



II. INTERNATIONAL APPLIED STATISTICS CONFERENCE (UYİK 2021)

In memory of Prof. Dr. Orhan DÜZGÜNEŞ



PROCEEDINGS BOOK

29 JUNE - 02 JULY 2021 TOKAT, TÜRKİYE



II. International
Applied Statistics
Conference **UYIK**
In memory of Prof. Dr. Orhan DÜZGÜNEŞ



II. INTERNATIONAL APPLIED STATISTICS CONFERENCE

Proceedings Book of the UYIK-2021

ISBN: 978-975-7328-80-3

HYBRID (ONLINE AND FACE TO FACE) - TOKAT / TÜRKİYE

Note: All administrative, academic and legal responsibilities of the departments belong to their authors

Proceedings Book of the II. International Applied Statistics Conference (UYIK-2021)

Publisher:

Tokat Gaziosmanpasa University

Editors:

Yalcin TAHTALI

E-Book Layout, Preparation and Composition:

Yalcin TAHTALI, Samet Hasan ABACI, Furkan YILMAZ

All published articles were peer-reviewed by Scientific Committee

The organizers do not have any legal liability for to contents of the presentation texts

Organized by

Tokat Gaziosmanpasa University, Tokat, Turkey

International Balkan University, Skopje, Macedonia

II. International Applied Statistics Conference (UYIK-2021)
Tokat / Turkey, 29 June - 2 July 2021

INVITED SPEAKERS

Common Mistakes in Statistics Applications and Solution Suggestions

Prof. Dr. Yuksel BEK (Retired Faculty Member)

Prof. Dr. Orhan KAVUNCU (Kastamonu University)

Prof. Dr. Zahide KOCABAS (Ankara University)

Prof. Dr. Zeynel CEBECI (Çukurova University)

Prof. Dr. Hayrettin OKUT (Kansas University)

Prof. Dr. Mehmet Ali CENGİZ (Ondokuz Mayıs University)

Assist. Prof. Dr. Emre DUNDER (Ondokuz Mayıs University)

Demonstrating the Central Limit Theorem with Simulation in Minitab

Prof. Dr. Tahsin KESİCİ (Retired Faculty Member)

Prof. Dr. Zahide KOCABAŞ (Ankara University)

Meeting of Sciences with Statistics; Historical Process

Prof. Dr. Hayrettin OKUT (Kansas University)

COMMITTEES

HONORARY CHAIRMAN

Prof. Dr. Mustafa Şentop, President of the Turkish Grand National Assembly

HONOURARY COMMITTEE

Dr. Ozan Balcı, Governor of Tokat

Prof. Dr. Bünyamin Şahin, Rector of Tokat Gaziosmanpaşa University

Prof. Dr. Mehmet Dursun Erdem, Rector of the International Balkan University

CONFERENCE CO-CHAIRMAN

Prof. Dr. İsa Gökçe, Tokat Gaziosmanpaşa University

Prof. Dr. Orhan Kavuncu, Kastamonu University

Prof. Dr. Rüstem Cangı, Tokat Gaziosmanpaşa University

Assoc. Prof. Dr. Yalçın Tahtalı, Tokat Gaziosmanpaşa University

CONFERENCE SECRETARIAT

Assoc. Prof. Dr. Samet Hasan Abacı, Ondokuz Mayıs University

Dr. Saniye Demir, Tokat Gaziosmanpaşa University

Dr. Nur İlkay Abacı, Ondokuz Mayıs University

PhD. Lütfi Bayyurt, Tokat Gaziosmanpaşa University

ORGANIZING COMMITTEE

Prof. Dr. Hasan Hüseyin Atar, Ankara University

Prof. Dr. Soner Çankaya, Ondokuz Mayıs University

Prof. Dr. Ecevit Eydurhan, Iğdır University

Prof. Dr. Aşir Genç, Necmeddin Erbakan University

Prof. Dr. Sedat Karaman, Tokat Gaziosmanpaşa University

Prof. Dr. Khalid Javed, University of Veterinary and Animal Sciences

Prof. Dr. İsmail Keskin, Selçuk University

Prof. Dr. Sıddık Keskin, Yüzüncü Yıl University

Prof. Dr. Zahide Kocabaş, Ankara University

Prof. Dr. İrfan Oğuz, Tokat Gaziosmanpaşa University

Prof. Dr. Hasan Önder, Ondokuz Mayıs University

Prof. Dr. Muhip Özkan, Ankara University

Prof. Dr. Şenay Sarıca, Tokat Gaziosmanpaşa University

Prof. Dr. İhsan Soysal, Namık Kemal University

Prof. Dr. Mustafa Şahin, Kahramanmaraş Sütçü İmam University

Prof. Dr. Arda Yıldırım, Tokat Gaziosmanpaşa University

Prof. Dr. Kenan Yıldız, Tokat Gaziosmanpaşa University
Prof. Dr. Kadri Yürekli, Tokat Gaziosmanpaşa University
Assoc. Prof. Dr. Melekşen Akın, Iğdır University
Assoc. Prof. Dr. Songül Akın, Dicle University
Assoc. Prof. Dr. Alper Serdar Anlı, Ankara University
Assoc. Prof. Dr. Şener Bıllal, International Balkan University
Assoc. Prof. Dr. Rukiye Dağalp, Ankara University
Assoc. Prof. Dr. Songül Gürsoy, Dicle University
Assoc. Prof. Dr. Dariusz Pıwczynski, UTP University, Poland
Assoc. Prof. Dr. Taner Tunç, Ondokuz Mayıs University
Assoc. Prof. Dr. Adnan Ünal, Niğde Ömer Halisdemir University
Dr. Yunus Akdoğan, Selçuk University
Dr. Yasin Altay, Eskişehir Osmangazi University
Dr. Emine Berberoğlu, Tokat Gaziosmanpaşa University
Dr. Yasemin Gedik, Eskişehir Osmangazi University
Dr. Kadir Karakaya, Selçuk University
Dr. Yeliz Kaşko Arıcı, Ordu University
Dr. Ahmet Pekgör, Necmettin Erbakan University
Dr. Dilek Sabancı, Tokat Gaziosmanpaşa University
Dr. Cem Tırınk, Iğdır University
Dr. Ömer Vanlı, İstanbul Teknik University
Dr. Fehmi Kiraz, Turkish Agricultural Engineers Association
Agr. Engineer Mete Türkoğlu, Iğdır Nature Conservation and National Parks Provincial Directorate

SCIENTIFIC COMMITTEE

Prof. Dr. Willy Arafah, Trisakti University, Indonesia
Prof. Dr. Hülya Atıl, Ege University
Prof. Dr. Mazroor Ahmad Bajwa, Balochistan University, Pakistan
Prof. Dr. Snezana Bilic, International Balkan University
Prof. Dr. Ömer Cevdet Bilgin, Atatürk University
Prof. Dr. Saim Boztepe, Selçuk University
Prof. Dr. Zeynel Cebeci, Çukurova University
Prof. Dr. Kabir Haruna Danja, Federal College of Education Zaira, Nigeria
Prof. Dr. Bejtulla Demiri, International Balkan University

Prof. Dr. Jacek Dlugosz, UTP University, Poland
Prof. Dr. Ercan Efe, Kahramanaraş Sütçü İmam University
Prof. Dr. Sait Ekinci, Kahramanaraş Sütçü İmam University
Prof. Dr. Ali Farajzadeh, Razi University, İran
Prof. Dr. Mehmet Ziya Fırat, Akdeniz University
Prof. Dr. Özkan Görgülü, Kırşehir Ahi Evran University
Prof. Dr. Wilhelm Grzesiak, West Pomeranian University of Technology, Poland
Prof. Dr. Asep Hermawan, Trisakti University, Indonesia
Prof. Dr. Özlem İlk Dağ, Orta Doğu Teknik University
Prof. Dr. Khalid Javed, University of Veterinary and Animal Sciences, Pakistan
Prof. Dr. Daniel Liberacki, Poznan University of Life Sciences, Poland
Prof. Dr. Violeta Madzova, International Balkan University
Prof. Dr. Mehmet Mendes, Çanakkale Onsekiz Mart University
Prof. Dr. Hikmet Orhan, Suleyman Demirel University
Prof. Dr. Emin Özköse, Kahramanaraş Sütçü İmam University
Prof. Dr. Renata Pilarczyk, West Pomeranian University of Technology, Poland
Prof. Dr. Levent Sangün, Çukurova University
Prof. Dr. Kire Sharlamanov, International Balkan University
Prof. Dr. Çiğdem Takma, Ege University
Prof. Dr. Mohammad Masood Tariq, Balochistan University, Pakistan
Prof. Dr. Muhammad Tariq, Agriculture University
Prof. Dr. Ali Reza Vaezi, Zanjan University, İran
Prof. Dr. Sanja Adzap Velickovska, International Balkan University
Prof. Dr. Abdulmojeed Yakubu, Nasarawa State University, Nigeria
Prof. Dr. Daniel Zaborski, West Pomeranian University of Technology, Poland
Prof. Dr. X.C. Zhang, Agricultural Research Center, United States
Assoc. Prof. Dr. Melekşen Akın, Iğdır University
Assoc. Prof. Dr. Ömer Alkan, Atatürk University
Assoc. Prof. Dr. Şenol Çelik, Bingöl University
Assoc. Prof. Dr. Sadiye Peral Eyduran, Iğdır University
Assoc. Prof. Dr. Yakut Gevrekçi, Ege University
Assoc. Prof. Dr. Farhat Iqbal, Balochistan University, Pakistan
Assoc. Prof. Dr. Abdulkarim Karaarslan, Atatürk University
Assoc. Prof. Dr. Natasha Kraveva, International Balkan University

Assoc. Prof. Dr. Dariusz Płwczynski, UTP University, Poland

Assoc. Prof. Dr. Roman Rolbiecki, UTP University, Poland

Assoc. Prof. Dr. Leman Tomak, Ondokuz Mayıs University

Assoc. Prof. Dr. Fatih Üçkardeş, Adıyaman University

Dr. Asım Faraz, Bahauddin Zakariya University, Pakistan

Dr. Emin Idrızı, International Balkan University

Dr. Abdurrahman Kara, Dicle University

Dr. Hande Küçükönder, Bartın University

Dr. Lujeta Sadıku, International Balkan University

Dr. Liza Halili Sulejmani, International Balkan University

Dr. Thobela Louis Tyasi, Limpopo University, North Africa

Dr. Abdul Waheed, Bahauddin Zakariya University, Pakistan

Note: Given in order of title and surname

PROF. DR. ORHAN DÜZGÜNEŞ'İN ÖZGEÇMİŞİ



Prof. Dr. Orhan DÜZGÜNEŞ: Büyük ama Mütevazı Bir İlim ve Dava Adamı,

Prof. Dr. Orhan DÜZGÜNEŞ 1917 yılında, İstanbul Rumeli Kavağı'nda doğdu. Babası polis memuru Aziz Efendi Bulgaristan'ın Lofça şehrinden, annesi ise Geredeliydi. İlkokulu Gerede'de, liseyi Kastamonu'da bitirdi. 1938 yılında Yüksek Ziraat Enstitüsünden birincilikle mezun oldu. Aynı yıl Yüksek Ziraat Enstitüsünün Zootečni Enstitüsüne asistan oldu.

1946 yılında doktorasını tamamladı ve Amerika Kaliforniya Üniversitesine genetik çalışmalarda bulunmak üzere gönderildi. Yüksek lisansı ve bilimsel araştırmalarıyla temayüz eden hocamız, ancak seçkin bilim adamlarının üye olduğu "Sigma XI" isimli araştırma derneğine üyelik teklifi aldı. 1950 yılında yurda döndükten sonra Ankara Üniversitesi Ziraat Fakültesi Zootečni Bölümünde verdiği Hayvan Islahı, Sığırcılık gibi dersler yanında Genetik ve Biyometri derslerini vermeye başlamış, bir süre sonra da Ziraat Genetik ve İstatistik Kürsüsünü kurmuştur. 1951'de doçent, 1957'de profesör olan Orhan DÜZGÜNEŞ, Türkiye'de Araştırma ve Deneme Metotları, Deneme Planlaması, Popülasyon Genetiği, Kantitatif Genetik gibi disiplinlerin Ziraat Fakültelerinin bilimsel çalışmalarında önemsenmesinde başlıca rol oynayan büyük bir ilim adamıdır.

Hayatı boyunca bilimsel çalışmalardan geri kalmayan Prof. Dr. Orhan DÜZGÜNEŞ, Macaristan, İskoçya, Amerika Birleşik Devletleri, İtalya, İspanya, Fransa, İsrail, Almanya ve Azerbaycan'da incelemeler yapmıştır. Hacettepe ve Ankara Üniversitesi Tıp Fakültelerinde İstatistik derslerini başlatan ve biyoistatistik bilim dalının kurulmasına öncülük eden bilim adamı Prof. Dr. Orhan DÜZGÜNEŞ'tir.

Bilimsel çalışmaların yanında, idari görevlerinden kaçınmayan DÜZGÜNEŞ, 1968-70 yıllarında Ankara Üniversitesi Ziraat Fakültesi Dekanlığı, 1984'te Ankara Üniversitesi Fen Bilimleri Enstitüsü Müdürlüğü yapmıştır. İki dönem Ziraat Mühendisleri Odası Başkanlığı yapmış, Türk Ziraat Mühendisleri Birliği Vakfı Başkanlığını vefat edinceye kadar uhdesinde taşımıştır. 1986 yılında TÜBİTAK "Hizmet Ödülü" verilmiştir. Tarım bakanlığında hayvan ıslahıyla ilgili çeşitli araştırma projesi, komisyon başkanlığı görevlerini hizmet şuuruyla bedelsiz olarak yapan Prpf. Dr. Orhan DÜZGÜNEŞ "Türkiye'deki Tabii Gen Kaynaklarının Korunması" çalışmalarına da hayatının son zamanlarına kadar aktif bir şekilde katılmıştır.

Prof. Dr. Orhan DÜZGÜNEŞ ciddi, alçak ve gani gönüllü bir insan, bilim ve dava adamıdır. Gerçek bir dost, büyüklerine saygılı, küçüklerine gerçek bir hami idi. Hocası Prof. Dr. İbrahim Yarkın'a sevgi dolu bir saygıyla muamelesi, hepimize örnek olacak bir nitelikteydi. Kin tutmazdı, affetmesini bilirdi. Dargınlığın müslümanlara yakışmadığını peygamber efendimizin hadislerinde ifade ettiklerini söyler, yüzde yüz haklı bile olsa veya doğrudan taraf dahi olmadığı, sonucu hakaret etmeye varan sert tartışmalardan bir müddet sonra, bu hadisin gereği olarak davranır, muhataplarından helallik ister ve büyüklük onda kalırdı. Neyse oydu. Hz. Mevlana'nın "olduğun gibi görün" ilkesi hoca ile bütünleşmişti. Hoca için küçük iş yoktu. Her işi ciddiye alırdı. Öğrenciler anlamadığı bir konuyu rahatlıkla sorabilirlerdi. Öğrenmek isteyen geri çevirmezdi.

İlmi araştırma zevkle yaptığı bir meslek idi. Bununla birlikte, hayvan ıslahı gibi doğrudan uygulamaya dayanan bir alanda çalışıyordu. Ayrıca günlük olaylardan, kısaca hayattan kopuk olmayan renkli bir bilim adamıydı. Bütün bu özellikleri, pratik zekasıyla da birleşince ortaya nazari bilgilerini ve bunlardan geliştirdiği kanaatlerini, hayatın her alanında kullanabilen meziyet sahibi, erdemli ve mütevazı bir insan profili ortaya çıkmıştır.

“İnsan” olmanın “alim” olmanın yanına “milliyetçi” olmayı da ekleyince, ortaya Orhan DÜZGÜNEŞ misali bir “münevver” çıkıyordu. Samimi bir Atatürk sevgisi, Atatürk’ü yakından görmüş, belki onunla sohbet eden gençler arasında yer almış olmaktan kaynaklanan doğal bir sevgiydi. Cumhuriyetin ilk neslindendi. Yıkılmış bir imparatorluğun enkazı içinden yeni bir devlet kurulmasını arzu eden bu ilk nesildeki Atatürk sevgisi berrak, saf ve samimiydi.

Almanca ve İngilizce bilen, çok sayıda ilim adamının yetişmesine yardımcı olan DÜZGÜNEŞ, konusuyla ilgili 14 kitap ve yüzün üzerinde bilimsel inceleme ve araştırma makalesi yayınlamıştır. Evli ve iki çocuk babası olan Prof. Dr. Orhan DÜZGÜNEŞ 27 Haziran 1996’da Ankara’da geçirmiş olduğu kalp krizi sonucu vefat etmiş ve Karşıyaka mezarlığına defnedilmiştir.

Hocam, sen gideli 25 yıl oldu, kocaman 25 yıl... Ve senin bizlere “evlatlarım akıllı olun ülkeye, birliğe, dirliğe ve tarıma sahip çıkın” nasihatini unutmamaya çalışıyor ve seni çok özledik...

DÜZENLEME KURULU

BIOGRAPHY OF PROF. DR. ORHAN DUZGUNES

Prof. Dr. Orhan Duzgunes: A Great but Modest Man of Science and a Man with a Cause

Orhan DUZGUNES was born in 1917 in Rumeli Kavağı in Istanbul. His father, Aziz Efendi, was a police officer from the Bulgarian city of Lovech (Lofça), and his mother was from Gerede. He finished primary school in Gerede and high school in Kastamonu in Turkey. He graduated from the Higher Agricultural Institute in 1938 with the first rank. In the same year, he became an assistant at the Animal Science Institute of the High Agricultural Institute.

He completed his doctorate in 1946 and was sent to the University of California, USA to study genetics. Our professor with the outstanding performance he displayed in his master's degree and scientific research, received a membership offer to the research association "Sigma Xi", where only distinguished scientists are members. After returning to Turkey in 1950, he started to give courses such as Animal Breeding and Cattle Breeding as well as Genetics and Biometrics in the Department of Animal Science at Ankara University Faculty of Agriculture. After a while, he established the Chair of Agricultural Genetics and Statistics. Orhan DUZGUNES, who became an Associate Professor in 1951 and a professor in 1957, was a great scholar who played a major role in ensuring that the disciplines of Research and Experimental Methods, Experimental Planning, Population Genetics, Quantitative Genetics are deemed important at the scientific studies of faculties of agriculture in Turkey.

Orhan DUZGUNES, who did not cease to perform scientific studies throughout his life, carried out research in Hungary, Scotland, and twice in America, Italy, Spain, France, Israel, Germany, and Azerbaijan. It is the scientist Orhan DUZGUNES who started the statistics courses in Hacettepe and Ankara University Medical Faculties and pioneered the establishment of the biostatistics branch.

DUZGUNES, who did not abstain from administrative roles, was the Dean of the Faculty of Agriculture at Ankara University in 1968-70 and the Director of the Institute of Science at Ankara University in 1984 in addition to his scientific studies. He was the President of the Chamber of Agricultural Engineers for two terms and the President of the Turkish Association of Agricultural Engineers Foundation until his death. In 1986, he was given TÜBİTAK "Service Award". Orhan DUZGUNES, who carried out various research projects on animal breeding and the chairmanship of the commission in the Ministry of Agriculture with the responsibility of service free of charge, actively participated in the studies of "Conservation of Natural Genetic Resources in Turkey" until the end of his lifetime.

Prof. Dr. Orhan DUZGUNES was a serious, modest, and generous person, scientist and a man with a cause. He was a true friend, respectful of elders, and true support for the young. His loving and respectful treatment of his teacher Prof. Dr. İbrahim YARKIN was an example for all of us. He would never hold grudges and knew how to forgive. He used to say that our Prophet stated that resentment was not suitable for Muslims in the hadiths. Even after harsh arguments that resulted in insults, he used to act as a requirement of this hadith and asked for people's blessings even if he was 100% right or not even a party and kept his greatness. He was what he seemed. Hz. Mevlana's principle of "Be as you seem" was reflected in him. No task was small for him. He took every task seriously. Students could easily ask a topic they did not understand. He would not refuse anyone who wanted to learn.

Scientific research was a profession that he performed with delight. Besides, he was working in a direct practice field like animal breeding. He was also a colourful scientist who was not detached from daily events, in short, from life. Therefore, unlike many, he was one of the rare people who internalized his knowledge and merged it with his life experience. Combining all these characteristics with his practical intelligence, a versatile, virtuous and modest person who could use his theoretical knowledge and the opinions he drew from them in every aspect of life emerged.

Adding being a "nationalist" to being a "human" in addition to being a "scholar," an "intellectual" like Orhan DUZGUNES emerged. A sincere love for Atatürk was maybe a natural love emerging due to being one of the young people who saw Atatürk closely and chatted with him. He was from the first generation of the republic. The love for Atatürk, who wanted to establish a new state from the ruins of a collapsed empire, was clear, pure, and sincere in this first generation.

Duzgunes, who spoke German and English and helped raise many scientists, published 14 books and over a hundred scientific studies and research articles on his subject. Prof. Dr. Orhan DUZGUNES, who was married and had two children, died on June 27, 1996, as a result of a heart attack he had in Ankara and was buried in Karşıyaka cemetery.

Dear Professor, it has been 25 years since you left, a huge 25 years... We are trying to heed your advice " My children, be wise, protect the country, unity, life and agriculture," we miss you very much...

ORGANIZING COMMITTEE

ÖNSÖZ

Çok Değerli Bilim İnsanları ve Değerli Dostlarımız,

Tokat Gaziosmanpaşa Üniversitesi ev sahipliğinde ve Uluslararası Balkan Üniversitesi iş birliği ile ikincisini düzenlediğimiz Uluslararası Uygulamalı İstatistik Kongresi, Türkiye Büyük Millet Meclisi Başkanı, kongre onursal başkanımız Sayın Prof. Dr. Mustafa ŞENTOP'un ve geniş bir akademisyen grubunun katkı ve katılımları ile açılmış, hem yüz yüze hem de çevrimiçi olarak gerçekleştirilmiştir.

Bilindiği üzere, 2020 Ekim ayında pandemi şartlarında birinci kongremizi Prof. Dr. Yüksel BEK adına gerçekleştirmiştik. Bu yıl ise, vefatının yirmi beşinci yılında, istatistik ve genetik bilimine ömrünü vermiş ve ülkemizde nice hocaların yetişmesinde büyük emekleri olan merhum Prof. Dr. Orhan DÜZGÜNEŞ hocamızın anısına kongremizi gerçekleştirmenin onurunu yaşadık.

Kongre düzenleme kurulu olarak, bu kongreleri düzenlerken, iki tema üzerinde duruyoruz; gençlere yatırım ve bu bilime emeği geçenlere vefa. Bu vesile ile kongrelerimizi, hatıralarını yaşatmak için hocalarımızın adına düzenlerken, kongre bünyesinde kıymetli hocalarımızın da emekleri ile muhtelif kurslar düzenleyerek genç araştırmacıların yanında olmaya gayret gösteriyoruz.

Bu amaçla kongre kapsamında 6 farklı kurs programı düzenlenmiş ve bu kurslara yaklaşık 156 kursiyer katılmıştır. Ayrıca kongreye 8 farklı ülkeden yaklaşık 98 araştırmacı 113 bildiri ile katıldı. Değerli hocalarımızın katılımıyla “İstatistik Çalışmalarındaki Hatalar ve Çözüm Önerileri” konulu panel düzenlendi.

Bu kongre hem yüz yüze hemde çevrim içi olarak gerçekleştirilmiştir. Çevrim içi salonlarımıza, ahde vefa göstergesi olarak yakın zamanda kaybettiğimiz Iğdır üniversitesi öğretim üyesi Prof. Dr. Ecevit EYDURAN'IN anısına “**Prof. Dr. Ecevit EYDURAN**” salonu ve Makendonya’da ikinci dünya savaşı sırasında önce Bulgarlara, daha sonra baskıcı Tito yönetimine karşı kurulan; Türklerin mevcudiyeti, kimlikleri ve inançları için mücadele eden bir “Türk mukavemet teşkilatının adı olan Yücelciler hareketine ithafen “**Yücelciler Salonu**” isimleri verilmiştir.

Bu kongre vesilesiyle, kongrenin onursal başkanı ve Türkiye Büyük Millet Meclisi Başkanı Sayın Prof. Dr. Mustafa ŞENTOP'a bu konferansın düzenlenmesine maddi ve manevi katkılarından dolayı; başta Tokat Gaziosmanpaşa Üniversitesi Rektörü Sayın Prof. Dr. Bünyamin ŞAHİN olmak üzere tüm üniversite yönetimine, Uluslararası Balkan Üniversitesi Rektörü Sayın Prof. Dr. Mehmet Dursun ERDEM ve Ankara Üniversitesi Ziraat Fakültesi Dekanı Sayın Hasan Hüseyin ATAR'A teşekkürlerimizi sunarız.

Ayrıca kongrenin ve kursların düzenlenmesinde emeği geçen, tüm hocalarımıza ve çalışma arkadaşlarımıza da teşekkürü borç biliriz.

Kongremizin üçüncüsünü “3. Uluslararası Uygulamalı İstatistik Kongresi”, yine sizlerin etkin katkı ve katılımlarıyla, kıymetli hocamız Prof. Dr. Tahsin KESİCİ adına Haziran 2022’de Uluslararası Balkan Üniversitesi ev sahipliği ve Tokat Gaziosmanpaşa Üniversitesi iş birliği ile Üsküp’te gerçekleştirmeyi planlıyoruz. Sizleri, yılın en güzel dönemlerinin birinde, Makedonya’da ağırlamaktan onur ve mutluluk duyacağız.

Kongremizin tüm uygulamalı istatistik alanlarına katkı sağlamasını, genç araştırmacılara faydalı olmasını ve bilime katkı sunanlar açısından ahde vefaya örnek oluşturmasını dileriz.

Kongre Düzenleme Kurulu

PREFACE

Dear Scientists and Dear Friends,

The second International Applied Statistics Conference, hosted by Tokat Gaziosmanpaşa University and in cooperation with the International Balkan University, was opened with the contributions and participation of the President of the Turkish Grand National Assembly and our honorary president of the congress Prof. Dr. Mustafa SENTOP and a large group of academics and was held both face-to-face and online.

As it is known, we held our first congress in the name of Prof. Dr. Yuksel BEK in October 2020 under pandemic conditions. This year, on the twenty-fifth anniversary of his death, we had the honor of holding our congress in memory of our late Prof Dr. Orhan DUZGUNES who gave his life to statistics and genetics and had great efforts in the training of many teachers in our country.

As the congress organizing committee, we focus on two themes while organizing these congresses; investment in young people and fidelity to those who contribute to this science. On this occasion, while we organize our congresses in the name of our professors to keep their memories alive, we try to stand by the young researchers by organizing various courses with the efforts of valuable teachers within the congress.

For this purpose, 6 different course programs were organized within the scope of the congress, and approximately 156 trainees participated in these courses. In addition, 98 researchers from 8 different countries participated in the congress with 113 proceedings. A panel on “Mistakes made in statistics studies and solution suggestions” was held with the participation of our valuable professors.

This congress was held both face-to-face and online. As a sign of loyalty, our online halls were named "**Prof. Dr. Ecevit EYDURAN Hall**" in memory of Prof. Dr. Ecevit EYDURAN, Iğdır University faculty member we lost recently and "**Yucelciler Hall**" in honor of Yücelciler movement, which is the name of a "Turkish resistance organization" that struggled for the existence, identity and beliefs of the Turks, established in Macedonia during the Second World War, first against the Bulgarians and then against the oppressive Tito administration.

On the occasion of the conference, we want to thank the honorary president of the conference and the President of the Turkish Grand National Assembly, Prof. Dr. Mustafa SENTOP, for his material and moral contributions to the organization of this conference; especially, Tokat Gaziosmanpasa University rector Prof. Dr. Bünyamin SAHİN and the entire university administration, the rector of the International Balkan University, Prof. Dr. Mehmet Dursun ERDEM and Ankara University Faculty of Agriculture Dean Hasan Hüseyin ATAR.

We would also like to thank all our professors and colleagues who contributed to the organization of the congress and courses.

We are planning to hold the third of our congress, “3. International Applied Statistics Conference”, in the name of our esteemed professor, Prof. Dr. Tahsin KESİCİ, again with your active contribution and participation, in Skopje in June 2022, hosted by the International Balkan University and with the cooperation of Tokat Gaziosmanpaşa University. We will be honored and happy to host you in Macedonia during one of the most beautiful times of the year.

We hope that our congress will contribute to all applied statistics fields, be beneficial to young researchers and set an example for those who contribute to science.

Congress Organizing Committee

TABLE OF CONTENT

COMMITTEES	i
PROF. DR. ORHAN DÜZGÜNEŞ'İN ÖZGEÇMİŞİ	v
BIOGRAPHY OF PROF. DR. ORHAN DUZGUNES	vi
ÖNSÖZ.....	viii
PREFACE	ix
TABLE OF CONTENT	x
ONLINE PRESENTATIONS	xi
ONLINE PRESENTATIONS CONTENT.....	xii
POSTER PRESENTATIONS	388
POSTER PRESENTATIONS CONTENT.....	389
NATIONALITY OF PRESENTER AUTHORS	392

ONLINE PRESENTATIONS

ONLINE PRESENTATIONS CONTENT

Şiddet-Süre-Frekans Bağıntısının Parçacık Sürü Optimizasyonu ve Genetik Programlama ile Belirlenmesi.....	1
Faktör Analizinde Faktör Puanlarının Hesaplanma Yöntemlerinin Karşılaştırılması	9
Üniversite Öğrencilerinin Bireysel Girişimcilik Algı Durumları ile İletişim Becerileri Arasındaki İlişkinin Yapısal Eşitlik Modeli ile İncelenmesi	10
On the Study of a Nonlinear Fractional-order Boundary Value Via Fixed Point	11
Boundedness and Stability For Generalized Schrödinger Equation.....	12
Hibrid Panel Veri Analiz Metotolojisi Üzerine Uygulama: Ulusal Harcamalar Örneği	13
Ankara İli Ekstrem Yağışlarının İstatistiksel Analizi.....	14
Evaluation of Glaucoma with Factor Analysis.....	21
Sağlık Bilimlerinde Seyrek Olumsuzluk Tabloları Üzerine Bir Çalışma	27
Quasi-Least Squares Regression Method with Veterinary Vaccination Data	28
Biyolojik Uygulamalarda Genlerin veya Özelliklerin Seçilmesi için Genetik Algoritmalarla Dayalı Makine Öğrenimi	29
Wine Veri Setinin Hiyerarşik Kümeleme Analizi Yöntemi ile İncelenmesi.....	30
Parametrik ve Parametrik Olmayan Korelasyon Analizinde Korelasyon Katsayısının Kullanımı ve Yorumlanması	37
Çoklu Durum Sistem Modellerinin Güvenilirlik Değerlendirmesi	38
Mean Estimation for Sensitive Study Variable Using Ln-type Estimators under Stratified Sampling. 39	
2-Boyutlu Ardıl-k Sistemlerin Güvenilirliği için Yaklaşımlar	40
Computational Insight into Analysis of the Effects of Deleterious SNPs of MAPT Gene in Alzheimer's Disease	41
Identification of Post-Translational Modifications Mediated by High-Risk Deleterious SNPs of Human DRD2 Gene Associated with Schizophrenia.....	42
Parametrik Olmayan Çok Değişkenli Varyans Analizi ve Beyaz Eşya Sektöründe Bir Uygulama.....	43
Daimî İkametgâh ve Suç Türüne Göre İllerin (İBBS 3. Düzey) Kümeleme Analizi ile İncelenmesi... 51	
Classification and Regression Tree as Data Mining Algorithm to Predict Body Weight of Boer Goats	58
Using QLS Method in Analysis of the Data Obtained from the Experimental FMD Vaccines.....	62
Determining Optimum Sample Size for Regression Tree and Automatic Linear Modeling.....	63
Tarla Silindirinin Farklı Ağırlıklarının Toprağı Sıkıştırma Olan Etkisinin Sonlu Elemanlar Yöntemiyle Analiz Edilmesi	68
Toprak İşlemesiz (Doğrudan Ekim) Tarımda Verim Üzerine Etkili Faktörlerin Araştırılmasına Yönelik Bir Meta-Analiz Çalışması	73
The Effect of Body Phenotypic Traits on Live Weight in Ovine Breeding: A Meta-Analysis.....	74
Farklı Aktivasyon Fonksiyonlarında Yapay Sinir Ağları Çalışması: Hayvansal Üretimde Bir Uygulama	83
Regression Analysis of Morphological Traits on Body Weight in Female Dorper Sheep Lambs.....	93
Classification of High Dimensional Imbalanced Data Using PCA, SMOTE Methods and Classification Algorithms.....	97
Denetimli İstatistiksel Öğrenme Algoritmaları ile Toprak Radon Gazının Tahmini	104
Fındıkta Yürütülen Seleksiyon Islahı Çalışmaları Özelliklerinin Değerlendirilmesi.....	105
Üretici Birliklerinin Örgütlenme Bilincine Katkıları; Telli Terbiye Sistemli Bağcılık Örneği.....	114
A New Generalized Pareto Distribution with Applications	115
On Evaluating Big and Complex Data Sets: Opportunities and Challenges	116

II. International Applied Statistics Conference (UYIK-2021)
Tokat / Turkey, 29 June - 2 July 2021

İyi Kaldıraç Noktalarından Etkilenmeyen Sağlam Varyans Artış Faktörü	117
Sıralı Lojistik Regresyon Yönteminin Sınıflama Performansının İncelenmesi	118
Makine Öğrenmesi Yöntemleri ve Algoritmalarının Kullanımı	119
Distance Based Regression Models.....	120
Moleküler Verilerde AMOVA'nın Kullanımı.....	127
Comparing the Performance of Eight Imputation Methods for Propensity Score Matching in Missing Data Problem.....	137
Bulanık Yapay Sinir Ağları ve Bir Uygulama	138
Comparison of Restricted and Unrestricted Biased Estimators and an Application	140
Hidrolojide Yapay Zekâ ve Gri Sistem Modeli ile Kuraklık Analizi.....	141
Lojistik Dağılım için Farklı Parametre Tahmin Metotlarının Karşılaştırılması	154
Türkiye’de Tüketici Fiyat Endeksi ile Yurt İçi Üretici Fiyat Endeksi, Yurt Dışı Üretici Fiyat Endeksi ve Tarım Ürünleri Üretici Fiyat Endeksi Arasındaki İlişki Analizi	162
Türkiye’de Çiğ Süt ve Yem Fiyatları Arasındaki Nedensellik İlişkisinin Analizi	169
Gaziantep İli İslahiye İlçesinde Bulunan Tahtaköprü Barajında Seviye Yükseltilmesiyle Sular Altında Kalacak Parsellerin Tesbiti.....	170
Sabit Kameralı Derin Öğrenme Fiziksel Şiddet Tespiti Sistemi için Aşırı Hareket Öncül Fonksiyonu	179
İlişki Ölçüleri Testleri	185
Taşkın Frekans Analizinde Kullanılan Yöntemler	186
Türkiye ve AB Ülkelerinin Covid-19’a Karşı Etkinliğinin Veri Zarflama Analizi ile Karşılaştırılması	210
Logistic Regression Analysis in Health Sciences Researches: An Application in SPSS.....	211
TOPSIS Yöntemi ile Öznitelik Seçimi.....	212
Evaluation of 366 Ridge Parameter Estimators in Terms of Their Normal Distribution, Generating Values of 0-10 and MSE Criteria.....	213
İşgücü Piyasası Güvensizliğini Etkileyen Faktörler: OECD Üyesi Ülkeler Üzerine Bir Değerlendirme	214
Information Entropies Which Are Based on Distance Matrix	215
A Survival Simulation Study on Usage of Train-Test Splitting, Cross Validation and Bagging Methods on Training Weibull, Exponential, Log-normal, Log-logistic, Quantile and Cox Regression Models.....	216
Ölümlülüğün Ufrs 17 Üzerindeki Etkisi.....	218
İnsani Gelişmişlik Endeksinin Çok Kriterli Karar Verme Yöntemlerine Dayalı Normal Karma Modeller ile Sınıflandırılması için Yeni Bir Yaklaşım	219
A Novel Approach for the Determination of Entropy and Compactivity of Granular Matter in a Tube	220
Genetic Defects Associated with Reproduction in Dairy Cattle	221
Residual-based Control Charts Under Liu Estimator For Monitoring Conway-maxwell-poisson Profiles	222
Bağlantı Dengesizliğinin Genotipik Varyans ve Unsurlarına Etkisi Üzerine bir Simulasyon Çalışması: I- Bir lokusta Dominans olan Modelde Populasyon Parametreleri	223
CfsSubsetEval ve MARS Yöntemleri ile BİST 100 Endeksi için Model Seçimi	224
A Unit - Cauchy Distribution: Definition, Properties and Application	225
On a Unit-eeg Quantile Regression Analysis.....	226
Varyasyon Katsayıları Arasındaki Farka Ait Örnekleme Dağılımının Şekli ve Güven Sınırları.....	227
A Simulation Study on the Comparison of the Methods Used in Comparing the Proportions and Hotelling T ² Method.....	228

II. International Applied Statistics Conference (UYIK-2021)
Tokat / Turkey, 29 June - 2 July 2021

Estimation and Prediction on the Half-Normal Distribution Based on Progressively Type-II Censored Samples	229
Bayesian, E-Bayesian and E-MSE of Half-Normal Distribution Parameter Under Different Loss Functions	230
Dengesiz Veri Setlerinde Farklı Dengeleme Algoritmalarının Optimum Denge Oranlarının Sınıflandırma ve Regresyon Ağaçları Yöntemi ile İncelenmesi: Simülasyon Çalışması.....	231
Bootstrap Confidence Intervals of Generalized Discrete process capability index C _{pyk}	233
Some Count Regression Models from Generalized Linear Models Approach to Investigate Pakistan Child Mortality Under-5 Data	234
Forecasting BIST 100 Index Using An Artificial Neural Network.....	235
Multiple Imputation Methods for Numerical Datasets.....	245
Türkiye’de Covid-19 Ölümünün İstatistiksel Süreç Kontrolü	246
Arctanh-Exponentiated Exponential Distribution and Applications	258
Lisansüstü Öğrenci Alım Mülakat Kriterlerinin Analitik Serim Süreci ile Ağırlıklandırılması	259
Öğrenci Elektronik Belge Yönetim Sistemi Tasarım/Seçim Kriterlerinin DEMATEL ile Değerlendirilmesi	260
Panel Veri Analizi ile Küresel Rekabet Endeksi, İhracat ve Büyüme Oranının İnovasyon Üzerine Etkisi.....	261
COVID-19 Hastalarına İlişkin Eksik Gözlem Durumunda Makine Öğrenimi Algoritmalarının Kullanımı.....	273
An Empirical Study on High Dimensional, Censored Survival Data Analysis with Supervised Principal Components, Extreme Learning Machine-Based and Penalized Cox Models	274
Kayıp Veri Analiz Yöntemleri Üzerine Bir İnceleme	275
Ülkelerin Endüstri 4.0 Etkilerinin İnsan, Makine ve Ürün Bağlamında Kümeleme Analizi	290
On the Existence and Uniqueness of Maximum likelihood Estimates for Unit Distributions	305
Yapısal Eşitlik Modellemesi ve Bir Uygulama	306
Investigating the Effect of Land Use and Hillslope Characteristics on Morphological Changes of Gullies in the West of Ardabil Province.....	323
Aşırı Yayılımlı Sayım Verilerinin Analizi	324
Petrol Fiyat Endeksi ve Doğrusal Olmayan Yaklaşım	325
Petrol Fiyat Endeksinin Model Yapısı: Markov Rejim Değişim Modeli	332
Some Factors Affecting Gully Erosion in Southwestern Hamadan, West of Iran	338
Comparison of Wood, Single and Two Knotted Cubic Spline Regression Models for Modeling Lactation Curves of Holstein Cattle in First Lactation.....	339
Genel Özellikleri ile İkili (Binary) Logit ve Probit Modellerin İncelenmesi	340
Performance Analysis of Electricity Distribution Companies in Turkey with Data Envelopment Analysis and (0,1) Regression Analysis	341
Çoklu Bağlantı Probleminde Alternatif Tahmin Edicilerin Performanslarının Karşılaştırılması.....	342
Samsun İli Ondokuzmayıs İlçesinde Son Yıllarda Çilek Üretiminin Artmasının Nedeni Kârlılığı Olabilir mi?	354
Prediction of Egg Production in Japanese Quails (<i>Coturnix coturnix japonica</i>) by Linear Regression, Linear Piecewise Regression, and Mars Algorithm	355
A New Quantile Regression Model and Application	356
Genelleştirilmiş Procrustes Analizi ve Üzümde Çeşit ile Mineral Madde Arasındaki İlişkinin İki Boyutlu Uzaydaki Konfigürasyonu	357
Tabakalı Medyan Sıralı Küme Örneklemesinde Örnek Çapının Paylaştırılması: En Uygun Paylaştırma Yöntemi.....	358
İstatistiksel Programlarda MARS Modelinin Elde Edilişleri Arasındaki Farklar	359

II. International Applied Statistics Conference (UYIK-2021)
Tokat / Turkey, 29 June - 2 July 2021

Increasing Sample Size Through Appropriate Use of Rolling Windows.....	360
Bayes Estimation of Stress-Strength Reliability for Unit Burn Type XII Distribution Based on Progressive Type-II Censoring.....	361
Permütasyonel Çok Değişkenli Varyans Analizi (PERMANOVA) ve Covid-19 Verisi ile Uygulama	362
Bayes Estimation for the Gompertz Gumbel Type-II Distribution Based on Progressive Type-II Censored Samples	363
Türkiye ve AB Ülkelerinin Covid-19' a Karşı Etkinliğinin Veri Zarflama Analizi ile Karşılaştırılması	364
Sağlık Bilimi Verilerinde Bir Uygulama ile Doğrusal Olmayan Temel Bileşen Analizinin Tanıtımı	376
Ülkelerin Sağlık Ekonomik ve Gelişmişlik Göstergeleri Arasındaki İlişki: Kanonik Bir Yaklaşım ..	377
Ölçek Geçerlilik ve Güvenilirliğinde Farklı Bir Yaklaşım	378
Goodness of Fit Tests Based on Kullback-Leibler Divergence under Censoring for Weibull Distribution.....	379
Investigating the Effect of Drainage Area Characteristics on Gully Density in North of Iran.....	387

Şiddet-Süre-Frekans Bağıntısının Parçacık Sürü Optimizasyonu ve Genetik Programlama ile Belirlenmesi

Eyyup Ensar BAŞAKIN^{1*}, Ömer EKMEKCİOĞLU¹, Mehmet ÖZGER¹, Hatice ÇITAKOĞLU²

¹İstanbul Teknik Üniversitesi İnşaat Mühendisliği Bölümü, İstanbul, Türkiye

²Erciyes Üniversitesi İnşaat Mühendisliği Bölümü, Kayseri, Türkiye

*Sunucu yazar e-posta: basakin@itu.edu.tr

Özet

Su yapılarının tasarım aşamasında, farklı frekanslarda gerçekleşmesi muhtemel yağış şiddeti değerlerine ihtiyaç duyulmaktadır. Bu değerlerin hesabı için şiddet-süre-frekans bağıntısının elde edilmesi büyük önem taşımaktadır. Bu çalışmada, Kayseri iline ait uzun süreli standart zamanlara ait en büyük yağış değerleri yardımıyla yağış şiddetinin hesaplanabileceği bir denklem elde edilmiştir. Yağış şiddeti, yağış süresi ve frekansın bir fonksiyonu olarak tasarlanmıştır. Denklem oluşturma işlemi ise Genetik Programlama (GP) yöntemi kullanılarak gerçekleştirilmiştir. Kayseri istasyonuna ait veriler ile oluşturulan denklem Nevşehir, Niğde ve Yozgat istasyonlarında test edilmiştir. Gözlenen ve tahmin edilen yağış şiddeti değerleri ortalama karesel hata (OKH) ve Verimlilik Katsayısı (NSE) ölçütleri ile değerlendirilmiştir. Kayseri istasyonuna ait OKH ve NSE değerleri sırasıyla 0.004 ve 0.99 olarak hesaplanmıştır. Nevşehir istasyonuna ait OKH ve VK değerleri sırasıyla 0.028 ve 0.91, Niğde istasyonuna ait OKH ve NSE değerleri sırasıyla 0.067 ve 0.76 son olarak Yozgat istasyonu için OKH ve NSE değerleri sırasıyla 0.027 ve 0.95 olarak hesaplanmıştır. Çalışmada ayrıca literatürde bulunan ampirik denklemlerden biri karşılaştırma için kullanılmıştır. Ampirik denklemin parametrelerinin optimizasyonu ise Parçacık Sürü Optimizasyonu (PSO) kullanılarak gerçekleştirilmiştir. GP ile elde edilen denklemin PSO yardımıyla optimize edilen denklemlere nazaran daha başarılı tahminlerde bulunduğu saptanmıştır.

Anahtar kelimeler: Genetik programlama, Şiddet-süre-frekans, Parçacık sürü optimizasyonu, Hidroloji

Determination of Intensity-Duration-Frequency Relation by Particle Swarm Optimization and Genetic Programming

Abstract

Rainfall intensity values at different frequencies are required to design the water structures. Thus, obtaining the intensity-duration-frequency relationship has vital importance. In this study, an equation was obtained in which the rainfall intensity can be calculated with the help of the intensity-duration-frequency values of long-term standard periods of Kayseri. Rainfall intensity is designed as a function of both rainfall duration and frequency. Genetic programming (GP) was used to generate the equation and obtained equation validated with the data taken from other stations in the region, such as Nevşehir, Niğde and Yozgat. The observed and predicted rainfall intensity values were evaluated with two different performance indicator, namely mean square error (MSE) and coefficient of efficiency (NSE) criteria. For Kayseri station, the MSE and NSE are calculated as 0.004 and 0.99, respectively. Also, the MSE and NSE values obtained for Nevşehir and Niğde stations are 0.028 and 0.91, and 0.067 and 0.76, respectively, while for Yozgat station, MSE was determined as 0.027 and NSE was found as 0.95. Furthermore, one of the empirical equations in the literature was used to perform a comparison. The parameter optimization for the empirical equation was carried out using particle swarm optimization

(PSO). Finally, it was concluded that the equation obtained with GP has higher accuracy than the obtained equation with the help of PSO.

Key words: Genetic programming, Intensity-duration-frequency, Particle swarm optimization, Hydrology

GİRİŞ

Yerleşim alanlarının planlanmasında alt yapı sistemleri büyük önem taşımaktadır. Alt yapı sistemlerinden en önemlisi ise yağmur sularının ve kullanılan suların tahliye edildiği kanalizasyon sistemleridir. Bu sistemlerin doğru şekilde planlanıp inşa edilmesi hayati önem arz etmektedir. Herhangi bir sağanak sonrası alt yapı sistemi optimum şekilde çalışmalıdır. Aksi durumlarda su baskınları meydana gelebilmekte, maddi hasara ve can kaybına neden olabilmektedir. Alt yapı su tahliye sistemlerinin tasarım aşamasında şiddet-süre-frekans (ŞSF) bilgileri sıklıkla kullanılmaktadır. Uzun yıllar kaydı tutulan, standart zamanlarda meydana gelen maksimum yağışlar vasıtasıyla ŞSF eğrileri oluşturulabilmektedir. Bu eğriler, tekerrür aralığına bağlı olarak yağış şiddeti değerinin yaklaşık olarak tespit edilmesi prensibine göre okunur. İnşa edilecek alt yapı elemanlarının boyutları bu bağıntılar kullanılarak elde edilen şiddet değerleri doğrultusunda inşa edilmektedir. Literatürde ŞSF değerleri arasındaki bağıntıyı açıklamak için kullanılan birçok ampirik denklem mevcuttur. Benzeden ve Bükey, (2001) Ege bölgesi için yaptıkları çalışmada yağış şiddeti değerlerini (i), yağış süresi(t) ve tekerrür yılına (T) göre; $i = \frac{A(T)}{(t+\theta)^n}$ şeklindeki bir denklem ile ifade etmiştir. Yüksek vd. (2016) Rize ili için yaptığı

çalışma sonunda ise; $i = \frac{a+b\ln(T)}{(t+c)^d}$ denklemini elde etmiştir. Ampirik denklemlerde bulunan $A, \theta, c, ve d$ parametreleri optimize edilerek ŞSF eğrileri oluşturulabilmektedir. Bu parametrelerin optimizasyonunun klasik optimizasyon yöntemleri ile optimize edilmesi zahmetli olabilmektedir. Karahan vd. (2008) ampirik denklemlerin optimizasyonu için Genetik Algoritmaları (GA) kullanarak bu alandaki çalışmalara yeni bir soluk kazandırmıştır.

Son yıllarda GA'ya ek olarak birçok sürü tabanlı optimizasyon algoritmaları geliştirilmiştir. Bu algoritmalarından en çok dikkat çeken parçacık sürü algoritmasıdır (Kennedy ve Eberhart, 1995). Su kaynakları ve hidroloji alanında en sık kullanılan sürü tabanlı optimizasyon algoritması olarak parçacık sürü optimizasyonu (PSO) örnek verilebilir.

Genetik algoritma tabanlı, Genetik Programlama (GP) yöntemi ile daha gelişkin problem çözümleri sunulmuştur. Literatürde yapılan çalışmalar incelendiğinde GP yönteminin su kaynakları ve hidroloji alanında birçok uygulamasının olduğu göze çarpmaktadır. Haddad vd. (2017) İran'ın Sefidrood Nehri'ndeki su kalitesi parametreleri modellemek için destek vektör makineleri, genetik algoritma ve genetik programlama yöntemlerini kullanmıştır. GP yönteminin diğer iki yöntemle göre daha başarılı sonuçlar verdiğini sunmuşlardır. Deo vd. (2017) su döngüsünün önemli parametrelerinden biri olan buharlaşmayı tahmin etmek için GP yöntemini kullanmışlardır. Hadi ve Tombul (2018) akım tahmini için dalgacık dönüşümlü genetik programlama yöntemini kullanmışlardır. Demirhan ve Atılğan (2015) solar radyasyon değerlerinin tahmini için GP yöntemi kullanmışlardır. Mattar, (2018) referans evapotranspirasyon (ET) değerlerinin tahmini için GP yöntemi kullanmıştır. Ampirik ET denklemlerine nazaran GP ile oluşturduğu denklemler ile daha başarılı tahminler elde etmiştir.

Bu çalışmada, Kayseri iline ait standart zamanlarda ölçülen yağış verileri kullanılarak GP yöntemi ile ŞSF eğri denklemleri elde edilmiştir. Elde edilen denklem, Karahan vd. (2008) tarafından önerilen ampirik denklem ile karşılaştırılmıştır. Ampirik denklem parametreleri parçacık sürü optimizasyonu (PSO) yöntemi ile optimize edilmiştir. Kayseri ili verileri ile elde edilen denklemler komşu illerden olan Yozgat, Nevşehir ve Niğde illeri değerlerine uygulanarak temsil kabiliyeti test edilmiştir.

MATERYAL VE YÖNTEM

Materyal

Çalışmada kullanılan veriler, standart zamanlarda ölçülen en büyük yağış değerlerinden oluşmaktadır. Yağış değerleri milimetre cinsinden su yüksekliklerini ifade etmektedir. Çalışma bölgesi olarak İç Anadolu bölgesinde bulunan Kayseri, Yozgat, Nevşehir ve Niğde illeri seçilmiştir (Şekil 1). Veri uzunlukları, Kayseri için 1950-2015, Nevşehir için 1965-2015, Niğde için 1959-2015 ve Yozgat için 1960-2015 yıllarını kapsamaktadır. İllere ait yağış değerlerinin genel istatistiki özellikleri Çizelge 1’de gösterilmiştir. Çalışmada kullanılan materyal açık bir şekilde kaynaklarına bağlı olarak verilmelidir. Ayrıca çalışmada kullanılan istatistiksel modeller açık ve tam olarak ifade edilmelidir.



Şekil 1. Çalışma alanı

Çizelge 1. Maksimum yağışların genel istatistiki değerleri

	Dakika												
Kayseri	5	10	15	30	60	120	180	240	300	360	480	1020	1080
Minimum(mm)	0.9	1.0	1.5	2.3	4.0	4.0	4.5	5.0	5.6	6.3	6.9	7.1	7.1
Maximum(mm)	11.4	21.6	21.9	32.5	34.8	40.9	41.9	42.0	42.2	50.9	51.7	51.8	51.8
Ortalama(mm)	5.3	8.2	9.9	12.8	14.8	17.9	19.0	19.8	20.5	21.3	22.4	23.5	25.1
Std. Sap. (mm)	2.57	4.19	4.65	6.83	8.02	8.95	8.75	8.85	8.67	9.08	9.11	8.91	8.77
Nevşehir	5	10	15	30	60	120	180	240	300	360	480	1020	1080
Minimum(mm)	0.9	1.8	1.9	3.5	5.0	6.0	7.0	7.5	8.1	8.8	8.8	8.8	8.8
Maximum(mm)	10.0	14.6	20.6	29.5	33.9	35.6	36.2	36.9	37.3	37.4	37.4	37.4	38.1
Ortalama(mm)	4.5	6.6	8.1	10.6	13.5	16.0	17.3	18.3	19.3	19.9	20.8	22.0	23.6
Std. Sap. (mm)	1.93	2.95	3.83	5.31	6.56	6.74	7.03	7.20	6.88	6.71	6.72	6.64	6.84
Niğde	5	10	15	30	60	120	180	240	300	360	480	1020	1080
Minimum(mm)	0.9	1.7	2.3	3.0	3.1	4.6	5.2	5.6	5.6	5.6	5.7	5.7	5.7
Maximum(mm)	10.1	14.3	17.6	26.7	33.4	39.2	42.6	42.6	42.6	42.6	42.6	42.6	45.0
Ortalama(mm)	4.5	6.5	7.6	10.0	12.1	14.8	16.0	17.2	18.0	18.8	19.6	20.6	22.6
Std. Sap. (mm)	1.93	2.55	3.29	5.13	6.29	6.40	6.60	6.54	6.42	6.43	6.70	7.13	7.76
Yozgat	5	10	15	30	60	120	180	240	300	360	480	1020	1080
Minimum(mm)	1.4	1.6	2.8	4.3	4.6	5	6.2	7.5	8.1	8.2	8.2	9.7	10.5
Maximum(mm)	15.0	20.8	31.1	39.9	51.2	53.9	54.0	54.0	54.0	54.0	54.0	56.4	67.2
Ortalama(mm)	6.0	8.5	11.0	14.4	17.3	19.8	21.6	22.8	23.6	24.0	26.0	28.2	31.6
Std. Sap. (mm)	2.93	4.41	6.25	8.04	9.73	9.95	9.62	9.10	8.84	8.71	7.85	8.74	10.56

Çizelge 1 incelendiğinde Kayseri, Nevşehir ve Niğde illeri standart zamanlarda ölçülmüş en büyük yağış değerlerinin karakterleri birbirine yakın seyretmektedir. Yozgat ili değerleri ise diğer üç ile

nazaran daha yüksek seyretmektedir. Yağış değerlerinin dağılımları incelendiğinde Kayseri ilinin Log-Pearson 3, Nevşehir ilinin Gumbel, Niğde ilinin Log-Normal 2 ve Yozgat ilinin Log-Normal 3 dağılımına uyduğu tespit edilmiştir.

Yöntem

Ampirik Denklem

Bu çalışmada yağış şiddeti değerlerinin tahmin edilmesi için, Karahan vd. (2008) tarafından önerilen ampirik bağlantı kullanılmıştır. Şiddet değerinin yağış süresi (t) ve yağış frekansı (T) arasındaki bağlantı denklem 1’de gösterildiği gibidir.

$$i = \frac{x_1 + x_2 \times \ln(T) + x_3 \times [\ln(T)]^2}{x_4 + x_5 \times \ln(t) + x_6 \times [\ln(t)]^2} \quad (1)$$

Denklemden bulunan i, yağış şiddetini, T yağış frekansını, t ise yağış süresini belirtmektedir. $x_1, \dots, x_6$ optimize edilmesi gereken denklem sabitleridir.

Genetik Programlama

Genetik algoritma tabanlı ve daha kullanışlı sürümü olan Genetik Programlama (GP) Koza (1992) tarafından geliştirilmiştir. GP, optimizasyon problemlerini ilkel çözüm yollarını gelişkin uyum kriterlerine göre uyarlaması (evirmesi) prensibine dayalı evrimsel optimizasyon yöntemidir. Genetik algoritmada çözüm gerçek sayılar ile ifade edilirken, GP’de çözüm genellikle ağaç tabanlı bilgisayar programları şeklinde ifade edilmektedir. GP’de ağaç sayısı ve uzunluğu işlem sırasında değişebilirken, GA’da bulunan gerçek sayılar sabittir.

GP işlemleri üç temel aşamada gerçekleşmektedir. İlk olarak başlangıç popülasyonu oluşturulur, ikinci olarak mutasyon ve çaprazlama ile yeni popülasyon oluşturulur, son olarak ise amaç fonksiyonu değeri kontrol edilerek sonlandırma kriteri dahilinde sonuca ulaşılır. Başlangıç popülasyonu mümkün olan çözüm uzayında olacak şekilde rastgele oluşturulur. Her popülasyon bireylerden oluşmaktadır ve bu bireyler fonksiyon ve terminaller içeren bir ağaçtır. Fonksiyon kümesi kullanıcı tarafından seçilmektedir ve +, -, /, × gibi aritmetik operatörler, sin, cos, tan, log, exp gibi matematiksel operatörler olabilmektedir. Çözüm için seçilen terminal ve fonksiyonlar ile rastgele ağaç benzeri bir yapı oluşturulur. Ağaçlarının ilk popülasyonu oluşturulduktan sonra GP, kullanıcı tanımlı amaç fonksiyonu seçilir. Daha sonra bu fonksiyon değerini minimum yapacak yeni bir nüfus oluşturmak için üreme ve varyasyon işlemi gerçekleştirilir. Bu işlem, kullanıcı tanımlı bir durma kriteri sağlanana kadar gerçekleştirilir. Popülasyon bir nesilden diğerine geliştikçe, yeni çözümler eskilerinin yerini alır ve daha iyi performans göstermesi beklenir. Çaprazlama, yeni çocuklar oluşturmak için seçilen üst ayrıştırma ağaçları arasındaki alt ağaçların rastgele değiştirilmesidir. Geçişler, evrimsel sürecin çözüm alanının vaat eden bölgelerine doğru ilerlemesine olanak tanır. Geçişin aksine, mutasyonda, tek bir ebeveyn ayrıştırma ağacı seçilir ve üzerinde rastgele değişiklikler yapılır. Mutasyon, yerel optimuma erken yakınsamayı önlemek için eklenmiştir

Parçacık Sürü Optimizasyonu

Parçacık Sürü Algoritması (PSO), sürü davranışlarından ilham alınarak geliştirilmiş bir algoritmadır (Kennedy ve Eberhart, 1995). Popülasyon tabanlı olan PSO, sürü içerisinde bilgi paylaşımını esas almaktadır. Sürüdeki bireyler, parçacık olarak adlandırılmakta ve her bir parçacığın, pozisyon ve hız bilgileri kaydedilmektedir. Her bir yinelemede bu hız ve pozisyon bilgilerinin güncellenmesi ile optimizasyon çalışması yürütülür. Bu güncellenme, parçacığın o anki hız bilgisine, konumuna, sürünün en iyi pozisyonuna ve o iterasyon içindeki en iyi pozisyona göre ayarlanmaktadır. “X” parçacık pozisyonunu (yön) ve “v” parçacık uçuş hızını temsil etmektedir. Parçacık sürüsündeki her x’e problemin çözüm yaklaşımına bağlı olarak bir puan verilir. Pbest, her bir parçacık için yerel en iyi

olanını; Gbest ise mevcut jenerasyondaki küresel en iyiyi temsil etmektedir. Parçacıklar, kendi en iyi pozisyonunu (Pbest) kaydeder ve sürüdeki tüm parçacıkların bildiği en uygun pozisyonu (Gbest) bulma yeteneğine sahiptir. Parçacığın bir adım sonraki pozisyonunun belirlenmesi amacıyla, hız ve yeni pozisyon vektörü, parçacığın o ana kadar ki bilgilerinden faydalanarak elde edilmektedir.

$$v_{id}^{t+1} = wv_i^t + c_1r_1(p_{id} - x_{id}^t) + c_2r_2(p_{gd} - x_{id}^t) \quad (2)$$

$$x_{id}^{t+1} = x_{id}^t + v_{id}^{t+1} \quad (3)$$

w atalet değeri, parçacığın hareketsizliğini önlemeye yarayan sabit, c1 ve c2 hızlanma sabitleri, r1 ve r2 ise [0, 1] aralığında rastgele üretilen bir değerdir. Algoritma rastgele bir hız ve konum bilgisi ile çözüm uzayında arama işlemine başlar. Her parçacığın konum bilgisi uygunluk fonksiyonun yardımıyla uygunluk değeri tespit edilir. Tespit edilen uygunluk değeri, parçacığın o ana kadar elde ettiği en iyi uygunluk değeri olan pbest değeri ile kıyaslanarak, pbest değeri güncellenir. Parçacığın o ana kadar sahip olduğu en iyi uygunluk değeri, sürünün o ana kadar elde ettiği en iyi uygunluk değeri olan gbest değeri ile karşılaştırılarak, gbest değeri güncellenir.

Performans Başarı Kriterleri

Bu çalışmada modellerin tahmin performansını değerlendirmek için ortalama karesel hata (OKH) ve Verimlilik Katsayısı (NSE) denklemleri kullanılmıştır.

$$OKH = \frac{1}{n} \sum_{i=1}^n (I_{gi} - I_{ti})^2 \quad (4)$$

$$NSE = 1 - \frac{\frac{1}{n} \sum_{i=1}^n (I_{ti} - I_{gi})^2}{\frac{1}{n} \sum_{i=1}^n (I_{oi} - I_{gi})^2} \quad (5)$$

Denklemdaki ifadeler; NSE=Verimlilik Katsayısı, n=toplam gözlem ve tahmin değeri sayısı, I_{ti} =i nci yağış şiddeti tahmin değeri, I_{gi} =i nci ölçülen yağış şiddeti değeri, I_{oi} = ölçülen yağış şiddeti değerleri ortalamasıdır.

BULGULAR

Çalışmada ilk olarak yağış şiddeti değerlerinin hesabı için ampirik denklem parametreleri parçacık sürü optimizasyonu (PSO) yöntemi ile optimize edilmiştir. Optimizasyon işlemi MATLAB paket programı vasıtasıyla gerçekleştirilmiştir. Ampirik denklemde optimize edilmesi gereken toplam 6 sabit mevcuttur. Şiddet-süre-frekans analizi yapılması için öncelikle en büyük yağış değerlerinin frekans analizi yapılmalıdır. Uzun süreli kayıtlar sayesinde uygun dağılım seçilebilmektedir. Seçilen dağılıma uygun 2, 5, 10, 25, 50, 100 yılda bir olması muhtemel en büyük yağış değerleri hesaplanmıştır. Yağış değerleri (mm) hesaplandıktan sonra yağış sürelerine bölünerek yağış şiddeti (mm/dakika) elde edilmiştir. Bu değerler ölçülen yağış değeri olarak isimlendirilmiştir ve istasyonun gerçek yağış şiddeti değerlerini temsil etmektedir. Elde edilen yağış şiddeti değerleri çıktı olarak, yağış süresi ve yağış frekansı girdi olarak sisteme tanıtılmıştır. Daha sonra ampirik denklem 6 adet sabit ile birlikte modele tanıtılmış ve modelden bu 6 adet sabitin optimize edilmesi istenmiştir. Optimizasyon işlemi sırasında amaç fonksiyonu olarak ortalama karesel hata (OKH) kullanılmıştır (denklem 4). Amaç fonksiyonunu minimum yapan değerler PSO algoritması tarafından hesaplanmıştır. Hesaplama işlemi sırasında PSO'ya ait popülasyon büyüklüğü değeri 100, atalet değeri 0.9 hızlanma sabitleri olan c1 ve c2 2 alınarak algoritma başlatılmıştır. Kayseri ili yağış değerleri ile optimize edilen ampirik denklem aşağıda verilmiştir.

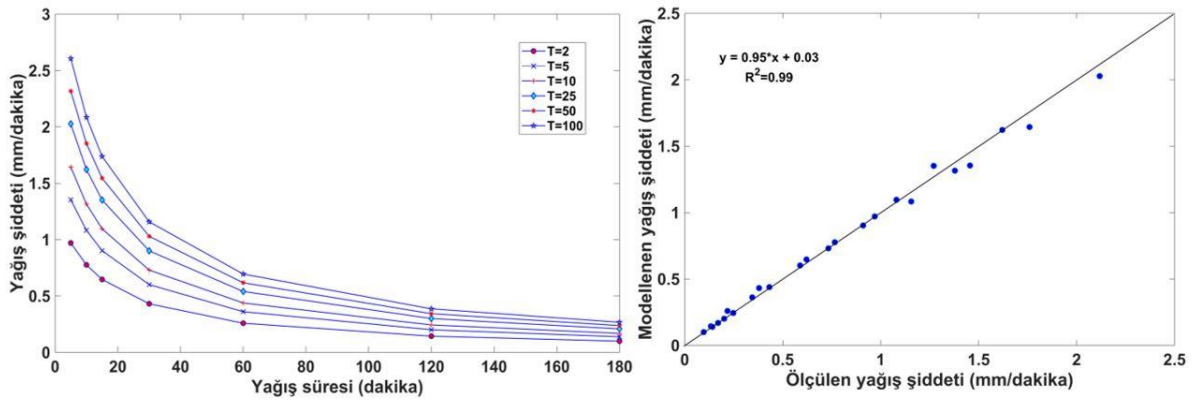
$$i = \frac{2384 + 1367 \times \ln(T) - 1113 \times [\ln(T)]^2}{-938 - 353 \times \ln(t) + 121 \times [\ln(t)]^2} \quad (6)$$

Optimizasyon işlemi sonrası Kayseri için elde edilen denklem vasıtasıyla Nevşehir, Niğde ve Yozgat illeri yağış şiddeti değerleri tahmin edilmiştir. Kayseri için OKH değeri 0.028, NSE değeri 0.94, Nevşehir için OKH 0.045, NSE 0.86, Niğde için OKH 0.071, NSE değeri 0.75, Yozgat için OKH değeri 0.044, NSE değeri 0.87 olarak hesaplanmıştır. Kayseri için çok iyi bir yakınsama mevcutken diğer illerde belirli sapmalar mevcuttur.

GP yöntemi ile model oluşturma işlemleri HeuristicLab paket programı vasıtasıyla gerçekleştirilmiştir. Program ücretsiz olup birçok GP modeli rahatlıkla oluşturulabilmektedir. GP tahmin modelleri oluşturma aşamasında popülasyon sayısı 50, iterasyon sayısı 1000, ağaç uzunluğu 15, ağaç derinliği 20 olarak seçilmiştir. Program işlem sonunda basitten karmaşığa doğru birçok denklem sunmaktadır. Bu çalışmada karşılaştırılan ampirik denklemden daha basit ve kullanışlı olması açısından aşağıdaki denklem 7 seçilmiştir.

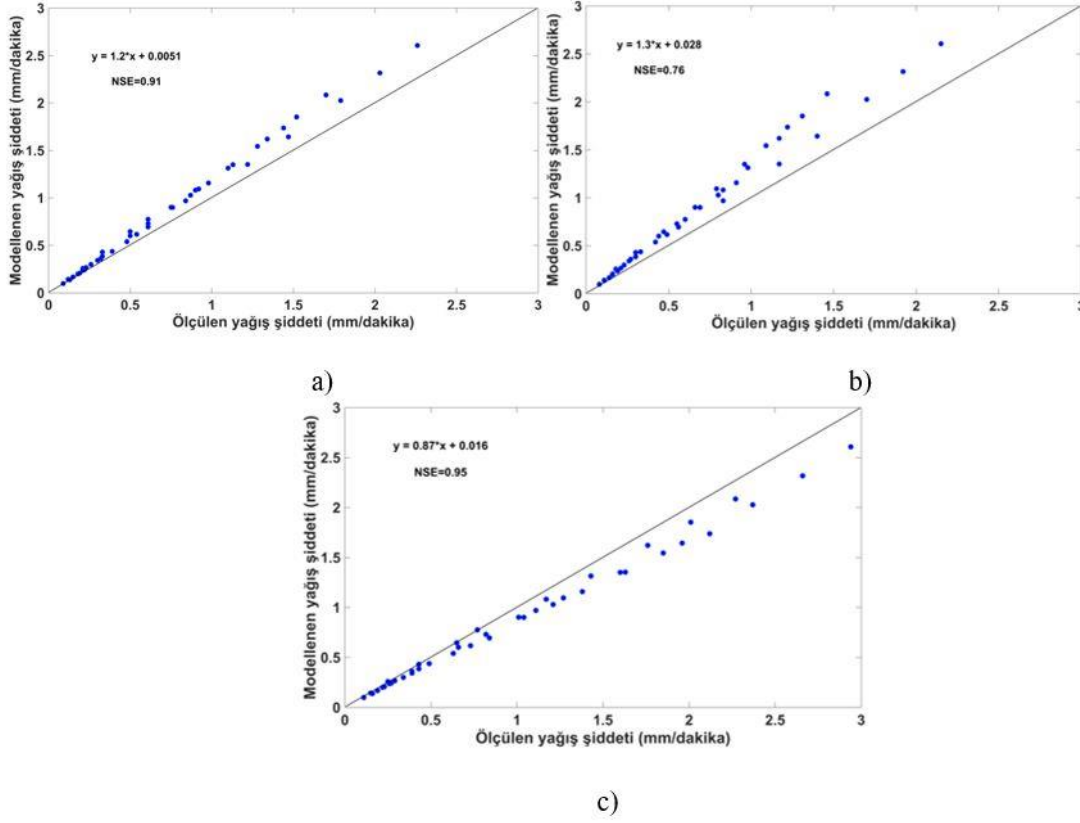
$$i = \frac{13.6 + 8.36 \times \ln(T)}{15 + t} \quad (7)$$

Kayseri yağış değerleri ve frekans değerleri kullanılarak oluşturulan denklemin OKH değeri 0.004 ve NSE değeri 0.99 olarak hesaplanmıştır. Uyum neredeyse mükemmel olarak adlandırılabilir. Kayseri ili için oluşturulan şiddet-süre-frekans eğrileri ve ölçülen yağış şiddeti ile tahmin edilen yağış şiddeti değerlerine ait saçılma grafikleri şekil3'te verilmiştir. Denklem 7 kullanılarak Nevşehir, Niğde ve Yozgat için hesaplanan yağış şiddeti değerleri ve ölçülen yağış şiddeti değerlerine ait saçılma grafikleri şekil4'te verilmiştir.



Şekil 3. Kayseri ili ŞSF eğrisi ve saçılma grafiği

Saçılma grafikleri incelendiğinde, en başarılı tahmin sonuçları Yozgat ili için elde edilmiştir. Yozgat ili için OKH değeri 0.027, NSE değeri ise 0.95 olarak hesaplanmıştır. Nevşehir ili için hesaplanan OKH ve NSE değerleri ise sırasıyla 0.028 ve 0.91'dir. Nevşehir ve Yozgat'a nazaran Niğde için hesaplanan OKH değeri 0.067 olarak daha yüksek, NSE ise 0.76 olarak daha düşük hesaplanmıştır. Kayseri ile Nevşehir ve Yozgat illerinin ŞSF eğrilerinin benzerlik gösterdiği söylenebilmektedir. Niğde ili için ise kullanışlı sayılabilir ama daha iyi sonuçlar almak için ikinci bir kalibrasyon gerekmektedir.



Şekil 4. Ölçülen ve modellenen yağış şiddeti değerleri (a) Nevşehir, (b)Niğde, (c) Yozgat

TARTIŞMA VE SONUÇ

Bu çalışmada Genetik Programlama yöntemi ile yağış şiddet-süre-frekans analizi yapılmıştır. Uzun yıllara ait en büyük yağış değerleri vasıtasıyla uygun dağılımlar belirlenmiş ve ölçülen yağış şiddeti değerleri elde edilmiştir. Genetik programlama yardımı ile Kayseri iline ait yağış değerleri kullanılarak bir denklem elde edilmiştir. Denklemin çıktısı yağış şiddeti, girdisi ise yağış süresi ve frekansdır. Elde edilen denklem yakın konumda olan farklı bölgelerde test edilmiş ve yüksek başarıda tahminler elde edilmiştir. Ayrıca literatürdeki bir ampirik denklem ile karşılaştırma yapılmıştır. Literatürdeki ampirik denklemde 6 adet sabitin optimizasyonu parçacık sürü optimizasyon algoritması ile yapılmıştır. Ampirik denklem, optimizasyonu sonrası çevre bölgeler üzerinde test edilmiştir. Sonuçlar istatistiksel olarak kabul edilebilir seviyelerde çıkmıştır. GP ile elde edilen tahmin sonuçları görece daha karmaşık olan ampirik denkleme nazaran daha başarılı tahmin sonuçları elde edilmiştir. Sonuçlar incelendiğinde Kayseri iline yakın bölgelerde enleme bağlı bir değişiklik olabileceği göze çarpmaktadır. İlerleyen çalışmalarda enlem değerleri de girdi olarak alınıp daha hassas modellerin kurulması düşünülmektedir.

Kaynaklar

- Benzeden, E. ve Bükey, B. (2001). Ege bölgesi örneğinde en uygun yağış şiddeti-süre tekerrür modelleri, III. Ulusal Hidroloji Kongresi, İzmir, 19-29.
- Yüksek, Ö., Anılan, T., Öztürk, H., Örgün, E., (2016). Şehir taşkınlarının tahmininde yağış şiddeti-süre-tekerrür analizi: Rize ili için bir uygulama, 765-771. 4.Ulusal Taşkın Sempozyumu, Kasım 23-25, 2016, Rize.
- Karahan, H., Ayvaz, M. T., Gürarslan, G., (2008). Şiddet-süre-frekans bağıntısının genetik algoritma ile belirlenmesi: GAP örneği. İMO Teknik Dergi, 290 (46):4393-4407.
- Kennedy, J. ve Eberhart, R. C, (1995) Particle Swarm Optimization. IEEE International Conference on Neural Networks, vol. IV, 1942-1948, Piscataway, NJ,

- Mirjalili, S. ve Lewis, A. (2014) Grey Wolf Optimizer, *Advances in Engineering Software*, Volume 69, 2014, Pages 46-61.
- Haddad, O., Soleimani, S., Loáiciga, H. A. (2017). Modeling Water-Quality Parameters Using Genetic Algorithm–Least Squares Support Vector Regression and Genetic Programming. *Journal of Environmental Engineering*, 143(7), 04017021.
- Deo, R. C., Samui, P. (2017). Forecasting Evaporative Loss by Least-Square Support-Vector Regression and Evaluation with Genetic Programming, Gaussian Process, and Minimax Probability Machine Regression: Case Study of Brisbane City. *Journal of Hydrologic Engineering*, 22(6), 05017003.
- Hadi, S. J. ve Tombul, M. (2018). Monthly streamflow forecasting using continuous wavelet and multi-gene genetic programming combination. *Journal of Hydrology*, 561, 674-687.
- Demirhan, H. ve Kayhan Atılgan, Y. (2015). New horizontal global solar radiation estimation models for Turkey based on robust coplot supported genetic programming technique. *Energy Conversion and Management*, 106, 1013-1023.
- Mattar, M. A. (2018). Using gene expression programming in monthly reference evapotranspiration modeling: A case study in Egypt. *Agricultural Water Management*, 198, 28-38.
- Koza, J. R., (1992) *Genetic programming: on the programming of computers by means of natural selection*. Vol. 1. MIT press.

Teşekkür

Veri temini konusunda yardımlarından ötürü Meteoroloji Genel Müdürlüğüne teşekkürü borç biliriz.

Çıkar Çatışması

“Yazarlar çıkar çatışması olmadığını beyan etmişlerdir” şeklinde yazılmalıdır.

Faktör Analizinde Faktör Puanlarının Hesaplanma Yöntemlerinin Karşılaştırılması

Turgut KARABULUT^{1*}, Selahattin YAVUZ¹

¹Erzincan Binali Yıldırım Üniversitesi

*Sunucu yazar e-posta: tkarabulut@erzincan.edu.tr

Özet

Sosyal bilimlerde en çok kullanılan yöntemlerden biri de faktör analizidir. Faktör analizi, birbirleriyle ilişkili çok sayıdaki değişkeni anlamlı ve birbirinden bağımsız olmak üzere daha az sayıda faktör haline getiren çok değişkenli istatistik tekniklerinden biridir. Faktör analizi özellikle araştırmalar için yapılan anket uygulaması sonrasında yaygın bir şekilde kullanılmaktadır. Literatürde faktör analizi sonrasında oluşturulan faktörlerin farklı biçimlerde puanlamasının yapıldığı görülmüştür. Genellikle ölçeklerde verilen puanların toplanması veya verilen puanların ortalaması alınarak faktör puanlarının oluşturulduğu görülmüştür. Faktör analizinde “dönüşümlü bileşen matris” inde faktörlerin hangi sorulardan oluştuğu tespit edilmektedir. Burada her bir sorunun faktörlerdeki ağırlıkları da verilmektedir. Faktör puanlarının bu ağırlıklara göre oluşturulması sonrasında sonuçlarda farklılıklara sebep olduğu ortaya çıkarılmaya çalışılmıştır. Çalışmada Spector’un İş Tatmin Düzeyini ölçen ölçeği kullanılarak, ortalama ve ağırlıklı ortalama yöntemleriyle oluşturulan faktörler arasında farklılıkların oluştuğu gözlemlenmiştir. Elde edilen bulgulara göre ağırlıklı ortalama yönteminin daha doğru sonuçlar verdiği düşünülmektedir.

Anahtar kelimeler: Faktör analizi, Faktör puanlama, Ağırlıklı ortalama yöntemi

Üniversite Öğrencilerinin Bireysel Girişimcilik Algı Durumları ile İletişim Becerileri Arasındaki İlişkinin Yapısal Eşitlik Modeli ile İncelenmesi

Nur İlkay ABACI^{1*}

¹Ondokuz Mayıs Üniversitesi, Ziraat Fakültesi, Tarım Ekonomisi Bölümü, 55139, Samsun, Türkiye

*Sunucu yazar e-posta: ilkaysonmez55@gmail.com

Özet

Bu araştırmanın amacı öğrencilerin girişimcilik algıları ile iletişim becerileri arasındaki ilişkiyi ortaya koymaktır. Bu amaç doğrultusunda Bireysel Girişimcilik Algı (BGA) Ölçeği ve İletişim Becerileri (İB) Ölçeği kullanılmıştır. Çalışmada Ondokuz Mayıs Üniversitesi, Ziraat Fakültesinde Girişimcilik dersi alan 304 öğrenci ile elektronik ortamda anketler yapılarak birincil veriler elde edilmiştir. Ölçeklerin güvenilirliklerini belirlemek için iç tutarlık katsayıları yapı geçerliliği için ise doğrulayıcı faktör analizi yapılmıştır. BAG ve İB ölçekleri arasındaki yapısal ilişkiyi belirlemek için Yapısal Eşitlik Modeli (YEM) kullanılmıştır. Verilerin analizinde SPSS ve Lisrel Paket programlarından yararlanılmıştır. Elde edilen YEM sonucuna göre BGA ölçeğinde en yüksek etkiye İletişim alt boyutu sahipken, en düşük etkiye Planlama ve Öz disiplin alt boyutlarının sahip olduğu belirlenmiştir. İB ölçeğinde ise en yüksek etkiye İletişim Kurmaya İsteklilik alt boyutu etkili iken en düşük etkiye İletişim İlkeleri ve Temel Beceriler alt boyutunun sahip olduğu belirlenmiştir. Ayrıca öğrencilerin bireysel girişimcilik algıları ile iletişim becerileri arasında %80 oranında bir ilişki tespit edilmiştir. Sonuç olarak bireylerin girişimcilik ruhuna sahip olmalarının önemli olduğu kadar, iş fikirlerini hayata geçirebilmeleri için yatırımcıları etkileyebilmeleri ve onlarla aynı dili konuşarak iş fikirlerini kabul ettirmeleri açısından iletişimin son derece önemli olduğu belirlenmiştir.

Anahtar kelimeler: Girişimcilik, Sosyal sermaye, İletişim, İş fikri

On the Study of a Nonlinear Fractional-order Boundary Value Via Fixed Point Theory

Habib DJOURDEM^{1*}

¹Laboratory of Fundamental and Applied mathematics of Oran (lmfao), University Of Oran

*Presenter author e-mail: djourdem.habib7@gmail.com

Abstract

We study a kind of nonlocal multipoint boundary value problem of nonlinear Riemann-Liouville fractional differential equations supplemented with multi-strip integral boundary conditions. Existence and uniqueness results for the given problem are obtained by means of standard fixed point theorems. Examples illustrating the main results are also discussed.

Key words: *Riemann-Liouville fractional derivative, Existence, Multi-point boundary conditions, Fixed point*

Boundedness and Stability For Generalized Schrödinger Equation

Nouredine BOUTERAA^{1*}

¹Laboratory of Fundamental and Applied mathematics of Oran (lmfao), University of Oran

*Presenter author e-mail: bouteraa-27@hotmail.fr

Abstract

In this paper, we study the boundedness and stability of a solutions for a class of nonlinear time-fractional differential equations with initial data by the help of fractional Duhamel principle. Well-known advantage of caputo derivatives with respect to the classical derivatives is its capability of taking into account the previous historical effects of model at each time step. This feature of fractional-order operators makes them more accurate and appropriate in modeling of the systems. An illustration how these are achieved and physical basis of fractional operators with different memory is presented in several research work recently

Key words: Fractional order, differential equation, Duhamel principle

Hibrid Panel Veri Analiz Metotolojisi Üzerine Uygulama: Ulusal Harcamalar Örneği

Selim DÖNMEZ^{1*}

¹Eskişehir Osmangazi Üniversitesi, Fen Edebiyat Fakültesi, İstatistik, 26040, Eskişehir, Türkiye

*Sunucu yazar e-posta: sdonmez3@gmail.com

Özet

Panel veri, zaman serilerinden oluşan çok birimli veri tipidir. Birimler arasındaki, birimlerin içindeki zamansal değişimi hem ayrı ayrı hem de ortak olarak açıklayıp betimleyici bir analiz gerçekleştirerek değişim dinamiklerini hakkında bilgi sağlar. Çalışmada model tabanlı kümeleme, ülkelerin gayrisafi yurtiçi hasılaya oranlanmış ulusal harcamalarını içeren panel veriye uygulandı ve bu uygulamada kullanılan panel veriye uygulanan doğrusal model uydurma yöntemi ile model tabanlı kümeleme yönteminden oluşan metodoloji ile ülkelerarası ulusal harcamaların davranış benzerlikleri açıklanmaya çalışıldı.

Anahtar kelimeler: Panel veri, Model tabanlı kümeleme, Değişim dinamiği

Ankara İli Ekstrem Yağışlarının İstatistiksel Analizi

Furkan YILMAZ^{1*}, Saniye DEMİR², Yunus AKDOĞAN³, Orhan KAVUNCU⁴, Yalçın TAHTALI¹, Kadir KARAKAYA³, Selma KÖKÇÜ⁴

¹Tokat Gaziosmanpaşa Üniversitesi, Ziraat Fakültesi, Zootekni Bölümü, Tokat, Türkiye

²Tokat Gaziosmanpaşa Üniversitesi, Ziraat Fakültesi, Toprak Bilimi ve Bitki Besleme Bölümü, Tokat, Türkiye

³Selçuk Üniversitesi, Fen Fakültesi, İstatistik Bölümü, Konya, Türkiye

⁴Kastamonu Üniversitesi, Mühendislik ve Mimarlık Fakültesi, Genetik ve Biyomühendislik Bölümü, Kastamonu, Türkiye

*Sunucu yazar e-posta: furkanyilmzz60@gmail.com

Özet

Ekstrem yağışlar ve iklim olayları (sıcak dalgalar, kuraklık, ağır yağışlar, hortumlar vs.) sürekli olarak insan hayatını tehdit etmektedir. İklim değişikliği ve buna bağlı olarak meydana gelen ekstrem yağışlar günümüzün en ciddi çevresel problemlerinden birisidir. Bu problemlerin azaltılması için yağışların önceden tahmin edilmesi gerekmektedir. Bunun için de yağışları uygun dağılımlarla ifade edebilmek gerekir. Bu çalışmada, Ankara yağış istasyonuna ait 1971-2020 yılları arasındaki bir, iki, üç, dört ve beş günlük yağışların Gumbel ve Gamma olasılık dağılımlarına uygunluğu araştırılmıştır. Çalışmanın sonucunda, her iki dağılımın sonucu elde edilen veriler arasında bir fark olmadığı görülmüştür.

Anahtar kelimeler: Ekstrem yağışlar, Hidroloji, Yüzey akış, Olasılık dağılımı

Statistical Analysis of Extreme Rainfall in Ankara Province

Abstract

Human life is continually threatened by extreme rains and climatic events (heat waves, droughts, heavy rains, tornadoes, and so on). Climate change and the resulting heavy rainfalls are one of the most important environmental issues of our time. It is important to know the pre-estimation of precipitation in order to mitigate these issues. For this the precipitation should be expressed with the convenient distributions. The compatibility of one, two, three, four, and five days of precipitation at the Ankara precipitation station between 1971 and 2020 with Gumbel and Gamma probability distributions was investigated in this report. As a result of the research, it was seen that there was no difference between the data obtained as a result of both distributions.

Key words: Extreme precipitation, Hydrology, Runoff, Probability distribution

GİRİŞ

Ekstrem yağışlar ve iklim olayları (sıcak dalgalar, kuraklık, ağır yağışlar, hortumlar vs.) sürekli olarak insan hayatını tehdit etmektedir. İklim değişikliği ve buna bağlı olarak meydana gelen ekstrem yağışlar günümüzün en ciddi çevresel problemlerinden birisidir. İklim değişikliği, küresel iklim sistemini oldukça yakından etkilemektedir. Ekstrem yağışların iklim değişikliğine yaptığı katkı göz ardı edilemeyecek kadar büyüktür. Büyük alanlarda, çok büyük felaketlere neden olan iklim değişikliği; bölgesel olarak topoğrafya ve arazi kullanımı gibi faktörlerin etkisi ile az veya şiddetli toprak kayıplarına

yol açmaktadır. Bundan dolayı, ekstrem yağışlardaki trendin belirlenmesi çok önemlidir (Mayooran ve Laheetharan, 2014).

Yağış ve yüzey akış gibi hidrolojik verileri toplamak ve düzenlemek bu alanda çalışan insanlar için oldukça zor bir iştir. Su ile ilgili yapıların tasarımında, su kanallarının yapımında bu hidrolojik veriler çok önem arz etmektedir. Pek çok durumda, bu veri seti içerisinde veriler sınırlı olmakla birlikte çoğu kez eksik olmaktadır. Verilerdeki bu eksikliğin giderilmesi, olasılık ve istatistiksel yöntemler ile mümkün olmaktadır. Frekans analizi, hidrolojik çalışmaların önemli bir parçası olarak kabul edilmektedir. Geçmişteki kayıtlar kullanılarak, gelecekteki olayların gerçekleşme ihtimalleri hesaplanabilmektedir. Bu durum, toprak ve su koruma çalışmalarında toprak kayıplarının ve yüzey akışın azaltılması açısından çok önemlidir.

Frekans analizinin öncelikle amacı; ekstrem yağışların şiddeti ile meydana gelme sıklıkları arasındaki ilişkiyi kurmaktır. Bu ilişkiyi, frekans dağılımlarını kullanarak belirlemek mümkündür. Frekans analizinde; belli bir zaman aralığında ölçülmüş yağış ve akım verilerinin bağımsız ve aynı dağılıma ait oldukları kabul edilmektedir. Frekans analizinde;

1. Veri setine uygun bir olasılık dağılım modelinin belirlenmesi,
2. Belirlenen dağılım modelinin parametrelerinin tahmini,
3. Yağış verilerine ait riskin uygun bir hassasiyet düzeyinde tahmin edilmesi şeklindeki sıra ile yapılmaktadır.

Türkiye’de ekstrem yağışların dağılımların belirlenmesine yönelik birçok çalışma bulunmaktadır. Öztekin (2011) Samsun, Sinop, Ordu ve Tokat illeri günlük en yüksek yağışlar için en uygun dağılımı belirlemiştir. Çalışmada 32 adet sürekli dağılım kullanılmıştır. FRANMOD modelinin kullanıldığı çalışmada, Samsun ve Sinop illeri için Wakeby, Ordu ili için beta-P, Tokat ili için Wakeby ve beta-P dağılımının uygun olduğu; gamma dağılımının uygun olmadığı ifade edilmiştir.

Anlı (2006) ise Giresun Aksu havzasında maksimum akım frekanslarının modellenmesi için olasılık dağılımlarını kullanmıştır. 39 yıllık süreli aylık ve yıllık maksimum akım verileri üzerinde yürüttükleri çalışmada; gamma dağılımının hiçbir veriye uygun olmadığı görülmüştür.

Nadarajah ve Withers (2001) ile Nadarajah (2005), Yeni Zelanda boyunca 16 ve Amerika Birleşik Devletlerinin Florida da dahil olmak üzere 14 doğu eyaletinde maksimum yağışların dağılımını araştırmışlardır. Diğer taraftan, Nadarajah ve Choi (2007), Güney Kore’de 5 lokasyona ait 1961-2001 yılları arasındaki ekstrem yağışları inceledikleri çalışmada, genelleştirilmiş ekstrem değerler dağılımından ziyade Gamma dağılımının yağışları tahmin etmede daha uygun olduğunu ifade etmişlerdir. Varathan ve ark. (2010), Kolombiya’da 110 yıllık yağış verilerini kullanarak maksimum yağışları incelemişlerdir. Yağışlar için en uygun dağılımın Gumbel dağılım olduğunu çalışmanın sonunda belirtmişlerdir. Bütün bu çalışmalarda farklı dağılımların uygun olarak belirlenmesi göstermektedir ki, ekstrem yağışların olasılık dağılımlarının modellenmesinde, çalışma alanının coğrafik özelliklerine uygun model seçilmesi gerekmektedir.

İklim değişikliğinin tartışıldığı ulusal ve uluslararası panellerde; uyum ve önleme çalışmalarının tüm paydaşlar ve hükümetler tarafından dikkatlice izlenilmesi gerektiği, değişikliğin olup olmadığının belirlenmesi ve olmuş ise ne kadar olduğunun belirlenmesinin gerektiği ifade edilmektedir. Türkiye, iklim değişikliğinden en fazla etkilenen ülkeler arasında yer almaktadır. Kuraklık, sel ve taşkın gibi felaketler ile çok sık karşılaşmaktadır. Bu felaketlerin azaltılması ya da önlenmesi için ekstrem yağış olaylarının olasılık dağılımlarının belirlenmesi gerekmektedir. Bu çalışmada; Ankara yağış istasyonuna ait 1971-2020 yılları arasındaki bir günlük, iki günlük, üç günlük, dört günlük ve beş günlük yağış verileri için 10, 20, 30, 50, 75, 100 ve 200 yıllık dönüşüm periyotlarının uyum iyiliği kriterleri kullanılarak olasılık dağılımlarının belirlenmesi amaçlanmıştır.

MATERYAL VE YÖNTEM

Materyal

Çalışma sahası olan Ankara, 39°97'27" K enlemi, 32°86' 37" D boylamı arasında yer almakta olup, 890.52m yüksekliğe sahiptir. İç Anadolu bölgesi ile kenar dağların bu bölgeye dönük bölümlerinin önünde kalan alanlarda hüküm süren Karasal İç Anadolu termik rejimi içerisinde yer almaktadır. Bu rejim tipinde Kış mevsimi batıya doğru artmak koşuluyla soğuk geçmekte; ortalama sıcaklık 00 ile - 30°C arasında değişmektedir. Ortalama 20°C olan yaz sıcaklıkları en sıcak aylarda 19⁰ - 23⁰ arasında değişmektedir (Koçman, 1993). Koçman (1993)'ın yağış rejimi tanımlamasına göre çalışma sahası içerisinde yer alan Ankara, İç Anadolu karasal geçiş tipi yağış rejimi etkisindedir. Bu rejimde denizlerin etkisini engelleyen yüksek kenar dağları ve artan karasallığın etkisiyle güney ve batıdaki Akdeniz, kuzeydeki Karadeniz ve doğudaki karasal rejimin özellikleri değişmekte, buraya özgü bir geçiş tipi ortaya çıkmaktadır. Bu bölgede termik nedenlerle oluşan konveksiyonel yağışlar yaz kuraklığını hafifletir. Ankara yöresinde Ekim ayında artmaya başlayan yağışlar Nisan ayına kadar devam etmekte; Mayıs ayında nispi nemin maksimuma ulaşmasına karşın Hazirandan itibaren azalmaya başlamakta ve Ağustos ayında minimuma ulaşmaktadır (Koçman, 1993).

Yöntem

Gumbel Dağılımı

Yıllık yağış serilerinin modellenmesinde en sık kullanılan yöntemlerden birisi Gumbel dağılımı yöntemidir. Hidrolojik veriler için en iyi frekans dağılımlarından birisidir. Ortak dağılım olarak da bilinen bu yöntem, tip 1 ekstrem değer dağılımı maksimum yağış değerlerinin analizi için uygulanmaktadır (Gumbel, 1942); (Salinas ve ark., 2014). Rastgele bir örneğin elemanlarının üstel bir dağılımı için kümülatif dağılım işlevi aşağıda gösterilmiştir.

Rastgele bir örneğin elemanlarının üstel dağılımı için kümülatif dağılım fonksiyonu aşağıda gösterilmiştir.

$$F(x) = 1 - e^{-\lambda x} \quad (1)$$

Sonra maksimum değer için marjinal dağılımdan rastgele örneklem için aşağıdaki denklem kullanılır:

$$F(x)_{max}(x) = [F(x)]^n \quad (2)$$

$$F(x)_{max}(x) = [1 - e^{-\lambda x}]^n \quad (3)$$

Konum ve ölçek parametreleri a ve b tanımlamasıyla,

$$F(x)_{max}(x) = \left[1 - \frac{1}{n} e^{-\frac{x-b}{a}}\right]^n \quad (4)$$

n sonsuza gittikçe, karşılık gelen olasılık yoğunluk fonksiyonu şu şekilde olacaktır:

$$F(x)_{max}(x) = \exp \left[e^{-\frac{x-b}{a}} \right] \quad (5)$$

$$F(x)_{max}(x) = \frac{1}{a} \exp \left[-\frac{x-a}{a} e^{-\frac{x-b}{a}} \right] - \infty < x < +\infty \text{ için} \quad (6)$$

Gamma Dağılımı

Bu dağılım birçok mühendislik alanında kurak bölgelerde aylık yağışların dağılımlarını belirlemede yaygın bir şekilde kullanılan yöntemidir (Şen ve Eljadid,1999). Hidrolojik frekans analizi için normal dağılım, log-normal, gamma, gumbel ve weibull dağılımları bilinen yöntemlerdir. Gamma dağılımının pozitif değerlere sahip olma özelliği vardır. Hidrolojik değişkenlerde maksimum veya minimum değerleri için gumbel ve gamma dağılımları yaygın olarak kullanılmaktadır. Farklı gamma dağılım türleri arasında 2 ve 3 parametrelili gamma dağılımı fonksiyonları genellikle frekans analizi için kullanılır. Genel olarak kullanılan gamma dağılımı yöntemi maksimum ve minimum yağış verileri için uygun bir dağılım yöntemidir. Gamma dağılımında kullanılan kümülatif olasılık aşağıdaki gibidir:

$$G(x) = \frac{1}{\beta^\alpha \Gamma(\alpha)} \int_0^x t^{\alpha-1} e^{-\frac{t}{\beta}} dt \text{ for } x > 0 \quad (7)$$

Burada; α - şekil parametresi, β - ölçek parametresi, x -yağış değerleri ve $\Gamma(\alpha)$ - gamma fonksiyonudur.

Akaike bilgi kriteri (AIC)

AIC (Akaike, 1998), bir modelin olabilirliğinin gelecekteki değerini tahmin etmeye yönelik oldukça popüler bir yöntemdir. Japon istatistikçi Akaike'nin önerdiği AIC gerçek veri çalışmalarında modellerin veriye olan uyumlarını karşılaştırmak için sıklıkla kullanılmaktadır. Veri modellemesinde bilgi kaybını en aza indiren model seçilmek istenir. Bu nedenle veriye uyum sağlayan modeller arasında hangi modelin veriye en iyi uyum gösterdiğini belirlemek için veriye uyumlarını kıyaslamak istenen modellerin AIC değerleri hesaplanarak en küçük değere sahip model veriye en iyi uyum sağlayan model olarak belirlenir. AIC aşağıdaki gibi tanımlanır.

$$AIC = -2\ell + 2k \quad (8)$$

Burada; ℓ olabilirlik fonksiyon değeri iken, k dağılıma ait parametre sayısını verir.

Bayes Bilgi Kriteri

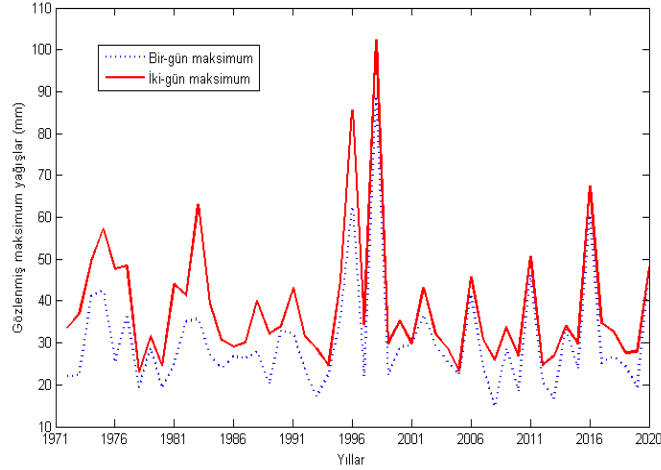
Bayes bilgi kriteri (BIC) Schwarz (1978) tarafından önerilmiştir. BIC, AIC ve olabilirlik fonksiyonunun aldığı değer ile yakından ilişkilidir. Model seçimi için kullanılan BIC aşağıdaki gibi tanımlanır.

$$BIC = -2\ell + k \log(n) \quad (9)$$

Burada; ℓ olabilirlik fonksiyon değeri iken, n örneklem hacmini (gözlem sayısı) ve k dağılıma ait parametre sayısını verir.

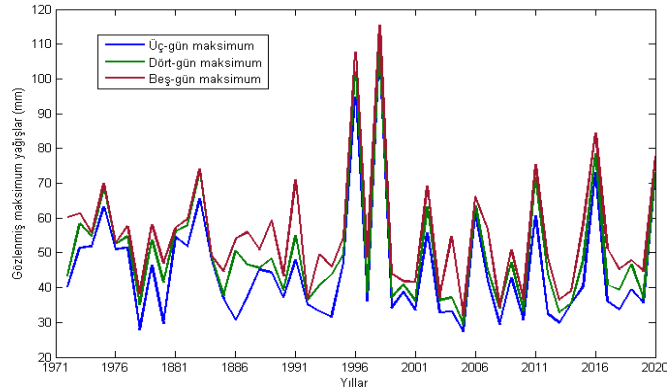
SONUÇLAR

Ankara ili yağış istasyonuna ait 1971-2020 yılları arasındaki bir, iki, üç, dört ve beş günlük yağışlar çalışılmıştır. İki gün için iki günlük toplam yağış, 5 gün için peş peşe 5 gün toplam yağıştır. Günlük yağışların dağılımları Şekil 1 ve Şekil 2'de verilmiştir. 1971-2020 yılları arasındaki bir ve iki günlük yağışlar incelendiğinde; Ankara ili'nin oldukça kurak bir iklime sahip olduğu görülmektedir (Şekil 1). 1996 ve 1998 yıllarında yağışların peak yaptığı ve 1999 yılında çok keskin bir şekilde düştüğü Şekil 1'de görülmektedir. 2016 yılında 60 mm civarında seyreden hem günlük hem de iki günlük yağışların tekrar azaldığı görülmektedir.



Şekil 1. Bir-gün ve iki-günlük gözlenen maksimum yağışlar

Bir-gün ve iki-gün yağışları için görülen durum, üç, dört ve beş gün yağışları için de geçerlidir (Şekil 2). 1971 den 1996 yılına kadar, çok küçük değerlerde artan ve azalan bir trend görülmektedir. 1986-1991 yılları arasındaki günlük yağışların dağılımı birbirinden farklıdır. 3-günlük yağışlarda bir artış gözlenirken, 4-günlük yağışlarda bir azalma söz konusudur. 1991-1996 yılları arasındaki zaman periyodu incelendiğinde; 1991-1994 yılları arasında 3-günlük yağışlarda azalan ve 1995 yılında ise artan; 4-günlük yağışlarda artan; 5-günlük yağışlarda artan ve azalan bir trend görülmektedir. 1996-1998 yılları arasında tüm günler için yağış değerleri peak yapmıştır. 1999 yılından sonra görülemeye başlanan azalan ve artan trend günümüze kadar devam etmektedir.



Şekil 2. Üç-gün, dört-gün ve beş-günlük gözlenen maksimum yağışlar

Hidrolojik yapıların dizaynında, günlük hatta saatlik yağış miktarları çok önemlidir. Çizelge 1 ve 2’de 1-günden 5-güne kadar yağışların Gumbel ve Gamma dağılımı kullanılarak bulunan dönüşüm periyotlarına ait sonuçlar verilmiştir. Çizelge 1 ve 2 incelendiğinde, her iki dağılımın da birbirine yakın tahminlerde bulundukları gözlenmiştir. Ankara ili çok kurak bir iklime sahip olduğundan, düzensiz bir yağış rejimi vardır.

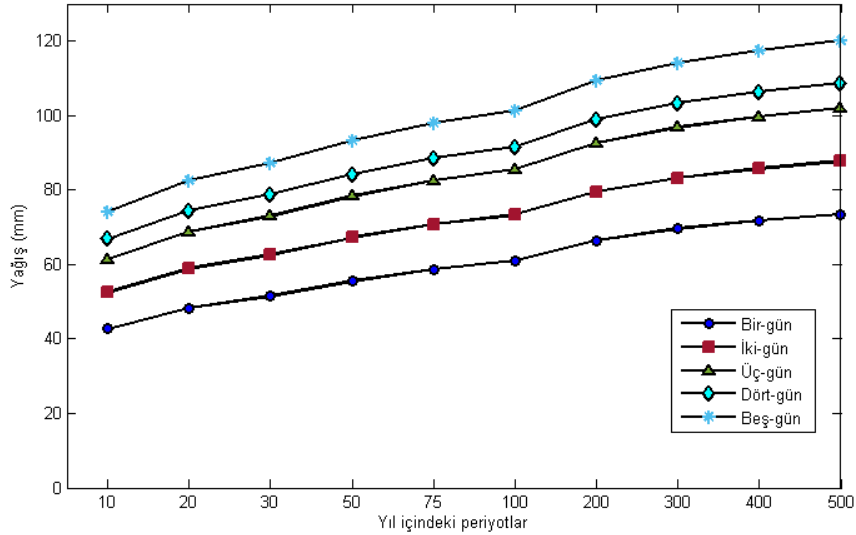
Çizelge 1. Gumbel dağılımına göre farklı dönüşüm periyotları için yağış miktarı

Yağış / dönüşüm periyotları	10 yıl	20 yıl	30 yıl	50 yıl	75 yıl	100 yıl	200 yıl
Bir-gün	42.66	48.28	51.50	55.54	58.72	60.97	66.40
İki-gün	52.51	58.91	62.59	67.19	70.82	73.39	79.57
Üç-gün	61.41	68.79	73.04	78.34	82.53	85.50	92.63
Dört-gün	66.88	74.47	78.84	84.30	88.61	91.66	99.00
Beş-gün	74.14	82.49	87.29	93.30	98.04	101.40	109.46

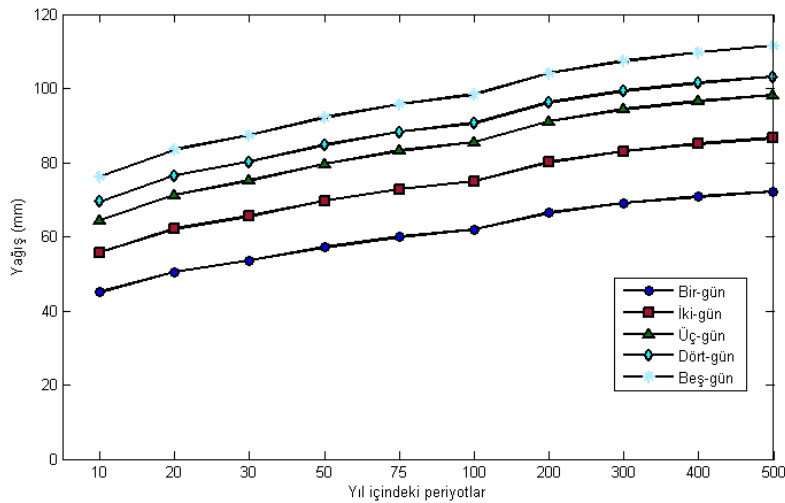
Çizelge 2. Gamma dağılımına göre farklı dönüşüm periyotları için yağış miktarı

Yağış / dönüşüm periyotları	10 yıl	20 yıl	30 yıl	50 yıl	75 yıl	100 yıl	200 yıl
Bir-gün	45.06	50.53	53.54	57.18	59.97	61.91	66.44
İki-gün	55.77	62.04	65.48	69.63	72.80	75.00	80.14
Üç-gün	64.37	71.29	75.07	79.63	83.11	85.52	91.14
Dört-gün	69.58	76.50	80.27	84.80	88.25	90.63	96.19
Beş-gün	76.16	83.45	87.41	92.18	95.80	98.30	104.14

En uygun dağılımları belirlemek için AIC ve BIC uyum iyiliği kriterleri kullanılmıştır. AIC ve BIC 'a ait en düşük değerler en iyi olasılık değeri olarak kabul edilir. Bir günlük Gamma dağılımı için AIC ve BIC değerleri sırasıyla 382.118 ve 385.943'tür. Bir günlük Gumbel dağılımı için AIC ve BIC değerleri sırasıyla 374.569 ve 378.393'tür. Her iki dağılım için AIC ve BIC değerlerine bakıldığında belirgin bir farklılık görülmediği için her iki dağılımda maksimum yağışları bulmak için kullanılabilir. Tüm günlerin dönüşüm periyotları Şekil 3 ve Şekil 4'de Gamma ve Gumbel dağılımlarına göre grafiksel olarak gösterilmiştir. Şekil 3 ve 4 incelendiğinde, her iki dağılımın arasında bir fark olmadığı ve 200 yıl periyotundan sonra yağışlarda bir değişim olmadığı da görülmektedir.



Şekil 3. Gumbel dağılımına göre farklı dönüş periyotları için yağış miktarı



Şekil 4. Gamma dağılımına göre farklı dönüş periyotları için yağış miktarı

TARTIŞMA

Ankara ili yağış istasyonu bir-gün ile beş-gün arasındaki yağışların toplamının çalışıldığı bu çalışmada, yağışlar için en uygun olasılık dağılımları belirlenmiştir. AIC ve BIC uyum iyiliği kriterleri kullanılarak yağışların Gumbel ve Gamma dağılımına göre dönüşüm periyotları tahmin edilmiştir. 10 yıldan başlayarak çok az miktarlarda olmak üzere 200 yıla kadar bir artış söz konusu iken; 200 yıldan sonra herhangi bir artış söz konusu değildir. Ayrıca, her iki dağılımdan elde edilen sonuçlar incelendiğinde aralarında önemli farkın olmadığı da görülmektedir. Özellikle hidrolojik yapıların dizaynı ve çalıştırılması ile tarımsal faaliyetlerde, yağışların dönüşüm periyotları çok önemlidir. Bu periyotların önceden tahmin edilmesi, meydana gelecek zararları önleyecek tedbirlerin alınmasını sağlayacak ve ekonomiye çok fazla katkı yapacaktır. Özellikle, iklim değişikliğinin çok fazla hissedildiği ülkemizde bu tür çalışmaların yapılması gerekmektedir. Karadeniz bölgesi gibi ekstrem yağışların çok fazla görüldüğü yerlerde toprak kayıplarını en aza indirmek, yapılan barajlarda sediment miktarını azaltmak, köprü ve yolların ömürlerini uzatmak açısından referans niteliğinde olacaktır.

Kaynaklar

- Anlı, A.S, 2006. Frequency Analysis for the Maximum Streamflows of Giresun Aksu Basin, Mediterranean Agricultural Sciences, Vol (19-1); 99-106.
- Gumbel, E.J, 1942. On The Frequency Distribution of Extreme Values in Meteorological Data, Bull. Am. Meteorol. Soc., 23, 95–104.
- Mayooran, T, Laheetharan, A, 2014. The Statistical Distribution of Annual Maximum Rainfall in Colombo District. Sri Lankan Journal of Applied Statistics, Vol (15-2).
- Nadarajah, S, Choi, D, 2007. Maximum Daily Rainfall in South Korea. J. Earth Syst. Sci. 116, No. 4, pp. 311–320, DOI:10.1007/s12040-007-0028-0.
- Nadarajah, S, 2005, The Exponentiated Gumbel Distribution with Climate Application, Environmetrics, 17(1), 13-23, DOI:10.1002/env.739.
- Nadarajah S., Withers, CS., 2001. Evidence of Trend in Return Levels for Daily Rainfall in New Zealand, Journal of Hydrology New Zealand, Volume 39, p.155-166.
- Salinas, J.L., Castellarin, A., Kohnová, S., Kjeldsen, TR., 2014. Regional Parent Flood Frequency Distributions in Europe—Part 2: Climate and Scale Controls. Hydrol. Earth Syst. Sci., 18, 4391–4401.
- Sen, Z, Eljadid, AG., 1999. Rainfall Distribution Function for Libya and Rainfall Prediction, Hydrol. Sci. J., 44(5), 665-680.
- Öztekin, T, 2011. Determination of The Best Fit Distributions for Daily Maximum Precipitations Measured at The City Centers Of Samsun, Sinop, Ordu and Tokat. Anadolu J Agr Sci., 26(3):194-202
- Varathan, N, Perera, K., Nalin, 2010. Statistical Modeling of Extreme Daily Rainfall in Colombo, M.Sc Thesis, Board of Study in Statistics and Computer Science of The Postgraduate Institute of Science, University of Peradeniya, Sri Lanka.

Evaluation of Glaucoma with Factor Analysis

Gülcan GENCER^{1*}, Kerem GENCER²

¹Karamanoglu Mehmetbey University, Karaman, Turkey

²Karamanoglu Mehmetbey University, Vocational School of Technical Sciences, Department of
Computer Programming, Karaman, Turkey

*Presenter author e-mail: gulcangencer@kmu.edu.tr

Abstract

Glaucoma, known as eye pressure, is the damage to the visual nerve due to the frequent increase in intraocular pressure. As a result, the person's field of vision gradually narrows. Glaucoma, an insidious disease that is noticed in the last stages of the disease, can cause serious damage to the optic nerve that cannot be repaired when diagnosed late. The main method used to reduce size in studies with interrelated variables is factor analysis. In this study, the variables that explain the glaucoma patients were divided into unrelated and fewer factors. A small number of conceptually meaningful new variables have been revealed, independent of the large number of variables used in the diagnosis of glaucoma, with the least loss of information.

Key words: Glaucoma, cup/disk, Eye pressure, SKK, RNFL.

INTRODUCTION

It is estimated that there are approximately 70 million glaucoma patients worldwide today. Almost half of these patients are unfortunately not aware that they have glaucoma. Due to the rapid increase in the elderly population, blindness due to glaucoma reaches approximately 6.7 million, and glaucoma ranks 2nd among the causes of blindness (Anonymous, 2019). People with higher than normal intraocular pressure have a higher risk of developing glaucoma; However, it does not mean that everyone with high intraocular pressure can have glaucoma. The risk of glaucoma increases in people over the age of 40. Glaucoma may be related to genetics. People with glaucoma in their family have a higher risk of developing it. In other words, there may be a defect in one or more genes and these individuals may become more susceptible to the disease. The risk of developing glaucoma is higher in patients with diabetes and hypothyroidism (goiter). Serious eye injuries can cause increased intraocular pressure. Other risk factors; retinal detachment, eye tumors, and eye inflammations such as chronic uveitis or iritis. Some eye surgeries can also trigger the development of secondary glaucoma. In myopia, which is generally known as nearsightedness, the frequency of glaucoma has increased approximately twice. Long-term use of cortisone may cause the development of secondary glaucoma. For such reasons, it is important that people with these features have regular eye exams for early detection of damage to the optic nerve. Pressure in the eye rises due to the inability to drain the intraocular fluid secreted in the eye and necessary for the eye to nourish.

The increased intraocular pressure also damages the eye nerve cells. As a result, glaucoma develops. Symptoms of glaucoma include headaches that become evident in the morning, blurred vision from time to time, light rings around lights at night, and pain around the eyes while watching television. Family history of glaucoma (genetic predisposition), being over 40 years old, Diabetes, Severe anemia or shocks, High-low systemic blood pressure (body blood pressure), High myopia, High hyperopia, Migraine, Long-term cortisone therapy, Eye injuries, Racial While family history is important, age,

gender, race, social life and habits are among the factors that increase the risk of glaucoma (Dünyagöz, 2020). Eye pressure is usually between 12 and 20 mm Hg in normal individuals.

In glaucoma patients, this value is generally above 20 mm Hg. Central corneal thickness (CCT) is the most important step in the diagnosis and treatment of glaucoma, in the correct measurement of intraocular pressure. Applanation tonometry measurements, which are accepted as the gold standard method in intraocular pressure measurement, are affected by central corneal thickness.

Goldmann tonometer gives the most accurate measurement with a 520 μm CCT. The central thickness of the cornea ranges from 537 to 554 μm in community studies (Whitacre, 1993). In normal eyes, the C / D (cup / disc) ratio is wider on the horizontal axis than the vertical axis. (Jonas, 1987) . In the early and middle stages of glaucoma, the vertical C / D ratio increases faster than the horizontal, and as a result, the horizontal / vertical C / D ratio is 1 ' to lower values than. Normally, the C / D ratio in healthy people is 0.25-0.30, while it is 0.5 in 10% of the population and 0.7 or more in 2% of the population.

An increase in this rate is accepted as an indicator of glaucomatous damage. The difference in C / D ratio between the two eyes is less than 0.2 in 99% of people and less than 0.1 in 92% (Armaly, 1967) in a study by Quigley et al. (1992). showed that the probability of developing glaucomatous visual field defects was higher in cases with a ratio of 0.55 and above than those with a C / D ratio of less than 0.55. Angle (irido-corneal angle) can be inherited and seen simultaneously in different members of the same family. It is more shallow in Asians and hyperops. In these people, the front camera is shallow compared to normal people. The angle at which the trabecular mesh is located between the cornea and iris is narrower than it should have been about 45 degrees. As the age gets older, this angle becomes narrower and the intraocular pressure increases due to the enlargement of the lens. When the angle is completely covered, acute glaucoma occurs (Anonymous, 2015).

Mean Deviation (MD): The mean deviation value is calculated by comparing the visual field measurement values, which are considered normal for the age of the patient, with the measurements obtained from the patient. It provides information about the magnitude of the visual field damage and is a parameter indicating extensive damage to the retinal nerve fiber layer (Wollstein, 2004). If the patient's values are outside the reference range of that age group in the society, p value is given next to the MD value. If this value is below 5%, it is considered as abnormal value. For example, $p < 1\%$ indicates that the test values of the patient are in the 1% of the individuals in that age group in the society (Werner, 2004).

Pattern Standard Deviation (PSD): Pattern standard deviation value indicates the irregularity of the patient's peak of vision surface. A low PSD value indicates that the hill is regular, and a high PSD value indicates that it is irregular. If the standard deviation value of the pattern is outside the reference ranges, it gives the p value in parentheses. A p value of less than 5% is considered abnormal (Stamper et al., 1999).

MATERIAL AND METHODS

Factor analysis; It is a set of multivariate statistical methods based on the correlation (or variance-covariance) matrix, as it is based on relationships between variables (Courtney, 2013). Factor analysis; It is a multivariate statistical method that aims to find a small number of conceptually meaningful new variables (factors), which are difficult to interpret, independent of many interrelated variables with minimal loss of information. The main purpose is; is to explain the total change with fewer factors than the number of variables (Alpar, 2017; Park, 2002). In addition to reducing the number of variables, this method enables the visualization of the data set, the associated data regression, etc. Biomedical research, epidemiological research for cancer diagnosis, clinical cardiovascular events determination studies in

cardiology, summarizing magnetic resonance images in radiology, pharmacology, etc. It is used in areas (Milewska, 2014).

Each variable was randomly generated by interviewing experts in the data set used in this study. The analysis of the data of 109 patients diagnosed with eye pressure; A total of 10 variables were defined: CCT, RNFL, MD, PSD, C / D, Angle, family history, accompanying systemic disease, age and gender. When the data obtained from each patient were examined, it was seen that the measurement units of the variables in the data set were different. In this case, the data have been standardized with the help of the following equation.

$$Z_i = \frac{x_i - \bar{x}}{\text{standart deviation}} \quad (1)$$

The data was analyzed by using Statistical Package for the Social Sciences (SPSS) -26.0 program. Factor Analysis were used. Table 1 shows descriptive statistics for variables used in glaucoma diagnosis.

Table 1. Descriptive Statistics of the data set

Variables	N	Minimum	Maximum	Mean	Std. Deviation
SKK	109	448,00	697,00	553,4679	43,50001
RNFL	109	38,60	223,00	91,3180	25,54361
C/D	109	0,10	1,00	0,4743	0,21534
Family history of glaucoma	109	0,00	1,00	0,0550	0,22912
Age	109	19,00	89,00	57,5046	14,10649
Sex	109	0,00	1,00	0,4037	0,49290
Concomitant Systemic Disease	109	0,00	1,00	0,3670	0,48421
Angle	109	16,10	135,91	36,1093	13,22806
MD	109	0,00	4,00	0,9817	0,99983
PSD	109	0,00	2,00	0,9450	0,57470

In Table 2, it is seen that the value of Kaiser-Mayer-Olkin measure is 0.697. From here, it is understood that factor analysis will be applied to the data.

Table 2. KMO and Bartlett's Test

Kaiser-Meyer-Olkin Measure of Sampling Adequacy.		0,697
Bartlett's Test of Sphericity	Approx. Chi-Square	146,782
	df	45
	Sig.	0,000

In Table 3, it is seen that the variable explained most by factor analysis is the accompanying systemic disease with 77.2%, the family history of glaucoma variable with 71.8%, the variable of age with 67.2%, the MD variable with 68.9%, and the variable of PSD with 66.4%.

Table 3. Communalities

Variables	Initial	Extraction
SKK	1,000	0,566
RNFL	1,000	0,536
C/D	1,000	0,415
Family history of glaucoma	1,000	0,672
Age	1,000	0,718
Sex	1,000	0,589
Concomitant Systemic Disease	1,000	0,772
Angle	1,000	0,509
MD	1,000	0,689
PSD	1,000	0,664

In Table 4, the number of factors is determined as 4, since the eigenvalue number is 4 components with greater than 1.

Table 4. Total Variance Explained

Component	Initial Eigenvalues		
	Total	% of Variance	Cumulative %
1	2,570	25,703	25,703
2	1,327	13,275	38,978
3	1,165	11,654	50,632
4	1,068	10,678	61,310
5	0,907	9,075	70,384
6	0,767	7,670	78,054
7	0,698	6,985	85,039
8	0,575	5,747	90,786
9	0,545	5,454	96,240
10	0,376	3,760	100,000

As can be seen from the scree plot in Figure 1, it is understood that the number of factors is 4.

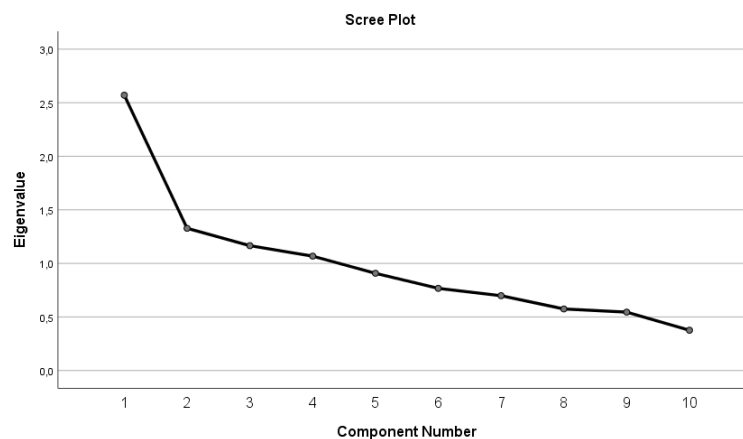


Figure 1. Total Variance Explained

Table 5 shows the total change explanation rates of the obtained factors. Accordingly, the first factor can explain 23.196% of the total change, the second factor 13.272% of the total change, the third factor 13.175% of the total change, and the fourth factor 11.667%.

Table 5. Total Variance Explained for Rotation Sums of Squared Loadings

Component	Total	% of Variance	Cumulative %
1	2,320	23,196	23,196
2	1,327	13,272	36,468
3	1,318	13,175	49,643
4	1,167	11,667	61,310
5			
6			
7			
8			
9			
10			

Since more than two factors were determined in the analysis of the data consisting of 10 variables related to 109 patients in Table 6, the main axes method, one of the factor finding techniques, was used. For the purpose of interpreting the factors more easily, the Varimax method, which enables the factors to be rotated so as to maximize the variances in the factor analysis, was used as a rotation technique. Here, the correlation value between variables and factors in absolute value is grouped under the factor of related variables consisting of 4 rotated factor axes with a value greater than 0.40.

Table 6. Rotated Component Matrixa

Variables	Component			
	1	2	3	4
SKK	-0,403			0,614
RNFL	-0,616			
C/D	0,632			
Family history of glaucoma				-0,780
Age		0,731		
Sex			-0,716	
Concomitant Systemic Disease		0,820		
Angle			0,672	
MD	0,829			
PSD	0,768			

The rotated factor matrix weights not only give the weights of the variables in the factors, but also show the direction of these weights in the factor with the +, - signs in front of them. It is in the same direction with the variables.

CONCLUSION

Glaucoma area ranks second among the causes of avoidable blindness in the world, that eye pressure, affecting close to 2.5 million people in Turkey, but due to late diagnosis of each patient treatment is performed only in one. In this study, glaucoma was collected under 4 unrelated factors. As a result, a

small number of conceptually meaningful new variables that are independent of the many variables used in the diagnosis of glaucoma with the least loss of information have been revealed.

REFERENCES

- Alpar R. Uygulamalı Çok Değişkenli İstatistiksel Yöntemler. Beşinci Baskı. Ankara: Detay Yayıncılık; 2017.
- Anonymous, (2019). Glokomun Körlükle Sonuçlanmasını Önlemek için Erken Tanı Çok Önemli, <http://www.yeditepehastanesi.com.tr/glokomun-korlukle-sonuclanmasini-onlemek-icin-erken-tani-cok-onemli>
- Anonymous, American Academy of Ophthalmology Glaucoma Panel (2015). Preferred Practice Pattern Guidelines. Primary Angle Closure. American Academy of Ophthalmology (www.aao.org/aboutpreferred-practicepatterns).
- Armaly MF.(1967). Genetic determination of cup/disc ratio of the optic nerve. Arch Ophthalmol, 78: 35–43.
- Courtney MGR. (2013). Determining the Number of Factors to Retain in EFA: Using the SPSS R-Menu v2.0 to Make More Judicious Estimations. Practical Assessment, Research, and Evaluation. 2013; 18(8).
- Dünyagöz, (2020), <https://www.dunyagoz.com.tr/tibbi-birimlerimiz/goz-tansiyonu-glokom/glokom-hakkinda#:~:text=Glokomun%20genetik%20ile%20ili%C5%9Fkisi%20olabilir,glokom%20geli%C5%9Fme%20riski%20daha%20fazlad%C4%B1r>.
- Jonas JB, Gusek GC, Guggenmoos-Holzmann I.(1987). Optic nerve head drusen associated with abnormally small optic disks. Int Ophthalmol, 11: 79-82.
- Milewska AJ, Jankowska D, Citko D, Wiesak T, Acacio B, Milewski R. (2014) The Use of Principal Component Analysis and Logistic Regression in Prediction of Infertility Treatment Outcome. Studies in Logic, Grammar, and Rhetoric., 39 (52).
- Park HS, Dailey R, Lemus D. (2002).The Use of Exploratory Factor Analysis and Principal Component Analysis in Communication Research. Human Communication Research, 28(4).
- Quigley HA, Katz J, Derick RJ, et al. (1992). An evaluation of optic disk and nerve fiber layer examinations in monitoring progression of early glaucoma damage. Ophthalmology, 99: 19–28.
- Stamper RL, Liberman MF, Drake MV. (1999).Visual field interpretation. Becker & Shaffer's Diagnosis and Therapy of the Glaucomas. 7 th Ed. St Louis-Missouri. Mosby; 144-171.,
- Werner EB. Visual field perimetry in glaucoma. Ophthalmology. Second Edition. Yanoff M, Duker JS, eds. Mosby. St Louis. 2004; 1441-1452.
- Whitacre M, Stein R, Hassanein K. (1993).The effect of corneal thickness on applanation tonometry. Am J Ophthalmol. 1993;115: 592-596.
- Wollstein G, Beaton S, Paunescu A, et al. (2004). Optical coherence tomography in glaucoma. Shuman JS, Puliafito GA, Fujimoto JG. Optical coherence tomography of ocular diseases. Second Edition. SLACK Inc. Thorofare, NJ.;483-8.

ACKNOWLEDGMENTS

The authors are grateful to Professor Ümit Kandaş from Medova Hospital, Konya, TURKEY for his valuable comments and contributions to the improvement of this paper.

CONFLICTS OF INTEREST

No conflict of interest was declared by the authors.

AUTHORS' CONTRIBUTIONS

Gülcan GENCER: Literature review, method design and planning, data collection and analysis, reporting. Kerem GENCER: Literature review, method design and planning, data collection and analysis, reporting.

Sağlık Bilimlerinde Seyrek Olumsallık Tabloları Üzerine Bir Çalışma

Gökçen ALTUN^{1*}

¹Bartın Üniversitesi, Fen Fakültesi, Bilgisayar Teknolojisi ve Bilişim Sistemleri Bölümü, 74100,
Bartın, Türkiye

*Sunucu yazar e-posta: gokcenefendioglu@gmail.com

Özet

Olumsallık tabloları iki ya da daha fazla kategorik değişkenin ortak sıklık dağılımıdır. Bir olumsallık tablosu iki ya da daha fazla değişkenin düzeylerinin çapraz sınıflandırılması olarak da düşünülebilir. Karesel olumsallık tabloları bağımlı örneklemelerde ortaya çıkan, satır ve sütun değişkenleri aynı düzeylere sahip olan çapraz tablolardır. Karesel olumsallık tabloları bir örnekleme yer alan bireylerin veya konuların herhangi bir özelliğe göre iki farklı değerlendirici tarafından birbirinden bağımsız şekilde derecelendirdiği sağlık çalışmalarında sıklıkla kullanılır. Örneğin, iki farklı doktor tarafından aynı hastaların hastalık derecelerine göre ayrı ayrı değerlendirmesi gibi. Karesel olumsallık tablolarının analizinde bazı özel modeller kullanılır. Bu modellerin temeli, ki-kare dağılımına yakınsayan simetri modelleri olarak tanımlanır. Özellikle mükemmel uyumun olması gereken sağlık çalışmalarında köşegen dışı elemanların sıfır veya sıfıra yakın değerler olması beklenir. Bir karesel olumsallık tablosu gözelerinde çok küçük sıklık veya örneklem sıfırı var ise bu tablo seyrek karesel olumsallık tablosu olarak adlandırılır. Seyrek olumsallık tablolarında ki-kare dağılımına yakınsama sağlanamadığı için bu modeller ile ilgili verinin analiz edilmesi istatistiksel olarak doğru sonuçlar vermemektedir. Bu nedenle, sıklıklardan etkilenmeyen, olasılıklar üzerinden hesaplanan sapma ölçülerini kullanarak bu tabloları yorumlamak daha güvenilir sonuçlar vermektedir. Bu çalışmada iki boyutlu karesel olumsallık tabloları için geliştirilen simetri modelleri ve sapma ölçüleri tanıtılmıştır. Ayrıca bu modeller ve sapma ölçüleri gerçek veri kümeleri üzerinde uygulanarak yorumlanmıştır.

Anahtar kelimeler: Sağlık bilimleri, Seyrek olumsallık tablosu, Simetri modeli

Quasi-Least Squares Regression Method with Veterinary Vaccination Data

Erdoğan ASAR^{1*}

¹Turkish Statistical Institute, Statistical Consultancy, 06100, Ankara, Turkey

*Presenter author e-mail: erdo6718@gmail.com

Abstract

Quasi-Least Squares Regression (QLS) is one of the methods used in longitudinal data analysis. Foot-and-Mouth disease (FMD) seen in cloven-hoofed animals is a highly contagious viral disease. Various vaccines and various routes of administration are used to protect animals from this disease. At the same time, studies on novel vaccination strategies and vaccines formulations have been studied. The main purpose of this study is to investigate the effect of vaccination methods used to protect animals from FMD and to introduce QLS method in data analysis for the veterinary vaccine studies. In addition, it is aimed to compare the results obtained with the regression model established with QLS with the results of the model established by Generalized Estimation Equations (GEE). In this study, which is an experimental study, FMD vaccine was administered in different routes and doses to three groups of cattle with similar characteristics. Vaccines were administered 0.5 ml intramuscularly to the first group, 0.5 ml intradermally to the second group, and 2 ml intramuscularly to the third group. Twenty-eight days after vaccination booster doses were administered using the same vaccine and routes. Blood samples were taken from cattle in these three groups at the specified six times to measure antibody titer levels (ATL). These measurements were first made before vaccination. Measurements were made five more times on the 7th, 14th, 28th, 58th and 88th days after vaccination. Virus neutralization test (VNT) for four different vaccine viruses were used to measure ATL. Analyzes on the longitudinal data obtained were performed with QLS and GEE methods using Stata (version 14) software. It was observed that the measurement times (days) in all ATL variables (response variables) were significant in both QLS and GEE methods ($p < 0.05$). On the other hand, the group effect was not found to be significant in both methods for all ATL variables ($p > 0.05$). For all response variables except one ATL variable, it was seen that the correlation value obtained with Markov working correlation structure was the highest correlation value among the correlation values obtained with all working correlation structures used in the analysis. The fact that the group variable effect on ATL variables was not found to be statistically significant means that the effects of 0.5 ml and 2 ml of intramuscular vaccine and 0.5 ml of intradermal vaccine on the groups were not different. Analyzed data is a type of data where there are missing measurements and measurements are made at unequal time intervals. Markov working correlation structure can only be applied to the QLS, which is suitable for such data. This gives QLS a privilege.

Key words: Quasi-least squares regression, Generalized estimating equations, Correlation structure, Foot-and-mouth disease

**Biyolojik Uygulamalarda Genlerin veya Özelliklerin Seçilmesi için Genetik
Algoritmalarla Dayalı Makine Öğrenimi**

Yalçın TAHTALI^{1*}, Lutfi BAYYURT¹

¹Gaziosmanpaşa Üniversitesi, Ziraat Fakültesi, Zootehni Bölümü, Tokat, Türkiye

*Sunucu yazar e-posta: yalcin.tahtali@gop.edu.tr

Özet

Gen veya özellik seçiminde, makine öğrenimi tekniklerini kullanarak model oluşturulması önemli ön işleme adımlarından birisidir. Bu adımlar, biyo belirteçlerin tanımlanması gibi farklı biyolojik uygulamalarda da önemli bir rol oynamaktadır. Yapılan bazı çalışmalarda özellik / gen seçme algoritması ve yöntemi üzerinde durulmuş olmasına rağmen, parametre ayarı veya düşük performans seviyesi gibi sorunlar ile karşılaşmıştır. Bu tür problemlerin üstesinden gelmek için, bu çalışmada, optimizasyon algoritmasına ve genetik algoritmaya (GA) dayalı bir yaklaşımı sunulmuştur. Bu çalışmada, yöntemi ortaya koymak amacıyla, kanser tedavisi için kullanılan ilaç, kanser teşhisi ve klinik uygulamalar gibi çeşitli biyolojik kapsamlardan farklı özelliklere sahip on üç sınıflandırma ve regresyon tabanlı veri kümesi seçilmiştir. Yapılan analiz sonucunda, elde edilen sonuçlar incelendiğinde, önerilen yöntemin diğer yöntemlere göre daha iyi performansa sahip olduğunu göstermiştir. Sonuç olarak, minimum sayıda gen ve maksimum ayrılabilirlik gücüne sahip bir biyo belirteç, tanı kiti elde edilmesi amacıyla biyolojik uygulamaları optimize edebilmektedir.

Anahtar kelimeler: Gen, Özellik seçimi, Makine öğrenmesi

Wine Veri Setinin Hiyerarşik Kümeleme Analizi Yöntemi ile İncelenmesi

Gülşah KEKLİK^{1*}

¹Çukurova Üniversitesi, Ziraat Fakültesi, Zootečni Bölümü, 01330, Adana, Türkiye

*Sunucu yazar e-posta: gulsahkeklk@gmail.com

Özet

Kümeleme analizi, öğelerin birbirine olan benzerliklerine göre gruplandırılarak verilerin birbiriyle olan ilişkilerini açıklamak için kullanılan çok değişkenli analiz yöntemlerinden biridir. Hiyerarşik kümeleme analizi, kümelerin iç içe geçerek alt-üst ilişkisine dayanan bir hiyerarşi söz konusu olduğunda kullanılan kümeleme analizi yöntemlerinden biridir.

Bu çalışmada, "wine" veri setinin R'de hiyerarşik kümeleme yöntemi ile analizi yapılmıştır. Çalışmada, NbClust paketindeki "NbClust" fonksiyonu ile optimum küme sayıları belirlenmiş olup "cutree" ile kesme işlemi gerçekleştirilmiştir. Kümeler arasındaki mesafe, bağlantı yöntemleri aracılığıyla hesaplanarak küme dendogramları oluşturulmuştur.

Anahtar kelimeler: Çok değişkenli analiz, Hiyerarşik kümeleme analizi

Examination of Wine Data Set by Hierarchical Clustering Analysis Method

Abstract

Clustering analysis is one of the multivariate analysis methods used to explain the relationships of data with each other by grouping items according to their similarities. Hierarchical clustering analysis is one of the clustering analysis methods used when there is a hierarchy based on sub-up relationship by intertwining clusters.

In this study, the "wine" data set was analyzed by hierarchical clustering method in R. In the study, optimum cluster numbers were determined with the "NbClust" function in the NbClust package, and cutting was performed with "cutree". Cluster dendograms were created by calculating the distance between the clusters using connection methods.

Key words: Multivariate analysis, Hierarchical cluster analysis

GİRİŞ

Kümeleme, nesneleri veya durumları birbirlerine olan benzerliklerine veya yakınlıklarına göre sınıflandırma işlemidir. Birbirine yakın öğeler bir gruba, yakın olmayan öğeler ise farklı bir gruba yerleştirilmektedir. Öğeler arasındaki benzerliğin yüksek ve kümeler arasındaki farklılıkların fazla olması, kümeleme işleminin ne kadar iyi bir şekilde yapıldığını ortaya koymaktadır. Bu sebeple, kaç küme oluşturulacağının tespiti için incelenen verinin doğru bir şekilde anlaşılması ve yorumlanması önemli bir husustur. Kümeleme analizi ise incelediğimiz bir veri setindeki öğeleri sınıflandırmak için faydalanan yöntemlerin bütünü ifade etmektedir.

YÖNTEM

Kümeleme Yöntemleri

Kümeleme analizi, üzerinde en fazla çalışılan çok değişkenli istatistik analiz yöntemlerinden biridir. Kümeleme yöntemi, genel olarak hiyerarşik ve hiyerarşik olmayan yöntemler olmak üzere

iki gruptan oluşmaktadır (Downs ve Barnard, 1995). Bu çalışmanın kapsamındaki analizler, hiyerarşik kümeleme yöntemi uygulanarak tamamlanmıştır.

Hiyerarşik Kümeleme Yöntemi

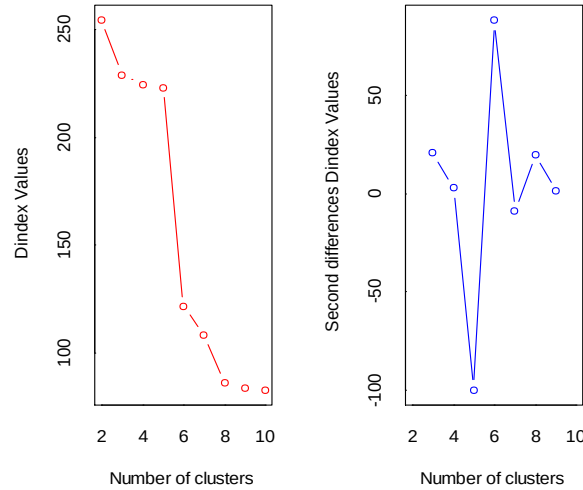
Hiyerarşik kümeleme yönteminde, kümeler birbiri ile iç içe geçecek şekilde oluşturulurlar. Kümeler oluşturulurken alt-üst ilişkisinden bahsedebileceğimiz bir hiyerarşi söz konusu olduğu için hiyerarşik kümeleme yöntemi denilmektedir.

R’de hiyerarşik kümeleme yöntemi için stats paketinde bulunan “hclust” nesnesi kullanılmaktadır. Bu çalışmanın kapsamında “wine” veri setinin hiyerarşik kümeleme yöntemi ile incelenmesi gerçekleştirilmiş olup bağlantı yöntemlerine göre kesim yükseklikleri arasındaki ilişkiler araştırılmıştır. NbClust paketindeki “NbClust” fonksiyonu ile optimum küme sayılarını belirlemek üzere elde edilen kümelerin ağaç grafikleri çizilerek hiyerarşik kümeleme yöntemi ile üretilen ağaçlara “cutree” ile kesme işlemi uygulanmıştır. Kümeler arasındaki mesafe, bağlantı yöntemleri aracılığıyla hesaplanarak küme dendogramları oluşturulmuş ve bağlantı yöntemlerine ait dendogramların bir arada incelenmesi olanak bulmuştur.

BULGULAR

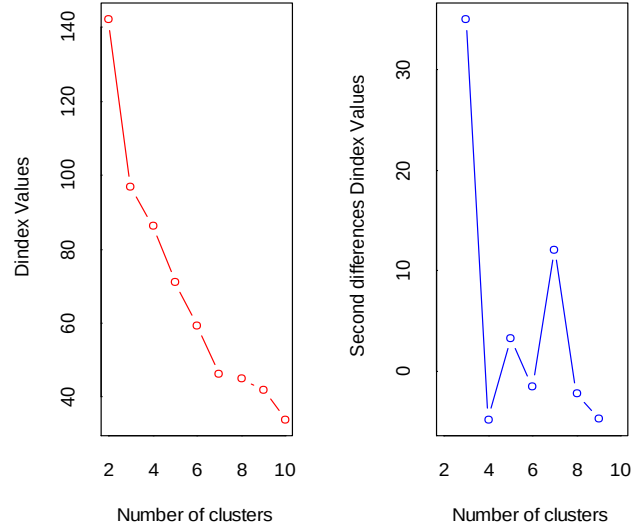
Uygun küme sayısının belirlenmesi için elde edilen kümelerin elverişli bir şekilde oluşturulması ve hangi elemanlardan meydana geldiğini bilmek önemlidir. Bunun için birtakım bağlantı yöntemleri geliştirilmiştir. Bu bağlantı yöntemlerinin oluşturulması, NbClust paketindeki “NbClust” nesnesinin döndürülmesiyle gerçekleştirilmiştir. R’deki NbClust, kümeler arasındaki uzaklık ölçülerini araştırarak en uygun kümeleme değerini göstermesi bakımından kümeleme analizinde tercih edilen bir pakettir.

Aşağıda verilen “tek” bağlantı yönteminde uygun küme sayısı (k) 2 olarak gözlemlenmektedir.



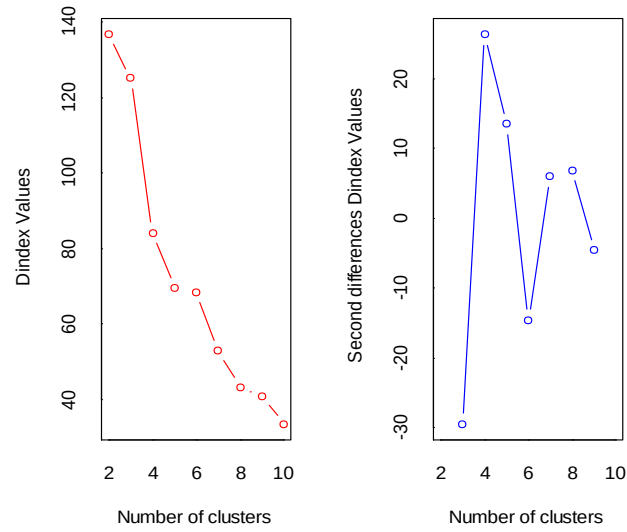
Şekil 1. Tek bağlantı yöntemine göre küme sayısının belirlenmesi

Aşağıda verilen “tam” bağlantı yönteminde uygun küme sayısı (k) 7 olarak gözlemlenmektedir.



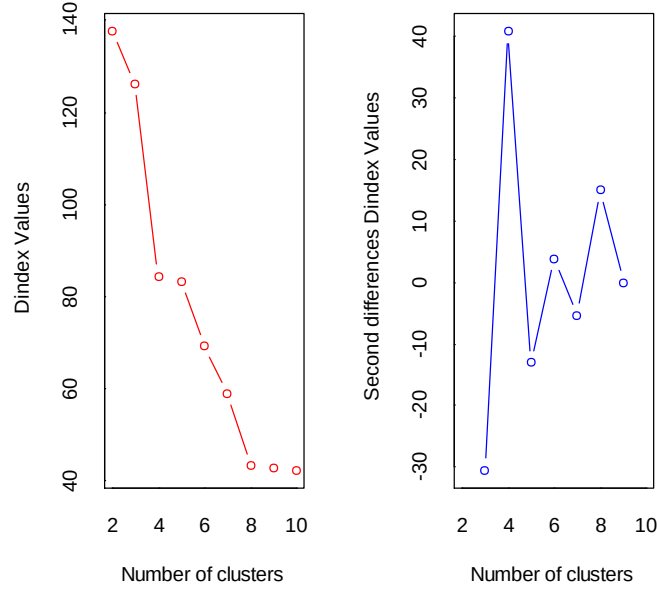
Şekil 2. Tam bağlantı yöntemine göre küme sayısının belirlenmesi

Aşağıda verilen “average” bağlantı yönteminde uygun küme sayısı (k) 2 olarak gözlemlenmektedir.



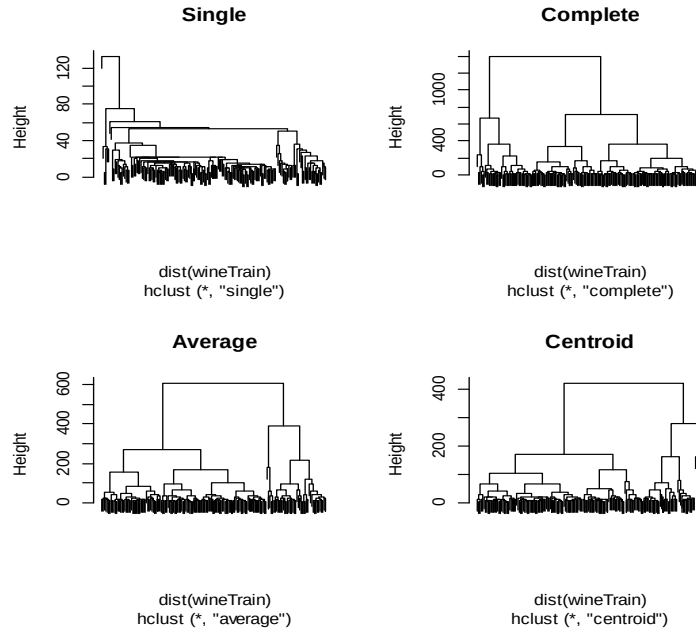
Şekil 3. Average bağlantı yöntemine göre küme sayısının belirlenmesi

Aşağıda verilen “centroid” bağlantı yönteminde uygun küme sayısı (k) 4 olarak gözlemlenmektedir.



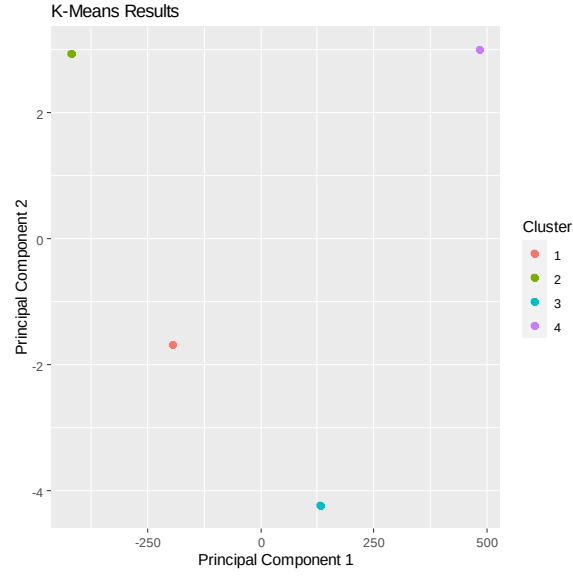
Şekil 4. Centroid bağlantı yöntemine göre küme sayısının belirlenmesi

Hiyerarşik kümeleme yöntemleri, “hclust” nesnesinde “method” parametresine single (tek), complete (tam), average (ortalama), centroid (merkez) bağlantı yöntemleri yazılarak farklı birleşme yüksekliklerine (h) sahip küme dendogramları oluşturulmaktadır. Böylelikle yöntemlere ait ağaç grafikleri arasındaki ilişkilerin bir arada gözlemlenmesi sağlanmaktadır. Ağaç grafikleri, “hclust” nesnesinde bulunan kümeleme sonuçlarından faydalanılarak R’deki “plot” fonksiyonu ile çizilmiştir.

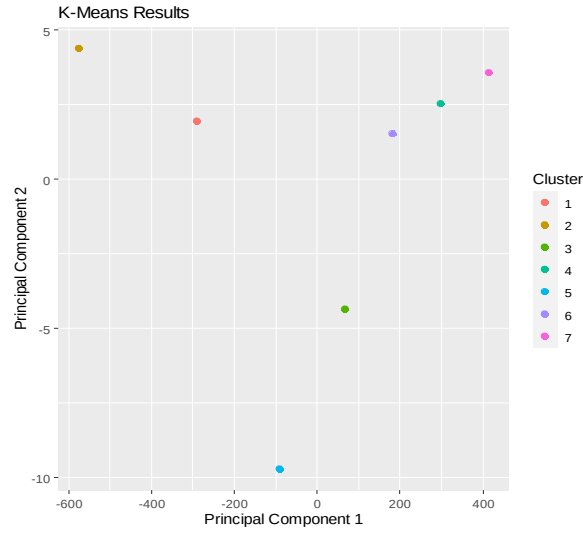


Şekil 5. Bağlantı yöntemlerine ilişkin küme dendogramları

Kümelerdeki elemanların dağılımını gözlemlemek ve görselleştirmek için serpilme grafiklerinden faydalanılır. Veriye ait elemanların hangi kümeye ait oldukları, kenarlarında verilen renk kodları ile belirlenebilmektedir. center=4 ve center=7 için oluşturulan serpilme grafikleri aşağıda verildiği gibidir.

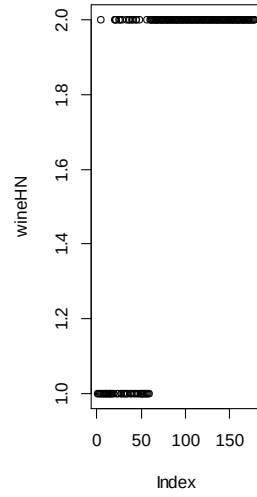


Şekil 6. Center=4 için kmeans sonuç grafiği

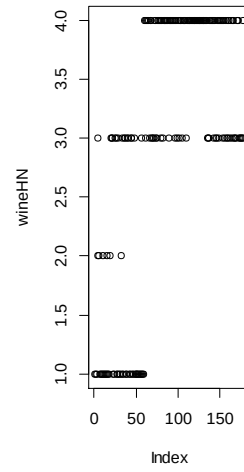


Şekil 7. Center=7 için kmeans sonuç grafiği

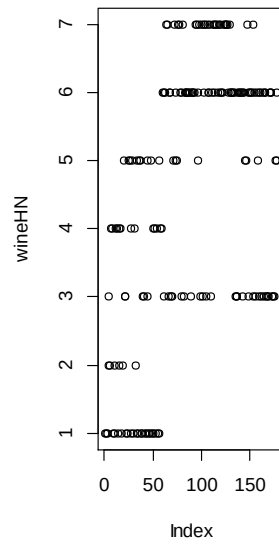
Kümeleme işleminde, uygun bir küme sayısının belirlenmesi ve dendogramın küme sayısına göre belirlenen herhangi bir yükseklikten kesilmesi gerekmektedir. Bunun için “cutree” ile kesme işlemi yapılmıştır. Alternatif olarak “rect.hclust” nesnesi de kullanılmaktadır.



Şekil 8. Kümelerin cutree nesnesi ile kesim işlemi (k=2 için)



Şekil 9. Kümelerin cutree nesnesi ile kesim işlemi (k=4 için)



Şekil 10. Kümelerin cutree nesnesi ile kesim işlemi (k=7 için)

SONUÇ VE TARTIŞMA

Kümeleme analizi, küme elemanlarının benzerliklerine göre gruplandırılması ve hangi kümeye ait olduklarının tespiti açısından önemli bir istatistik analiz yöntemidir. Bu çalışmadaki araştırmada, R’de “wine” veri seti üzerinde, kümeleme yönteminin analiz edilmesi sağlanmıştır.

Çalışmada, NbClust paketindeki “NbClust” fonksiyonu ile optimum küme sayıları belirlenmiş olup kümelerde “cutree” ile kesme işlemi gerçekleştirilmiştir. Bağlantı yöntemlerine ilişkin küme dendogramları oluşturularak kümeleme analiz yöntemine ilişkin sonuçlar yorumlanmıştır.

Bu çalışmaya ek olarak, karışıklık matrisleri, box-plot ve birleşme/ayırma yüksekliği grafikleri aracılığıyla kümeleme sonuçlarının karşılaştırılması işlemlerinin yapılması alternatif olarak önerilmektedir.

Kaynaklar

- Cebeci, Z., 2021. Biyoenformatik Veri Analizinde R ile Hiyerarşik Kümeleme, Papatya Yayıncılık Eğitim, İstanbul.
- Charrad, M. Ghazali, N., Boiteau, V., Niknafs, A. 2014. Package: ‘NbClust’: NbClust package for determining the best number of clusters. (url: <http://cran.r-project.org/web/packages/NbClust/NbClust.pdf>)
- Downs, G.M., Barnard, J.M. 1995. Hierarchical and non-Hierarchical Clustering. EUROMUG Meeting, UK.. k on 21.09.2014.
- Kaufman, L., Rousseeuw, P. J. 1990. Finding Finding Groups in Data: An Introduction to Cluster Analysis. Wiley Online Library.
- R Core Team, 2017. R: A language and environment for statistical computing. R Foundation for Statistical Computing, Vienna, Austria. URL <https://www.R-project.org>.
- Rajalingam, N., Ranjini, K., 2011. Hierarchical Clustering Algorithm-A Comparative Study. International Journal of Computer Applications, 19(3):42-46. (Accessed at <http://www.ijcaonline.org/volume19/number3/pxc3873052.pdf> on 22nd May, 2014).

Parametrik ve Parametrik Olmayan Korelasyon Analizinde Korelasyon Katsayısının Kullanımı ve Yorumlanması

Gülşah KEKLİK^{1*}

¹Çukurova Üniversitesi, Ziraat Fakültesi, Zootehni Bölümü, 01330, Adana, Türkiye

*Sunucu yazar e-posta: gulsahkeklk@gmail.com

Özet

Korelasyon analizi, modele alınacak değişkenler arasındaki ilişkinin büyüklüğünü ve yönünü açıklamak amacıyla sıklıkla başvurulan bir istatistiktir. Korelasyon analizinden bahsedildiğinde yaygın olarak Pearson korelasyonu ve Pearson korelasyonunun parametrik olmayan karşılığı olan Spearman korelasyonu algılanmaktadır. Parametrik olmayan korelasyon yöntemlerinin Spearman korelasyonu ile sınırlı olmadığı ve birçok alternatif korelasyon yöntemlerinin varlığı aşikardır. Değişkenler, normal dağılıma uygunluk gösterdiğinde Pearson korelasyonu, aksi durumda genel olarak Spearman korelasyonu yapılmaktadır. Korelasyon katsayısı(r), -1 ile $+1$ arasında değerler almakta olup korelasyon katsayısının -1 veya $+1$ olması, değişkenler arasında mükemmel bir ilişki olduğunu gösterirken 0 olması, herhangi bir ilişki olmadığını göstermektedir. Bu çalışmada, örnek bir veri seti üzerinden öncelikle normallik varsayımlarının kontrolü sağlanmış, sonrasında ise sırasıyla parametrik ve parametrik olmayan korelasyon analiz yöntemlerinden Pearson ve Spearman korelasyonunun uygulanması ve korelasyon katsayısının yorumlanması gerçekleştirilmiştir.

Anahtar kelimeler: Korelasyon katsayısı, Parametrik korelasyon katsayısı, Parametrik olmayan korelasyon katsayısı

Use and Interpretation of Correlation Coefficient in Parametric and Nonparametric Correlation Analysis

Abstract

Correlation analysis is a statistic that is frequently used to explain the magnitude and direction of the relationship between the variables to be modeled. When correlation analysis is mentioned, Pearson correlation and Spearman correlation, which is the non-parametric equivalent of Pearson correlation, are commonly perceived. It is obvious that non-parametric correlation methods are not limited to Spearman correlation and there are many alternative correlation methods. Pearson correlation is used when the variables conform to the normal distribution, otherwise Spearman correlation is generally used. The correlation coefficient (r) takes values between -1 and $+1$, and a correlation coefficient of -1 or $+1$ indicates a perfect relationship between the variables, while a 0 indicates no relationship. In this study, first of all, normality assumptions were checked on a sample data set, afterwards the application of Pearson and Spearman correlation, which is one of the parametric and nonparametric correlation analysis methods, respectively and the interpretation of the correlation coefficient were performed.

Key words: Correlation coefficient, Parametric correlation coefficient, Nonparametric correlation coefficient

Çoklu Durum Sistem Modellerinin Güvenilirlik Değerlendirmesi

Ahmet DEMİRALP^{1*}

¹Harran Üniversitesi, İİBF, Ekonometri Bölümü, 63050, Şanlıurfa, Türkiye

*Sunucu yazar e-posta: ahmt.dmrp@gmail.com

Özet

Mühendislik uygulamalarında kullanılan sistem modelleri için hem sistemin hem de sistemi oluşturan bileşenlerin iki mümkün durumundan söz edilebilir: ya çalışıyor ya da çalışmıyor. Bu şekilde kurduğumuz sistemler ikili sistemler olarak adlandırılmaktadır. Bu sistemler çeşitli mühendislik sistemlerinin güvenilirliklerini artırma konusunda önemli bir yere sahiptirler. Ama bazı mühendislik sistemleri için ikili varsayım her bir sistemin deneyimleyebileceği olası durumları doğru bir şekilde temsil etmekte yetersiz kalmaktadır. Bu yüzden ikiden daha fazla olası durumları içeren çoklu durum sistem modelleri geliştirilmiştir. Bu modelde, hem sistem hem de bileşenler ikiden fazla performans düzeyiyle karşılaşabilmektedir. Başka bir ifadeyle bir sistemin ve bileşenlerin $T+1$ mümkün durumu olsun $(0,1,2,...,T)$. Bu durumda 0 tamamen hatalı durumu, T tamamen çalıştığı durumu ve aradaki durumlar da diğer dereceli durumları göstermektedir. Bu çalışmada hem doğrusal hem de dairesel çoklu durum ardıl bağlı sistemlerin yapısını incelenecek ve birer örnekle güvenilirlik değerlendirmelerinin hesaplanması verilecektir.

Anahtar kelimeler: İkili sistem, Doğrusal (dairesel) çoklu durum ardıl sistem, Güvenilirlik

Mean Estimation for Sensitive Study Variable Using Ln-type Estimators under Stratified Sampling

Muhammad Nouman QURESHI^{1*}, Yousaf FAIZAN¹, Muhmmad HANIF¹

¹National College of Business Administration and Economics, Lahore, Pakistan

*Presenter author e-mail: nqureshi633@gmail.com

Abstract

In this paper, we have proposed ln-type estimators for estimating the mean of sensitive study variable in the presence of non-sensitive auxiliary variable under stratification. The mathematical equations of bias and mean square error of the proposed estimators are derived using first-order ln and Taylor expansions. Some special cases of proposed estimators are also obtained using different parameters associated with the auxiliary variable. Algebraic comparisons of the proposed estimators are also made with combined mean and ratio estimators. Simulation study is conducted to check the performance of the proposed estimators over the existing estimators using artificial populations. Real data application is also used to verify findings obtained from the simulation.

Key words: Auxiliary information, Ln-type estimators, Percentage relative efficiency, Randomized response technique, Sensitive study variable

2-Boyutlu Ardıl-k Sistemlerin Güvenilirliği için Yaklaşımlar

Ahmet DEMİRALP^{1*}

¹Harran Üniversitesi, İİBF, Ekonometri Bölümü, 63050, Şanlıurfa, Türkiye

*Sunucu yazar e-posta: ahmt.dmrp@gmail.com

Özet

Bir sistem, bileşen sayısı n ve en az ardıl k tane başarısız (çalışan) bileşenden oluşmaktadır. Yani bu sistemin hatalı (çalışan) olması için sistemdeki bileşenlerin en az ardıl k tanesinin hatalı (çalışan) olması yeterlidir. Bu tür sistemler, bileşenlerin tek boyutlu düzenlenmesi ile ilişkilendirilirler. Son zamanlarda ise tek boyutlu güvenilirlik modellerini iki veya üç veya d -boyutlu versiyonlarına genişletme çabalarını görmekteyiz ($d \geq 2$). Salvia ve Lasher tek boyutlu ardıl k sistemleri çok boyutlu sistemlere uyarlayan ilk yazarlardır. Tanımlamalarına göre, 2 ya da 3-boyutlu sistem n -kenarlı kare ya da kübik ızgaralar şeklindedir. Bu şekilde oluşan sistem ancak ve ancak tüm arızalı bileşenleri içeren en az bir kare veya k namlı küp varsa başarısız olur ($2 \leq k \leq n-1$). Bu çalışmada 2-boyutlu doğrusal ardıl k sistemleri inceledik. Bu sistemlerin genel özellikleri incelendikten sonra güvenilirlik değerlendirmeleri için bazı yaklaşık çözüm veren sınır yaklaşımlarını derledik. Derlediğimiz bu yöntemleri 2-boyutlu sistemlere uygun bir uygulama ile karşılaştırarak birbirlerine üstünlüklerini belirlemeye çalıştık.

Anahtar kelimeler: Doğrusal 2-boyutlu sistemler, Sınır yaklaşımlar, Güvenilirlik

Computational Insight into Analysis of the Effects of Deleterious SNPs of MAPT Gene in Alzheimer's Disease

Orcun AVŞAR¹, Aysenur SEN*

¹Hitit University, Faculty of Arts and Sciences, Department of Molecular Biology and Genetics,
19000, Corum, Turkey

*Presenter author e-mail: aysenursen97@yandex.com

Abstract

Alzheimer's disease is the most common cause of dementia, mostly seen in the elderly. Alzheimer's disease is a multifactorial neurodegenerative disorder characterized by cognitive impairment and memory loss. The disease can be categorized as early onset (EO, <65 age) and late onset (GO, > 65 age) Alzheimer's disease. Abnormal accumulation of microtubule-associated protein tau and inclusion bodies formation are seen in Alzheimer's disease. Mutations in the MAPT gene encoding the tau protein have been associated with Alzheimer's disease. In this study, we aimed to identify the effects of deleterious mutations in the MAPT gene. Deleterious SNPs of the MAPT gene associated with Alzheimer's disease was determined by the use of numerous computational servers -SIFT, Polyphen-2, PROVEAN, PhD-SNP, SNP&GO-. The effects of these SNPs on protein stability were analyzed by the I-Mutant 3.0 algorithmic program. Post-translational modifications (PTMs) affected by high-risk missense SNPs were identified by the MusiteDeep program. According to the results of the present study, five damaging missense SNPs in the human MAPT gene were shown to affect post translational modifications by modulating glycosylation at position 427 and ubiquitination at position 615. It was concluded that the deleterious missense SNPs can be implicated in the pathogenesis of Alzheimer's disease. Further experimental studies conducted with the reported deleterious SNPs in this study are required for confirmation of the findings.

Key words: Alzheimer' disease, MAPT, SNP, PTMs, Tau

Identification of Post-Translational Modifications Mediated by High-Risk Deleterious SNPs of Human DRD2 Gene Associated with Schizophrenia

Orcun AVSAR¹, Aysenur SEN^{1*}

¹Hitit University, Faculty of Arts and Sciences, Department of Molecular Biology and Genetics,
19000, Corum, Turkey

* Presenter author e-mail: aysenursen97@yandex.com

Abstract

Schizophrenia is a psychiatric disorder with a wide range of clinical symptoms and biological characteristics. Some patients experience a single psychotic episode, while others experience a remission period and gradual functional impairment with various negative symptoms. Although it is known that disturbances in dopaminergic neurotransmission is implicated in the pathogenesis of schizophrenia, the molecular mechanism of schizophrenia etiology is unclear yet. The aim of the present study was to investigate the deleterious single nucleotide polymorphisms (SNPs) and their effects in the DRD2 gene, which encodes the dopamine receptor D2 protein and is associated with schizophrenia. High-risk deleterious SNPs that have the potential to affect the structure and function of the dopamine receptor D2 and post-translational modifications (PTMs) mediated by these deleterious SNPs were determined by various bioinformatic methods. The four missense SNPs-rs148478679, rs199517364, rs200148328 and rs201724929- were found to have very damaging effects on structure, function, and stability of dopamine receptor D2 protein. However, any post-translational modifications at positions 120, 126, 150 and 219 were not affected by these high-risk deleterious SNPs. Based on the findings of bioinformatic analysis, our study supposes that reported high-risk SNPs may have the potential to serve as molecular targets for the treatment strategies of schizophrenia.

Key words: DRD2, Schizophrenia, SNP, PTMs

Parametrik Olmayan Çok Değişkenli Varyans Analizi ve Beyaz Eşya Sektöründe Bir Uygulama

Serhat YILMAZ^{1*}, Yüksel TERZİ¹

¹Ondokuz Mayıs Üniversitesi, Fen Bilimleri Enstitüsü, İstatistik Anabilim Dalı, 55139, Samsun, Türkiye

²Ondokuz Mayıs Üniversitesi, Fen Edebiyat Fakültesi, İstatistik Bölümü, 55139, Samsun, Türkiye

*Sunucu yazar e-posta:styilmazz@gmail.com

Özet

Perakende satış ve pazarlama sektörü veri analizlerinde etkenlerin birbirleriyle ilişkili birden fazla sonucu ve bu ilişkilerin ayrı ayrı analiz edilmesi I. Tip hatanın artmasına yol açmaktadır. Ancak, genellikle veri analizlerinde bu hata sıklıkla yapılmakta olup tek değişkenli analiz yöntemlerinden yararlanılmakta ve değişkenlerin önemli bir bölümü tek değişkenli model varsayımlarından olan normal dağılım ve varyansların homojenliği varsayımlarını sağlamamaktadır. Bu nedenle daha doğru ve güvenilir sonuç üretmesi açısından yeni metot ve yöntemlere ihtiyaç duyulmaktadır. Bu çalışmada, değişkenler arasındaki ilişkileri dikkate alan ve varsayımları az olan parametrik olmayan çok değişkenli varyans analiz yöntemlerinden, PERMANOVA (Parametrik Olmayan Çok Değişkenli Varyans Analizi) yönteminin teorik özellikleri anlatılmış ve perakende satış ve pazarlama üzerine elde edilen bir veri setine uygulanmıştır. Perakende satış mağazalarının bağlı oldukları satış grupları ile ürün grupları arasındaki ilişkileri, il-ilçe performansları, yıllık cirosal karşılaştırmalar incelenmiştir. Analiz sonucunda, ilgili değişkenler arasında farklılık olmamasına rağmen ürün grubu segmentlerinde yüksek oranda değişimler olduğu tespit edilmiştir.

Anahtar kelimeler: Permütasyon, Yapay F değeri, PERMANOVA

Non-Parametric Multivariate Analysis of Variance and an Application in the White Goods Industry

Abstract

In retail sales and marketing sector data analysis, more than one result of the factors related to each other and analyzing these relationships separately lead to an increase in Type I error. However, this error is often made in data analysis, univariate analysis methods are used and a significant part of the variables do not provide the assumptions of normal distribution and homogeneity of variances, which are univariate model assumptions. For this reason, new methods and methods are needed in order to produce more accurate and reliable results. In this study, the theoretical features of the PERMANOVA (Non-Parametric Multivariate Analysis of Variance) method, one of the non-parametric multivariate analysis of variance methods that take into account the relationships between the variables and have few assumptions, are explained and applied to a data set obtained on retail sales and marketing. The relations between the sales groups and product groups of the retail stores, the provincial-district performances, annual turnover comparisons were examined. As a result of the analysis, although there is no difference between the relevant variables, it has been determined that there are high changes in the product groups segments.

Key words: Permutation, Artificial F value, PERMANOVA

GİRİŞ

Veri analizlerinde deneysel faktör türlerinin, veriler üzerindeki bireysel etkinin yanı sıra birlikte etkisi araştırılıyor ve test edilmesine ihtiyaç duyulmaktadır. Çok değişkenli analizlerde ilgilenilen veri setlerinin birbirleriyle de mutlaka bir ilişkisi vardır. Bu ilişkiler göz ardı edilerek, tek değişkenli analizler ile değerlendirilmesi 1. Tip hata yapma olasılığını arttırmaktadır. Bu gibi araştırmalarda değişkenlerin birden fazla ilişkili sonuçlarının olması durumunda, tutarlı sonuçlar veren tek değişkenli varyans analizinin genelleştirilmiş hali olan çok değişkenli varyans analizi yönteminin kullanılması gerekmektedir. Bu yöntem sonuç değişkenlerinin çok değişkenli normal dağılım ve kovaryans matrislerinin homojenliği gibi varsayımların sağlanması gerekmektedir. Sonuç değişkenlerinin çok değişkenli normal dağılım göstermesi, varyans ve kovaryans matrislerinin homejenliği, değişken sayısı fazla olan birçok alanda karşılaşılabilecek bir durum değildir. Değişkenlerin çok değişkenli normal dağılıma uymaması veya varyans kovaryans matrisinin homojen olmaması durumunda çok değişkenli varyans analizi tutarlı sonuçlar vermeyecektir. Değişken sayısının örneklem sayısından büyük olduğu çalışmalarda ise genellikle hesaplanamamaktadır. Parametrik olan bu yöntemin varsayımları olduğu için kullanılıp kullanılmaması her çalışma için farklılık göstermektedir. Varsayımların sağlanmaması durumunda parametrik olmayan yöntemlerin kullanılması önerilir (Anderson, 2001).

Çok değişkenli varyans analizinde, parametrik ve parametrik olmayan testlerin arasındaki farkları incelediğimizde; parametrik testlerin tutarlı sonuçlar vermesi ve yapılabilmesi için normallik varsayımı ve grup varyansların eşitliği istenmektedir (Anderson, 2001). Bağımlı değişkenler arasındaki korelasyonlara hassas ve örneklemdeki birim sayısının değişken sayısından daha fazla olması şartını gerektirir, değişken sayısı örneklemdeki birim sayısından daha fazla olmamalıdır (Anderson, 2001).

Parametrik olmayan testlerde normal dağılım varsayımı olmayıp, gruplar arasındaki varyanslar homejen olmayabilir. Korelasyon hassaslığı gerektirmez. Kitle üzerindeki dağılım varsayımları için oldukça güçlüdür (Kıroğlu, 2001). Sonuç olarak, parametrik olan MANOVA testi gözlemlerin sıfır değer alması durumuna karşı aşırı derecede hassas iken parametrik olmayan PERMANOVA testi bu duruma karşı hassas ve duyarlı değildir (Pasin ve ark., 2016).

MATERYAL VE YÖNTEM

Bağımlı değişken sayısının birden fazla ve bağımsız değişken sayısı ikiden fazla olduğunda ise iki yönlü MANOVA testi kullanılır. Ancak satış pazarlama sektöründe varsayımların sağlandığı veri seti çok kısıtlı sayıdadır. Hatta birçok veri seti, çok değişkenli varyans analizi yapısına uymasına rağmen varsayımları sağlanmadığı durumda kullanılmamalıdır (Gower ve Krzanowski, 1999). Dolayısıyla varsayımların sağlanmadığı durumda alternatif olarak parametrik olmayan testlerden PERMANOVA testi tercih edilebilir. MANOVA testinde gözlemlerin sıfır değeri almasına çok hassas iken, Permaova bu duruma hassas ve duyarlı değildir (Anderson, 2001).

Yöntem, uzaklık veya benzerlik ölçülerinden yararlanmaktadır. Noktalar arasındaki uzaklıkların kareler toplamını, merkezlerin nokta sayısına bölünen ara noktaların uzaklıkların kareler toplamına eşit olmasıdır. F test istatistiği kullanılmakla beraber, p değeri hesaplanırken permütasyon teknikleri kullanılır. Bu sebeple bu yöntem dağılım varsayımı gerektirmemektedir (Anderson, 2001). Testin tek gerektirdiği varsayımı gözlemlerinin bağımsız olduğunu ve benzer dağılım gösterdiklerini varsaymaktadır. Bunun dışında herhangi bir varsayımı yoktur. Testin çok katı varsayımının olmaması kullanım alanlarını da arttırmaktadır.

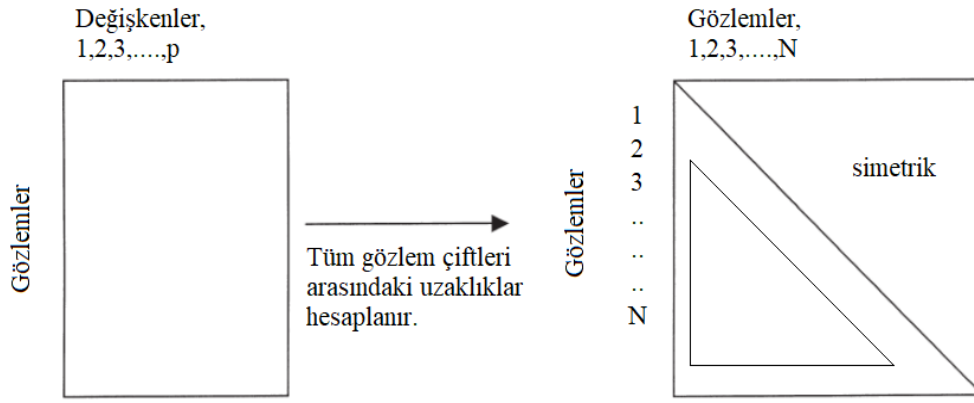
Tek yönlü modelde a tane grup ve her grupta n tane gözlem olduğunu düşünürsek toplam gözlem sayısı $N=an$ 'dir. Gözlemlerin i. ve j. bireylerinin gözlemler arasındaki uzaklık d_{ij} olmak üzere grup için kareler toplamı ve genel kareler toplamı aşağıdaki gibi olacaktır.

$$SS_T = \frac{1}{N} \sum_{i=1}^{N-1} \sum_{j=i+1}^N d_{ij} \quad (1)$$

$$SS_W = \frac{1}{N} \sum_{i=1}^{N-1} \sum_{j=i+1}^N d_{ij}^2 \varepsilon_{ij} \quad (2)$$

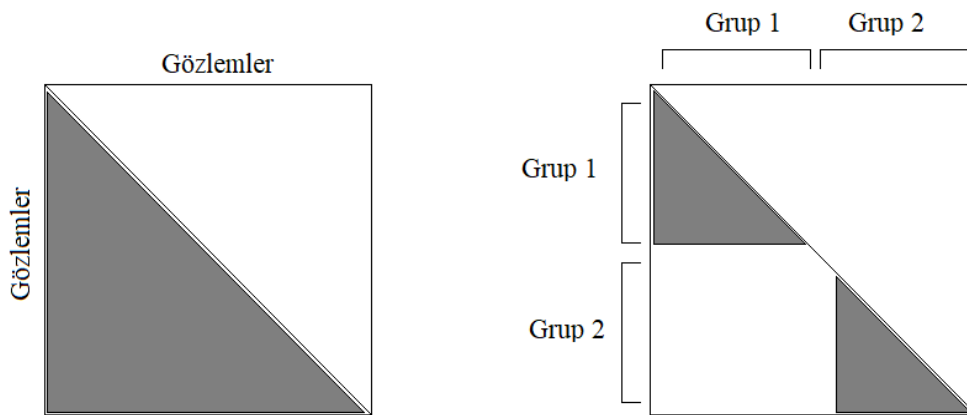
$$SS_A = SS_T - SS_W \quad (3)$$

- (1) numaralı denklem genel kareler toplamının formülüdür.
 (2) numaralı denklem grup için kareler toplamının formülüdür.
 (3) numaralı denklem gruplar arası kareler toplamının formülüdür.



Şekil 1. Genel kareler toplamının gösterimi

Genel kareler toplamı hesaplanırken uzaklık matrisinin alt köşegenine ait uzaklıkların kareleri toplanır ve N'ye bölünür (Pasin, 2016).



kareler toplamının gösterimi

Şekil 2. Grup içi Grup içi kareler toplamının hesaplanabilmesi için hata terimi dikkate alınmalıdır. Burada hata teriminin i. ve j. gözlemleri aynı grupta ise 1, farklı gruplarda ise 0 değerini alan bir katsayıdır. Bu katsayıya aynı grupta yer alan gözlemler arasındaki uzaklıkların kareleri eklenir Şekil 3.2.'deki gibi gösterilir (Pasin, 2016). Gruplar arasındaki farklılıklar için kareler toplamı hesaplanması ise genel kareler toplamından grup için kareler toplamının çıkarılması ile bulunur. F değeri

hesaplanırken gruplar arası serbestlik derecesi $a-1$, grup için serbestlik derecesi ise $N-a$ olarak alınır. Bu bilgiler ile hesaplanan yapay F test istatistiği, Fisher'in F dağılımına uygunluk göstermediği için p değerinin elde edilebilmesi için bilinen F tablosu kullanılmaz. Bu PERMANOVA yöntemini diğer yöntemlerden ayıran en önemli özelliktir (Pasin, 2016).

Klasik F tablosu kullanılamayacağı için I.Tip hata olmayacak olup p değerinin hesaplanmasında permütasyon yöntemlerinden yararlanılmaktadır. Permütasyon yöntemi ile F değeri hesaplanacağı için kurulan sıfır hipotezinin doğru olduğu varsayılmaktadır. Ancak ve ancak gruplar arasında farklılıklar yok ise gözlem değerleri gruplara rasgele olarak dağılabilecektir. Gözlemler her adımda gruplara rasgele dağıtılarak belirli bir gruba tanımlanır. Permütasyonlar sonucu elde edilen yapay F değeri F^π ile gösterilir. Bu rasgele elde edilen F^π değeri olası tüm sonuçlarla ilişkili satırların yeniden oluşturulmuş hali için hesaplanır. Grup sayısı a tane olmak üzere gruplarda n tekrar ile tek yönlü bir testte F^π istatistiğinin olası ayrı sonuçlarının sayısı $(an)!/(a!(n!)^a)$ kadardır. Sıfır hipotezinin doğruluğu varsayımı altında permütasyonlar sonucunda hesaplanan F değerinin dağılımı, orijinal satırlardan hesaplanan F değeri ile karşılaştırılır ve p değeri şu şekilde hesaplanır (Pasin, 2016).

$$p = \frac{(\text{Her bir } F^\pi \geq F \text{ koşulunun sağlanma sayısı})}{(\text{Toplam } F^\pi\text{'lerin sayısı})}$$

Denklemden F değeri permütasyonlar sonucunda elde edilen F^π dağılımının bir üyesi niteliğindedir. Tüm olası permütasyonların hesaplamak çok vakit alacağından pek tercih edilmemektedir. Bu nedenle p değeri olası permütasyonlardan rasgele seçilen bir alt seti kullanılarak da hesaplanabilir. Permütasyon sayısının fazla olması p değerlerinin duyarlılığını arttırdığı unutulmamalıdır.

PERMANOVA yönteminde üç farklı permütasyon yöntemi vardır. Bunlar satır veri, indirgenmiş model artıkları, tam model artıkları yöntemleridir. Örneklem hacmi, kullanılan veri çeşidi ve kullanım amacına göre dezavantaj veya avantajı vardır. Permütasyon yöntemlerinin açıklama ve avantajları şu şekildedir.

Satır veri yöntemi karmaşık ANOVA çalışmaları için tutarlı sonuçlar vermektedir. I.Tip hata α değerine yakındır. Örneklem hacmi büyük ise daha tutarlı sonuç vermektedir. Lakin çok büyük örneklem hacmine gerek yoktur (Manly, 1997). İndirgenmiş model artıkları yöntemi, veri seti karmaşık bir yapıya sahip ise I. Tip hatayı vermektedir (Pasin, 2016).

Uzaklık ölçülerinden metrik ve yarı metrik uzaklık ölçüleri kullanılmaktadır. Bunlardan metrik uzaklık ölçüleri Öklid, Manhattan ve Minkowski iken yarı metrik olanlar ise ki-kare, Bray Curtis, Jaccard'tır. Bu ölçülerin benzer özellikleri aşağıdaki gibidir. Farklı uzaklık ölçülerinin kullanımına izin veren bir yöntem olduğundan kullanımı kolay olup diğer yöntemlerden pozitif ayrılmaktadır (Kelly ve Gross, 2015).

Yarı metrik uzaklık özellikleri; i ve j herhangi iki nokta olmak üzere,

$i=j$ ise i ve j arasındaki mesafeyi gösteren $D_{(i,j)} = 0$ 'dır. Her zaman pozitifdir (eğer $i \neq j$ ise $D_{(i,j)} > 0$ 'dır). Simetriktir ($D(i,j)=D(j,i)$).

İki faktörlü bir tasarımı ele alalım a tane seviyesi olan birinci faktör A olarak adlandırılın. Aynı şekilde b tane seviyesi olan ikinci faktör B olarak adlandırılın. İki faktörün her ab kombinasyonundan n tekrar yapıldığı varsayılır ise gözlem sayısı $N=abn$ olacaktır (Pasin, 2016).

Çizelge 1. Kareler toplamı tablosu

Değişim Kaynağı	Kareler Toplamı
A (B faktörü göz ardı edildiğinde)	$SS_{W(A)} = \frac{1}{bn} \sum_{i=1}^{N-1} \sum_{j=i+1}^N d_{ij}^2 \varepsilon_{ij}^{(A)}$
B (A faktörü göz ardı edildiğinde)	$SS_{W(B)} = \frac{1}{an} \sum_{i=1}^{N-1} \sum_{j=i+1}^N d_{ij}^2 \varepsilon_{ij}^{(B)}$
A	$SS_A = SS_T - SS_{W(A)}$
B	$SS_B = SS_T - SS_{W(B)}$
AB	$SS_{AB} = SS_T - SS_A - SS_B - SS_R$
Artık	$SS_R = \frac{1}{n} \sum_{i=1}^{N-1} \sum_{j=i+1}^N d_{ij}^2 \varepsilon_{ij}^{(AB)}$

A faktörü göz ardı edilir ise B faktörleri için, B faktörü göz ardı edilir ise A faktörleri için hesaplama formülleri Çizelge 1’deki gibidir.

Formüller de yer alan $\varepsilon_{ij}^{(A)}$ ve $\varepsilon_{ij}^{(B)}$ terimleri i ile j terimleri sırasıyla A ve B faktörlerinin aynı grubunda yer alıyorsa “1” yer almıyorsa “0” değerini almaktadır. $\varepsilon_{ij}^{(AB)}$ terimi ise i ve j gözlemleri A ve B faktörlerinin aynı kombinasyonunda yer alıyorsa “1” yer almıyorsa “0” değerini almaktadır. Analizde her bir ana etki için karışık gelen kareler toplamı A ve B faktörleri için SS_A ve SS_B şeklinde gösterilmektedir. Artık kareler toplamı, A ve B faktörleri için her ab kombinasyonları içinde ara nokta uzaklıklar dikkate alınarak hesaplama yapılmakta ve SS_R şeklinde gösterilmektedir (Pasin, 2016).

Faktörlerin iç içe geçtiği modellerde etkileşim terimi olmadan aynı yaklaşım kuralları uygulanmaktadır. $SS_{B(A)} = SS_T - SS_A - SS_R$ olacaktır. Formülde yer alan $SS_{B(A)}$ terimi, B faktörünün A faktörü içine geçtiğini ifade etmektedir. Tek yönlü analiz yönteminde benzer şekilde hesaplanan F değerleri F tablo değerleri ile kullanılmamakta ve p değerlerinin hesaplanmasında permütasyon yöntemleri kullanılmaktadır. Permütasyon yöntemlerinin çok iyi seçilmesi faktör sayısı ile doğrudan orantılıdır. Yanlış seçilen bir permütasyon yöntemi sonuçları da etkileyeceği için yorumlamanın hatalı olmasına yol açmaktadır. Permütasyon yöntemleri ile kesin bir p değeri elde etmek için sıra faktör sayısının az olması gerekmektedir. Sıra faktör sayısı arttıkça kesin bir p değeri elde etmek zor olacak ve çok fazla vakit almaktadır. Bu durumun avantaja dönüşmesi için kesin permütasyon testleri yerine analizde yer alan bütün terimlere karşılık gelen satır verilerin ve artıkların beraberce permütasyonunu içeren yaklaşımsal permütasyon yöntemi kullanılmalıdır (Pasin, 2016).

Veri toplama araçları

Beyaz eşya sektöründe üretim yapan bir firmanın bölgelerine bağlı rasgele dört il seçilip, bu illerdeki perakende satışları ürün grubu ve dönemsel olarak seçilmiş ve bu değişkenlerin alt verileri için kodlama yapılmıştır. Bölge değişkeni dört farklı ili gösteriyor olup, 1 İç Anadolu, 2 Orta Anadolu, 3 Karadeniz, 4 Doğu Anadolu’dur. Yıl değişkeni iki yıl dönemini gösteriyor olup 1 2018, 2 2019’dur. Çeyrek değişkeni dört alt dönemi ifade ediyor olup, Q1 Ocak Şubat Mart, Q2 Nisan Mayıs Haziran, Q3 Temmuz Ağustos Eylül, Q4 Ekim Kasım Aralıktır. Segment değişkeni ise 1 orta segment 2 üst segment ürünlerin satışlarını göstermektedir. Çalışmada bölge ve segmentlere ait veriler kullanılmıştır.

BULGULAR

Bölge satış adetleri analiz edildiğinde verilerin normal dağılım göstermediği görülmüştür. Beyaz eşya perakende satışlarına yönelik değişkenler için tek ve iki yönlü modeller oluşturulmuş olup, uzaklık ölçülerinden p değeri en düşük olan uzaklık ölçü çıktıları ile hipotezler değerlendirilmiştir.

Bölgeler ve segmentler arasında anlamlı farklılık olduğu görülmüştür. İç Anadolu bölgesinin ürün grubu satış adetleri incelendiğinde diğer bölgelere kıyasla farklılık olduğu tespit edilmiştir. Tüm bölgelerde ürün gruplarının orta ve üst segment satışları incelendiğinde, anlamlı farklılıklar görülmüş ve farklılığın en yüksek olduğu bölge İç Anadolu bölgesi olmuştur. İç Anadolu bölgesinde üst segment satışları, orta segmentlere göre anlamlı derecede farklı çıkmıştır.

Veri grubumuzda bölgeler, yıllar, çeyrekler ve segmentler olmak üzere dört tane grup değişkenimiz mevcuttur. Bu değişkenlerin ürün grupları ile tek yönlü incelenmesi yapılarak farklılıkların olup olmadığına bakılmıştır.

Hipotezler,

H_0 : Grup değişkenleri arasında farklılık yoktur

H_1 : En az iki grup arasında farklılık vardır

Çizelge 2. Tek yönlü model için öklid uzaklığı çıktıları (N=10)

Açıklama	Bölge	Yıl	Çeyrek	Segment
Permüstasyon sayısı (N)	10	10	10	10
F istatistiği	16,20	0,52	1,09	12,39
p değeri	0,09	0,73	0,64	0,09

Permüstasyon sayısı 10 olarak belirlenmiştir. Tüm grup değişkenleri için $p>0,05$ olduğundan H_0 hipotezi red edilemez. Grup değişkenleri arasında anlamlı bir farklılık yoktur.

Çizelge 3. Tek yönlü model için öklid uzaklığı çıktıları (N=1000)

Açıklama	Bölge	Yıl	Çeyrek	Segment
Permüstasyon sayısı (N)	1000	1000	1000	1000
F istatistiği	16,20	0,52	1,09	12,39
p değeri	0,00	0,61	0,39	0,00

Permüstasyon sayısı 1000 olarak belirlenmiştir. Yıl ve çeyrek grupları için $p>0,05$ olduğundan H_0 hipotezi red edilemez. Bölge ve Segment grubu için $p<0,05$ olduğundan bu gruplar arasında anlamlı bir farklılık vardır.

Çizelge 4. Bölge ve Segment Öklid uzaklığı Analiz Çıktıları

Kaynak	S.D.	F istatistiği	p-değeri
Bölge	3	39,34	0,00
Segment	1	43,93	0,00
Etkileşim	3	15,27	0,00
Artıklar	56		
Toplam	63		

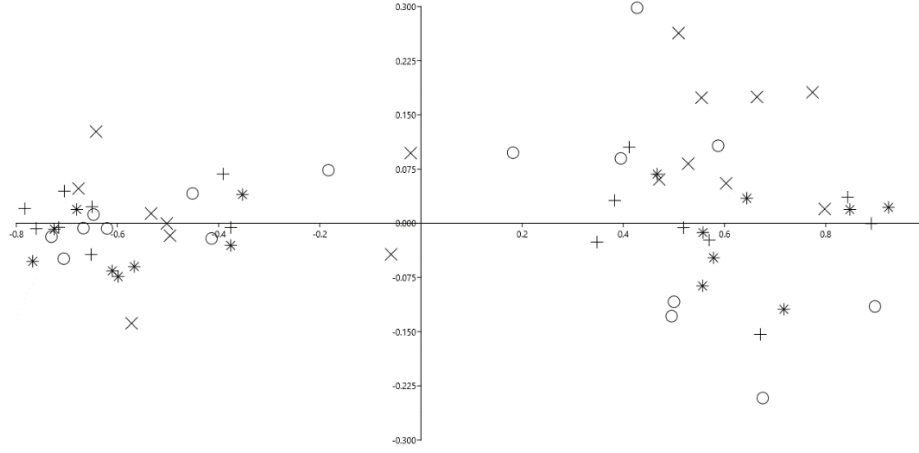
Çizelgelerdeki analiz sonuçları incelendiğinde uzaklık ölçüleri arasında sonucu değiştirecek bir fark olmadığı görülmektedir. Bölgeler ve segmentlerin p-değeri 0,05'den küçük olduğu için bölgeler arasında anlamlı bir farklılık vardır.

Etkileşim hipotezi şu şekilde kurulabilir,

H_0 : Bölge ve segmentlerin birlikte etkisinin ürün gruplarına bir etkisi yoktur

H_1 : En az iki grup arasında bir etkisi vardır

Bölgelerin ve segmentlerin ürün grupları üzerinde birlikte etkisi için etkileşim parametresine bakılmalıdır, p-değeri $0,00 < 0,05$ olduğundan H_0 hipotezi red edilir ve bölgeler ile segmentlerin birlikte etkisinin olduğu söylenebilir.



Şekil 3. Bölgelerin ürün grubu üzerindeki dağılımı

Şekil 3’de bölgelerin ürün gruplarına göre doğrudan satış adetlerinin dağılımı göstermektedir. Bölgelerin dağılışı ortak bir noktada birleşmemesinden dolayı bölgeler arasında bir farklılık olduğu yorumu yapılabilir.

TARTIŞMA VE SONUÇ

Birçok alanda olduğu gibi satış pazarlama alanında da bazı değişkenler normal dağılmamakta ve parametrik testlerin gerekli varsayımları sağlanamamaktadır. Literatürde Değişkenleri birbirleriyle ilişkili birden fazla sonucunu birlikte değerlendiren, normallik varsayımı gibi parametrik testlerin varsayımları sağlanmadığı için PERMANOVA yöntemi ile incelenmiştir. Veri setinde permütasyon sayısı arttıkça uzaklık ölçülerinin birbirine yakın sonuçlar verdiği gözlenmiştir. Permütasyon sayısı 50’nin altında iken uzaklık ölçü sonuçları birbirinden uzak ve seçilmesi durumunda sonucu değiştirdiği, uygun olan uzaklık ölçüsünün doğru seçilmesi gerektiği ve permütasyon sayısı artırılarak, hem tek yönlü hemde çift yönlü analizler yaparak karşılaştırmalara yer verilmiştir. Yarı metrik ve tam metrik olan uzaklık ölçülerinden Öklid çıktılarına yer vererek analiz sonuçları çizelgeler halinde verilmiştir.

Tek yönlü model sonuçları incelendiğinde, bölgeler ile ürün grupları arasında sonuçlar incelendiğinde, İç Anadolu, Orta Anadolu, Karadeniz ve Doğu Anadolu bölgeleri arasında anlamlı bir farklılık olduğu görülmüştür. Segmentler ile ürün grupları arasında sonuçlar incelendiğinde, orta ve üst segment arasında anlamlı bir farklılık olduğu görülmüştür.

İki yönlü model sonuçları incelendiğinde, bölgeler ve segmentlerin ürün grupları üzerine, anlamlı bir farklılık olduğu görülmüştür. Bu farklılığın İç Anadolu bölgesinin üst segment satışlarından kaynaklandığı Şekil 4’de grafiksel olarak gösterilmiştir. İç Anadolu bölgesinin ürün grubu satış adetleri incelendiğinde diğer bölgelere kıyasla farklılık olduğu tespit edilmiştir. Tüm bölgelerde ürün gruplarının orta ve üst segment satışları incelendiğinde, anlamlı farklılıklar görülmüş ve farklılığın en yüksek olduğu bölge İç Anadolu bölgesidir. İç Anadolu bölgesinde üst segment satışları, orta segmentlere göre anlamlı derecede farklılık olduğu görülmüştür.

Uygulamadan elde ettiğimiz sonuçlarda, üst segment satışlarının Orta Anadolu, Karadeniz ve Doğu Anadolu bölgelerinde gerçekleşen satışlar ile uyumlu olduğu ve bu bölgelerde üst segment ürün satışlarına yönelik üretim yapılması firma tarafında çalışmalar yapılabilir.

Çalışmalarda PERMANOVA analizinin çok kısıtlı alanda kullanıldığı görülmüş olup özellikle son yıllarda kullanım sıklığı artmaya başlamıştır. Peralta vd. (2012), bitki, mikrobiyal ve genel bakteri topluluklarına çevresel faktörlerin etkileri araştırılıyor. Elde edilen veri seti PERMANOVA analizi ile incelenmiştir. Tang vd. (2016), mikrobiyal toplulukların disbiyoz üzerindeki etkilerini incelemişler. Mesafeye dayalı yöntemlerin çoğu tek bir mesafe kullanır ve yanlış seçilir ise yanlış sonuçlar verdiğini söyleyerek, PERMANOVA analizine çoklu mesafeler dahil edilerek yeni bir mesafe analizi yöntemi olan PERMANOVA-S önermişlerdir. Zhu vd. (2020), Mikrobiyom çalışmalarında eşleşen küme verileri üzerine PERMANOVA ile doğrusal ayrışma modelini (LDM) kullanılarak iki analizden hangisinin daha esnek olduğunu çıkarmışlar. İki analiz yönteminde eşleşmiş küme verilerinde esneklik sağladığı görülmüştür.

Kaynaklar

- Anderson, M. J. 2001. A new method for non-parametric multivariate analysis of variance, *Austral Ecology*, 26, 32-46. doi: 10.1111/j.1442-9993.2001.01070. pp.x
- Anderson, M. J. 2001. Permutation tests for univariate or multivariate analysis of variance and regression. *Canadian Journal of Fisheries and Aquatic Sciences* 58, 626-639. doi: 10.1139/cjfas-58-3-626
- Gower, J. C. and Krzanowski, W. J. 1999. Analysis of distance for structured multivariate data and extensions to multivariate analysis of variance, *Journal of the Royal Statistical Society*, 48, 505-519. doi: 10.1111/1467-9876.00168
- Hammer, Ø., Harper, D. and Ryan, P. D. 2001. PAST: paleontological statistics software package for education and data analysis. *Palaeontol. Electron.* 4-9
- Kelly, B. J. and Gross, R. 2015. Power and sample-size estimation for microbiome studies using pairwise distances and PERMANOVA, *Bioinformatics*, 31:15, 2461-2468. doi: 10.1093/bioinformatics/btv183
- Kıroğlu G. 2001. Uygulamalı Parametrik Olmayan İstatistiksel Yöntemler (Altıncı Baskı), 148-187, Türkiye.
- Manly, BFJ., 1997. Randomization, Bootstrap and Monte Carlo Methods in Biology (2nd edition), Chapman and Hall, UK.
- Pasin, Ö., Ankaralı, H. ve Cangür Ş. 2006. Parametrik olmayan çok değişkenli varyans analizi ve sağlık alanında bir uygulaması, *Bilişim Teknolojileri Dergisi*, 9:1, 13-20. doi: 10.17671/btd.47787
- Peralta, A. L., Matthews, J. W., Flanagan, D. N. and Kent, A. D. 2012. Environmental factors at dissimilar spatial scales influence plant and microbial communities in restored wetlands, *Society of Wetland Scientists*, 32, 1125-1134. doi: 10.1007/s13157-012-0343-3
- Tang, Z. Z., Chen, G. and Alekseyenko, A. V. 2016. PERMANOVA-S: Association test for microbial community composition that accommodates confounders and multiple distances, *Human Genetics*, 96:5, 797-807.
- Zhu, Z., Satten, G., Caroline, M. and Hu, Y.J. 2020. Analyzing matched sets of microbiome data using the LDM and PERMANOVA, *Research Square*, 1-16. doi: 10.21203/rs.3.rs-17148/v1

Daimî İkametgâh ve Suç Türüne Göre İllerin (İBBS 3. Düzey) Kümeleme Analizi ile İncelenmesi

Ceren Yaman Yılmaz ^{1*}

¹Ankara Hacı Bayram Veli Üniversitesi, İktisadi ve İdari Bilimler Fakültesi, Adalet Yönetimi Bölümü,
06500, Ankara, Türkiye

*Sunucu yazar e-posta: ceren.yaman@hbv.edu.tr

Özet

Suç, insanlığın varlığıyla ortaya çıkmış ve pek çok araştırmaya konu olmuş bir olgudur. Suçla mücadele kapsamında atılacak adımların belirlenmesinde suçun gerçekleştiği ve suç türlerinin yoğunlaştığı coğrafyanın tespiti önem arz etmektedir. Bu sebeple, Türkiye İstatistik Kurumu'ndan temin edilen Türkiye İstatistiki Bölge Sınıflaması (İBBS) 3. Düzeye göre 2019 yılına ait suç göstergeleri kullanılarak, çok değişkenli istatistiksel analiz yöntemlerinden biri olan kümeleme analizi yapılmıştır. Suç türleri bakımından birine benzeyen iller farklı kümeleme yöntemleri uygulanarak kümelenebilir ve tablolar halinde sunulabilir yorumlanmıştır.

Anahtar kelimeler: Çok Değişkenli İstatistik, Kümeleme Analizi, Suç

Examination of Provinces by Permanent Residence and Type of Crime (NUTS 3rd Level) with Cluster Analysis

Abstract

Crime is a phenomenon that emerged with the existence of humanity and has been the subject of many studies. In determining the steps to be taken within the scope of the fight against crime, it is important to determine the geography where the crime occurs and the types of crimes are concentrated. For this reason, cluster analysis, which is one of the multivariate statistical analysis methods, was carried out by using crime indicators for the year 2019 according to the 3rd Level of the Turkey Statistical Region Classification (NUTS) obtained from the Turkish Statistical Institute. Provinces similar to one in terms of crime types were clustered by applying different clustering methods and presented in tables and interpreted.

Key words: Multivariate Statistics, Cluster Analysis, Crime

GİRİŞ

Dünyada ve Türkiye’de suç ile ilgili çalışmalar her geçen gün artmaktadır. Hukukun sosyal bilimler araştırmalarına konu edilmesi, dinamik toplum düzeninde suç ile ilgili dinamik bilgilere ulaşılmasını, toplumdaki aksaklıkların tespit edilmesini ve bu aksaklıkların çözümü için gerekli olan düzenlemelerin yapılabilmesi için zemin sağlayacaktır. Bu bağlamda suç, salt hukukçuları değil, sosyolog, kriminolog gibi diğer sosyal bilimcileri de ilgilendiren bir olgudur.

Suç teorisine ceza hukuku yönünden bakıldığında, *işlendiğinin subuta varması halinde bir cezanın ve çeşitli tedbirlerin uygulanmasını gerektiren suç* için oldukça fazla tanımlama yapılmış ve yazarlar tek bir tanımda birleşememişlerdir. Zaten suçun tüm unsurlarını tek bir tanımda işaret etmek ve suçu tümüyle anlatmak mümkün değildir. Buna rağmen suç için yapılan “hukuki nizamın netice olarak ceza gerektirdiği (terettüp) eylemdir (fiil)” tanımı kısa, basit ve kapsayıcı olması bakımından diğer tanımlardan üstündür (Alacakaptan, 1975).

“Suç” sözcüğünü, hukukçu olmayan kişiler teknik anlam yerine geniş anlamda kullanmaktadır. *Yalnızca ahlaka, töreye, hukuka aykırılıklar değil toplumdaki bütün “sapma”lar suç (ya da kabahat) olarak adlandırılmaktadır.* Ancak hukukçu olan ve olmayan kişilerin ortak tanımı; toplum içerisindeki kişilerin doğruluk, dürüstlük, haklılık, haksızlık gibi etik/ahlaksal boyutlu kavramlara göre değerlendirme yaptığı ve böylece haksızlığın ya da suçun, ahlaka aykırı ve alışılmıştan sapma olduğudur. (Selçuk, 2014).

Farklı disiplinlerin oluşturduğu farklı suç tanımlamaları bulunmaktadır. Ancak bu tanımlamalar, suçun evrensel ve toplumlara özgü bir olgu olması, hukuki bir yapıya sahip olması ve neden sonuç ilişkisi bakımından sosyolojik olması hususlarında birleşmişlerdir. Suç, her çağda ve her toplumda varolması nedeniyle evrenseldir. Davranışlar kendiliğinden suç olmaz, davranışı suç yapan toplumun o davranışı suç olarak kabul etmesidir. Bu bağlamda suç toplumsaldır. Toplumun olduğu yerde hukuk, hukukun olduğu yerde ise suç vardır. Örneğin, adada tek başına yaşayan roman kahramanı Robinson Crusoe’un uyması gereken bir kural bulunmaması nedeniyle suç işleme olasılığı yoktur (Güçlü - Akbaş, 2016; Özkan, 1998).

Suç analizi ile ilgili literatürde oldukça fazla uygulamalı çalışma bulunmaktadır. Bunlardan bazıları özetlenmiştir. Nath (2007), suçun jeo-mekansal yapısını orataya koymak, terörle mücadelede yardımcı olmak ve kolluk kuvvetlerine yol göstermek amacıyla k-means tekniğini kullanmıştır. Günay Atbaş (2008), “Kümeleme Analizinde Küme Sayısının Belirlenmesi Üzerine Bir Çalışma” başlıklı çalışmada, 11 başlık altında toplanan değişkenleri kullanarak kümeleme analizi uygulamış, farklı küme sayıları için illeri incelemiştir. Uittenbogaard ve Ceccato (2012), Stokholm şehrine ilişkin suç verilerini kullanarak, suçların önlenmesi ve kent güvenliğinin iyileştirilmesi açısından öneme sahip olabilecek bilgiler elde edilmiş, şehrin güvenli ve güvensiz olan bölgeleri ile hangi suçların hangi saatlerde işlendiği tespit edilmiştir. Agarwal ve ark. (2013), İngiltere ve Galler’de toplanan 1990-2012 yıllarına ait suç verisine k-means tekniği uygulanarak suç analizi yapılmıştır. Gera ve Vohra (2014), Delhi şehrine ait suç verisi kullanılarak kümeleme analizi kapsamında k-means tekniği uygulanmıştır. Öncelikle şehirde en çok görülen suçlar tespit edilmiş ve başlıklar altında toplanmıştır. Sonrasında ise şehrin hangi bölgelerinde hangi suçların yoğunlaşmakta olduğu araştırılmıştır. Afacan ve Bildirici (2017), işsizlik ve suç ilişkisi bakımından 81 ili incelemiş, yapılan kümeleme analizi sonucunda benzer özelliklere sahip 5 küme elde etmiştir. Altunkaynak ve ark. (2018), ikili kümelemede Bimax algoritması kullanarak suç türlerine göre illeri incelemiştir. Suç bölgeleri ile sosyo-ekonomik değişkenler arasındaki ilişki ortaya konmuştur. Sevil Ay (2018), Türk Ceza Kanunu’nda belirtilen dört temel suç (Uluslararası Suçlar, Kişilere Karşı Suçlar, Topluma Karşı Suçlar, Millete Karşı Suçlar) bağlamında 81 ili kümelemiştir. Farklı kümeleme teknikleri kullanılan çalışmada küme sayısı 7 olarak belirlenmiştir. Özgür Güler ve Keskin (2019), 2017 yılına ait 23 suç göstergesini kullanarak illeri kümeleme analizi ile alt gruplara ayırmıştır. Analiz sonucunda dört homojen küme elde edilmiştir. Arslan ve ark. (2018), kaçakçılıkta yakalanan malzeme türlerine göre suçlulukları, ikili kümeleme yöntemi kullanarak incelemiş ve suçlu profillerine ilişkin bilgileri ortaya koymuştur. Çil ve ark. (2019), Boston’daki Suçlar veri seti üzerinde farklı algoritmalar kullanmış, hangi algoritmanın daha üstün olduğunu ortaya koymuştur. Şen ve Yazıcı (2019), kümeleme analizini kullanarak cinsel suçlardaki suç oranının coğrafya ve sosyo-ekonomik koşullardan etkilenmediğini tespit etmiştir. Atmaja (2019), Yogyakarta şehrine ait suç verilerini kullanarak kümeleme analizi uygulamıştır. Suçu önleme çabasına katkı sağlama amacı güden çalışma, suç çeşitlerini kümeleyerek, 4 suç içeren yüksek suç seviyesi, 6 farklı suç içeren orta suç seviyesi ve 8 suç içeren düşük suç seviyesi olmak üzere üç grup elde etmiştir.

Günümüzde suç oranlarının hızla değiştiği düşünüldüğünde, suç verilerinin analiz edilmesi, kanun koyuculara suç eğilimleri hakkında bilgi vermek açısından önem arz etmektedir. Suç verileri kullanılarak yapılacak farklı analizler, suçun mekânsal dağılımının ve suç oranlarının tahmininin elde edilmesine vesile olacaktır. Bu çalışmanın temel amacı, toplumlar için çözüm bekleyen önemli bir

sosyal problem haline gelmiş suç olgusunun belirlenen suç türleri (öldürme, yaralama, cinsel suçlar, kişiyi hürriyetinden yoksun kılma, hakaret, hırsızlık, yağma, dolandırıcılık, uyuşturucu ve uyarıcı madde madde ile ilgili suçlar, kötü muamele, zimmet, rüşvet, kaçakçılık, trafik suçları, orman suçları, ateşli silahlar ve bıçaklar ile ilgili suçlar, icra-iflas kanununa muhalefet, askeri ceza kanununa muhalefet, tehdit, mala zarar verme, görevi yaptırmamak için direnme, ailenin korunması tedbirine aykırılık) ve bu suç türlerinin işlendiği illere göre incelenmesidir. Kümeleme analizi ile illerin homojen alt gruplara ayrıştırılması amaçlanmaktadır.

MATERYAL VE YÖNTEM

Materyal

İnsanoğlu tarih boyunca kendisini çevreleyen canlı ve cansız nesneleri sınıflandırmaya çalışmıştır. Nesneleri toplu kategoriler halinde sınıflandırmak, onları adlandırmak için bir ön koşuldur. Kümeleme, birimleri aynı gruba koymak için kâfi derecede benzer olduklarını kabul etmek ve birimlerden oluşan grupların arasındaki ayrımları tanımlayabilmek anlamına gelmektedir (Legendre - Legendre, 2012).

Kümeleme analizi, uzun bir geçmişe sahip yöntemlerden meydana gelmektedir. İlk önermeler 1900'lü yılların başında antropologlar tarafından yapılmış, analize ait yöntemler daha sonra psikoloji alanında kullanılmıştır. Popülerliği 1960'lı yıllarda artan kümeleme analizi zaman içerisinde geliştirilen algoritmalar sayesinde başta tıp (özellikle psikiyatri), biyoloji, sosyal bilimler, coğrafya/jeoloji ve mühendislik alanlarında yaygın biçimde kullanılan kullanılmıştır. Başlangıçta doğa bilimleri için geliştirilen yöntem sosyal bilimlerde daha çok rağbet görmüştür (Blashfield - Aldenderfer, 1988; Yim - Ramdeen, 2015).

Esas amacı büyük heterojen örneklerden, homojen özelliklere sahip gruplar oluşturmak olan kümeleme analizi, veri azaltma amacına hizmet ederek, araştırmacıların işe yarar, özet ve objektif bilgileri elde etmesine yardımcı olur. Buna ek olarak, aykırı değerlerin tespiti, hipotez oluşturma ve hipotezlerinin testi gibi kanıtlayıcı amaçlarla da kullanılmaktadır (Yılmaz - Doğu, 2021).

Günümüzde akla gelebilecek her alanda kullanılan kümeleme analizinde doğru sonuçlar elde edebilmek için, araştırmacıların oluşturulacak grupların (kümelerin) neden varolduğunu ve hangi birimlerin (değişkenlerin) gruplara neden gireceğine ilişkin kavramsal bir temele sahip olması gerekmektedir. Kümeleme analizi, verinin yapısı ne olursa olsun her zaman kümeleme işlemini gerçekleştireceğinden, araştırmacıların birimler arasındaki yapıya ilişkin sunduğu varsayımlar ve elde edilen küme yapıları kavramsal alt yapı olmadan anlamsız olacaktır. Analiz, kullanılan değişken, küme sayısı ve benzerlik ölçüsüne tamamen bağımlı olduğu için söz konusu hususların seçiminde hassas davranılmalıdır (Hair vd., 2014, 419-420).

Kümeleme analizinin aşamaları aşağıdaki gibi sıralanabilir (Hair vd., 2014; Hardle - Simar, 2015; Tatlıdil, 2002):

1. İlk aşama, araştırma probleminin tespiti ve gerekli verilerin temin edilmesidir.
2. İkinci aşama, benzerliğin ölçülmesi ve bunun ne şekilde yapılacağına karar verilmesidir. Söz konusu ölçüm farklı şekillerde yapılabilir. Kümeleme analizi uygulamalarında en çok kullanılan yöntemler, korelasyon ölçüleri (correlational measures), uzaklık ölçüleri (distance measures) ve ilişkilendirme ölçüleridir (association measure). Bunlardan uzaklık ölçüleri uygulamalarda daha fazla kullanılmakta ve genellikle Minkowski uzaklığı, Manhattan City-Block uzaklığı, Öklid (Euclidean) uzaklığı, Mahalanobis uzaklığı, Hotelling T^2 uzaklığı ve Canberra uzaklığının tercih edildiği bilinmektedir.
3. Üçüncü aşamada kümeleme algoritması seçilir. (Hiyerarşik yöntem, hiyerarşik olmayan yöntem veya her ikisinin birlikte kullanımı)

4. Seçilen algoritmaya göre küme sayısına karar verilir. Farklı küme sayıları farklı sonuçlar elde edilmesine neden olacağından, bu aşamada dikkatli olunmalıdır.

5. Elde edilen küme yapıları yorumlanarak anlamlılıkları araştırılır.

Kümeleme analizinde, hiyerarşik yöntemler grubu hiyerarşik olmayan yöntemlere göre daha fazla tercih edilmektedir. Tek bağlantı yöntemi, tam bağlantı yöntemi, ortalama bağlantı yöntemi ve Ward's yöntemi en çok kullanılanlar arasındadır. Yöntemlerin birbirine üstünlükleri incelendiğinde, Ward's tekniğinin araştırmacılar tarafından daha çok kullanıldığı görülmektedir. Hiyerarşik olmayan kümeleme yöntemlerinden en çok tercih edileni ise k-ortalamlar yöntemidir (Kalaycı, 2017; Legendre - Legendre, 2012; Tatlıdil, 2002).

Uygulama aşamasında, hiyerarşik ve hiyerarşik olmayan yöntemlerin birlikte kullanımı, bir yöntemin zayıflığının diğer yöntemin avantajı tarafından giderilmesi bakımından önerilmektedir. İlk olarak, uygun küme sayısı hiyerarşik kümeleme yöntemlerinden biriyle belirlenir. Sonrasında ise hiyerarşik olmayan yöntem kullanılır (Hair vd., 2014; Kalaycı, 2017).

Yöntem

Verilerin Elde Edilmesi

Ülkemizde araştırmacılar, suç ile ilgili istatistiki veriler ulaşmakta zorluk yaşamaktadır. Adalet Bakanlığı Ceza ve Tevkifevleri Genel Müdürlüğü, Türkiye Cumhuriyeti İçişleri Bakanlığı, Türkiye Statistik Kurumu (TÜİK), Adalet Bakanlığı Adli Sicil ve İstatistik Genel Müdürlüğü ve Hakimler Savcılar Yüksek Kurulu'ndan istatistiki veri temin edilebilmektedir. Ancak TÜİK dışında, diğer kurumlardan veri talep edebilmek için CİMER (Cumhurbaşkanlığı İletişim Merkezi) vasıtasıyla olmakta ve kurumların cevabı gecikebilmekte üstelik beklenen nitelikleri taşımayabilmektedir. Hem suç türlerini ayrıntılı sunması hem de erişim kolaylığı sebebiyle araştırmada kullanılan veriler TÜİK'in internet sayfasından alınmıştır.

Çalışmada kullanılan veriler 2019 yılına ait olup, Türkiye İstatistik Kurumu (TÜİK) Adalet İstatistikleri bölümünden temin edilmiştir. Bu veriler, 2019 yılında İBBS 3. düzeyde, daimî ikametgâh ve suç türüne göre ceza infaz kurumuna giren hükümlülere kapsamaktadır. Çalışma kapsamına alınan suç türleri Çizelge 1'de verilmiştir. Uyuşturucu ve uyarıcı madde ile ilgili suçlar tek değişken olarak alınmıştır. Söz konusu hükümlü sayıları illerin nüfusuna oranlanmıştır.

Çizelge 1. Suç türleri

Suç Türleri			
X_1 : Öldürme	X_7 : Yağma	X_{14} : Kaçakçılık	X_{19} : Askeri ceza
X_2 : Yaralama	X_8 : Dolandırıcılık	X_{15} : Trafik suçları	kanununa muhalefet
X_3 : Cinsel suçlar	X_9 : Uyuşturucu veya	X_{16} : Orman suçları	X_{20} : Tehdit
X_4 : Kişiyi	uyarıcı madde ile ilgili	X_{17} : Ateşli silahlar	X_{21} : Mala zarar verme
hürriyetinden yoksun	suçlar	ve bıçaklar ile ilgili	X_{22} : Görevi yaptırmamak
kılma	X_{10} : Sahtecilik	suçlar	için direnme
X_5 : Hakaret	X_{11} : Kötü muamele	X_{18} : İcra-iflas	X_{23} : Ailenin korunması
X_6 : Hırsızlık	X_{12} : Zimmet	kanununa muhalefet	tedbirine aykırılık
	X_{13} : Rüşvet		

İstatistiksel Analizler

Uygulama kapsamında, uzaklık matrisinin belirlenmesinde öklid uzaklığı (euclidian distance), hiyerarşik kümeleme yöntemlerinden Ward's tekniği, hiyerarşik olmayan kümeleme yöntemlerinden de k-ortalama tekniği kullanılmıştır. Analizler IBM SPSS (Statistical Package for the Social Sciences) v.25 ile yapılmıştır.

BULGULAR

Çalışmanın uygulama aşamasında hiyerarşik ve hiyerarşik olmayan iki teknik uygulanmıştır. İlk olarak hiyerarşik kümeleme yöntemlerinden Ward's tekniği, sonrasında ise hiyerarşik olmayan kümeleme yöntemlerinden de k-ortalama tekniği kullanılmıştır.

Hiyerarşik olmayan kümeleme yöntemlerinde süreç, araştırmacının önsel bilgisi ve deneyimi doğrultusunda belirlediği küme sayısı ışığında ilerlemektedir. Hiyerarşik olmayan yöntemlerden k-ortalama yöntem, birimleri araştırmacının belirlediği küme sayısına ayırmakta ve küme ayırt ediciliği ile bazı sayısal hedeflere ulaşana kadar iterasyona devam eden algoritmalar bütünüdür (Hair vd., 2014, 417). Hiyerarşik kümeleme süreci, her bir birimin tek bir küme kabul edilmesiyle başlamakta, sonraki adımlarda ise en benzer kümeler bir araya gelerek kümeleri oluşturmaktadır. Süreç, tüm birimlerin tek bir kümede birleşmesine kadar devam etmektedir (Hair vd., 2014, 416). Sürecin tamamı, dendrogramın (ağaç grafiği) soldan sağa okunması ya da yığılma tablosu (agglomeration schedule) incelenerek takip edilebilmektedir. Söz konusu inceleme yapıldığında, uygun küme sayısı 4 olarak belirlenmiştir. Analiz sonucunda elde edilen kümeler Çizelge 2.'de verilmiştir.

Çizelge 2. İllerin kümelenmesi

		Küme no			
		1	2	3	4
İller	İstanbul		Kocaeli (İzmit), Trabzon, Erzurum, Malatya, Aksaray, Kütahya, Tokat, Nevşehir, Çankırı, Kahramanmaraş, Gümüşhane, Erzincan, Kars, Bayburt, Iğdır, Rize, Ardahan, Adıyaman, Bingöl, Şanlıurfa, Siirt, Tunceli, Yalova, Artvin, Muş, Batman, Hakkari, Mardin, Bitlis, Ağrı, Van, Şırnak, Karaman, Sinop, Kırıkkale, Osmaniye, Kırklareli, Edirne, Niğde, Bilecik, Manisa	Düzce, Çanakkale, Kayseri, Uşak, Kastamonu, Kırşehir, Zonguldak, Ordu, Sivas, Bartın, Sakarya (Adapazarı), Bolu, Çorum, Giresun, Antalya, Tekirdağ, Yozgat, Eskişehir, Karabük, Elazığ, Samsun, Kilis, Isparta, Amasya, Denizli, Afyonkarahisar, Balıkesir, Aydın, Muğla, Burdur	İzmir, Bursa, Ankara, Adana, Gaziantep, İçel (Mersin), Konya, Diyarbakır, Hatay (Antakya)

Kümeleme analizi uygulamalarında dikkat edilmesi gereken en önemli noktalar, önemli değişkenlerin, kümeleme yönteminin, uzaklık ölçüsünün ve küme sayısının belirlenmesi olarak sıralanabilmektedir. Yöntem her halukarda bir kümeleme işlemi yapacağından, oluşan kümelerin doğru değerlendirmesi de önem arz etmektedir. Bu itibarla, İstanbul ilinin nüfus büyüklüğü ve tüm suç değişkenlerine ait yüksek derlere sahip olması nedeniyle tek başına bir küme oluşturması beklenen bir

durumdur. İstanbul ilinde en çok, hırsızlık (5888), yaralama (4463) ve uyuşturucu veya uyarıcı madde ile ilgili suçlar (4805) işlenmiştir.

Küme 2’de 41 ilin olduğu görülmektedir. Bu iller; Kocaeli (İzmit), Trabzon, Erzurum, Malatya, Aksaray, Kütahya, Tokat, Nevşehir, Çankırı, Kahramanmaraş, Gümüşhane, Erzincan, Kars, Bayburt, Iğdır, Rize, Ardahan, Adıyaman, Bingöl, Şanlıurfa, Siirt, Tunceli, Yalova, Artvin, Muş, Batman, Hakkâri, Mardin, Bitlis, Ağrı, Van, Şırnak, Karaman, Sinop, Kırıkkale, Osmaniye, Kırklareli, Edirne, Niğde, Bilecik ve Manisa’dır. Bu illerde en çok işlenen suçlar sırasıyla, hırsızlık, yaralama ve uyuşturucu veya uyarıcı madde ile ilgili suçlar olduğu tespit edilmiştir. Söz konusu suçların en çok işlendiği iller, Manisa, Kocaeli (İzmit) ve Şanlıurfa’dır.

Küme 3’te 30 il bulunmaktadır. Bu iller; Düzcce, Çanakkale, Kayseri, Uşak, Kastamonu, Kırşehir, Zonguldak, Ordu, Sivas, Bartın, Sakarya (Adapazarı), Bolu, Çorum, Giresun, Antalya, Tekirdağ, Yozgat, Eskişehir, Karabük, Elâzığ, Samsun, Kilis, Isparta, Amasya, Denizli, Afyonkarahisar, Balıkesir, Aydın, Muğla ve Burdur’dur. Nüfus bakımından daha kalabalık illerin bulunduğu bu kümedeki illerde çoğunlukla, hırsızlık, yaralama ve uyuşturucu veya uyarıcı madde ile ilgili suçların işlenmiştir. Antalya, Balıkesir ve Aydın bahsi geçen suçların en çok işlendiği illerdir.

Küme 4’te 9 il bulunmaktadır. Bu iller; İzmir, Bursa, Ankara, Adana, Gaziantep, İçel (Mersin), Konya, Diyarbakır ve Hatay (Antakya)’dır. Türkiye’nin nüfus bakımından en büyük illerinin bu kümede toplandığı görülmektedir. Diğer il kümelerinde olduğu gibi en çok işlenen suçlar, yaralama, hırsızlık ve uyuşturucu veya uyarıcı madde ile ilgili suçlar olduğu görülmektedir.

TARTIŞMA VE SONUÇ

Türkiye’de 1990’lı yıllardan itibaren özellikle iç göç ve neoliberal politikalar nedeniyle suç oranlarının arttığı bilinmektedir. Buna ek olarak kadın suç oranı ve kentlerdeki suç olgusu da istikrarlı biçimde artmaktadır. Genel afların ardından tutuklu sayısında 1974, 1991 ve 2000 yıllarında düşüş görülmüş, ardından bu düşüş yerini artışa bırakmıştır. 2016 yılından itibaren ise ceza evine giren hükümlü sayısında fark edilir bir artış gözlenmiştir (Demirel, 2017). Ceza evine giren hükümlü sayısındaki artış ve azalışlarda kuşkusuz ki ülkede yaşanan sosyal ve ekonomik olaylar ile uygulanan adalet politikalarının payı bulunmaktadır.

Suç olgusunu araştıran çalışmalarda çok değişkenli istatistiksel analiz tekniklerinden biri olan kümeleme analizi sıklıkla kullanılmaktadır. Bu çalışmada da kümeleme analizi tercih edilmiş, Türkiye İstatistik Kurumu’ndan elde edilen İBBS 3. Düzeyde, daimî ikametgâh ve suç türüne göre ceza infaz kurumuna giren hükümlüler verisi ışığında 23 değişken ile 81 il kümelendirilmiştir. Hiyerarşik ve hiyerarşik olmayan iki yöntem kombine edilerek kullanılmış ve uygun küme sayısı 4 olarak belirlenmiştir.

İstanbul ilinin tek başına bir küme oluşturduğu, nüfusu fazla olan illerin yine bir araya gelerek bir küme oluşturduğu gözlenmiştir. Kümelerin ortak özelliği, yaralama, hırsızlık ve uyuşturucu veya uyarıcı madde ile ilgili suçların en fazla işlenmiş olmasıdır.

Daha önce de belirtildiği gibi suç ile ilgili istatistiki verilerin elde edilmesinde çeşitli güçlükler yaşanmaktadır. Buna ek olarak, hükümlü sayılarının kadın-erkek-toplam şeklinde verilmesi yerine yaş gruplarına göre de verilmesi ve bu bilgilere çocuk hükümlülerin de eklenmesiyle illere göre hangi yaş gruplarının hangi illerde hangi suçların işlendiğinin tespiti gibi daha hassas analizler yapılabileceği düşünülmektedir.

Kaynaklar

Agarwal, J., Nagpal, R. ve Sehgal, R., 2013. Crime Analysis using K-Means Clustering: International Journal of Computer Applications, 83(4).

- Alacakaptan, U., 1975. Suçun Unsurları. Sevinç Matbaası, Ankara-Türkiye.
- Altunkaynak, B., Örlcü, H. H. ve Arslan, R., 2018. Şehirlerin Suç Türlerine Göre İkili Kümeleme Yöntemi ile Gruplandırılması: Türkiye Örneği, Süleyman Demirel Üniversitesi Fen Bilimleri Enstitüsü Dergisi, 22 (Özel Sayı): 110-120.
- Arslan, R., Örlcü, H. H. ve Altunkaynak, B., 2018. Kaçakçılıkta Yakalanan Malzeme Türlerine Göre Suçluların Kümelenmesi: İkili Kümeleme Yöntemi. 19. Uluslararası Ekonometri, Yöneyim Araştırması ve İstatistik Sempozyumu. 2018 EYİ Özel Sayısı, 883-896, İstanbul, Türkiye.
- Afacan, N. ve Bildirici, İ. Ö., 2017. Tematik Kartografyada Kümeleme Analizi. TMMOB Harita ve Kadastro Mühendisleri Odası, 16. Türkiye Harita Bilimsel ve Teknik Kurultayı. 3-6 Mayıs 2017, Ankara, Türkiye.
- Atbaş, A. C. G., 2008. Kümeleme Analizinde Küme Sayısının Belirlenmesi Üzerine Bir Çalışma, Yüksek Lisans Tezi, Ankara Üniversitesi, Fen Bilimleri Enstitüsü, Ankara, Türkiye.
- Atmaja, E. H. S., 2019. Implementation of k-Medoids Clustering Algorithm to Cluster Crime Patterns in Yogyakarta. International Journal of Applied Sciences and Smart Technologies, 1(1), 33-44.
- Ay, S., 2018. Türkiye'deki Ceza Davalarının İstatistiksel Analizi, İstanbul Ticaret Üniversitesi Sosyal Bilimler Dergisi, 33(1).
- Blashfield, R. K., Aldenderfer, M. S., 1988. The Methods and Problems of Cluster Analysis. N. J.R., & C. R.B. (Eds.), Springer. Handbook of Multivariate Experimental Psychology: 447-473. Boston-ABD.
- Çil, İ., Çakar, S. G., Sarı, N. ve Eydemir, O., 2019. İkili Kümeleme Yaklaşımıyla Suç Bölgelerinin Tespiti ve İkili Kümeleme Yöntemlerinin Karşılaştırılması. Sakarya University Journal of Computer and Information Science, 2(3).
- Demirel, P., 2017. Yetişkin Suçluluğuna Neden Olan Sosyoekonomik Faktörler, Doktora Tezi, Hacettepe Üniversitesi, Sosyal Bilimler Enstitüsü, Ankara, Türkiye.
- Gera, P. ve Vohra, R., 2014. City Crime Profiling Using Cluster Analysis: International Journal of Computer Science and Information Technologies, 5(4), 5145-5148.
- Güçlü, İ., ve Akbaş, H., 2016. Suç Sosyolojisi: Kavram, Teori, Uygulama. Seçkin, Ankara-Türkiye.
- Güler, E. Ö. ve Keskin, D., 2019. Türkiye'deki Suç Türlerinin ve Bu Suçların İşlendiği İllere Göre Ceza İnfaz Kurumuna Giren Hükümlülerin Kümeleme Analizi ile İncelenmesi. II. International Conference on Empirical Economics and Social Science (ICEESS' 19), 20-22 Haziran 2019, Bandırma, Türkiye.
- Hair, J. F., Black, W. C., Babin, B. J., ve Anderson, R. E., 2014. Multivariate Data Analysis. Pearson Education Limited, United States of America.
- Hardle, K. W., ve Simar, L., 2015. Applied Multivariate Statistical Analysis. Springer-Verlag, Berlin-Almanya.
- Kalaycı, Ş., 2017. SPSS Uygulamalı Çok Değişkenli İstatistik Teknikleri. Dinamik Akademi, Ankara-Türkiye.
- Legendre, P., ve Legendre, L., 2012. Cluster Analysis. P. Legendre, ve L. Legendre (Eds.), Developments in Environmental Modelling Book Series: Numerical Ecology. Canada.
- Nath, S.V., 2007. Crime Data Mining: Advances and Innovations in Systems, Computing Sciences and Software Engineering, 405-409.
- Özkan, Ş., 1998. Türkiye'de Yetişkin Suçluluğu Üzerine Cinsiyetler Arası Bir Karşılaştırma, Doktora Tezi. Ege Üniversitesi, Sosyal Bilimler Enstitüsü, İzmir, Türkiye.
- Selçuk, S., 2014. Prof. Dr. Feridun Yenisey'e armağan, 85 - 106. Beta Yayınevi. Ankara-Türkiye.
- Şen, H. ve Yazıcı, K., 2017. Türkiye'deki İllerin Cinsel Suçlar Açısından İncelenmesi, The Journal of Operations Research, Statistics, Econometrics and Management Information Systems, 5(2).
- Tatlidil, H., 2002. Uygulamalı Çok Değişkenli İstatistiksel Analiz. Akademi Matbaası, Ankara- Türkiye.
- Uittenbogaard, A. ve Ceccato, V., 2012. Space-time Clusters of Crime in Stockholm, Sweden: Article in Review of European Studies, 4(5).
- Yılmaz, C. Y., ve Doğu, G. A., 2021. Kümeleme Analizi ile İllerin Okul Sporları Faaliyetleri Bakımından İncelenmesi: Iğdır Üniversitesi Sosyal Bilimler Dergisi, 10(26): 182-205.
- Yim, O., ve Ramdeen, K. T., 2015. Hierarchical Cluster Analysis: Comparison of Three Linkage Measures and Application to Psychological Data: The Quantitative Methods for Psychology, 11(1): 8-21.

Classification and Regression Tree as Data Mining Algorithm to Predict Body Weight of Boer Goats

Madumetja Cyril MATHAPO¹, Thobela Louis TYASI^{1*}

¹School of Agricultural & Environmental Sciences, Department of Agricultural Economics and Animal Production, University of Limpopo, Private Bag X1106, Sovenga 0727, Limpopo, South Africa.

*Presenter author e-mail: louis.tyasi@ul.ac.za

Abstract

Classification and Regression Tree (CART) is a predictive algorithm method used to explain how the dependent variable can be predicted using independent variables (numerical and characters). The aim of the study was to identify importance of biometric traits on estimation of body weight of Boer goats. A total of seventy-one Boer goats between the age of one year and two years old were used. CART was used for data analysis. CART findings displayed that sex played a crucial role on body weight of Boer goats. CART model established in this study could be used by breeders to advise Boer goats' farmers who cannot afford to purchase weighing scale on which biometric traits they can use to select their animals in order to improve their herd. However, further studies need to be done to validate the use of CART in prediction of body weight from biometric traits of Boer goats using large sample size, different area or other goat breeds.

Key words: Body length, Ear length, Heart girth, Head width, Rump width

INTRODUCTION

Classification and regression tree have been used for estimation of body weight using morphological traits in Balochi's sheep (Huma and Iqbal, 2019), and in Beetal goats of Pakistan (Eyduan *et al.*, 2017). Classification and regression trees are a modern technique used to come up with the traits which have a statistical influence on the dependent variable without any assumption (Dariusz, 2009). However, there is a scarce information on prediction of body weight from biometric traits Boer goats using classification and regression tree. The objective of the current study was to establish a model to predict body weight from biometric traits using classification and regression tree. This study will aid poor rural farmers on which trait they can employ to improve and estimate body weight of their Boer goats, so that they can manage their animals well.

MATERIALS AND METHODS

Data collection

A total seventy-one Boer goats, fourteen males and fifty-seven females aged between one and two years old were used for data collection.

Biometric traits were measured using tailor measuring tape calibrated in centimeters (cm) while body weight was taken using a weighing scale calibrated in kilograms (kg). All measured traits were taken as explained in Lukuyu *et al.*, (2016). Figure below shows the measured traits.

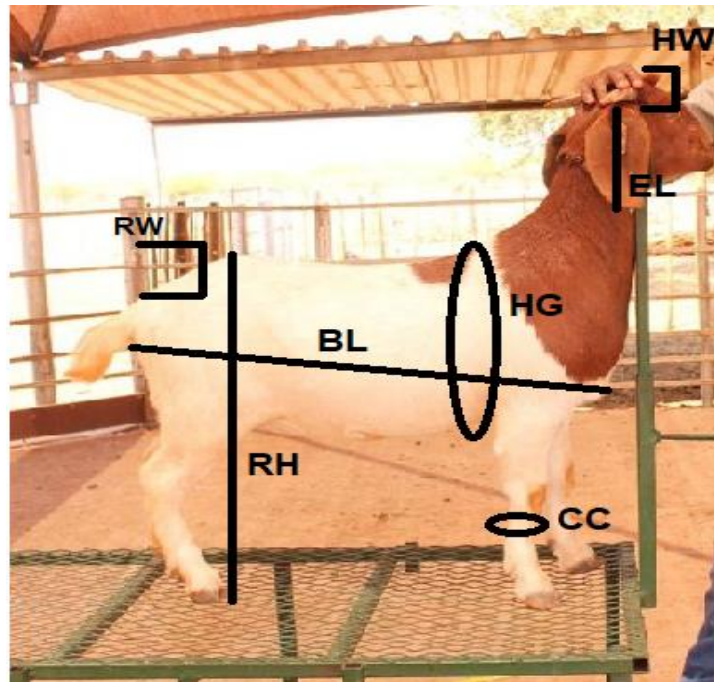


Figure 1. Boer goat showing biometric traits measured in the study. Source: Mathapo and Tyasi, (2021). HW: Head width, EL: Ear length, HG: Heart girth, CC: Cannon circumferences, BL: Body length, RH: Rump height, RW: Rump width.

Statistical Analysis

Data analysis was performed using Statistical Package for Social Sciences version 26 (IBM SPSS, 2019) software. CART was used to establish a model to predict body weight from biometric traits as explained by Eydurán *et al.* (2017).

RESULTS AND DISCUSSION

Classification and regression tree: CART model in the current study revealed that sex was the best influencer of body weight of Boer goats (Figure 1). The other traits which played a crucial role on body weight were, heart girth and age. The first node, which is node 0 predicted an overall body weight of 57.21 kg. Node 0 was subdivided based on sex into node 1 female ($n = 57$, predicted weight 46.51 kg) and node 2 male ($n = 14$, predicted weight = 100.79 kg). Node 1 was split into node 3 and node 4 based on heart girth. Yearling Boer goats with heart girth less or equals to 87.50 cm ($n = 45$) and those with heart girth greater than 87.50 cm ($n = 12$). These had predicted body weights of 42.52 kg and 61.41 kg, respectively. Node 3 was also split into node 5 and node 6 based on age. Node 5 Yearling Boer goats with thirteen, fourteen and eighteen months were ($n = 10$) with predicted body weight of 45.37 kg while node 6 showed yearling Boer goats with one year were ($n = 14$) with predicted body weight of 36.24 kg. Nodes 2,4,5 and 6 were terminal nodes as no partition was observed for providing homogeneousness in these nodes. In all the terminal nodes, node 2 was found to be the best node as it recorded the highest predicted mean (100.76 kg) as compared to node 4 (61.42 kg), node 5 (45.37 kg) and node 6 (36.24 kg). This model showed that node 6 had the lowest variance $[(3.246)^2 = 10.54]$ and the variance of the root node or dependent variable (body weight) was $S^2y = (24.149)^2 = 583.174$. The unexplained variation in the body weight was $S^2e = \text{risk value} \div S^2y = 49.534 \div 583.174 = 0.085$ and the variation in the model was explained as $S^2y = 1 - S^2e = 1 - 0.085 = 0.915$, respectively. This study revealed that sex is the most explanatory variable with significant effect on body weight of Boer goats. Our findings were supported by Eydurán *et al.* (2017) who showed sex as an explanatory variable which played significant role on body weight in Indigenous Beetal Goat of Pakistan. Celik. (2019) showed different results from current

study which revealed body length as trait that had a significant effect on the body weight. Mohammad *et al.* (2012) also indicated different results from the current study where withers height was found to be the best trait to estimate body weight in yearling sheep. This might be due to different breeds and environmental factors where the animals were raised.

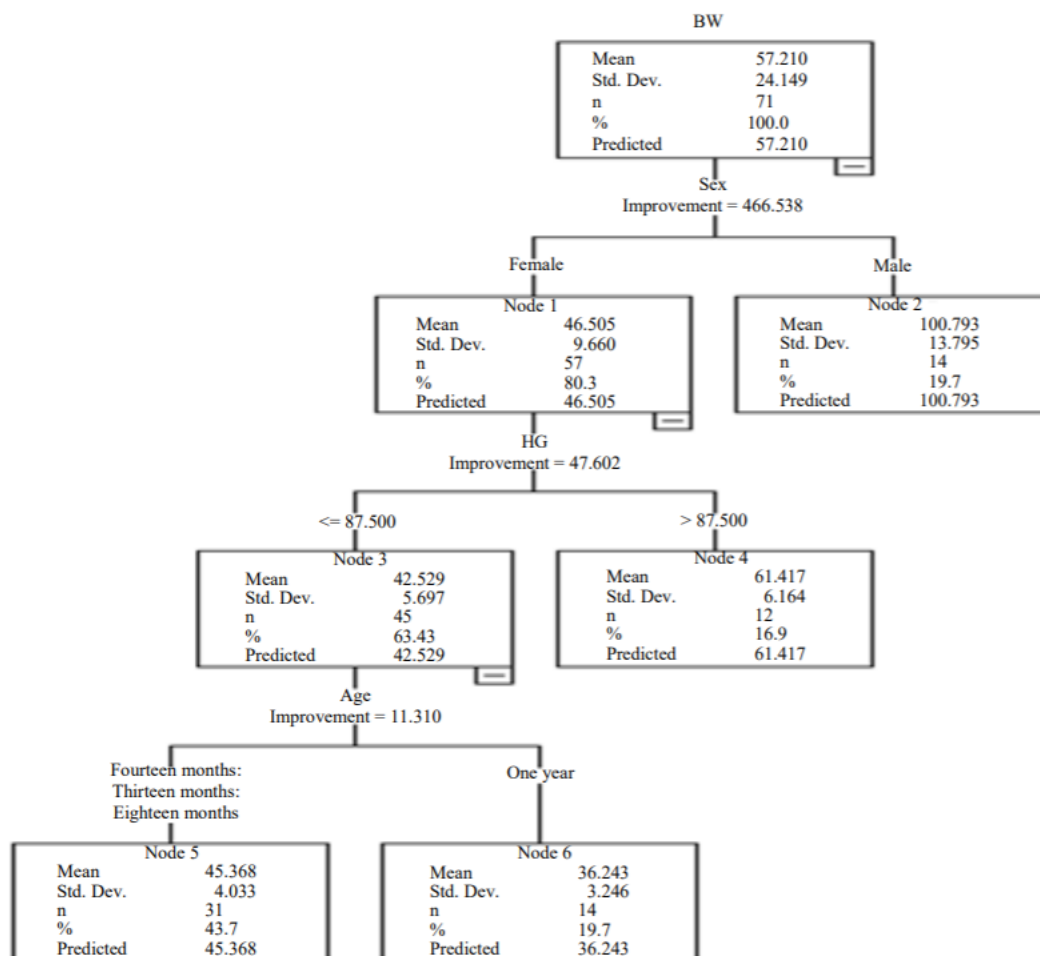


Figure 2. Showing classification and regression tree model

CONCLUSION

The current study suggests that sex is the one which play a significant effect on the body weight of Boer goats. Our findings might help poor farmers, researchers and the extension officers to know which trait to focus on when improving body weight of their goats and which explanatory variable to look when estimating body weight. More studies need to be conducted on developing a model to predict body weight of Boer goats, using more sample size or different animal.

References

- Celik, S., 2019. Comparing Predictive Performances of Tree Based Data Mining Algorithms and MARS Algorithm in the Prediction of Live Body Weight from Body Traits in Pakistan Goats. *Pakistan Journal of Zool.*, 51(4): 1447-1456. DOI: <http://dx.doi.org/10.17582/journal.pjz/2019.51.4.1447.1456>
- Dariusz, P., 2009. Using classification trees in statistical analysis of discrete sheep reproduction traits. *J Cent Eur Agric.*, 10 (3): 303-309
- Eyduran, E., Zaborski D., Waheed A., Celik S., Karadas K., Grzesiak, W., 2017. Comparison of the Predictive Capabilities of Several Data Mining Algorithms and Multiple Linear Regression in the Prediction of Body Weight by Means of Body Measurements in the Indigenous Beetal Goat of

- Pakistan. Pakistan Journal of Zoology., 49(1): 257-265. DOI: <http://dx.doi.org/10.17582/journal.pjz/2017.49.1.257.265>.
- Huma, Z. E., Iqbal, F., 2019. Predicting the body weight of Balochi sheep using a machine learning approach. Turkish Journal of Veterinary and Animal Sciences, 43(4), 500-506. <http://dx.doi.org/10.3906/vet-1812-23>
- Lukuyu, M. N., Gibson, J. P., Savage, D. B., Duncan, A. J., Mujibi, F. D. N., Okeyo, A.M., 2016. Use of body linear measurements to estimate live weight of crossbred dairy cattle in smallholder farms in Kenya. Springerplus., 5:63; <https://doi.org/10.1186/s40064-016-1698>.
- Mathapo, M.C., Tyasi, T.L., 2021. Prediction of body weight of yearling Boer goats from morphometric traits using classification and regression tree. American Journal of Animal and Veterinary Sciences, 16(2), 130-135.
- Mohammad, M. T., Rafeeq, M., Bajwa, M. A., Awan, M. A., Abbas, F., Waheed, A., Bukhari, F.A., Akhtar, P., 2012. Prediction of body weight from body measurements using regression tree (RT) method fir indigenous sheep breeds in Balochistan, Pakistan. The Journal of Animal & Plant Sciences., 22(1), 20-24.
- Statistical Package for Social Sciences., 2019. IBM SPSS statistics for windows, Version 26.0. IBM Corp., Armonk, NY

Acknowledgment

The authors would like to thank the Pieter Smith Boer goat farm workers for their support during data collection and financial contribution of University of Limpopo, Department of Agricultural Economics and Animal Production. We also thank a Boer goat stud, Mr Pieter Smith for allowing us to use his animals for data collection

Authors Contribution

Thobela Louis Tyasi: Designed the study and revised the manuscript and approved the final version.
Madumetja Cyril Mathapo: Collected data and analyzed the data and drafted the manuscript.

Ethics

Authors of this study followed the ethical principles of the University of Limpopo Animal Research Ethics Committee (AREC).

Conflict of Interest Declaration

The authors declare that there is no conflict of interest for this work.

Using QLS Method in Analysis of the Data Obtained from the Experimental FMD Vaccines

Erdoğan ASAR^{1*}

¹Turkish Statistical Institute, Statistical Consultancy, 06100, Ankara, Turkey

*Presenter author e-mail: erdo6718@gmail.com

Abstract

Foot and Mouth disease (FMD) is a highly contagious viral disease seen in cloven-hoofed animals. Vaccines are among the important measures against the disease. In addition, studies are carried out on more effective vaccine formulations against this disease. Quasi-Least Squares Regression (QLS) is one of the methods used in longitudinal data analysis. The main purpose of this study is to investigate the effect of a new FMD vaccine by introducing QLS method in data analysis for the veterinary vaccine studies. It was also aimed to compare the results obtained with the model established with QLS with the results of the model established with Generalized Estimating Equations (GEE).

Two water-in-oil-in-water (w/o/w) emulsion formulations of FMD vaccines were prepared with Montanide ISA 206 BVG (MISA206). Chitosan was incorporated in one of the formulations, and the other formulation was prepared without chitosan. The experimental FMD vaccines were administered to two groups of guinea pigs by subcutaneous route. Each group was consisted of ten animals. The blood samples were taken from the animals to measure the neutralizing antibody titer levels in serum ((NATL) by homologous virus neutralization test (VNT). The samples were taken seven times post vaccination (7, 14, 28, 60, 120, 180 and 220 days) under general anaesthesia. The results were statistically evaluated. Analyses on the longitudinal data were performed with QLS and GEE methods.

According to the analysis, the time effect on NATL was found to be significant in all study correlation structures in QLS and GEE ($p < 0.05$). However, when it comes to the effect of the group effect on NATL, it was found to be significant in QLS and GEE except for the regression model obtained using AR-1 working correlation structure ($p < 0.05$). Among the working correlation structures used in the analysis for the NATL response variable, it was also observed that the correlation value obtained with Markov working correlation structure was the highest correlation value.

The effect of the time variable on NATL leads to the conclusion that the new vaccine increases the antibody level. In addition, the finding that the group variable on NATL is significant except for AR-1 working correlation structure can be understood as the new vaccine is more effective. Contrary to AR-1 working correlation structure using measurement orders, according to the regression model established with Markov working correlation structure using real measurement times, it is an extremely important result that the group is effective on NATL and this structure is used in QLS.

Key words: Markov working correlation structure, Longitudinal data, Quasi-least squares regression

Determining Optimum Sample Size for Regression Tree and Automatic Linear Modeling

Serdar GENÇ^{1*}, Mehmet MENDEŞ²

¹Kirsehir Ahi Evran University - Faculty of Agriculture - Kirsehir, Turkey

²Canakkale Onsekiz Mart University - Faculty of Agriculture - Canakkale, Turkey

*Presenter author e-mail: serdargenc1983@gmail.com

Abstract

This study was carried out to determine optimum sample size for Regression Tree and Automatic Linear Modeling. A comprehensive Monte Carlo Simulation Study was designed for this purpose. Results of simulation study indicated that optimum sample size of Regression Tree and Automatic Linear Modeling methods was influenced by number of variables and structure of the variance-covariance matrix. Required sample size increased as the number of variables increased for both methods. For the RT, optimum sample size was determined as 10000 for P=5, 20000 for P=10 and more predictors. For ALM, on the other hand, for P=5 at least 2000 observations were needed while for P=10 and 15 at least 10000 observations were needed. As a result, the Regression Tree required much larger samples to make stable estimates when comparing to Automatic Linear Modeling.

Key words: Regression Tree, Automatic Linear Modeling, simulation, biased, data mining

INTRODUCTION

Sample size is one of the most important factors that affect the reliability and stability of estimates. Both Regression Tree (RT) and Automatic Linear Modeling (ALM) are multivariate techniques they require large samples to give reliable results. However, an excessive sample size is not a desirable situation as well since it will cause waste of resources. In this case, importance of studying with optimum sample size emerges (Breiman et. al., 1984; Cios et al., 2007; Lin et al., 2008; Kaur and Pulugurta 2008; Field 2013; Fan et al., 2014). In other word, it is needed to study with sufficient or optimum sample size to get the expected benefits from these methods. Therefore, determining proper or optimum sample size for both RT and ALM at the beginning of the experiment is an important issue. Optimum sample size may be defined as the minimum sample size which gives power value of 80 %. This simulation study aimed at determining optimum sample size for Regression Tree and Automatic Linear modeling techniques (Hill et al., 2004; Larose 2005; Lynda 2011; Kolaczyk 2013, Keskin and Mendes 2019). Thus, it is expected creating a simple and user-friendly route for the researchers who will use these techniques in analyzing their datasets.

MATERIAL AND METHODS

The material of this study is the random numbers generated from multivariate normal distribution by Monte Carlo simulation technique. RNMVN function of IMSL library of Microsoft FORTRAN Developer Studio were used in generating random numbers. R² and accuracy values were considered in determining optimum sample size (Anonim 1994, SPSS 2012; Stack Exchange 2011).

RESULTS FOR DETERMINING OPTIMUM SAMPLE SIZE

Results of the simulation study have been presented as graphically (Figure 1-6). Results of this simulation study showed that although optimum sample size for both methods was influenced by number of variables and structure of the variance-covariance matrix, optimum sample sizes were not the same

for RT and ALM. For the RT, optimum sample size was determined as 10000 for $P=5$, 20000 for $P=10$ and more predictors. As it can be noticed, as the number of variables increased, required sample size was also increased. Therefore, it is possible to suggest researchers study with at least sample sizes of 10000 when they have 5 predictors, 20000 when they have 10 and more predictors in order to get more reliable and stable estimates. As in the RT, the required sample size increases as the number of variables increased in the ALM. However, the ALM method required much less samples to give reliable results and make stable estimates when compared to the RT. For example, for $P=5$ at least 2000 observations were needed while for $P=10$ and 15 at least 10000 observations were needed.

Figure 1. Estimates versus actual values for RTM when $p=5$

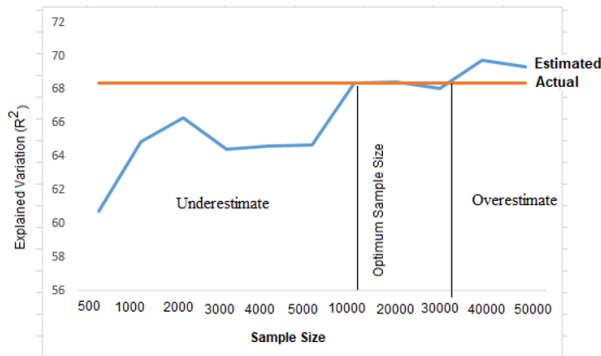


Figure 2. Estimates versus actual values for ALM when $p=5$

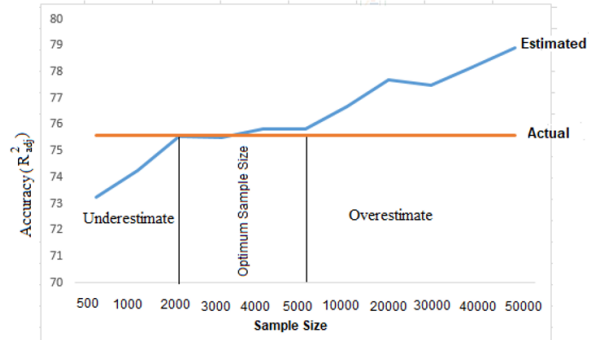


Figure 3. Estimates versus actual values for RTM when $p=10$

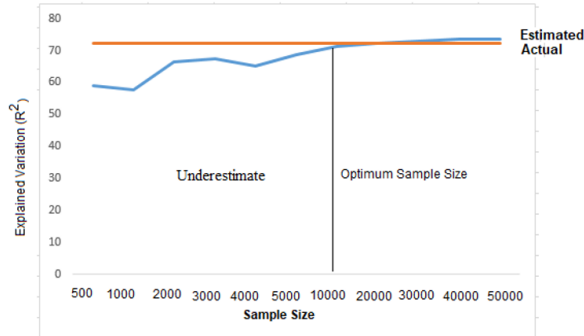


Figure 4. Estimates versus actual values for ALM when $p=10$

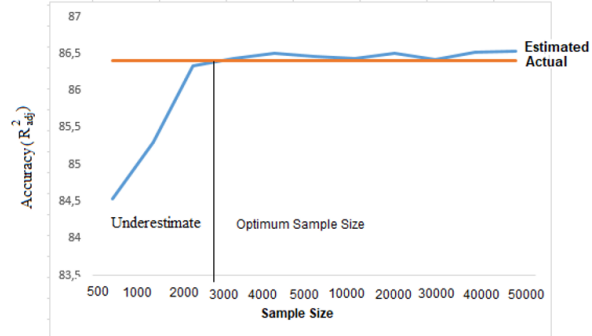


Figure 5. Estimates versus actual values for RTM when $p=15$

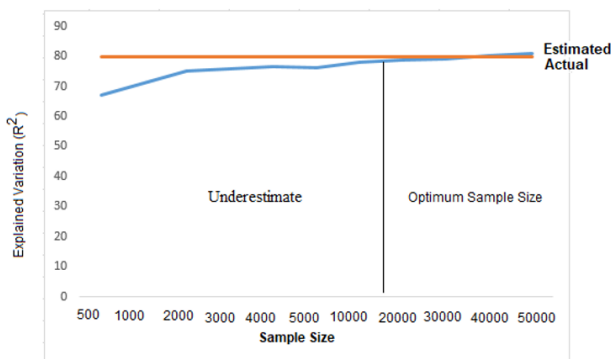
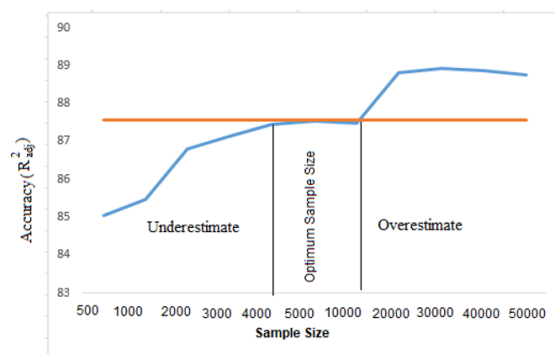


Figure 6. Estimates versus actual values for ALM when $p=15$



In this study, we tried to determine optimum sample size by using explained variation and accuracy values under different variable numbers and variance-covariance matrix structures. The minimum sample size, which gave the closest estimate to the reference value, was considered the optimal sample size. As a result, it was observed that sample sizes between 10000 and 30000 might be able to accept as optimum sample size for the RT when $P=5$ and 10. As it can be seen from figure 1, studying with the

sample size of smaller than 10000 caused underestimation problem while studying the sample sizes of larger than 30000 caused overfitting problem. Therefore, it can be concluded that in cases where the number of variables is between 5 and 10, it is necessary to work with samples of at least 10000 sizes in order to achieve reliable results or to make unbiased estimates. As it can be seen from figure 1, studying with the sample size of smaller than 10000 caused underestimation problem while studying the sample sizes of larger than 30000 caused overfitting problem. Therefore, the sample sizes between 10000 and 30000 might be able to accept as optimum sample size for the RT when $P=5$ and 10. Based on these findings, it can be concluded that in cases where the number of variables is between 5 and 10, it is necessary to work with samples of at least 10000 sizes in order to achieve reliable results or to make unbiased estimates under these conditions. For $P=15$ or more, it is possible to suggest that required minimum sample size should be at least 20000 in order to get reliable results or to make unbiased estimates. ALM method, on the other hand, has provided reliable results and stable estimates even when studying with much more smaller samples when compared to RT. As it can be seen from figure 2, the sample sizes between 2000 and 5000 might be accepted as optimum or sufficient sample sizes for the ALM method. As it can be noticed that the ALM method only produced unreliable results when sample size of smaller than 2000 (especially sample size was smaller than 1000). The ALM tended to give overestimate the actual value when sample size was larger than 5000. Domingos (1998) and Oates and Jensen (1997 and 1998) reported that decision tree based data mining tools were subject to over-fitting as the size of the data set increased. As it can be noticed from the results of this study (Figure 1-6), as the sample size increased overfitting problem occurred while underestimate problem occurred as the size of the data decreased. Therefore, the results of this study are compatible with the findings declared by Domingos (1998) and Oates and Jensen (1997). Morgan et al. (2003) reported that model accuracy improves at a decreasing rate with increasing sample size. When a power curve was fitted to accuracy estimates across various sample sizes, more than 80 percent of the time accuracy within 0.5 percent of the expected terminal (accuracy of a theoretical infinite sample) was achieved by the time the sample size reached 10,000 records. Although the idea that **"the larger the sample size is increased, the more reliable results are obtained"** in general, this idea is not particularly valid in practice. Because the time and cost allocated for a study must be taken into consideration. Therefore, studying with appropriate sample size will enable us to get reliable results and to make stable estimates by considering both factors of cost and time. Morgan et al. (2003) in their study reported that the relationship between sample size and model accuracy is an important issue for data mining and this relation should not be ignored since model accuracy improves at a decreasing rate with increasing sample size. Despite the increases in processing speeds and reductions in processing cost, applying data mining tools to analyze all available data is costly in terms of both economy and time required to generate and implement models. As it can be seen from the results of this study R^2 and accuracy values increased with sample size. Similar results were also reported by Morgan et al. (2003).

Oates and Jensen (1998) reported that increasing the amount of data used to build a model often results in a linear increase in model size, even when that additional complexity results in no significant increase in model accuracy. Frey and Fisher (1999) in their study they systematically examined the response of modeling accuracy to changes in sample size using the C4.5 decision tree algorithm applied to 14 datasets from the UCI repository. Although they did not focused on determining an optimal sample size, they found that the response of predictive accuracy to sample size was more accurately predicted by a regression based on a power law function than by regressions using linear, logarithmic, or exponential functions. As a result of this simulation study, it was found that the performance of the ALM method was higher and it could enable reliable estimates with samples with a smaller size compared to the RT method. The results of this study support the results of studies of Morgan (2003), Oates and Jensen (1998), Frey and Fisher (1999), and James (2013). Mannila (2000) suggested that the volume of

the data was probably not a very important difference. He reported that the number of variables often had a much more profound impact on the applicable analysis methods in discussing differences between statistical and data mining approaches. However, the best way for getting reliable results and making stable estimates is considering both effect of sample size and number of variables simultaneously along with other factors like correlations between predictors.

When a general evaluation is made based on simulation results, it is possible to conclude that although decision trees are generally known as one of the most successful data mining tools, they may not always be the best or not produce reliable results due to the fact that they are built based on some insatiable algorithms for limited or small data set. Therefore, comparison of the methods, algorithms or tools which may be able to use for the same purpose via a comprehensive simulation study will be beneficial for both evaluating their performance and to determine optimum sample size to get reliable results and stable estimates under different experimental conditions. As a result, it is possible to reach the following conclusions without ignoring the fact that differences in experimental conditions (i.e. number of variables, measurement types of variables, correlations between predictors etc.) may affect the reliability of the results to be obtained.

- a) The RT requires much more samples when comparing to ALM. Minimum required sample size for RT technique should be approximately twice that of the ALM test.
- b) Although it was observed that the ALM was much more powerful than RT for estimating continuous response and determining important factors, it should not be ignored a potential threat of misuse of the ALM due to its simplicity.

References

- Anonim, 1994. FORTRAN Subroutines for Mathematical Applications. IMSL MATH/LIBRARY. Vol.1-2, Houston: Visual Numerics, Inc.
- Breiman, L., Friedman, J., Olshen, R. and C. Stone. (1984). *Classification and Regression Trees*. Wadsworth International Group.
- Cios, K., Pedrycz, J., Swiniarski, W.R.W. and Kurgan, L.A. (2007). *Data mining: A knowledge discovery approach*. New York: Springer.
- Fan, J., Han, F. and Liu, H. (2014). *Challenges of Big Data Analysis*. National Science Review, 1 (2) :293-324.
- Field, A. (2013). *SPSS Automatic Linear Modeling*. Retrieved from <http://www.youtube.com/watch?v=pltr74IlxOg>
- Frey, L. and Fisher, D. (1999). *Modeling Decision Tree Performance with the Power Law*. *Proceedings of the Seventh International Workshop on Artificial Intelligence and Statistics*. San Francisco, CA: Morgan-Kaufmann, pp59-65.
- Hill, C. M., Malone, L.C. and Trocine, L. (2004). Data mining and traditional regression. In H. Bozdogan (Ed.), *Statistical data mining and knowledge discovery* (pp. 233–249). Boca Raton, FL: CRC Press LLC.
- IBM SPSS Inc. (2012). *IBM SPSS Statistics 21 Algorithms*. Chicago: Author.
- Kaur, D. and Pulugurta, H. (2008). Comparative Analysis of Fuzzy Decision Tree and Logistic Regression Methods for Pavement Treatment Prediction. *WSEAS Transactions on Computers*, 5 (6): 979-988.
- Keskin, E. and Mendeş, M. (2019). A Graphical Perspective for determining Factors Affecting Grade Success of Students. *III. International Conference on Awareness Proceedings*, pp:31-35, 5-7 December, Çanakkale/Turkey
- Kolaczyk, E. (2013). *Statistical Analysis of Network Data*. YES VI: Statistics for Complex and High Dimensional Systems, Jan 28-29, Eindhoven, The Netherlands.

- Larose, D.T., 2005. *Discovering knowledge in data: An introduction to data mining*. Hoboken, NJ: John Wiley & Sons, Inc.
- Lin, Z., Chen, Y., Liang, Y., 2008. Application of data mining classification algorithm for Customer membership Card Model. *College of Transportation and Management*, China.
- Lynda, P.C., 2011. *How to use automatic linear modeling*. Retrieved from <http://www.youtube.com/watch?v=JIs4sMxAFg0>
- Mannila, H., 2000. Theoretical Frameworks for Data Mining. *SIGKDD Explorations*, 1 (2):30-32, ACM SIGKDD.
- Morgan, J., Daugherty, R., Hilchie, A., Care, B., 2003. Sample Size and Modeling Accuracy of Decision Tree based Data Mining Tools. *Academy of Information and Management Sciences Journal*, 6 (2): 77-92.
- Morgan, J., Daugherty, R., Hilchie, A., Care, B., 2003. Sample Size and Modeling Accuracy of Decision Tree based Data Mining Tools. *Academy of Information and Management Sciences Journal*, 6 (2): 77-92.
- Oates, T., Jensen, D., 1997. The Effects of Training Set Size on Decision Tree Complexity. *Machine Learning: Proceedings of the Fourteenth International Conference*. Morgan Kaufmann, pp. 254-262.
- Oates, T., Jensen, D., 1998. Large Data Sets Lead to Overly Complex Models: an Explanation and a Solution. *Proceedings of the Fourth International Conference on Knowledge Discovery and Data Mining*. Menlo Park, CA: AAAI Press, pp. 294-298.
- Stack Exchange., 2011. *Is automatic linear modeling in SPSS a good or bad thing?* Retrieved from <http://stats.stackexchange.com/questions/7432/is-automatic-linear-modelling-in-spss-a-good-or-bad-thing>

Tarla Silindirinin Farklı Ağırlıklarının Toprağı Sıkıştırmaya Olan Etkisinin Sonlu Elemanlar Yöntemiyle Analiz Edilmesi

Hüsnüyusuf PINAR^{1*}, Songül GÜRSOY¹

¹Dicle Üniversitesi, Ziraat Fakültesi, Tarım Makinaları ve Teknolojileri Mühendisliği Bölümü, 21280, Diyarbakır, Türkiye

*Sunucu yazar e-posta: hsnuyusuf@hotmail.com

Özet

Tarımsal üretimde toprak işleme sonucunda kabaran toprağı bastırarak, toprak içerisinde oluşan yoğun hava boşluklarını giderilmesi ve tohumun toprakla temasını sağlayarak daha kolay çimlenmesine yardımcı olmak amacıyla ekimden önce, ekimle birlikte veya ekimden sonra çoğunlukla tarla silindirleri kullanılmaktadır. Fakat toprağın gereğinden fazla sıkıştırılması topraktaki oksijen miktarını ve bitki kök gelişimini azalttığı için ürün veriminde önemli ölçüde düşüslere neden olmaktadır. Bu çalışmada tarla silindirinin farklı ağırlık seviyelerinin silindir in topraktaki temas alanı ve gerilim-deformasyon özellikleri üzerine olan etkileri sonlu elemanlar metoduyla (FEM) araştırılmıştır. Sonlu elemanlar yöntemiyle belirlenen temas alanı değerleri, tarla koşullarında yürütülen deneme sonuçlarıyla karşılaştırılarak geliştirilen modelin geçerliliğı test edilmiştir. ANSYS 2019 R1 ticari paket programında toprağın lineer-elastik davranış sergilediğı varsayılarak geliştirilen modelin analiz sonuçlarında da silindire uygulanan yükün artışıyla hem silindir in toprakta meydana getirdiğı temas alanı hem de gerilim ve deformasyon değerlerinde artış meydana geldiğı görülmüştür. Modelin geçerliliğini test etmek amacıyla simülasyon çalışmalarının deneysel çalışmalarla karşılaştırılması sonucunda, Deneysel ve simülasyon çalışmalarındaki farklı yük seviyelerindeki silindir in toprakla temas alanı sonuçları arasındaki hata oranının %20'den düşük olduğı tespit edilmiştir.

Anahtar kelimeler: Tarla silindiri, Toprağın sıkıştırılması, Gerilim-deformasyon, Sonlu elemanlar yöntemi, ANSYS

Investigation of the Effect of Land Roller on Soil Compaction Using FEM

Abstract

The main purpose of the use of field rollers after sowing in agricultural production is to contribute to the gentle compaction of the soil without excessive compaction. The physical properties of the soil such as the soil type, the moisture content of the soil at the time of compaction, as well as the weight, diameter and working speed of the roller have a significant effect on the degree of compaction of any roller. Within the scope of this project, the effects of the different weight levels of the field roller used during the preparation of the seed bed or after sowing on the contact area of the roller in the soil and soil compaction were investigated by experimental and finite element method (FEM). The validity of the developed model was tested by comparing the simulation results which is obtained by the finite element method with experimental results carried out under field conditions. in the model developed by the finite element method, assuming that the soil is completely homogeneous and elastic, it has been observed that the stress and contact area in the soil increase in proportion to the load of the cylinder. In the experimental and simulation studies, the error rate between the soil contact area results of the roller at different roller load level was less than 20%.

Key words: Land roller, Soil compaction, Contact area, Finite element analysis

GİRİŞ

Tarımsal üretimde tohum yatağı hazırlığı esnasında veya tohumun ekiminden sonra toprak yüzeyinin bastırılmasında ve düzeltilmesinde genellikle düz merdaneler kullanılmaktadır. Düz silindir şeklinde olan bu tip merdaneler, toprak yüzeyini düzgünleştirilmesine katkı sağladığı gibi toprağın üst katmanının ekim derinliğinde sıkıştırdıkları için topraktaki nem kaybını önlemekle birlikte bitkinin kullanmadığı alt katmandaki nemin kapilarite ile yükselmesinde önemli rol oynamaktadırlar. Ayrıca, tohumun toprakla daha iyi temasına katkı sağlayarak, çimlenme ve çıkış oranında artış meydana getirmektedir (Klassen ve Watts, 1998; Berti ve ark., 2008; Tong ve ark., 2015). Fakat, uygun olmayan koşullarda toprağın aşırı şekilde sıkıştırılması, toprağın fiziksel ve kimyasal özelliklerini olumsuz yönde etkileyerek, bitki çıkış oranı, kök gelişimi, bitkinin besin maddelerinden yararlanma oranını ve sonuç olarak verimini olumsuz yönde etkilemektedir.

Tarla silindirlerinin toprağı sıkıştırma karakteristiğı, toprağın tekstürü ve fizikomekanik özellikleri, iklim koşulları, silindirin ağırlığı, silindir çapı, silindirin toprağı uyguladığı temas basıncı ve traktörün hızı gibi bir çok faktörün etkisi altındadır (Plaster, 1992). Ayrıca, silindirin toprakla meydana getirdiğı temas alanı, silindirin toprak yüzeyindeki temas basıncı miktarını değıştireceğı için toprağın gerilim-deformasyon ilişkisini önemli derecede etkileyecektir.

Tarımsal topraklara baskı uygulayarak sıkıştırılmasının toprak sıkışıklığına olan etkisinin belirlenmesinde, deneysel, analitik, sayısal yöntemler kullanılarak belirlenebilmektedir. Deneysel yöntemlerde, güvenilir bir test için çok tekerrürlü ve faktörlü denemelere ihtiyaç duyulmaktadır. Ayrıca deneysel yöntemde gerçek ölçümler kullanıldığı için maliyet ve zaman açısından dezavantajları vardır. Sayısal modelleme yönteminde, toprak-alet etkileşiminin toprak özelliklerine olan etkisini belirlemede, çalışma koşulları gerçeğine uygun şekilde simüle edilerek belirlenebilmektedir. Bilgisayar teknolojisindeki gelişmelere bağılı olarak ticari ve açık erişim paket programlarının geliştirilmesi, toprak işleme aletlerinin toprak sıkışıklığına etkisi daha kolay bir şekilde tahmin edilmesinde kullanılan modellerin oluşturulmasına katkı sağlamıştır. Bu konuda günümüzde en çok kullanılan yöntemler, Sonlu Elemanlar Yöntemi (FEM, Finite Element Method) ve ayrık elamanlar (DEM, Discrete Element Method) yöntemleridir (Smith, 2014).

Bu proje kapsamında tohum yatağının hazırlığı esnasında veya ekim sonrası kullanılan tarla silindirinin ağırlığının toprak sıkışıklığı üzerinde olan etkileri sonlu elemanlar metoduyla (FEM) araştırılmıştır. Yürütölen çalışmalarda, tarla silindiri farklı ağırlıklarının (687.28, 1181.83, 1676.38 kg) toprakla temas alanı incelenmiş, ayrıca silindirin topraktaki gerilme ve deformasyona olan etkisi araştırılmıştır. Sonlu elemanlar yöntemiyle geliştirilen modelin geçerliliğini test etmek amacıyla, simölasyon ortamında elde edilen silindir temas alanları değıerleri, tarla koşullarında yürütölen deneme sonuçlarıyla karşılaştırılmıştır

MATERYAL VE YÖNTEM

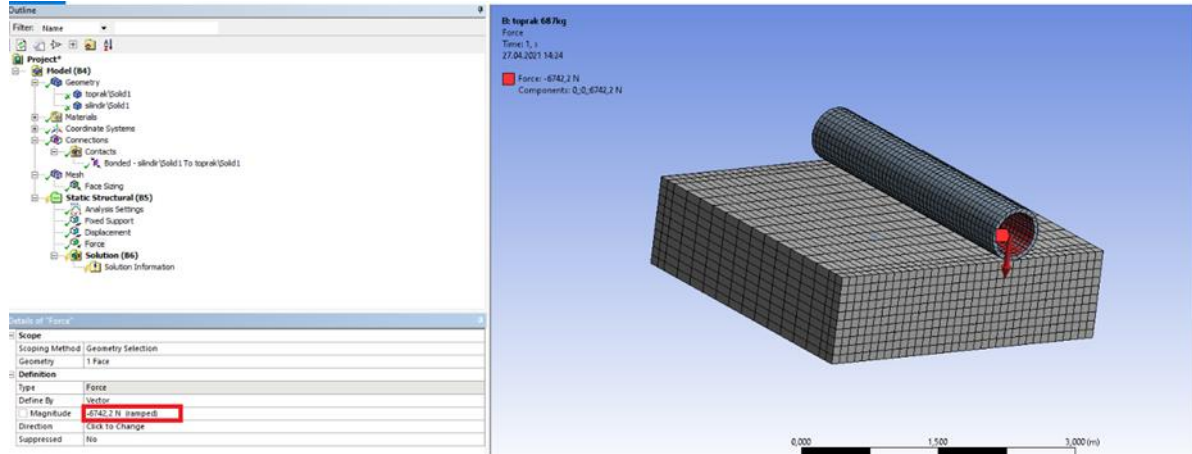
Silindirin statik olarak farklı ağırlıklarının toprakta meydana getirdiğı gerilme ve deformasyon özelliklerinin simölasyonu için, toprağın homojen-elastik olduğı varsayılarak ANSYS 2019 R1 ticari paket programında bir model oluşturulmuştur. Model oluşumunun ilk aşamasında SolidWorks 2019 R1 paket programı üzerinde çizimi ve montajı yapılan silindir ve toprak bloğunun geometrik modeli, ANSYS 2019 R1 Workbench paket programı üzerinde açılarak analiz edilmiştir. Oluşturulan geometrik toprak bloğı modelinin genişliği 3500 mm, uzunluğu 3500 mm ve derinliği 1000 mm, silindirin çapı 600 mm ve genişliği ise 3500 mm olarak tasarlanmıştır.

Toprak bloğunun malzeme özellikleri, tarla denemesindeki toprağın fiziksel özellikleri göz önünde bulundurularak literatürden seçilmiştir. Modelin oluşturulmasında kullanılmak üzere denemelerdeki toprak ve silindirin malzeme özellikleri Çizelge 1’de verilmiştir.

Çizelge 1. Modelin oluşturulmasında kullanılan toprak bloğu ve silindirin malzeme özellikleri (Lee ve Su, 1999; Cui et al (2007))

Model parametreleri	Sembol	Birim	Değer
Silindir (Çelik malzeme)-Yoğunluk	ρ_s	Kg/m ³	7800
Silindir (Çelik malzeme)-Elastiklik modülü	Es	Mpa	210 x 10 ³
Silindir (Çelik malzeme)-Poisson oranı	vs	-	0.29
Toprak bloku(killi toprak)-Yoğunluk	ρ_T	Kg/ m ³	1200
Toprak bloku (killi-tınlı)-Elastiklik modülü	Et	Mpa	0.3
Toprak bloku (killi-tınlı)-Poisson oranı	vt	-	0.3

Model analizi yapabilmek için sınır koşullarının tanımlanması amacıyla toprak bloğu alt yüzeyinden üç hareket ekseninde de sabit olacak şekilde mesnetlendirilmiştir. Silindir ve toprak zemini arasında bağlı (bonded) ilişki kurulmuştur. Daha sonra malzameye ait elastisite modülü ve poisson oranı tanımlanan model, sonlu elemanlara bölünerek ağ örgüsü (mesh) yapılmıştır. Sınır koşulları ve meshleme işlemlerinden sonra Static Structural sekmesinin altında silindirin iç yüzeyini seçerek uygulanacak kuvvet (force) miktarları modelde tanımlanmıştır. Şekil 1’de ANSYS programı yardımı ile sonlu elemanlar yöntemi kullanılarak tasarlanan modelin eleman ağı ve uygulanan kuvvet görülmektedir.



Şekil 1. Sonlu elemanlarla tasarlanan silindirin mesh yapılması ve kuvvet uygulaması

Silindirin toprakta meydana getirdiği gerilim-deformasyon analizi sağlamak için toprağın tamamen homojen ve elastik olduğu varsayılarak linear elastik kriter kullanılarak modellenmiştir. Oluşturulan modelin çözümü yapıldıktan sonra ANSYS programından çıktı olarak toprağın farklı derinliklerdeki gerilme dağılımı ve toprakta meydana gelen şekil değiştirme miktarı alınmıştır. Silindirin topraktaki temas alanı, toprakta meydana gelen çökme miktarının Solidworks programında silindir üzerinde kıyaslamasıyla silindirin temas eden yay uzunluğundan hesaplanmıştır.

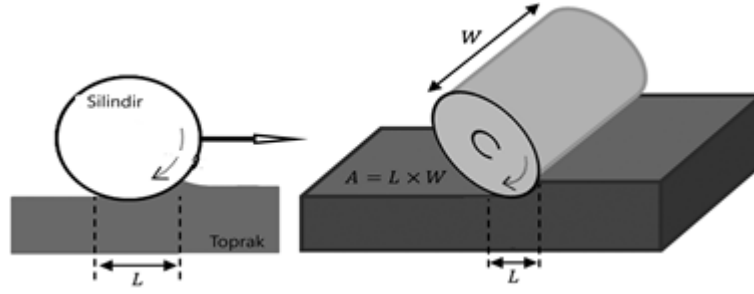
Geliştirilen modelin geçerliliğini test etmek amacıyla, Dicle Üniversitesi Ziraat Fakültesi’nin arpa ekiminin yapılmış üretim alanında 2020-2021 üretim sezonunda silindirin farklı ağırlıklarının (687.28, 1181.83, 1676.38 kg) toprakla temas alanına etkisini belirlemeye yönelik tarla denemesi çalışmaları yürütülmüştür.

Silindir ve toprak arasındaki temas alanı, silindirin toprakla temas ettiği yayın uzunluğu ölçülerek (Şekil 2’deki L uzunluğu), denklem 1’deki gibi hesaplanmıştır.

$$A = L \cdot W \quad 1$$

A= Silindirin toprakla temas alanı, m², L= Silindirin toprakla meydana getirdiği yay uzunluğu, m

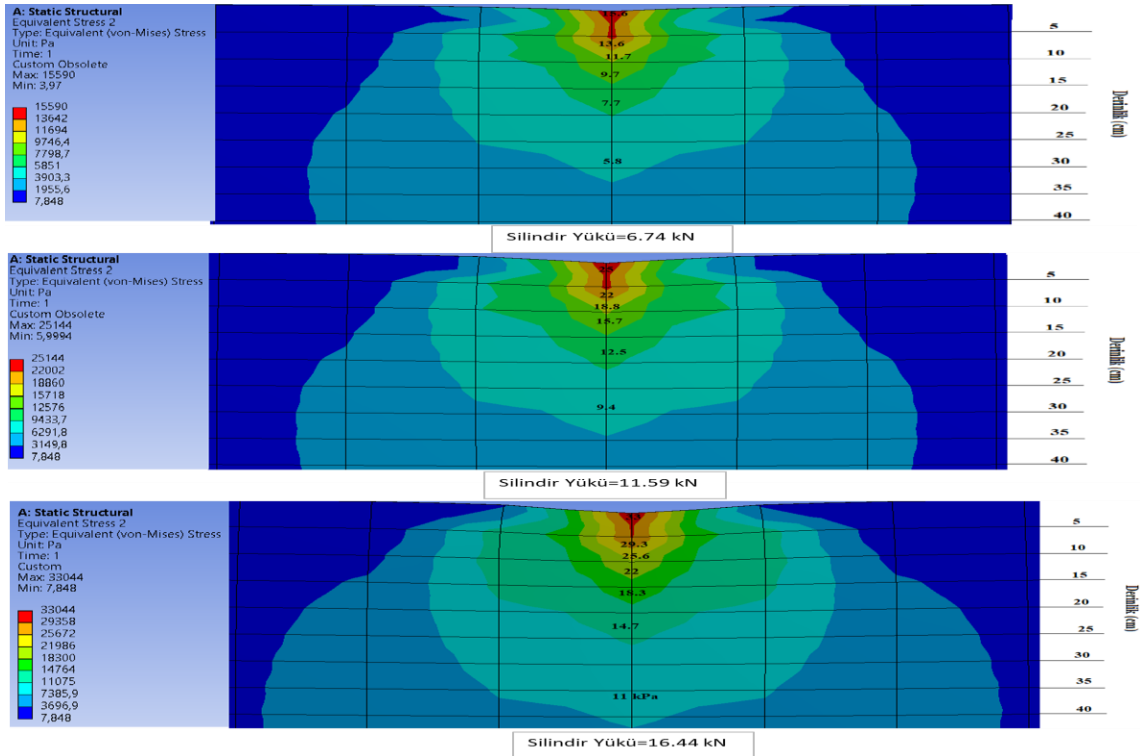
W= silindir genişliği, m



Şekil 2. Silindirin topraktaki temas alanının belirlenmesi

BULGULAR

Farklı yük miktarları (6.24, 11.59 ve 16.44 kN) altındaki silindir modelinin toprakta oluşturduğu gerilim dağılımları Şekil 3’de görülmektedir. Silindir modeli 6.24 kN düşey yük altındayken toprakta oluşturduğu maksimum gerilme 15590 Pa, minimum gerilme 3.97 Pa olurken, 11.59 kN yük altındaki topraktaki gerilme maksimum 25144 Pa, minumum 5.9994 Pa ve 16.44 kN yük altında ise maksimum 33044 Pa, minumum 7.848 Pa olarak belirlenmiştir. Silindirin tüm yük miktarları koşullarında toprağın derinliğinin artışıyla topraktaki gerilimin azaldığı Şekil 3’de görülmektedir. Aynı şekilde, silindirin yükünün artışıyla, silindirin etki derinliğinin daha fazla olduğu görülmüştür. Sonlu elemanlarla analiz edilen silindirin farklı düşey yük miktarları altında toprakta meydana getirdiği çökme miktarının Solidworks programında silindir üzerinde kıyaslamasıyla hesaplanan silindirin temas alanı değerleri Çizelge 2’de verilmiştir. Silindir yükündeki artışla silindirin temas alanının da önemli düzeyde arttığı Çizelge 2’de görülmektedir.



Şekil 3. Similasyon çalışmalarında farklı silindir yük miktarlarının toprağın gerilimine etkisi

Toprağın homojen ve elastik olduğu farzedilerek geliştirilen modelin geçerliliği, silindirin toprakla temas alanının simülasyon ve deneysel sonuçları arasındaki farktan bağıl hata oranı denklem 2'den hesaplanarak test edilmiştir. Silindirin farklı yük (6.74, 11.59, 16.44 kN) altındaki toprakla temas alanına ait simülasyon ve deneysel sonuçlarının bağıl hata oranları Çizelge 2'de verilmiştir.

$$RE = \left(\frac{A_d - A_s}{A_d} \right) \cdot 100 \quad 2$$

Çizelge 2 incelendiği zaman silindirin toprakla temas alanına ait deneysel ve simülasyon sonuçları arasındaki bağıl hata oranının %10'dan daha düşük olduğu görülmektedir. Bu da geliştirilen modelin silindirin toprakla temas alanının tahmin etmede kullanılabileceğini göstermektedir.

Çizelge 2. Silindirin toprakla temas alanına ait deneysel ve simülasyon sonuçları ve bağıl hata oranları

Silindir Yüğü, kN	A_d, m^2	A_s, m^2	RE, %
6.74	0.71	0.65	8.45
11.59	0.77	0.82	6.5
16.44	0.89	0.94	5.0

TARTIŞMA VE SONUÇ

Bu çalışma kapsamında, tohum yatağının hazırlığı esnasında veya ekim sonrası kullanılan tarla silindirinin farklı ağırlık seviyelerinin silindirin topraktaki temas alanı ve toprak sıkışıklığı üzerinde olan etkileri deneysel ve sonlu elemanlar metoduyla (FEM) araştırılmış ve silindir ağırlığının optimizasyonunu sağlamak amacıyla ANSYS 2019 R1 ticari paket programında bir model oluşturulmuştur. Geliştirilen bu modelin geçerliliği simülasyon çalışmalarında elde edilen silindirin temas alanı sonuçları, tarla koşullarında yürütülen deneme sonuçlarıyla karşılaştırılarak test edilmiştir.

Toprağın tamamen homojen ve elastik olduğu farzedilerek geliştirilen modelde, silindirin toprakla temas alanı %10'un altında bir hata payı ile tahmin edilebilmiştir. Fakat bilindiği gibi tarımsal topraklar homojen olmayıp, elasto-plastik yapıya sahiptirler. Dolayısıyla, toprak-alet etkileşiminin sonlu elemanlar yöntemiyle analizinde malzemenin elastisitesi kullanılarak modelinin kullanımından ziyade malzemenin plastisiteside göz önünde bulundurularak oluşturulan modelinin daha uygun olabileceği tavsiye edilmektedir. Toprak gibi malzemelerde en yaygın kullanılan Drucker-Prager plastisite modelidir. Ayrıca, bu çalışmada, statik haldeki silindir yükünün toprakta meydana getirdiği gerilim deformasyon analizi yapılmıştır. Gelecekteki çalışmalarda silindir modelinin hareketi sonucu toprakta meydana gelecek gerilim-deformasyon değişimlerinin araştırılması önem arz etmektedir.

Statik durumdaki silindirin toprakta meydana getirdiği gerilim analizinde, silindirin yük miktarındaki artışla toprakta meydana gelen gerilim miktarı da artmıştır. Ayrıca, silindirin yükünün artışıyla, silindirin etki derinliğinin daha fazla olduğu görülmüştür. Silindirin tüm yük miktarları koşullarında toprağın derinliğinin artışıyla topraktaki gerilimin azaldığı ve silindirin temas alanı üzerindeki gerilimin daha yüksek olduğu görülmüştür.

Kaynaklar

- Berti MT, Johnson BL Henson RA, 2008. Seeding depth and soil packing affect pure live seed emergence of Cuphea, Industrial Crops and Products, 27: 272–278.
- Klassen EP, Watts J, 1998. Effects of rolling a clay loam field, To press in stones on dry edible bean and dry pea production. Progress Reports on Pulse Research in Western Canada, 3:15.
- Plaster EJ, 1992. Soil Science and Management. Delmar Publishers Inc., New York, USA.
- Smith WC, 2014. Modeling of Wheel-Soil Interaction for Small Ground Vehicles Operating on Granular Soil. A dissertation submitted in partial fulfillment of the requirements for the degree of Doctor of Philosophy (Mechanical Engineering) in The University of Michigan
- Tong J, Zhang Q, Guo L, Chang Y, Guo Y, Zhu F, Chen D, Liu X, 2015. Compaction performance of biomimetic press roller to soil. Journal of Bionic Eng., 12:152–159.

**Toprak İşlemesiz (Doğrudan Ekim) Tarımda Verim Üzerine Etkili Faktörlerin
Araştırılmasına Yönelik Bir Meta-Analiz Çalışması**

Songül GÜRSOY^{1*}

¹Dicle Üniversitesi, Ziraat Fakültesi, Tarım Makinaları ve Teknolojileri Mühendisliği Bölümü, 21280,
Diyarbakır, Türkiye

*Sunucu yazar e-posta: songulgursoy@hotmail.com

Özet

Son yıllarda artan çevre bilinci, ekonomik üretim talepleri ve enerji kullanımında tasarrufa gitme zorunluluğu nedeniyle, Dünya’da ve Türkiye’de korumalı toprak işleme sistemleri arasında önemli bir yere sahip olan doğrudan ekim yöntemlerinin yaygınlaştırılmasına yönelik birçok çalışma yapılmasına rağmen, bu yöntemlerin adaptasyonlarının oldukça düşük olduğu görülmektedir. Dolayısıyla, bu çalışmada, ulusal/uluslararası nitelikli hakemli dergilerde yayımlanan ve Türk örneklemelerinden toplanan ampirik veriye dayalı makaleler üzerinden, doğrudan ekim yöntemlerinde verim üzerinde etkili faktörler meta-analitik yöntemle incelemiş ve verimi azaltan faktörlerin detaylı olarak analizi hedeflenmiştir.

Anahtar kelimeler: Toprak işleme, Doğrudan ekim, Verimi etkileyen faktörler, Meta-analiz

The Effect of Body Phenotypic Traits on Live Weight in Ovine Breeding: A Meta-Analysis

Senol CELİK^{1*}

¹Bingol University, Faculty of Agriculture, Department of Animal Science, Biometry and Genetic
Bingol, Turkey

*Presenter author e-mail: senolcelik@bingol.edu.tr

Abstract

A meta-analysis was conducted to determine the effect of various bodily variables on live weight in ovine animals such as sheep, goat and ram. The studies performed between 2006 and 2019 were reviewed on Google Scholar's and Web of Science's databases. 20 studies were included in the meta-analysis as a result of the literature review. The effect size of the studies included in the meta-analysis was calculated using correlation (r) using Jamovi statistical software program. Sample size (n) and correlation coefficient (r) statistics were used in order to obtain the effect size of the studies. In addition, the results have revealed that the overall effect sizes of correlation between live weight and various phenotypic traits of body in ovine animals is 0.87 based on the random effects model. Moreover, the results showed that there was no publication bias in the studies analyzed. It was found that various morphological traits of body were likely to have a significant effect on live weight in ovine husbandry, and that this was confirmed by the meta-analysis.

Key words: Meta analysis, effectiveness, live weight, body size

INTRODUCTION

The meta analysis is a statistical method also known as a secondary analysis conducted to have more reliable and accurate data about a study subject by combining certain rules for published studies performed by various researchers of the same discipline (Erol and Tekin, 2018). The meta analysis enables researchers to estimate reliable and viable parameters with the minimum number of variances for a study (Mosteller and Colditz 1996). The most outstanding aspect of the meta analysis, which has become more and more common in many disciplines including medicine, economics, sociology and engineering, is that it gathers all the studies performed before is not sufficient. As the number of researchers studying the same subject exponentially increases, almost each and every scientific discipline needs to adopt the meta-analysis (Yıldız and Tez, 2009).

Some popular studies have been performed in many disciplines concerning the meta analysis in recent years. The meta analysis is also performed in agriculture. Küçükönder and Efe (2014) adopted the Meta-Analysis of Mantel Haenszel method for eight subjects in agriculture, and interpreted the odds ratio based on the combination of results. Küçükönder et al. (2015) adopted the Meta Analysis to investigate the effect size of lactation order and calving season on a 305-day lactation yield. For another study, the meta analysis was adopted to identify factors affecting the surface roughness for various types of peaches. Whether there was any correlation between attacks of ants at trees and types of trees were investigated by the meta-analysis (Şelli and Doğan, 2011). The odds ratio was used for the effect measurement and the methods of meta-analysis based on Mantel Haenszel, Peto and inverse variance were used for the method in the above mentioned study. Doğan and Şahin (2003) investigated sex and mode of birth as a factor in terms of their effect on birth weight of sheep. The difference (d) between the means and Bare-Bones meta analysis were adopted as an effect size and a methodology respectively.

As a result, they concluded that the mode of birth was more effective on birth weight than sex for sheep. Jembere et al. (2017) investigated estimations of genetic parameters by means of meta analysis for growth, reproduction and milk yield in goats.

The aim of the present study was to examine the effect of various body measurements (height at withers, height at rump, chest girth, chest width, body length) and age on live weight in ovine animals such as sheep, goats and rams using meta-analysis.

MATERIAL AND METHOD

The studies included in the analysis were the ones performed between 2006 and 2019 and published with essential quantitative data and statistical analyses for estimations on live weight of ovine animals. The studies were reviewed on Google Akademik/Scholar website. The review was performed involving papers that contain the terms live weight, sheep, goat and ram in their titles and keywords. The review yielded many studies about the subject such as data mining, regression analysis and correlation analysis. Having been reviewed in line with the purpose of the study, 20 papers fulfilling the criteria were included in the analysis. The effect size was calculated based on the sample size and correlation coefficient.

Descriptive data regarding the 20 studies included in the meta-analysis are displayed in Table 1.

Table 1. Descriptive belong to the studies included in the meta-analysis.

Author(s) and Year	Trait	Sample size (n)	Correlation coefficient (r)
Gürçan and Akçapınar (2006)	Deutsche Merino sheep	68	0.82
Oral and Altınel (2006)	Hair goat	347	0.786
Afolayan et al. (2006)	Yankasa sheep	258	0.94
Sowande and Sobola (2008)	West African Dwarf sheep	210	0.97
Yakubu (2012)	Uda rams	499	0.787
Mohammad et al. (2012)	Sheep	239	0.849
Koncagül et al. (2012)	Zom sheep	211	0.75
Momoh et al. (2013)	Nigeria sheep	595	0.550
Khan et al. (2014)	Harnai sheep	730	0.917
Jafari et al. (2014)	Makuie sheep	2264	0.87
Menesatti et al. (2014)	Alpagota sheep	27	0.87
Ali et al. (2015)	Harnai lambs	161	0.918
Karakuş (2016)	Goat	77	0.754
Eyduran et al (2017)	Beetal goat	205	0.864
Karabacak et al. (2017)	Sheep	47	0.805
Celik et al. (2017)	Mengali rams	107	0.959
Akkol et al. (2017)	Hair goat	475	0.954
Çelik (2019)	Pakistan goats	130	0.950
Huma and Iqbal (2019)	Balochi sheep	131	0.994
Khorshidi-Jalali et al. (2019)	Raini Cashmere Goat	1389	0.927

The meta-analysis was used to calculate the effect size of correlation between live weight and other body traits for sheep, goats and rams. A meta-analysis includes the review of results from multiple studies performed about a certain subject (Şelli and Doğan 2011). A meta analysis is based on what statistical model is to be adopted and whether the effect size has a homogeneous distribution. P value of Q homogeneity test being above 0.05 points to homogeneous distribution, and consequently the random effects model (REM) is adopted when the fixed effect model (FEM) is below 0.05 (Ellis, 2010).

Jamovi 1.1.9.0 was used to calculate the individual effect size and overall effect size of the studies. The correlation coefficient itself is an effect size value (Kotrlik and Williams 2003). The effect size (ES) classification was as follows: $0 \leq ES < 0.10$ insignificant, $0.10 \leq ES < 0.30$ small, $0.30 \leq ES < 0.50$ moderate, $0.50 \leq ES < 0.80$ large, and $0.80 \leq ES < 1$ very large (Cohen et al., 2007). Next, these were revised and expanded by Cicchetti (2008), as: $ES \leq 0.10$ (trivial); $0.10 \leq ES \leq 0.29$ (small); $0.30 \leq ES \leq 0.49$ (moderate); $0.50 \leq ES \leq 0.69$ (large); and $EC \geq 0.70$ (very large).

Orwin's method and a funnel chart are adopted to establish the publication bias (Lipsey and Wilson 2001). In addition, a funnel plot can offer an insight into the publication bias. On the funnel plot, X axis signifies the effect size of each paper included in the study while Y axis signifies the sample size, variance or standard error of the papers. Based on the plot, a symmetric distribution of the papers included in the study points to the fact that they are reliable and therefore they do not have any publication bias (Üstün and Eryılmaz 2014).

Test of heterogeneity is performed as following (Higgins and Thompson 2002).

$$Q = \sum_{i=1}^k W * (ES_i - ES_{fixed})^2$$

to a χ^2 distribution with $k - 1$ degrees of freedom.

$$\tau^2 = \frac{Q - (k - 1)}{\sum_{i=1}^k w_i - \frac{\sum_{i=1}^k w_i^2}{\sum_{i=1}^k w_i}}$$

Where, τ^2 : Between-studies variance; Q: heterogeneity statistic, ES: effect size, EC_{fixed} : fitting effect size.

W^* is the weight calculated as shown following (Borenstein et al., 2009).

$$W * = \frac{1}{v_i}$$

where v_i is the within-study variance for study (i).

I^2 heterogeneity index from a collection of effect sizes by comparing the Q value to its expected value assuming homogeneity, that is, to its degrees of freedom ($df = k - 1$): (Higgins & Thompson, 2002).

$$I^2(\%) = \frac{Q - df}{Q} * 100$$

Here, df: degrees freedom, $df = k - 1$, (for $Q > k - 1$).

The studies composed of academic papers are from the period between 2006 and 2019. The distribution of studies performed in line with the original purpose by years is presented in Table 2.

Table 2. Distribution of studies by years

Year	Frequency	Percent (%)
2006	3	0.15
2008	1	0.05
2012	3	0.15
2013	1	0.05
2014	3	0.15
2015	1	0.05
2016	1	0.05
2017	4	0.20
2019	3	0.15
Total	20	100

Of 20 papers selected for the estimation of live weight in ovine animals, 6 of the them were about goats, 1 about lamb, 2 about rams and 11 about sheep (Table 1). The majority of the papers selected in line with the purpose of the meta analysis was focused on sheep.

The effect size for live weight-phenotypic traits of body in animals was calculated based on the sample size and correlation coefficient. Heterogeneity tests were performed to see whether their effect size had a normal distribution. The results of the heterogeneity test are presented in Table 3.

Table 3. Results on the effect size of studies based on the random effects model

Random-Effects Model (k = 20)								
	Estimate	se	Z	p	CI Lower Bound		CI Upper Bound	
Intercept	0.866	0.0227	38.1	<0 .001	0.821		0.910	
Note. Tau² Estimator: Restricted Maximum-Likelihood								
Heterogeneity Statistics								
Tau	Tau²		I²	H²	R²	df	Q	p
0.099	0.0098 (SE= 0.0033)		99.65%	287.078	.	19	1709.398	<0 .001

The heterogeneity figures in Table 3 show that the studies are heterogeneous (Tau=0.099, I²=%99.65, H²=287.078, Q=1709.398, p<0.001). Since the studies were heterogeneous, the analysis was performed based on the random effects model. Based on the random effects model, the standard error equaled to 0.0227, and the 95% confidence interval was between 0.55 and 0.99, and the mean effect size was 0.87. Performed for the purpose of statistically significant, the Z-test calculation revealed that Z value equaled to 38.1. Thus, one can argue that the result has a statistical significance with p=0.001 (Z=38.1; p=0.001).

Based on the results, it was concluded that the correlation between live weight and body traits had a perfect effect size, and was statistically significant. In other words, it could be noted that different body variables in ovine animals have a positive effect on live weight based on the study since the mean effect size turned out to be positive.

The effect sizes and weight of the studies regarding the variable of live weight are presented in Figure 1 (Forest Plot).

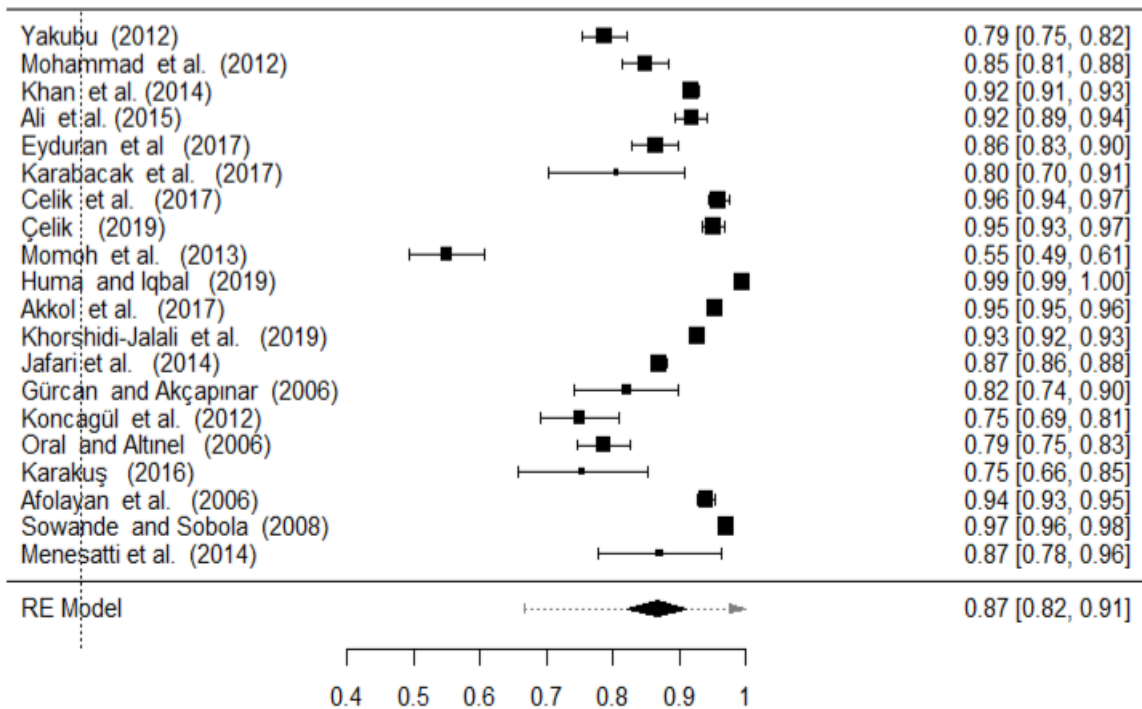


Figure 1. Forest plot (Effect size and weight of studies)

The squares in the forest plot of Figure 1 show the effect size of individual studies included in the analysis while the lines in both sides signify lower and upper limits of their effect size with a 95% confidence interval. The width of the square points to the weight of individual studies while the equilateral quadrangle signifies the overall effect sizes of the studies. The largest effect size turned out to be 0.99 while the smallest one corresponded to 0.55. All of the 20 papers selected in line with the purpose of the study had a positive effect. This suggests that various body sizes of ovine animals have a positive effect on live weight.

A funnel plot, fail-safe N analysis, rank correlation test for funnel plot asymmetry and regression test for funnel plot asymmetry were adopted as a methodology to confirm the reliability and validity of the meta analysis and to detect the publication bias.

The distribution of effect size estimates for the studies over live weight based on a funnel plot is presented in Figure 2.

Based on the Huni plot, the effect sizes of individual studies do not cause a publication bias should they be within the funnel lines and symmetrically distributed while the effect sizes of individual studies cause a publication bias should they be outside the funnel lines and asymmetrically distributed. When the Figure 2 is reviewed based on this data, it could be noted that the effect size estimate of the studies over live weight is nearly symmetric.

A nearly symmetric distribution shows that the publication bias is low. Performed for biasedness indicators of the funnel plot, the Rank Correlation Test for Funnel Plot Asymmetry and the Regression Test for Funnel Plot Asymmetry revealed that Kendall's tau equaled to 0.021, with $p=0.924$, $Z=-3.811$ and $p=0.001$. In this case, the p value is supposed to be over 0.05 for a significant difference, and it

equaled to over 0.05 (Table 4). The results of the analysis performed based on the aforementioned result show that there is no bias. It could be stated that one of the reasons behind it is the large quantity of primary studies. As a result, the funnel plot corroborates the argument that there is not any publication bias.

Table 4. Publication Bias Assessment

Fail-Safe N Analysis (File Drawer Analysis)	
Fail-safe N	p
6137014	<0 .001
Note. Fail-safe N Calculation Using the Rosenthal Approach	
Rank Correlation Test for Funnel Plot Asymmetry	
Kendall's Tau	p
0.021	0.924
Regression Test for Funnel Plot Asymmetry	
Z	p
-3.811	0.001

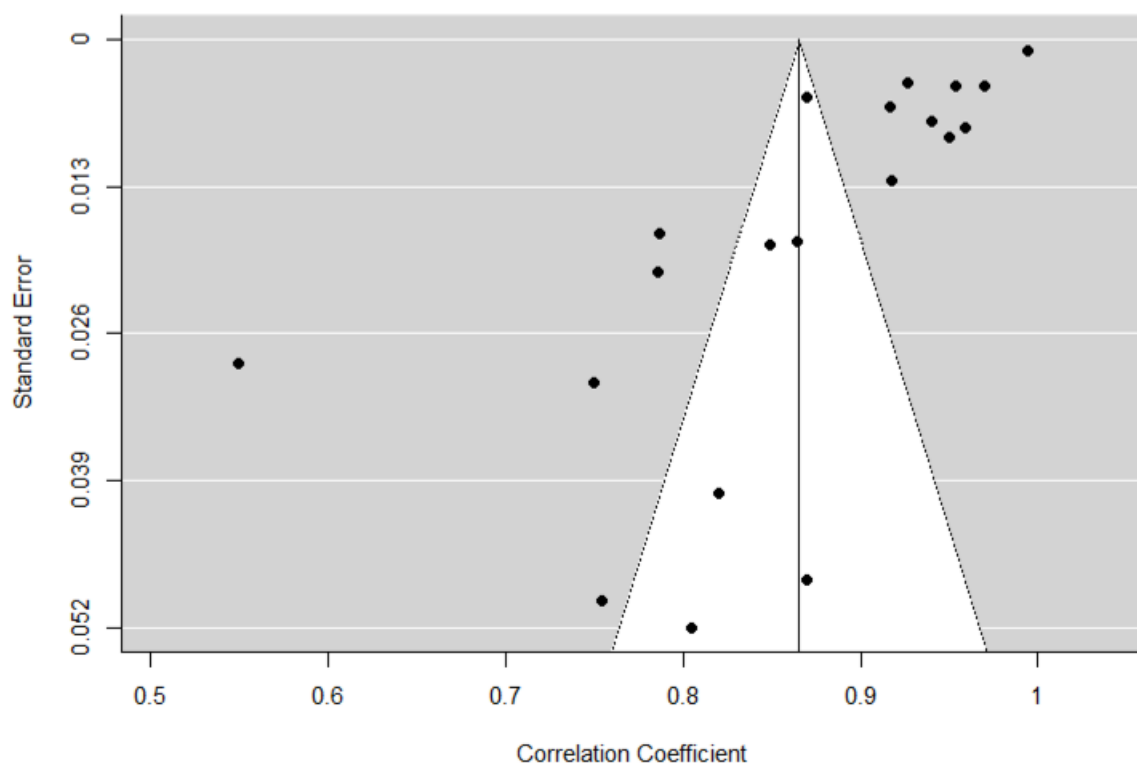


Figure 2. The funnel plot for the effect size estimates of the studies over live weight

A meta analysis of correlation coefficients was performed in this study. Since the confidence interval did not include zero (0) based on the Forest plots presented as a part of the analysis, the effect size estimates were statistically significant.

20 studies investigating the effect of phenotypic traits of body on live weight were selected. In most of those studies, the effect of body measurements such as height at withers, height at rump, chest girth, chest width and body length on live weight was analyzed by statistical methods such as data mining and regression analysis. All of the studies revealed that the body measurements analyzed as an independent variable were significant in terms of the effect on live weight.

The confidence interval for the effect size (ES) of 20 studies that investigate the effect of phenotypic traits of body on live weight ranged from 0.55 to 0.99, and never included zero (0). This result is similar to that of Küçükönder et al. (2015) with the ES confidence interval ranging from 0.596 to 0.967, and to that of Özder et al. (2004) from -0.50 to 0, and that of Aksoy et al. (2001) from 4 to 6. However, it was found to be different than that of Doğan and Şahin (2003) who reported that the ES confidence interval ranged from 0.315 to 1.117 for variables of sex and mode of birth having an effect on birth weight of lamb.

As it is heterogeneous, this study is similar to that of Bengtsson et al. (2005), Tuck (2014), Lean et al. (2009) and Nimer and Lundahl (2007). The study revealing that there is no publication bias, differs from that of Tuck (2014) in terms of unbiasedness. This study is solely focused on a discussion over meta analyses performed in agriculture.

CONCLUSION

As the result of the meta-analysis conducted in the study, it could be noted that the effect size of phenotypic traits of body on live weight for ovine animals is perfect. In fact, the studies show that the effect of body measurements on live weight is statistically significant. However, age and sex have also a significant effect on live weight. Of effect sizes of 20 studies analyzed, 6 of them were large, ranging from 0.50 to 0.80, and 14 of them were very large with the ES >0.80. It is believed that expanding the use of meta analyses for agricultural purposes and adopting different variables and parameters to achieve various purposes would be beneficial. It is expected that the unbiasedness of the results of this study and how perfect the effect size is will be an alternative statistical method to be adopted for further studies to analyze agricultural data sets.

References

- Afolayan, R. A., Adeyinka, I. A., and Lakpini, C. A. M.: The estimation of live weight from body measurements in Yankasa sheep. *Czech J. Anim. Sci.*, 51(8), 343–348, 2006
- Akkol, S., Akilli, A., and Cemal, I.: Comparison of Artificial Neural Network and Multiple Linear Regression for Prediction of Live Weight in Hair Goats. *YYU J. Agr. Science*, 27(1), 21-29, 2017.
- Aksoy, A. R., Saatcı, M., Özbey, M., and Dalcı, M. T.: Production traits of Tushin sheep I. reproductive ability and growth. *Vet Bil Derg*, 17(1), 73-77, 2001.
- Ali, M., Eydurani, E., Tariq, M. M., Tirink, C., Abbas, F., Bajwa, M. A., Baloch, M. H., Nizamani, A. H., Waheed, A., Awan, M. A., Shah, S. H., Ahmad, Z., and Jan, S.: Comparison of artificial neural network and decision tree algorithms used for predicting live weight at post weaning period from some biometrical characteristics in Harnai Sheep. *Pakistan Journal of Zoology*, 47, 1579-1585, 2015.
- Bengtsson, J., Ahnström, J., and Weibull A-C.: The effects of organic agriculture on biodiversity and abundance: A meta-analysis. *Journal of Applied Ecology*, 42, 261-269, 2005.
- Borenstein, M., Hedges, L. V., Higgins, J. P. T., and Rothstein, H. R.: *Introduction to Meta-Analysis*. Chichester: Wiley, 2009.

- Celik, S.: Comparing Predictive Performances of Tree Based Data Mining Algorithms and MARS Algorithm in the Prediction of Live Body Weight from Body Traits in Pakistan Goats. *Pakistan Journal of Zoology* (PJZ), 51(4), 1447–1456, 2019.
- Celik, S., Eydurhan, E., Karadas, K., and Tariq, M. M.: Comparison of predictive performance of data mining algorithms in predicting body weight in Mengali rams of Pakistan. *Revista Brasileira de Zootecnia*, 46(11), 863-872, 2017.
- Cicchetti, D. V.: From Bayes to the just noticeable difference to effect sizes: A note to understanding the clinical and statistical significance of Oenologic research findings. *Journal of Wine Economics*, 3, 185–193, 2008.
- Cohen, L., Manion, L., and Morrison, K.: *Research Methods in Education*. 2nd edition. New York. Routledge, 2007.
- Doğan İ, Şahin F. 2003. Investigation of birth type and sex factors which effect the birth weight in lambs by Bare-Bones META analysis. *Ankara Üniv. Vet. Fak. Derg.* 50: 135-140.
- Ellis, P. D.: *The essential guide to effect sizes: Statistical power, meta-analysis, and the interpretation of research results*. Cambridge: Cambridge University Press, 2010.
- Erol, E. B., and Tekin, M. E.: Meta-Analysis of the factors effecting the some growth characteristics of the lambs. *Eurasian J Vet Sci.*, 34(3), 150-155, 2018.
- Eyduran, E., Zaborski, D., Waheed, A., Celik, S., Karadas, K., and Grzesiak, W.: Comparison of the Predictive Capabilities of Several Data Mining Algorithms and Multiple Linear Regression in the Prediction of Body Weight by Means of Body Measurements in the Indigenous Beetal Goat of Pakistan. *Pakistan Journal of Zoology*, 49(1), 273-282, 2017.
- Gürçan, I. S., and Akçapınar, H.: Estimation of Live Weight by Statistical Methods Using Body Measurements in Merino Sheep. *Lalahan Hay. Araşt. Enst. Derg.* 46(1), 7-17, 2006.
- Higgins, J. P. T., and Thompson, S. G.: Quantifying heterogeneity in a meta-analysis. *Statistics in Medicine*, 21: 1539-1558. DOI: 10.1002/sim.1186, 2002.
- Huma, Z. E., and Iqbal, F.: Predicting the body weight of Balochi sheep using a machine learning approach. *Turkish Journal of Veterinary and Animal Sciences*, 43, 500-506, 2019.
- Jafari, S., Hashemi, A., Darvishzadeh, R., and Manafiazar, G.: Genetic parameters of live body weight, body measurements, greasy fleece weight, and reproduction traits in Makuie sheep breed. *Spanish Journal of Agricultural Research*, 12(3), 653-663, 2014.
- Jembere, T., Dessie, T., Rischkowsky, B., Kebede, K., Okeyo, A. M., and Haile, A.: Meta-analysis of average estimates of genetic parameters for growth, reproduction and milk production traits in goats. *Small Ruminant Research*, 153, 71-80, 2017.
- Karabacak, A., Celik, S., Tatliyer, A., Keskin, I., Ertürk, Y. E., Eydurhan, E., Javed, Y., and Tariq, M. M.: Estimation of Cold Carcass Weight and Body Weight from Several Body Measurements in Sheep through Various Data Mining Algorithms. *Pakistan Journal of Zoology*, 49(5), 1731-1738, 2017.
- Karakuş, F.: The Effect of Body Condition Score on Fertility, Live Weight and Some Body Measurements in Goats. *YYU J Agr. Sci.*, 26(3), 372-379, 2016.
- Khan, M. A., Tariq, M. M., Eydurhan, E., Tatliyer, A., Rafeeq, M., Abbas, F., Rashid, N., Awan, M. A., and Javed, K.: Estimating body weight from several body measurements in Harnai sheep without multicollinearity problem. *The Journal of Animal Plant Science*. 24(1), 120-126, 2014.
- Khorshidi-Jalali, M., Mohammadabadi, M. R., Esmailizadeh, A., Barazandeh, A., and Babenko, O. I.: Comparison of Artificial Neural Network and Regression Models for Prediction of Body Weight in Raini Cashmere Goat. *Iranian Journal of Applied Animal Science*, 9(3), 453-461, 2019.
- Koncagül, S., Akça, N., Vural, M. E., Karataş, A., and Bingöl, M.: Morphological Characteristics of Zom Sheep. *Kafkas Univ Vet Fak Derg.*, 18(5), 829-837, 2012.
- Kotrlík, J. W., and Williams, H. A.: The incorporation of effect size in information technology, learning, and performance research. *Information Technology, Learning, and Performance Journal*, 21(1), 1-7, 2003.
- Küçükönder, H., and Efe, E. 2014. Meta Analysis and Agricultural Applications. *KSU J. Nat. Sci.* 17(1): 1-7, 2014.
- Küçükönder, H., Üçkardeş, F., and Efe, E. Estimation of the effect size of lactation sequence and calving season on 305-day lactation milk yield with Meta Analysis approach. *Türk Tarım – Gıda Bilim ve Teknoloji Dergisi*, 3(1), 17-21, 2015.

- Lean, I. J., Rabiee, A. R., Duffield, T. F., and Dohoo, I. R.: Use of meta-analysis in animal health and reproduction: Methods and applications. *J. Dairy Sci.* 92, 3545–3565, doi: 10.3168/jds.2009-2140, 2009.
- Lipsey, M., and Wilson, D.: Practical meta-analysis. Beverly Hills, CA: Sage Publications, 2001.
- Menesatti, P., Costa, C., Antonucci, F., Steri, R., Pallottino, F., and Catillo, G.: 2014. A low-cost stereovision system to estimate size and weight of live sheep. *Computers and Electronics in Agriculture*, 103, 33-38, 2014.
- Mohammad, M. T., Rafeeq, M., Bajwa, M. A., Awan, M. A., Abbas, F., Waheed, A., Bukhari, F. A., and Akhtar, P.: Prediction of body weight from body measurements using regression tree (RT) method for indigenous sheep breeds in Balochistan, Pakistan. *The Journal of Animal Plant Science*, 22(1), 20-24, 2012.
- Momoh, O. M., Rotimi, E. A., and Dim, N. I.: Breed effect and non-genetic factors affecting growth performance of sheep in a semi-arid region of Nigeria. *Journal of Applied Biosciences*, 67, 5302–5307, 2013.
- Mosteller F., and Colditz G. A.: Understanding Research Synthesis (Meta-Analysis). *Annual Review Public Health*, 17, 1-23, 1996.
- Nimer, J., and Lundahl, B.: Animal-Assisted Therapy: A Meta-Analysis. *Anthrozoös*, 20(3), 225-238, 2007.
- Oral, H. D., and Altinel, A.: The Phenotypic Correlations Among Some Production Traits of the Hair Goats Bred on the Private Farm Conditions in Aydın Province. *J. Fac. Vet. Med. Istanbul Univ.*, 32(3), 41-52, 2006.
- Özder, M., Kaymakçı, M., Taşkın, T., Köycü, E., Karaağaç, F., and Sönmez, R.: Türkgeldi koyun tipinin gelişme ve süt verim özellikleri. *Türk J. Vet. Anim. Sci.*, 28, 195-200, 2004.
- Sowande, O. S., and Sobola, O. S.: Body measurements of West African dwarf sheep as parameters for estimation of live weight. *Tropical Animal Health and Production*, 40, 433-439, 2008.
- Şelli, M., and Doğan, Z.: Meta analiz ile tarımsal verilerin değerlendirilmesi. *Harran Üniversitesi, Ziraat Fakültesi Dergisi*, 15(4), 45-56, 2011.
- Tuck, S. L., Winqvist, C., Mota, F., Ahnstrom, J., Turnbull, L. A., and Bengtsson, J.: Land-use intensity and the effects of organic farming on biodiversity: a hierarchical meta-analysis. *Journal of Applied Ecology*, 51, 746–755, 2014.
- Üstün, U., and Eryılmaz, A.: Etkili araştırma sentezleri yapabilmek için bir araştırma yöntemi: Meta-analiz. *Eğitim ve Bilim*, 1-32, 2014.
- Yıldız N, and Tez M.: Statistical methods used to combine categorical data in a Meta-Analysis: An evaluation of active and passive smokers. *Istanbul University Journal of the School of Business Administration*, 38(2), 134-146, 2009.
- Yakubu, A.: Application of regression tree methodology in predicting the body weight of Uda sheep. *Scientific papers. Animal Science and Biotechnologies*, 45(2), 484-490, 2002.

Farklı Aktivasyon Fonksiyonlarında Yapay Sinir Ağları Çalışması: Hayvansal Üretimde Bir Uygulama

Senol ÇELİK^{1*}

¹Bingöl Üniversitesi, Ziraat Fakültesi, Zootekni Bölümü, Biyometri ve Genetik ABD, Bingöl, Türkiye

*Sunucu yazar e-posta: senolcelik@bingol.edu.tr

Özet

Bu çalışmada koyunlarda süt üretimini modellemek ve tahmin edecek bir Yapay Sinir Ağı (YSA) modeli geliştirilmiştir. YSA modelinin geliştirilmesinde girdi parametresi olarak zaman değişkeni, çıkış parametresi olarak süt üretim miktarı kullanılmıştır. Çalışmada Türkiye’de 1991-2019 yılları arası yerli ve merinos koyunlarında süt üretim miktarı verileri kullanılmıştır. Yerli ve merinos koyunlarda süt üretim modeli için iki ayrı YSA yöntemi uygulanmıştır. En iyi modeli oluşturmak için semi lineer, sigmoid, bipolar sigmoid ve hiperbolik tanjant gibi 4 farklı aktivasyon fonksiyonları kullanıldı. En uygun aktivasyon fonksiyonunu belirlemek için modelin etkinliğini belirleyen Hata Kareler Ortalaması (MSE) ve Hatanın Mutlak Ortalaması (MAE) istatistikleri kullanılmıştır. Farklı aktivasyon fonksiyonların performansları karşılaştırıldığında en uygun aktivasyon fonksiyonu her iki modelde de hiperbolik tanjant fonksiyonu olmuştur. Çünkü hiperbolik tanjant fonksiyonu kullanıldığında en küçük MSE ve MAE değerlerine ulaşılmıştır. YSA modeli tek gizli katmanlı, 10 işlem elemanlı (10-10-1) ve öğrenme algoritması olarak Levenberg–Marquardt geri yayılım algoritması (trainlm) olarak kullanılan bir ağ mimarisine sahiptir. YSA ile yerli ve merinos koyun süt üretimine ait 2020-2025 yılları arası öngörü yapılmıştır. Öngörü sonuçlarına göre, 2020-2025 yılları arasında yerli koyun süt üretim miktarının 1 472 557-1 490 614 ton arasında, merinos koyunu süt üretim miktarının ise 72535-75027 ton arasında olacağı beklenmektedir.

Anahtar kelimeler: Yapay sinir ağları, aktivasyon fonksiyonu, koyun sütü

Study of artificial neural networks in different activation functions: An application in animal production

Abstract

In this study, an Artificial Neural Network (ANN) model was developed to model and predict milk production in sheep. In the development of the ANN model, time variable was used as the input parameter and milk production amount as the output parameter. In the study, milk production amount data of domestic and merino sheep between 1991 and 2019 were used. Two different ANN methods were applied for the milk production model in domestic and merino sheep. Four different activation functions, such as semi-linear, sigmoid, bipolar sigmoid, and hyperbolic tangent, were used to create the best model. Mean Squares Error (MSE) and Absolute Mean Error (MAE) statistics, which determine the efficiency of the model, were used to determine the most appropriate activation function. When the performances of different activation functions were compared, the most appropriate activation function was the hyperbolic tangent function in both models. Because the smallest MSE and MAE values were reached when the hyperbolic tangent function was used. ANN model has a network architecture with one hidden layer, 10 processing elements (10-10-1), and Levenberg-Marquardt backpropagation algorithm as learning algorithm (trainlm). A prediction was made for domestic and merino sheep milk

production between the years 2020-2025 with ANN. According to the prediction results, it is expected that the domestic sheep milk production amount will be between 1 472 557-1 490 614 tones, and the milk production amount of merino sheep will be between 72535-75027 between the years 2020-2025.

Key words: Artificial neural networks, activation function, sheep milk

GİRİŞ

Yapay sinir ağları (YSA); insan beyninden esinlenerek, öğrenme sürecinin matematiksel olarak modellenmesi ile ortaya çıkan (Kabalcı, 2014), paralel dağıtılmış ağlar, bağlantılı ağlar, nuromorfik ağlar olarak tanımlanır (Keskenler, 2017).

YSA, doğrusal olmama, paralel çalışma, öğrenme, genelleme, eksik verilerle çalışma, hata toleransı, uyarlanabilirlik, çoklu değişken ve parametre kullanımı gibi özelliklere sahiptir. YSA modelleri tek katmanlı algılayıcılar, çok katmanlı algılayıcılar, ileri beslemeli yapay sinir ağları ve geri beslemeli yapay sinir ağları olarak dört grupta incelenir.

Bir yapay sinir ağında, birbirleriyle bağlantılı sinir hücrelerinin yer aldığı girdi katmanı (input layer), çıktı katmanı (output layer) ve gizli katmanı (hidden layer) gibi üç katman bulunur. Girdi katmanı, dışarıdan gelen verilerin yapay sinir ağına alınmasını sağlar. Girdi katmanı probleme etki eden parametrelerden oluşur ve girdi katmanındaki nöron sayısı parametre sayısına göre şekillenir. Çıktı katmanı, bilgilerin dışarıya iletilmesini sağlar. Gizli katman ise girdi katmanı ile çıktı katmanı arasında yer alır. Gizli katmanda bulunan nöronların dış ortamla bağlantıları yoktur. Yalnızca girdi katmanından gelen sinyalleri alırlar ve çıktı katmanına sinyal gönderirler (Benli, 2002). Ayrıca toplama fonksiyonu ve aktivasyon fonksiyonu da yapay sinir ağlarının diğer önemli öğeleridir (Öztemel, 2012).

YSA yöntemi ile çeşitli alanlarda yapılmış çalışmalara rastlanılmıştır. Bu çalışmaların birinde özel bir şirkette ele alınan üç ürünün satış rakamlarına YSA yöntemi uygulanarak öngörü yapılmıştır ve uygun performans değerlendirmesi ile iyi sonuçlar elde edilmiştir (Ataseven, 2013). Ergene Havzası'nda Sodyum Absorbsiyon Oranı Yapay Sinir Ağları ile tahmin edilmiştir (Arkoç ve ark., 2016). Lineer Olmayan Dinamik Sistemlerin Yapay Sinir Ağları ile Modellenmesinde çok katmanlı yapay sinir ağları ile radyal taban fonksiyonlu yapay sinir ağlarının yapılarının model performansları karşılaştırılmıştır. Aynı model üzerinde ve benzer ağ yapısında yapılan modelleme çalışmasında radyal taban fonksiyonlu yapay sinir ağlarının diğerine göre daha başarılı olduğu görülmüştür (Kılıç ve ark., 2012). Aydınlatma kavramı yapay sinir ağları ile farklı açılardan incelenerek, iç mekânlardaki aydınlık kalitesi belirlenmiştir. Klasik hesaplama yöntemiyle mümkün olmayan, ölçme cihazlarıyla uzun zaman alan ölçümler, yapay sinir ağı sayesinde kısa bir sürede, minimum hata oranı ile tahmin edilmiş ve tahminler üç boyutlu olarak modellenmiştir. Böylece aydınlatma sistemlerinin karmaşık ve uzun zaman alan bakımları daha kısa sürede ve daha güvenilir bir şekilde yapılmıştır (Şahin ve ark., 2013).

Bu çalışmada, yapay sinir ağlarının yerli ve yabancı koyun süt üretim miktarının modellenmesi ve tahmin edilmesi amaçlanmıştır.

MATERYAL VE YÖNTEM

Materyal

Çalışmanın materyalini Türkiye İstatistik Kurumu (TÜİK)'nin www.tuik.gov.tr web adresinden alınan Türkiye'de süt üretimi başlığı altında yerli ve merinos koyun üretim miktarı değerleri oluşturmıştır. Çalışmada 1991-2019 yılları arası veriler kullanılmış, yapay sinir ağları (YSA) yöntemi uygulanmıştır. Uygun modeller belirlendikten sonra 2020-2025 yılları arası yerli ve merinos koyun üretimi öngörüsü yapılmıştır.

Yöntem

Yapay sinir ağları (YSA), deneme yolu ile öğrenme ve genelleştirme yapabilmektedir. YSA'nın kullanıldığı önemli alanlardan biri de geleceği tahmindir. YSA, veriler arasındaki bilinmeyen ve fark edilmesi güç ilişkileri ortaya çıkartabilir. YSA'nın eğitilebilmesi ve hedef çıktılara ulaşılması için çok sayıda girdi ve girdilere ilişkin çıktı dizisine gereksinim duyulur. YSA, insan beyninin fonksiyonel özelliklerine benzer şekilde öğrenme, optimizasyon, analiz, sınıflandırma, geneleme ve ilişkilendirme gibi konularda başarılı olarak uygulanmaktadır (Öztemel, 2012).

YSA'nın çalışmasına esas teşkil eden en küçük birimler, yapay sinir hücresi ya da işlem elemanı olarak adlandırılır. Yapay sinir hücresi girdiler, ağırlıklar, toplama fonksiyonu, aktivasyon fonksiyonu ve çıkış olmak üzere beş ana bileşenden oluşmaktadır.

Girdiler, bir yapay sinir hücresine dış dünyadan gelen bilgilerdir. Dış dünyadan veya bir önceki katmandan alınan bilgiler giriş olarak yapay sinir hücrelerine gönderilir (Özveren, 2006).

Ağırlıklar bir yapay hücreye gelen bilginin önemini ve nöron üzerindeki etkisini gösterir (Öztemel, 2012). Ağırlıklar ($w_1, w_2, w_3, \dots, w_i$), yapay sinir tarafından alınan girişlerin sinir üzerindeki etkisini belirleyen uygun katsayılardır (Elmas, 2003). Toplama fonksiyonu bir hücreye gelen net girdiyi hesaplar. Bu fonksiyon aşağıdaki gibi formüle edilir.

$$z_i = \sum_{i=1}^n (w_{ij} x_i + b_j)$$

Burada w girdiler, x ağırlıklar, n ise girdi (proses elemanı) sayısıdır.

Toplam fonksiyonu sonucunda elde edilen değer, doğrusal ya da doğrusal olmayan türevlenebilir bir aktivasyon fonksiyonundan geçirilerek işlem elemanının çıktısı hesaplanır. Bu durum aşağıdaki gibidir (Yavuz ve Deveci, 2012).

$$y = f(z_i) = f\left(\sum_{i=1}^n (w_{ij} x_i + b_j)\right)$$

Aktivasyon fonksiyonu hücreye gelen net girdiyi işleyerek hücrenin bu girdiye karşılık üreteceği çıktıyı belirler. Aktivasyon fonksiyonu genellikle doğrusal olmayan bir fonksiyon seçilir (Çayiroğlu, 2015). Genel olarak en yaygın kullanılan aktivasyon fonksiyonları doğrusal (lineer), sigmoid, bipolar sigmoid ve hiperbolik tanjant aktivasyon fonksiyonlarıdır.

Doğrusal (lineer) aktivasyon fonksiyonu, doğrusal problemler çözmek amacıyla aktivasyon fonksiyonu doğrusal bir fonksiyon olarak seçilebilir. Toplama fonksiyonundan çıkan sonuç, belli bir katsayı ile çarpılarak hücrenin çıktısı olarak hesaplanır.

$$F(NE) = A * NE$$

Burada A , sabit bir sayıdır.

Sigmoid aktivasyon fonksiyonu sürekli ve türevi alınabilir bir fonksiyondur. Bu fonksiyon girdi değerlerinin her biri için 0 ile 1 arasında bir değer üretir ve aşağıdaki gibi ifade edilir (Çayiroğlu, 2015).

$$F(NE) = \frac{1}{1 + e^{-NE}}$$

Bipolar Sigmoid aktivasyon fonksiyonu hiyerarşik kademeli bir ağ yapısı oluşturmak için geliştirilmiş bipolar aktivasyon fonksiyonunu kullanma fırsatı sunar. Kısaca geliştirilmiş bipolar

sigmoid aktivasyonu kullanılmıştır. Bu etkinleştirme işlevi, ağırlıkların +1 ve -1 aralığında tutan bir gösterim seçerek hatayı en aza indirmek için tasarlanmıştır. Bu fonksiyon,

$$F(NET) = \frac{1 - e^{-2NET}}{1 + e^{-2NET}}$$

şeklindedir (Kaur ve Gupta, 2020).

Hiperbolik tanjant aktivasyon fonksiyonu, sigmoid fonksiyonuna benzer bir fonksiyondur. Sigmoid fonksiyonunda çıkış değerleri 0 ile 1 arasında değişirken hiperbolik tanjant fonksiyonunun çıkış değerleri -1 ile 1 arasında değişmektedir (Çayıroğlu, 2015). Bu fonksiyon,

$$F(NET) = \frac{e^{NET} - e^{-NET}}{e^{NET} + e^{-NET}}$$

şeklinde hesaplanır (Öztemel, 2012; Alp ve Öz, 2019).

Performans ölçütü

YSA model performansı genellikle Hata Kareler Ortalaması (MSE) ve Ortalama Mutlak Hata (MAE) ile tespit edilir. MSE aşağıdaki gibi hesaplanır (Singh ve ark., 2009).

$$MSE = \frac{\sum_{i=1}^n (Y_i - \hat{Y}_i)^2}{n}$$

MAE ise aşağıdaki gibi hesaplanır.

$$MAE = \frac{\sum_{i=1}^n |Y_i - \hat{Y}_i|}{n} = \frac{\sum_{i=1}^n |\varepsilon_i|}{n}$$

Burada Y_i : Bağımlı değişkenin gözlenen değerleri, $\hat{Y}_i$: Bağımlı değişkenin tahmini değerleri, n ise gözlem sayısıdır.

BULGULAR

YSA giriş, gizli ve çıktı tabakalarının sayıları sırasıyla 10-10-1 olarak belirlenmiş olup, geri yayılma öğrenimi (Back Propagation Learning) ile 1000 iterasyonlu olarak uygulanmıştır. YSA yönteminde gerek yerli koyun süt üretimi gerek se merinos koyun sütü üretimi için 4 farklı aktivasyon fonksiyonları kullanılmıştır. Bunlar semi lineer, sigmoid, bipolar sigmoid ve hiperbolik tanjant aktivasyon fonksiyonlarıdır. En uygun aktivasyon fonksiyonun kullanıldığı modeli belirlemek için MSE ve MAE istatistikleri kullanılmıştır. Yerli ve merinos koyun süt üretimi için kullanılan farklı aktivasyon fonksiyonlarına göre hesaplanan MSE ve MAE istatistikleri Tablo 1’de verilmiştir.

Tablo 1. Aktivasyon fonksiyonlarında hesaplanan MSE ve MAE değerleri

Aktivasyon fonksiyonu	Yerli		Merinos	
	MSE	MAE	MSE	MAE
Yarı lineer fonksiyon	8 820 508 975	74 945	41 875 609	5540
Sigmoid fonksiyon	58 311 030 255	21 2646	296 634 471	15100
Bipolar sigmoid fonksiyon	4 015 476 076	48 390	14 539 513	3034
Hiperbolik tanjant fonksiyonu	2 808 416 786	42 606	8 939 994	2407

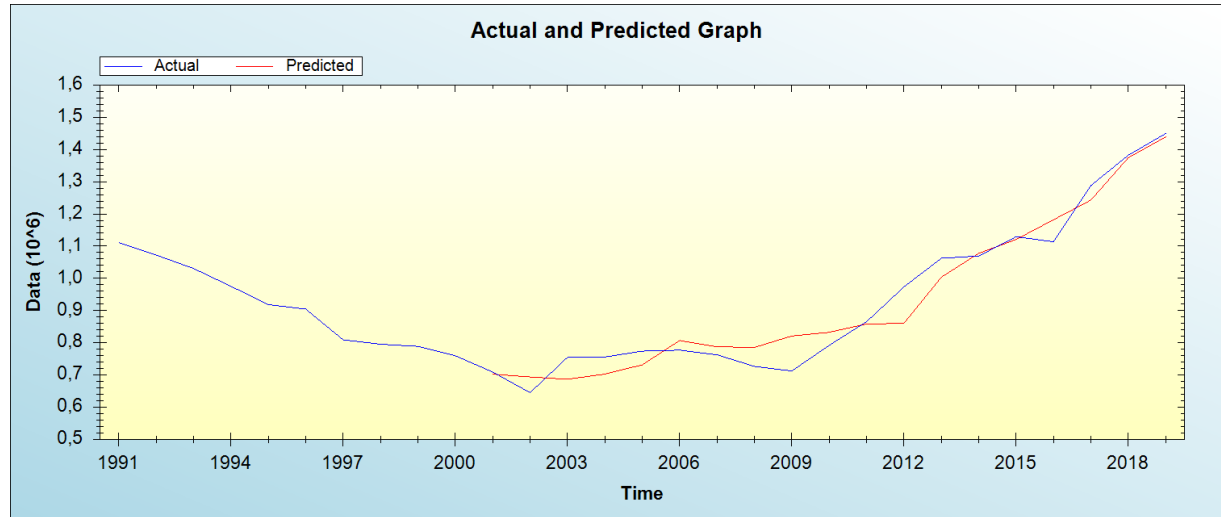
Yerli koyun sütü üretim modellemesi için YSA yönteminde kullanılan aktivasyon fonksiyonlarında en düşük MSE ve MAE değeri hiperbolik tanjant aktivasyon fonksiyonunda elde edilmiştir. MSE=2 808 416 786 ve MAE=42 606 bulunmuştur (Tablo 1). Dolayısıyla hiperbolik tanjant aktivasyon fonksiyonu kullanılarak YSA analizi yapılmıştır.

YSA yöntemi sonucunda tahmin edilen ve gözlenen değerlerle birlikte hata terimleri değerleri Tablo 2’de sunulmuştur.

Tablo 2. Gözlenen, tahmini ve hata terimleri (residual)

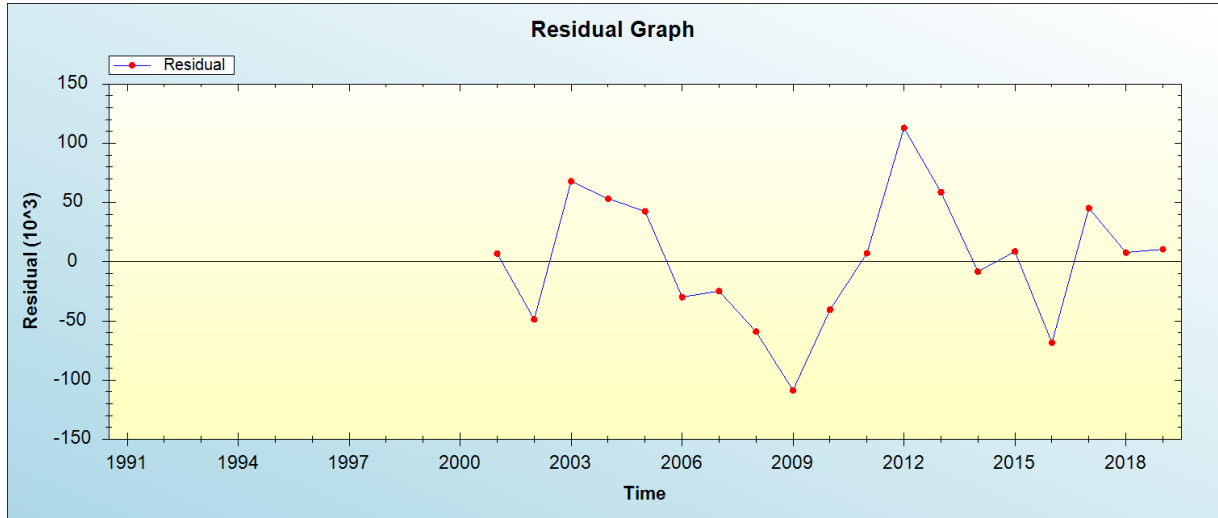
Dönem	Gözlenen	Tahmin Değeri	Hata terimi
2005	774344.0000	731760.0107	42583.9893
2006	777385.0000	807212.4217	-29827.4217
2007	762930.0000	787737.4559	-24807.4559
2008	726894.0000	785879.6085	-58985.6085
2009	712784.0000	821327.3218	-108543.3218
2010	792122.0000	832708.8385	-40586.8385
2011	865577.0000	858412.7468	7164.2532
2012	973619.0000	860683.2364	112935.7636
2013	1062274.0000	1003589.2779	58684.7221
2014	1069441.0000	1077746.4328	-8305.4328
2015	1129237.0000	1120685.6610	8551.3390
2016	1113469.0000	1181859.7501	-68390.7501
2017	1288041.0000	1242787.1432	45253.8568
2018	1382026.0000	1374270.5030	7755.4970
2019	1449351.0000	1438834.4784	10516.5216

Yerli koyun süt üretim miktarı tahmini için YSA uygulaması sonucu gerçek ve tahmini değerlerin seyri ve dağılımı grafiği Şekil 1’de verilmiştir.



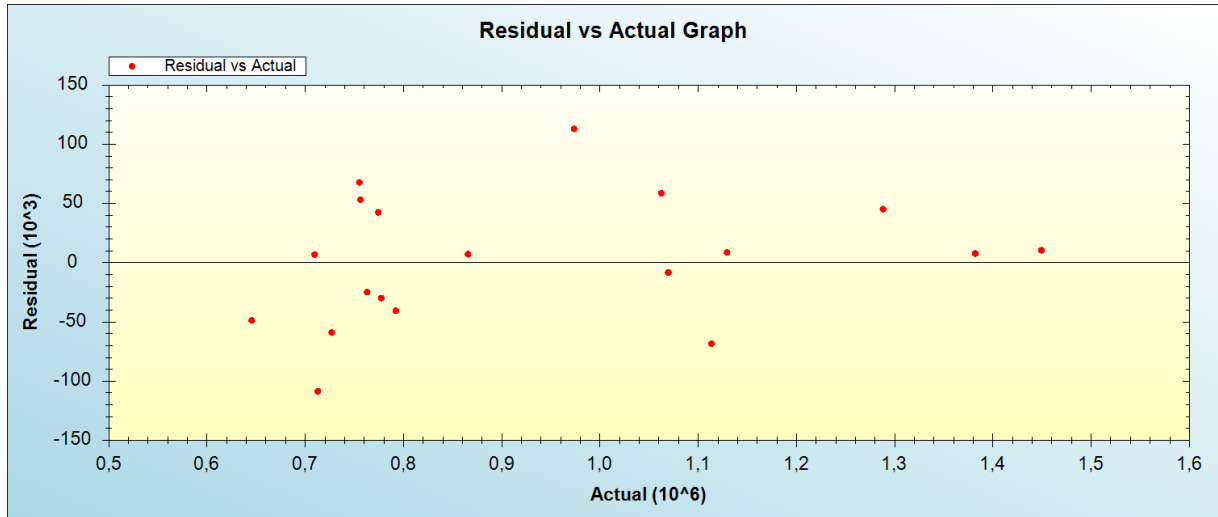
Şekil 1. Gözlenen ve tahmin edilen değerlerin grafiği

YSA uygulaması sonucu elde edilen hata terimlerinin grafiği Şekil 2’te sunulmuştur.



Şekil 2. Hata terimleri grafiği

Şekil 2’de hata terimlerinin rasgele dağıldıkları görülmüştür. Yerli koyun süt üretimi miktarı gerçek değerleri ile hata terimlerinin grafiği Şekil 3’te verilmiştir. Gerçek değerler ile hata terimleri birbirinden bağımsız olup rasgele dağılmışlardır.

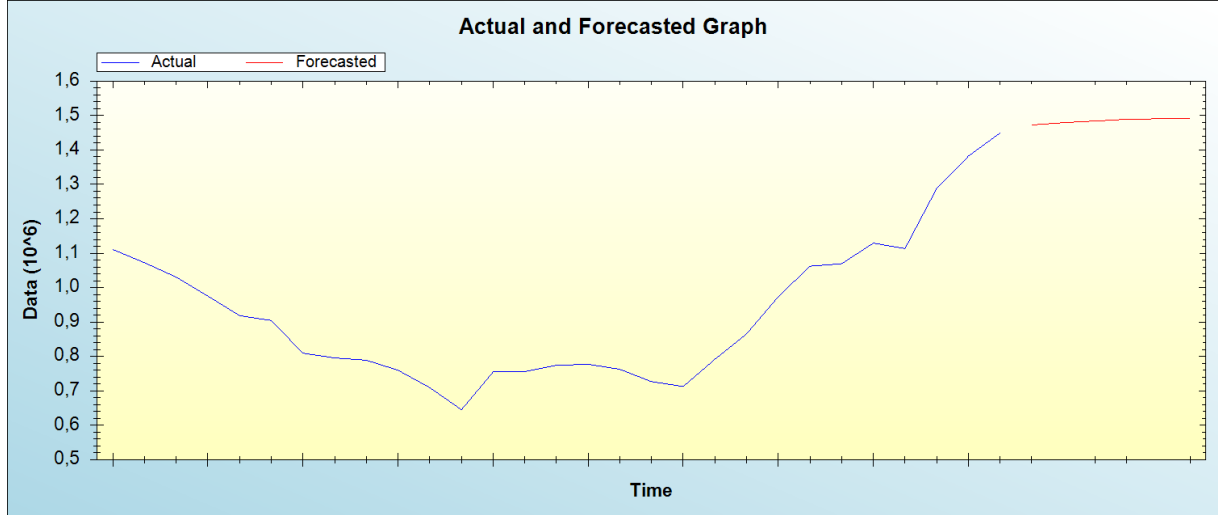


Şekil 3. Gerçek değerler ve hata terimlerinin grafiği

Bu aşamadan sonra 2020-2025 yılları arası yerli koyun süt üretim miktarı (ton) öngörüsü Tablo 3’de ve Şekil 4’te verilmiştir.

Tablo 3. Gelecek dönem için yerli koyun süt üretimi öngörüsü

Yıllar	Koyun sütü üretim miktarı (yerli)
2020	1472557
2021	1479281
2022	1483965
2023	1489177
2024	1490598
2025	1490614



Şekil 4. Gelecek dönem için gerçekleşen değerler ve öngörü

Benzer şekilde merinos koyun üretim miktarı modellenmesi de YSA yöntemi ile yapılmış olup aşağıdaki gibi özetlenmiştir.

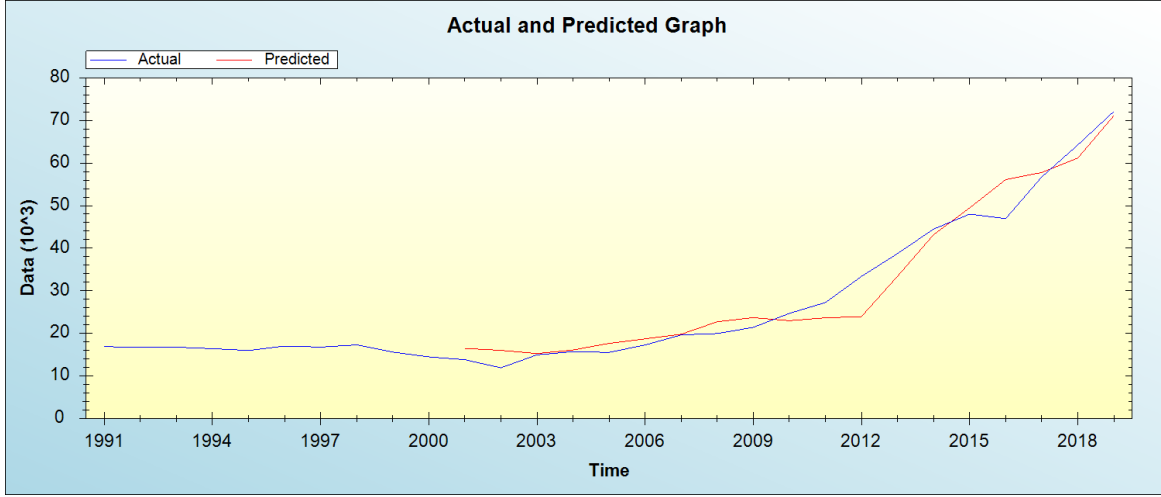
Merinos koyun sütü üretim modellenmesinde YSA yönteminde kullanılan aktivasyon fonksiyonlarında en düşük MSE ve MAE değerleri hiperbolik tanjant aktivasyon fonksiyonunda bulunmuştur. MSE=8 939 994 ve MAE=2407 bulunmuştur (Tablo 1). Bu nedenle hiperbolik tanjant aktivasyon fonksiyonu kullanılarak YSA analizi yapılmıştır.

Merinos koyunlarda süt üretim tahmini ile ilgili YSA yöntemi sonucunda tahmin edilen ve gözlenen değerlerle birlikte hata terimleri değerleri Tablo 4’de sunulmuştur.

Tablo 4. Gözlenen, tahmini ve hata terimleri

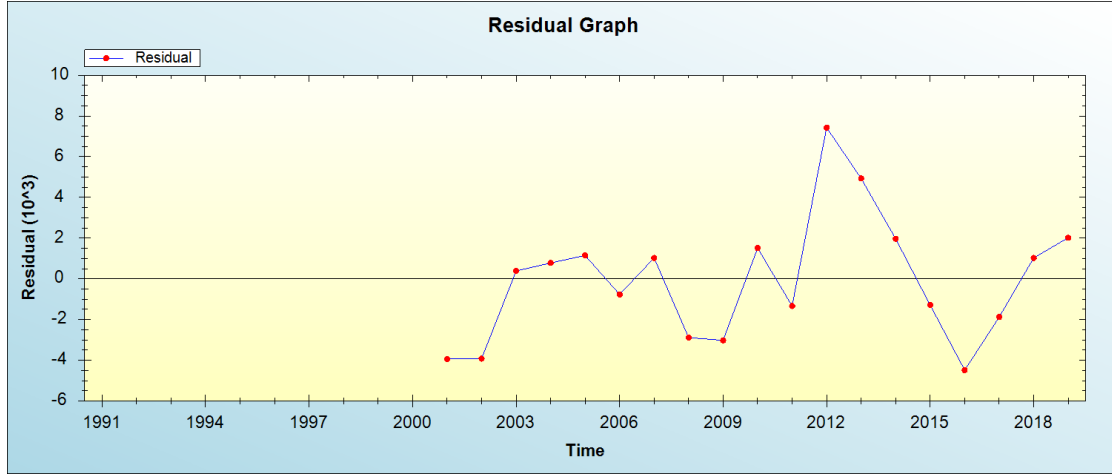
Dönem	Gözlenen	Tahmin değeri	Hata terimi
2005	15533.0000	14383.0717	1149.9283
2006	17296.0000	18060.9859	-764.9859
2007	19657.0000	18629.4729	1027.5271
2008	19978.0000	22862.0512	-2884.0512
2009	21435.0000	24456.2659	-3021.2659
2010	24710.0000	23199.6644	1510.3356
2011	27245.0000	28577.5867	-1332.5867
2012	33388.0000	25959.5098	7428.4902
2013	38739.0000	33806.3993	4932.6007
2014	44496.0000	42523.2964	1972.7036
2015	47990.0000	49274.8269	-1284.8269
2016	46943.0000	51426.0506	-4483.0506
2017	56738.0000	58611.7672	-1873.7672
2018	64245.0000	63216.3272	1028.6728
2019	72105.0000	70086.1672	2018.8328

Merinos koyun süt üretim miktarı tahmini için YSA uygulaması sonucu gerçek ve tahmini değerlerin dağılımının grafiği Şekil 5’de verilmiştir.



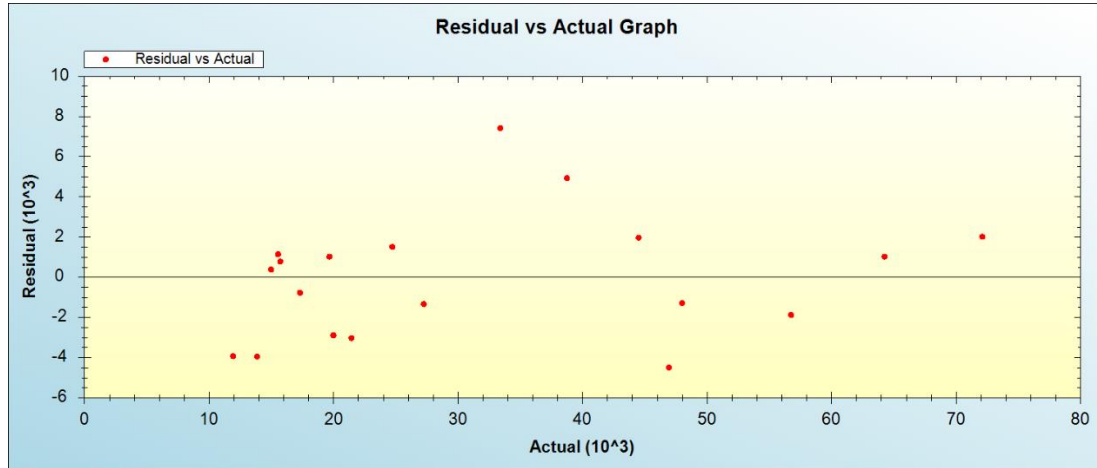
Şekil 5. Merinos koyunlarında süt üretiminin gerçek ve tahmini değerleri

Merinos koyun süt üretim miktarı modellemesinde YSA uygulaması sonucu elde edilen hata terimlerinin grafiği Şekil 6'da sunulmuştur.



Şekil 6. Hata terimleri grafiği (Merinos koyun süt üretimi için)

Şekil 6'da görüldüğü gibi hata terimleri rasgele dağılmıştır. Merinos koyun süt üretimi miktarı gerçek değerleri ile hata terimlerinin grafiği Şekil 7'de verilmiştir. Gerçek değerler ile hata terimleri birbirinden bağımsız olup rasgele dağılmışlardır.



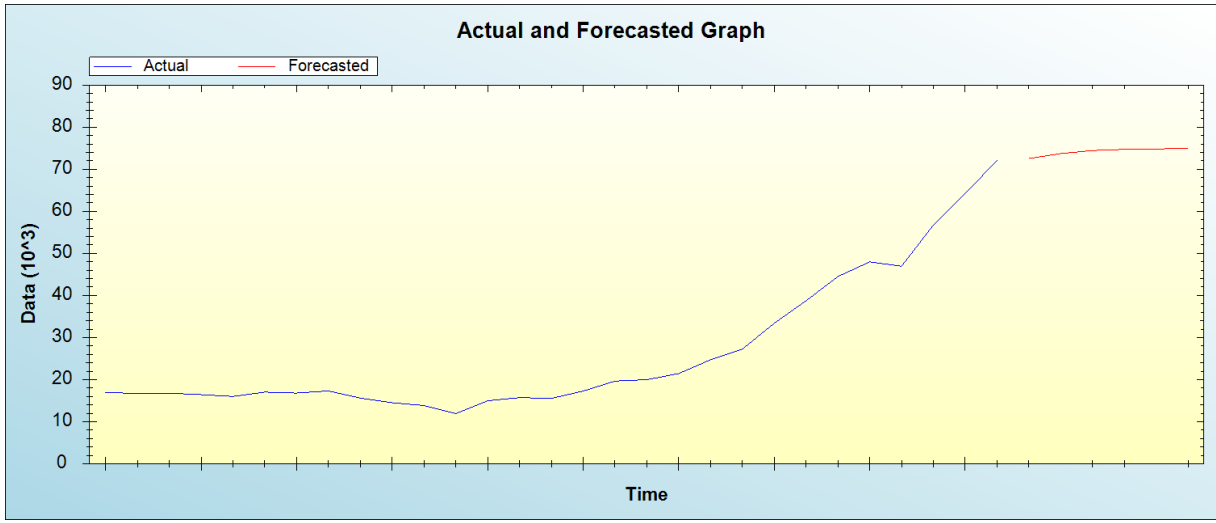
Şekil 7. Merinos koyunlarında süt üretiminde gözlenen değerler ve hata terimleri

2020-2025 yılları arası olması beklenen Merinos koyu sütü üretim miktarı öngörüsü Tablo 5’de verilmiştir.

Tablo 5. 2020-2025 dönemi öngörüsü

Yıllar	Koyun sütü üretim miktarı (merinos)
2020	72535
2021	73715
2022	74435
2023	74728
2024	74813
2025	75027

Tablo 5’de görüldüğü gibi Merinos koyunlarında süt üretim miktarı artış eğilimindedir. 2025 yılında 75027 ton süt üretimi olacağı beklenmektedir. Bununla ilgili gözlenen ve ilgili dönemdeki öngörüü gösteren grafik Şekil 8’de gösterilmiştir.



Şekil 8. Merinos koyunların gerçek süt üretim miktarı ve gelecek dönem öngörüsü

SONUÇ

Bu çalışmada yapay sinir ağları ile Türkiye’de yerli ve Merinos koyunlarda süt üretim miktarı modellenmiştir. Girdi değişkeni olarak yıllar (1961-2019) olup, çıktı değişkeni olarak ise süt üretim miktarı değerleri kullanılmıştır. Aktivasyon fonksiyonlarından semi lineer, sigmoid, bipolar sigmoid ve hiperbolik tanjant fonksiyonları kullanılmıştır ve bunlar karşılaştırılmıştır. En uygun aktivasyon fonksiyonu ise hem yerli koyun hem de merinos koyun sütü üretim miktarı modellerinde hiperbolik tanjant fonksiyonu olmuştur. Çünkü bu aktivasyon fonksiyonunda en küçük MSE ve MA değerleri elde edilmiştir.

Sonraki aşamada ağı eğitimi, test ve doğrulama işlemleri yapılmış ve tahmin işlemi gerçekleştirilmiştir. Elde edilen sonuçlar, kurulan YSA yönteminin iyi tahminler verdiğini ortaya koymuştur. Eğitim, test ve doğrulama aşamalarındaki düşük MSE ve MAE değerleri de bunu göstermektedir.

Türkiye’de 2020-2025 yılları arasında yerli koyun süt miktarının 1 472 557-1 490 614 ton arasında, merinos koyunu süt miktarının ise 72 535-75 027 ton arasında olacağı öngörülmüştür.

Süt üreticilerinin bu hususu dikkate almalarında yarar vardır. Genel olarak yapay sinir ağlarının mevcut verileri tahmin etmede iyi sonuçlar verdiği görülmüştür. Geleceğe ait tahmin çalışmalarında

yapay sinir ağıları ve alternatif teknikleri kullanarak hayvansal üretimde iyi sonuçlar vereceği beklenmektedir.

Kaynaklar

- Alp, S., Öz, E. 2019. Makine Öğrenmesinde Sınıflandırma Yöntemleri ve R Uygulamaları. Nobel Akademik Yayıncılık, Ankara.
- Arkoç, O., Akıncı, T. Ç., Nogay, H. S. 2016. Yapay Sinir Ağları Yardımı ile Yeraltı Suyunda Sodyum Absorbsiyon Oranı (SAR) Tahmini: Ergene Havzası Doğu Akiferi Örneği. Jeoloji Mühendisliği Dergisi, 40(2):177-188.
- Ataseven, B. 2013. Yapay Sinir Ağları ile Öngörü Modellemesi. Öneri Dergisi, 10(39):101-115.
- Benli, Y. 2002. Finansal Başarısızlığın Tahmininde Yapay Sinir Ağı Kullanımı ve İMKB’de Bir Uygulama. Muhasebe Bilim Dünyası Dergisi, 4(4): 17-30.
- Çayıroğlu, İ. (2015). İleri Algoritma Analizi-5 Yapay Sinir Ağları.
<http://www.ibrahimcayiroglu.com/Dokumanlar/IleriAlgoritmaAnalizi/IleriAlgoritmaAnalizi-5.Hafta-YapaySinirAglari.pdf>
- Elmas, Ç. 2003. Yapay Sinir Ağları, Birinci Baskı, Ankara: Seçkin Yayıncılık
- Kabalıcı, E. 2014. Yapay Sinir Ağları. Ders Notları.
<https://ekblc.files.wordpress.com/2013/09/ysa.pdf>
- Kaur, J., Gupta, N. 2020. Bipolar Sigmoid Algorithm for Designing Constructive Neural Network. International Journal on Emerging Technologies, 11(2):991–996
- Keskenler, M. F., Keskenler, E. F. 2017. Geçmişten günümüze Yapay Sinir Ağları ve tarihçesi. Takvim-i Vekayi, 5(2):8-18
- Kılıç, E., Özbacı, Ü., Özçalık, H. R. 2012. Lineer Olmayan Dinamik Sistemlerin Yapay Sinir Ağları ile Modellenmesinde MLP ve RBF Yapılarının Karşılaştırılması. ELECO '2012 Elektrik Elektronik ve Bilgisayar Mühendisliği Sempozyumu, 29 Kasım-01 Aralık 2012, Bursa.
- Öztemel, E. 2012. Yapay Sinir Ağları. Papatya Yayıncılık, İstanbul.
- Özveren, U. 2006. Pem Yakıt Hücrelerinin Yapay Sinir Ağları İle Modellenmesi. Yayınlanmamış Yüksek Lisans Tezi, İstanbul: Yıldız Teknik Üniversitesi Fen Bilimleri Enstitüsü
- Singh, K. P., Basant, A., Malik, A., Jain, G. 2009. Artificial neural network modeling of the river water quality—A case study. Ecological Modelling, 220(6): 888-895.
- Şahin, M., Büyüktürk, F., Oğuz, Y. 2013. Yapay Sinir Ağları ile Aydınlik Kalitesi Kontrolü. Afyon Kocatepe Üniversitesi Fen ve Mühendislik Bilimleri Dergisi, 13:1-10.
- TÜİK, 2019. Hayvansal Üretim İstatistikleri. Tür ve ırklarına göre sağılan hayvan sayısı ve süt üretim miktarı. Türkiye İstatistik Kurumu, www.tuik.gov.tr
- Yavuz, S., Deveci, M. 2012. İstatiksel Normalizasyon Tekniklerinin Yapay Sinir Ağı Performansına Etkisi. Erciyes Üniversitesi İktisadi ve İdari Bilimler Fakültesi Dergisi, 40:167-187

Regression Analysis of Morphological Traits on Body Weight in Female Dorper Sheep Lambs

Lebelo Joyceline SELALA^{1*}, Thobela Louis TYASI¹

¹School of Agricultural and Environmental Sciences, Department of Agricultural Economics and Animal Production, University of Limpopo, Private Bag X1106, Sovenga 0727, Limpopo, South Africa

*Presenter author e-mail: lebeloselala1997@gmail.com

Abstract

Regression is a very influential statistical technique that uses independent variables to describe the dependent variable. The present study was performed to determine the best fitted model from different regression techniques to predict body weight (BW) from morphological traits viz. heart girth (HG), rump height (RH), body length (BL), withers height (WH) and sternum height (SH). Total of twenty-eight female Dorper sheep lambs at birth were used for data collection. The simple and multiple regression were used for data analysis. The coefficients of determination (R^2) and mean square error (MSE) were used to determine the best-fitted regression model. The results indicated that in simple regression the best fitted regression model for estimation of BW in female Dorper lamb was the model including BL ($R^2 = 0.79$, $MSE = 1.43$), in multiple regression, the best fitted regression model for estimation of body weight in female Dorper lamb was model including HG, RH, BL, WH, SH ($R^2 = 0.89$, $MSE = 0.90$) in the study. The results of simple regression suggest that BL can truly estimate the body weight in female Dorper sheep lambs. Multiple regression findings suggest that two or more morphological traits can truly estimate body weight in the female Dorper lambs. The study will help the farmers to accurately predict the body weight of the Dorper sheep lambs using the morphological traits.
Key words: Coefficients of determination, Mean square error, Heart girth, Sternum height, Body length

INTRODUCTION

Assessing different live body weight of different sexes of animals help in improving body weight of the animals (Karim and Karym, 2018). According to Melesse et al. (2013) the regression equations are mostly used as techniques for predicting body weight using body dimensions of many breeds. Body dimensions namely: heart girth, rump height, ear length, head length, body length, sternum height and withers height are considered to be the easiest way of predicting body weight of animals as compared to the use of weighing scale because are not easily accessible and are expensive to most communal farmers (Taye et al., 2016). However, there are scarce readings on comparison of different regression techniques for estimation of body weight using morphological traits in female Dorper sheep lambs. The objective of the study is to determine the best fitted regression models to estimate body weight using morphological traits amongst two regression techniques namely: simple and multiple regressions. The current paper will assist the farmers to know which morphological traits contribute more variation to the body weight of female Dorper sheep lambs which will help farmers to improve strategies used in the breeding programs.

MATERIALS AND METHODS

The study was conducted at the University of Limpopo experimental farm, South Africa. Twenty-eight female Dorper sheep lambs at birth were used. Sick female Dorper lambs were excluded during data collection for accurate results. The body weight (BW) and morphological traits viz. heart girth (HG), rump height (RH), body length (BL), withers height (WH) and sternum height (SH) of female Dorper

lambs were recorded as described by Raja et al. (2013) using weighing scale (kg) and measuring tape (cm). The data was analyzed using Statistical Package for Social Sciences (SPSS, 2019) software, version 26 for windows. The data of descriptive of female Dorper sheep lambs were obtained. Live body weight was regressed on morphological traits using simple regression and multiple regression as it explained by Tropal et al., (2003) to design the best fitted regression model. Below is the simple and multiple regression equation that was formed:

$$Y = a + b_1X_1 + b_2X_2 + b_3X_3 + b_4X_4 + b_5X_5$$

Where;

Y= dependent variable (Body weight),

$X_1 - X_5$ = Independent variables (Morphological traits),

$b_1 - b_5$ = Coefficient of regression and

a = Constant (regression intercept).

RESULTS

Table 1 shows the results of descriptive statistics of female Dorper sheep lambs between body weight and morphological traits viz heart girth (HG), rump height (RH), body length (BL), withers height (WH), and sternum height (SH). The results indicates that all the morphological traits had higher numerical mean values than the body weight of female Dorper lambs which are indicated as follows HG (40.83 ± 1.13), RH (33.67 ± 0.94), BL (37.33 ± 1.24) and WH (34.56 ± 0.76), and SH (26.71 ± 0.50) respectively, while the BW had (6.05 ± 0.49) which is lower than all the morphological traits.

Table 1: Descriptive statistics of female Dorper sheep lambs between BW and morphological traits

Traits	Female lambs	n = 28
	Mean $\pm$ SE	CV (%)
BW (kg)	6.05 ± 0.49	8.10%
HG (cm)	40.83 ± 1.13	2.77%
RH (cm)	33.67 ± 0.94	2.79%
BL (cm)	37.33 ± 1.24	3.32%
WH (cm)	34.56 ± 0.76	2.20%
SH (cm)	26.71 ± 0.50	1.87%

SE: standard error; CV: coefficient of variance; BW: body weight; WH: withers height; RH: rump height; BL: Body length; BD: Body depth; RW: rump width; RL: rump length; HG: heart girth; n: 28.

Simple regression model on female Dorper lambs using morphological traits on body weight

The model of simple linear regression was formed using dependent variable (BW) and independent variables (HG, RH, BL, WH and SH) to determine the regression model that is the best fitted for estimation of body weight in the female Dorper lambs. The findings indicate that all the equation models were statistical significant ($p < 0.05$) showed in Table 2 below. The best fitted regression model with BL has considered to be the best because of highest R^2 and lowest MSE of 0.79 and 1.43 respectively. The results indicate that 79 per cent of variation in the body weight is contributed by BL.

Table 2: Simple linear equations for estimation of body weight using different morphological traits

Model of female Dorper lambs sheep	Variables	Equations	R ² values	MSE
1	HG	BW = -8.17 + 0.35HG	0.65	2.42
2	RH	BW = -7.52 + 0.40RH	0.60	2.76
3	BL	BW = -7.09 + 0.35BL	0.79	1.43
4	WH	BW = -13.42 + 0.56WH	0.76	1.67
5	SH	BW = -10.09 + 0.60SH	0.39	4.26

R²: coefficient of determination; MSE: mean square error; BW: body weight; HG: heart girth; RH: rump height; BL: body length; WH: withers height; SH: sternum height; EL: ear length and HL: head length.

Multiple regression model on the female Dorper lambs using morphological traits on body weight

Table 3 shows the summary of the multiple regression analysis in female Dorper sheep lambs. The results revealed that the SH had highest variation contribution of 0.09 followed BL with variation contribution of 0.05 ($p < 0.05$) to the body weight with coefficients of determination of 0.89 and mean square error of 0.90 in female Dorper lambs. The results indicates that 89 per cent of the variation in the BW of the female Dorper lambs is described by this model. The equation of multiple regression was developed as followed: $BW = -14.72 + 0.53HG + 0.10RH + 0.16BL + 0.12WH + 0.20SH$, whereby HG, RH and WH were not statistically significant ($p > 0.05$).

Table 3: Multiple regression for estimation of BW using morphological traits

Multiple regression	Morphological traits				
Parameters	HG	RH	BL	WH	SH
P value	0.38	0.15	0.01	0.27	0.04
B	0.05	0.10	0.16	0.12	0.20
SE	0.06	0.07	0.05	0.11	0.09

R² = 0.89, MSE = 0.90, constant = -14.723

HG = heart girth, RH = rump height, BL = body length, WH = withers height, SH = sternum height, EL = ear length, HL = head length, SE = Standard error and R² = coefficient of determination, b = regression coefficient.

DISCUSSION

According to Taye et al. (2016) body dimensions are perfect instruments used to assess the body weight of the animals. The results of descriptive statistics indicated that the numerical mean value of BW is lower than the rest of morphological traits. The findings are dissimilar with the study of Mathapo and Tyasi. (2021) in yearling Boar goats. The different may be due to the use of different breeds. The simple regression was initially used to determine the best-fitted model in the estimation of body weight using morphological traits viz. HG, RH, BL, WH and SH in female Dorper lambs sheep. The highest R² and the lowest MSE was achieved from body length in female Dorper lamb sheep. The results indicate that body length is the one that pay more to the variation in body weight of female Dorper lamb sheep. The study is in disagreement with the study of Tropal et al. (2003) in Awassi sheep; Raja et al. (2013) in atteppady black goats. Secondly multiple regression was used to determine which morphological traits contribute more to variation of body weight in female Dorper lambs. The study is in oppose with the study of Tyasi et al. (2020) in Nguni cattle.

CONCLUSION

The results of the current study suggest that body length is the most important body dimensions that can evaluate body weight using simple regression and again, the sternum height and body length contribute more variation to the body weight of female Dorper lambs using multiple regression. The sternum height and body length can be included in the selection criteria of to improve body weight of female Dorper

lambs. Therefore, the study will benefit Dorper sheep farmers to assess body weight using model of regression that is established. Additional readings are required for increasing body weight in Dorper sheep using simple and multiple regression.

References

- Karim, G.M. and Karym, C.M., 2018. Prediction of body weight from body dimensions in Karadi sheep. *Journal of Zankoy Sulaimani*. JZS, 2ndInt. Conference of Agricultural Science.
- Mathapo, M.C. and Tyasi, T.L., 2021. Prediction of Body Weight of Yearling Boer Goats from Morphometric Traits using Classification and Regression Tree. *American Journal of Animal and Veterinary Sciences*. 16 (2): 130-135.
- Melesse, A., Banerjee, S., Lakew, A., Mersha, F., Hailemariam, F., Tsegaye, S., Makebo, T., 2013. Variations in linear body measurements and establishing prediction equations for live weight of indigenous sheep populations of southern Ethiopia. *Journal of Animal Science*. 2(1):15-25.
- Raja, T.V., Venkatachalapathy, R.T., Kannan, A., Bindu, K.A., 2013. Determination of best-fitted regression model for prediction of body weight in attappady black goats. *Global Journal of Animal Breeding and Genetics*. 1 (1): 020-025.
- Taye, M., Yilma, M., Rischkowsky, B., Dessie, T., Okeyo, M., Mekuriaw, G., Haile, A., 2016. Morphological characteristics and linear body measurements of Doyogena sheep in Doyogena district of SNNPR, Ethiopia. *African Journal of Agricultural Research*. 11(48): 4873-4885.
- Tropal, M., Yildiz, N., Esenbuba, N., Aksakal, V., Macit, M., Ozdemir, M. 2003. Determination of best fitted regression model for estimation of body weight in Awassi sheep. *Journal of Applied Animal Research*. 23: 201-208.
- Tyasi, T.L., Mathye, N.D., Danguru, L.W., Rashijane, L.T., Mokoena, K., Makgowa, K.M., Mathapo, M.C., Molabe, K.M., Bopape, P.M., Maluleke, D., 2020. Correlation and path analysis of body weight and biometric traits of Nguni cattle breed. *Journal of Advanced Veterinary and Animal Research*. 7(1):148-55.

Acknowledgement

The writers would like to thank the farmworker of University of Limpopo during data collection for their support and University of Limpopo, Department of Agricultural Economics and Animal Production for financial involvement during this study.

Conflict of interest

There is no conflict of interest stated by authors.

Classification of High Dimensional Imbalanced Data Using PCA, SMOTE Methods and Classification Algorithms

Guhdar A. A. MULLA^{1*}, Yıldırım DEMİR²

¹Van Yuzuncu Yil University, Institute of Natural and Applied Sciences, Department of Statistics, 65080, Van, Turkey

²Van Yuzuncu Yil University, Faculty of Economics and Administrative Sciences, Department of Statistics, 65080, Van, Turkey

*Presenter author e-mail: guhdar.abdulaziz@gmail.com

Abstract

Data mining is a technique for extracting possible trends and correlations from data obtained from various sources in order to uncover secret information. For data visualization, high-dimensionality statistics and dimensionality reduction methods are often used. If the classification groups are not roughly evenly represented, the dataset is imbalanced. In this study, we used SMOTE method to rebalanced data set with Principal Component Analysis (PCA) is proposed to solve the high dimensional data. Four classification algorithms are Linear Discriminant Analysis (LDA), Artificial Neural Network (ANN), Stochastic Gradient Descent (SGD) and Random Forest Analysis (RFA). In this study we used Cancer dataset were used to check the efficiency of the proposed method and select the result of the classifiers. Respectively, raw datasets, converted datasets by PCA, SMOTE methods, were analyzed with the given algorithms. Analyzes were made using WEKA.

Key words: Classification, Imbalance, Principal component analysis

INTRODUCTION

Data mining is the process of looking through vast amounts of data or archives for patterns and then using those patterns to forecast future events. One of the data mining techniques for categorizing a specific category of objects into targeted categories is classification. The main aim of classification is to predict the disposition of an object or data based on the groups of objects that are available (Adebayo and Chaubey, 2019). There are several different algorithms for classification in the literature. Linear Discriminant Analysis (LDA), Artificial Neural Network (ANN), Stochastic Gradient Descent (SGD) and Random Forest Analysis (RFA) are the most important and commonly machine learning algorithms for classification process. The classification of imbalanced data is one of the most complex problems encountered in classification. Because problems arise in binary grouping when there are many instances of one class (majority) and a small number of instances of the other (minority). However, if there is no disparity between samples from the positive and negative classes, this problem may not be so serious (Al-Rousan et al., 2012). In this paper, it is aimed to solve the classification problem for high-dimensional imbalanced data. In the study, the classification problem was investigated with simultaneous rebalancing and dimension reduction methods (using PCA and SMOTE methods). In accordance with this purpose, four well-known classification algorithms (LDA, ANN, SGD and RFA) were used for different imbalanced datasets where the number of samples in one class (majority) is substantially higher than the number of samples in the other class (minority).

MATERIALS AND METHODS

Machine learning deals with different types of learning, and classification algorithms in detail. In this paper discusses LDA, ANN, SGD, and RFA data mining and machine learning algorithms to deal with the class imbalance problem in high dimensional data. In addition, PCA and SMOTE methods were used.

In order to investigate the effectiveness and reliability of the proposed methods, the methods used were applied to the Cancer disease dataset. The dataset was selected based on the imbalanced percentage values between negative (37.25%) and positive (62.75%) class; the imbalance between negative and positive classes is 25.5%. The Wisconsin breast cancer diagnostic dataset consists of 357 benign and 212 malignant cases. Also, the number of variables in the dataset is 32. This data This dataset was taken from <https://www.kaggle.com/uciml/breast-cancer-wisconsin-data.on> 12.06.2020.

The imbalanced data set is balanced with the SMOTE method and then the dimension is reduced by the PCA method. Of the rebalanced and the dimension reduction dataset, 30% was used as test data and 70% as training data. By the end of this process, the data set was classified using four classification algorithms (LDA, ANN, SGD, RFA). Accuracy, Precision and ROC area measurements were used as evaluation measures. Using tenfold cross-validation for evaluation, the information was automatically divided into ten equal parts, testing and training procedure was carried out ten times. After the data are prepared for classification and evaluation, the method appoints each test record to the most likely class. In the last stage, the classification performances before and after the PCA and SMOTE processes were compared separately to determine whether there was an improvement in the performance of the classification models and the efficiency of the methods.

Classification

The classification and prediction tasks deal with predicting the value of one field (the target) based on the values of other fields (attributes or features). The assigned assignment is called classification whether the objective is discrete (e.g. nominal or ordinal). Classification is typically a supervised process in which the model learns to correctly identify new unseen instances based on a previously correctly identified collection of training instances. Predicting whether not to award a credit to a customer is an example of a classification challenge. So, a set of yes or no reflecting a positive and negative judgment respectively, could shape the values of the class C in in this problem. Details about a consumer will be the input to the classification system (that is, for a classifier). If the theory space is made up of rules, the output could be made up of a series of learned rules. (Ławrynowicz and Tresp, 2014).

The Random Forest Algorithm (RFA)

Random Forest, as the name suggested, is a tree-based ensemble in which each tree is dependent on a set of random variables. Assuming that there is an undefined co-distribution $P_{XY}(X, Y)$, for a p-dimensional random vector $X = (X_1, \dots, X_p)^T$ representing the real-valued input or predictor variables and a random variable Y representing the real-valued response. The aim is to find $f(X)$, which is a prediction function to predict Y. The prediction function is determined by the loss function $L(Y, f(X))$ and is defined to minimize the expected value of equation 1.

$$E_{XY} (L(Y, f(X))) \quad (1)$$

where the subscripts denote expectation with respect to the co-distribution of X and Y (Cutler et al., 2012).

Stochastic Gradient Descent (SGD)

The stochastic gradient descent (SGD) algorithm is a significant reduction in complexity. Rather than computing the gradient of $E_n(f_w)$ fully, each iteration estimates this gradient to the basis of a single randomly chosen example z_t :

$$w_{t+1} = w_t - \gamma_t \nabla_w Q(z_t, w_t) \quad (2)$$

The stochastic process $\{w_t, t = 1, \dots\}$ depends on the samples randomly selected at each iteration (Pham et al., 2016). The stochastic algorithm can process examples on the fly in a deployed system because it does not need to recall which examples were visited during previous iterations. In this case, as the examples are randomly attracted from the basis truth distribution, stochastic gradient descent directly optimizes the inevitable risk (Bottou, 2012).

Artificial Neural Network (ANN)

The first artificial neural network was created using a very basic concept of neural connections. The neurobiologist Frank Rosenblatt (1957), inventing the "Mark I Perceptron" machine, suggested that the mismatch between the real and predicted performance of the artificial connections between neurons could be reduced by a supervised learning mechanism. The mismatch between the real and predicted responses of the network contains crucial information for optimizing learning outcomes. Predicted performance is obtained using training data (Dell'Aversana, 2019).

Linear Discriminant Analysis (LDA)

Linear Discriminant Analysis finds an ideal set of discriminant projection vectors W , to map the original data space onto a lower dimensional feature space, by maximizing the fisher criterion $J(W)$ which ensure that the overlap between the classes in lower dimensional feature space is minimum. (Qin et al., 2005). For example, consider two classes in a two-dimensional feature space. W, W' , which reflects the distribution of classes and shows two spaces, is shown in figure 1. Here W' indicates significantly class overlap in the projected space, while the W projection shows greatly improved class separation. Thus, LDA is described for dimensionality reduction. Assume that $X = \{x_1, x_2, \dots, x_p\}$ is a data set of p -dimensional vectors. Each data point is associated with one of the C object classes $\{X_1, X_2, \dots, X_i, \dots, X_C\}$ (Shashoa et al., 2016). LDA, which is a function of W , is given by equation (3).

$$J(W) = \frac{|W^T S_B W|}{|W^T S_W W|} \quad (3)$$

Here, W must be chosen such that it maximizes $J(W)$. S_B shows the between-class distribution matrix, while S_W shows the within-class distribution matrix (Poston and Marchette, 1998).

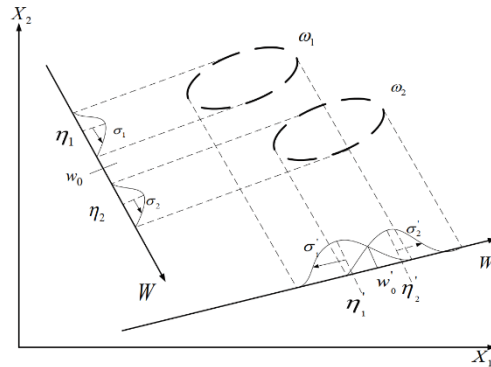


Figure 1. An example linear mapping (Shashoa et al., 2016).

Confusion matrix

The representation of the following parameters in the form of a simple matrix is expressed as the Confusion matrix. Confusion matrix, in the classical binary classification problem, the classifier units are labeled as positive and negative, and the matrix has four outcomes (Al-Rousan, 2012).

Table 1. Confusion matrix

Actual	Predicted	
	True Positive (TP)	False Negative (FN)
	False Positive (FP)	True Negative (TN)

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} * 100 \quad (4)$$

This measure is proposed as a fitness measure in evaluating each subset generated by data mining classification algorithms.

$$Precision = \frac{TP}{TP + FP} \quad (5)$$

Another one measure used in classification is precision, and the denominator in precision consists only of the total predicted positive.

Another measure is the Roc area, a graph that shows how well a classification model performs over all classification thresholds. This curve is based on two parameters, true and false positive rate. Equations (6) and (7) are given for True Positive Rate (TPR) and False Positive Rate (FPR), respectively (Mulla, 2021).

$$TPR = \frac{TP}{TP + FN} \quad (6)$$

$$FPR = \frac{FP}{FP + TN} \quad (7)$$

RESULT

In Table 2 according to Accuracy (ACC), Precision (PRE), and ROC area (ROC) measure, the analysis results of the raw data set and the data sets obtained by the methods (PCA and SMOTE) are given in determining the effectiveness of methods and performance of the four classification algorithms.

Table 2. Measurement rates for classification algorithms and methods

Algorithms	Raw Data			PCA			SMOTE		
	ACC	PRE	ROC	ACC	PRE	ROC	ACC	PRE	ROC
LDA	96.4	96.6	99.6	97.0	97.1	99.1	97.4	97.5	99.4
ANN	97.6	97.7	97.8	97.7	97.7	99.2	99.1	99.2	99.3
SGD	98.2	98.3	97.7	97.7	97.8	97.0	97.8	97.9	97.9
RFA	94.7	94.8	99.6	95.4	95.4	99.0	97.8	97.9	99.0

Table 2. show that, the accuracy, precision and ROC area measures were calculated for checking the performance of the four classification algorithms with and without using dimensionality reduction and rebalancing data methods. For the raw data, the highest Accuracy, Precision rate of 98.2%, 98.3% were obtained in SGD respectively; but the highest Roc area rate of 99.6%, was obtained in LDA and RFA; According to the Accuracy and Precision measurements calculated, the lowest rates as 94.7%, 94.8% were obtained in RFA. On the other hand, when we used PCA the result was improved significant because the number of variables decreased to 12 variables. However, when the PCA dimensionality reduction method was used, we can see that the highest Accuracy rate of 97.7%, were obtained in ANN, SGD respectively; The highest Precision rate of 97.8% was obtained in SGD; on the other hand, the highest Roc area rate of 99.2% was obtained in ANN algorithm. When the SMOTE oversampling method was applied, the highest Accuracy, Precision rate of 99.1% 99.2%, were obtained in ANN, respectively. When Table 1 is examined in more detail, it is seen that the performance of the classification algorithms is very good, although the data is rebalanced and the dimensionality is reduced. Thus, it can be said that the given methods are important to deal with high-dimensional data in the presence of imbalance problem.

Figure 2 shows the performance of each classification algorithms, before and after using PCA and SMOTE, according to Accuracy, Precision and ROC area.

When Figure 2 is examined, when PCA and SMOTE methods are used, a good improvement was observed in the performances of classification algorithms except for SDG, according to almost all three measures. namely, the algorithms gave very good results although 32 variables were reduced to 12 variables with PCA and 25.5% imbalance rate was eliminated with SMOTE. Therefore, this study indicates that both methods are very useful.

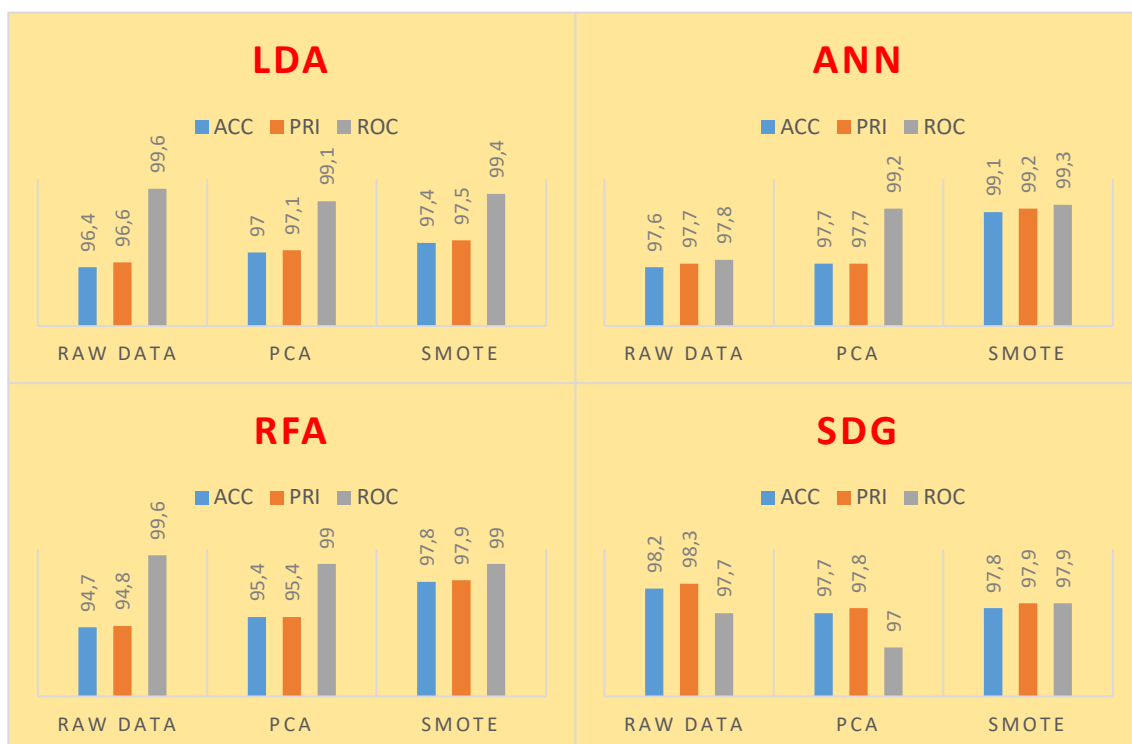


Figure 2. Performance of each classification algorithms, before and after using methods.

According to Figure 2, the LDA figure showed that accuracy and precision increased after using PCA and SMOTE methods, but Roc area was slightly decreased. It was observed that the ANN algorithm gave very good results according to all three measurements in the SMOTE method, and the Roc field in the PCA method gave very good results compared to the first case and the other two measurements.

Considering the ROC area measure, the RFA algorithm performed poorly with both methods compared to the raw data set. However, according to Accuracy and Precision measurements, the algorithm performed very well in both methods, especially in the SMOTE method. According to the three measurement criteria, it was determined that the SDG algorithm gave better results in the raw data set. In general, it can be said that the SDG algorithm does not give bad results, but it does worse than the initial state and other algorithms. This may be due to the fact that the SDG algorithm already gives very good results before the methods are used.

DISCUSSION AND CONCLUSION

The problem of dealing with high dimensional imbalanced data is resolved by removing redundant features (via the PCA method of dimensionality reduction) and rebalancing the results (using the SMOTE method). By defining a faster and more effective model, more ideal solutions and more successful results can be obtained in many different fields, especially in the field of health (Naseriparsa and Kashani, 2014).

Experimental results on the four different classification algorithms for imbalanced high-dimensional data showed that all classification algorithms have enhanced the classification performance of datasets using either PCA, SOMTE methods. For the dataset, the SGD algorithm has provided the best results for raw data before applying any rebalancing or dimensionality reduction methods. On the other hand, when the SMOTE method is used, the ANN classifier gave better results than other algorithms. It was observed that PCA and SMOTE methods did not affect the SGD algorithm results, but these two methods improved the results of LDA, ANN and RFE algorithms.

As a result, it has been observed that classification performances increase in general when PCA and SMOTE methods are used. When these methods are used, it can be said that the best results are obtained with the ANN algorithm and although good results are obtained with the SDG algorithm, no improvement is achieved with the methods.

References

- Adebayo, A. O., Chaubey, M. S. 2019. Data Mining Classification Techniques on the Analysis of Student's Performance. *Global Scientific Journal*, 7(4), 79-95.
- Al-Rousan, N., Haeri, S., Trajković, L. 2012. Feature selection for classification of BGP anomalies using Bayesian models. 2012 International Conference on Machine Learning and Cybernetics, Xi'an, China, 15-17 July 2012, pp. 140-147.
- Bottou, L. 2012. Stochastic Gradient Descent Tricks. Grégoire Montavon. Geneviève B. Orr and Klaus-Robert Müller (Eds.) *Neural Networks: Tricks of the Trade (2nd ed.)*, pp. 421-436. Springer-Verlag, Berlin.
- Cutler, A., Cutler, D. R., Stevens, J. R. 2012. Random Forests. Cha Zhang and Yunqian Ma (Eds.) *Ensemble Machine Learning: Methods and Applications*, pp. 157-175. Springer-Verlag, New York. DOI 10.1007/978-1-4419-9326-7
- Dell'Aversana, P. 2019. A Global Approach to Data Value Maximization: Integration, Machine Learning and Multimodal Analysis. Cambridge Scholars Pub., Newcastle.
- Ławrynowicz, A., Tresp, V. 2014. Introducing Machine Learning. Jens Lehmann and Johanna Voelker (Eds.) *Perspectives on Ontology Learning*, pp. 35, 50, IOS Press, Germany.
- Mulla, G.A.A. 2021. Combination of PCA with SMOTE Oversampling for Classification of High-Dimensional Imbalanced Data [M.Sc. Thesis]. Van Yuzuncu Yil University, Institute of Natural and Applied Sciences, Van.
- Naseriparsa, M., Kashani, M. M. R. 2014. Combination of PCA with SMOTE resampling to boost the prediction rate in lung cancer dataset. *International Journal of Computer Applications*, 77(3): 33-38.
- Qin, A. K., Shi, S. Y. M., Suganthan, P. N., Loog, M. 2005. Enhanced direct linear discriminant analysis for feature extraction on high dimensional data. *Proceedings of the 20th National Conference on Artificial Intelligence (AAAI-05)* (pp. 851-855).

- Pham, T. S., Hoang, T. H., Vu, V. C. 2016. Machine learning techniques for web intrusion detection – a comparison. 2016 Eighth International Conference on Knowledge and Systems Engineering (KSE 2016), pp. 291-297, 6-8 October 2016, Hanoi, Vietnam.
- Poston, W. L., Marchette, D. J. 1998. Recursive Dimensionality Reduction Using Fisher's Linear Discriminant. *Pattern Recognition*, 31(7), 881-888.
- Shashoa. N. A. A., Salem, N. A., Jleta, I. N., Abusaeeda. O. 2016. Classification Depend on Linear Discriminant Analysis Using Desired Outputs. *17th International Conference on Sciences and Techniques of Automatic Control and Computer Engineering (STA'2016)* (pp. 328-332), Sousse, Tunisia, December 19-21.

Denetimli İstatistiksel Öğrenme Algoritmaları ile Toprak Radon Gazının Tahmini

Çağla ÖZTÜRK ZAN^{1*}, Neslihan DEMİREL²

¹Dokuz Eylül Üniversitesi-Fen Bilimleri Enstitüsü, İstatistik Ana Bilim Dalı, 35390, İzmir, Türkiye

²Dokuz Eylül Üniversitesi-Fen Fakültesi, İstatistik Bölümü, 35390, İzmir, Türkiye

*Sunucu yazar e-posta: cerencaglaozturk@gmail.com

Özet

Radon, kayalarda doğal uranyum ve radyum elementlerinin bozunmasından oluşan tatsız, kokusuz, renksiz asal bir gazdır. Radon, yeryüzünde bulunan tüm radyasyon kaynakları içerisinde en yüksek doza maruz kalan doğal radyasyon kaynağıdır. Radonun hareketini şekillendiren faktörler arasında; Radon izotoplarının bozunma hızı, gözenekleri dolduran akışkanlar (hava, su ve diğer gazlar), atmosferik basınç gibi meteorolojik faktörler yer alır. Bu çalışmanın amacı bazı meteorolojik faktörlerin, Radon gazı üzerindeki etkilerinin makine öğrenme algoritmaları ile değerlendirilmesi ve Radon gazı değerlerinin bu faktörlere göre tahmin edilmesidir. Bu amaçla 30 Ekim 2006 - 04 Haziran 2007 tarihleri arasında saatlik periyodlarla Seferihisar bölgesinden elde edilen Radon gazı seviyesine ek olarak T.C. Tarım ve Orman Bakanlığı Meteoroloji Genel Müdürlüğü'nden Saatlik Aktüel Basınç(hPa), Saatlik 50 cm Toprak Sıcaklığı(°C), Saatlik Nispi Nem(%), Saatlik Sıcaklık(°C), Saatlik Rüzgâr Derecesi(°), Saatlik Rüzgar Hızı(m/sn) ve Saatlik Rüzgar Yönü parametreleri ölçümleri elde edilmiştir. Denetimli istatistiksel öğrenme algoritmalarından; Çoklu Doğrusal Regresyon, K en yakın komşu, Destek Vektör Makineleri, Regresyon Ağaçları, Torbalama(Bagging), Rassal Ormanlar, XGBOOST yöntemleri kullanılmıştır. k-çapraz geçerlilik (k=5) ile doğrulama testleri gerçekleştirilmiştir. Algoritmaların performanslarını karşılaştırmak üzere Belirtme Katsayısı (R²), Hata Kareler Ortalaması (Mean Square Error-MSE), Kök Hata Kareler Ortalaması (Root Mean Square Error-RMSE), Ortalama Mutlak Hata (Mean Absolute Error-MAE) değerleri kullanılmıştır. Performans kriterleri dikkate alındığında en iyi sonucu Rassal Ormanlar Regresyonu vermiştir. Bu yöntemi, birbirlerine çok yakın sonuçlar veren XGBOOST ve K en yakın komşu algoritmaları takip etmiştir.

Anahtar kelimeler: Denetimli istatistiksel öğrenme algoritmaları, Toprak radon gazı, Meteorolojik faktörler.

Fındıkta Yürütülen Seleksiyon İslahı Çalışmaları Özelliklerinin Değerlendirilmesi

Ali TURAN^{1*}

¹Giresun Üniversitesi, Teknik Bilimler Meslek Yüksekokulu, Bitkisel ve Hayvansal Üretim Bölümü,
28100, Giresun

*Sunucu yazar e-posta: ali.turan@giresun.edu.tr

Özet

Ülkemizde seleksiyon çalışmaları 1969 yılında Doğu Karadeniz Bölgesi'nde survey şeklinde başlamış ve çalışma sonucu 522 çeşit ve tip tespit edilerek Fındık Araştırma Enstitüsü genetik kaynaklar parseline dikilmiştir. Genetik kaynakları parseline dikilen bireylerin UPOV 1969, IBGR 1981 ve 1988'e göre 79 karakter yönünden karakterizasyonu yapılmıştır. Karakterizasyon çalışması ile mevcut çeşit ve tiplerin özellikleri belirlenerek ıslahçıların daha çok ve bilinçli kullanması amaçlanmış ancak istenen düzeyde kullanılamamıştır. Daha sonraları da bu ve benzeri pek çok seleksiyon çalışması yürütülmüştür. Ancak bu çalışmalarda eleme için kullanılan özellikleri değerlendiren bir çalışma bulunmamaktadır. Bu çalışma seleksiyon ıslahı çalışmalarında kullanılan verim, çotanaktaki meyve sayısı, meyve ve iç büyüklüğü, kabuk kalınlığı, beyazlama oranı ve buruşuk iç oranı gibi temel özellikleri çeşitler bazında değerlendirmek için yürütülmüştür.

Anahtar kelimeler: Çeşit, Fındık, Özellik, Seleksiyon

Evaluation of Traits of Selection Breeding Studies Conducted in Hazelnut

Abstract

Selection studies in our country started in 1969 in the Eastern Black Sea Region as a survey and as a result of the study, 522 cultivars and types were determined and planted in the genetic resources area of the Hazelnut Research Institute. The genotypes planted in the genetic resources area were characterized in terms of 79 traits according to UPOV 1969, IBGR 1981 and 1988. With the characterization study, the traits of the existing cv and types were determined and it was aimed to be used more and more consciously by the breeders, but it could not be used at the desired level. Afterwards, this and many similar selection works were carried out. But, there are no studies that evaluate the traits used for elimination. This study was carried out to evaluate the basic traits used in selection breeding studies such as yield, number per cluste, nut and kernel size, shell thickness, pellicle removal and shrivel kernel on the basis of cultivars.

Key words: Cultivar, Hazelnut, Traits, Selection

GİRİŞ

Ülkemizde yetiştiriciliği çok eskilere dayanan fındık, en uygun yetiştirme ekolojisini Karadeniz Bölgesinde bulmuş, Türkiye'nin önemli tarım ürünlerinden biridir. Karadeniz Bölgesinde yer alan Ordu, Giresun, Trabzon, Bolu, Sakarya ve Samsun'da Türkiye fındık üretiminin yaklaşık % 90'ı gerçekleştirilmektedir. Bu bölge bol yağış almakta olup kıyılarda yazlar serin (23-24°C), kışlar ılık (5-7°C) Karadeniz iklimi, iç kesimlerde ise daha çok karasal iklim görülmektedir. Dağların kıyıya paralel uzanması tarım alanlarını kısıtlamakta olup iklimin de farklı olmasına sebep olmaktadır. Ülkemizin en fazla yağış alan bölgesi olan Karadeniz Bölgesi'nde dağların kıyı kesimin nemli olan havasının iç

kesimlere geçmesini engellemesi bitki örtüsünün de farklılık göstermesine neden olmaktadır (Öztürk ve Serttaş, 2018; Turan ve İslam, 2020) .

Bu iklim özelliklerinden dolayı da dünyanın en kaliteli ve nitelikli fındıkları bu alanda yetiştirilmektedir. Fındık, gıda üretimine hem doğrudan girmekte hem de fındık yağı olarak kullanılmaktadır. İçerisinde çeşitli amino asitler ve yağ asitleri belirlenmiş, besin değerleri tespit edilmiş (Dönmez ve ark., 2016) olan fındık aynı zamanda, tekli ve çoklu doymamış oleik ve linoleik yağ asitleri bakımından, steroller, temel mineraller, serbest fenolik asitler, fenolik bileşikler ve organik asitler bakımından zengin ve sağlıklı bir üründür (Shafiei ve ark., 2020). Bu yüzden de fındık meyvesi çerez olarak tüketiminin yanı sıra, bütün, doğranmış ya da un olarak gıda sanayinde geniş çapta kullanılmaktadır. Ayrıca fındık yağı da yağ endüstrisi için önem arz etmektedir. Sanayi ürünü olan fındığın %80'i çikolata imalatında, %15'i şekerleme, bisküvi ve pasta imalatı, %5'i ise işlem yapılmaksızın tüketilmektedir (Cansev ve ark., 2018).

Bunlara ilave olarak, fındık içeriğindeki antioksidanlar sayesinde kanser ve damar tıkanıklığı hastalıklarını önlemekte, fitokimyasallar ve fenolik bileşikler sayesinde de kanserin ve oksidatif stresin zararlı etkilerine karşı koruma sağlamaktadır (Yılmaz ve ark., 2019). Bu zengin besin içeriği çeşit, hasat zamanı, kurutma, kolonal farklılık, besleme, rakım ve ekolojiye göre farklılık göstermektedir (Turan, 2018). Ama en önemli faktörün çeşidin ve/veya klonun genetik yapısı olduğu bilinmektedir (Turan, 2007; İslam, 2019). Ekonomik değeri olan fındık çeşitleri çotanakdaki meyve sayısı, meyve ve iç ağırlığı, randıman, kabuk kalınlığı, beyazlama oranı, buruşuk iç oranı ve çıtlama oranı gibi pek çok özellik bakımından farklılık göstermektedir. Ama çeşit ıslahı için yüksek verim ve sanayiye yönelik üstün kalite özellikleri gösteren bireyler tercih edilmektedir.

Islahçı bu tercihini yapabilmek için, birden fazla gen (multiple factor) tarafından kontrol edilen verim gibi karmaşık özelliği etkileyen birden fazla karakteri eş zamanlı olarak araştırılması ve değerlendirmesi gerekmektedir. Bu durumda, farklı genler tarafından kontrol edilen özelliklerin fizyolojileri ile ilgili ilişkilerin detaylı araştırılması zorunludur. Bu nedenle, seleksiyon ıslahı programlarında farklı ekonomik karakterler arasındaki korelasyonlara ait bilgiler de ayrıca çok önem arz etmektedir. İki veya daha fazla özellik arasındaki pozitif korelasyon birden fazla değişkenin aynı anda ıslahını mümkün kılarken, negatif korelasyon ise arzu edilen karakterler arasında bir bağlantıya ihtiyaç duyulduğunu gösterir (İşbakan ve Bostan, 2020). Ayrıca, bu kantitatif karakterler çevresel faktörlerden çok fazla etkilenebilmektedir. Bu durum da aslında, tek bir gen tarafından kontrol edilen özelliklere oranla çok daha esnek davranacağı anlamına gelmektedir. O nedenle bir ıslahçı, herhangi bir çevre koşullarında verimdeki değişkenlik nedenlerini çok iyi bir şekilde tanımlayabilmelidir. Aksi durumlarda büyük bir yanlış olabilir ve yapılan çalışma ve emekler karşılığını da aynı zamanda bulamayabilir.

1969 yılından günümüze kadar çok sayıda seleksiyon çalışması yürütülmüş (İslam, 2000; Turan, 2007; Turan ve Beyhan, 2009; Balık ve ark., 2013; Göğüs, 2015; Pekdemir, 2019; Şahin, 2019; Kan, 2019; İslam ve Çayan, 2019) ve hala yürütülmektedir. Bitkiler arasında meydana gelen mutasyon ve doğal melezlemeler aynı çeşit içerisinde geniş bir varyasyona sebebiyet verebilir. Bundan dolayı önemli bir kaynak olan varyasyonlar içerisinde istenilen özellikleri taşıyan bireylerin seçimi ıslahçılar için çok önemlidir (İslam, 2019). Çünkü yeni oluşan varyasyonlar ekonomik açıdan esas çeşide oranla daha iyi olabileceği gibi bu durum tam tersine de dönüşebilir. Bu sebeple de seleksiyon ıslahı çalışmaları süreklilik arz etmektedir.

Seleksiyon ıslahı çalışmalarının çoğunluğu maalesef çok geniş alanlarda ve çok materyalle devam etmiş ve/veya etmektedir. Oysa ki, dar alanda detaylı çalışma yürütmenin çok daha verimli ve sonuca yönelik olacağı arazi çalışmalarıyla ortaya konulan bir gerçektir. Ayrıca bu çalışmaların önemli bir kısmında seleksiyon kriterleri benzerlik göstermekte ve maalesef çeşide ait özellikler iyi tanımlanarak yapılan seleksiyon çalışması sayısı çok fazla değildir. Diğer yandan da günümüze kadar fındıkta

seleksiyon çalışmalarını detaylı değerlendiren bir çalışmaya rastlanılmamıştır. O nedenle bu çalışma mevcut seleksiyon kriterlerini çeşitlere göre değerlendirmek amacıyla yürütülmüştür. Değerlendirme sonucunda ortaya konulacak veriler bundan sonraki ıslah çalışmalarına katkı sağlayacağı öngörülmektedir.

FINDIKTA SELEKSİYON ISLAHI KRİTERLERİ

Seleksiyon ıslahı çalışmalarında pek çok meyve ve iç özelliği kriter olarak kullanılmaktadır. Şöyleki, verim, erken hasada gelme, geç yapraklanma, çotanakdaki meyve sayısı (çms), zuruf boyu, zurufun meyveyi sarma durumu, meyve ve iç büyüklüğü, şekil, şekil indeksi, çıtlak meyve oranı, kabuk kalınlığı, randıman, liflilik durumu, beyazlama oranı, göbek boşluğu, buruşuk oranı, abortif iç oranı, çift iç oranı, boş meyve, göbek boşluğundaki kahverengileşme durumu, urlu iç, küflü iç, gizli küflü, gizli çürük, limonlaşma, ekşi limonlu ve çürük meyve olmak üzere çok sayıda meyve ve iç özelliği kullanılmaktadır (İslam, 2000; İslam ve ark., 2005; Turan, 2007; Turan ve Beyhan, 2009; Turan, 2017; Turan, 2019). Bu özelliklerin önemli bir kısmı seleksiyon çalışmalarında eleme kriteri olarak kullanılmaktadır. Bu seçim işlemleri arazi gözlemleri ve fiziksel özelliklerden oluşmaktadır. Seçilen bireyler daha sonra kontrollü verim denemesine alınarak seleksiyon II aşamasında tekrar değerlendirilmektedir. Bu seleksiyon II aşamasında daha çok fiziksel özellikler, çok az sayıda moleküler tanımlama ve yakın akrabalık seviyesinin belirlenmesi üzerinden yapılmaktadır. Moleküler tanımlamada ise amaç uzak akrabalar arasında melezleme çalışması ile varyasyonu artırarak yeni bireylerin ortaya çıkmasını sağlamaktır.

Ancak bu çalışmalarla önceden başlamış ve henüz tamamlanan çalışmalar haricinde (Örnek: Giresun melezi) henüz pratiğe aktarılacak sonuç alınamamıştır. Seleksiyon II aşaması aslında, yukarıda bahsedilen özelliklere ilaveten fındığın besin içeriğini de içermelidir. Şöyleki, son dönem çalışmaları fındık ve diğer sert kabuklu meyvelerin insan sağlığı üzerine özellikle kalp-damar ve beyin fonksiyonları üzerine çok önemli etki yaptığını ortaya koymaktadır (Hammad ve ark., 2016; Marzocchi ve ark., 2017; Pelvan ve ark., 2018; Alaşalvar ve ark., 2020; Turan, 2021). Dolayısıyla kontrollü verim aşamasında eleme kriteri olarak diğer özelliklere ilaveten yağ asitleri kompozisyonu gibi pek çok özelliği içermesinde büyük yarar görülmektedir. Mesela, nervonik asit fındıkta çok düşük olmasına rağmen beyin fonksiyonları üzerine olağanüstü etkisi olduğu bilinmektedir (Tang ve ark., 2013; Cömlekcioglu ve Mutlu, 2021). Bu nedenle seleksiyon ıslahı aşamasında fiziksel özelliklerde belli eşik değeri yakaladıktan sonra mutlaka insan sağlığı yönüyle eliminasyon yapılmalıdır. Çünkü son dönem salgın hastalıkları insan sağlığından daha önemli bir şeyin olmadığını çok ağır bir şekilde göstermiştir.

Seleksiyon kriterlerinin değerlendirilmesi

Verim

Verim genel olarak seleksiyon ıslahı çalışmalarında üç şekilde değerlendirilmektedir. Birinci olarak, aynı bahçe içerisinde, üreticinin yönlendirmesiyle bulunan en verimli ocaktan, en verimli dal seçilerek bu daldaki bütün meyveler toplanır. Daha sonra meyveler ayıklanır ve kurutulur. Kurutma işlemi sırasında kabuklu meyvede nem oranının %7'nin, iç meyvedeki nem oranı % 5'in; toplamda %12'nin altına düşürülmesi sağlanır ve devamında meyve örnekleri tartılır. Tartım işlemi 0.01g'a duyarlı hassas terazide yapılmalı ve g/bitki olarak ifade edilmelidir (Turan, 2007). İkinci olarak, dal verim etkinliğini belirlemek için seçilen her dal önce topraktan en az 10 cm yükseklikten (Her bahçe ve ocak için aynı yükseklik olmayabilir. Bu nedenle 30 cm yüksekliğe kadar çıkabilir) her dalda kuzey-güney ve doğu-batı doğrultusunda 2 çap ölçümü yapıp ortalaması alınmaktadır. Bulunan değer yarısı (r) $\pi \cdot r^2$ formülünde yerine konulup gövde kesit alanı belirlenmekte ve gr/cm² olarak ifade edilmektedir. Üçüncü olarak ise, gözlemlerle belirlenmekte ve 1-5 skalasında değerlendirilmektedir. En düşük verimi olan klon 1, en yüksek verimi olan klon 5 puan verilerek yarışdırılmaktadır.

Her üç değerlendirme kriterinde de bazı noktalarda sorunlar yaşanabilmektedir. Şöyle ki, birinci verim değerlendirmesinde bahçeler arasındaki farklılık tam olarak değerlendirilememektedir. Çünkü toprak yapısı, bitki yaşı, ekolojik farklılık ve yöney verimi önemli düzeyde etkilemektedir. Ölçüm metrik olabilir ancak bu farklılık zaman zaman göz ardı edilebilmektedir. Yüksek rakım ve iklimi biraz daha sert olan tepe/sırt bölgelerde meyveler küçük olacağından haliyle verimde de azalma olmuş gibi değerlendirilebilmektedir. İkinci değerlendirme de ise, bitkilerin gelişimi, kök yaşı ve köklerin toprak yüzeyine yakınlığı haliyle gövde kalınlığını etkileyecektir. Her ne kadar ortak belli bir yükseklikten ölçümde anlaşılabilir bu farklılıkların araştırmaya olumsuz etki edeceği değerlendirilmektedir. Son verim değerlendirme kriteri bana göre en uygun seçenek gibi görünmektedir. Metrik olmadığı için görsel olarak insanların yanılacağı düşünülerek bazı araştırmacılar tarafından tercih edilmemektedir. Oysa ki, arazi gözlemi iyi ve tecrübeli bir ıslahçı bitkinin gelişimi, tacı, rakımı ve bahçenin genel durumunu hızlı bir şekilde gözden geçirerek en iyi değerlendirmeyi yapabilir. Arazi tecrübesi olmayan ıslahçıların belli bir olgunluk seviyesine gelmeden ve usta çırac ilişkisi içerisinde yetişmeden kesinlikle seçim işlemlerinde belirleyici olmaması gerekir. Yoksa öngörülemez hatalar zinciri oluşabilir ve tüm emekler karşılıksız kalabilir. Sonuç olarak, verim gözlemle değerlendirilmeli ve diğer yöntemler sadece destekleyici olarak kullanılmalıdır. İlla bazı yöntemler metrik olduğu için tercih edilmemelidir.

Çotanaktaki meyve sayısı (ÇMS)

Hasat edilen bitkilerden alınan 100 çotanak ve meyveler sayılarak çotanaktaki meyve sayısı tespit edilmektedir (Turan, 2007; TMS: Toplam meyve sayısı; TÇS: Toplam çotanak sayısı).

$$\text{ÇMS} = \frac{\text{TMS}}{\text{TÇS}}$$

Tamamı hasat edilen bitkiden çotanaklar tesadüfen seçilmeli ve sağlam meyveler değerlendirilmelidir. Çotanaktaki meyve sayısı aslında bir çeşit özelliğidir. Dolayısıyla meyvenin büyüklüğü ve şekli üzerinde doğrudan etkisi bulunmaktadır. En ideal çms 4 adet/çotanak olmalıdır. Çünkü çotanakta meyve sayısı arttıkça meyve küçülmekte ve birbirine baskı oluşturduğundan şekil değişkenliğine ve/veya meyve üzerinde deformasyona neden olmaktadır. Meyvenin az olması ise meyve büyüklüğünü ve dolayısıyla kabuk kalınlığını, en nihayetinde meyve standardını bozmaktadır. Değerlendirme aşamasında en yüksek puan 4 adet/çotanak sayısına verilmeli (Örnek: Tombul), 3-5 adet/çotanak değerleri de kabul edilebilir olarak değerlendirilmelidir. Diğer sayılara sahip olan klonlar değerlendirilmeye alınmadan baştan elenmelidir.

Meyve ve iç ağırlığı (g)

Meyve ve iç ağırlığı tesadüfen seçilen 30 meyve 0.01 g'a duyarlı terazide tek tek tartılıp aritmetik ortalaması alınarak hesaplanmaktadır (Turan, 2017).

$$\text{MA} = \frac{\sum X_i}{n}$$

Meyve ve iç ağırlığı çeşit, verim, bitkinin beslenme durumu, rakım ve ekoloji gibi pek çok faktöre göre değişkenlik göstermektedir. Yao ve Mehlenbacher (2000) ve Turan (2019), meyve ağırlığının kalıtım değerinin $h^2=0.63$, İA'nın ise kalıtım derecesinin $h^2=0.67$ olduğu bildirmişlerdir. Kalıtım derecelerinin yüksek olması teorik olarak özellikle iç ağırlığının çevre şartlarından daha az etkilenebileceğinin göstergesi olabilir. Sanayici ve ihracat ve/veya taze tüketim yönüyle bakıldığında, değerlendirilen kısım iç meyve olması nedeniyle daha da çok ön plana çıkmaktadır. O nedenle eleme kriteri olarak, meyve ağırlığından ziyade iç ağırlığı değerlendirilmesi daha büyük katkı sağlayacağı öngörülmektedir.

İç Oranı (%)

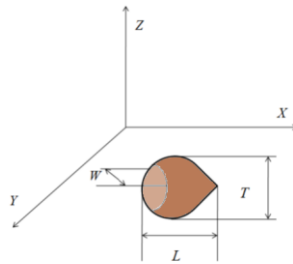
Toplam meyve ağırlığının toplam iç (dolgun ve kusurlu içler) ağırlığına oranlaması yoluyla yüzde (%) olarak ve tesadüfen alınan 50 meyve kullanılarak hesaplanmaktadır (Turan, 2019).

$$\text{İç Oranı (İO; \%)} = \frac{\text{İç Ağırlığı}}{\text{Meyve Ağırlığı}} \times 100$$

İO özelliğinin çeşitlere göre farklılık gösterdiği bilinmekte birlikte (Turan, 2017), verim, ekoloji ve bakım şartlarından aşırı derecede etkilenmediği, bunun nedeninin ise kalıtım derecesinin yüksek ($h^2=0.87$; [Yao ve Mehlenbacher, 2000; Turan, 2019]) olmasından kaynaklandığı bildirilmiştir. Ayrıca İO aynı çeşit içerisinde yıl ve lokasyon farkına göre de çok az düzeyde varyasyon göstermektedir (Hosseinpour ve ark., 2013; Turan, 2019). Bu durum aslında seçim işleminde eleme kriteri olarak kullanılmasının önemini ortaya koymaktadır. İç oranı zaman zaman farklı şekillerde de değerlendirilen yada çok anlam yüklenen bir özellik olarak karşımıza çıkmaktadır. Şöyle ki, fındık piyasasında çeşit farklılığının yanında ve hatta önüne geçen özellik olarak görülmekte ve değerlendirilmektedir. Çünkü fiyatlandırma işlemi %50 randıman üzerinden yapılmaktadır. Çok yüksek randıman olması istenir mi? Yada en yüksek randımanı olan Uzunmusa çeşidi neden en kaliteli fındık olarak değerlendirilmiyor? Aslında fındık çeşitlerini birbirinden ayıran ve değerli kılan düzenli verim ve standart büyüklükte meyve oluşturmalarıdır. Diğer yandan çok yüksek iç dolgunluğu sert kabuk ve iç arasındaki boşluğu azaltması sebebiyle kırma aşamasında vurguna neden olmaktadır. Bu nedenle %50-54 arasındaki randıman değerinin daha uygun olacağı anlaşılmaktadır. Randımanı en yüksek fındık en kalitelidir algısı doğru değildir. Aslolan fındığın kimyasal bileşiminin yanında en az hasarla kırma ve işleme sektöründe kullanılabilmesidir. Hasar ve/veya kusurlu iç oranının artması maliyetin artışına yol açması nedeniyle fındık sektörü tarafından tercih edilmemektedir.

Meyve ve iç iriliği (mm)

Ortalama meyve ve iç boyutlarını belirlemek için tesadüfen seçilmiş meyveler üzerinden, 0.01 mm hassasiyette dijital kumpas ile fındıkların (kabuklu ve iç fındıklar) uzunluk (MU), genişlik (MG) ve kalınlık (MK) boyutları ölçülerek (Şekil 1) ve 30 meyvenin aritmetik ortalaması alınarak hesaplanmaktadır. Meyve büyüklüğü (MB), tesadüfen alınan 30 meyvenin meyve uzunluğu, meyve genişliği ve meyve kalınlığının geometrik ortalaması alınarak hesaplanmaktadır (Turan, 2017; Turan, 2019).



Şekil 1. Kabuklu ve iç fındığın meyve boyutları (L: Meyve uzunluğu, T: Meyve kalınlığı, W: Meyve genişliği).

Meyve şekil indeksi (MŞİ) aşağıdaki formül ile hesaplanmaktadır (Turan ve Beyhan, 2009; Turan, 2019).

$$\text{MŞİ} = \frac{\text{MU}}{\text{MG}}$$

Meyve boyutları üzerine çeşit, besleme, verim, ekoloji ve hasat zamanı gibi pek çok faktör etki etmektedir (Xu ve Hanna, 2010; Ercişli ve ark., 2011; Turan, 2017; Turan ve İslam, 2018). Ayrıca meyve

uzunluğu (MU), meyve genişliği (MG) ve meyve kalınlığının (MK) kalıtım derecesinin yüksek olduğu ($h^2=0.68, 0.78$ ve 0.89 , sırasıyla) bilinmektedir (Yao ve Mehlenbacher, 2000). Kalıtım derecesinin yüksek olmasından çevre şartlarından çok fazla etkilenmediklerini anlayabiliriz. Ancak fındıkta verim, meyve büyüklüğünü değiştiren en önemli etmenlerden birisi olarak bilinmektedir. Ayrıca tek bir çotanak içindeki meyveler arasında bile fiziksel farklılık görülmektedir (Turan, 2017). Çünkü verimin yüksek olduğu sezonlarda çotanaktaki meyve sayısı ve randıman artış gösterirken meyve büyüklüğü ve kabuk kalınlığı azalış göstermektedir. Hatta çotanaktaki meyve sayısının artışı meyvenin şekil değerinin değişmesine de neden olmaktadır. Bu nedenle de bir çeşit içinde bile meyvenin fiziksel özelliklerinde yıldan yıla değişkenlik gösteren farklılıklar görülmektedir. O nedenle çalışmayı yürütecek kişi çeşidi çok iyi tanımalı, meyve büyüklüğü ve şekil değerini hızlıca gözlemleyip değerlendirmeli ve seçimini yapmalıdır. Aksi durumda çok zaman kaybedilerek arzu edilen hedefe ulaşma süresi çok uzayabilmektedir.

Kabuk kalınlığı (mm)

Kabuk kalınlığı (KK), fındık tablasından yukarıya doğru orta ve/veya ortaya yakın kısmından şişkin yerin en kalın yerinden 0.01mm 'ye duyarlı kumpas kullanılarak tesadüfen seçilen toplam 30 meyve üzerinden yapılmaktadır (İslam ve ark., 2005). Fındıkta kabuk kalınlığı ortalama 1 mm etrafında şekillenmektedir. 1 mm 'nin çok üzerine çıkması kalın kabuklu olarak değerlendirilmektedir. O nedenle kabuk kalınlığının 1 mm altında olması tercih edilmelidir. Örneklerde bu değer altındaki bireyler seçilmeli ve diğer özellikleri bakımından tartılı derecelendirmeye alınmalıdır. Bu değer üzerinde kabuk kalınlığına sahip olan bireyler doğrudan elenerek değerlendirilmeye alınmamalıdır. Çünkü kabuk kalınlığı randımanı etkileyen en önemli özelliklerin başında yer almaktadır.

Göbek boşluğu (mm)

Göbek boşluğu (GB), birleşen iki kotiledon arasında kalan boşluk göbek boşluğu olarak ifade edilmektedir. Göbek boşluğunun en geniş çapı 0.01 mm 'ye hassas kumpas ile ölçülmeli ve mm olarak ifade edilmeli ve ölçümler 30 meyvede yapılmalıdır. Göbek boşluğu kararsız özellik olarak bilinmekte ve çeşitlere göre de değişkenlik göstermektedir (Turan, 2007). Bunlara ilave olarak bir çotanaktaki meyveler arasında da göbek boşluğu boyutlarında farklılık olduğu ve meyve büyüklüğüne göre de değişkenlik gösterdiği bilinmektedir. Bu yüzden, çok değişkenlik gösteren bu özelliğin tartılı derecelendirmede kullanılmaması daha uygun olacaktır. Çünkü yanılgıya yol açarak hatalı değerlendirmelere neden olabilmektedir.

Buruşuk iç oranı (%)

Buruşuk iç oranı (BRŞ), genellikle ürünün bol olduğu yıllarda veya kuraklık ve beslenme yetersizliği gibi etkenler nedeniyle veya kalıtsal olarak meydana gelen ve bir meyvenin dış yüzeyinin yaklaşık %50'sinden fazla bir kısmının buruşuk olması olarak tanımlanmalı ve %50 değerinin altında kalanlarda da buruşuk iç olarak değerlendirilmelidir. Bu özellik buruşuk içlerin yüzdesi (%) olarak belirlenmeli ve 50 meyve üzerinden yapılmalıdır (Turan, 2007). Buruşuk iç oranını toprak yapısı, rakın, erken hasat, klonal farklılık, çeşit ve iklim gibi pekçok faktör etkilemektedir (Turan ve Beyhan, 2009, Kalkışım ve ark., 2016). Bu nedenle değerlendirme de kullanılsa bile tartılı derecelendirmede eleme kriteri olarak dikkate alınmaması daha uygun olacaktır.

Çift İç Oranı (%)

Diğer sert kabuklu meyvelerde olduğu gibi fındıkta da yumurtalık içerisinde iki adet tohum taslağı bulunur. Normal şartlarda bunlardan birisi döllenip gelişerek tohumu oluşturur. Ancak bazen tohum taslaklarından ikisi birden gelişebilir. Böyle durumlarda endokarp içerisinde iki tohum yan yana bulunur. Bunlara çift iç denilmektedir ve bir tohumun gelişebileceği boşlukta iki tohumun gelişmesi her

iki tohumunda bozuk şekilli olmasına neden olmaktadır. Bu iç fındıklar gerek görünüşleri ve gerekse boylamada yarattıkları güçlük nedeniyle yetiştiricilikte istenmezler. Bu oran gelişmiş iki içe sahip meyvelerin oranı olarak hesaplanmalı, yüzde (%) olarak ifade edilmeli ve 50 meyve üzerinden hesaplanmalıdır (Turan, 2007; Turan, 2017). Çift iç oranı olduğu tespit edilen klonlar değerlendirmeye alınmadan elenmelidir. Çünkü çift iç oluşturma özelliği bir kusurdur ve bir taraftan meyve standardını bozarken diğer taraftan da ürünün işleme maliyetini arttırmaktadır.

Beyazlama oranı (%)

Beyazlama oranı (BO), sağlam iç fındıklar fırında 175°C’ de 15 dakika bekletilmeli, daha sonra el ile 15–20 saniye ovularak testa çıkarılmalı ve aşağıdaki formülü ile hesaplanmalıdır (Turan ve Beyhan, 2009).

$$BO (\%) = \frac{\text{Beyazlamış iç}}{\text{Toplam iç}} \times 100$$

Beyazlama oranı (BO), ihraç edilecek ürünlerde aranan başlıca özelliklerden biridir ve yüksek olması arzu edilir (Turan, 2017). Toprak yapısı, ekoloji ve çeşit gibi pek çok özellik tarafından etkilenen BO’nın kalıtım derecesinin $h^2=0.64$ olduğu bilinmektedir (Yao ve Mehlenbacher, 2000). Bu değer bize, beyazlama oranının çevre şartlarından aşırı etkilenmeyeceğini göstermektedir. Tam beyazlayanlar değerlendirmeye alınmalı, üzerinde biraz testa kalan bireyler dahi elenmelidir. Çünkü Tombul çeşidinde bu oran ortalama %97, o nedenle bu çeşitte yürütülen çalışmalarda %95 ve/veya altında beyazlama oranına sahip olanlar elenmelidir. Diğer çeşitler için ise %90 eşik değeri belirlenmeli ve bu değer altına düşenler elenerek devam edilmelidir. Çünkü Türk fındık çeşitlerinde beyazlama oranı yabancı çeşitlere göre çok yüksek ve ihraç edilen ürünlerde tercih sebebi olmaktadır.

Çıtlak meyve oranı (%)

Antepfıstıklarında yapılan çalışmalarda hasat zamanı, sulama durumu, bitki besleme, budama ve anacın çıtlak meyve üzerine etkili olduğu bildirilmiştir (Ertürk ve ark., 2015). Ayrıca başka bir çalışmada çıtlamanın bir çeşit özelliği olduğu ve bu özelliğin yerli çeşitlerde düşük olduğu bildirilmiştir (Özçağırın ve ark., 2005). Diğer sert kabuklu meyvelerden kestanede bu özelliğin kalıtım derecesinin $h^2=0.48$ olduğu bilinmektedir (Nishio ve ark., 2014). Fındıkta ise, çıtlama özelliği ile ilgili günümüze kadar yürütülen çalışma sayısı son derece sınırlıdır. Ancak bazı çeşitlerde bu özelliğin yüksek olduğu bilinmektedir. Uzun yıllar fındıkta yürütülen çalışmalarda içini dolduran fındıkların yanı sıra içini doldurmayanlarda da çıtlak meyve olduğu bilinmektedir. Buradan da, fındığın genetik yapısında böyle bir özelliğin bulunabileceği düşünülebilir. Buna ilave olarak, çıtlama eğilimi olan çeşitlerde ekoloji, besleme ve iç dolgunluğu ile çıtlamanın daha belirgin hale gelebileceğini, ancak çıtlamaya eğilimli olmayan çeşitlerde bu tespiti yapmanın çok zor olacağını bildirilmektedir (Turan, 2017). Bu nedenle de çıtlama özelliği üzerinde çalışılmalı ve kalıtım derecesi mutlaka tespit edilmelidir. Seleksiyon çalışmasında ise çıtlak meyve tespit edilen klonlar değerlendirmeye alınmadan survey aşamasında elenmelidir.

SONUÇ

Bu derleme, fındıkta seleksiyon kriterlerini ele alarak değerlendiren literatürdeki ilk çalışmadır. Önceki araştırmaların önemli bir kısmının birbirine benzer ve geniş alanda çok bireyle yürütüldüğü, bu nedenle aşırı emek harcanmasına rağmen detaya pek çok çalışmada inilemediği anlaşılmaktadır. Sonuç olarak, seleksiyon ıslahı çalışmalarının geniş alandan ziyade dar alanda ve çok detaylı olmasının amaca yönelik ıslah çalışmalarına büyük katkı sağlayacağı anlaşılmaktadır. Bu yüzden Tombul gibi çeşitlerde beyazlama oranı ve randıman gibi özelliklerin değerlendirilmesinden ziyade verim, Çakıldak’ta ise kabuk kalınlığı ve meyve büyüklüğünün arka plana atılarak beyazlama oranı ve iç büyüklüğü üzerinde yoğunlaşmak gereklidir. Palaz ve Foşa çeşitlerinde de meyve özelliklerinde belli bir standart

yakalanması durumunda değerlendirme sisteminin kabuk kalınlığı üzerinden yürütülmesinde yarar görülmektedir. Tabiki bu fiziksel özellikler üzerinden değerlendirme yapmak seleksiyon I aşaması içi uygun görülmektedir. Bu yüzden seleksiyon II aşamasında tartılı derecelendirmede kullanılan temel özelliklere ilaveten, fındığın insan sağlığı üzerine etkisi olan kimyasal özelliklerinde mutlaka değerlendirmede kullanılması gerekir.

Kaynaklar

- Alaşalvar C, Salvado JS, Ros E., 2020. Bioactives and health benefits of nuts and dried fruits. Food Chem, 314: 126192.
- Balık Hİ, Balık SK, Köse ÇB, Duyar Ö, Sıray E, Sezer A, Turan A, Beyhan N, Erdoğan V, İslam, A, Kurt H, Ak K, Kalkışım Ö, 2014. Development of the new cultivars of hazelnut by selection from 'Tombul' hazelnut populations in Giresun and Trabzon provinces. International Mesopotamia Agriculture Congress, Proceedings, 22-25 September 2014, pp: 172-179, Diyarbakır, Turkey.
- Cansev A, Tüccar M, Turhan Ş., 2018. Sakarya İli Kocaali İlçesi'nde faaliyette bulunan fındık işletmelerinin mevcut yapısı ve sorunları. BAHÇE 47(2): 23-31.
- Comlekcioglu N, Kutlu M., 2021. Fatty acids, bioactive substances, antioxidant and antimicrobial activity of *Ankyropetalum* spp., a novel source of nervonic acid. Grasas Aceites, 72 (1): e399.
- Dönmez, İE, Selçuk S, Sargın S, Özdeveci H., 2016. Kestane, fındık ve antepfıstığı meyve kabuklarının kimyasal yapısı. Turkish Journal of Forestry, 17(2): 174- 177.
- Ercisli S, Öztürk I, Kara M, Kalkan F, Seker H, Duyar O, Ertürk Y., 2011. Physical properties of hazelnuts. International Agrophysics, (25): 115-121.
- Ertürk YE, Geçer MK, Gülsoy E, Yalçın S., 2015. Antepfıstığı üretimi ve pazarlaması. Iğdır Üniversitesi Fen Bilimleri Enstitüsü Dergisi, 5(2): 43-62.
- Göğüs A, 2015. Tirebolu Karakaya vadisinde Tombul fındık klon seleksiyonu, Yüksek Lisans Tezi. Ordu Üniversitesi, Fen Bilimleri Enstitüsü, Ordu, Türkiye.
- Hammad S, Pu S, Jones PS., 2016. Current evidence supporting the link between dietary fatty acids and cardiovascular disease. Lipids, (51):507-517.
- Hosseinpour A, Seifi, E, Javadi, D, Ramezanpour SS, Molnar TJ., 2013. Nut and kernel characteristics of twelve hazelnut cultivars grown in İran. Scientia Horticulturea, (150): 410-413.
- İslam A, 2000. Ordu ilinde yetişen Türk fındık çeşitlerinde klon seleksiyonu, Doktora Tezi. Çukurova Üniversitesi, Fen Bilimleri Enstitüsü, Adana, Türkiye.
- İslam A, Turan A, Kurt H, 2005. Effect of ocak and single trunk training systems on yield and nut quality. Sixth International Congress on Hazelnut, 14-18 June 2004, Book of Proceeding, Page:259-262, Tarragona-Reus, Spain.
- İslam A., 2019. Fındık ıslahında gelişmeler. Akademik Ziraat Dergisi, 8 (Özel Sayı): 167-174.
- İslam A, Çayan M., 2019. Ordu ili Gürgentepe ilçesinde yetiştirilen Çakıldak fındık çeşidinde klon seleksiyonu. Akademik Ziraat Dergisi, 8 (Özel Sayı): 1-8.
- İşbakan H, Bostan SZ., 2020. Fındıkta bitki morfolojik özellikleri ile verim ve meyve kalite özellikleri arasındaki ilişkiler. Ordu Üniv. Bil. Tek. Derg, 10 (1): 32-45.
- Kalkışım O, Turan A, Okcu Z, Özdeş D., 2016. Evaluation of the effect of different harvest time on the fruit quality of foşa nut. Erwerbs-Obstbau, (58): 89-92.
- Kan E, 2019. Trabzon'un bazı ilçelerinde yetiştirilen Trabzon sivrisi fındık popülasyonunda klon seleksiyonu, Yüksek Lisans Tezi. Ordu Üniversitesi, Fen Bilimleri Enstitüsü, Ordu, Türkiye.
- Marzocchi S, Pasini F, Verardo V, Cierniewska-Żytikiewicz H, Caboni MF, Romani S., 2017. Effects of different roasting conditions on physical-chemical properties of Polish hazelnuts (*Corylus avellana* L. var. *Kataloński*). LWT Food Sci Tehnol, 77: 440-468.
- Nishio S, Yamada M, Takada N, Kato H, Onoue N, Sawamura Y, Saito, T., 2014. Environmental variance and broad-sense heritability of nut traits in japanese chestnut breeding. Hortscience 49(6): 696-700.
- Özçağırın R, Ünal A, Özeker E, İsfendiyoğlu M., 2005. Ilıman İklim Meyve Türleri, Sert Kabuklu Meyveler Cilt-III, Ege Üniversitesi Ziraat Fakültesi Yayın No:566, İzmir, Türkiye.

- Öztürk A, Serttaş S., 2018. Karadeniz Bölgesi Meyveciliğinin Mevcut Durumu ve Potansiyeli. Journal of the Institute of Science and Technology, (8): 4, 11-20.
- Pekdemir E, 2019. Piraziz (Giresun) İlçesi Tombul fındık populasyonunun verim ve kalite özelliklerinin belirlenmesi, Yüksek Lisans Tezi. Ordu Üniversitesi, Fen Bilimleri Enstitüsü, Ordu, Türkiye.
- Pelvan E, Olgun EÖ, Karadağ A, Alaşalvar C., 2018. Phenolic profiles and antioxidant activity of Turkish Tombul hazelnut samples (natural, roasted, and roasted hazelnut skin). Food Chem, 244: 102-108.
- Shafie G, Ghorbani M, Hosseini H, Mahoonak AS, Maghsoudlou Y SM., 2020. Estimation of oxidative indices in the raw and roasted hazelnuts by accelerated shelf-life testing. Food Sci Technol, 57: 2433-2442.
- Şahin N, 2019. Giresun ili merkez ilçede yetiştirilen sivri fındık çeşidinde klon seleksiyonu, Yüksek Lisans Tezi. Ordu Üniversitesi, Fen Bilimleri Enstitüsü, Ordu, Türkiye.
- Tang TFan, Liu XM, Ling M, Lai F, Zhang L, Zhou YH, Sun RR., 2013. Constituents of the essential oil and fatty acid from Malania oleifera. Industrial Crops and Products 43, 1– 5.
- Turan A, 2007. Giresun ili Bulancak ilçesi Tombul fındık klon seleksiyonu, Yüksek Lisans Tezi. Ordu Üniversitesi, Fen Bilimleri Enstitüsü, Ordu, Türkiye.
- Turan A, Beyhan N, 2009. Investigation of the pomological characteristics of selected Tombul hazelnut clones in the Bulancak area of Giresun province. VI. International Congress on Hazelnut, Book of Proceeding, 23–27 June, Page:61-66,Viterbo, Italy.
- Turan A, 2017. Fındıkta Kurutma Yöntemlerini Meyve Kalitesi ve Muhafazası Üzerine, Etkileri Doktora Tezi. Ordu Üniversitesi, Fen Bilimleri Enstitüsü, Ordu, Türkiye.
- Turan A., 2018. Effect of drying methods on fatty acid profile and oil oxidation of hazelnut oil during storage. European Food Research and Technology, 12: 2181–2190.
- Turan A, İslam A, 2018. Postharvest differences between ‘Tombul’ and ‘Palaz’. IX International Congress on Hazelnut, 15–18 August 2017, pp: 351–358, Samsun, Türkiye.
- Turan A., 2019. Kurutma yöntemlerinin fındığın fiziksel özellikleri üzerine etkisi. Anadolu J Agr Sci, (34): 296-303.
- Turan A, İslam A, 2020. Fındığın ekolojik istekleri. Ali İslam (Editör). Fındık Yetiştiriciliği. Yeşiller grafik tasarım reklam ve matbaacılık, I. Baskı, 27-29s, Fatsa, Ordu, Türkiye.
- Turan A., 2021. Effect of the damages caused by the green shield bug (*Palomena prasina* L.) on the qualitative traits of hazelnuts. Grasas Aceites, 72 (1): e391.
- Yao Q, Mehlenbacher, SA., 2000. Heritability, variance components and correlation of morphological and phenological traits in hazelnut. Plant Breeding, 119: 369–381.
- Yılmaz M, Karakaya O, Balta MF, Balta F, Yaman İ., 2019. Çakıldak fındık çeşidinde iç meyve iriliğine bağlı olarak biyokimyasal özelliklerin değişimi. Akademik Ziraat Dergisi Cilt: (8:Özel Sayı) 61-70.
- Xu YX, Hanna MA., 2010. Evaluation of Nebraska hybrid hazelnuts: Nut/kernel characteristics, kernel proximate composition, and oil and protein properties. Industrial Crops and Products, (31): 84–91.

Çıkar Çatışması

Yazarın herhangi bir çıkar çatışması yoktur.

Üretici Birliklerinin Örgütlenme Bilincine Katkıları; Telli Terbiye Sistemli Bağcılık Örneği

Songül AKIN^{1*}

¹Dicle Üniversitesi, Ziraat Fakültesi, Bahçe Bitkileri Bölümü, 21200, Diyarbakır, Türkiye

*Sunucu yazar e-posta: sakin@dicle.edu.tr

Özet

Kırsal alanda tarım sektöründe faaliyet gösteren organizasyonlarından biride üretici birlikleridir. Tarımsal üretici birlikleri, üreticilere geniş bir hizmet yelpazesi sunmakta ve aynı zamanda örgütsel farkındalığın yaratılması ve yaygınlaştırılmasında rol model olarak hareket etmektedir. Bu çalışmada, Diyarbakır'ın Dicle, Han, Kulp, Silvan ilçelerinde İFAD desteğiyle tesis edilen 79 Telli Terbiye Sistemli Bağcılık yapan işletmelerinin örgütlenme konusundaki bakış açıları birliğe üye olan ve birliğe üye olmayan üreticiler açısından incelenmiştir. Söz konusu işletmelerin 45 adedi Dicle Organik Meyve Yetiştirici Birliğine üyedir. Verilerin analizinde tanımlayıcı istatistik yöntemler, corelasyon analizi ve bağımsız örneklem T testi kullanılmıştır. Likert ölçeği kullanılarak elde edilen sıralı (ordinal) verilerde uyum analizi yapılmış bu amaçla Cronbach's Alpha istatistiği hesaplanmıştır. Çalışmada organik meyve yetiştiricileri birliğine üyelik ve telli terbiye sistemli bağcılığa devam etme isteği arasında pozitif fakat $r=0.095$ düzeyinde çok zayıf bir ilişki olduğu görülmüştür. Üyelik durumuna göre örgütlenmenin önemi ve gerekliliği ile ilgili olumlu ifadelerle katılım ölçeğine bağımsız örneklem t testinde iki grup arasındaki fark pozitif ve anlamlı olarak ($p<0,05$) bulunmuştur.

Anahtar kelimeler: Tarımsal örgütlenme, Telli terbiye sistemli bağ, Bilinçlenme

Contribution of Agricultural Producer Unions to the Awareness of Organization; Trellis System Vineyard Example

Abstract

One of the organizations operating in the agricultural sector in the rural area is the producer unions. Agricultural producer organizations offer a wide range of services to producers and also act as role models in creating and spreading organizational awareness. In this study, the organizational perspectives of 79 trellis system vineyard businesses established with the support of IFAD in Diyarbakır's Dicle, Han, Kulp and Silvan districts were examined in terms of producers who are members of the union and not members of the union. 45 of these enterprises are members of the Dicle Organic Fruit Growers Organization. Descriptive statistical methods, correlation analysis and independent sample T test were used in the analysis of the data. The ordinal data obtained by using a Likert scale were analyzed and Cronbach's Alpha statistics were calculated for this purpose. In the study, it was observed that there was a positive but very weak relationship at the $r = 0.095$ level between the membership of the organic fruit growers association and the desire to continue viticulture with the trellis system. The difference between the two groups was found to be positive and significant ($p < 0.05$) in the independent sample t test according to the scale of participation in positive statements about the importance and necessity of organization according to membership status.

Key words: Agricultural organization, Wire training system, Awareness

A New Generalized Pareto Distribution with Applications

Najma SALAHUDDIN^{1*}, Alamgir KHALIL²

¹Shaheed Benazir Bhutto Women University, Peshawar, KP, Pakistan

²University of Peshawar, Peshawar, KP, Pakistan

*Presenter author e-mail: najma.fwu@gmail.com

Abstract

In this paper, a new generalization of the generalized Pareto distribution is proposed using the “Khalil generator” [1], named as Khalil Generalized Pareto (KGP) distributions. Various mathematical properties of the suggested model have been derived including moments, quantile function, moment generating function, entropy measures, order statistics. The proposed distribution is very flexible in its nature covering several hazard rates shapes (symmetric and asymmetric). Parametric estimation of the model is conferred using method of maximum likelihood. For assessing the performance of the maximum likelihood estimates in terms of their bias and mean squared error using simulated samples, a simulation study is performed. Two real data sets are used from survival and reliability theory for the purpose of illustration. The practical applications show that the suggested model provides better fits than the other considered models.

Key words: Khalil generalized family, Generalized Pareto, moment generating function, hazard function, maximum likelihood

On Evaluating Big and Complex Data Sets: Opportunities and Challenges

Serdar GENÇ^{1*}, Mehmet MENDEŞ²

¹Kirsehir Ahi Evran University - Faculty of Agriculture - Kirsehir, Turkey

²Canakkale Onsekiz Mart University - Faculty of Agriculture - Canakkale, Turkey

*Sunucu yazar e-posta: serdargenc1983@gmail.com

Abstract

In the past, one of the main problems for the researchers was collecting data. Today, thanks to developments in science and technology it is possible to measure different variables of many experimental units. In other words, nowadays the data is flowing from everywhere. Although having many measurements of different variables is a great advantage to obtain more detailed and reliable information about the effect of interested factor(s), this case may cause to run into different challenges especially at the stage of statistical analysis. Since there will be many measures of experimental units in terms of interested factor(s), the researchers will have to work with a massive and complex data sets. In this study, it has been focused on a general evaluation about opportunities and challenges studying with big and complex datasets.

Key words: Big Data, data analysis, Vomplex datasets, New statistical thinking

İyi Kaldıraç Noktalarından Etkilenmeyen Sağlam Varyans Artış Faktörü

Osman Ufuk EKİZ^{1*}

¹Gazi Üniversitesi, Fen Fakültesi, İstatistik Bölümü, 06500, Ankara, Türkiye

*Sunucu yazar e-posta: ufukekiz@gazi.edu.tr

Özet

Regresyon analizinde açıklayıcı değişkenler arasında yüksek korelasyon bulunması çoklu bağlantı problemi olarak bilinmektedir. Bu problem regresyon doğrusu parametrelerinin en çok olabilirlik (EÇOB) tahminlerinin yanlış işaretli ve varyanslarının olduğundan daha büyük çıkmasına neden olmaktadır. Bu durumda ridge, sağlam ve ridge-tipi sağlam gibi farklı tahmin edicilerden hangisinin kullanılacağı çoklu bağlantının varlığına göre farklılık göstermektedir. Bu nedenle çoklu bağlantının teşhisi oldukça önemlidir. Bilinen bir çoklu bağlantı ölçüsü, regresyon katsayılarının EÇOB tahmin edicisine dayalı, açıklayıcı değişkenler arasındaki belirleme katsayısının ($R^2_{EÇOB}$) bir fonksiyonu olarak tanımlanan varyans artış faktörüdür ($VAF_{EÇOB}$). Kaldıraç gözlemler, $R^2_{EÇOB}$ 'un dolayısıyla $VAF_{EÇOB}$ 'un üzerinde oldukça etkili olabilmektedir. Veride şiddetli çoklu bağlantı varken yokmuş, yokken de varmış gibi gösterebilirler. Bu durumda regresyon katsayılarının herhangi bir sağlam tahmin edicisine dayalı sağlam belirleme katsayısının (RR^2) bir fonksiyonu olan sağlam VAF 'ın (SVAF) kullanılması önerilmektedir. Ancak iyi kaldıraç gözlemlerden bazıları SVIF değerlerinin de yanlış sonuçlar vermesine neden olmaktadır. Bu çalışmada iyi kaldıraç gözlemlerden de etkilenmeyen ve dolayısıyla daha iyi sonuçlar veren yeni bir SVAF önerilmektedir. Bu yöntemde, öncelikle sağlam tahmin edicilere dayalı Mahalanobis uzaklık değerlerine bakılarak kaldıraç gözlemler veriden çıkarılmaktadır. Daha sonra kalan veri üzerinden RR^2 'ye bağlı SVAF değerleri elde edilmektedir. Dolayısıyla elde edilen bu değerler iyi kaldıraç gözlemlerin olası negatif etkilerinden de etkilenmemektedir. Gerçek veri ve simülasyon çalışmalarında yeni SVAF 'ın iyi sonuçlar verdiği görülmüştür.

Anahtar kelimeler: Çoklu bağlantıyı-şiddetlendiren kaldıraç, Çoklu bağlantıyı-maskeleyen kaldıraç, Doğrusal regresyon, Aykırı gözlem, Sağlam istatistikler

Sıralı Lojistik Regresyon Yönteminin Sınıflama Performansının İncelenmesi

Ali Vasfi AĞLARCI^{1*}, Cengiz BAL²

¹Eskişehir Osmangazi Üniversitesi, Tıp Fakültesi, Biyoistatistik, 26040, Eskişehir, Türkiye

²Eskişehir Osmangazi Üniversitesi, Tıp Fakültesi, Biyoistatistik, 26040, Eskişehir, Türkiye

*Sunucu yazar e-posta: aglarci32@hotmail.com

Özet

Bu çalışmada sıralı lojistik regresyon yönteminin (OLR) çeşitli korelasyon yapıları, örnek büyüklüğü, yanıt değişkeninin kategori sayısı ve eşit dağılıp dağılmama durumlarına göre sınıflama performansının incelenmesi amaçlanmıştır. Sıralı lojistik regresyon (OLR) yöntemi bağımlı değişkenin en az üç kategoriden oluştuğu ve sıralı ölçekle ölçüldüğü durumlarda bağımlı ve bağımsız değişkenler arasındaki ilişkiyi incelemek için kullanılan bir yöntemdir. Son yıllarda bu yöntemin yanıt değişkeninin sıralı-ordinal yapıda olduğu durumlarda sınıflama amaçlı da kullanılmaya başlandığı gözlenmiş olup, farklı durumlar altında sınıflama performansını değerlendirmek için simülasyon çalışması yapılmıştır. R programı kullanılarak yapılan simülasyon çalışması ile 3 farklı korelasyon yapısı (düşük, orta, yüksek), 3 farklı yanıt değişkeni kategori sayısı (3, 4 ve 5 kategorili), 5 farklı örnek büyüklüğü (100,250,500,1000,2500) ve yanıt değişkeninin eşit dağılım gösterdiği ve göstermediği 90 farklı durum için 10 adet değişken (1 sıralı, 3 ikili, 6 sürekli) üretilerek sınıflama performansı değerlendirilmiştir. Performans ölçümü için doğruluk oranı ve kappa değeri hesaplanmış olup, doğruluk oranının bazı durumlarda yanıltıcı olabileceği düşünülerek çeşitli senaryolardaki performans kıyaslamasında kappa değeri dikkate alınarak değerlendirmeler yapılmıştır. Herbir durum için 1000 tekrar yapılarak ortalama değerler sunulmuştur. Yapılan değerlendirmeler sonucunda sıralı lojistik regresyon yönteminin kappa değerine göre sınıflama performansı incelendiğinde örnek büyüklüğü arttıkça, düşük korelasyon yapısı ve yanıt değişkeninin dağılımının eşit olmadığı durum hariç tüm durumlarda sınıflama performansı artma eğilimi göstermiştir. Düşük korelasyon yapısı ve yanıt değişkeninin eşit dağılım göstermediği durumda özellikle 4 ve 5 kategorili durum için örneklem büyüklüğü arttıkça sınıflama performansının düştüğü gözlenmiştir. Yanıt değişkeni kategori sayısının performansa etkisi incelendiğinde ise kategori sayısındaki artış tüm durumlarda performansı düşürdüğü görülmüştür. Korelasyon seviyesindeki artış 3, 4 ve 5 kategorili durumlar için sınıflama performansını artırmıştır. Yanıt değişkeninin eşit dağılım gösterip göstermeme durumunun sınıflamaya etkisine baktığımızda ise düşük korelasyon yapısı hariç diğer durumlarda sınıflama performansının değişmediği söylenebilir. Düşük korelasyon yapısında ise yanıt değişkeninin eşit dağılım gösterdiği durumdaki performans değerleri, eşit dağılımın olmadığı duruma göre daha yüksek bulunmuştur.

Anahtar kelimeler: Sınıflandırma, Ordinal veri, Makine öğrenmesi

Makine Öğrenmesi Yöntemleri ve Algoritmalarının Kullanımı

Özgür KOŞKAN¹, Emre ASLAN^{1*}

¹Isparta Uygulamalı Bilimler Üniversitesi, Ziraat Fakültesi, Zootečni Bölümü, 32000, Isparta, Türkiye

*Sunucu yazar e-posta: emreaslan0314@gmail.com

Özet

Bu çalışmada son yıllarda çalışılmakta olan makine öğrenmesinin işleyişi hakkında bilgiler verilmiştir. Programlama ve bilgisayar teknolojisinin gelişmesi ile birlikte istatistik metotlar da gelişmektedir. Makine öğrenmesi gelişen bu yöntemlerin bilgisayar teknolojisi ile birleştirilerek bir takım işlemlerin gerçekleştirilmesiyle pek çok farklı durumda kullanılması sağlanmaktadır. Makine öğrenmesi 1959 yılında geliştirilmiş ve bilgisayar biliminin yapay zekâ da model tanımlama ve sayısal öğrenme çalışmalarının bir alt dalıdır. Yapısal fonksiyon olarak pekiştirilen ve veriler üzerinden varsayım yapabilen algoritmaların çalışma ve yapısını araştıran bir sistemdir. Bu yapıda ki algoritmalar veri tabanlı varsayımları ve kararları uygulayabilmek için bir model kurarak çalışırlar. Çok büyük miktarlardaki verinin elle işlenmesi ve analizinin yapılması mümkün değildir. Amaç eldeki mevcut veriyi kullanarak gelecek için tahminlerde bulunmaktır. Bu problemleri çözmek için Makine öğrenmesi (machine learning) yöntemleri geliştirilmiştir. Makine öğrenmesi yöntemleri, eldeki mevcut veriyi kullanarak sonuç için en uygun modeli bulmaya çalışmaktadır.

Anahtar kelimeler: Makine öğrenmesi, Algoritma, Büyük veri

Distance Based Regression Models

Burcu Kurnaz^{1*}, Hasan Önder¹

¹Ondokuz Mayıs University, Faculty of Agriculture, Department of Animal Science, Samsun, Turkey

*Presenter author e-mail: burcu2039@hotmail.com

Abstract

Distance-based regression (DBR) is an alternative method to parameter estimation in regression models when dealing with mixed-type of exploratory variables. The method of DBR is similar to classical linear regression, but the explanatory variables are measured based on distance instead of raw values. Euclidean, Manhattan and Gower dissimilarity measures was examined with using two different explanatory variable structure (Normal and Binomial). To execute the analysis the R package “dbstats” was used. Results showed that the Gower dissimilarity can be used for the binomial explanatory variables and Euclidean distance measure should be used for normal data according to lower AIC and BIC values.

Keywords: Distance measures, Regression, Package dbstats, R software

INTRODUCTION

Statistical analyzes are applied when it is desired to compare or establish a relationship on the data used in many scientific kinds of research. Regression analysis methods interested in the relational part of the statistics and examine a cause and response relationship. It is a relationship method that explores how one or more causes affect the result, and these relationships are one-way.

With regression analysis; Is there a relationship between dependent and independent variables? If there is a relationship, what is the strength of that relationship? What kind of relationship is there between the variables? Is it possible to predict the future values of the dependent variable and how should it be estimated? What is the effect of a particular variable or group of variables on the other variable or variables and how does it change if certain conditions are controlled? Attempts are made to seek answers to questions such as (URL1);

Regression analysis general use to estimate:

- Variables can be obtained in a long time from early measured variables,
- From the data set of a variable that is easy to measure, the values of a variable that is difficult to measure
- Variable expensive to measure from variable cheaper to measure (Huber and Dutter, 1974).

Linear regression analysis is used to create a model that predicts the variable to be determined based on the variable(s) that can be detected more easily or earlier than the variable to be determined (Alpar, 2010). Linear regression analysis is examined under two headings as simple linear regression and multiple regression analysis.

Simple regression analysis explains the linear relationship between the response variable and a single explanatory variable. If a linear or curvilinear relationship between a single response variable and more than one explanatory variable is desired, the relationship is examined by multiple linear regression analysis (Okur, 2009; Weisberg, 2005).

In order for the parameter estimations of the regression model, which will be obtained as a result of both simple and multiple linear regression analysis, to be reliable, some assumptions about the model must be provided. In order to use the regression equation obtained in the simple linear regression analysis for estimation; error terms ($\epsilon_i = Y_i - \hat{Y}$) show random normal distribution, the mean of the expected value of the errors is 0 and the variance is homogeneous and equal to σ^2 , residuals should be independent [$\text{Cov}(\epsilon_i, \epsilon_j) = 0$], some assumptions need to be made, such as the absence of correlation between error terms and explanatory variable(s) (Alma ve Vupa, 2008).

In multiple linear regression, in addition to the assumptions in simple linear regression, the assumption that the explanatory variables are independent from each other must also be provided (Vural, 2007). This assumption, which can also be explained as the condition that the simple linear correlation coefficients between the explanatory variables are zero or very close to zero, is expressed as the absence of "multiple linear correlation" in statistics (Orhunbilge, 2017). Least Squares (Least Squares) estimation method loses its power in case of multicollinearity (Vural, 2007).

Therefore, when selecting explanatory variables, it is recommended that the response variable and simple linear correlation coefficients of these variables be high (close to 1), and the simple linear correlation coefficients between each other should be low (close to 0) (Damodar, 2001).

One of the few models developed as a solution for the above-mentioned situation in parameter estimation methods is Distance Based Regression methods. The purpose of these methods is to properly address problems with estimators of real values, including categorical or a mix of real-valued and categorical explanatory variables (Arenas and Cuadras, 2002).

The distance-based regression model has many applications in analysis of multivariate response regression in various fields, such as ecology, genomics, genetics, human microbiomics and neuroimaging. Multiple outcomes are generally correlated and can be analyzed by a wide range of multivariate methods, such as the traditional multivariate analysis of variance, multivariate regression, and canonical correlation analysis. A good alternative is distance-based regression, which is a nonparametric approach proposed by McArdle and Anderson for analyzing multivariate ecological data based on the distance or dissimilarity measures among the observations. The distance-based regression model uses a pseudo F test statistic to detect differences among different groups in terms of some dissimilarity measure (Li et al., 2019).

Recently, distance-based regression model has been successfully applied in many areas. In genomics, Xu et al. ranked genes based on their aggregated associations by a distance-based regression model in the presence of a driver mutation, and selected the significant genes. In neuroscience, Shehzad et al. detected the voxels associated with brain-phenotypes by a distance-based regression model. In human microbiome research, Chen et al. identified the factors affecting the microbiome composition by a regression-based approach. In all of these applications, the significance values of the pseudo F test statistic derived from the distance-based regression model were calculated numerically by the permutation procedure, which has proven superior for this purpose (Li et al., 2019).

According to the information obtained, the theoretical distribution properties of the pseudo F test statistic have not been discussed in the existing literature. Therefore, based on spectral decomposition and matrix normal assumption, we study the asymptotic properties of the pseudo F test statistic when the distance measure is the Euclidean or Mahalanobis distance. For the Euclidean distance measure, we find that the pseudo F test statistic has the same distribution as a random variable formed as the quotient of two Chi-squared-type mixtures. Such an approximation is useful because it broadens the application range of distance-based regression models to higher dimensional cases, and enables further inference on the pseudo F test statistic (Li et al., 2019).

In this study, it is aimed to compare the distance based regression models created by using Euclid, Manhattan and Gower distance measures. For this aim two different data set were used; 1) regression of crude cellulose (CC) and NDF variables on dry matter consumption (DMC), 2) regression of birth type (BT) and gender on birth weight (BW).

MATERIAL AND METHOD

In the Distance Based Regression Model statistics and data analysis, the concept of geometric distance between individuals or populations has been applied in fields such as biology, genetics, and psychology.

For a selection of reference input points $R = \{\mathbf{m}_k\}_{k=1}^K$ with $R \subseteq X$ and corresponding outputs $T = \{t_k\}_{k=1}^K$ with $T \subseteq Y$, define $\mathbf{D}_x \in \mathbb{R}^{N \times K}$ in such a way that its k th column contains the distances $d(\mathbf{x}_i, \mathbf{m}_k)$ between the $i = 1, \dots, N$ input points $\mathbf{x}_i$ and the k th reference point $\mathbf{m}_k$. Analogously, define $\Delta_y \in \mathbb{R}^{N \times K}$ in such a way that its k th column contains the distances $\delta(\mathbf{y}_i, t_k)$ between the N output points $\mathbf{y}_i$ and the output t_k of the k th reference point. The mapping g between the input distance matrix $\mathbf{D}_x$ and the corresponding output distance matrix Δ_y can be reconstructed using the multiresponse regression model

$$\Delta_y = g(\mathbf{D}_x) + \mathbf{E}.$$

The columns of the matrix $\mathbf{D}_x$ correspond to the K input vectors and the columns of the matrix Δ_y correspond to the K response vectors, the N rows correspond to the observations. The columns of the $N \times K$ matrix $\mathbf{E}$ correspond to the K residuals.

Assuming that mapping g between input and output distance matrices has a linear structure for each response, the regression model has the form

$$\Delta_y = \mathbf{D}_x \mathbf{B} + \mathbf{E}. \quad (1)$$

The columns of the $K \times K$ regression matrix $\mathbf{B}$ correspond to the coefficients for the K responses. The matrix $\mathbf{B}$ can be estimated from data through a minimization of the multivariate residual sum of squares as loss function:

$$\text{RSS}(\mathbf{B}) = \text{tr}((\Delta_y - \mathbf{D}_x \mathbf{B})' (\Delta_y - \mathbf{D}_x \mathbf{B})). \quad (2)$$

Under the normal conditions where the number of equations in Eq. (1) is larger than the number of unknowns, the problem is overdetermined and, usually, with no solution. This corresponds to the case where the number of selected reference points is smaller than the number of available points (i.e. $K < N$). In the case, we must rely on the approximate solution provided by the usual least squares estimate of $\mathbf{B}$,

$$\hat{\mathbf{B}} = (\mathbf{D}_x' \mathbf{D}_x)^{-1} \mathbf{D}_x' \Delta_y. \quad (3)$$

If in Eq. (1) the number of equations equals the number of unknowns (i.e. $K = N$ because all the learning points are also reference points), then the problem is uniquely determined and has a single solution if the matrix $\mathbf{D}_x$ is full-rank. In this case,

$$\hat{\mathbf{B}} = \mathbf{D}_x^{-1} \Delta_y. \quad (4)$$

Clearly less interesting is the case where in Eq. (1) the number of equations is smaller than the number of unknowns (i.e. for $K > N$, corresponding to the situation where, after selecting the reference points, only a smaller number of learning points used). This case leads to an underdetermined problem with, usually, infinitely many solutions (de Souza Jr et al., 2015).

The sums of squares associated with any term in any linear model (i.e. for MANOVA, MANCOVA, or multivariate regression) can be calculated directly from a distance matrix. This is because, for any

centered data matrix $Y_{(n \times p)}$ (of n sample units for each of p variables), the relevant information contained in the $(p \times p)$ inner product matrix $Y'Y$ (used in classical multivariate analysis) is also contained in the $(n \times n)$ outer product matrix YY' . In addition, an outer product matrix can be obtained from any $(n \times n)$ distance matrix (Gower 1966), thus allowing the analysis to be based on a distance measure of choice, including semimetric measures like Bray-Curtis.

Let $X(n \times m)$ be a model (aka design or regression) matrix, with m the number of parameters. Traditional multivariate analysis proceeds through partitioning of the $(p \times p)$ total sum of squares and cross products (SSCP) matrix, which is the inner product $Y'Y$. The total sum of squares (S_T) is the trace, or sum of diagonal elements (sums of squares for each variable) in this matrix, which we will symbolize by $tr(Y'Y)$. Partitioning can be done according to the linear model $Y=X\beta + \epsilon$, where β is the matrix of model parameters, ϵ is the matrix of errors, and we wish to test $H_0 : \beta = 0$. The least-squares solution for β is $B=(X'X)^{-1}X'Y$. The matrix of fitted values is $\hat{Y}=XB =HY$, where H is the idempotent ‘‘hat’’ matrix $X(X'X)^{-1}X'$ (e.g., Johnson and Wichern 1992). So the matrix of residuals is $R=Y-\hat{Y}=(I-H)Y$. The total SSCP matrix is thus partitioned into predicted (model) and residual SSCP matrices in the following manner: $Y'Y=\hat{Y}'\hat{Y} + R'R$. Also, S_T is partitioned into hypothesis sums of squares $S_H = tr(\hat{Y}'\hat{Y})$ and residual sums of squares $S_R = tr(R'R)$, as follows:

$$tr(Y'Y) = tr(\hat{Y}'\hat{Y}) + tr(R'R) \quad (5)$$

An appropriate statistic to test the null hypothesis of no effect of the model parameters is a pseudo F statistic:

$$F = \frac{tr(\hat{Y}'\hat{Y})/(m-1)}{tr(R'R)/(n-m)} \quad (6)$$

Note that if there is only one variable, then Eq. 6 reduces to Fisher’s univariate F statistic. For a nonparametric test, the P value may be obtained as $P = P(F^* > F)$ where F^* is the value of F obtained by a random equiprobable permutation across the n units. The degrees of freedom $(m-1)$ and $(n-m)$ are not necessary in Eq. 6 for the test by permutation, as they remain constant (McArdle and Anderson, 2001).

Euclidean Distance

Euclidean distance is the most commonly used distance measure to measure similarity between two units and is based on the length of a straight line to be drawn between two units.

By using the Euclidean distance measure, the distance between two units, with n being the number of units and p being the number of variables; $i, j = 1, 2, 3, \dots, n$, i and j distance from each other

$$(d_{i,j}) = \sqrt{(x_{i1} - x_{j1})^2 + (x_{i2} - x_{j2})^2 + \dots + (x_{ip} - x_{jp})^2}$$

it is calculated with the formula (Dinler, 2014).

It has been proven that this method is compatible with the classical linear regression model when the Euclidean distance measure is used.

Manhattan Distance

It is the sum of the distances between units according to their dimensions. In the measurement of the distance between two objects in two-dimensional space, the hypotenuse of the triangle shown inside shows the Euclidean distance. The sum of the lengths of the sides of this triangle outside the hypotenuse gives the Manhattan City Block distance. It is generally recommended to be used for variables which are discrete quantitative data. It is calculated as (Alpar, 2013):

$$d_{i,j} = \sum_{k=1}^p |x_{ik} - x_{jk}|$$

Manhattan City Block distance is a measure of distance that is less sensitive to outliers (Timm, 2002).

Gower Distance

The Gower distance was proposed by Gower in 1971. The most basic feature of the Gower distance is that it can be used in a data set with both categorical and continuous data. The Gower distance is calculated using standardized data. The Gower distance is calculated with a separate formula when only continuous data are used. The distance used for the data set with both categorical and continuous data is called the Gower general similarity measure adlandırılır (URL2). Gower's general similarity measure for categorical variables,

$$S_{ij} = \frac{\sum_{k=1}^p W_{ijk} S_{ijk}}{\sum_{k=1}^p W_{ijk}}$$

expressed in the form. Here S_{ijk} , is the similarity measure between i . and j . observations according to k . variable. W_{ijk} is takes 0 when there is no variable value and 1 other situations in comparison i . and j . observations for variable k . Gower (1971), the similarity measure for continuous variables in the data;

$$S_{ij} = 1 - \frac{|x_{ik} - x_{jk}|}{R_k}$$

defined in the form. Here R_k , is determined as rank of i . and j . observations of variable k (Servi, 2009).

In this study, the R program was used to analyze the data. Codes (Boj et al., 2016) and data sample for this analysis are given below.

Example for continuous data;

	DMC	CC	NDF
2	41	60	
1.87	52	64	
1.96	47	61	

Example for discrete data;

	BT	Gender	BW
1	2	3,6	
1	1	4	
2	2	2,7	

The sample code sequence used in the analysis;

```
library("dbstats")
library("cluster")
mydata <- read.delim(file.choose())
mydata.active <- mydata
analiz <- dbglm(formula = DMC ~ CC + NDF, data = mydata, family = gaussian(),
```

```
method = "GCV", full.search = TRUE, metric = "euclidean", weights = NULL,  
range.eff.rank = c(1, 2))  
  
summary(analiz)  
  
glm1 <- glm(DMC ~ CC + NDF, data = mydata, family = gaussian())
```

RESULTS

The results obtained from the analysis using Euclidean, Manhattan and Gower distance measures and continuous variables are given in Table 1.

Table 1. Comparison criteria from distance measures for continuous data

Distance	AIC	BIC	GCV
Euclidean	-254.16	-249.09	0
Manhattan	-169.52	-164.45	0
Gower	-153.92	-150.54	0

According to the results, it was determined that the model obtained by using the Euclidean distance measure, which produced the lowest AIC and BIC values, was the best model. Generalized cross validation (GCV) was calculated as zero in all models and could not be used for comparison purposes.

The model;

$$\text{DMC} = 3.649 - 0.0006 \cdot \text{CC} - 0.027 \cdot \text{NDF}$$

was obtained. The reason for the higher fit of the model obtained by using the Euclidean distance, which is based on the length of a straight line to be drawn between two units, may be because the data is continuous. This may be due to the linearity of the model.

The results obtained from the analysis using Euclidean, Manhattan and Gower distance measures and discrete variables are given in Table 2.

Table 2. Comparison criteria from distance measures for discrete data

Distance Measure	AIC	BIC	GCV
Euclidean	142.68	149.9	0.32
Manhattan	142.96	150.18	0.32
Gower	142.68	149.9	0.32

According to the results obtained, it was determined that the model obtained using Euclidean and Gower distance measures, which produced the lowest AIC and BIC values, were the best. Generalized cross validation (GCV) was calculated as 0.32 in all models and could not be used for comparison purposes.

The model;

$$\text{BW} = 5.93 - 1.14 \cdot \text{BT} - 0.42 \cdot \text{Gender}$$

was obtained. Since this dataset contains discrete data, although it produced the same results as the Euclidean distance, it may be useful to prefer the Gower distance. This result may be dependent on the data set. For this reason, when it is desired to establish a model with discrete data, the preference of Gower distance can be recommended.

CONCLUSION

The comparison of Euclid, Manhattan and Gower distance measures within the scope of distance based regression aimed to this study showed that use of distance based regression may be give more reliable results especially existence of discrete explanatory variables.

References

- Alma GÖ, Vupa Ö, 2008. Regresyon analizinde kullanılan en küçük kareler ve en küçük medyan kareler yöntemlerinin karşılaştırılması. *SDÜ Fen Edebiyat Fakültesi Fen Dergisi*, 3(2): 2019-229.
- Alpar R. 2010. Basit doğrusal regresyon çözümlemesi: spor, sağlık ve eğitim bilimlerinden örneklerle uygulamalı istatistik ve geçerlik-güvenirlilik. 285-304s, Detay Yayıncılık, Ankara, Turkey.
- Alpar R, 2013. Uygulamalı Çok Değişkenli İstatistiksel Yöntemler. 858s, Detay Yayıncılık, Ankara, Turkey.
- Arenas C, Cuadras CM, 2002. Recent statistical methods based on distances. *Contributions to Science*, 2(2): 183-191.
- Boj E, Caballe A, Delicado P, Esteve A, Fortiana J. 2016. Global and local distance-based generalized linear models. *TEST*, 25: 170-195.
- Damodar, N. 2001. Temel Ekonometri, 2. Basım, 192s, Literatür Yayıncılık, İstanbul, Turkey.
- de Souza Jr AH, Corona F, Barreto GA, Miche Y, Lendase A, 2015. Minimal learning machine: A novel supervised distance-based approach for regression and classification. *Neurocomputing*, 164: 34-44.
- Dinler M, 2014. Kümeleme analizi yöntemlerinin hayvancılık verilerinde karşılaştırılması olarak incelenmesi. Yüksek Lisans Tezi, Bingöl Üniversitesi, Fen bilimleri Enstitüsü, Bingöl, Turkey.
- Li J, Zhang W, Zhang S, Li Q, 2019. A theoretic study of a distance-based regression model. *Sci China Math*, 62: 979-998.
- McArdle BH, Anderson MJ, 2001. Fitting multivariate models to community data: a comment on distance-based redundancy analysis. *Ecology*, 82(1): 290-297.
- Okur S, 2009. Parametrik ve parametrik olmayan doğrusal regresyon analiz yöntemlerinin karşılaştırılması olarak incelenmesi. Yüksek Lisans Tezi, Çukurova Üniversitesi, Fen Bilimleri Enstitüsü, Adana, Turkey.
- Orhunbilge N, 2017. Uygulamalı regresyon ve korelasyon analizi, 3. Basım, 386s, Nobel Yayıncılık, Ankara, Turkey.
- Servi T, 2009. Çok değişkenli karma dağılım modeline dayalı kümeleme analizi. Doktora Tezi, Çukurova Üniversitesi, Fen bilimleri Enstitüsü, Adana, Turkey.
- Timm NH, 2002. Applied Multivariate Analysis. 185-309s, Springer-Verlag, New York, USA.
- URL1: http://istatistikanaliz.com/regresyon_analizi.asp (Access date: 18.06.2021).
- URL2: <http://emredunder.blogspot.com/2011/06/kumeleme-analizinde-kullanilan-baz.html> (Access date: 18.06.2021).
- Vural A, 2007. Aykırı değerlerin regresyon modellerine etkileri ve sağlam kestiriciler. Yüksek Lisans Tezi, Marmara Üniversitesi, Sosyal Bilimler Enstitüsü, İstanbul, Turkey.
- Weisberg S, 2005. Applied linear regression. 317s, John Wiley & Sons, Inc. Hoboken, New Jersey, USA.

Moleküler Verilerde AMOVA'nın Kullanımı

Malik ERGİN^{1*}, Özgür KOŞKAN¹

¹Isparta Uygulamalı Bilimler Üniversitesi Ziraat Fakültesi Zootečni Bölümü, 32200, Isparta, Türkiye

*Sunucu yazar e-posta: malikergin@isparta.edu.tr

Özet

Populasyonlar arasındaki genetik farklılıkların belirlenmesi ve elde edilen bulguların karşılaştırılması populasyon genetiği çalışmalarının amaçlarından biridir. Populasyonlarda görülen genetik farklılaşma mutasyon, göç, seleksiyon ve şans gibi evrimsel süreçlerin etkisiyle meydana gelmektedir. Birçok türde olduğu gibi çiftlik hayvanlarında da geçmişten günümüze kadar geçen süreçte şekillenen genetik çeşitlilik türlerin ve ırkların sürdürülebilir kullanımı ve gelecekte değişmesi muhtemel çevre koşullarında yapılacak ıslah çalışmaları için oldukça önem taşımaktadır. Bir türdeki genetik çeşitlilik o türü oluşturan tüm ırklar ve alt gruplardaki (tip, ekotip vb.) genetik çeşitliliğin toplamını yansıtmaktadır. Populasyonlar arasındaki genetik farklılaşmanın ölçüsü olarak sabitleme indeksleri (fixation index) kullanılmaktadır. Geçmişten günümüze birçok genetik farklılaşma istatistikleri öne sürülmüş olmakla birlikte, Wright tarafından bildirilen F_{ST} indeksi en popüler ve ilk kez öne sürülen farklılaşma indeksidir. F_{ST} istatistiğine benzer şekilde G_{ST} , R_{ST} , D vb. indeksler de yer almaktadır. Bu indeksler genel olarak allel frekanslarıyla tahmin edilmekte ve populasyonlarda meydana gelen evrimsel süreçlerin sebep olduğu mutasyonel durumlarda yetersiz kalabilmektedir. Bundan dolayı moleküler varyans analizi olarak da bilinen AMOVA'nın kullandığı Φ (Phi) indeksi, birçok moleküler marker verisi (AFLP, RFLP, Mikrosatellit, SNP) ile uyum gösteren önemli bir genetik farklılaşma ölçüsüdür. Bu derlemede G_{ST} , R_{ST} , D indekslerinden de kısaca bahsedilmiş olup daha çok AMOVA testine ve Φ (Phi) indeksine vurgu yapılmıştır.

Anahtar kelimeler: Moleküler varyans analizi, genetik farklılaşma, genetik çeşitlilik

Usage of AMOVA in Molecular Data

Abstract

Identifying genetic differences between populations and comparing the results is one of the aims of population genetics studies. Genetic differentiation observed on the populations occurs with the effect of evolutionary processes such as mutation, migration, selection and chance. As with many species, genetic diversity in livestock that have taken shape in the from past to the present period is essential for the sustainable use of species, breeds and for breeding studies in environmental conditions that are likely to change in the future. Genetic diversity in a species reflects the sum of genetic diversity in all breeds and subgroups (type, ecotype, etc.) that comprise species. Fixation indexes are used as a measure of genetic differentiation between populations. Although many genetic differentiation statistics have been suggested from past to the present, the F_{ST} index reported by Wright is the most popular and first time differentiation index. Similar to F_{ST} statistics, G_{ST} , R_{ST} , D etc. indexes are also included. These indexes are generally estimated with allele frequencies and may be inadequate in mutational cases caused by evolutionary processes occurring in populations. Therefore, the Φ (Phi) index used by AMOVA, also known as analysis of molecular variance, is an important measure of genetic differentiation that align with many molecular marker data (AFLP, RFLP, Microsatellite, SNP). In this

review, G_{ST} , R_{ST} , D indexes are also briefly mentioned and more emphasis is placed on AMOVA test and index Φ (Φ).

Key words: Molecular variance analysis, genetic differentiation, genetic diversity

GİRİŞ

Birçok canlı türünde grup veya populasyon içindeki allel frekansları, genetik sürüklenme veya bölgesel adaptasyonlar sebebiyle farklı izler taşımaktadır ve oluşacak yeni mutasyonlar bulundukları populasyonların ötesine serbestçe yayılamazlar. Bundan dolayı gruplar veya populasyonlar arası farklılık bu durumdan kaynaklanmaktadır (Mona ve Bertorelle, 2013). Bunun yanı sıra populasyonların arasındaki genetik farklılıkların sebeplerine genetik sürüklenme, seleksiyon ve göç de etki etmektedir. Populasyonlar arasındaki genetik farklılıkların ortaya çıkarılmasında günümüze kadar birçok farklı yöntem önerilmiştir. Wright'ın F istatistiği, genetik farklılığın tespit edilmesini sağlayan en popüler ve ilk kez geliştirilen yöntemdir. F istatistiğine benzer şekilde çalışan ve çalışma prensibinde küçük farklılıklar bulunduran yöntemler de farklı durumlar ve çeşitli genetik markerler için önerilmektedir (Mona ve Bertorelle, 2013). Bazı istatistiksel yöntemler yalnızca alel frekanslarına göre işlem yaparken bazıları da aleller arasındaki evrimsel süreçlerin sebep olduğu uzaklıklara göre işlem yapmaktadır (Mona ve Bertorelle, 2013). Bu sebeple mutasyon olayları sonucunda oluşan genetik farklılıkların tespit edilmesinde ikinci kategoride belirtilen özellikteki istatistiklerin kullanılması daha uygun olacaktır.

Birçok türde olduğu gibi çiftlik hayvanlarında da geçmişten günümüze kadar geçen süreçte şekillenen genetik çeşitlilik türlerin ve ırkların sürdürülebilir kullanımı ve gelecekte değişmesi muhtemel çevre koşullarında yapılacak ıslah çalışmaları için oldukça önem taşımaktadır. Bir türdeki genetik çeşitlilik o türü oluşturan tüm ırklar ve alt gruplardaki (tip, ekotip vb.) genetik çeşitliliğin toplamını yansıtmaktadır. Genetik çeşitlilik; belirli türler, alt türler, populasyonlar veya ırklardaki genlerin çeşitliğidir ve bu genlerin var olup olmaması ile ölçülmektedir. Bir tür veya populasyon içerisindeki bireyler birbirinden az veya çok farklılık göstermektedirler. Eğer bu bireyler arasında akrabalık seviyeleri azalmaktaysa, birbirleri arasındaki genetik farklılık artmaktadır (Fadhil ve ark., 2018). Bu genetik uzaklığın ölçülmesinde kullanılan terim akrabalık katsayısı (F), bir bireydeki söz konusu gende yer alan iki alelin homozigot olup olmadığı yani bir diğer ifadeyle aynı atadan gelip gelmediğinin olasılığıdır. Rastgele çiftleşen populasyonlarda iki aleli bulunan bir genin (örnek olarak, A ve a) frekansları sırasıyla p ve q ($p + q = 1$) olarak bilinmektedir. Yani iki aleli bulunan bir lokus için birinci alel frekansı p , diğer alternatif alel frekansı $(1-p)$ olmaktadır (Hartl, 1988).

Büyük populasyonların daha küçük populasyonlara ayrılması sebebiyle küçük populasyonlarda homozigotluk seviyesinde artış, diğer bir ifade ile bireyler arasındaki akrabalık seviyesinde yükselme meydana gelmektedir. Küçülmeye gitmiş, yani alt populasyonlarda ana populasyonlara göre örneklem hatası daha büyük olmaktadır (Keskin, 2019). Tüm bu sebeplerden dolayı, alt populasyonlarda genetik sürüklenme sonucunda alel frekanslarındaki değişimler çok daha hızlı meydana gelmektedir.

Genetik Farklılığın Tespit Edilmesinde Kullanılan Bazı Yöntemler

Wright'ın F İstatistikleri

F istatistiği, bir populasyondaki alt gruplar arasında ortalama heterozigotluk ile populasyonun tüm bireylerinin mevcut heterozigotluk frekansı arasındaki farkın bir ölçüsüdür. Başka tanımlamalar yapmak gerekirse F_{ST} , gametler arasındaki genetik olarak bir korelasyon ölçüsü veya söz konusu lokusta populasyonlar arasındaki alel frekanslarının varyansı ve atasal benzerlik olasılığı olarak da ifade edilebilmektedir. F değeri aşağıdaki formüldeki gibi ifade edilmektedir.

$$F_{ST} = \frac{\sigma_S^2}{\sigma_T^2} = \frac{\sigma_S^2}{\bar{p}(1-\bar{p})}$$

σ_s^2 : Alt populasyonlarda frekansların toplanmış varyansı

σ_T^2 : Toplam populasyondaki (T) allel varyansını temsil etmektedir.

F değeri teorik olarak 0 ile 1 arasında değişmektedir. Eğer F değeri 0 ise bütün populasyon ve onun alt populasyonları arasında bir farklılığın bulunmadığı anlaşılmaktadır. Fakat, çoğu durumda birbirinden çok yüksek seviyede ayrılma gösteren popülasyonlarda bile F değeri 1'den az çıkmaktadır. Fiksasyon indekslerinin (F istatistikleri) hesaplanabilmesi için öncelikle populasyon yapısında bulunan her seviye grupları için alel frekanslarının belirlenmesinin ardından ortalama heterozigotluk seviyeleri hesaplanmalıdır.

Çizelge 1. Populasyonun farklı grup seviyeleri için ortalama heterozigotluk

Populasyondaki Grup Seviyesi	Heterozigotluk
Alt populasyonlar için	$H_S = \sum_i \sum_j 2p_{ij} (1 - p_{ij})$
Alt populasyonların grupları için	$H_G = \sum_i 2 \frac{\sum_j p_{ij}}{n_i} \left(1 - \frac{\sum_j p_{ij}}{n_i}\right)$
Genel populasyon için	$H_T = \sum 2 \frac{\sum_i \sum_j p_{ij}}{n_i} \left(1 - \frac{\sum_i \sum_j p_{ij}}{n_i}\right)$

Genel populasyon içerisindeki alt populasyon gruplarının akrabalık derecesi,

$$F_{SG} = \frac{H_G - H_S}{H_G}$$

Gruplar içerisindeki alt populasyonların akrabalık derecesi,

$$F_{GT} = \frac{H_T - H_G}{H_T}$$

Genel populasyondaki alt grupların akrabalık derecesi,

$$F_{ST} = \frac{H_T - H_S}{H_T} \quad \text{olarak hesaplanmaktadır.}$$

NEI'n GST İstatistiği

Nei (1977) tarafından ortaya çıkarılan G istatistiği, F istatistiğine benzer şekilde populasyonlar arasındaki farklılaşmanın bir ölçüsüdür. G istatistiği çoklu aleller için daha uygun bir farklılaşma ölçüsüdür. S sayıda alt populasyonlara ayrılmış olan bir diploid populasyon olduğunu ve r kadar allel ($A_1, A_2, \dots, A_r$) bulunduğunu varsayalım. Bu durumda;

p_{ik} : A_k allelinin frekansı

p_{ikl} : $A_k A_l$ genotipinin frekansı olacaktır.

Nei (1977) G_{ST} istatistiğini aşağıdaki gibi tanımlamıştır.

$$G_{ST} = 1 - \frac{H_S}{H_T}$$

Bu formülde,

$$H_S = 1 - \sum_{k=1}^r \bar{p}_k^2$$

$$\bar{p}_k^2 = \sum_{i=1}^s w_i p_{ik}^2 \bar{p}^k = \sum_{i=1}^s w_i p_{ik}$$

$$H_T = 1 - \sum_{k=1}^r \bar{p}_k^2$$

w_i : alt populasyonların nispi büyüklüğü olarak tanımlanmıştır.

w_i genellikle $\frac{1}{S}$ olarak varsayılmaktadır.

H_S : Populasyon Hardy-Weinberg Dengesinde olduğu durumda alt populasyonlardaki heterozigotluk

H_T: Populasyon Hardy-Weinberg Dengesinde olduğu durumda genel populasyondaki heterozigotluk

SLATKIN R_{ST} İstatistiği

Slatkin'in R_{ST} istatistiği F_{ST}'ye benzer biçimde populasyonların göç oranlarını veya populasyonların farklılaşmasından bu yana geçen süreleri tahmin etmek için kullanılabilir (Slatkin, 1995). Mikrosatellitler oldukça yüksek seviyede polimorfizm gösteren markerlerdir. DNA sekansı içerisinde tekrar eden nükleotid dizilerinden oluşan mikrosatellitlerin analizleri oldukça kullanışlı olmakla birlikte populasyonların yapılarını tanımlamada oldukça bilgi vericidir. Slatkin (1994) R_{ST} istatistiğini ortaya ilk çıkardığında mikrosatellit verilerin populasyonlardaki genetik farklılaşmaları tespit etmekte oldukça kullanışlı bir yöntem olduğunu bildirmiştir.

$$R_{ST} = \frac{\bar{S} - S_W}{\bar{S}}$$

$$\bar{S} = \frac{2n-1}{2nd_s-1} S_W + \frac{2n(d_s-1)}{2nd_s-1} S_B$$

Burada d_s kadar populasyon içerisinde n sayıda birey örneklediğimizi varsayalım ve a_{ij}, j. populasyonda (j=1,2,...,d_s) i. bireyin (i=1,2,...,2n) allel büyüklüğünü göstermektedir.

Her bir populasyon içindeki allel büyüklükleri arasındaki farklılıkların ortalama kareler toplamı,

$$S_W = \frac{1}{d_s} \sum_{j=1}^{d_s} \frac{2}{2n(2n-1)} \sum_{i < i'} (a_{ij} - a_{i'j})^2$$

Her bir birey çifti arasındaki allel büyüklüğü farklılıklarının ortalama kareler toplamı,

$$S_B = \frac{2}{(2n)^2 d_s (d_s - 1)} \sum_{j < j'} \sum_{i < i'} (a_{ij} - a_{i'j'})^2$$

Jost'n D İstatistiği

Genetik farklılaşma ölçülerinden incelediğimiz yöntemler heterozigotluk seviyelerine dayanmaktayken Jost (2008) D istatistiğinde etkili allel sayısını kullanmaktadır. Bu yöntem çok büyük populasyonlarda genetik farklılaşmanın tespit edilmesinde kullanışlı bir yöntem olmakla birlikte aynı zamanda küçük ve büyük alt populasyonlardaki farklı genetik çeşitlilik dinamiklerinin dikkate alınmaması açısından bir dezavantaj olarak ifade edilmektedir (Doğan ve Doğan, 2016). Jost D İstatistiği bir lokus için aşağıdaki gibi ifade edilmektedir.

$$D = \left(\frac{k}{k-1} \right) \left(\frac{H_T - H_S}{1 - H_S} \right)$$

k: Alt populasyonların sayısı

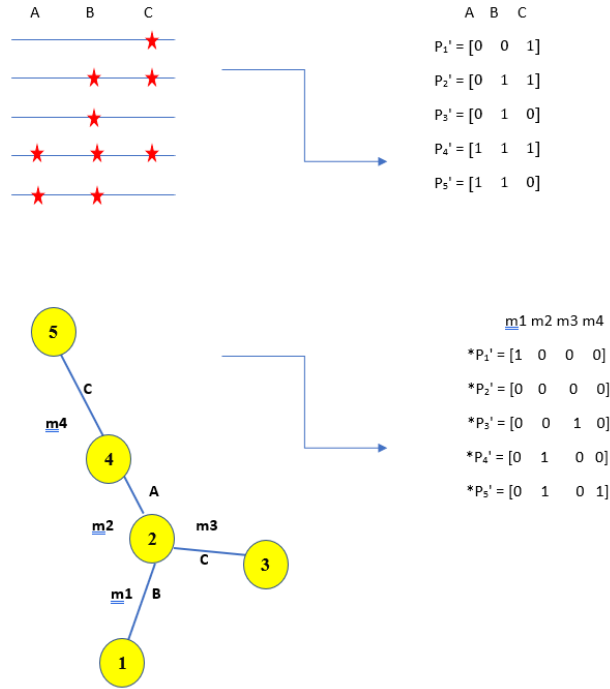
$$H_S = 1 - \sum_{k=1}^r \bar{p}_k^2$$

$$H_T = 1 - \sum_{K=1}^r \bar{p}_K^2$$

AMOVA

Moleküler Varyans Analizi olarak bilinen AMOVA, populasyonlar içi bireyler ve populasyonlar arası bireylerde grup ortalamalarından sapmaların genetik mesafe olarak hesaplandığı ve karesel sapmaların varyans olarak alındığı bir yöntemdir. AMOVA türler içerisinde moleküler olarak farklılığın tespit edilmesinde kullanılmaktadır. Genetik farklılıkların tespit edilmesinde önerilen ve yukarıda ifade ettiğimiz birçok yöntem (F_{ST}, D_{ST}, R_{ST} ve diğer farklı yöntemler) bireyler arasındaki allel frekanslarına göre karşılaştırma yapmasından dolayı genler arası mutasyon farklılıkları gibi durumlarda yetersiz kalabilmektedir. Genler arasında mutasyon farklılıklarının tespit edilmesinde bu yöntemlerin

yetersizliğinden dolayı Excoffier (1992) tarafından moleküler verilerin varyans analizi olarak bilinen AMOVA geliştirilmiştir. Moleküler varyans analizinde birçok çeşit moleküler veri – moleküler marker verileri (RFLP, AFLP gibi) veya moleküler verilere göre oluşturulmuş filogenetik ağaçlar- analiz edilebilmektedir. Mikrosatellitler, Tek Nükleotid Polimorfizmleri (SNPs) ve tüm sekans verileri de AMOVA’da kullanılmaktadır. AMOVA, her bir moleküler veriyi bir Boolean vektörü (p_i) olarak algılar. Boolean vektörü 1 ve 0’lardan oluşan ve markerin yokluğunu 0, varlığını 1 olarak belirten bir vektördür (Excoffier ve ark., 1992). Boolean vektörünün çalışma prensibi Şekil 1’deki gibi gösterilmiştir.



Şekil 1. Boolean vektörünün oluşturulma aşamasına bir örnek. (Şekil Excoffier ve ark. (1992) tarafından yapılan çalışmadan uyarlanmıştır).

Öncelikle her bir haplotip bir boolean vektörüne (p_i) dönüştürülür. Bu dönüşümde markerin varlığı veya yokluğuna göre 1 ve 0 sayılarına göre değerler alırlar. Söz konusu grafikte restriksiyon bölgelerinin varlığı veya yokluğuna göre dönüşüm yapılmıştır. Yapılacak ikinci aşama ise her bir haplotipin diğer haplotiplerle ortak mutasyon ilişkilerini ortaya çıkaran bir filogenetik ağ ortaya çıkarmaktır. En son adım ise bir referans haplotip (h_j) referans olarak seçilir ve mutasyon durumlarının (m_i) oluşumları her bir haplotip için boolean vektörü ($*p_j$) olarak kodlanır.

Her bir haplotipin ikili formattaki (binary) vektörünü diğer haplotip vektöründen ($p_j - p_k$) formülü kullanarak çıkarılmasıyla vektör çiftleri arasındaki Öklid uzaklıkları hesaplanmaktadır. İki vektör arasındaki karesel Öklid uzaklığı aşağıdaki formül aracılığıyla hesaplanır;

$$p' = [p_1 p_2 p_3 \dots p_s]$$

p_1 : sekansın mutasyonun bulunduğu bölgede yer alması durumunda 1, yer almaması durumunda 0 olmasıdır.

2 sekans arasındaki farklılık,

$$(p_j - p_k)' = [(p_{1j} - p_{1k})(p_{2j} - p_{2k}) \dots (p_{sj} - p_{sk})] \text{ olarak gösterilmiştir.}$$

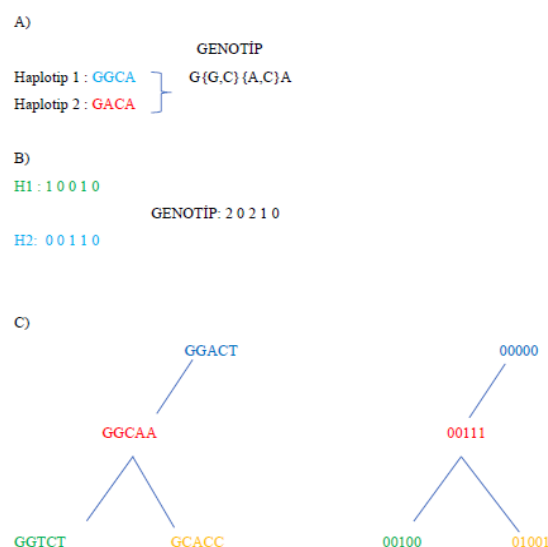
İki sekans arasındaki öklid uzaklığı,

$$\delta_{jk}^2 = (p_j - p_k)' W (p_j - p_k) \text{ formülü ile ifade edilmektedir.}$$

Eşitlikte yer alan W , bir ağırlık matrisidir. Birim matrise eşit olduğu varsayıldığından sonucu değiştirmemektedir. Fakat, bazı durumlarda ağırlık matrisi (W), transisyon, transversiyon ve indellere (insersiyon ve delesyon) göre farklı değerler alabilmektedir (Excoffier, 2004). Boolean vektörlerinin tüm ikili dönüşümleri için karesel öklid uzaklıkları hesaplanır ve bunlar daha sonra her birinin populasyon içinde bir alt grubu temsil ettiği alt matrislere dönüştürülerek genel bir matris içerisinde birleştirilir. Bireyler arası moleküler sekans bilgileri ile oluşturulan vektörler 1 ve 0 olarak kodlanmış ve karesel Öklid uzaklıkları bu vektörlere göre hesaplandığı için bir normal dağılım söz konusu olmamaktadır. Hesaplanan bu uzaklık değerleri normal dağılım göstermemesi sebebiyle tekrar tekrar seçilerek bir parametrik olmayan permütasyona tabi tutulurlar. Her bir permütasyonda örnek sayısı sabit kalmak şartıyla populasyonun her bir bireyi rastgele olarak seçilmekte ve farklı farklı populasyonlara dağıtılmaktadır. Bu aşamada permütasyon testi gerçekleştirilmekte ve boş bir dağılım elde edilmektedir. Daha sonra hipotez testi, yeniden örneklemeye sonucunda oluşturulan boş dağılıma göre test edilmektedir (Keskin, 2019). AMOVA uygulanırken populasyon hakkında birtakım varsayımların kabul edilmesi söz konusudur. Bunlardan ilki, bireylerin haplotiplerinin rastgele olarak seçildiği, birbirinden bağımsız olduğudur. İkincisi ise populasyondan örneklenen bireyler arasında akrabalık olmadığı ve rastgele çiftleşmenin söz konusu olduğu varsayımdır (Excoffier ve ark., 1992). Hiyerarşik bir varyans analizi, populasyon seviyelerinin bulunması ve bu seviyeler arasındaki farklılıklar nedeniyle toplam varyansı kovaryans bileşenlerine bölmektedir. Kovaryans bileşenleri (σ_i^2) fiksasyon indekslerinin hesaplanması için kullanılmaktadır (Excoffier ve Lischer, 2010).

Diploid ve Haploid Veriler

Haplotipik veriler ile anlatılmak istenen, allellerin haplotip formda bulunmasıdır. Bu haplotipik veriler haploid genomlardan (mtDNA, Y Kromozomu, Prokaryotlar) elde edilebileceği gibi diploid genomlardan da elde edilebilmektedir. Genotipik veri yapısı ile anlatılmak istenen ise diploid genomlardan elde edilmiş olmasıdır (Excoffier ve Lischer, 2010). Diploid genomlarda 2 haplotipin bir genotipi oluşturması söz konusu olduğu için bu durumda veriler uygun bir formata dönüştürülmelidir. Her bir haplotip sekansı ikili (binary) veri formatı şeklinde matris vektörlerini oluşturacaktır. Şekil 2’de haplotip ve genotip sekansları arasındaki ikili kodlama ilişkileri basit olarak anlatılmaya çalışılmıştır. Şekildeki örnek gen dizilimleri tamamen birbirinden bağımsız olup sadece örnek vermek amacıyla seçilmiştir.



Şekil 2. A) Bir bireyin genotipini oluşturan iki haplotip dizisi. B) Haplotiplerin birbirleri arasındaki farklılıklara göre ikili (binary) formatta kodlanması ve genotipin kodlanma düzeneği C) Her bir Haplotipin filogenetik ağaç üzerindeki ilişkisi.

Örnek vermek gerekirse Şekil 3'te gösterildiği gibi üç genotipten oluşan bir A matrisi oluşturulduğunu varsayalım. Bu genotiplerin meydana geldiği haplotip sekansları ikili formatta (1 ve 0'lardan oluşan) vektörler ile matriste yerleştirilmelidir.

$$A = \begin{bmatrix} 2 & 0 & 2 & 1 & 0 \\ 2 & 2 & 0 & 2 & 0 \\ 0 & 2 & 2 & 2 & 2 \end{bmatrix} \quad B = \begin{bmatrix} 1 & 0 & 0 & 1 & 0 \\ 0 & 0 & 1 & 1 & 0 \\ 1 & 0 & 0 & 1 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 1 & 0 \end{bmatrix}$$

Şekil 3. Farklı sekanslara sahip genotipleri oluşturan bir A matrisi, her satırda 1 ve 0'lardan oluşan ikili formdaki vektörlerin oluşturduğu bir B matrisine dönüştürülür.

Bu şekilde oluşturulan vektörlerin karesel Öklid uzaklıkları ile hesaplanan unsur nükleotidler arasındaki değil bireylerin sekansları arasındaki mutasyon sayısıdır. Çizelge 1,2 ve 3'te haplotipik ve genotipik veriler için varyans analiz tablolarının genel yapısı gösterilmiştir. Moleküler varyans analizi için kullanılacak olan matematiksel model aşağıdaki gibi düzenlenebilmektedir

$$x_{ijk} = x + a_k + b_{jk} + c_{ijk}$$

x_{ijk} = k. Grubun içindeki j. Populasyonda bulunan i. Haplotip frekans vektörü

x : x_{ijk} 'nın vektörü

a : Grup etkisi

b : Populasyon etkisi

c : Haplotip etkisi

σ_a^2 = Grup etkisi için kovaryans

σ_b^2 = Populasyon etkisi için kovaryans

σ_c^2 = Haplotip etkisi için kovaryans

σ^2 = Toplam moleküler varyans,

$$\sigma^2 = \sigma_a^2 + \sigma_b^2 + \sigma_c^2$$

Belirli sayıdaki birey için toplam karesel sapmaların toplamı geleneksel varyans analizine benzer şekilde olmasına rağmen, moleküler varyans analizinde toplam karesel sapmalar populasyonun alt seviyelerine indirgenmektedir.

$$SSD (\text{Toplam}) = SSD (\text{Gruplar Arası}) + SSD (\text{Gruplar İçi})$$

$$SSD (\text{Toplam}) = SSD (\text{Populasyon İçi}) + SSD (\text{Grup İçi Populasyonlar Arası}) + SSD (\text{Gruplar Arası})$$

$$SSD (\text{Populasyon İçi}) = \sum_g \sum_i \frac{\sum_j \sum_k S_{jk}^2}{2N_{ig}}$$

$$SSD (\text{Grup İçi Populasyonlar Arası}) = \sum_{g=1}^G \left(\frac{\sum_{i=1}^{I_g} \sum_{i'=1}^{I_g} \sum_{j=1}^{N_{ig}} \sum_{k=1}^{N_{ig}} S_{jk}^2}{2N_{ig}} - \sum_{i=1}^{N_{ig}} \frac{\sum_j \sum_k S_{jk}^2}{2N_{ig}} \right)$$

$$SSD \text{ (Gruplar Arası)} = \frac{1}{2N} \sum_{j=1}^N \sum_{k=1}^N S_{jk}^2 - \sum_{g=1}^G \frac{\sum_{i=1}^{I_g} \sum_{i'=1}^{I_g} \sum_{j=1}^{N_{ig}} \sum_{k=1}^{N_{ig}} S_{jk}^2}{2N_{ig}}$$

SSD (T): Toplam karesel sapmalar toplamı

SSD (AG): Populasyondaki Gruplar Arasında toplam karesel sapmalar

SSD (AP): Populasyonlar arasındaki toplam karesel sapmalar

SSD (AI): Bireyler arasındaki toplam karesel sapmalar

SSD (WP): Populasyon içi toplam karesel sapmalar

SSD (WI): Bireyler içi toplam karesel sapmalar

SSD (AP/WG): Grup İçi Populasyonlar Arası toplam karesel sapmalar

SSD (AI/WP): Grup İçi Bireyler Arası toplam karesel sapmalar

G: Grup Sayısı

P: Toplam populasyon sayısı

N: Genotipik veya haplotipik veri için toplam birey sayısı

N_p: Genotipik veya haplotipik veri için p. populasyondaki birey sayısı

N_g: Genotipik veya haplotipik veri için g. gruptaki birey sayısı

$$n = \frac{2N - \sum_P \frac{2N_p^2}{N}}{P-1}$$

$$F_{ST} = \frac{\sigma_a^2}{\sigma_T^2}, F_{IT} = \frac{\sigma_a^2 + \sigma_b^2}{\sigma_T^2}, F_{IS} = \frac{\sigma_b^2}{\sigma_b^2 + \sigma_c^2}$$

$$n = \frac{2N - \sum_{g \in G} \sum_{p \in P} \frac{2N_p^2}{N_g}}{P-G}$$

$$n' = \frac{\sum_{g \in G} \frac{(N-N_g)}{N_g} \sum_{p \in g} 2N_p^2}{N(G-1)}$$

$$n'' = \frac{2N - \sum_{g \in G} \frac{2N_g^2}{N}}{G-1}$$

$$F_{CT} = \frac{\sigma_a^2}{\sigma_T^2}$$

$$F_{IT} = \frac{\sigma_a^2 + \sigma_b^2 + \sigma_c^2}{\sigma_T^2}$$

$$F_{IS} = \frac{\sigma_c^2}{\sigma_c^2 + \sigma_d^2}$$

$$F_{SC} = \frac{\sigma_b^2}{\sigma_b^2 + \sigma_c^2 + \sigma_d^2}$$

$$S_G = \sum_{g \in G} \sum_{p \in g} \frac{N_p^2}{N_g}$$

$$n = \frac{N - S_G}{P-G}$$

$$n' = \frac{S_G - \sum_P \frac{N_p^2}{N}}{G-1}$$

$$n'' = \frac{N - \sum_G \frac{N_g^2}{N}}{G-1}$$

$$F_{CT} = \frac{\sigma_a^2}{\sigma_T^2}, F_{SC} = \frac{\sigma_b^2}{\sigma_b^2 + \sigma_c^2}, F_{ST} = \frac{\sigma_a^2 + \sigma_b^2}{\sigma_T^2}$$

Çizelge 1. Genotipik verilere göre tek seviyeye sahip populasyonlarda bireyler içi seviyeli AMOVA tablosu

VARYASYON KAYNAĞI	SERBESTLİK DERECELERİ	KARELER TOPLAMI(KT)	KARELER ORTALAMASI
Populasyonlar Arası	P-1	KT(AP)	$n\sigma_a^2 + 2\sigma_b^2 + \sigma_c^2$
Populasyon içi Bireyler Arası	N-P	KT(AI / WP)	$2\sigma_b^2 + \sigma_c^2$
Bireyler İçi	N	KT(WI)	σ_c^2
Genel	2N-1	KT(T)	σ_T^2

Çizelge 2. Genotipik verilere göre birden fazla seviyeye sahip populasyonlarda bireyler içi seviyeli AMOVA tablosu

VARYASYON KAYNAĞI	SERBESTLİK DERECELERİ	KARELER TOPLAMI(KT)	KARELER ORTALAMASI
Populasyon İçi Gruplar Arası	G-1	KT(AG)	$n'' + \sigma_a^2 + n'\sigma_b^2 + 2\sigma_c^2 + \sigma_d^2$
Grup İçi Alt Populasyonlar Arası	P-G	KT(AP / WG)	$n\sigma_b^2 + 2\sigma_c^2 + \sigma_d^2$
Alt Populasyonlar İçi Bireyler Arası	N-P	KT(AI/WP)	$2\sigma_c^2 + \sigma_d^2$
Bireyler İçi	N	KT(WI)	σ_d^2
Genel	2N-1	KT(T)	σ_T^2

Çizelge 3. Haplotipik verilere göre, birden fazla gruba sahip populasyonlar için AMOVA tablosu

VARYASYON KAYNAĞI	SERBESTLİK DERECELERİ	KARELER TOPLAMI(KT)	KARELER ORTALAMASI
Populasyon içi Gruplar arası	G-1	KT(AG)	$n'' + \sigma_a^2 + n'\sigma_b^2 + \sigma_c^2$
Grup içi Alt Populasyonlar Arası	P-G	KT(AP/WG)	$n\sigma_a^2 + \sigma_c^2$
Alt Populasyonlar İçi Bireyler Arası	N-P	KT(WP)	σ_c^2
Genel	N-1	KT(T)	σ_T^2

Varyans bileşenleri F_i istatistikleri (Φ , ϕ_i) kullanarak hesaplanabilmektedir. Φ istatistikleri F_i istatistiklerine benzer olup, populasyonun seviyeleri arasındaki farklılıkların dereceleri hakkında bilgi vermektedir. Aşağıdaki tabloda görüldüğü üzere, örneğin Φ_{ST} populasyonun ve populasyon alt grupları arasında bir hipotez kurulmasını sağlamaktadır. Örneğin, alt populasyonların varyansı genel populasyon varyansından önemli ölçüde farklılık göstermediği sonucuna varılırsa, alt populasyonların genel populasyondan farklı çıkacağı hipotez reddedilecektir.

Çizelge 4. Φ (Φ) istatistiklerinin elde edilişi

Populasyon Seviyeleri	Φ İSTATİSTİKLERİ
Grup içindeki alt populasyonlar arası	$\Phi_{SC} = \frac{\partial_b^2}{\partial_b^2 + \partial_c^2}$
Genel populasyondaki gruplar arası	$\Phi_{CT} = \frac{\partial_a^2}{\partial^2}$
Populasyon içindeki alt populasyonlar arası	$\Phi_{ST} = \frac{\partial_a^2 + \partial_b^2}{\partial^2}$

SONUÇ

Moleküler varyans analizi (AMOVA) genetik varyasyonun farklı seviyelerdeki populasyon gruplarında dağılımını ve bu populasyon gruplarını test etmek için kullanılabilecek bir yöntemdir. Standart varyans analizi (ANOVA) ile moleküler varyans analizi (AMOVA) arasındaki fark, populasyonun farklı seviyeleri için grup içi ve gruplar arası farklılıkların test edilebilmesini sağlamasıdır. Dolayısıyla standart ANOVA'dan farklı olarak kareler toplamı ve kareler ortalaması populasyonun farklı seviyeleri için özel olarak hesaplanmaktadır. Bu çalışma moleküler varyans analizinin çalışma prensiplerini incelemek ve ileride farklı yaklaşımların getirilebilmesi veya oluşacak olan analiz hatalarının nerelerden kaynaklanabileceği hakkında bilgi verici olması amacıyla gerçekleştirilmiştir. Moleküler varyans analizi (AMOVA) uygulamak için günümüzde Arlequin ve GenAlEx paket programları oldukça popüler olarak

kullanılmaktadır. Bu paket programların yanı sıra son zamanlarda biyoinformatik alanındaki gelişmelerin de etkisiyle birlikte R ve Python programlama dilleri içerisinde birçok kütüphane paketleri geliştirilmiş olup, bu paketler aracılığıyla moleküler varyans analizleri gerçekleştirilebilmektedir.

Kaynaklar

- Doğan İ, Doğan N, 2016. Statistical Measures for Genetic Differentiation: Review. *Turkiye Klinikleri J Biostat* 2016;8(2):180-6
- Excoffier L, Smouse P.E., Quattro M.J., 1992. Analysis of Molecular Variance Inferred From Metric Distances Among DNA Haplotypes: Application to Human Mitochondrial DNA Restriction Data. *Genetics*,131:479-491.
- Excoffier L, 2004. Analysis of Population Subdivision. *Handbook of Statistical Genetics*, 2nd Edition. Editörler: Balding D.J., Bishop M., Cannings C., John Wiley & Sons, Ltd., 713-750s.
- Hartl D, 1988. *A Primer of Population Genetics*. Sinauer Associates, 2.Baskı, 66s, Sunderland, Massachusetts, USA.
- Excoffier L, Lischer H, 2010. Arlequin suite ver 3.5: A new series of programs to perform population genetics analyses under Linux and Windows. *Molecular Ecology Resources*. 10: 564-567.
- Fadhil M, Zülkadir U, Aytekin İ, 2018. Genetic Diversity in Farm Animals. *Elixir Hormones and Signaling*, 117: 50032-50037
- Keskin S, 2019. Moleküler Varyans Analizi (AMOVA). II. HASAT ULUSLARARASI TARIM VE ORMAN KONGRESİ, 8-9 Kasım 2019, Sayfa:411-419, İzmir, Türkiye.
- Mona S, Bertorelle G, 2013. *Population Differentiation: Measures*. John Wiley & Sons, Ltd: Chichester.
- Nei M, 1977. F-statistics and analysis of gene diversity in subdivided populations. *Ann Hum Genet* 41(2):225-33
- Slatkin M, 1995. A Measure of Population Subdivision Based on Microsatellite Allele Frequencies. *Genetics* 139: 457-462.
- Jost L, 2008. GST and its relatives do not measure differentiation. *Mol Ecol* 17(18):4015-26s

Comparing the Performance of Eight Imputation Methods for Propensity Score Matching in Missing Data Problem

İmran KURT ÖMÜRLÜ¹, Buğra VAROL^{2*}, Mevlüt TÜRE¹

¹Adnan Menderes University, Faculty of Medicine, Division of Biostatistics, Aydın, Turkey

²Adnan Menderes University, Institute of Health Sciences, Division of Biostatistics, Aydın, Turkey

*Presenter e-mail: bugravarol87@gmail.com

Abstract

Propensity score (PS) is a popular method to control for covariates in observational studies. A challenge in PS analyses are missing values in covariates. Several strategies for handling missing values exist, but guidance in choosing the best method is needed. The aim of this study is to investigate how different imputation methods of handling missing values of covariates in a PS analysis can affect average treatment on the treated (ATT) estimates.

In this study, missing data imputation methods were evaluated using different data sets, whose covariates were low, medium, and high ($r=0.10, 0.50, 0.85$) correlated with each other, for $n=200$ units and 1000 times running simulation. Missing data structures were created according to the MAR mechanism and different missing rates. Different datasets were obtained after having imputed the missing values separately by eight imputation methods including mean, median, mode, hot deck, last observation carried forward (LOCF), next observation carried backward (NOCB), regression and predictive mean matching (PMM). Then the PS nearest neighbor matching was implemented to estimate ATT scores using the imputed data sets. Predictive performance of imputation methods was compared using hierarchical clustering analysis with Euclidean distance complete linkage.

ATT scores that obtained from mean, median, mode, hot deck, LOCF, NOCB, regression and PMM methods ranged from 3.223 to 3.366 for low correlated data sets, 2.920 to 2.996 for medium correlated data sets and, 2.539 to 2.647 for high correlated data sets. ATT scores of regression and PMM methods were closer to each other and these methods showed the best predictive performance especially in low and medium correlated data sets. The regression method outperformed slightly the PMM method, when the correlation among the covariates were high. Additionally, almost all imputation methods tended to produce more accurate ATT estimates when less data was missing and when there were larger amounts of missing data, the PMM was a best method of choice.

Ignoring missing values on covariates for PS analyses causes information loss significantly and this information loss becomes greater as the rate of missing data increases. PS analyses might be biased if missing data on covariates are ignored also. To prevent this information loss and bias, PS analyses should be performed after solving the problem of missing data with MAR mechanism on covariates by regression and PMM methods, which showed statistical superiority compared to other methods in this study.

Key words: Missing Data, Imputation, Simulation, Propensity Score, Hierarchical Clustering

Bulanık Yapay Sinir Ağları ve Bir Uygulama

Derviş TOPUZ^{1*}

¹Niğde Ömer Halisdemir Üniversitesi, Niğde Zübeyde Hanım Salık Hizmetleri MYO, Tıbbi hizmetler ve Teknikler Bölümü, Niğde, Türkiye

*Sunucu yazar e-posta: topuz@ohu.edu.tr

Özet

Doğru tahmin yöntemlerini belirlemek her sektör için en iyi kararların verilmesinde büyük öneme sahiptir. Sağlık hizmetleri, küresel bir sağlık ortamında rekabet gücünü korumak için gelişmiş bilgi işleme teknolojilerine giderek daha fazla bağımlı hale geliyor. Bu nedenle, veri madenciliği tekniklerini kullanarak teşhis ve tedavide tahmin problemi sağlıktaki en önemli konulardan biridir. Bu alan büyük bilimsel ilgi görmüş ve daha kesin bir tahmin süreci sağlamak için çok önemli bir araştırma alanı haline gelmiştir. Bulanık mantık (FL) ve Yapay Sinir Ağı (YSA), tahmin uygulamaları için geniş bir kapsamı olan heyecan verici ve umut verici bir tekniktir. Bulanık mantık, belirsiz veya kesin olmayan bilgilerden sonuçlar çıkarmanın bir yolunu sağlar. Yapay Sinir Ağı, doğrusal olmayan değişkenler arasındaki ilişkileri öğrenme ve tespit etme yeteneği nedeniyle sağlık alanında yaygın olarak kabul gören veri madenciliği tekniklerinden biridir. YSA, istatistiksel regresyon modellerinden daha iyi performans gösterir ve ayrıca, özellikle kısa bir süre içinde değişim gösterme eğilimi olan büyük veri kümelerinin daha derin analizine olanak tanır. Bu çalışmada, bir YSA türü olan bulanık mantık ve çok katmanlı algılayıcının (MLP) sağlıkla ilgili tahmin probleminin üstesinden gelme yeteneği araştırılmaktadır. Bu çalışmada, yapay sinir ağlarının (YSA) ve bulanık mantığın (FL) teorik temellerini açıklayarak gerekli bulanık istatistik değerlerin hesaplanması ve yorumlanması aşamalarının örnek bir veri kümesi üzerinde sistematik olarak gösterimleri amaçlanmıştır.

Anahtar Kelimeler: Nöro-bulanık Sistemler, Nöral ağlar, fonksiyon, Fuzzy logic ANFIS

Fuzzy Artificial Neural Networks and an Application

Abstract

Determining the right forecasting methods is of great importance in making the best decisions for each industry. Healthcare is increasingly dependent on advanced information processing technologies to remain competitive in a global healthcare environment. Therefore, the problem of prediction in diagnosis and treatment using data mining techniques is one of the most important issues in health. This area has received great scientific interest and has become a very important research area to provide a more precise forecasting process. Fuzzy logic (FL) and Artificial Neural Network (ANN) is an exciting and promising technique with a wide scope for prediction applications. Fuzzy logic provides a way to draw conclusions from uncertain or imprecise information. Artificial Neural Network is one of the widely accepted data mining techniques in the healthcare field due to its ability to learn and detect relationships between nonlinear variables. ANN outperforms statistical regression models and also allows deeper analysis of large datasets that tend to change over a short period of time. In this study, the ability of fuzzy logic and multilayer perceptron (MLP), which is a type of ANN, to overcome the health-related prediction problem is investigated.

In this study, it is aimed to explain the theoretical foundations of artificial neural networks (ANN) and fuzzy logic (FL) and to systematically show the calculation and interpretation stages of the necessary fuzzy statistical values on a sample data set.

Key Words: *Neuro-fuzzy Systems, Neural Networks, function, Fuzzy logic, ANFIS*

Comparison of Restricted and Unrestricted Biased Estimators and an Application

Israa Mahmoud SHALTOOT^{1*}, Nihal İNCE¹, Sevil SENTURK¹

¹Eskisehir Technical University, Faculty of Science, Department of Statistics, 26470, Eskisehir, Turkey

*Presenter author e-mail: israashaltoot@eskisehir.edu.tr

Abstract

The linear regression model is frequently used in almost all scientific fields to provide more sense of the relationships between variables, fields such as psychology, economics, physics, management, genetics, and sociology. The problem of multicollinearity and its statistical consequences on a linear regression model are very well known in practical applications. One of the major consequences of multicollinearity on the ordinary least squares (OLS) estimator is that the estimator exhibits a large and increasing mean squared error (MSE) loss, particularly when there are near-linear dependencies among the regressor variables. This study mainly aims to examine and compare between some of the restricted and unrestricted biased estimators theoretically in sense of matrix mean square error criterion, in order to identify which linear model is better in estimating the parameter vector in the presence of the multicollinearity, the restricted or the unrestricted. Finally, two numerical examples are computed to demonstrate the behavior of estimators in view of the study's pre-discussed theories. These examples are conducted on the well-known Portland cement data and the total national research and development expenditures data by using the MATLAB program.

Key words: Multicollinearity, Mean square error, Biased estimators, Two parameter estimator, Linear restrictions.

Hidrolojide Yapay Zekâ ve Gri Sistem Modeli ile Kuraklık Analizi

Sunda ŞEN^{1*}, Alper Serdar ANLI¹

¹Ankara Üniversitesi Ziraat Fakültesi Tarımsal Yapılar ve Sulama Bölümü 06110, Ankara, Türkiye

*Sunucu yazar e-posta: sundasen08@gmail.com

Özet

Atmosfer ve iklimde meydana gelen birçok olayı model kullanarak geleceğe yönelik tahmin etme, analiz yapma ve önlem alınabilmesi gibi çalışmaların yapılması mümkün hale gelmiştir. Hidrolojik süreçlerin ve kuraklığın analizinde bu durum modelleme ve simülasyon çalışmalarıyla yapılabilmektedir. Günümüzün en büyük problemlerinden biri olan kuraklık yalnızca fiziksel bir olay veya bir doğa olayı olarak görülmemelidir. Çünkü insan ve faaliyetlerinin su kaynaklarına olan bağımlılığı nedeniyle toplum üzerinde çeşitli etkileri vardır. Kuraklığı erken tahmin edebildiğimiz sürece olumsuz etkilerini azaltabilme imkânımız doğmuş olur. Bu amaçla çalışmada kuraklığın tahmin edilmesinde yapay zekâ ve gri sistem modeli kullanılmıştır. Yapay zekâ; tutarlı ve hızlı tahmin yapması, düşük hata miktarı elde etmesi ve değişkenler arasındaki ilişkilerin en doğru şekilde gözlemlenmesi imkânı sunar. Çalışma sonunda yapay zekâda python programlama dili ve gri sistem modelinde de gri sistem tahmini kullanılmış olup yöntemlerin uygulanabilir olduğu belirlenmiştir.

Anahtar kelimeler: Kuraklık, Yapay zekâ, Python, Gri sistem modeli, Gri tahmin

Drought Analysis with Artificial Intelligence and Gray System Model in Hydrology

Abstract

It has become possible to predict, analyze and take precautions for the future by using models of many events occurring in atmosphere and climate. In the analysis of hydrological processes and drought, this can be done with modelling and simulation studies. Drought, one of the biggest problems of our time, should not be seen as just a physical event or a natural phenomenon. Because it has various effects on society due to its dependence on water resources of humans and their activities. As long as we can predict the drought early, we have the opportunity to reduce its negative effects. For this purpose, artificial intelligence and gray system model were used to predict drought in the study. Artificial intelligence; it provides consistent and fast estimation, low error quantity and optimal observation of relationships between variables. At the end of the study, python programming language in artificial intelligence and gray system prediction were used in gray system model and methods were determined to be applicable.

Key words: Drought, Artificial intelligence, Python, Gray system modelling, Gray prediction

GİRİŞ

Yeryüzündeki en küçük canlı organizmadan en büyüğüne kadar, insanın tüm faaliyetlerini ve biyolojik yaşam ve döngüyü ayakta tutan da sudur. 21. yüzyıl'da meydana gelme olasılıkları artan iklim değişikliği ve kuraklık, gıda güvenliğini tehdit etmektedir. Bu çalışmanın amacı, kuraklık izleme ve etkilerini azaltarak çalışmalarda başlangıç aşamasının tamamlanmasına katkı sağlamaktır.

Atmosfer ve iklimde meydana gelen birçok olayı model kullanarak geleceğe yönelik tahmin etme, analiz yapma ve önlem alınabilmesi gibi çalışmaların yapılması mümkündür. Hidrolojik süreçleri ve kuraklığı analiz yaparken bu durum modelleme ve simülasyon çalışmalarıyla mümkün hale gelebilmektedir.

Kuraklık

Kuraklık genel olarak, yağışın normalin altına düşmesi olarak tanımlanır. Bununla beraber, bu eksikliğin zaman ve süresine göre kuraklıkla ilgili çeşitli tanımlar ortaya çıkmıştır. Kuraklığın başlangıç ve bitişinin belirsiz oluşu, kümülatif olarak artması, aynı anda birden fazla kaynağa etkisi ve ekonomik boyutunun yüksek olması onu diğer doğal afetlerden ayıran en önemli özelliklerdir (Willeke vd. 2004).

Çok yavaş gelişerek belirli bir süreçte oluşan bu doğal olayın süresi uzadıkça sonuçları da çok tehlikeli boyutlara ulaşmaktadır. Esas olarak yağış yetersizliğine bağlı olarak su azlığıyla ortaya çıkan kuraklık, üretimde azalmaya, yetersiz beslenmeye, sonuçta kıtlık, açlık ve ölümlere neden olabildiğinden çok önemli sosyal ve ekonomik sorunların yaşanmasına neden olmaktadır (Kömüşçü 2001). Ancak bütün tanımlarda, iklim dalgalanmalarına bağlı yağış yetersizliği, bu olayın temel nedeni olarak gösterilmektedir. Bunun için "genellikle yağış yetersizliği nedeniyle, yeraltı ve yerüstü doğal su varlığının belli bir süreçte, bölgesel boyutta ve önemli ölçüde ortalama değerlerin altına düşmesiyle oluşan su açığı" şeklindeki kuraklık tanımı, bugün için en yaygın ve en geçerli olanıdır (Kömüşçü vd. 2000).

Kuraklık tabiatın gizli bir tehlikesidir. Genellikle herhangi bir mevsim veya bir zaman diliminde yağış miktarındaki azalmadan dolayı meydana gelir.

Kuraklık hesaplamalarında bir bölgedeki yağış ve evapotranspirasyon arasındaki dengenin uzun süreli ortalaması göz önünde bulundurulmalıdır. Kuraklık zamana bağlı bir parametredir (Graedel vd. 2007). Yağışların başlangıcındaki gecikmeler, yağış-zaman ilişkisi, yağışların yoğunluk ve şiddeti kuraklık üzerine etkili bazı faktörler arasında yer alır.

KURAKLIK ETKİLERİ

Kuraklık yavaş gelişen tehlikeli bir meteorolojik olay olmasına rağmen, insanlara ve çevreye en fazla zarar veren doğal afetlerin başında gelmektedir. Çünkü kuraklığın bir yerde ne zaman başladığı ve ne zaman biteceği tam olarak bilinmemektedir. İlk aşamada daha önce alınan bazı önlemlerle ortaya çıkan olumsuzluklar atlatılabilecek gibi görünse de, kuraklığın uzaması halinde bu önlemler yetersiz kalmaktadır.

Bir yörede, bir bölgede veya bir ülkede görülen kuraklık sadece orada yaşayanlar için değil, gelişmiş ve gelişmekte olan bütün ülkeler, dolayısıyla dünyadaki bütün insanlar için önemli sorunlar doğuran bir tehlikedir. Çünkü bir yerde, iklim dalgalanmalarına bağlı olarak görülen kuraklık, mutlaka diğer ülkelerde de bir iklim anomalisinin yaşanmasına neden olmaktadır (Willeke vd. 2004).

GRI SİSTEM TEORİ

Gri sistem teorisi ilk kez 1982 yılında Prof. Julong Deng tarafından önerilen, geleneksel istatistik yöntemlerin özünde olan kusurları ortadan kaldırır ve tek belirsiz bir sistemin davranışını tahmin etmek için veya analizini gerçekleştirmek için sınırlı bir veri gerektirir. Gri sistem teorisinde araştırma alanı iki bölüme ayrılır, biri ilişkisel analiz ve diğeri tahmindir.

Geleneksel Gri İlişkisel Kalitesinin Geliştirilmesi

Geleneksel gri ilişkisel sınıfının iki türü vardır ve aşağıda tüm kavramları açıklanmıştır.

Deng'in Gri İlişkisel Analizi

Geleneksel yöntemler kullanılarak, Deng öncelikle dört temel öneriyi iki bölüme ayırmıştır ki ilk olarak gri ilişkisel derece formülünü elde etmiştir (Wen 2004).

Gri İlişkisel Katsayı

$$\gamma(X_0(k), X_i(k)) = \Delta_{\min.} + \zeta \Delta_{\max.} / \Delta_{0i}(k) + \zeta \Delta_{\max} \quad (1)$$

Burada;

$$i=1,2,3,\dots,m. \quad k=1,2,3,\dots,n. \quad j \in i$$

X_0 : Referans sırası

X_i : İncelenen sıra

$$\Delta_{0i} = |X_0(k) - X_i(k)| : X_0 \text{ ve } X_i \text{ arasındaki fark.}$$

$$\Delta_{\min.} = \min_{j \in i} V_{j \in i}^{\min.} k |X_0(k) - X_i(k)|,$$

$$\Delta_{\max.} = \max_{j \in i} V_{j \in i}^{\max.} k |X_0(k) - X_i(k)|$$

$$\zeta : \text{Ayırt edici katsayısı ve } \zeta \in [0,1]$$

Gri İlişkisel Derece

Gri ilişkisel katsayısının ortalaması

$$\Gamma_{0i} = 1/n \sum_{k=1}^n \gamma(X_0(k), X_i(k)) \quad (2)$$

Genel olarak her bir faktörün ağırlığı eşit değildir, bu nedenle (2) eşitliği genişletilirse

$$\gamma(X_i, X_j) = \sum_{k=1}^n \beta_k \gamma(X_i(k), X_j(k))$$

burada β_k ; her bir faktörün ağırlığı ve $\sum_{k=1}^n \beta_k = 1$ 'dir.

ζ : Ayırt edici katsayısı

Ayırt edici katsayısının (ζ) temel amacı Δ_{0i} ve $\Delta_{\max}$ arasındaki farkı ayarlamaktır. Ayırt edici katsayısı aralığına istediğimiz değeri verebiliriz ancak genellikle bu 0.5'e eşittir. Matematiksel ispata göre, ζ değerinin değişimi, gri ilişkisel derece sıralamasını değiştirmez.

Gri İlişkisel Derecesini Sıralamak

Gri ilişkisel derece hesaplandıktan sonra, değerine göre dereceyi sıralarız ve bu prosedüre gri ilişkisel derece denir.

Referans dizisi X_0 ve incelenen dizi X_i için, burada

$$X_0 = (X_0(k)), \quad X_i = (X_i(k)) \quad (3)$$

$$k=1,2,3,\dots,n, \quad i=1,2,3,\dots,m$$

eğer

$$\gamma(X_0, X_i) \geq \gamma(X_0, X_j)$$

Daha sonra X_0 referans dizisi altında, X_i 'nin gri ilişkisel sıralaması, X_j 'nin gri ilişkisel sıralamasından daha yüksek olduğunu bulduk.

Gri İlişkisel Katsayı

$$\gamma(X_0(k), X_i(k)) = \{ \Delta_{\max} - \Delta_{0i}(k) / \Delta_{\max} - \Delta_{\min} \}^{\zeta} \quad (4)$$

Gri İlişkisel Derece

Deng'in gri ilişkisel derecesiyle aynıdır.

$$r_{0i} = 1/n \sum_{k=1}^n \gamma(X_0(k), X_i(k)) \quad (5)$$

Gri ilişkisel analiz, gri sistem teorisinin alt başlıklarından birisi olarak bilimsel çalışmalarda yerini almış olup bir derecelendirme, sınıflandırma ve karar verme tekniğidir (Lin vd. 2004). Julong Deng tarafından geliştirilen gri sistem teorisinin en önemli konusu olan gri ilişkisel analizin karmaşık hesaplamalar ve formüllere ihtiyaç duymaması, belirli ve net hesaplama süreci ve adımlarından oluşması bu tekniğin kolay ve uygulanabilir olmasını sağlamaktadır (Deng 1982, 1989, Lu ve Wevers 2007, Baş 2010).

YAPAY ZEKÂ

Teknolojik gelişmelerin hızla ilerlediği günümüzde farklı disiplinler tarafından ele alınan ve üzerinde sıkça durulan konulardan biri de yapay zekâdır. Çeşitli alanlarda yapay zekâ ile ilgili yapılan çalışmalardan son derece olumlu sonuçlar elde edilmektedir.

Yapay zekâ (Artificial Intelligence) bilgisayar biliminin bir dalı olarak kabul edilen insanlar tarafından geliştirilen modeller ile geçmiş veriler üzerinden öğrenme süreci sonrasında makinelerin öğrendikleri bilgiyi insanlar gibi sunabilen bir teknoloji modelidir.

Yapay zekâ nedir?

Yapay zekânı ne olduğu konusunda ortak bir tanım bulunmamaktadır. Bunun en önemli nedeni yapay zekânın birçok farklı uygulama alanı olmasıdır ve doğal olarak her alan kendine göre bir tanım yapmaktadır. Aşağıda bu tanımların bazıları yer almaktadır.

"Yapay zekânın amacı, normal olarak insan zekâsını gerektiren görevleri yapabilecek makineler yapmaktır.

Yapay zekâ araştırmalarının amacı, insan varlığından gözlemlediğimiz ve akıllı davranış olarak adlandırdığımız davranışları gösterebilen bilgisayarlar yapmaktır,

Yapay zekâ, genel olarak insan tarafından yapıldığında, doğal zekâyı gerektiren görevleri yapabilecek mekanizmanın oluşturulması çabalarının tümü olarak değerlendirilebilir.

İlerleyen zamanlarda yapay zekâyı sahip robotlarla yaşamak zorunda kalacağız. Hayatımızın her anında yapay zekâyı yapılmış çalışmalar, tarımda kullanılacak robotlar, hidrolojik olaylarda ve kuraklık bilgilerini artık yapay zekâ ile analiz edebileceğiz.

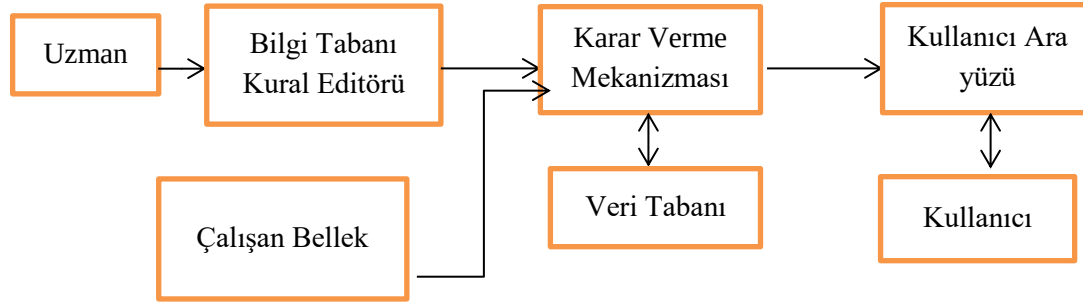
YAPAY ZEKÂ UYGULAMALARI

Doğadaki varlıkların akıllı davranışlarını yapay olarak üretmeyi amaçlayan (Charniak ve McDermot 1985, Akt. Nahiye 2012), bu meyanda işini mükemmel yapan canlı sistemlerini ve insan beynini model alan yapay zekâ çalışmaları; günlük hayatın farklı alanlarında ürünler vermesinin yanında, tahmin, sınıflandırma, kümeleme gibi amaçlar için de kullanılmaktadır. Başlıca olarak uzman sistemler, genetik algoritmalar, bulanık mantık, yapay sinir ağları, makine öğrenmesi gibi teknikler, genel olarak yapay

zekâ teknolojileri olarak adlandırılmaktadır. Bu tekniklerin yanı sıra doğanın taklidi amacıyla da canlılar incelenmekte ve benzeri akıllı yöntemler önerilmektedir. Karınca kolonisi, parçacık sürü ve yapay arı gibi algoritmalar, yapay zeka optimizasyon teknikleri olarak kullanılmaktadır.

Uzman Sistemler

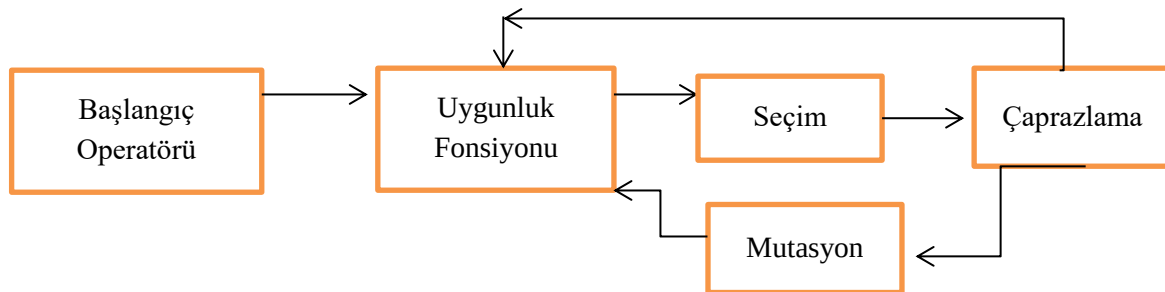
Uzman sistemler, çözümü bir uzmanın bilgi ve yeteneğini gerektiren problemleri, bilgi ve mantıksal çıkarım kullanarak o uzman gibi çözeblen sistemlerdir. Yani problemi çözmede uzman kişi veya kişilerin bilgi ve mantıksal çıkarım mekanizmasının modellenmesi amaçlanmaktadır. (Edward Feigenbaum'dan aktaran: Harmon ve King, 1985:5). Uzman sistemlerde, bilgiler depolanıp daha sonra bir problemle karşılaştırıldığında bu bilgi üzerinden yapılan çıkarımlarla sonuçlara ulaşmaya çalışılmaktadır. Uzman sistemler şu temel öğelerden teşekkül edilir: Bilgi tabanı (kural tabanı), veri tabanı, çalışan bellek (yardımcı yorumlama modülü), çıkarım motoru (karar verme mekanizması, mantıksal çıkarım modülü) ve kullanıcı ara yüzü. Bilgi tabanı, bilgilerin tutulduğu ve tutulan bilgilerden yeni bilgiler üretilmesine imkan sağlayan birim olup uzman sistemin beyni ve yapı taşıdır. Veri tabanı ise bilgi tabanı ile sürekli ilişki halinde olmalıdır (Nabiyev 2012). Kullanıcı ara yüzü; bilgi kazanma, bilgi tabanı ile hata ayıklama deneme, test durumlarını çalıştırma, özet sonuçlar üretme, sonuca götüren nedenleri açıklama ve sistem performansını değerlendirme gibi görevleri yerine getirmektedir. Çıkarım motoru, uygun bilgi için bilgi tabanını taramakta ve mevcut problem verisine dayanarak çıkarımda bulunmaktadır. Çalışan bellekte, problem ile ilgili soruların cevapları ve tanısal testlerin sonuçları gibi mevcut problem verisi saklanmaktadır (Balaban ve Kartal 2015). Çıkarım motorunda mantıksal sonuçlar elde edilmesine yardımcı olur.



Şekil 1. Uzman Sistemin Genel Yapısı

Genetik Algoritmalar

Genetik algoritmalar, bilinen yöntemlerle çözülemeyen veya çözüm süresi problemin büyüklüğüne göre oldukça fazla olan problemlerde, kesin sonuca çok yakın sonuçlar verebilen bir yöntemdir.



Şekil 2. Genetik Algoritmanın Genel Akış Şeması (Nabiyev 2012)

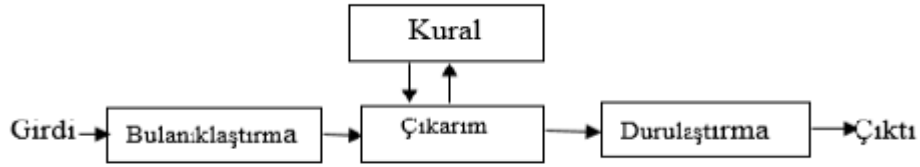
Bulanık Mantık

İki değerli mantıkta her şey ya doğru ya yanlıştır. Bulanık sistem bulanık kümelerin/ bulanık mantığın kullanılması birkaç şekilde olabilir. Bunlar (Baykal ve Beyan, 2004);

1-) Sistem "eğer-o halde" şeklinde kurallarla tanımlanabilir. Bu şekilde tanımlanan sistemlere kural tabanlı bulanık sistemler adı verilir.

2-) Sistem parametreleri gerçek sayılar yerine bulanık sayılar kullanılarak parametre değerlerindeki belirsizlik tanımlanabilir.

3-) Sistemin girdi, çıktı ve durum değişkenleri insan algısı ile ilişkili nicelikleri ifade ediyor veya sözel bilgiyi taşıyorsa bu değişkenler bulanık küme ile tanımlanabilir.



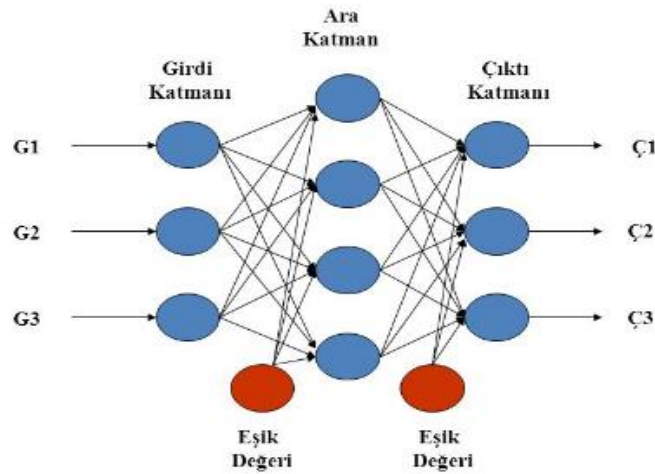
Şekil 3. Bulanık Sistemin Genel Yapısı

Makine Öğrenmesi

Makine öğrenmesi, bir problemi o probleme ait veriye göre modelleyen bilgisayar algoritmalarının genel adıdır. Mevcut veri seti ve kullanılan algoritma ile oluşturulan model, en yüksek performansı vermek üzere kurulmaktadır. Bu sebeple pek çok makine öğrenmesi yöntemi geliştirilmiş olup, bunlardan bazıları; k-en yakın komşu algoritması, basit (naive) bayes sınıflandırıcı, karar ağaçları, lojistik regresyon analizi, k-ortalamlar algoritması, destek vektör makineleri ve yapay sinir ağlarıdır. Bu yaklaşımların bir kısmı tahmin ve kestirim, bir kısmı kümeleme ve bir kısmı da sınıflandırma yapabilme yeteneğine sahiptir.

Yapay Sinir Ağları

İnsan beyninin temel işlem elemanı ve sinir sisteminin en basit elemanı olan nöron ve bu nöronlar arası bağlantılara şekilsel ve işlevsel olarak benzeyen bir yapay sinir ağı, bu haliyle adeta biyolojik sinir sisteminin basit bir simülasyonudur.



Şekil 4. Çok Katmanlı Bir Yapay Sinir Ağının Genel Yapısı

YAPAY ZEKÂ PROGRAMLAMA DİLLERİ

Yapay zekâ çalışmalarında en çok kullanılan programlama dilleri aşağıda açıklamasıyla birlikte verilmiştir.

Python

Yapay zekâda en çok tercih edilen bir programlama dilidir. Geliştirilmiş kütüphanesi sayesinde öğrenilmesi de kolaydır. Matematiksel işlemler için Numpy, ileri derecede programlama için Scipy, yapay zekâ işlemleri için kullanılan Pybrain gibi onlarca kütüphaneye sahiptir.

Bizim de bu çalışmamızda yapay zekâ adı altında python programlama dili kullanılacaktır.

Lisp

En eski programlama dillerinden birisidir. Eski olmasına rağmen tercih edilenler arasında yerini alabilmiştir. Lisp'in ana amaçlarından biri bilgisayar sistemleri için matematiksel gösterimler yapabilmektedir.

Prolog

Kullanım kolaylığı ve rahatlığı açısından Lisp'e yakındır. Yazılımların ilişkiler ile ifade edildiği bildirime dayalı bir programlama dilidir.

C++

Genellikle oyun programlarında kullanılıyor. Yıllardır eskimeyen bu dil, yapay zekâ konularında da hat safhada kullanım alanı elde ediyor. C ++ sadece yazılıma değil, donanıma da müdahale edebildiğimizden uygulamaların çok daha hızlı olmasını sağlayabiliriz.

Java

Lisp ve Prolog programlama dilleri gibi üst düzey olmayan ve Python ve C ++ kadar hızlı çalışmasa da yapay zekâ konularında kullanılan başka bir popüler programlama dilidir.

R

Temeli 1976 yılından bu yana Bell Laboratuvarları'nda istatistiksel programlama dili olarak geliştirilen S dilidir.

R; veri işleme, hesaplama ve grafiksel gösterimi için bir dil ve çevre sağlar. Geniş bir yelpazede istatistiki ve grafiksel teknikleri içerir. Yani; doğrusal ve doğrusal olmayan modelleme, klasik istatistik testleri, zaman-serileri analizi, sınıflandırma, kümeleme,...

Açık kaynak kodlu olması itibariyle geliştirilmeye çok yatkındır. R; fonksiyonel bir programlama dili olan R istatistikçiler için kod yazmayı kolaylaştıran fonksiyonlara sahiptir. Veri çözümlemede kullanılabilecek grafiksel araçlara sahiptir.

MATERYAL VE YÖNTEM

Materyal

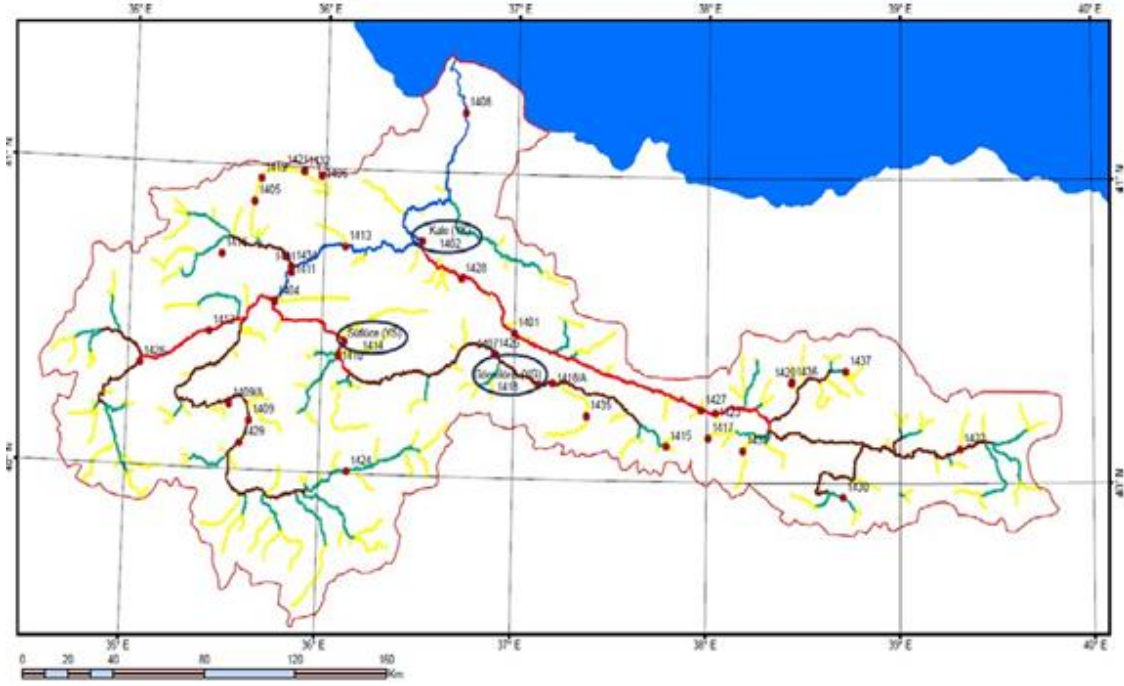
Çalışma Alanı

Ülkemizin en önemli akarsularından bir tanesi olan Yeşilırmak, ülkemiz topraklarında doğup ülkemiz topraklarında denize dökülen akarsulardan bir tanesidir. Ülkemizde ki Kızılırmak, Fırat ve Dicle akarsuları ile birlikte en çok bilinen akarsuların başında gelmektedir.

Yeşilırmak nehri uzunluk itibariyle bir Kızılırmak bir Fırat veya bir Dicle nehri kadar uzun olmasa da ülkemizdeki diğer akarsuların birçoğundan uzun bir akarsudur.

Yeşilırmak; Tokat, Amasya ve Samsun illerinden geçerek çeşitli akarsularla birleşir. Nehir üç ana kolun birleşmesinden oluşur. Nehrin en büyük kolu Kelkit Çayı ve uzunluğu 320 km'dir. Nehrin diğer kolları Çekerek ve Yeşilırmak'tır. Çekerek ve Kelkit ırmaklarının birleştiği yerden sonra Yeşilırmak başlar. Kelkit'in katılmasıyla Yeşilırmak nehrinin debisi artar.

Yeşilırmak nehri akarsu rejimi bakımından düzensiz olan akarsularımızdan bir tanesidir. Yeşilırmak nehrinin toplam yağış alanının genişliği 35.000 km²'dir. Bu büyüklük ülkemiz topraklarının yaklaşık olarak 1/22'lik kısmına tekabül etmektedir.



Şekil 5. Yeşilırmak Havzası ve ele alınan akım istasyonları

Çalışmamızda Yeşilırmak-Gömelönü, Yeşilırmak-Sütluce ve Yeşilırmak-Kale akım istasyonları göz önüne alınmıştır. İstasyonların koordinatları ve gözlem yılları Çizelge 1'de verilmiştir.

Çizelge 1. Yeşilırmak nehri üzerindeki akım istasyonlarının coğrafik konumları

İstasyon Adı	Enlem	Boylam	Gözlem Süresi
Yeşilırmak-Gömelönü (YG)	37°7'43"	40°18'42"	1963-2013
Yeşilırmak-Sütluce (YS)	36°7'5"	40°26'3"	1955-2013
Yeşilırmak-Kale (YK)	36°30'45"	40°46'18"	1939-2013

Yöntem

Gri tahmin yöntemi, elde mevcut verileri kullanarak gri model GM(1,1) yardımıyla geleceğe ilişkin tahminlerde bulunmak için geliştirilmiştir. GM(1,1), bir grup türevlenebilir eşitliği barındıran zaman serisi tahmin modelidir. GM(1,1) gösterimi, tek değişkene sahip birinci dereceden türevlenebilir eşitliklerin yer aldığı gri modeli ifade etmek için kullanılmaktadır. Gri tahmin yöntemi aşağıda detaylı olarak açıklanan temel adımlarından oluşur (Liu ve Lin, 2006).

Adım-1: Başlangıç verilerini kullanarak aşağıda gösterilen $X^{(0)}$ ham veri setini oluşturuyoruz.

$$X^{(0)} = (x^{(0)}(1), x^{(0)}(2), x^{(0)}(3), \dots, x^{(0)}(n))$$

Adım-2: Birinci dereceden toplam üretim operatörü (AGO) kullanılarak $X^{(1)}$ i oluşturuyoruz.

$$\sum_i^k = 1 \quad X_{(0)}(i) \quad (i = 1, 2, 3, \dots, n) \quad (6)$$

$$X^{(1)} = (x^{(1)}(1), x^{(1)}(2), x^{(1)}(3), \dots, x^{(1)}(n))$$

Adım-3: Birinci dereceden ortalama değer üretim operatörünü kullanarak $Z^{(1)}$ i oluşturuyoruz.

$$Z^{(1)}(k) = 0.5x^{(1)}(k) + 0.5x^{(1)}(k-1) \quad (7)$$

$$Z^{(1)} = (Z^{(1)}(2), Z^{(1)}(3), Z^{(1)}(4), \dots, Z^{(1)}(n))$$

Adım-4: GM(1,1) 'nin parametrelerini bulma;

$$\hat{a} = \begin{bmatrix} a \\ b \end{bmatrix} \quad (8)$$

$$C = \sum_k^n = 2 \quad Z_{(1)}(k) \quad (9)$$

$$D = \sum_k^n = 2 \quad X_{(0)}(k) \quad (10)$$

$$E = \sum_k^n = 2 \quad Z_{(1)}(k) \cdot X_{(0)}(k) \quad (11)$$

$$F = \sum_k^n = 2 \quad Z_{(1)}(k)^2 \quad (12)$$

n = Veri sayısı

$$a = CD - (n-1)E / (n-1)F - C^2 \quad (13)$$

$$b = DF - CE / (n-1)F - C^2 \quad (14)$$

Adım-5: $dx^{(1)}(k) / dk + ax^{(1)}(k) = b$ şeklinde gösterilen birinci dereceden türevlenebilir eşitliği çözüp, tahmin modelini elde edersek;

$$\hat{X}^{(0)}(k+1) = (1 - e^a)(X^{(0)}(1) - b/a)e^{-ak} \quad k=1, 2, 3, \dots \quad (15)$$

Adım-6: GM(1,1) modelinin hata formülü;

$$e(k) = \left| X^{(0)}(k) - \hat{X}^{(0)}(k) / X^{(0)}(k) \right| \cdot \%100, \quad k \in N \quad (16)$$

burada,

1. $X^{(0)}(k)$: Gerçek değer

2. $\hat{X}^{(0)}(k)$: Tahmin edilen değer

Adım-7: Ortalama hata;

$$e_{ort.} = e(1)+e(2)+e(3)+..... / \text{Veri sayısı} \quad (17)$$

Yapay zekâ yöntemlerinden de makine öğrenmesi yöntemleri kullanılarak modellenmiş olup, bir,üç ve altı ay sonraki kuraklık değerleri tahmin edilmiştir. Destek vektör makineleri (SVM) ve K-en yakın komşuluk (KNN) algoritmaları kullanılarak oluşturulan modeller ile yapılan tahminlerin başarı oranı istatistiksel olarak değerlendirilmiştir.

İstatistiksel Analizler

Çalışmada göz önüne alınan Yeşilırmak Havzası'nda bulunan akım istasyonlarından elde edilen verilerin Gri Sistem Teorisi (GST) ile tahminleri yapılmıştır. Tahminde kullanılan verilerin tahmindeki başarısı Hataların Karesinin Ortalamasının Karekökü (RMSE), Hataların Mutlak Değerlerinin Ortalaması (MAE), Nash-Sutcliffe Katsayısı (E) ve Pearson Korelasyon Katsayısı (r) ile belirlenecektir (Jachner ve ark., 2007). Bu ilişkiler aşağıda açıklanmıştır.

Hataların Karesinin Ortalamasının Karekökü (RMSE)

Gözlenmiş (dörtlü veri seti) ve GST ile tahmin edilmiş her veri seti için farklar alınarak hata miktarları (gözlenen değerden olan sapma) belirlenir. Bunların kareleri alınarak ortalaması bulunur. Ortalama değer karekökünün alınmasıyla RMSE değeri saptanır. Hesaplanan değer sıfıra yaklaşması gözlenen değerlere yakın tahmin yapıldığını göstermektedir. Hataların karesinin ortalamasının karekökü aşağıda verilen ilişki ile saptanmaktadır.

$$RMSE = \sqrt{n^{-1} \sum_{i=1}^n \{Q_{obs}(i) - Q_{pred}(i)\}^2} \quad (18)$$

İlişkide verilen Q_{obs} ve Q_{pred} sırasıyla gözlenmiş ve tahmin edilmiş akım değerlerini, n ise gözlem sayısını vermektedir.

Hataların Mutlak Değerlerinin Ortalaması (MAE)

Bu yöntemde hataların (gözlenen ve tahmin edilen değerlerin farkı) mutlak değerlerinin ortalaması alınarak MAE değeri hesaplanır. Bu yöntemde elde edilen MAE değerinin sıfıra yakın değerler alması tahminin başarılı olduğunu, sıfırdan uzaklaşması ise modelin performansında azalmanın olduğunu belirtir.

$$MAE = n^{-1} \sum_{i=1}^n |Q_{obs}(i) - Q_{pred}(i)| \quad (19)$$

Nash-Sutcliffe Katsayısı (E)

Bu yöntem hidrolojik değişkenlerin modellenmesinde tahminin başarısının değerlendirilmesinde yaygın olarak kullanılan bir yöntemdir. Nash-Sutcliffe katsayısı, gözlenmiş ve tahmin edilmiş değerler ile gözlenmiş değerlerin ortalamasına göre elde edilmektedir. Bu katsayı $-\infty$ ile $+1$ arasında değişen değerler alır. Bu ilişkiden elde edilen sonucun $+1$ olması gözlenen ve tahmin edilen değerler arasında mükemmel bir eşleşmenin olduğunu gösterir. Buradan elde edilen sonucun sıfır çıkması ise tahmin değerlerinin gözlenmiş değerlerin ortalamasıyla benzeşim gösterdiğini belirtir. Bu katsayının $-\infty$ 'a doğru gitmesi ise tahmin değerlerinin gözlenmiş değerlerin ortalamasından daha büyük değerler olduğunu belirtir.

$$E = 1 - \frac{\sum_{i=1}^n \{Q_{obs}(i) - Q_{pred}(i)\}^2}{\sum_{i=1}^n \{Q_{obs}(i) - \overline{Q_{obs}}\}^2} \quad (20)$$

İlişkide verilen $\overline{Q_{obs}}$ gözlenmiş akım değerlerinin ortalamasını belirtmektedir.

Pearson Korelasyon Katsayısı

Bu ilişki iki değişken (bu çalışmada gözlenmiş ve tahmin edilmiş akım değerleri) arasındaki doğrusal ilişkinin gücünü ve yönünü göstermektedir. Çalışmada göz önüne alınan iki değişkenin arasındaki ilişkinin ölçülmesinde kullanılan ve aşağıda ilişkisi verilen Pearson Korelasyon Katsayısının (r) +1 ve -1 yakın değerler alması iki değişken arasındaki ilişkinin güçlü olduğunu belirtmektedir. Ancak elde edilen katsayı değerinin pozitif olması iki değişkenin verilerinin artan bir eğilime sahip olduğunu, bu katsayının negatif olması ise değişkenlerin azalan yönde bir trend gösterdiğini belirtmektedir. Ancak bu katsayı değerinin sıfıra yaklaşması iki değişken arasındaki ilişkinin zayıflığını göstermektedir.

$$r = \frac{\sum_{i=1}^n [Q_{obs}(i) - \overline{Q_{obs}}][Q_{pred}(i) - \overline{Q_{pred}}(i)]}{\sqrt{\sum_{i=1}^n [Q_{obs}(i) - \overline{Q_{obs}}]^2 \cdot [Q_{pred}(i) - \overline{Q_{pred}}(i)]^2}} \quad (21)$$

İlişkide verilen $\overline{Q_{pred}}$ tahmin edilen akım değerlerinin ortalamasını belirtmektedir.

Homojenlik Analizi

Homojenlik, hidrolojik verilerdeki değişkenliğinin ortaya çıkarılması anlamında önemlidir. Hidrolojik verinin homojen olması koşulunda, veri ölçümünün zaman içerisinde gerek çevre gerekse alet anlamında hiçbir değişikliğin olmadığı durumlarda gerçekleştirildiği anlaşılmaktadır. Bundan dolayı ölçülen değerlerin aynı popülasyondan geldiği ifade edilmektedir. Bu çalışmada homojenlik analizi gözlenmiş ve tahmin edilen verilerin benzer dağılım karakteristiğine yani aynı popülasyondan geldiğiyle veya gelmediğiyle ilgili yapılan analizdir. Homojenliğin istatistiki anlamda test edilemesinbir çok yöntem bulunmasına karşın yaygın olarak kullanılan yöntem Mann-Whithney U (MWU) testidir. Bu test ile ilgili açıklamalar aşağıda verilmiştir (Nachar, 2008). Bu yöntemde n_1 ve n_2 gözlem sayılı iki örneğin aynı medyanlıpopülasyondan alınıp alınmadığı test edilir. Bu istatistik analizde H_0 ve H_1 hipotezi aşağıda verildiği şekilde oluşturulur.

H_0 : n_1 ve n_2 gözlem sayılı örnekler aynı dağılım karakteristiği göstermektedir.

H_1 : n_1 ve n_2 gözlem sayılı örnekler farklı dağılıma sahiptir.

Bu test uygulanırken öncelikle her iki örneğin gözlemleri birlikte göz önüne alınarak her bir gözlemin sırası (rank) saptanır. Her bir örneğin rankları toplanarak R_1 ve R_2 rank toplamaları elde edilir. Elde edilen bu değerlere bağlı olarak MWU istatistik değeri aşağıda verilen ilişkiler yardımıyla belirlenir.

$$U_1 = n_1 \cdot n_2 + n_1 \cdot (n_1 + 1) / 2 - R_1 \quad (22)$$

$$U_2 = n_1 \cdot n_2 + n_2 \cdot (n_2 + 1) / 2 - R_2 \quad (23)$$

Aşağıda verilen test istatistiği (Z_{MWU}) hesaplanırken U_1 ve U_2 'nin sayıca küçük olanı alınır ve bu değer U 'ya karşılık gelir.

$$U_{ort.} = U_1 \cdot U_2 / 2 \quad (24)$$

$$SS_U = \left[\frac{n_1 \cdot n_2 \cdot (n_1 + n_2 + 1)}{12} \right]^{1/2} \quad (25)$$

$$Z_{MWU} = U - U_{ort.} / SS_U \quad (26)$$

Elde edilen Z_{MWU} değeri, α (0.05) önem seviyesinde standart normal dağılımın tablo değeri olan ± 1.96 değeri ile karşılaştırılır. Hesaplanan Z_{MWU} değeri kritik tablo değerinin arasında kalırsa yukarıda verilen H_0 hipotezi kabul edilir. Yani iki gözlemin aynı popülasyondan geldiğine karar verilir.

BULGULAR

Bu çalışmada yapay zekâ yöntemleri olarak adlandırılan makine öğrenmesi yönteminin ve gri sistem tahmin modelinin kuraklık analizinde ve indislerinin tahmin edilmesinde uygulanabilirliği araştırılmıştır. Elde edilen sonuçlar kuraklık alanında yapay zekâ ve gri sistem tahmin modelinin önemli katkı sağladığı sonucuna varılmıştır.

Gri Sistem Teorisinin en önemli avantajlarından birisi olan az sayıda ve belirsiz bilgiler içeren veri kümeleri ile çalışabilmek birçok gerçek hayat uygulamasına zemin hazırlamaktadır. Gri Sistem Teorisi çeşitli alanlarda başarılı analizlerin yapılmasına zemin oluşturmuş güçlü bir teoridir. Geleneksel metotlar genellikle büyük veri kütleleri ile çalışmada başarılı olurken Gri Sistem Teorisi ile küçük sayıda veri içeren ve eksik bilgi tablolarından sağlıklı bilgilerin çıkarılması işlemi yapılabilmektedir.

Çalışmada, farklı uzunluktaki akım verileri kullanılarak Gri Sistem tahmin modeli yaklaşımıyla akım tahminleri yapıldı. Daha sonra RMSE, MAE, E ve Pearson (r) ile tahmin edilen ortalama akım verilerinin hata analizi incelenmiştir. Ardından MWU testi ile gözlem verilerinin aynı dağılım gösterip-göstermediğini belirlenmeye çalışılmıştır.

Çalışmada kullanılan tüm istatistik analizleri sonuçlarına göre; Gri Sistem tahmin modelinin diğer tahmin yöntemlerine göre görsel açıdan anlaşılabilirliği daha kolay olup ileride yapılacak olan tüm tahmin çalışmalarında bu yöntemden yararlanılması önerilmektedir.

Gelecekteki çalışmalarda araştırmacıların önışlem yöntemlerini kullanarak kuraklık analizi yapmalarının faydalı sonuçlar doğuracağı kanaatine varılmaktadır. Karmaşık yapıya sahip modeller görece iyi sonuçlar verse de, basit modellerin de önışlem yöntemleri ile birlikte kullanımının göz ardı edilemeyeceği de bu çalışma ile varılan önemli bir diğer sonuçtur. İfade edildiği üzere kuraklık belirsizlikler içeren rasgele bir olaydır ve önceden tahmin edilmesi oldukça zordur. Kuraklık yönetiminin etkili bir şekilde planlanması ve işletilmesi için ileriye dönük uzun dönemli kuraklık tahminlerinin önemli bir girdi parametresi olarak değerlendirilmesi gerekir. Karar vericiler gelecekteki kurak devreleri ne kadar önceden bilebilirlerse zararları da o kadar iyi bertaraf etme imkanı bulacaklardır. Ayrıca kuraklık karar destek sisteminin vazgeçilmez bir parçası olan kuraklık tahminlerinin başarısı, karar vericilerin de başarısında önemli rol oynayacaktır.

TARTIŞMA VE SONUÇ

Hidrolojik olaylar birçok değişkenin etkisi altında meydana geldiğinden bu değişkenlerin gelecekteki miktarlarının önceden tahmin edilmesi mümkün olmamaktadır. Bu bakımdan hidrolojik değişkenlerin gelecekteki miktarları hakkında bilgi sahibi olmak için istatistik yöntemlerden ve simülasyon tekniklerinden yararlanılmaktadır. Burada da istatistik yöntemin doğru seçilmesi önemli olmaktadır. Bu gibi sorunların çözümünde birçok doğrusal ve doğrusal olmayan simülasyon teknikleri mevcuttur. Bu çalışmada da yeni geliştirilmiş ve ülkemizde hidrolojik anlamda yeterince çalışılmamış olan gri sistem teorisi yaklaşımı ile ve yapay zekâ yöntemiyle Yeşilırmak Havzasındaki kuraklık analiz edilmeye çalışılmıştır.

Çalışmamızda kuraklık, yapay zekâ adı altında makine öğrenmesi ve gri sistem modeli kullanılarak çalışmalardan son derece olumlu sonuçlar elde edilmiş ve yapay zekânın tarımda ve hidrolojide de çalışmaların her geçen gün artacağını söyleyebiliriz.

Kaynaklar

- Akarçay Havzası Kuraklık Yönetim Planı.* (2015). Kuraklık İndisleri, İndikatörleri ve Eşik Değerlerinin Tespiti Raporu, Orman ve Su İşleri Bakanlığı Su Yönetimi Genel Müdürlüğü Taşkın ve Kuraklık Yönetimi Dairesi Başkanlığı .
- Bacanlı, Ü. G., Dikbaş, F., & Fırat, M. (2011). *Yapay Sinir Ağları ve Bulanık Mantık Yöntemleri ile Kuraklık Tahmini.* Proje Çalışması, Pamukkale Üniversitesi, Denizli.
- Başakın, E., Ekmekcioğlu, Ö., & Özger, M. (2019). Makine Öğrenmesi Yöntemleri ile Kuraklık Analizi. *Pamukkale Üniversitesi Mühendislik Bilimleri Dergisi*, 25(8), 985-991.
- Dabanlı, İ. (2019). Kuraklık Riskinin Bulanık Mantık Yardımıyla Türkiye Geneline Değerlendirilmesi. *Dicle Üniversitesi Mühendislik Fakültesi Dergisi*, 10(1), 359-372.
- Efe, B., & ve Özgür, E. (2014). Standart Yağış İndeksi (SPI) ve Normalin Yüzdesi Metodu (PNI) ile Konya ve Çevresinin Kuraklık Analizi. *II. Uluslararası Katılımlı Kuraklık ve Çölleşme Sempozyumu.*
- İçen , D., & Günay, S. (2014). Uzman Sistemler ve İstatistik. *İstatistikçiler Dergisi: İstatistik & Aktüerya*, 7, 37-45.
- Kömüşçü, A., Erkan, A., & Turgu, E. (19-21 Mart 2003). Normalleştirilmiş Yağış İndeksi Metodu ile Türkiye'de Kuraklık Oluşum Oranlarının Bölgesel Dağılımı. *III. Atmosfer Bilimleri Sempozyumu Bildirileri*, (s. 268-275). İstanbul.
- Özelkan, E. (2019). Uzaktan Algılama ile Belirlenen Baraj Gölü Alanının Zamansal Değişiminin Meteorolojik Kuraklık ile Değerlendirilmesi: Atıkhisar Barajı (Çanakkale) Örneği. *Türk Tarım ve Doğa Bilimleri Dergisi*, 6(4), 904-916.
- Özfidaner, M., & Topaloğlu, F. (2020). Standart Yağış İndeksi Yöntemi ile Güneydoğu Anadolu Bölgesinde Kuraklık Analizi. *Toprak Su Dergisi*, 9(2), 130-136.
- Özfidaner, M., Gönen, E., & Kartal, S. (2020). Adana İlinde Topsis Yöntemi ile Kuraklık Analizi. *Mediterranean Agricultural Sciences*, 33(1), 101-106.
- Topçu , E., & Karaçor, F. (2020). Erzurum İstasyonunun Standartlaştırılmış Yağış Evapotranspirasyon İndeksi ve Bütünleşik Kuraklık İndeksi Kullanılarak Kuraklık Analizi. *Journal of Polytechnic*.
- Tüfekçi, A., & Köse, U. (2013). Bilgisayar Programlama Öğretiminde Yapay Zekâ Tabanlı Bir Yazılım Sisteminin Geliştirilmesi ve Değerlendirilmesi. *Hacettepe Üniversitesi Eğitim Fakültesi Dergisi*, 28(2), 469-481.
- Türkeş, M. (1990). Palmer Kuraklık İndisine Göre İç Anadolu Bölgesi'nin Konya Bölümündeki Kurak Dönemler ve Kuraklık Şiddeti. *Coğrafi Bilimler Dergisi*, 129-144.
- Willeke, G., Hosking, J., J.R. and , & Guttman, N. (1994). Institute for Water Resources Report, The National Drought Atlas.
- Yılmaz, M. (2018). Kuraklık Analizleri, Kuraklık Analizleri ve Metodolojileri Konulu Hizmet İçi Eğitim Programı.

Çıkar Çatışması

“Yazarlar çıkar çatışması olmadığını beyan etmişlerdir” .

Lojistik Dağılım için Farklı Parametre Tahmin Metotlarının Karşılaştırılması

Asuman YILMAZ^{1*}, Mahmut KARA²

^{1,2}Van Yüzüncü Yıl Üniversitesi, İktisadi ve İdari Bilimler Fakültesi, Ekonometri Bölümü, 65080,
Van, Türkiye

*Sunucu yazar e-posta: asumanduva@yyu.edu.tr

Özet

Bu çalışmada, lojistik dağılımın bilinmeyen parametrelerinin tahmin edicileri ele alınmıştır. Lojistik dağılımın bilinmeyen parametrelerini tahmin etmek için en olabilirlik tahmin edicileri (MLE), L-moment tahmin edicileri (LME), moment tahmin edicileri (ME), en küçük kareler tahmin edicileri (LSE), ağırlıklı en küçük kareler tahmin edicileri (WLSE), yüzdelik tahmin ediciler (PE) ve TL-momentler (TLME) tahmin edicileri kullanılmıştır. Bu tahmin yöntemlerinin performansları, bir simülasyon çalışması ile yanları ve ortalama kare hataları açısından karşılaştırılmıştır.

Anahtar kelimeler: Parametre tahmini, en çok olabilirlik, L-moment, TL-moment, simülasyon.

Comparison of Different Parameter Estimation Methods for Logistic Distribution

Abstract

In this paper, the estimators of the unknown parameters of the logistic distribution are considered. The maximum likelihood estimators (MLEs), L-moments estimators (LMEs), moment's estimators (MEs), least squares estimators (LSEs), and weighted least squares estimators (WLSEs), percentile estimators (PEs) and Trimmed L-moments (TLMEs) are used in parameter estimation. The performances of these estimation methods are compared with respect to their biases and mean square errors through a simulation study.

Key words: Parameter estimation, maximum likelihood estimator, L-moments, TL-moments, simulation.

GİRİŞ

Bu çalışmada lojistik dağılımın bilinmeyen parametrelerinin tahmini üzerinde durulacaktır. Parametrelerin tahmininde farklı metotlar kullanılacaktır. Bu metotlardan momentler yöntemi (ME) ve en çok olabilirlik yöntemi (MLE) literatürde oldukça sık kullanılmaktadır. Fakat bu metotların bazı dezavantajlara sahip olduğu bilinmektedir. Örneğin, en çok olabilirlik tahmin edicilerinin hesaplanması bazı durumlarda zordur ve bu yöntem genel olarak küçük örneklem için iyi performans göstermemektedir (Dey ve ark., 2018). Moment tahmin edicileri kolaylıkla elde edilebilmektedir. Ancak bu yöntemin etkinliği bazı durumlarda düşük çıkmaktadır. Bu nedenle bu tahmin edicilere alternatif olarak literatürde oldukça sık kullanılan başka parametre tahmin yöntemleri bulunmaktadır. Bu çalışmada MLE ve ME metotlarının yanı sıra en küçük kareler (LSE), ağırlılandırılmış en küçük kareler (WLSE), yüzdelik tahmin edici (PE), L- moment tahmin edicisi (LME) ve TL-moment tahmin edicisi (TLME) kullanılmıştır.

MATERYAL VE YÖNTEM

Materyal

Lojistik dağılımının dağılım fonksiyonu ve olasılık yoğunluk fonksiyonu sırasıyla;

$$f(x) = \frac{e^{-\left(\frac{x-\mu}{\sigma}\right)}}{\sigma \left(1 + e^{-\left(\frac{x-\mu}{\sigma}\right)}\right)^2}, \quad -\infty < x < \infty, -\infty < \mu < \infty, \sigma > 0,$$

$$F(x) = \frac{1}{\left(1 + e^{-\left(\frac{x-\mu}{\sigma}\right)}\right)}.$$

Burada μ konum ve σ ölçek parametresini temsil etmektedir. Lojistik dağılım beklenen değer ve varyansı

$$E(X) = \mu \text{ ve } Var(X) = \frac{\sigma^2 \pi^2}{3}, \text{dür.}$$

L- moment yöntemi

Hosking (1990) tarafından önerilen L- Moment yöntemi sıra istatistiklerinin lineer birleşimi olarak elde edilmektedir. (Dey ve ark., 2017). Bu yöntemin en önemli avantajı veri setindeki aykırı değerlerden etkilenmemesidir. $\theta = (\theta_1, \theta_2, \dots, \theta_r) \in \Theta$ olmak üzere $r(r \geq 1)$ bileşenli θ parametre vektörünü tahmin etmek isteyelim. $X_1, X_2, \dots, X_n$ olasılık yoğunluk fonksiyonu $f(\cdot; \theta), \theta \in \Theta$ olan dağılımdan bir örneklem ve $X_{1:n} \leq X_{2:n} \leq X_{n:n}$ 'ler bu örnekleme karşılık gelen sıra istatistikleri olmak üzere, var olması halinde, kitle dağılımın L-momentleri,

$$L_k = k^{-1} \sum_{j=0}^{k-1} (-1)^j \binom{k-1}{j} E(X_{k-j:k}) \quad k=1, 2, 3, \dots \quad (1)$$

ve örneklem L-momentleri,

$$l_k = \frac{1}{k \binom{n}{k}} \sum_{i=1}^n \sum_{j=0}^{k-1} (-1)^j \binom{k-1}{j} \binom{i-1}{k-j-1} \binom{n-i}{j} X_{i:n} \quad k=1, 2, 3, \dots \quad (2)$$

ile verilir (Hosking,1990; Elamir ve Seheult,2003).

Buna göre kitle L-momentleri ile örneklem L-momentlerinden ilk r tanesinin eşitlenmesi ile elde edilen ve $\theta_1, \theta_2, \dots, \theta_r$ bilinmeyenlerine göre r tane denklemden oluşan denklem sistemlerinin çözümü olan,

$$\theta_1^o(X_1, X_2, \dots, X_n), \theta_2^o(X_1, X_2, \dots, X_n), \dots, \theta_r^o(X_1, X_2, \dots, X_n)$$

istatistiklerine L-momentler yöntemi tahmin edicileri denir.

TL-moment tahmin yöntemi

TL-moment yöntemi Elamir ve Seheult (2003) tarafından L-moment yönteminin bir genellemesi olarak önerilmiştir. L-moment yönteminde olduğu gibi bu yöntem ile elde edilen tahmin edicilerde veri setindeki aykırı değerlerden etkilenmemektedir. $\theta = (\theta_1, \theta_2, \dots, \theta_r) \in \Theta$ olmak üzere $r(r \geq 1)$ bileşenli θ parametre vektörünü tahmin etmek isteyelim. $X_1, X_2, \dots, X_n$ olasılık yoğunluk fonksiyonu $f(\cdot; \theta), \theta \in \Theta$ olan dağılımdan bir örneklem ve $X_{1:n} \leq X_{2:n} \leq X_{n:n}$ 'ler bu örnekleme karşılık gelen sıra istatistikleri olmak üzere, var olması halinde, kitle dağılımın TL-momentleri,

$$\lambda_k^{(s,t)} = k^{-1} \sum_{j=0}^{k-1} (-1)^j \binom{k-1}{j} E(X_{k+s-j:k+s+t}) \quad k=1,2,3,\dots \quad (3)$$

ve örneklem TL-momentleri,

$$l_k^{(s,t)} = \frac{1}{k \binom{n}{k+s+t}} \sum_{j=s+1}^n \sum_{i=0}^{k-1} (-1)^i \binom{k-1}{i} \binom{j-1}{k+s-i-1} \binom{n-j}{t+i} X_{j:n} \quad k=1,2,3,\dots \quad (4)$$

ile verilir (Elamir ve Seheult,2003; Hosking,2007).

Burada s ve t pozitif tam sayılardır. Buna göre kitle TL-momentleri ile örneklem TL-momentlerinden ilk r tanesinin eşitlenmesi ile elde edilen ve $\theta_1, \theta_2, \dots, \theta_r$ bilinmeyenlerine göre r tane denklemden oluşan denklem sistemlerinin çözümü olan,

$$\theta_1(X_1, X_2, \dots, X_n), \theta_2(X_1, X_2, \dots, X_n), \dots, \theta_r(X_1, X_2, \dots, X_n)$$

istatistiklerine TL-momentler yöntemi tahmin edicileri denir.

Yöntem

Verilerin Elde Edilmesi

Bu bölümde lojistik dağılımın bilinmeyen parametreleri MLE, ME, LSE, WLSE, PE, LME, TLME, tahmin yöntemleri ile incelenecektir.

En çok olabilirlik yöntemi

$X_1, X_2, \dots, X_n$, μ ve σ parametrelili lojistik dağılımından bir örneklem olsun. Bu durumda olabilirlik fonksiyonu;

$$L(\mu, \sigma; x) = \prod_{i=1}^n \frac{e^{-\left(\frac{X_i - \mu}{\sigma}\right)}}{\sigma \left(1 + e^{-\left(\frac{X_i - \mu}{\sigma}\right)}\right)^2} = \frac{e^{-\sum_{i=1}^n \left(\frac{X_i - \mu}{\sigma}\right)}}{\sigma^n \left(\prod_{i=1}^n \left(1 + e^{-\left(\frac{X_i - \mu}{\sigma}\right)}\right)\right)^2}$$

olarak bulunur.

$$\ln L(\mu, \sigma; x) = -\sum_{i=1}^n \left(\frac{X_i - \mu}{\sigma}\right) - n \ln \sigma - 2 \sum_{i=1}^n \ln \left(1 + e^{-\left(\frac{X_i - \mu}{\sigma}\right)}\right)$$

ile verilir. Burada $\frac{X_i - \mu}{\sigma}$ olarak ele alınsın. Bu durumda

$$\ln L(\mu, \sigma; x) = -\sum_{i=1}^n z_i - n \ln \sigma - 2 \sum_{i=1}^n \ln(1 + e^{-z_i})$$

olmak üzere bu fonksiyonu maksimum yapacak μ ve σ parametrelerini bulmak için μ ve σ 'ya göre kısmi türevler alınarak,

$$\frac{\partial \ln L}{\partial \mu} = \frac{n}{\sigma} - \frac{2}{\sigma} \sum_{i=1}^n \frac{1}{1 + e^{z_i}} = 0$$

$$\frac{\partial \ln L}{\partial \sigma} = \frac{-n}{\sigma} + \frac{1}{\sigma} \sum_{i=1}^n z_i - \frac{2}{\sigma} \sum_{i=1}^n \frac{z_i}{1+e^{z_i}} = 0$$

denklemlerine ulaşılır. Bu denklemlerin açık çözümü olmadığından Newton-Raphson gibi herhangi bir iterasyon yöntemi kullanılarak $\hat{\sigma}_{MLE}$ ve $\hat{\mu}_{MLE}$ elde edilir.

Momentler yöntemi

Lojistik dağılımının beklenen değer ve varyansının,

$$E(X) = \mu \text{ ve } Var(X) = \frac{\sigma^2 \pi^2}{3}$$

olduğu bilinmektedir.

Örneklem beklenen değeri ve örneklem varyansı,

$$\bar{X} = \frac{\sum_{i=1}^n X_i}{n} \text{ ve } S^2 = \frac{\sum_{i=1}^n (X_i - \bar{X})^2}{n-1}$$

olmak üzere lojistik dağılımın beklenen değer ve varyansının örneklem beklenen değer ve varyansına eşitlenmesiyle,

$$\sigma = \bar{X} \text{ ve } \frac{\sigma^2 \pi^2}{3} = \frac{\sum_{i=1}^n (X_i - \bar{X})^2}{n-1}$$

denklemlerini elde edilir. Bu denklemler sisteminin çözümünden μ ve σ 'nın tahmin edicisi sırasıyla;

$$\hat{\mu} = \bar{X} \text{ ve } \hat{\sigma} = \frac{\sqrt{3}}{\pi} S$$

olarak elde edilir.

L- moment yöntemi

Eşitlik (1) ve Eşitlik (2) ifadelerinin kullanılmasıyla lojistik dağılımın ilk iki kitle ve örneklem L-momentleri;

$$L_1 = E(X) \text{ ve } L_2 = \frac{1}{2} E(X_{2:2} - X_{1:2})$$

ve

$$l_1 = \frac{1}{n} \sum_{j=1}^n X_{j:n} = \bar{X} \text{ ve } l_2 = \frac{1}{n(n-1)} \sum_{j=1}^n (2j-n-1) X_{j:n}$$

olur. Bu ifadelerin kullanılmasıyla lojistik dağılımın ilk iki kitle L-momentleri

$$L_1 = \mu \text{ ve } L_2 = \sigma$$

olur. Bu durumda kitle L-momentlerinin örneklem L-momentlerine eşitlenmesi ile

$$\hat{\mu}_{LME} = \bar{X} \text{ ve } \hat{\sigma}_{LME} = \frac{1}{n(n-1)} \sum_{j=1}^n (2j-n-1) X_{j:n}$$

olarak elde edilir.

T-L moment yöntemi

Bu çalışmada, lojistik dağılım simetrik bir dağılım olduğundan $s=1$ ve $t=1$ alınmıştır. Böylece Eşitlik (3) ve Eşitlik (4)'den bu dağılım için ilk iki kitle ve örneklem TL-momentleri

$$\lambda_1^{(1,1)} = E(X_{2:3}) \text{ ve } \lambda_2^{(1,1)} = \frac{1}{2} [E(X_{3:4}) - E(X_{2:4})]$$

ve

$$l_1^{(1,1)} = \frac{6}{n(n-1)(n-2)} \sum_{j=1}^{n-1} (n-j)(j-1) X_{j:n} \quad l_2^{(1,1)} = \frac{6 \sum_{j=1}^{n-1} (j-1)(n-j)(2j-n-1) X_{j:n}}{n(n-1)(n-2)(n-3)}$$

olur. Bu ifadelerin kullanılmasıyla lojistik dağılımın ilk iki kitle TL-momentleri

$$\lambda_1^{(1,1)} = \mu \text{ ve } \lambda_2^{(1,1)} = \frac{\sigma}{2}$$

olarak elde edilir. Kitle TL-momentlerinin örneklem TL-momentlerine eşitlenmesiyle

$$\mu = \frac{6}{n(n-1)(n-2)} \sum_{j=1}^{n-1} (n-j)(j-1) X_{j:n}$$

$$\frac{\sigma}{2} = \frac{6}{n(n-1)(n-2)(n-3)} \sum_{j=1}^{n-1} (j-1)(n-j)(2j-n-1) X_{j:n}$$

denklem sistemlerine ulaşılır. Bu denklem sistemlerinin çözümünden μ ve σ parametrelerinin TL-moment tahmin edicileri sırasıyla;

$$\hat{\mu}_{TLME} = \frac{6}{n(n-1)(n-2)} \sum_{j=1}^{n-1} (n-j)(j-1) X_{j:n}$$

ve

$$\hat{\sigma}_{TLME} = \frac{12 \sum_{j=1}^{n-1} (j-1)(n-j)(2j-n-1) X_{j:n}}{n(n-1)(n-2)(n-3)}$$

olur.

En küçük kareler yöntemi

$X_1, X_2, \dots, X_n$ lojistik dağılımdan bir örneklem ve $X_{1:n} \leq X_{2:n} \leq \dots \leq X_{n:n}$ ler bu örnekleme karşılık gelen sıra istatistikleri olsun. Bu durumda lojistik dağılım için

$$\sum_{i=1}^n \left(F(X_{i:n}) - \frac{i}{n+1} \right)^2 = \sum_{i=1}^n \left(\left(1 + e^{-\frac{(X_{i:n} - \mu)}{\sigma}} \right)^{-1} - \frac{i}{n+1} \right)^2$$

ifadesi elde edilir. Bu ifadenin μ ve σ 'ya göre minimize edilmesiyle μ ve σ 'nın en küçük kareler tahmin edicisi olan $\hat{\mu}_{LSE}$ ve $\hat{\sigma}_{LSE}$ tahmin edicileri elde edilir.

Ağırlıklandırılmış en küçük kareler yöntemi

$X_1, X_2, \dots, X_n$ lojistik dağılımdan bir örneklem ve $X_{1:n} \leq X_{2:n} \leq \dots \leq X_{n:n}$ ler bu örnekleme karşılık gelen sıra istatistikleri olsun. Bu durumda lojistik dağılım için

$$\sum_{i=1}^n W_i \left(F(X_{i:n}) - \frac{i}{n+1} \right)^2 = \sum_{i=1}^n W_i \left(\left(1 + e^{-\left(\frac{X_{i:n} - \mu}{\sigma} \right)} \right)^{-1} - \frac{i}{n+1} \right)^2$$

denklemini elde edilir. Bu denklemin μ ve σ 'ya göre minimize edilmesiyle μ ve σ 'nın ağırlıklandırılmış en küçük kareler tahmin edicisi olan $\hat{\mu}_{WLSE}$ ve $\hat{\sigma}_{WLSE}$ tahmin edicileri elde edilir.

Burada

$$W_i = \frac{1}{\text{Var}(F(X_{i:n}))} = \frac{(n+1)^2(n+2)}{i(n-i+1)}, \quad i=1, 2, \dots, n$$

dır.

Yüzdelik yöntem

$X_1, X_2, \dots, X_n$ lojistik dağılımdan bir örneklem ve $X_{1:n} \leq X_{2:n} \leq \dots \leq X_{n:n}$ ler bu örnekleme karşılık gelen sıra istatistikleri olsun. Bu durumda lojistik dağılım için

$$\sum_{i=1}^n \left(X_{i:n} - F^{-1} \left(\frac{i}{n+1} \right) \right)^2 = \sum_{i=1}^n (X_{i:n} - \mu - \sigma z_i)^2$$

denklemini elde edilir. Bu denklemin minimize edilmesiyle μ ve σ 'nın yüzdelik tahmin edicileri

$$\hat{\sigma}_{PE} = \frac{n \sum_{i=1}^n X_{i:n} z_i - \sum_{i=1}^n z_i \sum_{i=1}^n X_{i:n}}{n \sum_{i=1}^n z_i^2 - \left(\sum_{i=1}^n z_i \right)^2}, \quad i=1, 2, \dots, n \quad \text{ve} \quad \hat{\mu}_{PE} = \frac{\sum_{i=1}^n X_{i:n} - \hat{\sigma} \sum_{i=1}^n z_i}{n}, \quad i=1, 2, \dots, n$$

olur. Burada $z_i = \ln \left(\frac{n+1}{n+1-i} \right)$, dır.

İstatistiksel Analizler

Bu bölümde μ ve σ parametreleri için önerilen tahmin edicilerin performansları bir simülasyon çalışması ile incelenmiştir. Tahmin edicilerin performanslarını değerlendirmek için yan ve MSE değerleri kullanılmıştır. Yapılan simülasyon çalışmasında $\sigma=1$ ve $\mu=0$ olarak alınmıştır. Ayrıca tahmin edicilerin küçük ve büyük örneklem özelliklerini incelemek için n örneklem hacmi sırasıyla 20, 30, 50 ve 100 olarak alınmıştır. Sonuçlar Çizelge 1'de verilmiştir. Tablolar 10000 simülasyon üzerine kurulmuştur.

Çizelge 1. Lojistik dağılım ($\mu = 1, \sigma = 0$) durumunda μ ve σ parametreleri için önerilen klasik tahmin edicilerin tahmin ve MSE'leri.

n	Tahminci	$\hat{\mu}$	MSE	$\hat{\sigma}$	MSE
		Tahmin		Tahmin	
20	MLE	-0.0172	0.1560	0.9632	0.0358
	ME	-0.0201	0.1710	0.9764	0.0390
	LSE	-0.0145	0.1568	1.0552	0.0532
	WLSE	-0.0178	0.1544	1.0380	0.0488
	PE	-0.0214	0.1472	1.1239	0.0639
	LME	-0.0201	0.1710	0.9985	0.0373
	TLME	-0.0170	0.1555	0.9964	0.0466
30	MLE	-0.0077	0.1066	0.9849	0.0222
	ME	-0.0066	0.1147	0.9904	0.0242
	LSE	-0.0073	0.1080	1.0350	0.0351
	WLSE	-0.0079	0.1044	1.0204	0.0285
	PE	-0.0063	0.1001	1.1001	0.0385
	LME	-0.0066	0.1147	1.0021	0.0229
	TLME	-0.0078	0.1068	0.9989	0.0287
50	MLE	-0.0068	0.0604	0.9812	0.0147
	ME	-0.0068	0.0646	0.9869	0.0161
	LSE	-0.0057	0.0615	1.0108	0.0201
	WLSE	-0.0054	0.0589	1.0099	0.0170
	PE	-0.0066	0.0504	1.0638	0.0218
	LME	-0.0066	0.0646	0.9936	0.0148
	TLME	0.0076	0.0605	0.9901	0.0173
100	MLE	0.9999	0.0278	0.9966	0.0072
	ME	0.9956	0.0250	0.9978	0.0080
	LSE	0.9977	0.0190	1.0144	0.0099
	WLSE	1.0003	0.0171	1.0067	0.0082
	PE	1.0104	0.0278	1.0461	0.0107
	LME	1.0071	0.0176	1.0022	0.0074
	TLME	1.0107	0.0104	1.0030	0.0084

BULGULAR

Çizelge 1'deki simülasyon sonuçlarına göre tüm durumlarda tahmin edicilerin yan ve MSE'leri n artıkça azalmaktadır. Bu durum tahmin edicilerin asimptotik yansız ve tutarlı olduğunu göstermektedir.

Çizelge 1'den TLME, MLE ve PE tahmin edicilerinin MSE'leri diğer tahmin edicilere nazaran küçük örnek hacimlerinde daha düşük çıkmıştır. Ayrıca tüm örnek hacimlerinde LME, TLME, LSE ve WLSE tahmin edicilerinin performanslarının iyi olduğu görülmektedir. LME ve TLME tahmin edicileri kendi aralarında karşılaştırıldığında TLME tahmin edicisinin performansının daha iyi olduğu anlaşılmaktadır. Benzer şekilde LSE ve WLSE tahmin edicileri kendi aralarında karşılaştırıldığında her ne kadar LSE tahmin edicisinin yanı sıra düşük çıksada yan ve MSE birlikte düşünüldüğünde WLSE tahmin edicisinin daha iyi olduğu söylenebilir.

TARTIŞMA VE SONUÇ

Literatürde lojistik dağılım finans, yenilenebilir enerji ve mühendislik gibi birçok alanda normal dağılım ve lognormal dağılıma alternatif olarak sıklıkla kullanılmıştır. Bakınız (Braserton ve ark., 1999; Prince ve ark., 1988; Jeon ve Kang 2020; Chai ve ark., 2017).

Parametre tahmini istatistiğin en önemli amaçlarından biri olmasının yanı sıra pek çok alanda yaygın olarak kullanılmaktadır. Bu nedenle bu çalışmada, lojistik dağılımının farklı parametre tahmin yöntemleri kullanılarak μ ve σ parametreleri için tahmin ediciler elde edilmiştir. Genel olarak

simülasyon sonuçlarına göre tahmin edicilerin yansız ve tutarlı olduğu ve MSE' lerinin n arttıkça azaldığı görülmüştür. Böylece tahmin edicilerin asimptotik yansız ve tutarlı olduğu simülasyon ile gösterilmiştir. Ayrıca, LSE ve WLSE yöntemleri ile elde edilen tahmin ediciler kendi aralarında karşılaştırıldıklarında beklenildiği gibi WLSE yöntemi ile elde edilen tahmin edicilerin daha iyi performans gösterdiği gözlemlenmiştir. Son olarak LME ve TLME yöntemleri ile elde edilen tahmin ediciler kendi aralarında karşılaştırıldıklarında genel olarak TLME tahmin edicisinin performansının daha iyi olduğu anlaşılmıştır.

Kaynaklar

- Braselton, J., Rafter, J., Humphrey, P., & Abell, M.,1999. Randomly walking through Wall Street: Comparing lump-sum versus dollar-cost average investment strategies. *Mathematics and computers in simulation*, 49(4-5), 297-318.
- Chai, H., He, R., Ma, C., Dai, C., & Zhou, K.,2017. Path planning and vehicle scheduling optimization for logistic distribution of hazardous materials in full container load. *Discrete Dynamics in Nature and Society*, 2017.
- Dey, S., Mazucheli, J., Nadarajah, S., 2018. Kumaraswamy distribution: different methods of estimation, *Computational and Applied Mathematics*, 3: 2094-2111.
- Dey, S., Raheem, E., Mukherjee, S., 2017. Statistical properties and different methods of estimation of transmuted Rayleigh distribution. *Revista Colombiana de Estadística*, 40: 165-203.
- Elamir, E.A.H., Seheult, A.H, 2003. Trimmed L-moments, *Computational statistics & Data analysis*,43: 299-314.
- Hosking, J.R.M., 1990. L-moments analysis and estimation of distributions using linear combinations, *F statistics. J.Roy. Statistics. Soc. B.*,52:105-124.
- Hosking, J.R.M., 2007. Some theory and practical uses of trimmed L-moments, *Journal Of Statistical Planning And Inference*,137: 3024-3039.
- Jeon, Y. E., & Kang, S. B.,2020. Estimation for the half-logistic distribution based on multiply Type-II hybrid censoring. *Physica A: Statistical Mechanics and its Applications*, 550, 124501.
- Prince, J. R., Mumma, C. G. and Kouvelis, A. (1988). Applications of the logistic distribution to 125 I sensitometry in nuclear-medicine. *Journal of Nuclear Medicine*, 29, 273{273.

Türkiye’de Tüketici Fiyat Endeksi ile Yurt İçi Üretici Fiyat Endeksi, Yurt Dışı Üretici Fiyat Endeksi ve Tarım Ürünleri Üretici Fiyat Endeksi Arasındaki İlişki Analizi

İlhami MİNTEMUR^{1*}

¹Türkiye İstatistik Kurumu, Ankara

*Sunucu yazar e-posta: ilhamimintemur.@gmail.com

Özet

Bu çalışmada 2010-2019 yıllarına ait aylık bazda veriler kullanılarak, Türkiye’de fiyat değişimlerini gösteren bazı endekslerin (TÜFE, YİÜFE, YDÜFE, TARÜFE), TÜFE endeksi üzerindeki etkileri, bu değişkenlerin zaman içinde TÜFE değişkeni ile nasıl hareket ettikleri birim kök testleri, Var gecikme uzunluğu, eşbütünleşme, Johansen Eşbütünleşme testi, değişkenler arasında kısa ve uzun dönem ilişkileri ve Vektör Hata Düzeltme Modeli sonuçları kullanılarak incelenmiştir. Yapılan analizler sonucunda, TÜFE ile diğer endeksler arasında uzun dönemli bir denge ilişkisi bulunurken, kısa dönemde TARÜFE değişkeninin, TÜFE değişkeninin nedeni olduğu, YİÜFE ve YDÜFE değişkenlerinin ise TÜFE değişkeninin nedeni olmadığı sonucuna ulaşılmıştır.

Anahtar kelimeler: Birim kök, Eşbütünleşme, Johansen Eşbütünleşme, Vektör Hata Düzeltme Modeli

GİRİŞ

Bu çalışmada fiyat seviyesini gösteren enflasyonu ifade etmek üzere kullanılan Tüketici Fiyat Endeksi (TÜFE) ve Yurt İçi Üretici Fiyat Endeksi (YİÜFE), Yurt Dışı Üretici Fiyat Endeksi (YDÜFE) ve Tarım Ürünleri Üretici Fiyat Endeksi (TARÜFE) olmak üzere dört endeks kullanılmıştır.

Buna göre, TÜFE nihai mal ve hizmet fiyatlarındaki değişimi ifade ederken, YİÜFE belirli bir referans döneminde ülke ekonomisinde üretimi yapılan ve yurt içine satışa konu olan ürünlerin fiyat değişimlerini, TARÜFE çiftçinin üreterek piyasaya satışını yaptığı ürünlerin ilk el satış fiyatlarındaki değişimi göstermektedir.

Bu çalışmada, TÜFE, YİÜFE, YDÜFE ve TARÜFE değişkenleri arasındaki ilişkiler incelenmiştir. Bununla birlikte, enflasyonun niteliğini belirleme konusunda TÜFE, YİÜFE ile TARÜFE arasındaki ilişkinin incelenmesi önem arz etmektedir. Buna göre, girdi fiyatlarındaki yükselişten kaynaklanan maliyet enflasyonu önce YİÜFE’yi artıracak, YİÜFE-TÜFE arasındaki aktarma mekanizmasının nitelik ve hızına göre gecikmeli olarak TÜFE’ye de yansımaya (Saraç ve Karagöz, 2010). Aynı şekilde TARÜFE’de meydana gelen maliyet artışları gecikmeli olarak TÜFE’ye yansımaya. Diğer taraftan, gelir artışı veya beklentilerden kaynaklanan bir talep enflasyonu önce TÜFE’yi yine ilk durumda olduğu gibi gecikmeli olarak YİÜFE ve TARÜFE’yi etkileyebilmektedir. Bu nedenle, TARÜFE, YİÜFE ile TÜFE arasındaki ilişkinin belirlenmesi ekonomi politikaları açısından önem arz etmektedir (Saraç ve Karagöz, 2010).

TARÜFE ile TÜFE arasında direkt bir ilişki mevcuttur. Tarım ürünlerinin ilk elden satışları, toptancı veya perakendeciler yolu ile direkt TÜFE’ye yansımaktadır. Tarım ürünlerinin ilk elden satış fiyat değişimleri, önce toptancı veya direkt sanayi girdilerini etkileyecek, sanayi girdileri maliyet değişimlerinden dolayı YİÜFE’yi etkileyecektir (Saraç ve Karagöz, 2010). YDÜFE endeksi ile YİÜFE arasındaki ilişki önemlidir. Ayrıca YDÜFE döviz kurlarından, global ekonomiden ve jeopolitik risklerden etkilendiği için öncü gösterge olarak değerlendirilebilir.

Diğer bir görüş, bu endeksler arasında herhangi bir istatistiksel bağlantının bulunmadığını öne sürmektedir. Endeksleri farklı metodoloji, farklı ana kitleler için farklı örneklemelerden elde edilen fiyatlara dayalı olarak hesaplanmaktadır (Saraç ve Karagöz, 2010). YİÜFE ve TÜFE her ayın 3'ünde, TARÜFE her ayın 14-16'sında ve YDÜFE ise her ayın 20-22 'sinde Türkiye İstatistik Kurumu (TÜİK) tarafından yayınlanmaktadır.

VERİ SETİ VE YÖNTEM

Bu çalışma, dahil edilen değişkenlerin (endekslerin) zaman serileri analizine göre dört bölümde incelenmiştir. Birinci bölümde serilerin durağanlıkları, ikinci bölümde uygun gecikme uzunluğu incelenmiştir. Üçüncü bölümde Johansen Eşbütünleşme analizi ve son bölümde Vektör Hata Düzeltme Modeli (VECM) uygulanarak sonuçlar yorumlanmıştır. Modelde kullanılan seriler, Türkiye İstatistik Kurumu'nun web sayfalarından alınmıştır. Tüm endeksler 2010 Ocak- 2019 Aralık dönemi (2010=100), aylık bazda ve logaritma formuna dönüştürülerek ve sırasıyla LNTUFE, LNYDUFE, LNYIUFE ve LNTARUFE simgeleriyle kodlanmıştır.

Durağanlık Analizi

Zaman serileri tesadüfi değişkenlerle çalışmaktadır. Tesadüfi bir süreç izleyen zaman serilerinde serinin durağan olup olmadığı önem kazanmaktadır. Tesadüfi bir değişkenin zaman içindeki ortalamasının, varyansının ve kovaryansının sabit olması durumuna durağanlık denilmektedir. Eğer bir zaman serisi durağan değilse, seri ile yapılan çalışma, ele alınan tahmin dönemi için geçerli olacaktır. Diğer dönemler için genelleme yapılamaz. Açıklayıcı değişkenlerden herhangi birisi durağan olmadığında regresyon modeli bozulur. Zaman serisi verileri ile çalışırken serilerin aynı seviyede durağan olması önemlidir. Aksi takdirde değişkenler arasında gerçekten bir ilişki olmamasına rağmen durağan olmayan değişkenlerin trendinden kaynaklı bir ilişki görülebilir ve bu da sahte regresyon sorununa yolaçabilmektedir. Ayrıca serilerin aynı seviyede durağan olmaları eşbütünleşme analizi için ön koşuldur (Emrah ve Kanat, 2018).

Zaman serisi analizlerinde çeşitli durağanlık sınamaları olmasına rağmen bu çalışmada durağanlık testi olarak Düzeltilmiş Dukey Fuller (ADF) testi kullanılmıştır.

Tablo 1. Serilerin ADF Birim Kök Testi Sonuçları

Değişkenler	Düzyey (%1 anlamlılık)		1. Sıra Fark (%1 anlamlılık)	
	Sabit	Sabit+Trend*	Sabit	Sabit+Trend*
TÜFE	2.672 (prob.1.00)	-0.358 (prob.0.98)	-6.267 (prob.0.00)	-7.023 (prob.0.00)
YDÜFE	0.7413 (prob.0.99)	-1.559 (prob.0.80)	-8.530 (prob.0.00)	-8.642 (prob.0.00)
YİÜFE	1.171 (prob.0.99)	-1.004 (prob.0.94)	-7.11 (prob.0.00)	-7.315 (prob.0.00)
TARÜFE	-1.392 (prob.0.58)	-3.390 (prob.0.06)	-10.395 (prob.0.00)	-10.348 (prob.0.00)

*:Bu çalışmada sabit+trend modeli kullanılmıştır.

Tablo 1 incelendiğinde, ADF test istatistiğine göre TÜFE, YDÜFE, YİÜFE ve TARÜFE endeks serileri düzeyde durağan değilken, birinci farkları alındığında ADF test istatistiğinin değeri %1 anlamlılık düzeyinde serilerin durağanlaştığını göstermektedir. Serilerin birinci fark değerleri için birim

kök olduğunu ifade eden sıfır hipotezi reddedilmiştir. Çalışmada kullanılan tüm değişkenlerin durağanlık derecesi $I(1)$ 'dir. Yani birinci farkı alınmış serilerin (sabit (1), sabit ve trend (2)) durağan olduğu saptanmıştır.

$$\Delta Y_t = \alpha + \rho Y_t - 1u_t \quad (1) \text{ sabit}$$

$$\Delta Y_t = \alpha + \beta t + \rho Y_t - 1u_t \quad (2) \text{ sabit+trend}$$

Tablo 2. Uygun Gecikme Uzunluğunun Belirlenmesi

Lag	LogL	LR	FPE	AIC	SC	HQ
0	551.5353	NA	5.58e-10	-9.955186	-9.856987	-9.915356
1	1243.172	1320.398	2.58e-15	-22.23949	-21.74850	-22.04034
2	1289.403	84.89711	1.49e-15*	-22.78915*	-21.90536*	-22.43068*
3	1298.429	15.91782	1.70e-15	-22.66234	-21.38575	-22.14455
4	1307.643	15.58113	1.93e-15	-22.53897	-20.86958	-21.86186
5	1327.764	32.55866	1.81e-15	-22.61389	-20.55171	-21.77746
6	1347.735	30.86461*	1.70e-15	-22.68610	-20.23111	-21.69034
7	1362.776	22.15039	1.77e-15	-22.66865	-19.82087	-21.51357
8	1379.504	23.41942	1.79e-15	-22.68189	-19.44131	-21.36749
9	1397.595	24.01133	1.78e-15	-22.71990	-19.08653	-21.24618
10	1413.025	19.35819	1.87e-15	-22.70954	-18.68337	-21.07651

*: Seçilen gecikme düzeyini gösterir. Her ardışık test için %5 anlamlılık düzeyinde LR test kriteri, FPE: Final Prediction kriteri, AIC: Akaike Bilgi Kriteri, SC: Schwarz bilgi kriteri, HQ: Hannan Quin bilgi kriteridir.

Tablo 2 incelendiğinde kriterlerin çoğunluğu (LogL, LR, FPE, AIC, SC, HQ) hangi gecikme sayısını işaret ediyorsa o gecikme sayısını kullanmak modeli daha güvenilir hale getirmektedir (Kocabıyık, 2016). Modelde en uygun gecikme uzunluğu tüm bilgi kriterlerine göre 2 olarak seçilmiştir. TÜFE ile ÜFE arasındaki ilişki çalışmasında TÜFE, ÜFE modelinde gecikme sayısını 3, ÜFE, TÜFE modelinde gecikme sayısını 1 olarak bulunmuştur (Saraç ve Karagöz, 2010). Serilerin aynı seviyede durağanlaştıkları tespit edildikten sonra ve VAR gecikme sayısı belirlenerek, eşbütünleşme testine geçilebilmektedir (Kocabıyık, 2016). Değişkenlerin (TÜFE, YDÜFE, YİÜFE ve TARÜFE) birinci sıra farklarının durağan çıkmaları, yani düzeyde durağan olmaları nedeniyle aralarındaki uzun ve kısa dönem ilişkilerini incelemek amacıyla Johansen eşbütünleşme testi uygulanmıştır.

Johansen Eş Bütünleme Analizi

Zaman serileri arasında eşbütünleşme ilişkisini ortaya koyabilmek için öncelikle serilerin hangi seviyede durağan olduklarının incelenmesi gerekmektedir. Seriler eğer aynı seviyede durağan hale geliyorsa eşbütünleşme analizine dahil edilebilir (Kocabıyık, 2016). Serilerin aynı seviyede durağanlaştıkları tespit edildikten ve VAR gecikme sayısı belirlendikten sonra eşbütünleşme testine geçilebilmektedir. Tablo 1'de görüldüğü gibi tüm seriler 1. farkları %1 anlamlılık düzeyinde durağan ve Tablo 2'de uygun uzunluk 2 olarak bulunmuştur.

Tablo 3. Johansen Eş Bütünleme Test Sonuçları

Örneklem (Düzeltilmiş): 2010M04 2019M12							
Dahil edilen gözlem sayısı: 117							
Trend varsayımı: Doğrusal Deterministik trend (Kısıtlı)							
Seriler: TÜFE YDÜFE YİÜFE TARÜFE							
Geçikme aralığı (1. farkları): 1- 2							
Maximum Özdeğer Testi (Maximum Eigen value)				İz Testi (Trace Test)			
H ₀ Hipotez	Alternatif Hipotez	Test İstatistiği	%5 kritik Değeri	H ₀ Hipotez	Alternatif Hipotez	Test İstatistiği	%5 kritik Değeri
$r=0$	$r=1$	69.706	63.876	$r=0$	$r=1$	34.620	32.118
$r=1$	$r=2$	35.086	42.915	$r=1$	$r=2$	21.091	25.823
$r=2$	$r=3$	13.995	25.872	$r=2$	$r=3$	10.486	19.387

Tablo 3'teki sonuçlar incelendiğinde hem maksimum öz değer testi hem de İz testi açısından ele alınan seriler arasında uzun dönemli bir ilişkinin varlığı görülmektedir. Herhangi bir eşbütünleşik vektörün bulunmadığı ifade edilen (H_1 : *Hiç Eş bütünleşme Vektörü yoktur*) hipotez için maksimum özdeğer 69.706, %5 anlamlılık düzeyindeki kritik değer olan 63.876 değerinden ve İz testi açısından da değerler sırasıyla 34.620 ve 32.118 değerinde büyük olduğu için H_1 hipotezi reddedilmektedir. Her iki istatistik de eşbütünleşik vektörü göstermektedir. Yani 1 tane eşbütünleşme denklemi bulunmuştur.

Değişkenlerin kısa dönemdeki dengeden sapma eğilimlerinin vektör hata düzeltme modeli çerçevesinde ele alınabileceğini göstermektedir. Böylece uzun dönem de değişkenlerin eşbütünleşik oldukları, uzun dönem birlikte hareket ettikleri ve birbirlerini etkilediklerini söylemek mümkündür.

Vektör Hata Düzeltme Modelleri

Değişkenler arasında eşbütünleşme olduğunun belirlenmesi halinde, kısa dönemli dengesizlikleri gideren bir vektör hata düzeltme modeli olduğunu gösterilmiştir (Kocabıyık,2016). Vektör hata düzeltme modeli, uzun dönem ilişki yani dengeden sapmayı gösterir. Değişkenler arasındaki uzun dönem ilişkiyi eşbütünleşme olarak görebiliriz. Ancak serilerin aynı dereceden durağan olması gerekir. Durağanlık testi için serilere fark uygulanmaktadır. Uzun dönem serilerinde fark uygulanması veri kayıplarına sebep olmaktadır. Bu nedenle hata düzeltme modelleri kullanılmaktadır.

Tablo 4. Vektör Hata Düzeltme Tahminleri

Örneklem (Düzeltilmiş): 2010M04 2019M12	
Dahil edilen gözlem sayısı: 117	
Standart hata () ve t -istatistiği []	
Eşbütünleşme denklemi	CointEq1
TÜFE (-1)	1.0000
YDÜFE (-1)	-0.379000 (0.17658) [-2.14634]
YİÜFE	0.029760 (0.22029) [0.13510]
TARÜFE (-1)	0.288638 (0.07445) [3.87687]
@TREND (10M01)	-0.005689 (0.00060) [-9.49377]
C	-4.252265

Aşağıdaki denklem tüm değişkenler arasında uzun dönem ilişkisi ortaya koyan denklemdir, e_t ise hata terimini göstermektedir.

$$D(TÜFE) = C(1) * (TÜFE(-1) - 0.3790 * YDÜFE(-1) + 0.0297 * YİÜFE(-1) + 0.2886 * TARÜFE(-1) - 0.0056 * @TREND(10M01) - 4.2522) + e_t \quad (3)$$

Denklemden bağımlı değişken TÜFE endeksi ile diğer bağımsız değişkenlerin ilişkileri görülmektedir. Bu denkleme göre hata düzeltme katsayısı negatif ve istatistiksel olarak anlamlıdır (hata düzeltme katsayısı=-0.156602, p değeri=0.000). Endekslerin uzun dönemde birlikte hareket ettiklerine işaret etmektedir. TÜFE ile ÜFE arasındaki ilişki çalışmada hata düzeltme kat sayısı pozitif ve istatistiksel olarak anlamlı bulunmamıştır (Saraç ve Karagöz, 2010). Buna göre kısa dönemde meydana gelen dengesizliklerin uzun dönemde düzeleceği görülmektedir. Uzun dönemde YİÜFE ile TARÜFE değişkenleri TÜFE endeksi ile aynı yönde hareket ederken, YDÜFE endeksi ile negatif yönde hareket etmektedir. YDÜFE aynı zamanda döviz kuru olarak yorumlanabilir. Çünkü ihracat genellikle döviz kuru üzerinden yapılmaktadır. Döviz kurunun artması, TÜFE değişkeni üzerinde negatif etkiye sahiptir.

Tablo 5. Vektör Hata Düzeltme Model Sonuçları

Değişkenler	Katsayılar	Standard Hata	T-İstatistikleri	Olasılık Değerleri
ECM _{t-1}	-0.156602	0.028628	-5.470223	0.0000*
D (TÜFE (-1))	-0.009347	0.097636	-0.095734	0.9239
D (TÜFE (-2))	-0.187545	0.099853	-1.878206	0.0631
D (YDÜFE (-1))	0.062197	0.047726	1.303200	0.1953
D (YDÜFE (-2))	-0.053118	0.044799	-1.185697	0.2384
D (YİÜFE (-1))	0.133069	0.121550	1.094770	0.2761
D (YİÜFE (-2))	-0.037071	0.098647	-0.375796	0.7078
D (TARÜFE (-1))	0.043282	0.020209	2.141669	0.0345*
D (TARÜFE (-2))	0.033918	0.020844	1.627225	0.1066
Sabit terim	0.008028	0.001064	7.545663	0.0000*

$R^2 = 0,50$, F-istatistik değeri =11.954, p (0,000), Durbin-Watson: 1,938

*: İlgili değişkenleri %5 önem düzeyinde istatistiksel açıdan anlamlı olduğunu ifade etmektedir.

Hata düzeltme modelin kısa ve uzun dönem denklemi (4)'te verilmiştir.

$$D(TÜFE) = C(1) * (TÜFE(-1) - 0.379000 * YDÜFE(-1) + 0.029760 * YİÜFE(-1) + 0.288637 * TARÜFE(-1) - 0.005688 * @TREND(10M01) - 4.252265) + C(2) * D(TÜFE(-1)) + C(3) * D(TÜFE(-2)) + C(4) * D(YDÜFE(-1)) + C(5) * D(YDÜFE(-2)) + C(6) * D(YİÜFE(-1)) + C(7) * D(YİÜFE(-2)) + C8 * D(TARÜFE(-1)) + C(9) * D(TARÜFE(-2)) + C(10) \quad (4)$$

C (2)'den C (9)'ye kadar katsayılar kısa dönem ilişkisi gösteren katsayılardır. Burada diğer endekslerin kısa dönemde D(TÜFE) bağımlı değişkeni etkileyip etkilemediği araştırılmaktadır.

Tablo 6 - Walt Test Sonuçları

Yokluk hipotezleri	P değerleri	
	F - Olasılık değeri	Ki-Kare-Olasılık değeri
H ₀ : [D (YDÜFE (-1)) = D (YDÜFE (-2))] = [C (4) = C (5) = 0]	0.2096	0.2048
H ₀ : [D (YİÜFE (-1)) = D (YİÜFE (-2))] = [C (6) = C (7)] = 0	0.5435	0.5416
H ₀ : [D (TARÜFE (-1)) = D (TARÜFE (-2))] = [C (8) = C (9)] = 0	0.0359	0.0323*

*: %5 önem düzeyinde istatistiksel açıdan anlamlı olduğunu ifade etmektedir.

Tablo 6'da TÜFE ile diğer endekslerin kısa dönem ilişkileri görülmektedir. Seçilen endekslerin kısa dönem ilişkilerine bakıldığında, YİÜFE ve YDÜFE değişkenlerinin birinci ve ikinci gecikmeleri kısa dönemde TÜFE değişkeninin nedenselliği olarak görülmemektedir. Fakat TARÜFE değişkeninin birinci ve ikinci gecikmeleri kısa dönemde TÜFE değişkeninin nedenselliği olarak görülmektedirler. 1991-2009 yılları kapsayan TÜFE ve ÜFE çalışmasında TÜFE ve ÜFE oranları arasında nedensellik ilişkisi bulunmamıştır (Saraç ve Karagöz, 2010).

Tablo 7. Tanısal testler

Hata terimlerine ilişkin testler		p-değeri
Otokorelasyon (Breusch-Godfrey LM Testi)	F-istatistiği: 1.5009 Obs*R-squared: 3.2519	Prob: 0.2277 Prob: 0.1967
Değişen varyans (White test)	F-istatistiği: 1,2618 Obs*R-squared: 61.2557	Prob: 0.1875 Prob: 0.2318
Normal dağılım (Jarque-Bera)	Jarque-Bera:1.9608	Prob: 0.3751

Vektör hata düzeltme modeliyle ortaya koyduğumuz sonuçların kabul edilebilmesi için, modelim artıklarının normal dağılması, otokorelasyona sahip olmaması ve değişen varyans problemi olmaması gerekmektedir (Kocabıyık,2016). Yapılan testlere göre, Vektör Hata Düzeltme Modeli'nin varsayımları uygun olduğu sonucuna varılmıştır.

BULGULAR

TÜFE, YDÜFE, YİÜFE ve TARÜFE değişkenlerinin, düzeyde %1 anlamlılıkla, ADF test istatistikleri sonucuna göre durağan olmadıkları görülmektedir. Birinci farkları alınmış tüm değişkenler durağan olduğu saptanmıştır. VAR modeli için uygun gecikme uzunluğu 2 olarak bulunmuştur. Hem maksimum öz değer testi hem de İz testi açısından ele alınan seriler arasında uzun dönemli bir ilişkinin varlığı tespit edilmiştir. Seriler arasında 1 eşbütünleşik vektör olduğu sonucuna ulaşılmıştır. Vektör Hata Düzeltme Model'ine göre hata düzeltme katsayısı negatif ve istatistiksel olarak anlamlıdır. (Hata düzeltme

katsayısı=-0.156602, p değeri=0.000) Endeksler uzun dönemde birlikte hareket etmektedir. Buna göre kısa dönemde meydana gelen dengesizliklerin uzun dönemde düzeleceğini göstermektedir. Uzun dönemde, YİÜFE ve TARÜFE ve değişkenleri TÜFE endeksi ile pozitif yönde hareket ederken, YDÜFE ile negatif yönde hareket etmektedir.

YİÜFE ve YDÜFE değişkenlerinin birinci ve ikinci gecikmeleri kısa dönemde TÜFE değişkeninin nedenselliği olarak görülmemektedir. Fakat TARÜFE değişkeninin birinci ve ikinci gecikmeleri kısa dönemde TÜFE değişkeninin nedenselliği olarak görülmektedirler

Kaynaklar

- Saraç, T. B., & Karagöz, K., 2010. Türkiye’de tüketici ve üretici fiyatları arasındaki ilişki: Yapısal kırılma ve sınır testi. Maliye Dergisi, 159, 220-232.
- Türkiye İstatistik Kurumu, 2021. "Haber Bültenleri". Erişim Adresi: http://www.tuik.gov.tr/PreTablo.do?alt_id=1014; Erişim Tarihi:22.04.2021.
- Türkiye İstatistik Kurumu, 2021. "Haber Bültenleri". Erişim Adresi: http://www.tuik.gov.tr/PreTablo.do?alt_id=1076; Erişim Tarihi:22.04.2021.
- Türkiye İstatistik Kurumu, 2021. "Haber Bültenleri". Erişim Adresi: http://www.tuik.gov.tr/PreTablo.do?alt_id=1099; Erişim Tarihi:22.04.2021.
- Türkiye İstatistik Kurumu, 2021. "Haber Bültenleri". Erişim Adresi: http://www.tuik.gov.tr/PreTablo.do?alt_id=1101; Erişim Tarihi:22.04.2021.
- Türkiye Cumhuriyeti Merkez Bankası, 2019. <https://evds2.tcmb.gov.tr/index.php?/evds/portlet/hIdR20CDwM4%3D/tr>; Erişim Tarihi:15.03.2019.
- Emrah, Ö. G. E. T., & Kanat, E., 2018. Bitcoin ile Türkiye ve G7 Ülke Borsaları Arasındaki Uzun ve Kısa Dönemli İlişkilerin İncelenmesi. Finans Ekonomi ve Sosyal Araştırmalar Dergisi (FESA), 3(3), 601-614.
- Kocabıyık, T., 2016. Johansen Eşbütünleme Testinde Karar Aşamalarının Analizi. Sosyal Bilimler Enstitüsü Dergisi
- Petek, A., & Çelik, A., 2017., Türkiye’de Enflasyon, Döviz Kuru, ihracat ve ithalat Arasındaki İlişkinin Ekonometrik Analizi (1990-2015). Finans Politik & Ekonomik Yorumlar, 54(626), 69.
- Mucuk, M., & Alptekin, V., 2008. Türkiye’de Vergi ve Ekonomik Büyüme İlişkisi: VAR Analizi (1975-2006). Maliye Dergisi, 155(2), 159-174.
- Bozdağlıoğlu, E. Y. U., 2007. Türkiye’nin İthalat ve İhracatının Eşbütünleşme Yöntemi ile Analizi (1990-2007). Gazi Üniversitesi İktisadi ve İdari Bilimler Fakültesi Dergisi, 9(3), 213-224.
- Arı, E., & Yıldız, A., 2017. Eşbütünleşme Analizi İle Genç İşsizliği Etkileyen Değişkenlerin Araştırılması, Alphanumeric Journal, 5 (2), 309-316.

Türkiye’de Çiğ Süt ve Yem Fiyatları Arasındaki Nedensellik İlişkisinin Analizi

Kaan KAPLAN^{1*}, Adnan ÇİÇEK²

¹Tokat Gaziosmanpaşa Üniversitesi, Ziraat Fakültesi, Tarım Ekonomisi, 60200, Tokat, Türkiye

²Tokat Gaziosmanpaşa Üniversitesi, Ziraat Fakültesi, Tarım Ekonomisi, 60200, Tokat, Türkiye

*Sunucu yazar e-posta: kaan.kaplan@gop.edu.tr

Özet

Bu çalışmanın amacı Türkiye’de çiğ süt fiyatlarına etki eden faktörlerden yem fiyatlarının etkilerini değerlendirmektir. Bu amaçla 1993-2020 yılları arasında Türkiye’de çiğ süt, süt yemi ve besi yemi fiyatları veri olarak kullanılmıştır. Fiyatların enflasyondan etkisinin arındırılması amacıyla cari fiyatlar reel fiyatlara çevrilmiş ve reel fiyatlar üzerinden analiz edilmiştir. Çalışmada çiğ süt fiyatı, süt yemi fiyatı ve besi yemi fiyatı arasındaki eşbütünleşme ve nedensellik ilişkisinin tespitinde kullanılan johansen kointegrasyon ve granger nedensellik modelleri için ön koşul olarak ADF birim kök testi uygulanmıştır. Araştırma sonuçlarına seriler birinci fark düzeyinde birim kök içermektedirler. Sonuç olarak seriler arasında uzun dönemde eşbütünleşik bir ilişki olduğu ve süt yemi fiyatı ile besi yemi fiyatı arasında çift yönlü bir nedensellik ilişkisinin olduğu, süt yemi fiyatı ve besi yemi fiyatlarının çiğ süt fiyatı ile aralarında bir nedensellik ilişkisinin saptanmadığı belirlenmiştir.

Anahtar kelimeler: Granger nedensellik analizi, Çiğ süt fiyatı, Süt yemi fiyatı, Besi yemi fiyatı

Gaziantep İli İslahiye İlçesinde Bulunan Tahtaköprü Barajında Seviye Yükseltilmesiyle Sular Altında Kalacak Parsellerin Tesbiti

Ömer VANLI^{1*}

¹İTÜ, Bilişim Enstitüsü, Coğrafi Bilgi Teknolojileri, 34469, İstanbul, Türkiye

*Sunucu yazar e-posta: vanli@itu.edu.tr

Özet

İnsanın doğaya müdahalesi, doğanın kendini yenileme kapasitesini aştığı için doğadaki bozulma bazen telafi edilemez durumlara sebebiyet verebilmektedir. Barajlar, genellikle bir akarsu vadisini kapatarak su biriktiren insan yapımı tesislerdir. Bu yapıların tarım alanlarının sulanması, hidroelektrik enerji üretilmesi, içme ve kullanma suyu olarak kullanmak gibi faydalarının yanında, yeraltı, yerüstü ve atmosfer alanlarındaki doğal döngünün bozulması, flora ve fauna başta olmak üzere doğadaki biyoçeşitlilik unsurlarına olan birçok olumsuz etkileri de kaçınılmaz olmuştur. Doğaya yapılan müdahalelerin etkilerinin izlenmesi ve tesbiti için son zamanlarda teknolojiye önemli gelişmeler sağlanmıştır. Bunlardan uzaktan algılama ile (UA), fiziksel temas olmadan çok geniş alanların kısa zamanda görüntülenerek bilgi elde edilebilirken; Coğrafi Bilgi Sistemleri (CBS) ile de bu tür verilerin veritabanı ortamında depolanması, işlenmesi, sorgulanması, analizi ve kullanıcıya sunulması işlemleri sonucu ilgili mekansal alan hakkında en uygun kararlar almamızı sağlarlar. Bu çalışmada, Gaziantep ili, İslahiye ilçesi sınırları içerisinde bulunan Tahtaköprü barajının, komşu bölgelerdeki su taşkınlarını önleme ve geniş ovaların sulanması gibi amaçlarla barajdaki su seviyesinin yükseltilmesi planlanırken, baraj çevresinde bulunan zengin bitki çeşitliliği ve tarımsal alanların sular altında kalması sonucu ekosisteme, bölge ekonomisine ve sosyal yapıya olumsuz etkileri irdelenecektir. Tüm bu etkileri değerlendirmek üzere, bölgenin uydu görüntüsü ve sayısal yükseklik Modeli (SYM) görüntülerinden, Erdas ve ArcGIS programlarını kullanarak, farklı analiz ile yapılacak ön işlemlerle su kotunun planlanan yüksekliğe getirilerek, ortaya çıkacak muhtemel durumda sular altında kalacak araziler sınıflandırılarak değerlendirilmiştir. Ayrıca sonuçlar ve alınacak önlemler ile ilgili öneriler sıralanmıştır.

Anahtar Kelimeler; Tahtaköprü Barajı, UA, CBS, SYM, Erdas, ArcGIS

Determination of the Parcels to be Inundated by Leveling Up the Tahtaköprü Dam in the İslahiye District of Gaziantep Province

Abstract

Since human intervention in nature exceeds the capacity of nature to renew itself, deterioration in nature can sometimes cause irreparable situations. Dams are man-made facilities that collect water, usually by closing a stream valley. In addition to the benefits of these structures such as irrigation of agricultural areas, production of hydroelectric energy, and use as drinking and utility water, the disruption of the natural cycle in underground, aboveground and atmospheric areas, many negative effects on biodiversity elements in nature, especially flora and fauna, have been inevitable. Significant advances have been made in technology recently to monitor and determine the effects of interventions to nature. With Remote Sensing (RS), it is possible to obtain information by viewing very large areas in a short

time without physical contact; With Geographic Information Systems (GIS), they enable us to make the most appropriate decisions about the relevant spatial area as a result of storing, processing, querying, analyzing and presenting such data in the database environment. In this study, while Tahtaköprü Dam, located in Gaziantep province, İslahiye district, is planned to increase the water level in the dam for purposes such as preventing floods in the neighboring regions and irrigation of large plains; The negative effects on the ecosystem, regional economy and social structure as a result of the rich plant diversity around the dam and the flooding of agricultural areas will be examined. In order to evaluate all these effects, using Erdas and ArcGIS programs from the satellite image and digital elevation Model (DEM) images of the region; The lands that will be inundated in the possible scenarios that will arise in case the water level is brought to the planned height with the preliminary processes to be made with different analyzes have been evaluated by classifying them. In addition, the results and suggestions for the measures to be taken are also discussed.

Key words; Tahtaköprü Dam, RS, GIS, DEM, Erdas, ArcGIS

GİRİŞ

Dünyada yeterli miktarda su potansiyeli bulunmasına karşın suyun bölgesel ve mevsimsel dağılımı çok büyük düzensizlikler göstermektedir. Dünyadaki ve ülkemizdeki nüfusun hızlı bir şekilde artmasına paralel olarak içme-kullanma suyu talebi, sanayi ve zirai faaliyetlerdeki gelişim suya olan ihtiyacın devamlı olarak artmasına yol açmaktadır.

Barajların, Tarım alanlarının zamanında ve yeterli olarak sulanmasını sağlaması, Hidroelektrik enerji üretmesi, İçme, kullanma ve endüstri için gerekli suyu düzenli ve sürekli sağlaması ve Yerleşim ve tarım alanlarını taşkınlardan koruması gibi faydaları bulunmaktadır. Bunlarla birlikte Barajlar, Doğal dengenin bozulması, Göl sahası içinde kalan tarım alanlarının ve doğal kaynakların kullanılamaması, Göl sahası içinde kalan yerleşim merkezlerinin nakli, Göl sahası içinde kalan yol, köprü vb. yatırımların yerine yeni yatırım yapma ihtiyacı, Yeraltı su seviyesinin yükseltilmesi dolayısı ile meydana gelen olumsuz etkiler, Su yükünün artması dolayısı ile meydana gelebilecek tehlikeli heyelanlar ve diğer jeolojik olaylar, Artan su buharlaşması dolayısı ile kullanılabilir su miktarının azaltılması, Akarsuların taşkın mevsimlerinde birlikte getirdikleri toprak gücünü artıran besleyicilerden bilhassa delta ovalarının mahrum kalması, Suyun içinde taşınan maddelerin azalması nedeni ile baraj mansabında daha fazla yatak oyulması, Kıyı erozyonunun artması, Göl alanında kalan tarihi eserler, gibi konulara da dikkat edilmesi gerekmektedir. (1)

Uzaktan algılama, cisimlere fiziksel temasta bulunmadan cisimlerin yansıttığı elektromanyetik enerjinin ölçülmesi yoluyla onlar hakkında bilgi elde etme yöntemi olarak tanımlanmaktadır. Bu teknik, cisimler tarafından yansıtılan elektromanyetik enerjinin kaynağının algılayıcı platformun kendisi olması durumunda aktif uzaktan algılama, kaynağın Güneş gibi doğal bir enerji kaynağı olması durumunda ise pasif uzaktan algılama olarak adlandırılır. Yeryüzü hakkında uydularca toplanan ve yer istasyonlarına aktarılan bilgi, yeryüzü parçasına ait görüntülerdir. Alıcılar yeryüzü parçasının elektromanyetik enerjiyi yansıtma oranını ölçerler. (2)

Dünyada daha önce yapılan bir kısım çalışmalardan birinde, çok-zamanlı Landsat verileri 18 yıllık çalışma döneminde Alqueva Bölgesinde arazi kullanım değişiklikleri ölçmek için kullanıldı. Arazi örtüsü haritaları üretmek ve arazi kullanımı değişiklikleri değerlendirmek için benimsenen metodoloji (piksel tabanlı sınıflandırma yoluyla nesne tabanlı sınıflandırma) diğer çalışma alanlarına adapte edildi. Burdan da Piksel tabanlı yaklaşım (kontrollü sınıflandırma) ve nesne yönelimli yaklaşım (kontrolsüz sınıflandırma) ile ilgili benzer sonuçlar üretilmiştir. Ancak, nesne tabanlı yaklaşım görüntülerin önemli bir bölümü sınıflandırma vermedi. Bu nedenle, piksel tabanlı yaklaşım arazi örtüsü analizinde kabul edildi. Değişim tespiti istatistik ve CBS analizlerine dayanarak, iki arazi kullanımı dönüşümleri

planlandı. Gelecekte, bu çalışma arazi örtüsü sınıflarının belirlenmesinde doğruluğunu artırmak amacıyla saha araştırmalarına devam edilmesi önerildi(4).

Yapılan başka bir çalışmada Ayamama Deresi için taşkın risk analizi çalışmaları uzaktan algılama ve CBS ile gerçekleştirilmiştir. TRA ile yapılan risk analizleri sonucunda, Ayamama Deresi için AHP yöntemi ile olası risk bölgeleri belirlenmiştir. Oluşturulan risk haritası değerlendirildiğinde “Çok Yüksek” taşkın riskine sahip alanların yerleşim bölgesi üzerinde kaldığı ve bu alanlar olası bir taşkında zarar görme olasılığı en yüksek yerler olarak belirlenmiştir. (5)

Bir başka çalışmada ise, Burdur gölünün kıyı kenar çizgisi değişimleri farklı yıllara ait çok zamanlı uydu görüntüleri üzerinde uzaktan algılama yöntemleri kullanılarak araştırılmıştır. Burdur gölünün kıyı kenar çizgileri ve üç boyutlu göl topoğrafyası yardımıyla, 1975 yılından 2002 yılına kadar farklı yıllarda göl alanı ve hacim değişimleri belirlenmiştir. 1975 yılında 210 km² olan göl alanı, 2002 yılında 153 km²'ye kadar düşmüştür. Göl hacminde de 1.68 km³'lük azalma meydana gelmiştir. 27 yıllık dönem içerisinde, göl seviyesi yaklaşık 10 m düşmüştür. Bu dönemde göl alanı ve hacmi yaklaşık % 27 oranında azalmıştır.(6)

Son örnek çalışmada ise, Porsuk Nehri su toplama alanının belirlenmesi, Bahçelik Barajı drenaj alanına düşen ortalama yıllık toplam yağış miktarının hesaplanması, maksimum su kotluna göre aktif hacmin belirlenmesi, rezervuara ait kot-alan-hacim grafiği ve drenaj alanı kot-alan grafiği için veri hazırlanması ArcGIS 10 kullanılarak gerçekleştirilmiştir. İlk bölümde zamana ve mekâna ait verilerin yönetimi gerçekleştiren CBS'nin önemli bir bileşeni olan ArcHydro kullanılarak su toplama havzalarının belirlenmesine çalışılmıştır. Hidrolojik verileri içeren zaman serilerinin yönetiminde büyük kolaylık sağlayan ArcHydro ile zamana bağlı veriler mekânsal özelliklerle ilişkilendirilerek yeryüzünde suyun hareketini belirlemek mümkün olabilmektedir. (7)

Gaziantep İslahiye-Hatay Amik bölgesinde bulunan 32.000 ha'lık alanın sulanması ve öngörülemeyen su taşkınlarının önlenmesi amacıyla Klavuzlu barajından Tahtaköprü barajına su taşınması ve bu yolla baraj gölü seviyesinin 12 m yükseltilmesi amaçlanmaktadır. Bu nedenle su altında kalacak tarımsal arazilerin belirlenmesi neticesinde arazi sahiplerine gerekli rayiç bedelleri üzerinden ödemelerin yapılması sağlanmış olunacaktır. Yapılması planlanan bu aşamalar ışığında problemi şu şekilde tanımlamak mümkün olabilecektir. Gaziantep ili, İslahiye ilçesi'nde bulunan Tahtaköprü barajında 12 m'lik su seviyesi yüksekliği gerçekleştiğinde hangi tarımsal araziler sular altında kalacaktır? Su seviyesi yükselmesi sonucu sular altında kalacak köyler var mıdır? Biber bitkisinin tarımsal üretiminde önemli bulunan bölgenin Sosyoekonomik durumunu etkileyecek ciddi oluşumlar yaşanacak mıdır?

MATERYAL VE YÖNTEM

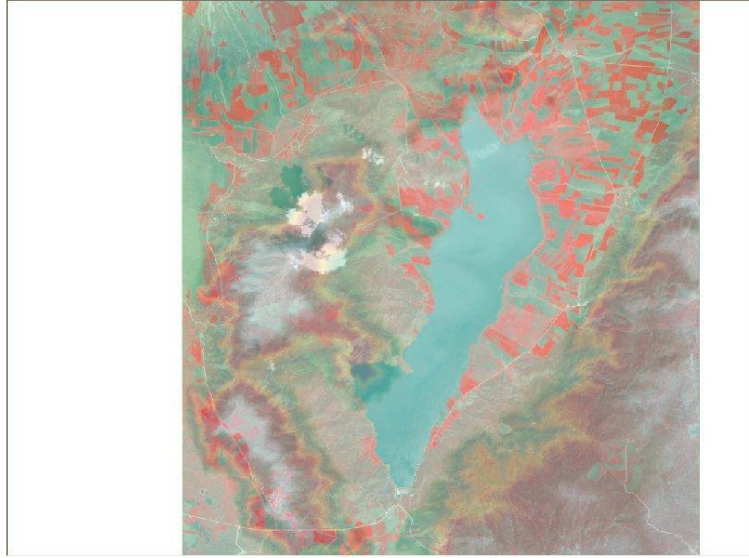
Materyal

İslahiye ilçe merkezinin güneydoğusunda bulunan Tahtaköprü barajı genel amaçları dışında bölgedeki kırsal alanda yaşayanların balıkçılık ve mesire alanları olarak kullanılmasını da sağlamaktadır. Aşağıdaki şekil'de barajın teknik özellikleri gösterilmektedir.

Barajın Yeri	Hatay
Akarsuyu	Karasu
Amacı	S+T
İnşaatın (başlama-bitiş) yılı	1968 - 1976
Gövde dolgu tipi	Zonlu
Gövde hacmi	1,95 hm ³
Yükseklik (talvegden)	33,50 m
Normal su kotunda göl hacmi	200 hm ³
Normal su kotunda göl alanı	23,40. km ²
Sulama alanı	11 900 ha
Güç	- MW
Yıllık Üretim	- GWh

Şekil 1. Tahtaköprü barajının teknik özellikleri

İlgili bölgenin uzaktan algılama görüntüleri olarak ise, Tahtaköprü barajı ve çevresinin uydu görüntüsünün (SPOT5) İTÜ UHÜZAM dan temin edilmesi; Spot 5 görüntüsünün, geometrik düzeltilmesinin yapılmış olduğundan tekrar yapılmamıştır. Aşağıdaki şekilde bölgenin uzaktan algılama görüntüsü bulunmaktadır.



Şekil 2. Tahtaköprü barajı SPOT5 Uydu Görüntüsü

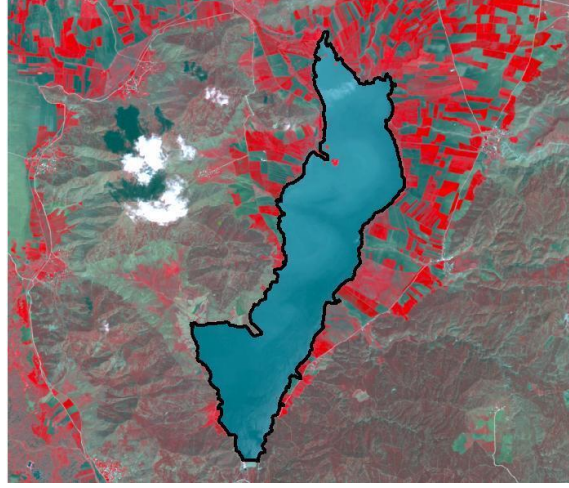
Görüntünün Teknik özellikleri ise aşağıda belirtilmektedir.

Çizelge 1. Spot 5 uzaktan algılama görüntüsünün teknik özellikleri

Algılama tarihi	Uydu	Mekansal Çözünürlük	Radyometrik Çözünürlük	Spektral Çözünürlük
13.04.2010	Spot 5	10 m	8 bit	4 bant

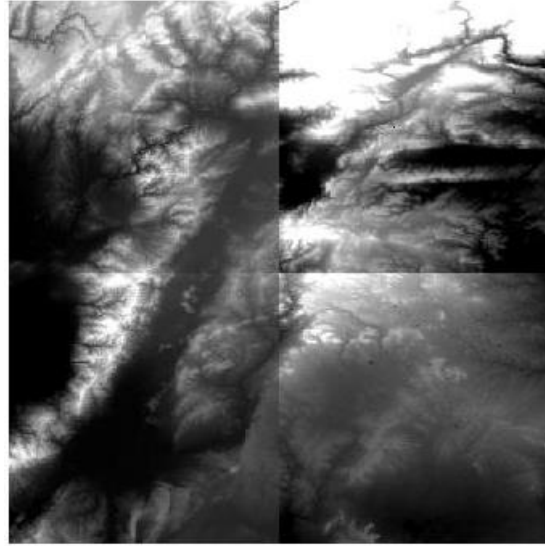
BULGULAR

Öncelikle uydu görüntüsü üzerinde ERDAS programında baraj gölü sınırlarının editleme yapılarak sınırları çizilmiştir.



Şekil 3. Sınırları belirlenmiş uydu görüntüsü

Ardından Sayısal Yükseklik Modeli oluşturmak için, ASTER-GDEM'den Tahtaköprü barajı gölü ve çevresinin 4 parça görüntü olarak alınarak mozaikleme işlemine başlanmıştır.



Şekil 4. Aster Gdem den alınan 4 parça SYM

ASTER-GDEM sitesinden alınan (ASTGTM2_N36E036, N36E037, N37E036, N37E037) parçalı Sayısal Yükseklik Modeli görüntüleri ERDAS programında birleştirilerek Mozaik görüntü elde edilmiştir.



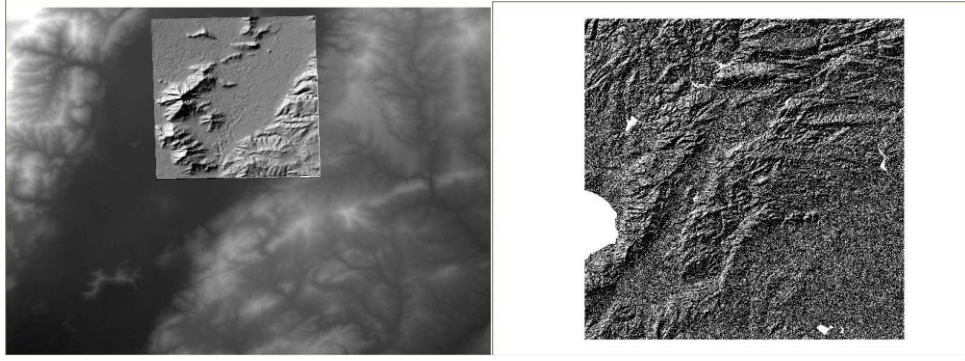
Şekil 5. Bölgenin mozaik görüntüsünün oluşturulması

Baraj gölünün iki farklı yıldaki özellikleri de aşağıdaki tabloda gösterilmiştir.

Çizelge 2. DSI'den alınan 13 Nisan 2010 ve Güncel (14 Aralık 2012) su kotları

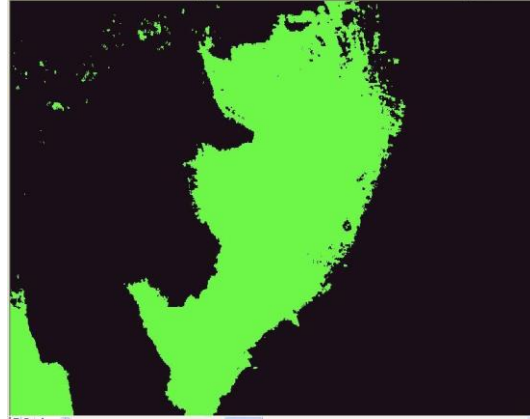
Baraj Alan özellikleri	2010	2012
Göl Kotu (m)	404.41	398.85
Göl Hacmi (hm ³)	164.111	75.866
Göl Alanı (km ²)	21.139	11.121

Tahtaköprü barajının uydu görüntülerinden ArcGIS programından yapılan yüzey analizleri ile Hillshade görüntüleri de elde edilmiştir.



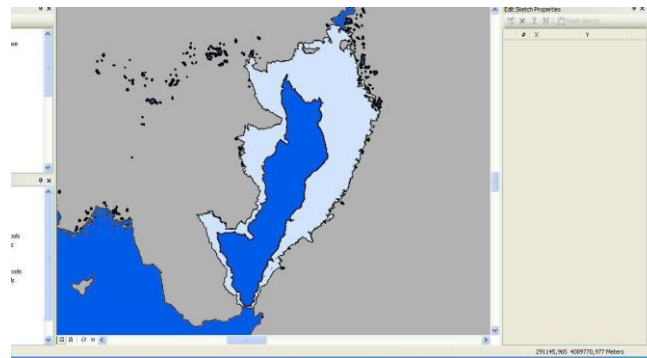
Şekil 6. Hillshade görüntüleri

ArcGIS programında toolbox'lar içinde raster calculator hesaplaması için 405 m olan göl seviyesinin 12 m yükseltilerek 417 m den büyük-eşit yapılarak yeni göl sınırlarımız için çeşitli analizler de yapılmıştır.



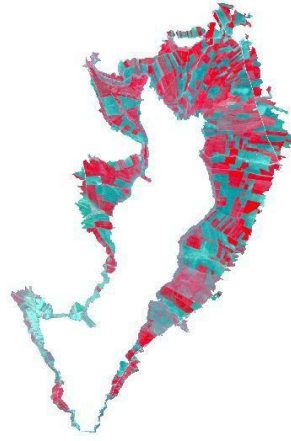
Şekil 7. Raster calculator sonrası göl görüntüsü

Daha sonra yine ArcGIS programında “export output” işlemi yani select by attribute yapılarak “0” ları çıkartıp, 1 olanları (417 m) göstermesi istendi.



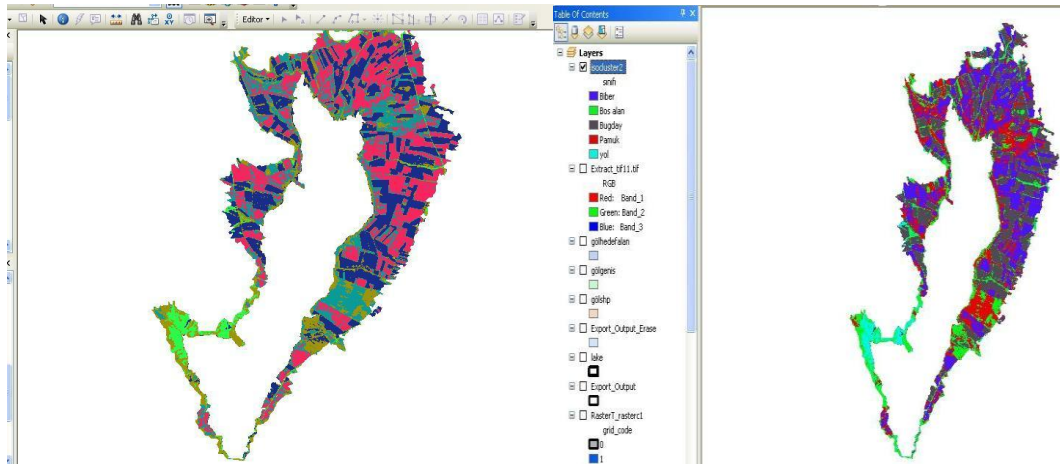
Şekil 8. Export Output sonrası görüntü

Extract işlemi olarak da çalışma yapılacak alanlar diğer yerlerden ayırma işlemi uygulandı.



Şekil 9. Extract işlemi ile ayrılan çalışma hedef alanı

ArcGIS programında, spatial analyst / Multivariate / isocluster Unsupervised Classification işlemi yapıldı. Bunun sonucunda da 5 adet sınıf üretildi ve alan büyüklüklerine göre sınıflandırma yapıldı.



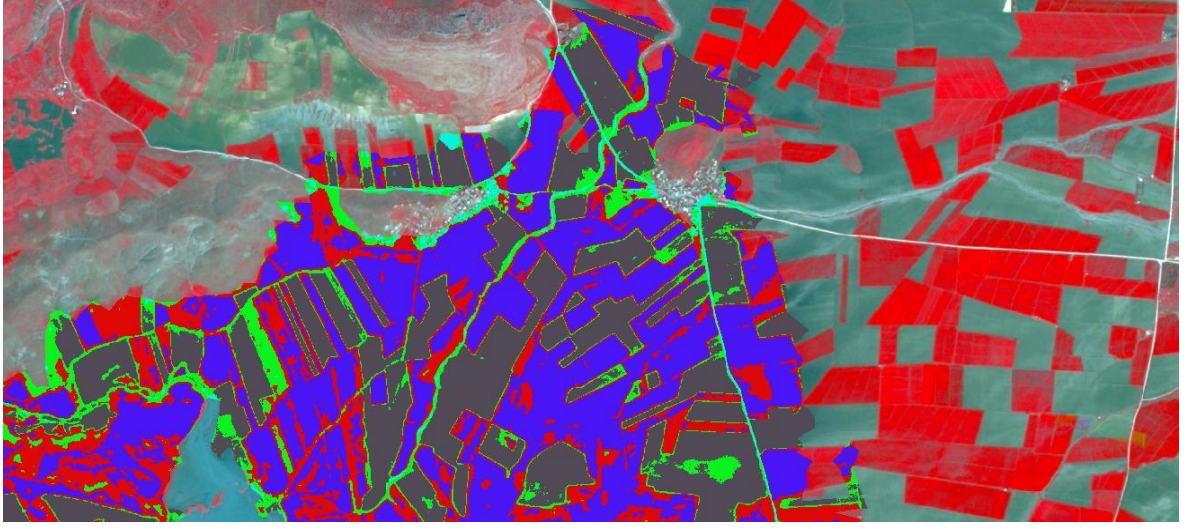
Şekil 10. Arazideki bitki sınıfı grupları

Arazi gruplarından 5 adet bitki türü üretimleri alanları ile birlikte tesbit edilerek sınıflandırma tablosu oluşturuldu.

Table				
isocluster2				
OBJECTID	Value	Count	sınıfı	
3	3	386698	Biber	
1	1	376584	Bugday	
4	4	288039	Pamuk	
2	2	208581	Bos alan	
5	5	55710	yol	

Şekil 11. Bitki türlerinin alanlara göre dağılımı

Arazi ve köylerin son durumu yönüyle ise yukarıdaki yapılan tüm işlemlerden sonra Şahmaran ve Akınyolu köylerinin sınırlarına kadar göl sularının gelebileceği görüldü.



Şekil 12. Göl yeni sınırları ve köylerin genel görünümü

TARTIŞMA VE SONUÇ

Yapılan modellememizle ortaya çıkan bu katmanlarda su seviyesinin 12 m arttırılması ile hangi arazilerin ve köy yerleşim bölgelerinin sular altında kalabileceklerinin tespiti yapılmaya çalışılmıştır. Yapılan bu bilimsel ve teknolojik ön çalışmalarla arazi ve köylerin öngörülebilir durumu daha da netleşmiş olmakla, zaman ve mali kaynaklardan önemli ölçüde tasarruf sağlanmış olacaktır.

DSİ tarafından uygulanacak olan bu proje ile amaç, ilgili bölgedeki 32 bin hektar alanın sulanması ve su taşkınlarının önüne geçilmesidir. Fakat bununla birlikte su altında kalacak arazilerin ekolojik habitatındaki değişim ile, ÇED (Çevre Etki Değerlendirmesi) hesaplanarak gerekli önlemlerin minimuma indirilmesinin sağlanması ve bölge halkının sosyoekonomik durumlarına karşı da refah seviyesinin ve genel ekonomik seviyenin düşürülmemesi için önlemler alınması gerekmektedir.

Biber bitkisi, İslahiye ovasının ana ürünü olarak bölgenin %60'ının geçim kaynağını oluşturmaktadır. Bölge, tüm ülkenin ise Kırmızıbiber üretiminin %35 ini karşılamaktadır. 387 bin hektar alandan kaybedilecek alan başta bölge ekonomisi olmak üzere kırmızıbiber sektöründe de önemli bir daralmaya neden olacaktır.

Ayrıca özellikle verimli arazilerin gerekirse kurtarılması için alternatif göl kotu yönlendirilmesi de yapılabilir. Böylece atıl durumdaki araziler veya 4.sınıf tarım arazilerinin kullanılmasıyla ekosistem yönünden olmasa da ekonomik açıdan zarar indirgenebilir. Çünkü değerli arazilerin getirisi ile göl alanının sulama ile ortaya çıkacak getirisi konusunda çok doğru tercihlerde bulunması gerekmektedir.

Kaynaklar

Helvacı Deniz, DİM Barajının Dinamik Analizi, Süleyman Demirel Üniversitesi, Fen Bilimleri Enstitüsü, Yüksek Lisans Tezi Isparta 2009.

Kandemir Egemen, Uzaktan algılama tekniğinde NDVI Değerleri ile doğal bitki örtüsü tür Dağılımı arasındaki ilişkilerin Belirlenmesi üzerine araştırmalar, Yüksek Lisans Tezi, Ege Üniversitesi, Fen bilimleri Enstitüsü.

Dam-break flood modeling for Tangjiashan Quake Lake Juan Dai, Shihui Zhang, Chongsheng Xue, Faculty of Geoscience, China University of Geosciences, Wuhan, China Dept. of Physics.

Land use/land cover changes and flooding surface estimation in Alqueva (Portugal) using 18 years of Landsat data, Ana Teodoro, Sofia Riosb, Dário Ferreira, Geo-Space Sciences Research Centre,

Faculty of Science, University of Porto, Rua do Campo Alegre, 4169-007 Porto, Portekiz.
Çok kriterli karar verme ve bilgi difüzyonu yöntemleri İle taşkın risk analizi Aybike Saral 1 , Nebiye Musaoğlu 2 İTÜ, İstanbul Teknik Üniversitesi, Bilişim Enstitüsü ,Uydu Haberleşmesi ve Uzaktan Algılama Bölümü, Maslak, İstanbul, saralaybik@itu.edu.tr İTÜ, İstanbul Teknik Üniversitesi, Geomatik Mühendisliği Bölümü, Maslak, İstanbul, musaoglune@itu.edu.tr
Burdur Gölü Seviye Değişimlerinin Çok Zamanlı Uydu Görüntüleri İle İzlenmesi Erhan Şener, Ayşen Davraz, Tefik İsmailov
Coğrafi bilgi sistemleri ile hidroloji uygulamaları T.C. orman ve su işleri bakanlığı, DSİ genel müdürlüğü teknoloji dairesi başkanlığı.
Çokal Barajı (Çanakkale) çökme modeli ve taşkın risk analizi, Hasan Özdemir, Cengiz Akbulak, Hasan Özcan.

Sabit Kameralı Derin Öğrenme Fiziksel Şiddet Tespiti Sistemi için Aşırı Hareket Öncül Fonksiyonu

Muhammet Fatih POLAT^{1*}, Aylin ALIN²

¹Dokuz Eylül Üniversitesi, Fen Bilimleri Enstitüsü, 35390, İzmir, Türkiye

²Dokuz Eylül Üniversitesi, Fen Fakültesi, İstatistik, 35390, İzmir, Türkiye

*Sunucu yazar e-posta: 2019900077@ogr.deu.edu.tr

Özet

Son yıllarda video görüntülerindeki fiziksel şiddetin tespitine yönelik farklı çözüm önerileri sunulmuştur. Bu çalışmalar arasında en yüksek doğruluk oranına sahip yöntemlerin başında derin sinir ağı kombinasyonları gelmektedir. Yüksek öğrenme kapasitesine sahip bu yapay sinir ağlarının işlem hızı donanım türüne, derinliğe ve algoritmaya bağlı olarak görece yavaş ve yapay sinir ağının hesaplama için ihtiyaç duyduğu enerji miktarı da görece fazla olabilmektedir. Böyle bir sistemin sürekli olarak aktif tutulmasının getirebileceği negatif etkiler bulunmaktadır. Bu çalışmada bu etkileri azaltmak amacı ile video görüntülerindeki piksel lokasyonları üzerinde “zamana bağlı mod değişkenliği” olarak tanımladığımız bir mod merkezli varyans işlemi gerçekleştirilerek video içerisindeki aşırı hareketi bulan bir yöntem önerisinde bulunulmakta ve önerilen yöntemin performansı, UCF-Crime veri setinden sabit kamera ile elde edilmiş farklı video görüntüleri üzerinde incelenmektedir.

Anahtar Kelimeler: Aşırı hareket tespiti, Derin sinir ağları, Fiziksel şiddet tespiti

Extreme Motion Antecedent Function for Deep Learning Physical Violence Detection System with Fixed Camera

Abstract

In last decades, different solutions were proposed for physical violence in videos. Heads of the best methods are combinations of deep neural networks. Deep neural network combinations are one of the most accurate methods in these studies. The processing speed of these artificial neural networks with high learning capacity could be relatively slow, depending on the type of hardware, depth and algorithm, and the amount of energy needed by the artificial neural network for calculation could be relatively high. There are negative effects that could be brought about by keeping that kind of system active continuously. In this study, in order to reduce these effects, a mode-centered variance process, which we define as "time-dependent mode variation" is performed on the pixel locations in videos, is proposed to find excessive motion in the video. The performance of the proposed method is examined on different video images obtained with a fixed camera from the UCF-Crime data set.

Key words: Excessive motion detection, Deep neural networks, Physical violence detection

GİRİŞ

Çağlar boyunca şiddet ögesi içeren eylemler insanlık tarihinin ayrılmaz bir parçası ve aynı zamanda sorunu olmuştur, olmaya da devam etmektedir. Bu eylemlere, ülkelerin anayasaları ile suç olarak kabul edilmesine rağmen toplumda sık rastlanılmaktadır.

Günümüzde gelişen teknoloji ile birlikte birçok problem, otonom hale getirilmiş araçlarla çözülebilmektedir. Bu gelişmelerden bir tanesi de günümüzde kullanımı ve kapasitesi arttırılmış yapay sinir ağlarıdır. Yapay sinir ağları ile birlikte karmaşık işlemler, yüksek doğruluk oranları ile rahatlıkla başarılabilmektedir. Bu karmaşık işlemler arasında resimlerin ve/veya resimlerdeki nesnelerin sınıflandırılması ve işaretlenmesi de yer almaktadır. Bunlar gibi karmaşık işlemlerin başarılmasından sonra son yıllarda yapılan çalışmalar ile birlikte karma yapay sinir ağı sistemlerinin video görüntülerinin sınıflandırılmasında kullanılabileceği görülmüştür.

Video görüntülerinin sınıflandırılması ile birlikte bu çalışmalar daha da özelleştirilerek video görüntüleri içerisindeki şiddeti tespit edebilen yapay sinir ağı sistemleri geliştirilmiştir. Bunlar arasında evrişimli sinir ağları ile uzun-kısa süreli hafıza (*Long-short term memory*) kombinasyonlarını içeren (Hanson vd. 2019) ya da üç boyutlu evrişimli sinir ağlarından oluşan sistemler bulunmaktadır (Li vd. 2019).

Bu ve birçok alanda kullanılabilen yapay sinir ağları, donanıma, ağ derinliğine, ağ genişliğine ve kullanılan algoritmaya bağlı olarak yüksek enerji ihtiyacına sahip olabilmektedir.

Video görüntülerindeki şiddeti tespit etmesi amacıyla kurulan bir yapay sinir ağı sisteminin sürekli olarak aktif tutulması kurulan donanım ile birlikte akış halindeki görüntülerle senkronizasyonu sağlama noktasında zorlanması muhtemeldir. Bunun önüne geçebilmek amacıyla bu sistemler için hesaplaması farklı algoritmalarla hızlandırılabilen bir yöntem önerisinde bulunmaktayız. Bu yöntemi “zamana bağlı mod değişkenliği” olarak tanımlamaktayız.

Sonraki bölümlerde yöntem, bulgular, tartışma ve sonuç kısımları bulunmaktadır.

YÖNTEM

Mod, bir sayı dizisinde en fazla tekrar eden sayı anlamına gelmektedir. Mod işleminin, uzun süreli bir video klibi üzerine, zamana bağlı olarak uygulanması ile arka plana ait görüntü matrisi bulunabilmektedir.



Şekil 1. Zamana bağlı mod ile arka plan görüntüsünün bulunması (Zheng vd., 2006)

Yöntemimiz kısa süreli video kliplerinde en çok tekrar eden değeri, o klibin arka planı kabul ederek o lokasyonda arka plan değerinden ne kadar değişkenlik gösterdiğini bulmaya dayanmaktadır. Kavga gibi eylemlerin hızlı yer değiştirme içermesi bu yöntemin dayanak noktasıdır. Önerdiğimiz yöntemi, diğer hareket tespiti yöntemlerinden ayıran temel fark fonksiyonumuzun tek yönlü ve az değişim gösteren eylemlerin (yürüme eylemi gibi) etkilerini filtrelemesidir. Kare farkı, hareketli kenar, arka plan çıkartımı gibi algoritmalar ise doğrudan hareketi tespit etmek için oluşturulmuştur (Zhan vd. 2007).

Biz, bu çalışma kapsamında zamana bağlı dizilerdeki modu, kısa süreli video kliplerinde (30 – 60 – 90 kare) bulmaktayız. Bulunan mod matrisindeki değerler, her x - y piksel lokasyonu için referans değer kabul edilmektedir. Referans değerler, zamana bağlı olarak alınan piksel dizilerinde varyans hesaplamasında ortalama yerine kullanılmaktadır. Bu işlem sonucunda da x – y piksel lokasyonları için aşırı hareketlilik matrisi oluşturulmaktadır.

M, 3 boyutlu video matrisi olarak kabul edilirse (M_{mnz});

$$1. \quad M_z = \begin{bmatrix} x_{11} & \cdots & x_{m1} \\ \vdots & \ddots & \vdots \\ x_{1n} & \cdots & x_{mn} \end{bmatrix}$$

$$2. \quad mod(M_{zaman}) = Y$$

$$3. \quad M_{aşırıhareket} = \frac{\sum (M_{mnz} - Y)^2}{N}$$

$$4. \quad M_{aşırıhareket} = \begin{bmatrix} x_{11} & \cdots & x_{m1} \\ \vdots & \ddots & \vdots \\ x_{1n} & \cdots & x_{mn} \end{bmatrix}$$

M_{mnz} : 3 boyutlu video matrisi

m, n, z : Matris boyutları

Y : Mod matrisi

$M_{aşırıhareket}$: Aşırı hareket matrisi

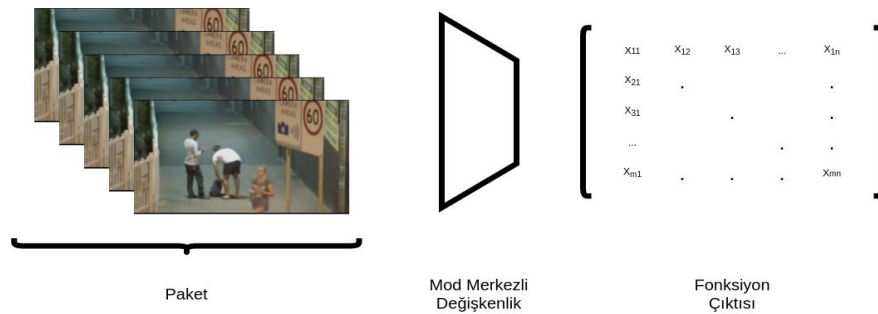
Aşırı hareketlilik fonksiyonunun farklı eşik değerleri ya da farklı karar mekanizmalarıyla yapay sinir ağı sistemlerine öncül fonksiyon olarak entegresi sağlanması mümkün olacaktır.

Önerdiğimiz fonksiyona ait algoritma aşağıdaki gibidir;

Algoritma

Boş bir "paket" 3 boyutlu dizi oluşturun.
Video kareleri döngü başlangıcı:
Video karesini gri ölçeğe çevir.
Video karesini pakete ekle.
Eğer paketteki kare sayısı belirlenen sayıya ulaştı ise:
Piksel lokasyonları için mod bul.
Zaman ekseninde mod ile piksel farklarını al.
Karelerini topla.
Kare sayısına böl.

Algoritmadaki fonksiyonun dönüşüm işlemi Şekil 2 ile görselleştirilmiştir.

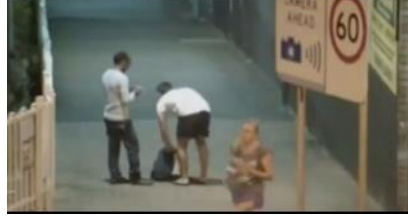


Şekil 2. Fonksiyon işleyişi

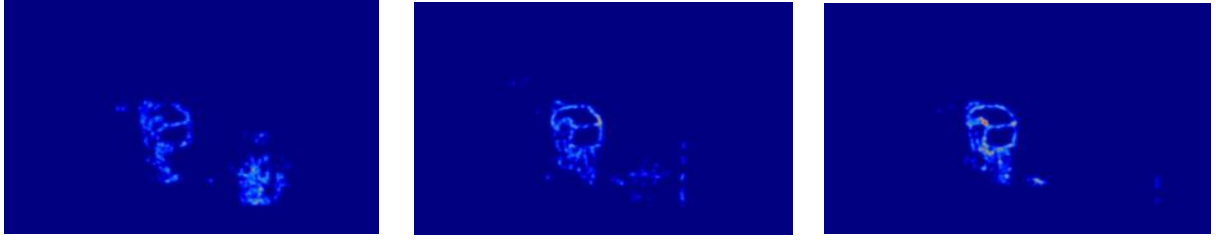
Bu algoritmadaki mod işlemi farklı şekillerde optimize edilerek fonksiyon daha hızlı hale getirilebilir.

BULGULAR

Algoritma, UCF-Crime veri setinden sabit kamera ile kayıt altına alınmış video görüntüleri üzerinde test edilmiştir(Sultani vd. 2018). Bu testlerde paketler 30-60-90 kare uzunluğunda alınmıştır.



Şekil 3. Orjinal kare

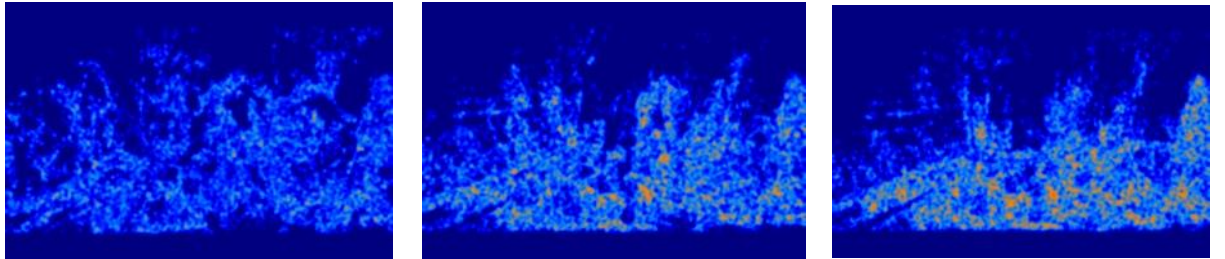


Şekil 4. Mod merkezli değişkenlik çıktıları (Soldan sağa 30-60-90 karelik paketler)

Bu video görüntüsünde 30 ve 60 kare içeren paketlerde fonksiyon çıktısında sadece tek yönde hareket eden bireye ait pikseller aşırı hareketli olarak görülmektedir. 90 kare içeren pakette ise tek yönde sıradan bir eylem gerçekleştiren bireye ait aşırı hareket olarak algılanan piksel miktarında azalma meydana gelmiştir.



Şekil 5. Orjinal kare

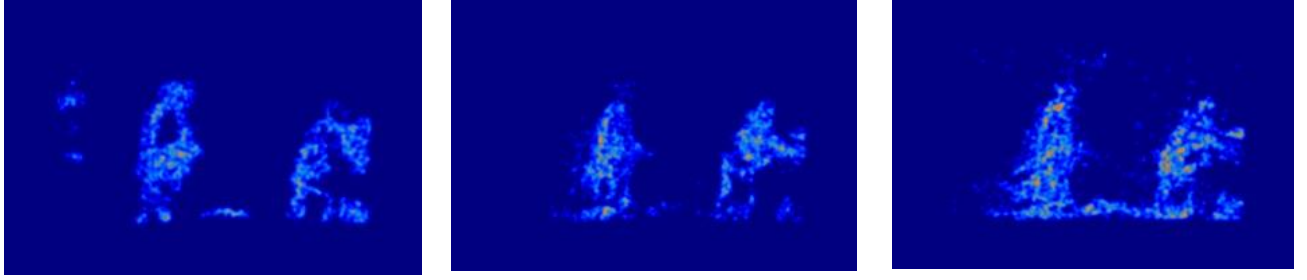


Şekil 6. Mod merkezli değişkenlik çıktıları (Soldan sağa 30-60-90 karelik paketler)

Bu video görüntüsünde aşırı hareketli eylemlerde bulunan bireyler sahnenin geneline yayıldığı görülmektedir. Paket büyüklüğünün artışıyla beraber aşırı hareket olarak belirtilen pikseller daha belirgin hale gelmiştir.



Şekil 7. Orjinal kare

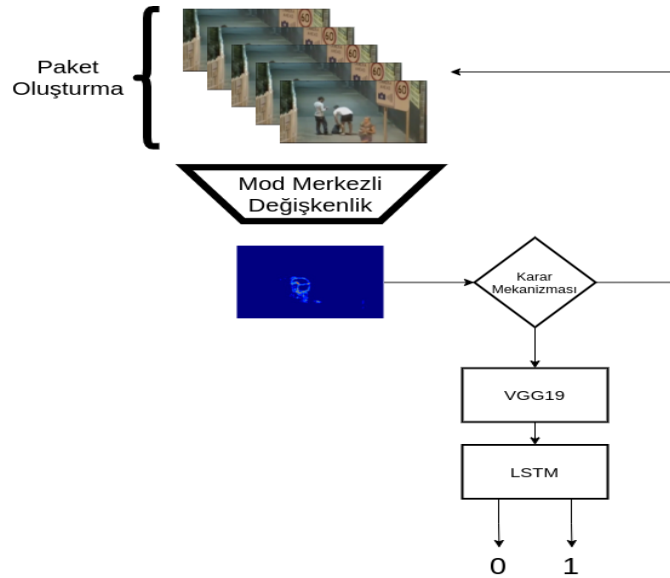


Şekil 8. Mod merkezli değişkenlik çıktıları (Soldan sağa 30-60-90 karelik paketler)

Bu video görüntüsünde ise 30 karelik pakete ait çıktının sol tarafındaki aşırı hareket olarak belirtilen pikseller dışında çok büyük farklılıklar görülmemiştir.

TARTIŞMA VE SONUÇ

Yapmış olduğumuz çalışmada, algoritmamızın uygulandığı videoların birçoğunda amaca uygun sonuçlar ürettiği gözlemlenmiştir. Bununla beraber bu algoritmanın sahne yapısına göre performansında değişiklik yaşanabileceğini düşünmekteyiz. Algoritmanın kullanımı için sahne yapısına uygun belirleyici ve sınıflandırıcı fonksiyonların olması gerekmektedir. Örnek bir akış şeması aşağıda gösterilmiştir (Şekil 9). Daha sonraki çalışmamızda bu fonksiyonlar da algoritmaya entegre edilerek algoritma geliştirilecektir.



Şekil 9. Örnek uygulama şeması

Kaynaklar

- Sultani, W., Chen, C., & Shah, M. (2018). Real-world anomaly detection in surveillance videos. In *Proceedings of the IEEE conference on computer vision and pattern recognition* (pp. 6479-6488).
- Hanson, A., Pnvr, K., Krishnagopal, S., & Davis, L. (2019). Bidirectional convolutional LSTM for the detection of violence in videos. In *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)* (Vol. 11130 LNCS, pp. 280–295). Springer Verlag. https://doi.org/10.1007/978-3-030-11012-3_24
- Li, J., Jiang, X., Sun, T., & Xu, K. (2019). Efficient violence detection using 3D convolutional neural networks. In *2019 16th IEEE International Conference on Advanced Video and Signal Based Surveillance, AVSS 2019*. Institute of Electrical and Electronics Engineers Inc. <https://doi.org/10.1109/AVSS.2019.8909883>
- Zheng, J., Wang, Y., Nihan, N. L., & Hallenbeck, M. E. (2006). Extracting Roadway Background Image: Mode-Based Approach. *Transportation Research Record*, 1944(1), 82–88. <https://doi.org/10.1177/0361198106194400111>
- Zhan, C., Duan, X., Xu, S., Song, Z., & Luo, M. (2007, August). An improved moving object detection algorithm based on frame difference and edge detection. In *Fourth international conference on image and graphics (ICIG 2007)* (pp. 519-523). IEEE.

İlişki Ölçüleri Testleri

Murat SÖZEYATARLAR¹, Mustafa ŞAHİN², Esra YAVUZ^{3*}

^{1,3}Kahramanmaraş Sütçü İmam Üniversitesi, Ziraat Fakültesi, Zootehni Bölümü, 46050,
Kahramanmaraş, Türkiye

²Kahramanmaraş Sütçü İmam Üniversitesi, Ziraat Fakültesi, Tarımsal Biyoteknoloji Bölümü, 46050,
Kahramanmaraş, Türkiye

*Sunucu yazar e-posta: yavuz7346@gmail.com

Özet

Çalışmada, araştırmacıların yöntem ve çalışmalarına başlarken verilerin hangi ölçekle ele alınacağı ve bu ölçek doğrultusunda verilere hangi istatistiksel ilişki ölçülerinin uygulanacağı hakkında bilgi verilmektedir. Bu amaç doğrultusunda istatistiksel ilişki ölçüleri testlerinden oluşan toplam 18 test incelemeye alınmıştır. Ayrıca Kahramanmaraş' ta eğitim veren üç farklı lise türünde 2148 öğrenciden alınan veri seti üzerinde istatistiksel ilişki ölçüleri uygulamalı olarak incelenmiştir. Bu amaçla Kahramanmaraş' ta Fen lisesi, Anadolu lisesi ve Meslek lisesinde okuyan öğrencilerden alınan bilgiler kullanılmıştır. Çalışmada, cinsiyet, babanın eğitim durumu, özel oda durumu, kardeş sayısı, bilgisayar durumu, yaşadığı yer, özel ders alma durumu, öğrencinin yaşı, öğrencinin vücut ağırlığı ve ailenin gelir durumu gibi faktörlerin lise türüne etkisi istatistiksel ilişki ölçüleri ile değerlendirilmiştir. Aralıklı veya oran ölçekli veri setlerine parametrik testler, sınıflayıcı veya sıralayıcı ölçekli veri setlerine ise parametrik olmayan testler uygulanmıştır. Sonuç olarak bu tez çalışmasında ilişki ölçüleri testleri ile öğrencinin lise türünü etkileyebileceği düşünülen faktörlerin, etki derecesi ve yönünün ne olduğu belirlenmeye çalışılmıştır.

Anahtar kelimeler: İstatistik, İlişki Ölçüleri, Testler

Taşkın Frekans Analizinde Kullanılan Yöntemler

Ayşe DOĞANÜLKER^{1*}, Alper Serdar ANLI¹

¹Ankara Üniversitesi, Ziraat Fakültesi, Tarımsal Yapılar ve Sulama Bölümü, 06110, Ankara, Türkiye

*Sunucu yazar e-posta: doganulkerayse@gmail.com

Özet

Taşkınlar doğal gerçekleşen bir olay olmasının yanı sıra toplum yaşamı ve ekonomi üzerinde fazla olumsuz etki yaratan bir afettir. Ülkemizde depremlerden sonra en büyük ekonomik kayıplara neden olan yaşanmış taşkınlara bakıldığında 1975-2010 yılları arasında 695 taşkın olayı yaşanmış ve 634 kişi bu olaylarda yaşamını yitirmiştir. Ülkemizde plansız şehirleşme, akarsu yataklarındaki düzensiz yapılaşmalar, iklim değişikliği gibi nedenlerden dolayı taşkın olayı artmıştır. Bu nedenle hidrolik yapıların doğru planlanması ve tasarlanmasıyla birlikte taşkınların büyüklüklerinin ve meydana gelme frekanslarının güvenilir biçimde tahmin edilmesi büyük önem taşımaktadır. Bu çalışmada taşkın frekans analizi için kullanılan yöntemler incelenmiş ve bu yöntemlerle taşkınların verdiği zararı en aza indirmenin mümkün olduğu sonucuna varılmıştır.

Anahtar kelimeler: Taşkın frekans analizi, Olasılık dağılımları, Parametre tahminleri

Methods Used in Flood Frequency Analysis

Abstract

Floods are disaster that have a very negative impact on community life and the economy, as well as being a naturally occurring event. When looking at the floods that caused the biggest economic losses after earthquakes in our country, there occurred 695 flood incidents between 1975 and 2010, and 634 people died in these incidents. In our country, due to reasons such as unplanned urbanization, irregular constructions in stream beds, climate change increased floods. Therefore, it is very important to reliably predict the size and frequency of floods with the correct planning and design of hydraulic structures. In this study, the methods used for flood frequency analysis were examined and it was concluded that it is possible to minimize the damage caused by floods by these methods.

Key words: Flood frequency analysis, Probability distributions, Parameter estimates

GİRİŞ

Taşkın, akarsu veya dere yataklarının şiddetli yağışlar sonucu debilerin yüksek değerlere ulaşmasıyla yataklarından taşarak çevresinde can ve mal kaybına, arazilere ve yapılara zarar verecek hale gelmesi olayına denilmektedir. Taşkın dünyanın birçok yerinde, gelişmiş ülkelerde bile önemli derecede can ve mal kayıplarına neden olmaktadır.

Dünyanın pek çok yerinde olduğu gibi, ülkemizde de taşkınlar çok önemli zararlara yol açmaktadır. Geçmişten bugüne yaşanan taşkınlar pek çok insanın ölümüne, yaralanmasına ve çeşitli şekillerde sağlığının bozulmasına neden olmaktadır. Her yıl milyonlarca TL taşkınlardan kaynaklanan zararın azaltılmasına ve yaraların kapatılmasına harcanmaktadır (Yüksek vd. 2013).

İklim değişikliği daha fazla yağışa ve daha fazla buharlaşmaya neden olacaktır. Hidrolojik çevrimin hızlanması, genelde daha fazla yağışla sonuçlanır. Buradaki sorun, bu yağışın ne kadarının yere

düşeceğidir. Yağışlar muhtemelen bazı bölgelerde artacak, bazı bölgelerde ise azalacaktır. İklim modelleri bölgesel tahminlerin yapılmasında yetersiz olup hidrolojik çevrim son derece karmaşıktır. Yağışlarda meydana gelen değişiklik, yerdeki yağış miktarı, yansıma ve bitki örtüsünü, bunlar da buharlaşma ve bulut formasyonunu etkileyebilir. Ayrıca ormanların yok edilmesi, şehirleşme ve su kaynaklarının aşırı kullanımı gibi insani faaliyetler de hidrolojik sistemi etkiler (Kalaycı 2003).

Taşkınlara çeşitli doğal etkiler ve rastgele iklimsel faktörler neden olduğu için rastgele olaylar olarak da kabul edilmektedir. Oluşma zamanları rastgele karakterli olaylar olduğundan kesin tahminlerde bulunmak mümkün olmamaktadır. Bu nedenle geçmiş tecrübelerden yararlanılarak belirli olasılıklarla meydana gelecek gelecekteki olayların tahminine başvurulur.

Hidrolojik çevrimin her bir kısmında meydana gelen hidrolojik olaylar pek çok değişkenden etkilendikleri için, her bir olaydaki değişkenler arasındaki bağıntılar kesin bir şekilde elde edilemediğinden, hidrolojik olayların deterministik bağıntıları çoğunlukla belirlenemezler. Bu nedenle rastgele karaktere sahip olan hidrolojik olayları incelemek olasılık ya da istatistik yaklaşımla mümkün olabilmektedir. Olasılık teorisi rastgele karakterdeki olayların olasılıklarını inceleyen bir matematik dalı olarak tarif edilirken, istatistik de rastgele nitelikteki bir değişkene ait gözlenmiş örnekleri inceleyerek bu değişkenin toplumu hakkında yargılara varan bir bilim olarak tanımlanmaktadır (Yılmaz 1995).

Herhangi bir akarsuda aynı yılda birden fazla değişen büyüklükte taşkın olaylarının gerçekleşebilmesi gösterilebilir. Dünyanın güneş etrafında bir yıllık sürede dönmesinden dolayı tekrarlanan hidrolojik olaylar için bir yıllık periyot, birim zaman değeri olarak alınabilir. Bir yıl içerisinde meydana gelen taşkın piklerinin en büyük değeri yıllık taşkın piki olarak tanımlanır ve bu değer hidrolojide önemli bir rastgele değişkendir. Çünkü en büyük debi değeri, en büyük taşkın zararını doğurur. Gerçek güvenilir taşkın pik serilerinde seri bağımlılığının ihmal edilecek mertebede oldukları istatistiksel analizlerle doğrulanabilir. Böylece, yeterince uzun zaman serisine sahip yıllık taşkın pikleri, bir doğal akarsuyun belirli bir kesitinde ölçülmüş gerçek akım değerleri ise, klasik bir frekans analizinde yapıldığı gibi, uygun bir teorik olasılık dağılımı uygulanabilir (Haktanır 1990).

Günümüzde birçok su yapısının tasarımında, taşkın kontrolünün planlama ve projelendirilmesi önemli bir yer tutmaktadır. Bu tür çalışmalarda maksimum yağışların ve akışların analizi önem kazanmaktadır. Bir yerleşim alanında tasarlanacak drenaj çalışmaları için 2 ve 5 yıl tekrerrür süreli maksimum yağış şiddetine ihtiyaç varken, erozyon ve sediment kontrol çalışmalarında 25 yıllık, baraj, gölet ve sulama tesisleri gibi yapılar için 100 yıllık tekrerrür periyoduna sahip maksimum yağış ve akış değerlerine ihtiyaç vardır. Bu nedenle ülkemizde Devlet Su İşleri (DSİ) ve Elektrik İşleri Etüd İdaresinin (EİEİ) akarsular üzerinde kurmuş olduğu akım gözlem istasyonlarından elde edilen maksimum yıllık akış değerlerinden yola çıkarak gelmesi muhtemel taşkınının büyüklüğü istatistik yöntemler kullanılarak tahmin edilir. Analizde kullanılan istasyonların gözlem süreleri ne kadar uzun olursa tahmin edilen tekrerrür değerleri gerçek değerlere o kadar yakın olur. Böylece gelebilecek taşkınının büyüklüğü önceden tahmin edilerek zararın minimuma indirilmesi için önlemler alınır (Seçkin 2009).

Bunun için istatistiksel yöntemlere ve frekans dağılımlarının hesabına gidilir. Meydana gelecek taşkın büyüklüğü ve frekansı su yapısının planlama ve projelendirmesinde, yer seçiminde, boyutlandırılmasında gereklidir ve taşkın yatağının kullanımındaki riskin belirlenmesinde önemli bilgiler verir. Yapılacak tahminlerin güvenilir olmasıyla da taşkın anında oluşabilecek zararlar olabildiğince minimum düzeye indirilmiş olacaktır.

TAŞKIN FREKANS ANALİZİ

Hidrolik yapıların projelendirilmesinde çeşitli yineleme aralıklarına karşılık gelen taşkın debilerin tahmini önemli bir yer tutmaktadır. Yanlış bir proje debisi seçilmesi durumunda istenmeyen iki durum ortaya çıkabilir. Bunlardan birincisi gereğinden büyük seçilen bir proje debisi nedeniyle yapı ekonomik

olmayacaktır. İkincisi ve daha tehlikelisi de küçük seçilen proje debisi nedeniyle projelendirilen yapı proje debisinden daha büyük gelebilecek taşkın debisi nedeniyle, yapı gelebilecek taşkın debileri durumunda yıkılma riskine maruz kalarak büyük can kayıplarına ve maddi zararlara neden olacaktır. Dolayısıyla su yapılarının projelendirilmesinde debilerin büyüklüğünün yeterince doğru olarak tahmin edilmesi oldukça önemlidir (Aydoğan 2008).

Özellikle taşkın koruma yapıları ve diğer su yapılarının projelendirilmesi aşamasında belirli tekerrür aralığında meydana gelebilecek olan taşkın debileri kullanılarak hesaplamalar yapıldığından taşkın debilerinin doğru bir şekilde tahmin edilmesi büyük önem taşımaktadır (Aydın 2015, Seçkin ve Topçu 2016).

Eldeki akım gözlemleri yardımıyla, çeşitli yineleme (tekerrür) süreli taşkın debilerinin tahmin edilmesi çalışmalarına taşkın frekans analizi denir (Anılan 2014). Taşkın frekans analizlerinde genelde iki tür seri kullanılır. Bunlar:

1. Yıllık taşkın serileri
2. Kısmi süre serileri

Yıllık taşkınlar, bir su yılında bir defa oluşan maksimum akıştır. Her bir yılda bir maksimum taşkının kullanılması, yıllık taşkınların analizinde esas alınır. Diğer yılların herhangi bir akışı başka bir yılı n maksimum akışından büyük olsa dahi dikkate alınmaz. Kısmi süre serileri ise daha pik değerlerle analiz yapabilmek için her yılın maksimum değeri yerine belli bir değerden daha büyük değerlerle oluşturulan serilerdir (Seçkin 2009). Taşkın frekans analizi ise noktasal taşkın frekansı ve bölgesel taşkın frekansı olmak üzere iki farklı yöntem ile yapılabilir.

TAŞKIN FREKANS ANALİZİ YÖNTEMLERİ

Noktasal Taşkın Frekans Analizi

Tek bir noktadaki hidrolojik verilerin (akım, yağış miktarı veya şiddeti) analizi noktasal taşkın frekans analizi ile elde edilir (Anılan 2014). Noktasal taşkın frekans analizinde seçilen istasyona ait veriler kullanılarak bu verilere en uygun olan olasılık yoğunluk fonksiyonları uyumun iyiliği testleri sonucu belirlenir ve bu olasılık yoğunluk fonksiyonu kullanılmak suretiyle belirli bir tekerrür aralığında meydana gelebilecek olan taşkın debileri hesap edilebilir (Aydın 2017).

Noktasal taşkın frekans analizindeki olasılık yoğunluk fonksiyonu belirlenirken Kolmogorov Smirnov, Anderson Darling ve Ki Kare testleri gibi uygunluk testleri kullanılarak seçilen veri setine ait en uygun olasılık yoğunluk fonksiyonu belirlenir. Bu şekilde belirlenen olasılık yoğunluk fonksiyonu kullanılarak taşkın debileri hesaplanır.

Nam vd. (2005) noktasal frekans analizinin tahmin edilecek tekrarlanma periyodundan daha kısa veri olduğu durumda uygun olmadığını, ancak gözlem süresinin tekrarlanma periyodunun iki katından fazla olduğu durumda ise uygun olduğunu belirttikleri çalışmalarında, Kore'deki yıllık maksimum yağmur verisinin genellikle 50 yıldan daha az olduğunu söylemişler ve 100 yıllık bir tekrarlanma periyodu tahmini için bölgesel frekans analizinin kullanılması gerektiğini savunmuşlardır. Güney Kore'de saatlik yağmur verisi ile gerçekleştirdikleri çalışmada kümeleme analizi yaparak homojen bölgelere karar vermişler, uygun dağılımı ise uygunluk ölçüsü testi ile genel lojistik dağılımı olarak belirlemişlerdir. Sonuç olarak kullandıkları frekans analizi tekniğinin performansını simülasyon deneyleri ile araştırmışlar ve uyumsuz olan istasyonların tekrarlanma miktarlarına etkilerini araştırmışlardır (Anlı 2009).

Nam vd. (2008) çok değişkenli tekniklerle yaptıkları bölgesel yağmur frekans analizinde homojen bölgeleri oluşturmada değişkenlerin tipi ve sayısının etkili olduğunu savunmuşlar ve Procrustes analizini

uygulayarak 21 değişken kullanmışlardır. Seçtikleri bu değişkenlere faktör analizi uygulayarak 5 faktör oluşturmuşlar ve Fuzzy-c tekniği ile 68 istasyonu 6 bölgeye ayırmışlardır. Daha sonra genel lojistik dağılımına bağlı olarak bölgesel yağmur miktarlarını noktasal frekans analizi, gösterge taşkın ve bölgesel şekil tahmin yöntemleri ile tahmin etmişler ve bunların karşılaştırmalarını yapmışlardır (Anlı 2009).

Bölgesel Taşkın Frekans Analizi

Farklı istasyonlar için benzer frekanslar hesaplanıyorsa, noktasal taşkın frekans yönteminde olduğu gibi tek bir istasyondan örnek kullanmakla çok sağlıklı sonuçlar elde edemeyebiliriz. Tüm veriler toplanarak analiz yaparak çok daha sağlıklı sonuçlara ulaşabiliriz. Bölgesel taşkın frekans analizi olarak adlandırılan bu yöntem, genellikle ölçüm yapılmamış veya yetersiz miktarda ölçümün bulunduğu havzalardaki taşkın debilerinin tahmininde kullanılır.

Kahya ve Gürses (1998), Konya Havzasındaki Yıllık Pik Akımların U-Momentler Analizi adlı çalışmada, bağımsız ve homojen karakterdeki yıllık pik akım serilerine bölgesel frekans analiziyle birlikte homojenlik analizi yapmak için, L-momentler yöntemi kullanmışlardır. Bu amaçla, bölgesel frekans analizinde kullanılacak olan istasyon grubunda Hosking (1986) ve Wallis'in (1993) L-momentlere dayalı ifade ettikleri üç faydalı istatistik bu çalışmada ele almışlardır. Heterojenlik ölçüsü ve tasarlanan olasılık dağılımının verilere uyup uymadığını belirlemek amacıyla uygunluk ölçüsü olmak üzere tipik bölgesel frekans analizini gerçekleştirmişlerdir. Analiz yöntemleriyle birlikte; LogNormal, Gumbel, Logistic, Ekstrem Değerler, Pearson-3, Log-Pearson-3 frekans dağılımları Konya Havzası'nda bulunan akarsulara ait yıllık pik akım serilerine uygulanmış ve sonuçları değerlendirmişlerdir.

Castellarin vd. (2001), bölgesel taşkın frekans analizi homojenlik belirlenmesi için hidrolojik benzer ölçülerin etkilerini araştırmışlardır. Analizde kullanılan hidrolojik benzerlik ölçümleri; ekstrem olayların zamanlamasını tanımlayan mevsimsel göstergeler, yağış frekans dağılımların havza ölçeğinde karakteristikleri, önceki yağış istatistikleri ve havza geçirimsizlik değerlerinden faydalanılmaktadır. Araştırmacılar Kuzey İtalya'da büyük bir alanda Monte Carlo yönteminden yararlanarak yaptıkları uygulama sonucunda, bölgesel göstergelere dayalı ekstrem değer tahminlerinin daha etkileyici sonuçlar verdiğini göstermişlerdir.

Heo vd. (2001), bölgesel frekans analizinde kullanılan Weibull modelinin hesap ve asimptotik varyans değeri üzerinde bir çalışma yapmayı amaçlamışlardır. İki parametreliliğin parametrelerini ve her bir parametrenin asimptotik varyansını, momentler ve olasılık ağırlıklı momentler yöntemleriyle hesaplamışlardır. Asimptotik varyansları, Cramer Rao testiyle mukayese etmişlerdir.

McCuen ve Beighley (2003), yıllık maksimum ve mevsimsel maksimum akımlar arasında verilen bir yağış dağılımı içindeki mevsimsel değişimi herhangi bir T tekerrür periyodu için araştırmışlardır. Genel bir problem olan ölçülmüş bir veriye mevsimsel frekans analizi uygulandığı zaman, eksik akım değerleriyle karşılaşma sorununu aşmanın iki yolu vardır. Birincisi; maksimum olabilirlik yöntemi kullanılarak eksik verileri tahmin etmek, diğeri günlük akımların aritmetik ortalamasını hesaplayıp eksik veri yerine ortalama değerini koymaktır. Yağış kayıtları bulunmayan bölgelerde frekans analizinin yağış-akım modeli veya regresyon eşitlikleri yardımıyla elde edilen frekans eğrilerinin kullanılmasıyla yapılabileceği kanısına varmıştır.

Özen ve Benzeden (2005), Batı Akdeniz akarsu havzalarında yıllık zirve akışlarını bölgesel analizi ile incelemişlerdir. 48 ölçüm istasyonundaki verilerin inceleme sonucunda 48132 km² civarındaki proje alanının "Yukarı-Batı Akdeniz", "Aşağı-Batı Akdeniz" ve "Antalya" isimli üç alt bölgeye ayrılması gerektiği görmüşlerdir. Yıllık ortalama taşkın verimi-yağış alanı, çarpıklık değişkenlik katsayıları arasındaki bölgesel ilişkiler incelemişlerdir. Bu incelemeler, istatistikî açıdan anlamlı, birbirine hemen hemen paralel iki ayrı taşkın verimi-yağış alanı bağıntısı ve Gamma dağılımına oldukça yakın tek bir

çarpıklık-değişkenlik bağıntısının bulunduğunu göstermiştir. Bölgesel homojenlik koşulunu gerçekleştirmeyen istasyonlar çıkarıldıktan sonra, değişkenlik ve çarpıklık katsayılarının bölgesel değerleri hesaplanmış ve üç alt bölge için bölgesel frekans eğrileri elde etmişlerdir.

Bölgesel frekans analizinde izlenen aşamalar

Bölgesel taşkın frekans analizi dört aşamadan meydana gelmektedir. Bunlar; verilerin gözden geçirilmesi, homojen bölgelerin belirtilmesi, en uygun bölgesel olasılık dağılımının belirtilmesi ve bölgesel olasılık dağılımının değerlendirilmesidir.

Verilerin gözden geçirilmesi:

Herhangi bir istatistiksel analiz için ilk gerekli adım verinin analiz için uygun olup olmadığını kontrol etmektir. Frekans analizi için bir istasyondan toplanan veri, ölçülen miktarları doğru bir şekilde temsil etmeli ve bütün miktarların aynı olasılık dağılımından çekilmesi gereklidir. Bu nedenle verinin derlenerek incelenmesi, büyük hataların, tutarsızlıkların, aykırı değerlerin giderilmesi ve zaman içinde meydana gelen değişimlerden dolayı verinin istatistiksel karakterinin değişip değişmediğinin araştırılması çok önemlidir (Anlı 2009).

Uyumsuzluk ölçüsü, istasyon verilerinin L-moment oranları ile hesaplanır. İstasyonun L-moment oranları (L-Cv, L-çarpıklık, L-basıklık) bir noktanın üç boyutlu koordinatları olarak tanımlanır. Bu tanımlanan noktaların L-cv ve L-çarpıklık değerleri grafikte karşılıklı olarak noktalandığında bir grup oluşturur ve bu grup bir merkeze yani orta noktaya sahiptir. Uyumsuz olarak adlandırılan herhangi bir nokta, bu merkezden oldukça uzaktır.

$$D_i = \frac{1}{3} N(u_i - \bar{u})^T A^{-1} (u_i - \bar{u}) \quad (1)$$

D_i , bölgedeki istasyon sayısına bağlı olarak tanımlanır. Eğer hesaplanan D_i değeri kritik D_i değerinden büyük ise o istasyon uyumsuzdur denir (Hosking ve Wallis 1997).

Homojen bölgelerin belirlenmesi:

Homojen yağmurların etkili olduğu alana, hidrolojik homojen bölge denir. Frekans analizlerine göre belirtilen miktarların etkili olduğu alanın bilinmesi gerekir (Anlı 2006). Bu bakımdan su toplama havzası veya göz önüne alınan alan, hidrolojik homojen bölgelere ayrılır. Hidrolojik yönden homojen olan bir bölgede farklı istasyonlarda ölçülen yağmurlar bir istasyonda ölçülmüş olarak alınabilir. Böylece her istasyonda ölçülen yağmurlardan meydana gelen sayıda bir veri elde edilir (Anlı 2009).

Heterojenlik ölçüsü özellikle homojen olması muhtemel bölgelerin istasyonları arasında örnek L-momentlerin varyasyonlarını karşılaştırır. Homojen bir bölgede bulunan tüm istasyonlar, aynı toplam L-moment oranlarına sahiptir.

$$H_i = \frac{(V_i - \mu_v)}{\sigma_v} \quad (2)$$

Hosking ve Wallis (1993), eğer $H_i < 1$ ise bölgenin kabul edilebilir derecede homojen olduğunu, $1 < H_i < 2$ ise bölgenin muhtemelen heterojen olduğunu, $H_i > 2$ ise bölgenin kesinlikle heterojen olduğunu söylemişlerdir. Eğer bölge yeterince homojen değil ise, bölge daha alt bölgelere ayrılarak homojen hale getirilmeye çalışılır (Seçkin ve Yurtal 2008).

En uygun bölgesel olasılık dağılımının belirtilmesi:

Olasılık dağılımlarının seçilmesi büyük ölçüde, dağılımların veriye ne kadar uygun olduğuna bağlıdır. Mevcut veriye birden fazla dağılım uygunluk gösterdiğinde, bunların arasından en iyi tekrarlanma tahminlerini verecek olan güçlü dağılımın seçilmesi bölgesel frekans analizinin hedefidir. Dağılımın uygunluğunun belirtilmesinde kullanılan önemli testlerden birkaçı Kolmogorov-Smirnov, ki-kare,

moment ve L momentlere dayanan testler olarak sayılabilir. Varsayılan bir bölgesel frekans dağılımının her istasyon verisine yeterliliği, her istasyon için hesaplanan uygunluk istatistiği (noktasal) ile değerlendirilir ve elde edilen sonuçlar bölgesel uygunluk istatistiği ile birleştirilir (Rao ve Hamed 1997).

Bölgesel olasılık dağılımının değerlendirilmesi:

Bölgesel frekans analizi için kullanılabilir verinin olduğu istasyonların, yaklaşık olarak homojen ve aynı olasılık dağılımlarına sahip olan bir bölgeye atanması ve her bir bölge verisine uygun bir olasılık dağılımı seçilerek tekrarlanma miktarlarının elde edilmesi bu kısımda gerçekleştirilmiştir. Farklı istasyonlardaki olasılık dağılımlar arasındaki ilişki, bölgesel frekans analizi için bir karar verme anlamındadır. Farklı istasyonların verisinin birleştirilmesi ile elde edilen dağılım parametrelerini ve tekrarlanma miktarlarını belirten bu ilişki, her bir istasyona ayrı olarak uygulanan dağılımlar ile elde edilenden daha doğru tahminlere olanak sağlar (Anlı 2009).

İleride yapılacak benzer çalışmalara yardımcı olması açısından bazı öneriler aşağıda verilmiştir (Anlı 2009):

- Bölgesel frekans analizi, bölgelerin bir miktar heterojen olması, istasyonlar arasında bağımlılık olması ve olasılık dağılımının yanlış belirtilmesi durumunda bile noktasal frekans analizinden daha doğru sonuçlar verebilir.
- Bölgeselleştirme işlemi, özellikle dağılımın alt ve üst kuyruklarındaki büyüme eğrisi ve tekrarlanma miktarlarının tahmininde önemlidir.
- Heterojen bölgelerdeki tekrarlanma ve büyüme eğrisi tahminlerindeki hatalar, bölgelerdeki istasyon sayısının artmasıyla oldukça azalmaktadır. Ancak istasyon sayısının 20'den fazla olduğu bölgelerde ise, söz konusu tahminlerde hatalar meydana gelebilmektedir.
- Uzun gözlem süreleri, bölgesel tahminlerin noktasal tahminlere göre tamamen doğru olabileceğini göstermemektedir. Ancak uzun gözlem süreleri heterojenliği belirtmede kolaylık sağlamaktadır. İstasyonlardaki gözlem süreleri uzunsa, heterojenliği belirtmek için bölgelerin birkaç istasyon içermesi düşünülebilir.
- Bu çalışmada bölgesel frekans analizinde iki parametrelili dağılımlar kullanılmamıştır. Çünkü dağılımların L-moment oranlarının bölgelerde bulunan istasyonlardaki olasılık dağılımlarının L-moment oranlarına yakınlığı hakkındaki eksik bilgi, tekrarlanma miktarlarında büyük taraflılık doğuracaktır.
- Olasılık dağılımının yanlış belirtilmesi özellikle yüksek olasılıklardaki ($F > 0.99$) tekrarlanma tahminlerinde önemli bulunmaktadır.
- İstasyonlarda hesaplanan tahminlerdeki taraflılığı ortaya çıkaran heterojenlik tüm bölge için geçerli olmayabilir, bu taraflılık tahmin edilen miktarlarda ve büyüme eğrisindeki hatalardan kaynaklanabilmektedir.
- Bölgesel tahminlerde istasyonlar arası küçük miktarda bağımlılık pek göz önüne alınmamaktadır.

MATERYAL VE YÖNTEM

PARAMETRE TAHMİN YÖNTEMLERİ

Bir rastgele değişkenin toplumunun tümünü belirlemek mümkün olmadığından, bu rastgele değişkene ait bir örnekten bu toplumun olasılık dağılım fonksiyonunun parametrelerinin tahmin edilmesi gerekmektedir. Bir üstel fonksiyon olan olasılık dağılım fonksiyonunun şekli, biçimi, simetriği, yaygınlığı ve sivriligi, analitik ifadesine ve parametrelere bağlıdır. Söz konusu parametreler, belirli bir dağılıma uyan rastgele değişkene ait istatistiksel karakteristiklerin belirlenmesinde kullanılır (Önöz 1994).

Rastgele değişkenin parametrelerinin gerçek değerleri tam bir şekilde belirlenemeyeceğinden elimizdeki örnekten tahmin edilmesi gerekir. Bu tahminlerde iyi bir yöntem kullanılmasıyla parametrelere yaklaşık değerler elde edilebilir. Parametre tahmininden beklenen tahminlerin yansız olmasıdır. Yansız bir tahmin beklenen değerlerin parametrenin gerçek değerleri ile eşit olan tahmindir.

Momentler Yöntemi

Bir rastgele değişkenin olasılık yoğunluk fonksiyonunun çeşitli parametreleriyle merkezsiz momentler arasında ilişkiye dayanan yöntemdir. Buna göre i . dereceden moment,

$$\mu_i = \int_{-\infty}^{\infty} x^i f(x) dx \quad (3)$$

ilişkisiyle hesaplanır. $i=1$ iken $x=\mu_1$ noktasına göre alınan momente merkezsiz moment adı verilir,

$$\mu_i = \int_{-\infty}^{\infty} (x - m_1)^i f(x) dx \quad (4)$$

Birinci dereceden moment bizi aritmetik ortalamaya diğer bir deyişle rastgele değişkene ait olan merkezsiz değere ulaştırır.

$$\mu_x = \int_a^b x \cdot f(x) dx \quad (5)$$

ortalamayı veren bu eşitlikte a alt sınır ve b üst sınır değerlerini temsil eder. Elde edilen ortalama rastgele değişkenin bu değer çevresindeki yayılmasının büyüklüğüyle ilgili bilgi vermez. Rastgele değişkenin bu değer çevresindeki yayılmasının büyüklüğüne ikinci dereceden denklem merkezsiz moment değeri olan varyans ile ulaşılır. Varyans,

$$Var(x) = \sigma^2 = \int_a^b (x - m_x)^2 f(x) dx \quad (6)$$

eşitliğiyle hesaplanır. Varyansın karekökünü ifade eden standart sapma, rastgele değişkenin aldığı değerler ile aynı boyutta olmasından dolayı dağılım yayılımını hesaplamada sıklıkla kullanılan bir parametredir. Standart sapmanın büyümesiyle değişken dağılımının yaygınlığı artmaktadır.

Rastgele değişkenin olasılık dağılımının merkez etrafında simetrik olması durumunda tek sayılı merkez momentler sıfırdır. Bu nedenle 3. dereceden merkezsiz moment çarpıklığının iyi bir ölçüsüdür. Çarpıklık katsayısı,

$$G = \frac{\mu_3}{\sigma_x^3} \quad (7)$$

eşitliğiyle hesaplanır. Rastgele değişkenlerin dağılımları genelde çarpık olduğundan çarpıklık katsayısı önemli bir parametredir. Çarpıklık katsayısı pozitif ise dağılım sağa doğru, negatif ise sola doğru kuyruğu olduğunu gösterir. Çarpıklık katsayısının sıfır olması olasılık yoğunluk fonksiyonun simetrik olduğunu gösterir. Olasılık yoğunluk fonksiyonunun tepesini düz ya da sivri oluşu 4. dereceden merkezsiz moment olan Kurtosis (basıklık) katsayısı ile belirlenir. Kurtosis katsayısı,

$$K = \frac{\mu_4}{\sigma_x^4} \quad (8)$$

eşitliğiyle hesaplanır. Kurtosis katsayısının değeri büyüdükçe sivrilmesi artar ve bu normal dağılımın sivrilmesine göre ölçülür.

Yukarıda bahsi geçen momentlerin toplumdaki gerçek değerleri tam olarak bilinmemesi nedeniyle örnekten bu değerlere en iyi yaklaşımlar yapılabilir. Diğer bir deyişle istatistiksel karakteristikler örnekten hesaplanır. Formüller aşağıda verilmiştir.

Ortalama:

$$\bar{x} = \frac{\sum_{i=1}^N X_i}{N} \quad (9)$$

Varyans:

$$S_x^2 = \frac{\sum_{i=1}^N (X_i - \bar{x})^2}{N - 1} \quad (10)$$

Çarpıklık katsayısı:

$$C_s = \frac{m_x^3}{S_x^3} \quad (11)$$

$$m_x^3 = \frac{N \cdot \sum_{i=1}^N (X_i - \bar{x})^3}{(N - 2) \cdot (N - 1)} \quad (12)$$

Kurtosis Katsayısı:

$$C_{sx} = \frac{m_x^4}{S_x^4} \quad (13)$$

$$m_x^4 = \frac{N^2 \cdot \sum_{i=1}^N (X_i - \bar{x})^4}{(N - 3) \cdot (N - 2) \cdot (N - 1)} \quad (14)$$

Yukarıdaki formüller iki parametrelili dağılımlar için kullanılır. Dağılıma ait parametre sayısı üç ise sol taraflarının integralleri alınarak sağ taraflarındaki örnek seriden elde edilen tarafsız tahminlere eşitlenip parametreler bulunur (Benjamin ve Cornell 1970). Momentler yöntemi etkin tahminler vermesi ve uygulanabilirliğinin kolay olmasından dolayı sıklıkla kullanılan bir yöntemdir.

Maksimum Olabilirlik Yöntemi

Olasılık yoğunluk fonksiyonu $f(x; a, B, \dots)$ bir rastgele değişkene ait örnekteki $x = x_i$ olayının meydana gelme olasılığı $f(x_i; a, B, \dots)$ ile orantılı olduğundan $x = x_1, x = x_2, \dots, x = x_N$ bağımsız olaylarının meydana gelme olasılığında,

$$L = \prod_{i=1}^N f(x_i; a, B, \dots) \quad (15)$$

ifadesi ile orantılıdır. L'ye olabilirlik fonksiyonu denir ve bu fonksiyonu maksimum yapan $\hat{a}, \hat{B}, \dots$ değerleri $a, B, \dots$ parametrelerinin tahminleri olmaktadır. Pratikte L için tanımlanan çarpımı toplam haline getirmek için L yerine $\ln L$ ile çalışılır. Bu taktirde $\hat{a}, \hat{B}$ tahminleri için aşağıdaki denklemler çözülmelidir.

$$\frac{\partial \ln L}{\partial \hat{a}} = \frac{\partial \ln L}{\partial \hat{B}} = \dots = 0 \quad (16)$$

Bu yöntemle yapılan tahminler daha etkin tahminler olmakla birlikte denklemlerin çözümü ancak ardışık yaklaşımlarla mümkündür (Bayazıt 1981).

Olasılık Ağırlıklı Momentler Yöntemi

Greenwood vd (1979) tarafından Wakeby dağılımının parametre tahmini için geliştirilen olasılık ağırlıklı momentler yöntemi, daha sonra Hosking (1986) tarafından incelenmiş ve bu momentlerin merkezel istatistik momentlerle eş değer özelliklere sahip olduğu sonucuna varılmıştır. Bu momentlerin örnek tahminleri özellikle küçük örnekler için tarafsızdır ve aykırı değerlere karşı hassas sonuçlar vermemektedirler. Aykırı değerlerle kasıt, mevcut verinin trendinden önemli derecede ayrılan veri noktalarıdır. Bu aykırılıkların modifikasyonları ve bozulmaları özellikle küçük örnekler için datadan hesaplanan istatistik parametrelere mühim bir derecede tesir edebilir. Bunun yanında bu yöntemle, verinin lineer fonksiyonu olmaları sebebiyle diğer momentlere göre örnekleme değişimlerinden daha az etkilenmektedir. Bu özellik, diğer yöntemlere göre bir avantajdır (Önöz 1994).

Olasılık ağırlık momentler,

$$M_{1,j,k} = E[X^j F^j (1 - F)^k] \quad (17)$$

$$M_{1,j,k} = \int_0^1 x(F)^j = F^j \cdot (1 - F)^k d \quad (18)$$

şekliyle ifade edilmiştir. Eşitliklerdeki $F = F(x) = P(X \leq x)$ ve $1, j, k$ pozitif çift sayılardır. $j=k=0$ ve 1 de tam sayı olduğu durumda $M_{1,0,0}$ moment 1 . dereceden merkezel momente eşit olur. Buna göre,

$$M_{1,j,k} = \frac{1}{N} \sum_{i=1}^N \frac{\binom{I-1}{j}}{\binom{N-1}{j}} x(i) \quad (19)$$

$$M_{1,0,k} = \frac{1}{N} \sum_{i=1}^N \frac{\binom{I-1}{k}}{\binom{N-1}{k}} x(i) \quad (20)$$

tarafsız denklemleri elde edilir. Eşitliklerde, $j=0, 1, \dots, N-1$; $k=0, 1, \dots, N-1$; $i=1, 2, \dots, N$ dir. $x(i)$ örneği temsil etmektedir ve $j=1$ örnekteki en küçük değeri gösterir (Hosking 1986).

Olasılık ağırlık momentleri kullanılarak parametre tahmininde bir başka yol da şu şekildedir. Rastgele değişkenin i 'inci değere eşit veya küçük kalma frekansı Landwehr formülleri ile belirlenir (Haktanır ve Bozduman 1995).

$$F(i) = \frac{i - 0.35}{n} \quad (20)$$

$$F(i) = \frac{i - 1}{n - 1} \quad (21)$$

$$F(i) = \frac{(i - 1)(i - 2)}{(n - 1)(n - 2)} \quad (22)$$

eşitliklerine dayanır. Dağılım veriye iyi bir şekilde uyarsa yöntem başarıya ulaşabilmektedir. Olasılık fonksiyonu aşağıdaki eşitliklerle tanımlanır:

$$M_{1,j,0} = \frac{1}{N} \sum_{i=1}^N x(i) f(i)^j \quad (23)$$

$$M_{1,0,k} = \frac{1}{N} \sum_{i=1}^N x(i) [1 - F(i)]^k \quad (24)$$

Bu yöntem maksimum olabilirlik yönteminden farklı olarak, $x = x(F)$ şeklinde ters formu açıkça belirlenen dağılımların parametre tahminlerinde kullanılır ve tahminleri oldukça kolaydır.

$$M_{1,0,k} \sum_{j=0}^k j (1)^j M_{1,j,0} \quad (26)$$

$$M_{1,j,0} \sum_{k=0}^j k (1)^k M_{1,0,k} \quad (27)$$

ifadeleri elde edilir. Buna göre ilk iki moment,

$$M_{1,0,0} = M_{1,0,0} \quad (28)$$

$$M_{1,1,0} = M_{1,0,0} - M_{1,0,1} \quad (29)$$

$$M_{1,2,0} = M_{1,0,0} - 2M_{1,0,1} + M_{1,0,2} \quad (30)$$

Örnekteki gözlenmiş elemanların büyüklükleri gibi istatistikleri hesaba katması ve üç parametrelili dağılımlar için, tüm olasılık ağırlık momentleri için örnekten hesaplanan x ifadelerinin birinci kuvvetlerinin hesaplanması gerekirken, momentler yönteminin üçüncü kuvvet ifadesinin istenmesi, bu yöntemin en önemli iki avantajıdır. Özellikle küçük örnekler için, üçüncü kuvvet ifadesi yukarıda açıklanan aykırı gözlemler çıkmasına sebep olmaktadır. Bunun için momentler yöntemine göre stabil bir yöntem olduğu söylenebilir (Haktanır 1991b).

Karışık Momentler Yöntemi

Bu metot, Rao (1983) tarafından özellikle Log-Pearson 3 dağılımının parametrelerini hesaplamak üzere tasarlanmıştır. Log-Pearson 3 parametrelerinin tahmini için momentler ve genelleştirilmiş momentler gibi dağılımın ortalama varyans ve çarpıklık değerlerini kullanan yöntemler bulunmaktadır. Bunlar arasında en basit işlem; Log-Pearson-3 dağılımını logaritmik dönüştürme işlemi yapılmış veriye uygulayan dolaylı moment yöntemidir. Bobee (1975), dönüştürme uğramamış verinin orijinine ait ilk üç momenti kullanan direkt moment metodunu uygulamıştır. Rao, üçüncü örnek momentinin uygulamada istenmeyen sonuçlar verdiğini fark etmiş ve orijinal verinin örnek ortalama ve varyansı ile logaritmik dönüşüme uğramış verinin örnek büyüklüğünü kullanan karışık moment metodunu tasarlamıştır (Koutrovelis ve Canavos 2000).

$f(x)$ dağılımının olasılık yoğunluk fonksiyonuna ve y değerlerinin ortalama, varyans ve çarpıklık katsayıları momentler yöntemine göre,

$$i_y = c + b.a \quad (31)$$

$$\sigma^2 = b.a^2 \quad (32)$$

$$\tilde{a}_y = 2a(a\sqrt{b}) \quad (33)$$

eşitliklerinden elde edilir. Pearson-3 dağılımının moment değerlerinin genelleştirilmesiyle daha elverişli temsil fonksiyonu aşağıdaki şekilde yazılabilir:

$$g(y)^t = E[\exp(t.y)] = \exp(c.t)(1 - a.t)^b \quad (t=1, 2, 3) \quad (34)$$

$$h g(y)^t = c.b.In(1 - a.t) \quad (t=1, 2, 3) \quad (35)$$

Yapılan çalışmalarda, büyük tekerrür periyotları için taşkın tahminlerinde ve özellikle tamsayıların moment değerlerinin hesaplanmasında uygun sonuçlar verdiği gözlemlenmiştir (Rao 1983). Bu yöntem şekil parametresi (a)nın 0.5'den küçük değerleri için geçerlidir. Olasılık dağılımın a parametresi,

$$P(a) = \frac{2\ln(1-a) \ln(1-2a)}{a + \ln(1-a)} = P(n) \quad (36)$$

eşitliğiyle bulunur. $P(n)$,

$$P(n) = \frac{h g_n(y) \cdot 3 \cdot \beta - 2h g_n(y) f(x)}{y - h g_n(y) f(x)} \quad (37)$$

ile hesaplanan b ve c 'nin eliminasyonundan sonra üç eşitliğin çözümünden elde edilen eşitliklerdir. Daha sonra $t=1$ değeri için, b ve c parametreleri elde edilir. Sistem ancak iteratif yaklaşımla çözülebilir. İşlemleri kolaylaştırmak amacıyla Arora ve Sign tarafından $a < 0.5$ değerleri için $P(a)$ değerlerini veren tablo ortaya konmuştur (Koutrouvelis ve Canovos 2000).

Maksimum Entropi Yöntemi

Maksimum olabilirlik yöntemine alternatif olarak geliştirilmiştir. Yani temel düşüncesi, eldeki serinin, olması muhtemel benzeri örnek serilere göre meydana gelme ihtimalinin büyük olması durumudur. Maksimum olabilirlik yönteminde parametre tahminleri oldukça zordur. Bu yaklaşımda bilinen en büyük problem eldeki veriyi iyi temsil etme zorunluluğudur. Maksimum entropi yöntemi bu sakıncayı gidermek amacıyla tasarlanmıştır. Bu yöntem özellikle sınırları verinin parametrelerinin büyük değer almaları durumunda iyi sonuçlar vermektedir (Howitt ve Reynaud 2003).

Olasılık kütle fonksiyonu P_k ile gösterilen bir rastgele değişken x ; $x_1, x_2, \dots, x_n$ değerleri almış olsun. Değişkenin olasılık fonksiyonu,

$$\sum_{k=1}^k g(x) \cdot P_k = c \quad (38)$$

eşitliği elde edilmiş olur. Eğer $g(x) = X$ ise, c değeri X' in ortalamalarına eşittir. x_i değerlerini maksimum yapan entropi fonksiyonu denklemi,

$$|H(x) = \sum_{i=1}^N p(X = x_i) \cdot \log(p(X = x_i)) \quad (39)$$

Fonksiyonun çözümü için P_k değeri hesaplanmalıdır. P_k değerinin hesaplanması için örneğin olabilirliğini maksimum yapan m parametresinin değerleri bulunur. Maksimum olabilirlik yöntemine benzer olarak olabilirlik fonksiyonunu maksimum yapan $\theta_1, \theta_2, \dots, \theta_m$ değerlerini bulma işlemi ile karşılaşılır. θ değeri,

$$P_e(O|x) \propto P(x|O)P_e(\theta) \quad (40)$$

ile elde edilir. X incelenen veriyi temsil eder. Denklemdaki P_e değeri,

$$P_e(\theta) \propto e^{H(\theta)} \quad (41)$$

eşitliğiyle elde edilir. Buradaki $H(\theta)$ entropi değeridir. Olması mümkün en büyük olasılığı tahmin amacıyla verilen fonksiyonu maksimum yapan C değeri,

$$C = b \cdot \log P(O) + (1 - b)H(T) \quad (42)$$

eşitlikten bulunur. b , “0” ile “1” arasında değişen denge katsayısıdır. T , olması mümkün değerleri gösterir.

Özgün Olasılık Ağırlıklı Momentler Yöntemi

Bu yöntem olasılık ağırlıklı momentler yönteminin farklı bir düzenlemesi şeklinde tasarlanmıştır. Bu yöntemde, toplumun j ' inci olasılık ağırlıklı moment değeri, sonlu sayıda örnek serileri için,

$$pwm_j = \sum_{i=1}^N x_i (Pn x_i)^j / N, \quad (j = 0, 1, 2, \dots) \quad (43)$$

$$pwm_j = \sum_{i=1}^N x_i (1 - Pn x_i)^j / N, \quad (j = 0, 1, 2, \dots) \quad (44)$$

ile hesaplanır. Yöntemin analitik uygulaması için öncelikle, $Pn x_i$ değeri olasılık dağılımının eklenik dağılım fonksiyonunda yerine konması durumunda,

$$pwm_j = \left[\sum x_i \cdot F(x_i, parametreler)^j \right] / N \quad (45)$$

$$pwm_j = \left[\sum x_i \cdot \{1 - F(x_i, parametreler)\}^j \right] / N \quad (46)$$

olur. Eşitliklerden hesaplanan olasılık dağılım parametre değerlerinin, tam anlamıyla orijinal parametre değerlerine uyduğu görülmüştür. Parametreler eklenik olasılık dağılım fonksiyonunun içinde mevcut olmalarının yanında, herhangi bir dağılımın parametrelerinin sonuçlarını hem bağımlı hem de bağımsız değişken olarak vermektedir. Bu nedenle, parametre değerlerinin hesaplanmasında bazı kısa iteratif sayısal metotlar kullanılmalıdır.

Yöntemde, ölçülmüş örnek seri elemanlarının umulan olasılık tahmin değerleri grafik üzerinde işaretlenir. Yöntemin uygulanabilmesi için öncelikle örnek elemanlarının bir kurala göre sıralanması gerekir.

Sınırlı uzunluktaki örnekler için de iyi yaklaşımlar vermemektedirler. Bu metodun iyi sonuç verebileceği örnek uzunluğu sayısı olan N , Landwehr formülüyle hesaplanabilir. Örneğin, eleman sayısı 50 olan bir örnek üzerinde yapılan parametre tahminin % 99.3 güven aralığında sonuç verdiği görülmüştür. Bu yönüyle olasılık ağırlıklı moment yöntemine göre daha iyi sonuç vermektedir (Haktanır 1997).

L-momentler Yöntemi

Olasılık dağılımı, dağılımın momentleriyle tanımlanmaktadır. Momentler, dağılımın merkezi eğilimini gösteren ortalama ve daha yüksek dereceli momentler olarak ele alınabilmektedir. Bu teknik Hosking (1990) tarafından tanımlanan, L-moment istatistikleri, gözlem verisinin karesinin ve küpünün alınmadan elde edilen doğrusal bileşenleridir. L-momentler Olasılık dağılımların şekillerini tarif eden bir sistem olmakla birlikte, normal çarpım momentlerine göre uzun süreli veride daha az duyarlılığa sahiptir. Bir X verisinin L-momentleri olasılık ağırlıklı momentlerin fonksiyonu olarak tarif edilmiş ve buradan sıralanmış gözlemlerden $X(j:n)$ elde edilen olasılık ağırlıklı momentlerin yansız örnek tahmini olarak Greenwood vd. (1979) tarafından aşağıdaki gibi ifade edilmiştir,

$$b_r = n^{-1} \sum_{j=1}^n \frac{(j-1)(j-2) \dots (j-r)}{(n-1)(n-2) \dots (n-r)} x_{j:n} \quad (47)$$

Daha sonra b_r değerlerinin ilk dördü ($r=0, 1, 2, 3$) olasılık ağırlıklı momentler (b_0, b_1, b_2 ve b_3) bulunduktan sonra, herhangi bir dağılım için L-moment istatistikleri,

$$\ell_1 = b_0, \quad (48)$$

$$\ell_2 = 2b_1 - b_0, \quad (49)$$

$$\ell_3 = 6b_2 - 6b_1 + b_0, \quad (50)$$

$$\ell_4 = 20b_3 - 30b_2 + 12b_1 - b_0 \quad (51)$$

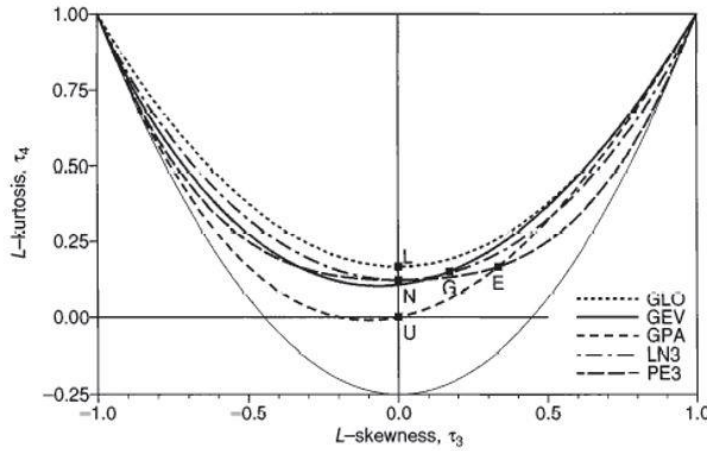
ilişkilerinde saptanır. İlk L-moment olan, merkezi eğilim ölçüsü olmasının yanı sıra dağılımın ortalamasına eşittir. Dağılımın ölçüsü ise ℓ_2 dir. Buradan boyutsuz L-moment oranları,

$$t = \frac{\ell_2}{\ell_1} (L \text{ değişim katsayısı}) \quad (52)$$

$$t_3 = \ell_3 / \ell_2 (L \text{ çarpıklık}) \quad (53)$$

$$t_4 = \frac{\ell_4}{\ell_2} (L \text{ basıklık}) \quad (54)$$

tahmin edilmiştir. Yukarıda bahsedildiği gibi, b_1 ve b_2 , L-CV (L değişim katsayısı, t) ve t_3 ve t_4 yani L-moment oranları, olasılık dağılımlarını özetleyen en kullanışlı değerlerdir. Farklı dağılımların L-momentlerini temsil etmenin en uygun yolu L-moment oran diyagramıdır. Sadece yer ve ölçü parametreleri ile tanımlanan iki parametrelili dağılımların sadece iki değeri rastgele dağılmakta, L-çarpıklık ve L-basıklık değerleri aynı olmakta ve sadece bir nokta ile diyagramda gösterilmektedir. Üç parametrelili bir dağılımda ise, yer, ölçü ve şekil parametreleri olmaktadır. Bu sebeple farklı şekil parametreleri ile farklı noktalar işaretlenmekte ve bir çizgi ifade edilebilmektedir. Bu sebeple aşağıdaki diyagram özellikle üç parametrelili dağılımlar için kullanılmaktadır (Hosking ve Wallis 2005).



Şekil 1. Bazı dağılımlara ait L moment oran diyagramı

Bu metodun uygulanması için aşağıda sıralanan hususlar aranır (Şorman 2004):

1. Frekans analizi hatasız ve güçlü olmalıdır. Bunun anlamı: Bir modelleme yönteminin güçlü olabilmesi için yöntemin gerçek fiziksel sürecin modelin kabullerinden farklılaşma göstermesi durumunda dahi tahmin edilen değerin (Q_T) gerçekten randımanı iyi ve oldukça hatasız olmasıdır. Aksi halde Q_T nin bulunmasındaki yöntem ile zayıf tahmin yapılıyorsa bunun güçlü olmayacağı anlaşılmalıdır.
2. Frekans analiz yöntemi simülasyona dayandırılmalıdır.
3. Bölgeselleştirmenin önemli ölçüde bu tür çalışmada katkısı vardır.
4. Bölgelerin coğrafik baza dayandırılma zorunluluğu yoktur.
5. L-moment istatistik parametreleri dağılımın geniş bir alanını kapsar ve hatalı olma özelliği azdır.

Klasik parametre tahmin yöntemleri ile karşılaştırıldığında L-momentler yönteminin avantajları şunlardır (Şorman ve Okur 2000):

1. L-momentler yöntemi ile bulunan varyasyon, çarpıklık ve basıklık katsayıları hemen hemen hatasızdır ve yaklaşık normal bir dağılıma sahiptir. Aynı çarpım momentleri küçük örneklerde oldukça değişken ve hatalıdır.
2. L-momentler, çarpım momentlerinden daha hatasız oldukları için moment diyagramları oluşturulmasında kullanımları daha uygundur.
3. L-momentler, dağılım ortalamasının bulunabildiği her durumda hesaplanır. Bu özellik, bazı çarpım momentlerinin hesaplanmadığı durumlar için de geçerlidir.
4. Çarpım momentlerinde bir sınırlama yoktur. L-moment oranları -1 ile 1 arasında değiştiğinden doğal bir sınıra sahiptir. Bu sınırlama, bu değerlerin yorumlanmasını kolaylaştırır.
5. Çarpım momentlerinde cebirsel sınırlamalar vardır. n örneklemin uzunluğu, g çarpıklık katsayısı ve k basıklık katsayısı olmak üzere; $|g| \leq n/2$ ve $k \leq n+3$ olmalıdır. Oysa L-momentlerde örnekleme bağlı bu türlü sınırlamalar yoktur. Popülasyon değerlerinin alabileceği tüm değerleri örneklem L-momentleri de alabilir.
6. L-momentlerin aksine çarpım momentleri dağılımın uçlarına daha fazla ağırlık verirler ve uçlardaki gözlemlerden daha fazla etkilenirler.
7. Klasik tekniklerle kıyaslandığında, L-momentler daha fazla sayıda dağılımın parametrelerinin bulunmasında kullanılabilir.
8. L-momentler bir örneklemeden tahmin edildiğinde, örnekleme bulunan uç değerlere karşı daha doğru ve etkin sonuçlar verir.
9. L-momentler kullanılarak elde edilen dağılım parametreleri küçük örneklerde genellikle daha doğru sonuçlar verir.
10. L-momentler, verilerin doğrusal fonksiyonları oldukları için örneklemin değişkenliğinin etkisi fazla değildir (Hosking 1990).
11. L-momentler bölgeselleştirme tekniklerinde kolaylıkla kullanılır. L-momentler tekniği ilgili istasyonlardan bölgesel parametrelerin elde edilmesi için en üstün tekniktir (WMO 1989).
12. L-momentleri temel alan analizler, model varsayımlarından sapmalar olduğunda ve/veya uygun dağılım fonksiyonunun seçilmediği durumlarda daha etkin sonuçlar vermektedir.

OLASILIK DAĞILIMLARI

İki ve Üç Parametrelili Log-Normal Dağılımı

Simetrik bir yoğunluk fonksiyonuna sahip olan standart normal dağılım, olasılık dağılımlarının bir ana modelidir. Merkez limit teoremine göre, bir x değişkeninin normal dağılmış olmasının kabul edilmesi için, x değişkenini etkileyen çok sayıda küçük etkenlerin birbiriyle toplanacak şekilde etkilenmemelidir. Benzer şekilde, eğer x , birçok küçük etkenlerin sonucu eşitse, $\ln x$ normal dağılmış olarak kabul edilebilir. Bu $Y = \ln x$ 'in hesaplanmasıyla, $Y = \ln(x_1, x_2, \dots, x_n) = \ln x_1 + \ln x_2, \dots, \ln x_n$ şeklinde gerçekleştirilir. x_i 'ler rastgele değişkendirler ve $\ln x_i$ 'ler de rastgele değişkendirler (Haan 1977).

Log-normal dağılımın analitik ifadesi:

$$f(x) = \frac{1}{\sigma\sqrt{2\pi}} e^{-1/2[(x-\mu)/\sigma]^2} \quad -\infty < x < \infty \quad (55)$$

eşitlikteki σ ve μ sırasıyla, logaritması alınmış serinin standart sapması ve ortalamasıdır. Yıllık taşkın piklerine Log-normal dağılım uygulama sebebi, yıllık taşkın serilerinin genellikle sağa çarpık olan histogramlara sahip olmaları ve değişkenlerin logaritmaları alındığında frekans diyagramının daha dar ve simetrik halle dönüşmesindendir (Bayazıt 1981). Log-normal dağılımına uyan bir x serisinin üç parametrelilik olasılık yoğunluk fonksiyonu,

$$f(x) = \frac{1}{\sqrt{2\pi} \cdot a(x-c)} \cdot \exp\left(-\frac{\left(\frac{\ln(x-c)}{b}\right)^2}{2a^2}\right) \quad (56)$$

ilişkisi ile ifade edilir. Çarpıklık katsayısının sıfırdan büyük olduğu durumda, rastgele değişkenin alt sınırı biçim parametresi c'dir. Bundan dolayı örnekteki en küçük elemandan daha küçük değere sahip olur ve iki parametrelilik Log-normal dağılımına bağlı olarak,

$$y = \ln(x - c) \quad (57)$$

şeklinde ifade edilir. Çarpıklık katsayısının sıfırdan küçük ise c, örnekteki en büyük değerlerden daha büyüktür ve

$$y = \ln(c - x) \quad (58)$$

şeklinde tanımlanır. Ancak gözlenen yıllık taşkın pik serilerinin çoğunun pozitif çarpıklık göstermelerinden dolayı, genellikle çarpıklık katsayısının sıfırdan büyük olduğu durumdaki formülü geçerlidir. Buna göre dağılımın ortalaması, standart sapması ve diğer parametreleri;

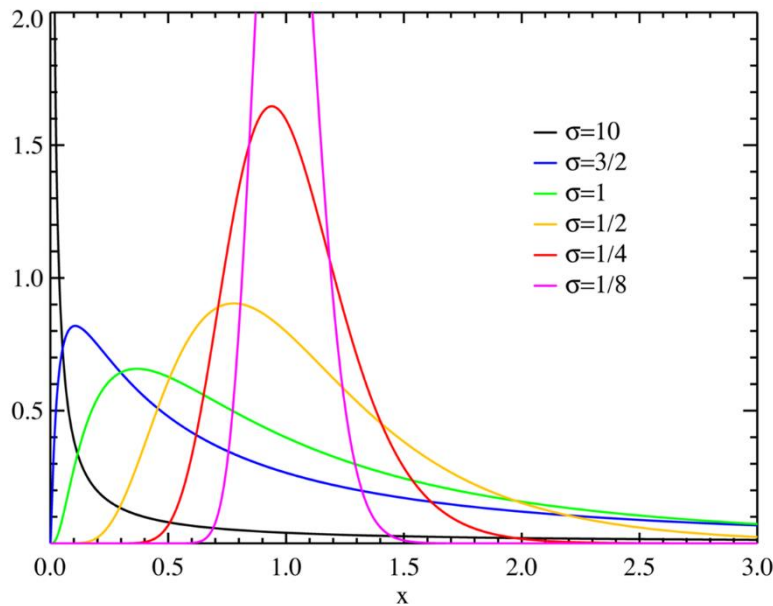
$$\bar{y} = \frac{\sum y_i}{n} \quad (59)$$

$$S = \sqrt{\frac{\sum (y_i - \bar{y})^2}{(n-1)}} \quad (60)$$

$$b = \exp(\bar{y}) \quad (61)$$

$$a = S_y \quad (62)$$

eşitlikleriyle bulunabilir.



Şekil 2. Log-normal dağılımı

Gumbel Dağılımı

Çeşitli örneklerin bir dizi maksimum veya minimum dağılımını modellemek için kullanılır. Maksimum gözlemlenen değerlerin belirli bir yıl içinde maksimum düzeyde dağılımını göstermek için kullanılır. Bu Dağılım deprem, sel veya diğer doğal afetlerin meydana gelme olasılığını tahmin etmede kullanılabilir. (Kumanlıoğlu ve Ersoy 2018).

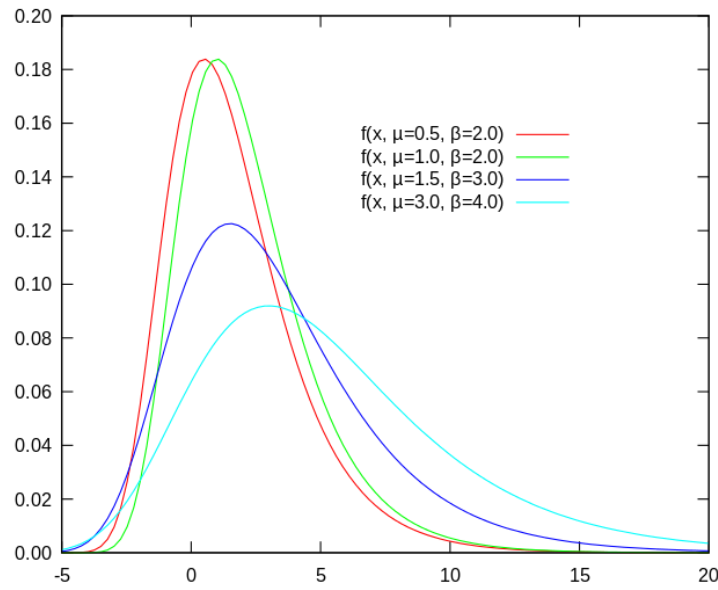
Gumbel dağılımının olasılık yoğunluk fonksiyonu,

$$f(x) = a \cdot e^{[-a(x-\beta)-a(x-\beta)]} \quad (63)$$

olarak ifade edilirken, eklenik olasılık fonksiyonu da,

$$f(x) = e^{-e^{-a(x-\beta)}} = \exp\{-\exp[-a(x-\beta)]\} \quad (64)$$

eşitliğiyle ifade edilebilir. a ve β dağılımın ölçek ve yer parametreleridir.



Şekil 3. Gumbel Dağılımı

Gamma Dağılımı

Bu dağılım iki parametrelili sürekli bir olasılık dağılımıdır. Gamma dağılımının olasılık yoğunluk fonksiyonu,

$$P(x) = \frac{1}{\Gamma(a)} x^{a-1} e^{-x} \quad x \geq 0 \quad (65)$$

şeklindedir. Burada $\Gamma(a)$ tablolatırılmış gamma fonksiyonu olup $a > 0$ için tanımlanır. Dağılımın özellikleri;

$$Ex = a \quad (66)$$

$$Var_x = a \quad (67)$$

$$Cs = \frac{2}{\sqrt{a}} \quad (68)$$

$$Ek = \frac{6}{a} \quad (69)$$

yukarıda verilen formüllerde x yerine $\beta > 0$ olmak üzere x/β yerleştirilirse 2 parametrelili damma dağılımının ihtimal yoğunluk fonksiyonu,

$$P(x) = \frac{1}{\beta^a} x^{a-1} e^{-x/\beta} \quad x \geq 0 \quad (70)$$

elde edilir. Bu dağılımın özellikleri;

$$Ex = a \cdot \beta \quad (71)$$

$$Var_x = a \cdot \beta^2 \quad (72)$$

$$Cs = \frac{2}{\sqrt{a}} \quad (73)$$

$$Ek = \frac{6}{a} \quad (74)$$

ilk denklemde yapılan gibi x yerine $(x-x_0)/\beta$ yerleştirilirse 3 parametrelili gamma dağılımı elde edilir (Pearson-3 dağılımı):

$$P(x) = \frac{1}{\beta^a} (x - x_0)^{a-1} e^{-(x-x_0)/\beta} \quad x \geq x_0 \quad (75)$$

Bu dağılımın özellikleri;

$$Ex = x_0 + a \cdot \beta \quad (76)$$

$$Var_x = a \cdot \beta^2 \quad (77)$$

$$Cs = \frac{2}{\sqrt{a}} \quad (78)$$

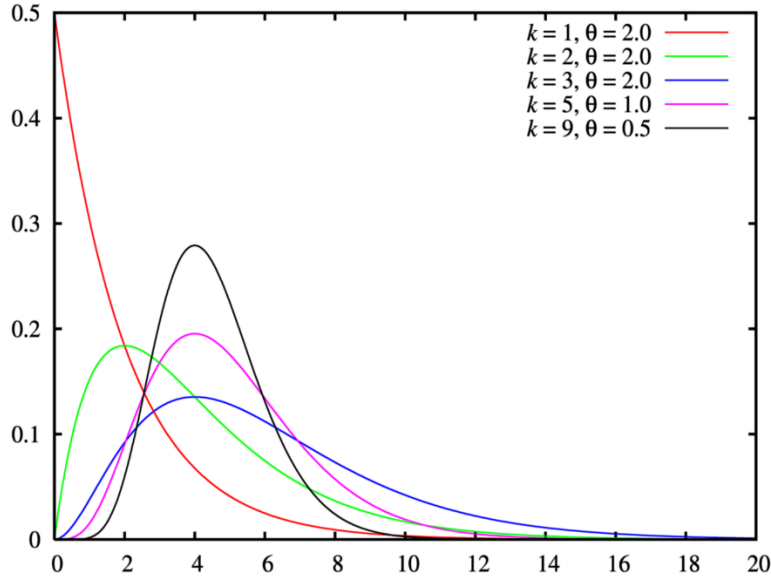
$$Ek = \frac{6}{a} \quad (79)$$

Bir parametrelili gamma dağılımının α parametresi;

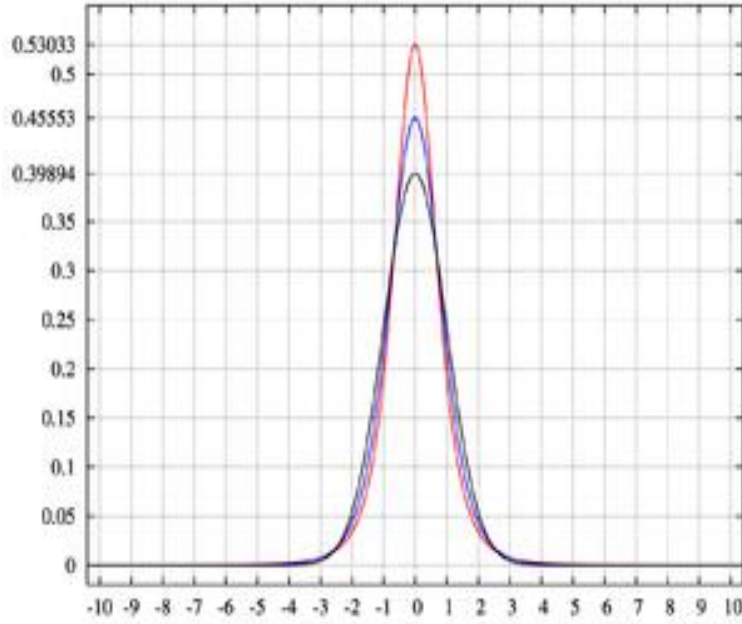
$$L = \prod_{i=1}^N p(x_i; \alpha) = \prod_{i=1}^N \frac{1}{\Gamma(a)} x_i^{a-1} e^{-x_i} = [\Gamma(a)]^{-N} \prod_{i=1}^N x_i^{a-1} e^{-x_i} \quad (80)$$

$$\frac{d(\ln L)}{da} = -N \frac{d[\ln \Gamma(a)]}{da} + \sum_{i=1}^N \ln x_i = 0 \quad (81)$$

eşitliklerle elde edilir.



Şekil 4. Gamma Dağılım



Şekil 5. Pearson-3 Dağılımı

Pareto Dağılımı

Büyük debilerin dağılımlarının büyük dönüş aralıklarında genelleştirilmiş Pareto dağılımına uyduğu kabul edilmektedir (Smith 1989). Yönteme göre; bir istasyonda, u belli bir aşılma olasılığına karşı gelen alt sınıryken w ise üst sınırı göstermek üzere u debisini aşma koşuluyla Q debisinden küçük kalma koşullu olasılığı:

$$F(x|u) = p(X \leq Q|x > u) = p(X \leq u + x|uX > u) \quad (82)$$

$$F(x|u) = (F(u + x) - \frac{F(u)}{1 - F(u)}) \quad (83)$$

p_0 'ın büyük değerleri için yukarıdaki ifade genelleştirilmiş Pareto dağılımına eşit olur ve olasılık fonksiyonu aşağıdaki eşitlikle ifade edilir:

$$F(x|u) = F(x|s, k) \quad (84)$$

Genelleştirilmiş Pareto dağılımının eklenik dağılım fonksiyonu,

$$F\{x|s, k\} = 1 - \sqrt[k]{1 - k \cdot s^{-1} \cdot x} \quad , k \neq 0 \text{ için} \quad (85)$$

$$F\{x|s, k\} = 1 - \exp(-s^{-1} \cdot x) \quad , k = 0 \text{ için} \quad (86)$$

Burada, k dağılımın biçim parametresi ve s ise ölçek parametresidir. Dağılımda k'nın bölge içinde değişmediği s'in ise çeşitli havza özelliklerin bir fonksiyonu olduğunu kabul edilmektedir. Bu durumda, k=0 ise 0=x ve k>0 ise 0=x=k⁻¹s dir. Sınır değerleri olmak üzere, dağılımın olasılık yoğunluk fonksiyonu;

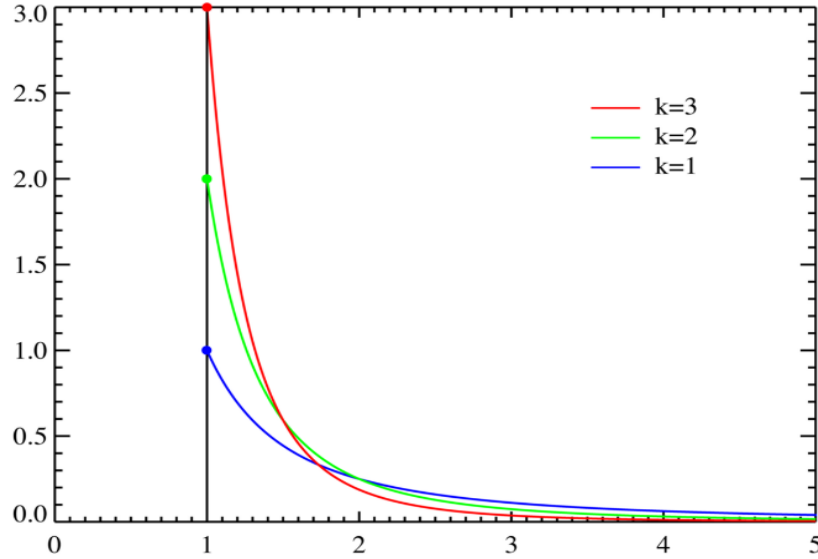
$$f\{x|s, k\} = s^{-1} \cdot \frac{(k-1)}{x} \sqrt[k-1]{1 - k \cdot s^{-1} \cdot x} \quad , k \neq 0 \text{ için} \quad (87)$$

$$f\{x|s, k\} = s^{-1} \cdot \exp(-s^{-1} \cdot x) \quad , k = 0 \text{ için} \quad (88)$$

şeklinde ifade edilir. k' nın pozitif olduğu durumlarda dağılımın üstten sınırı;

$$y = u + k^{-1} \cdot s \quad (89)$$

eşitliği ile hesaplanır. k'nın sıfıra eşit olduğu durumlarda dağılımın üst sınırı üstel, k'nın negatif değerlerinde ise dağılım sınırsızdır.



Şekil 6. Pareto Dağılımı

Ekstrem Değer Dağılımı

Yıllık pik akım serilerinin örnek büyüklükleri genellikle 30 ile 10 yıl arasında olduklarından, bunlardan ekstrem değerlere ait elde edilebilecek örneklerdeki eleman sayıları da en fazla 30-10 kadar olur. Bu örneklerin frekans analizi ile bulunan frekans dağılım çizgisini uzatıp, örnek büyüklüğünden daha uzun aralıklar için yapılacak tahminlerde hatalar olabilir. Bunun için de teorik ekstrem dağılımlar kullanmak daha uygundur (Yılmaz 1995).

Taşkın frekans analizlerinde meşhur dağılım modellerinden birisi olan ekstrem değerler dağılımının eklenik olasılık fonksiyonu;

$$F(x) = \exp\{-1[1 - a(x - c)/b]^{1/a}\} \quad (90)$$

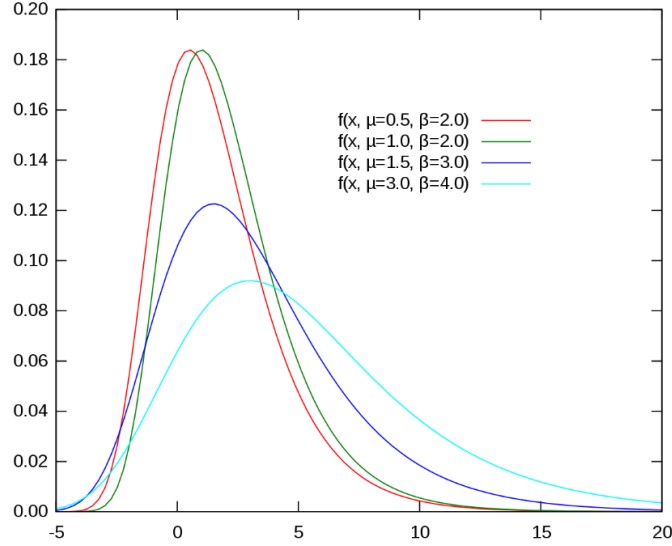
$$x_T = c + b\{1 - [-\ln(1 - 1/T)]^a\}/a \quad (91)$$

eşitliğiyle tanımlanır. Burada a yer parametresiyken b ölçek ve c biçim parametresidir. a genellikle -0.5 ve +0.5 arasında veya daha dar sınır değerler alır. b ise pozitif bir değer almak zorundadır. Ekstrem değer olasılık dağılımının 3 tipi vardır. Bunlar;

$$1.14 < G < +\infty, \quad a < 0, \quad c + b/a < x < +\infty \quad (\text{Tip 2}) \quad (92)$$

$$0 < G < 1.14, \quad a > 0, \quad -\infty < x < c + b/a \quad (\text{Tip 3}) \quad (93)$$

$$G = 1.14, \quad a = 0, \quad -\infty < x < +\infty \quad (\text{Tip 1, Gumbel}) \quad (94)$$



Şekil 7. Gumbel Dağılımı (Tip 1)

Log-Lojistik Dağılımı

Dağılım Ahmad tarafından İngiltere’de 1988 yılında taşkın frekans analizinde yeni bir yöntem olarak kullanılması önerilmiştir. Log-Lojistik dağılımının olasılık yoğunluk ve eklenik olasılık yoğunluk fonksiyonları sırasıyla,

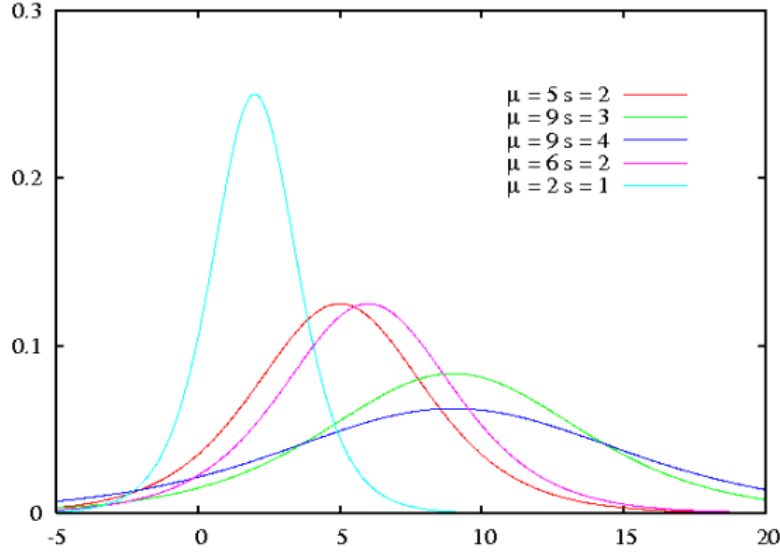
$$F(x) = \frac{\left(\frac{x-c}{b}\right)^{-1/a}}{a(x-c) \cdot \left[1 + \left(\frac{x-c}{b}\right)^{-1/a}\right]^2} \quad (95)$$

$$f(x) = \frac{1}{1 + \left(\frac{x-c}{b}\right)^{-1/a}} \quad (96)$$

eşitlikler böyleyken ileride yapılacak hesaplarda kullanılacak olan F(x),

$$X = c + b \cdot \frac{F(X=x)}{1 - F(X=x)^a} \quad (97)$$

a şekil, b ölçek, c şekil parametreleridir. Burada, $a > 0$, $b > 0$ ve $c < X_{\min}$ olmalıdır.



Şekil 8. Lojistik Dağılımı

Wakeby Dağılımı

1978 yılında Houghton tarafından daha çok bölgesel taşkın frekans analizi için tasarlanmıştır. Wakeby dağılımının olasılık yoğunluk fonksiyonu ve eklenik dağılım fonksiyonunun açık olarak ifadesi yoktur, büyüklüğü küçük kalma ihtimaline iliştiren bir analitik ifadesi vardır. 1980’li yıllarda tanınmaya başlayan Wakeby dağılımı beş parametrelidir olduğundan çeşitli olumlu eleştiriler almasına rağmen, yakın zamanlarda yapılan bazı detaylı çalışmalarda hesap yükü fazla olduğundan, üç parametrelilik kadar iyi olmadığı sonucuna varılmıştır (Haktanır ve Horlacher 1993). Greenwood vd. (1979) tarafından verilen şekliyle ters formu x,

$$x = m + a[1 - (1 - F)^b] - c[1 - (1 - F)^{-d}] \quad (98)$$

Wakeby dağılımının sol ucu b, sağ ucu ise genelde d parametresinden etkilenmektedir. Gerçekte sağ uç, d yeteri kadar büyük bir pozitif değerse b parametresinden bağımsızdır. Ancak b değeri kabul edilen b_{max} dan küçük alınarak, $(1-F)^b$ teriminin yukarıdaki denkleme katkısının tamamen yok edilmemesi uygun olmaktadır (Önöz 1994).

Wakeby dağılımı, bilinen parametre tahmin yöntemlerine iyi uyum sağlamamaktadır. Maksimum olabilirlik yöntemi çok fazla iterasyon gerektirirken momentler yöntemi de, dağılımın parametreleri momentlerin bir fonksiyonu olarak tanımlanamadığından iteratif çözümler gerektirmektedir. Bundan dolayı parametrelerin tahmini parametrelerin açıkça tanımlandığı olasılık ağırlıklı momentler yardımıyla mümkündür.

Wakeby parametrelerinin tahmininde kullanılan bu yöntem ile geçerli parametre kombinasyonları her zaman sağlanamadığı için bu gibi durumlarda çözüm bulunamamaktadır. Ancak dağılımın alt sınırının ($m=0$) bulunduğu kabulünün yapıldığı dört parametrelilik için çözümün bulunamadığı durumlar ile çok seyrek karşılaşmaktadır. Dağılımın alt sınırının sıfır olarak alınması dağılımın yapısında bir kayıp meydana getirmemektedir (Önöz 1994).

UYGUNLUK TESTİ

Uygunluk testlerinden olan ve sadece sürekli rastgele değişkenlere tatbik edilebilen Kolmogorov-Smirnov (K-S) testi ve hem sürekli hem de kesikli değişkenlere uygulanabilen Cramer Von Misses (CvM) testi, belirli bir önem seviyesinde hipotez dağılımın kabulü veya reddinin tespiti için kullanılır. Diğer bir istatistik test olan Khi-kare testi ise hem kesikli hem de sürekli rastgele değişkenlere

uygulanabilir ve diğer iki testte olduğu gibi eklenik yoğunluk fonksiyonları yerine, olasılık yoğunluk fonksiyonlarının mukayesesi esasına dayanır. Bununla beraber, özellikle küçük örnekler için, hem K-S testi hem de Khi-kare testi, gerçekte hipotez yanlışken, doğru kabul etme bakımından güçlü değildir. CvM testinin güvenilirliği ise örnek büyüklüğü 20' nin üzerinde olduğu durumlarda artmaktadır. Örneğin küçük olması durumunda histogram ile teorik ihtimal yoğunluk fonksiyonunu karşılamak yerine, örnekten elde edilen eklenik frekans dağılımını, öngörülen eklenik frekans dağılımıyla karşılaştırmak uygun olabilir (Haktanır 1990).

Kolmogorov-Smirnov Testi

Kolmogorov-Smirnov testi ile gözlenen verilerin herhangi bir dağılıma uygun olup olmadığı test edilmektedir. Bu test ile çekilen örneğin, sözü edilen dağılıma sahip ana kitleden gelip gelmediğini araştırmak için beklenen dağılımın altında geçerli olan kümülatif frekans dağılımını gözlenen frekans dağılımı ile karşılaştırmak gerekmektedir. Bu karşılaştırma sırasında iki kümülatif dağılım arasında en büyük sapma, en büyük mutlak fark olarak belirlenir. Daha sonra örnek dağılımına göre bu ölçüdeki farkın şans eseri elde edilip edilmeyeceği araştırılır (Bayazıt 2008).

$$D_{maxi} = |F(x_i) - F^*(x_i)| \quad (99)$$

Burada $F(x_i)$ seçilen dağılım fonksiyonun aynı x_i 'lere karşılık gelen ordinatları, $F^*(x_i)$ ise gözlenen örnekten hesaplanan eklenik frekans dağılımı ordinatıdır.

$$F^*(x_i) = \frac{i}{N} \quad (100)$$

D istatistiği gözlenen ve teorik eklenik dağılımların arasındaki farkların en büyüğüdür. D istatistiğinin dağılımı rastgele değişkenin dağılımından bağımsız olup sadece örnekteki N eleman sayısına bağlıdır (Uçar 2010).

Cramer-Yon Mises Testi

Bu testin uygulanması dağılımın toplam olasılık fonksiyonu ile örnek seriden pratik olarak elde edilen toplam olasılık frekans eğrisinin karşılaştırılması şeklindedir. Örnek büyüklüğünün 20'nin üzerinde olduğu durumlarda daha iyi sonuç verdiği görülmüştür (Park ve Kim 1992). Cramer-von Mises istatistiği,

$$CvM = \sum_{i=1}^M U_i \frac{(2i-1)^2}{2M} + \frac{1}{12M} (1, 2, 3, \dots, M) \quad (101)$$

eşitliğiyle hesaplanabilir. Burada uniform dağılımından elde edilen sıralı istatistik değeri olan U_i , ve β ,

$$U_i = \frac{X_i^\beta}{X_N^\beta} \quad (i = 1, 2, 3, \dots, N) \quad (102)$$

$$\beta = \frac{N}{\sum_{i=1}^M \ln\left(\frac{X_N}{X_i}\right)} \quad (103)$$

eşitliğinden elde edilir. Bulunan CvM istatistiği seçilen önem seviyesine göre Crow tarafından oluşturulan kritik tablo değerinden küçük ise dağılımın söz konusu seriye uyduğu kabul edilir (Park ve Kim 1992).

Ki Kare Testi

Gözlenen N elemanlı örnek m sınıfa ayrılarak her bir sınıftaki eleman sayısı N_i (veya O_i), seçilen dağılımın teorik eleman sayısı N_i' (veya e_i) ise, χ^2 değeri şöyle hesaplanır:

$$\chi^2_h = \sum_{i=1}^m \left[\frac{(N_i - N_i')^2}{N_i'} \right] = \sum_{i=1}^m \left[\frac{(O_i - e_i)^2}{e_i} \right] \quad (104)$$

hesaplanan χ_h^2 değeri, α anlamlılık düzeyi için, tablo değerinden küçükse veya eşitse ($\chi_h^2 \leq \chi^2$ ise) uygunluk var, büyükse uygunluk yok, $\chi_h^2 = 0$ ise tam uygunluk var demektir. χ^2 testinde bir rastgele değişkene ait N elemanlı bir örnek, m adet sınıfa ayrılır (sınıf seçimi için birden fazla formül vardır ve belirli durumlar için optimal bir seçim bulunmamaktadır) ve her bir sınıftaki eleman sayısı (N_i) hesaplanır. Seçilen olasılık yoğunluk fonksiyonuna göre aynı sınıf aralıklarında olma olasılıkları (p_i) hesaplanır. Bu sınıftaki beklenen eleman sayısı, bu olasılık değeri veri sayısı ile çarpılarak bulunur. Hesaplanan χ_h^2 değeri, α anlamlılık düzeyi için tablo değerinden küçükse ($\chi_h^2 < \chi^2$), gözlenen verilerin ilgili dağılıma uygun olduğuna karar verilir. Tablo değeri okunurken χ^2 dağılımının serbestlik derecesi m-3 olarak hesaplanır (Bayazıt 1981, Bayazıt ve Oğuz 1994, Bayazıt 1996).

TARTIŞMA VE SONUÇ

Bu araştırma sonucunda bir hidrolojik yapının planlanmasında ve taşkın yataklarındaki riskin belirlenmesinde taşkın frekans analizi önemli bilgiler sağlamaktadır. Belirli bir dönüş aralığında karşı gelen taşkın debisinin tahmini taşkın frekans analizinde kullanılan yöntemlerle belirlenmektedir. Bölgesel ve noktasal analizlerde taşkın riski hesaplanarak taşkınların toplum yaşamına ve ekonomiye verdiği zararı en aza indirmenin mümkün olduğu sonucuna varılmıştır. Ayrıca bölgesel ve noktasal frekans analizi değerlendirmesinde bölgeselleştirme işlemi heterojen bölgelerde noktasal analize göre daha güvenilir sonuçlar vermektedir.

Kaynaklar

- Anılan T. 2014. Doğu Karadeniz Havzası'nın L-Momentlere Dayalı Taşkın Frekans Analizinde Yapay Zeka Yöntemlerinin Uygulanması, Doktora Tezi, Karadeniz Teknik Üniversitesi, Fen Bilimleri Enstitüsü, Trabzon.
- Anlı A. S. 2009. Ankara'da Meydana Gelen Yağmurların L-Momentler Yöntemi İle Bölgesel Frekans Analizi, Doktora Tezi, Ankara Üniversitesi, Fen Bilimleri Enstitüsü, Ankara.
- Aydın M. 2015. Hidrolojik Verilerdeki Aykırı Değerlerin Taşkın Frekans Analizi Üzerine Etkisinin İncelenmesi: Van Gölü Havzası Örneği, Yüzüncü Yıl Üniversitesi Fen Bilimleri Enstitüsü Dergisi, 2015, 20(1-2), 47-55.
- Aydın M. 2017. Batı Akdeniz Havzası Taşkın Debilerinin L-momentler Yöntemi Ve Noktasal Taşkın Frekans Analizi İle Belirlenmesi. El-Cezerî Fen ve Mühendislik Dergisi 2018, 5(1); 117-125.
- Aydoğan, D. 2008. L-momentleri Yöntemiyle Çoruh Havzası'nın Bölgesel Frekans Analizinin Yapılması, Yüksek Lisans Tezi, Karadeniz Teknik Üniversitesi, Fen Bilimleri Enstitüsü, Trabzon.
- Bayazıt M. ve Önöz B. Taşkın ve Kuraklık Hidrolojisi. Nobel Yayın Dağıtım. 259 s. Ankara. 2008.
- Bayazıt, M. 1981. Hidrolojide İstatistiksel Metotlar, İTÜ Matbaası, İstanbul.
- Bayazıt, M. ve Oğuz, B. 1994. Mühendisler için İstatistik. Birsan Yayınevi, İstanbul, 197s.
- Benjamin, J. R. and C. A. Cornell, 1970. Probability, Statistics, and Desicion for Civil Engineers. McGraw-Hill Bode Company, New York.
- Büyükkaracıoğlu N. 2009. Türkiye Akarsuları Yıllık Pik Akım Serilerinin Frekans ve Değişiklik Analizi, Doktora Tezi, Selçuk Üniversitesi, Fen Bilimleri Enstitüsü, Konya
- Greenwood, J.A., Landwehr, J.M. and Matalas, N.C., 1979. Probability Weighted Moments: Definition and Relation of Parameters of Several Distributions Expressible In Inverse Form. Water Resources Research, 15: 1049-1054.
- Haan, C.T. 1977. Statistical Methods in Hydrology. The Iowa State University Press, Ames, Iowa, USA, 377.

- Haktanır, T and Bozduman, A. 1995. A study on Sensitivity of the ProbabilityWeighted Method on the Choice of the Plotting Position Formula. *Journal of Hydrology* 168: 265-281.
- Haktanır, T. 1990. A Few Distributions Complied Together For Flood Frequency Analysis. *Doğa Dergisi*, 14: 146-165.
- Haktanır, T. 1991b. Pratical Computation of Gamma Frequency Factors. *Hydrological Sciences* 36,6: 559-610.
- Haktanır, T. 1997. Self Determined Probability-Weighted Moments Method and its Application to Various Distributions, *Journal of Hydrology*. 97 May: 34-60.
- Haktanır, T. and Horlacher, H.B. 1993. Evaluation of Various Distributions for Flood Frequency Analysis. *Hydrological Sciences Journal* 136 (1-4), 15-32.
- Hosking J.R.M. and Wallis, J.R., “Regional Frequency Analysis An Approach Based on L-moments”, Cambridge University Press, 1997.
- Hosking, J. R. M. 1990. L-moments: Analysis and estimation of distributions using linear combinations of order statistics. *Journal of the Royal Statistical Society. Series B* 52(1):105-124.
- Hosking, J. R. M. 2005. Fortran routines for use with the method of L-moments, Version 3.04. Research Report RC 20525, IBM Research Division, T.C. Watson Research Center, Yorktown Heights, N.Y.
- Howitt, R. and Reynaud, A., 2003. Spatial Dissagregation of Agricultural Data Using Maximum Entropy, *European Review of Agricultural Economics*, 30(3): 1-29.
- Kalaycı, S. 2003. Türkiye’ deki Nehir Debisi Değerlerinin Değişkenlik Analizi, Doktora Tezi, Selçuk Üniversitesi Fen Bilimleri Enstitüsü, Konya.
- Koutrouvelis, I.A. and Canavos, G.C. 2000. A Comparison of Moment-Based Methods of Estimation for The Log Pearson Type 3 Distribution. *Jornal of Hydrology* 234: 71-81.
- Kumanoğlu, A. A. ve Ersoy S.B. 2018. Akım Gözlemi Olmayan Havzalarda Taşkın Akımlarının Belirlenmesi: Kızıldere Havzası. *Dokuz Eylül Üniversitesi-Mühendislik Fakültesi Fen ve Mühendislik Dergisi Cilt 20, Sayı 60, Eylül, 2018.*
- Önöz, B. 1994. Yeni Bir Parametre Tahmin Yöntemi Olasılık Ağırlıklı Momentler Yöntemi. *DSİ Teknik Bülteni*, 81: 49-54.
- Park, W.J. and Kim, Y.G. 1992. Goodness-of-fit Tests for The Power-Law Process. *Institute of Electrical and Electronics Engineers* 41: 107-111, New York, NY, ETATS-UNIS.
- Rao, A. R. and Hamed, K. H. 1997. Regional frequency analysis of Wabash River flood data by L-moments, *Journal of Hydrologic Engineering*. 2, 169-179.
- Seçkin N. ve Topçu E. 2016. Adana ve Çevre İllerde Gözlenen Yıllık Maksimum Yağışların Bölgesel Frekans Analizi, *Journal of the Faculty of Engineering and Architecture of Gazi University*, 2016, 31(4), 1049-1062.
- Seçkin N. ve Yurtal R. L-Momentlere Dayalı Gösterge-Sel Metodu İle Bölgesel Taşkın Frekans Analizi. *Ç.Ü Fen Bilimleri Enstitüsü*, 2008. Cilt:19-4.
- Seçkin, N. 2009. L-Momentlere Dayalı Gösterge-sel Metodu ile Bölgesel Taşkın Frekans Analizi, Doktora Tezi, Çukurova Üniversitesi, Fen Bilimleri Enstitüsü, Adana.
- Smith J.A. 1989. Regional Flood Frequency Analysis Using Extreme Order Statistics of the Annual Peak Records. *Water Resources Research*, 25: 3113-17.
- Şorman, A. Ü. 2004. Bölgesel frekans analizindeki son gelişmeler ve Batı Karadeniz’de bir uygulama. *İMO Teknik Dergi*, Cilt 15, Sayı 2.
- Şorman, A. Ü. ve Okur, A. 2000. L-moment tekniği kullanılarak noktasal ve bölgesel frekans analizinin uygulanması. *İMO Teknik Dergi*, Cilt 11, Sayı 3.
- Uçar İ. 2010. Trabzon Değirmendere Havzası’nda Coğrafi Bilgi Sistemleri ve Bir Hidrolik Model Yardımıyla Taşkın Analizi Yapılması.” *Gazi Üniversitesi, Fen Bilimleri Enstitüsü, Yüksek Lisans Tezi*, 158s, Ankara. 2010.
- Yılmaz, T. 1995. Hidrolojide İstatistik Yöntemler Yüksek Lisans Ders Notları, Selçuk Üniversitesi, Konya.
- Yüksek, Ö., Kankal, M. Üçüncü, O. 2013. Assessment of Big floods in the Eastern Black Sea Basin of Turkey, *Environmental Monitoring and Assessment*, 185:797–814.

Çıkar Çatışması

Yazarlar çıkar çatışması olmadığını beyan etmişlerdir.

Türkiye ve AB Ülkelerinin Covid-19'a Karşı Etkinliğinin Veri Zarflama Analizi ile Karşılaştırılması

Gonca ÇETİNKAYA EROĞLU^{1*}, Hasan BAL¹

¹Gazi Üniversitesi

*Sunucu yazar e-posta: goncacetinkaya@hotmail.com

Özet

2020 yılının başında, tüm halk sağlığı camiasının onlarca yıldır korktuğu senaryo gerçek olmuş ve hayvanlardan insanlara bulaşarak hastalığa neden olan korona virüslerin yeni bir tipi Çin’ den tüm dünyaya hızlıca yayılmaya başlamıştır. Dünya Sağlık Örgütü tarafınca “Covid-19” olarak adlandırılan virüs 24 Ocak 2020 tarihinde Avrupa kıtasına ulaşmış ülkemizde ise ilk vaka 11 Mart 2020 tarihinde görülmüştür. Bu çalışmanın amacı Aralık 2019’da Çin’in Wuhan şehrinde ortaya çıkan ve kısa sürede dünyayı saran COVID-19 salgını ile mücadelede Avrupa Birliği ülkelerinin ve Türkiye’nin etkinlik düzeyinin Veri Zarflama Analizi ile araştırılmasıdır. Sağlık sistemlerinin performans ölçümünde sıklıkla kullanılan Veri Zarflama Analizi, pek çok girdi ve çıktının kullanılabildiği parametrik olmayan etkinlik ölçüm yöntemidir. Analizde girdi değişkenleri olarak; onbin kişi başına doktor, hemşire, hastane ve yoğun bakım yatak sayıları ile sağlık harcamalarının Gayri Safi Yurtiçi Hâsıla (GSYH) içindeki oranı, çıktı olarak; ülkelerin 04 Mart 2021 tarihine ait milyon kişi başına Covid-19 test sayısı, vaka sayısı ve ölüm sayısı ele alınmıştır. Etkinlik analizinde sabit ölçekli Charnes, Cooper, Rhodes (CCR) ve değişken ölçekli Banker, Charnes, Cooper (BCC) yöntemleri kullanılmış, etkin olan ülkelerin etkinlik skorları belirlenmiş ve etkin olmayan ülkeler için potansiyel iyileştirme önerilerinin geliştirilmesi amaçlanmıştır.

Anahtar kelimeler: Covid-19, Veri zarflama analizi, Avrupa birliği

Comparison of the Effectiveness of Turkey and EU Countries Against Covid-19 with Data Envelope Analysis

Abstract

At the beginning of 2020, the scenario that the entire public health community feared for decades came true, and a new type of corona viruses that cause disease by transmitting from animals to humans began to spread rapidly from China to the whole world. The virus, called "Covid-19" by the World Health Organization, reached the European continent on January 24, 2020, and the first case was seen on March 11, 2020 in our country. The purpose of this study is to investigate the efficiency level of the European Union countries and Turkey in the fight against the COVID-19 pandemic, which emerged in Wuhan, China in December 2019 and spread around the world in a short time, using Data Envelopment analysis. Data Envelopment Analysis, which is frequently used in the performance measurement of health systems, is a non-parametric efficiency measurement method where many inputs and outputs can be used. As input variables in the analysis; The number of doctors, nurses, hospitals and intensive care beds per ten thousand people, and the ratio of health expenditures to Gross Domestic Product (GDP), as output; COVID-19 pandemic data of countries from March 04, 2021, Number of Tests, Number of Cases and Number of Deaths per million population were used. In the study, output oriented Charnes, Cooper and Rhodes (CCR) and Banker, Charnes and Cooper (BCC) methods were used, scale efficiency scores were determined, and it was aimed to develop potential improvement suggestions for inefficient countries.

Key words: Covid-19, Data envelopment analysis, European union

Logistic Regression Analysis in Health Sciences Researches: An Application in SPSS

Adnan ÜNALAN^{1*}

¹Niğde Ömer Halisdemir University, Faculty of Medicine, Department of Biostatistics, 51240, Niğde, Turkey

*Presenter author e-mail: unalanadnan@gmail.com

Abstract

In a significant part of health sciences research, it is aimed to accurately reveal the relationship between some characteristics measured or defined from the research units and also considered as dependent (result) and independent (reason) variables. Correlation and/or regression analysis methods are used for this purpose. While correlation analysis gives the direction and degree of the linear relationship between the variables examined, regression analysis also provides the opportunity to see the estimated values of the dependent (result) variable for certain values of the independent (reason) variable with the help of the regression model to be created. One of the variables to be included in the regression models are used as dependent variable and one/some as independent variable/s. Today, there are many regression analysis methods developed for different purposes. In the selection of the regression analysis method for the analysis of the data, the number of dependent and independent variables (such as one dependent and one independent, one dependent and more than one independent), the type of variables (categorical or numerical, discrete or continuous) and whether the data show a normal distribution or not which are the determining factors. If both of the dependent and independent variables are continuous numerical variables with normal distribution and if more than one independent variable is taken and there is no multicollinearity between these independent variables, simple linear regression or multiple linear regression analysis methods are used; If the dependent variable is categorical or discrete (binary, multiple/multinomial, ordinal etc.), logistic regression analysis (binary, multiple/multinomial, ordinal etc.) methods are used. In this study, SPSS application of binary logistic regression analysis method was given on an appropriate data set related to health sciences and the points to be considered in the interpretation of the analysis results are emphasized.

Key words: Health sciences researches, Binary logistic regression, An application in SPSS

TOPSIS Yöntemi ile Öznitelik Seçimi

Serkan AKOGUL^{1*}

¹Pamukkale Üniversitesi, Fen Edebiyat Fakültesi, İstatistik Bölümü, 20000, Denizli, Türkiye

*Sunucu yazar e-posta: sakogul@pau.edu.tr

Özet

Günümüzde, verilerdeki boyut sayısının artması, öznitelik (değişken) seçim yöntemlerine olan ilgiyi de arttırmıştır. Küçük boyutlu verilerde bile, sınıflandırma ve kümeleme sonuçlarının yorumlanmasını kolaylaştırmak için öznitelik seçiminin önemli olduğu savunulmaktadır. Bu çalışmada, öznitelik seçim problemi çok kriterli karar verme (ÇKKV) problemi olarak düşünülmüş olup bazı filtreleme tabanlı öznitelik seçim algoritmaları ile karar matrisi oluşturulmuş. ÇKKV, birden fazla karar kriterin optimizasyonu ile alternatif çözümlerin seçimi, sıralanması veya sınıflanmasını sağlayan yöntemleri içerir. Bu çalışmanın amacı, ÇKKV yöntemlerinden birisi olan TOPSIS (Technique for Order Preference by Similarity to Ideal Solution) yöntemi kullanarak özniteliklerin sıralanması ve seçimi için yeni bir yaklaşım önermektir. Önerilen yöntemin uygulaması, ücretsiz ve açık kaynak veri analizi programı ve istatistiksel programlama olan RStudio da gerçekleştirilmiştir.

Anahtar kelimeler: Çok kriterli karar verme (ÇKKV), Filtreleme tabanlı öznitelik seçimi, Öznitelik (değişken) seçimi, TOPSIS.

**Evaluation of 366 Ridge Parameter Estimators in Terms of Their Normal Distribution,
Generating Values of 0-10 and MSE Criteria**

Selman MERMI^{1*}, Özge AKKUŞ¹, Atila GÖKTAŞ¹

¹Muğla Sıtkı Koçman University, Faculty of Science, Department of Statistics, 48000, Muğla, Turkey

*Presenter author e-mail: selmanmermi@mu.edu.tr

Abstract

Ridge regression is a method developed first by Hoerl and Kennard (1970) as an alternative to the Ordinary Least Squares in the case of multicollinearity problem. Many ridge parameter estimators have been suggested for ridge regression analysis which plays a crucial role in correctly estimating the coefficients of the multiple linear regression model. In most research, ridge estimators' performance was evaluated solely in terms of the Mean Square Error of regression coefficients and in a limited scope in terms of the number of independent variables (p), sample size (n), correlation coefficient among the variables (ρ), and variance of the random error (σ^2). In this study, the performances of 366 different ridge estimators encountered in the literature so far have been evaluated based on the study of Mermi et al. (2021), as well as the MSE criterion, in terms of the normal distribution of the values they produced and being between 0-10. For this purpose, a Monte Carlo simulation study with $N=10000$ replications, in which p , n , ρ and σ^2 are handled in a wide range, was carried out. The results showed that only a few of them have been an effective performance in terms of all the criteria.

Key words: Monte Carlo simulation, Multicollinearity, Ridge parameters, Ridge regression

İşgücü Piyasası Güvensizliğini Etkileyen Faktörler: OECD Üyesi Ülkeler Üzerine Bir Değerlendirme

Emrah AKDAMAR^{1*}

¹Bandırma Onyedü Eylül Üniversitesi, Denizcilik Fakültesi, Denizcilik İşletmeleri Yönetimi, 10200, Balıkesir, Türkiye

*Sunucu yazar e-posta: eakdamar@bandirma.edu.tr

Özet

İşgücü piyasası güvensizliği, bir ülkenin iş gücü performansının belirlenmesinde oldukça önemlidir. Bu amaçla hesaplanan işgücü piyasası güvensizliği göstergesi, işsiz kalan bireylerin kazançlarındaki beklenen kayıpları ve ülkelerin bu kayıplara karşı alabildiği önlemleri yansıtır. Bu çalışmada, işgücü piyasası güvensizliğini etkileyen diğer işgücü piyasası faktörleri OECD üyesi ülkeler üzerinden araştırılmış ve ülkeler bu faktörler bağlamında irdelenmiştir. Çalışmada, işgücü piyasası güvensizliğini etkileyen faktörleri araştırmak amacıyla çoklu doğrusal regresyon analizi, anlamlı bulunan faktörler bağlamında ülkeleri irdelenmek için ise çok boyutlu ölçekleme analizi kullanılmıştır. Çoklu doğrusal regresyon analizinde, işgücü piyasası güvensizliği göstergesi bağımlı değişken, uzun dönem işsizlik oranı ve istihdam oranı ise bağımsız değişken olarak ele alınmıştır. Söz konusu değişkenlere ilişkin veriler OECD tarafından geliştirilen “Better Life Index” veri tabanından elde edilmiştir. Çalışmada, verilerine ulaşılamayan, Şili, Kolombiya, Norveç, İsviçre ve Litvanya analiz dışı bırakılmış ve toplam 37 OECD üyesi ülke arasından 32’si analize dahil edilmiştir. Bulgular, uzun dönem işsizlik oranı değişkeninin iş gücü piyasası güvensizliği üzerinde pozitif yönlü ($\beta = 1.559$, $p < 0.01$) ve istihdam oranı değişkeninin iş gücü piyasası güvensizliği üzerinde negatif yönlü ($\beta = -0.171$, $p < 0.05$) ve istatistiksel olarak anlamlı etkiye sahip olduğunu göstermektedir. Ayrıca, söz konusu bağımsız değişkenlerin, işgücü piyasası güvensizliğindeki değişimin %82’sini açıkladığı belirlenmiştir. Ülkeler, işgücü piyasası güvensizliği üzerinde anlamlı etkisi olduğu görülen uzun dönem işsizlik oranı ve istihdam oranı göstergeleri bağlamında çok boyutlu ölçekleme analizi ile görselleştirilmiştir. Elde edilen bulgular, Yunanistan’ın tüm ülkelerden olumsuz anlamda ayrıştığını, İspanya, İtalya ve Türkiye’nin de olumsuz ayrışan ülkeler grubunda yer aldığını göstermektedir. Bu ülkelerden Yunanistan, hem uzun dönem işsizlik hem de istihdam oranı bakımından olumsuz anlamda aykırı değer olarak görülmekte iken; İspanya uzun dönem işsizlik bakımından, Türkiye ise istihdam oranı bakımından olumsuz anlamda aykırı değer olarak görülmüştür. Her ne kadar iş gücü piyasası güvensizliği göstergesi, devletin işsiz kalan bireylere olan/olmayan desteği çerçevesinde tanımlansa da; Bulgular, istihdam oranını arttırmanın ve uzun dönem işsizliği azaltmanın iş gücü piyasası güvensizliğini açıklamada ve kontrol etmede oldukça önemli olduğunu göstermektedir.

Anahtar kelimeler: İşgücü piyasası güvensizliği, OECD, Çoklu doğrusal regresyon analizi, Çok boyutlu ölçekleme analizi

Information Entropies Which Are Based on Distance Matrix

Bünyamin ŞAHİN^{1*}

¹Selcuk University, Science Faculty, Department of Mathematics, 42130, Konya, Turkey

*Sunucu yazar e-posta: bunyamin.sahin@selcuk.edu.tr

Abstract

The first entropy measure was defined by Shannon in 1948. Bonchev and Trinajstić introduced some entropy measures which is based on distances to interpret the molecular branching of molecular graphs. These entropy measures are based on the distance matrix. These entropies are obtained by the frequency of informations in distance matrix. We follow two ways. Firstly, we investigate the frequency of informations in full of the distance matrix. After that since the distance matrix is symmetric matrix, we use upper triangular matrix for this informations.

Keywords: Entropy, Distance, Distance Matrix, Information functional

A Survival Simulation Study on Usage of Train-Test Splitting, Cross Validation and Bagging Methods on Training Weibull, Exponential, Log-normal, Log-logistic, Quantile and Cox Regression Models

Fulden CANTAS^{1*}, İmran KURT ÖMÜRLÜ¹, Mevlüt TÜRE¹, Hakan ÖZTÜRK¹

¹Adnan Menderes University, Faculty of Medicine, Division of Biostatistics, Aydın, Turkey

*Presenter author e-mail: fulden35cantas@gmail.com

Abstract

It is of great importance to increase the performance of statistical methods, and accordingly to estimate the models with minimum error. The goal of this study is to compare the performance of dividing the data set into training-test sets, using cross validation and bagging methods in estimating the survival time of the models during the training of Weibull, exponential, log-normal, log-logistic, quantile and Cox regression models on simulated survival data with different sample sizes. In the simulation study conducted for this purpose, survival data with different sample sizes are derived in the R software language. During the training of the simulated data with Weibull, exponential, log-normal, log-logistic, quantile and Cox regression models, three methods are used: splitting the data set into training-test sets, cross validation and bagging. At the end of the 1000 times repetitive simulation, the root mean square error (RMSE) of the models is calculated to evaluate the performance of the models in predicting the survival time. The performances of the estimation models are examined by hierarchical clustering analysis. The median range of RMSE obtained by training the models after the data set is divided into training-test set is 56.67 - 129.06; 57.27 – 130.30 for cross validation and 54.94 – 135.00 for bagging. According to the results obtained, the methods that give similar results in all conditions are found as log-normal, log-logistic, quantile regression models; and also log-logistic, Cox regression models. Quantile regression is found to be the best performing method while the exponential regression model is the worst. The estimation errors obtained by splitting the data set into training-test sets and training it with log-normal, log-logistic and Cox regression models are not affected by the sample size; however; it is found that the errors of the quantile regression model increase, while the errors of the exponential model decrease by increasing sample size. When the model prediction errors obtained by applying cross validation to the data set during the training of the models are examined; as the sample size increases, the errors of the models other than quantile regression decrease; on the other hand, the estimation errors of the quantile regression model increases; however, it has been found that it gives similar results when $n=250$ and $n=500$. When the effect of the bagging algorithm on the performance of the models according to the varying sample size is examined, it is seen that the estimation errors of the Cox and quantile regression models increase as the sample size increases; it is observed that the estimation errors of other methods decrease. As a consequence; while quantile regression gives the opportunity to choose the model with the least error among the models created for each quantile value in estimating the survival time, other models reveal a single predictive model. It has been observed that the performance of quantile regression stands out in all conditions due to this flexible structure, and its performance is slightly higher when it is used together with the bagging algorithm.

Key words: Bagging, Cross Validation, Survival, Quantile, Simulation

Eğitim-Test Ayrımı, Çapraz Geçerlilik ve Bagging Yöntemlerinin Weibull, Üstel, Log-normal, Log-lojistik, Kantil ve Cox Regresyon Modellerinin Eğitilmesinde Kullanımı Üzerine Bir Sağkalım Simülasyon Çalışması

Özet

İstatistiksel yöntemlerin performanslarının artırılması, buna bağlı olarak en az hatalı modellerin tahmin edilmesi büyük önem taşımaktadır. Bu çalışmanın amacı; farklı örneklem genişliğine sahip sağkalım verilerinde Weibull, üstel, log-normal, log-lojistik, kantil ve Cox regresyon modellerinin eğitilmesi aşamasında veri setinin eğitim-test setlerine ayrılmasının, k-fold çapraz geçerlilik ve bagging yöntemlerinin kullanılmasının modellerin sağkalım süresini tahmin etmedeki performanslarının karşılaştırılmasıdır.

Bu amaca göre yapılan simülasyon çalışmasında R yazılım dilinde farklı örneklem genişliğine sahip sağkalım verileri türetilmiştir. Türetilen verilerin Weibull, üstel, log-normal, log-lojistik, kantil ve Cox regresyon modelleri ile eğitilmesi aşamasında veri setinin eğitim-test setlerine ayrılması, k-fold çapraz geçerlilik ve bagging olmak üzere üç yöntem kullanılmıştır. 1000 döngü ile yürütülen simülasyon sonunda modellerin sağkalım süresini tahmin etmedeki performanslarının değerlendirilmesinde hata kareler ortalamasının karekökü (HKOK) hesaplanmıştır. Elde edilen modellerin performansları aşamalı kümeleme analizi ile incelenmiştir.

Veri setinin eğitim-test setine ayrıldıktan sonra modellerin eğitilmesiyle elde edilen HKOK medyan aralığı 56,67 - 129,06; çapraz geçerlilik için 57,27 – 130,30; bagging için ise 54,94 – 135,00 olduğu belirlenmiştir. Bu sonuçlara göre tüm koşullarda birbirine benzer sonuç veren yöntemlerin log-normal, log-lojistik, kantil; bunun dışında log-lojistik ve Cox regresyon modelleri olduğu belirlenmiştir. Kantil regresyonun en yüksek performans; üstel regresyon modelinin ise en düşük performans sergileyen yöntem olduğu bulunmuştur. Veri setinin eğitim-test setlerine ayrılarak log-normal, log-lojistik ve Cox regresyon modelleriyle eğitilmesiyle elde edilen tahmin hataları örneklem büyüklüğünden etkilenmediği; örneklem büyüklüğü arttıkça kantil regresyon modelinin hatalarının artarken üstel modelin hatalarının azaldığı belirlenmiştir. Modellerin eğitilmesi sırasında veri setine çapraz geçerlilik uygulanmasıyla elde edilen model tahmin hataları incelendiğinde; örneklem büyüklüğü arttıkça kantil regresyon dışındaki modellerin hatalarının azaldığı; buna karşılık kantil regresyon modelinin tahmin hatalarının arttığı; ancak $n=250$ ve $n=500$ olduğunda benzer sonuçlar verdiği tespit edilmiştir. Bagging algoritmasının değişen örneklem büyüklüğüne göre modellerin performanslarına etkisi incelendiğinde ise örneklem büyüklüğü arttıkça Cox ve kantil regresyon modellerine ilişkin tahmin hatalarının arttığı; diğer yöntemlerin tahmin hatalarının ise azaldığı gözlenmiştir.

Sonuç olarak; kantil regresyon sağkalım süresinin tahmininde her bir kantil değeri için oluşturulan modeller arasında minimum hatalı olan modeli seçme imkanı sunarken, diğer modeller tek bir tahmin modeli ortaya koyar. Kantil regresyonun bu esnek yapısından dolayı her koşulda performansının öne çıktığı, bagging algoritması ile birlikte kullanılmasıyla ise performansının biraz daha yükseldiği gözlenmiştir.

Anahtar kelimeler: Bagging, Çapraz doğrulama, Sağkalım, Kantil, Simülasyon

Ölümlülüğün Ufrs 17 Üzerindeki Etkisi

Ciğdem LAZOĞLU^{1*}, Uğur KARABEY¹

¹Hacettepe Üniversitesi, Fen Fakültesi Aktüerya Bilimleri Bölümü, Beytepe/ANKARA

*Sunucu yazar e-posta: cigdemkobal@hacettepe.edu.tr

Özet

Sigorta sektöründe primler tazminat ödemelerinden önce ya da sonra gerçekleşebilir. Herhangi bir sigortalı grubuna ait maliyetler genellikle tam olarak bilinmemektedir. Bu nedenle hem primin zaman içinde nasıl değerlendirileceği hem de prime ilişkin maliyetlerin değerlendirilmesinde karmaşıklık ortaya çıkmaktadır. UFRS 17 standardı, bu karmaşıklığı ortadan kaldırmayı amaçlamaktadır. Bu çalışmada UFRS 17 standardının düzenleyici metinlerinin matematiksel ifadeleri incelenmektedir. Tayvan ve Japonya'nın 1985-2018 yılları arasındaki ölümlülük verileri kullanılmıştır. Ölümlülük olasılığı için Poisson log-bilinear Lee-Carter modeli, mortalite trendi için ARIMA modeli kullanılarak simülasyon yapılmıştır. Sözleşme sayısı iç içe (nested) binom süreci ile modellenerek hayat sigortası üzerine bir uygulama yapılmıştır. Ölümlülüğün sözleşme hizmet marjı, hasar bileşeni ve kar&zarar üzerindeki etkisi incelenmiştir.

Anahtar kelimeler: UFRS 17, Poisson log-bilinear Lee-Carter Modeli, ARIMA modeli

İnsani Gelişmişlik Endeksinin Çok Kriterli Karar Verme Yöntemlerine Dayalı Normal Karma Modeller ile Sınıflandırılması için Yeni Bir Yaklaşım

Maruf GOGEBAKAN^{1*}

¹Bandırma Onyedi Eylül Üniversitesi

*Sunucu yazar e-posta: mgogebakan@bandirma.edu.tr

Özet

İnsani Gelişmişlik Endeksi (İGE), Birleşmiş Milletler Kalkınma Programı (BMKP) tarafından ülkeleri gelişmişlik sınıflarına ayırmak amacıyla geliştirilen bir yöntemdir. Ülkelerin insani gelişmişlikte hangi kriterlere göre hangi sınıfta bulundukları önemli bir göstergedir. İGE ülkelerin Eğitim, Sağlık ve Ekonomik verilerine göre sınıflandırılmasını önermektedir. İGE veri seti kullanılarak her bir ülkenin ölçülen kriterlere (değişkenlere) göre sınıflandırılmasında oluşacak küme yapısını belirlemede Normal Karma Modeller (GMM) kullanılabilir. Bu çalışmada Birleşmiş Milletlere üye 188 ülke için HDI veri setinin Normal karma dağılımlar ile modele dayalı kümelenebilirliği önerilmiştir. Önerilen kümeleme yönteminde, çok kriterli karar verme yöntemleri (ÇKKVY) ile değişkenlerdeki heterojen yapılar ortaya çıkarılarak bir kümeleme algoritması oluşturulmuştur. Kümelerin sayısına ve yapısına etki etmeyen homojen (tek bileşenli) yapıdaki değişkenler belirlenmiş ve sınıflandırmada bu kriter etkisiz bırakılmıştır. Böylece oluşturulan endeks için belirleyici değişkenler ortaya çıkarılmıştır. HDI veri setini oluşturan değişkenlere ile 188 ülke için bileşen sayıları belirlenmiş ve bileşenlere düşen ülkelerin oluşturduğu küme sayıları için aralık tahmini yapılmıştır. GMM için bileşen sayılarına dayalı karma modeller oluşturulmuş, her bir aday modelin parametreleri Beklenti ve Ençoklama (EM) algoritması kullanılarak tahmin edilmiştir. Ençok olabilirlik kestirimi (Maximum Likelihood) ile bilgi kriterleri hesaplanmış, aday karma modeller arasından en iyi küme sayısı ve yapısı bilgi kriterlerine göre belirlenmiştir. BM tarafından çeyrekliklere göre belirlenen 4 sınıf üyelikleri ile önerilen modele dayalı kümeleme sınıfları ile elde edilen 3 sınıf arasındaki ilişki incelenmiştir. Ülkelerin değişkenler baz alınarak sınıflandırılması alternatif bir sınıflandırma yöntemi olarak önerilmiştir.

Anahtar kelimeler: *İnsani gelişmişlik endeksi (İGE), Modele dayalı kümeleme, Normal karma modeller (GMM), Çok kriterli karar verme yöntemleri, Bilgi kriterleri*

A Novel Approach for the Determination of Entropy and Compactivity of Granular Matter in a Tube

Yahya ÖZ^{1*}

¹Turkish Aerospace, R&D Directorate, Advanced Materials, Processes and Energies Technology Center, 06980, Ankara, Turkey

*Presenter author e-mail: yahya@unam.bilkent.edu.tr

Abstract

Granular matter in a tube is studied in this work. For this purpose, we consider identical disks confined in a tube that has one open and one closed end whereby all surfaces are assumed to be hard and frictionless. Thus, a one-dimensional system is studied. Jamming occurs due to the local pressure in a uniform gravitational field parallel to the tube and rotation. The framework of configurational statistics with a generalized Pauli principle is used, where jammed microstates are encoded in particle configurations while macrostates are produced by uniform arbitrary agitations and described by entropy and volume. Moreover, mass density, compactivity and population correlations of quasiparticles are studied analytically.

Key words: Entropy, Granular matter, Jamming, Quasiparticles, Statistical mechanics

Genetic Defects Associated with Reproduction in Dairy Cattle

Habiba MOHSHINA^{1*}, Yasemin GEDİK¹

¹Eskisehir Osmangazi University, Faculty of Agricultural- Institute of science, Animal Science
Department, 26160 Odunpazari, Eskisehir, Turkey

*Presenter author e-mail: habiba41712@bau.edu.bd

Abstract

Reproductive performance is a key goal in dairy cattle breeding programs since it determines milk output and provides a source of replacement animals, and it is one of the most important factors determining long-term viability of animal agriculture. In tandem with the tremendous increase in dairy cattle output reproductive performances have decreased in all of the major dairy breeds over the last sixty years, and this decline has been attributed to a number of factors including increased milk yield, improvements in management and herd structure, inbreeding, and a lack of reproduction selection. Despite the fact that reproduction is not generally thought to be a highly heritable trait, genetics has evidently played a role in the decrease in fertility. Furthermore, fertility is a complex trait determined by a number of genes. A haplotype is a group of alleles that are inherited together at close intervals on a chromosome and believed to harbour lethal recessives; in dairy cattle evidence showing these affect fertility and related to early embryonic failure. For using a narrow parent stock exclusively in most dairy cattle breeding, fertility and reproductive efficiencies are decreased. In many cases loss of functions (LOF) mutation occurs in specific breeding bulls and through artificial insemination and embryo transfer those recessive genes are spread out throughout the world. In present times haplotype related to reproduction is growing concerns to researchers in findings of reduced fertility. Here we will discuss different defects sourced from haplotypes discovered recently, causing reproductive failure in Brown Swiss, Holstein Friesian and Jersey cattle breeds, which require emergency screening among the cattle population for preventing reproductive lethal effects and economic loss.

Keywords: Reproductive efficiency, Recessive defects, Screening

Residual-based Control Charts Under Liu Estimator For Monitoring Conway-maxwell-poisson Profiles

Ulduz MAMMADOVA^{1*}, M. Revan ÖZKALE¹

¹Çukurova Üniversitesi, Fen-Edebiyat Fakültesi, İstatistik Bölümü

*Presenter author e-mail: ulduzozel@gmail.com

Abstract

A control chart is the widely used simple tool for monitoring, identification, and evaluation of shifts that may occur in the production process. In this study, we modified traditional control charts for monitoring Conway-Maxwell-Poisson profiles under multicollinearity. We proposed an iterative Liu estimator for the Conway-Maxwell-Poisson regression model to eliminate the adverse implications of multicollinearity and obtained Liu deviance residuals. Then, we defined Liu deviance residual-based Shewhart, CUSUM, and EWMA control charts to monitor the Conway-Maxwell-Poisson profiles. The proposed method is tested by using the data set from the bike rental application. The performances of Liu deviance residual-based Shewhart, CUSUM, and EWMA control charts are analyzed and, results are compared to the deviance residual-based Shewhart, CUSUM, and EWMA in terms of out-of-control average run length and first signal.

Key words: Liu estimation, COM-Poisson distribution, profile monitoring, multicollinearity, control charts

**Bağlantı Dengesizliğinin Genotipik Varyans ve Unsurlarına Etkisi Üzerine bir
Simulasyon Çalışması: I- Bir lokusta Dominans olan Modelde Populasyon
Parametreleri**

**Selma KÖKÇÜ^{1*}, Orhan KAVUNCU¹, Furkan YILMAZ², Kadir KARAKAYA³, Yunus
AKDOĞAN³, Yasemin GEDİK⁴**

¹Kastamonu Üniversitesi, Mühendislik ve Mimarlık Fakültesi, Genetik ve Biyomühendislik Bölümü.

²Tokat Gaziosmanpaşa Üniversitesi, Ziraat Fakültesi, Zootečni Bölümü, Tokat, Türkiye

³Selçuk Üniversitesi, Fen Fakültesi, İstatistik Bölümü, Konya, Türkiye

⁴Eskişehir Osmangazi Üniversitesi, Ziraat Fakültesi, Zootečni Bölümü, Eskişehir, Türkiye

*Sunucu yazar e-posta: selma14110396@gmail.com

Özet

Bağlantı dengesizliği katsayısının genetik varyans ve unsurları üzerine etkisini sayısal olarak somut bir şekilde göstermek için biri dengede diğeri dengede olmayan 2 populasyon tanımlandı ve parametreleri deterministik olarak hesaplandı. Buna göre, iki lokuslu bir modelde üç allelli birinci lokusta genler arasında dominanslık sapması yoktur. İki allelli ikinci lokusta ise bir allel diğere tam dominanttır. Hesaplamalar sonucunda parametrelerin bağlantı dengesizliğinden etkilendiği açıkça görüldü. Bağlantı dengesizliği söz konusu olduğunda eklemeli lokusta dominanslık sapması görüldüğü gibi genetik varyansta ve unsurlarında bağlantı dengesizliğinden kaynaklanan kovaryans teriminin etkisinin ortaya çıktığı görüldü. Daha sonraki çalışmalarda bağlantı dengesizliğinin bu etkileri analitik olarak çalışılacak ve genotip değerler üzerine bir çevre unsuru eklenerek simulasyonla üretilecek örneklerde varyans unsurlarının örnekleme dağılımı çalışılacaktır.

Anahtar kelimeler: Bağlantı dengesizliği katsayısı, Varyans, Genotipik varyans

CfsSubsetEval ve MARS Yöntemleri ile BİST 100 Endeksi için Model Seçimi

Dilek SABANCI^{1*}, Serhat KILIÇARSLAN²

¹Gaziosmanpaşa Üniversitesi, Fen-Edebiyat Fakültesi, Matematik Bölümü, 60250, Tokat/Türkiye

²Gaziosmanpaşa Üniversitesi, Rektörlük, Enformatik Bölümü, 60250, Tokat/Türkiye

*Sunucu yazar e-posta: dilek.kesgin@gop.edu.tr

Özet

Bu çalışmada, BİST 100 endeksi tahminlemesi için boyut indirgemesinde Korelasyon Tabanlı Özellik Altkümesi Seçimi (CfsSubsetEval) yöntemi altında Genetik (Genetic), Açgözlü (Greedy), Harmoni (Harmony) ve Parçacık Sürü Optimizasyonu (PSO) arama algoritmaları; model tahminlemesinde ise Çok Değişkenli Uyarlamalı Regresyon Şeritleri (MARS) yöntemi kullanılarak hibrit bir model önerilmiştir. Veri seti, pandemi dönemini (01 Ocak 2020 – 03 Haziran 2021) kapsayacak şekilde BİST 100 endeksini etkileyen Türk Hisseleri, Emtia ve Döviz Kuru başlıkları altından seçilen 120 özellik kullanılarak oluşturulmuştur. Arama algoritmaları ile boyut indirgeme yapılarak elde edilen alt boyutlara ait veri setine ve boyut indirgeme yapılmamış veri setine MARS yöntemi uygulanmıştır. Böylece beş farklı tahmin modeli kurularak BİST 100 endeksini etkileyen özellikler belirlenmiştir. Bu modellerden istatistiksel olarak istikrar ve doğruluk açısından en iyi olanı hata ölçüm değerlerine bakılarak incelenmiş ve belirlenmiştir.

Anahtar kelimeler: MARS, CfsSubsetEval, Özellik seçimi, Arama algoritmaları, Model seçimi

A Unit - Cauchy Distribution: Definition, Properties and Application

Ako Rajab QADER^{1,3*}, Talha ARSLAN^{1,2}

¹Department of Statistics, Institute of Natural and Applied Sciences, Van Yüzüncü Yıl University,
65080, Van, Turkey

²Department of Econometrics, Van Yüzüncü Yıl University, 65080, Van, Turkey

³School of Medicine, Koya University, Koya, Iraq

*Presenter author e-mail: ako.najar@gmail.com

Abstract

Unit distributions play a crucial role in modeling data on unit interval (0,1), defining the families of continuous distributions, and constructing a regression model for bounded response. Therefore, researchers have proposed a new unit distribution to model data and generate a new family of distributions. This study introduces a unit-Cauchy distribution by following the similar lines given in Arslan (2021). Then, the maximum likelihood method is used for estimating its parameter. The probability density function, cumulative distribution function, and hazard rate function of the unit-Cauchy distribution are obtained in closed form like the Kumaraswamy distribution. Also, the unit-Cauchy distribution has only one shape parameter different from the beta and Kumaraswamy distributions with two shape parameters. A real data set is used to show the implementation of the unit-Cauchy distribution. The result indicates that the unit-Cauchy distribution has better goodness-of-fit performance than the beta and Kumaraswamy distribution.

Key words: Beta distribution, Cauchy distribution, Kumaraswamy distribution, Maximum likelihood

On a Unit-eeg Quantile Regression Analysis

**Mohamed El Moustapha SIDI^{1*}, Coşkun KUŞ¹, M. Çağatay KORKMAZ², Ahmet PEKGÖR³,
Kadir KARAKAYA¹**

¹Selçuk University, Faculty Of Science, Department Of Statistics, 42130, Konya, Turkey

²Artvin Çoruh University, Department Of Measurement And Evaluation, 08100, Artvin, Turkey

³Necmettin Erbakan University, Faculty Of Science, Department Of Statistics, 42090, Konya, Turkey

*Presenter author e-mail: moustafabtt@gmail.com

Abstract

In this talk, we consider a quantile regression analysis. Let X be an EEG random variable and use $X/(X+1)$ transformation. Then we have a unit range distribution. This distribution is used to construct a new quantile regression analysis when the dependent data lies in the interval $(0,1)$. The maximum likelihood estimation is discussed and a new method is proposed to achieve the maximum likelihood estimates with arbitrary initial values. A numerical example is also provided.

Key words: Unit distribution, Quantile regression, Maximum likelihood estimation

Varyasyon Katsayıları Arasındaki Farka Ait Örnekleme Dağılımının Şekli ve Güven Sınırları

Rabia ALBAYRAK DELİALİOĞLU^{1*}

¹Ankara Üniversitesi, Ziraat Fakültesi, Zootekni Bölümü, Biyometri ve Genetik Anabilim Dalı, 06110, Ankara, Türkiye

*Sunucu yazar e-posta: ralbayrak@ankara.edu.tr

Özet

Üzerinde çalışılan özellikler veya gruplar değişim bakımından karşılaştırılırken genellikle standart sapmadan yararlanılmaktadır. Ancak çalışılan veri setinde grupların veya değişkenlerin ortalamaları arasında fark olduğunda standart sapma bu farktan etkileneyeceği için, standart sapma üzerinden ortalamanın etkisini kaldıran bir değişim ölçüsü olan varyasyon katsayısına bakmak çok daha güvenilir olacaktır. Ayrıca varyasyon katsayısı birimden bağımsız bir değişim ölçüsü olup değişimi % olarak ifade eder. Normal dağılım varsayımı altında varyasyon katsayısına ait güven aralığının hesaplanmasında birçok araştırmacı tarafından çalışmalar yapılmıştır. Bu simülasyon çalışmasında farklı örnek genişlikleri ve farklı dağılımlar için iki varyasyon katsayısı arasındaki farka ait örnekleme dağılımı elde edilmiş ve varyasyon katsayıları arasındaki farka ait örnekleme dağılımının şekli ile %95 ve %99 güvenilirlik ile alt ve üst güven sınırları bulunmuştur. Çalışmanın materyalini Microsoft Power Station Developer Studio ve IMSL Library yardımıyla FORTRAN'da simülasyon (benzetim) tekniği ile $Z(0,1)$ ve $\chi^2(3)$ dağılımlarından çeşitli genişliklerde ($n=3-10, 15, 20, 25, 30, 50, 100, 500$) üretilen tesadüf sayıları oluşturmaktadır. Hem $Z(0,1)$ hem de $\chi^2(3)$ dağılımları için ayrı ayrı yapılan simülasyon çalışmasında popülasyondan $n_1=n_2=n$ olacak şekilde iki örnek alınmıştır. Örneklerin Z-dağılımından alındığı durumlarda negatif ve sıfır olan değerler denemeye dahil edilmemiştir. Alınan örneklerde varyasyon katsayıları hesaplanmış ve aralarındaki fark bir dosyaya kaydedilmiştir. Bu işlem 1.000.000 defa tekrarlanmış ve böylece varyasyon katsayıları arasındaki farka ait örnekleme dağılımı elde edilmiştir. Sonra bu farklar büyüklüklerine göre sıralanmış ve %99 ve %95 güvenilirlikle alt ve üst güven sınırlarına denk gelen değerler bulunmuştur. Ayrıca hesaplanan farklara ilişkin tanıtıcı istatistikler ve histogramlar da çalışmada verilmiştir.

Anahtar kelimeler: Simülasyon, Varyasyon katsayısı, Güven sınırları

A Simulation Study on the Comparison of the Methods Used in Comparing the Proportions and Hotelling T^2 Method

Muslu Kazım KÖREZ^{1*}, Neriman AKDAM¹, Ahmet PEKGÖR², Coşkun KUŞ³

¹Selçuk University, Faculty of Medicine, Division of Biostatistics, 42130, Konya, Turkey

²Necmettin Erbakan University, Faculty of Science, Department of Statistics, 42090, Konya, Turkey

³Selçuk University, Faculty of Science, Department of Statistics, 42130, Konya, Turkey

*Sunucu yazar e-posta: mkkorez@gmail.com

Abstract

Probably the most common quantitative problem in medical research involves the question of the presence or absence of a particular attribute. Such a question requires testing the difference between two or more proportions. In order to test the difference between two or more proportions, proportion tests and chi-square tests are generally recommended. Also, in the case of paired samples, the McNemar test is used. Here, we discuss the usage of the one sample Hotelling T^2 test to compare the two or more proportions unlike traditional tests. Aim of the study is to calculate and compare the simulated significance level and the power for different test statistics by including the Hotelling T^2 for testing equality of two and more proportions. For this purpose, simulation study is performed both in the case of uncorrelated and correlated samples. It is observed that the Hotelling T^2 test is easy to use and easy to enter data, works in two or more groups, and has a good performance (stability of significance level and higher power) when the groups are correlated. Consequently, considering all possible situations, we recommend using the one sample Hotelling T^2 test instead of the tests currently used to compare proportions.

Key words: Proportion tests, Monte-Carlo simulation, Significance level, Power of tests, Hotelling T^2

**Estimation and Prediction on the Half-Normal Distribution Based on Progressively
Type-II Censored Samples**

Sümeýra SERT^{1*}, İhab A.S. ABUSAİF¹, Ertan AKGENÇ¹, Kadir KARAKAYA¹, Coşkun KUŞ¹

¹Selçuk University, Science Faculty, Department of Statistics, 42130, Konya, Turkey

*Presenter author e-mail: sumeyra.sert@selcuk.edu.tr

Abstract

In this study, the estimation and prediction problems are discussed for the half-normal distribution under a progressively Type-II censoring scheme. Several predictors and predictive intervals for removed failure times are derived. An extensive Monte Carlo simulation study is performed to discuss the mean squared error (mean squared prediction errors) and bias of estimates (predictors). The coverage probabilities and average length of the confidence and predictive intervals are simulated. A numerical example is also provided.

Key words: Confidence interval, Half-normal distribution, Predictive interval, Progressively censoring.

**Bayesian, E-Bayesian and E-MSE of Half-Normal Distribution Parameter Under
Different Loss Functions**

Sümeýra SERT^{1*}, Kadir KARAKAYA¹, Coşkun KUŞ¹

¹Selçuk University, Science Faculty, Department of Statistics, 42130, Konya, Turkey

* Presenter author e-mail: sumeyra.sert@selcuk.edu.tr

Abstract

In this study, Bayesian, E-Bayesian and E-MSE of the Half-normal distribution parameter is obtained under different loss functions. A comprehensive Monte Carlo simulation study is conducted to analyze the performance of different loss functions using different parameter settings and different sample sizes. Several prior distributions are also considered to examine the effect on E-Bayesian estimations. A numerical data analysis is also provided.

Key words: Half-normal distribution, E-Bayesian estimation, E-MSE, SELF, ELF.

Dengesiz Veri Setlerinde Farklı Dengeleme Algoritmalarının Optimum Denge Oranlarının Sınıflandırma ve Regresyon Ağaçları Yöntemi ile İncelenmesi: Simülasyon Çalışması

Hakan ÖZTÜRK^{1*}, Mevlüt TÜRE¹, İmran KURT ÖMÜRLÜ¹

¹Aydın Adnan Menderes Üniversitesi, Tıp Fakültesi, Biyoistatistik Anabilim Dalı,
09100, Aydın, Türkiye

*Sunucu yazar e-posta: hakan.ozturk@adu.edu.tr

Özet

İki sınıflı bir veri setinde, sınıflardaki örnek sayıları arasında önemli bir fark varsa, sınıf dengesizliği problemi ortaya çıkar. Sınıf dengesizliği problemine çözüm olarak çok sayıda dengeleme algoritması önerilmiştir. Literatürde, veri dengeleme algoritmalarının performanslarının değerlendirildiği birçok çalışma yapılmıştır. Bu çalışmaların çoğunda dengeleme algoritmalarının performansları sadece gerçek veri setleri üzerinden ya da gerçek veri setlerinden elde edilen dengesiz veri setleri üzerinden değerlendirilmiştir. Kullanılan performans ölçüsüne göre en yüksek değere ulaşan yöntemlerin en başarılı yöntemler olduğu ifade edilmiştir. Önceki çalışmalardan farklı olarak, çalışmamızda, 7 farklı dengeleme algoritması, orijinal bir simülasyon çalışması ışığında, hem bilinen bir toplum parametresini tahmin performansları bakımından değerlendirildi hem de optimal azınlık-çoğunluk sınıfı denge oranları belirlendi. Çalışmamızda, simüle edilen bir toplum veri setinden örneklenen dengesiz veri setleri, random over sampling (ROS), synthetic minority over sampling technique (SMOTE), majority weighted minority oversampling technique (MWMOTE), adaptive synthetic sampling approach (ADASYN), random under sampling (RUS), random under boosting (RUSBoost) ve under bagging (UB) algoritmaları ile kademeli olarak dengelendi ve her kademede sınıflandırma ve regresyon ağaçları (CART) yöntemi ile sınıflandırıldı. Dengeleme algoritmalarının performanslarının değerlendirilmesinde ROC eğrisi altında kalan alan (AUC) kullanıldı. Sonuç olarak, RUSBoost ve UB algoritmalarının, belirli denge oranları aşıldığında iyimser sonuçlar ürettiği, diğer yöntemlerin ise iyimser sonuçlar üretmediği ve en iyi sonuçları tam denge durumunda (50:50) ürettiği gözlemlendi.

Anahtar kelimeler: Dengesiz veri, Dengeleme algoritmaları, Sınıflandırma ve regresyon ağaçları, Simülasyon

Examining the Optimal Balance Ratios of Different Balancing Algorithms in Imbalanced Data Sets by Classification and Regression Trees: Simulation Study

Abstract

In a two-class dataset, the class imbalance problem occurs if there is a considerable difference between the number of samples in the classes. Many balancing algorithms have been proposed as a solution to the class imbalance problem. There have been many studies in the literature evaluating the performance of data balancing algorithms. In most of these studies, the performances of balancing algorithms were evaluated only on real data sets or on imbalanced data sets obtained from real data sets. It is stated that the methods that reach the highest value according to the performance criteria used are the most successful methods.

Differently from the previous studies, in our study, 7 different balancing algorithms were evaluated in terms of their predictive performance of a known population parameter and optimal minority-majority class distribution were determined in the light of an original simulation study. In our study, unbalanced datasets sampled from a simulated population dataset were gradually balanced with random over sampling (ROS), synthetic minority over sampling technique (SMOTE), majority weighted minority oversampling technique (MWMOTE), adaptive synthetic sampling approach (ADASYN), random under sampling (RUS), random under boosting (RUSBoost) and under bagging (UB) algorithms and classified by the classification and regression trees (CART) method at each step. The area under the ROC curve (AUC) was used to evaluate the performance of balancing algorithms.

In conclusion, it was observed that RUSBoost and UB algorithms produce optimistic results when certain balance ratios are exceeded, while other methods do not produce optimistic results and produce the best results in fully balance (50:50).

Key words: *Imbalanced data, Balancing algorithms, Classification and regression trees, Simulation*

Bootstrap Confidence Intervals of Generalized Discrete process capability index Cpyk

Yunus AKDOĞAN¹, Tenzile ERBAYRAM^{2*}

¹Selçuk Üniversitesi, Fen Fakültesi, İstatistik, 42130, Konya, Türkiye

²Selçuk Üniversitesi, Fen Fakültesi, İstatistik, 42130, Konya, Türkiye

*Sunucu yazar e-posta: tenzile.erbayram@selcuk.edu.tr

Abstract

In this paper, we consider estimation methods to estimate the parameters and the discrete process capability index Cpyk for the discrete Lindley distribution. It is well known that confidence interval provides much more information about the population characteristic of interest than does a point estimate. Hence, next, we consider six bootstrap confidence intervals namely, standard, normal, basic, percentile, student and bias-corrected percentile bootstrap for obtaining confidence intervals of discrete process capability index Cpyk using the methods of estimation. A Monte Carlo simulation has been used to investigate the estimated coverage probabilities and average widths of the bootstrap confidence intervals. Finally, real data sets are analyzed for illustrative purposes.

Key words: Coverage probabilities, process capability index, estimation, bootstrap confidence intervals, discrete Lindley distribution, Monte Carlo simulation.

Some Count Regression Models from Generalized Linear Models Approach to Investigate Pakistan Child Mortality Under-5 Data

Mohamad ALNAKAWA^{1*}, Neslihan İYİT¹

¹Selcuk University, Faculty of Science, Statistics Department, 42130, Konya, Turkey

*Presenter author e-mail: abosleem.syria@yahoo.com

Abstract

In generalized linear models (GLMs), the distribution of the response variable is considered to come from the exponential family. In this study, some count regression models such as Poisson, Negative Binomial and Geometric Regression Models as special cases of GLMs have been addressed in terms of model building process and parameter estimation to analyze the count data. In Poisson regression model, the conditional probability function of the response variable follows the Poisson distribution with the parameter λ . Negative binomial regression model is based on negative binomial distribution which is a mixture of Poisson and Gamma distributions. On the other hand, geometric regression model is a special case of negative binomial regression model when the dispersion parameter is equal to zero. Geometric regression model is also a generalization of Poisson regression model removes the requirement to satisfy the assumption that the variance is equal to the mean of the response variable in Poisson regression model. The aim of this study is to determine the possible risk factors affecting "child mortality under-5" in Pakistan selected from South Asia region lower-middle income countries by using the Multiple Indicators Cluster Survey 5 (MICS5) data conducted by UNICEF and GLMs approach in the aspect of Poisson, Negative Binomial and Geometric Regression Models.

Key words: Poisson regression, Negative binomial regression, Geometric regression, Generalized linear model, Child mortality under-5, UNICEF

Forecasting BIST 100 Index Using An Artificial Neural Network

Yüksel Akay ÜNVAN¹, Cansu ERGENÇ^{1*}

¹Ankara Yıldırım Beyazıt University, Faculty of Management, Finance and Banking, 06760 Ankara, Turkey

*Presenter author e-mail: cansuergenc7@gmail.com

Abstract

Making reliable forecasts is very important for financial analysis. For this reason, financial analysts make analyzes using different models. Financial analysts try to make the most accurate estimation in these analyzes. The artificial neural network model is a widely used method in the field of finance. In this study, BIST 100 index was estimated using artificial neural networks and regression model. By using the closing prices of the BIST 100 index between 2010 and 2020, the closing values of the BIST 100 index for 2021 were estimated. Moreover, the regression model and artificial neural network model predictions were obtained. The mean square error of the neural networks and regression model was also found. Finally, according to the result of the mean of error squares, the performance of the models was compared and seen that the artificial neural network model was better.

Key words: Artificial neural network, Regression, BIST 100

INTRODUCTION

The concept of forecasting, which is a large part of the financial and economic world, is crucial to investors and governments. Financial and economic forecasting has been around for centuries. Economic forecasting uses historical data published by countries or geographic regions in previous economic reports. (Lenel et al., 2020)

Economic forecasting is the process used to try to predict or anticipate future economic conditions by using various economic variables and indicators. Indeed, the methods for selecting a forecasting model, estimating its parameters, communicating the resulting forecasts, and evaluating their accuracy have improved in many ways over the past 20 years. Nevertheless, economic forecasting models are far from perfect, and understanding their limitations is important for interpreting economic forecasts. (Elliott and Timmermann, 2008)

The fundamental importance of economic forecasting is to help policymakers make better decisions. For example, if the economy forecast is to enter a recession, the government might consider expansionary fiscal policy (increased spending financed by borrowing) to maintain demand in the economy and prevent a severe downturn in the economy (Wieland and Wolters, 2013)

Economic forecasts are critical in setting monetary policy/fiscal policy. If the economy is indeed expected to recover, inflation may pick up and the bank may need to raise interest rates. If the economy is likely to contract further, the Bank may need to undertake further quantitative easing.

Companies also use forecasts to develop business strategies. Financial and operational decisions are made based on economic conditions and what the future will look like, even if it is uncertain. Data from the past is collected and analyzed so patterns can be found. Today, big data and artificial intelligence have changed the methods of business forecasting. Business forecasting is vital for companies as it allows them to plan production, financing, and other strategies.

There is a long history of research in finance and economic modeling. Time series analysis and regression are the most commonly used traditional approaches in financial and economic modeling. However, in some cases, these models are not sufficient. Recent evidence shows that financial markets are nonlinear; however, these linear methods still provide good opportunities to describe nonlinear systems found in time series analysis of the financial market (Maciel and Ballini, 2008). For this reason, financial analysts started to look for other models. In recent years, artificial neural networks have been widely used as a powerful modeling technique in financial and economic forecasting.

Various variables such as gold prices, bonds, exchange rates, daily US dollar returns and oil prices have been used in many different studies to predict the returns of investment instruments through artificial neural network models (Zhang and Berardi, 2001; Farahani and Mehralian, 2013; Khodayari et al., 2020; Diaz and Nguyen, 2021; Mir et al., 2021; Le et al. 2021).

There are studies that make financial forecasts by combining regression models and artificial neural network models. (Pradeepkumar and Ravi, 2017; Siami-Namini and Namin, 2018; Cao and Li, 2019). Other studies have combined time series models and artificial neural network models for financial data analysis. (Patra et al., 2017; Altan and Karasu, 2019; Benrhmach, 2020) Some studies have compared the results of the different methods, such as regression and time series models, with the neural network by using the same data. (Carvalho and Ribeiro, 2007; Kyung et al., 2008).

In the literature, there are studies that predict the index BIST 100 using the artificial neural network model and compare it with other models (Kılıc et al., 2014; Ozbey et al., 2020; Molla et al., 2020; Karakul, 2020).

In this study, the BIST 100 index was estimated by using artificial neural networks and a regression model. Using the closing prices of the BIST 100 index between 2010 and 2020, the closing values of the BIST 100 index for 2021 were estimated. Also, the predictions of the regression model and the artificial neural network model were obtained. The mean square errors of the neural networks and the regression model were also obtained. In the analysis, estimates were first made with the regression model, and then the artificial neural network model was applied.

MATERIAL AND METHOD

The regression method was proposed by the English statistician Francis Galton in the 19th century. He proposed this method for the purpose of a biological study (Galton, 1877). Udney Yule and Karl Pearson applied and developed this method to broader general statistical fields (Yule, 1897). In this period, the dependent and independent variables are assumed to have normal distributions. This assumption was extended with Fisher's publications in 1922 and 1925 to apply only to cases in which the conditional distribution of the dependent variable is normal (Fisher, 1922).

Regression analysis is a set of statistical processes for estimating the relationships between a dependent variable and one or more independent variables. In simple linear regression analysis, it is assumed that there is a linear relationship between two variables. In multiple regression analysis, it is assumed that there is a linear relationship between more than two variables (Schroeder et al., 2016).

The general linear regression model can be stated by the equation below (Montgomery, 2021).

$$y_i = \beta_0 + \beta_1 X_{1i} + \beta_2 X_{2i} + \dots + \beta_k X_{ki} + \varepsilon_i$$

where,

y_i =dependent variable

x_i =explanatory variables

β_0 = constant term

β_k =slope coefficients for each explanatory variable

ε_i =the model's error term

The multiple regression model is based on the following assumptions (Seber and Lee, 2012):

- There is a linear relationship between the dependent variables and the independent variables.
- The observations are independently and randomly selected from the population.
- $\varepsilon_i \sim N(0, \sigma^2)$

Sometimes data can have a non-linear relationship. One way to try to explain such a relationship is to use a polynomial regression model (Aiken, 1991);

$$y_i = \beta_0 + \beta_1 X + \beta_2 X^2 + \dots + \beta_h X^h + \varepsilon_i$$

Here h is called the degree of the polynomial. For lower degree, the relation has a specific name (i.e. h=2 means quadratic, h = 3 means cubic, h = 4 means quartic, etc.). Although this model enables a nonlinear relationship between Y and X, polynomial regression is still considered linear regression as the regression coefficients $\beta_1, \beta_2, \dots, \beta_h$ are linear.

While applying the regression model in time-dependent series, various transformations are applied. Time is considered as an independent variable in time series regression analysis. Models such as cubic regression model, exponential regression model, quadratic regression model and logistic regression model can be used to transform time series. The model with the smallest mean squared error is selected and cubic regression was used in this study.

In 1997, Ramon Lawrence tested the performance of artificial neural networks in price prediction in his article "Using Neural Networks to Forecast Stock Market Prices". As a result of the study, he said that price prediction with artificial neural networks is not perfect, but outperforms regression and other mathematical models. In 2003, Chen et al. estimated stock returns using Taiwan Stock Exchange data for January 1982 and August 1992. As a result of the study, it is emphasized that the returns can be predicted with artificial neural networks and investment decisions can be made with this analysis.

In artificial neural networks, there are 3 layers. These are the input layer, the output layer and the hidden layers. The input layer is the first layer in the model. In statistics, independent variables can be used as the input layer. The output layer is the last layer in the model. In statistics, dependent variables can be used as the output layer. Hidden layers transfer the information from the input layer to the output layer.

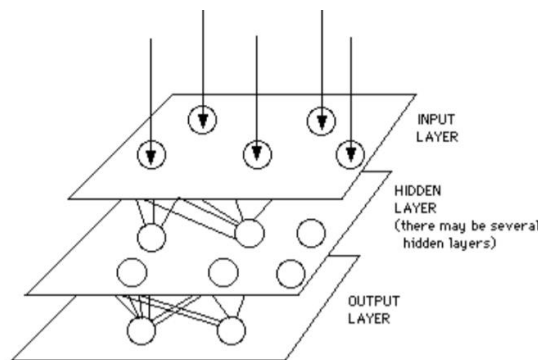


Figure 1. Artificial neural networks layers (Maind et al., 2014)

The following are the three most commonly used types of neural networks in artificial intelligence. (Patel and Gaurav, 2020):

1. **Feedforward Neural Networks:** Feedforward neural networks are the first type of artificial neural networks created and can be considered the most used today. These neural networks are called feedforward neural networks because the flow of information through the network is unidirectional and does not loop.
2. **Recurrent Neural Networks:** Recurrent neural networks (RNN), as the name suggests, involve the repetition of operations in the form of loops. These are much more complicated than feedforward networks and can perform more complex tasks than simple image recognition.
3. **Convolutional Neural Networks:** Convolutional neural networks have been associated almost exclusively with computer vision applications since their inception. This is because their architecture is specifically suited for performing complex visual analysis.

RESEARCH FINDINGS AND DISCUSSION

In the analysis, first the regression model and then the artificial neural network model estimation were made. BIST100 index was used as the dependent variable and gold price, USD rate, inflation rate, exchange rate and treasury bond were used as the independent variables. At the beginning of the study, the descriptive statistics were calculated. Then, it was tested whether the variables have normal distribution or not. Finally, regression and artificial neural network models were developed. The descriptive statistics used in the study are given in Table 1 and the normality test is given in Table 2.

Table 1. Descriptive statistics

	Mean	Std. Deviation	Minimum	Maximum	Median	Skewness	Kurtosis
BIST100	824,55	186,55	595,67	1133,56	773,15	0,355	-1,303

Table 2. Tests of Normality

Tests of Normality						
	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
BIST100	,196	11	,200*	,909	11	,237

When the results given in Table 2 are examined, we can say that the data are normally distributed. "Curve Estimation" was applied to decide which method to use in the regression analysis. The model with the smallest mean square error is selected. According to the results of the calculations, it was decided to use cubic regression. The results obtained by using cubic regression for the year 2021 are as follows:

II. International Applied Statistics Conference (UYIK-2021)
Tokat / Turkey, 29 June - 2 July 2021

Table 3 The results for the year 2021 obtained by using the cubic regression

DATE	BIST 100	ESTIMATION WITH CUBIC REGRESSION
2020-01	1191,4	1143,1
2020-02	1059,9	1152,0
2020-03	896,4	1160,9
2020-04	1011,1	1170,0
2020-05	1055,2	1179,2
2020-06	1165,2	1188,5
2020-07	1126,9	1197,9
2020-08	1078,6	1207,4
2020-09	1145,2	1217,0
2020-10	1112,3	1226,7
2020-11	1283,5	1236,5
2020-12	1476,7	1246,5
2021-01	1473,4	1256,5
2021-02	1471,3	1266,7
2021-03	1391,7	1277,0
2021-04	1397,8	1287,4
2021-05	1420,4	1297,9

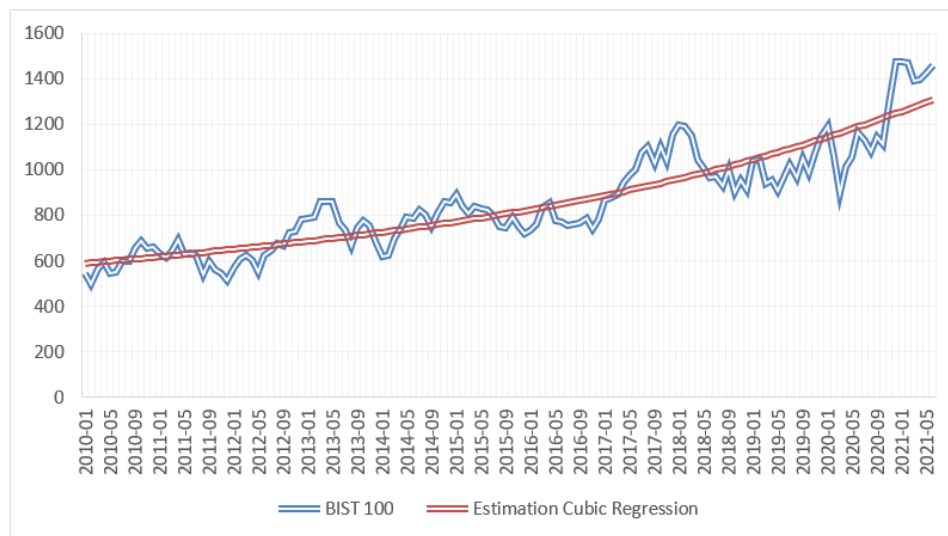


Figure 2. Forecast chart with BIST 100 and Cubic regression model

When Table 3 and Figure 2 are examined, it is seen that there is a difference between BIST 100 index values and forecast values. It was found that the estimated values for the year 2021 data were lower than the index values of BIST 100. When the years 2020-2021 were examined, there was an increase in the results of both models until February. However, despite the decrease in BIST 100 index values after February, an increase in forecast values was observed. In May 2021, the BIST 100 index values also began to increase. After testing the results of the cubic regression model, the artificial neural network model was constructed. 20% of the data was used as test data and 80% as training data. These data were selected randomly. The mean square error of the artificial neural network model was found to be 601.36. The mean square error of the cubic regression model was found to be 1120.17. Based on these results, it was found that the artificial neural network model has better performance. The results for 2021 obtained by using the artificial neural network model are as follows:

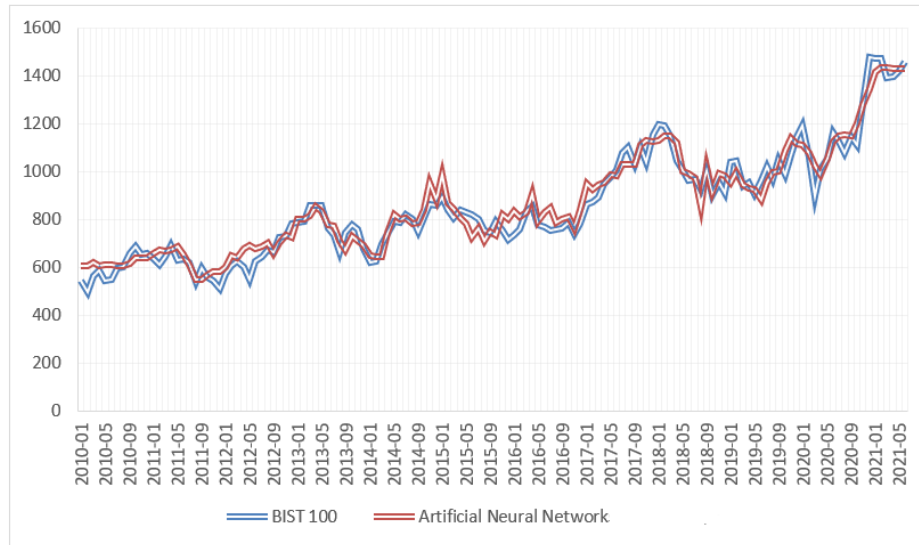


Figure 3. Forecast chart with BIST 100 and artificial neural network model

Table 4 The results for the year 2021 obtained by using the artificial neural network

DATE	BIST 100	ESTIMATION WITH ARTIFICIAL NEURAL NETWORK
2020-01	1191,4	1112,2
2020-02	1059,9	1076,6
2020-03	896,4	1026,6
2020-04	1011,1	988,6
2020-05	1055,2	1055,9
2020-06	1165,2	1125,4
2020-07	1126,9	1149,5
2020-08	1078,6	1151,8
2020-09	1145,2	1149,2
2020-10	1112,3	1201,5
2020-11	1283,5	1278,3
2020-12	1476,7	1345,8
2021-01	1473,4	1416,3
2021-02	1471,3	1435,4
2021-03	1391,7	1433,9
2021-04	1397,8	1431,5
2021-05	1420,4	1429,3

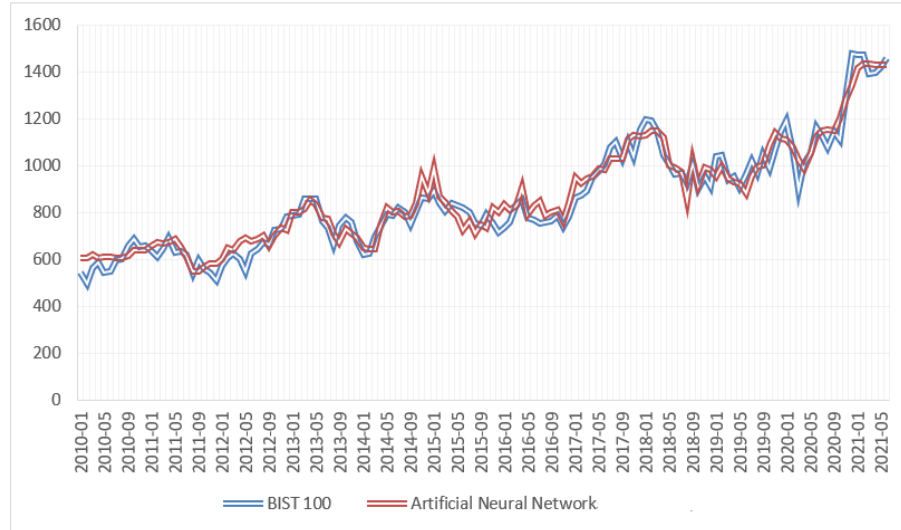


Figure 4. Forecast chart with BIST 100 and artificial neural network model

When Table 4 and Figure 3 were examined for the dataset obtained with the artificial neural network model, it could be seen that the index BIST 100 and the prediction results were very close. Similar to the regression model, an increase was observed in the artificial neural network model until the end of the period. However, in the real data of BIST 100 index, a decrease was observed. When the artificial neural network model and the regression model were compared, it was seen that the predictions obtained by the artificial neural network were closer to the real values. The prediction values obtained by the artificial neural network model and the regression model are given below.

Table 5 The results for the year 2021 obtained by using the artificial neural network

DATE	BIST 100	ESTIMATION WITH CUBIC REGRESSION	ESTIMATION WITH ARTIFICIAL NEURAL NETWORK
2020-01	1191,4	1143,1	1112,2
2020-02	1059,9	1152,0	1076,6
2020-03	896,4	1160,9	1026,6
2020-04	1011,1	1170,0	988,6
2020-05	1055,2	1179,2	1055,9
2020-06	1165,2	1188,5	1125,4
2020-07	1126,9	1197,9	1149,5
2020-08	1078,6	1207,4	1151,8
2020-09	1145,2	1217,0	1149,2
2020-10	1112,3	1226,7	1201,5
2020-11	1283,5	1236,5	1278,3
2020-12	1476,7	1246,5	1345,8
2021-01	1473,4	1256,5	1416,3
2021-02	1471,3	1266,7	1435,4
2021-03	1391,7	1277,0	1433,9
2021-04	1397,8	1287,4	1431,5
2021-05	1420,4	1297,9	1429,3

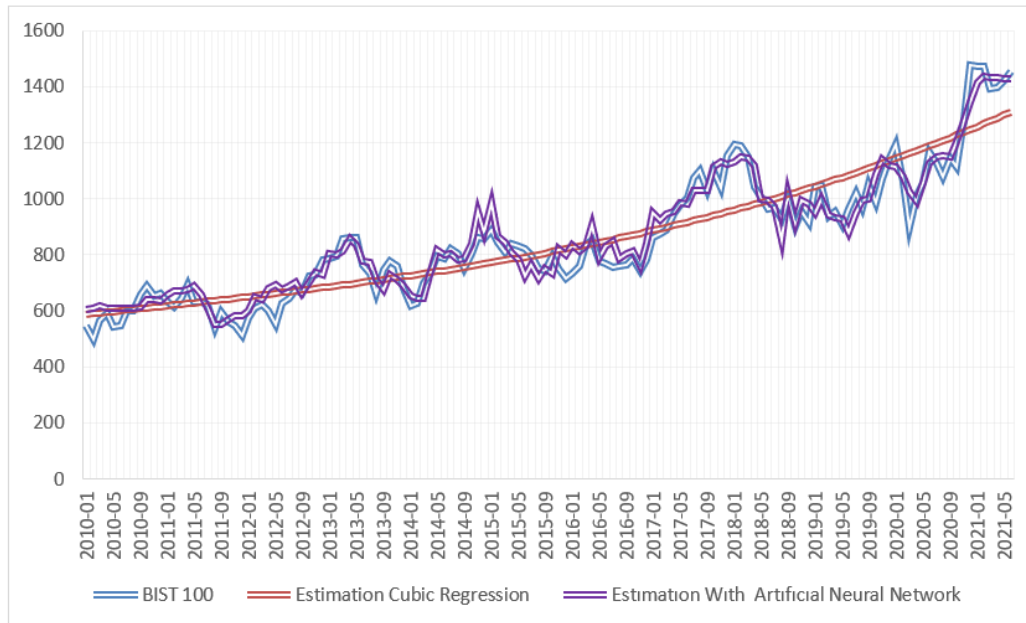


Figure 5. Forecast chart with BIST 100 and artificial neural network model and cubic regression model

CONCLUSIONS AND RECOMMENDATIONS

In this study, cubic regression and artificial neural network models were developed by using BIST 100 index data between 2010-2020. First, the mean square error values of the models were applied to decide which regression model to use. According to these results, it was decided to select the cubic regression model. Then, the results of the cubic regression model and the actual values of the index BIST 100 were compared. There was a difference between the estimated values and the actual values. Then the artificial neural network model was created.

As a result of the study, it can be seen that the artificial neural network model gives better results and performs better than the cubic regression model. On examining the tables and graphs, a decrease in BIST 100 values is observed in March 2020. The reason for this decline is believed to be the Covid 19 epidemic that started in Wuhan, China in late 2019 and spread around the world. This is a variable that is not added to the model. Therefore, both artificial neural networks and regression models were not affected by this situation and did not decrease.

A similar decrease was observed in February 2021. The reason for this is probably the quarantine process that occurred due to the Covid19 outbreak. Since unexpected situations are not taken into account in the model, there are discrepancies between the estimates and the actual values. However, it is shown that the predictions obtained with the artificial neural network model give better results than the cubic regression model. If it is assumed that there will be no factors that are not added to the model in the future; it can be said that the BIST 100 index can be estimated more accurately with the artificial neural network model.

References

- Aiken, L. S., West, S. G., & Reno, R. R. (1991). *Multiple regression: Testing and interpreting interactions* sage.
- Altan, A., & Karasu, S. (2019). The effect of kernel values in support vector machine to forecasting performance of financial time series. *The Journal of Cognitive Systems*, 4(1), 17-21.

- Benrhmach, G., Namir, K., Namir, A., & Bouyaghroumni, J. (2020). Nonlinear Autoregressive Neural Network and Extended Kalman Filters for Prediction of Financial Time Series. *Journal of Applied Mathematics*, 2020.
- Cao, J., Li, Z., & Li, J. (2019). Financial time series forecasting model based on CEEMDAN and LSTM. *Physica A: Statistical Mechanics and its Applications*, 519, 127-139.
- Cao, Q., Parry, M. E., & Leggio, K. B. (2011). The three-factor model and artificial neural networks: predicting stock price movement in China. *Annals of Operations Research*, 185(1), 25-44.
- Carvalho, A., & Ribeiro, T. (2008). Do artificial neural networks provide better forecasts? Evidence from Latin American stock indexes. *Latin American Business Review*, 8(3), 92-110.
- Diaz, J. F., & Nguyen, T. T. (2021). Application of grey relational analysis and artificial neural networks on corporate social responsibility (CSR) indices. *Journal of Sustainable Finance & Investment*, 1-19.
- Elliott, G., & Timmermann, A. (2008). Economic forecasting. *Journal of Economic Literature*, 46(1), 3-56.
- Farahani, K.M. and Mehralian, S. (2013), Comparison between artificial neural network and neuro fuzzy for gold price prediction, 13th Iranian Conference on Fuzzy Systems (IFSC), Tehran.
- Fisher, R. A. (1922). The goodness of fit of regression formulae, and the distribution of regression coefficients. *Journal of the Royal Statistical Society*, 85(4), 597-612.
- Gannous, A. S., & Elhaddad, Y. R. (2011). Improving an artificial neural network model to predict thyroid bending protein diagnosis using preprocessing techniques. *WASET*, 50, 124-128.
- Gauss, C. F. (1809). *Theoria motus corporum coelestium in sectionibus conicis solem ambientium auctore Carolo Friderico Gauss.* sumtibus Frid. Perthes et IH Besser.
- Karakul, A. K., (2020) Forecasting Stock Market Index With Artificial Neural Networks. *Mehmet Akif Ersoy Üniversitesi İktisadi ve İdari Bilimler Fakültesi Dergisi*, 7(2), 497-509.
- Kazem, A., Sharifi, E., Hussain, F. K., Saberi, M., & Hussain, O. K. (2013). Support vector regression with chaos-based firefly algorithm for stock market price forecasting. *Applied soft computing*, 13(2), 947-958.
- Khodayari, M. A., Yaghobnezahad, A., & Khalili Eraghi, K. E. (2020). A Neural-Network Approach to the Modeling of the Impact of Market Volatility on Investment. *Advances in Mathematical Finance and Applications*, 5(4), 569-581.
- Lawrence, R. (1997). Using neural networks to forecast stock market prices. *University of Manitoba*, 333, 2006-2013.
- Le, T. L., Abakah, E. J. A., & Tiwari, A. K. (2021). Time and frequency domain connectedness and spill-over among fintech, green bonds and cryptocurrencies in the age of the fourth industrial revolution. *Technological Forecasting and Social Change*, 162, 120382.
- Lee, K. J., Chi, A. Y., Yoo, S., & Jongdae Jin, J. (2008). Forecasting Korean Stock Price Index (Kospi) Using Back Propagation Neural Network Model, Bayesian Chiao's Model, And Sarima Model. *Academy Of Information & Management Sciences Journal*, 11(2).
- Lenel, L., Köster, R., & Fritsche, U. (2020). Introduction (Futures Past. Economic Forecasting in the 20th and 21st Century). *Futures Past. Economic Forecasting in the 20th and 21st Century*.
- Maciel, L. S., & Ballini, R. (2008). Design a neural network for time series financial forecasting: Accuracy and robustness analysis. *Anales do 9º Encontro Brasileiro de Finanças, Sao Pablo, Brazil*.
- Maind, S. B., & Wankar, P. (2014). Research paper on basic of artificial neural network. *International Journal on Recent and Innovation Trends in Computing and Communication*, 2(1), 96-100.
- Mir, H., Zaratgari, R., & Sotoudeh, R. (2021). Comparison of Risk Factors for Investing in Tehran Stock Exchange Using Smart Neural Network (Forecasting Tehran Stock Exchange with Neural Networks). *Agricultural Marketing and Commercialization Journal*, 5(1), 43-57.
- Molla, B., Cagil, G., and Uyaroglu, Y. (2021). Chaotic Analysis Of BIST 100 Return Time Series And Short-Term Predictability With ANFIS.
- Montgomery, D. C., Peck, E. A., & Vining, G. G. (2021). *Introduction to linear regression analysis*. John Wiley & Sons.
- Morelli, D. (2002). The relationship between conditional stock market volatility and conditional macroeconomic volatility: Empirical evidence based on UK data. *International Review of Financial Analysis*, 11(1), 101-110.

- Ozbey, F., Paksoy, S. (2020). GARCH Ailesi Modelleri ve ANN Entegrasyonu İle BİST 100 Endeks Getirisinin Volatilite Tahmini 1. *Business And Economics Research Journal*, 11(2), 385-396.
- Patel, L., & Gaurav, K. A. (2020). Introduction to Machine Learning and Its Application. In *Applications of Artificial Intelligence in Electrical Engineering* (pp. 262-290). IGI Global.
- Patra, A., Das, S., Mishra, S. N., & Senapati, M. R. (2017). An adaptive local linear optimized radial basis functional neural network model for financial time series prediction. *Neural Computing and Applications*, 28(1), 101-110.
- Pradeepkumar, D., & Ravi, V. (2017). Forecasting financial time series volatility using particle swarm optimization trained quantile regression neural network. *Applied Soft Computing*, 58, 35-52.
- Schroeder, L. D., Sjoquist, D. L., & Stephan, P. E. (2016). *Understanding regression analysis: An introductory guide* (Vol. 57). Sage Publications.
- Seber, G. A., & Lee, A. J. (2012). *Linear regression analysis* (Vol. 329). John Wiley & Sons.
- Siami-Namini, S., & Namin, A. S. (2018). Forecasting economics and financial time series: ARIMA vs. LSTM. *arXiv preprint arXiv:1803.06386*.
- Wieland, V., & Wolters, M. (2013). Forecasting and policy making. In *Handbook of economic forecasting* (Vol. 2, pp. 239-325). Elsevier.
- Xiong, T., Bao, Y., & Hu, Z. (2014). Multiple-output support vector regression with a firefly algorithm for interval-valued stock price index forecasting. *Knowledge-Based Systems*, 55, 87-100.
- Yule, G. U. (1897). On the theory of correlation. *Journal of the Royal Statistical Society*, 60(4), 812-854.
- Zhang, G. P., & Berardi, V. L. (2001). Time series forecasting with neural network ensembles: an application for exchange rate prediction. *Journal of the operational research society*, 52(6), 652-664.

Multiple Imputation Methods for Numerical Datasets

Burak DİLBER^{1*}, A. Fırat ÖZDEMİR¹

¹Dokuz Eylül University, Department of Statistics, 35390, İzmir, Turkey

*Sunucu yazar e-posta: burakdilber91@gmail.com

Abstract

Missing value is a common problem experienced by researchers, statisticians and data scientists. One of the conventional and effective way of solving this problem is imputation of missing observations by generating multiple substitutes. There are various multiple imputation methods in the literature. In this study, eight current methods are compared by applying corresponding R packages, functions and numerical data sets with different characteristics. These methods are predictive mean matching, classification and regression trees, Bayesian linear regression, non - Bayesian linear regression, random sample from observed values, quantile regression, hot deck and linear model by Bayesian bootstrap methods. Performance of these methods is assessed by using Normalized Root Mean Square Error (NRMSE) and the results show that predictive mean matching and classification and regression trees outperformed the others.

Keywords: *Missing value, Multiple imputation methods, Predictive mean matching, Classification and regression trees*

Türkiye’de Covid-19 Ölümlerinin İstatistiksel Süreç Kontrolü

Sümeyye ALTUN^{1*}, Zeynep ÖZTÜRK²

¹Artvin Çoruh Üniversitesi, Lisansüstü Eğitim Enstitüsü, İşletme Anabilim Dalı, 08100, Artvin, Türkiye

²Artvin Çoruh Üniversitesi, Hopa İİBF, İşletme Bölümü, 08200, Artvin, Hopa, Türkiye

*Sunucu yazar e-posta: sumeyyealtun008@gmail.com

Özet

2020 yılında başlayan ve tüm dünyayı etkisi altına alan Covid-19 salgını insan yaşamını olumsuz yönde etkilemiştir. Bulaşıcı olan bu hastalık ateş, öksürük, halsizlik, nefes darlığı gibi pek çok semptomlar göstermektedir. Daha çok öksürme, hapşırma ve nefes verme yoluyla veya temas yoluyla bulaştığı bilinmektedir. Bireyden bireye farklılık gösteren hastalığın şiddeti, insan yaşamını sonlandırmaya kadar gidebilmektedir. Tüm dünyada olduğu gibi Türkiye’de de vakaların görüldüğünü ve bu vaka sayılarının artan ve azalan şekilde olduğu gözlemlenmektedir. Bu çalışma, Türkiye’deki Covid-19 salgınından ölen insan sayılarının R programı “qcc” paketi yardımı ile istatistiksel süreç kontrolü teknikleri kullanılarak analiz edilmesi üzerine yapılmıştır. İstatistiksel süreç yaklaşımı, Covid 19 nedeni ile ölen insan sayısında ve yüzdesinde aylara göre değişimin ve istatistiksel olarak sürecin hangi zamanlarda kontrolden çıktığının kolayca görülmesini sağlar. Yapılan analizler sonucunda ölüm sayılarının kontrol altında olmadığı görülür.

Anahtar kelimeler: İstatistiksel süreç kontrolü, Covid-19, R paket programı, “qcc”

Statistical Process Control of Covid-19 Deaths in Turkey

Abstract

The covid-19 epidemic, which began in 2020 and affected the entire world, negatively affected human life. This disease, which is contagious, shows many symptoms such as fever, cough, weakness, shortness of breath. It is known to be transmitted mainly through coughing, sneezing and exhaling, or through contact. The severity of the disease, which differs from individual to individual, can go as far as death. As in the whole world, it is observed that cases are observed in Turkey, and the number of these cases is increasing and decreasing. In this study, the people who died from the Covid-19 epidemic in Turkey were analyzed using statistical process control techniques with the help of the R program “qcc” package. The statistical process approach allows you to easily see the change in the number and percentage of people who die due to Covid 19 by month, and statistically at what times the process gets out of control. As a result of the analysis, it is seen that the number of deaths is not under control.

Key words: Statistical process control, Covid-19, R package program, “qcc”

GİRİŞ

COVID-19, bugünlerde modern dünyanın en önemli küresel krizidir. Dünya çapında yayılan bu salgın, neredeyse tüm modern zaman hastalıklarının hızını ve boyutunu aşmıştır. Pandemi, sosyal kargaşa ile birlikte beklenmedik küresel ekonomik krizi tetiklemiş ve ayrıca alışkanlıklarımızı, ticaret yöntemimizi, sosyal ve ekonomik ilişkilerimizi, çalışma biçimlerimizi ve siyasi organizasyonlarımızı değiştirmiştir (WHO). Her devlet, COVID-19’un vakalarının yayılması ve neden olduğu ölümler için önleyici

politikalarla vatandaşlarını korumaya çalışmaktadır. Türkiye önleyici faktörler olarak, sosyal faktörler ve sokağa çıkma yasağı, sosyal mesafe, maske ve el yıkama alışkanlıkları, uzaktan eğitim ve uzaktan çalışma vb. gibi bazı önleyici tedbirler almıştır.

İstatistiksel Süreç Kontrolü (SPC), süreçleri izlemek ve kontrol etmek için yaklaşık bir asırdır kullanılan bir araçtır. Amerikalı istatistikçi Walter A. Shewhart, bir sürecin istatistiksel kontrol içinde mi yoksa dışında mı olduğunu belirlemeye izin veren kontrol çizelgelerini geliştirmiştir (WEC, 1956). İstatistiksel Kalite Kontrol ya da İstatistiksel Süreç Kontrolü, verileri bilimsel yöntemler kullanarak analiz eder ve bu verileri mühendislik, yönetim, bakım ve az çok sayılarla ifade edilebilecek diğer faaliyetlerdeki sorunları çözmek için uygular (WEC, 1956). İstatistik, sayılardan sonuç çıkaran yöntemleri ifade eder, kalite, incelenen şeyin özellikleriyle ilgilidir ve kontrol, gözlemlenen belirli sınırlar içinde tutmakla ilgilidir. İstatistiksel Süreç Kontrolünde gözlemlenen ve analiz edilen nesne bir süreçtir. Bir süreç, bir sonuç üretmek için birlikte çalışan bir dizi koşul veya neden olarak tanımlanabilir (WEC, 1956). Süreç değişimini analiz ederek, bu değişimin nedenlerini bularak ve sebebi ortadan kaldırarak süreç performansı iyileştirilebilir. İstatistiksel Süreç Kontrol Çizelgeleri zaman içindeki performansın kullanımı kolay görsel bir sunum imkânı sağlar.

İstatistiksel süreç kontrolü başlangıçta endüstriyel operasyonların ve özelliklerin izlenmesi ve kontrolü için geliştirilmiş olsa da, çevresel izleme sırasında mikrobiyolojik kalite, su kalitesi, cerrahi alan enfeksiyonu (SSI) çalışması ve inceleme özelliklerinin veya fenomenlerinin trendlenmesi ve salgınların veya pandemilerin analizi izlenmesinde de kullanılmıştır (Eissa, 2019). Yani, süreç kontrol çizelgeleri farklı alanlarda geniş bir uygulamaya sahiptir. Eissa (2019a)'de beslenme kaynağı olarak kullanılan belediye içme suyunun artırılması ve sağlık tesisi besleyen su arıtma istasyonları üzerinden veri oluşturularak analizi yapılmıştır. Eissa (2019b) ise, 20 yıllık izleme ve kayıt süresi boyunca Siklosporiosis salgının incelenerek kontrol süreçleri içinde analiz edilmesiyle yapılmıştır.

Shewhart kontrol çizelgeleri, bir sürecin istatistiksel kontrol içinde mi yoksa dışında mı olduğunu belirlemek için süreç verilerini kullanır. İstatistiksel kontrol dışındaki bir süreç savurgan ve verimsizdir. X-Bar ve R çizelgeleri, süreç değişkeni ortalamalarının ve varyansının zaman içinde nasıl davrandığını ölçer ve çizelgelerdeki varyansın ortak neden mi, kontrol limitleri içindeki varyans mı yoksa limitlerin dışında özel bir neden mi olduğunu analiz etmek için üst ve alt kontrol limitlerini kullanır. Bu aynı zamanda normal veya anormal varyans olarak da adlandırılabilir (WEC, 1956). Shewhart kontrol çizelgeleri, bir süreçteki hataları belirlemek için bir araç olarak hizmet edebilir, süreç davranışının temel nedenlerini belirlemeye yardımcı olabilir ve süreci iyileştirmeye yönelik herhangi bir eylem için istatistiksel bir temel oluşturabilir. Shewhart çizelgelerinin yorumlanması istatistiksel ölçümlerdeki rastgele dalgalanmalar ile gerçek "sinyaller" arasında ayırım yapmanıza olanak tanır ve bir şeyler değiştiğinde tepki vermenize yardımcı olur. Ek olarak, çizelgeler, bir sürecin ayarlanmaya mı yoksa bakıma mı ihtiyacı olduğu veya yalnız bırakılabileceği konusunda kararlar verebilir. Bu çizelgelerde oluşturulan kalıplara bağlı olarak, bunlar iki veya daha fazla değişken arasındaki döngüleri, eğilimleri, kaymaları, istikrarsızlığı veya etkileşimleri tanımlamaya yardımcı olabilir (WEC, 1956). Shewhart kontrol çizelgeleri büyük ölçüde normal dağılımın analizine dayanır ve sıfır hipotezi, X-Bar ve aralık ortalamasının süreç ortalamalarına eşit olduğudur. α 'nın %0,3 ve anlamlılık düzeyinin %99,7 olduğu y eksenine ± 3 sigma kontrol limitleri uygulandığında, kontrol limitlerinin dışındaki noktalar sıfır hipotezinin reddedilmesine yol açacaktır. Ek olarak, x eksenini zamanı temsil eder ve hipotez testine başka bir boyut ekler. Bu sıfır hipotezinin reddedilip reddedilmeyeceğini değerlendirmek için karmaşık bir kural setidir. (WEC, 1956). İstatistiksel Süreç Kontrolünün amacı, gelecekteki performansı iyileştirmek için geçmiş verileri kullanmaktır ve geleneksel olarak analiz, süreç veya süreçlerden ziyade çıktıya odaklanma eğilimindedir.

Covid 19 salgını zamanında nicel olarak incelemenin ilginç yollarından biri de istatistiksel süreç kontrolüdür. Üst orta gelirli ülkelerden birinde COVID-19'dan bildirilen günlük vakalar ve ölümler bir Excel veri tabanından filtrelendirilmiş, katmanlara ayrılmış ve kronolojik olarak düzenlenmiştir. Veri işleme, viral salgının modelini ve davranışını izlemek için istatistiksel programlar kullanılarak gerçekleştirilmiştir. Veri toplama, 12 Mart 2020'den 29 Temmuz 2020'ye kadar, bildirilen vakaların ve ölümlerin ana sapmalarının kaydedildiği belirli bir dönemi kapsamıştır (Eissa, 2020). Bu çalışmada, tüm dünyayı saran ve insan sağlığını etkileyen Covid-19 virüsünün aylık ölüm toplamalarının Türkiye genelini ne kadar etkilediğini, salgının kontrol altında olup olmadığını ve Nisan 2020-Nisan 2021 dönemleri arası Covid 19 nedeniyle ölüm sayıları için sürecin izlenmesi amaçlanmıştır. İstatistiksel Süreç kontrol tekniklerinden Shewart, Cusum ve Ewma kontrol grafik teknikleri İstatistiksel süreç kontrolü bilgisayar programları içerisinde önemli yeri bulunan ve sürekli gelişime açık olan R programı kullanılarak yapılmıştır. Bu üç yöntem karşılaştırmalı olarak verilmiştir. Çalışmada, İstatistiksel Süreç kontrolü ve bazı tekniklerden bahsedilmiştir. R programı “qcc” paketi tanıtılmıştır. Shewart, Ewma ve Cusum kontrol grafikleri ve sonuçları karşılaştırmalı olarak verilmiştir. Ayrıca, Türkiye de Covid 19 Pandemisi ve alınan önlemler anlatılmıştır. Bulgular kısmında elde edilen sonuçlar ve yorumlar yapılmıştır.

MATERYAL VE YÖNTEM

İstatistiksel Süreç kontrolü

Bireylerin ihtiyaçları ve talepleri gün geçtikçe artmaktadır. Tüketime artmasıyla sanayi devriminden günümüze gelen seri üretimin gelişmesiyle denetimin zorlaştığı ve maliyetlerin yüksek üretimi nedeniyle istatistiksel olarak bu durumu çözümlemek amacıyla yöntemler geliştirilmeye başlanmıştır. Böylelikle İstatistiksel Süreç Kontrolü olarak devamlı denetim ve sürecin her anının kontrolünün sağlayan bir yöntem şekli ortaya çıkmıştır (Can, 2007:15; Patır, 2009:232).

Geçmişte üretilen ürünlerdeki kalitenin, tamamen muayene ve ölçme teknikleri ile sağlanırken günümüze doğru kalitenin muayene ve yüzde yüz oluşamaması sebebiyle istatistiksel süreç kontrolü gibi yöntemler var olmuştur. İstatistiksel Süreç Kontrolü, belirlenen istatistik teknikleri kullanarak veri toplamak, toplanan verilerle analiz etmek, yorumlama ve çözümlerle hatalı üretimin azaltılmasıyla birlikte kaliteli ürün üretimini gerçekleştirip ileride oluşabilecek problemleri önceden tahmin edip verimliliği yükselten bir metottur (Yıldırım, Karaca, 2013:78; Aydın, Kargı, 2018:42).

İstatistiksel Süreç Kontrol Teknikleri

Kaliteli üretim sağlamakla birlikte denetim ve kontrol amacıyla istatistiksel tekniklerden yararlanma İstatistiksel Süreç Kontrolü ve tekniklerini oluşturmaktadır. Bu teknikleri şu şekilde sıralanabilir:

- Sınıflandırma: Verilerin değişkenlik durumlarına göre veya kategorilere ve özelliklerine göre sınıflandırma (gruplandırma) yöntemi olarak kullanılmaktadır. Problem araştırmasında hatalı ve hatasız ürünlerin hangi makinalardan veya ürünlerden kaynaklandığını saptamak için kullanılır (Can, 2007:24; Ndong Ovono Ekouma, 2019:33).
- Çetele: Basit bir istatistik süreç kontrolü tekniği olan çetele, kullanımında verilerin toplanması, toplanan verilerin analiz edilmesi ve kayıt edilmesi konusunda pek çok imkân sunar (Yıldırım, Karaca, 2013:78).
- Histogram: Nicel verilerin dağılım aralıklarını grafiksel olarak gösterimini sağlayan istatistiksel bir yöntemdir. Bununla beraber süreç iyileştirme işlemlerinin gösterildiği aynı zamanda çan şeklinde bir eğri ile gösterilen istatistiksel süreç kontrolü tekniğidir (Yıldırım, Karaca, 2013:78).
- Pareto Analizi: İtalyan bir ekonomist olan Vilfredo Pareto, problemleri oluşturan kaynağın ne

derecede önemli olduğu açısından sıralamayı oluşturan aynı zamanda bu pareto analizi 80'e 20 kuralı şeklinde de tanımlanan histogram türü olarak bilinmektedir (Ndong Ovono Ekouma, 2019:24).

- Neden- Sonuç Diyagramı: Balık kılçığı olarak da bilinen bu yöntem, bir problemin sonuçları ile nedenleri arasında ki karşılıklı ilişkiyi düzenleyerek sunmaktır. Diyagram çizilirken, oluşan problem ayrıntılarıyla açıklanıp ve olabilecek nedenler için bir ana kategori oluşturulur (Patır, 2009:237).

- Serpilme Diyagramları: İki değişken arasındaki ilişki ile birlikte aralarında herhangi bir korelasyon bir bağlantı olup olmadığını karar vermek ve neden- sonuç ilişkisiyi kanıtlamak veya kanıtlamamak için oluşan bir grafiksel diyagramdır. Bu grafiksel yöntemde bağımsız değişken X, bağımlı değişken Y ekseninde gösterilir (Patır, 2009:237; Yıldırım, Karaca, 2013:78).

- Kontrol Grafikleri: Kontrol grafiklerinin önemi üretimin gerçekleşmesi aşamasında dışsal çevrenin kaliteyi etkileyecek sorunların oluşmasını analiz edip değerlendirilmesini bunlar sonucunda ortaya çıkacak zararların engellenmesiyle verimliliğin yüksek düzeylere çıkarılmasını sağlamaktadır. Kontrol grafikleri, istatistik süreç kontrolü teknikleri içinde ekonomik ve güvenilirlik açısından en etkili yöntemlerdendir (Patır, 2009:238).

Kontrol grafikleri üç kontrol limitine bağlı olarak gösterilir. Merkezi Değer (MS), üst kontrol limiti (ÜKL) ve alt kontrol limiti (AKL) şeklinde sıralanır. Bunlar istatistiksel süreçten elde edilen veriler ile oluşturulur. Bu veriler kontrol limitlerinin dışına çıkması ile sürecin kontrol dışına çıktığını gösterir (Ndong Ovono Ekouma, 2019:35-36).

- Shewhart kontrol grafikleri: X-bar grafiği, bir kontrol grafiğindeki ortalamaların bir grafiğidir. R- çubuğu, zamana göre grup aralıklarının veya yanıtların bir grafiği şeklindedir.(Shomran, 2013:15).

- Ewma kontrol grafikleri: Katlanarak ağırlıklı Hareketli Ortalama (EWMA) kontrol grafiği, süreçteki değişimleri tespit etmek için kullanılan kontrol grafiklerinden biridir. İlk olarak Roberts (1959) tarafından tanımlanan bu şemanın özelliklerinden biri, süreçteki küçük ve orta ölçekli kaymaların hızlı bir şekilde tespit edilmesidir (Shomran, 2013:24).

- Cusum kontrol grafikleri: Küçük değişimleri tespit edebilen daha gelişmiş bir alternatif olan kontrol grafiği kümülatif toplam kontrol(CUSUM) grafikleri kullanılmaktadır (Stijn, 2018:8).

R Paket Programı (“qcc”)

R programı 1990 yıllarda Yeni Zelanda bulunan Auckland Üniversitesi, Ross Ihaka ve Robert Gentleman tarafından ortaya çıkmıştır. S diline benzer bir dil olan R dilinin ilk sürümü 29 Şubat 2000 tarihinde yayınlanmıştır. R, istatistiksel hesaplama ve grafikler için yazılım ortamı olup aynı zamanda programlama dilidir. Veri işleme, hesaplama, grafik gösterimi yapmayı sağlayan bir programdır. R, çok geniş istatistiki (doğrusal ve doğrusal olmayan modelleme, klasik istatistik testleri, zaman serileri analizi, sınıflandırma, kümeleme ve diğer) ve grafik çizim teknikleri sunmaktadır (Beydoğan, Çetin, Orbay, 2014:4-6).

“qcc” paketi, istatistik kalite kontrol süreçlerinde kullanılan, Shewhart, EWMA, CUSUM kontrol grafikleri, temel süreç kontrolleri grafiklerinin oluşturulmasında kullanılan bir pakettir (Orçanlı, 2019:1395).

“qcc” paketini R programında kullanmak için ilk önce paketin yüklenmesi ve kütüphaneden çağırılması gerekmektedir.

```
> install.packages("qcc")
WARNING: Rtools is required to build R packages but is not currently installed. Please
download and install the appropriate version of Rtools before proceeding:

https://cran.rstudio.com/bin/windows/Rtools/
Installing package into 'C:/Users/bilgisayar/Documents/R/win-library/4.0'
(as 'lib' is unspecified)
URL 'https://cran.rstudio.com/bin/windows/contrib/4.0/qcc_2.7.zip' deneniyor
Content type 'application/zip' length 3557347 bytes (3.4 MB)
downloaded 3.4 MB

package 'qcc' successfully unpacked and MD5 sums checked

The downloaded binary packages are in
C:\Users\bilgisayar\AppData\Local\Temp\RtmpERUojo\downloaded_packages
> library(qcc)

  /--  /--  /--  Quality Control Charts and
 | ( ) | ( ) | ( ) Statistical Process Control
  \--  \--  \--
      |
      version 2.7
Type 'citation("qcc")' for citing this R package in publications.
```

Covid-19 Pandemisi

Covid-19, Çin'in Wuhan şehrinde, 2020 yılında ortaya çıkmakla beraber tüm dünyayı etkisi altına alan bir salgın haline gelmiştir. Öksürük, ateş, nefes almada zorluk gibi belirtiler gösteren bir virüs türüdür. Dünya Sağlık Örgütü'nün verilerine göre "salgın 212 ülkede yaklaşık 1.9 milyon kişi enfekte olmuş ve 120.000'den fazla kişinin ölümüne yol açmıştır." (Atay, 2020:168).

Bu salgının geçmişte yaşanan salgınlardan farklı olarak insanların hareketlilik derecesinin yüksek olmasıyla yayılma ve hastalık kapma oranının fazla olmasından kaynaklanmaktadır. Bu durumla Çin'de başlayan bu virüs kısa zamanda tüm dünyayı etkisi altına almıştır. Ekonomik anlamda sıkıntıların doğması, salgından kaynaklı tedbir amacıyla sınırların kapatılması, ülkelerin veya şehirlerin yoğun salgından dolayı karantinaya alınması gibi etkenleri beraberinde getirmiştir. Vaka sayılarının giderek artması, ölüm oranlarının yükselmesine neden olmuştur (Turan, Çelikyay, 2020:4).

Covid-19 virüsüne yakalanan bireylerin öksürmesi, hapşırması ile etrafa saçılan su damlacıkları veya ellerin, temas edilen yerlere dokunulması ile göz, burun, ağıza götürülmesi virüs bulaşmasına neden olabilir. Hastalığı yakalanan bireylerin bazıları hafif atlatırken, bazıları ise ciddi şekilde ağır geçirebildiği görülmektedir. Hastalıktan, kronik rahatsızlığı olanlar, kalp, diyabet, hipertansiyon hastalığı olanlar, kanser ve sağlık çalışanları en çok etkilenen grup içerisinde (Covid19 Aşı, 2021).

Korunmak için;

- Ellerin temizliğine dikkat etmek ve yıkamadan yüze temas ettirmemek,
- Hastalığın yoğun olduğu ve kalabalık ortamlarda bulunmamak,
- Hapşırma ve öksürük esnasında ağız ve burnu tek kullanımlık mendil ile kapatmak,
- Çiğ veya az pişmiş etleri yemeden kaçınmak,
- Seyahatten sonra oluşabilecek belirtiler için 14 gün içinde en yakın sağlık kuruluşuna başvurmak gerekmektedir (Covid19 Aşı, 2021).

Salgının artması vaka sayılarının arttırırken Türkiye bu artan vaka sayıları için genelgeler ve tedbirler uygulamıştır. İlk vakanın görülmesiyle Türkiye genelinde kısıtlamalar başlamış, maske takmak zorunlu

hale gelmiş ve iş yerleri geçici olarak faaliyetlerine ara vermek durumunda kalmışlardır. Bununla beraber 65 yaş ve üstü ile kronik rahatsızlığı olanlar için sokağa çıkma yasağı uygulanmıştır. Bu yaş grubu için Vefa Sosyal Destek Grupları tarafından ihtiyaçları giderilmiştir. Şehir girişi ve çıkışlara belli sürelerle aralıklarında yasaklamalar getirilmiştir. Artan vaka sayısı ile 20 yaş ve altı yaş grubuna da sokağa çıkma yasağı getirilmiş ve bazı il, ilçelerde karantina uygulamaları başlamıştır. 65 yaş ve üstü ile 20 yaş ve altı yaş grubundaki bireylerin sokağa çıkmaları belirli saat aralıklarına gerçekleştirilmiştir. Uygulanan tedbirlerin vaka sayılarını düşürmesiyle giriş çıkışı yasak olan ve seyahat engelinin bulunması kaldırılmış normalleşme sürecine girilmeye başlamıştır. Belirtilen yaş grupları izin belgesi almadan il değişikliği yapmalarında engel kalmamıştır. Fakat hafta sonları Türkiye genelinde zorunlu hal ve izin belgesi haricinde olanlar için sokağa çıkma yasağı uygulanmaya devam edilmiştir. Sınavlar sorunsuz bir şekilde yapılmaya başlanmıştır. “Hayat Eve Sığar” uygulaması üzerinden HES kodu alınarak tüm kamu ve özel kurum ve kuruluşlarda, otobüslerde, konaklama alanlarında, alışveriş merkezlerinde neredeyse tüm kapalı alanlara girişte HES kodu uygulaması getirilerek denetim sağlanmıştır. Alınan yeni kararlar ile alışveriş merkezlerinin saat aralıkları 10.00 ile 20.00 olmuş, restoran, lokanta, pastane, kafe gibi yeme içme yerleri sadece paket servis yapmaları, 65 yaş ve üstü

10.00 ile 13.00, 20 yaş ve altı ise 13.00’den 16.00 saat aralıklarında sokağa çıkabilecekler ve hafta sonları 10.00 ile 20.00 saatleri dışında sokağa çıkmak yasaklanmıştır. Kamu ve özel kuruluşlar, eczane, veteriner gibi pek çok iş yeri faaliyetlerine devam edebilmişlerdir. Aralık 2020 yılında vakaların artış göstermesiyle yeni kararlar alınmaya başlamıştır. Hafta sonları tamamen sokağa çıkma yasağı gelmiş ve bu durumdan “üretim, imalat, tedarik ve lojistik zincirlerinin aksamaması, sağlık, tarım ve orman faaliyetlerinin sürekliliğini sağlamak” amacıyla yasaktan muaf tutulmuştur. Esnek çalışma saatleri belirlenmiştir. Buna göre illere göre bir harita çıkartılmış ve renklere göre değerlendirme yapılmıştır. Kırmızı olan en yüksek risk taşıyan iller, mavi olan iller ise vaka sayısının düşük olan illeri temsil etmektedir. Renklere göre tedbirler oluşturulmuştur (İçişleri Bakanlığı, 2021).

Vaka sayılarında görülen patlama, ölüm sayıları Türkiye genelini “Tam Kapanma”ya götürmüştür. 29 Nisan 2021 ile 17 Mayıs 2021 tarihleri arasında tam kapanma uygulanmıştır. Zorunlu olmadıkça sokağa çıkmak yasaklanmış ve en yakın market olmak koşuluyla temel ihtiyaçların giderilmesi için olanak tanınmıştır. Yeme içme yerleri sadece paket servis yapmak için faaliyetlerine devam edebilmişlerdir. Şehirlerarası seyahat tamamen yasaklanmış, sadece görev ve izin belgesi bulunanlar muaf olmuşlardır. “Tam kapanma sürecinde sağlık, güvenlik, acil çağrı gibi kritik görev alanları hariç olmak üzere kamu kurum ve kuruluşlarında hizmetlerin sürdürülebilmesi için gerekli olan asgari personel seviyesine düşürülerek uzaktan veya dönüşümlü çalışma gerçekleştirmişlerdir.” Pazar yerleri sadece sebze, meyve ve fide ihtiyaçların giderilmesi için kurulmuştur. Tam kapanma sonu vakaları az da olsa düşmesi ile kademeli normalleşme süreci başlamış ve Haziran ayında normalleşmeye geçilmiştir. Pazar hariç haftasonu sokağa çıkma yasağı kalmış, hafta içi ve cumartesi günü saat 22.00’ten saat 05.00’a kadar yasak getirilmiştir. Seyahat yasakları kalmış, yeme-içme yerleri faaliyetlerine kaldığı yerden devam etme imkânı tanınmıştır (İçişleri Bakanlığı, 2021).

Bu süreçte aşılama yöntemi de etkili olmaktadır. Türkiye’de 03 Haziran 2021 tarihine kadar toplam 29.665.125 aşılama gerçekleşmiştir (Covid19 Aşı, 2021).

BULGULAR

Analizde, Türkiye’de 2020 Nisan ayından 2021 Nisan ayını kapsayan günlük ölen kişi sayısı dikkate alınarak oluşturulan veriler kullanılmıştır. 13 ayı kapsayan günlük ölüm sayısı verileri Çizelge 1’de verilmiştir. Veriler Sağlık Bakanlığı web sitesinden alınmıştır.

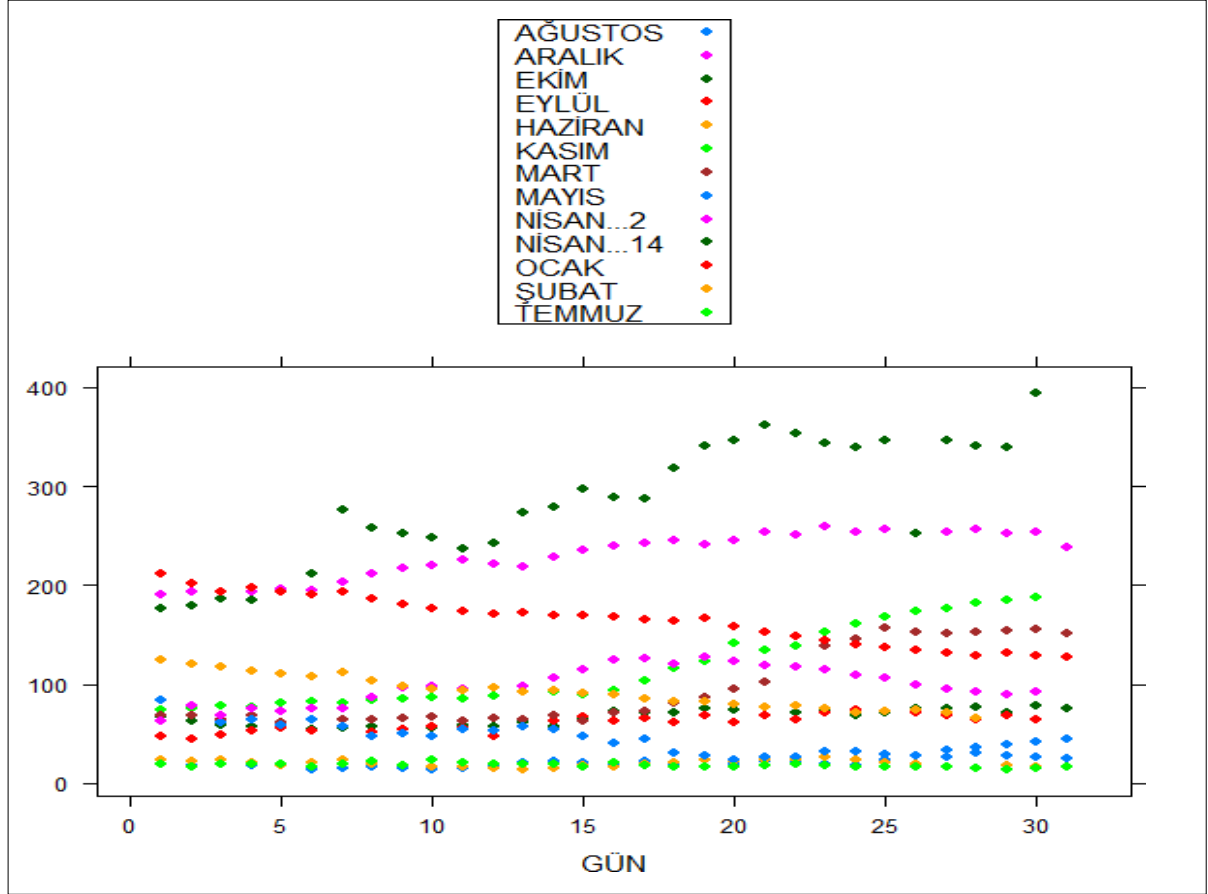
Çizelge 1. Günlük Ölüm Sayısı Verileri

GÜN	2020 Yılı Vefat Edenler									2021 Yılı Vefat Edenler			
	NİSAN	MAYIS	HAZİRAN	TEMMUZ	AĞUSTOS	EYLÜL	EKİM	KASIM	ARALIK	OCAK	ŞUBAT	MART	NİSAN
1	63	84	23	19	19	47	67	74	190	212	124	69	176
2	79	78	22	17	18	45	63	76	193	202	120	68	179
3	69	61	24	19	19	49	59	79	187	193	117	65	186
4	76	64	21	20	18	53	57	77	193	197	113	68	185
5	73	59	18	19	19	56	57	81	196	194	110	62	193
6	75	64	21	16	14	53	55	83	195	191	108	64	211
7	76	57	23	19	15	57	56	81	203	194	112	65	276
8	87	48	19	22	16	52	58	84	211	186	103	64	258
9	96	50	18	18	15	55	55	85	217	181	98	66	253
10	98	47	17	23	14	58	56	87	220	176	95	67	248
11	95	55	17	21	15	56	59	86	226	174	94	63	237
12	97	53	15	19	18	48	58	88	222	171	97	66	243
13	98	58	14	19	21	57	62	93	218	173	93	65	273
14	107	55	15	20	22	63	57	92	229	170	94	68	279
15	115	48	18	17	21	67	66	89	235	169	91	63	297
16	125	41	17	21	19	63	73	94	240	168	90	71	289
17	126	44	19	18	22	66	71	103	243	165	86	73	288
18	121	31	21	17	20	62	72	116	246	164	83	81	318
19	127	28	23	16	23	68	75	123	241	167	82	87	341
20	123	23	22	17	19	61	74	141	246	159	80	95	346
21	119	27	23	18	22	68	68	135	254	153	77	102	362
22	117	27	24	19	22	65	71	139	251	149	78	117	354
23	115	32	27	18	19	72	74	153	259	144	75	138	343
24	109	32	24	17	18	74	69	161	254	140	72	146	339
25	106	29	21	16	24	73	72	168	256	137	73	157	347
26	99	28	19	17	20	71	75	174	253	134	74	153	253
27	95	34	17	17	26	68	76	177	254	132	71	151	346
28	92	30	15	15	36	65	77	182	257	129	66	153	341
29	89	28	18	14	39	68	72	185	253	131		154	339
30	93	26	16	15	42	65	78	188	254	129		155	394
31		25		17	44		75		239	128		152	

Çizelge 1’ de 2020 Nisan ayından Ağustos ayının sonlarına doğru vaka sayılarında azalış gözlenirken Ağustos ayının sonralarından sonra giderek artan vefat sayıları kısıtlamaların ve yasakların gelmesine neden olmuştur.

Günlük ölüm verilerinin aylara göre dağılımını elde etmek için R programında “Rcmdr” paketi kullanılarak yazılan komutlar ve elde edilen günlük ölüm verilerinin aylara göre dağılım grafiği Şekil 1’ de gösterilmiştir.

```
Rcmdr> xyp1ot(AĞUSTOS + ARALIK + EKİM + EYLÜL + HAZİRAN + KASIM + MART + MAYIS +  
Rcmdr+ NİSAN...2 + NİSAN...14 + OCAK + ŞUBAT + TEMMUZ ~ GÜN, type="p", pch=16,  
Rcmdr+ auto.key=list(border=TRUE), par.settings=simpleTheme(pch=16),  
Rcmdr+ scales=list(x=list(relation='same'), y=list(relation='same')), data=cvd19)
```



Şekil 1. Günlük verilerin aylara göre dağılımı

Şekil 1’ de ölüm sayısı en yüksek olan ay 2021 yılının Nisan ayıdır ve en düşük ölüm sayısı 2020 yılının Haziran ve Temmuz aylarında olmuştur. Nisan 2021 sonuçları üzerine Türkiye’de devlet önleme politikası olarak tam kapanma kararı ile vaka ve ölüm sayılarının kontrol altına alınması amacı sağlanmıştır. Çizelge 2’ de aylık toplam ölüm sayısı verileri verilmiştir. Bu veriler kullanılarak ölüm sayılarının kontrol sınırları içerisinde olan ve olmayan ayları saptamak amacıyla istatistiksel süreç kontrol teknikleri ile analiz edilmesi sağlanmıştır.

Çizelge 2. Aylık Toplam Ölüm Verileri

2020-2021 Ayları	COVID-19 Ölüm Sayısı
NİSAN	2960
MAYIS	1366
HAZİRAN	591
TEMMUZ	560
AĞUSTOS	679
EYLÜL	1825
EKİM	2057
KASIM	3494
ARALIK	7135
OCAK	5112
ŞUBAT	2576
MART	2968
NİSAN	8494

Analiz

R programı üzerinden “qcc” paketi ile Shewhart kontrol grafiklerinden “x-bar” grafiğini elde etmek için verinin boyutu ve ortalaması yazılarak yazılan komutlar, sonuçlar ve Shewhart grafiği Şekil 2’de gösterilmiştir.

```
> cvd19=qcc(cvd19,type = "xbar",sizes = 13)
> summary(cvd19)

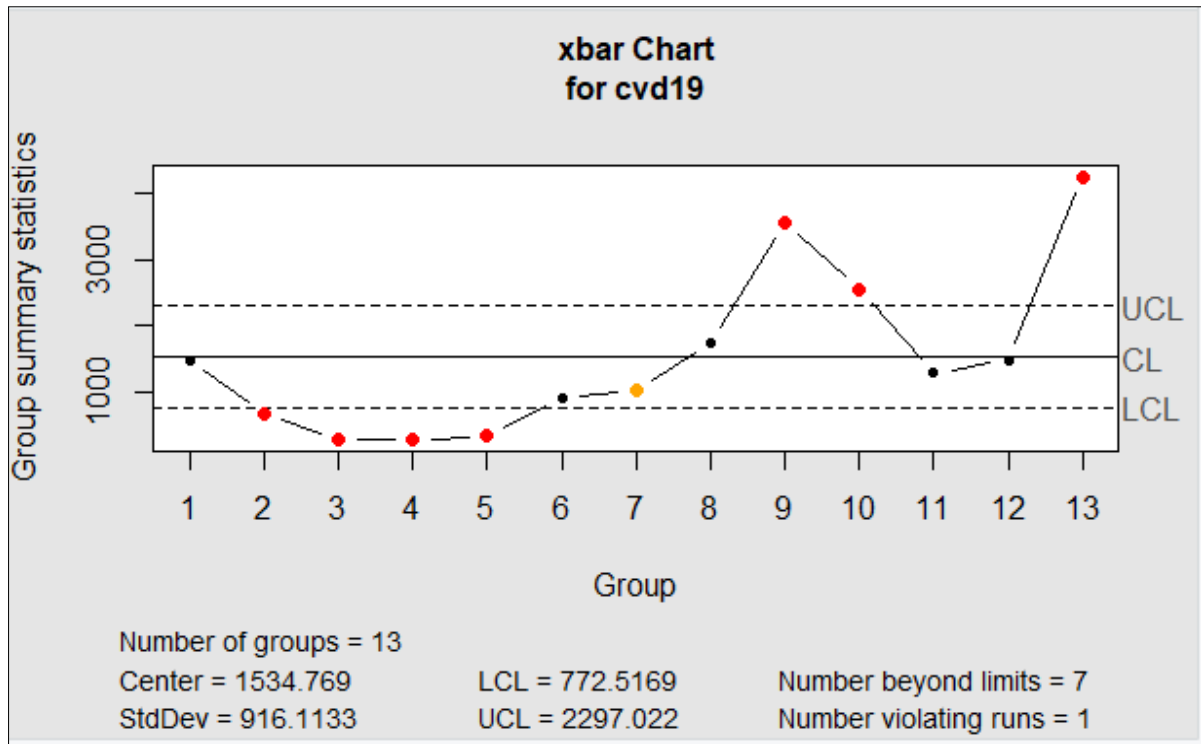
Call:
qcc(data = cvd19, type = "xbar", sizes = 13)

xbar chart for cvd19

Summary of group statistics:
      Min. 1st Qu.  Median    Mean 3rd Qu.    Max.
 286.000  687.000 1293.500 1534.769 1750.000 4251.500

Group sample size: 13
Number of groups: 13
Center of group statistics: 1534.769
Standard deviation: 916.1133

Control limits:
      LCL      UCL
772.5169 2297.022
```

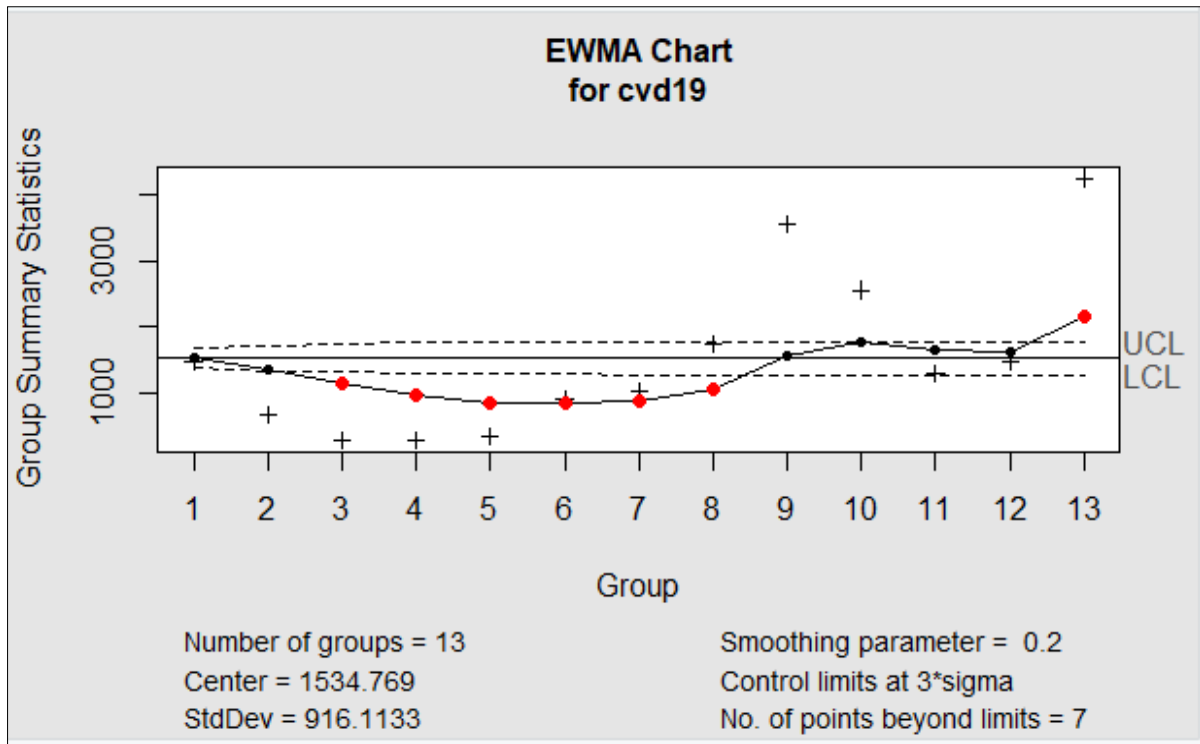


Şekil 2. Shewhart X-bar Grafiği

R programında “qcc” paketinin yardımı ile Türkiye’deki Covid-19 ölen kişi sayısı ile oluşturulan veriden Shewhart grafiklerinden X-bar kontrol grafiğini kullanarak Covid-19 sürecinin kontrol altında olup olmadığı incelenmektedir. Üst kontrol limitinin 772.5169, ortalamasının 1534.769 ve alt kontrol limitinin 2297.022 olduğu gözlemlenmektedir. X-bar grafiğine göre sonuçların kontrol sınırları dışında olan değerlerin var olması nedeniyle Covid-19 sürecinin kontrol altında olmadığını ifade etmektedir.

Şekil 3’ te Ewma kontrol grafiği için kullanılan komut ve grafik elde edilmiştir.


```
> ewma(cvd19,sizes = 13)
List of 15
 $ call      : language ewma(data = cvd19, sizes = 13)
 $ type      : chr "ewma"
 $ data.name : chr "cvd19"
 $ data      : num [1:13, 1:2] 9 8 5 12 1 4 3 6 2 10 ...
 ..- attr(*, "dimnames")=List of 2
 $ statistics: Named num [1:13] 1484 687 298 286 340 ...
 ..- attr(*, "names")= chr [1:13] "1" "2" "3" "4" ...
 $ sizes     : num [1:13] 13 13 13 13 13 13 13 13 13 13 ...
 $ center    : num 1535
 $ std.dev   : num 916
 $ x         : int [1:13] 1 2 3 4 5 6 7 8 9 10 ...
 $ y         : Named num [1:13] 1525 1357 1145 973 847 ...
 ..- attr(*, "names")= chr [1:13] "1" "2" "3" "4" ...
 $ sigma     : num [1:13] 50.8 65.1 72.8 77.3 80 ...
 $ lambda    : num 0.2
 $ nsigmas   : num 3
 $ limits    : num [1:13, 1:2] 1382 1340 1317 1303 1295 ...
 ..- attr(*, "dimnames")=List of 2
 $ violations: Named int [1:7] 3 4 5 6 7 8 13
 ..- attr(*, "names")= chr [1:7] "3" "4" "5" "6" ...
 - attr(*, "class")= chr "ewma.qcc"
```

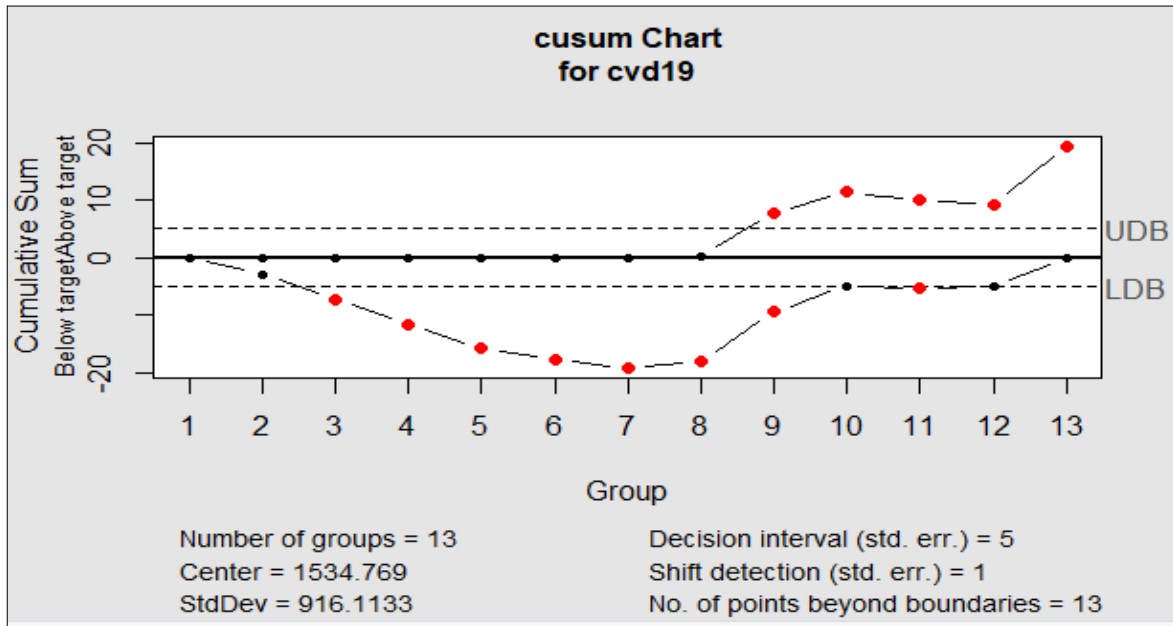


Şekil 3. Ewma Grafiği

Ewma grafikleri küçük ve orta ölçekli kaymaları tespit edilmek için kullanıldığı için üst ve alt sınırlar önemlidir. Shewhart grafiğinden farkı Şekil 3'te görüldüğü gibi sigma değeri önemlilik göstermektedir. Şekil 3 incelendiğinde kırmızı noktaların bulunduğu yerlerin alt ve üst sınırın dışında olduğu görülmektedir. Bu da yapılan analiz sonucunda kontrol altında olmadığını ifade etmektedir.

Şekil 4' te Cusum kontrol grafiği için kullanılan komut ve grafik elde edilmiştir.

```
> cusum(cvd19,sizes = 13)
List of 14
 $ call      : language cusum(data = cvd19, sizes = 13)
 $ type      : chr "cusum"
 $ data.name  : chr "cvd19"
 $ data      : num [1:13, 1:2] 9 8 5 12 1 4 3 6 2 10 ...
 ..- attr(*, "dimnames")=List of 2
 $ statistics : Named num [1:13] 1484 687 298 286 340 ...
 ..- attr(*, "names")= chr [1:13] "1" "2" "3" "4" ...
 $ sizes     : num [1:13] 13 13 13 13 13 13 13 13 13 13 ...
 $ center    : num 1535
 $ std.dev   : num 916
 $ pos       : num [1:13] 0 0 0 0 0 ...
 $ neg       : num [1:13] 0 -2.84 -7.2 -11.62 -15.82 ...
 $ head.start : num 0
 $ decision.interval : num 5
 $ se.shift   : num 1
 $ violations : List of 2
 ..- attr(*, "class")= chr "cusum.qcc"
```



Şekil 4. Cusum Grafiği

Shewhart grafiği küçük değişimlere dikkat etmezken Cusum grafiği küçük değişimleri tespit etmek için kullanılan kontrol grafiğidir. Diğer grafik türlerinde olduğu gibi alt ve üst sınır arasında olmak zorundadır. Bunun için, Şekil 4 incelendiğinde kırmızı noktaların bulunmasının kontrol sınırları dışında olduğunu ve sonuç olarak kontrol altında olmadığı gözlemlenmektedir.

SONUÇ

İstatistiksel süreç kontrolleri teknolojinin gelişmesiyle yeni, gelişen teknikler ve sistemler ortaya çıkmasına neden olmuştur. Analizlerin, ölçümlerin yapılması kolaylaşmış ve zaman kaybını aza indirecek sistemler oluşturulmuştur. Bu sistemler R gibi önemli ölçüde üretim aşamasında istatistiksel süreçleri kontrol etmek için yararlanılan programlar haline gelmiştir. R programı bunlardan biri olarak istatistiksel süreç tekniklerinin kullanımı ile birlikte analiz edilme kolaylığı sağlamıştır.

Bu çalışma için veri olarak kullanılan Covid-19 salgını, tüm dünyayı etkisi altına almıştır. Zorlaşan koşullar ile pek çok sektörü etkilemiştir. İnsan sağlığını tehlikeye sokan bu salgın halen daha devam etmektedir. Ölüm sayılarının değişkenlik göstermesi alınan tedbirler ve izlenecek süreçleri değiştirmektedir. Bu çalışmada da istatistiksel süreç kontrol teknikleri ile birlikte Türkiye'deki Covid-19 virüsünün ortaya çıkardığı ölüm sayılarının aylık toplamının analiz edilmesi için kullanılan R programının yardımı, sonuçların kısa zamanda ve hatasız bir şekilde hesaplanmasını sağlamıştır. Sonuç

olarak yapılan analizler sonucunda ölüm sayılarının kontrol altında olmadığı gözükmemektedir. Türkiye de artan vaka sayıları ve ölüm sayıları nedeniyle dönem dönem farklı önlemler alınmıştır. Sürecin iyileştirmeye sonulanıp sonulanmadığını belirlemek için istatistiksel süreç kontrolü ve kontrol çizelgeleri ölüm sayılarının artması durumunda mümkün olan en kısa sürede tepki verilebilmesi için süreçlerin izlenmesine ve böylece performanslarını sürdürebilmeleri ve aynı zamanda performanstaki rastgele dalgalanmalar olabilecek şeylere aşırı tepki vererek zaman ve enerji israfından korunmasını sağlayabilir.

Kaynaklar

- Atay, L. (2020). COVID-19 Salgını ve Turizme Etkileri. 3, 168–172.
- Aydın, Z. B., & Arıkan Kargı, V. S. (2018). İstatistiksel Kalite Kontrol Teknikleri İle Otomotiv Sektöründe Bir Uygulama. *Yönetim Ve Ekonomi Araştırmaları Dergisi*, March, 41–63. <https://doi.org/10.11611/yead.332129>
- Baydoğan, Çetin ve Orbay, (2014). R, *inet-tr'14*, İzmir - 27/11/2014
- Can, M. (2007). İstatistiksel Süreç Kontrolünde Deney Tasarımlı Süreç Optimizasyonu. *Procedia Manufacturing*, 30(22 Jan), 588–595.
- Covid19 Aşı. (2021, haziran 5). Covid19 Aşı: <https://covid19asi.saglik.gov.tr/> adresinden alındı
- Eissa M.(2020) SPC and Evaluation of Analyst Performance | *Pharma Articles [Internet]. Pharmafocusasia.com*. 2020 [cited 9 August 2020]. Available from: <https://www.pharmafocusasia.com/articles/spc-and-evaluation-of-analyst-performance>
- Eissa, M. (2019a). Application of control charts for non-normally distributed data using statistical software program: A technical case study. *World Journal of Advanced Research and Reviews*. 2019; 1(1):039-48.
- Eissa ME,(2019b). Surveillance of Cyclospora cayetanensis Epidemics in USA from Long-Term National Outbreaks Reporting System-Based Monitoring: An Observational Study Using Statistical Process Control Methodologies. *Iberoamerican Journal of Medicine*. 2019 Nov 21;2(1):4-9.
- İçişleri Bakanlığı. (2021, Haziran 5). İçişleri Bakanlığı: <https://www.icisleri.gov.tr/> adresinden alındı
- Ndong Ovono Ekouma, M. G. (2019). İstatistiksel Süreç Kontrolü: Bir Çağrı Merkezinde Uygulama, *Sakarya Üniversitesi İşletme Enstitüsü, Yüksek Lisans Tezi*.
- Orçanlı, K. (2019). Kalite Kontrol Grafiklerinde R Programlama Dilinin Kullanımı İle İlgili İçerik Analizi. *OPUS Uluslararası Toplum Araştırmaları Dergisi*, 12–14. <https://doi.org/10.26466/opus.589423>
- Patır, S. (2009). İstatistiksel Proses Kontrol Teknikleri Ve Kontrol Grafiklerinin Malatya'daki Bir Tekstil (İplik Dokuma) İşletmesinde Bobin Sarım Kontrolüne Uygulanması, *Sosyal Ekonomik Araştırma Dergisi*, 231-250
- Shomran, H. T. (2013). Simulation Study Of Statistical Process Control Schemes Performance For Monitoring Variation, *Faculty of Mechanical and Manufacturing Engineering University Tun Hussein Malaysia, Masters of Engineering in Mechanical*
- Stijn, M. Van. (2018). Change Detection In System Parameters Of Lithography Machines, *Technische Universiteit Eindhoven University Of Technology, Department Of Mathematics And Computer Science*
- Turan, A., & Çelikyay, H. (2020). Türkiye’de Covid-19 İle Mücadele: Politikalar Ve Aktörler. *Uluslararası Yönetim Akademisi Dergisi*, 1–25. <https://doi.org/10.33712/Mana.733482>
- Yıldırım, H., & Karaca, E. (2013). Üretim Sürecinde İstatistiksel Proses Kontrol (İpk) Uygulamaları Ve Elektronik Sektöründe Bir İnceleme, *Marmara Üniversitesi, İşletme Fakültesi, İşletme Bölümü*, 77–87.
- WEC (1956). Statistical Quality Control Handbook. Western Electric Company.
- World Health Organization (WHO). World situation report; 2020.

Arctanh-Exponentiated Exponential Distribution and Applications

Mustafa Ç. KORKMAZ¹, Kadir KARAKAYA², Yunus AKDOĞAN^{2*}

¹Artvin Çoruh University Department of Measurement and Evaluation, Artvin, Turkey

²Selçuk University, Science Faculty, Department of Statistics, Konya, Turkey

*Sunucu yazar e-posta: yakdogan@selcuk.edu.tr

Abstract

In this study, a new unit interval distribution called inverse hyperbolic tangent exponentiated exponential distribution, whose cumulative distribution function is $F(x) = (1 - \exp(-\lambda \arctan h(x)))^\alpha$, is proposed with the inverse hyperbolic tangent transform. We consider five frequentist estimation methods, namely, the method of maximum likelihood, the method of least square, the method of weighted least square, the method of Anderson-Darling, and the method of Cramer von-Mises to estimate the parameters inverse hyperbolic tangent exponentiated exponential distribution. An extensive Monte Carlo simulation study is conducted to observe the performance of the five estimators using bias and mean square error. We showed the resilience of the introduced distribution over existing well-known distributions by using real dataset applications.

Key words: Inverse hyperbolic tangent, exponentiated exponential distribution, Monte Carlo simulation, estimation, data analysis

**Lisansüstü Öğrenci Alım Mülakat Kriterlerinin Analitik Serim Süreci ile
Ağırlıklandırılması**

Abdurrahman YILDIZ^{1*}

¹Kütahya Dumlupınar Üniversitesi, Mühendislik Fakültesi, Endüstri Mühendisliği Bölümü, Kütahya,
Türkiye

*Sunucu yazar e-posta: abdurrahman.yildiz@dpu.edu.tr

Özet

Lisansüstü programlara öğrenci alımında, ALES Puanı, Yabancı Dil Puanı ve Not Ortalaması, genelde kullanılan ortak kriterler olarak karşımıza çıkmaktadır. Ancak bazı üniversitelerde bu kriterlerin yanında mülakat puanı ve/veya bilim sınav puanı kriterleri de kullanılmaktadır. Kütahya Dumlupınar Üniversitesi'nde ALES Puanı, Yabancı Dil Puanı, Not Ortalaması ile birlikte Mülakat Puanı aday puanı hesaplanmasında kullanılmaktadır.

Bu çalışmada Kütahya Dumlupınar Üniversitesi Endüstri Mühendisliği Anabilim Dalı Yüksek Lisans Programına öğrenci alımı esnasında, mülakat puanının belirlenmesinde standartlaştırılmış bir yapının oluşturulabilmesi amacıyla, kullanılacak kriterlerin ve bu kriterlerin ağırlık değerlerinin belirlenmesi gerçekleştirilmiştir. Buna göre beş adet kriter belirlenmiş olup, bu kriterler;

- 1. Akademik çalışma yapmaya ilgi düzeyi,*
- 2. Bir işte çalışıp çalışmadığı ve iş yoğunluğu,*
- 3. Kodlama bilgisi ve düzeyi,*
- 4. Yapay Zeka ve yeni tekniklere ilgi ve bu teknikler hakkındaki bilgi düzeyi,*
- 5. Lisans tez çalışmasının kalite düzeyidir.*

Belirlenen kriterlerin ağırlıklarının belirlenmesi için Analitik Serim Süreci (Analytic Network Process-ANP) kullanılmıştır. ANP, Çok Kriterli Karar Verme yaklaşımlarından biri olup, ilgilenilen unsur grupları (kriterler, alt kriterler vd.) arasında ve içinde her iki yönlü etkileşimi dikkate alabilen ve bu unsurların ağırlıklarının (önem derecelerinin) belirlenmesini sağlayan bir yaklaşımdır. Önce ilgilenilen unsurlar arasındaki etkileşimlerin belirlenmesi, sonrasında bu unsurlar arasında ikili karşılaştırma yaparak değerlendirme esasına dayanır. Yapılan değerlendirmenin geçerli olabilmesi için ise, oluşturulan karşılaştırma matrislerinin tutarlılık kontrolünden geçirilmesi ve tutarlı olmaları temel şarttır. Çalışmada önerilen kriterler ve hesaplanan kriter ağırlıkları kullanılarak, adayların mülakat puanları aynı ölçütlere göre değerlendirilerek standartlaştırılmış bir şekilde belirlenebilecektir.

Anahtar kelimeler: Öğrenci alımı, Analitik serim süreci, Çok kriterli karar verme

Öğrenci Elektronik Belge Yönetim Sistemi Tasarım/Seçim Kriterlerinin DEMATEL ile Değerlendirilmesi

Abdurrahman YILDIZ^{1*}

¹Kütahya Dumlupınar Üniversitesi, Mühendislik Fakültesi, Endüstri Mühendisliği Bölümü, Kütahya, Türkiye

*Sunucu yazar e-posta: abdurrahman.yildiz@dpu.edu.tr

Özet

COVID-19 pandemisi nedeniyle, üniversitelerde birçok faaliyet ve süreç kökten değişmek durumunda kalmıştır. Bunlardan bir tanesi de öğrencilerin başvuru yapma, dilekçe vb. evraklarını teslim etme süreçleridir. Öğrencilerin şehir dışında olmaları, evrak alışverişinde virüs bulaşma riski vb. nedenlerle, başvuru ve evrak teslimi doğal olarak elektronik ortama aktarılmıştır. Elektronik ortamda evrak teslimi için, e-posta'nın en fazla kullanılan yaklaşım olduğu gözlenmekte olup, hazırlanan online formlara evrakların yüklenmesi ve Öğrenci Bilgi Sistemi (OBS) mesaj sistemi de kullanılan diğer yaklaşımlar olarak gözlemlenmektedir. Ancak bu yaklaşımların her biri pratikte iş görüyor olmasına karşın, evrak takibi ve uygulama birliği açısından bazı sorunlar da oluşturabilmektedir.

Bunun yanında Pandemi nedeniyle üniversitelerde yaşanan bu zorunlu değişimlerden bazılarının ise pandemi öncesine göre daha faydalı ve kullanışlı olduğu, pandemi sonrasında da kullanımına devam edilmesi isteği, gözlenen bir gerçektir. Bunlardan biri de evrak teslimi ve başvuruların elektronik ortamda yapılmasına devam edilmesidir. Ancak pandeminin sona ereceği ve şimdilik kullanılan bu sistemin geçici olduğu düşüncesiyle uygulama birliğinin oluşmaması ve buna bağlı olarak oluşan durumlar, sorun olarak karşımıza çıkmaktadır. Bu sorunların ortadan kaldırılması için standardize edilmiş bir yapının oluşturulması gerekmektedir.

Oluşturulması gereken yapı, halihazırda üniversite akademik ve idari yönetim sisteminde kullanılan Elektronik Belge Yönetim Sistemi (EBYS) benzeri bir yapı olmalıdır. Bilindiği üzere, EBYS'de her türlü resmi yazışma ve evrak gönderimi, elektronik ortamda yapılmakta, evraklarda e-imza kullanılmaktadır. Ancak üniversitedeki mevcut EBYS'lerin, hem yüksek maliyetli olması, hem de binlerce öğrenci için e-imza oluşturulmasının uygulanabilir olmaması nedeniyle, bu amaçla kullanımı için öğrencilere de açılması mümkün görünmemektedir. Bunun, ancak OBS'ye entegre edilecek bir ara sistemle mümkün olabileceği değerlendirilmektedir.

Bu çalışmada OBS'ye entegre edilecek "Öğrenci EBYS"nin tasarımında veya seçiminde dikkate alınacak kriterlerin oluşturulması ve bu kriterlerin birbirleri ile etkileşimlerinin DEMATEL (Decision Making Trial And Evaluation Laboratory) Yaklaşımı ile belirlenmesi ile ilgilenilmiştir. Yazılım özellikleri (kullanıcı arayüzü, evrak yönetimi, resmi evrak yükleme-alma prosedürü, evrak takibi, evrak yedekleme, evrak arama, raporlama), Veritabanı özellikleri ve Yönetimi, Güvenlik, OBS vb. diğer sistemlerle entegrasyon ve Maliyet temel kriterleri ve alt kriterlerinin değerlendirmede kullanımı önerilmektedir. DEMATEL, ilgilenilen unsurların (kriterler vd.) birbirleriyle etkileşimini (etkileyen-etkilenen), bu etkileşimin derecesini ve ağırlıklarını belirlemeye imkân veren bir Çok Kriterli Karar Verme yaklaşım olduğu için kullanımı tercih edilmiştir. Önerilen kriterler ve DEMATEL ile belirlenen kriter ağırlıkları kullanılarak, "Öğrenci EBYS" tasarımı yapılabileceği veya seçimi gerçekleştirilebileceği düşünülmektedir.

Anahtar kelimeler: Elektronik belge yönetim sistemi, DEMATEL, Çok kriterli karar verme

Panel Veri Analizi ile Küresel Rekabet Endeksi, İhracat ve Büyüme Oranının İnovasyon Üzerine Etkisi

Zeynep ÖZTÜRK¹, Aysel ÖZAYDIN^{2*}

¹Artvin Çoruh Üniversitesi, Hopa İİBF, İşletme Bölümü, 08200, Artvin, Hopa, Türkiye

²Artvin Çoruh Üniversitesi, Lisansüstü Eğitim Enstitüsü, İşletme Bölümü, 08200, Artvin, Türkiye

*Sunucu yazar e-posta: ayselozaydin@hotmail.com

Özet

Çağımızda ülkelerin inovasyon konusundaki başarıları, rekabet avantajı elde etmesini sağlayan önemli unsurlardan biri olarak görülmektedir. Küresel inovasyon endeksi (KİE) 2007 yılında ülkelerin inovasyon düzeylerini belirlemek amacıyla Avrupa İşletme Yönetimi Enstitüsü (INSEAD) tarafından geliştirilmiştir. Küresel rekabet endeksi, ülkelerin rekabetçilik durumunu takip etmek için Dünya Ekonomik Forumu tarafından 1979 yılından beri belirli dönemlerde hesaplanmakta ve sürekli olarak güncellenmektedir. Bireyler, ülkeler, firmalar ve hane halkları gibi birimlere ait ya da yatay kesit gözlemlerinin belli bir zaman döneminde bir araya getirilerek ekonomik ilişkilerin tahmin edilmesi yöntemine panel veri analizi denilmektedir. Bu çalışmanın amacı, panel veri analizi ile Küresel Rekabet Endeksi, İhracat ve büyüme oranının Küresel İnovasyon Endeksi üzerine etkisini araştırmaktır. 2011 ile 2017 yılları arasında 25 ülkeye ait veriler alınmış olup elde edilen veriler R yazılımı ile R Studio programı kullanılarak panel veri analizi ile analiz edilmiştir. Sabit etkili model en uygun model olarak bulunmuştur. Modelde Değişen Varyans ve Otokorelasyon sorunu olduğundan robust tahmin edicilerinden Arellano ile standart hata düzeltilmesi yapılarak tahminler elde edilmiştir. Küresel Rekabet Endeksinin, Küresel İnovasyon Endeksi üzerinde istatistiksel olarak anlamlı etkisi olduğu sonucuna varılmıştır.

Anahtar kelimeler: Küresel inovasyon endeksi, Küresel rekabet endeksi, Panel veri analizi

The Effect of Global Competitiveness Index, Export and Growth Rate on Innovation with Panel Data Analysis

Abstract

In our age, the success of countries in innovation is seen as one of the important factors that enable them to gain competitive advantage. The global innovation index (CII) was developed by the European Institute of Business Management (INSEAD) in 2007 to determine the innovation levels of countries. The global competitiveness index has been calculated periodically since 1979 by the World Economic Forum in order to monitor the competitiveness of countries and is constantly updated. Panel data analysis is the method of estimating economic relations by bringing together horizontal section observations of units such as individuals, countries, firms and households in a certain period of time. The aim of this study is to investigate the effect of Global Competitiveness Index, Export and growth rate on Global Innovation Index with panel data analysis.

Data from 25 countries were taken between 2011 and 2017, and the obtained data were analyzed by panel data analysis using R software and R Studio program. The fixed effect model was found to be the most suitable model. Since there is a problem of varying variance and autocorrelation in the model, the

estimations were obtained by correcting the standard error with Arellano, one of the robust estimators. It is concluded that the global competitiveness index has a statistically significant effect on the Global Innovation index

Key words: Global innovation index, Global competitiveness index, Panel data analysis

GİRİŞ

Çağımızda yaşanan teknolojik gelişmeler, toplumların gerçekleştirmeyi planladıkları ekonomik ve sosyal hedefler konusunda önemli rol oynamaktadır. İngiltere’de 18. yüzyılın ikinci yarısından itibaren başlayan Endüstri Devriminden sonra Ar-Ge laboratuvarları uzmanlaşmaya başlamış daha kurumsal bir yapı haline dönüşmüştürler. Endüstri devriminden sonra ülkelerin büyüüp gelişmesi, diğer ülkelerle rekabet edebilmesi için bilim, teknoloji ile inovasyon önem arz etmeye başlamıştır. Zaman içerisinde bilim ve teknoloji de yaşanan gelişmeler iktisadi analizlerde de kullanılmaya başlanmıştır.

Ülkeler tarafından iyi şekilde geliştirilmiş ulusal inovasyon sistemleri, uluslararası rekabetin ve ülkelerin kalkınma düzeylerinin önemli göstergelerindendir. Zaman içerisinde gelişen bilim ve teknoloji sayesinde ülkeler hakkında yeterince veri toplanabilmiş ve bu sayede ülkelerin inovasyon düzeylerini gösteren çalışmalar yapılmaya başlanmıştır. Genelde yapılan çalışmalar gelişmiş ülkelerde gerçekleşmekte az gelişmiş ya da gelişmekte olan ülkelerde çok az sayıda yapılmaktadır (Zuhal,2020).

25 ülkeye ait Küresel İnovasyon Endeksinin üzerine, Küresel Rekabet Endeksi, İhracat ve GSYİH değerlerinin etkisi olup olmadığının ampirik analizinin yapılması amaçlanmaktadır. Çalışmanın birinci bölümünde KİE, KRE, İhracat ve GSYİH hakkında genel bilgi verilerek genel bir çerçeve oluşturulacaktır. İkinci bölümünde ülkelerin iki endeks değerini karşılaştırılacaktır. Son bölümde ise çalışmadan elde edilen sonuçlar ile veri analizi yapılacaktır.

Literatür Taraması

Zuhal (2020) bu çalışmada, ülkelere ait ulusal inovasyon kapasitesi belirleyicisinin gelişmiş ülkelerle gelişmekte olan ülkelerin karşılaştırmasını yapmayı amaçlamıştır. Bu çalışmalardaki önermelerden ve yöntem denemelerinden yola çıkılarak ulusal inovasyon düzeyi hesaplanmakta ve temsili olarak bu değişken kullanılmaktadır. Ülkelerin 1996 ile 2016 yıllarına ait verileri ele alınmıştır. Panel veri analizi ile 31 tanesi gelişmiş, 18 tanesi gelişmekte olan ülkenin inovasyon kapasitesi endeksi seçilerek analiz edilmiştir. Ulusal teknolojik yetenek ve altyapı faktörleri ulusal inovasyon düzeyi belirleyicilerindendir ve bu faktörün gelişmekte olan ülkelerde daha zayıf bir etkiye sahip olduğu sonucuna ulaşılmıştır.

Ay Türkmen ve Aynaoglu (2017) bu çalışmada, 2009 yılı KİE sıralamasında yer alan ilk 29 ülkenin, regresyon analizi yöntemi ile KRE parametrelerinin KİE üzerinde ne yönde bir etkiye sahip olduğunu tespit etmeyi amaçlamıştır. KRE ve KİE’nin 2009-2017 dönemine ait veriler kullanılmıştır. İnovasyon, yüksek eğitim ve öğretim ve emek piyasasının etkinliği, KRE ve KİE’nin parametreleri olup bu değişkenler arasında tam pozitif yönlü bir ilişki olduğu sonucuna ulaşılmıştır. Makroekonomik çevre parametresi ile aralarında ise en düşük etkinin olduğu görülmektedir.

Torun (2016) bu çalışmanın temel amacı, inovasyon algılarının, inovasyon sürecindeki liderlik tarzlarına etkisini ve inovasyon sürecindeki liderlik tarzlarının da işletmenin inovasyon performansına etkisini incelemektir. Araştırmaya katılan yönetici ve işletme sahipleri Düzce illinde çalışan KOBİ’lerden seçilmiştir. İstatistiksel analiz olarak; betimleyici istatistikler, korelasyon analizi ve regresyon analizi kullanılmıştır. Elde edilen sonuçlara göre en yaygın inovasyon algısı, inovasyonun, ürün kalitesini artıracığı algısı olmuştur. En yaygın görülen inovasyon sürecindeki liderlik tarzı, dönüştürücü liderlik tarzıdır.

Aynaoglu (2018), bu çalışmada, KRE ile AR-GE/GSYİH, işgücü verimliliği, kişi başı GSYİH, patent sayıları ve İhracat/GSYİH arasındaki ilişkiyi incelemiştir. 2017 ile 2018 arası dönemi kapsayan

Avrupa Birliği üye ülkelerinin verileri, sabit ekiler modeli ve panel veri seti kullanılarak analiz edilmiştir. Araştırma sonucuna göre ihracatın GSYH'deki payının, kişi başı GSYH'nin Küresel Rekabet Endeksi üzerinde pozitif yönde etkisi olduğu tespit edilmiştir. Patent sayılarının, AR-GE'nin GSYH'deki payının ve işgücü verimliliğinin rekabetçilik üzerinde negatif yönlü etkisi olduğu ortaya çıkmıştır.

Çetin ve Süt' ün (2018) yapmış olduğu bu çalışmada, ekonomik büyüme üzerinde Ar-Ge, araştırmacı ve patent sayısı gibi tek bileşenli değişkenler ile toplam faktör verimliliği ile inovasyon endeksleri gibi birçok bileşenden oluşan değişkenlerin etkilerini panel veri yöntemiyle analiz etmiştir. OECD, The Conference Board, EC, WEF ve Dünya Bankası veri tabanlarından 2006-2015 yıllarını kapsayan döneme ait veriler kullanılmıştır. Bu çalışmada, ekonomik büyüme üzerinde inovasyonun önemli etkisi olduğu tespit edilmiştir. İnovasyon endeksleri ile faktör verimliliğinin inovasyonu temsil eden en iyi göstergeler olduğu sonucuna ulaşılmıştır.

Çalıpınar ve Baç (2007) bu çalışmasında KOBİ'lerin karakteristik özelliklerinin inovasyon yapma sayılarına etkisi olup olmadığını araştırmayı amaçlanmıştır. İnovasyon faaliyetlerine etki eden reklam harcamaları, çalışan sayıları, ihracaat gelirleri, patent harcamaları, Ar-Ge harcamaları, kurulan dış ortaklıklar ve işletme yaşı verileri anket ile toplanmış ve veriler arasındaki ilişki doğrusal regresyon modeli kullanılarak analiz edilmiştir. İşletmenin yaşı ve çalışan sayısı ile inovasyon değişkenleri arasında negatif yönde ilişki bulunduğu tespit edilmiştir. Ortalama Ar-Ge harcamaları dışında bulunan diğer faktörler ile inovasyon değişkenleri arasında ise, istatistiksel olarak anlamlı bir ilişki bulunmuştur.

Küresel İnovasyon Endeks (KİE)

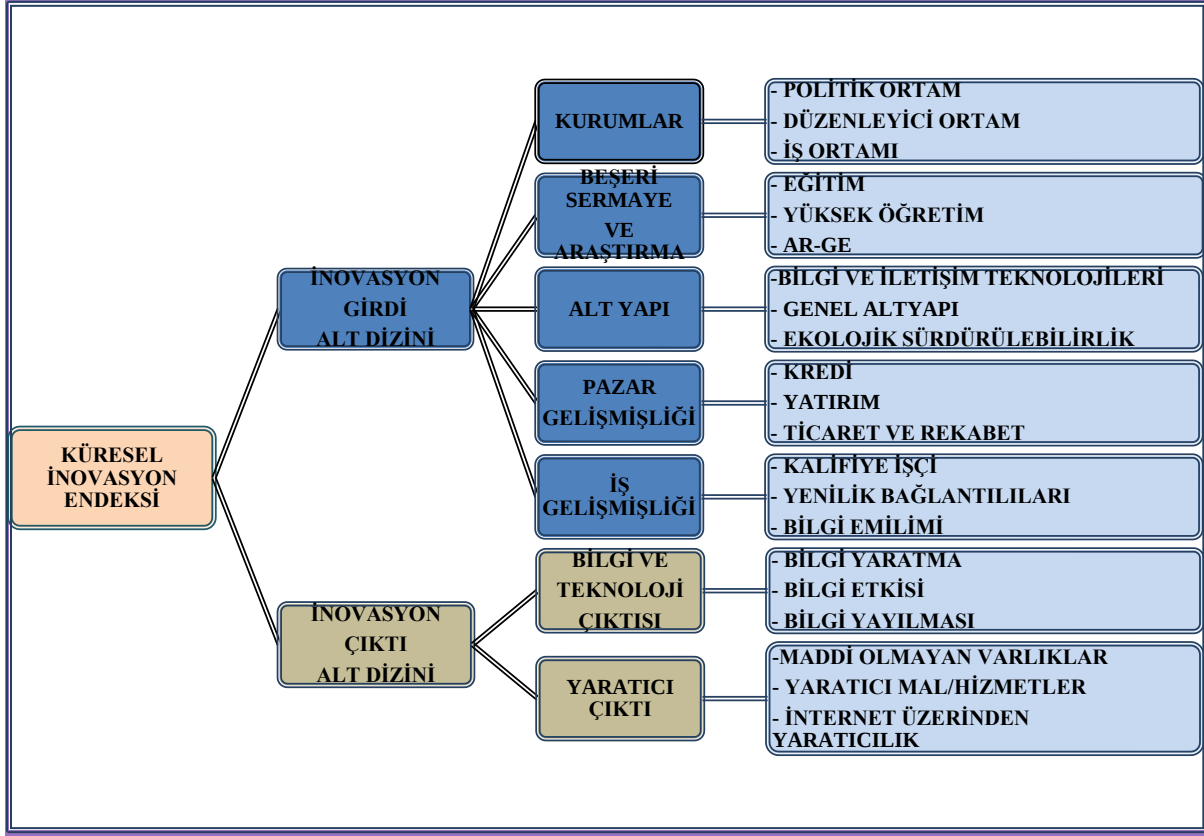
İnovasyon, yeni bir şeyin ortaya konulması suretiyle mevcut durumla ilgili bir takım değişikliklerin gerçekleştirilmesi sürecidir (Yılmaz, 2015). İnovasyon endeksi, araştırmacıların, kurumların, işletmelerin ve araştırma için seçilen bölgelerin inovasyon kapasitelerinin özet olarak sayısal göstergesidir. İnovasyon endeksi, ülkelerin inovasyon performanslarını, teknoloji düzeyini ve ekonomi seviyesini gösterdiği gibi ulusal rekabet edebilirliği de ölçer (Hancioğlu,2016). İnovasyon rekabet avantajı ve kalkınmanın önemli bir anahtarıdır. KİE, ülkelerin inovasyon politikaları ve inovasyon uygulamaları ile hem zayıf hemde güçlü yönlerini ortaya koymak için tasarlanmıştır.

İnovasyon; sadece bölgeler ve ülkeler için değil bireyler ve işletmeler içinde rekabetçiliği ve verimliliği artıran bir etkiye sahiptir. İnovasyon, ülkelerin verimliliği, rekabet edebilirliği, istihdam performansları ve üretim kapasitelerinin altında yatan ana faktörlerden olduğundan, ülkelerin gelişmişlik düzeylerini belirlemede kullanılmaya başlamıştır (Ay Türkmen, Aynaoglu,2017).

2007 yılında Avrupa İşletme Yönetimi Enstitüsü (INSEAD) tarafından, ülkelerin inovasyon seviyelerini belirlemek için Küresel İnovasyon Endeksi (KİE) geliştirilmiştir. INSEAD'ın yaptığı bu çalışmaya Dünya Fikri Mülkiyetler Örgütü (WIPO) 2011 yılında ve Cornell Üniversitesi'nde 2013 yılında da dâhil olmuştur. Küresel İnovasyon Endeksi hesaplaması, 140'tan fazla ülke verilerinden yararlanır ve 81 göstergeden yola çıkarak elde edilir (Ay Türkmen, Aynaoglu,2017).

KİE'nin iki alt endeksi bulunmaktadır. Birincisi girdi alt endeksi, ikincisi çıktı alt endeksidir. İnovasyon girdi alt endeksi, beş temel değişken, inovasyon çıktı alt endeksi ise, iki temel değişkenden oluşmaktadır. Burada her bir değişken de üç bileşen vardır. Değişkenlerle ilgili veriler; kamu ve özel sektör kuruluşlarından elde edilen nicel veriler, alanında uzman kuruluşların yayınladıkları diğer endekslerin verileri, Dünya Ekonomik Forumu yönetici anketi verileri kullanılarak elde edilmektedir (Ay Türkmen, Aynaoglu,2017). Küresel İnovasyon Endeksi bileşenleri Çizelge-1'de gösterilmektedir.

Çizelge 1. Küresel inovasyon endeksi bileşenleri



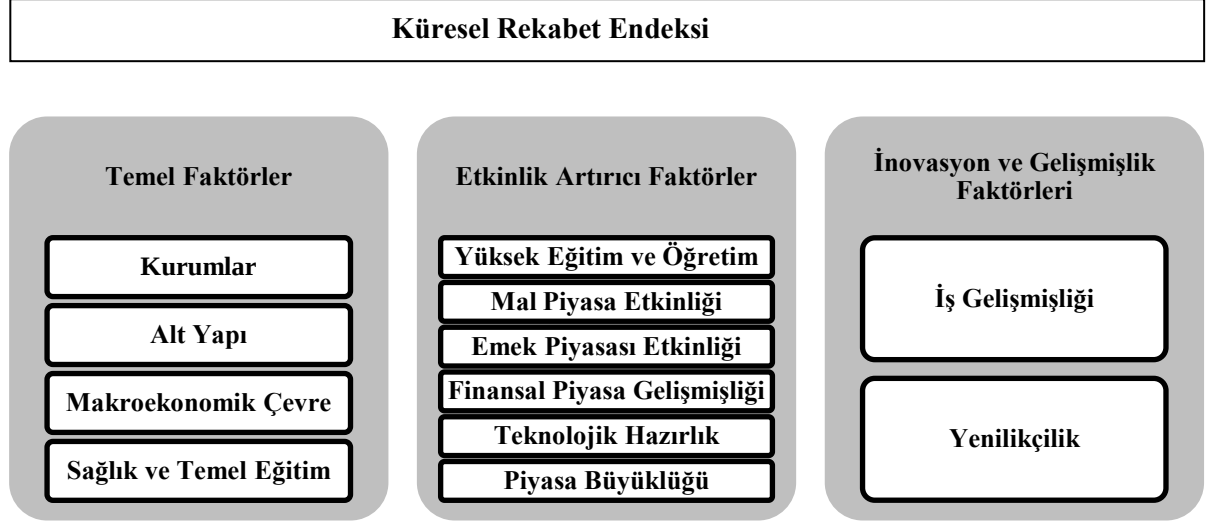
Küresel Rekabet Endeks (KRE)

Küresel Rekabet Endeksi 1979 yılından itibaren belirli aralıklarla Dünya Ekonomik Forumu (WEF) tarafından hesaplanarak yayınlanmaktadır. Dünya Ekonomik Forumu tarafından, bir ülkenin verimliliğini belirleyen kurumların, politikaların ve faktör setlerinin, rekabetçilik olarak tanımlaması yapılmaktadır (Ay Türkmen ve Aynaoglu, 2017).

Küresel Rekabet Endeksi hesaplamalarında yaklaşık olarak 120 ile 148 arasında ülke ele alınmaktadır. Bu ülkelere ait yaklaşık olarak 20.000 civarında veri kullanılmaktadır. Verilerin temin edilmesi aşamasında, Dünya Bankası, Ekonomik İstihbarat Birimi, bakanlıklar, bölgesel kalkınma bankaları, enstitüler, ajanslar, ulusal istatistik birimlerinin kaynaklarından faydalanılmaktadır. KRE, ülkelerin rekabet düzeylerinin karşılaştırılabilmesi açısından önemlidir. Kurumlardan ve kuruluşlardan küresel anlamda veri sağlanamadığından WEF anket yoluyla veri elde edebilmektedir. Ülkemizde bu çalışma Sabancı Üniversitesi Rekabetçilik Forumu ile Türk Sanayicileri ve İşadamları Derneği (TÜSİAD) tarafından yapılmaktadır (Ay Türkmen, Aynaoglu, 2017).

Küresel Rekabet Endeksi, temel faktörler, etkinlik artırıcı faktörler, inovasyon ve gelişmişlik faktörleri olmak üzere üç ana faktöre ayrılmıştır. Temel faktör bileşeni, kurumlar, altyapı, makroekonomik çevre, sağlık ve temel eğitimi kapsamaktadır. Etkinlik artırıcı faktör bileşenleri, yüksek eğitim ve öğretim, emek piyasalarında etkinlik, mal piyasalarında etkinlik, finansal piyasalarda gelişmişlik, teknolojik hazırlık ve piyasa büyüklüğünden oluşmaktadır. İnovasyon ve gelişmişlik faktörü bileşenleri ise, yenilikçilik ve iş gelişmişliğidir (Ay Türkmen, Aynaoglu, 2017). Bilim ve teknolojinin hızla ilerlediği günümüzde ülkelerin gelişip büyümelerine katkı sağlayacak politikalar geliştirmesi aşamalarında Küresel Rekabet Endeksi bir yol göstericisi olmaktadır. KRE ana faktör grupları aşağıda yer alan Çizelge-2 de gösterilmiştir.

Çizelge 2. Küresel rekabet endeksi 3 temel faktör etkisi



KRE'nin alt endeksleri tüm ülkeler üzerinde geçerlidir ancak ülkeler bu bileşenlere aynı derecede önem vermemektedirler. Hedeflenen rekabet düzeyine ulaşmak için Amerika Birleşik Devletlerinin izleyeceği yol ile Türkiye'nin izleyeceği yol farklı olabilmektedir. Ülkelerin izlediği yolların farklı olması nedeniyle ülkelerin rekabetçilik düzeyleri sınıflandırılmaktadır. Faktör çeşitli, etkinlik çeşitli ve yenilikçilik çeşitli olmak üzere üç aşamadan oluşmaktadır. Bu sınıflandırmalar Çizelge 3'de gösterilmiştir.

Çizelge 3. Küresel rekabetçilik endeksi raporu

1. Düzey	Faktör Çeşitli	Bu düzeyde ekonomi faktör çeşitli olup ülkeler faktör kaynaklarına (doğal kaynaklar ve temel olarak vasıfsız işgücü) göre rekabet etmektedir. Faktör çeşitli modeli tercih eden ekonomilerde kurum ve kuruluşlar, makroekonomik çevre, alt yapı şartları, temel eğitim ve sağlık alanları temel belirleyicidir.
2. Düzey	Etkinlik Çeşitli	Bu düzeyde gelişmişlik artmıştır. Üretkenlik artmıştır. Ücretler yüksektir. Ülkeler rekabetçi gelişmişlikte etkinlik çeşitli düzeye geçmiştir. Yüksek eğitim ve öğretim, gelişmiş finans piyasaları, etkin mal piyasaları, teknoloji, iyi işleyen emek piyasaları ile geniş iç ve dış piyasa düzeyleri de artmıştır.
3. Düzey	Yenilikçilik Çeşitli	Yenilikçilik çeşitli aşama en son aşamadır. Bu aşamada kurum ve kuruluşlar en gelişmiş makine vb teçhizatlar ile üretim için yeni ve farklı ürünler üreterek rekabet etmektedirler.

Uluslararası arenada ülkelerin birbirleri ile rekabet ederken kendilerini ortaya çıkaracak rekabet güçlerini artıracak faktörlere yönelerek onlara daha fazla önem vermeleri gerektiği aşikardır. Rekabet gösterge ağırlıklarına göre ülkeler, birinci aşama faktör çeşitli, birinci aşamadan ikinci aşamaya geçiş, ikinci aşama etkinlik çeşitli, ikinci aşamadan üçüncü aşamaya geçiş ve üçüncü aşama yenilikçilik çeşitli olmak üzere 5 gruba ayrılmaktadır. Bir ülkenin rekabet gücü hesaplanırken, temel faktörlerin ağırlığı, etkinliği artıran faktörlerin ağırlığı, inovasyon ve uzmanlaşmayı artıran faktörlerin ağırlığı ve ülkelerin kişi başına düşen gelir düzeyleri kategorileri oranlarına bakılmaktadır. Rekabet gösterge ağırlıkları ve oranları ayrıntılı olarak Çizelge-4'de gösterilmiştir (Ay Türkmen ve Aynaoglu,2017).

Çizelge 4. Rekabet gösterge ağırlıkları

	1.Aşama Faktör Çekişli	1.Aşamadan 2.Aşamaya Geçiş	2.Aşama Etkinlik Çekişli	2.Aşamadan 3.Aşamaya Geçiş	3.Aşama Yenilikçilik Çekişli
Kişi başına GSYH (Amerikan Doları)	< 2000	2000-2999	3000-8999	9000-17000	>17000
Temel faktörlerin ağırlığı	%60	%40-60	%40	%20-40	%20
Etkinliği artıran faktörlerin ağırlığı	%35	%35-50	%50	%50	%50
İnovasyon ve uzmanlaşmayı artıran faktörlerin ağırlığı	%5	%5-10	%10	%10-30	%30
Toplam (%)	%100	%100	%100	%100	%100

Gayrisafi Yurtiçi Hasıla (GSYH)

Bir ülkenin sınırları içerisinde belli bir dönemi kapsayan zamanda üretilen mal ve hizmetlerin tamamının para birimi cinsinden değerine Gayrisafi Yurtiçi Hasıla denmektedir. Ülkelerin gelişmişlik seviyelerinin belirlenebilmesi adına kullanılan en önemli ölçütlerden biridir. Bir ülkenin para birimiyle sabit fiyatlarla kişi başına GSYH kullanılarak yıllara göre gelişmişlik düzeyi ölçülürken, ülkelerin uluslararası gelişmişlik seviyelerini karşılaştırmak için ise satın alma gücü paritesine göre düzenlenmiş Gayrisafi Yurtiçi Hasıla kullanılır (Aynaoglu, 2018).

Bir ülkenin üretim artışı ile birlikte, ihracatında, istihdamında ve GSYH’inde, sanayileşme oranının artması ve bu arada tarıma ait oranında azalması ülkenin ekonomik olarak kalkınmasının en önemli göstergelerindendir (Yılmaz ve İncekaş,2018).

GSYH İçindeki İhracat Yüzdesi

Bir ülkede kişi ve kuruluşlar tarafından üretilen serbest dolaşımda bulunan mal ve hizmetlerin yabancı ülkelere satılması ya da gönderilmesine ihracat denir. İhracat yapabilme yeteneği, ülkenin sahip olduğu üretimi ve doğal kaynaklarını en etkin biçimde değerlendirmesi, ülkedeki yaşam düzeyini arttırması ve ayrıca verimlilik kapasitesini yükseltmesini gerektirir (Aynaoglu, 2018).

Ülkelerin ihracat gelirleri arttığı zaman buna bağlı olarak ithalatın ihracatı karşılama yüzdesi de yükselecektir. Ülkenin ekonomik büyümesi hız kazanacak ve GSYH ve dış ticaret dengesi gerçekleşebilecektir (Aynaoglu, 2018).

BULGULAR

Veri Seti ve Örneklem Kümesi

Ülkelerin inovasyon, teknoloji ve bilgi düzeyleri ile temel makroekonomik göstergeler arasındaki ilişkinin boyutunu bilmek geleceğe dair planlama yapmak konusunda önem arz etmektedir. Ülkelerin temel makroekonomik göstergelerinin gelişmişlikleri ve rekabet düzeyleri gelişmelerine yönelik birer göstergedirler. Bu çalışmada, 2011-2017 dönemini kapsayan 25 ülkenin küresel rekabet endeks puanı (Rek), makroekonomik gösterge olarak da GSYH içindeki ihracat yüzdesi (Ihr) ve GSYH büyüme oranı (GSYH) bağımsız değişkenler ile Küresel İnovasyon Endeks(Inv) puanı bağımlı değişken olarak alınarak panel veri seti olarak düzenlenmiştir. Küresel Rekabet Endeks puan ve Küresel İnovasyon Endeks puan(Inv) verileri Dünya Fikri Mülkiyet Örgütü sayfasından ve diğer veriler Dünya Bankası

(data.worldbank.org) 'ndan alınmıştır. Çalışmanın analizleri R yazılım paket programları kullanılarak yapılmıştır. Çalışmada yer alan 25 ülke çizelge-5'te gösterilmiştir.

Çizelge 5. Çalışmada ele alınan ülkeler

SIRA	ÜLKE	SIRA	ÜLKE
1	ALMANYA	14	İSRAİL
2	AMERİKA BİRLEŞİK DEVLETLERİ	15	İSVEÇ
3	ARJANTİN	16	İSVİÇRE
4	AZERBAYCAN	17	İTALYA
5	BREZİLYA	18	JAPONYA
6	BULGARİSTAN	19	KAZAKİSTAN
7	ÇİN	20	KENYA
8	ERMENİSTAN	21	RUSYA FEDERASYONU
9	FİNLANDİYA	22	SUUDİ ARABİSTAN
10	FRANSA	23	SLOVENYA
11	GÜRCİSTAN	24	TÜRKİYE
12	HİNDİSTAN	25	YUNANİSTAN
13	İRAN		

Panel Veri Modeli

Çalışmanın bağımlı değişkeni inovasyon endeks puanını etkileyen bazı faktörleri belirleyebilmek için oluşturulan regresyon modeli aşağıdaki gibidir:

$$\ln Inv_{it} = \beta_0 + \beta_1 \ln Re k_{it} + \beta_2 \ln I hr_{it} + \beta_3 GSYH_{it} + \varepsilon_{it}$$

Bu denklemde, $i = 1, 2, 3, \dots, N$ yatay kesit birimlerini ifade eder, $t = 1, 2, 3, \dots, t$ zaman boyutunu ve panel hata terimini ise ε ifade eder. Parametre tahminlerini elde etmek için R yazılım programı “plm” paketinden yararlanılmıştır.

Birim Kök Analizi

Belli bir dönemi kapsayan serinin istatistiksel analizi yapılmaya başlanmadan önce, serinin durağan olup olmadığına yani zaman içerisinde sabit olup olmadığına incelenmesi gerekir. Birim kök analizi yapabilmek için öncelikle panel verileri kullanıldığında yatay kesit bağımlılığı olup olmadığının tespit edilmelidir. Panel veri setinde yatay kesit bağımlılığı varlığı reddedilirse, 1. nesil birim kök testleri kullanılır. Eğer panel verilerinde yatay kesit bağımlılığı varsa 2. nesil birim kök testlerini kullanılır. Yatay kesit bağımlılığının test edilebilmesi için Breusch-Pagan (1980) ve Pasaran CD (kesitsel bağımlılık) testleri kullanılacaktır.

Yatay Kesit Bağımlılığı testi:

Breusch-Pagan LM bağımsızlık testi ve Pasaran CD (kesitsel bağımlılık) testleri, artıkların girdiler arasında korelasyon gösterip göstermediğini test etmek için kullanılır. Kesitsel bağımlılık, test sonuçlarında yanlılığa yol açabilir (aynı zamanda eşzamanlı korelasyon olarak da adlandırılır). Yatay kesit bağımlılık test sonuçları Çizelge-6'da gösterilmektedir.

Çizelge 6. Yatay Kesit Bağımlılık Test Sonuçları

Değişkenler	Breusch-Pagan LM Bağımsızlık Testi		Pasaran CD (Kesitsel Bağımlılık) Testi	
	İstatistik Değeri	Prob. değeri	İstatistik Değeri	Prob. değeri
Inv	566.43	0.000	8.5686	0.000
Rek	795.3	0.000	11.301	0.000
Ihr	940.2	0.000	0.732	0.464
GSYH	--	-	-	-

H_0 : Kesitler arasında bağımlılık yoktur. H_1 : Kesitler arasında bağımlılık vardır.

Panel bazda yapılan yatay kesit bağımlılığı (YKB) test sonuçlarına göre, Breusch-Pagan LM bağımsızlık testi tüm değişkenler için olasılık değerinin 0.05'ten küçük çıktığı tespit edilmiştir. Böylece, Breusch-Pagan LM bağımsızlık gereğince sıfır hipotez “kesitler arasında bağımlılık yoktur” şeklindedir ve sıfır hipotez reddedilmektedir. Ancak, Pasaran CD (kesitsel bağımlılık) testinde, sadece Ihr değişkeni için sıfır hipotez kabul edilmektedir. Bu durumda diğer değişkenler için panel veride kesitler arasında yatay kesit bağımlılığı söz konusudur. Yani, ülkelerden birine gelen bir birimlik etki diğer ülkeleri de etkilemektedir. Çalışmada ele alınan değişkenler için yatay kesit bağımlılığı tespit edildiği için durağanlık, ikinci nesil birim kök testlerinden Pesaran (2007) tarafından geliştirilen CADF testi ile ele alınmıştır. CADF testi, hem $T > N$ ve hem de $N > T$ durumlarında tutarlıdır. CADF testinde; sıfır hipotezi “Seriler durağan değildir.” şeklindedir. R programında “CADFtest” paketi kullanılarak elde edilmiştir.

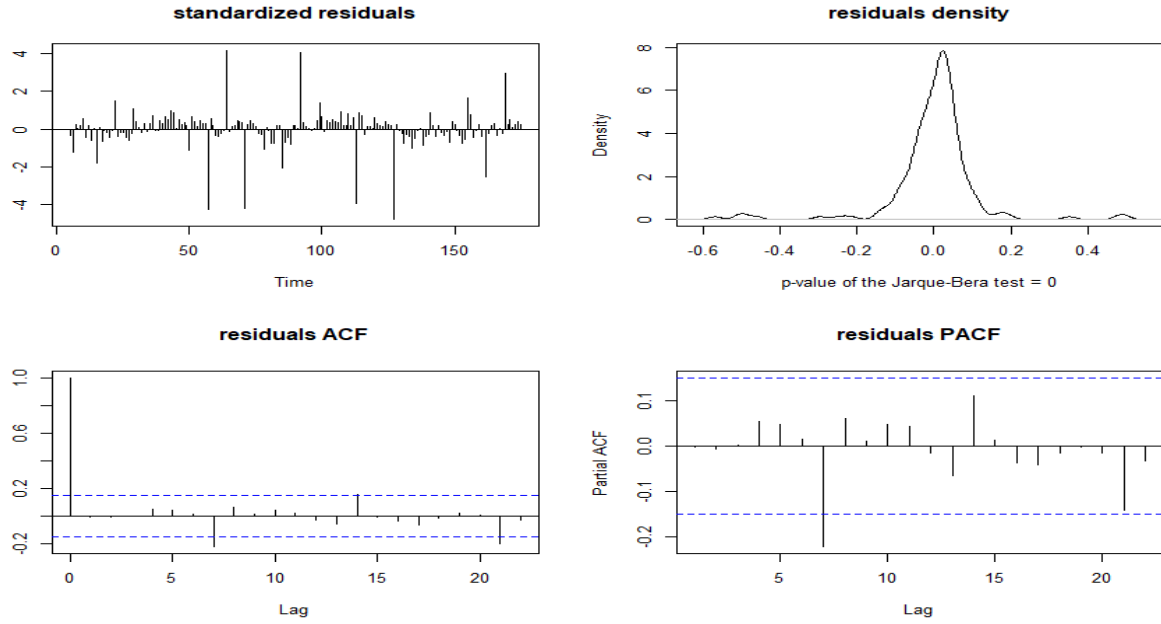
CADF Birim Kök Test Sonuçları

CADF Birim kök testi yapılmış olup sonuçlar çizelge 7’de sunulmuştur.

Çizelge 7. CADF birim kök testi.

Değişkenler	İstatistik Değeri	Prob. değeri
Inv	-3.011	0.132
Rek	-3.526	0.039
Ihr	-4.225	0.005
GSYH	-5.504	0.000
Dif(Inv)	-4.625	0.002

CADF Birim kök testi sonuçlarına göre, Rek, Ihr, GSYH değişkenleri temel düzeyde $I(0)$ durağan iken, Inv değişkeni 1.düzeyde $I(1)$ durağan sonucuna varılmıştır. Aşağıdaki Grafik-1’de Inv değişkeninin temel düzeyde durağan olmadığı görülmektedir.



Şekil-1: Grafikler

Model Seçimi

Klasik panel model de, sabit parametler ile eğim parametrelerinin hem birimlere hem de zamana göre sabit olduğu varsayılır. Havuzlanmış en küçük kareler yöntemi ile elde edilen modelin analiz sonuçları Çizelge-8’de verilmiştir.

Çizelge 8. Havuzlanmış En Küçük Kareler Yöntemi İle Elde Edilen Parametre Tahminleri

Değişkenler	Tahmin	Standart Hata	T değeri	P değeri
Sabit Terim	1.0532593	0.1269394	8.2973	0,000
lRek	1.7057316	0.0788641	21.6287	0,000
Xllhr	0.0285680	0.0199872	1.4293	0,155
XGsyh	-0.0083420	0.0030333	-2.7502	0,007
F-istatistik değeri 172.896			p değeri 0,000	

Klasik panel modelde elde edilen sonuçlara göre, sabit terim, Rek ve Gsyh değişkenleri istatistiksel olarak anlamlı sonuç bulunmuştur.

Bir panel veri modelinde hata terimi ile ilgili temel varsayımlar olarak değişen varyans ve otokorelasyon varsayımları bulunmaktadır. Değişen varyans, varyansın homojenlik varsayımının geçerli olmadığını ifade eder. Yani, hata terimlerinin varyanslarının tüm birimler için farklı olduğu ve kovaryanslarının sıfıra eşit olmama durumudur. Bu çalışmada değişen varyans varsayımı, Breusch-Pagan değişen varyans testi ile yapılmıştır

Hata terimleri arasındaki ilişki otokorelasyon varsayımını gösterir. Birim değerlerinin birbirini etkilemesi modelde ilişkiye neden olur ve model tahminlerinde sapmalar ve tutarsızlık oluşur. Bu çalışmada otokorelasyon varsayımı, Breusch-Godfrey/Wooldridge test ile yapılmıştır Sıfır hipotez, “Serinin otokorelasyonunun olmamasıdır.” Şeklinde.

Breusch-Pagan Değişen Varyans Testi

Breusch-Pagan Değişen Varyans testi yapılarak sonuçları Çizelge 9’da verilmiştir.

Çizelge 9. Breusch-Pagan Test Değişen Varyans

BP = 17.887	p-value = 0.0004642
-------------	---------------------

Otokorelasyon Testi

Otokorelasyon testi yapılarak elde edilen sonuçlar Çizelge 10’da gösterilmiştir.

Çizelge 10. Otokorelasyon

chisq = 115.36	p-value = 2.2e-16
----------------	-------------------

Çizelge 9 ve Çizelge 10 a göre modelde değişen varyans ve otokorelasyon sorunu olduğundan robust standart hatalı tahminler yapılması tahminlerin daha dirençli ve etkin olmasını sağlar. Bu nedenle aşağıda Arellano yöntemi ile elde edilen tahmin sonuçları Çizelge 11’de gösterilmiştir.

Çizelge 11. Arellano Yöntemi ile Elde Edilen Tahmin Sonuçları

Değişkenler	Tahmin	Standart Hata	T değeri	P değeri
Sabit Terim	1.05326	0.22659	4.64834	0.00001
lRek	1.70573	0.14320	11.91133	0.00000
XlIhr	0.02857	0.04717	0.60563	0.54556
XGsyh	-0.00834	0.00458	-1.82167	0.07025

Panel modelin robust tahmin sonuçlarına göre ise, sabit terim ve Rek değişkenlerinin istatistiksel olarak anlamlı olduğu sonucu bulunmuştur.

Modelin havuzlanmış modelle mi yoksa sabit etkiler modeli ile mi tahmin edilmesi gerektiğini belirlemek için F testi çıktıkları göz önüne alınır. Çizelge 12’ de birim etkiler için F testi sonucu verilmiştir.

Çizelge 12. Birim Etkiler için F testi

F = 69.856	p-value < 2.2e-16
------------	-------------------

Çizelge 12’deki F testi sonuçlarına göre olasılık değeri kritik değerin altında olup sıfır hipotezi reddedilmiş bulunmaktadır. Sonuç olarak modelde sabit etkiler modeli ile tahmin edilmesi daha etkin olacaktır. Sıfır hipotezi ile alternatif hipotez aşağıdaki şekilde kurulmuştur; $H_0: \mu_i = 0$ Birim etki yoktur. Havuzlanmış model uygundur. $H_1: \mu_i \neq 0$ Birim etki vardır. Sabit etkiler modeli kullanılmalıdır. (FE) p olasılık değerine göre H_0 reddedilebilir. Sabit etkili model daha uygun model olacağı sonucuna varılır.

Sabit etkiler modeli tahmincisi ile rassal etkiler modeli tahmincileri arasında seçim yapmak için Hausman testi kullanılmaktadır. H_0 :Rassal Etkiler Modeli Geçerlidir. H_1 :Sabit Etkiler Modeli Geçerlidir. Hausman testini uygulamadan önce ilk olarak sabit ve rassal etkili modellerin tahmin edilmesi gerekmektedir. Çizelge 13’deki Hausman test sonucu ki-kare test ve p değeri istatistiklerini içermektedir. Hausman test istatistiği sonucuna göre ($0.0004529 < 0.05$) sıfır hipotezi reddedilmektedir. Sabit etkiler modelinin varsayımlarının geçerli olduğunu, parametreler arasındaki farkın sistematik olduğu ve açıklayıcı değişkenler ile birim etki arasında korelasyon olduğu sonucuna ulaşılır.

Çizelge 13. Hausman Test

chisq = 17.938	p-value = 0.0004529
----------------	---------------------

Sabit Etkili Model

Sabit etkiler modeli esas alınarak hesaplanmış değişen varyans ve otokorelasyon test istatistikleri Çizelge 14’te ve Çizelge 15’de gösterilmektedir

Çizelge 14. Değişen Varyans Tespiti

Breusch-Pagan test	Test İst. Değeri:32,522	P değeri:0,000
---------------------------	-------------------------	----------------

Breusch-Pagan testi ile modelde değişken varyans varsayımı sınanmış ve yapılan analizlerde değişen varyanslılık sorunu ortaya çıkmıştır.

Çizelge 15. Otokorelasyon Tespiti

Breusch-Godfrey/ Wooldridge test	Test İst. Değeri: 46.525	P değeri: 0,000
---	--------------------------	-----------------

Modelde otokorelasyon varsayımı sınanmış ve elde edilen sonuçlara göre otokorelasyon sorunu olduğu tespit edilmiştir.

Değişen varyans ve otokorelasyon sorunlarından en az bir tanesinin varlığında, tahminler tutarsız fakat etkindir. Bu durumda, parametre tahminlerine dokunulmadan standart hatalar düzeltilmeli (dirençli standart hatalar elde edilmeli) ya da varlıkları halinde uygun yöntemlerle tahmin yapılmalıdır. Modelde değişen varyans ve otokorelasyon sorunlarının varlığından sabit etkili modeller için önerilen Arellano tarafından geliştirilen Arellano robust yöntemi ile panel standart hataların düzeltilmesi yoluyla tahminleme kullanılacaktır.

Sabit etkiler modelinin ekonometrik varsayımlara uyumluluğuna dair otokorelasyon, değişen varyans testleri yapılmıştır. Elde edilen sonuçlara göre sabit etkili panel veri modelinde otokorelasyon ve değişen varyans problemi vardır. Son olarak varsayımlardan sapmaları düzeltmek için sabit etkili model robust standart hatalar ile tahmin edilmiştir ve sonuçlar Çizelge 16’da gösterilmiştir;

Çizelge 16. Robust Standart Hatalar İle Sabit Etkili Model Parametre Tahminleri

Değişkenler	Tahmin	Standart Hata	T değeri	P değeri
lRek	0.75373032	0.15360668	4.9069	2.428e-06
XlIhr	0.04295638	0.03744584	1.1472	0.2532
XGsyh	0.00080366	0.00209173	0.3842	0.7014

Çizelge 16’da robust standart hatalar ile sabit etkili model parametre tahminleri elde edilmiştir. KRE nin KİE üzerinde istatistiksel olarak anlamlı etkisi olduğu sonucuna varılmıştır.

SONUÇ

İnovasyon konusu dünyanın gündemine oturduğundan birçok ülke bu konuda stratejiler geliştirme çabasına girmiştir. Ülkelerin refahını artırmak, ekonomik açıdan büyümek, belirlenen hedeflerin gerçekleştirilmesini sağlamak ve uluslararası rekabete karşı mücadele edip ayakta kalabilmek adına inovasyon kritik bir kavram olarak görülmektedir. Bir ülkede inovasyon yapan kurum ve kuruluşların

sayısının fazla olması, o ülkede yaşayan insanların yaşam şartlarının da o kadar yüksek olduğu anlamına gelmektedir.

Bu çalışmada 25 ülkenin 2011 ile 2017 yılları arasındaki, Küresel İnovasyon Endeksi, Küresel Rekabet Endeksi, GSYH ve GSYH içinde İhracat oranı alınarak elde edilen verileri panel veri analizini R yazılımı ile R Studio programı kullanılarak yapılmıştır. Burada sabit etkili model en uygun model olarak bulunmuştur. Modelde Değişen Varyans ve Otokorelasyon sorunu olduğundan robust tahmin edicilerinden Arellano ile standart hata düzeltilmesi yapılarak tahminler elde edilmiştir. Küresel Rekabet Endeksinin, Küresel İnovasyon Endeksi üzerinde istatistiksel olarak anlamlı etkisi olduğu sonucuna varılmıştır.

Kaynaklar

- Ay Türkmen,M.; Aynaoglu,Y., (2017),“Küresel Rekabet Endeksi Göstergelerinin Küresel İnovasyon Endeksi Üzerindeki Etkisi”,*Business&Management Studies:An International Journal*, 5(4):257-282
- Aynaoglu,Y.,(2018),“Küresel Rekabet Endeksi İle İnovasyon Ve Makroekonomik Göstergeler Arasındaki İlişkinin Analizi”, Pamukkale Üniversitesi Sosyal Bilimler Enstitüsü Yüksek Lisans Tezi İşletme Ana Bilim Dalı Üretim Yönetimi ve Pazarlama Programı, Denizli.
- Çalıpınar,H.,Baç,U.,(2007),“Kobi’lerde İnovasyon Yapmayı Etkileyen Faktörler ve Bir Alan Araştırması”, Ege Akademik Bakış,
- Çetin, A.K.;SÜT.,E.,(2018), “Alternatif İnovasyon Göstergelerinin Büyüme Üzerinde Etkileri: Panel Veri Analizi”, *Social, Mentality and Researcher Thinkers Journal*, Cilt:4 , Sayı: 12
- Hancıoğlu,Y., (2016), “Küresel İnovasyon Endeksini Oluşturan İnovasyon Girdi ve Çıktı Göstergeleri Arasındaki İlişkinin Kanonik Korelasyon Analizi İle İncelenmesi: Oecd Örneği”,*AİBÜ Sosyal Bilimler Enstitüsü Dergisi*, Cilt:16, Yıl:16, Sayı: 4, 16:
- Torun, B.,(2016),“İnovasyon Algısı, İnovasyon Sürecindeki Liderlik Tarzları Ve İşletmenin İnovasyon Performansı Arasındaki İlişkiler:Düzce’deki Kobi’ler Üzerinde Bir Araştırma”, Düzce Üniversitesi, Yüksek Lisans Tezi, Düzce.
- Yılmaz, H., “Stratejik İnovasyon Yönetimi”, Aralık 2015, İstanbul, 1.Baskı.
- Yılmaz,Z.,İncekaş,E.,(2018), “Türkiye’de İnovasyon ve Bölgesel Kalkınma”, *Kırklareli Üniversitesi Sosyal Bilimler Dergisi* , Cilt No2(1)
- Zuhal, M.,(2020), “Gelişmekte Olan Ülkelerde Ulusal İnovasyon Kapasitesinin Belirleyicileri: Panel Veri Analizi Yöntemi”, *Gümüşhane Üniversitesi Sosyal Bilimler Enstitüsü Elektronik Dergisi*, Yıl: 2020 / Cilt: 11 / Sayı: Ek

**COVID-19 Hastalarına İlişkin Eksik Gözlem Durumunda Makine Öğrenimi
Algoritmalarının Kullanımı**

Senem KOC^{1*}, Soner ÇANKAYA², Özgür ENGİNYURT³

¹Nişantaşı Üniversitesi, Tıp Fakültesi, Temel Bilimler Bölümü, Tıbbi İstatistik ABD, 34398, İstanbul, Türkiye

²Ondokuz Mayıs Üniversitesi Yaşar Doğu Spor Bilimleri Fakültesi, Spor Yöneticiliği Bölümü, 55139, Samsun, Türkiye

³Ordu Üniversitesi Tıp Fakültesi Aile Hekimliği ABD, 52200, Ordu, Türkiye

*Sunucu yazar e-posta: senemk55@outlook.com

Özet

Tıp alanında yürütülen çalışmalarda toplanan veriler içerisinde eksik gözlemler ile sıklıkla karşılaşmaktadır. Bu gözlemlerin çalışmadan çıkartılması ise bilgi kaybına sebep olmaktadır. Bu nedenle alanda çalışan istatistikçiler bu tür durumlarla karşılaştıklarında eksik veri probleminin üstesinden gelmek için regresyon ile atama yöntemi, ortalama atama, çoklu atama yöntemleri gibi literatürde çeşitli yöntemler geliştirilmiştir. Makine öğrenme algoritmaları bu tipteki tıbbi verileri analiz etmek için sıkça kullanılmaktadır. Dolayısı ile bu çalışmanın amacı Covid-19 hastalarına (hem yatan hem ayakta tedavi gören) ait eksik veri probleminde naive bayes ve rastgele orman algoritmalarını kullanarak elde edilen sonuçları karşılaştırmaktır. Araştırma sonuçlarına göre, gözlem değerleri içerisinde eksik veri mevcut iken doğruluk %47 olarak tahmin edilmiş iken, ortalama atama yöntemi kullanarak yerine atama ile elde edilen doğruluk naive bayes yöntemi için %60, rastgele orman yöntemi için ise %91 olarak tahmin edilmiştir. Sonuç olarak ortalama atama yöntemi ile eksik veriler tamamlanması durumunda kullanılan algoritmaların performansı arttığı ve en iyi performansı kullanılan yöntemler içerisinde rastgele orman algoritmasının gösterdiği belirlenmiştir.

Anahtar kelimeler: Eksik gözlem, Covid-19, Naive bayes, Rasgele orman, Sınıflandırma

An Empirical Study on High Dimensional, Censored Survival Data Analysis with Supervised Principal Components, Extreme Learning Machine-Based and Penalized Cox Models

Fulden CANTAS^{1*}, İmran KURT ÖMÜRLÜ¹, Mevlüt TÜRE¹

¹Adnan Menderes University, Faculty of Medicine, Division of Biostatistics, Aydın, Turkey

*Presenter author e-mail: fulden35cantas@gmail.com

Abstract

In the analysis of high-dimensional ($p > n$) data, some difficulties arise such as the high correlation of some variables with each other, the high processing load, and the difficulty of interpreting the results. In order to solve this problem, the dimension of the data is reduced by using dimension reduction methods or methods developed for the analysis of these kind of data are used. In this study, it is aimed to compare the performances of supervised principal components analysis (SuperPC), extreme learning machine (ELM)-based (ELMCox, ELMCoxEN, ELMCoxBAR, ELMCoxBoost, ELMmBoost) and penalized Cox models in high-dimensional survival data with varying censoring rates in predicting survival risk. In the simulation study conducted for this purpose, high-dimensional survival data with varying censoring rates are derived in the R software language. In order to analyze generated survival data, SuperPC, ELMCox, ELMCoxEN, ELMCoxBAR, ELMCoxBoost, ELMmBoost and L2-norm regularized Cox regression (Cox-L2) models are used. At the end of the 1000 times repetitive simulation, Harrell's Concordance Index (C-index) and Integrated Brier Score (IBS) are calculated in order to evaluate the performance of the models in estimating the survival risk. The performances of the estimation models are examined by hierarchical clustering analysis. According to the results obtained, ELMCoxBoost is found to be the best performing method while the ELMmBoost is the worst. It has been observed that ELMmBoost and SuperPC methods differ from other methods as the censoring rate in the data set changes. When the performances of the methods are examined according to varying censoring rates; it is observed that the SuperPC method is not affected by the censoring rate, and the performance of the other methods decrease with the increase in the censoring rate. As a result; during the analysis of high-dimensional, censored survival data, the SuperPC method reduces the independent variables to sub-dimensions; ELM-based and penalized Cox models analyze all of the independent variables in the data set simultaneously. It has been determined that ELM-based Cox with likelihood-based boosting algorithm (ELMCoxBoost) has the highest estimation performance among the methods that stand out in this aspect.

Key words: Extreme learning machine, Supervised principal component analysis, Penalized cox regression, Survival, simulation

Kayıp Veri Analiz Yöntemleri Üzerine Bir İnceleme

Abdullah DOĞAN^{1*}, Aydın KARAKOCA¹

¹Necmettin Erbakan University, Faculty of Science, Department of Statistics, 42090, Konya, Turkey

*Sunucu yazar e-posta: dgnapo@hotmail.com

Özet

Kayıp veri sorunu, araştırmacıları yıllardır rahatsız eden bir sorundur. Bir veri setinde eksik değerler varsa bu işleme eksik veri süreci denir. Kayıp veri sürecinde önemli olan kayıp değerlerin varlığı değil, bunların veri setinde nasıl dağıldığı ve bununla nasıl başa çıkıldığıdır. Eksik verilerle başa çıkmanın iki yolu vardır, bunlar silme ve atama yöntemleridir. Silme yöntemleri, liste bazında silme ve çiftler bazında silme yöntemi olarak ikiye ayrılır. Atama yöntemleri, ortalama atama, medyan atama, regresyon atama, stokastik regresyon atama ve beklenti-maksimizasyon atama (EM) yöntemleridir. Bu çalışmada, Matlab istatistiksel analiz programından farklı gözlem sayıları ile çok değişkenli normal dağılımdan 2 değişkenli veri setleri oluşturulmuş ve değişkenler arasındaki korelasyonlar düşük, orta, yüksek olarak üretilmiştir. N = 30, 50, 100, 200 olarak oluşturulan veri setlerinin her bir değişkeni %10, %20, %30 ve %40 eksik oranları içerecek şekilde oluşturulmuştur. Bu veriler MCAR (Missing Completely at Random), MAR (Missing at Random) ve NMAR (Not Missing at Random) eksik veri yapılarında üretilmiştir. Simülasyonda üretilen bu verilerin silme ve atama performansları ortalama hata karesi açısından karşılaştırılmıştır.

Anahtar kelimeler: Ortalama atama, Liste bazında silme, Çiftler bazında silme, Regresyon atama

A Study on Missing Data Analysis Methods

Abstract

The problem of missing data is a problem that has plagued researchers for years. If there are missing values in a data set, this process is called the missing data process. The important thing in the missing data process is not the presence of missing values, but how they are distributed in the data set and dealing with it. There are two ways to deal with missing data, they are deletion and imputation methods. The deletion methods is the listwise and pairwise methods. The imputation methods are mean imputation, median imputation, regression imputation, stochastic regression imputation, and expectation-maximization imputation (EM) methods. In this study, 2-variable data sets were generated from the multivariate normal distribution with different observation numbers from the Matlab statistical analysis program, and the correlations between the variables were produced as low, medium, high. Each variable of the data sets generated as N = 30, 50, 100, 200 was generated to contain 10%, 20%, 30% and 40% missing rates. These data were generated in MCAR (Missing Completely at Random), MAR (Missing at Random), and NMAR (Not Missing at Random) missing data structures. The deletion and imputation performances of these data produced in the simulation were compared in terms of the mean square error.

Key words: MCAR (Missing Completely at Random), MAR (Missing at Random) and NMAR (Not Missing at Random), Mean imputation, Regression imputation

GİRİŞ

Verinin önemi ve büyüklüğünün ön plana çıktığı bir dönemde veri bilimi, makine öğrenmesi, yapay zeka gibi uygulamaları popülerliğini giderek artırırken verinin değeri yadsınamaz. Veriyle bu kadar haşır neşir olunan bir ortamda ise çeşitli sebeplerden dolayı kayıp veriyle karşılaşmak doğal bir sonuçtur. Böyle bir durumda önemli olan kayıp veriyle karşılaşmak değil, kayıp veriyle karşılaşıldığında izlenecek yoldur. Çünkü çoğu istatistiksel yazılım eksik verili durumlarla karşılaştıklarında ya sonuç vermemekte ya da yanlış sonuçlar vermektedir.

Özellikle anket çalışmalarında görülen, sorulara yanıt alamama, uzun anketlerde soruları atlama, sorulara cevap bulamama, cevabının bilinmediği durumlarda soruların boş bırakılması ya da veri girişi sırasında dikkatsizlikten ya da hatalı kodlamadan dolayı eksik veri problemiyle karşılaşmaktadır. Bazen bu süreç ile karşılaşılacağı öngörülebilir ve önlem alınabilir. Ancak yine de eksik veriyle karşılaşmak mümkündür. Bu problemle başa çıkmak araştırmanın genellenebilirliği ve güvenilirliği açısından hayati bir öneme sahiptir.

Kayıp veriler çalışmanın sonuçlarını çeşitli şekilde etkilemektedir. Eğer kayıp veri miktarı fazlaysa elde edilen sonuçların güvenilirliği, genellenebilirliği ve yapılacak istatistiksel çıkarımlar bundan önemli derecede etkilenecektir. Bunun sonucunda da yapılan istatistiksel çıkarımlar yanıltıcı olacaktır. Ayrıca kayıp veriler çalışmanın geçerliliğini de olumsuz yönde etkileyecektir (McKnight ,2007).

Eksik veriyle karşılaşıldığında izlenecek yollar ana hatlarıyla belirlidir. İlk olarak veri setindeki eksikliğin tamamen rastsal olarak kayıp (MCAR), rastsal olarak kayıp (MAR) veya rastsal olarak kayıp olmayan (NMAR) veri mekanizmasına uyup uymadığına bakılır. İkinci olarak da eksik verinin ne sıklıkta olduğuna bakılır. Kullanıcılar veri setindeki eksik gözlemleri göz ardı ederek silme yöntemlerini veya eksik verileri uygun gözlemlerle tamamlayarak atama yöntemlerini kullanırlar. Her bir yöntemin kendi içlerinde artıları ve eksileri mevcuttur. Eksik veriler ortadan kaldırıldığında başta yeterli olarak görülen veri seti eksik veriler silindikten sonra yeterli olmayabilir bu da araştırmanın sonucunda tahmin edilenden çok farklı bir sonuca sebep olabilir. Diğer bir durum ise çeşitli yöntemler kullanılarak eksik gözlemlerin değerlerini tahmin etmektir. Her iki durumda da yanlış ve istenmeyen sonuçlar görülebilir.

Kayıp veri çalışmaları 1970’lerde ad-hoc prosedürlerine kadar dayanmaktadır. Bugünlerde de sıkça kullanılan kayıp veri probleminin teorik temelleri 1976 yılında Rubin tarafından atılmıştır. Little tarafından ortaya atılan MCAR testi de verideki rastgeleliği araştırmada kullanılan ve günümüzde hala geçerliliğini koruyan yöntemlerden biridir. Özellikle 1980’lerde benzerliğe dayalı yöntemler daha da popülerleşmiş ve yoğunluk EM (Expectation Maximization) algoritmasında toplanmıştır. 1990’larda ise çoklu atama (Multiple Imputation) kullanılmıştır.

Bu çalışmada literatürde yer alan silme ve tekli atama yöntemlerinin farklı kayıp veri oranlarına, gözlem sayılarına, değişkenler arası ilişki düzeylerine ve kayıp veri yapılarına olan duyarlılıkları detaylı bir Monte Carlo simülasyon çalışmasıyla incelenmiştir.

Kayıp Veri Yapıları

$X=(y_{ij})$, $n \times p$ boyutlu veri seti ve $E=(m_{ij})$ gözlem gösterge matrisi olmak üzere, eksik veri mekanizmaları X verilmişken E ’nin koşullu dağılımıdır ve $f(X/E, \emptyset)$ şeklinde gösterilir. Burada $\emptyset$ bilinmeyen parametreleri göstermektedir. Eğer eksiklik E ’nin değerlerine bağlı değil ise,

$$f(X/E, \emptyset) = f(X/E)$$

şeklinde ifade edilen durum herhangi bir değişkende kayıp veri olmasının diğer bir değişkene bağlı olmadığı MCAR kayıp veri yapısını temsil eder.

$E_{göz}$, E veri setinde gözlenen kısmı, E_{eks} , E veri setinde eksik olan kısmı göstermek üzere,

$f(X/E, \emptyset) = f(X/E_{göz}, \emptyset)$ tüm $B_{eks}, \emptyset$ 'ler için şeklinde tanımlanan eşitlik herhangi bir değişkendeki kayıp veri olması olasılığının diğer bir değişkenden kaynaklandığı MAR kayıp veri yapısını ifade eder.

$f(X/E, \emptyset) = f(X/E_{göz}, E_{eks}, \emptyset)$ tüm $B_{göz}, B_{eks}, \emptyset$ 'ler için şeklinde tanımlanan herhangi bir değişkendeki kayıp veri olması olasılığının diğer değişkenlere değil de yalnız kayıp veri bulunan değişkenin kendi değerinden etkilenmesi sonucu da NMAR kayıp veri yapısını ifade eder.

Kayıp veriyle başa çıkmada kullanılan yöntemler iki ana başlıkta incelenir. Bunlardan ilk silme yöntemleri diğeri ise atama yöntemleridir.

Kayıp veri silme yöntemleri

Listwise deletion (Liste bazında silme): Veri setinde bir ya da daha fazla eksik veri içeren gözlem birimlerinin veri setinden çıkartılması durumudur.

Pairwise deletion (Çiftler bazında silme): Uygulanacak analize göre mümkün olan tüm kayıp veri içermeyen gözlemleri kullanmaya dayalı silme yöntemidir. Herhangi bir değişkenin ortalamasını hesaplamak için yalnız o değişkendeki tam veriler kullanılırken, iki değişken arasındaki korelasyonu hesaplamak için yalnız ilgili iki değişkenin tam olan gözlemlerinin kullanılması bu yöntemle örnek verilebilir.

Kayıp veri atama yöntemleri

Literatürde sıklıkla kullanılan tekli atama yöntemlerinden adhoc yöntemleri bu çalışmada kullanılmıştır.

Ortalama değer atama yöntemi: Bu yöntemde veri setinin ilgili değişkenine ait ortalamayı kullanarak eksik gözlemlerinin tamamlanmasına dayanan bir yöntemdir.

Medyan atama yöntemi: İlgili değişkenin medyan değerinin bulunup kayıp gözlemlerin yerine atanması durumudur.

Regresyon atama yöntemi: Herhangi bir kayıp gözlem içeren değişkenin bağımlı değişken, diğer değişkenlerinde bağımsız değişken alınarak oluşturulan regresyon modeli yardımıyla eksik gözlemlerin tamamlanmasına dayanan bir yöntemdir.

Stokastik regresyon atama yöntemi: Regresyon atama modeli kullanılarak atanan verilere sıfır ortalamalı, sabit σ^2 varyanslı hata dağılımından elde edilen değerlerin eklenmesiyle elde edilen atama yöntemidir.

Beklenen değer atama yöntemi (EM): En çok olabilirlik yöntemi temel alınarak geliştirilen bu yöntem, iki aşamalı bir yöntemdir. E adımında tam gözlemleri kullanarak eksik gözlemleri tahmin etmek için oluşturulan ortalama vektörü ve kovaryans matrisi ile stokastik regresyon modeli oluşturulur ve model yardımı ile eksik gözlemler tahmin edilir. Algoritmanın E adımının sonunda, eksik gözlemlerin tahmini ile veri seti tamamlanmış veri durumuna getirilir. Algoritmanın M adımında, tamamlanmış veri durumundaki prosedürler uygulanarak ortalama vektörü ve kovaryans matrisi tahminleri yenilenir. Bu süreç tahmin değerlerindeki ya da logaritmik olabilirlik değerlerinin toplamı arasındaki değişimin önemsenmeyecek derecede azalmasına kadar devam eder (Sarı, 2012).

SİMÜLASYON ÇALIŞMASI

Kayıp veri analizinde farklı kayıp veri yapılarının ve kullanılan atama yöntemlerinin performanslarını detaylı bir şekilde inceleyebilmek için Monte Carlo simülasyon çalışması aşağıdaki şekilde planlanmıştır.

Araştırma örneklemini, Matlab istatistiksel analiz programı kullanılarak oluşturulan standart çok değişkenli normal dağılımdan elde edilen veri setleri kullanılarak gerçekleştirilmiştir. Oluşturulan veri

setleri atama yöntemlerinin performanslarını detaylı bir şekilde karşılaştırmak açısından farklı parametre seçimlerine göre incelenmiştir. Gözlem sayısının düşük, orta ve yüksek değerleri için $N = 30, 50, 100$ ve 200 değerleri incelenmiştir. Veri setlerindeki değişkenler arasındaki ilişkinin düşük, orta ve yüksek değerlerini temsil edecek şekilde herhangi iki değişken arasındaki korelasyon katsayısının sırasıyla $-0.30 < \rho_{xy} < 0.30$, $\pm 0.30 < \rho_{xy} < \pm 0.70$ ve $\rho_{xy} > 0.95$ olduğu durumlar ele alınmıştır. Her bir değişkendeki kayıp veri sayısı %10, %20, %30 ve %40 oranında eksik veri barındıracak şekilde incelenmiştir. Ayrıca herbir veri seti MCAR (Missing Completely at Random), MAR (Missing at Random) ve NMAR (Not Missing at Random) kayıp veri yapılarında üretilerek yöntemlerin performansları değerlendirilmiştir.

Simülasyon çalışması 1000 tekrarlar yürütülmüştür. Her bir veri seti yukarıda belirtilen parametre değerleri için her bir silme ve atama yöntemiyle elde edilen ortalama değerleri kullanılarak eksik veriler tamamlanmış ve elde edilen ortalamaların gerçek veri setinin bilinen parametre değerleriyle uyumu hata kareler ortalamasıyla (HKO) karşılaştırılarak değerlendirilmiştir. Herhangi bir değişken için tam verinin ortalaması μ , θ yöntemiyle yapılan atama sonunca elde edilen ortalama değeri θ_μ olmak üzere θ yöntemine ilişkin hata kareler ortalaması,

$$HKO_\theta = \left(\frac{\sum_{i=1}^{1000} (\theta_\mu - \mu)^2}{1000} \right)$$

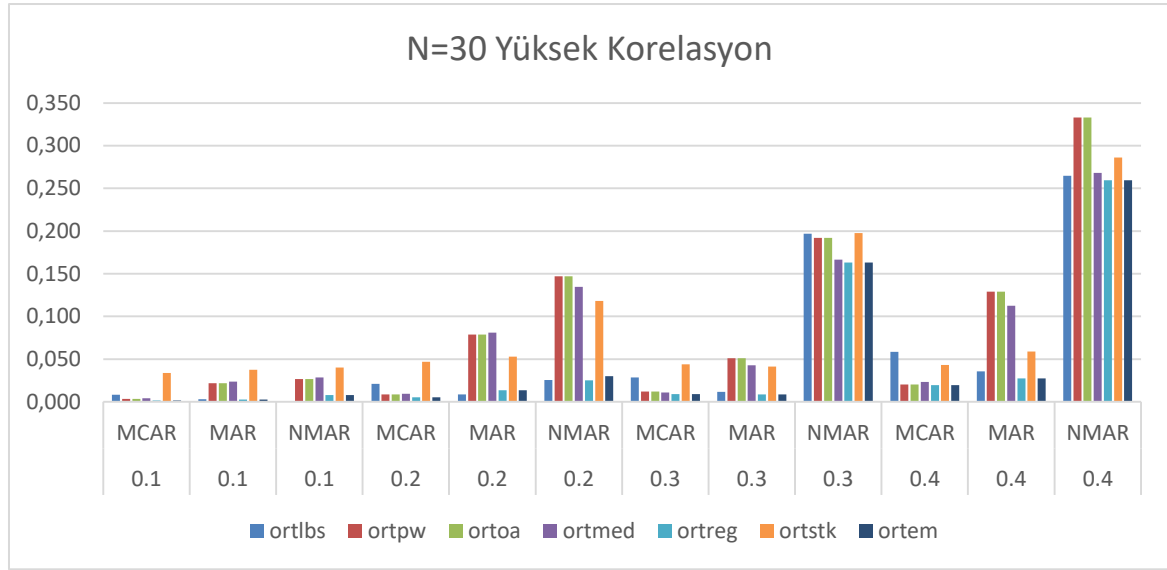
şeklinde hesaplanmıştır.

BULGULAR

Simülasyon çalışmasıyla her bir yönteme ilişkin HKO değerleri Çizelge 1-12’de verilmiştir. Çizelge 1’de 30 birimli, yüksek korelasyonlu veri seti için değerler verilmiştir.

Çizelge 1. N=30 ve yüksek korelasyon durumunda yöntemlerin HKO performansları

	0.1	0.1	0.1	0.2	0.2	0.2	0.3	0.3	0.3	0.4	0.4	0.4
	MCAR	MAR	NMAR	MCAR	MAR	NMAR	MCAR	MAR	NMAR	MCAR	MAR	NMAR
ortlbs	0,008	0,003	0,000	0,021	0,008	0,026	0,028	0,012	0,197	0,059	0,036	0,265
ortpw	0,004	0,022	0,027	0,009	0,079	0,147	0,012	0,051	0,192	0,020	0,129	0,333
ortoa	0,004	0,022	0,027	0,009	0,079	0,147	0,012	0,051	0,192	0,020	0,129	0,333
ortmed	0,004	0,024	0,028	0,009	0,081	0,135	0,011	0,043	0,166	0,023	0,113	0,268
ortreg	0,002	0,003	0,008	0,005	0,014	0,025	0,009	0,009	0,163	0,020	0,027	0,259
ortstk	0,034	0,038	0,040	0,047	0,053	0,118	0,044	0,041	0,198	0,043	0,059	0,286
ortem	0,002	0,003	0,008	0,005	0,014	0,030	0,009	0,009	0,163	0,020	0,027	0,259



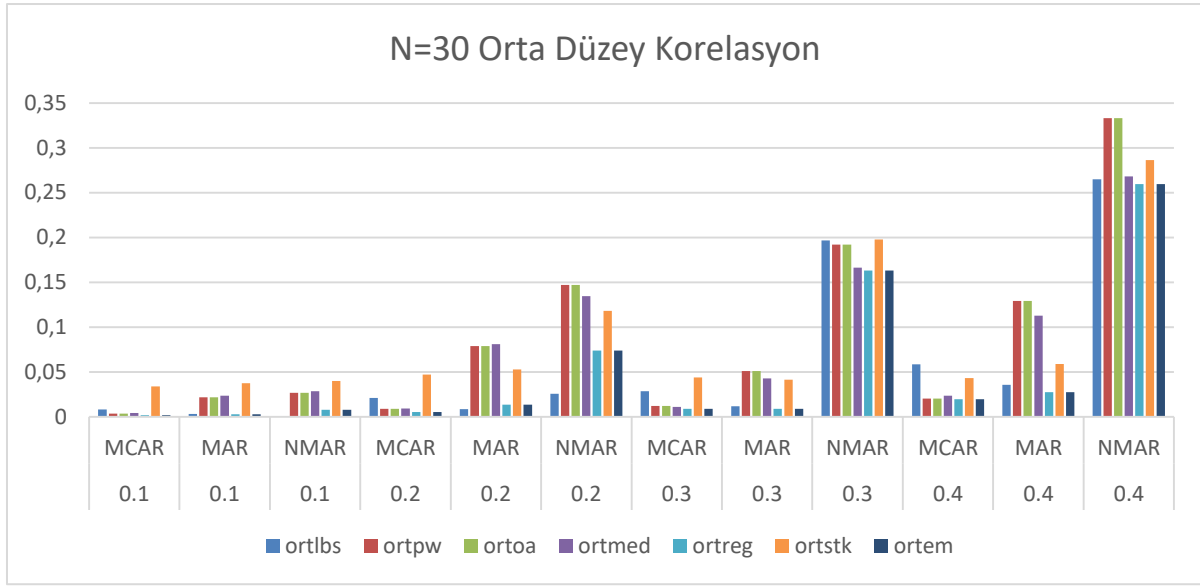
Şekil 1. 30 birimlik yüksek korelasyona sahip veri seti için yöntemlerin performansları

30 birimlik ve yüksek korelasyon durumunda veri yapısı incelendiğinde kayıp oranı arttıkça HKO artmaktadır. Tüm kayıp oranları incelendiğinde en düşük hata kareler ortalaması MCAR durumunda olduğu söylenebilir. MCAR ve MAR yapısında regresyon ve EM atama, NMAR liste bazında silme yöntemi en düşük HKO'ya sahiptir.

Çizelge 2'de N=30 orta düzey korelasyona sahip veri seti incelenmeye çalışılmıştır. MCAR değerlerinin her yöntemde daha düşük değerlere sahip olduğu göze çarpmaktadır.

Çizelge 2. N=30 ve orta korelasyon durumunda yöntemlerin HKO performansları

	0.1	0.1	0.1	0.2	0.2	0.2	0.3	0.3	0.3	0.4	0.4	0.4
	MCAR	MAR	NMAR	MCAR	MAR	NMAR	MCAR	MAR	NMAR	MCAR	MAR	NMAR
ortlbs	0,008	0,003	0,000	0,021	0,008	0,026	0,028	0,012	0,197	0,059	0,036	0,265
ortpw	0,004	0,022	0,027	0,009	0,079	0,147	0,012	0,051	0,192	0,020	0,129	0,333
ortoa	0,004	0,022	0,027	0,009	0,079	0,147	0,012	0,051	0,192	0,020	0,129	0,333
ortmed	0,004	0,024	0,028	0,009	0,081	0,135	0,011	0,043	0,166	0,023	0,113	0,268
ortreg	0,002	0,003	0,008	0,005	0,014	0,074	0,009	0,009	0,163	0,020	0,027	0,259
ortstk	0,034	0,038	0,040	0,047	0,053	0,118	0,044	0,041	0,198	0,043	0,059	0,286
ortem	0,002	0,003	0,008	0,005	0,014	0,074	0,009	0,009	0,163	0,020	0,027	0,259



Şekil 2. 30 birimlik orta düzey korelasyona sahip veri seti için yöntemlerin performansları

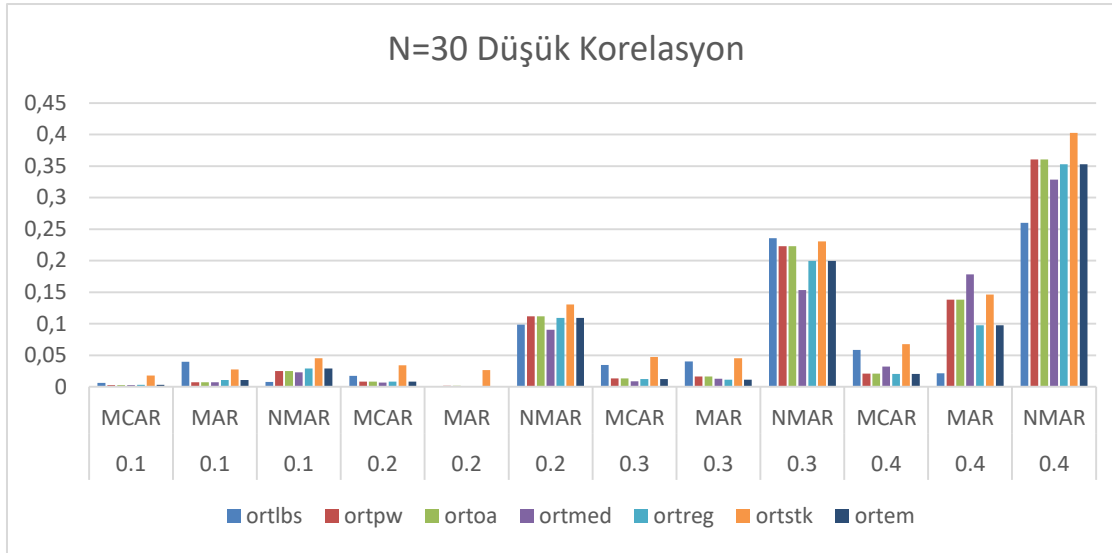
30 gözlemlilik orta düzey korelasyona sahip veri yapısında ise kayıp oranı arttıkça hata kareler ortalaması artmakta ve MCAR veri yapısına sahip değerler en düşük hata kareler ortalamasına sahip değerlerdir. NMAR durumundaki veri seti ise kayıp veri oranından en çok etkilenen veri seti olarak göze çarpmaktadır. MCAR ve MAR durumunda EM atama daha düşük HKO'na sahipken, NMAR'da liste bazında silme değerlerinin HKO değeri daha düşüktür.

Çizelge 3'te veri seti arasındaki ilişki düşük düzeyde korelasyon yapısına sahip 30 gözlemlilik bir veri seti incelenmiştir.

Çizelge 3. N=30 ve düşük korelasyon durumunda yöntemlerin HKO performansları

	0.1	0.1	0.1	0.2	0.2	0.2	0.3	0.3	0.3	0.4	0.4	0.4
	MCAR	MAR	NMAR	MCAR	MAR	NMAR	MCAR	MAR	NMAR	MCAR	MAR	NMAR
ortlbs	0,006	0,040	0,008	0,017	0,000	0,098	0,035	0,040	0,236	0,059	0,022	0,260
ortpw	0,003	0,007	0,025	0,008	0,002	0,112	0,013	0,016	0,223	0,021	0,138	0,361
ortoa	0,003	0,007	0,025	0,008	0,002	0,112	0,013	0,016	0,223	0,021	0,138	0,361
ortmed	0,003	0,007	0,023	0,007	0,001	0,090	0,009	0,013	0,154	0,032	0,178	0,328
ortreg	0,003	0,011	0,029	0,008	0,001	0,109	0,012	0,011	0,199	0,020	0,098	0,353
ortstk	0,018	0,027	0,045	0,034	0,026	0,130	0,047	0,045	0,230	0,068	0,146	0,402
ortem	0,003	0,011	0,029	0,008	0,001	0,109	0,012	0,011	0,199	0,020	0,098	0,353

Çizelge 3'te incelenen veri seti setinde diğer veri setlerinde farklı olarak MAR değerlerinin belirli bir orana kadar hata kareler ortalamasının düşük olması göze çarpmıştır.



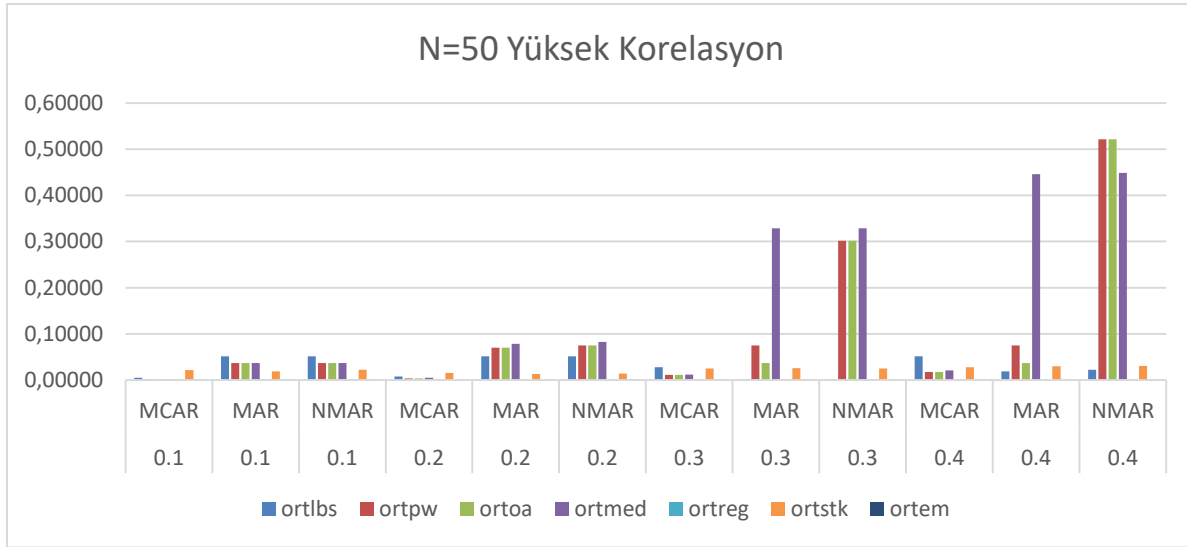
Şekil 3. 30 birimlik düşük düzey korelasyona sahip veri seti için yöntemlerin performansları

30 gözlemlili düşük korelasyon yapısına sahip veri setinde kayıp veri oranına bağlı olarak hata kareler ortalama değişmektedir. MCAR veri yapısına sahip değerler daha küçük hata kareler ortalama değerine sahip iken, NMAR veri yapısındaki değerler en yüksek hata kareler ortalamasına sahip değerlerdir. MCAR medyan ortalama MAR ve NMAR durumunda liste bazında silme yönteminin HKO'ları daha düşüktür.

Veri setinde gözlem sayısı arttığında yöntemler arası performans değerlerindeki değişiklikleri gözlemlemek adına Çizelge 4'te 50 birimlik bir veri seti incelenmiştir.

Çizelge 4. N=50 ve yüksek korelasyon durumunda yöntemlerin HKO performansları

	0.1	0.1	0.1	0.2	0.2	0.2	0.3	0.3	0.3	0.4	0.4	0.4
	MCAR	MAR	NMAR	MCAR	MAR	NMAR	MCAR	MAR	NMAR	MCAR	MAR	NMAR
ortlbs	0,005	0,051	0,051	0,008	0,051	0,051	0,028	0,000	0,000	0,051	0,019	0,022
ortpw	0,002	0,037	0,037	0,004	0,070	0,075	0,011	0,075	0,301	0,018	0,075	0,521
ortoa	0,002	0,037	0,037	0,004	0,070	0,075	0,011	0,037	0,301	0,018	0,037	0,521
ortmed	0,002	0,037	0,037	0,005	0,079	0,082	0,012	0,329	0,329	0,021	0,446	0,449
ortreg	0,000	0,000	0,000	0,000	0,000	0,002	0,000	0,000	0,000	0,000	0,000	0,000
ortstk	0,022	0,019	0,022	0,015	0,013	0,014	0,025	0,026	0,025	0,028	0,030	0,030
ortem	0,000	0,000	0,000	0,000	0,000	0,002	0,000	0,000	0,000	0,000	0,000	0,000



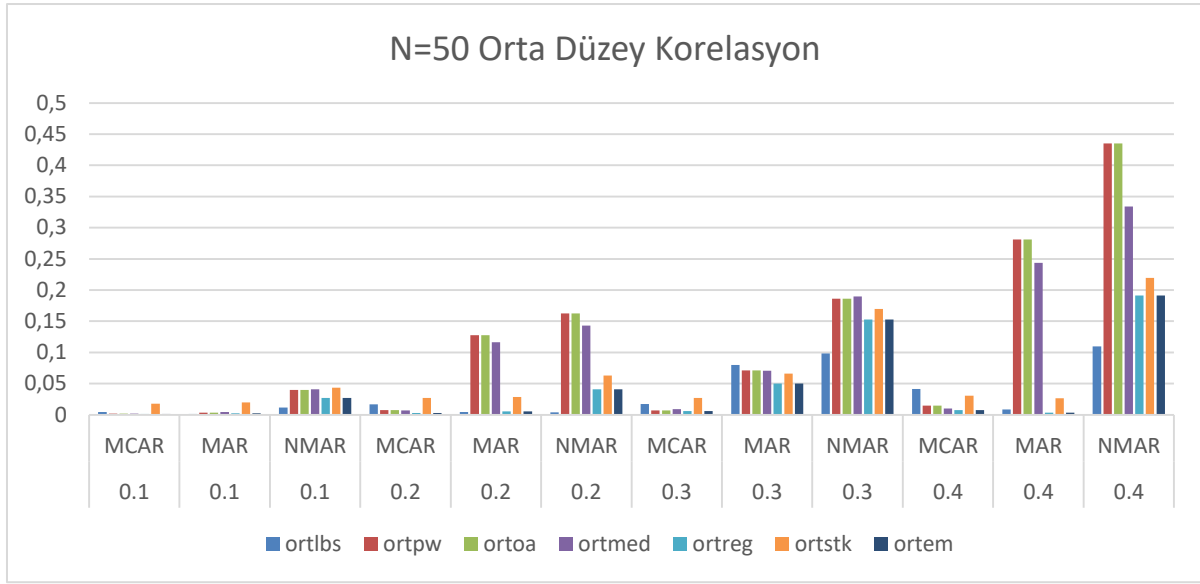
Şekil 4. 50 birimlik yüksek korelasyona sahip veri seti için yöntemlerin performansları

50 birimlik yüksek korelasyona sahip veri seti incelendiğinde 0.1 ve 0.2 kayıp değer oranına sahip değerler neredeyse aynı hata kareler ortalamasına sahip iken 0.3’den sonra özellikle MAR veri yapısında medyan değeriyle ortalama atama yapılan değerlerin hata kareler ortalamasındaki artışlar göze çarpmaktadır. MCAR veri türü en düşük hata kareler ortalamasına sahip veriler üretse de 0.3 ve 0.4 kayıp yüzdeli veri setlerinde NMAR durumundaki verilerin hata kareler ortalaması çok artmıştır. MCAR regresyon atama, MAR durumunda EM, NMAR durumunda ise liste bazında silme yönteminin HKO’ları daha düşüktür.

Orta düzey korelasyon seviyesine sahip 50 birim içeren tablo ya kayıp veri yöntemlerinin performanslarını incelemek amacıyla çeşitli yöntemler uygulanmış ve bu yöntemler Çizelge 5’te gösterilmiştir.

Çizelge 5. N=50 ve orta düzey korelasyon durumunda yöntemlerin HKO performansları

	0.1	0.1	0.1	0.2	0.2	0.2	0.3	0.3	0.3	0.4	0.4	0.4
	MCAR	MAR	NMAR	MCAR	MAR	NMAR	MCAR	MAR	NMAR	MCAR	MAR	NMAR
ortlbs	0,004	0,001	0,012	0,016	0,004	0,004	0,017	0,080	0,098	0,041	0,008	0,110
ortpw	0,002	0,003	0,040	0,007	0,127	0,163	0,007	0,071	0,186	0,015	0,281	0,435
ortoa	0,002	0,003	0,040	0,007	0,127	0,163	0,007	0,071	0,186	0,015	0,281	0,435
ortmed	0,002	0,005	0,041	0,007	0,116	0,143	0,009	0,071	0,190	0,010	0,244	0,334
ortreg	0,001	0,002	0,027	0,003	0,005	0,041	0,006	0,050	0,153	0,007	0,003	0,191
ortstk	0,018	0,020	0,043	0,027	0,029	0,063	0,027	0,066	0,170	0,031	0,027	0,219
ortem	0,001	0,002	0,027	0,003	0,005	0,041	0,006	0,050	0,153	0,007	0,003	0,191



Şekil 5. 50 birimlik orta düzey korelasyona sahip veri seti için yöntemlerin performansları

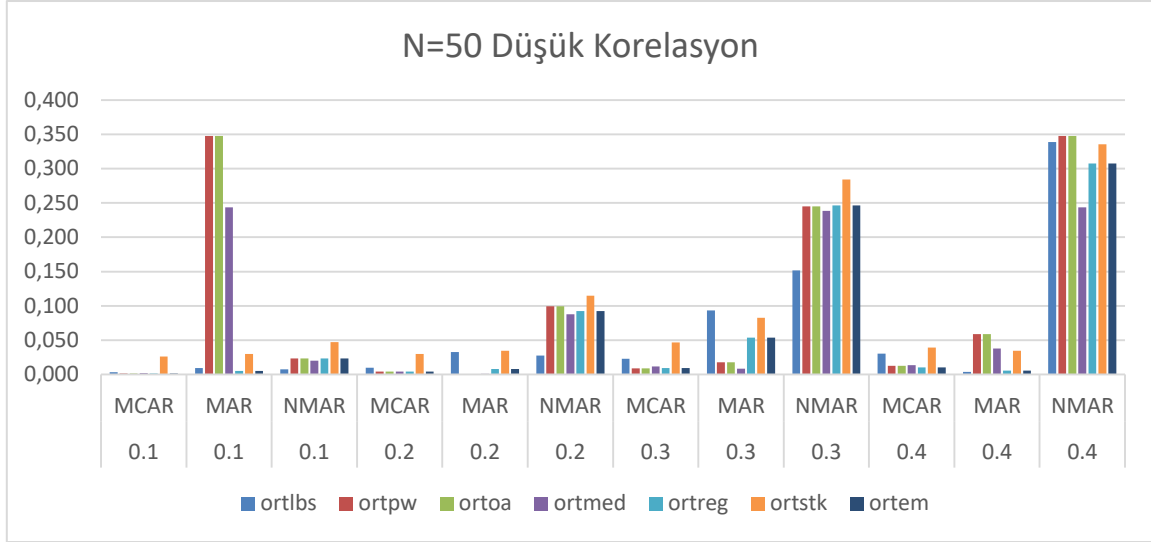
50 birimlik orta düzey korelasyon ilişkisine sahip veri setinin grafiği incelendiğinde MCAR veri yapısına sahip birimlerin hata kareler ortalaması düşük iken, MAR ve NMAR veri yapısında yüzde yirmilik kayıp oranda dahil olmak bu orandan itibaren artan bir yapıya sahiptir. MCAR regresyon atama ve EM atama, MAR durumunda liste bazında silme, NMAR durumunda ise çiftleri bazında silme yönteminin HKO'ları daha düşüktür.

50 birim içeren düşük düzeyde korelasyona sahip veri setine silme ve atma yöntemleri uygulanmış ve bu uygulamanın sonuçları Çizelge 6'da gösterilmiştir.

Çizelge 6. N=50 ve düşük korelasyon durumunda yöntemlerin HKO performansları

	0.1	0.1	0.1	0.2	0.2	0.2	0.3	0.3	0.3	0.4	0.4	0.4
	MCAR	MAR	NMAR	MCAR	MAR	NMAR	MCAR	MAR	NMAR	MCAR	MAR	NMAR
ortlbs	0,004	0,010	0,008	0,010	0,033	0,028	0,023	0,093	0,152	0,031	0,004	0,339
ortpw	0,002	0,347	0,024	0,005	0,000	0,100	0,009	0,018	0,245	0,013	0,059	0,347
ortoa	0,002	0,347	0,024	0,005	0,000	0,100	0,009	0,018	0,245	0,013	0,059	0,347
ortmed	0,002	0,244	0,020	0,005	0,001	0,088	0,012	0,008	0,238	0,013	0,038	0,244
ortreg	0,002	0,005	0,023	0,004	0,008	0,093	0,009	0,054	0,246	0,011	0,006	0,307
ortstk	0,026	0,030	0,047	0,030	0,035	0,115	0,047	0,083	0,284	0,039	0,034	0,336
ortem	0,002	0,005	0,023	0,004	0,008	0,093	0,009	0,054	0,246	0,011	0,006	0,307

50 birimlik düşük düzey korelasyonlu veri setine atama ve silme yöntemleri uygulandıktan sonra hata kareler ortalama sonuçları Şekil 6'te gösterilmiştir.



Şekil 6. 50 birimlik düşük düzey korelasyona sahip veri seti için yöntemlerin performansları

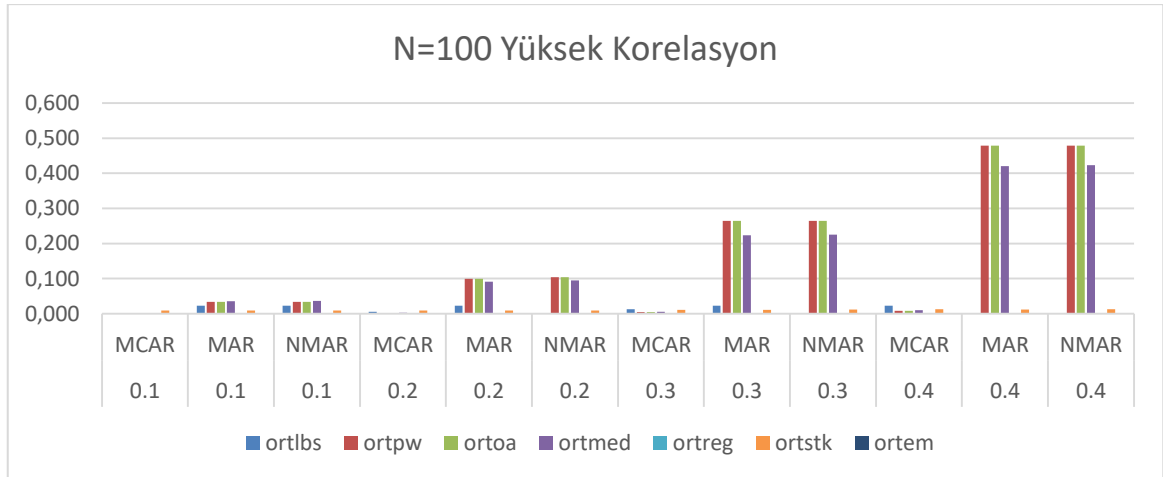
50 birimli düşük korelasyonlu veri yapısına bakıldığında MCAR yapısında da özellikle 0.3'dan itibaren artışlar gözlenmektedir. Hata kareler ortalaması NMAR'da 0.3 ve 0.4 kayıp veri oranında en yüksek değerlere sahip olduğu görülmektedir. MCAR EM atama, MAR durumunda çiftler bazında silme ve ortalama atama, NMAR durumunda ise çiftleri bazında silme yönteminin HKO'ları daha düşüktür.

100 birimden oluşan yüksek korelasyona sahip veri setinde sırasıyla %10, %20, %30 ve %40 değerler eksiltilerek elde edilen veri setinin hata kareler ortalaması Çizelge 7'de verilmiştir.

Çizelge 7. N=100 ve yüksek korelasyon durumunda yöntemlerin HKO performansları

	0.1	0.1	0.1	0.2	0.2	0.2	0.3	0.3	0.3	0.4	0.4	0.4
	MCAR	MAR	NMAR	MCAR	MAR	NMAR	MCAR	MAR	NMAR	MCAR	MAR	NMAR
ortlbs	0,002	0,023	0,023	0,005	0,023	0,000	0,013	0,023	0,000	0,023	0,000	0,001
ortpw	0,001	0,033	0,034	0,002	0,099	0,104	0,005	0,265	0,265	0,008	0,479	0,479
ortoa	0,001	0,033	0,034	0,002	0,099	0,104	0,005	0,265	0,265	0,008	0,479	0,479
ortmed	0,001	0,036	0,036	0,002	0,092	0,095	0,006	0,223	0,226	0,010	0,420	0,423
ortreg	0,001	0,001	0,001	0,001	0,001	0,000	0,001	0,000	0,001	0,000	0,000	0,001
ortstk	0,009	0,009	0,009	0,009	0,009	0,009	0,011	0,011	0,012	0,013	0,011	0,012
ortem	0,000	0,000	0,000	0,000	0,000	0,000	0,000	0,000	0,000	0,000	0,000	0,000

Çizelge 7 incelendiğinde MCAR değişkenlerin en düşük hata kareler ortalamalarına sahip olduğu görülmektedir. Çizelgedeki yapıyı daha iyi anlamak adına Şekil 7'de histogram grafiği verilmiştir.



Şekil 7. 100 birimlik yüksek düzey korelasyona sahip veri seti için yöntemlerin performansları

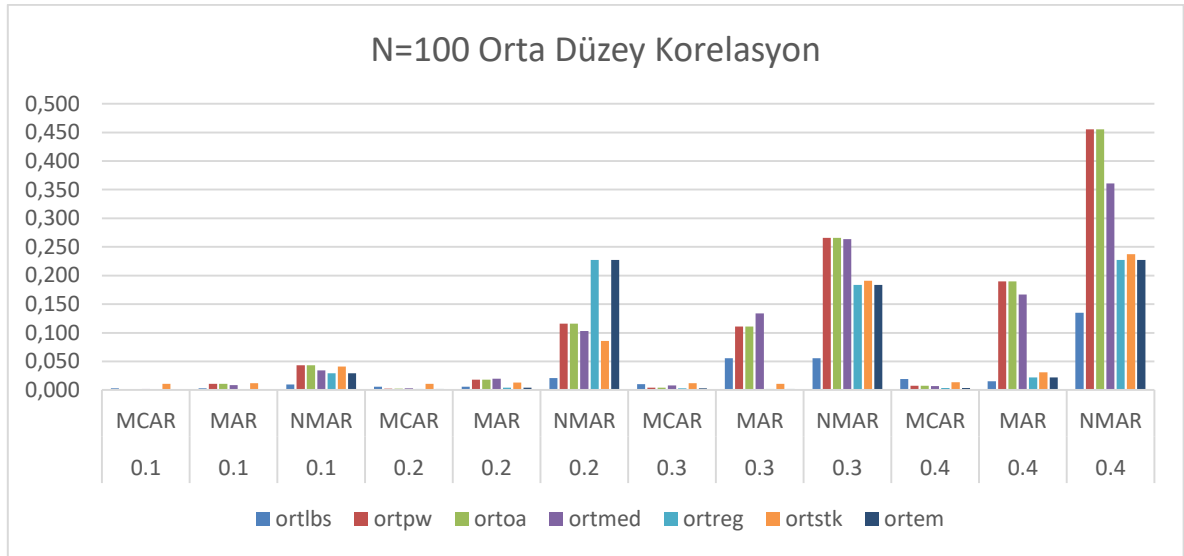
100 birimlik yüksek korelasyon ilişkisine sahip bu veri setini incelediğimizde ise MCAR yapıda hata kareler ortalaması neredeyse tüm eksik veri oranlarında en düşük seviyelerde seyretmiştir. MAR ve NMAR veri yapısında ise ortalama atama, medyan atama ve pairwise silme yöntemlerinde yüksek değerlere sahip olduğu görülmektedir.

100 birimlik ortalama korelasyon yapısındaki eksik veri atama ve silme yöntemlerinin performansları Çizelge 8’de karşılaştırılmıştır.

Çizelge 8. N=100 ve orta düzey korelasyon durumunda yöntemlerin HKO performansları

	0.1	0.1	0.1	0.2	0.2	0.2	0.3	0.3	0.3	0.4	0.4	0.4
	MCAR	MAR	NMAR	MCAR	MAR	NMAR	MCAR	MAR	NMAR	MCAR	MAR	NMAR
ortlbs	0,003	0,003	0,009	0,006	0,006	0,021	0,010	0,055	0,055	0,019	0,015	0,135
ortpw	0,001	0,011	0,043	0,002	0,018	0,116	0,004	0,111	0,266	0,007	0,190	0,455
ortoa	0,001	0,011	0,043	0,002	0,018	0,116	0,004	0,111	0,266	0,007	0,190	0,455
ortmed	0,002	0,008	0,034	0,003	0,020	0,103	0,008	0,134	0,264	0,007	0,167	0,361
ortreg	0,001	0,001	0,029	0,002	0,004	0,227	0,003	0,002	0,184	0,004	0,022	0,227
ortstk	0,011	0,012	0,041	0,011	0,013	0,086	0,012	0,011	0,191	0,014	0,031	0,237
ortem	0,001	0,001	0,029	0,002	0,004	0,227	0,003	0,002	0,184	0,004	0,022	0,227

Orta düzeyde korelasyon yapısına sahip tablo incelendiğinde MCAR değerlerin en küçük hata kareler ortalamasına sahip olduğu görülmektedir. Bu tablodaki yapıyı incelemek amacıyla Şekil 8’de histogram grafiği oluşturulmuştur.



Şekil 8. 100 birimlik orta düzey korelasyona sahip veri seti için yöntemlerin performansları

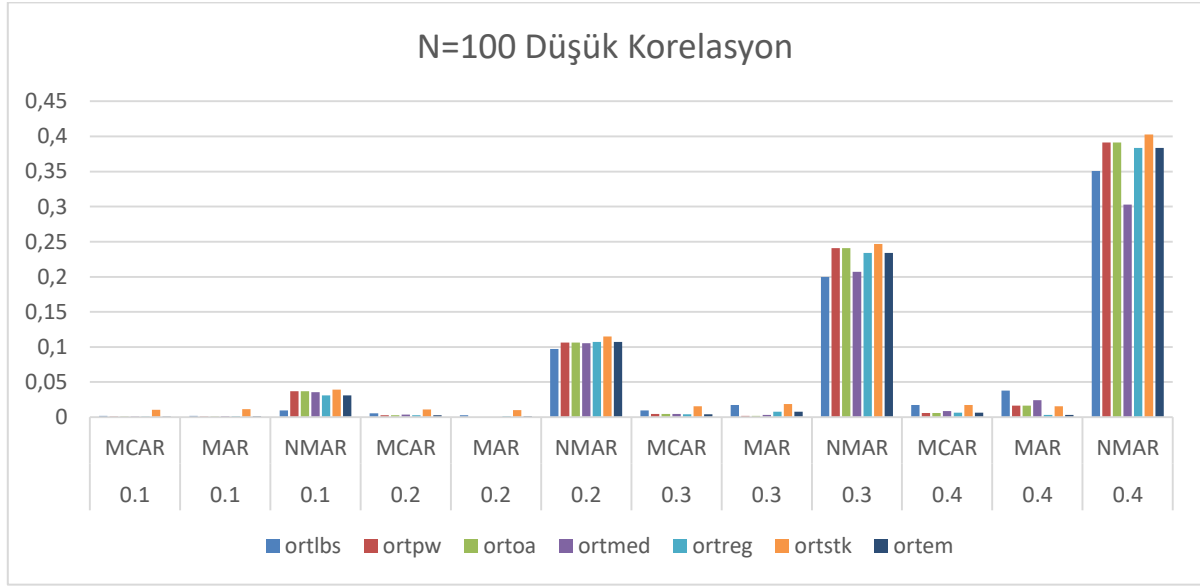
Orta düzeyde korelasyona sahip 100 birimlik veri setine sahip grafiği incelediğimizde MCAR veri yapısına sahip birimler en düşük hata kareler ortalama değerine sahip iken MAR veri yapısında 0.3 oranından sonra artışlar gözlenmektedir. MAR’da %40 eksik veri oranında liste bazında silme, regresyon atama, stokastik regresyon atama ve em atama yöntemlerinin hata kareler ortalamasının düşük değere sahip oldukları görülmektedir. NMAR veri yapısına sahip veriler ise %10’dan itibaren en yüksek hata kareler ortalamasına sahip değerler olarak gözlemlenmişlerdir.

100 birim içeren tam veri setinde çeşitli oranlarda veri silme işlemleri uygulanmıştır. Silme işleminden sonra silme yöntemlerinin performansları karşılaştırılmıştır. Bu karşılaştırma Çizelge 9’da gösterilmiştir.

Çizelge 9. N=100 ve düşük düzey korelasyon durumunda yöntemlerin HKO performansları

	0.1	0.1	0.1	0.2	0.2	0.2	0.3	0.3	0.3	0.4	0.4	0.4
	MCAR	MAR	NMAR	MCAR	MAR	NMAR	MCAR	MAR	NMAR	MCAR	MAR	NMAR
ortlbs	0,002	0,002	0,010	0,006	0,003	0,097	0,010	0,018	0,200	0,018	0,038	0,351
ortpw	0,001	0,001	0,037	0,003	0,391	0,106	0,004	0,002	0,241	0,006	0,016	0,391
ortoa	0,001	0,001	0,037	0,003	0,391	0,106	0,004	0,002	0,241	0,006	0,016	0,391
ortmed	0,001	0,001	0,036	0,004	0,391	0,105	0,005	0,003	0,207	0,009	0,024	0,303
ortreg	0,001	0,001	0,031	0,003	0,001	0,107	0,004	0,008	0,234	0,006	0,003	0,384
ortstk	0,011	0,011	0,039	0,011	0,010	0,115	0,015	0,019	0,247	0,017	0,015	0,403
ortem	0,001	0,001	0,031	0,003	0,001	0,107	0,004	0,008	0,234	0,006	0,003	0,384

Çizelge 9’da verilen tabloyu daha iyi anlamak adına histogram grafiği Şekil 9’da verilmiştir.



Şekil 9. 100 birimlik düşük düzey korelasyona sahip veri seti için yöntemlerin performansları

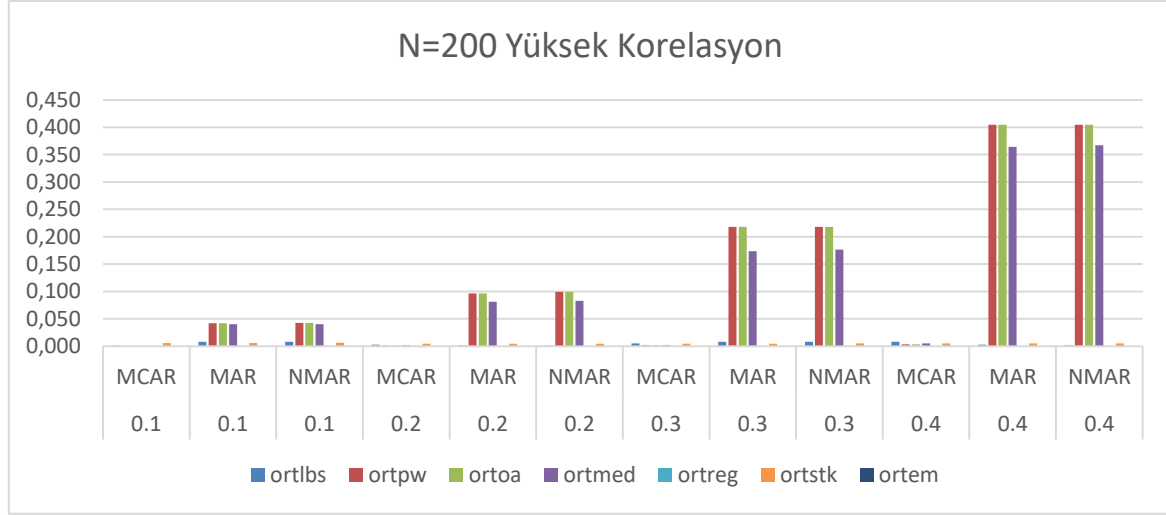
Düşük korelasyon yapısına sahip 100 gözlemlilik veri seti incelendiğinde MCAR ve MAR veri yapılarının hata kareler ortalamaları düşük iken, NMAR veri yapısına sahip gözlemlilerin ortalamaları yüksek değerlere sahiptir.

200 birim içeren tam veri yapısına %10, %20, %30 ve %40 oranında silme işleminden sonra eksik veriyi gidermeye yönelik işlemler uygulanmıştır bu işlemlerin sonuçları Çizelge 10’da verilmiştir.

Çizelge 10. N=200 ve yüksek korelasyon durumunda yöntemlerin HKO performansları

	0.1	0.1	0.1	0.2	0.2	0.2	0.3	0.3	0.3	0.4	0.4	0.4
	MCAR	MAR	NMAR	MCAR	MAR	NMAR	MCAR	MAR	NMAR	MCAR	MAR	NMAR
ortlbs	0,001	0,008	0,008	0,003	0,002	0,001	0,005	0,008	0,008	0,008	0,003	0,002
ortpw	0,001	0,042	0,043	0,001	0,097	0,099	0,002	0,218	0,218	0,003	0,404	0,404
ortoa	0,001	0,042	0,043	0,001	0,097	0,099	0,002	0,218	0,218	0,003	0,404	0,404
ortmed	0,001	0,040	0,041	0,002	0,081	0,083	0,002	0,173	0,177	0,006	0,364	0,367
ortreg	0,001	0,001	0,001	0,001	0,001	0,001	0,001	0,001	0,001	0,001	0,001	0,000
ortstk	0,006	0,006	0,006	0,005	0,005	0,005	0,005	0,005	0,005	0,005	0,005	0,005
ortem	0,001	0,001	0,001	0,001	0,001	0,001	0,001	0,001	0,001	0,001	0,001	0,000

Çizelge sonucunda ortalama medyan atama sonuçlarının yüksekliği göze çarpmaktadır. Bu sonuçları değerlendirebilmek adına Şekil 10’da grafik verilmiştir.



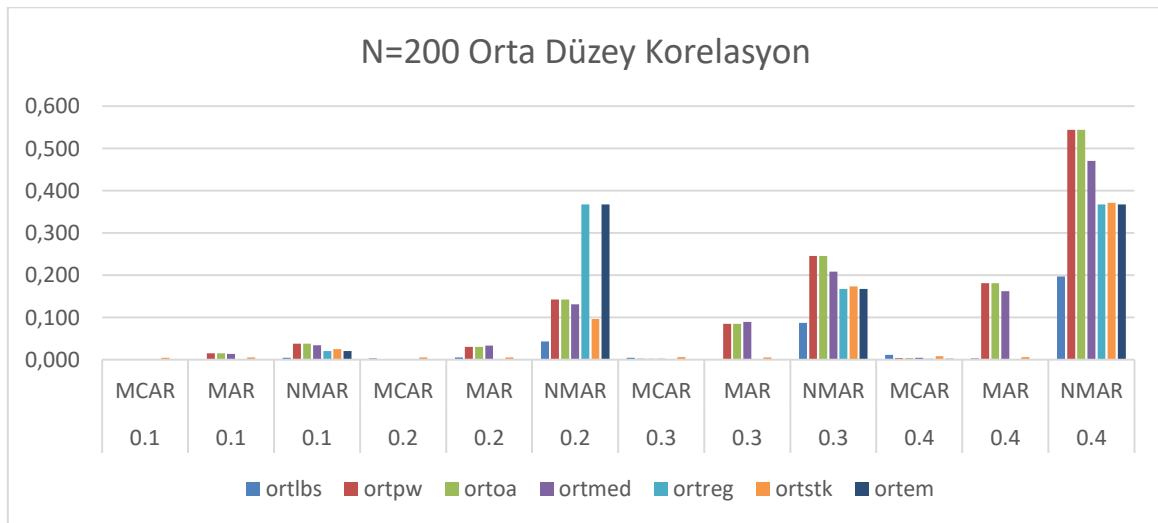
Şekil 10. 200 birimlik yüksek düzey korelasyona sahip veri seti için yöntemlerin performansları

200 birimlik, yüksek korelasyon yapısına sahip bu veri setinin grafiğini incelediğimizde MAR veri yapısının özellikle medyan atama değerlerinin yüksekliği göze çarpıyor. MAR ve NMAR veri yapısına sahip gözlem birimleri en yüksek hata kareler ortalama değerine sahip iken MCAR veri yapısına sahip gözlem birimleri en düşük değerlere sahiptir.

200 birimlik ortalama korelasyon yapısındaki eksik veri atama ve silme yöntemlerinin performansları Çizelge 11’de karşılaştırılmıştır.

Çizelge 11. N=200 ve orta düzey korelasyon durumunda yöntemlerin HKO performansları

	0.1	0.1	0.1	0.2	0.2	0.2	0.3	0.3	0.3	0.4	0.4	0.4
	MCAR	MAR	NMAR	MCAR	MAR	NMAR	MCAR	MAR	NMAR	MCAR	MAR	NMAR
ortlbs	0,001	0,001	0,005	0,003	0,006	0,044	0,005		0,087	0,011	0,004	0,197
ortpw	0,001	0,015	0,038	0,001	0,030	0,143	0,002	0,085	0,245	0,004	0,181	0,544
ortoa	0,001	0,015	0,038	0,001	0,030	0,143	0,002	0,085	0,245	0,004	0,181	0,544
ortmed	0,001	0,014	0,035	0,001	0,033	0,131	0,003	0,090	0,208	0,005	0,162	0,471
ortreg	0,000	0,000	0,021	0,001	0,002	0,368	0,002	0,001	0,168	0,003	0,000	0,368
ortstk	0,005	0,005	0,026	0,005	0,005	0,097	0,007	0,006	0,174	0,009	0,006	0,371
ortem	0,000	0,000	0,021	0,001	0,002	0,368	0,002	0,001	0,168	0,003	0,000	0,368

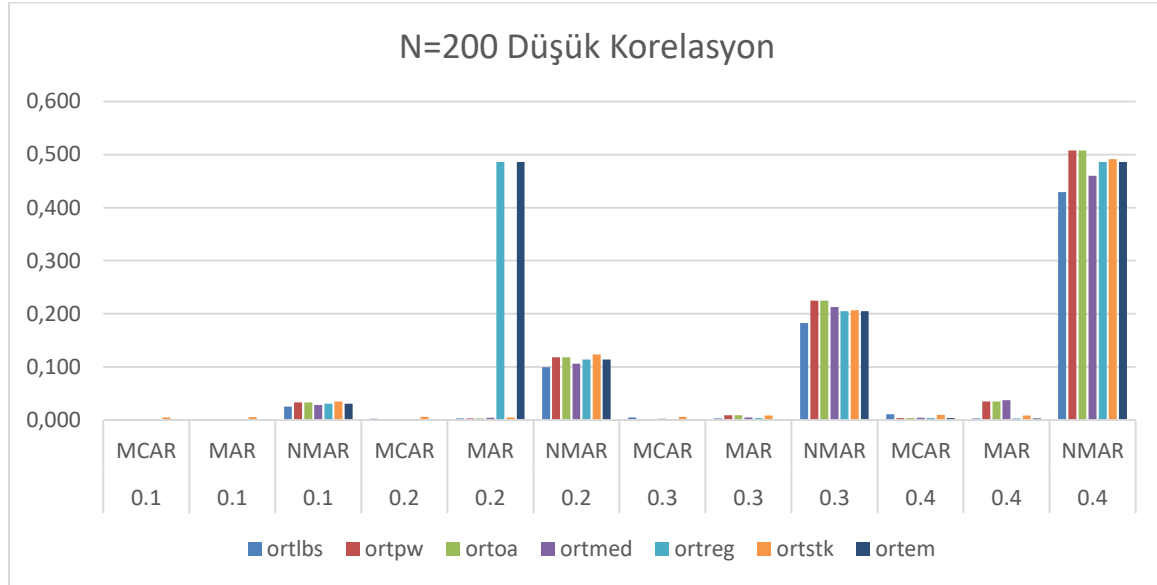


Şekil 11. 200 birimlik orta düzey korelasyona sahip veri seti için yöntemlerin performansları

Orta düzey korelasyon yapısı incelendiğinde MAR yapısında %40 eksik gözlem yapısına sahip verilerin liste bazında silme, regresyon atama, stokastik regresyon atama ve em atama değerlerinin düşük olduğu görülmektedir. Orta düzey korelasyon yapısında 200 birimlik veri setini incelediğimizde MAR yapısının hata kareler ortalaması %30 eksik oranından itibaren artarken, NMAR bu artış yüzde 20'den itibaren görülmektedir. MCAR veri yapısı ise en düşük ortalama değere sahiptir.

Çizelge 12. N=200 ve düşük korelasyon durumunda yöntemlerin HKO performansları

	0.1	0.1	0.1	0.2	0.2	0.2	0.3	0.3	0.3	0.4	0.4	0.4
	MCAR	MAR	NMAR	MCAR	MAR	NMAR	MCAR	MAR	NMAR	MCAR	MAR	NMAR
ortlbs	0,001	0,002	0,025	0,003	0,003	0,099	0,005	0,003	0,183	0,011	0,003	0,429
ortpw	0,001	0,001	0,034	0,001	0,003	0,118	0,002	0,010	0,225	0,004	0,035	0,508
ortoa	0,001	0,001	0,034	0,001	0,003	0,118	0,002	0,010	0,225	0,004	0,035	0,508
ortmed	0,001	0,001	0,029	0,001	0,004	0,107	0,003	0,005	0,213	0,004	0,037	0,460
ortreg	0,000	0,001	0,031	0,001	0,486	0,114	0,002	0,004	0,205	0,004	0,003	0,486
ortstk	0,005	0,005	0,035	0,006	0,005	0,124	0,006	0,009	0,207	0,010	0,009	0,491
ortem	0,000	0,001	0,031	0,001	0,486	0,114	0,002	0,002	0,205	0,004	0,003	0,486



Şekil 12. 200 birimlik düşük korelasyon durumunda yöntemlerin performanları seviyesine sahip veri seti

Şekil 12 incelendiğinde ise MAR ve MCAR veri yapılarının hata kareler ortalaması en düşük değerlere sahip iken NMAR veri yapısındaki ortalamalar en yüksek değerlere sahiptir.

SONUÇLAR VE ÖNERİLER

Günümüzde verinin değeri yadsınamaz veriyi kullanan işletmeler ve yapılar rakiplerine göre her zaman bir adım önde olacaktır. Verinin bu denli önemli bir konumda olması eksik problemini de ön plana çıkarmaktadır. Eksik veri araştırmacılar için günün sonunda kaçınılmaz bir sonuçtur. Burada tartışılması gereken asıl problem eksik veriden kaçınmak değil, eksik veriyle nasıl mücadele edilebileceğine dair olmalıdır.

Simülasyon çalışması sonucunda üretilen verilerin analizlerin performanslarını değerlendirebilmek için çeşitli eksik veriyi gidermeye yönelik çalışmalar (Listwise Deletion, Pairwise Deletion, Mean Imputation, Median Imputation, Regression Imputation, Stochastic Regression Imputation, Expectation-Maximization Imputation) yöntemleri uygulanmıştır. Hata kareler ortalamasına

bakıldığında en düşük hata karelerin MCAR (Missing Completely at Random) yapılarında olduğu gözlenmiştir. MAR ve NMAR yapılarındaki verilerde de yanlış sonuçların ortaya çıktığı görülmektedir.

Kaynaklar

- Rubin, D. B. (1976). Inference and missing data. *Biometrika*, 63(3), 581-592.
- McKnight, P. E., McKnight, K. M., Sidani, S., & Figueredo, A. J. (2007). *Missing data: A gentle introduction*. Guilford Press.
- Little, R. J. (1988). A test of missing completely at random for multivariate data with missing values. *Journal of the American statistical Association*, 83(404), 1198-1202.
- Sarı İsmet (2012). Karma ve ayrıştırma analizinde kayıp gözlem tahmin yöntemlerinin değerlendirilmesi
- Forbes, 2017. Data is the new oil and that's a good thing
<https://www.forbes.com/sites/forbestechcouncil/2019/11/15/data-is-the-new-oil-and-thats-a-good-thing/?sh=649d7d227304>

Ülkelerin Endüstri 4.0 Etkilerinin İnsan, Makine ve Ürün Bağlamında Kümeleme Analizi

Sennur YILDIRIM^{1*}, Zeynep ÖZTÜRK¹

¹Artvin Çoruh Üniversitesi, Lisansüstü Eğitim Enstitüsü, İşletme Bölümü, 08000, Artvin, Türkiye

*Sunucu yazar e-posta: sennuryildirim00@gmail.com

Özet

Endüstri 4.0'ın sanayi alanına getirmiş olduğu yeniliklerin etkisiyle, verimliliğin arttırılarak kaynak kullanımının en yüksek seviyede yeralması, iş akış süreçlerinin hızlanıp, maliyetlerin düşmesi, üretimde yer alan insan gücü yerine makineleri kullanarak, insandan kaynaklı hataları en aza indirerek, üretim hızını arttırma çalışmaları günümüzde hızla devam etmektedir. Endüstri 4.0'ın getirmiş olduğu yenilikler her ne kadar endüstriyel bazda bir gelişme olarak değerlendirilse de, insan hayatını her alanda etkileyen uygulamalar ile öne çıkmaktadır. Bu çalışmada, Şangay, OECD ve diğer ülkelerin kümeleme analizi yöntemiyle Endüstri 4.0 etkilerinin insan, makine ve ürün bağlamında kümelenmesi amaçlanmıştır. Bu bağlamda Makine ve Ulaşım Ekipmanları, Orta ve Yüksek Teknoloji Endüstrisi, Küreselleşme Endeksi, Dünya Ekonomik Özgürlükler Endeksi, Küresel İnovasyon Endeksi, Üretimin Geleceğine Hazırlık Raporu, Küresel Rekabet Endeksi, İnsani Gelişme Endeksleri değişken olarak ele alınmıştır. Endüstri 4.0 bağlamında ürün, insan ve makine kapsamında ele alınan değişkenler sayesinde Hiyerarşik kümeleme Ward Yöntemi ile ülkeler 4 kümeye ayrılmıştır. Bu kümeler Endüstr 4.0 da ülkelerin gelişmişliklerinin dağılımlarını yansıtmaktadır.

Anahtar kelimeler: Endüstri 4.0, Yapay zeka, İnavasyon, Kümeleme analizi

Cluster Analysis of Countries' Industry 4.0 Effects in the Context of Human, Machine and Product

Abstract

With the effect of the innovations that Industry 4.0 has brought to the industrial field, efforts to increase the production speed by increasing the efficiency and maximizing the use of resources, accelerating the work flow processes and reducing the costs, by using machines instead of manpower in production, minimizing human-induced errors, are rapidly increasing today. continues. Although the innovations brought by Industry 4.0 are considered as an industrial development, they stand out with applications that affect human life in every field. In this study, it is aimed to cluster the effects of Industry 4.0 in the context of human, machine and product with the cluster analysis method of Shanghai, OECD and other countries. In this context, Machinery and Transportation Equipment, Medium and High Technology Industry, Globalization Index, World Economic Freedoms Index, Global Innovation Index, Preparedness Report for the Future of Production, Global Competitiveness Index, Human Development Indices are considered as variables. In the context of Industry 4.0, countries are divided into 4 clusters with the Hierarchical Clustering Ward Method, thanks to the variables considered within the scope of product, human and machine. These clusters reflect the distribution of countries development in Industry 4.0.

Key words: Industry 4.0, Artificial Intelligence, Innovation, Cluster Analysis

GİRİŞ

Su ve buhar enerjili, mekanik üretim sistemlerinin ortaya çıkmasıyla başlayan 18. yy'daki "Sanayi 1.0" 20. yy başlangıcında elektriğin ve kitle üretim sistemlerinin etkisiyle "Sanayi 2.0," 1970'lerle birlikte "üretim otomasyonunu daha üst düzeye taşıyan elektronik ve bilgi teknolojilerinin devreye girmesiyle" de "Sanayi 3.0" (Şahin, 2019) yenilikleri getirmiştir.

2011 yılında internetin de yaygın kullanılmasıyla iletişim ve ulaşımın ilerlemesiyle hem üretimde hem de tüketimde hızlı bir şekilde yeni bir dönemin başlangıcı olan Sanayi 4.0 kavramı ile (Taş & Küçükoğlu, 2019) günümüzde de halen devam eden yenilikler hız kazanmıştır.

Nesnelerin interneti, Siber-Fiziksel Sistemler, Big Data gibi kavramların getirdikleri yenilikler sonucunda üretim sürecini hızlandırabilen hata oranını ortadan kaldıran 3 Boyutlu (3B) baskı, gibi üretimdeki verimliliği artıran faktörlerin etkisiyle ortaya çıkan Endüstri 4.0, başlangıcında üretim esaslı ortaya çıksa da etki alanının sadece üretim olmayacağı hayatın hemen her alanında etkisinin (Şahin, 2019) görüleceği öngörülmüştür. Günümüzde insanların bireysel kullandığı uygulamalar da bu öngörünün gerçekleştiğini göstermektedir.

Endüstri 4.0 kısa bir sürede teknik bir tanım olmaktan çıkıp milyon dolarlık yatırım sahası haline gelmiştir. Ülkelerin küresel sahada başka ülkelerle rekabet edebilmeleri için Endüstri 4.0 ile birlikte gelen yenilikleri yakından takip edip uygulamada öncü olmaları gerektiği, büyük verileri doğru kullanarak müşteri beklentilerini en iyi şekilde analiz etmelerinin (Soylu, 2018) önemi her geçen gün daha fazla artmaktadır.

Endüstri 4.0'ın getirmiş olduğu yenilik ve faydaların yanısıra gelişmekte olan ülkelerin, bu yeniliklere adapte olma noktasında yaşadıkları zorlukların yanında ileri teknolojilere sahip ülkelerin ürün yelpazesinin geniş olması ve insanların bu ürünlere ulaşımının çok daha kolay olması gelişmekte olan ülkelerin bu erişimi ve çeşitliliği aynı kolaylıkta halkına sağlayamaması bu ülkeleri oldukları yerden daha fazla geriye götüreceği anlamına gelebilmektedir.

Ülkelerin endüstri 4.0'ın etkileri nedeniyle karşılaştıkları değişime karşı üretim sistemleri ile ilgili gelişen teknolojilere açık olarak diğer doğru adımları atmalarının halklarının refah düzeyine de etkisi olacağı düşünülmektedir.

Ülkeler için sadece Endüstri 4.0 ile birlikte gelişen ve gelişmekte olan kavramların getirdiği yeniliklerden bahsetmek eksik kalacaktır. Çünkü Endüstri 4.0'ın uygulama alanlarının, yapay zeka uygulamalarıyla desteklendiği gerçeği göz ardı edilemez boyuttadır. Hemen her sektörde etkisi gözlemlenen, Endüstri 4.0 ile yapay zeka uygulamalarının getirdiği yeniliklere karşı uyum gösteren ülkeler, sürdürülebilir olmalarını sağlamamanın yanında kendi bulundukları bölgelerde rekabet yönünden de güçlerini arttırdığı bilinmektedir.

Bu çalışmanın amacı, Şangay, OECD ve diğer ülkelerin kümeleme analizi yöntemiyle Endüstri 4.0 etkilerinin insan, makine ve ürün bağlamında kümelenmesi amaçlanmış olup bu bağlamda Makine ve Ulaşım Ekipmanları, Orta ve Yüksek Teknoloji Endüstrisi, Küreselleşme Endeksi, Dünya Ekonomik Özgürlükler Endeksi, Küresel İnovasyon Endeksi, Üretimin Geleceğine Hazırlık Raporu, Küresel Rekabet Endeksi, İnsani Gelişim Endeksi değişkenlerini içine alan Endüstri 4.0 anlamında benzerlikleri bir araya getirilerek kümelere ayrılmış ve birbirleriyle etkileşimleri incelenmiştir.

Literatür Taraması

Ünlü ve Atık (2018), çalışmada Endüstri 4.0 kapsamında 28 Avrupa Birliği ülkesi ve Türkiye ile birlikte toplam 29 ülkenin, Monitoring the Digital Economy & Society 2016-2021 adlı Avrupa Komisyonu tarafından yayınlanan verilerinden, 2015-2016 dönemini kapsayan farklı dönemlere ait Eurostat tarafından yayınlanan dijital ekonomi göstergelerini içeren veri tabanı ve TÜİK'ten alınan veriler

çerçevesinde belirlenen 10 değişken ; Kurumsal kaynak planlamasını (ERP) kullanan firmaların payı, Müşteri ilişkileri yönetimini (CRM) kullanan firmaların payı, Tedarik zincirinde bilgi paylaşımını gerçekleştiren firmaların payı, İnternete mobil bağlantı için taşınabilir cihazlar kullanan firmaların payı, İnternet üzerinden sipariş alan firmaların payı, Müşteri İlişkileri Yönetimi (CRM) gibi yazılımlar kullanan firmaların payı, Farklı fonksiyonel alanlar arasında bilgi paylaşmak için kurumsal kaynak planlama (ERP) yazılım paketine sahip olan firmaların payı, Farklı fonksiyonel alanlar arasında bilgi paylaşmak için kurumsal kaynak planlama (ERP) yazılım paketine sahip olan firmaların payı, Ücretli bulut bilişim uygulamalarını kullanan firmalar, Kamu kurumları ile iletişimde interneti kullanan firmaların oranı, şeklinde, Endüstri 4.0 kapsamında performans düzeylerini belirleyebilmek için faktör analizi, birbirlerine benzeyen ülkeler ile Türkiye'nin hangi ülkelerin birbiri ile benzerlik gösterdiği ayrıca Türkiye'nin hangi ülkeler ile aynı grupta yer aldığını belirlemek amacı ile de kümeleme analizi yapılmıştır. Çalışma sonucunda Endüstri 4.0 çerçevesinde en iyi performansa sahip ülkenin Almanya olduğu, 15. Sıradaki Türkiye'nin Macaristan, Letonya ve Polonya ile aynı küme içerisinde yer alırken, 29 ülkenin Endüstri 4.0 perspektifinde homojen bir görünüm sergilemediği belirlenmiştir.

Duman (2020), Endüstri 4.0'ın Teknoloji Bileşenlerinin Örgütsel Performansa Etkilerini Belirlemeye Yönelik Araştırmasında, Endüstri 4.0 teknoloji bileşenlerinin işletmeler üzerindeki performansı ve işletmelere sağlayacağı faydaları ortaya koymaya çalışmıştır. Pilot uygulamasını Malatya Organize Sanayi Bölgesi'nde yaparken, Endüstri 4.0 teknolojilerini kullanan ve Endüstri 4.0 adına stratejik hedefleri olan işletmeler için Kayseri Organize Sanayi Bölgesi'ndeki işletmeler araştırmada ana kütle olarak belirlenmiştir. Üç bölümden oluşan anket formu her soru birebir araştırmacı tarafından açıklanarak çoğunlukla üst yöneticilerden biriyle ve azınlık olarak mühendislerle görüşülerek alınan cevaplar doğrultusunda doldurulmuştur. Anket formu üç bölümden oluşmaktadır. Birinci bölümde işletmelere ait kurumsal bilgilere yer verilmiştir. İkinci bölümde olgunluk düzeyi ölçülmüş olup teknoloji bileşenleri içerisinde yer alan yedi bileşenin işletme içerisinde uygulanıp uygulanmadığı beşli likert ölçek ile sorulmuş. Endüstri 4.0 konusunda işletmeleri zorlayan sebepleri öğrenmek için iki soru sorulmuş bu sorular kapsamında siber güvenlik düzeyleri ölçümlenmeye çalışılmıştır. Üçüncü bölümünde ise örgüt performans değerlerini ölçümlemek amacıyla sorular sorulmuştur. Elde edilen veriler IBM SPSS 21 programı ile istatistiki analizler yapılmıştır. Analiz sonucunda Endüstri 4.0 teknoloji bileşenlerini kullanan işletmelerin örgütsel performanslarını, işletmelerin kişi başına üretim miktarlarını, kârlılıklarını, maliyetlerini, satışlarını, üretim hızlarını ve ürün kalitesini olumlu yönde etkilediğini görülmüştür.

Gürler (2020), Endüstri 4.0 ve Bir Kümeleme Analizi Uygulaması çalışmasında, Dünya Ekonomik Forumu tarafından yayınlanan Küresel Rekabetçilik Endeksi ile Huawei tarafından yayınlanan Küresel Bağlanabilirlik Endeksinde yer alan toplam 15 değişken analizde benzer ülkelerin kümelenmesi amaçlanmıştır. Analizde kullanılan değişkenler; Küresel Rekabetçilik Endeksinde Kurumsal yapılanma, Altyapı, Makroekonomik ortam, Sağlık, Beceriler, Mal piyasası, İşgücü piyasası, Mali sistem, Pazar büyüklüğü, İş dünyasının dinamizmi, İnovasyon kabiliyeti olarak 11 değişken yer alırken, Küresel Bağlanabilirlik Endeksi içerisinde Geniş bant, Bulut, Yapay zeka, Nesnelerin interneti 4 değişken ile birlikte büyümenin önemli bir göstergesi olan GSYİH verileri de analize eklenmiştir. Elde edilen verilerle benzer ülkeleri bir araya getirebilmek için dirsek yöntemi ile 78 ülkeden 7 küme şeklinde olacak şekilde belirlendikten sonra hiyerarşik olmayan kümeleme yöntemlerinden k-ortalamalar yöntemi kullanılarak her bir kümenin endeks değerleri hesaplanarak kümeler arasında karşılaştırmalar yapılmıştır. Türkiye'ninde içerisinde yer aldığı, Küme 5'te Cezayir, Arjantin, Ekvador, Mısır, Ürdün, Lübnan, Sırbistan ve Ukrayna'da bulunmaktadır. Kurumsal yapılanma ve alt yapıda Türkiye'nin, Rusya'dan daha iyi bir skora sahip olduğu, sağlık konusunda ABD'den çok daha iyi durumda olduğu, pazar büyüklüğünde ise ülkeler arasında 13. ülke içerisinde yer aldığı, makroekonomik ortam da analiz içerisinde yer alan 78 ülke içerisinde 72. sırada yer bulduğu görülmüştür.

Şahin (2019), çalışmasında Ülkemizde içerisinde bulunan G-20 ülkelerinin Endüstri 4.0 ve dijital dönüşüm için sahip olması gereken kriterlerinin belirlenmesi için Entropi tabanlı COPRAS yöntemi kullanılarak, endüstri 4.0 ve dijital dönüşüm seviyelerinin ölçülmesi ve sıralamaları için belirlenmiş olan değişkenler; açık pazar endeksi, bilgi ve iletişim teknolojileri gelişmişlik endeksi, dünya ekonomik özgürlükler endeksi, e-devlet kalkınma endeksi, kof küreselleşme endeksi, küresel rekabet endeksi, şebekeleşmiş hazır bulunuşluk endeksi, e-katılım endeksi, orta-ileri teknoloji üretimin toplam üretim içerisindeki payı, global inovasyon endeksi, küresel girişimcilik monitörü, internet penetrasyon bağlanma oranı (toplam nüfus içerisindeki internet kullanıcı sayısı), toplam çalışan içerisinde tarım sektöründe çalışan sayısı oranı, 15-64 yaş aralığının toplam nüfusun içerisindeki payı, toplam ihracat içerisinde yüksek teknoloji ihracatıdır. Çalışmada COPRAS yönteminin uygulama aşamaları anlatılmış olup G-20 içerisinde yer alan ülkelere ait 2015-2018 yıllarından elde edilmiş verilerle bu ülkelerin Endüstri 4.0 seviyelerinin kendi aralarında karşılaştırılması yapılmış olup, bu çalışmada ülkemiz 15. olmuştur

Barutçu (2019) çalışmasında, Endüstri 4.0'ın öncü firmalarından olan 2019 Ocak ayında Bosch Sanayi ve Ticaret Anonim Şirketi'nin Bursa'da bulunan fabrikasında üretim süreçlerinde kullanılan Endüstri 4.0 uygulamaları ve üretim süreçleri incelenmiş ve etkileri ortaya konulmuştur. Bu etkilerin verileri derinlemesine mülakat yöntemi kapsamında sözel iletişim teknikleriyle toplanmıştır. Bosch Bursa Fabrikası'nda görev yapan bir Endüstri 4.0'dan Sorumlu Üretim Mühendisi, bir Proje Mühendisi, bir Kıdemli Üretim Mühendisi ve bir Üretim Mühendisi tarafından cevaplanmıştır. Araştırma kapsamında görüşmecilere derinlemesine mülakat kapsamında; Üretim süreçlerinde kullanılan Endüstri 4.0 uygulamaları nelerdir? Endüstri 4.0 Uygulamalarının Üretim Sistemlerine bir etkisi var mıdır? Endüstri 4.0 Uygulamalarının Tedarik Zinciri Yönetimine bir etkisi var mıdır? Endüstri 4.0 Uygulamalarının Kaynak Planlamasına bir etkisi var mıdır? Endüstri 4.0 Uygulamalarının Stok Yönetimine bir etkisi var mıdır? Endüstri 4.0 Uygulamalarının Süreç Yönetimine bir etkisi var mıdır? Endüstri 4.0 Uygulamalarının Toplam Kalite Yönetimine bir etkisi var mıdır? Kullanılan Endüstri 4.0 Uygulamalarının süreç iyileştirme ve çözüm üretme gibi konularda ne gibi etkileri vardır? Şeklinde sorular 15 tane soru yöneltilmiştir. Araştırmanın sonucunda Endüstri 4.0'ın fabrikanın her aşamasında olumlu etkilerinin olduğu, yüksek verimlilik ve kalite oranlarına ulaştığı ortaya çıkmıştır.

Kabaca (2019) çalışmasında, Endüstri 4.0 dijital dönüşümün gerekliliğinin ve etkilerinin belirlenmesine yönelik bir çalışma olup araştırma bu kavramının çıkış noktası olan Almanya'da gerçekleştirilmiştir.

MATERYAL VE YÖNTEM

Kümeleme Analizi

Kümeleme analizi çok değişkenli istatistiksel analiz yöntemlerinden biri olup, değişik yapıdaki verilerin küme sayı ve yapısını araştırarak, veri seti içerisinde birimleri, değişkenleri birbiri ile benzer alt kümeler (grup, sınıf) ayıran analiz yöntemidir (Yalçın, 2013).

Gruplanmış verileri benzerliklerine göre sınıflandırma da kullanılan bu yöntem, birey ya da nesnelerin temel özelliklerini inceleyerek onları gruplandırmaktır. Ancak kümeleme analizi ile veri seti içerisindeki gerçek tiplerin belirlenmesi, oluşturulan gruplar için ön tahmin, verilerin tek tek değil küme halinde değerlendirilmesi, hipotez testi ve veri seti içerisinde varsa aykırı değerlerin bulunması gibi amaçlarla da kullanılmaktadır (Kalaycı, 2014).

Kümeleme analizinin temeli, veri seti içerisindeki değişkenler arasındaki benzerlik ölçümlerinin hesaplanmasıdır. Bir değişkenin hangi kümeye ait olduğunu benzerlik ölçümlerini gösteren geometrik modeller aracılığıyla belirlenirken, Kümeleme analizinin; “mesafeye, her bir kümenin merkez noktasını

bulmaya ve dağılıma dayanan istatistiksel teknikler” kullandığı yöntemlerdendir (Uğuz, 2019). Özdamar (2018) ise, kümeleme analizinin kullanım amaçlarını;

1. n sayıda birimi, nesneyi ve oluşumu, p değişkene yönelik özelliklerine göre, kendi içinde oldukça homojen ve kendi aralarında da heterojen alt kümelerle ayırmak.
2. p sayıda olan değişkenin, n sayıda birimde saptanarak elde edilen değerlere göre ortak özellikleri açıkladığı varsayılan alt kümelerle ayırarak ortak faktör yapıları ortaya çıkarmak,
3. Birimleri ve değişkenleri birlikte ele alarak n birimi p değişkene göre ortak özellikli alt kümelerle ayırmak,
4. Birimlerin, P değişkene göre saptanan yapılar aracılığıyla toplumdaki biyolojik ve taksonomik sınıflandırmayı yapabilmek olarak açıklamıştır.

Kümeleme Analizinin uygulama aşamaları aşağıdaki gibidir;

- a) Veri içerisinde yer alan değişkenlerin seçilerek veri matrisinin belirlenmesi,
- b) Birimlerin, değişkenlerin veya her birinin birbirleriyle olan benzerlik veya uzaklık ölçülerinin oluşturulması,
- c) Veri seti üzerinde uygulanacak kümeleme tekniği belirlenerek uygun sayıda kümelerin oluşturulması,
- d) Üzerinde çalışılan veri seti içerisinde elde edilen kümelerin yorumladıktan sonra, bu küme yapıları baz alınarak kurulan hipotezlerin doğruluğunu gösterebilmek amacıyla gerekli analitik yöntemleri uygulamak, denetlemek ve test etmek (Alpar, 2013),
- e) Örneklem büyüklüğünün çalışma için yeterli olup olmadığı,
- f) Veri seti içerisinde uç değerlerin olup olmadığı, uç değerlerin olması halinde kaldırılıp kaldırılamayacağını belirlenmesidir (Çokluk, Büyüköztürk, & Şekercioğlu, 2018).

Benzerlik ve Uzaklık Ölçüsü

Kümeleme analizi, gözlem ve değişkenlerle elde edilmiş ham veri setlerinin, birimleri, değişkenleri veya birim ve değişkenleri, yakınlık, uzaklık, benzerlik ya da farklılıkları baz alarak kendi içinde homojen gruplara ayrılırken aralarında heterojen, farklı farklı kümelerle bölünmesi işlemidir. Kümeleme analizinin temel amacı, değişkenler arasındaki benzerlikleri ya da uzaklıkları tespit etmektir. Kümeleme analizinde iki değişken arasındaki ilişkinin kuvveti benzerlik ölçüsü ile açıklanırken alınan sonuç ölçüğe ve veri tipine göre farklı şekillerde elde edilir. Uzaklık ölçüsü ise, değişkenler arasındaki zıtlık veya uyumsuzluğun sonucundaki farklılıkları ölçerek değişkenler içerisinde yer alan verilerin nicel veya karışık veriler olmamasına göre farklılık gösterir (Yaz, 2014). Kümeleme analizinde uzaklık ve benzerlik ölçüleri amaca göre gözlem ve değişken için hesaplanabilmektedir. Gözlemler için “n×n” boyutlu uzaklık veya benzerlik matrisi ile hesaplanırken değişkenleri ise “p×p” uzaklık veya benzerlik matrisi ile hesaplanır (Alpar, 2013)

Aşamalı (Hiyerarşik) Kümeleme Yöntemi

Hiyerarşik kümeleme yöntemleri veri seti içerisinde yer alan gözlem veya değişkenleri kümelemek amacıyla birbirlerine uygun uzaklık veya benzerlik ölçülerini kullanarak değişik aşamalarda birleştiren yöntemdir (Özdamar, 2018).

Hiyerarşik kümeleme yöntemi, veri seti içerisinde kaç grup olduğu bilinmeyen durumlarda kullanımının uygun bir yöntem olduğu yine bu yöntemin veri seti içerisinde daha önce gözlenmemiş birbirleriyle olan ilişkileri gözleme olanağı sağlaması nedeniyle de tercih edilen bir yöntem olduğu belirtilmektedir (Çokluk, Büyüköztürk, & Şekercioğlu, 2018).

İki çeşit hiyerarşik kümeleme algoritması vardır bunlardan biri Birleştirici (Agglomerative) Algoritmalar, diğeri ayrıştırıcı (divisive) Algoritmalar. Birleştirici hiyerarşik kümeleme yöntemi,

başlangıçta her bir gözlemin tek başlarına ayrı birer küme olduğu kabul edilerek sonrasında en yakın iki kümenin tek bir küme içerisinde birleştirilerek devam etmesi sonucunda tüm gözlemlerin bir küme içinde gruplanmasıdır. Ayrıştırıcı hiyerarşik kümeleme yöntemi ise birleştirici kümeleme yöntemlerindeki sürecin tersi işleyerek tüm gözlemleri içerisine alan tek bir küme bulundurmaktadır. Daha sonraki adımlarda birbirlerine benzemeyen uzak gözlemler birbirinden ayrılarak daha küçük kümeler oluşturulur (Alpar, 2013)

Tek Bağlantı Yöntemi, tam bağlantı yöntemi, ortalama bağlantı yöntemi ve Ward yöntemi, sıklıkla kullanılan hiyerarşik kümeleme yöntemleridir. Tek Bağlantı Yöntemi; en yakın komşuya ait en kısa mesafede ve en çok benzerliği olan iki kümenin birleştirilmesi olarak tanımlanır. Bu yöntemle uzaklıklar matrisi kullanılarak birleştirme işlemi gerçekleştirilmekte ve peş peşe tekrar edilerek devam ettirilmektedir (Mercan & Atalay, 2021). Tam bağlantı yönteminin, tek bağlantı yönteminden tek farkı en uzak iki gözlemden başlaması olarak tanımlanmaktadır. Ortalama Bağlantı Yöntemi, aşırı uç gözlemlerden başlamayarak bir kümenin ortasına düşen gözlemi baz alması şeklindedir. Ward yöntemindeki amaç, küme içindeki varyansı en aza indirecek şekilde durumları kümeler halinde birleştirmektir. Bunu yapmak için, her gözlem kendi kümesi olarak başlar. Kümeler daha sonra bir küme içindeki değişkenliği azaltacak şekilde birleştirilir. Daha kesin olmak gerekirse, bu birleşme karelerin hata toplamında minimum artışla sonuçlanırsa iki küme birleştirilir. Temel olarak bu, her aşamada kümenin ortalama benzerliğinin ölçüldüğü anlamına gelir. Bir küme içindeki her durum ve bu ortalama benzerlik arasındaki fark hesaplanır ve karesi alınır. Kare sapmaların toplamı, bir küme içindeki hata ölçüsü olarak kullanılır. Kümeye dahil edilmesi hatada en az artışa neden olan durum ise, kümeye girmek için bir durum seçilir. Ward Yöntemi Benzerlik ölçüsü olarak Öklid Uzaklık Ölçüsü kullanır.

Öklid Uzaklık Ölçüsü

En çok kullanılan uzaklık ölçülerinden biri olan Öklid uzaklığı, daha önce üzerinde işlem yapılmamış ham verilerde kullanılırken, verilerin farklı ölçeklerden elde edilmesinden oldukça fazla etkilenip aykırı değerlerden aynı ölçüde etkilenemediği bilinmektedir. Büyük değerlere sahip veri içerisinde yer alan değişkenler öklid uzaklığını maksimize etme eğilimindedir (Deniz, 2020).

$$d_{ij} = \sqrt{\sum_{k=1}^p w_k^2 (x_{ik} - x_{jk})^2}$$

Hiyerarşik Kümelemede Küme Sayısını Belirlemek için Aşamalı (hiyerarşik) Kümeleme analizinde Ward yönteminde Öklid Uzaklık ölçüsü göz önüne alınarak yapılan kümelendirmelerin kolay anlaşılabilirliği için dendrogram (ağaç diyagramı) grafiğinden faydalanılır.

BULGULAR

Bu çalışmada, teknoloji, inovasyon ve bilginin rol aldığı Endüstri 4.0 ın etkilerini insan kapsamında İnsani Gelişme Endeksi, Küreselleşme Endeksi; makine kapsamında Makine ve Ulaşım Ekipmanları, Orta ve Yüksek Teknoloji Endüstrisi; ürün kapsamında Dünya Ekonomik Özgürlükler Endeksi, Küresel İnovasyon Endeksi, Üretimin Geleceğine Hazırlık Raporu, Küresel Rekabet Endeksi ele alınarak 52 ülkenin Endüstri 4.0 noktasında kümelenebilirliği amaçlanmaktadır. 52 ülke ve 8 değişken ile 2018 yılına ait veriler ile sınıflandırma işlemi Hiyerarşik Kümeleme analizi Ward Yöntemi kullanılarak yapılmıştır. Kümeleme analizinin en doğru sonucu verebilmesi için birimlerin kayıp gözlem içermemesi gerekmektedir. Bu çalışmaya ait veri setinde kayıp gözlem bulunmamaktadır. Çizelge 1’de analizde kullanılan değişkenler, değişkenlerin tanımlanması ve her değişkene ait ortalama ve standart sapma değerleri verilmiştir.

II. International Applied Statistics Conference (UYİK-2021)
Tokat / Turkey, 29 June - 2 July 2021

Çizelge 1. Analizde kullanılan değişkenler ve istatistikleri

Sayı	Değişken Adı	Açıklama	Ortalama (Standart Sapma)
1	Makine Ve Ulaşım Ekipmanları (İmalatta Katma Değerin Yüzdesi) 2018	Makine ve Ulaşım Ekipmanları (İmalatta Katma Değerin Yüzdesi) hakkındaki veriler, dünya bankasının kaynak olarak gösterdiği Birleşmiş Milletler Sınai Kalkınma Örgütü (UNIDO), verileridir. Bu sektörde gerçekleşen teknolojik gelişmelerin diğer sektörlerle yayılmasının sağladığı avantajlardan dolayı ülkelerin gelişmesinde çok önemli bir role sahiptir (URL-1, 2021).	21,475 (13,330)
2	Orta Ve Yüksek Teknoloji Endüstrisi (İnşaat Dahil) (İmalat Katma Değerinin Yüzdesi) 2018	Dünya bankası verilerince de, orta ve yüksek teknoloji sanayi katma değerinin toplam imalat katma değeri içindeki oranı olarak tanımlanmaktadır. İmalat sanayisinin bir ülkenin gelişmişliğinin göstergelerinden biri olarak görülen imalat katma değer kişi başına düşen miktar ne denli az ise o ülkede o denli az gelişmiş olduğu belirlenmektedir (URL-2, 2021).	38,070 (15,650)
3	Küreselleşme Endeksi	Küreselleşme, insanlar, bilgi ve fikirler, sermaye ve mallar dahil olmak üzere çeşitli akışlar aracılığıyla kıtalar arası mesafelerde aktörler arasında bağlantı ağları oluşturma süreci olarak tanımlanmaktadır. Küreselleşme Endeksi, dünyanın dört bir yanındaki ülkelerin küreselleşen ekonomik, sosyal ve politik boyutlarını ölçmeyi hedefliyor. Bu endekste, anket yapılan ülkelerdeki mevcut ekonomik akışları, ekonomik kısıtlamaları, bilgi akışlarına ilişkin verileri, kişisel iletişim verilerini ve kültürel yakınlık verilerini değerlendirmeye çalışır(URL-3, 2021).	77,96 (8,754)
4	Dünya Ekonomik Özgürlükler Endeksi	Kanada'nın önde gelen düşünce kuruluşu Fraser Institute tarafından 162 ülkenin verileri ile yayımlanan Dünya Ekonomik Özgürlükler Endeksi, "Devletin Büyüklüğü, Hukuk Sistemi ve Mülkiyet Hakları, Sağlam Para, Uluslararası Ticaret Yapma Özgürlüğü, Yönetmelik'ten oluşan beş ana alan içinde, 26 bileşeni bulunmaktadır. Bu bileşenlerin çoğu, birkaç alt bileşenden oluşur. Toplamda, 44 farklı değişkenden oluşmaktadır (URL-4, 2021).	7,395 (0,714)
5	Küresel İnovasyon Endeksi	Dünya Fikri Haklar Örgütü (WIPO), Cornell Üniversitesi ve INSEAD işbirliğinde hazırlanan 126 küresel ekonominin 7 ana alan ve 80 alt başlığıyla inovasyon performansı incelenmiştir. İnovasyon, yaygın olarak ekonomik büyüme ve kalkınmanın itici bir gücü olarak kabul edilmektedir (URL-5, 2021).	46,949 (11,144)
6	Küresel Rekabet Endeksi	Rekabet gücünün karşılaştırılmasında deneyime dayanan endeks, 141 ekonominin rekabet gücünü ele alan 12 tema halinde düzenlenmiş olup, 103 gösterge aracılığıyla haritalandırmıştır. Kullanan her gösterge, bir ekonominin ideal duruma veya rekabet edebilirliğin sınırına ne kadar yakın olduğunu göstermektedir. Geniş sosyo-ekonomik unsurları kapsayan ana başlıklar şunlardır; Kurumlar, Altyapı, BİT'in benimsenmesi, Makroekonomik istikrar, Sağlık, Beceriler, Ürün Piyasası, İşgücü Piyasası, Finansal Sistem, Pazar Büyüklüğü, İş Dinamizmi ve İnovasyon kabiliyetidir (URL-6, 2021).	71,292 (8,270)
7	İnsani Gelişim Endeksi	İnsani Gelişim Endeksi insani gelişmenin uzun ve sağlıklı yaşam, bilgiye erişim, ve özet bir ölçüm yöntemidir. "2019 İnsani Gelişim Raporu, 189 ülke ve BM tarafından tanınan toprak için 2018 İGE'yi (değerler ve sıralamalar), 150 ülke için EUİGE(Eşitsizliğe Uyarlanmış İnsani Gelişim Endeksi), 166 ülke için CDGE (Cinsiyete Dayalı Gelişim Endeksi) , 162 ülke için TCEE (Toplumsal Cinsiyet Eşitsizliği Endeksi) ve 101 ülke için ÇBYE (Çok Boyutlu Yoksulluk Endeksi)" verilerini vermektedir. Endüstri 4.0 ile ülkelerin insani gelişim endeksleri arasındaki etkileşim olup olmadığı irdelenmiştir (URL-7, 2021).	,858 (0,076)
8	Üretimin Geleceğine Hazırlık Raporu	Dünya Ekonomik Forumu ve A.T. Kearney iş birliğiyle hazırlanan Üretimin Geleceğine Hazırlık raporu, ülkelerin Dördüncü Sanayi Devrimi hakkındaki yaklaşımlarını ana tema olarak almıştır (URL-8, 2021).	6,019 (1,129)

II. International Applied Statistics Conference (UYİK-2021)
Tokat / Turkey, 29 June - 2 July 2021

Çizelge 2. Hiyerarşik küme liste

Sıra	Birleştirilmiş Küme		Katsayılar	İlk Görüldüğü Aşama Kümesi		Gelecek Aşama
	Küme 1	Küme 2		Küme 1	Küme 2	
1	26	40	7,465	0	0	9
2	5	7	18,805	0	0	19
3	9	20	30,677	0	0	27
4	29	44	47,964	0	0	10
5	30	41	65,309	0	0	13
6	2	19	83,410	0	0	19
7	36	45	102,636	0	0	30
8	1	15	123,826	0	0	27
9	26	32	146,191	1	0	33
10	29	37	177,392	4	0	17
11	11	14	210,794	0	0	43
12	49	50	247,155	0	0	26
13	30	46	284,883	5	0	46
14	12	43	325,305	0	0	33
15	6	8	372,761	0	0	44
16	22	33	420,784	0	0	39
17	29	51	472,369	10	0	22
18	16	23	526,983	0	0	31
19	2	5	585,411	6	2	28
20	13	17	646,396	0	0	37
21	42	47	716,888	0	0	23
22	29	35	789,496	17	0	25
23	18	42	863,899	0	21	37
24	3	38	940,994	0	0	32
25	27	29	1026,395	0	22	36
26	48	49	1116,438	0	12	42
27	1	9	1210,308	8	3	31
28	2	39	1308,808	19	0	38
29	10	25	1418,439	0	0	40
30	21	36	1535,010	0	7	34
31	1	16	1674,812	27	18	41
32	3	24	1825,761	24	0	38
33	12	26	1978,420	14	9	34
34	12	21	2145,576	33	30	45
35	31	52	2332,311	0	0	39
36	27	28	2611,238	25	0	48
37	13	18	2900,516	20	23	41
38	2	3	3202,793	28	32	47
39	22	31	3511,626	16	35	43
40	4	10	3852,776	0	29	47
41	1	13	4328,679	31	37	44
42	34	48	4895,256	0	26	48
43	11	22	5465,609	11	39	45
44	1	6	6207,136	41	15	46
45	11	12	7008,619	43	34	49
46	1	30	8206,127	44	13	51
47	2	4	9471,103	38	40	49
48	27	34	10789,799	36	42	50
49	2	11	13346,781	47	45	50
50	2	27	19522,104	49	48	51
51	1	2	35378,234	46	50	0

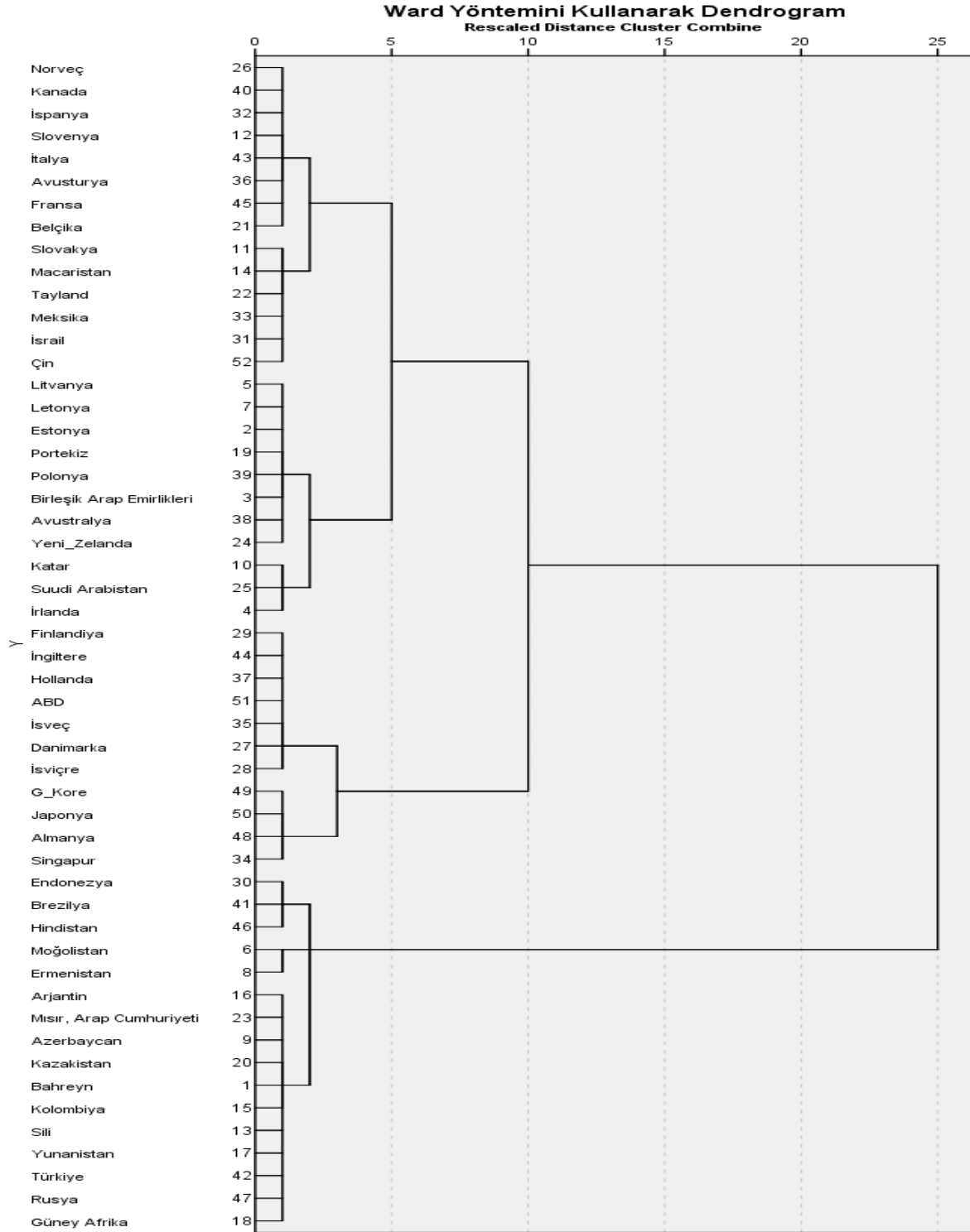
Çizelge 1’de değişkenlerin tanımları ve alındığı adresler belirtilmiştir. Ülkelerin Makine Ve Ulaşım Ekipmanları ortalaması 21.475, Ülkeler arasındaki değişim oranı 13.33 olup, ülkeler arasındaki fark yüksektir. İnsani Gelişim Endeksi ortalaması 0.858, standart sapması 0.076’dır. Bu sonuç ülkelerin İnsani Gelişim Endeks değerlerinin değişim oranlarının farklı olmadığını gösterir. Aşağıdaki Çizelge 2’de Kümeleme analizi yapılmış ve hiyerarşik küme liste sonucu verilmiştir.

Yukarıdaki Çizelge 2’deki tablo ile hangi ülkenin hangi aşamada hangi ülke ile kümelendiği de görülebilir. Hiyerarşik tabloda, 8 değişkene bağlı katsayıya göre birbirlerine en çok benzeyen ülkeler eşleştirilir. Çizelge 2’de kümeleme analizinin 51 aşaması olduğu görülmektedir. İlk aşamada 26. gözlem ile (Norveç) ve 40. gözlem (Kanada) arasındaki uzaklığın en az olduğu görülmektedir. İkinci aşamada ise, 5. (Litvanya) ve 7. (Letonya) ülkelerin birbirine yakın olduğu görülmektedir. 3. Aşamada ise 9. (Azerbaycan) ve 20. (Kazakistan) ülkelerin birbirine yakın olduğu görülmektedir. 4.aşamada ise, 29. (Finlandiya) ve 44. (İngiltere) ülkelerin birbirine yakın olduğu görülmektedir. 10. aşamaya kadar ülkeler farklı kümelerde yer almaktadır. Ancak 10. aşamada 29. (Finlandiya), 37. (Hollanda) ülkelerin aynı kümede yer aldığı görülmektedir. Bu yüzden 29. Ülke (Finlandiya), 44. gözlem (İngiltere) ve 37. ülke (Hollanda) bir küme oluşturmaktadır. Bu şekilde kümeleme analiz algoritması 51. aşamaya kadar devam eder. kümeler oluşur.

Kümeleri belirlemek için elde edilen Dendrogram grafiği Şekil 1’de yandaki gibidir. Grafiğin yatay ekseninde Ülkeler, dikey ekseninde ise ülkelerin birbirlerine olan uzaklıkları ve oluşturdukları kümeler arasındaki bağlantılar görülmektedir. Grafiğe bakıldığında 3 birim uzaklıkta 9 küme ve 5 birim uzaklıkta 4 kümenin olduğu görülmektedir. En uygun küme sayısının ise 4 küme ile 5 birim uzaklıkta gerçekleştiği görülür.

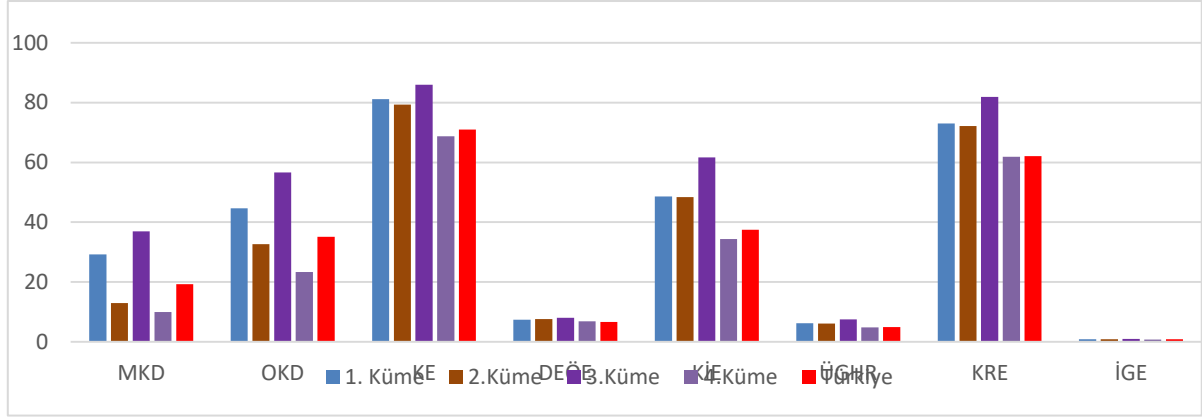
Çizelge 3. Ülkeleri gösteren küme grupları

1.Küme	2.Küme	3.Küme	4.Küme
Norveç	Litvanya	Finlandiya	Endonezya
Kanada	Letonya	İngiltere	Brezilya
İspanya	Estonya	Hollanda	Hindistan
Slovenya	Portekiz	ABD	Moğolistan
İtalya	Polonya	İsveç	Ermenistan
Avusturya	Birleşik Arap Em.	Danimarka	Arjantin
Fransa	Avustralya	İsviçre	Mısır
Belçika	Yeni_Zellanda	G_Kore	Azerbaycan
Slovakya	Katar	Japonya	Kazakistan
Macaristan	Suudi Arabistan	Almanya	Bahreyn
Tayland	İrlanda	Singapur	Kolombiya
Meksika			Sili
İsrail			Yunanistan
Çin			Türkiye
			Rusya
			Güney Afrika



Şekil 1. Ward yöntemi ile dendrogram

Şekil 2’de Ülkelerin değişkenler doğrultusunda elde edilen dağılımlar gösterilmektedir.

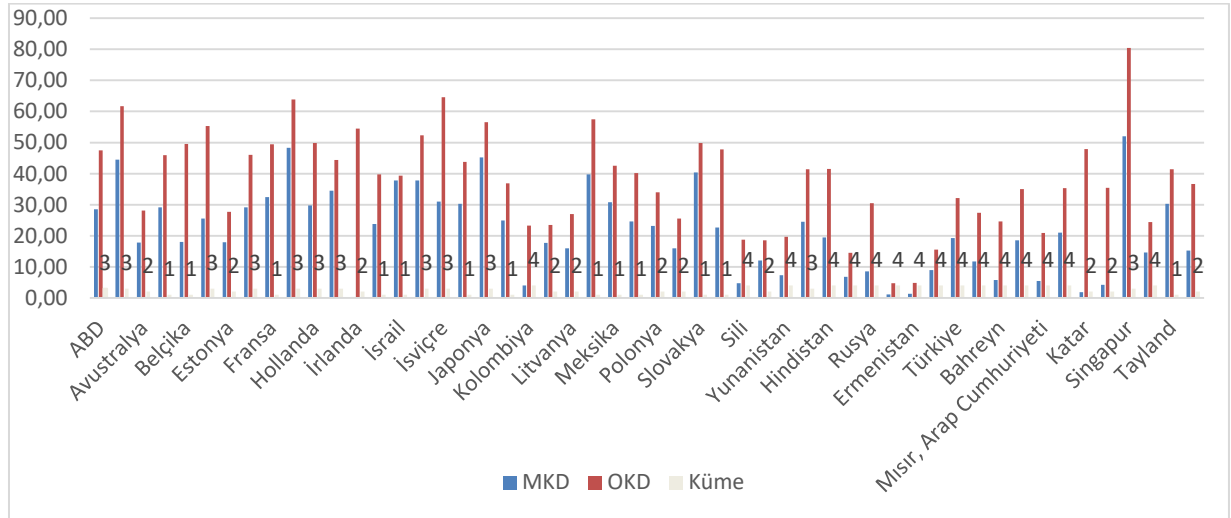


Şekil 2 Kümelerin ve Türkiye'nin değişkenler bazında dağılımı

Grafikte, çalışmamız içerisinde bulunan 4 küme içerisinde 1. Küme diğer küme grupları içerisinde 3. kümeden daha düşük 2. ve 4. Kümeden daha yüksek ortalamaya sahiptir. 2. küme 1. ve 3. Kümeden daha düşük 4. Kümeden daha yüksek ortalamaya sahiptir. 3. Küme tüm kümeler içerisinde değişken ortalamaları ile en yüksek kümedir. 4. küme ise değişken değerleri en düşük küme ortalamasına sahiptir.

3.Kümenin Endüstri 4.0 noktasında bu değişkenlere göre gelişmişliklerinin iyi olduğu en düşük değerlere ise 4. Kümenin sahip olduğu insani gelişim endeksi ve makine endeksi ve Küreselleşme Endeksi (KE) ile hem insan hem ürün anlamında ortalamasının düşük olması en düşük kümenin 4. Küme olduğunu gözlemlemekteyiz.

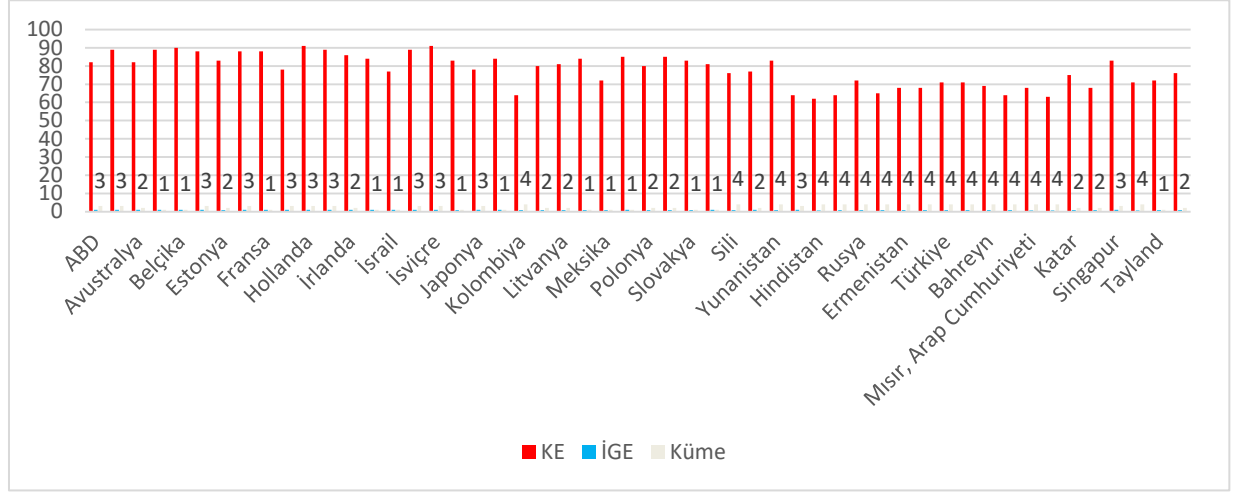
Türkiye kümeler içerisinde yer alan en düşük küme olan 4. Küme içinde yer almaktadır. Diğer küme gruplarıyla karşılaştırıldığında diğer kümelere yakın olduğu, ancak ortalamasıyla başka küme içerisinde giremediği Dünya Ekonomik Özgürlükler Endeksi (DEÖE) değişkeni dışında kalan tüm değişkenlerde kendi kümesinden daha yüksek ortalamaya sahip olduğu görülmektedir.



Şekil 3. Makine kapsamında; makine ve ulaşım ekipmanları, orta ve yüksek teknoloji endüstrisi değişkenlerinin ülkeler bazlı küme tablosu.

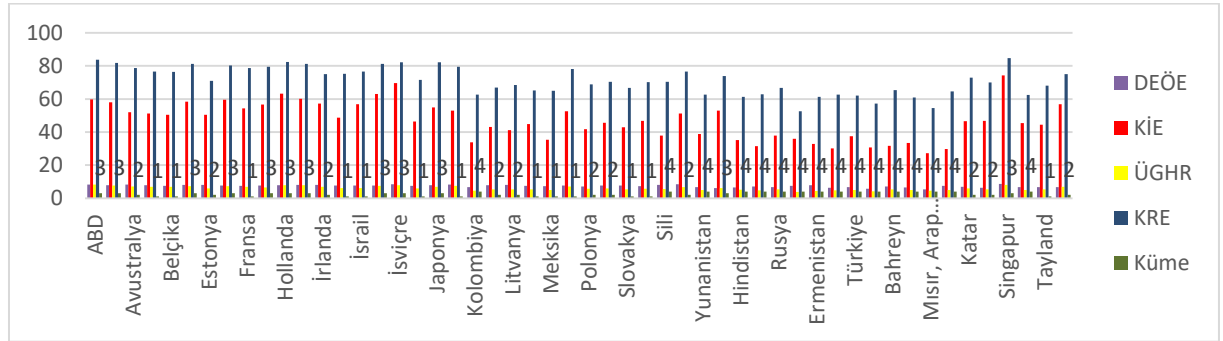
Değişken ortalamalarını tek tek incelenirse, Şekil 3 te Makine Ve Ulaşım Ekipmanları (MKD) değişkeninde en iyi performansı Singapur (Küme3) sergilerken, en kötü performansı İrlanda (Küme 2) göstermiştir. Makine Ve Ulaşım Ekipmanları (MKD) değişkeninde Küme 3, Küme 1'den daha iyi bir ortalamaya sahipken, Türkiye'nin de içerisinde bulunduğu 4. Küme en düşük ortalamaya sahiptir. Bu değişkende Türkiye kendi küme ortalamasından daha yüksek bir değere sahiptir. Şekil 3' te görüldüğü

üzere Orta Ve Yüksek Teknoloji Endüstrisi (OKD) değişkeninde en iyi performansı Singapur (Küme3) sergilerken, en kötü performansı Moğolistan (Küme 4) göstermiştir. Orta Ve Yüksek Teknoloji Endüstrisi(OKD) değişkeninde Küme 3, Küme 1’den daha iyi bir ortalamaya sahiptir. Bu değişkende Türkiye kendi küme ortalamasından daha yüksek bir değere sahiptir.



Şekil 4. İnsan kapsamında; insani gelişme endeksi, küreselleşme endeksi, değişkenlerinin ülkeler bazlı küme tablosu

Şekil 4’te, Küreselleşme Endeksi (KE) değişkeninde en iyi performansı İsviçre ve Hollanda (Küme3) sergilerken, en kötü performansı Hindistan (Küme 4) göstermiştir. Küreselleşme Endeksi (KE) değişkeninde kümeler birbirine oldukça yakın değerdedirler. Bu değişkende Türkiye kendi küme ortalamasından daha yüksek bir değere sahiptir. Şekil 4’te, İnsani Gelişme Endeksi (İGE) değişkeninde en iyi performansı Norveç (Küme 1) sergilerken, en kötü performansı Hindistan (Küme 4) göstermiştir. İnsani Gelişme Endeksi (İGE) değişkeninde Küme 3, birbirlerine yakın değerlere sahip Küme 1 ve Küme 2’den daha iyi bir ortalamaya sahipken, Türkiye’nin de içerisinde bulunduğu 4. Küme en düşük ortalamaya sahiptir. Bu değişkende Türkiye kendi küme ortalamasından daha yüksek bir değere sahiptir.



Şekil 5. Ürün kapsamında; Dünya Ekonomik Özgürlükler Endeksi, Küresel İnovasyon Endeksi, Üretimin Geleceğine Hazırlık Raporu, Küresel Rekabet Endeksi, değişkenlerinin ülkeler bazlı küme tablosu.

Şekil 5’te; Dünya Ekonomik Özgürlükler Endeksi (DEÖE) değişkeninde en iyi performansı Singapur (Küme3) sergilerken, en kötü performansı Mısır (Küme 4) göstermiştir. Dünya Ekonomik Özgürlükler Endeksi (DEÖE) değişkeninde Küme 3, Küme 2’den daha iyi bir ortalamaya sahipken, Türkiye’nin de içerisinde bulunduğu 4. Küme en düşük ortalamaya sahiptir. Bu değişkende Türkiye kendi küme ortalamasından daha yüksek bir değere sahiptir. Küresel İnovasyon Endeksi (KİE) değişkeninde Şekil 5’te en iyi performansı Singapur (Küme3) sergilerken, en kötü performansı Mısır (Küme 4) göstermiştir. Küresel İnovasyon Endeksi (KİE)değişkeninde Küme 3, birbirlerine yakın değerlere sahip Küme 1 ve

Küme 2’den daha iyi bir ortalamaya sahipken, Türkiye’nin de içerisinde bulunduğu 4. Küme en düşük ortalamaya sahiptir. Bu değişkende Türkiye kendi küme ortalamasından daha yüksek bir değere sahiptir. Şekil 5’te, Küresel Rekabet Endeksi (KRE) değişkeninde en iyi performansı Singapur (Küme 3) sergilerken, en kötü performansı Moğolistan (Küme 4) göstermiştir. Küresel Rekabet Endeksi (KRE) değişkeninde Küme 3, birbirlerine yakın değerlere sahip Küme 1 ve Küme 2’’den daha iyi bir ortalamaya sahipken, Türkiye’nin de içerisinde bulunduğu 4. Küme en düşük ortalamaya sahiptir. Bu değişkende Türkiye kendi küme ortalamasından daha yüksek bir değere sahiptir. Üretimin Geleceğine Hazırlık Raporu (ÜGHR) değişkeninde Şekil 5’te en iyi performansı ABD (Küme 3) sergilerken, en kötü performansı Moğolistan (Küme 4) göstermiştir. Üretimin Geleceğine Hazırlık Raporu (ÜGHR) değişkeninde Küme 3, birbirlerine yakın değerlere sahip Küme 1 ve Küme 2’den daha iyi bir ortalamaya sahipken, Türkiye’nin de içerisinde bulunduğu 4. Küme en düşük ortalamaya sahiptir. Bu değişkende Türkiye kendi küme ortalamasından daha yüksek bir değere sahiptir.

SONUÇ

Bu çalışmada; Bu çalışmada; Endüstri 4.0 etkilerinin insan, makine ve ürün bağlamındaki etkilerini Şangay, OECD ve bazı ülkelerin toplamını oluşturan 52 ülkenin yer aldığı, Makine ve Ulaşım Ekipmanları, Orta ve Yüksek Teknoloji Endüstrisi, Küreselleşme Endeksi, Dünya Ekonomik Özgürlükler Endeksi, Küresel İnovasyon Endeksi, Üretimin Geleceğine Hazırlık Raporu, Küresel Rekabet Endeksi, İnsani Gelişme Endekslerinin yer aldığı değişkenlerle Hiyerarşik Kümeleme Ward Yöntemi ile ülkeler 4 kümeye ayrılmış olup bu kümeler Endüstri 4.0’da ülkelerin gelişmişliklerinin dağılımlarını yansıtmaktadır.

Küme 1’de, Norveç, Kanada, İspanya, Slovenya, İtalya, Avusturya, Fransa, Belçika, Slovakya, Macaristan, Tayland, Meksika, İsrail, Çin ile birlikte 14 ülke yer almaktadır. Küme 1’in bu değişkenlerde aldığı ortalama değerleri ile Küme 3’e oldukça yakın, Küme 2 ve 4’ e uzaktır. Küme içerisinde bulunan her bir değişkende farklı farklı ülkeler sıralamanın önünde veya arkasında olmaktadır. Küme 1’in ortalama değeri en yüksek olan Küme 3’ten sonra en iyi ortalamaya sahip olduğu görülmektedir.

Küme 2’de 11 ülke: Litvanya, Letonya, Estonya, Portekiz, Polonya, Birleşik Arap Emirlikleri, Avustralya, Yeni Zelanda, Katar, Suudi Arabistan, İrlanda yer almaktadır. Küme 4’ten çok daha iyi bir ortalamaya sahipken Küme 1 ile oldukça yakın ortalamalara sahip olduğu görülmektedir.

Küme 3 ’de 11 ülke: Finlandiya, İngiltere, Hollanda, ABD, İsveç, Danimarka, İsviçre, G_Kore, Japonya, Almanya, Singapur yer almaktadır. Kümeler içerisinde en iyi ortalamaya sahip kümedir. Tüm değişkenler içerisinde diğer kümelerin üzerinde ortalamaya sahip olup en yakın olduğu küme grubu isel. Kümedir.

Küme 4’te 16 ülke: Endonezya, Brezilya, Hindistan, Moğolistan, Ermenistan, Arjantin, Mısır, Azerbaycan, Kazakistan, Bahreyn, Kolombiya, Sili, Yunanistan, Türkiye, Rusya, Güney Afrika yer almaktadır. Değişken ortalamaları içerisinde en düşük ortalamaya sahip küme grubudur. İçerisinde yer alan Türkiye’nin “Dünya Ekonomik Özgürlükler Endeksi” dışındaki tüm değişken ortalama değerleri daha yüksektir.

Küme 4’te yer alan Türkiye Makine ve Ulaşım Ekipmanları, Orta ve Yüksek Teknoloji Endüstrisi değişken ortalama değerleri ile hem içerisinde bulunduğu küme 4 hemde Küme 2’den daha iyi bir ortalamaya sahiptir. Küresel Rekabet Endeksi, Küresel İnovasyon Endeksi, Üretimin Geleceğine Hazırlık Raporu, Küresel Rekabet Endeksi, İnsani Gelişme Endeks değişkenleriyle çok daha iyi ortalamaya sahip olduğu görülmektedir.

İnsan aklını modelleyerek yine insan hayatını kolaylaştırma yönünde hızlı bir değişimin oluşumundaki başlıca aktör olan Endüstri 4.0 insan hayatı ile ilgili kolay pratik bir yaşam yaratma hedefi ile ülkelerin rekabet alanı genişlemektedir.

Endüstri 4.0'ın etkileri sadece üretimden kaynaklı ülke ekonomilerini etkilemekle kalmayıp o ülke vatandaşının hayatının bir çok alanında yer alarak en nihayetinde insan hayatını yakından etkilemektedir. Hızla gelişen teknoloji sebebiyle insan hayatı bir taraftan kolaylaşırken başka taraflardan zorlaştırmakta olduğu varsayılsada Endüstri 4.0 etkilerinin olumlu yönde fazla olan ülkelerdeki insani gelişim endeksinin de yüksek olduğu bu çalışmada da gözlemlenmektedir.

Endüstri 4.0 ile birlikte, iş gücünün değişim yaşaması, mavi yakalıların yerini robot, beyaz yakalıların yapacağı iş yerine ise yapay zekanın geleceği, bu sebeple teknoloji okur yazarlığın gelişmesi gerektiği, endüstri 4.0 ile yapay zeka kavramlarının etkileri sonucunda akıllı fabrikaların getirdikleri yenilik ve ucuzluk ile ürün çeşitliliği ve kolay ulaşılabilirliği sonucunu, makinelerin insan hayatını kolaylaştırma yolculuğu olarak varsayabiliriz. Tüm ülkelerin Endüstri 4.0 'ın etkileri nedeniyle gelişen bu süreci en iyi şekilde analiz ederek yeniliklerin getirdiği avantajlardan en üst düzeyde faydalanabileceklerin öncü ülkeler olabileceği düşünülmektedir.

Çalışmada incelenen Makine ve Ulaşım Ekipmanları, Orta ve Yüksek Teknoloji Endüstrisi, Küreselleşme Endeksi, Dünya Ekonomik Özgürlükler Endeksi, Küresel İnovasyon Endeksi, Üretimin Geleceğine Hazırlık Raporu, Küresel Rekabet Endeksi, ortalama değeri yüksek olan küme gruplarının İnsani Gelişim Endeksi ortalama değeri de yüksektir. Ülkelerin endüstri 4.0 ın getirdiği teknolojileri yakalama yönündeki başarısı ülke insanının hayatının refahını da yakından etkilemektedir. Bu en yüksek değişken ortalamaya sahip 3. Küme'nin İnsani gelişim endeksinin 0,93 olması ve Endüstri 4.0 noktasında en gelişmiş küme olması ile görülmektedir.

Kaynaklar

- Alpar, R. (2013). Çok Değişkenli İstatistiksel Yöntemler. Ankara: Detay Yayıncılık.
- Barutcu, H. C. (2019). Endüstri 4.0 Uygulamalarının Üretim Süreçlerine Etkisi: Bosch Sanayi Ve Ticaret Anonim Şirketi Örneği. Yüksek Lisans Tezi. İstanbul.
- Duman, M. Ç. (2019). Endüstri 4.0 Teknoloji Bileşenlerinin Örgütsel Performansa Etkilerini Belirlemeye Yönelik Bir Araştırma. Malatya.
- Gürler, C. (2020). Endüstri 4.0 Ve Bir Kümeleme Analizi Uygulaması. Doktora Tezi. İstanbul.
- Çokluk, O., BÜYÜKÖZTÜRK, Ş., & ŞEKERCİOĞLU, G. (2018). Sosyal Bilimler İçin Çok Değişkenli İstatistik: SPSS ve LISREL Uygulamaları. Ankara: Pegem Akademi.
- Deniz, S. S. (2020). Kayıp Veri İçeren Veri Setlerinde Kümeleme Uygulamaları. Doktora Tezi. Van.
- Kabaca, D. Y. (2019). Endüstri 4.0 Devrimi Ve Uygulamaları İle Etkilerinin İncelenmesi. Yüksek Lisans. İstanbul.
- Kalaycı, Ş. (2014). Spss Uygulamalı Çok Değişkenli İstatistik Teknikleri. Ankara: Asil Yayın Dağıtım Ltd. Şti.
- Mercan, T., & Atalay, M. (2021). Türkiye'deki Havalimanlarının Performanslarının Kümeleme ve Topsis Yöntemleriyle Değerlendirilmesi. Adıyaman Üniversitesi Sosyal Bilimler Enstitüsü Dergisi, 666-711.
- Özdamar, K. (2018). Paket Programlar İle İstatistiksel Veri Analizi. Eskişehir: nisan.
- Ünlü, F., & Atik, H. (2018). Türkiye'deki İşletmelerin Endüstri 4.0'a Geçiş Performansı: Avrupa Birliği Ülkeleri İle Karşılaştırmalı Ampirik Analiz. Ankara Avrupa Çalışmaları Dergisi, 431-463.
- Şahin, C. (2019). Ülkelerin endüstri 4.0 düzeylerinin copras yöntemi ile analizi: g-20 ülkeleri ve türkiye. Yüksek lisans tezi. Bartın.
- Taş, H. Y., & Küçükoglu, M. (2019). Endüstri 4.0 (Dördüncü Sanayi Devrimi) Sürecinde Yeni Nesil Girişimler Hakkında Bir İnceleme. 4.Uluslararası Gap Sosyal Bilimler Kongresi, (S. 205-212). Şanlıurfa.
- Yalçın, N. (2013, OCAK). Kümeleme Analizi Ve Uygulaması. Yüksek Lisans Tezi. Elazığ.

II. International Applied Statistics Conference (UYIK-2021)
Tokat / Turkey, 29 June - 2 July 2021

- Yaz, H. F. (2014, Ocak). Çok Değişkenli İstatistiksel Tekniklerden Kümeleme Analizi; Spss İle Bir Uygulama.
- URL-1, <https://data.worldbank.org/indicator/NV.MNF.TECH.ZS.UN>. Erişim tarihi: (2021,03.18).
- URL-2, <https://data.worldbank.org/indicator/NV.MNF.TECH.ZS.UN>. Erişim tarihi : 15.03.2021.
- URL-3, <https://kof.ethz.ch/en/forecasts-and-indicators/indicators/kof-globalisation-index.html>. Erişim tarihi: 15.03.2021.
- URL-4, <https://www.fraserinstitute.org/economic-freedom/dataset?geozone=world&year=2018&page=dataset&filter=0&date-type=single>. Erişim tarihi: 15.03.2021.
- URL-5, https://www.globalinnovationindex.org/userfiles/file/reportpdf/gii_2018-report-new.pdf. Erişim tarihi: 15.03.2021.
- URL-6, <https://www.weforum.org/reports/how-to-end-a-decade-of-lost-productivity-growth>. Erişim tarihi: 10.04.2021.
- URL7, https://www.tr.undp.org/content/turkey/tr/home/library/human_development/hdr2019.html. Erişim tarihi: 15.04.2021.
- URL-8, <https://www.weforum.org/reports/readiness-for-the-future-of-production-report-2018>. Erişim tarihi: 08.04.2021.

On the Existence and Uniqueness of Maximum likelihood Estimates for Unit Distributions

Yunus AKDOĞAN¹, Tenzile ERBAYRAM^{1*}

¹Selçuk University, Faculty of Science, Statistics, 42130, Konya, Turkey

*Presenter author e-mail: tenzile.erbayram@selcuk.edu.tr

Abstract

The unit distributions have been studied for many fields recently. However, most importantly, the existence and uniqueness of estimates is a big problem in these studies based on parameter estimations. It is not always possible to show that the estimator exists and that it is unique. In the maximum likelihood and other estimator finding methods where iterative methods are used, the initial value problem is encountered and the obtained estimate may not be a global estimate. In this study, a very simple but unique estimator method that both reaches the global maximum value and is applied for unit distributions. In addition, two real data applications were made to show that the method works correctly.

Key words: Unit Weibull distribution, Unit Burr-XII distribution, Graphical methods, Cauchy–Schwarz inequality, Data analysis

Yapısal Eşitlik Modellemesi ve Bir Uygulama

Yuke BAI^{1*}, İlkay ALTINDAĞ², Aşır GENÇ¹

¹Necmettin Erbakan Üniversitesi, Fen Fakültesi, İstatistik Bölümü, 42370, Konya, Türkiye

²Necmettin Erbakan Üniversitesi, Uygulamalı Bilimler Fakültesi, Finans ve Bankacılık Bölümü, 42370, Konya, Türkiye

*Sunucu yazar e-posta: yukebai94@163.com

Özet

Müşteri sadakatini geliştirmek, ticari operasyonların en önemli hedeflerinden biridir. Literatür taramasından, aralarında daha yaygın olan sınıflandırmanın tutum sadakati ve davranış sadakatine ayrılabilceği birçok sadakat sınıflandırma yöntemi olduğu bulunmuştur. Geleneksel sadakat araştırması, tutum sadakati ile davranışsal sadakat arasında ayırım yapmaz. Bu nedenle, bu çalışma, farklılıkları keşfetmek için iki sadakat modelini ayrı ayrı karşılaştırmak için doğrusal bir yapı modeli kullanır. Bu çalışmada, kurumsal imaj, değiştirme maliyetleri ve memnuniyetin tutumsal sadakat modeli ve davranışsal sadakat modeli üzerindeki etkisinin farklı olup olmadığını araştırmak için 318 mağaza müşterisinden elde edilen anket verileri kullanılmıştır. Karşılaştırma yaparak, iki sadakat modelinin oldukça iyi olduğunu görebiliriz, bu nedenle her iki model de kullanılabilir.

Anahtar kelimeler: Doğrulamalı faktör analizi, Yapısal eşitlik modeli, Tutumsal sadakat modeli, Davranışsal sadakat modeli

GİRİŞ

Yapısal Eşitlik Modellemesi (YEM), ikinci nesil veri analiz tekniği olarak (Bagozzi ve Fornell, 1982), regresyon gibi birinci nesil istatistiksel tekniklere kıyasla, birçok bağımlı ve bağımsız değişkenler arasındaki ilişkilerin modellenmesi ile karmaşık bir araştırma problemini tek bir süreçte, sistematik ve kapsamlı bir şekilde ele almayı sağlamaktadır (Anderson ve Gerbing, 1988). Özellikle karmaşık modellerin testinde başarılı olduğu, birçok analizi bir defada yaptığı, incelenen modeldeki ilişkiler ağına yönelik varsa yeni düzenlemeler tavsiye ettiği, aracılık ve düzenleyicilik (moderasyon) etkilerini incelemeyi kolaylaştırdığı, ölçüm hatalarını hesaba katıyor olması gibi nedenlerle yapısal eşitlik modellemesi yöntemi, birçok teoremin test edilmesinde ve yeni modellerin geliştirilmesi sürecinde kullanılmakta olan bir yöntemdir. Regresyon analizi ise temel olarak bağımlı değişkendeki değişimin ne kadarlık kısmının bağımsız değişkenler tarafından açıklandığını ortaya koyan birinci nesil veri analiz tekniklerindendir. Bağımlı değişken ile bağımsız değişkenler arasındaki doğrudan ilişkilerin yanı sıra dolaylı ilişkilerin varlığının söz konusu olduğu çok basamaklı bir modelde, regresyon analiziyle doğrudan etkiler tespit edilebilirken, değişkenlerin dolaylı etkileri göz ardı edilmektedir. Dolayısıyla doğrusal regresyon gibi geleneksel yöntemlerde bağımlı ve bağımsız değişkenler arasındaki bağlantıların sadece tek bir düzeyde ele alınması, yapısal eşitlik modellemesinde ise, her bir ilişki düzeyinin eş anlamlı olarak değerlendirilmesi yöntemleri arasındaki farklılıklardan yalnızca birisidir.

YEM yöntemiyle analiz edilen bir model, geleneksel regresyon analizi yöntemleriyle de yapılabilir de, regresyon analizlerinde her bir ilişki için bir regresyon analizine gerek duyulurken, Lisrel, AMOS vb. programlarla gerçekleştirilen analizlerde, değişkenler arasında belirlenen tüm ilişkiler tek bir analizle ortaya konmakta, ayrıca ek olarak path analizinde ölçmeden kaynaklanan hata miktarı elimine edilebilmektedir. Hatanın devre dışı bırakılması, yapısal eşitlik modellemesine dayalı olan tüm analiz yöntemlerinin en önemli avantajlarından birisidir (Tatlıdil, 1992'den aktaran, Yener, 2007).

Bu çalışmada, yapısal eşitlik modeli incelemek amaçlanmıştır. Ajzen (1985, 1991) tarafından geliştirilmiş olan planlanmış davranış teorisi modeliyle kadınların satın alma davranışı için geliştirilmiş üç model, her iki yöntemle analiz edilmiştir.

Çok değişkenli model analizlerinde, amaç bir yandan bağımsız değişkenlerin her birindeki bir birimlik değişimin bağımlı değişkende ne kadarlık bir değişime yol açacağını görmek iken, diğer yandan söz konusu bağımsız değişkenlerin bağımlı değişkendeki değişimin ne kadarını açıklıyor olduğunu tespit edebilmektedir. Bu çalışma ile çok değişkenli bir modelde araştırmacının taşıdığı bu iki amaç açısından karşılaştırılan yöntemler nasıl bir performans sergiliyor sorusuna yanıt aranmıştır.

Yapısal Eşitlik Modeli

Yapısal Eşitlik Modeli, aynı zamanda kovaryans yapı modeli olarak da bilinen Yönelimli bir çok değişkenli analiz aracıdır. YEM, özellikleri ve özellikleri analiz etmek için karakteristik değişkenlerin kovaryans matrisini temel alır. Sosyal bilimler, ekonomi ve finans, psikoloji ve yönetim alanındaki birçok araştırmada, birçok durumda öğrenme motivasyonu, kullanıcı memnuniyeti vb. gibi doğrudan gözlemlenemeyen örtük değişkenler bulunmaktadır ve geleneksel istatistiksel yöntemler bu sorunları iyi çözmemektedir. YEM 1980'lerde olgunlaşmış ve geleneksel hesaplama yöntemlerinin eksikliklerini telafi edebilmiştir. YEM aynı anda birden fazla bağımlı değişkeni, yani endojen değişkenleri işleyebilir. Geleneksel regresyon modellerinde regresyon katsayıları ve path analizindeki path katsayıları her bağımlı değişken için tek tek hesaplanır ve diğer faktörler göz ardı edilir. Yapısal Eşitlikte, diğer faktörlerin varlığı tamamen dikkate alınacak, yani her bir faktörün içindeki yapı, birlikte var olan diğer değişkenler dikkate alınarak ayarlanacak ve değiştirilecek, böylece sadece faktörler arasındaki ilişki değişmeyecek. Fakat faktörlerin iç yapısı da değişiklikleri değiştirecektir.

YEM gözlenen değişkenler (observed variable) ve örtük değişkenler (latent variable) arasındaki nedensel ilişkilerin ve korelasyon ilişkilerinin bir arada bulunduğu modellerin test edilmesi için kullanılan istatistiksel bir teknik olup bağımlılık ilişkilerini tahmin etmek için, varyans, kovaryans analizleri, faktör analizi ve çoklu regresyon gibi analizlerin birleşmesiyle meydana gelen çok değişkenli bir yöntemdir. Yapısal eşitlik modellemesi özellikle psikoloji, pazarlama vb. bilimlerde değişkenler arasındaki ilişkilerin değerlendirilmesinde ve modellerin testinde kullanılmaktadır (Tüfekçi ve Tüfekçi, 2006).

Örtük değişkenler bu yöntemin en önemli kavramlarından biridir. Pazarlamada, müşteri memnuniyeti, kalite algılanışı, tutumlar gibi kavramlar örtük değişkenlere örnek olarak verilebilir. Bu değişkenler gözlenmediği için doğrudan ölçülemezler. Bu yüzden, araştırmacı, örtük değişkeni işlemsel olarak tanımlamak için, örtük değişkeni gözlenebilir değişkenlerle ilişkilendirir (Yılmaz, 2004a).

Yöntemin temel özelliği, tamamen teoriye dayalı olmasıdır ve örtük değişkenler seti arasında bir nedensellik yapısının var olduğunu kabul etmesidir (Yılmaz, 2004b). Yapısal eşitlik modellemesinin uygulanmasındaki en kritik konu, oluşturulan modelin oldukça sağlam bir teorik alt yapıya sahip olmasıdır

Verilerin modeli destekleyip desteklemediğini değerlendirmek amacıyla yapısal eşitlik modellemesi literatüründe kullanılan en yaygın yöntem, iki aşamalı yöntemdir (Anderson ve Gerbing, 1988). Analizlerde birinci aşama olarak önce ölçme modeli test edilerek (Huchting vd., 2008) modelde yer alan yapılara ait ölçümlerin ilgili yapıları doğru ölçüp ölçmediğine bakılır, ikinci aşamada ise yapısal modeller incelenir.

i. Ölçüm Modeli: Gizil değişkenler ile gözlenen değişkenler arasındaki ilişkiyi gösterir.

ii. Yapısal Model: Gizil yapılar veya endojen (içsel) değişkenler arasındaki ilişkileri gösterir.

Araştırmacının elinde doğru bir ölçüm yoksa, yapıları ölçtüğünü varsaydığı ifadeler söz konusu yapıyı yeterince ölçmüyorsa, yapısal modeli analiz etmenin bir anlamı olmayacaktır.

YEM analizi sonuçlarını değerlendirmek için birçok uyum iyiliği test istatistiği vardır. Literatürde sıkça kullanılan test istatistikleri; Ki-kare istatistiği, Yaklaşık Hataların Ortalama Karekökü (RMSEA), Uyum İyilik İndeksi (GFI), Düzeltilmiş Uyum İyilik İndeksi (AGFI), Normlandırılmış Uyum İndeksi (NFI), Standartlaştırılmış Ortalama Hataların Karekökü (SRMR) ve Karşılaştırmalı Uyum İndeksi (CFI) olarak ele alınmaktadır. Uyum iyiliği değerleri ve uyum referans aralıkları Tablo 1'de verilmiştir (Schermelleh-Engel, Moosbrugger ve Müller, 2003; Çelik ve Yılmaz, 2016).

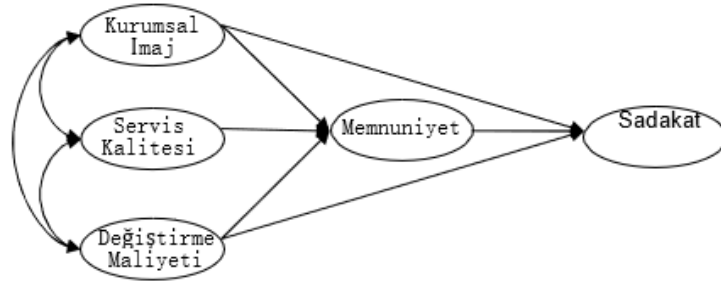
Tablo 1. Uyum İndeksleri Referans Aralığı

Uyum İndeksleri	İyi Uyum	Kabul Edilebilir Uyum
χ^2/sd	$\chi^2/sd \leq 2sd$	$\chi^2/sd \leq 3sd$
RMSEA	$RMSEA \leq 0.05$	$RMSEA \leq 0.08$
SRMR	$SRMR \leq 0.05$	$SRMR \leq 0.10$
NFI	$0.95 \leq NFI$	$0.90 < NFI$
CFI	$0.97 \leq CFI$	$0.95 \leq CFI$
GFI	$0.95 \leq GFI$	$0.90 \leq GFI$
AGFI	$0.90 \leq AGFI$	$0.85 \leq AGFI$

sd: Serbestlik Derecesi

Araştırmanın Amacı ve Araştırma Modelleri

Bu çalışmada, veri analizi ve model doğrulama için Amos 22.0'ı kullanılarak YEM analizi gerçekleştirilmiştir. Literatüre uygun olarak, iki rekabetçi model oluşturulmuş, en uygun modeli bulmak için modeller karşılaştırılmıştır. Şekil 1 bu çalışmanın kavramsal modelidir. Sadakat analizi modeli 1: Tutumsal sadakat Modeli ve modeli 2: Davranışsal sadakat modelinin iki tane modeli.



Şekil 1. Çalışmada ele alınan kavramsal model

Araştırmada test edilecek hipotezler:

Hipotez 1: Tutumsal sadakat modelinin beklenen kovaryans matrisi ile örnek kovaryans matrisi arasında fark yoktur.

Hipotez 2: Davranışsal bağlılık modelinin beklenen kovaryans matrisi ile örnek kovaryans matrisi arasında fark yoktur.

Hipotez 3: Memnuniyetin tutum sadakati ve davranışsal sadakat üzerindeki etkisinde fark yoktur.

Hipotez 4: Kurumsal imajın tutumsal sadakati ve davranışsal sadakat üzerindeki etkisinde bir fark yoktur.

Hipotez 5: Değişim maliyetlerinin tutum bağlılığı ve davranışsal bağlılık üzerindeki etkisinde hiçbir fark yoktur.

YÖNTEM

Araştırma modeli

Literatürü okuyarak sadakat ile ilgili iki yapısal eşitlik model elde edilmiş, iki farklı modeli farklı sadakat modeli ile karşılaştırmak için yapısal eşitlik modeli kullanılmıştır. Bu çalışmada kullanılan boyutlar 5'li Likert ölçeği ile değerlendirilmiştir. Anket tasarımı için referans materyalleri aracılığıyla, sadakati etkileyebilecek faktörleri belirlendikten sonra çeşitli etkili faktörler için anketler tasarlanmıştır. Anketteki sorular, kurumsal imaj, servis kalitesi, somuluk, garanti, güvence, değiştirme maliyetinin altı yönünü birleştirir ve ankete katılan kişinin bir memnuniyet puanı vermesine olanak tanır. Skor beş seviye içerir: (1: kesinlikle katılmıyorum, 2: katılmıyorum, 3: emin değilim, 4: katılıyorum 5: kesinlikle katılıyorum)

Veri toplama

Anket yöntemi: Çin'deki AVM'lerinde, kullanmak istediğimiz verileri toplamak için alışverişe gelen müşterilere anket uygulanacaktır. Bu anket anketinde, 318'i geçerli anket anketi sonucu olmak üzere 388 katılımcı yer almıştır.

SPSS Programında

Bu makale, her bir gizli değişkenin deneysel verilerinin iç tutarlılık gereksinimlerini karşılayıp karşılamadığını belirlemek için her boyutun güvenirlik katsayısını, CITC değerini ve madde silindikten sonra güvenirlik katsayısını hesaplamak için SPSS.21'i kullanır.

AMOS Programında

Çalışmanın özünü oluşturacak bazı yapısal eşitlik modelleri kurulacak ve hesaplanmalar yapılacaktır.

Analizlere alınan örnek grubunun büyüklüğü 318'dur. Analizlerde yapısal eşitlik modellemesi kullanıldığından ve bu yöntem için örnek hacminin olması gereken büyüklüğü çeşitli araştırmacılar tarafından en az 200 olmak kaydıyla 200-500 aralığında ifade edildiğinden (Kline, 1998; Möbius, 2003), söz konusu örnek hacmi yeterli bulunmuştur. YEM analizleri AMOS programı, regresyon analizleri ise SPSS programı kullanılarak yapılmıştır.

BULGULAR

Katılımcıların Demografik Özellikleri

Ele alınan 318 örneklem arasından, cinsiyet açısından bakıldığında 184 erkek ve 134 kadın vardır. Oranladığımızda ise sırasıyla % 57,9'u erkek, %42,1'i kadındır. Erkek örneklem kadın örnekleminden biraz fazladır. Eğitim düzeyi açısından ise %22,1 oranındaki 70 kişi ortaokul mezunu, %23,9 orandaki 76 kişi lise mezunu, %28,3 oranındaki 90 kişi lisans mezunu, %25,8 oranındaki 82 kişi yüksek lisans mezunudur. Akademik niteliklerin örneklem seçimi nispeten dengelidir, meslekler açısından, %14,2 oranında 45 memur, %18,2 oranında 58 beyaz yakalı işçi, %11,46 oranında 35 öğrenci, %14,5 oranında 46 çiftçi, %9,4 oranında 3 askeri personel, 30'unu, %18,2'sini oluşturuyor. Bu oran %9,4, 12 emekli, %3,8 ve 92 kişi, 28,9, aylık gelir açısından 48 kişi 1.000-1.500 yuan örneklemde, %15,1, 1.500-3.000 yuan hesap Örneklemde 101 kişi var, 31,8, 57 kişi 3000-5000, 17,9, 44 kişi 5000-7000 ve 68 13,8 veya 7000'den fazla olan kişiler, 21,4'ü oluşturuyor. Yukarıdaki sonuçlar görülebilir, bu çalışmanın orijinalliğini sağlamak için seçilen örnekler daha kapsamlıdır.

Tablo 2. Çalışmaya katılan kişilere ilişkin demografik bilgiler

Temel Bilgiler	Sınıflandırma	Sıklık	Yüzde
Cinsiyet	Erkek	184	57.9
	Kadın	134	42.1
Eğitim Seviyesi	Orta okul	70	22
	Lise	76	23.9
	Lisans	90	28.3
	Yüksek Lisans	82	25.8
	Memurluk	45	14.2
	Beyaz yakalı işçiler	58	18.2
	Öğrencilik	35	11
Meslek	Çiftçilik	46	14.5
	Askerlik	30	9.4
	Emekli	12	3.8
	Diğerler	92	28.9
	1000-1500	48	15.1
Aylık Geliri	1500- 3000	101	31.8
	3000-5000	57	17.9
	5000-7000	44	13.8
	7000 +	68	21.4

Güvenirlilik Analizi

Güvenirlilik analizi, ölçek tarafından toplanan sonuçların tutarlılık düzeyinin bir ölçüsüdür. Bu çalışmanın güvenirlik analizinde ağırlıklı olarak Cronbach's Alpha katsayısı kullanılmaktadır. İlgili literatüre göre, Cronbach's Alpha değeri genellikle 0.7'den büyük olduğunda ölçek gereksinimlerini karşılar. İç tutarlılık nispeten yüksektir ve daha fazla analiz edilebilir. 0.7'den küçükse, anketin ayarlanması veya örneklem büyüklüğünün artırılması gerekir. Bu çalışmada, SPSS 21 programı ile elde edilen Cronbach's Alpha değerleri aşağıdaki tabloda verilmiştir.

Tablo 3. Ölçeklere ilişkin güvenirlik testi sonuçları

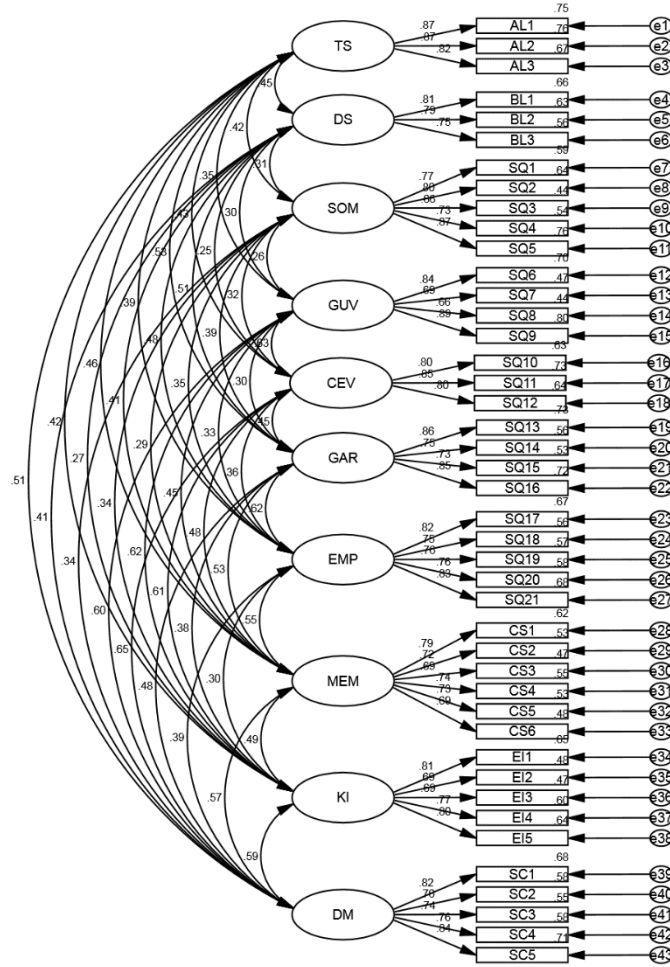
Değişken	Sorun Kodu	Değişken Sayısı	Cronbach's Alpha
Tutumsal Sadakat	AL1-AL3	3	0.883
Davranışsal Sadakat	BL1-BL3	3	0.82
Somutluk	SQ1-SQ5	5	0.876
Güvence	SQ6-SQ9	4	0.874
Cevaplanabilirlik	SQ10-SQ12	3	0.859
Garanti	SQ13-SQ16	4	0.869
Empati	SQ17-SQ21	5	0.891
Memnuniyet	CS1-CS6	6	0.882
Kurumsal İmaj	EI1-EI5	5	0.868
Değiştirme Maliyeti	SC1-SC5	5	0.880

Tablo 3'ten bu çalışmada kullanılan anketin iyi bir güvenirliğe sahip olduğu söylenebilir.

Doğrulayıcı faktör analizi

Yapısal eşitlik modelinde, ölçüm modelinin verilerle uyumunu tespit etmek için doğrulayıcı faktör analizi (DFA) yöntemi kullanılmaktadır. Bu bağlamda çalışmada yer alan tutum ifadelerinin yapısal

geçerliliğini test etmek, maddelerin temsil gücünü ve boyutlar arasındaki ilişkiyi belirlemek için AMOS 22.0 paket programı ile doğrulayıcı faktör analizi uygulanmıştır. DFA sonucunda elde edilen ölçüm modeline ilişkin path diyagramı Şekil 2’de gösterilmiştir. uyum iyiliği indeksleri ise Tablo 4’te gösterilmiştir. Kolaylık sağlamak için aşağıdaki kelimeler için kısaltmalar kullanıyoruz. : Tutumsal Sadakat : TS ; Davranışsal Sadakat : DS ; Somutluk : SOM ; Güvence : GUV; Cevaplanabilirlik : CEV ; Garanti : GAR ; Empati : EMP ; Memnuniyet : MEM ; Kurumsal İmaj : KI ; Değiştirme Maliyeti : DM ; Servis Kalitesi:SK.



Şekil 2. Birinci düzey doğrulayıcı faktör analizi modeli diyagramı

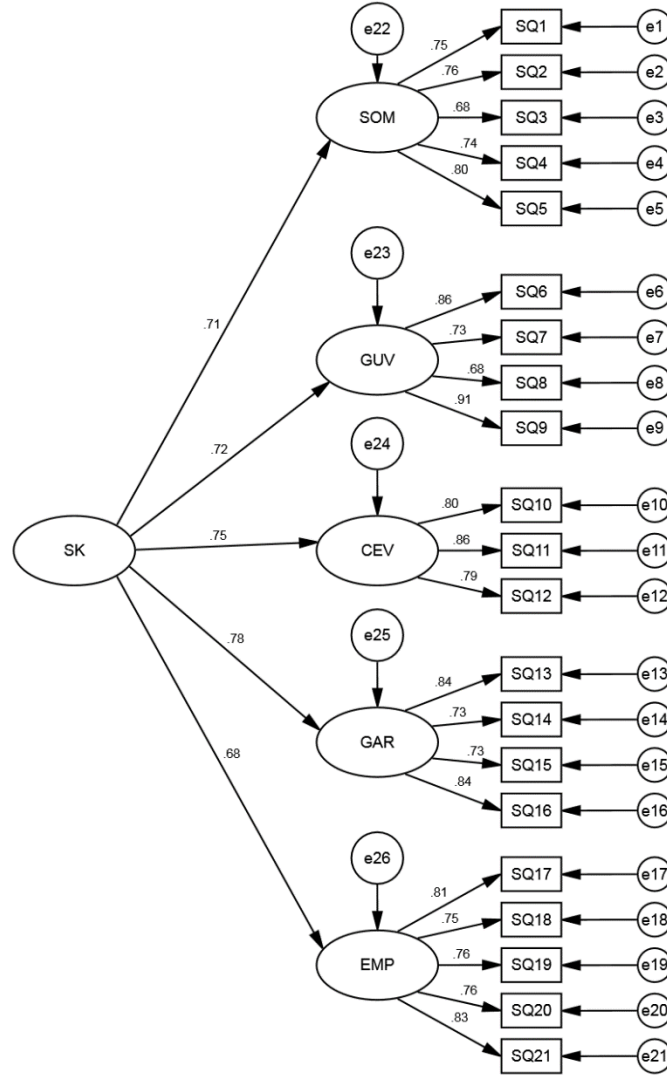
Elde edilen DFA sonucun uyum iyiliği indeksleri Tablo 4’te gösterilmektedir.

Tablo 4. Birinci düzey doğrulayıcı faktör analizi modeli uyum indeksi sonuçları

Uyum İndeksleri	χ^2/sd	GFI	AGFI	TLI	IFI	CFI	RMSEA	SRMR
	1.669	0.836	0.810	0.926	0.934	0.933	0.046	0.0445

Yukarıdaki tabloya göre, χ^2/sd değeri 1.669'dur, bu 3'ten küçüktür; RMSEA, 0.046'dır, bu da 0.08'lik standart seviyeden daha düşüktür ve iyi bir uyumu gösterir. GFI=0.836, AGFI=0.810, TLI=0.926, IFI=0.934, CFI=0.933, tüm uyum iyiliği göstergeleri genel standardı karşılamaktadır, bu da bu araştırmada kurulan birinci derece doğrulayıcı faktör analizi modelinin etkili ve tutarlı olduğunu göstermektedir.

Birinci Düzey DFA sonuçları uygun bulunduktan sonra ilişkin ikinci düzey DFA analizi yapılmıştır. AMOS programından elde edilen path diyagramı Şekil 3'te gösterilmektedir.



Şekil 3. İkinci düzey doğrulayıcı faktör analizi modeli diyagramı

İkinci düzey DFA analizinden elde edilen sonuçlar Tablo 5'te gösterilmektedir.

II. International Applied Statistics Conference (UYİK-2021)
Tokat / Turkey, 29 June - 2 July 2021

Tablo 5. İkinci Düzey Doğrulayıcı faktör analizi sonuçları

Örtük Değişken	Gözlenen Değişken	Standart Yük	Standart Hata	t	P	R ²	CR	AVE
Servis Kalitesi	Somutluk	0.708				0.501		
	Güvence	0.72	0.123	8.507	***	0.518		
	Cevaplanabilirlik	0.753	0.123	8.399	***	0.567	0.851	0.533
	Garanti	0.783	0.158	8.753	***	0.613		
	Empati	0.682	0.153	8.068	***	0.465		
Tutumsal Sadakat	AL1	0.866				0.750		
	AL2	0.874	0.052	18.642	***	0.764		
	AL3	0.798	0.054	16.728	***	0.637	0.884	0.717
Davranışsal Sadakat	BL1	0.794				0.630		
	BL2	0.785	0.073	12.962	***	0.616		
	BL3	0.75	0.069	12.557	***	0.563	0.820	0.603
Somutluk	SQ1	0.759				0.576		
	SQ2	0.794	0.069	14.176	***	0.630		
	SQ3	0.674	0.072	11.865	***	0.454		
	SQ4	0.741	0.07	13.159	***	0.549		
	SQ5	0.858	0.072	15.32	***	0.736	0.877	0.589
Güvence	SQ6	0.864				0.746		
	SQ7	0.734	0.055	15.127	***	0.539		
	SQ8	0.69	0.059	13.855	***	0.476		
	SQ9	0.899	0.051	20.224	***	0.808	0.877	0.642
	SQ10	0.807				0.651		
Cevaplanabilirlik	SQ11	0.856	0.065	16.226	***	0.733		
	SQ12	0.792	0.064	14.982	***	0.627	0.859	0.670
	SQ13	0.852				0.726		
Garanti	SQ14	0.743	0.055	14.852	***	0.552		
	SQ15	0.74	0.052	14.779	***	0.548		
	SQ16	0.831	0.057	17.324	***	0.691	0.871	0.629
Empati	SQ17	0.814				0.663		
	SQ18	0.76	0.059	14.824	***	0.578		
	SQ19	0.767	0.061	15.005	***	0.588		
	SQ20	0.758	0.058	14.782	***	0.575		
	SQ21	0.843	0.059	16.98	***	0.711	0.892	0.623
Memnuniyet	CS1	0.796				0.634		
	CS2	0.739	0.061	13.893	***	0.546		
	CS3	0.71	0.061	13.241	***	0.504		
	CS4	0.762	0.061	14.417	***	0.581		
	CS5	0.744	0.065	14	***	0.554		
	CS6	0.719	0.059	13.448	***	0.517	0.882	0.556
Kurumsal İmaj	EI1	0.808				0.653		
	EI2	0.705	0.062	13.157	***	0.497		
	EI3	0.72	0.063	13.501	***	0.518		
	EI4	0.75	0.062	14.209	***	0.563		
	EI5	0.793	0.061	15.209	***	0.629	0.869	0.572
Değiştirme Maliyeti	SC1	0.832				0.692		
	SC2	0.746	0.063	14.756	***	0.557		
	SC3	0.723	0.063	14.163	***	0.523		
	SC4	0.731	0.061	14.363	***	0.534		
	SC5	0.828	0.059	17.045	***	0.686	0.881	0.598

Tablo 6. İkinci düzey doğrulayıcı faktör analizi uyum indeksi sonuçları

Uyum İndeksleri	χ^2/df	GFI	AGFI	TLI	IFI	CFI	RMSEA	SRMR
	2.380	0.885	0.855	0.925	0.935	0.935	0.066	0.0642

Yukarıdaki tabloya göre χ^2/df değeri 3'ten küçük olan 2.380'dir; RMSEA, 0.066'dır, bu da 0.08'lik standart seviyeden daha düşüktür ve iyi bir uyuma işaret eder. GFI=0.885, AGFI=0.855, TLI=0.925, IFI=0.935, CFI=0.935, SRMR=0.0642 tüm uyum iyiliği göstergelerinin genel standardı karşılaması, bu araştırmada kurulan ikinci dereceden doğrulayıcı faktör analizi modelinin etkili ve tutarlı olduğunu göstermektedir.

Yukarıdaki tablodan, ikinci derece değişken hizmet kalitesi altındaki beş boyutun path katsayılarının hepsinin 0,5'ten büyük olduğu görülebilir. Her bir maddenin her bir boyut altındaki standartlaştırılmış faktör yükü 0,5'ten büyüktür, bu da her maddenin kendi boyutunu iyi açıklayabildiğini göstermektedir.

Birleşik güvenilirlik (CR), her bir örtük değişkendeki tüm öğelerin örtük değişkeni tutarlı bir şekilde açıklayıp açıklamadığını yansıtan, modelin doğal kalitesi için kriterlerden biridir. Tablodan, birleşik güvenilirlik CR'sinin 0.7'den büyük olduğu görülebilir; bu sonuçlar her gizli değişkendeki tüm ölçüm öğelerinin gizli değişkeni tutarlı bir şekilde açıklayabildiğini gösterir.

Her boyutun yakınsak geçerliliği, ortalama varyans çıkarımı (AVE değeri) ile yansıtılır. Değer genellikle ölçeğin yakınsak geçerliliğini yansıtmak için kullanılır. Gizil değişken tarafından açıklanan varyansın ne kadarının ölçümden geldiğini doğrudan gösterebilir. AVE değeri ölçüm değişkeni ne kadar büyük olursa, gizli değişken tarafından açıklanan varyasyon yüzdesi o kadar büyük olursa, bağıl ölçüm hatası o kadar küçük olur. Genel olarak, değer gereksinimi 0,5'in üzerindedir. Yukarıdaki tablodan, AVE değerlerinin hepsinin standart 0,5 değerinin üzerinde olduğu görülebilir, bu sonuçlar çalışmanın ölçeğinin iyi yakınsak geçerliliğe sahip olduğunu göstermektedir.

Tablo 7. Ayırt edici geçerlilik sonuçları

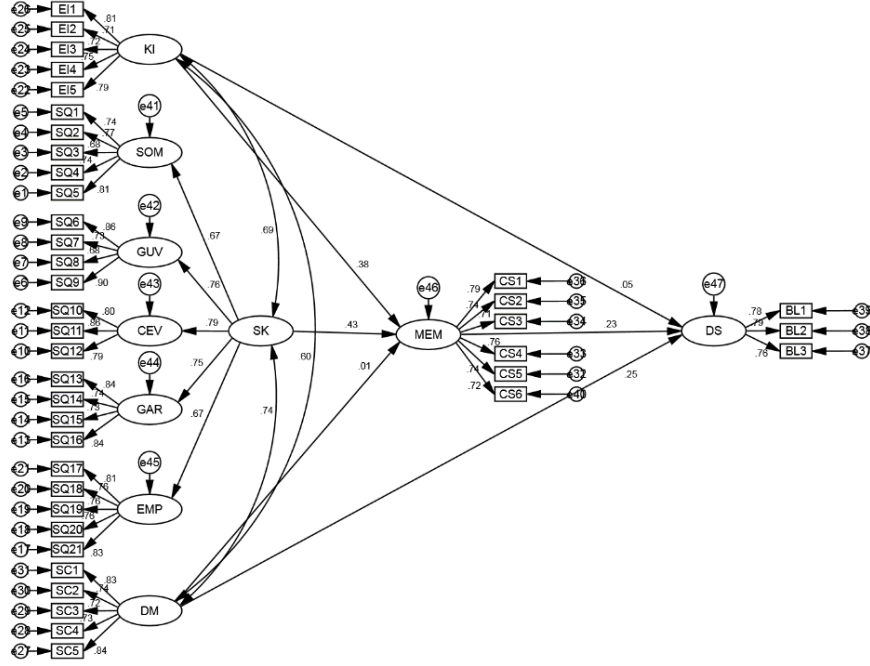
	Değişirme Maliyeti	Kurumsal İmaj	Memnuniyet	Empati	Garanti	Cevaplanabilirlik	Güvence	Somutluk	Davranışsal Sadakat	Tutumsal Sadakat
Değişirme Maliyeti	0.773									
Kurumsal İmaj	0.6	0.756								
Memnuniyet	0.559	0.68	0.745							
Empati	0.374	0.417	0.552	0.789						
Garanti	0.475	0.442	0.53	0.617	0.793					
Cevaplanabilirlik	0.661	0.551	0.469	0.342	0.486	0.818				
Güvence	0.603	0.464	0.452	0.31	0.297	0.625	0.801			
Somutluk	0.323	0.372	0.31	0.347	0.392	0.316	0.265	0.767		
Davranışsal Sadakat	0.4	0.356	0.4	0.477	0.495	0.221	0.293	0.304	0.789	
Tutumsal Sadakat	0.468	0.495	0.477	0.366	0.528	0.427	0.353	0.402	0.388	0.846
AVE	0.598	0.572	0.556	0.623	0.629	0.670	0.642	0.589	0.603	0.717

Not: Kalın yazı tipi AVE'nin aritmetik karekökünü göstermektedir.

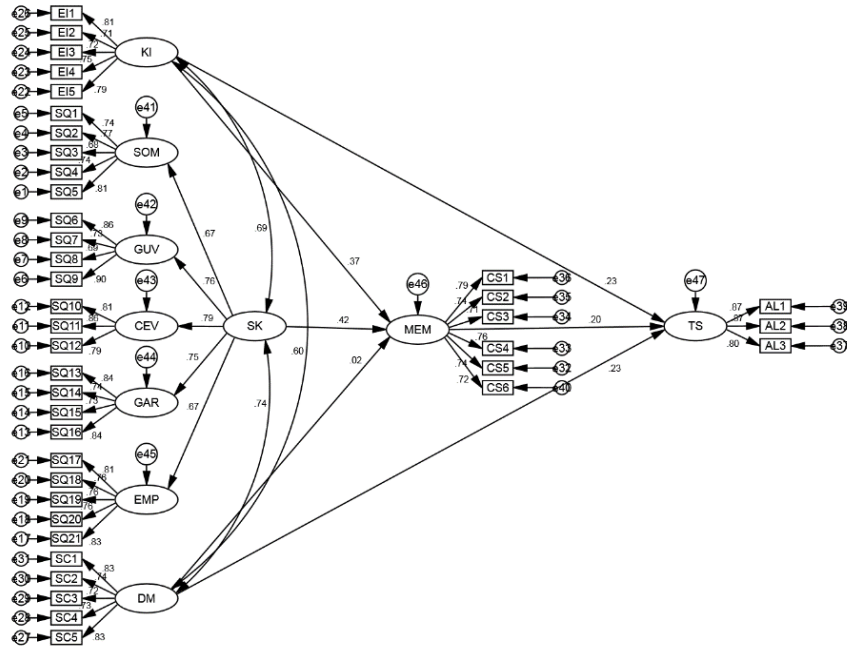
Yukarıdaki tablodan, her bir boyutun AVE'sinin 0,5'ten büyük olduğu ve AVE'nin karekökünün korelasyon katsayısından büyük olduğu görülebilir, bu da ölçeğin iyi bir yakınsama geçerliliğine ve ayırt edici geçerliliğe sahip olduğunu gösterir.

Önerilen Yapısal Eşitlik Modeli İçin Analiz Sonuçları

Araştırmada ele alından hipotezleri test etmek için YEM analizi gerçekleştirilmiştir. Davranışsal Sadakat modeli için yapısal model uygulaması sonucu elde path diyagramı Şekil 4'te, Tutumsal Sadakat Modelinin Yapısal Eşitlik Modeli sonuçları ise Şekil 5'te gösterilmektedir.



Şekil 4. Davranışsal Sadakat Modelinin Yapısal Eşitlik Modeli Path Diyagramı



Şekil 5. Tutumsal Sadakat Modelinin Yapısal Eşitlik Modeli Path Diyagramı

Tablo 8 incelendiğinde davranışsal sadakat modelinin path analizi sonuçlarından hizmet kalitesinin memnuniyete giden standart path katsayısı 0.431'dir (t değeri=3.979, $p=0.000<0.05$), bu da hizmet kalitesinin memnuniyet üzerinde anlamlı bir pozitif etkiye sahip olduğunu gösterir. Bu nedenle, hipotez geçerlidir; dönüşüm maliyetinin doyuma giden standartlaştırılmış path katsayısı 0,013'tür (t değeri = 0,152, $p = 0,879>0,05$), bu da dönüştürme maliyetinin memnuniyet üzerinde önemli bir etkisi olmadığını, dolayısıyla hipotezin geçerli olmadığını gösterir; kurumsal imajın memnuniyet üzerindeki etkisi standartlaştırılmış path katsayısı 0.376'dır (t değeri=4.774, $p=0.000<0.05$), kurumsal imajın memnuniyet üzerinde anlamlı bir pozitif etkisi olduğunu gösterir, bu nedenle hipotez geçerlidir; memnuniyetin R^2 değeri boyut 0,561 olup, yorumlama derecesinin %56,1'e ulaştığını göstermektedir.

Standartlaştırılmış yol katsayısının tutum bağlılığına dönüşüm maliyeti 0,235 (t değeri=3,126, $p=0,002<0,05$) olması, dönüşüm maliyetinin tutum bağlılığı üzerinde anlamlı bir pozitif etkiye sahip olduğunu gösterir, bu nedenle hipotez geçerlidir; kurumsal imajın standartlaştırılmış yolu Katsayı 0.226'dır (t değeri=2.525, $p=0.012<0.05$), bu da kurumsal imajın tutum bağlılığı üzerinde anlamlı bir pozitif etkiye sahip olduğunu gösterir, bu nedenle hipotez geçerlidir; standartlaştırılmış memnuniyetten tutum bağlılığına giden yol katsayısı 0.203 (t değeri=2.395, $P=0.017<0.05$), memnuniyetin tutum bağlılığı üzerinde anlamlı bir pozitif etkiye sahip olduğunu, dolayısıyla hipotezin geçerli olduğunu belirtirken; tutum bağlılığı boyutunun SMC değeri 0.330 olup, yorumlanma derecesine ulaştığını göstermektedir. %33.

Tablo 8. Davranışsal Sadakat Modelinin Path Analizi Sonuçları

Varsayımsal path		Standartlaştırılmış path katsayısı	Standart hata	T	P	R^2
Servis Kalitesi	→	Memnuniyet	0.431	0.159	3.979	***
Değiştirme Maliyeti	→	Memnuniyet	0.013	0.084	0.152	0.879
Kurumsal İmaj	→	Memnuniyet	0.376	0.082	4.774	***
Değiştirme Maliyeti	→	Davranışsal Sadakat	0.253	0.073	3.024	0.002
Kurumsal İmaj	→	Davranışsal Sadakat	0.051	0.089	0.522	0.602
Memnuniyet	→	Davranışsal Sadakat	0.232	0.082	2.458	0.014

Tablo 9 incelendiğinde tutum bağlılığı modelinin path analizi sonuçlarından hizmet kalitesinin memnuniyete giden standartlaştırılmış path katsayısının 0,427 (t değeri=3.899, $p=0.000<0.05$) olduğu, hizmet kalitesinin memnuniyet üzerinde pozitif etkisi vardır. Bu nedenle, hipotez geçerlidir; dönüşüm maliyetinden memnuniyete standartlaştırılmış path katsayısı 0.017'dir (t değeri=0.208, $p=0.835>0.05$), dönüşüm maliyetinin memnuniyet üzerinde anlamlı bir etkisi olmadığını gösterir, dolayısıyla hipotez geçerli değildir; kurumsal imajın memnuniyet üzerindeki etkisi Standartlaştırılmış path katsayısı 0.374'tür (t değeri=4.712, $p=0.000<0.05$), kurumsal imajın memnuniyet üzerinde anlamlı bir pozitif etkisi olduğunu gösterir, dolayısıyla hipotez geçerlidir; Memnuniyet boyutunun ÖBY değeri 0,561 olup yorumlama derecesinin% 56,1'e ulaştığını göstermektedir.

Tutum bağlılığına dönüşüm maliyetinin standartlaştırılmış path katsayısı 0.235'tir (t değeri=3.126, $p=0.002<0.05$), bu, dönüşüm maliyetinin tutum bağlılığı üzerinde anlamlı bir pozitif etkiye sahip olduğunu gösterir, bu nedenle hipotez geçerlidir; kurumsal imajın standartlaştırılmış pathu Katsayı

0.226'dır (t değeri=2.525, $p=0.012<0.05$), bu da kurumsal imajın tutum bağlılığı üzerinde anlamlı bir pozitif etkiye sahip olduğunu gösterir, bu nedenle hipotez geçerlidir; standartlaştırılmış memnuniyetten tutum bağlılığına giden path katsayısı 0.203 (t değeri=2.395, $P=0.017<0.05$), memnuniyetin tutum bağlılığı üzerinde anlamlı bir pozitif etkiye sahip olduğunu, dolayısıyla hipotezin geçerli olduğunu; tutum bağlılığı boyutunun R^2 değeri 0.330 olup, yorumlanma derecesine ulaştığını göstermektedir. %33.

Tablo 9. Tutum Bağlılık Modeli Path analizi sonuçları

Varsayımsal path			Standartlaştırılmış path katsayısı	Standart hata	T	P	R^2
Servis Kalitesi	→	Memnuniyet	0.427	0.160	3.899	***	0.561
Değiştirme Maliyeti	→	Memnuniyet	0.017	0.085	0.208	0.835	
Kurumsal İmaj	→	Memnuniyet	0.374	0.083	4.712	***	
Değiştirme Maliyeti	→	Tutumsal Sadakat	0.235	0.074	3.126	0.002	0.330
Kurumsal İmaj	→	Tutumsal Sadakat	0.226	0.091	2.525	0.012	
Memnuniyet	→	Tutumsal Sadakat	0.203	0.082	2.395	0.017	

Davranışsal bağlılık modeli ile tutum bağlılığı modeli arasındaki karşılaştırmaya göre, davranışsal bağlılık modelinde, değiştirme maliyetinin memnuniyet üzerinde anlamlı bir etkisi olmadığı ve kurumsal imajın davranışsal bağlılık üzerinde anlamlı bir etkisi olmadığı görülmektedir. Tutum bağlılığı modelinde, değiştirme maliyetinin de memnuniyet üzerinde önemli bir etkisi yoktur ve diğer değişkenlerin önemli bir etkisi vardır.

Test edilen modellerin uygunluğunun karşılaştırılması

Araştırmada test edilen Davranışsal bağlılık modeli ve tutum bağlılığı uyum indeksi sonuçları aşağıdaki Tablo 10'da gösterilmiştir. : iki modelin uyum indeksi genel standarda ulaşmıştır. Davranışsal bağlılık $\chi^2=1291.899$, $\chi^2/sd=1.779$, GFI=0.831, AGFI=0.809, RMSEA=0.050, TLI=0.918, IFI=0.924, CFI=0.924, ECVI=4.668, AIC=1479.899, BIC=1833.532; tutum Sadık $X^2=1308.637$, χ^2/sd , GFI=0.832, AGFI=0.810, RMSEA=0.050, TLI=0.918, IFI=0.925, CFI=0.924, ECVI=4.721, AIC=1496.637, BIC=1850.27. İki model arasında pek bir fark yok.

Tablo 10. Test edilen modellerin genel model uyumlarının karşılaştırılması

Adaptasyon indeksi	Genel standart	Davranışsal Sadakat Modeli	Tutum sadakat modeli
χ^2	Ne kadar küçükse o kadar iyi	1291.899	1308.637
χ^2/sd	<3	1.779	1.803
GFI	>0.8	0.831	0.832
AGFI	>0.8	0.809	0.810
RMSEA	<0.08	0.050	0.050
TLI	>0.9	0.918	0.918
IFI	>0.9	0.924	0.925
CFI	>0.9	0.924	0.924
ECVI	Ne kadar küçükse o kadar iyi	4.668	4.721
AIC	Ne kadar küçükse o kadar iyi	1479.899	1496.637
BIC	Ne kadar küçükse o kadar iyi	1833.532	1850.270

Tablo 10'a bakıldığında genel test edilen her iki modelin de uyum iyiliği indekslerinin iyi uyum sergilediği ve elde edilen uyum iyiliği değerlerinin birbirine oldukça yakın değerler aldığı görülmektedir.

SONUÇ

Araştırma hipotezinin sonucu

Bu çalışmanın hipotezi1 ve hipotezi 2 için yani tutumsal sadakat modeli ve davranışsal sadakat modeli oldukça iyi bir uyuma sahiptir. Davranışsal sadakat modelinde, değiştirme maliyetlerinin memnuniyet üzerinde önemli bir etkisi yoktur ve kurumsal imajın davranışsal sadakat üzerinde önemli bir etkisi yoktur. Tutum bağlılığı modelinde, değiştirme maliyetinin de memnuniyet üzerinde önemli bir etkisi yoktur ve diğer değişkenlerin önemli bir etkisi vardır.

Tablo 11. Test edilen modellerin genel model uyumlarının karşılaştırılması

	Hipotez	Sonuç
H1	Tutumsal sadakat modelinin beklenen kovaryans matrisi ile örnek kovaryans matrisi arasında fark yoktur.	Kabul ediyor
H2	Davranışsal bağlılık modelinin beklenen kovaryans matrisi ile örnek kovaryans matrisi arasında fark yoktur.	Kabul ediyor
H3	Memnuniyetin tutum sadakati ve davranışsal sadakat üzerindeki etkisinde fark yoktur.	Kabul ediyor
H4	Kurumsal imajın tutumsal sadakati ve davranışsal sadakat üzerindeki etkisinde bir fark yoktur.	Reddediyor
H5	Değişim maliyetlerinin tutum bağlılığı ve davranışsal bağlılık üzerindeki etkisinde hiçbir fark yoktur.	Kabul ediyor

İki sadakat modelin karşılaştırılması

İki sadakat modelinin uygunluk göstergelerinin sonuçlarından iki modelin uygunluk göstergelerinin genel standarda ulaştığı görülmektedir. Davranışsal bağlılık $X^2=1291.899$, $X^2/df=1.779$, $GFI=0.831$, $AGFI=0.809$, $RMSEA=0.050$, $TLI=0.918$, $IFI=0.924$, $CFI=0.924$, $ECVI=4.668$, $AIC=1479.899$, $BIC=1833.532$; tutum Sadık $X^2=1308.637$, $X^2/df=1.803$, $GFI=0.832$, $AGFI=0.810$, $RMSEA=0.050$, $TLI=0.918$, $IFI=0.925$, $CFI=0.924$, $ECVI=4.721$, $AIC=1496.637$, $BIC=1850.27$. İki model arasında pek bir fark yok. Bu çalışmanın sonuçları Bowen ve Chen'den (2001) farklıdır. Bowen ve Chen (2001) tutum bağlılığı ve davranış bağlılığının bileşik ölçümünün daha iyi tahmin gücü elde edeceğine inanmaktadır. Farklı sonuçların olası nedeni, Bowen ve Chen'in çalışmasında sadakatin alt boyutunun yüksek düzeyde ilişkili olması, bu çalışmada sadakatin alt boyutunun ise sadece orta-düşük korelasyon göstermesidir, yani sadece Tutum sadakati ve davranışsal sadakatin her ikisi de yüksek düzeyde ilişkiliyse, yani sonraki araştırmacılar sadakat modelini incelemek istiyorsa, bileşik ölçüm kullanıp kullanmama konusunda öncelikle alt boyutların korelasyon katsayılarını dikkate almalıdırlar.

Rekabetçi tutum sadakati ve davranış sadakati modelinin karşılaştırılmasından, tutum sadakati ve davranış sadakati etkisinde sadece kurumsal imajın anlamlı bir farklılığa sahip olduğu ve kurumsal imajın davranış sadakati üzerindeki etkisinin tutumdan daha fazla olduğu bulunmuştur. sadakat. Kurumsal imaj, tutum bağlılığı ve davranış bağlılığı üzerinde önemli bir etkiye sahiptir ve her ikisinin de etkisi olduğunu gösterir. Çalışmanın bu bölümünün sonuçları çoğu çalışma ile tutarlıdır (Jones ve diğerleri, 2000; Sharma ve Patterson, 2000).

Değiştirme maliyetinin tutum bağlılığı ve davranışsal bağlılık üzerindeki etkisi açısından, değiştirme maliyetinin tutum bağlılığı üzerinde anlamlı bir pozitif etkiye sahip olması, Josee ve Gaby'nin (2002) araştırma sonuçlarıyla tutarlı olarak, değiştirme maliyetinin sadakati etkileyeceğini göstermektedir.

Memnuniyetin tutum bağlılığı üzerindeki etkisi olumlu ve anlamlıdır. Geçmiş çalışmalarda sadakati ölçmek için sentetik bir yöntem kullanıldığından bu sonuç çoğu çalışma ile aynıdır. Ancak, bu çalışmada

memnuniyetin davranışsal sadakat üzerindeki etkisi anlamlı değildir. Bu araştırma sonucu Fornell ve diğerlerinin (1996) araştırma sonucundan farklıdır: müşteri memnuniyeti yüksektir ve patronaj olasılığı daha yüksek olacaktır. farklı.

Tutumsal Sadakat Modeli ve Davranışsal Sadakat Modelinde Değişkenlerin Yönetimdeki Önemi

Müşteri sadakatini geliştirmek, ticari operasyonların en önemli hedeflerinden biridir. Müşteri sadakati kabaca tutum sadakati ve davranış sadakati olarak ikiye ayrılabilir. Tutum sadakati modelinin analizinde memnuniyet, kurum imajı ve dönüşüm maliyetinin tutum bağlılığı üzerinde anlamlı ve olumlu etkileri olduğu tespit edilmiştir. Tüketim yapmak Bir kişi başkalarına tavsiye etmeye istekliyse, müşteri memnuniyetini artırmak, kurumsal imajı artırmak için daha fazla etkinlik yapmak ve müşteri değiştirme maliyetlerini artırmak mümkündür. Davranış sadakati modelinin analiz sonuçlarında, kurumsal imajın davranış sadakati üzerindeki etkisinin önemli olmadığı, memnuniyetin davranış sadakati üzerindeki etkisinin ise önemli bir marjın kenarında olduğu bulunmuştur. davranış sadakati anlamlıdır ve pozitif ilişkilidir. Bu nedenle, tüketicilerin yeniden satın alma niyetlerini teşvik etmek için müşterinin değiştirme maliyetini artırmak en büyük önceliktir. Örneğin, sayaç personeli müşterilerle ilişkiler kurmak, özelleştirilmiş hizmet öğeleri sağlamak ve eski müşterilere daha fazlasını sağlamak için daha fazla zaman harcamak zorundadır. değer indirimleri veya indirimleri, diğer büyük mağazalardan farklı hizmetler sağlama vb., tümü müşteri değiştirme maliyetlerini artırmaya yardımcı olur. Böyle bir yaklaşım, tüketicilerin yeniden satın alma istekliliği olasılığını artırabilir. Kurumsal imajı artırmanın satın alma duyarlılığını teşvik etme üzerinde çok az etkisi olduğu açıktır. Başka bir deyişle, müşterilerin gerçek satın alma davranışlarını veya niyetlerini yalnızca kurumsal imajın iyileştirilmesi yoluyla göstermelerini sağlamak çok zor olacaktır. çok katlı mağaza işletmecilerinin ortak markalı kartlar, fiyat indirimleri, cömert hediyeler gibi promosyon faaliyetleri ile cirolarını artırmaları hala kaçınılmazdır. Ayrıca bu araştırma sayesinde, dönüşüm maliyeti düşük olduğunda gerçek müşteri sadakatinin abartılmayacağını anlayabiliriz. Bir mağaza pazar bölümlendirmesi yaparken, eğer değiştirme maliyetleri bir bölümleme değişkeni olarak kullanılabilirse, o zaman şirket, farklı memnuniyet seviyeleri ve müşteriler için farklı değiştirme maliyetlerinin kombinasyonuna dayalı olarak şirketin çabasının oranını belirleyebilir. müşteri gruplarını yüksek memnuniyet ancak düşük değiştirme maliyetleri ile birleştirirseniz, bu müşteri grubu yüksek derecede sadakate sahip olacak ve şirketin gelirin katkıda bulunacaktır.

Kaynaklar

- Anderson, E., Fornell, C., Lehmann, D. R. (1994), "Customer satisfaction, market share and profitability: findings from Sweden," *Journal of Marketing*, 58(3), 53-66.
- Auh, S., Johnson, M.D. (2005). Compatibility effects in evaluations of satisfaction and loyalty. *Journal of Economic Psychology* 26(1), 35-57.
- Baumgartner and C., Homburg (1996), "Applications of Structural Equation Modeling in Marketing and Consumer Research: a review." *International Journal of Research in Marketing* 13 (2), 139-161.
- Bolton, R. N., Drew, J. H. (1991), A multistage model of customers' assessments of service quality and value," *Journal of Consumer Research*, 17(4), 375-384.
- Boomsma, A. (2000), "Reporting analyses of covariance structures." *Structural Equation Modeling*, 7, 461-483.
- Bowen, J. T., Chen, S. L. (2001), "The relationship between customer loyalty and customer satisfaction," *International Journal of Contemporary Hospitality Management*, 13(5), 213-217.
- Byrne, B. B. (2010), *Structural equation modeling using AMOS. Basic concepts, applications, and programming*, 2nd Ed, New York: Routledge. 23
- Chang, C. H. and Tu, C. Y. (2005), "Exploring store image, customer satisfaction and customer loyalty relationship: evidence from Taiwanese hypermarket industry," *Journal of American Academy of Business*, 7, 197-202.

- Cheung, G. W., and Rensvold, R. B. (2002), "Evaluating goodness-of-fit indexes for testing measurement invariance," *Structural Equation Modeling*, 9(2), 233-255.
- Chin, W. W. (1998), "Issues and Opinion on Structural Equation Modeling," *MIS Quarterly* 22(1), vii-xvi.
- Chumpitaz R., N. G. Paparoidamis. (2004). Service quality and marketing performance in business-to-business markets: exploring the mediating role of client satisfaction. *Managing Service Quality* 14(2/3), 235-248
- Churchill, G. A., Surprenant, C. (1982), "An investigation into the determinants of customer satisfaction," *Journal of Marketing Research*, 19(4), 491-504.
- Cronin, J. J., Brady, M. K., Hult, T. M. (2000), "Assessing the effects of quality, value, and customer satisfaction on consumer behavioral intentions in service environments," *Journal of Retailing*, 76(2), 193-218.
- Cronin, J.J. and Taylor, S.A. (1992). Measuring Service Quality: Reexamination and Extension, *Journal of Marketing*, 56, 55-68.
- Nguyen N., G. LeBlanc (1998) "The mediating role of corporate image on customers' retention decisions: an investigation in financial services", *International Journal of Bank Marketing*, 16(2), 52 – 65.
- Nguyen, N. and Leblanc G. (2001) "Corporate image and corporate reputation in customers' retention decisions in services, " *Journal of Retailing and Consumer Services*, 8(4), 227-236,.
- Oliver, R. L. (1999), "Whence customer loyalty?" *Journal of Marketing*, 63(special issue), 33-44.
- Parasuraman, V.A., Zeithaml, V.A. and Berry, L.L., (1988). SERVQUAL: a multiple-item scale for measuring consumer for perceptions of service quality, *Journal of Retailing*, 64, 12–40.
- Robertson, T.S. and Gatignon, H. (1986), "Competitive effects on technology diffusion," *Jouranal of Marketing*, 50, 1-12.
- Schreiber, J. B., Nora, A., Stage, F. K., Barlow, E. A., and King, J. (2006), "Reporting structural equation modeling and confirmatory factor analysis results: Areview," *The Journal of Educational Research*, 99, 323–337.
- Schreiber, J.B.(2008). Core reporting practices in structural equation modeling. 26 *Administrative Pharmacy*, 4, 83-97.
- Sharma, N. and P. G. Patterson. (2000). "Switching Costs, Alternative Attractiveness and Experience as Moderators of Relationship Commitment in Professional, Consumer Services." *International Journal of Service Industry Management* 11(5), 470-490.
- Torkzadeh, Koufteros, and Pflughoeft (2003), "Confirmatory analysis of computerself-efficacy, " *Structural Equation Modeling*, 10(2): 263-275.
- Tse D. K. and Peter C. Wilton (1988). Models of Consumer Satisfaction Formation: An Extension. *Journal of Marketing Research*, 25(2), 204-212.
- Wong A., A. Sohal. (2003) Service quality and customer loyalty perspectives on two levels of retail relationships. *Journal of Services Marketing* 17(5), 495-513.
- Zeithaml, V. A., Berry, L. L., and Parasuraman, A. (1996), "The behavioral consequences of service quality, " *Journal of Marketing*, 60(2), 31-46.
- Dick, A. S., Basu, K. (1994), "Customer loyalty: toward an integrated conceptual framework," *Journal of the Academy of Marketing Science*, 22(2), 99-113.
- Doll, W. J., Xia, W., and Torkzadeh, G. (1994), "A confirmatory factor analysis of the end-user computing satisfaction instrument," *MIS Quarterly*, 18 (4), 357–369.
- Duncan, O. D. (1975), *Introduction to structural equation models*, New York: Academic Press.
- Fornell, C. (1992), "A national customer satisfaction barometer: the Swedish experience," *Journal of Marketing*, 56(1), 6-21.
- Fornell, C., and Larcker, D. F. (1981), "Evaluating structural equation models with unobservable variables and measurement error," *Journal of Marketing Research*, 18 , 39-50.
- Fornell, C., Johnson, M. D., Anderson, E. W., Cha, J., Bryant, B. E. (1996), "The American customer satisfaction index: nature, purpose, and findings," *Journal of Marketing*, 60(4), 7-18.
- Hellier, Phillip K., Geursen, Gus M., Carr, Rodney A., and Rickard, John A. (2003). "Customer repurchase intention: A general structural equation model". *European Journal of Marketing*, 37(11/12), 1762.

- Host, V., Knie-Andersen, M. (2004), "Modeling customer satisfaction in mortgage credit companies", The International Journal of Bank Marketing, 22 (1), 26-42.
- Hoyle, R. H., and Panter, A. T. (1995), "Writing about structural Equation models". In R. H. Hoyle (Ed.), Structural equation modeling: Concepts, issues, and applications, Thousand Oaks, CA: Sage, 158-176.
- Jones, T. O., Sasser, W. E. (1995), "Why satisfied customers defect," Harvard Business Review, 73(6), 88-99.
- Kumar V., D. Shah, and R. Venkatesan (2006), "Managing Retailer Profitability — One Customer at a Time!" Journal of Retailing, 82 (4), 277–294.
- Lam, S. Y., Shankar, V., Erramilli, M. K., Murthy, B. (2004), "Customer value, satisfaction, loyalty, and switching costs: an illustration from a business-to-business service context," Journal of the Academy of Marketing Science, 32(3), 293-311.
- MacCallum, R.C., Browne, M.W., and Sugawara, H.M. (1996), "Power analysis and determination of sample size for covariance structure modeling," Psychological Methods 1(2), 130-149.
- McDonald, R.P., and Ho, M. H. R. (2002), "Principles and practice in reporting structural equation analyses," Psychological Methods, 7, 64–82.
- Çelik, H. E., & Yılmaz, V. (2016). LISREL 9.1 ile yapısal eşitlik modellemesi: Temel kavramlar-uygulamalar-programlama (3. Baskı). Ankara: Anı Yayıncılık.
- Schermelleh-Engel, K., Moosbrugger, H., Müller, H. (2003). Evaluating the Fit of Structural Equation Models: Tests of Significance and Descriptive Goodness-of-Fit Measures. Methods of Psychological Research Online, 8(2): 23-74.

Anket

Cinsiyet:

Erkek ☐ kadın ☐

Eğitim Seviyesi:

Ortaokul ☐ lise ☐ lisans ☐ lisansüstü ☐

Meslek:

Memur ☐ Beyaz yakalı işçi ☐ Öğrenci ☐ Çiftçi ☐ Asker ☐ Emekli ☐ Diğer ()

Aylık Gelir:

1000-1500 ☐ 1500- 3000 ☐ 3000-5000 ☐ 5000-7000 ☐ 7000+ ☐

1’den 5’e sıralayacak şekilde tercihlerinizi seçin:

(1: kesinlikle katılmıyorum, 2: katılmıyorum, 3: emin değilim, 4: katılıyorum 5: kesinlikle katılıyorum)

1. AVM’nin avantajlarından başkalarına da bahsedeceğim.

1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐

2. Biri benden tavsiye etmemi istediğinde, AVM’yi tavsiye edeceğim.

1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐

3. Bu AVM’yi aktif olarak arkadaşlarıma ve aileme tavsiye edeceğim.

1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐

4. Gelecekte alışveriş yaparken, AVM’yi alışveriş için ilk tercihim olarak kullanmaya hazırım.

1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐

5. Gelecekte, bu AVM’de ihtiyacım olan mal veya hizmetleri almaya hazırım.

1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐

6. İleride AVM ile ticaret yapmaya devam edeceğim.

1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐

7. AVM'nin tabelaları açıkça belirtilmiştir ve tanımlanması kolaydır.

1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐

8. AVM'nin müziği, dekorasyonu ve atmosferi rahat.

1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐

9. AVM'nin konumu park etmek kolay veya bir park yeri var.

1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐

10. AVM'nin ürün sınıflandırma ekranı ve mağaza akışı iyi planlanmış.

1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐

11. AVM'nin servis personeli özenle giyinmiş ve bakımlıdır.

1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐

12. Bu AVM'da satılan ürünler kaliteli.

1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐

13. AVM'de söz verilen hizmetleri (garanti süresi, satış sonrası bakım vb.) sağlar.

1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐

14. AVM'den satın alınan ve uygun olmayan ürünler hızlı bir şekilde değiştirilebilir.

1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐

15. AVM tarafından sağlanan ürün ve hizmetler tutarlıdır.

1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐

16. AVM'nin servis personeli müşterilere yardımcı olmaktan mutluluk duyar.

1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐

17. AVM'nin servis personeli, müşteri taleplerine hızlı bir şekilde cevap verecektir.

1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐

18. AVM'nin servis personeli, müşteri sorunlarını çözme yeteneğine sahiptir.

1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐

19. AVM, gelecekte sunacağı hizmetleri bilgilendirmek için inisiyatif alacaktır.

1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐

20. AVM servis personeli tarafından bana verilen bilgilerin yeterli olduğuna inanıyorum.

1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐

21. AVM'nin çalışanlara uygun eğitim ve öğretim verdiğini düşünüyorum.

1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐

Investigating the Effect of Land Use and Hillslope Characteristics on Morphological Changes of Gullies in the West of Ardabil Province

Ali Reza VAEZI¹, Saeid SAFARI^{1*}, Saniye DEMIR²

¹Department of Soil Science and Engineering, Faculty of Agriculture, University of Zanjan

²TOGU University, Faculty of Agriculture Department of Soil Science, Tokat, 60100

*Presenter author e-mail: s1366saeed@gmail.com

Abstract

Gully erosion is one of the most important types of soil erosion and land degradation, especially in drylands in semi-arid regions. The purpose of this study was to investigate the effect of land use and hillslope characteristics on morphological changes of permanent gullies in the west of Ardabil province, north west Iran. First, twenty seven hillslopes affected by gully erosion were identified in an area of 250 km² and twenty gullies were considered for field measurements. Gully length extension was determined along with morphological characteristics of hillslopes including surface area, slope gradient, shape and surface area of rangeland and rainfed lands in each hillslopes. Based on the results, the gully length in the area ranges from 80.20 m to 830 m. There is a positive correlation between the gully length and the surface area of rainfed ($r = 0.40$ and $p < 0.05$), whereas gully length is negative correlated with rangeland surface area ($r = -0.38$ and $p < 0.05$). The change of rangelands to rainfed uses and overgrazing in rangelands can extend gully formation in the area.

Key words: Conservation tillage, Hillslopes length, land use, soil loss

Aşırı Yayılımlı Sayım Verilerinin Analizi

Gazel SER^{1*}

¹Van Yüzüncü Yıl Üniversitesi, Ziraat Fakültesi, Zootehni Bölümü, 65040, Van, Türkiye

*Sunucu yazar e-posta: gazelser@gmail.com

Özet

Bu çalışma, sayıma dayalı olarak elde edilen verilerde sıkça karşılaşılan aşırı yayılım sorununun incelenmesi amacıyla gerçekleştirilmiştir. Sayım verilerinin varsayımsal olarak poisson dağılımı gösterdiği kabul edilir. Bu dağılımın önemli özelliği ortalamasının varyansının eşit olmasıdır. Ancak, sayım verilerinde bu eşitliği sağlamak oldukça güçtür. Çünkü, sayıma dayalı bağımlı değişkenin varyansı, sayım aralığının genişlemesine bağlı olarak artma eğilimi gösterir. Değişkenliğin artmasına bağlı olarak, varyansın ortalamadan büyük olması durumu ortaya çıkar ve aşırı yayılım sorunuyla karşılaşılır. Aşırı yayılım, varsayılan modelin yanlış olduğunun bir göstergesidir ve modifikasyonlar gereklidir. Bu nedenle çalışmada, sayımla elde edilen uygulama verisinde genelleştirilmiş doğrusal karışık model yaklaşımı kullanılarak poisson ve alternatif olarak negatif binomial dağılımlardan elde edilen sonuçlar karşılaştırılmış ve aşırı yayılım sorununa çözüm önerileri sunulmuştur. Sonuç olarak, sayıma dayalı verilerde aşırı yayılım durumu doğru yaklaşımlar kullanılarak incelenmelidir. Bu sorunun göz ardı edilmesi ya da doğru yaklaşımların kullanılmaması durumunda, oldukça küçük standart hatalar, şişirilmiş test istatistikleri ve I. tip hata oranının beklenenden çok daha büyük elde edilmesine yol açabilir. Bu durumda, sonuçlar yanlı ve güvenilmezdir.

Anahtar kelimeler: Sayım verisi, Aşırı yayılım, Poisson dağılım, Negatif binom dağılım

Petrol Fiyat Endeksi ve Doğrusal Olmayan Yaklaşım

Buket TAŞTAN^{1*}, M. Kenan TERZİOĞLU²

¹Trakya Üniversitesi, Sosyal Bilimler Enstitüsü, Ekonometri Anabilim Dalı, 22100, Edirne, Türkiye

²Trakya Üniversitesi, İktisadi ve İdari Bilimler Fakültesi-Ekonometri Bölümü, 22100, Edirne, Türkiye

*Sunucu yazar e-posta: bukettastan2@trakya.edu.tr

Özet

Enerji tüketiminde önemli paya sahip olan petrol fiyatlarında meydana gelen değişim ile ülkelerin makro-ekonomik değişkenleri doğrudan ya da dolaylı olarak etkilenmektedir. Ekonomik yapıda yarattığı belirsizlikler ve ani değişimler, ülkelerin ithalatçı/ihracatçı yapıya sahip olması ve uluslararası ticaretteki konumlarına göre farklı düzeyde etkilere sahip olmaktadır. Çalışma kapsamında, New York borsasında işlem gören West Texas Intermediate ham petrolüne ilişkin veriler ele alınarak, petrol fiyatlarında meydana gelen değişimin ortalamada doğrusal olmayan modellere ortaya çıkartılması amaçlanmaktadır. Bu çerçevede, kendinden eşik değerli otoregresif TAR ailesine ait olan SETAR, MTAR (Momentum TAR) ve yumuşak geçişli otoregresif STAR modellerden LSTAR ele alınarak modeller arasında karşılaştırma yapılmıştır ve en uygun model olarak SETAR seçilerek rejimler arası geçiş yorumlanmıştır.

Anahtar kelimeler: TAR, SETAR, MTAR, Petrol fiyat endeksi

Oil Price Index and Non-Linear Approach

Abstract

The change in oil prices, which has a significant share in energy consumption, directly or indirectly affects the macro-economic variables of the countries. Uncertainties and sudden changes created by oil prices in the economic structure has different effects depending on the importer/exporter structure of the countries and their position in international trade. Within the scope of the study, it is aimed to reveal the change in oil prices using non-linear models by considering the data on West Texas Intermediate crude oil traded in the New York stock exchange. In this framework, SETAR and MTAR (Momentum TAR) models belonging to the self-threshold autoregressive TAR family, and LSTAR, soft-transition autoregressive STAR models, were compared and transition between regimes were interpreted by choosing SETAR the most suitable model.

Keywords: TAR, SETAR, MTAR, Oil price index

GİRİŞ

Küreselleşme ile birlikte artan sermaye aktarımı ülkelerin finansal etkileşimini arttırırken ülkelerin makro-iktisadi değişkenleri üzerinde sapmalara neden olmaktadır. Enerjide ithalata bağımlı ve döviz kurunun ekonomik sistem üstünde baskın olduğu ülkelerde, küresel petrol fiyat değişimlerinde ortaya çıkan ani hareketler piyasaları, özellikle gelişmekte olan ve enerji bağımlı ülkelerde, olumlu/olumsuz etkilere maruz bırakmaktadır. Küresel çapta enerji tüketiminde en fazla paya sahip olan petrolün ikamesinin kısıtlı olması ekonomik büyüme başta olmak üzere uluslararası ticareti ve gelişmişlik düzeyini kısa ve uzun vadede farklılaştırabilmektedir (Moshiri, 2015). Kısa ve uzun vadede petrol

fiyatlarında meydana gelen artış petrolü ithal eden ülkelerde ticaret haddi şokları yaratarak, petrol ithal eden ülkenin para birimi üzerinde azaltıcı etki yaratabilmektedir (Killian, 2010).

Zaman serilerinde ortalamada asimetri etkisini dikkate alan, doğrusal olmayan yapıyı inceleyen, modeller kullanılarak seri yapıları doğru bir şekilde ortaya konulabilmektedir (Barca ve Arabacı, 2020). Çalışma kapsamında, petrol fiyat serisinde doğrusal olmayan yapının varlığının belirlenmesi ve doğrusal olmama bulgusu altında seriye ilişkin model yapısının oluşturulması için doğrusal olmayan yöntemler ele alınmaktadır. Tong (1978), Tong ve Lim (1980) ve Tong (1983), eşik değerli otoregresif (TAR) modeli ortaya konmaktadır. Eşik değişkeninin ve eşik değerinin bilinmemesi sorununun ortadan kaldırılması için kendinden uyarımlı eşik değerli otoregresif (SETAR) model yapısı geliştirilmektedir. Enders ve Granger (1998) ve Enders ve Siklos (2001), Momentum TAR modelini önermektedir. Ek olarak, rejimler arası geçişin çok sert olmasının yanıltıcı sonuçlar ortaya çıkartabilmesinden dolayı rejimler arası geçişte fonksiyonel yapı kullanılarak yumuşaklığın sağlanmasına imkân sağlayan yumuşak geçiş eşik değerli otoregresif (STAR) model yapısı kullanılmaktadır.

Dufrenot vd. (2005), hisse senedi ve bireysel varlık fiyatları ile birinci rejimde uzun hafıza süreci, ikinci rejimde ise kısa hafıza süreci elde ederek iki rejimli SETAR modelini sunmaktadır. Akgül ve Özdemir (2012), enflasyonun ekonomik büyüme üzerindeki etkilerini iki rejimli TAR modeli ile incelemektedir. Aydın ve İşçi (2012), ihracat birim endeks verisini ele alarak, parametrik yöntemlerden AR ve SETAR, parametrik olmayan yöntemlerden toplamsal regresyon modelini (ARM) ele alarak uygun model yapısını belirlemeye çalışmaktadır. Mohammadi ve Jahan-Parvar (2012), reel petrol fiyatları ve döviz kurları arasındaki uzun ve kısa dönemli ilişkiyi TAR ve MTAR modelleri kapsamında incelemektedir. Gemici ve Polat (2019), kripto para fiyat davranışını otoregresif birim kökü olan iki rejimli TAR modeli kullanarak araştırmaktadır. Alp (2008), reel ücretlerin yapısını TAR modeli ile incelemektedir. Goo ve Chen (2020), bakır vadeli işlem getirilerinin oynaklığı üzerindeki asimetrik momentum eşiği etkilerini gözlemlemek amacıyla MTAR-GARCH modelini kullanmaktadır.

Çalışma kapsamında, petrol fiyat endeks serisi (WTI) 24 Ocak 1986-21 Mayıs 2021 dönem aralıkları ele alınarak en uygun ortalamada doğrusal olmayan model yapısının tahmin edilmesi amaçlanmaktadır. İlk bölümde konuya genel bir giriş yapıldıktan sonra, ikinci bölümde ekonometrik metodolojiye yer verilmektedir. Üçüncü bölümde, petrol fiyat endeksine ilişkin en uygun model yapısı belirlenmektedir. Son bölümde ise ampirik bulgulardan yola çıkarak sonuç ve önerilere yer verilmektedir.

YÖNTEM

Konjonktürel dalgalanmaların asimetrik ve doğrusal olmayan etkileri rejim değişikliği olarak tanımlanmaktadır. TAR modelin özel bir durumu olan SETAR modelinde eşik değer, modelin bağımlı değişkeninin gecikmeli değerleri arasından seçilmektedir. Eşik değer bu özelliği ile rejim değişimine izin vermesi modelin yapısını doğrusal olmayan hale getirmektedir (Güriş, 2020).

Eşik değerle karşılaştırılan $q_t = y_{t-d}$ terimi ve $d = 1$ olmak üzere iki rejimli SETAR modeli, otoregresif parametreler ϕ , eşik parametre c ve bağımsız ve özdeş dağılan hata terimi ε_t olmak üzere,

$$y_t = \begin{cases} \phi_{1,0} + \phi_{1,1}y_{t-1} + \varepsilon_t & y_{t-1} \leq c \\ \phi_{2,0} + \phi_{2,1}y_{t-1} + \varepsilon_t & y_{t-1} > c \end{cases} \quad (1)$$

şeklinde tanımlanmaktadır. Modelde rejim sayısı ikiden fazla olması durumunda SETAR model yapısı,

$$y_t = \begin{cases} \phi_{1,0} + \phi_{1,1}y_{t-1} + \dots + \phi_{1,p_1}y_{t-p_1} + \varepsilon_t & y_{t-d} \leq \gamma \\ \phi_{2,0} + \phi_{2,1}y_{t-1} + \dots + \phi_{2,p_2}y_{t-p_2} + \varepsilon_t & y_{t-d} > \gamma \end{cases} \quad (2)$$

şeklinde gösterilmektedir (Brooks, 2002).

MTAR model, standart Dickey Fuller testinin asimetrik şekilde düzenlenmesi ile tanımlanmaktadır. $I(.)$ gösterge fonksiyonu ve eşik parametresi τ olmak üzere,

$$I_t = \begin{cases} 1 & \text{eğer } y_{t-1} \geq \tau \\ 0 & \text{eğer } y_{t-1} < \tau \end{cases}$$

$$I_t = \begin{cases} 1 & \text{eğer } \Delta y_{t-1} \geq \tau \\ 0 & \text{eğer } \Delta y_{t-1} < \tau \end{cases}$$

MTAR modeli,

$$\Delta y_t = I_t \rho_1 y_{t-1} + (1 - I_t) \rho_2 y_{t-1} + \sum_{i=1}^p \gamma_i \Delta y_{t-i} + \varepsilon_t \quad (3)$$

şeklinde ifade edilmektedir. Δy_{t-1} , eşik değerden büyük olduğunda uyarılma katsayısı ρ_1 , Δy_{t-1} eşik değerden küçük olduğunda ise uyarılma katsayısı ρ_2 olmaktadır. MTAR modelde ilk olarak TAR modeline ait y_{t-1} değerleri ve sonrasında ise MTAR modeline ait Δy_{t-1} değerleri küçükten büyüğe doğru sıralanmaktadır. Bu sıralama gerçekleştirilirken serinin maksimum ve minimum değerlerinin %15'i inceleme dışı (trim) bırakılmaktadır. Eşik değer seçimi en küçük hata terimleri karesine sahip olan modele göre belirlenmektedir (Bildirici, 2020).

Eşik değerin belirlenmesinden sonra rejimler arasında geçişin ani olduğu (TAR) modeller yerine rejimler arasında geçişin yumuşak olduğu ve lojistik yumuşak geçiş otoregresif (LSTAR) ile üstel yumuşak geçiş otoregresif (ESTAR) olmak üzere ikiye ayrılan STAR modelleri tercih edilmektedir. LSTAR modeli ekonominin daralma ve genişleme dönemlerinin farklı dinamiğe sahip olduğunu ve bunlar arası geçişin yumuşak olması gerektiğini savunmaktadır (Güriş, 2020). STAR modelinin genel biçimi, sıfır ve bir arasında değer alan geçiş fonksiyonu $G(y_{t-1}, \gamma, c)$ ve parametreye göre bir rejimden diğerine geçiş sağlayan düzgünleştirme parametresi γ olmak üzere,

$$y_t = (\phi_{1,0} + \phi_{1,1} y_{t-1})(1 - G(y_{t-1}, \gamma, c)) + (\phi_{2,0} + \phi_{2,1} y_{t-1})(1 - G(y_{t-1}, \gamma, c)) + \varepsilon_t \quad (4)$$

şeklinde verilmektedir (Franses ve Dijk, 2000). Geçiş fonksiyonu LSTAR modele göre,

$$G(s_t; \gamma, c) = (1 + \exp\{-\gamma(s_t - c)\})^{-1} \quad \gamma > 0 \quad (5)$$

şeklinde tanımlanmaktadır. ($\gamma \rightarrow \infty$) olduğunda lojistik fonksiyon bire yaklaşmakta ve LSTAR modeli iki rejimli TAR modeline dönüşürken; ($\gamma = 0$) olduğunda LSTAR modeli doğrusal AR modeline dönüşmektedir.

BULGULAR

Doğrusal olmama olgusundan yola çıkılarak incelenen 24 Ocak 1986 ve 21 Mayıs 2021 dönemleri kapsamında Energy Information Administration veri tabanından elde edilen WTI (West Texas Intermediate) ham petrol veri seti kullanılarak petrol fiyat endeksine ilişkin model yapısının belirlenmesi amaçlanmaktadır. Varyansın kararlılığının sağlanması ve aykırı gözlemlerin mevcut olması durumunda etkilerinin azaltılması, diğer bir ifadeyle serilerde mevcut olan aşırı değişiklerin dengelenmesi, için petrol fiyat endeks serisine ilişkin logaritmik dönüşüm yapılmaktadır. Tablo 1.'de WTI petrol verisine ait tanımlayıcı istatistiklere yer verilmektedir.

Tablo 1. WTI petrol verisine ait tanımlayıcı istatistikler

Tanımlayıcı İstatistikler							
	Ortalama	Max.	Min.	St. Sapma	Çarpıklık	Basıklık	J-B
WTI	3.580	4.978	2.987	0.647	0.195	-1.301	685.58*

***, **, * sırasıyla %10, %5 ve %1 önem düzeyini göstermektedir.

Tablo 1.'de yer alan tanımlayıcı istatistikler incelendiğinde, WTI petrol serisinin çarpıklı katsayısının pozitif olması sağa çarpık dağılımı ve basıklık katsayısının negatif olması ise basık dağılımı işaret etmektedir. Jarque-Bera test istatistiği de serinin normal dağılmadığını göstermektedir. Değişkenliğin bir göstergesi olarak standart sapma incelendiğinde, WTI petrol fiyatlarında değişkenliklerin olduğu gözlemlenmektedir. WTI petrol serisinin doğrusallığının tespit edilmesinde, Tablo 2'de yer alan, SETAR tipi Hansen (1999) testi, SETAR tipi LR testi, Keenan (1985) testi ve Keenan testinin geliştirilmiş hali olan çapraz çarpımların kullandığı Tsay (1986) testlerinden faydalanılmaktadır.

Tablo 2. WTI petrol değişkeninin doğrusallık sınaması

	Hansen Testi		
	1vs2	1vs3	2vs3
Test İstatistiği	42.275*	45.435*	3.1450
	LR Testi	Keenan Testi	Tsay Testi
	421.3869*	3.5603***	6.899*

***, **, * sırasıyla %10, %5 ve %1 önem düzeyini göstermektedir.

Tablo 2.'de Hansen (1999) testi çerçevesinde doğrusallık, iki ve üç rejime karşı sınanmaktadır. Doğrusallığın iki rejime karşı test edilmesinde olasılık değeri anlamlı olduğu için yokluk hipotezi red edilmekte ve serinin iki rejime sahip olduğuna karar verilmektedir. Doğrusallığın üç rejime karşı test edilmesinde ise olasılık değeri anlamlı olduğu için yokluk hipotezi red edilerek serinin üç rejimli olduğuna karar verilmektedir. İki rejimin üç rejime karşı sınanmasında istatistiksel olarak anlamlı sonuç elde edilememektedir. Tablo 2.'de yer alan LR testi ve Tsay (1986) testi çerçevesinde olasılık değeri anlamlı olduğu için doğrusallığı savunan yokluk hipotezi red edilmektedir. Keenan (1985) testinde ise olasılık değeri anlamlı olduğu için doğrusallığı savunan yokluk hipotezi red edilmektedir. Bu sonuçlara göre modellerin doğrusal olmayan bir yapı gösterdiği kabul edilmektedir.

WTI petrol serisinin doğrusal olmadığı belirlendikten sonra Tablo 3.'de yer alan SETAR model yapısına dayanan doğrusal olmayan birim kök testi Caner ve Hansen (2001) testi, MTAR model yapısına dayanan ait doğrusal olmayan birim kök testi Enders ve Granger (1998) testi, LSTAR model yapısına dayanan doğrusal olmayan birim kök testi Sollis (2004) testi kullanılarak doğrusal olmayan birim kök sınaması yapılmaktadır. SETAR tipi Caner ve Hansen (2001) testi, tek taraflı test ve birimci rejime göre durağan iken ikinci rejimde durağan değildir. MTAR tipi Enders ve Granger (1998) testine göre WTI petrol serisinin tüm düzeylerde durağan olduğu, LSTAR tipi Sollis (2004) testine göre WTI petrol serisinin Durum 2 ve Durum 3'de durağan olduğu gözlemlenmektedir. Elde edilen sonuçlara göre, tüm test istatistikleri göz önüne alınarak seri durağan kabul edilmektedir.

Tablo 3. WTI petrol değişkenine ait doğrusal olmayan birim kök sınaması

Caner ve Hansen (2001) Testi			
	Tek Taraflı Wald Testi	t₁	t₂
Test İstatistiği	4.2*	1.560***	1.340
Enders ve Granger (1998) Testi			
	Durum 1	Durum 2	Durum 3
Test İstatistiği	5.678*	7.581*	11.297*
Sollis (2004) Testi			
	Durum 1	Durum 2	Durum 3
Test İstatistiği	11.453	12.335**	14.876**

*Caner ve Hansen testine ait değerler Bootstrap değerleri dikkate alınarak hesaplanmaktadır. k=36, m=4 gecikme parametreleridir. t₁=Birinci rejime ait birim kök sınaması, t₂=İkinci rejime ait birim kök sınaması olarak ifade edilmektedir. Enders ve Granger (1998) testine ait tablo değerleri %1 için Durum 1=4.85, Durum 2=6.91, Durum 3=8.74'dür. Sollis (2004) testine ait tablo değerleri %5 için Durum 1=8.845, Durum 2=11.309, Durum 3=12.558'dir. Durum 1: ham veri, Durum 2: ortalamadan arındırılmış veri ve Durum 3: ortalama ve trendden arındırılmış veri tanımlanmasında kullanılmaktadır.

Tablo 4.'de, petrol fiyat endeks serisinin durağan bulunması sonucunda, SETAR, MTAR ve LSTAR eşik değerli modeller karşılaştırılmakta ve modellere ilişkin eşik değerler verilmektedir. SETAR, MTAR ve LSTAR model yapıları karşılaştırıldığında alt ve üst rejim parametrelerinin tüm gecikmelerde anlamlı olması ile en uygun modelin SETAR eşik modeli olduğu sonucuna varılmaktadır. SETAR modeline ilişkin F istatistiği ele alındığında, modelin doğrusal olduğu sonucuna varılmaktadır.

Tablo 4. WTI petrol değişkenine ait eşik değerli modellerin karşılaştırılması

SETAR Model					
		Sabit	WTI(-1)	WTI(-2)	
Üst Rejim		0.082*	0.854*	0.116*	
Alt Rejim		0.0005	0.972*	0.027**	
F İstatistiği	2.485*	Eşik Değeri		2.8937	
MTAR Model					
		Sabit	WTI(-1)	WTI(-2)	WTI(-3)
Üst Rejim		-0.009*	0.809*	0.237*	-0.045*
Alt Rejim		0.005*	0.955*	0.002	0.040*
			Eşik Değeri	-0.010	
LSTAR Model					
		Sabit	WTI(-1)	WTI(-2)	WTI(-3)
Üst Rejim		0.047*	0.865*	0.005	0.113*
Alt Rejim		-0.047*	0.106*	0.050***	-0.141*
F İstatistiği	3.120**	Eşik Değeri	2.936*	Geçiş Hızı	100

***, **, * sırasıyla %10, %5 ve %1 önem düzeyini göstermektedir.

MTAR modelinde birinci rejime ilişkin ikinci gecikme istatistiksel olarak anlamlı iken, ikinci rejime geçildiğinde seriye ilişkin ikinci gecikme istatistiksel olarak anlamsızlaşmaktadır. LSTAR modeli ele alındığında ise, birinci rejimde ikinci gecikme istatistiksel olarak anlamsız iken, ikinci rejime geçildiğinde ikinci gecikme istatistiksel olarak anlamlı olarak gözlemlenmektedir. Bununla birlikte, WTI serisine göre SETAR model tahmini;

$$WTI_t = \begin{cases} 0.082 + 0.854_{WTI_{t-1}} + 0.116_{WTI_{t-2}} + \varepsilon_t, & WTI_{t-2} \leq 2.8937 \\ 0.0005 + 0.972_{WTI_{t-1}} + 0.027_{WTI_{t-2}} + \varepsilon_t, & WTI_{t-2} > 2.8937 \end{cases}$$

şeklinde gerçekleşmektedir. Eşik değerine göre, petrol fiyat endeksinde 2.8937 oranında bir değişim gerçekleştiğinde alt rejimden üst rejime bir kayma gerçekleşmektedir. Rejim değişimin üç dönem öncesinden öngörülebildiği belirlenmektedir. Diğer bir ifadeyle, WTI serisinde meydana gelmesi muhtemel değişim üç dönem önceden tahmin edilebilmektedir. Şekil 1'de SETAR modele ait rejim değişimi ve eşik değeri (kırmızı çizgi) gösterilmektedir. İki rejimin gözlemlendiği grafikte, SETAR modelin özelliği olarak rejimden diğer rejime geçişin çok sert olduğu belirlenmektedir.



Şekil 1: WTI Petrol Değişkenine ait SETAR Rejim Değişim ve Eşik Değer Grafiği

TARTIŞMA VE SONUÇ

Petrol fiyat endeksinde meydana gelen değişimin makro-iktisadi değişkenler ve uygulanan/belirlenen politikalar üzerinde dolaylı/doğrudan etkileri bulunmaktadır. WTI serisinin doğrusal yapısı hakkında bilgi sahibi olunmadan, doğrusal modeller kullanılarak yapılan tahminleme çalışmalarında doğru sonuçlara ulaşılamamaktadır. Doğrusal olmayan modeller, ortalamada ve varyansta doğrusal olmayan modeller olarak ikiye ayrılmaktadır. Çalışma kapsamında, ortalamada doğrusal olmayan modellerden SETAR, MTAR ve LSTAR modelleri kullanılarak en uygun model seçimine karar verilmektedir. Bu amaçla, New York borsasında işlem gören WTI petrol fiyat endeks serisi ele alınarak, petrol fiyatlarında meydana gelen artış ya da azalışların etkileri SETAR, MTAR ve LSTAR model ile incelenmektedir. Elde edilen bulgulara göre, WTI serisinin ortalama doğrusal olmayan seri olduğu tespit edilmekte ve tüm gecikmelere ait katsayıların anlamlı çıkması dolayısıyla SETAR model yapısının en uygun model olduğu belirlenmektedir. SETAR model yapısına göre, fiyat endeksi belirli bir seyir halinde iken 2.8937 oranında değişim gerçekleştiğinde diğer rejime geçiş sağlandığı tespit edilmektedir.

Kaynaklar

- Akgül I., Özdemir S., 2012. Enflasyon Eşiği ve Ekonomik Büyümeye Etkisi, İktisat, İşletme ve Finans, 27(313): 85-106.
- Alp, E.A., 2008. Türkiye’de Reel Ücretlerin TAR Modeli İle Analizi ve Birim Kök Sınaması, Tartışma Metni No.2008/10, www.tek. org.tr.
- Aydın D., İşçi Ö., 2012. Doğrusal Olmayan Otoregresif Zaman Serileri Modellerinin Kestirimi, Muğla Üniversitesi Sosyal Bilimler Enstitüsü Dergisi, 28: 205-218.
- Barca O., Arabacı Ö., (2020). BİST Altın Fiyatları Serisinin Markov Rejim Değişim Modeli İle Analizi, Muhasebe ve Finansman Dergisi, (85): 209-222.
- Bildirici M., Alp, A.E., Ersin Ö.Ö., Bozoklu Ü., 2010. İktisatta Kullanılan Doğrusal Olmayan Zaman Serisi Yöntemleri, Türkmen Kitabevi, Yayın No:357, 1.Baskı, 384s, İstanbul, Türkiye.
- Brooks C., 2002. Introductory Econometrics for Finance. Cambridge University Pres, Cambridge.
- Caner M., Hansen B. E., 2001. Threshold Autoregression with a Unit Root. Econometrica, 69(6): 1555-1596.
- Dufrenot G., Guegan D., Feisolle A.P., (2005). Long-Memory Dynamics in A SETAR Model – Applications to Stock Markets, Journal of International Financial Markets, Institutions and Money, 15(5): 391-406.
- Enders W., Granger C., 1998. Unit Root Tests and Asymmetric Adjustment with an Example Using the Term Structure of Interest Rates, Journal of Business & Economic Statistics, 16(3): 304-311.
- Enders W., Siklos P.L., 2001. Cointegration and Threshold Adjustment, Journal of Business & Economic Statistics, 19(2): 166-176.
- Franses P.H., Dijk V.D., 2000. Non-linear Time Series Models in Empirical Finance. Cambridge University Press.
- Gemici E., Polat M., 2019. Bitcoin Fiyatlarında Eşik Değer Etkisi, Mehmet Akif Ersoy Üniversitesi İktisadi ve İdari Bilimler Fakültesi Dergisi, 6(3): 669-681.
- Goo Y.J., Chen C.C., 2020. Asymmetric Momentum Threshold Effect of Copper Futures Returns on Spot Returns Volatility in London Metals Exchange under High Volatility, Modern Economy, 11(1): 51-61.
- Güriş B., 2020. R Uygulamalı Doğrusal Olmayan Zaman Serileri Analizi. Der Yayınları No: 0293, 1. Baskı, 290s, İstanbul, Türkiye.
- Hansen B.E., 1996. Inference When a Nuisance Parameter is not Identified Under the Null Hypothesis, Econometrica: Journal of The Econometric Society, 413-430.
- <http://www.econstor.eu/handle/10419/57602>, Date of Access: 15.06.2013.
- Kapetanios G., Shin Y., 2006. Unit Root Tests in Three-Regime SETAR Models, Econometrics Journal, 9(2): 252-278.
- Keenan D.M. 1985. A Tukey Nonadditivity-Type Test for Time Series Nonlinearity, Biometrika, 72: 39-44.

- Killian, L., (2010), Oil Price Volatility: Origins and Effects, World Trade Organization Staff Working Paper, ERSD-2010-02, Internet Address:
- Mohammadi H., Jahan-Parvar M.R., 2012. Oil Prices and Exchange Rates in Oil-Exporting Countries: Evidence From TAR and M-TAR Models, *Journal of Economics and Finance*, 36: 766-779.
- Moshiri S. (2015). Asymmetric Effects of Oil Price Shocks in Oil-Exporting Countries: The Role of Institutions, *OPEC Energy Review*, 39 (2): 222- 246.
- Sollis R., 2004. Asymmetric Adjustment and Smooth Transitions: A Combination of Some Unit Root Tests, *Journal of Time Series Analysis*, 25(3): 409-417.
- Tong H., 1978. On a Threshold Model, *Pattern Recognition and Signal Processing*, Sijthoff and Noordhoff, Amsterdam.
- Tong H., 1983. *Threshold Models in Non-Linear Time Series Analysis* 1.bs. New York: Springer-Verlag.
- Tong H., Lim K.S., (1980). Threshold Autoregression, Limit Cycles and Cyclical Data. *Journal of Royal Statistical Society B*, 42(3): 245-292.
- Tsay R.S., 1989. Nonlinearity Tests for Time Series, *Biometrika*, 73(2): 461-466.

Petrol Fiyat Endeksinin Model Yapısı: Markov Rejim Değişim Modeli

Buket TAŞTAN^{1*}, M. Kenan TERZİOĞLU²

¹Trakya Üniversitesi, Sosyal Bilimler Enstitüsü, Ekonometri Anabilim Dalı, 22100, Edirne, Türkiye

²Trakya Üniversitesi, İktisadi ve İdari Bilimler Fakültesi, Ekonometri Bölümü, 22100, Edirne, Türkiye

*Sunucu yazar e-posta: bukettastan2@trakya.edu.tr

Özet

Kısıtlı ikameye sahip olan ve enerji tüketiminde önemli paya sahip olan petrol, ülkelerin gelişmişlik düzeyleri ve sanayileşme oranlarına bakılmaksızın ülke ekonomilerinde belirleyici rol oynamaktadır. Bu çerçevede, petrol fiyatlarında ortaya çıkan muhtemel dalgalanmalar, arz ve talep ilişkisi ve ülkelerin ekonomik istikrarı göz önünde bulundurularak, makro iktisadi değişkenleri etkilemektedir. Çalışma kapsamında, West Texas Intermediate ham petrolüne ait 24 Ocak 1986-24 Mayıs 2021 dönemleri arası veriler göz önüne alınarak petrol fiyatlarında meydana gelen yapısal değişikliklerin doğrusal olmayan yapısının Markov rejim değişim modeliyle ortaya konması amaçlanmaktadır.

Anahtar kelimeler: Markov rejim değişim modeli, Petrol fiyat endeksi, WTI

Oil Price Index Model Structure: Markov Regime Switch Model

Abstract

Petroleum, which has limited substitution and has a significant share in energy consumption, plays a decisive role in the country's economy, regardless of the development level and industrialization rates of the countries. In this context, possible fluctuations in oil prices, affects macroeconomic variables according to supply-demand relations and the countries' economic stability. Within the scope of the study, it is aimed to reveal the nonlinear structure of the structural changes in oil prices with the Markov regime switching model, taking into account the data of West Texas Intermediate crude oil between January 24, 1986 and May 24, 2021.

Keywords: Markov regime switch model, Oil price index, WTI

GİRİŞ

Ekonomik yapı içerisinde enerji tüketim talebinin karşılanması için petrol ve türevlerine ayrılan kısmın fazla olması, petrol fiyat değişimlerinin, özellikle ithalatçı yapıda olan ülkelerde döviz kuru ile de senkronize olması sonucunda, makro-ekonomik değişkenler üzerindeki etkilerinin hissedilir derecede olmaktadır. İkamesi kısıtlı olan petrolün fiyatlarındaki artış ile birlikte ülkelerin diğer mallara ayırdığı oranlarda azalma gözlemlenmektedir. (Sadorsky, 2006).

Zaman serilerinde meydana gelen değişimlerin ortaya çıkarılmasında ve uygulanması planlanan politika önerilerin geliştirilmesinde doğrusal modeller ile ortaya konan bulgular serilerin özelliklerini tam olarak yansıtamadığında asimetrik etkiyi dikkate alan rejim değişim modellerine geçiş yapılmaktadır. Bu çerçevede, Neftçi (1984) ve Hamilton (1989), iktisadi değişkenlerin özelliklerinin ortaya çıkarılmasında doğrusal olmama durumu ve asimetrik etki yaklaşımını kullanmaktadır. Hamilton (1989), konjonktürel dalgalanmanın genişleme-daralma dönemlerinde (farklı rejimler) sergilediği

asimetriyi birinci dereceden Markov zinciri (MS-AR) ile modellemektedir. Hamilton (1990), farklı rejimler arasındaki geçişi olasılıksal olarak ifade etmektedir. Markov rejim değişim modelleri, diğer rejim modellerinden farklı olarak, olasılıksal çıkarım yapılmasını sağlamaktadır (Bildirici, 2010). Krolzig (1997), rejimler arasındaki geçişin sabit ve tek değişkenli olarak modellenmesindeki kısıtları kaldırmak için çok değişkenli Markov rejim değişim otoregresif modellerini ortaya koymaktadır.

Stale, Kayhan ve Koçyiğit (2013), işsizlik oranının asimetrik etkilerini Markov rejim değişim modeli ile açıklamaktadır. Özdemir ve Akgül (2015), ham petrol fiyatları ve benzin fiyatlarındaki ani değişimlerin sanayi üretimine etkisini Markov değişim vektör otoregresif modelleri (MS-VAR) ile incelemektedir. Aktaş vd. (2016), petrol arz-talebinin petrol fiyatları üzerindeki etkisini Markov rejim değişim modeli kullanarak ortaya koymaktadır. Evci vd. (2016), altın piyasasına ilişkin rejim etkisini gözlemlemek amacıyla Markov rejim değişim modellerini kullanmaktadır. Koy (2017), fiyat endeksi ve vadeli işlem sözleşmelerindeki fiyat değişimlerini Markov rejim değişim vektör otoregresif model ile açıklamaktadır. Çil ve Yılmaz (2018), ham petrol fiyatının doğrusal olmayan yapısını Markov rejim değişim otoregresif modellerle test etmektedir. Maupin (2019), kripto para serilerinin gelecek fiyatlamasını AR, MS-AR, VAR, MS-VAR yöntemlerini kullanarak tahmin etmektedir.

Çalışma kapsamında, Markov rejim değişim modeli kullanılarak WTI (West Texas Intermediate) petrol fiyat endeksine ilişkin simetrik/asimetrik yapının varlığının araştırılması ve asimetrik yapının ortaya çıkması durumunda ise rejim değişim olasılıklarının belirlenmesi amaçlanmaktadır.

Markov Rejim Değişim Modeli

Doğrusal olmama durumu, parametrelere ilişkin ifadelerde farklı yapıların kullanılması (kuvvet alımı, bölme/çarpma, vb. gibi) durumunda ortaya çıkan parametrelere göre doğrusal olmama durumu ve/veya değişkenlerin matematiksel yapısından dolayı ortaya çıkan değişkenlere göre doğrusal olmama durumu olmak üzere iki durum altında sınıflandırılmaktadır. Modeller ele alındığında, dönüşüm teknikleri uygulanarak doğrusallaştırma söz konusu olurken; bazı durumlarda ise dönüşümler doğrusallaştırma işlemini sağlayamamakta ve doğrusal olmayan model yapılarının analize edilemesi gerekliliğini ortaya koymaktadır.

Hamilton (1989) ve Hamilton (1990), ekonomideki konjonktüre ilişkin genişleme ve daralma dönemlerini farklı rejimler olarak inceleyen Markov rejim değişim otoregresif (MS-AR) modelini ortaya koymaktadır. Rejimler arasındaki geçişin olasılıksal olarak modellendiği Markov rejim değişim modeli, tek değişkenli dördüncü dereceden otoregresif model, iki durum ($M=2$) arasında değişen koşullu ortalama $\mu(s_t)$ ve $\mu_t \sim NID(0, \sigma^2)$ olmak üzere,

$$\Delta y_t - \mu(s_t) = \phi_1(\Delta y_{t-1} - \mu(s_{t-1})) + \dots + \phi_4(\Delta y_{t-4} - \mu(s_{t-4})) + \mu_t \quad \mu_t \sim NID(0, \sigma^2) \quad (1)$$

şeklinde ifade edilmektedir. Duruma ve rejime bağlı olarak μ daralma döneminde negatif değer alırken; genişleme döneminde pozitif değer almaktadır. MS-AR modeli, gözlemlenemeyen durum değişkeni olarak tanımlanan rejimi ifade eden s_t 'ye bağlı olmaktadır. s_t 'nin olasılık değerinin bir önceki döneme bağlı olması durumu, i durumunda j durumuna geçiş olasılığı p_{ij} olmak üzere,

$$\Pr(s_t = j | s_{t-1} = i, s_{t-2} = k, \dots, s_0 = h) = \Pr(s_t = j | s_{t-1} = i) p_{ij} \quad (2)$$

şeklinde gösterilmektedir (Bildirici vd., 2010)

Ekonomide genişleme ($s_t = 2$) ve daralma ($s_t = 1$) olmak üzere iki rejim $s_t = \{1, 2\}$ olduğu kabul edildiğinde, iki rejim arasındaki geçiş olasılığı

$$\begin{aligned} \Pr[s_1 = 1 | s_{t-1} = 1] &= p_{11} = p \\ \Pr[s_1 = 2 | s_{t-1} = 1] &= p_{12} = 1 - p \end{aligned} \quad (3)$$

$$Pr[s_1 = 1 | s_{t-1} = 2] = p_{21} = 1 - q$$

$$Pr[s_1 = 2 | s_{t-1} = 2] = p_{22} = q$$

şeklinde verilmektedir (Bildirici vd., 2010). y_t serisi ele alındığında, birinci rejimde kalma olasılığı (genişlemede iken bir dönem sonra genişlemede bulunma olasılığı) p , birinci rejimden ikinci rejime geçiş olasılığını (genişlemede iken bir dönem sonra daralmada bulunma olasılığı) $1 - p$, ikinci rejimde kalma olasılığını (daralmada iken bir dönem sonra yine daralmada kalma olasılığı) q , ikinci rejimden birinci rejime geçiş olasılığı (daralmada iken genişlemeye geçiş olasılığı) $1 - q$ olarak ifade edilmektedir. Birinci rejimde kalma süresi $(1/1 - p_{11})$ ve ikinci rejimde kalma süresi ise $(1/1 - p_{12})$ olarak hesaplanmaktadır (Hamilton, 1989).

Markov rejim değişim modelinde yer alan gözlemlenebilen y_t serisinin istatistiksel özelliği rejimin değişmesi ile birlikte değişim göstermesi durumunda, $E[y_t | s_t = 1] = \mu_1$ ya da $E[y_t | s_t = 2] = \mu_2$ elde edilmektedir. Ekonominin daralma dönemi μ_1 ve genişleme dönemi μ_2 olarak ele alındığından koşullu ortalama

$$\mu_{s_t} = \begin{cases} s_t = 2 \text{ ise } \mu_2 > 0 \text{ (genişleme)} \\ s_t = 1 \text{ ise } \mu_1 < 0 \text{ (daralma)} \end{cases}$$

şeklinde ifade edilmektedir (Bildirici vd., 2010). Markov rejim değişim modeli ise, genel haliyle,

$$y_t = \mu_{s_t} + \mu_t, \quad \mu_t \sim iid(0, \sigma^2) \quad (4)$$

şeklinde gösterilmektedir.

BULGULAR

Doğrusal olmama olgusundan yola çıkılarak incelenen 24 Ocak 1986 ve 21 Mayıs 2021 dönemleri arasında petrol arz/talebine ve fiyatına bağlı olarak farklılaşmaların/yapısal değişimlerin gözlemlendiği söylenebilmektedir. Bu nedenle, çalışma kapsamında, Energy Information Administration veri tabanından elde edilen WTI (West Texas Intermediate) ham petrol veri seti kullanılarak petrol fiyat endeksindeki dönemler itibarıyla yapısal farklılıklarının/dönüşümlerinin ortaya çıkartılması amaçlanmaktadır. Varyansın kararlılığının sağlanması ve aykırı gözlemlerin mevcut olması durumunda etkilerinin azaltılması, diğer bir ifadeyle serilerde mevcut olan aşırı değişikliklerin dengelenmesi, için petrol fiyat endeks serisine ilişkin logaritmik dönüşüm yapılmaktadır. Tablo 1.'de WTI petrol verisine ait tanımlayıcı istatistiklere yer verilmektedir.

Tablo 1. WTI petrol verisine ait tanımlayıcı istatistikler

	Tanımlayıcı İstatistikler						
	Ortalama	Max.	Min.	St. Sapma	Çarpıklık	Basıklık	J-B
WTI	3.580	4.978	2.987	0.647	0.195	-1.301	685.58*

***, **, * sırasıyla %10, %5 ve %1 önem düzeyini göstermektedir.

Tablo 1.'de yer alan tanımlayıcı istatistikler incelendiğinde, WTI petrol serisinin çarpıklık katsayısının pozitif olması dağılımların sağa çarpık olduğunu, basıklık katsayısının negatif olması ise basık bir dağılıma sahip olduğunu işaret etmektedir. Jarque-Bera test istatistiği incelendiğinde ise serinin normal dağılmadığı gözlemlenmektedir. Değişkenliğin bir göstergesi olarak standart sapma incelendiğinde, WTI petrol fiyatlarında değişkenliğin olduğu ifade edilebilmektedir. Markov rejim değişim modeli oluşturulmadan önce WTI petrol değişkeninin doğrusal bir yapıya sahipliğinin ortaya çıkarılmasında doğrusal olmayan testlerden Brock, Dechert, Scheinkman (BDS) (1987) testi ele alınarak, serinin bağımsız ve eş dağılıma (iid) sahipliği Tablo 2.'de gösterilmektedir. Serideki gözlem sayısı 500'den fazla ise boyut değeri (m) altıdan (6) küçük ve ε değeri standart sapmanın 0.5 ile 2 katı

arasında seçilmektedir (Sümer ve Hepsağ, 2007). BDS testinin yanı sıra Keenan Testi (1985) ve Keenan testinden farklı olarak değişkenlerin çapraz çarpımlarının kullanıldığı Tsay Testi (1986) sonuçlarına da Tablo 2.'de gösterilmektedir.

Tablo 2. WTI petrol değişkenine ait doğrusallık sınaması

Boyut	BDS Test			
	m=2	m=3	m=4	m=5
$\varepsilon = 0.5$	1492.222*	2876.860*	6098.895*	14113.027*
$\varepsilon = 1$	1734.983*	2315.411*	3207.957*	4664.218*
$\varepsilon = 1.5$	566.7738*	634.8978*	721.5764*	843.9393*
$\varepsilon = 2$	342.7166*	341.4685*	341.5307*	341.5307*
Keenan Test				
Test İstatistiği	3.560***			
Tsay Testi				
Test İstatistiği	6.899*			

***,**, * sırasıyla %10, %5 ve %1 önem düzeyini göstermektedir

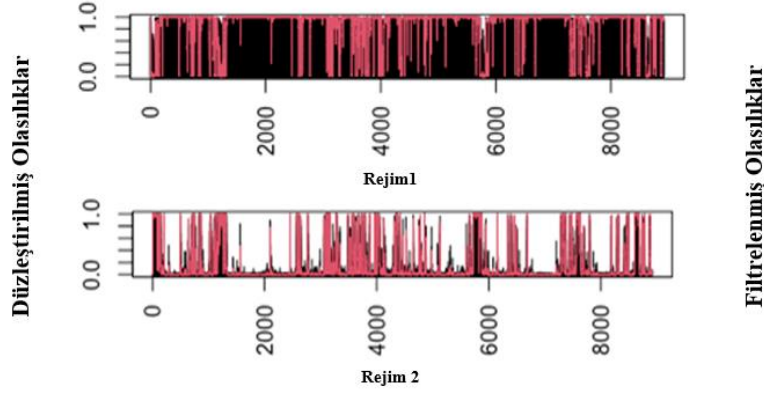
Tablo 2.'de hata terimlerine uygulanan BDS testi sonuçlarına göre, tüm ε değerlerinde ve m boyutlarında hesaplanan değerlerin hata terimlerinin bağımsız ve eş dağılıma sahip olduğunu savunan hipotez kabul edilmemektedir. Keenan test istatistiğine ve Tsay test istatistiğine göre de serilerin doğrusal olduğunu savunan hipotez kabul edilmemektedir. Bu çerçevede, modellerin doğrusal olmayan bir yapı gösterdiği tespit edildikten sonra petrol fiyat endeksi serisinin yapısının ortaya çıkarılması için doğrusal olmayan zaman serisi yöntemlerinden olan Markov rejim değişim modeline geçiş yapılmaktadır.

Tablo 3. WTI petrol değişkenine ait markov rejim değişim modeli

Değişkenler	Rejim 1		
	Katsayı	St. Hata	t Değeri
C	0.0017	0.0012	1.4167
WTI (-1)	0.9884*	0.0051	193.8039
WTI(-2)	0.0113**	0.0051	2.2157
Değişkenler	Rejim 2		
	Katsayı	St. Hata	t Değeri
C	0.0274**	0.0128	2.1406
WTI (-1)	0.9023*	0.0360	25.0639
WTI(-2)	0.0886**	0.0358	2.4749
Rejim Geçiş Olasılıkları			
	Rejim 1	Rejim 2	
Rejim 1	0.9838	0.1199	
Rejim 2	0.0161	0.8800	

***,**, * sırasıyla %10, %5 ve %1 önem düzeyini göstermektedir

Tablo 3.'de WTI petrol değişkenine ait Markov rejim değişim katsayılarının istatistiksel olarak anlamlı olduğu gözlemlenmektedir. Elde edilen bulgulara göre, WTI petrol değişkeninin iki rejimli doğrusal olmayan bir yapı sergilediği sonucuna ulaşılmaktadır. Rejim geçiş olasılıklarına göre, Rejim 1'de ve Rejim 2'de kalıcılık söz konusu olurken; Rejim 1'den sonraki dönemde Rejim 2'ye geçiş olasılığı oldukça düşük (0.1199) ve Rejim 2'den sonraki dönemde Rejim 1'e geçiş olasılığı oldukça düşük (0.0161) hesaplanmaktadır.



Şekil 1. Rejim geçiş olasılıkları grafiği

Şekil 1.'de Rejim 1 ve Rejim 2 dönemlerine ilişkin geçiş ve bulunma (kalma) olasılıkları gösterilmektedir. WTI petrol fiyatlarında meydana gelen dalgalanmaların hem Rejim 1'de hem Rejim 2'de yakın olasılıklar gösterdiği ifade edilebilmektedir. Bu durum dalgalanmaların az olduğu durumlarda rejimlerin birbirlerine geçiş olasılıklarının düşük olduğu sonucunu ortaya çıkarmaktadır.

TARTIŞMA VE SONUÇ

Doğrusallık varsayımını ele alan model yapıları, asimetrik davranışa sahip serilerin özelliklerini çıkarmada yetersiz kalmaktadır. Bu durumda konjonktürel dalgalanmada meydana gelen etkileri modellemeyi sağlayan Markov rejim değişim modelleri kullanılabilir. Markov rejim değişim modelleri konjonktürel dalgalanmayı genişleme ve daralma olmak üzere iki farklı rejim olarak ele alabilmektedir. Asimetrik etkiyi ortaya çıkaran farklı doğrusal olmayan model yapıları olmakla birlikte, Markov rejim değişim modeli rejimden rejime geçiş olasılıklarını da sağlamaktadır. Çalışma kapsamında, WTI petrol değişkeninin iki rejimli doğrusal olmayan yapı sergilediği sonucuna ulaşılmaktadır. WTI petrol fiyatlarında herhangi bir değişim meydana gelmeden önce birinci rejimde yaklaşık 62 hafta geçirildiği ve meydana gelen bir değişim ile ikinci rejime geçilerek yaklaşık 8 hafta ikinci rejimde kaldığı bulgusu elde edilmektedir. Aynı rejimde kalma olasılığının yüksek olması ve diğer rejime geçme olasılıklarının düşük olması Markov rejim değişim modelinin güçlü ve uyumlu sonuçlar verdiğinin göstergesi olmaktadır.

Kaynaklar

- Aktaş H., Kayaşlıdere K., Güleç T.C., 2016. Petrol Arz ve Talebindeki Değişimlerin Petrol Fiyatları Üzerindeki Etkisi. 3rd International Congress on Economics and Business. 26-27 Ocak 2016.
- Bildirici M., Alp, A.E., Ersin Ö.Ö., Bozoklu Ü., 2010. İktisatta Kullanılan Doğrusal Olmayan Zaman Serisi Yöntemleri, Türkmen Kitabevi, Yayın No:357, 1.Baskı, 384s, İstanbul, Türkiye.
- Box G.E.P, Jenkins G.M. 1970. Time Series Analysis, Forecasting and Control, San Francisco: Holden-Day.
- Brock W.A., Dechert W., Scheinkman J., 1987. A Test for Independence Based on the Correlation. Dimension Working paper, University of Wisconsin at Madison, University of Houston, and University of Chicago.
- Çil N., Yılmaz Ç., 2018. Markov Switching Autoregressive Model for WTI Crude Oil Price, Ekonometrik ve İstatistik e-Dergisi, 14(28): 45-56.
- Evci S., Şak N., Karaağaç Adana G., 2016. Altın Fiyatlarındaki Değişimin Markov Rejim Değişim Modelleriyle İncelenmesi, Business and Economics Research Journal, 7(4): 67-77.
- Hamilton J.D., 1989. A New Approach to the Economic Analysis of Nonstationary Time Series and The Business Cycle, Econometrica, 57(2): 357-384.
- Koy A., 2017. Spot ve Vadeli Piyasa İlişkilerine Markov Rejim Değişim Yaklaşımı, Bankacılar Dergisi, 101:70-87.

- Krolzig, H.M., 1997. Markov Switching Vector Autoregressions Modelling: Statistical Inference and Application to Business Cycle Analysis, Berlin: Springer.
- Maupin T., 2019. Can Bitcoin, and Other Cryptocurrencies, be Modeled Effectively with a Markov-Switching Approach? Royal Institute of Technology School of Engineering Sciences, Stocholm, Sweden
- Neftçi S., 1984. Are Economic Time Series Asymmetric over the Business Cycles, Journal of Political Economy, 92:307-328.
- Özdemir S., Akgül I., 2015. Ham Petrol ve Benzin Fiyatlarının Sanayi Üretimine Etkisi: MS-VAR Modelleri ile Analizi, Ege Akademik Bakış, 15(3): 367-378.
- Sümer K., Hepsağ A. 2007. Finansal Varlık Modelleri Çerçevesinde Piyasa Risklerinin Hesaplanması: Parametrik Olmayan Yaklaşım, Bankacılar Dergisi, 62: 3–24.
- Terzioğlu, M.K., 2018. Ham Petrol Fiyatları Ve Döviz Kuru: Markov-Geçiş Hata Düzeltme Modeli, Finans Ekonomi ve Sosyal Araştırmalar Dergisi, 3(1): 339-347.
- Tsay R.S., 1989. Nonlinearity Tests for Time Series, Biometrika, 73(2): 461-466.

Some Factors Affecting Gully Erosion in Southwestern Hamadan, West of Iran

Ali Reza VAEZI¹, Hossein SHAMKHANI^{1*}, Saniye DEMIR²

¹Department of Soil Science Engineering, Faculty of Agriculture, University of Zanjan, Zanjan, Iran.

²TOGU University, Faculty of Agriculture, Department of Soil Science, Tokat, 60100, Turkey.

*Presenter author e-mail: hossein.shamkhani@yahoo.com

Abstract

Gully erosion is one of the most important types of water erosion, which causes strongly soil degradation in hillslopes, especially in semi-arid regions. Very little information is available on how to develop this erosion in different land uses such as rainy lands and rangelands. The present study was conducted to investigate the relationships between morphometric features of the hillslopes on gully erosion in a semi-arid region located in the southwest of Hamadan province, west of Iran. A topographic area with 9 km² was selected using the Google Earth, where there are signs of permanent gully erosion. Twenty one grids with dimensions of 400 m × 300 m were considered on the area map and a gully was investigated in each grid. The results showed that the density of gullies in the area varies from 15.36 to 40.19 km km⁻¹. There are significant relationships between gully density and the characteristics of drainage area such as slope gradient (($r= 0.87$ and $p<0.05$)), slope length ($r= 0.95$ and $p<0.05$), Gravelius's coefficient ($r= 0.90$ and $p<0.05$) and Miller's coefficient ($r= 0.63$ and $p<0.05$). The density of gully increases in the area with increasing slope gradient, length, Gravelius's coefficient and Miller's coefficient of drainage area. Hillslope length is the most important morphological factor influencing gully erosion in the area. The application of soil and water conservation practices in the gullies along with using biological methods on uplands can be effective strategies in preventing the expansion of gully erosion in the area.

Key words: Gully Density, Gully Erosion, Hillslope characteristics, Morphometric, Rangelands

Comparison of Wood, Single and Two Knotted Cubic Spline Regression Models for Modeling Lactation Curves of Holstein Cattle in First Lactation

Samet Hasan ABACI^{1*}, Ertugrul KUL²

¹Ondokuz Mayıs University, Faculty of Agriculture, Department of Animal Science, 55139, Samsun, Turkey

²Kırşehir Ahi Evran University, Faculty of Agriculture, Department of Animal Science, 40100, Kırşehir, Turkey

*Presenter author e-mail: shabaci37@gmail.com

Abstract

The lactation curve, which expresses the graphical representation of milk production during the lactation period, can be examined using different mathematical models. With mathematical models, milk yield can be predicted in cattle under different conditions. For this, it is important to choose the most suitable model. This study, it is aimed to compare the Wood model, which is widely used in modeling the lactation curve, and the Cubic Spline Regression model (single knot: CSR1 and two-knot: CSR2), which is useful when the relationship between a response and a set of covariates is not known beforehand. For this purpose, monthly test-day milk yield (TDMY) records (minimum 7, maximum 10 TGSV) of 84 head Holstein cattle raised in a private cattle farm in Kastamonu province İnebolu district were used. The SAS statistical package program was used to estimate the models. In the comparison of the models, the mean squares error (MSE) and coefficient of determination (R^2) values were used. According to the findings obtained for Wood, CRS1 and CRS2 models, the MSE values were 5.05, 3.24 and 2.86, respectively, and the R^2 values were 0.706, 0.860 and 0.915. Accordingly, it was determined that the CSR2 model gave more reliable estimates in modeling the lactation curves of Holstein cattle in the first lactation, due to the lower PCR value and higher R^2 value than the other models used in this study.

Key words: Modeling, Milk yield, Lactation, Holstein

Genel Özellikleri ile İkili (Binary) Logit ve Probit Modellerin İncelenmesi

Sıddık KESKİN¹, Duygu KORKMAZ^{2*}

¹Van Yüzüncü Yıl Üniversitesi, Tıp Fakültesi, Biyoistatistik, 65000, Van, Türkiye

²Van Yüzüncü Yıl Üniversitesi, Tıp Fakültesi, Tıp Eğitimi ve Bilişimi, 65000, Van, Türkiye

*Sunucu yazar e-posta: duyukorkmaz@yyu.edu.tr

Özet

Cevap değişkeninin “evet-hayır”, “katılıyorum-katılmıyorum”, “evli-bekar”, “ölü-sağ” gibi ikili olması durumunda, logit veya probit model yaygın olarak kullanılmaktadır. Özellikle ilaç uygulamalarında, bir ilacın akut toksisitesi, %50'sini öldürecek dozun (ölümcül doz veya letal doz, LD50) hesaplanmasıyla belirlenebilmekte ve letal (öldürücü) doz hesaplamalarında probit veya logit model kullanılmaktadır. Bu modeller, ölümcül zaman veya konsantrasyon hesaplamak için de kullanılabilir. Bu çalışmada, genel özellikleri ile logit ve probit modellere değinilerek, uygulama yapılması ve elde edilen sonuçların yorumlanması amaçlanmıştır. Çalışmada uygulama materyali olarak, R paket programı (ver≥3.3.0) “ecotox” paketi içerisinde yer alan serbest erişimli “lamprey_time” (n=44) ve “lamprey_tox” (n=64) veri setleri kullanılmıştır. Lamprey_time veri seti; 1 saatten 12 saate kadar balık öldürücüye (TFM) maruz kalan lampreyler, toksisite testinin yapıldığı ay (Mayıs, Haziran, Ağustos, Eylül) maruziyet süresi, ortalama TFM dozu, maruziyet süresinde yaşayan ve ölen lamprey sayılarından oluşmuştur. Lamprey_tox veri setinde ise maruziyet süresi yerine tank değişkeni yer almaktadır. Lamprey_tox veri setinden ölümcül konsantrasyon, Lamprey_time veri setinden ise ölümcül zaman hesaplanmıştır. Lamprey tox veri setinden logit modelle hesaplanan Mayıs ayına ait ölümcül konsantrasyon LD₅₀ ve LD₉₉ sırası ile 1.26, 2.24 bulunurken; probit modelle hesaplanan ölümcül konsantrasyon; LD₅₀ ve LD₉₉ sırası ile 1.25, 2.11 bulunmuştur. Lamprey time veri setinden logit modelle hesaplanan Mayıs ayına ait ölümcül zaman ise LD₅₀ ve LD₉₉ sırası ile 9.97, 47.6 bulunurken; probit modelle hesaplanan ölümcül zaman LD₅₀ ve LD₉₉ için sırası ile 9.99, 38.1 bulunmuştur. Lamprey tox veri setinde tüm aylarda ölçülen öldürücü konsantrasyonlar hem probit modelle hem de logit modellerde istatistik olarak anlamlı bulunmuştur (p<0.05). Lamprey time verisinde ise Haziran-Eylül dışındaki zaman karşılaştırmaları istatistik olarak anlamlı bulunmuştur (p<0.05). Logit ve probit modellerin her ikisinin de öldürücü doz ve zamanları tahmin etmede benzer sonuçlar verdiği gözlenmiştir. Alt grup (ay) düzeyinde yapılan karşılaştırmalarda da her iki yöntem birbirleri ile tutarlı sonuçlar vermiştir.

Anahtar kelimeler: LD50, Lojistik dağılım, Bağlantı fonksiyonu, Eklemeli olasılık

Performance Analysis of Electricity Distribution Companies in Turkey with Data Envelopment Analysis and (0,1) Regression Analysis

Ammar Mohammed Salih ALBASRI^{1*}, Coşkun KUŞ¹, Hasan ÖRKÇÜ², Mustafa Ç. KORKMAZ³

¹Selçuk University, Science Faculty, Department of Statistics, Konya, Turkey

²Gazi University, Department of Statistics, Ankara, Turkey

³Artvin Çoruh University Department of Measurement and Evaluation, Artvin, Turkey

* Presenter author e-mail: ammar_mohammedsalih@yahoo.com

Abstract

Electricity distribution companies play a crucial role for both households and industries. In this study, the performance of electricity distribution companies was examined with the data envelopment analysis model at the first stage. Data Envelopment Analysis (DEA) is a non-parametric methodology for evaluating the relative efficiency of decision-making units (DMUs) with common inputs and outputs. In the DEA modeling, the number of staff, length of cables (km), installed power (MVA), energy losses (MWh), and investment amount (TL) are taken as inputs. As output, gross electricity generation (MWh), the urban area served (km²), net consumption (MWh), the number of customers is used to measure of service efficiency of Turkish electric distribution companies.

In the second stage, the efficiency score, that is, the performance of the companies is accepted as a dependent variable, and the effect of external variables affecting performance is examined by beta and Kumaraswamy regression analysis. In addition, a new quantile type regression analysis is also proposed, and it is applied to our data besides beta and Kumaraswamy regression analysis. At this stage, the effects of the variables of the number of villages and towns, number of poles, distribution company region population, annual breakdown and interruption, number of transformers, and the number of fixtures and lamps on the efficiency score were examined. The data used in the study belong to 2019. The results show that 14 of 21 are efficient companies, and also, the effect of all of the external variables on the scores is statistically significant.

Key words: Data envelopment analysis, Electricity distribution companies, Beta regression, Kumaraswamy regression

Çoklu Bağlantı Probleminde Alternatif Tahmin Edicilerin Performanslarının Karşılaştırılması

Fahreddin KALKAN^{1*}, Aydın KARAKOCA²

¹Selçuk Üniversitesi, Fen Fakültesi, Aktüerya Bilimleri Bölümü, 42130, Konya, Türkiye

²Necmettin Erbakan Üniversitesi, Fen Fakültesi, İstatistik Bölümü, 42090, Konya, Türkiye

*Sunucu yazar e-posta: fahreddin.kalkan@selcuk.edu.tr

Özet

Çoklu doğrusal regresyon analizinde en önemli varsayımlardan biri bağımsız değişkenlerin ilişkisiz olmasıdır. Bu varsayımın bozulması durumunda çoklubağlantı problemi ortaya çıkar. Çoklubağlantı probleminin çözümü için farklı yöntemler önerilmiş olsa da bunlardan en popüler olanı Ridge Regresyon yöntemidir. Bu çalışmada literatürde önerilen, 50 farklı ridge tahmin edicisinin hata kareler ortalaması açısından performansları detaylı olarak incelenmiştir. Çalışmada örneklem büyüklükleri, korelasyon düzeyleri, hata varyansları ve bağımsız değişken sayısı için belirlenen farklı değerler için hata kareler ortalaması kriterine göre karşılaştırmalar yapılmıştır. Ayrıca bu tahmin edicilerin performansları, simülasyon çalışmalarında kullanılan parametre seçim yöntemlerine göre de ayrıca karşılaştırılmıştır. Elde edilen Monte Carlo simülasyon sonuçlarına göre incelenen her durum için en iyi performansı gösteren tahmin edicilerin değiştiği dolayısıyla her veri seti için doğru ridge tahmin edicisinin seçiminin önemine ilişkin sonuçlar incelenmiş ve öneriler sunulmuştur.

Anahtar kelimeler: Çoklubağlantı, Çoklu doğrusal regresyon, Ridge regresyon, Ridge tahmin edicisi, Monte Carlo simülasyonu

GİRİŞ

Çoklu doğrusal regresyon modeli bir bağımlı değişken ile iki ya da daha fazla bağımsız değişken arasında mevcut olan doğrusal ilişkiyi matematiksel olarak modellemek ve özellikle ileriye dönük öngörü ve tahminlerde bulunmak amacıyla birçok araştırmacı tarafından kullanılan istatistiksel bir yöntemdir. Çoklu doğrusal regresyon analizinin uygulanabilmesi için gerekli varyasyonların sağlanması gerekir. Bu varsayımlardan biri de bağımsız değişkenlerin birbiriyle ilişkisiz olmasıdır. Bağımsız değişkenler arasında doğrusal bir ilişki yoksa bağımsız değişkenler ortogondur. Bağımsız değişkenler ortogonal olduğunda çıkarımlar nispeten kolay bir şekilde yapılabilir ancak çoğu regresyon uygulamalarında bağımsız değişkenler ortogonal değildir. Bağımsız değişkenler arasında neredeyse doğrusal bağımlılıklar olduğunda, çoklubağlantı probleminin var olduğu söylenir. Bu durumda, regresyon modeline dayalı çıkarımlar yanıltıcı veya hatalı olabilir (Montgomery ve Peck, 1991). Literatürde çoklubağlantı problemini ortadan kaldırmak için önerilen birçok yöntem mevcuttur. Bu yöntemlerden en popüler olanı Hoerl ve Kennard (1970a, 1970b) tarafından önerilen Ridge Regresyon yöntemidir.

Çoklu doğrusal regresyon denkleminin matris notasyonuyla gösterimi;

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\varepsilon} \quad (1.1)$$

Burada $\mathbf{y}$, $n \times 1$ boyutlu bağımlı değişken vektörü, $\mathbf{X}$, $n \times p$ boyutlu stokastik olmayan açıklayıcı değişkenlerin gözlem matrisi, $\boldsymbol{\beta}$, $p \times 1$ boyutlu bilinmeyen parametreler vektörü ve $\boldsymbol{\varepsilon}$ sıfır ortalamalı, σ^2 varyanslı, $n \times 1$ boyutlu hata vektörüdür. $\boldsymbol{\beta}'$ 'nin en küçük kareler tahmini;

$$\hat{\boldsymbol{\beta}} = (\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{y} \quad (1.2)$$

eşitliği ile elde edilir. Ridge regresyon yöntemi $X'X$ korelasyon matrisinin köşegen elemanlarına küçük bir yanlılık sabiti ekleyerek, regresyon katsayılarının varyans ve kovaryanslarını azaltmaya ve hata kareler ortalamasını (MSE) minimum yapmayı hedefler. Hoerl ve Kennard (1970) tarafından önerilen;

$$\hat{\beta}(k) = (X'X + kI_p)^{-1}X'y, \quad k > 0 \quad (1.3)$$

eşitliği β^* 'nın ridge tahmin edicisidir. Burada, k sabiti ridge parametresi veya yanlılık parametresi olarak bilinir (Kibria, 2003). Ridge tahmin edicileri modele ait parametreleri yanlı tahmin etse de, varyans ve kovaryansı azaltmakla birlikte en küçük kareler yönteminin hata kareler ortalamasından daha küçük hata kareler ortalamasına sahiptir. Ridge tahmininde kritik nokta hata kareler ortalamasını minimum yapan k değerini belirlemektir. k 'yı tahmin etmek için literatürde önerilen bir çok ridge tahmin edicisi mevcuttur. Literatür de Hoerl ve Kennard (1970a), Hoerl ve Kennard (1970b), Hoerl ve ark. (1975), Lawless ve Wang (1976), Hocking ve ark. (1976), Schaeffer ve ark.(1984), Nomura ve Masuo (1988), Kibria (2003), Khalaf ve Shukur (2005), Alkhamisi ve ark. (2006), Alkhamisi ve Shukur (2007), Muniz ve Kibria (2009), Kibria ve ark. (2011), G.Khalaf (2012), Muniz ve ark. (2012), Khalaf (2013) ve Khalaf ve ark. (2013) tarafından farklı tahmin ediciler önerilmiştir.

Bu çalışmada literatürde öne çıkan 50 farklı ridge tahmin edicisinin MSE yönünden seçilen farklı parametre değerlerine göre performansları incelenmiştir. Çalışmada örneklem büyüklükleri, korelasyon düzeyleri, hata varyansları, bağımsız değişken sayısı ve simülasyon çalışmalarında parametre değerlerinin seçim yöntemleri MSE'yi etkileyen faktörler olarak ele alınmıştır.

BAZI RIDGE TAHMİN EDİCİLERİ

Bu bölümde çalışmada incelenen ve literatür de dikkat çeken Ridge tahmin edicilerinden seçilen 50 tanesi Çizelge 1'de verilmiştir.

Çizelge 1. Literatürde önerilen bazı ridge tahmin edicileri

$k_1 = \frac{\hat{\sigma}^2}{\hat{\alpha}_i^2}$	Hoerl ve Kennard (1970a)
$k_2 = \frac{\hat{\sigma}^2}{\hat{\alpha}_{max}^2}$	Hoerl ve Kennard (1970b)
$k_3 = \frac{p\hat{\sigma}^2}{\hat{\beta}'\hat{\beta}}$	Hoerl, Kennard ve Baldwin (1975)
$k_4 = \frac{p\hat{\sigma}^2}{\hat{\alpha}'X'X\hat{\alpha}}$	Lawless ve Wang (1976)
$k_5 = \hat{\sigma}^2 \frac{\sum_{i=1}^p (\lambda_i \hat{\alpha}_i)^2}{(\sum_{i=1}^p \lambda_i \hat{\alpha}_i^2)^2}$	Hocking, Speed ve Lynn (1976)
$k_6 = \frac{1}{\hat{\alpha}_{max}^2}$	Schaeffer, Roi ve Wolfe (1984)

Çizelge 1. devamı

$k_7 = \frac{p\hat{\sigma}^2}{\sum_{i=1}^p \left[\frac{\hat{\alpha}_i^2}{1 + \sqrt{1 + \lambda_i \left(\frac{\hat{\alpha}_i^2}{\hat{\sigma}^2} \right)}} \right]}$	Nomura ve Masuo (1988)
$k_8 = \sum_{i=1}^p \frac{\hat{\sigma}^2}{\hat{\alpha}_i^2}$	Kibria (2003)
$k_9 = \frac{\hat{\sigma}^2}{(\prod_{i=1}^p \hat{\alpha}_i^2)^{1/p}}$	Kibria (2003)
$k_{10} = \text{Median} \left\{ \frac{\hat{\sigma}^2}{\hat{\alpha}_i^2} \right\} i = 1, 2, \dots, p$	Kibria (2003)
$k_{11} = \frac{\lambda_{\max} \hat{\sigma}^2}{(n - p - 1) \hat{\sigma}^2 + \lambda_{\max} \hat{\beta}_{\max}^2}$	Khalaf ve Shukur (2005)
$k_{12} = \max \left(\frac{\hat{\sigma}^2}{\hat{\beta}_i^2} \right), i = 1, 2, \dots, p$	Alkhamisi ve ark. (2006)
$k_{13} = \frac{1}{p} \sum_{i=1}^p \left(\frac{\lambda_i \hat{\sigma}^2}{(n - p) \hat{\sigma}^2 + \lambda_i \hat{\beta}_i^2} \right)$	Alkhamisi ve ark. (2006)
$k_{14} = \max \left(\frac{\lambda_i \hat{\sigma}^2}{(n - p) \hat{\sigma}^2 + \lambda_i \hat{\beta}_i^2} \right)$	Alkhamisi ve ark. (2006)
$k_{15} = \hat{k}_{md}^{KS} = \text{median} \left(\frac{\lambda_i \hat{\sigma}^2}{(n - p) \hat{\sigma}^2 + \lambda_i \hat{\beta}_i^2} \right)$	Alkhamisi ve ark. (2006)
$k_{16} = k_3 + \frac{1}{\lambda_{\max}}$	Alkhamisi ve Shukur (2007)
$k_{17} = \max \left(\frac{\hat{\sigma}^2}{\hat{\beta}_i^2} + \frac{1}{\lambda_i} \right)$	Alkhamisi ve Shukur (2007)
$k_{18} = \frac{1}{p} \sum_{i=1}^p \left(\frac{\hat{\sigma}^2}{\hat{\beta}_i^2} + \frac{1}{\lambda_i} \right)$	Alkhamisi ve Shukur (2007)
$k_{19} = \text{median} \left(\frac{\hat{\sigma}^2}{\hat{\beta}_i^2} + \frac{1}{\lambda_i} \right)$	Alkhamisi ve Shukur (2007)
$k_{20} = \frac{p\hat{\sigma}^2}{\sum_{i=1}^p \lambda_i \hat{\beta}_i^2} + \frac{1}{\lambda_{\max}}$	Alkhamisi ve Shukur (2007)
$k_{21} = \left(\prod_{i=1}^p \frac{\lambda_i \hat{\sigma}_i^2}{(n - p) \hat{\sigma}_i^2 + \lambda_i \hat{\alpha}_i^2} \right)^{1/p}$	Muniz ve Kibria (2009)

Çizelge 1. devamı

$k_{22} = \max \left(\frac{1}{\sqrt{\frac{\hat{\sigma}_i^2}{\hat{\alpha}_i^2}}} \right)$	Muniz ve Kibria (2009)
$k_{23} = \max \left(\sqrt{\frac{\hat{\sigma}_i^2}{\hat{\alpha}_i^2}} \right)$	Muniz ve Kibria (2009)
$k_{24} = \left(\prod_{i=1}^p \frac{1}{\sqrt{\frac{\hat{\sigma}_i^2}{\hat{\alpha}_i^2}}} \right)^{\frac{1}{p}}$	Muniz ve Kibria (2009)
$k_{25} = \left(\prod_{i=1}^p \sqrt{\frac{\hat{\sigma}_i^2}{\hat{\alpha}_i^2}} \right)^{\frac{1}{p}}$	Muniz ve Kibria (2009)
$k_{26} = \text{median} \left(\frac{1}{\sqrt{\frac{\hat{\sigma}_i^2}{\hat{\alpha}_i^2}}} \right)$	Muniz ve Kibria (2009)
$k_{27} = \text{median} \left(\sqrt{\frac{\hat{\sigma}_i^2}{\hat{\alpha}_i^2}} \right)$	Muniz ve Kibria (2009)
$k_{28} = \max \left(\frac{1}{\frac{\lambda_{\max}}{(n-p)\hat{\sigma}^2 + \lambda_{\max}\hat{\alpha}_i^2}} \right)$	Kibria, Mansson ve Shukur (2011)
$k_{29} = \max \left(\frac{\lambda_{\max}}{(n-p)\hat{\sigma}^2 + \lambda_{\max}\hat{\alpha}_i^2} \right)$	Kibria, Mansson ve Shukur (2011)
$k_{30} = \prod_{i=1}^p \left(\frac{1}{\frac{\lambda_{\max}}{(n-p)\hat{\sigma}^2 + \lambda_{\max}\hat{\alpha}_i^2}} \right)^{\frac{1}{p}}$	Kibria, Mansson ve Shukur (2011)

Çizelge 1. devamı

$k_{31} = \prod_{i=1}^p \left(\frac{\lambda_{max}}{(n-p)\hat{\sigma}^2 + \lambda_{max}\hat{\alpha}_i^2} \right)^{\frac{1}{p}}$	Kibria, Mansson ve Shukur (2011)
$k_{32} = median \left(\frac{1}{\frac{\lambda_{max}}{(n-p)\hat{\sigma}^2 + \lambda_{max}\hat{\alpha}_i^2}} \right)$	Kibria, Mansson ve Shukur (2011)
$k_{33} = median \left(\frac{\lambda_{max}}{(n-p)\hat{\sigma}^2 + \lambda_{max}\hat{\alpha}_i^2} \right)$	Kibria, Mansson ve Shukur (2011)
$k_{34} = \frac{\hat{\sigma}^{*2}}{\hat{\beta}_i^2}, \hat{\sigma}^{*2} = \frac{RRS}{n-p+2}$	G.Khalaf (2012)
$k_{35} = \frac{p\hat{\sigma}^{*2}}{\hat{\beta}'\hat{\beta}}, \hat{\sigma}^{*2} = \frac{RRS}{n-p+2}$	G.Khalaf (2012)
$k_{36} = max \left(\frac{1}{\sqrt{\frac{\lambda_{max}\hat{\sigma}^2}{(n-p)\hat{\sigma}^2 + \lambda_{max}\hat{\alpha}_i^2}}} \right)$	Muniz, Kibria ve Shukur (2012)
$k_{37} = max \left(\sqrt{\frac{\lambda_{max}\hat{\sigma}^2}{(n-p)\hat{\sigma}^2 + \lambda_{max}\hat{\alpha}_i^2}} \right)$	Muniz, Kibria ve Shukur (2012)
$k_{38} = \left(\prod_{i=1}^p \frac{1}{\sqrt{\frac{\lambda_{max}\hat{\sigma}^2}{(n-p)\hat{\sigma}^2 + \lambda_{max}\hat{\alpha}_i^2}}} \right)^{\frac{1}{p}}$	Muniz, Kibria ve Shukur (2012)
$k_{39} = \left(\prod_{i=1}^p \sqrt{\frac{\lambda_{max}\hat{\sigma}^2}{(n-p)\hat{\sigma}^2 + \lambda_{max}\hat{\alpha}_i^2}} \right)^{\frac{1}{p}}$	Muniz, Kibria ve Shukur (2012)
$k_{40} = median \left(\frac{1}{\sqrt{\frac{\lambda_{max}\hat{\sigma}^2}{(n-p)\hat{\sigma}^2 + \lambda_{max}\hat{\alpha}_i^2}}} \right)$	Muniz, Kibria ve Shukur (2012)
$k_{41} = \frac{(\lambda_{max} + \lambda_{min})}{2 \sum_{i=1}^p \hat{\beta}_i } \frac{p\hat{\sigma}^2}{\sum_{i=1}^p \hat{\beta}_i^2}$	Khalaf (2013)

Çizelge 1. devamı

$k_{42} = \frac{(\lambda_{max} + \lambda_{min})\lambda_{max}\hat{\sigma}^2}{2 \sum_{i=1}^p \hat{\beta}_i ((n-p)\hat{\sigma}^2 + \lambda_{max}\hat{\beta}_{max}^2)}$	Khalaf (2013)
$k_{43} = k_2 * \frac{\lambda_{max}}{\sum_{j=1}^p \hat{\alpha} _j}$	Khalaf, Mansson ve Shukur (2013)
$k_{44} = k_9 * \frac{\lambda_{max}}{\sum_{j=1}^p \hat{\alpha} _j}$	Khalaf, Mansson ve Shukur (2013)
$k_{45} = k_{10} * \frac{\lambda_{max}}{\sum_{j=1}^p \hat{\alpha} _j}$	Khalaf, Mansson ve Shukur (2013)
$k_{46} = k_{11} * \frac{\lambda_{max}}{\sum_{j=1}^p \hat{\alpha} _j}$	Khalaf, Mansson ve Shukur (2013)
$k_{47} = k_{14} * \frac{\lambda_{max}}{\sum_{j=1}^p \hat{\alpha} _j}$	Khalaf, Mansson ve Shukur (2013)
$k_{48} = k_{15} * \frac{\lambda_{max}}{\sum_{j=1}^p \hat{\alpha} _j}$	Khalaf, Mansson ve Shukur (2013)
$k_{49} = k_{21} * \frac{\lambda_{max}}{\sum_{j=1}^p \hat{\alpha} _j}$	Khalaf, Mansson ve Shukur (2013)
$k_{50} = k_{22} * \frac{\lambda_{max}}{\sum_{j=1}^p \hat{\alpha} _j}$	Khalaf, Mansson ve Shukur (2013)

SİMÜLASYON ÇALIŞMASI

Simülasyon çalışmasında, bir çok çalışmada standart veri üretim tekniği olarak kullanılan McDonald ve Galarneau (1975) tarafından önerilen (3.1) eşitliği yardımıyla ilişkili yapıda bağımsız değişkenler üretilmiştir.

$$x_{ij} = (1 - \gamma^2)^{1/2} z_{ij} + \gamma z_{ip}, \quad i = 1, 2, \dots, n, \quad j = 1, 2, \dots, p \quad (3.1)$$

Burada, z_{ij} bağımsız standart normal dağılan bir rastgele vektör ve γ^2 herhangi iki bağımsız değişken arasındaki korelasyon katsayısıdır (Kibria, 2003). Simülasyon çalışması bağımsız değişkenler arasında seçilen $\rho = 0.90, 0.95$ ve 0.99 korelasyon değerleri, $n = 30, 50$ ve 100 örneklem büyüklükleri, $\sigma^2 = 0.25, 0.50, 1, 2$ ve 4 hata varyans değerleri, $p = 4$ ve 8 bağımsız değişken sayısı ve model parametrelerinin seçimi için 2 farklı yöntemle tahmin edicilerin MSE performanslarını karşılaştırmak üzere tasarlanmıştır. Newhouse and Woman (1971) β 'nin $X'X$ matrisinin en büyük özdeğerine karşılık gelen normalleştirilmiş özvektörler olduğunda hata kareler ortalamasının en aza indirildiğini belirtmiştir. Simülasyon çalışmasında parametre seçiminde Newhouse and Woman (1971) tarafından önerilen yöntemi ve β 'nin manuel olarak girildiği yöntem karşılaştırılmıştır. Newhouse and Woman (1971)'in önerdiği yöntemle kanonik, diğer yöntemle ise standartlaştırılmamış ismi verilmiştir. Bağımlı değişken y belirlenen parametre seçim yöntemine göre oluşturulan regresyon modelinde ilgili hata varyansı ve oluşturulan tasarım matrisi kullanılarak elde edilmiştir. Ayrıca tahmin edicilerin değişkenliğine bakmak amacıyla tahmin edicilere ait standart hatalarda hesaplanmıştır. Tahmin edicilerin MSE değeri,

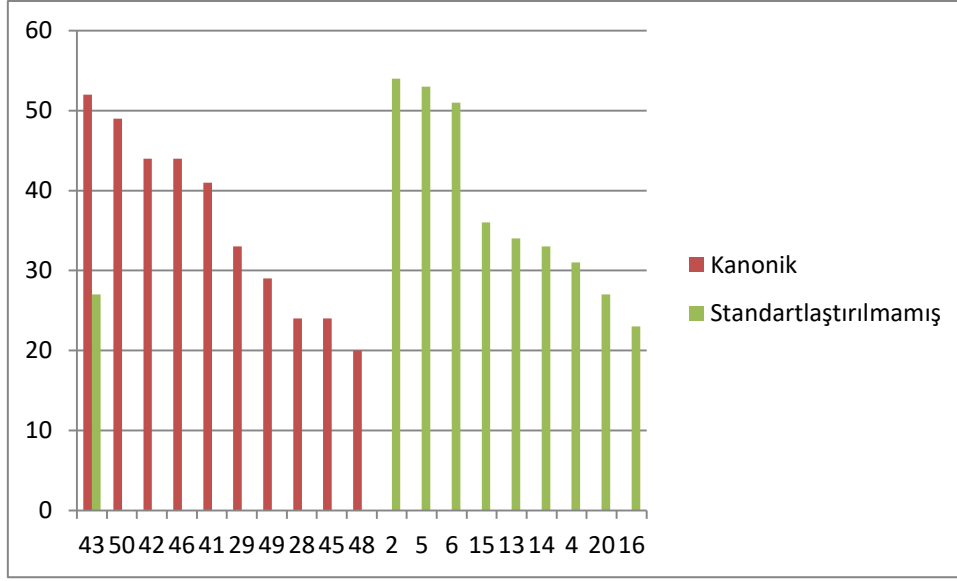
$$MSE[\hat{\beta}(k)] = \frac{1}{5000} \sum_{r=1}^{5000} (\hat{\beta}_r - \beta)' (\hat{\beta}_r - \beta) \quad (3.2)$$

eşitliği ile değerlendirilmiştir.

Simülasyon çalışmasında örneklem büyüklüğü için 3, korelasyon değerleri için 3, hata varyansı için 5, bağımsız değişken sayısı için 2 ve model parametresi seçimi için 2 farklı yöntemin incelenmesinden kaynaklanan toplam $3 \times 3 \times 5 \times 2 \times 2 = 180$ farklı durum incelenmiştir. İncelenen tahmin edicilerin her bir durum için MSE performanslarına göre en küçük ilk 5 MSE değerine sahip tahmin edici değerlendirmeye alınmıştır. Bu değerlendirmelerde kullanılan simülasyon sonuçlarından örnek bir durum Çizelge 2’de verilmiştir. Çizelge sayısı fazla olduğu için, sonuçların değerlendirmeleri Şekil 1-5 üzerinden yapılmıştır.

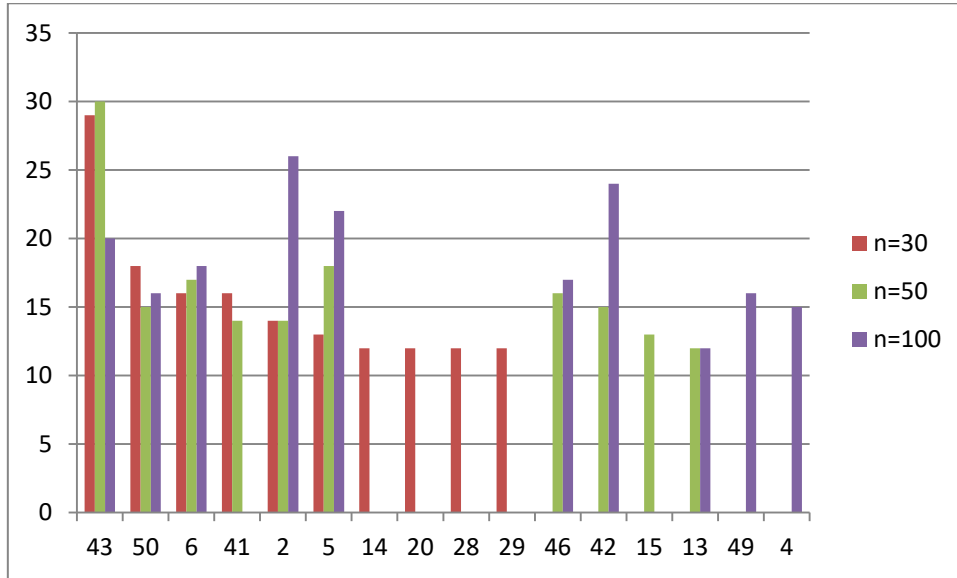
Çizelge 2. Kanonik parametre seçimi için MSE ve tahmin edicilere ait s.s. değerleri ($n = 30, p = 4$)

$n = 30, p = 4$														
$\rho = 0.90$														
$\sigma^2 = 0.25$			$\sigma^2 = 0.5$			$\sigma^2 = 1$			$\sigma^2 = 2$			$\sigma^2 = 4$		
MSE	k	σ	MSE	k	σ	MSE	k	σ	MSE	k	σ	MSE	k	σ
0.026	29	6.456	0.055	43	14.028	0.105	46	35.051	0.078	50	10.459	0.091	50	7.801
0.027	43	6.824	0.056	49	2.158	0.107	43	29.385	0.162	46	42.709	0.276	47	73.703
0.029	49	1.811	0.057	42	8.744	0.107	50	15.355	0.184	41	74.266	0.298	46	46.830
0.030	48	1.832	0.057	29	3.228	0.111	42	13.408	0.192	43	57.665	0.301	41	128.079
0.032	42	5.227	0.067	48	1.878	0.120	49	2.454	0.226	47	80.960	0.361	43	103.045
$\rho = 0.95$														
$\sigma^2 = 0.25$			$\sigma^2 = 0.5$			$\sigma^2 = 1$			$\sigma^2 = 2$			$\sigma^2 = 4$		
MSE	k	σ	MSE	k	σ	MSE	k	σ	MSE	k	σ	MSE	k	σ
0.026	29	6.874	0.059	43	13.681	0.103	50	14.619	0.076	50	10.940	0.078	50	8.310
0.026	43	6.633	0.060	41	18.705	0.115	41	33.489	0.226	28	4.839	0.199	28	9.765
0.030	41	9.448	0.066	46	20.424	0.122	46	26.477	0.241	41	54.537	0.279	45	60705.735
0.038	46	14.277	0.071	29	3.437	0.133	43	26.214	0.242	47	67.956	0.292	44	4289.413
0.047	42	4.787	0.097	42	7.660	0.181	47	75.553	0.248	45	65430.098	0.442	17	88729755.094
$\rho = 0.99$														
$\sigma^2 = 0.25$			$\sigma^2 = 0.5$			$\sigma^2 = 1$			$\sigma^2 = 2$			$\sigma^2 = 4$		
MSE	k	σ	MSE	k	σ	MSE	k	σ	MSE	k	σ	MSE	k	σ
0.014	19	72.131	0.023	19	41.440	0.034	28	13.852	0.059	28	27.728	0.071	50	8.956
0.016	36	1.596	0.027	28	6.896	0.039	19	38.735	0.072	19	91.973	0.092	30	4.695
0.016	22	1.608	0.032	36	1.687	0.065	36	1.719	0.084	50	12.241	0.117	28	55.466
0.027	28	3.385	0.033	22	1.702	0.067	22	1.737	0.102	30	2.609	0.123	32	9.744
0.036	40	0.621	0.068	18	74420504.070	0.084	18	138303.051	0.109	18	829182.111	0.139	19	96.612



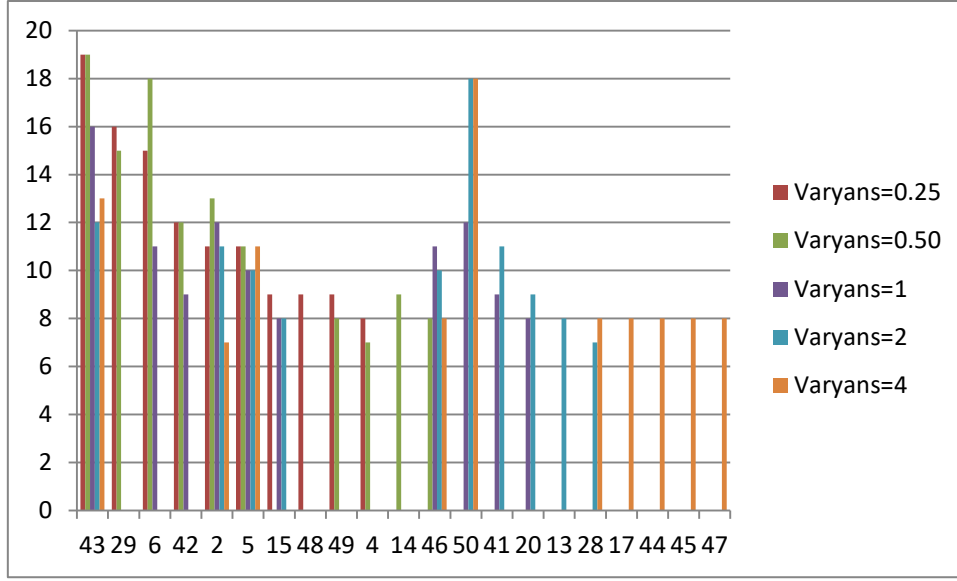
Şekil 1. Parametre seçim yöntemine göre tahmin edicilerin performansları

Farklı parametre seçim yöntemine göre ridge tahmin edicilerinin performansları incelenen 180 farklı durumda en çok ilk 5'e giren tahmin edicilerden en yüksek frekansa sahip ilk 10 tahmin edici için Şekil 1'de verilmiştir. Simülasyonda kullanılan 2 farklı parametre seçim yöntemine göre sadece k_{43} tahmin edicisinin her iki durumda da ilk 10 içerisinde yer aldığı gözlemlenmiştir. Kanonik parametre seçimi için k_{43} , standartlaştırılmamış parametre seçimi için k_2 tahmin edicisi en çok ilk 5'e giren tahmin ediciler arasında yer almıştır. Dolayısıyla tahmin edicilerin parametre seçim yönteminden etkilendiği görülmektedir.



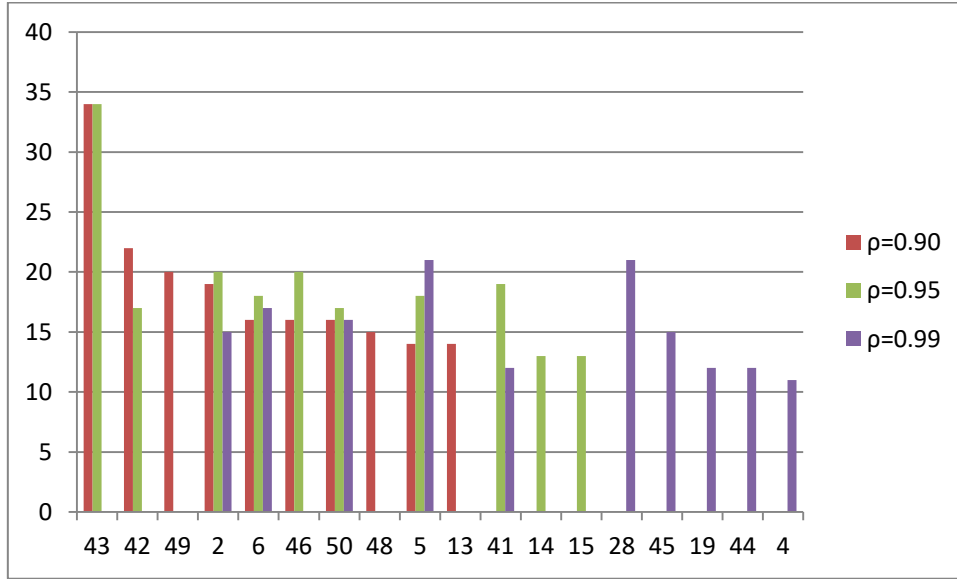
Şekil 2. Farklı örneklem büyüklükleri için tahmin edicilerin performansları

Farklı örneklem büyüklüklerine göre tahmin edicilerin performanslarını karşılaştırmak amacıyla 3 farklı örneklem büyüklüğü için en çok ilk 5'e giren ilk 10 tahmin edici Şekil 2'de verilmiştir. Gözlem sayıları $n=30$ ve $n=50$ için k_{43} , $n=100$ için k_2 tahmin edicisi en iyi performansı göstermiştir. İncelen tahmin edicilerin gözlem sayısından etkilendiği Şekil 2'den görülmektedir.



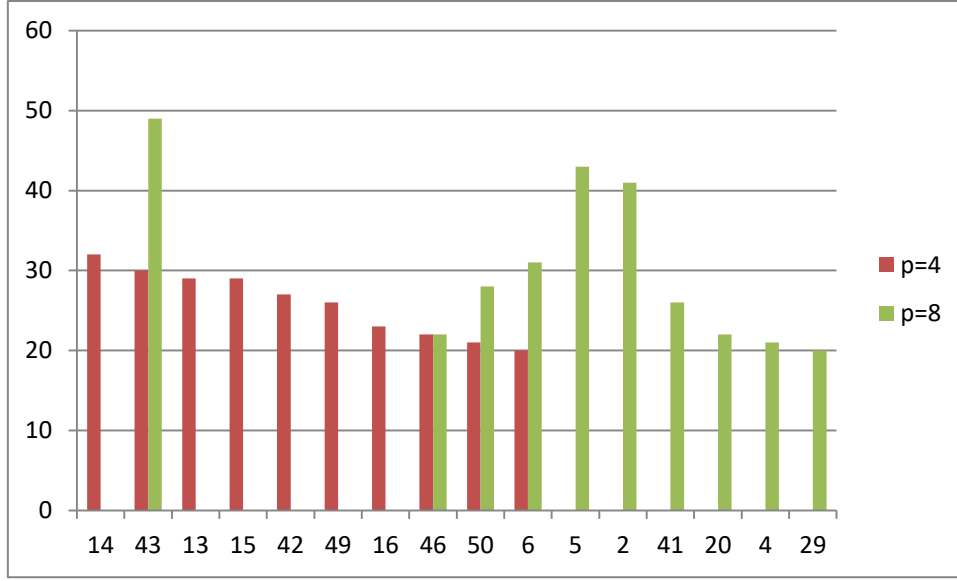
Şekil 3. Farklı varyans değerleri için tahmin edicilerin performansları

Varyans değerleri için tahmin edicilerin performansları Şekil 3’de verilmiştir. Şekil 3 incelendiğinde düşük varyans değerleri için ($\sigma^2 = 0.25, 0.50$ ve 1) k_{43} , yüksek varyans değerleri için ($\sigma^2 = 2$ ve 4) k_{50} tahmin edicisi en iyi performansı göstermiştir. Hata varyanslarının değişimine göre tahmin edicilerin performansının etkilendiği Şekil 3’den görülmektedir.



Şekil 4. Farklı korelasyon düzeyleri için tahmin edicilerin performansları

Bağımsız değişkenler arasındaki korelasyon düzeyine göre tahmin edicilerin performansları Şekil 4’te verilmiştir. Şekil 4 incelendiğinde $\rho = 0.90$ ve 0.95 için k_{43} , $\rho = 0.99$ için k_{28} tahmin edicisinin en iyi performansı gösterdiği gözlenmiştir. İncelen tahmin edicilerin korelasyon düzeyinden etkilendiği Şekil 4’den görülmektedir.



Şekil 5. Farklı değişken sayıları için tahmin edicilerin performansları

Farklı değişken sayıları için tahmin edicilerin performansı Şekil 5’te verilmiştir. Şekil 5 incelendiğinde $p = 4$ için k_{14} , $p = 8$ için k_{43} tahmin edicisinin en iyi performansı gösterdiği saptanmıştır. İncelen tahmin edicilerin değişken sayısından etkilendiği Şekil 5’den görülmektedir.

TARTIŞMA VE SONUÇ

Çalışma sonucunda literatürde incelenen ridge tahmin edicilerinin, örneklem büyüklüğü, bağımsız değişken sayısı, korelasyon düzeyi, hata varyansı ve model parametrelerinin yapısına göre incelenen her durumdan etkilendikleri gözlenmiştir. Sonuç olarak literatürde birçok ridge tahmin edicisi olmasına karşın her durumda en iyi sonucu veren bir tahmin edicinin gözlenmemesi bu alandaki çalışmaların devam edeceğini göstermektedir.

Kanonik parametre seçimi için örneklem sayısı arttıkça k_{17} , k_{18} , k_{19} , k_{22} ve k_{36} tahmin edicilerinin hata kareler ortalamalarının arttığı gözlemlenirken, k_{47} , k_{48} ve k_{49} tahmin edicilerinin hata kareler ortalamalarının azaldığı gözlemlenmiştir. Ayrıca k_{28} , k_{29} , k_{41} , k_{42} , k_{43} , k_{44} , k_{45} , k_{46} ve k_{50} tahmin edicilerinin örneklem sayısı arttığında hata kareler ortalamalarının gösterdiği tepkinin düzensiz olduğu gözlemlenmiştir. Kanonik parametre seçimi için örneklem sayısı arttığında tahmin edicilere ait standart sapmalarındaki değişim incelendiğinde ise k_{17} , k_{41} , k_{42} , k_{43} , k_{44} , k_{46} , k_{47} , k_{48} , k_{49} ve k_{50} tahmin edicilerinin standart sapmalarının arttığı saptanırken k_{22} , k_{28} ve k_{36} tahmin edicilerinin standart sapmalarının azaldığı saptanmıştır. Örneklem sayısı arttığında standart sapmalarının gösterdiği tepkinin düzensiz olduğu tahmin ediciler ise k_{18} , k_{19} , k_{29} ve k_{45} tahmin edicileri olmuştur. Kanonik parametre seçimi için varyans arttıkça k_7 , k_{17} , k_{18} , k_{19} , k_{22} , k_{29} , k_{36} , k_{41} , k_{42} , k_{43} , k_{46} , k_{47} , k_{48} ve k_{49} tahmin edicilerin hata kareler ortalamaları artarken, k_{30} tahmin edicisinin hata kareler ortalaması azaldığı gözlemlenmiştir. Buna ilaveten varyans arttığında k_{28} , k_{44} , k_{45} ve k_{50} tahmin edicilerinin hata kareler ortalamalarının verdiği tepkinin düzensiz olduğu saptanmıştır. Ayrıca varyans arttığında tahmin edicilerin değişkenliğini incelediğimizde ise k_7 , k_{22} , k_{28} , k_{30} , k_{36} , k_{41} , k_{42} ve k_{43} tahmin edicilerinin değişkenliğinin arttığı, k_{29} , k_{45} , k_{47} , k_{50} tahmin edicilerinin değişkenliğinin ise azaldığı saptanmıştır. Varyansdaki değişime karşılık k_{17} , k_{18} , k_{19} , k_{44} , k_{46} , k_{48} ve k_{49} tahmin edicilerin değişkenliğinin düzensiz olduğu gözlemlenmiştir. Kanonik parametre seçimi için bağımsız değişkenler arasındaki korelasyon arttıkça k_{47} , k_{48} ve k_{49} tahmin edicilerinin hata kareler ortalamaları artarken, k_{17} , k_{28} ve k_{44} tahmin edicilerinin hata kareler ortalamalarının azaldığı gözlemlenmiştir. Ayrıca korelasyon arttıkça k_{17} , k_{28} ve k_{44} tahmin edicilerinin hata kareler ortalamalarının gösterdiği tepkinin düzensiz

olduğu gözlemlenmiştir. Korelasyon arttıkça k_{28} ve k_{29} tahmin edicilerine ait standart sapmalarının arttığı, k_{41} , k_{44} , k_{45} , k_{46} , k_{47} , k_{48} ve k_{49} tahmin edicilerinin standart sapmalarının ise azaldığı saptanmıştır. k_{17} , k_{42} , k_{43} ve k_{50} tahmin edicilerinin standart sapmalarının gösterdiği tepkinin ise düzensiz olduğu gözlemlenmiştir.

Kanonik parametre seçimi için bağımsız değişken sayısı arttığında k_{19} , k_{22} , k_{30} ve k_{36} tahmin edicilerinin hata kareler ortalaması artarken, k_{17} , k_{18} , k_{28} , k_{29} , k_{41} , k_{42} , k_{43} , k_{44} , k_{45} , k_{46} , k_{47} , k_{48} , k_{49} ve k_{50} tahmin edicilerinin hata kareler ortalamalarının gösterdiği tepkinin düzensiz olduğu gözlemlenmiştir. Ayrıca değişken sayısı arttığında k_{22} , k_{28} , k_{36} , k_{47} , k_{48} ve k_{50} tahmin edicilerine ait standart sapmalarının arttığı, k_{19} , k_{30} , k_{44} ve k_{45} tahmin edicilerine ait standart sapmalarının azaldığı görülmüştür. Buna ek olarak değişken sayısı arttığında k_{17} , k_{18} , k_{29} , k_{41} , k_{42} , k_{43} , k_{46} ve k_{49} tahmin edicilerine ait standart sapmalarının gösterdiği tepkinin düzensiz olduğu saptanmıştır.

Parametre seçiminin standartlaştırılmamış olması durumunda gözlem sayısı arttıkça tüm tahmin edicilerin hata kareler ortalamasının azaldığı gözlemlenmiştir. Gözlem sayısı arttığında k_{27} , k_{39} , k_{47} , k_{48} ve k_{49} tahmin edicilerine ait standart sapmaların arttığı gözlemlenirken, k_2 , k_3 , k_5 , k_6 , k_{13} , k_{15} , k_{16} , k_{20} , k_{21} , k_{34} , k_{35} ve k_{43} tahmin edicilerine ait standart sapmaların azaldığı gözlemlenmiştir. Ayrıca gözlem sayısı arttığında k_4 , k_{14} ve k_{33} tahmin edicilerinin standart sapmalarının gösterdiği tepkinin düzensiz olduğu saptanmıştır. Parametre seçimi standartlaştırılmamış için varyans arttıkça tüm tahmin edicilerin hata kareler ortalamasının arttığı gözlemlenmiştir. k_{15} tahmin edicisine ait standart sapmanın gösterdiği tepkinin düzensiz olduğu gözlemlenirken, diğer tahmin edicilere ait standart sapmanın arttığı gözlemlenmiştir. Bağımsız değişkenler arasındaki korelasyon arttıkça tüm tahmin edicilerin hata kareler ortalamasının arttığı gözlemlenmiştir. Parametre seçimi standartlaştırılmamış için korelasyon arttıkça, k_2 , k_3 , k_5 , k_6 , k_{20} , k_{33} , k_{34} , k_{35} ve k_{43} tahmin edicilerin değişkenliğinin arttığı, k_{13} ve k_{21} tahmin edicilerin değişkenliğinin azaldığı saptanmıştır. Ayrıca k_4 , k_{14} , k_{15} ve k_{16} tahmin edicilerinin değişkenliğinin göstermiş olduğu tepkinin düzensiz olduğu görülmüştür. Parametre seçimi standartlaştırılmamış için bağımsız değişken sayısı arttığında tahmin edicilere ait hata kareler ortalamasının arttığı ve tahmin edicilere ait standart sapmanın azaldığı gözlemlenmiştir.

Yapılan bu çalışmada incelenen durumlara göre $n=30$ ve $n=50$ için k_{43} , $n=100$ için k_2 , düşük varyans değerleri için ($\sigma^2 = 0.25, 0.50$ ve 1) k_{43} , yüksek varyans değerleri için ($\sigma^2 = 2$ ve 4) k_{50} , $\rho = 0.90$ ve 0.95 için k_{43} , $\rho = 0.99$ için k_{28} ve $p = 4$ için k_{14} , $p = 8$ için k_{43} tahmin edicilerinin kullanılması önerilmektedir.

Kaynaklar

- Alkhamisi, M., Khalaf, G., & Shukur, G., 2006. Some Modifications for Choosing Ridge Parameters. Communications in Statistics - Theory and Methods, 35(11), 2005–2020.
- Alkhamisi, M. A., & Shukur, G., 2007. A Monte Carlo Study of Recent Ridge Parameters. Communications in Statistics - Simulation and Computation, 36(3), 535–547.
- Hoerl, A. E., & Kennard, R. W., 1970. Ridge Regression: Biased Estimation for Nonorthogonal Problems. Technometrics, 12(1), 55–67.
- Hoerl, A. E., & Kennard, R. W., 1970. Ridge Regression: Applications to Nonorthogonal Problems. Technometrics, 12(1), 69–82.
- Hoerl AE, Kennard RW, Baldwin KF 1975. Ridge Regression: Some Simulations. Communications in Statistics- Theory and Methods. 1975;4:105-123.
- Hocking, R. R., Speed, F. M., & Lynn, M. J, 1976. A Class of Biased Estimators in Linear Regression. Technometrics, 18(4), 425–437.
- Kibria, B. M. G., 2003. Performance of Some New Ridge Regression Estimators. Communications in Statistics – Simulation and Computation. 2003;32(2):419-435.

- Khalaf, G., & Shukur, G., 2005. Choosing Ridge Parameter for Regression Problems. *Communications in Statistics - Theory and Methods*, 34(5), 1177–1182.
- Kibria, B. M. G., Månsson, K., & Shukur, G. 2011. Performance of Some Logistic Ridge Regression Estimators. *Computational Economics*, 40(4), 401–414.
- Khalaf G. A, 2012. Proposed Ridge Parameter to Improve the Least Squares Estimator. *Journal of Modern Applied Statistical Methods*. 2012;11(2):443-449.
- Khalaf G. A, 2013. Comparison between Biased and Unbiased Estimators in Ordinary Least Squares Regression. *Journal of Modern Applied Statistical Methods*. 2013;12(2):293-303.
- Khalaf G, Månsson, K., Shukur, G, 2013. Modified Ridge Regression Estimators, *Communications in Statistics-Theory and Methods*. 2013;42:8,1476-1487.
- Lawless JF, Wang P, 1976. A Simulation Study of Ridge and other Regression Estimators. *Communications in Statistics- Theory and Methods*. 1976;14:1589-1604.
- McDonald, G. C., & Galarneau, D. I., 1975. A Monte Carlo Evaluation of Some Ridge-Type Estimators. *Journal of the American Statistical Association*, 70(350), 407.
- Montgomery, D, Peck, E, 1991-91. *Introduction to Linear Regresyon Analysis*. John Wiley and Sons Inc., 5, 645, New York.
- Muniz, G., & Kibria, B. M. G., 2009. On Some Ridge Regression Estimators: An Empirical Comparisons. *Communications in Statistics - Simulation and Computation*, 38(3), 621–630.
- Muniz G, Kibria BMG, Shukur G., 2012. On Developing Ridge Regression Parameters: a Graphical Investigation. *SORT (Stat Oper Res T)*. 2012;36(2):115-138.
- Nomura M., 1988. On the Almost Unbiased Ridge Regression Estimation. *Communications in Statistics – Simulation and Computation* 1988;17(3):729-743.
- Schaefer, R. L., Roi, L. D., & Wolfe, R. A, 1984. A ridge logistic estimator. *Communications in Statistics - Theory and Methods*, 13(1), 99–113.

Çıkar Çatışması

Yazarlar çıkar çatışması olmadığını beyan etmişlerdir.

Samsun İli Ondokuzmayıs İlçesinde Son Yıllarda Çilek Üretiminin Artmasının Nedeni Kârlılığı Olabilir mi?

Erhan Ersin DEMİR^{1*}, Göktuğ ÇAYIROĞLU², Esin HAZNECİ²

¹Ondokuz Mayıs Tarım ve Orman İlçe Müdürlüğü, 55420, Samsun, Türkiye

²Ondokuz Mayıs Üniversitesi, Ziraat Fakültesi, Tarım Ekonomisi Bölümü, 55139, Samsun, Türkiye

*Sunucu yazar e-posta: erhanersindemir@hotmail.com

Özet

Bu çalışma yörede son yıllarda üretimi artan çilek yetiştiriciliğinin Samsun ili Ondokuzmayıs ilçesinde kârlı bir üretim faaliyeti olarak yapıp yapılmadığını ortaya koymak amacıyla gerçekleştirilmiştir. Araştırmanın ana materyalini 2020 üretim döneminde çiftçilerden elde edilen üretim verileri oluşturmuştur. Araştırma verileri Ondokuzmayıs ilçesinde çilek yetiştiriciliği yapan 54 tarım işletmesinden gayeli olarak seçilmiş 30 işletmeden anket yoluyla toplanmıştır. Elde edilen veriler SPSS paket programı kullanılarak analiz edilmiştir. Çileğin üretim maliyetinin belirlenmesinde tek ürün bütçe analiz yöntemi kullanılmış, çilek üretiminde brüt kâr, net kâr, oransal kâr ve 1 kg çileğin maliyeti hesaplanmıştır. Araştırma sonucunda incelenen işletmelerde destekleme gelirleri de dikkate alındığında çilek üretiminden dekara 26784,61 TL brüt kâr ve 21697,57 TL net kâr elde edildiği hesaplanmıştır. Sonuç olarak işletmelerin üretim maliyetlerini karşıladıktan sonra, destekleme gelirleri olmasa bile çilek yetiştiriciliğinden oldukça tatminkar bir kâr seviyesi elde edeceği saptanmıştır. Üreticiler çilek üretimine yatırdıkları her 1 TL karşılığında, 1,85 TL gelir elde etmektedirler. Yörede başlangıçta çilek yetiştiriciliğinin artmasının asıl sebebinin verilen fide ve malç naylonu dseteği olduğu, sonrasında da kârlılığı nedeniyle üreticilerin çilek üretimine yöneldiği tespit edilmiştir. Çilek üretiminin bölgesel olarak desteklenmeye devam edilmesi ve üreticilere üretim teknikleri ile ilgili modern uygulamaların tarım teşkilatına bağlı uzmanlarca aktarılması, yörede çilek üretiminin sürdürülebilir olması bakımından önerilmektedir.

Anahtar kelimeler: Çilek, Üretim maliyetleri, Kârlılık analizi, Samsun, Türkiye

**Prediction of Egg Production in Japanese Quails (*Coturnix coturnix japonica*) by
Linear Regression, Linear Piecewise Regression, and Mars Algorithm**

Özgür KOŞKAN¹, Yasin ALTAY^{2*}, Sedat AKTAN¹

¹Isparta University of Applied Sciences, Faculty of Agriculture, Department of Animal Science,
32260, Isparta, Turkey

²Eskisehir Osmangazi University, Faculty of Agriculture, Department of Animal Science, 26160,
Eskisehir, Turkey

*Presenter author e-mail: yaltay@ogu.edu.tr

Abstract

The study aims to predict the egg production by linear regression, linear piecewise regression, and multivariate adaptive regression splines algorithms using age of sexual maturity, sexual maturity weight, weight average of the first 10 eggs, 20th, 30th, 40th, 60th, 80th, 100th, and 150th-day partial egg yields in Japanese quails. In this context, the material of the study consists of 128 female Japanese quails. In the comparison of estimation methods, the fit criteria of 15 different models were examined and the models were compared with the most common criteria as R^2 , $AdjR^2$, RMSE, RRMSE, RAE, CV, MRAE, MAPE, MAD, AIC, and cAIC. In the estimation of egg production, the best performance of the models with partial egg production on the 20th and 30th days included in the model as an independent variable belongs to the linear piecewise regression. All estimation methods showed similar results in models with partial egg production at 40th, 60th, and 80th days as independent variables. Although the linear regression and multivariate adaptive regression splines methods showed good performance on the 100th and 150th days of partial egg production, the linear piecewise regression method gave a worse predictive performance than these two methods. As a result, a successful prediction of partial egg production as an early selection criterion in egg production was made from the 100th day with the independent variables and linear regression - multivariate adaptive regression splines methods.

Key words: Egg production, Partial egg production, Linear regression, Piecewise Linear regression, MARS

A New Quantile Regression Model and Application

Yunus AKDOĞAN¹, Kadir KARAKAYA^{1*}, Aşır GENÇ², Mustafa Ç. KORKMAZ³, Fatih ŞAHİN¹

¹Selçuk University, Science Faculty, Department of Statistics, Konya, Turkey

²Necmettin Erbakan Üniversitesi, Fen Fakültesi, İstatistik, Konya, Turkey

³Artvin Çoruh University Department of Measurement and Evaluation, Artvin, Turkey

* Presenter author e-mail: kkarakaya@selcuk.edu.tr

Abstract

The alpha logarithmic Weibull distribution is introduced by Akdoğan et. al (2017). In this study, a new quantile regression model is introduced based on the alpha logarithmic Weibull distribution as an alternative to other existing quantile regression models. The maximum likelihood estimation method is conducted to estimate the parameters of the new quantile regression model. A practical example is studied to show the capacity of the new quantile regression models.

Key words: Alpha logarithmic Weibull distribution, Estimation, quantile regression, Gamma regression model, Real data analysis.

Genelleştirilmiş Procrustes Analizi ve Üzümde Çeşit ile Mineral Madde Arasındaki İlişkinin İki Boyutlu Uzaydaki Konfigürasyonu

Nurhan KESKİN¹, Birhan KUNTER², Sıddık KESKİN^{1*}

¹Van Yüzüncü Yıl Üniversitesi, Ziraat Fakültesi, Bahçe Bitkileri Bölümü, 65000, Van, Türkiye

²Ankara Üniversitesi, Ziraat Fakültesi, Bahçe Bitkileri Bölümü, 06110, Dışkapı, Ankara, Türkiye

³Van Yüzüncü Yıl Üniversitesi, Tıp Fakültesi, Biyoistatistik Anabilim Dalı, 65000, Van, Türkiye

*Sunucu yazar e-posta: skeskin@yyu.edu.tr

Özet

Genelleştirilmiş Procrustes Analizi kısaca; N adet ürün veya nesne; M adet değişken, özellik veya nitelik; K adet uzman, hakem, faktör veya koşula bağlı olarak elde edilen veri matrisini analiz etmek ve bunların birlikte konfigürasyonunu iki boyutlu uzayda göstermek üzere yapılan işlem aşamalarını içerir. Yöntem, daha çok duyuşal veriler (sensory data) üzerinde kullanılmakla birlikte, diğer alanlarda rahatlıkla kullanılabilir. Sağlıklı beslenme açısından üzüm önemli olmakla birlikte, sofralık üzümlerde önemli bir kalite bileşeni olan mineral madde ile çeşit arasındaki ilişkinin de doğru belirlenmesi oldukça önemlidir. Bu çalışmada, Genelleştirilmiş Procrustes Analizi yönteminin incelenmesi ve sofralık üzümlerde çeşit ile mineral madde arasındaki ilişkinin iki boyutlu uzaydaki konfigürasyonu amaçlanmıştır. Çalışmanın materyalini, 25 adet çekirdekli sofralık üzüm çeşidi (Alphonse Lavallée, Amasya Beyazı, Atasarı, Cardinal Çavuş, Erenköy Beyazı, Güllüzümü, Hafızali, Hamburg Misketi, İskenderiye Misketi, Italia, Kabarcık, Kadın Parmağı, Kalecik Karası, Karagevrek, Kozak Siyahı, Kömüş Memesi, Muscat, Müşküle, Perlette, Razakı, Safranbolu Çavuşu, Tahannebi, Uslu ve Yalova İncisi) oluşturmuştur. Bu çeşitlerin, çekirdeksiz ve çekirdekli ekstraktlarında; Sodyum (Na), Potasyum (K), Kalsiyum (Ca), Mangan (Mn), Demir (Fe), Çinko (Zn), Bakır (Cu) ve Magnezyum (Mg) içerikleri ölçülmüştür. Ölçülen bu değerler, Genelleştirilmiş Procrustes Analizi yöntemi ile analiz edilmiştir. Genelleştirilmiş Procrustes Analizi, N adet ürün veya nesne; M adet değişken, özellik veya nitelik; K adet uzman, hakem, faktör veya koşula bağlı olarak elde edilen veri matrislerini birlikte analiz ederek, bunların iki boyutlu uzaydaki en uygun konfigürasyonunu bulmayı amaçlar. Bunun için bireysel veri matrislerine; öteleme (translation), döndürme/yansıtma (rotation/reflection) ve izotropik ölçekleme (isotropic scaling) işlemlerini içeren transformasyonları (dönüşümleri) yapar. Analiz sonucunda, elde edilen PANOVA (Procrustes Analysis of Variance) tablosuna göre transformasyonların (öteleme, döndürme/yansıtma ve izotropik ölçekleme) etkileri önemli bulunmuştur ($p < 0.001$). Birinci boyut varyansın %70.38'ini açıklarken, ikinci boyut %29.15'ini açıklamış ve toplamda açıklanan varyans oranı %99.53 olarak bulunmuştur. Sonuç olarak, çalışmada Genelleştirilmiş Procrustes analizi incelenmiş ve çeşitlerle mineral maddeler arasındaki ilişkinin, Genelleştirilmiş Procrustes analizi ile değerlendirilebileceği ve yöntemin birçok alanda kullanılabileceği vurgulanmıştır.

Anahtar kelimeler: İzotropik ölçekleme, PANOVA, boyut, üzüm, mineral

Tabakalı Medyan Sıralı Küme Örneklemesinde Örnek Çapının Paylaştırılması: En Uygun Paylaştırma Yöntemi

Ayşe BOZKURT^{1*}, Sinem Tuğba ŞAHİN TEKİN², Yaprak Arzu ÖZDEMİR²

¹T.C. Hazine ve Maliye Bakanlığı, Gelir İdaresi Başkanlığı, 06110, Ankara, Türkiye

²Gazi Üniversitesi, Fen Fakültesi, İstatistik, 06374, Ankara, Türkiye

*Sunucu yazar e-posta: aysee.bozkurt@gmail.com

Özet

Araştırmalarda örnek seçim işlemi gerçekleştirilirken; seçilen örneğin yığını iyi temsil edip etmediği, örnek çapının doğru yöntemle belirlenip belirlenmediği ve ilgilenilen istatistiğin varyansının alternatif yöntemlere göre daha düşük olup olmadığı irdelenmelidir. Özellikle ilgilenilen değişken bakımından birimleri ölçmenin maliyetli veya çok zaman alıcı olması durumunda, sıralı küme örnekleme kullanılarak basit tesadüfi örnekleme ve tabakalı örneklemeyle göre daha etkin tahmin edicilerin elde edilmesi mümkün olmaktadır. Bu çalışmada, öncelikle tabakalı örneklemeyle alternatif olarak önerilen tabakalı medyan sıralı küme örnekleme ele alınmıştır. Daha sonra örnek çapını tabakalara paylaştırmak için doğrusal olmayan maliyet kısıtı altında bu yöntemin varyansı minimize edilmeye çalışılarak, tabakalardan seçilecek örnek çapı formülü teorik olarak elde edilmiştir. Uygulama olarak, belirlenen küme çapları, her tabakadan bir birim seçmenin maliyeti ve tabaka ağırlıklarına göre doğrusal olmayan maliyet kısıtı altında gerekli örnek çapları ve yığın ortalamasına ilişkin tahmin edicinin varyansı elde edilmiştir. Elde edilen sonuçlara göre, tabakalı medyan sıralı küme örneklemesinden elde edilen tahmin edicinin tabakalı tesadüfi örneklemeden elde edilen tahmin ediciden daha etkin olduğu gözlenmiştir.

Anahtar kelimeler: Sıralı küme örnekleme, Tabakalı medyan sıralı küme örnekleme, En uygun paylaştırma, Doğrusal olmayan maliyet fonksiyonu

İstatistiksel Programlarda MARS Modelinin Elde Edilişleri Arasındaki Farklar

Dilek SABANCI^{1*}

¹Gaziosmanpaşa Üniversitesi, Fen-Edebiyat Fakültesi, Matematik Bölümü, 60250, Tokat/Türkiye

*Sunucu yazar e-posta: dilek.kesgin@gop.edu.tr

Özet

Regresyon analizi, araştırmacılar tarafından bağımlı ya da tahmin edilecek değişkeni modellemek amacıyla oldukça fazla tercih edilen istatistiksel bir analiz tekniğidir. Bünyesinde birçok alternatif modeli barındırmaktadır. Bunlar arasında en güçlü olan tekniklerden birisi de Çok Değişkenli Uyarlamalı Regresyon Şeritleri (MARS) yöntemi olarak karşımıza çıkmaktadır. MARS makine öğreniminde denetimli bir öğrenme modelidir. “MARS” terimi yalnızca Minitab bünyesindeki Salford Predictive Modeler (SPM) ürününe ait ticari markadır ve lisanslanmıştır. Grafikselleştirilmiş bir kullanıcı arayüzü vardır. Dolayısıyla, bu marka bünyesindeki MARS yazılımı ücretli olarak kullanıcılara sunulmaktadır. MARS yöntemini kullanan StatSoft, SAS vb. bir çok farklı yazılım bulunmaktadır. Ancak güncel ilerlemeler düşünüldüğün de açık kodlu, ücretsiz ve hızlı geliştirilebilir olması münasebeti ile R-project programı araştırmacılar tarafından oldukça talep görmektedir. Bu nedenle istatistik analiz uygulamalarında alışılmışın dışına çıkmak isteyen araştırmacılar için bu çalışmada, MARS modelinin elde edilişindeki farklılıklar ticari olan ve olmayan farklı program üzerinden örneklerle karşılaştırılarak sonuçları yorumlanmıştır.

Anahtar kelimeler: Regresyon analizi, MARS, Minitab SPM MARS, R-project

Increasing Sample Size Through Appropriate Use of Rolling Windows

Virgilio PÉREZ^{1*}, Jose M. PAVÍA¹, Cristina AYBAR¹

¹University of Valencia, Faculty of Economics, Department of Applied Economics, 46022, Valencia, Spain

*Presenter author e-mail: virgilio.perez@uv.com

Abstract

When conducting research, it is often impossible, or too expensive, to collect data from all the people or items we are interested in. Instead, we take a sample of the population of interest and predict how the population behaves from that sample.

There are lot of aspects that influence how much (or little) a sample resembles the population and thus the validity and reliability of our conclusions. One of the key concepts to consider when conducting a research is sample size. It is well known that as more and more information become available (larger the sample size), the uncertainty reduces. For that reason, in this study we have focused on how to extend the sample size without the need to obtain new data, and without barely increase the associated error levels.

As we say, sometimes we cannot validate our research because we have little data (small sample sizes), a circumstance that causes a lot of volatility. To solve this problem, we propose to use adjacent information, i.e., to admit that close subjects behave similar. To verify this assumption, we used a large database with more than 150 variables and more than 700,000 observations, over 30 years. We have found that there are indeed certain time variables that evolve very smoothly, such as political ideology, among others, so it is not unreasonable to establish data clusters (or data pools). The proposed method describes the use of the rolling windows, achieving the described objective of increasing the sample size without barely increase the estimation error.

Key words: Rolling windows, Sample size, Estimation error, Predictions

**Bayes Estimation of Stress-Strength Reliability for Unit Burn Type XII Distribution
Based on Progressive Type-II Censoring**

Fatma ÇİFTÇİ^{1*}, Neriman AKDAM², Yunus AKDOĞAN³

¹Konya PTT General Directorate, 42040, Konya, Turkey

²Selçuk University, Faculty of Medicine, Division of Biostatistics, 42130, Konya, Turkey

³Selçuk University, Faculty of Science, Department of Statistics, 42130, Konya, Turkey

*Presenter author e-mail: fatma_ifc@hotmail.com

Abstract

In this paper, the stress–strength reliability parameter, $R=P(Y<X)$, is estimated based on progressive type II censored samples when X and Y are two independent Unit Burn Type XII distributions with different parameters. The maximum likelihood estimator of R is obtained and, in addition, the approximate Bayes estimator under squared error loss function with the help of Lindley's approximations of R is derived. A Monte Carlo simulation study is carried out to compare the performance of the approximate Bayes estimators with the maximum likelihood estimator. Also, one data set is analyzed for illustrative purposes.

Key words: Unit Burn Type XII distribution, Bayes estimator, Lindley's approximation, Monte Carlo simulation, Squared-error loss function

Permütasyonel Çok Değişkenli Varyans Analizi (PERMANOVA) ve Covid-19 Verisi ile Uygulama

Yeliz KAŞKO ARICI^{1*}, Sıddık KESKİN², Şeyda Tuba SAVRUN³

¹Ordu Üniversitesi, Tıp Fakültesi, Biyoistatistik ve Tıbbi Bilişim AD., 52200, Ordu, Türkiye

²Van Yüzüncü Yıl Üniversitesi, Tıp Fakültesi, Biyoistatistik AD., 65080, Van, Türkiye

³Ordu Üniversitesi, Tıp Fakültesi, Acil Tıp AD., 52200, Ordu, Türkiye

*Sunucu yazar e-posta: ykaskoarici@odu.edu.tr

Özet

Sürekli bir bağımlı değişkenin, faktör değişken(ler)in seviyelerine göre değişiminin incelenmesinde en yaygın kullanılan istatistik yöntem tek değişkenli varyans analizidir (ANOVA). ANOVA ile faktörlerin esas etkilerinin yanı sıra, çalışmada birden fazla faktör varsa aralarındaki etkileşimler (etkileşimler) da incelenebilmektedir. Ancak birden fazla bağımlı değişkenin birlikte incelenmek istendiği durumlarda ANOVA, bağımlı değişkenler arası ilişkileri göz ardı etmektedir. Doğada birçok değişkenin birbiri ile ilişkili olabileceği düşünüldüğünde, değişkenlerin ayrı ayrı ANOVA ile analiz edilmesi bilgi kaybına neden olabilir. Bununla birlikte, bağımlı değişkenlerin ayrı ayrı analizi de I. tip hatayı artırabilir. Çok değişkenli varyans analizi (MANOVA), bağımlı değişkenler arasındaki ilişkileri dikkate alarak, faktörlerin etkilerini ve olası etkileşimleri (etkileşimleri) eşzamanlı incelemeye olanak sağlamaktadır. Ancak literatürde bağımlı değişkenlerin çoğunlukla ayrı ayrı ANOVA ile analiz edildiği görülmektedir. Bu durumun, özellikle sağlık bilimlerinde daha yaygın olduğu söylenebilir. Bu nedenle, MANOVA gibi birden fazla değişkenin birlikte analiz edilmesine olanak sağlayan çok değişkenli istatistik analiz yöntemlerinin kullanımının yaygınlaştırılması gerektiği düşünülmektedir. Parametrik bir test olması nedeniyle, MANOVA'nın kullanımı da ANOVA gibi bazı varsayımların yerine getirilmiş olmasına bağlıdır. Değişkenlerin normal dağılımlı olmaması ve grupların homojen varyanslara sahip olmaması, ANOVA gibi MANOVA'nın da kullanımını sınırlamaktadır. MANOVA'nın parametrik olmayan karşılığı olan Permütasyonel çok değişkenli varyans analizi (Permutational Multivariate Analysis of Variance, PERMANOVA) bu varsayımların ihlalden etkilenmemektedir. Dolayısıyla, hem bağımlı değişkenler arasındaki ilişkileri dikkate alan hem de faktörler arasındaki etkileşimleri inceleme imkanı sağlayan PERMANOVA, varsayımlarının daha az olmasıyla da kullanım kolaylığı sunmaktadır. PERMANOVA, özellikle birçok değişkenin parametrik test varsayımlarını yerine getirmediği bilinen sağlık bilimleri için önem kazanmaktadır. PERMANOVA yaygın istatistik programlarının menülerinde henüz yer almamış olmasa da bazı istatistik yazılımlar ile yapılabilmektedir. Bu çalışmada, çok değişkenli analiz yöntemlerinden birisi olan PERMANOVA'nın açıklanması, temel kavramlarına değinilmesi ve konunun anlaşılmasına katkı sağlamak üzere Covid-19 veri seti ile bir uygulama yapılması amaçlanmıştır. Bu amaçla Covid-19 tanısı alan 435 hastaya ait veriler kullanılmıştır. Covid-19 tanısı için laboratuvar biyobelirteçleri olarak kabul edilen Ferritin, D-Dimer ve WBC'nin cinsiyet ve hastaneye yatıp yatmama durumuna göre değişimi sabit etkili tam faktöriyel modele uygun olarak incelenmiştir. XLSTAT 2021 programı kullanılarak gerçekleştirilen analizde uzaklık ölçüsü olarak Bray-Curtis kullanılmış ve 999 adet permütasyon yapılmıştır. Analiz sonucunda "cinsiyet×yatış durumu" etkileşimi istatistik olarak anlamlı bulunmuştur ($p<0.05$). Sonuç olarak, birden fazla bağımlı değişkenin birlikte ele alınması gereken araştırmalarda daha az varsayıma sahip olan PERMANOVA'nın kullanılabileceği önerilmiştir.

Anahtar kelimeler: PERMANOVA, Bray-Curtis, Çok değişkenli analiz, Parametrik test varsayımları, Covid-19

Bayes Estimation for the Gompertz Gumbel Type-II Distribution Based on Progressive Type-II Censored Samples

Neriman AKDAM^{1*}, Muslu Kazım KÖREZ¹, Fatma ÇİFTÇİ²

¹Selçuk University, Faculty of Medicine, Division of Biostatistics, 42130, Konya, Turkey

²Konya PTT General Directorate, 42040, Konya, Turkey

* Presenter author e-mail: nkaradayi@selcuk.edu.tr

Abstract

In this study, the maximum likelihood estimators (MLEs) and Bayes estimators for four parameters of Gompertz Gumbel Type-II (GGTT) distribution based on progressive type-II censored samples are derived. Because the Bayes estimators cannot be obtained in closed forms, the approximate Bayes estimator is computed using the idea of Lindley under squared-error loss function. Then, the approximate Bayes estimates are compared with the maximum likelihood estimates in terms of the estimated risks (ERs) using Monte Carlo simulation. Also, the analysis of real data set is presented for illustrative purposes.

Key words: Gompertz Gumbel type-II distribution, Bayes estimator, Lindley's approximation, Monte Carlo simulation, Squared-error loss function

Türkiye ve AB Ülkelerinin Covid-19' a Karşı Etkinliğinin Veri Zarflama Analizi ile Karşılaştırılması

Gonca CETİNKAYA EROĞLU¹, Hasan BAL¹

¹Gazi Üniversitesi, Fen Fakültesi, İstatistik Bölümü, 06560, Ankara, Türkiye

*Sunucu yazar e-posta: goncacetinkaya@hotmail.com

Özet

2020 yılının başında, tüm halk sağlığı camiasının onlarca yıldır korktuğu senaryo gerçek olmuş ve hayvanlardan insanlara bulaşarak hastalığa neden olan korona virüslerin yeni bir tipi Çin' den tüm dünyaya hızlıca yayılmaya başlamıştır. Dünya Sağlık Örgütü tarafınca "Covid-19" olarak adlandırılan virüs 24 Ocak 2020 tarihinde Avrupa kıtasına ulaşmış ülkemizde ise ilk vaka 11 Mart 2020 tarihinde görülmüştür. Bu çalışmanın amacı Aralık 2019'da Çin'in Wuhan şehrinde ortaya çıkan ve kısa sürede dünyayı saran COVID-19 salgını ile mücadelede Avrupa Birliği ülkelerinin ve Türkiye'nin etkinlik düzeyinin Veri Zarflama Analizi ile araştırılmasıdır. Sağlık sistemlerinin performans ölçümünde sıklıkla kullanılan Veri Zarflama Analizi, pek çok girdi ve çıktının kullanılabildiği parametrik olmayan etkinlik ölçüm yöntemidir. Analizde girdi değişkenleri olarak; onbin kişi başına doktor, hemşire, hastane ve yoğun bakım yatak sayıları ile sağlık harcamalarının Gayri Safi Yurtiçi Hâsıla (GSYH) içindeki oranı, çıktı olarak; ülkelerin 10 Haziran 2021 tarihine ait milyon kişi başına Covid-19 test sayısı, vaka sayısı ve ölüm sayısı ele alınmıştır. Etkinlik analizinde sabit ölçekli Charnes, Cooper, Rhodes (CCR) ve değişken ölçekli Banker, Charnes, Cooper (BCC) yöntemleri kullanılmış, etkin olan ülkelerin etkinlik skorları belirlenmiş ve etkin olmayan ülkeler için potansiyel iyileştirme önerilerinin geliştirilmesi amaçlanmıştır.

Anahtar kelimeler: Covid-19, Veri zarflama analizi, Avrupa birliği, Sağlık sistemleri

Comparison of the Effectiveness of Turkey and EU Countries Against Covid-19 with Data Envelope Analysis

Abstract

At the beginning of 2020, the scenario that the entire public health community feared for decades came true, and a new type of corona viruses that cause disease by transmitting from animals to humans began to spread rapidly from China to the whole world. The virus, called "Covid-19" by the World Health Organization, reached the European continent on January 24, 2020, and the first case was seen on March 11, 2020 in our country. The purpose of this study is to investigate the efficiency level of the European Union countries and Turkey in the fight against the COVID-19 pandemic, which emerged in Wuhan, China in December 2019 and spread around the world in a short time, using Data Envelopment analysis. Data Envelopment Analysis, which is frequently used in the performance measurement of health systems, is a non-parametric efficiency measurement method where many inputs and outputs can be used. As input variables in the analysis; The number of doctors, nurses, hospitals and intensive care beds per ten thousand people, and the ratio of health expenditures to Gross Domestic Product (GDP), as output; COVID-19 pandemic data of countries from June 10, 2021, Number of Tests, Number of Cases and Number of Deaths per million population were used. In the study, output oriented Charnes, Cooper and Rhodes (CCR) and Banker, Charnes and Cooper (BCC) methods were used, scale efficiency

scores were determined, and it was aimed to develop potential improvement suggestions for inefficient countries.

Key words: Covid-19, Data envelopment analysis, European union, Health systems

GİRİŞ

2020 yılının başında, tüm halk sağlığı camiasının onlarca yıldır korktuğu senaryo gerçek olmuş ve zoonotik virüsler olarak tanımlanan, hayvanlardan insanlara bulaşarak hastalığa neden olan korona virüslerin yeni bir tipi Çin’ den tüm dünyaya hızlıca yayılmaya başlamıştır. İlk olarak 31 Aralık 2019’ da Çin’ in Hubei Eyaleti, Vuhan şehrinde farklı hayvan türleri satan toptan balık ve canlı hayvan pazarı çalışanlarında ortaya çıkmış daha sonra sırasıyla Tayland ve Japonya’ da görülmüştür. Dünya Sağlık Örgütü tarafınca “Covid-19” olarak adlandırılan virüs 24 Ocak 2020 tarihinde Avrupa kıtasına ulaşmış ilk olarak Fransa’ nın Bordeaux şehrinde tespit edilmiştir. Ülkemizde ise ilk vaka 11 Mart 2020 tarihinde görülmüş ve aynı gün virüsün dünyaya yayılma hızının endişe verici bir hal alması sonucu Covid-19 salgını, Dünya Sağlık Örgütü tarafından pandemi olarak ilan edilmiştir.

Bilgi ve sermaye akışının hızlanmasını sağlayan küreselleşme, sunduğu seyahat imkânları neticesinde Covid-19 salgının yayılım hızını arttırmış ve dünya, daha önce yaşamadığı kadar büyük çapta sağlık krizi ile yüz yüze kalmıştır. Yaşanan gelişmelerin ardından ülkeler peş peşe sınır kontrolü uygulamalarına başlamış ve salgını kontrol altına almak için kısıtlama stratejileri geliştirmeye çalışmıştır. Bu noktada ülkelerin sağlık sistemlerinin performansı salgınla mücadelede etkin bir rol oynamıştır. Her ülkenin hastane ve sağlık personeli yoğunluğu, yapılan test sayısı, kişi başına düşen sağlık harcamaları gibi kriterler hastalığın teşhis hızı, bulaşıcılığın azaltılması ve tedavinin etkinliği ile ilişkilendirilmiştir. Yaşlı ve sistematik rahatsızlıkları olan kişilerde ölümle sonuçlanma olasılığı daha yüksek olan Covid-19’ la mücadele etmek, yaşlı nüfusa sahip Avrupa Birliği’ ne üye ülkeler için de çok zor olmuştur. 13 Mart 2020 tarihinde Dünya Sağlık Örgütü’ nün Avrupa’ nın, salgının merkez üssü haline geldiğini duyurmasının üzerine AB üye ülkeleri, virüsün yayılma hızını engelleme amaçlı birbirlerinden bağımsız olarak farklı önlemler almışlardır. Ülkelerin bazıları kısmi kısıtlamalara giderken bazı ülkeler sağlık sistemini felç etmemek için sosyal hayatı ciddi şekilde kısıtlayan sokağa çıkma yasakları uygulamışlardır.

Avro krizi ve Brexit’ in etkilerini üzerinden atmaya çalışan Avrupa Birliği, salgının ilk aylarında ortak eylem planı oluşturamamış tıbbi teçhizat yetersizlikleri, sağlık personeli ve yatak sayısı eksikliği olan ülkelere yardım konusunda yeterli bulunmamıştır. Salgının ilk aylarında hastalığın tedavisi için önem arz eden solunum cihazı ve yayılım hızını azaltan maske gibi tıbbi enstrümanları, her ülkenin kendisi için stok yapması ve İtalya örneğinde görüldüğü gibi yardım bekleyen ülkelere ulaştırmaması Birlik’ in varlığının dahi tartışılmasına yol açmıştır. Sağlık alanında olduğu gibi ekonomik ve siyasi alanlarda da olumsuz sonuçlara neden olan Covid-19 salgını, ilk vakanın ortaya çıktığı günden bugüne kadar 208 ülkeye yayılmış, 178.618.383 kişiye bulaşarak 3.798.151 kişinin ölümüne sebep olmuştur (WHO)

Pandemi yönetimi süreçlerinin değerlendirilmesi, salgının ilerleyen dönemde yönetimine ve gelecek olası salgınlara rehber olması amacıyla hayati önem taşımaktadır (Varol ve Tokuç, 2020). Salgın boyunca her ülke tanı, tedavi, alınması gereken önlemler, sağlık yatırımları vb. gibi pek çok açıdan birbirinden farklı uygulamalar getirmişlerdir. Bir ülkede bir salgın durumunda enfekte ve ölüm vakalarının sayısı ve bunların artma oranı büyük ölçüde o ülkenin hazırlık durumuna ve sağlık sistemine bağlıdır (Selamzade ve Özdemir, 2020). Bu nedenle Covid-19 salgınında ülkelerin sağlık sistemlerinin performanslarının karşılaştırılması ve etkin süreçlerin belirlenmesi, salgınla mücadelenin en önemli adımlarından biridir.

Başta Avrupa ülkeleri olmak üzere tüm dünyada son yirmi yıldır gündeme gelen sağlık reformlarının temel amaçlarından biri de, sağlık sistemlerinin verimlilik performanslarının nasıl geliştirilebileceği ile ilgilidir. Başka bir ifade ile ülkeler; sağlık sistemlerinin kalite, beklentilerin karşılanması, adil finansman, etkililik, verimlilik vb performanslarını artırma konusunda baskı altındadırlar. Birçok ülke ve bu arada da uluslararası kuruluşlar (DSÖ, DB ve OECD gibi) sağlık sistemlerinin performanslarını ölçmek için çeşitli performans kriterleri ve stratejileri geliştirmekte, sağlık sistemlerinin performanslarını yükseltmeyi amaçlayan sağlık reformu çalışmalarına girişmektedir. Bu nedenle AB ölçeğinde üye ve aday ülke sağlık sistemlerinin verimlilik performanslarının ortaya konulması önem arz etmektedir(Yıldırım, 2005).

Bu çalışmanın amacı AB üye ülkeler ile aday olan ülkemizin Covid-19 krizinde, sağlık sistemlerinin etkinliklerinin karşılaştırılmasını sağlamak ve elde edilen analiz sonuçlarının gelecekte meydana gelebilecek olası salgınlarla mücadelede ülkelerin sağlık sistemlerinin performanslarını arttırmak üzere değerlendirilmesine katkı sağlamaktır. Ayrıca analiz sonucunda her ülkenin etkinlik skorları belirlendikten sonra yapılan sıralamaya göre etkinlik skoru 1 olan ülkeler arasından seçilmiş referans kümeleri etkin olmayan ülkeler için yol gösterici olmaktadır. Böylece araştırma kapsamında ele alınan girdi ve çıktılar, etkin olmayan ülkelere gerçekleştirilmesi gerekli sağlık reformları açısından Avrupa Birliği ülkelerine ve Türkiye'ye tavsiye niteliğinde bilgiler sunmaktadır.

MATERYAL VE YÖNTEM

Materyal

Çalışma kapsamında Avrupa Birliği ülkeleri ve Türkiye'nin Covid-19 salgınına karşı etkinliklerini karşılaştırmak için Veri Zarflama Analizinden yararlanılmıştır. Veri zarflama analizi için gerekli olan girdi ve çıktılar belirlenmiş; girdi sayısı 5 adet ve çıktı sayısı 3 adetle sınırlandırılmıştır. Avrupa Birliği'ne üye 27 ülke ve aday ülke konumundaki Türkiye'nin on bin kişi başına düşen doktor sayısı, hemşire sayısı, hastane yatak sayısı, yoğun bakım yatak sayısı ve Gayri Safi Yurtiçi Hâsıla (GSYH) içindeki sağlık harcamalarının oranları girdi olarak belirlenmiştir. Ülkelerin 10 Haziran 2021 tarihine ait milyon kişi başına düşen test sayısı, vaka sayısı (tersi) ve ölüm sayısı (tersi) ise çıktı olarak ele alınmıştır. Girdilere ait veriler Avrupa Birliği İstatistik Ofisi (Eurostat) veri tabanından elde edilmiş Dünya Bankası (World Bank) veri tabanı, Dünya Sağlık Örgütü (WHO) veri tabanından kontrol edilmiştir. Çıktılara ait veriler, <https://www.worldometers.info> internet adresinde ülkelerin 10 Haziran 2021 gününe ait Covid-19 verilerinden alınmıştır. Etkinlik analizinde sabit ölçekli Charnes, Cooper, Rhodes (CCR) ve değişken ölçekli Banker, Charnes, Cooper (BCC) yöntemleri kullanılmıştır.

Yöntem

Verilerin Elde Edilmesi

Girdiler için gerekli olan veri seti Dünya Bankası (World Bank) veri tabanı, Dünya Sağlık Örgütü (WHO) veri tabanı ve Avrupa Birliği İstatistik Ofisi (Eurostat) ' ın internet sayfalarından alınarak verilerin tutarlılığı karşılaştırılmıştır. Her üç sitede de girdilere ait en güncel veriler 2018 yılına ait olarak bulunmaktadır. Fakat girdi değişkenlerinden biri olan yoğun bakım yatak sayısı bahsedilen üç sitede de bulunmamaktadır. Avrupa Birliği ülkelerine ait en yakın tarihli yoğun bakım sayısına <https://www.statistica.com> internet sitesinden ulaşılmış olup veriler 2012 yılına aittir. Türkiye'nin yoğun bakım yatak sayıları güncel olmasına rağmen analizde hata olmaması açısından T.C. sağlık Bakanlığının 2012 yılına ait verileri uygulamada yer almıştır (<https://rapor.saglik.gov.tr/istatistik/rapor>). Çıktılar için kullanılan veri seti ise tüm ülkelerin güncel koronavirüs verilerinin yer aldığı <https://www.worldometers.info> internet sitesinden alınmıştır. Çalışmada analizi yapılan çıktılar bahsedilen internet sitesinin 10 Haziran 2021 gününe ait koranavirüs verileridir. Derlenen veri kümesi, Veri Zarflama Analizi yapmak üzere EMS 1.3.0 paket programında kullanılmıştır.

İstatistiksel Analizler

Veri Zarflama Analizi (VZA), 1978 yılında Charnes ve ark.'nın bir dizi homojen karar verme biriminin girdi ve çıktıları, göreceli etkinliklerini değerlendirerek karşılaştırmak için yaygın olarak kullanılan ve paratmetrik olmayan matematiksel programlama yaklaşımıdır.

VZA kullanan deneysel analizde, verimli olarak derecelendirilen her gözlem, bir verimlilik sınırı tanımlamak için kullanılır. Çok verimli olmayan ve aynı çıktı veya girdi karışımına sahip olan gözlemler, sınırdaki varsayımsal bir gözlem ile karşılaştırılarak değerlendirilir. Daha sonra, karşılaştırılan gözlemin verimliliği, gerçek çıktı seviyelerinin varsayımsal gözleme oranıdır. Göreceli verimlilik ölçüsü, her bir gözlemin potansiyeline göre ne kadar iyi performans gösterdiğine dair bir gösterge verir. En iyi gözlemlerin 0 ile 1 arasında 1 puan alması gerektiğinden, puanlardaki farklılık, politika yapıcılara olası iyileştirmenin kapsamı hakkında bir fikir verir. Veri zarflama analizi (VZA), işletme düzeyinde verimliliği hesaplamak ve çeşitli kaynakların kullanımındaki durgunluğu ölçmek için kullanılır (Bowlin, 1998).

VZA'nın başlangıcı, 1957 yılında Journal of the Royal Statistical Society'de yayınlanan, M.J. Farrell'in "Üretken Verimliliğin Ölçümü" başlıklı makalesi olarak kabul görmektedir. Daha önceki çalışmaların birden çok girdinin ölçümünü herhangi bir tatmin edici genel verimlilik ölçüsü ile birleştirmede başarısız olduklarını savunan Farrell, sorunu daha yeterli bir şekilde ele alabilecek "etkinlik analizi" yaklaşımını önermiştir. İlk VZA modeli, 1978'de A. Charnes, W.W. Cooper ve E. Rhodes'un (CCR) sunduğu şekliyle, Farrell'in 1957'de yayınlanan önceki çalışması üzerine inşa edilmiştir (Charnes ve ark., 1978). Bu model, VZA'daki diğer varyantları formüle etmek için temel formüldür ve bu model ölçeğe göre sabit getirilidir. Birçok durumda ölçeğe göre değişken getiri durumu söz konusu olduğuna göre Banker ve ark., tarafından 1984 yılında CCR geliştirilmiş ve ölçeğe farklı bir getiri düşünülmüş ve bu nedenle, çoklu girdi yönelimli bir form olarak BCC'yi tanıtmıştır (Cooper ve ark., 2011).

CCR modelinin iki farklı türü bulunmaktadır. Bunlardan ilki mevcut bulunan çıktı seviyesini karşılayabilecek şekilde girdileri minimize etmeyi amaçlayan girdi odaklı model, bir diğeri ise, mevcut girdilerden daha fazlasını talep etmeyecek şekilde çıktıları maksimize etmeyi amaçlayan çıktı odaklı modeldir (Kıran, 2008).

$$E_k = \text{Max} \frac{\sum_{r=1}^p u_r Y_{rk}}{\sum_{i=1}^m v_i X_{ik}} \quad (1)$$

Burada CCR modeli, m sayıda girdi ve p sayıda çıktı üreten n tane Karar Verme Biriminin (KVB) her biri için ağırlıklı girdilerin ağırlıklı çıktılarına oranını maksimize eder. Girdi yönelimli analiz yapılırken (1) formülünde bulunan $\sum_{i=1}^m V_i X_{ik}$ ifadesi, çıktı yönelimli analiz yapılırken ise $\sum_{r=1}^p u_r Y_{rk}$ ifadesi 1'e eşitlenerek maksimize edilir.

$$\text{Kısıtlar: } \sum_{j=1}^n X_{ij} \lambda_j \leq X_{i0}; \sum_{j=1}^n Y_{rj} \lambda_j \geq Y_{r0}; \lambda_j \geq 0$$

Girdi odaklı modelin çözümü sonucunda elde edilen değerler göreceli etkinlik ölçüleridir ve oranın 1'e eşit olması durumunda KVB'nin etkin olduğu, 1'den küçük olması durumunda ise etkin olmadığı söylenebilir. Çıktı odaklı modelde ise 1'e eşit olması durumunda KVB'nin etkin olduğu, 1'den büyük olması durumunda ise etkin olmadığı söylenebilir.

Ölçeğe göre değişen getiri durumuna sahip sistemlerin etkinliklerini belirleyebilmek için, 1984 yılında Banker, Charnes ve Cooper kendi isimlerinin baş harfleri ile anılan BCC modelini geliştirmişlerdir (Banker ve ark., 1984). Bunun için CCR modellerinin dualine konvekslik kısıtı denilen $\sum_{j=1}^n \lambda_j = 1$ kısıtı eklenmiştir. Buna göre; bir KVB için hesaplanan λ_j 'lerin (ağırlıkların) toplamı birden büyük ise KVB ölçeğe göre azalan getiriye; birden küçük ise artan getiriye ve bire eşit ise sabit getiriye

göre faaliyet gösteriyor anlamına gelmektedir(Yıldız, 2006).

$$E_k = \text{Max} \frac{\sum_{r=1}^p u_r Y_{rk} - \mu_0}{\sum_{i=1}^m v_i X_{ik}} \quad (2)$$

Kısıtlar: $\sum_{j=1}^n X_{ij} \lambda_j \leq X_{i0}$; $\sum_{j=1}^n Y_{rj} \lambda_j \geq \beta Y_{r0}$; $\sum_{j=1}^n \lambda_j = 1$; $\lambda_j \geq 0$; $\mu_0 = 0$. KVB ‘ye ait serbest işletli değişken.

Çalışmada kullanılan girdi ve çıktıları ait istatistiksel değerler Çizelge 1’ de gösterilmiştir. AB üye ülkeleri ve Türkiye’nin girdi değişkenleri incelendiğinde onbin kişi başına doktor sayısı en çok olan ülke Yunanistan (61,04), en az olan ülke Türkiye (18,81) olmuştur. Onbin kişi başına hemşire sayısı en çok olan ülke Belçika (72,81), en az olan ülke İsveç (0,01) olmuştur. Onbin kişi başına hastane yatak sayısı en çok olan ülke Almanya (80,02), en az olan ülke İsveç (21,38) olmuştur. Onbin kişi başına yoğun bakım yatak sayısı en çok olan ülke Almanya (2,92), en az olan ülke Portekiz (0,42) olmuştur. GSYH içinde cari sağlık harcamaları en çok olan ülke Almanya (11,43), en az olan ülke Türkiye (4,12) olmuştur.

Aynı şekilde AB üye ülkeleri ve Türkiye’nin çıktı değişkenleri incelendiğinde bir milyon kişi başına en çok test yapılmış ülke Danimarka (10.391.780), en az test yapılmış ülke ise Bulgaristan (426.572) olmuştur. Bir milyon kişi başına vaka sayısı en çok olan ülke Çekya (155.174), en az olan ülke de Finlandiya (16.871) olmuştur. Bir milyon kişi başına ölüm sayısı Macaristan (3.102) en az olan ülke de Finlandiya (174) olmuştur.

Çizelge 1. Değişkenlerin genel istatistiği

Değişkenler	Girdiler			
	En çok	En az	Ortalama	Standart Sapma
Doktor	61,04	18,81	38,09	8,77
Hemşire	72,81	0,001	39,39	21,71
Hastane	80,02	21,38	48,85	17,34
Yoğun Bakım	29,2	0,42	11,95	6,75
Harcama	11,43	4,12	8,03	1,93
Değişkenler	Çıktılar			
	En çok	En az	Ortalama	Standart Sapma
Test	10.391.780	426.572	1.794.170	2198950,4
Vaka	155.174	16.871	77.975	27954,31
Ölüm	3.102	174	1.511	735,43

BULGULAR

Çizelge 2’de 28 ülkenin CCR ve BCC modellerine göre etkinlik skorları sıralanmış ayrıca ölçek etkinlik skorları da hesaplanmıştır. KVB lerin ölçek etkinliklerinin hesaplanması CCR yöntemi ile elde edilen etkinlik skorlarının BCC yöntemi ile elde edilen etkinlik skorlarına bölünmesi ile elde edilir (Kutlar ve Babacan, 2008).

Çizelge 2’de ölçeğe göre sabit getirili CCR yöntemi ile yapılmış etkinlik analizine göre Danimarka, Finlandiya, Kıbrıs (GKRY), İsveç, Malta ve Slovakya olmak üzere 6 Avrupa Birliği ülkesi etkin olmuştur. Bu yöntemle göre ortalama etkinlik skoru %55,42’dir. Lüksemburg (0,8525), Türkiye (0,8028) ve Romanya (0,6536) ortalamasının üzerinde etkinlik skoru almışlardır. Türkiye aldığı % 80,28’

lik etkinlik skoru ile KVB' leri içerisinde 8. sırayı almıştır. CCR yöntemine göre en az etkin olan üç ülke İspanya (0,3033) , Hollanda (0,2669) ve Slovenya (0,2587) 'dır.

CCR yöntemi ile yapılan etkinlik analizinde Danimarka 22, Finlandiya 1, Kıbrıs (GKRY) 1, İsveç 1, Malta 1 ve Slovakya 2 ülkeye referans olmuştur. Örneğin etkin olan Slovakya; etkin olmayan Romanya ve Hırvatistan tarafından referans alınmıştır. Etkinlik skoru 0,5400 olan Avusturya, % 49 oranında Danimarka' yı, % 7 oranında Finlandiya'yı referans alarak etkinliğini arttırabilmektedir. Aynı şekilde etkinlik skoru 0,6536 olan Romanya, % 25 oranında Finlandiya' yı, %2 oranında Danimarka 'yı ve % 18 Slovakya' yı referans alarak etkinliğini arttırabilmektedir.

Çizelge 2'den teknik etkinlik olarak da adlandırılan ölçeğe göre değişken getirili BCC modelinden elde edilen skorlar da belirtilmiştir. CCR yöntemine göre etkin olan 6 ülkenin dışında İrlanda, Polonya, Romanya, Lüksemburg, Slovenya, Yunanistan ve Türkiye olmak üzere toplam 13 ülkenin BCC yöntemine göre etkin olduğu görülmektedir. CCR modeli ile etkin olan 6 ülke, BCC modeli ile de etkin çıkmaktadır. $\theta_{BCC}^* \geq \theta_{CCR}^*$ kısıtı sağlanmış olup tersi geçerli değildir. BCC yöntemine göre ortalama etkinlik skoru % 0,8873'dir. Portekiz (0,9458), Hollanda (0,9416) ve Letonya (0,9380) ortalamanın üzerinde etkinlik skoru almışlardır. BCC yöntemine göre en az etkin olan üç ülke Belçika (0,6987) , Avusturya (0,6431) ve Almanya (0,5352)'dir.

BCC yöntemi ile yapılan etkinlik analizinde Türkiye etkin olmayan karar verme birimleri tarafından 15 kez referans gösterilmiştir. Danimarka 8, Malta 7, Polonya 5, Kıbrıs (GKRY), İrlanda, İsveç ve Slovakya 4, Finlandiya 2, Romanya ve Malta 1 kez referans gösterilmiştir. Örneğin etkin olan Türkiye; etkin olmayan 15 ülke (Almanya, Avusturya, Belçika, Bulgaristan, Çekya, Estonya, Fransa, Hırvatistan, Hollanda, İspanya, Letonya, Litvanya, Macaristan ve Portekiz) tarafından referans alınmıştır. Etkinlik skoru 0,6431 olan Avusturya; % 43 oranında Kıbrıs (GKRY), % 32 oranında Türkiye, % 23 oranında Danimarka, % 2 oranında Lüksemburg'u referans alarak etkinliğini arttırabilmektedir.

Çizelge 2'de ölçek etkinlikleri de hesaplanmıştır. KVB lerin ölçek etkinliklerinin hesaplanması CCR yöntemi ile elde edilen etkinlik skorlarının BCC yöntemi ile elde edilen etkinlik skorlarına bölünmesi ile elde edilir (Kutlar ve Babacan, 2008). Ölçek etkinlik ve teknik etkinlik değerlerinin bilinmesi, toplam etkin olmayan bir KVB'nin etkinsizliğinin nedeninin teknik etkinsizlikten mi yoksa ölçek etkinsizlikten mi, ya da her ikisinden de mi kaynaklandığının belirlenmesini de sağlamaktadır.

$$\text{Ölçek etkinlik} = \frac{\text{Toplam Etkinlik}}{\text{Teknik Etkinlik}}$$

Ölçek etkinliği analizi sonuçlarına göre 14 ülkenin ölçeğe göre etkin olduğu ve sabit getiri (C) özelliği sağladığı görülmüştür. CCR modeline göre etkin olan 6 ülkenin dışında 8 ülke daha sabit getirili ve ölçeğe göre etkin olduğu gözlenmiştir. Bu ülkeler Almanya, Avusturya Fransa, Estonya, Belçika, Bulgaristan, Çekya ve İtalya'dır. Bu ülkelerin ölçeklerini değiştirmeden aynı şekilde COVID-19 pandemisi ile mücadele etmesi gerekmektedir. Aynı zamanda kalan 14 ülkenin ölçeğe göre artan getiri (I) özelliğinin olduğu tespit edilmiş, bu ülkelerin ölçeklerini azaltarak Covid-19 ile mücadele etmesinin gerektiği söylenebilmektedir. Ülkelerin toplam etkinlik ortalaması 0,5542 iken ölçek etkinlik ortalaması 0,6149 olmuştur.

II. International Applied Statistics Conference (UYIK-2021)
Tokat / Turkey, 29 June - 2 July 2021

Çizelge 2. CCR ve BCC modellerine göre etkinlik skorları

Ülkeler	CCR Etkinlik Skoru	Ülkeler	BCC Etkinlik Skoru	Ölçek Etkinlik Skoru	Ölçeğe Göre Getiri Türü
Danimarka	1	Danimarka	1	1	C
Finlandiya	1	Finlandiya	1	1	C
Kıbrıs (GKRY)	1	Kıbrıs (GKRY)	1	1	C
İsvec	1	İsvec	1	1	C
Malta	1	Malta	1	1	C
Slovakya	1	Slovakya	1	1	C
Luksemburg	0,8525	İrlanda	1	0,501	I
Türkiye	0,8028	Polonya	1	0,4581	I
Romanya	0,6536	Romanya	1	0,6536	I
Yunanistan	0,5472	Luksemburg	1	0,8525	I
Avusturya	0,54	Slovenya	1	0,2587	I
İrlanda	0,501	Yunanistan	1	0,5472	I
Polonya	0,4581	Türkiye	1	0,8028	I
Letonya	0,4529	Portekiz	0,9458	0,3642	I
Almanya	0,4337	Hollanda	0,9416	0,2834	I
Fransa	0,3773	Letonya	0,938	0,4828	I
Litvanya	0,3699	İspanya	0,8667	0,3499	I
Belçika	0,3597	İtalya	0,8157	0,4162	C
Bulgaristan	0,3567	Cekya	0,8043	0,4253	C
Portekiz	0,3445	Hirvatistan	0,7891	0,3948	I
Cekya	0,3421	Litvanya	0,7819	0,4730	I
İtalya	0,3395	Bulgaristan	0,778	0,458	C
Estonya	0,3308	Macaristan	0,7774	0,4069	I
Macaristan	0,3164	Fransa	0,7718	0,4888	C
Hirvatistan	0,3116	Estonya	0,7597	0,4354	C
İspanya	0,3033	Belçika	0,6987	0,5148	C
Hollanda	0,2669	Avusturya	0,6431	0,8396	C
Slovenya	0,2587	Almanya	0,5352	0,8103	C
Ortalama	0,5542		0,8873	0,6149	

Çizelge 3' te etkinlik analizi sonucuna göre etkin olmuş ülkelerin CCR ve BCC yöntemleri ile elde edilmiş süper etkinlik skorları verilmiştir. CCR yöntemi ile elde edilen tahmin sonuçlarına göre Malta 23,6809 ile en yüksek, Finlandiya 5,2991 ve İsveç 5,30915 ile en yüksek süper etkinlik skoru almışlardır. CCR yöntemi ile en düşük etkinlik skoru ile İspanya (0,3033), Hollanda (0,2669) ve Slovenya (0,2587) olmuştur. BCC yöntemi ile elde edilen tahmin sonuçlarına göre en yüksek süper etkinlik skoru olan ülkeler Danimarka, Finlandiya ve İsveç olmuştur. BCC yöntemi ile en düşük etkinlik skoru ile Belçika (0,6987), Avusturya (0,6431) ve Almanya (0,5352) olmuştur.

Çizelge 3. CCR ve BCC modellerine göre süper etkinlik skorları

Ülkeler	CCR Süper Etkinlik Skoru	Ülkeler	BCC Etkinlik Skoru
Malta	23,6809	Danimarka	Big
Finlandiya	5,2991	Finlandiya	Big
Isvec	5,0915	Isvec	3743,4782
Danimarka	2,8377	Slovakya	992,2090
Slovakya	2,4264	Malta	45,6403
Kıbrıs (GKRY)	1,2717	Türkiye	1,6920
		Kıbrıs (GKRY)	1,2981
		Luksemburg	1,1735
		Polonya	1,1569
		Romanya	1,1057
		İrlanda	1,0860
		Yunanistan	1,0237
		Slovenya	1,0185

Çizelge 4'te CCR yöntemi ile tahmin edilmiş etkinlik analizi sonucunda tam etkin olmayan ülkeler için önerilen potansiyel iyileştirme oranları sunulmuştur. Etkin olmayan 22 ülke gösterilerek etkin hale gelebilmeleri için girdi ve çıktı değişkenlerini ne yönde ve ne oranda değiştirmesi gerektiği belirtilmiştir.

Çıktı odaklı CCR modeli sonuçlarına göre, 22 etkin olmayan ülkeden Avustralya, Bulgaristan, Çekya, Estonya, Letonya, Litvanya, Luxemburg, Macaristan, Portekiz, Romanya ve Yunanistan'ın onbin kişi başına düşen doktor sayısının % 0,19 ile % 17,25 arasında değişen oranlarda fazla olduğu tespit edilmiştir. On bin kişi başına düşen hemşire sayısında ise Romanya, Hırvatistan ve Slovenya hariç kalan 19 ülkede % 0,34 ile % 47,39 arasında değişen oranlarda fazla olduğu tespit edilmiştir. Onbin kişi başına düşen hastane yatak sayısı Portekiz hariç tüm ülkelerde %0,9 ile % 47,29 arasında değişen oranlarda fazla olduğu görülmektedir. . Onbin kişi başına düşen yoğun bakım yatak sayısında yine Portekiz hariç tüm ülkelerde % 0,20 ile % 2,34 arasında değişen oranlarda fazla olduğu görülmektedir. Girdi değişkenlerinden GSYH içerisindeki cari sağlık harcamalarının oranında etkin olmayan ülkelere 11' inde (Almanya, Belçika, Fransa, Hollanda, İrlanda, İtalya, İspanya, Polonya, Portekiz, Slovenya, Türkiye) % 0,3 ile % 4,51 arasında değişen oranlarda fazla olduğu tespit edilmiştir.

Girdi değişkenlerinin azaltılmasının mümkün olmadığı durumlarda, etkin olmayan ülkelerin çıktı değişkenlerinin değiştirilmesi de mümkün olmaktadır. Fakat Çizelge 4 incelendiğinde çıktı odaklı CCR yöntemine göre çıktılar arasında milyon kişi başına düşen ölüm oranlarının % 0,01 ile % 0,03 arasında değişen oranlarda daha az olması beklenmektedir. Türkiye'nin hastane yatak sayısını % 14,25 oranında, hemşire sayısını % 13,34 ve yoğun bakım yatak sayısını % 1,11 azaltması durumunda CCR yöntemine göre etkin ülkeler arasında olması mümkün olacaktır. Türkiye'nin çıktı değişkenlerinde değişiklik önerilmemiştir.

II. International Applied Statistics Conference (UYIK-2021)
Tokat / Turkey, 29 June - 2 July 2021

Çizelge 4. CCR Yöntemleri ile Yapılan Tahmin Sonucunda İyileştirme Öneriler

CCR								
GİRDİLER			ÇIKTILAR					
ÜLKELER	Hastane	Doktor	Hemşire	Yoğun Bakım	Harcama	Vaka	Ölüm	Test
Almanya	47,22	0	40,72	2,34	2,88	0	0,03	0,01
Avusturya	45,99	8,28	3,52	1,49	0	0	0,01	0
Belçika	34,32	0	47,39	1,14	3,67	0	0,02	0
Bulgaristan	47,07	4,85	20,38	0,72	0	0	0,03	0
Çekya	46,5	7,77	3,72	0,65	0	0	0,01	0
Estonya	22,67	2,57	21,33	0,97	0	0	0	0
Fransa	36,97	0	31,22	0,71	4,51	0	0,01	0
Hırvatistan	26,77	0,19	0	0,95	0	0	0,02	0
Hollanda	5,19	0	21,57	0,13	2,35	0	0	0
İrlanda	5,32	0	37,33	0,2	0,3	0	0,02	0
İtalya	2,68	0	18,47	0,7	0,43	0	0,03	0
İspanya	0,9	0	8,06	0,41	0,6	0	0,02	0
Letonya	33,1	1,99	5,99	0,53	0	0	0,01	0
Litvanya	44,61	17,25	26,18	1,14	0	0	0,01	0
Lüksemburg	32,3	7,82	27,3	2,13	0	0	0	0
Macaristan	45,5	0,73	14,52	0,93	0	0	0,03	0
Polonya	47,29	0	26,57	0,37	1,61	0	0,02	0
Portekiz	0	9,56	0,34	0	0,42	0	0,02	0
Romanya	39,44	1,89	0	1,64	0	0	0,01	0
Slovenya	21,89	0	0	0,2	1,59	0	0,01	0
Yunanistan	12,75	22,36	6,93	0,08	0	0	0,03	0
Türkiye	14,25	0	13,34	1,11	0,37	0	0	0

Çizelge 5` te BCC yöntemi ile tahmin edilmiş etkinlik analizi sonucunda tam etkin olmayan bölgeler için önerilen potansiyel iyileştirme oranları sunulmuştur. Etkin olmayan 15 ülke gösterilerek etkin hale gelebilmeleri için girdi ve çıktı değişkenlerini ne yönde ve ne oranda değiştirmesi gerektiği belirtilmiştir.

Çıktı odaklı BCC modeli sonuçlarına göre, 15 etkin olmayan ülkeden Avustralya, Bulgaristan, Litvanya ve Portekiz`in onbin kişi başına düşen doktor sayısının % 4,89 ile % 12,52 arasında değişen oranlarda azaltması gerektiği tespit edilmiştir. On bin kişi başına düşen hemşire sayısında ise Letonya, Hırvatistan, Çekya ve İspanya hariç kalan 11 ülkede % 1,96 ile % 41,95 arasında değişen oranlarda fazla olduğu tespit edilmiştir. Onbin kişi başına düşen hastane yatak sayısı Hollanda, İtalya, İspanya ve Portekiz hariç tüm ülkelerde % 4,45 ile % 46,98 arasında değişen oranlarda fazla olduğu görülmektedir. Onbin kişi başına düşen yoğun bakım yatak sayısında Hırvatistan, Hollanda, Letonya, Litvanya ve Portekiz hariç tüm ülkelerde % 0,02 ile % 2,22 arasında değişen oranlarda fazla olduğu görülmektedir. Girdi değişkenlerinden GSYH içerisindeki cari sağlık harcamalarının oranında etkin olmayan ülkelere 9`unda (Almanya, Avusturya, Belçika, Fransa, Hırvatistan, Hollanda, İrlanda, İtalya ve İspanya) % 0,14 ile % 4,37 arasında değişen oranlarda fazla olduğu tespit edilmiştir.

Aynı şekilde girdi değişkenlerinin azaltılmasının mümkün olmadığı durumlarda, etkin olmayan ülkelerin çıktı değişkenlerinin değiştirilmesi de mümkün olmaktadır. Çizelge 5 incelendiğinde çıktı odaklı BCC yönteminde çıktılar arasında milyon kişi başına düşen ölüm oranlarının % 0,01 ile % 0,03 arasında değişen oranlarda daha az olması beklenmektedir. Ayrıca milyon kişi başına test sayısının

İtalya'da %27576,9 ve İspanya'da %157997,5 oranda daha fazla olması etkin hale gelmeleri için gerekmektedir.

Çizelge 5. BCC Yöntemleri ile Yapılan Tahmin Sonucunda İyileştirme Öneriler

ÜLKELER	BCC							
	GİRDİLER				ÇIKTILAR			
	Hastane	Doktor	Hemşire	Yoğun Bakım	Harcama	Vaka	Ölüm	Test
Almanya	45,67	0	39,31	2,22	2,84	0	0,03	0
Avusturya	46,98	9,92	6	1,52	0,39	0	0,01	0
Belcika	26,91	0	41,95	0,58	3,52	0	0,02	0
Bulgaristan	41,13	4,89	14,43	0,26	0	0	0,03	0
Cekya	35,48	0	0	0,15	0	0	0,02	0
Estonya	8,92	0	9,12	0,02	0	0	0	0
Fransa	29,8	0	26	0,17	4,37	0	0,01	0
Hirvatistan	4,45	0	0	0	0,14	0	0,01	0
Hollanda	0	0	1,96	0	2,22	0	0	0
İtalya	0	0	14,78	0,48	0,32	0	0,03	27576,94
İspanya	0	0	0	0,22	0,23	0	0,02	150997,5
Letonya	23,27	0	0	0	0	0	0,01	0
Litvanya	27,23	12,52	12,31	0	0	0	0,01	0
Macaristan	35,32	0	4,92	0,17	0	0	0,02	0
Portekiz	0	8,64	2,5	0	0	0	0,01	0

TARTIŞMA VE SONUÇ

Son yılların en büyük küresel problemi olarak tanımlanan Covid-19 salgını, Dünya Sağlık Örgütü tarafından pandemi ilan edimesinden çalışmanın yapıldığı 10 Haziran 2021 gününe dek dünya çapında yaklaşık 178 milyon kişinin hastalanmasına ve 3,8 milyon kişinin ölümüne sebep olmuştur. Hayatı durma noktasına getiren pandemi süresince tüm devletler sağlık sistemlerinin ayakta kalması ile ilgili ciddi sınavlar vermiş birçoğu da yetersiz kalmıştır. Bu nedenle tüm dünyada olduğu gibi Avrupa Birliği'ne üye ülkeler ve Türkiye için de sağlık sistemlerinin performansını ölçerek Covid-19 mücadelesinde ne derece etkin olduklarının belirlenmesi önem taşımaktadır. Başlangıçta kar amacı gütmeyen kuruluşların verimliliğini ölçmek için kullanılan Veri Zarflama Analizi, son yıllarda sağlık sistemlerinin performansını ölçmek üzere yaygın olarak kullanılmaktadır. Çalışmamızda VZA' nın ölçeğe göre sabit getirili CCR ve değişken getirili BCC yöntemleri ile ülkelerin etkinlik skorları belirtilmiş, etkin olmayan ülkeler için potansiyel iyileştirme oranları gösterilmiştir.

Analiz sonuçlarına göre Avrupa Birliği'ne üye 27 ülke ve aday olan Türkiye olmak üzere toplam 28 ülke arasında 6 tanesi (% 21'i) CCR modeline göre Covid-19 salgınına karşı etkin çıkmıştır. BBC modeline göre bahsedilen 6 ülkeye 7 ülke daha eklenerek toplam 13 ülke (%46'sı) Covid-19 salgınına göre etkin olarak belirtilmiştir. Danimarka, Finlandiya, Kıbrıs (GKRY), İsveç, Malta ve Slovakya CCR modelinde etkin olan ülkeler olup İrlanda, Polonya, Romanya, Lüksemburg, Slovenya, Yunanistan ve Türkiye sadece BCC modeline göre etkin çıkmışlardır. Etkin ülkeler içerisinde ise Malta 23,6809, Finlandiya 5,2991 ve İsveç 5,0915 'lük süper etkinlik skoru ile CCR modelinde, Danimarka, Finlandiya ve İsveç ise süper etkinlik skorları ile BCC modelinde ilk üç sırayı almıştır.

Etkin olmayan ülkeler için tahmin edilen potansiyel iyileştirme oranlarına göre çıktı yönelimli CCR analizi sonucunda tam etkinlik skoruna en yakın skor almış Lüksemburg' un mevcut girdiler arasında on bin kişi başına hastane yatak sayısını %32,3 oranında, on bin kişi başına doktor sayısını % 7,82 oranında ve on bin kişi başına hemşire sayısını %27,30 oranında azaltması gerektiği söylenebilir. Ayrıca çıktılar arasında değişiklik yapması gerekmemektedir. BCC analizi sonucunda ise tam etkinlik skoruna

en yakın skor almış Portekiz'in mevcut girdiler arasında on bin kişi başına doktor sayısını % 8,64 oranında ve on bin kişi başına hemşire sayısını % 2,5 oranında azaltması gerektiği söylenebilir.

Selamzade ve Özdemir (2020) çalışmasında, OECD ülkelerinin Covid-19'a karşı etkinliklerini Veri Zarflama Analizi (VZA) üzerinden değerlendirmiştir. 24 Nisan 2020 tarihine göre ele aldığı verilerle CCR yöntemini kullanarak bulduğu sonuçlara göre süper etkinlik skoru en yüksek olan ülkeler Slovakya (9,007), Meksika (7,068), İzlanda (3,444) olmuştur. BCC yöntemine göre ise süper etkinlik skoru en yüksek olan ülkeler İzlanda (4,183), Japonya (1,732) ve Lüksemburg (1,249) olmuştur.

Chen, Yu ve Chia-Yu Su (2020) çalışmasında G20 ülkelerinden 19 ülke ile İran, İspanya, Nijerya ve Pakistan'ın bulunduğu 23 ülkenin Covid-19'a karşı etkinliklerini Veri Zarflama Analizi (VZA) üzerinden değerlendirmiştir. Pandeminin başından itibaren geçen 120 günlük süreçte veriler değerlendirmeye alınmıştır. Analiz sonucunda Avustralya ve Kore en yüksek skoru alırken Brezilya, Rusya ve Amerika en düşük skoru almışlardır.

Asanduluia, Romanb ve Fatulescu (2014) çalışmasında Avrupa'daki 30 ülkenin sağlık sistemlerinin etkinliklerini Veri Zarflama Analizi (VZA) üzerinden değerlendirmiştir. CCR yöntemine göre çıkan sonuçlarda 6 tane ülke (Kıbrıs, Malta, Bulgaristan, Romanya, UK ve İsveç) etkin çıkmıştır. BCC yöntemine göre ise 7 adet ülke (Kıbrıs, Malta, Litvanya, Romanya, İspanya, UK ve İsveç) etkin çıkmıştır.

Tarihte örneklerini gördüğümüz salgıların, küreselleşen dünya düzeni nedeniyle bulaşma hızını arttırması gün geçtikçe daha büyük kitleleri tehdit etmektedir. Bu nedenle ülkelerin sağlık sistemlerinin pandemilere hazırlıklı olmalarının önemi Covid-19 salgını ile açıkça görülmüştür. Bu çalışmada hayati risklerin dışında ekonomik, sosyal ve eğitim vb. gibi pek çok olumsuz etkilerini yaşadığımız Covid-19 salgını ile mücadele yönetiminde, ülkelerin sağlık sistemlerinin performansları değerlendirilmeye çalışılmış, bu konuda çalışan yöneticiler, karar vericiler ve politika yapıcılar için yol gösterici olması hedeflenmiştir.

Kaynaklar

- Asanduluia L, Romanb R, Fatulescu F, 2014. The efficiency of healthcare systems in Europe: a Data Envelopment Analysis Approach, *Procedia Economics and Finance*, 10, 261 – 268.
- Avrupa İstatistik Ofisi, Eurostat, 2018. Erişim adresi: https://ec.europa.eu/eurostat/statistics-explained/index.php?title=Healthcare_personnel_statistics/_nursing_and_caring_professionals&oldid=355980; Erişim tarihi: 10.06.2021.
- Avrupa İstatistik Ofisi, Eurostat, 2018. Erişim adresi: https://ec.europa.eu/eurostat/statistics-explained/index.php?title=Healthcare_personnel_statistics/_physicians&oldid=355980; Erişim tarihi: 10.06.2021.
- Avrupa İstatistik Ofisi, Eurostat, 2018. Erişim adresi: https://ec.europa.eu/eurostat/statistics-explained/index.php?title=Healthcare_expenditure_statistics; Erişim tarihi: 10.06.2021.
- Avrupa İstatistik Ofisi, Eurostat, 2018. <https://ec.europa.eu/eurostat/web/products-datasets/-/tps00046>; Erişim tarihi: 10.06.2021.
- Bowlin W, 1998. Measuring Performance: An introduction to data envelopment analysis (DEA), *The Journal of Cost Analysis*, 27, 918.
- Charnes A., Cooper WW, Rhodes E, 1978. Measuring the efficiency of decision making units, *European Journal of Operation Research*, 2, 429-444.
- Chen Y, Yu S, Chia-Yu Su E, 2020. An examination of COVID-19 mitigation efficiency 2 among 23 countries, *MedRXiv*, 20 (25).
- Cooper WW, Seiford LM, Zhu J, 2011. *Handbook on Data Envelopment Analysis*, Second Edition, Springer: New York.
- Kıran B, 2008. Kalkınmada öncelikli illerin ekonomik etkinliklerinin veri zarflama analizi yöntemi ile değerlendirilmesi. Yüksek lisans Tezi, Çukurova Üniversitesi, Sosyal Bilimleri Enstitüsü, Adana, Türkiye.

- Kutlar A, Babacan A, 2008. Türkiye’deki kamu üniversitelerinde CCR etkinliği-ölçek etkinliği analizi: DEA tekniği uygulaması. Kocaeli Üniversitesi Sosyal Bilimler Dergisi, 15, 148- 172.
- Rhodes A, Ferdinande P, Flatten H, Kılavuz B, Metnitz, PG, Moreno RP, 2012. The variability of critical care bed number in Europe . Intensive Care Med. 38 (10), 647- 653.
- Selamzade F, Özdemir Y, 2020. Covid-19’a Karşı OECD Ülkelerinin Etkinliğinin VZA ile Değerlendirilmesi, Turkish Studies, 15 (4): 977-991.
- T.C. Sağlık Bakanlığı, 2012. Erişim adresi: <https://rapor.saglik.gov.tr/istatistik/rapor/>; Erişim tarihi:12.06.2021.
- Varol G, Tokuç B, 2020. Halk Sağlığı Boyutuyla Türkiye’de Covid-19 Pandemisinin Değerlendirmesi, Namık Kemal Tıp Dergisi, 8 (3): 579-594.
- WHO, 2018. Erişim adresi: [https://www.who.int/data/gho/data/indicators/indicator-details/GHO/hospital-beds-\(per-10-000-population\)](https://www.who.int/data/gho/data/indicators/indicator-details/GHO/hospital-beds-(per-10-000-population)); Erişim tarihi:10.06.2021.
- WHO, 2021. Erişim adresi: <https://covid19.who.int/>; Erişim tarihi: 10.06.2021.
- Worldometers, 2021. Erişim adresi: <https://www.worldometers.info/coronavirus/>; Erişim tarihi: 10.06.2021.
- Yıldırım HH, 2005. Avrupa Birliği’ne üye ve aday ülke sağlık sistemlerinin karşılaştırmalı performans analizi: veri zarflama analizine dayalı bir uygulama. Verimlilik Dergisi.

Çıkar Çatışması

Yazarlar çıkar çatışması olmadığını beyan etmişlerdir.

Sağlık Bilimi Verilerinde Bir Uygulama ile Doğrusal Olmayan Temel Bileşen Analizinin Tanıtımı

Canan DEMİR^{1*}, Sıddık KESKİN¹

¹Van Yüzüncü Yıl Üniversitesi

*Sunucu yazar e-posta: canandemir@yyu.edu.tr

Özet

Doğrusal Olmayan Temel Bileşen Analizi, açıklayıcı boyut indirgeme tekniklerinden biridir ve doğrusal veya doğrusal olmayan ilişkileri içeren değişken kümesi için sayısal ve grafiksel sonuçlar sunar. Bu çalışmada Doğrusal Olmayan Temel Bileşenler Analizi tanıtılmış ve bu yöntemle travma (cinsel ve fiziksel) ile öğrencilerdeki demografik özellikler arasındaki ilişki incelenmiştir. Veriler, 548 öğrenciden anket yoluyla elde edildi. Çalışmaya travma (cinsel ve fiziksel) ile birlikte dokuz değişken dahil edildi. NLPCA, 9 değişken ve travma arasındaki doğrusal ve doğrusal olmayan ilişkileri belirlemek için kullanıldı. Birinci ve ikinci boyutun özdeğerleri sırasıyla 1.766 ve 1.504 olarak bulundu. Bu özdeğerlere göre, toplam varyansın %17.656'sı birinci boyut tarafından, %15,044'ü ise ikinci boyut tarafından açıklanmaktadır. Açıklanan toplam varyans %28,550 olarak bulunmuştur. Cinsiyet, yaş, medeni durum ve intihar değişkenlerinin travma üzerinde etkili olduğu bulundu. NLPCA ile kategorik değişkenler en uygun şekilde istenilen boyuta ölçeklendirildi ve böylece değişkenler arasındaki doğrusal olmayan ve doğrusal ilişkiler modellenildi. Sonuç olarak NLPCA, başta sağlık olmak üzere birçok alanda farklı türde değişkenleri ve değişken yapıları içeren çalışmalarda faydalı olabilir.

Anahtar kelimeler: Bileşen yükleri, Boyut küçültme, Doğrusal olmayan temel bileşen analizi, Optimal ölçekleme

Ülkelerin Sağlık Ekonomik ve Gelişmişlik Göstergeleri Arasındaki İlişki: Kanonik Bir Yaklaşım

Taner TUNC^{1*}, Firdevs Berrak CABI¹

¹Ondokuz Mayıs Üniversitesi Fen Edebiyat Fakültesi İstatistik Bölümü, 55200 Samsun, Türkiye

*Sunucu yazar e-posta: ttunc@omu.edu.tr

Özet

Dünyada gelişmiş ve gelişmekte olan ülkelerin bazı göstergelerinin birbirleri ile ilişkili olup-olmadığı önemli bir araştırma konusudur. Gelişmemiş ülkelere model olacak bu durum hakkında bilgi sahibi olmak ve strateji geliştirmek oldukça gereklidir. Bu çalışmada öncelikle dünyadaki bazı ülkelerin belirli alanlardaki 2019 yılına ait güncel değişken değerleri faktör analizi ile geçerli ve güvenilir bir alt boyut oluşturulmuştur. Oluşturulan alt boyutlar barındırdıkları değişkenler itibarıyla ülkelerin sağlık, ekonomi ve gelişmişlik özelliklerini yansıtmaktadır. Oluşturulan alt boyutların birbirleriyle olan ilişkileri ise veri setleri arasındaki doğrusal ilişkiyi analiz eden kanonik korelasyon analizi ile çözümlenmiştir. Buna göre ülkelerin sağlık ekonomi ve gelişmişlik düzeyi göstergelerine ait veri setleri arasında istatistiksel olarak anlamlı ilişkiler olduğu gözlenmiştir.

Anahtar kelimeler: Faktör analizi, Kanonik değişken, Açıklanan varyans, Gereksizlik ölçütü, Gösterge

Ölçek Geçerlilik ve Güvenilirliğinde Farklı Bir Yaklaşım

Taner TUNÇ^{1*}, Erdinç KOLAY²

¹Ondokuz Mayıs Üniversitesi Fen Edebiyat Fakültesi İstatistik Bölümü, 55200 Samsun, Türkiye

²Sinop Üniversitesi Fen Edebiyat Fakültesi İstatistik Bölümü, 57000 Sinop, Türkiye

*Sunucu yazar e-posta: ttunc@omu.edu.tr

Özet

Tüm bilim dallarında bilgiye ulaşmak için belirli yollar vardır. Bunlardan bir tanesi ise belli bir olguyu ölçen, belirli bir örüntüyü ortaya çıkarmak için geliştirilen ölçeklerdir. Bunlar için çok önemli bir kavram olan geçerlilik ve güvenilirlik ise ölçme aracında yer alan maddelerin istenilen olguyu veya tutumu her defasında tutarlı ölçüp ölçmediğinin tespitini yapmaktadır. Geçerlilik, ölçme aracında yer alan maddelerin incelenen olguyu bir araya gelerek tutarlı bir şekilde ölçüp-ölçmediğinin bir ölçüsüdür. Bununla beraber ölçme aracının her tekrarında elde edilen sonuçların tutarlılığı ise güvenilirlik anlamına gelir. Bu çalışmada, 2012 yılında geliştirilmiş İstatistik Tutum Ölçeği, 2014 yılında 477 üniversite öğrencisine ve 2019 yılında 1034 üniversite öğrencisine uygulanarak maddelerin alt boyutlara yüklenmelerinin sabitleşip-sabitleşmediği gözlenmiştir. Buna göre geçerlilik ve güvenilirlik için geliştirilen yöntemler ve ölçütlerin yanı sıra ölçeklerdeki maddelerin dağılımının analizini yapan, yıllar içinde değişmezliğini ve eğilimini tespit eden RİDİT analizi, ilk defa ölçekler için geçerlilik ve güvenilirlik anlamında kullanılmıştır. Bu bağlamda literatürde ilk kez Ridit Analizi ölçek geçerliliği ve güvenilirliği için kullanılmıştır. Her madde için hesaplanan “ridit” değerleri itibarıyla İstatistik Tutum Ölçeği maddelerinin iç tutarlılığı hem 2014 hem de 2019 yılında incelendiğinde iyi düzeyde ve uyumlu olduğu söylenilebilir. Bu sonuç ise ölçeğin tutarlı bir hale geldiğinin göstergesidir.

Anahtar kelimeler: Ridit analizi, Ortalama ridit, İstatistik tutum ölçeği, Geçerlilik, Güvenilirlik

Goodness of Fit Tests Based on Kullback-Leibler Divergence under Censoring for Weibull Distribution

Gülcan GENCER^{1*}, İsmail KINACI²

¹ Karamanoglu Mehmetbey University, Karaman, TURKEY,

² Selçuk University, Department of Actuarial Science, Konya, TURKEY

*Presenter author e-mail: gulcangencer@kmu.edu.tr

Abstract

Measuring the fit between a known theoretical distribution and the distribution containing sample value is called the goodness of fit test. In this thesis study, critical values and powers were obtained based on the best and worst bandwidth tests for Weibull distribution based on Kullback-Leibler divergence in case of censored and complete sample. As a result, in the case of complete and censored sampling, based on Kullback-Leibler divergence, critical value and power comparison was made under different parameters based on the best and worst bandwidths. In all comparisons, it was seen that the power of the test increased as the sample size increased. It has been understood that the selection of bandwidth w is important by giving the best and the worst bandwidths in terms of the goodness of fit tests problem. Since it is understood that the test cannot be performed in the case where bandwidth w is selected as the worst for TV test, the importance of w for this test stands out. For the tests with bandwidth comparison, it was seen that the strongest test was mostly the TC test, then the TV test, and the lowest test was the TVE test.

Keywords: Bandwidth, Kullback-Leibler Divergence, Censoring, Power of Test, Goodness of Fit

Weibull Dağılımı için Sansürleme Altında Kullback-Leibler Uyuşmazlığına Dayalı Uyum İyiliği Testleri

Özet

Bilinen bir teorik dağılım ile örnek değer içeren dağılım arasındaki uyumun ölçülmesi uyum iyiliği testi olarak adlandırılır. Bu çalışmada, sansürlü ve tam örneklem durumunda Kullback-Leibler uyumsuzluğuna dayalı Weibull dağılımı için uyum iyiliği testleri en iyi ve en kötü bant genişliği baz alınarak kritik değerler ve güçler elde edilmiştir. Sonuç olarak tam ve sansürlü örneklem durumunda Kullback-Leibler uyumsuzluğuna dayalı olarak, en iyi ve en kötü band genişlikleri baz alınarak farklı parametreler altında kritik değer ve güç karşılaştırması yapılmıştır. Tüm karşılaştırmalarda örnek hacmi arttıkça testin gücünün arttığı görülmüştür. Uyum iyiliği testleri problemi bakımından en iyi ve en kötü band genişlikleri verilerek bant genişliği w 'nin seçiminin önemli olduğu anlaşılmıştır. TV testi için genellikle bant genişliği w 'nin en kötü seçildiği durumda testin gerçekleştirilemediği anlaşıldığından bu test için w 'nin önemi öne çıkmaktadır. Bant genişliği karşılaştırması yapılan testler için çoğunlukla en güçlü testin TC testi daha sonra TV testi, en düşük testin ise TVE testi olduğu görülmüştür.

Anahtar kelimeler: Bant genişliği, Kullback-Leibler Uyuşmazlığı, Sansürleme, Testin Gücü, Uyum İyiliği.

INTRODUCTION

Goodness of fit tests are used to determine the model owned by the audience, namely the probability distribution, at least to determine whether a considered probability distribution is suitable for the population. If a hypothesis test measures the fit between a known theoretical distribution and a distribution containing sample values, it is called a goodness of fit test (Massey, 1951). From the earliest days when statistics were used, statisticians started their analysis by proposing a distribution for their observations, and then checked whether the distribution they used was correct for the data they observed. Thus, many test procedures have emerged over the years and studying these procedures has become known as goodness of fit (D'agostino and Stephens, 1986). In the literature, using the Kullback-Leibler divergence measure, goodness of fit tests based on complete and censored samples have been suggested for various distribution families. (Choi et al., 2004) proposed an exponential test based on the Kullback-Leibler divergence measure. As an estimator of Shannon entropy, Van Es and Correa's entropy was used. Comparison has been made with various goodness of fit tests and it has been found to have more power than other tests.

(Lim and Park, 2007) made a comparison of power for different distributions with monotone decreasing hazard function, monotone increasing hazard function and non-monotone hazard function for Exponential and Normal distribution based on the partial Kullback-Leibler divergence measure in the case of a type-2 censored sample. For the exponential test of the test statistics, it was observed that the distributions with monotonous increasing hazard function had more power than other tests, and Tukey (5) had better power in the comparison for normality test. (Balakrishnan et al., 2007) proposed a goodness-of-fit test based on Kullback-Leibler knowledge in progressive type-2 censored data for exponentiality test. He compared some alternatives with different hazard functions, and as a result, the test was found to be stronger at times for alternatives with non-monotonous hazard functions.

(Rad et al., 2011) proposed a goodness of fit test based on Kullback-Leibler information for progressive type-2 censored data, using the approximate maximum likelihood estimators and maximum likelihood estimators of the model parameters for Pareto, Log-normal and Weibull distributions. They examined the strength of the test for different alternatives under different sample sizes. The Weibull distribution is a 2-parameter distribution with α shape and β scale parameter. It was obtained by Weibull (1951) by generalizing the exponential distribution. This distribution is known as β averaged exponential distribution when $\alpha = 1$ and Rayleigh distribution when $\alpha = 2$. The probability density function and distribution function of a X random variable with α and β parameter Weibull distribution are as follows.

$$f(x) = \alpha\beta^{-\alpha}x^{\alpha-1}e^{-\left(\frac{x}{\beta}\right)^{\alpha}}, \quad \alpha > 0, \beta > 0, x > 0 \quad (1)$$

$$F(x) = 1 - e^{-\left(\frac{x}{\beta}\right)^{\alpha}}, \quad \alpha > 0, \beta > 0, x > 0 \quad (2)$$

METHODOLOGY

Complete Sample Status

Let $X_1, X_2, \dots, X_n$ be independent random variables having $Weibull(\alpha, \beta)$ distribution with α and β parameters. Then the log-likelihood function is given by

$$\ln L(\alpha, \beta) = \sum_{i=1}^n \ln f(x_i) \\ \propto n \ln(\alpha) - n\alpha \ln(\beta) - \left(\sum_{i=1}^n \left(\frac{x_i}{\beta} \right)^\alpha \right) + (\alpha - 1) \sum_{i=1}^n \ln x_i \quad (3)$$

The hypothesis to be tested here,

H_0 : The population probability distribution is Weibull

H_1 : is not

or is given by

$$H_0 : F_0(x) = 1 - e^{-\left(\frac{x}{\beta}\right)^\alpha} \\ H_1 : F_0(x) \neq 1 - e^{-\left(\frac{x}{\beta}\right)^\alpha} \quad (4)$$

Using the log-likelihood function given in Equation (3), Vasicek's test (TV), VanEs' test (TVE) and Correa's test (TC) statistics to be used for testing the mentioned hypothesis are as follows, respectively (Vasicek, 1976; Van Es, 1992; Correa, 1995).

$$TV_{mn} = -HV_{mn} - \frac{1}{n} \left(n \ln(\alpha) - n\alpha \ln(\beta) - \left(\sum_{i=1}^n \left(\frac{x_i}{\beta} \right)^\alpha \right) + (\alpha - 1) \sum_{i=1}^n \ln x_i \right) \\ = -HV_{mn} - \ln(\alpha) + \alpha \ln(\beta) + \frac{1}{n} \left(\sum_{i=1}^n \left(\frac{x_i}{\beta} \right)^\alpha \right) - \frac{(\alpha - 1)}{n} \sum_{i=1}^n \ln x_i \quad (5)$$

$$TVE_{mn} = -HVE_{mn} - \frac{1}{n} \left(n \ln(\alpha) - n\alpha \ln(\beta) - \left(\sum_{i=1}^n \left(\frac{x_i}{\beta} \right)^\alpha \right) + (\alpha - 1) \sum_{i=1}^n \ln x_i \right) \\ = -HVE_{mn} - \ln(\alpha) + \alpha \ln(\beta) + \frac{1}{n} \left(\sum_{i=1}^n \left(\frac{x_i}{\beta} \right)^\alpha \right) - \frac{(\alpha - 1)}{n} \sum_{i=1}^n \ln x_i \quad (6)$$

$$TC_{mn} = -HC_{mn} - \frac{1}{n} \left(n \ln(\alpha) - n\alpha \ln(\beta) - \left(\sum_{i=1}^n \left(\frac{x_i}{\beta} \right)^\alpha \right) + (\alpha - 1) \sum_{i=1}^n \ln x_i \right) \\ = -HC_{mn} - \ln(\alpha) + \alpha \ln(\beta) + \frac{1}{n} \left(\sum_{i=1}^n \left(\frac{x_i}{\beta} \right)^\alpha \right) - \frac{(\alpha - 1)}{n} \sum_{i=1}^n \ln x_i \quad (7)$$

Progressive Type Censored Sample Status

Let $X_{1:m:n}^R, X_{2:m:n}^R, \dots, X_{m:m:n}^R$ be independent random variables having progressive censored sample taken from $Weibull(\alpha, \beta)$ distribution with α and β parameters. Then likelihood and log-likelihood function are given by

$$L(\alpha, \beta) \propto \prod_{i=1}^m f(x_i) [1 - F(x_i)]^{R_i}$$

$$= \alpha^m \beta^{-\alpha m} e^{-\left(\left(\sum_{i=1}^m x_i\right)/\beta\right)^\alpha} \prod_{i=1}^m x_i^{\alpha-1} \left[e^{-\left(\left(\sum_{i=1}^m x_i\right)/\beta\right)^\alpha} \right]^{R_i} \quad (8)$$

$$\ln L(\alpha, \beta) \propto m \ln(\alpha) - m\alpha \ln(\beta) + (\alpha - 1) \sum_{i=1}^m \ln x_i - \left(\alpha \sum_{i=1}^m (R_i + 1) \left(\frac{x_i}{\beta} \right) \right) \quad (9)$$

Maximum likelihood estimators of the α and β parameters are obtained from the solution of the likelihood equations with respect to $\hat{\alpha}$ and $\hat{\beta}$.

$$\frac{d \ln L(\alpha, \beta)}{d\alpha} = \frac{m}{\alpha} - m \ln(\beta) + \sum_{i=1}^m \ln(X_{i:m:n}^R) - \left(\sum_{i=1}^k \left(\frac{(R_i + 1) X_{i:m:n}^R}{\beta} \right) \right) \bigg|_{\substack{\alpha=\hat{\alpha} \\ \beta=\hat{\beta}}} = 0 \quad (10)$$

$$\frac{d \ln L(\alpha, \beta)}{d\beta} = -\frac{m\alpha}{\beta} - \alpha \left(\sum_{i=1}^m \left(-\frac{(R_i + 1) X_{i:m:n}^R}{\beta^2} \right) \right) \bigg|_{\substack{\alpha=\hat{\alpha} \\ \beta=\hat{\beta}}} = 0 \quad (11)$$

Since analytical solutions of these equations are not available, using numerical methods and maximum likelihood estimates of their parameters can be obtained. Here, the hypotheses for the test of conformity to the Weibull distribution are as follows.

$$H_0 : F_0 = 1 - e^{-\left(\frac{x}{\beta}\right)^\alpha}$$

$$H_A : F_0 \neq 1 - e^{-\left(\frac{x}{\beta}\right)^\alpha} \quad (12)$$

Kullback-Leibler information on whether it is suitable for Weibull distribution for progressive types of censored samples is obtained as follows.

$$I_{1...m:m:n}(f : f^0) = -n\bar{H}_{1...m:m:n} - \left(m \ln(\alpha) - m\alpha \ln(\beta) + (\alpha - 1) \sum_{i=1}^m \ln x_i - \left(\alpha \sum_{i=1}^m (R_i + 1) \left(\frac{x_i}{\beta} \right) \right) \right)$$

$$= -n\bar{H}_{1...m:m:n} - m \ln(\alpha) + m\alpha \ln(\beta) - (\alpha - 1) \sum_{i=1}^m \ln x_i + \left(\alpha \sum_{i=1}^m (R_i + 1) \left(\frac{x_i}{\beta} \right) \right) \quad (13)$$

In the above Kullback-Leibler information, derivatives are taken according to α and β parameters and estimators are reached. The test statistic is obtained when the most likelihood estimators for the α and β parameters are replaced in the test statistic below. The Kullback-Leibler information test statistics under the progressive type censored sample is calculated as follows.

$$TA(w, n, m) = -\frac{1}{n} \sum_{i=1}^m \log \left\{ \frac{G(X_{(i+w:m:n)}; \hat{\theta}) - G(X_{(i-w:m:n)}; \hat{\theta})}{X_{i+w:m:n} - X_{i-w:m:n}} \right\} + \frac{1}{n} \sum_{i=1}^m R_i \log \left\{ \frac{1 - m/n}{1 - G(X_{(i:m:n)}; \hat{\theta})} \right\} \quad (14)$$

SIMULATION STUDY

Here, if the real distribution is Weibull, the best and worst powers are examined under different distributions and parameter states according to the choice of bandwidth w under the complete sample. When Table 1 is examined, for testing the suitability of a population with an $EE(2, 2)$ distribution to the Weibull distribution. The strengths of TV, TVE and TC tests against the best bandwidth w are almost equal and all of them are higher than the power of the K-S test, even if a little TC test is stronger than other tests, TV test cannot be performed against the worst w and this situation is important to choose w for TV test, The powers of TVE and TC tests are close to each other in all sample volumes examined, when bandwidth w is selected as the best and worst, i.e. the choice of w is not very important for these tests, in TC test the best bandwidth and the worst bandwidth is found to be 49.

Table 1. If the true distribution is $EE(2, 2)$, the best bandwidth w and corresponding powers (Complete Sample)

n	K-S	TV			TVE			TC		
	Power	w	Critical Value (%95)	Power	w	Critical Value (%95)	Power	w	Critical Value (%95)	Power
10	0,1101	4	-4,44328	0,3967	4	-6,07689	0,4056	4	-4,7538	0,4106
20	0,1840	6	-13,3113	0,6684	5	-15,0653	0,6674	1	-13,7446	0,6873
30	0,3014	12	-22,5728	0,8291	12	-23,8991	0,8310	4	-22,8763	0,8437
40	0,4369	12	-31,4763	0,9141	12	-32,9644	0,9099	3	-31,8961	0,9174
50	0,5629	22	-40,9008	0,9559	22	-42,1424	0,9569	2	-40,8942	0,9608
60	0,6752	21	-49,9107	0,9791	26	-51,1370	0,9797	6	-50,0415	0,9810
70	0,7684	25	-58,9897	0,9897	25	-60,3975	0,9896	2	-59,2969	0,9906
80	0,8420	22	-67,9241	0,9967	22	-69,5363	0,9960	22	-68,3531	0,9967
90	0,8901	32	-77,4760	0,9984	8	-79,0139	0,9989	8	-77,5309	0,9991
100	0,9287	49	-87,2728	0,9994	49	-88,2933	0,9995	49	-87,1407	0,9995

w: bandwidth

When Table 2 is examined, in testing the suitability of a population with $EE(2, 2)$, $EP(2, 2)$ or $EG(0.4, 2)$ distribution under the censored sample to the Weibull distribution, For $EE(2, 2)$ distribution, especially in the 13th and later censorship schemes, the choice of bandwidth w is important, since there is a difference between the best and worst powers, results similar to those obtained for distribution $EE(2, 2)$ are also obtained for the distribution, the powers for the $EE(2, 2)$ and $EP(2, 2)$ distributions are generally low, for the $EG(0.4, 2)$ distribution, there are significant changes in the powers according to the censorship schemes, and the powers are higher in the $R_m = 12$, $R_i = 0, i=1, 2, \dots, m-1$ shaped censorship schemes where all censorship is made when the last observation is obtained, Since there is a difference between the best and the worst powers for the $EG(0.4, 2)$ distribution, comments can be made such that the choice of bandwidth w is important. The progressive type censorship scheme used in Monte Carlo simulations is given in Table 3.

II. International Applied Statistics Conference (UYIK-2021)
Tokat / Turkey, 29 June - 2 July 2021

Table 2. Best bandwidth w and corresponding power obtained for different alternative distributions (Progressive Type Censored Sample)

Censoring Scheme No	EE(2,2)			EP(2,2)			EG(0.4,2)		
	w	Critical Value (%95)	Power	w	Critical Value (%95)	Power	w	Critical Value (%95)	Power
1	3	-0,2583	0,0828	3	-0,2583	0,0737	3	-0,2583	0,2369
2	2	0,0676	0,0618	2	0,0676	0,0553	2	0,0676	0,5336
3	3	-0,0142	0,0679	3	-0,0142	0,0590	2	0,0211	0,5745
4	4	-0,3198	0,0914	4	-0,3198	0,0775	5	-0,3341	0,1821
5	2	0,1010	0,0749	2	0,1010	0,0652	2	0,1010	0,6608
6	2	-0,1255	0,0761	2	-0,1255	0,0698	5	-0,1663	0,1133
7	6	-0,2955	0,0944	6	-0,2955	0,0818	7	-0,3044	0,1362
8	7	0,0280	0,0833	6	0,0423	0,0731	7	0,0280	0,5821
9	3	-0,1814	0,0747	3	-0,1814	0,0683	3	-0,1814	0,0580
10	4	-0,1866	0,0988	4	-0,1866	0,0790	4	-0,1866	0,4246
11	4	0,0131	0,0610	4	0,0131	0,0545	2	0,0487	0,7305
12	2	0,0242	0,0678	2	0,0242	0,0584	2	0,0242	0,5070
13	8	-0,3236	0,1123	7	-0,3195	0,0886	8	-0,3236	0,3062
14	8	0,0198	0,0762	8	0,0198	0,0626	2	0,0908	0,9493
15	3	-0,0415	0,1127	3	-0,0415	0,0945	3	-0,0415	0,1498
16	14	-0,3340	0,1204	14	-0,3340	0,0888	14	-0,3340	0,2148
17	14	0,0198	0,1109	14	0,0198	0,0892	14	0,0198	0,9469
18	13	-0,0783	0,1301	13	-0,0783	0,1097	14	-0,0804	0,8972
19	8	-0,2504	0,1220	6	-0,2430	0,0899	8	-0,2504	0,4713
20	7	0,0188	0,0710	4	0,0412	0,0592	2	0,0623	0,9768
21	4	-0,0003	0,0755	6	-0,0133	0,0630	2	0,0182	0,7337
22	16	-0,3536	0,1465	14	-0,3482	0,1078	17	-0,3543	0,2769
23	19	0,0149	0,1134	19	0,0149	0,0861	3	0,0818	0,9977
24	5	-0,1093	0,1286	6	-0,1127	0,1037	5	-0,1093	0,1294
25	23	-0,2920	0,1356	18	-0,2836	0,0954	23	-0,2920	0,2007
26	24	0,0195	0,1400	24	0,0195	0,1147	24	0,0195	0,9820
27	24	-0,0779	0,1514	24	-0,0779	0,1258	24	-0,0779	0,9227

EE:Exponentiated Exponential

EP:Exponential Poisson

EG:Exponential Geometric

Table 3. Progressive type censorship scheme used in Monte Carlo Simulation

Censoring Scheme No	n	m	$R = (R_1, R_2, \dots, R_m)$
[1]	20	8	$R_1 = 12, R_i = 0, i \neq 1$
[2]	20	8	$R_8 = 12, R_i = 0, i \neq 8$
[3]	20	8	$R_1 = 6, R_8 = 6, R_i = 0, i \neq 1, 8$
[4]	20	12	$R_1 = 8, R_i = 0, i \neq 1$
[5]	20	12	$R_{12} = 12, R_i = 0, i \neq 12$
[6]	20	12	$R_3 = R_5 = R_7 = R_9 = 2, R_i = 0, i \neq 3, 5, 7, 9$
[7]	20	16	$R_1 = 4, R_i = 0, i \neq 1$
[8]	20	16	$R_{16} = 4, R_i = 0, i \neq 16$
[9]	20	16	$R_5 = 4, R_i = 0, i \neq 5$
[10]	40	10	$R_1 = 30, R_i = 0, i \neq 1$
[11]	40	10	$R_{10} = 30, R_i = 0, i \neq 10$
[12]	40	10	$R_1 = R_5 = R_{10} = 10, R_i = 0, i \neq 1, 5, 10$
[13]	40	20	$R_1 = 20, R_i = 0, i \neq 1$
[14]	40	20	$R_{20} = 20, R_i = 0, i \neq 20$
[15]	40	20	$R_i = 1, i = 1, 2, \dots, 20$
[16]	40	30	$R_1 = 10, R_i = 0, i \neq 1$
[17]	40	30	$R_{30} = 10, R_i = 0, i \neq 30$
[18]	40	30	$R_1 = R_{30} = 5, R_i = 0, i \neq 1, 30$
[19]	60	20	$R_1 = 40, R_i = 0, i \neq 1$
[20]	60	20	$R_{20} = 40, R_i = 0, i \neq 20$
[21]	60	20	$R_1 = R_{20} = 10, R_{10} = 20, R_i = 0, i \neq 1, 10, 20$
[22]	60	40	$R_1 = 20, R_i = 0, i \neq 1$
[23]	60	40	$R_{40} = 20, R_i = 0, i \neq 40$
[24]	60	40	$R_{2i-1} = 1, R_{2i} = 0, i = 1, 2, \dots, 20$
[25]	60	50	$R_1 = 10, R_i = 0, i \neq 1$
[26]	60	50	$R_{50} = 10, R_i = 0, i \neq 50$
[27]	60	50	$R_1 = R_{50} = 5, R_i = 0, i \neq 1, 50$

CONCLUSION

In the case of complete sample, it was understood that the power of TV, TVE and TC tests were almost equal to the best bandwidth w in the test of conformity to the Weibull distribution under $EE(2, 2)$ distribution, and the strongest test of the TC test was the test with the least power. It was found that the power of all the tests compared was higher than the power of the K-S test. The TV test cannot be performed against the worst w and this situation is important for the selection of w for the TV test, the powers of the TVE and TC tests are close to each other in all sample volumes examined, when bandwidth w is selected as the best and worst. It has been determined that the choice of w is not very important for these tests. In the case of a censored sample, in testing the suitability of a population with an $EE(2, 2)$, $EP(2, 2)$ and $EG(0.4, 2)$ distribution and a population with a $EE(2, 10)$, $EP(3, 0.7)$ and

$EG(0.2, 0.7)$ distribution to the Weibull distribution, the powers for the $EE(2, 2)$, $EP(2, 2)$, $EE(2, 10)$ and $EP(3, 0.7)$ distributions are generally low, and the censorship schemes for the $EG(0.4, 2)$ and $EG(0.2, 0.7)$ distributions are It was observed that there were significant changes in the powers, the powers were higher when the censoring was made when the last observation was obtained, and also the choice of bandwidth w was important because there was a difference between the best and the worst powers for these distributions. As a result, in the case of complete and censored sampling, based on Kullback-Leibler divergence, critical value and power comparison was made under different parameters based on the best and worst bandwidths. In all comparisons, it was seen that the power of the test increased as the sample size increased.

References

- Van Es B, 1992. Estimating Functionals Related to a Density by a Class of Statistics Based on Spacings. Scandinavian Journal of Statistics, 19, 1, 61-72.
- Vasicek O, 1976. A Test for Normality Based on Sample Entropy. Journal of the Royal Statistical Society. Series B (Methodological), 38, 1, 54-9.
- Balakrishnan N, Aggarwala R, 2000. Progressive Censoring: Theory, Methods, and Applications, Boston: Birkhauser, p.
- Massey FJ, 1951. The Kolmogorov-Smirnov Test for Goodness of Fit. Journal of the American Statistical Association, 46, 253, 68-78.
- D'agostino RB, Stephens MA, 1986. Goodness of Fit Techniques, Newyork, ABD, MDI Press, p.
- Choi B, Kim K, Heun Song S, 2004. Goodness-of-Fit Test for Exponentiality Based on Kullback-Leibler Information. Communications in Statistics - Simulation and Computation, 33, 2, 525-36.
- Lim J, Park S, 2007. A Goodness of Fit Tests Based on the Partial Kullback-Leibler Information with the Type II Censored Data. Journal of Applied Statistics, 34, 9, 1051-64.
- Balakrishnan N, Rad AH, Arghami NR, 2007. Testing exponentiality based on Kullback-Leibler information with progressively type-II censored data. IEEE Transactions on Reliability, 56, 2, 301-7.
- Rad AH, Yousefzadeh F, Balakrishnan N, 2011. Goodness-of-Fit Test Based on Kullback-Leibler Information for Progressively Type-II Censored Data. IEEE Transactions on Reliability, 60, 3, 570-9.
- Weibull W, 1951. A Statistical Distribution Function of Wide Applicability. Journal of Applied Mechanics, 18, 293-7.

Acknowledgment

This study is presented as a PhD thesis by G.G, Selcuk University, Konya, TURKEY, 2020, December. Advisor: Prof. Dr. İsmail KINACI.

Conflicts of Interest

No conflict of interest was declared by the authors.

Authors' Contributions

Gülcan GENCER: Literature review, method design and planning, data collection and analysis, reporting. İsmail KINACI: Literature review, method design and planning, data collection and analysis, reporting.

Investigating the Effect of Drainage Area Characteristics on Gully Density in North of Iran

Ali Reza VAEZI¹, Hosein PANDI¹, Saniye DEMIR^{2*}

¹Department of Soil Science and Engineering, Faculty of Agriculture, University of Zanjan,

²TOGU University, Faculty of Agriculture Department of Soil Science, Tokat, 60100

*Presenter author e-mail: saniye.demir@gop.edu.tr

Abstract

Gully erosion is one of the most destructive types of erosion in soil resources, which causes a lot of damage by producing a large volume of sediment in watersheds. In order to manage soil resources and controlling soil erosion and sedimentation, identifying the areas with gullies in the watershed is very important. The present study was conducted to investigate the relationship between the gully density and the characteristics of the drainage area in a mountainous area of 28 km² in the west of the Rudbar township, located in the south of Guilan province in north Iran. Tward this, eighty gullies were identified in the area and some morphological characteristics of drainage area including slope gradient, shape indicators (Miller's coefficient and Gravius coefficient) and land use type (garden and rangeland) were determined. The results indicated that the gully density is positively correlated with the surface drainage area ($r= 0.28$, $p<0.05$), Miller's coefficient ($r= 0.29$, $p<0.01$) and the Gravius's coefficient ($r= 0.37$, $p<0.01$), while negative correlation was found for slope gradient ($r=-0.52$, $p<0.01$). No significant correlation was fund between the gully density and the land use type in each drainage area i.e. garden surface ($r=0.14$) and rangeland surface ($r= 0.11$). The general results show that the use of satellite images to monitor gully erosion reduces the difficulties of field study at the macro level. Also, in the drainage area with gentle slope and circular shape, conservation practices should be taken to prevent the development of gullies

Keywords: Gully erosion, Slope gradient, Shape coeffiicient, Land use, Conservation practices

POSTER PRESENTATIONS

POSTER PRESENTATIONS CONTENT

Differentiating Between Girls and Boys in Transition to Smoking Stages: A Sex-specific Growth Mixture Modeling.....	390
Applications of Proportional Odds Ordinal Logistic Regression Models and Continuation Ratio Models in Examining the Association of Physical Inactivity with Erectile Dysfunction Among Type 2 Diabetic Patients.....	391

Differentiating Between Girls and Boys in Transition to Smoking Stages: A Sex-specific Growth Mixture Modeling

Mohammad Asghari JAFARABADI^{1,2*}, Nasrin JAFARI², Asghar Mohammadpour ASL¹

¹Road Traffic Injury Research Center, Tabriz University of Medical Sciences, Tabriz, Iran

²Department Of Statistics And Epidemiology, Tabriz University of Medical Sciences, Tabriz, Iran

*Corresponding author e-mail: m.asghari862@gmail.com

Abstract

This study was done to differentiate between girls and boys in the transition to smoking stages utilizing the sex-specific growth mixture models (GMMs) among 4903 adolescents. By a two-point cohort study design and a multistage sampling procedure, illustrative samples of 10th-grade students were assessed in Tabriz, Iran. Smoking status, intention to start smoking, smoking during the past week/month were measured as the primary outcomes, which were collected by a valid and reliable instrument. GMMs with an optimal number of classes were fitted to estimate the intercepts and slopes along with covariates. Conclusively, GMMs lead in a 2-class optimal model: the "Occasional/Intending Smokers" and the "Non-Smokers" classes. Considering the boys in the second class as the referent, GMMs lead in significantly negative intercepts (range over -8.5 to -0.6), indicating that girls in both classes had lower outcome levels than the referent. Importantly, according to the GMMs' significantly and positive slopes (range over 0.2 to 2.7) for boys in both classes and girls in the second class, the positive transition occurred in the outcomes. and finally, the results of the study highlighted the importance of stopping the initiation and transition in smoking stages with special sex-specific planning for adolescent girls and boys.

Key words: Sex-Specific, Growth mixture models, Smoking stages, Adolescents, Transition

Applications of Proportional Odds Ordinal Logistic Regression Models and Continuation Ratio Models in Examining the Association of Physical Inactivity with Erectile Dysfunction Among Type 2 Diabetic Patients

Elbin SİBYI^{1*}

¹Inductive Quotient Analytics Pvt. Ltd. India, St. Thomas College Palai, Kerala, India

²Inductive Quotient Analytics Pvt. Ltd, India

*Corresponding author e-mail: elbinsiby97@gmail.com

Abstract

Many studies have observed a high prevalence of erectile dysfunction among individuals performing physical activity in less leisure-time. However, this relationship in patients with type 2 diabetic patients is not well studied. In exposure outcome studies with ordinal outcome variables, investigators often try to make the outcome variable dichotomous and lose information by collapsing categories. Several statistical models have been developed to make full use of all information in ordinal response data, but they have not been widely used in public health research. In this paper, we discuss the application of two statistical models to determine the association of physical inactivity with erectile dysfunction among patients with type 2 diabetes. A total of 204 married men aged 20-60 years with a diagnosis of type 2 diabetes at the outpatient unit of the Department of Endocrinology at PSG hospitals during the months of May and June 2019 were studied. We examined the association between physical inactivity and erectile dysfunction using proportional odds ordinal logistic regression models and continuation ratio models. The proportional odds model revealed that patients with diabetes who perform leisure time physical activity for over 40 minutes per day have reduced odds of erectile dysfunction (odds ratio = 0.38) across the severity categories of erectile dysfunction after adjusting for age and duration of diabetes. The present study suggests that physical inactivity has a negative impact on erectile function. We observed that the simple logistic regression model had only 75% efficiency compared to the proportional odds model used here; hence, more valid estimates were obtained here.

Key words: Physical inactivity, Erectile dysfunction, Ordinal regression, Diabetic patients, Continuation ratio model, Asymptotic efficiency

NATIONALITY OF PRESENTER AUTHORS

Country	n	%
Algeria	2	2.04
India	1	1.02
Iranian	3	3.06
Iraq	1	1.02
Pakistan	2	2.04
South Africa	2	2.04
Spain	1	1.02
Turkey	86	87.76
Total	98	100
