



#### **Sahibi / Owner**

Doç. Dr. M. Hanefi CALP

#### **Baş Editör / Editor in Chief**

Doç. Dr. M. Hanefi CALP

#### **Yardımcı Editörler / Co-Editors**

Doç. Dr. Ahmet DOĞAN

Doç. Dr. Serkan SAVAŞ

Dr. Öğr. Üyesi Ömer Çağrı YAVUZ

#### **Alan Editörleri / Field Editors**

Prof. Dr. Alptekin ERKOLLAR

Prof. Dr. Latif ÖZTÜRK

Prof. Dr. Türksel KAYA BENSGHİR

Prof. Dr. Tülay İLHAN NAS

Prof. Dr. Üstün ÖZEN

#### **Yayın Kurulu / Editorial Board**

Prof. Dr. Abdulkadir PEHLİVAN, Karadeniz Teknik Üniversitesi, Yönetim Biliřim Sistemleri

Prof. Dr. Ali HALICI, Bařkent Üniversitesi, Yönetim Biliřim Sistemleri

Prof. Dr. Aslıhan TÜFEKÇİ, Gazi Üniversitesi, Yönetim Biliřim Sistemleri

Prof. Dr. Bilal GÜNEŞ, Gazi Üniversitesi, Fizik Eğitimi

Prof. Dr. Bilal TOKLU, Gazi Üniversitesi, Endüstri Mühendislięi

Prof. Dr. Birgöl Kutlu BAYRAKTAR, Boğaziçi Üniversitesi, Yönetim Bilişim Sistemleri  
Prof. Dr. Bogdan PATRUT, Alexandru Ioan Cuza Üniversitesi, Matematik ve Bilgisayar Bilimleri  
Prof. Dr. Bünyamin ER, Karadeniz Teknik Üniversitesi, Yönetim Bilişim Sistemleri  
Prof. Dr. Cevriye GENCER, Gazi Üniversitesi, Endüstri Mühendisliği  
Prof. Dr. Cihan TANRIÖVEN, Ankara Hacı Bayram Veli Üniversitesi, İşletme  
Prof. Dr. Efendi NASİBOĞLU, Dokuz Eylül Üniversitesi, Bilgisayar Bilimleri  
Prof. Dr. Erdoğan DOĞDU, Çankaya Üniversitesi, Bilgisayar Mühendisliği  
Prof. Dr. Erman COŞKUN, Bakırçay Üniversitesi, Yönetim Bilişim Sistemleri  
Prof. Dr. Hadi GÖKÇEN, Gazi Üniversitesi, Endüstri Mühendisliği  
Prof. Dr. Halil İbrahim OKUMUŞ, Karadeniz Teknik Üniversitesi, Elektrik-Elektronik Mühendisliği  
Prof. Dr. Hamdi Tolga KAHRAMAN, Karadeniz Teknik Üniversitesi, Yazılım Mühendisliği  
Prof. Dr. Hasan Erdiñ KOÇER, Selçuk Üniversitesi, Elektrik-Elektronik Mühendisliği  
Prof. Dr. İly LevİN, Tel Aviv Üniversitesi, Bilim ve Teknoloji Eğitimi  
Prof. Dr. İsmail SARITAŞ, Selçuk Üniversitesi, Elektrik-Elektronik Mühendisliği  
Prof. Dr. İsmail ŞAHİN, Gazi Üniversitesi, Endüstriyel Tasarım Mühendisliği  
Prof. Dr. Kürşad ZORLU, Ankara Hacı Bayram Veli Üniversitesi, Yönetim Bilişim Sistemleri  
Prof. Dr. Latif ÖZTÜRK, Ankara Hacı Bayram Veli Üniversitesi, Yönetim Bilişim Sistemleri  
Prof. Dr. M. Ali AKCAYOL, Gazi Üniversitesi, Bilgisayar Mühendisliği  
Prof. Dr. M. Nihat SOLAKOĞLU, Çankaya Üniversitesi, Yönetim Bilişim Sistemleri  
Prof. Dr. Mehmet AKTAN, Necmettin Erbakan Üniversitesi, Endüstri Mühendisliği  
Prof. Dr. Mehmet BAŞ, Ankara Hacı Bayram Veli Üniversitesi, İşletme  
Prof. Dr. Meltem ÖZTURAN, Boğaziçi Üniversitesi, Yönetim Bilişim Sistemleri  
Prof. Dr. Metehan TOLON, Ankara Hacı Bayram Veli Üniversitesi, İşletme  
Prof. Dr. Murat DENER, Gazi Üniversitesi, Bilgisayar Mühendisliği  
Prof. Dr. Murat Paşa UYSAL, Başkent Üniversitesi, Yönetim Bilişim Sistemleri  
Prof. Dr. Nicu BIZON, Pitesti Üniversitesi, Elektronik, İletişim ve Bilgisayar Bilimleri  
Prof. Dr. Nursal ARICI, Gazi Üniversitesi, Yönetim Bilişim Sistemleri  
Prof. Dr. Oğuz KAYNAR, Sivas Cumhuriyet Üniversitesi, Yönetim Bilişim Sistemleri  
Prof. Dr. Rahmi CANAL, İnönü Üniversitesi, Biyomedikal Mühendisliği  
Prof. Dr. Sabri KOÇER, Necmettin Erbakan Üniversitesi, Bilgisayar Mühendisliği  
Prof. Dr. Selçuk KARAMAN, Ankara Hacı Bayram Veli Üniversitesi, Yönetim Bilişim Sistemleri  
Prof. Dr. Selçuk Kürşat İŞLEYEN, Gazi Üniversitesi, Endüstri Mühendisliği  
Prof. Dr. Sevinç GÜLSEÇEN, İstanbul Üniversitesi, Enformatik  
Prof. Dr. Shadi A. ALJAWARNEH, Jordan Üniversitesi, Bilim ve Teknoloji  
Prof. Dr. Suat ÖZDEMİR, Hacettepe Üniversitesi, Bilgisayar Mühendisliği  
Prof. Dr. Süleyman ERSÖZ, Kırıkkale Üniversitesi, Endüstri Mühendisliği  
Prof. Dr. Şükrü ÖZŞAHİN, Karadeniz Teknik Üniversitesi, Endüstri Mühendisliği  
Prof. Dr. Talip KELLEGÖZ, Gazi Üniversitesi, Endüstri Mühendisliği  
Prof. Dr. Tülay İlhan NAS, Karadeniz Teknik Üniversitesi, İşletme  
Prof. Dr. Türksel KAYA BENSGHIR, Ankara Hacı Bayram Veli Üniversitesi, İşletme  
Prof. Dr. Uğur YAVUZ, Atatürk Üniversitesi, Yönetim Bilişim Sistemleri

Prof. Dr. Üstün ÖZEN, Atatürk Üniversitesi, Yönetim Bilişim Sistemleri  
Prof. Dr. Yılmaz GÖKŞEN, Dokuz Eylül Üniversitesi, Yönetim Bilişim Sistemleri  
Doç. Dr. Gürcan ÇETİN, Muğla Sıtkı Koçma Üniversitesi, Bilişim Sistemleri Mühendisliği  
Doç. Dr. Ekrem BAHÇEKAPILI, Karadeniz Teknik Üniversitesi, Yönetim Bilişim Sistemleri  
Doç. Dr. Hakan ÖZKÖSE, Bartın Üniversitesi, Yönetim Bilişim Sistemleri  
Doç. Dr. Muhammet BERİGEL, Karadeniz Teknik Üniversitesi, Yönetim Bilişim Sistemleri  
Doç. Dr. Murat DÖRTERLER, Gazi Üniversitesi, Bilgisayar Mühendisliği  
Doç. Dr. Osman ÖZKARACA, Muğla Sıtkı Koçma Üniversitesi, Bilişim Sistemleri Mühendisliği  
Doç. Dr. Paolo TORRONI, Bologna Üniversitesi, Bilgisayar Bilimleri ve Mühendisliği  
Doç. Dr. Utku KÖSE, Süleyman Demirel Üniversitesi, Bilgisayar Mühendisliği  
Doç. Dr. Ümit ATİLA, Gazi Üniversitesi, Bilgisayar Mühendisliği  
Doç. Dr. Ahmet DOĞAN, Osmaniye Korkut Ata Üniversitesi, Yönetim Bilişim Sistemleri  
Dr. Öğr. Üyesi Bilgehan İMAMOĞLU, Karadeniz Teknik Üniversitesi, Yönetim Bilişim Sistemleri  
Dr. Öğr. Üyesi Emin Sertaç ARI, Osmaniye Korkut Ata Üniversitesi, Yönetim Bilişim Sistemleri  
Dr. Öğr. Üyesi Güler KARAMAN, Ankara Hacı Bayram Veli Üniversitesi, Yönetim Bilişim Sistemleri  
Dr. Öğr. Üyesi Mevlüt UYSAL, Gazi Üniversitesi, Yönetim Bilişim Sistemleri  
Dr. Öğr. Üyesi Mustafa TANRIVERDİ, Gazi Üniversitesi, Yönetim Bilişim Sistemleri  
Dr. Öğr. Üyesi Ömer Çağrı YAVUZ, Trabzon Üniversitesi, Yönetim Bilişim Sistemleri  
Dr. Iulian FURDU, Vasile Alecsandri Üniversitesi, Bilişim ve Eğitim Bilimleri  
Dr. Pandian VASANT, Teknoloci Petronas Üniversitesi, Bilişim Sistemleri  
Dr. Tomayess ISSA, Curtin Üniversitesi, Bilişim Sistemleri  
Arş. Gör. Berat TAHTABİÇEN, Ankara Hacı Bayram Veli Üniversitesi, Yönetim Bilişim Sistemleri  
Arş. Gör. Nadide Gizem GÜRSÖN DOLAR, Ankara Hacı Bayram Veli Üniversitesi, Yönetim Bilişim Sistemleri  
Arş. Gör. Mahmud Zahid MUTLU, Ankara Hacı Bayram Veli Üniversitesi, Yönetim Bilişim Sistemleri

#### **Teknik Koordinatör / Technical Coordinator**

Arş. Gör. Mahmud Zahid MUTLU

#### **Sekreterlik / Secretarial**

[mahmud.mutlu@hbv.edu.tr](mailto:mahmud.mutlu@hbv.edu.tr), [hanefi.calp@hbv.edu.tr](mailto:hanefi.calp@hbv.edu.tr)

Ankara Hacı Bayram Veli Üniversitesi, İktisadi ve İdari Bilimler Fakültesi, Yönetim Bilişim Sistemleri Bölümü,  
Emniyet Mahallesi, Muammer Bostancı Caddesi, No:4, 06500 Beşevler/Ankara, Türkiye

#### **İÇİNDEKİLER (Cilt: 7 / Sayı: 1 - Haziran 2025)**

|                                                                                                                                                            |       |
|------------------------------------------------------------------------------------------------------------------------------------------------------------|-------|
| Topic Modelling of Postgraduate Theses in the Field of Management Information Systems in Turkey<br>with LDA Algorithm<br><i>Göktaş İLİSU, Nursal ARICI</i> | 1-13  |
| Quality Management System Based Blockchain Applications<br><i>Hilal YILMAZ, Özlem ŞENVAR</i>                                                               | 14-28 |

|                                                                                                                                                                                         |         |
|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------|
| <i>MLP Temelli Yapay Zeka Modellerinde Çıktıları Etkileyen Giriş VeriSeti Niteliklerinin Model Mimarisine Göre Davranışlarının Araştırılması</i><br><i>Muhammer İLKUÇAR</i>             | 29-37   |
| <i>RSO Kullanıcılarının RSO'ya İlişkin Algısı</i><br><i>Selçuk KIRAN, Sena SÜZÜK</i>                                                                                                    | 38-51   |
| <i>An Investigation of International Field Indexes in the Discipline of Management Information Systems: A Bibliometric Analysis</i><br><i>Fatma AKGÜN, Cevriye GENCER, Hakan ÖZKÖSE</i> | 52-73   |
| <i>COVID-19'un Yayılım Modelini Tahmin Etmek İçin Hibrit GA-ConvLSTM Modeli: Şanghay İşbirliği Örgütü İçin Bir Vaka Çalışması</i><br><i>Anıl UTKU</i>                                   | 74-89   |
| <i>Kuantum Yazılım Test Teknikleri Ve Araçları Üzerine Bir İnceleme</i><br><i>Fatma Betül ÖZDEMİR, Savaş ÖZTÜRK</i>                                                                     | 90-104  |
| <i>Domates Yapraklarında Hastalık Tespiti İçin Transfer Öğrenme Kullanılması Ve Mobil Uygulamaya Entegre Edilmesi</i><br><i>Tahir Çağrı ÖZBEN, Osman GÜLER</i>                          | 105-113 |
| <i>Can Law Deter in Cyberspace? Türkiye's Experience in the Context of the Turkish Penal Code</i><br><i>Mehmet Onur ÖZER, Mehmet KAHRAMAN</i>                                           | 114-130 |



## Türkiye’de Yönetim Bilişim Sistemleri Alanında Yapılan Lisansüstü Tezlerin LDA Algoritması ile Konu Modellemesi\*

Göktuğ İLİSU<sup>a</sup>, Nursal ARICI<sup>\*.b</sup>

<sup>a</sup> Gazi Üniversitesi, Bilişim Enstitüsü, Bilişim Sistemleri, ANKARA, 06500, TÜRKİYE

<sup>b\*</sup> Gazi Üniversitesi, Uygulamalı Bilimler Fakültesi, Yönetim Bilişim Sistemleri Bölümü, ANKARA, 06500, TÜRKİYE

\* This article was produced from the Master’s thesis prepared by Göktuğ İlisu under the supervision of Prof. Dr. Nursal Arıcı.

### MAKALE BİLGİSİ

Alınma: 29.09.2024  
Kabul: 19.06.2025

#### Anahtar Kelimeler:

Yönetim bilişim sistemleri, konu modelleme, gizli dirichlet tahsisi algoritması

#### \*Sorumlu Yazar

e-posta:  
goktugilisul@gmail.com

### ÖZET

Yönetim Bilişim Sistemleri, işletmelerin ve kurumların stratejik, yönetsel ve operasyonel düzeylerdeki bilgi ihtiyaçlarını karşılamak için bilgi teknolojisi çözümlerini ve iş süreçlerini entegre etmeye odaklanan bir bilim alanıdır. Bu yönüyle bilgisayar bilimi, yönetim bilimi, istatistik, organizasyon teorisi, karar teorisi gibi çeşitli referans disiplinlerden beslenen çok disiplinli bir araştırma alanıdır. Bu çalışmanın temel amacı, Yönetim Bilişim Sistemleri bilim dalının Türkiye’deki lisansüstü tez konularına yansımalarını ve gelişimini incelemektir. Bu amaçla, 2002-2023 yılları arasında yönetim bilişim sistemleri alanında hazırlanan ve YÖK Ulusal Tez Merkezi web sitesi üzerinden erişilebilen 1070 lisansüstü tez (Yüksek Lisans-f: 951 ve Doktora-f: 119) inceleme kapsamına alınarak Gizli Dirichlet Tahsisi algoritmasıyla konu modellemesi gerçekleştirilmiştir. Konu modellemesinde kullanılan veri seti, lisansüstü tezlerin İngilizce özetleridir. Tez özetlerine öncelikle metin ön işleme ve kök çözümlemesi uygulanmıştır. Ortaya çıkan tüm kelimeler iç içe listelere dönüştürülüp LDA algoritması uygulanarak konu modelleri elde edilmiştir. Veri görselleştirme ile kelime bulutları, konu kümeleri, kelime sıklık histogramları ve belge- konu dağılımları oluşturulmuştur. Konu modellerinde çoğunlukla “data”, “model”, “research”, “technology”, “system” kelimelerinin yer aldığı tespit edilmiştir. Bu kelimelerin sıklıkla kullanıldığının tespit edilmesi yönetim bilişim sistemleri bilim dalında çalışılan konular için beklenen bir sonuç olarak değerlendirilmektedir.

DOI: 10.59940/jismar.1557818

## Topic Modelling of Postgraduate Theses in the Field of Management Information Systems in Turkey with LDA Algorithm

### ARTICLE INFO

Received: 29.09.2024  
Accepted: 19.06.2025

#### Keywords:

Management information systems, topic modelling, latent dirichlet allocation algorithm

#### \*Corresponding Authors

e-mail:  
goktugilisul@gmail.com

### ABSTRACT

Management Information Systems is a branch of science that deals with integrating information technology solutions and business processes to meet the information requirements of businesses and organizations at strategic, managerial, and operational levels. In this respect, it is a multidisciplinary research field that draws upon several different disciplines, including computer science, management science, statistics, organization theory and decision theory. The primary objective of this study is to investigate the reflections and improvement of the field of Management Information Systems as evidenced by graduate thesis topics in Turkey. To this end, 1,070 graduate theses (951 MSc-f and 119 PhD-f) prepared in the field of Management Information Systems between 2002 and 2023 and accessible via the CoHE National Thesis Centre website were included in the scope of the study. Topic modeling was performed with the Latent Dirichlet Allocation algorithm. The data set employed in the topic modelling process comprised the English abstracts of graduate theses. The thesis abstracts were subjected to text preprocessing and stemming. The resulting words were converted into nested lists and topic models were obtained by applying the LDA algorithm. Data visualization was employed to create word clouds, topic clusters, word frequency histograms and document-topic distributions. It was established that the terms “data,” “model,” “research,” “technology,” and “system” were predominantly incorporated

into the topic models. The fact that these words are frequently used is considered as an expected result for the subjects studied in the field of management information systems.

DOI: 10.59940/jismar.1557818

## 1. INTRODUCTION (GİRİŞ)

Management information systems is a scientific discipline that centers upon providing solutions to the knowledge requirements of states, societies, organizations, groups, and individuals at strategic, managerial, and operational levels through information systems and technologies. It is a multidisciplinary field that develops methodologies to find a way out to real-life problems in an application-oriented framework and to rule information technology resources most appropriately by drawing on various reference study fields such as computer science, management science, statistics, organization theory, and operations research. The field is also concerned with issues in sociology, economics, and psychology, such as the utilization and influence of information technology [1].

Management Information Systems (MIS) is a study field that focuses on the strategic and practical use of technology to improve organizational performance. It lies at the intersection of business and technology, aiming to facilitate the flow of information within an institution to promote decision-making, coordination, control, analytics, and visualization of data [2].

Text mining significantly enhances the capabilities of management information systems by providing the required tools to analyze unstructured text data, leading to better decision-making, increased efficiency and a deeper understanding of both internal operations and external environments. MIS platforms facilitate the collection and storage of large scales of text data from several sources. Integrating text mining tools and techniques into MIS permits the processing and analysis of text data to generate actionable insights. Combining text mining with real-time data processing talents in MIS can provide up-to-date insights for timely decision-making [3].

### 1.1. Text Mining (Metin Madenciliği)

Scientific literature and documents from marketing and economic sectors are frequently gathered as extensive text data. Additionally, large datasets can be collected in semi-structured formats, such as log files from servers and networks. In this scenario, text mining analysis is highly useful for both unstructured and semi-structured textual data. Although text mining is akin to data mining, it specifically concentrates on text analysis rather than structured data [4].

Text mining, in other words, knowledge discovery, involves the extraction of valuable information from textual data. This field, called text analytics or natural language processing, integrates computer science, linguistics, statistics, and machine learning techniques to derive meaningful insights from unstructured text. Unstructured text data encompasses any text that lacks a predefined format or structure, such as emails, social media content, articles, and customer feedback [5].

#### 1.1.1. Text preprocessing (Metin ön işleme)

Before analysis, text data is typically subjected to preprocessing procedures to enhance its quality and ensure uniformity across different datasets. Text preprocessing represents a fundamental aspect of numerous text mining algorithms. A conventional text classification framework typically involves four main stages: preprocessing, feature extraction, feature selection, and classification. The research indicates that the effectiveness of the classification process is heavily influenced by the methods used in feature extraction, feature selection, and the choice of classification algorithm. However, the preprocessing stage has also been found to have a noticeable impact on the success of this process. Uysal et al. investigated the impact of preprocessing tasks, with a particular focus on their influence in the field of text classification. The preprocessing step typically comprises a series of tasks, including tokenization, filtering, lemmatization and stemming [6].

Tokenization is dividing a sequence of characters into discrete units, called tokens, which may include words or sentences. Punctuation marks may also be discarded. The resulting list of tokens is then used for further processing. The main aim here is to pinpoint individual words within a sentence. Effective text classification and mining rely heavily on a robust parser capable of accurately tokenizing documents.

The process of filtering typically involves the removal of specific words or phrases from documents. A common practice in text filtering is the removal of so-called "stop words." Stop words are lexical items that occur with high frequency in a text but lack significant content-bearing information. Examples of such words include prepositions and conjunctions. Similarly, words that appear with considerable frequency in the text are identified as having minimal information content and thus being unable to distinguish between different documents. Furthermore, words that appear infrequently are likely to be irrelevant and can be excluded from the documents [7].

Different capitalization patterns are used in the creation of text and document data points, forming sentences. Since documents contain many sentences, inconsistent capitalization can significantly complicate the classification of extensive documents. One standard method to tackle this issue is to convert all letters to lowercase, which effectively brings all words in the text and document into a consistent property space. Punctuation marks and private symbols are also excluded from the sentences because they can challenge classification algorithms.

Lemmatization examines words based on their morphological makeup, consolidating several inflexive forms of a word into a unique entity for analysis. In essence, stemming techniques strive to normalize verbs to their root forms and standardize nouns into a consistent format [5].

The objective of stemming methods is to identify the stem of derived words. The specific stemming algorithms employed vary depending on the language in question. In the case of English, the stemmer algorithm is a commonly utilized approach. Text stemming involves modifying words to generate different word forms by applying diverse linguistic operations like affixation (the attachment of prefixes and suffixes) [7].

### 1.1.2. Feature selection (Özellik seçimi)

After preprocessing, the text must be transformed into a numerical pattern suitable for machine learning analysis. A promising approach proposes that incorporating both syntactic and semantic features into text representations can be very effective for sentence selection, particularly in technical genomic texts. Another method to tackle syntactic challenges is to use the n-gram technique for feature extraction [8].

The n-gram technique involves identifying sequences of n-letters appearing in a specific order within a given text corpus. This method does not only serve as a direct representation of the text, but also rather functions as a feature for text representation. The Bag of Words (BOW) model represents text by using individual words non-sequentially. This model is simple to implement, and the text is represented by a vector, typically with a manageable dimensionality. An n-gram, in this context, is a BOW feature used to represent text through sequences of words. The use of two-letter and three-letter combinations is common. This approach allows the extracted text feature to detect more information than a single word [9].

In natural language processing, term frequency (TF) is a statistical metric used to determine how frequently a specific term or word appears in a corpus. The simplest form of weighted feature extraction involves TF, where each term is assigned a value based on its count throughout the corpus. More advanced approaches that build on TF often apply binary or logarithmic scaling to word frequencies for weighting. In these methods, documents are converted into vectors that represent word frequencies. While this technique is easy to understand, it can be limited by the overrepresentation of common words in the feature vectors [10].

The bag-of-words (BoW) model provides a streamlined and simplistic depiction of a text document by extracting key features such as word frequency. This approach finds applications in many areas, including document classification, information retrieval, computer vision, natural language processing (NLP), Bayesian spam filtering, and machine learning. In the BoW framework, a text—whether a document or a sentence—is represented as a collection of individual words. During the BoW process, word lists are generated. These words are not the elements that make up sentences and grammar; they are simply listed in a matrix without considering their semantic relationship. While the order of appearance and grammatical structure are ignored, the focal points of documents are still identified [11].

K. Sparck Jones proposed the concept of Inverse Document Frequency (IDF) to mitigate the influence of words that are inherently prevalent in a given sentence [11]. IDF dedicates higher weights to words that are either very frequent or infrequent across documents. The integration of Term Frequency (TF) and IDF is known as Term Frequency-Inverse Document Frequency (TF-IDF). The mathematical formula for calculating the weight of a term in a document using TF-IDF is expressed in equation (1).

$$W(d, t) = TF(d, t) * \log\left(\frac{N}{df(t)}\right) \quad (1)$$

In this framework, N denotes the overall documents, while df(t) indicates the number of documents that feature the term t. The initial component of the equation improves recall, whereas the latter component enhances the precision of the term's representation. Although TF-IDF helps reduce the impact of frequently occurring terms, it has its drawbacks. It fails to account for the semantic relationships between words within a document, treating each word in isolation. However, recent advancements in modeling, such as word embeddings,

offer new methods for capturing word similarities and integrating part-of-speech information [12].

## 1.2. Topic Modelling (Konu Modelleme)

A significant proportion of the literature is digitized and stored electronically in databases, either through digital libraries or social network databases. It is therefore necessary to have access to powerful automated tools to read this data and to realize the underlying themes. A significant pivotal objective of data evaluation is to discern the attributes that data entries exhibit in common. In the field of text analytics, this frequently entails the identification of the situations or constructs that a given document addresses. Although this information is intuitively understood by human readers, computer programs interpret text in its literal form. To address this challenge in programming, data scientists utilize topic modeling. Topic modeling is a widely adopted technique in text mining that reveals underlying patterns within large datasets. Though it is particularly effective for analyzing textual data, it is also valuable in fields such as bioinformatics, social science, and environmental studies. This approach helps structure extensive datasets, facilitating easier navigation and analysis [13].

The ability to derive valuable statistics and features from a dataset depends heavily on selecting the right methods. While contemporary topic modeling techniques greatly surpass earlier algorithms, they still need to be fine-tuned and optimized to ensure accurate results. Various topic modeling approaches are tailored to handle specific types of data relationships and structures, including short texts, long sequences, highly correlated information, and data with intricate structural patterns. To develop a topic modeling process that effectively meets the needs of a data analysis project, it is essential to grasp the distinctions between different models and the foundational algorithms of them [14].

### 1.2.1. Classification of Topic Modelling (Konu Modelleme Sınıflandırması)

Topic modelling is a statistical technique used to uncover the latent "topics" within a collection of documents. As a subset of unsupervised machine learning and natural language processing (NLP), it seeks to classify and structure extensive text corpora by identifying recurring patterns, themes, and structures. By applying a topic modelling algorithm to preprocessed data, one can discover these underlying patterns and topics. Notable algorithms in this field include Latent Dirichlet Allocation (LDA) and Non-negative Matrix Factorization (NMF).

Latent Dirichlet Allocation (LDA) is a highly regarded method in topic modeling. It operates on the premise that documents are composed of a combination of topics, each of which is a collection of words. LDA functions by iteratively dedicating words to topics based on their co-occurrence patterns within the documents, gradually refining the association between words and topics [15].

Non-Negative Matrix Factorization (NMF) presents an alternative approach for uncovering topics within a corpus. It achieves this by decomposing the document-term matrix into two distinct lower-dimensional matrices: one matrix that captures the underlying topics and another that reflects how documents relate to these topics [16].

### 1.2.2. Topic modelling with LDA algorithm (LDA algoritmasıyla konu modelleme)

In the realm of natural language processing (NLP) and machine learning, topic modelling has emerged as a powerful tool for uncovering hidden themes and patterns within large text corpora. Among the various algorithms used for topic modelling, Latent Dirichlet Allocation (LDA) stands out as one of the most influential and largely used methods.

Latent Dirichlet Allocation (LDA) is a generative probabilistic model that postulates that each document in a corpus is a mixture of distinct topics, with each topic itself a mixture of words. The fundamental idea behind LDA is to reverse-engineer this generative process to uncover the hidden topic structure within the documents [17].

The Dirichlet distribution is a key component of LDA, serving as a prior distribution over the topic distributions in documents and the word distributions in topics. It is parameterized by a vector of positive reals and ensures that the resulting distributions are proper probability distributions. For document-topic distributions, the Dirichlet prior is denoted by  $\alpha$ , and for topic-word distributions, it is denoted by  $\beta$ . These hyperparameters influence the sparsity of the distributions, with smaller values leading to sparser distributions.

The generation process essentially models how a set of words in a document can be generated given a set of topics. By applying Bayesian inference, LDA aims to reverse this process to discover the hidden topic structure. LDA creation process is shown in Figure 1 [18].

Text preprocessing is applied as tokenization, stop words removal, and lemmatization to the text dataset. Following preprocessing, textual data is typically transformed into a numerical format using a document-term matrix. Here, each document corresponds to a row, and each word corresponds to a column in the matrix.

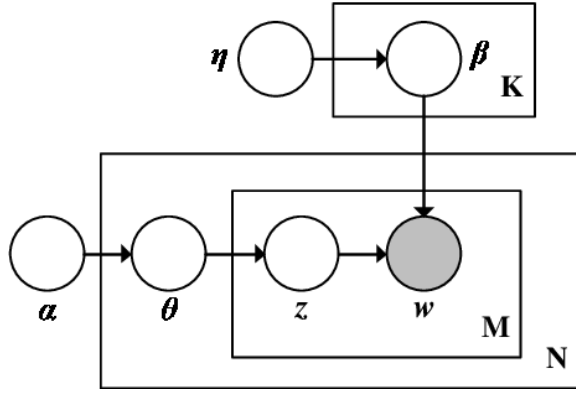


Figure 1. LDA creation process [18]  
(LDA üretim süreci)

The number of topics is determined and LDA is applied to the preprocessed data by using the Gensim library of Python which provides an efficient and easy-to-use interface for running LDA.

The quality of the topics is evaluated using metrics such as perplexity (a metric for the model's sample prediction) and coherence score (an evaluation of the semantic similarity among high-probability words within a topic). Based on the evaluation results, the hyperparameters and the number of topics are also tuned to achieve the best performance.

Once the model converges, the topics are interpreted by examining the most probable words in each topic and the topic distribution is also interpreted for each document. Finally, the topics are visualized using tools like word clouds, topic distribution graphs, or interactive plots (e.g., LDAvis) to gain insights into the underlying themes in the corpus [19].

Latent Dirichlet Allocation (LDA) is a robust method for uncovering hidden topics in extensive text collections, offering insights into the underlying structure and themes of the data.

By leveraging the principles of Bayesian inference and Dirichlet distributions, LDA provides a robust framework for topic modelling, with wide-ranging applications in content analysis, recommendation systems, information retrieval, and beyond. Through careful implementation, evaluation, and

interpretation, LDA can transform unstructured text data into meaningful and actionable knowledge [20].

## 2. LITERATURE REVIEW (LİTERATÜR TARAMASI)

A review of the literature revealed the existence of several studies on topic modelling with the LDA algorithm. However, there is a paucity of studies that have been conducted with the objective of conducting a quantitative and qualitative evaluation of postgraduate theses in terms of subject matter. The following section will provide an overview of these studies.

The research conducted by Çallı et al. [21] analyzed 574 graduate thesis abstracts completed between 2002 and 2020 in the MIS department in Turkey. The abstracts were examined using a text mining technique called the Latent Dirichlet Allocation algorithm. The analysis yielded 11 clusters, which were identified as follows: e-commerce and Marketing, System Development and Effects, Effects of Information Systems on Organizations, Data Mining, Human Resources Management, Organizational Change, Field Specific Studies I, Field Specific Studies II, Security, Education and Training, Prediction and Decision Support. In the context of this research, the similarities and differences between the estimation results and those presented in the national and international literature were discussed. This study aims to offer researchers in the field of management information systems insights and direction.

Parlina and Kusumarani [22] sought to employ bibliometric analysis to examine the intellectual structure and thematic development of the field of management information systems (MIS). A comprehensive analysis of the characteristics of publications in the three most prominent MIS journals in the SCOPUS database (IJIM, JSIS, and MIS Quarterly: Management Information Systems) was conducted, spanning the period from 1980 to 2021. In this study, the latent Dirichlet distribution (LDA) is incorporated into the approach to extend and improve the scientific research on MIS, resulting in a more comprehensive and up-to-date analysis. The indications of the study demonstrate the trend of publishing articles, the scientific structure, and the prominent issues in the top three journals.

Özköse and Gencer [23] conducted a comprehensive analysis of the field of Management Information Systems (MIS) through the use of bibliometric mapping. To achieve this objective, 222 journals that are indexed in the Science Citation Index Expanded (SCI-E) and the Social Science Citation Index (SSCI) were selected from the Web of Science and Scopus

databases. To determine the corpus of journals, 24 journals were selected for analysis, with the input of experts who could provide a more nuanced interpretation of the field. Initially, 20,497 English-language articles from these journals were gathered from the Web of Science (WoS) Core Collection between the years 1980 and 2015. Following text mining, the most influential organizations, authors and countries are displayed on graphs with statistical analysis using BibExcel. Furthermore, the development per annum of published articles is illustrated, and a trend analysis of these articles is presented. Additionally, the most cited articles are provided. Subsequently, using VosViewer, the most pertinent terms in this field were extracted through co-occurrence analysis from abstracts and keywords. The terms and their clusters are displayed on a graph. Density maps were also employed. The graphs and density maps are interpreted in detail, respectively.

The objective of this study is to analyze the master's and doctoral theses published in the field of MIS in universities in Turkey and the TRNC by text mining. Topic modelling with the LDA algorithm is employed as a method, although a multitude of data visualization techniques are utilised.

### 3. METHODOLOGY (YÖNTEM)

The aim of this study is to perform topic modelling of postgraduate theses prepared in the field of management information systems in universities in Turkey by text mining. Latent Dirichlet Allocation (LDA) algorithm is used in topic modelling for this purpose.

#### 3.1. Dataset Creation Process (Veri Seti Oluşturma Süreci)

The data set consists of 1070 postgraduate theses (Master's-f: 951 and Doctorate-f: 119) prepared in the field of Management Information Systems between 2002 and 2023 and accessible through the YÖK National Thesis Centre website [24].

The YÖK National Thesis Centre website was accessed and the detailed search section was selected via the link <https://tez.yok.gov.tr/UlusalTezMerkezi/>. Subsequently, the term "Management Information Systems" was entered into the primary discipline, branch of science and subject sections. This process was repeated for each of these occasions. The dataset was created by this way as an Excel table. To make the dataset suitable for use in Python, the relevant dataset was organized so that Turkish and English topics, Turkish and English thesis names, Turkish and

English keywords, and thesis abstracts were included in separate columns of the Excel table.

#### 3.2. Text Preprocessing (Metin Önışleme)

For the post graduate theses published in the field of MIS, there are thesis abstracts in both Turkish and English on the YÖK National Thesis Centre website. Within the scope of the study, text preprocessing was performed based on the English abstracts of these theses and the Subject (English) column of the data structure related to the organized data set. The text preprocessing process was applied with Python before the subject models, which are intended to be created by using only thesis abstracts as a data structure. Through the Python Re library, all punctuation marks and numbers in the English abstracts of the post graduate theses were removed. The content of all remaining texts was converted into lower case letters.

Stopwords of the English language are accessible through NLTK library of Python. In addition to the aforementioned stop words, the NLTK library also includes a list of words that have been evaluated as having lost their meaning in English thesis abstracts. These include:

['In', 'using', 'used', 'also', 'however', 'since', 'via', 'within', 'although', 'among', 'besides', 'whereas', 'dont', 'u', 'can', 'non', 'thus', 'may', 'towards', 'according', 'study', 'thesis', 'one', 'result', 'obtained', 'different', 'many', 'first', 'second', 'third', 'important', 'use', 'along', 'therefore', 'around', 'moreover', 'furthermore', 'nevertheless', 'whether', 'with', 'without', 'could', 'would', 'should', 'often', 'fourth', 'fifth', 'sixth', 'always', 'generally', 'sometimes', 'never', 'whenever', 'hence', 'across', 'thereby', 'thesis', 'before', 'after', 'meanwhile']

The Python Gensim library offers a simple preprocess function that can be used to break down text into words. In this case, the function was used to tokenize the abstracts. The stop words were then extracted from the English thesis abstracts using the simple preprocess function. A nested list was created for all remaining words. Each sub-list in the nested list consists of the words in a thesis abstract that have been removed from all stop words.

The English natural language processing model, provided by the Spacy library, was initially loaded using the command `nlp = spacy.load('en_core_web_sm')`. Subsequently, the words within each sub-list, which constituted the nested list, underwent lemmatization.



various LDA models. It provides insight into the extent to which the topics produced by the models are analogous or disparate. The corpus distance is typically calculated using measures such as the Kullback-Leibler (KL) divergence or the Jensen-Shannon divergence. These measures quantify the discrepancy between two probability distributions. For instance, the KL divergence between the topics of two LDA models indicates the extent of their divergence.

The formula for KL divergence, where P and Q are the probability distributions obtained from two different LDA models, is as given in equation (3).

$$D_{KL}(P \parallel Q) = \sum_i \log \frac{P(i)}{Q(i)} \quad (3)$$

The combined use of perplexity, compatibility, exclusivity and corpus distance measures helps to select the best model in subject modelling processes and to comprehensively evaluate the performance of the model [17].

#### 4. TOPIC MODELLING WITH LDA ALGORITHM (LDA ALGORİTASIYLA KONU MODELLEME)

In this section, the perplexity, corpus distance, coherence and exclusivity values were calculated using the Python programming language, and graphs were created for each measure. To calculate these values, the CoherenceModel, LdaModel, similarities and TfidfModel functions of the Python gensim library were installed. Subsequently, the corpus was transformed with TF-IDF values. The minimum number of topics to be formed was determined to be two, while the maximum number of topics was determined to be twenty. The LDA model was trained for different numbers of topics with the LDA model function. Finally, graphs for perplexity, corpus distance, compatibility and exclusivity values were created. Related graph is presented in Figure 3.

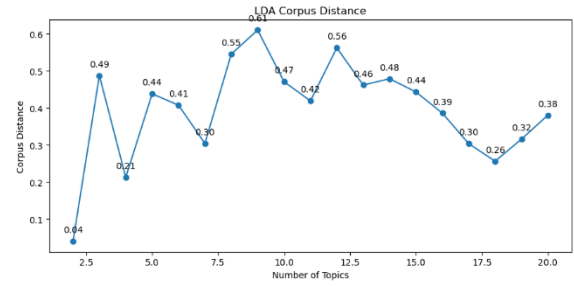
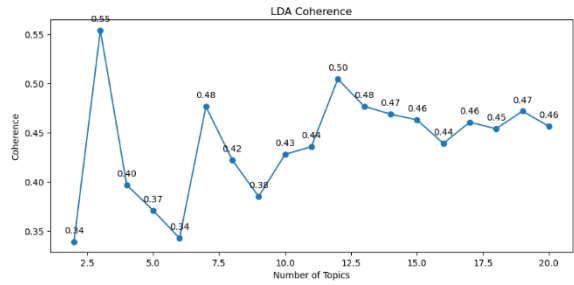
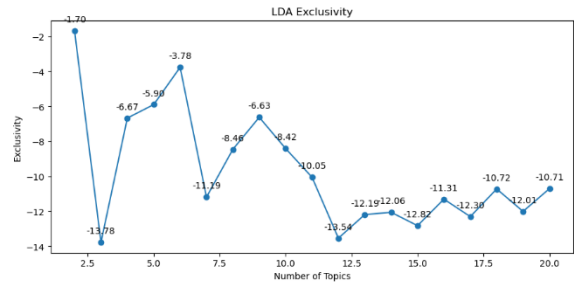
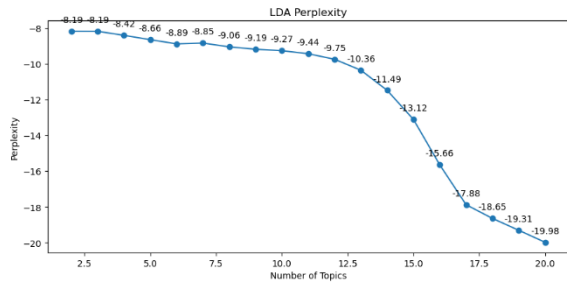


Figure 3. LDA topic modelling measures  
(LDA konu modelleme ölçüleri)

Once the measures of perplexity, corpus distance, compatibility and exclusivity have been determined, the coherence values of the LDA models for a given set of topics are analyzed in order to ascertain the number of topics with the lowest compatibility value. This process is employed to ascertain the optimal number of topics that most accurately represent the text data. The coherence value is typically a measure of the consistency and meaningfulness of the identified topics. Consequently, determining the optimal number of topics represents a crucial step in the text analysis and model evaluation process. At this stage, the optimal number of topics was automatically determined by Python code based on the coherence value. The number of topics is assigned as “2” automatically by Python. Thus, there exist two topics by using English abstracts of postgraduate thesis as a dataset.

#### 4.1. Creating LDA Topic Model (LDA Konu Modelinin Oluşturulması)

In the LDA topic model, which is suitable for the data structure in which postgraduate English thesis abstracts were used, the number of topics was automatically determined as 2. Subsequently, the corpus, number of topics and word identification number were used through the Python pprint library and the LdaMulticore function of the gensim library, resulting in the creation of the LDA topic model. A visual representation of the resulting topic model is presented in Figure 4.

```
[('0,
'0.013*datum" + 0.011*information" + 0.011*technology" + 0.010*analysis" ,
'+ 0.010*method" + 0.010*system" + 0.008*model" + 0.008*process" + '
'0.008*research" + 0.007*application'),
(1,
'0.012*datum" + 0.011*research" + 0.009*system" + 0.009*technology" + '
'0.009*information" + 0.008*model" + 0.008*process" + 0.008*method" + '
'0.008*analysis" + 0.007*development"')]
```

Figure 4. LDA Topic Model  
(LDA Konu Modeli)

As illustrated in Figure 4, the LDA topic model identifies two primary topics: Topic 0, which encompasses “Data and Information Technology Analysis” and Topic 1, which pertains to “Research and Development in Information Systems”.

#### 4.2. Visualization of LDA Topic Model (LDA Konu Modelinin Görselleştirilmesi)

Following the generation of the topic model, a visual representation of the LDA topic model was produced using the gensimvis function of the Python pyLDAvis library. This is presented in Figure 5.

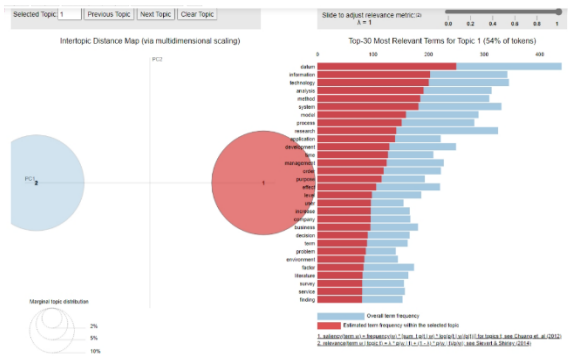


Figure 5. LDA topic model visualization when  $\lambda = 1$   
( $\lambda = 1$  iken LDA konu modeli görselleştirilmesi)

When the scroll bar is completely to the right ( $\lambda = 1$ ), the words in Topic 0 (Data and Information Technology Analysis) are shown in Figure 5 as an example.

The lambda ( $\lambda$ ) used in the PyLDAvis visualization is a tool for effectively exploring the results of the LDA topic model. In the PyLDAvis interface, the variable  $\lambda$ , which is used to determine how relevant a particular topic is to a particular word, can be set between 0 and 1. The properties of the variable  $\lambda$  are as follows [25]:

- $\lambda = 1$  means that the general frequency of words is more prominent.
- If  $\lambda = 0$ , the specific relevance of words to a particular topic is more prominent.

A shift to the left of the scroll bar ( $\lambda < 1$ ) results in alterations to the position and frequency of words within the topics. This phenomenon is exemplified in Figure 6.

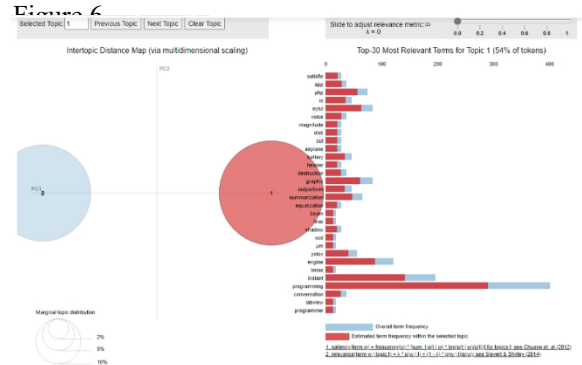
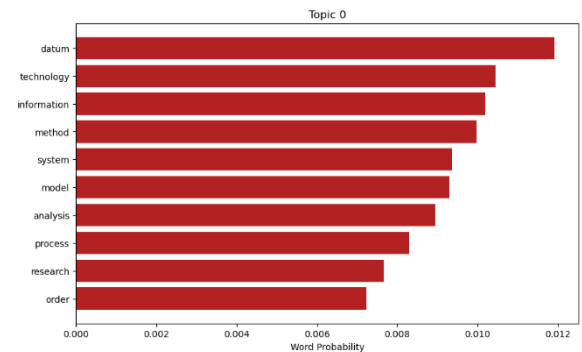


Figure 6. LDA topic model visualization when  $\lambda = 0$   
( $\lambda = 0$  iken LDA konu modeli görselleştirilmesi)

In the LDA topic model, which is represented as a cluster, Topic 0 is the predominant topic within the model, as it is the cluster with the largest area (54% of tokens are represented).

#### 4.3. Displaying Topics with a Bar Graph (Konuların sütun grafiğiyle görüntülenmesi)

Matplotlib library of Python is used to create bar graphs to express the most used ten words in topics of the model. These bar graphs are illustrated in Figure 7 as a visualization of Figure 4.



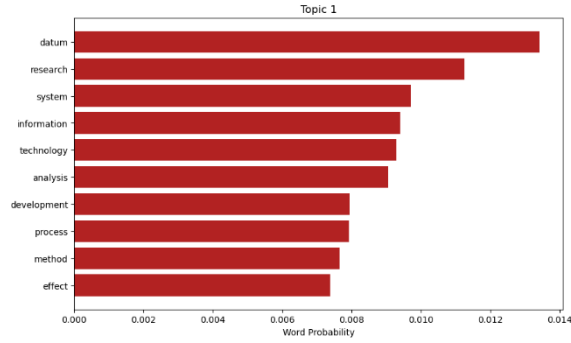


Figure 7. LDA topic model bar graphs  
(LDA konu modeli bar grafikleri)

#### 4.4. TF, IDF, TF-IDF values for postgraduate thesis abstracts (Lisansüstü Tezlerle ilişkin TF, IDF ve TF-IDF Değerleri)

The TF, IDF and TF-IDF values are presented in Figure 8, with the graph displaying the 20 words with the highest term frequency, inverse document frequency and term frequency-inverse document frequency values, as observed in postgraduate thesis abstracts.

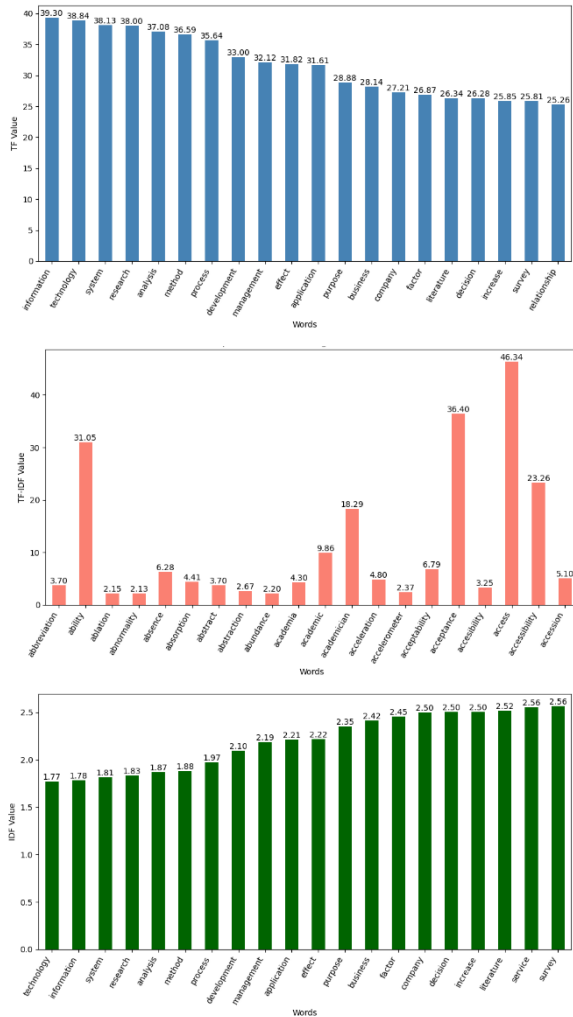


Figure 8. Words with the highest TF, IDF, TF-IDF values in postgraduate thesis abstracts  
(Lisansüstü tez özetlerinde en yüksek TF, IDF, TF-IDF değerlerine sahip olan kelimeler)

Figure 8 illustrates that the word with the highest term frequency is 'information', with a TF value of 39,29. Conversely, the word with the highest inverse document frequency is 'service' and 'survey', with an IDF value of 2,56. The term 'access' had the highest term frequency-inverse document frequency (TF-IDF) score, recorded at 46,34.

#### 4.5. Creating Word Clouds for the Topics in the LDA Topic Model (LDA Konu Modelindeki Konular için Kelime Bulutlarının Oluşturulması)

Figure 4 presents the word clouds containing the 10 most frequently mentioned words in the topics generated by the LDA topic model. These word clouds were created using Python libraries, specifically the wordcloud and matplotlib packages. The word clouds for the topic model are presented in Figure 9. It is seen that the words in the bar graphs in Figure 7 are the same as the words in the word clouds in Figure 9.



Figure 9. Word Clouds for LDA Topic Model  
(LDA Konu Modeli Kelime Bulutları)

#### 4.6. Creating a Heatmap of the LDA Topic Model

(LDA Konu Modeli için Isı Haritası Oluşturulması)

The output of the Latent Dirichlet Allocation (LDA) topic model can be used to create a visual representation of the distribution of documents (thesis abstracts) across topics, in the form of a heatmap. The LDA topic model takes the topic probabilities in each document and stores these probabilities in a matrix, which is then visualised as a heatmap. The resulting heatmap is presented in Figure 10.

The map provides a more accessible and analytically tractable representation of the output of the LDA topic model. The degree of colour intensity in each cell is indicative of the relevance of the document in question to the topic in question. In creating this heatmap, the colour palette was set to 'YlGnBu'. The use of dark colours to represent high probability and light colours to represent low probability allows for a clear visualisation of the distribution of topics among documents and the degree of relatedness between documents and topics. The heatmap demonstrates that both topics exhibit a high document probability, as indicated by the high density of dark colors. In this context, the term 'document' refers to the number of words in the original word list from which repeated words have been extracted for each thesis.

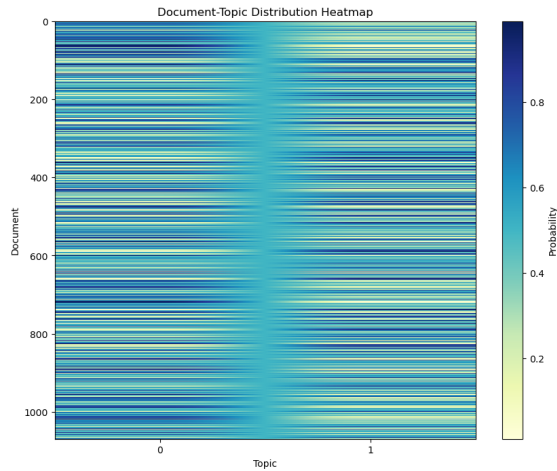


Figure 10. LDA Topic Model Document-Topic Distribution Heatmap

(LDA Konu Modeli Belge- Konu Dağılımı Isı Haritası)

#### 5. CONCLUSION (TARTIŞMA)

In latent Dirichlet allocation (LDA) topic modelling, fit values serve as a means of evaluating the model's efficacy and the interpretability of the topics it identifies. A higher fit value indicates that the model performs better and generates more meaningful topics. The application of the LDA algorithm to the data set yielded topic models in which the words

'data', 'model', 'research', 'technology', and 'system' were identified as predominant.

In topic modelling, the size of the data set and its compatibility within itself are important factors. Despite the topic models being created from master's and doctoral theses in the field of management information systems, the different subjects and contents affect the topic model performance.

A comparison of the results obtained in the studies presented in Section 2 with the results obtained in this thesis can be expressed as follows:

In the study [21], the Latent Dirichlet Allocation algorithm, a text mining method, was employed to analyze 574 graduate thesis abstracts completed between 2002 and 2020 in the MIS department. In this thesis, the same method was used to analyze 1170 graduate thesis abstracts completed between 2002 and 2023. This indicates that the number of theses published during the three years between 2020 and 2023 is higher than the number of theses published during the 18 years between 2002 and 2020. Furthermore, subject differentiation was observed in the results obtained by the subject models. The analysis conducted with the LDA algorithm revealed the prevalence of topics such as data analysis, decision-making, system analysis, system development, and information management in postgraduate theses.

In the study [22], the application of topic modelling with the LDA algorithm led to the conclusion that the most frequently obtained topics in the first three MIS journals in the SCOPUS database were business performance, value management, data analysis, training, knowledge management and model use. Similarly, the topic modelling with the LDA algorithm applied in this thesis study yielded results that highlighted data analysis, knowledge management and modelling as prominent topics.

In the study [23], words such as "study", "research", "analysis", "use", "method", "algorithm" were identified as prominent in the density maps within the scope of bibliometric analysis studies conducted on leading journal articles in the field of management information systems in WoS. In light of the aforementioned findings, a comparison of the subject model presented in Figure 5, which is among the article's key findings, reveals a similar trend. This suggests that graduate theses and international journal articles share a significant overlap in terms of content.

In the future studies, topic models and model performances can be evaluated by including

postgraduate theses prepared in 2024. The processes can be improved by using different algorithms such as Top2Vec, LSA, Bertopic instead of the LDA algorithm. Apart from this study, an examination of the articles published in the field of MIS can be determined as another study topic.

## REFERENCES (KAYNAKLAR)

- [1] Baskerville, R. L., & Myers, M. D. (2002), "Information Systems as a Reference Discipline". *MIS Quarterly*, 26(1), 1-14. <https://doi.org/10.2307/4132338>
- [2] Laudon, K. & Laudon, J. (2006), "Management Information Systems: Managing the Digital Firm", 9th ed. Prentice Hall.
- [3] Berry, M. W., & Castellanos, M. (2007). "Survey of Text Mining: Clustering", *Classification, and Retrieval*.
- [4] O'Mara-Eves, A., Thomas, J., McNaught, J. *et al.* "Using text mining for study identification in systematic reviews: a systematic review of current approaches", *Syst Rev* 4, 5 (2015). <https://doi.org/10.1186/2046-4053-4-5>
- [5] Peersman, C., Edwards, M., Williams, E. & Rashid, A. (2022), "A Survey of Relevant Text Mining Technology", *10.48550/arXiv.2211.15784*.
- [6] Parlak, B., & Uysal, A. K. (2015, May). Classification of medical documents according to diseases. In *2015 23rd signal processing and communications applications conference (siu)* (pp. 1635-1638). IEEE.
- [7] Kowsari, K., Meimandi, K. J., Heidarysafa, M., Mendu, S., Barnes, L., Brown, D. (2019), "Text Classification Algorithms: A Survey", *Information* 2019(10), 150; doi:10.3390/info10040150.
- [8] Liang, H., Sun, X., Sun, Y., Gao, Y. (2017), "Text feature extraction based on deep learning: a review", *Liang et al. EURASIP Journal on Wireless Communications and Networking*, (2017)211. doi: 10.1186/s13638-017-0993-1.
- [9] Cavnar, W. B., & Trenkle, J. M. (1994, April), "N-gram-based text categorization", In *Proceedings of SDAIR-94, 3rd annual symposium on document analysis and information retrieval* (Vol. 161175, p. 14).
- [10] Biricik, G., Diri, B., Sönmez, A.C. (2012), "Abstract feature extraction for text classification", *Turk J Elec Eng & Comp Sci*, Vol.20, No. Sup.1, 2012, TÜBİTAK doi:10.3906/elk-1102-1015.
- [11] Sakai, T., & Sparck-Jones, K. (2001, September), "Generic summaries for indexing in information retrieval", In *Proceedings of the 24th annual international ACM SIGIR conference on Research and development in information retrieval*, 190-198.
- [12] Vayansky, I & Kumar, S. (2020), "A review of topic modeling methods", *Information Systems*, 94. 101582. 10.1016/j.is.2020.101582.
- [13] Jurafsky, D., & Martin, J. H. (2000), "Speech and Language Processing: An Introduction to Natural Language Processing", *Computational Linguistics, and Speech Recognition*.
- [14] Garoufallou, E., & Gaitanou, P. (2021), "Big data: opportunities and challenges in libraries, a systematic literature review", *College & Research Libraries*, 82(3), 410.
- [15] Alghamdi, R., & Alfalqi, K. (2015), "A survey of topic modeling in text mining", *Int. J. Adv. Comput. Sci. Appl. (IJACSA)*, 6(1).
- [16] Blei, D. M., Ng, A. Y., & Jordan, M. I. (2003), "Latent dirichlet allocation", *Journal of machine Learning research*, 3(Jan), 993-1022.
- [17] Onan, A., Korukoglu, S., & Bulut, H. (2016), "LDA-based topic modelling in text sentiment classification: An empirical analysis", *Int. J. Comput. Linguistics Appl.*, 7(1), 101-119.
- [18] Kuang, D., Choo, J., & Park, H. (2015), "Nonnegative matrix factorization for interactive topic modeling and document clustering", *Partitional clustering algorithms*, 215-243.
- [19] Sievert, C., & Shirley, K. (2014). "LDAvis: A method for visualizing and interpreting topics", *Proceedings of the Workshop on Interactive Language Learning, Visualization, and Interfaces*.
- [20] Hasan, M., Rahman, A., Karim, M. R., Khan, M. S. I., & Islam, M. J. (2021), "Normalized approach to find optimal number of topics in Latent Dirichlet Allocation (LDA)", In *Proceedings of International Conference on Trends in Computational and Cognitive Engineering: Proceedings of TCCE 2020*, 341-354. Springer Singapore.
- [21] Çallı, L., Çallı, F., & Alma Çallı, B. (2021), "Yönetim Bilişim Sistemleri Disiplininde Hazırlanan Lisansüstü Tezlerin Gizli Dirichlet Ayrımı Algoritmasıyla Konu Modellemesi", *MANAS Sosyal*

*Araştırmalar Dergisi*, 10(4), 2355-2372.  
<https://doi.org/10.33206/mjss.894809>

[22] Parlina, A., & Kusumarani, R. (2023), "A Latent Dirichlet Allocation-Based Bibliometric Exploration of Top-3 Journals in Management Information Systems", *Jurnal Studi Komunikasi dan Media*, 27(1), 77-92.

[23] Özköse, H., & Gencer, C. T. (2017). "Bibliometric analysis and mapping of management information systems field", *Gazi University Journal of Science*, 30(4), 356-371.

[24] Internet: YÖK National Thesis Center Web Site. [Accessed: 02/05/2024.]

<https://tez.yok.gov.tr/UlusalTezMerkezi/>

[25] Sievert, C., & Shirley, K. (2014), "LDavis: A method for visualizing and interpreting topics." *Proceedings of the Workshop on Interactive Language Learning, Visualization, and Interfaces*.



## Quality Management System Based Blockchain Applications



Hilal YILMAZ,<sup>a</sup>



Ozlem SENVAR<sup>\*b</sup>

<sup>a</sup> Marmara Üniversitesi, Mühendislik Fakültesi, Endsütri Mühendisliği Bölümü, İstanbul, 34854, TÜRKİYE

<sup>b\*</sup> Marmara Üniversitesi, Mühendislik Fakültesi, Endsütri Mühendisliği Bölümü, İstanbul, 34854, TÜRKİYE

### ARTICLE INFO

Received: 07.10.2024  
Accepted: 05.02.2025

#### Keywords:

Quality Management  
System (QMS),  
Blockchain, Distributed  
Ledger Technology  
(DLT)

#### \*Corresponding Authors

e-mail:  
ozlemsenvar@gmail.co  
m

### ABSTRACT

This study investigates blockchain technologies and blockchain related researches from various sectors considering sectoral applications including food, healthcare, automotive, supply chain, information security, banking and quality management issues associated with these sectors. This study provides comparisons of various industries considering blockchain technology features. The aim of this study is to present an overview to intelligent quality management system based blockchain. This study examines standards for blockchain and distributed ledger technologies and discusses quality challenges for blockchain applications.

DOI: 10.59940/jismar.1563163

## Kalite Yönetim Sistemi Tabanlı Blok Zincir Uygulamaları

### MAKALE BİLGİSİ

Alınma: 07.10.2024  
Kabul: 05.02.2025

#### Anahtar Kelimeler:

Kalite Yönetim Sistemi  
(KYS), Blok Zincir,  
Dağıtık Defter Teknolojisi  
(DLT)

#### \*Sorumlu Yazar

e-posta:  
ozlemsenvar@gmail.com

### ÖZET

Bu çalışma, gıda, sağlık, otomotiv, tedarik zinciri, bilgi güvenliği, bankacılık ve bu sektörlerle ilişkili kalite yönetimi sorunlarını sektörel uygulamaları göz önünde bulundurarak çeşitli sektörlerden blok zinciri teknolojilerini ve blok zinciriyle ilgili araştırmaları incelemektedir. Bu çalışma, blok zinciri teknolojisi özelliklerini göz önünde bulundurarak çeşitli endüstrilerin karşılaştırmalarını sunmaktadır. Bu çalışmanın amacı, zeki kalite yönetim sistemine dayalı blok zincirine genel bir bakış sunmaktır. Bu çalışma, blok zincir ve dağıtık defter teknolojileri için standartları inceler ve blok zincir uygulamaları için kalite zorluklarını tartışır.

DOI: 10.59940/jismar.1563163

### 1. INTRODUCTION (GİRİŞ)

With the developing technologies in recent years, people's consumption habits and consumers' perceptions of quality are changing day by day. In the rapidly evolving digital environment,

organizations across various industries are looking for innovative ways to improve their processes and increase overall efficiency, effectiveness and productivity. Traditional quality management systems are required to be adapted with new technologies. In recent years, the importance of data,

data analysis and real-time data tracking have been increased.

Blockchain technology has the potential to streamline and optimize complex processes, removing intermediaries, reducing costs, and enhancing efficiency. Blockchain technologies have grabbed attention and continue to increase in the near future. Simultaneous traceability, transparency, immutability and reliability are important features of blockchain technology.

The integration of these advanced technologies into Quality Management System (QMS) with the promise of streamlining and enhancement of outranking by increasing transparency, traceability, reliability, immutability and real-time visibility, accountability. With the enrichment of product diversity, new consumption habits increase the supply-demand balance, which reveals the necessity of adaptations in supply chain management, which involves a wide range of operations such as production, transportation, storage and delivery. It has to be taken into account that supply chain management has many benefits, but there are also drawbacks and shortcomings to consider. With the use of new technology, these problems were resolved and the system was aimed to work more effectively.

This study aims to investigate the main advantages, disadvantages and weaknesses along with limitations of blockchain. Moreover, this study examines blockchain related research from different sectors. This study provides an overview to incorporating blockchain technologies into quality management. This article explores the key components of this technological convergence and its implications for the future of quality management. To the best of our knowledge, blockchain technology based quality management has not been handled, yet.

The rest of this study is organized as follows: Section 2 explains blockchain technology concepts. In section 3, blockchain sectoral applications including food, healthcare, automotive, supply chain, information security, banking and quality management issues associated with these sectors are reviewed. In section 4, an overview to intelligent quality management system based blockchain is given. In section 5, quality challenges for blockchain applications are discussed. The last section presents conclusions and recommendations for future directions.

## 2. BLOCKCHAIN TECHNOLOGY (BLOK ZİNCİR TEKNOLOJİSİ)

Blockchain is the technology used to build the cryptocurrency Bitcoin, which is called the next generation digital currency. Bitcoin was first introduced in an article by Satoshi Nakamoto in 2008[1]. A blockchain is fundamentally a distributed, decentralized digital ledger that securely and openly records transactions and uses cryptography to ensure the security and transparency of transactions (Figure 1). It is made up of a series of blocks, each of which has a list of transactions on it. A block becomes a chain once it has all of the transactions from the prior block linked to it. With its decentralized and transparent nature, combined with cryptographic security, opens up new possibilities for enhancing trust, efficiency, and security in various domains.

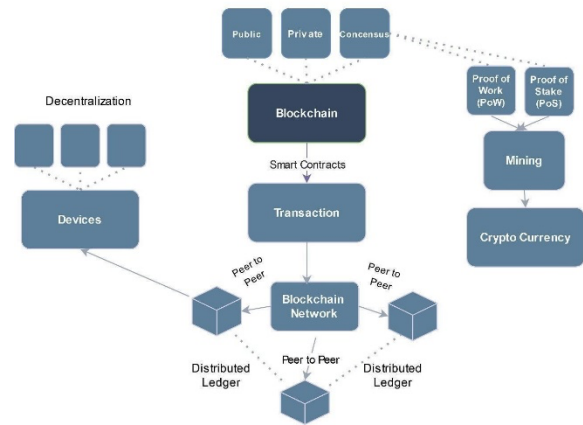


Figure 1. Blockchain technology architecture  
(Blok zincir teknolojisi mimarisi)

In 2014, Buterin [2] published a white paper proposing Ethereum, in which smart contracts are integrated into the blockchain system. The content of these contracts is written into the code and allowed to run on the network. Ethereum has paved the way for the development of decentralized applications (dapps) with its innovation [3].

Blockchain technology builds a secure structure and provides transparent, real-time monitoring and unchangeable properties of data records. Blockchain technology has various features, which are briefly explained below:

- **Decentralization** is the division of power and authority among several actors in a network so that no one party controls the entire system. Thus, it increases security and trust between network users by eliminating a single point of control. It also operates on a peer-to-peer network where multiple

computers or nodes are involved in recording and verifying transactions.

- Immutability** This provides a tamper-proof ledger as it is nearly impossible to remove or alter data added to the blockchain. This is one of the most known and trusted features of blockchain technology.

- Transparency** All network users can view the open and transparent ledger that the blockchain provides. All network users can view the open and transparent ledger provided by the blockchain. When a user makes a transaction in the ledger, other users also have instant access to this information. This provides real-time visibility into blockchain technology.

- Traceability** In a distributed ledger, a decentralized network of nodes keeps an unchangeable and transparent record of transactions. In a distributed ledger, a decentralized network of nodes keeps an immutable and transparent record of transactions. It is not possible to change previously made operations.

- Anonymity** A certain degree of anonymity is achieved in blockchain technology by providing users with cryptographic identifiers rather than revealing personal information.

- Dis-intermediation** By eliminating the need for intermediaries in transactions, blockchain reduces costs and increases efficiency.

- Security** Transactions on the blockchain are secured using cryptographic techniques, making them highly resistant to fraud and manipulation.

- Reliance** The blockchain eliminates the need for users to put their faith in a central authority in order to execute and verify transactions.

- Smart contracts** by automating and enforcing agreements according to predetermined rules, self-executing contracts reduce the need for middlemen. Ethereum and Turing Complete smart contract platforms make it easy to generate complex, programmable contracts.

- Credibility** The transparency of blockchain improves transaction credibility and lowers the possibility of fraud.

- Scrutiny** The blockchain makes all of its transactions publicly available, facilitating thorough examination and auditing.

- Performance** Different blockchain networks have different performance capabilities; some are made for low latency and high throughput.

- Execution** the quantity of transactions or operations that, at a given asset level, can be completed every second.

- Scalable** It is the capacity of blockchain technology to manage increasing transaction volume without sacrificing efficiency.

- Cryptography** These techniques secure data, ensuring that transactions are tamper-resistant and identities are protected.

Nowadays, there are problems with mutual trust in collaborations due to fraud in records. When intermediaries are used in transactions, time delays and costs increase, which prevents the work from being completed efficiently and at low cost in a short time. With this technology, data is not stored in a single center, but is kept safe against possible data loss and malicious attacks by being kept on more than one server on the same network [4,5].

Blockchain technology started to develop after the Bitcoin cryptocurrency [1]. Blockchain structure combined with distributed ledger technology is the simultaneous distribution of multiple copies of data to different servers. There are different structures in blockchain: public, private and consensus mechanisms (Figure 1).

## 2.1. Distributed Ledger (Dağıtık Defter)

Blockchain processes transactions using distributed ledger technology (DLT). With this technology, unlike traditional databases, information and transactions are recorded in multiple sources instead of a single source. With distributed ledger technology, data is recorded securely, creating a transparent, traceable and unchangeable structure.

## 2.2. Public Blockchain (Genel Blok Zincir)

Public blockchain can make any user to access blockchain network. This is also known as permissionless blockchain. In a public blockchain, users verify transactions, which are then shared publicly through a timestamped consensus mechanism. Face challenges in scaling to handle a large number of transactions quickly. Public blockchains are truly decentralized, democratic and independent of authority. The most important disadvantage of public blockchains is that they require a lot of energy consumption.

## 2.3. Private Blockchain (Özel Blok Zincir)

Private blockchains, which are not decentralized, are structures in which access to the system is given only to authorized users. Only those with permission can make transactions or verify changes on the blockchain. Private blockchains are specifically designed for enterprise applications. In private blockchain, data privacy is inevitable so that secure data sharing is ensured.

## 2.4. Hybrid Blockchain (Hibrit Blok Zincir)

Hybrid blockchain is an integrated structure that includes private and public blockchain features to get the best efficiency from the system. It is generally used by private institutions. It has a closed structure and provides a safer networking environment.

## 2.5. Consensus Mechanism (Konsensus Mekanizması)

In a consensus mechanism, nodes in the network agree on the accuracy of the transactions made and ensure that these transactions are included in the blockchain. In this way, errors that may occur in the system are prevented and a safe environment is provided for transactions. Several consensus methods are used to confirm transactions and maintain the integrity of the blockchain, such as Proof of Work (PoW) and Proof of Stake (PoS).

•*Proof of Work (PoW)* Users (crypto miners) must solve complex cryptographic puzzle to add a new block to the system. The user who finds the solution, which needs to be verified, receive crypto coin as a reward. It is the most decentralized and secure of all authentication mechanisms [6]. However, low transaction rates, excessive energy usage and expensive operating fees are the disadvantages of PoW. Cryptocurrency Bitcoin uses the POW mechanism.

•*Proof of Stake (PoS)* The authority for confirming transactions and creating new blocks is determined according to the amount of funds held by users (stakers). PoS is an alternative consensus mechanism with low cost and low energy consumption [6]. For this reason, cryptocurrency Ethereum has switched from PoW mechanism to PoS mechanism.

The foundation of cryptocurrencies like Bitcoin and Ethereum, blockchain technology, has proven to be a transformative force that has the potential to completely transform a number of industries [3,7]. Its special properties enable a decentralized, transparent and secure solution to long-standing problems in conventional systems.

## 2.6. Web 3.0 (Web 3.0)

Web 3.0, which represents the process that has been ongoing since 2010, is also called the Semantic Web. In this new era, computers imitate human analysis with the help of artificial intelligence or machine learning, use data in an autonomous structure and build user-specific results. Unlike traditional web servers, Web 3.0 applications use blockchain servers that communicate with each other in decentralized

networks rather than a single server. The foundations of Web 3.0 protocols come from cryptocurrency systems. While transactions are made with tokens in crypto currencies, there is no such structure in Web 3.0.

## 2.7. Internet of Things (IoT) (Nesnelerin İnterneti (IoT))

The Internet of Things (IoT) refers to a set of devices that are connected to the internet connected devices through a network or other communication networks and exchange data among themselves. With blockchain, manufacturers can make agreements with secure transactions in production and sales processes through smart contracts whose content cannot be changed [8,9]. In this way, transactions are confirmed quickly and it can be easily monitored whether the content of the contract is complied with. To record data, IoT devices collect data at every stage of production and this data is securely stored with DLT technology. Since the data is not stored in a single center, real-time traceability of the data is ensured by all users in the network

## 3. Sectoral Applications via Blockchain (Blok Zincir Sektör Uygulamaları)

Blockchain technology is becoming increasingly widespread in a wide variety of sectors due to its accessibility, useful structure, and effective and efficient solutions. In this study, various sectors, given in Figure 2 were examined. Figure 2 shows sectoral application fields via blockchain based systems. Table 1 shows blockchain technology related literature review considering various sectors involving more specifically food, healthcare, automotive, supply chain, information security and banking.

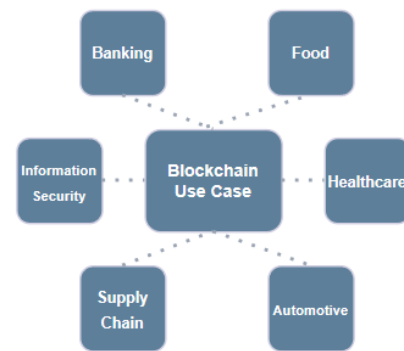


Figure 2. Sectoral application fields via blockchain based systems

(Blok zincir tabanlı sistemler üzerinden sektörel uygulama alanları)

Table 1. Blockchain technology related literature review considering various sectors

(Farklı sektörleri göz önünde bulundurarak blok zincir teknolojisi ile ilgili literatür taraması)

| Sector           | Author – Year               | Topics                     | Technology       |
|------------------|-----------------------------|----------------------------|------------------|
| Food&Agriculture | Tian, 2016                  | Agri-food products         | RFID, Blockchain |
| Banking          | Guo & Liang, 2016           | Banking industry           | Blockchain       |
| Banking          | Cocco et al., 2017          | Cost savings               | Blockchain       |
| Banking          | Wang & Kang, 2017           | Account information        | Blockchain       |
| Supply Chain     | Chen et al., 2017           | Quality management         | Blockchain       |
| Security         | Kumar & Mallick, 2018       | Security, IoT              | Blockchain       |
| Healthcare       | Jamil et al., 2019          | Pharmaceutical supply      | Blockchain       |
| Food&Agriculture | Musah et al., 2019          | Cocoa farms in Ghana       | Blockchain       |
| Supply Chain     | Tönnissen & Teuteberg, 2020 | Logistic supply            | Blockchain       |
| Food&Agriculture | Stranieri et al., 2021      | Food sustainability        | Blockchain       |
| Automotive       | Xu et al., 2022             | German automotive supplier | Blockchain       |
| Automotive       | Yasmin & Devi, 2023.        | Automotive suppliers       | Blockchain       |
| Healthcare       | Abdallah& Nizamuddin,2023   | Online sale                | Blockchain       |

### 3.1. Food Industry (Gıda Endüstrisi)

#### 3.1.1 Blockchain Technology in Food Industry

(Gıda Endüstrisinde Blok Zincir Teknolojisi)

It's known that there are inefficiencies, data discrepancies, and a lack of transparency within traditional food traceability systems. There are various problems in the food industry supply chain, such as food fraud and food spoilage due to improper storage. Thus, economic costs increase, delivery of the product at the desired time is delayed, and productivity decreases [10,11]. Blockchain technology offers many features such as simultaneous data tracking throughout the entire supply chain from producer to consumer, ability to monitor data by users on the same network, and immutability of data. Blockchain technology produces innovative solutions for the food industry and provides a safe, transparent and immutable system structure. To prevent food from spoiling, it must be stored in appropriate environmental conditions. By controlling the ambient temperature and humidity with the help of IoT devices, it can be checked whether the food is stored properly or not, and problems can be solved immediately. Blockchain technology, with its many features, meets the needs of the food sector and allows the establishment of a more efficient system [12].

The main features of applying blockchain technology in food supply chain management can be listed as follows: With its transparency and traceability feature, the process from the production of the products to their delivery to the consumer can

be followed in real time by all participants involved in the process. In this way, it is easily checked whether the requirements such as storing and t

transporting food in appropriate environments are met and whether the necessary precautions are taken [11]. With its immutability feature, it provides a safer environment by preventing the recorded data from being changed, thus preventing possible fraud in food products. With the feature of simultaneous data recording and monitoring, product data records can be viewed by all people in the network and transactions can be carried out more efficiently and faster.

#### 3.1.2. Blockchain Based Intelligent Food Quality Management System

(Blok Zincir Tabanlı Zeki Gıda Kalite Yönetim Sistemi)

Nowadays, food and agricultural implementations Information and Communications Technologies are incorporated with blockchain technology. Blockchain unique digital identifiers to food products provide traceability within the food supply chain including information such as agri-food growth conditions, lot numbers, and expiry dates preventing food waste and fraud, monitoring ecological footprint along with registering transactions of immutable food enabling source identification of foodborne illness. Digital nature of blockchain technologies can track on-farm data sharing [13].

According to [14], real time quality management and control systems with IoTs in the food supply chain increase security. Radio frequency identification

(RFID) and blockchain technology in the agri-food supply chain traceability system ensures the authenticity of the food safety and quality [11]. In [15] presented blockchain based quality management architecture developed for short food supply chains providing unique ability to store specific quality related data and supporting non-repudiation.

### **3.2. Healthcare System** *(Sağlık Sistemi)*

#### **3.2.1. Blockchain Technology in Healthcare Sector** *(Sağlık Sektöründe Blok Zincir Teknolojisi)*

Health services are of vital importance, and the medicines and materials used in the treatment processes and the duration of the treatment must be used in favor of the patient. In parallel with this situation, the healthcare supply chain also has a wide network, and there are many providers, from manufacturers to hospitals, who will ensure the supply of the requested materials within the specified period. Due to various negativities that may occur in this process, supply processes may be disrupted and negative effects may occur for patients.

Blockchain technology offers innovative solutions to various problems in the healthcare industry. Patient data should be stored in a secure environment and access by unauthorized persons should be prevented. The supply of medicines and other equipment should be carried out under appropriate environmental conditions. It is a structure that can be used for problems in the health sector by storing health data in a secure environment, preventing access by unauthorized persons and at the same time not being able to change it [16,17].

Medicines and other medical supplies must be original and imitation products must not be used. However, due to fraud in some drugs, different products are put on the market instead of real products, and this causes various problems in terms of treatment. Since the process from production to purchase of products produced with Blockchain technology is recorded simultaneously and in an unalterable way, any fraud in the products is prevented.

#### **3.2.2. Healthcare Quality Management** *(Sağlık Kalite Yönetim Sistemi)*

The health sector is a comprehensive system and there are many factors in the functioning of the health quality management system; Corporate management, accurate diagnosis and treatment of

diseases, patient satisfaction and supply management constitute the system as a whole. Implementation of a quality management system in healthcare improves the efficient use of resources by increasing the reliability and trustworthiness of all personnel. Institutions adopt quality management criteria in a competitive environment and ensure continuity with systems with new technologies in line with the principle of continuous improvement. There are various studies [18,19] in the literature to make a better structure in the quality management system in health.

### **3.3. Automotive Industry** *(Otomotiv Endüstrisi)*

#### **3.3.1. Blockchain Technology in Automotive Industry** *(Otomotiv Endüstrisinde Blok Zincir Teknolojisi)*

Blockchain technology offers convenient and efficient solutions to the problems experienced in automotive industry supply chain management. Since the materials in automotive production are supplied from many suppliers, they can also be from different countries, and accordingly, the supply chain scope of the sector is wide. The system is expected to operate effectively, efficiently and safely so that production can be completed on time [20,21].

With Blockchain technology, all processes in the automotive industry, from the supplier to the manufacturer, are recorded, ensuring that authorized users on the same network can access the data accurately and securely [22]. In this way, possible delays and disruptions are detected and a faster and more effective process is made. Stock, price and delivery information of the products are recorded completely and unchangeably, preventing possible fraud, preventing cost and time loss, and providing users with a traceable, transparent, low cost and fast infrastructure [23].

#### **3.3.2. Automotive Quality Management System** *(Otomotiv Kalite Yönetim Sistemi)*

The automotive industry is competitive and complex, and materials for production are sourced from many different global suppliers. Any delay, damage, or disruption of planned production in products supplied due to reasons such as a global crisis or epidemic can cause great losses for companies. In these processes, with the existence of an effective quality management system, timely solutions can be provided to problems that may arise. IATF 16949:2016 [24] is a standard specific to the automotive industry. It is intended for organizations that manufacture automobiles and related parts and

contains certain additional requirements. It is based on ISO 9001:2015 [25] quality management system.

### **3.4. Supply Chain Management** *(Tedarik Zinciri Yönetimi)*

#### **3.4.1. Blockchain Technology in Supply Chain Management** *(Tedarik Zinciri Yönetiminde Blok Zincir Teknolojisi)*

The importance of supply chain management continues to increase due to increasing production types and consumption habits with new technologies. The production supply shortage experienced worldwide due to sudden events (such as the pandemic period of Covid-19, chip crisis in Taiwan) has necessitated improvements in supply chain management. Features of blockchain technology eliminate the problems experienced in the supply chain due to the transparent, reliable, real-time traceable and unchangeable [26]. The processes of the products from the supplier to the manufacturer are seamlessly tracked via blockchain. In this way, it ensures that the products are delivered on time and under favorable conditions. Using blockchain smart contract technology, a secure trading environment is provided with agreements made by all participants in the process [27]. As a matter of fact, there can be conflicts between supply chain individuals, specifically in coordination of forecasting demand. If forecasting demand is uncoordinated, distorted demand forecasts result in bullwhip having forecast error and variance. Blockchain technology provides an information system and consensus formation mechanism that can intermediate supply chain network behavior.

#### **3.4.2. Supply Chain Quality Management** *(Tedarik Zinciri Kalite Yönetimi)*

Supply chain quality management (SCQM) operationalize and understand the impact of supply chain management on quality management [28]. Blockchain based Supply Chain Quality Management, which involves management mechanisms along with new Information Technologies (IT) systems used in supply chain quality management, solves the issues of distrust on the basis of unchanged information and traceable records through standardized norms and agreements. In blockchain based SCQM, blockchain technology, is integrated to new supply chain system in which information sharing and quality control are assured. Framework of blockchain based SCQM covers enterprises on the supply chain, blockchain, smart contracts and various IoT sensors. Blockchain technology adopts the governance model of human

society in IT systems, and further develops the decentralized system that provides different interest groups to share power in the same IT system, which improves the qualities of products and services in supply chains by contracts. Establishing automated executions of quality management contracts, it is possible to develop an auto-run intelligent system. Blockchain and smart contract establish more reliable quality track and control system, more agile ultimate customer demand catching system and more secure and convenient supply chain financial system [29].

### **3.5. Information Security Systems** *(Bilgi Yönetim Sistemleri)*

#### **3.5.1. Blockchain Technology in Information Security Systems** *(Bilgi Yönetim Sisteminde Blok Zincir Teknolojisi)*

The use of developing technologies also increases the amount of data obtained. Adequate technologies and storage centers are needed to store huge amount of data, protect integrity and confidentiality. Compared to traditional databases, blockchain technology offers a more secure and robust platform for storing and accessing data than other platforms. While data is kept in a single center in traditional databases, through blockchain's distributed ledger technology, copies of the data are stored in the relevant network. Whenever a new file is saved or any update is made to the file, it is instantly copied to the storage. With this and the automatic data backup feature, data loss is prevented. Data are stored securely through smart contracts allowing users to carry out their transactions without sharing important personal information.

Applications of blockchain technology to strengthen cybersecurity include: Blockchain technology ensures the security of the systems it operates on and the devices connected to it. It prevents unauthorized persons from accessing the data with end-to-end encryption methods used on the data [30]. By storing data in a decentralized structure, it minimizes malicious attacks and data loss that may occur in possible attacks.

#### **3.5.2. Information Security System Quality Management** *(Bilgi Güvenliği Sistemi Kalite Yönetimi)*

Protection of information and data breaches are important. There are various regulations regarding this issue. These regulations provide benefits for organizations in ensuring the security and continuity of their information assets. As an international standard, ISO/IEC 27001 [31] is one of them. The

features that form the basis of information security systems are included in the ISO/IEC 27001 standard as follows: confidentiality, availability and integrity. *Confidentiality* refers to protection of information against unauthorized access. *Integrity* refers to prevention of information from being modified by unauthorized persons. *Accessibility* refers to availability and usability of information by authorized persons.

### 3.6. Banking (Bankacılık)

#### 3.6.1. Blockchain Technology in Banking

(Bankacılıkta Blok Zincir Teknolojisi)

The banking sector, like other sectors, has to adapt its systems to new technologies in order to remain strong in the face of rapidly developing technologies. Various limitations of the current system cause various disruptions in transactions. Blockchain technology has emerged as a robust, reliable and versatile new technology that brings a different perspective to existing systems. Banking also started to integrate its systems with this technology to benefit from this [32]. The contributions of blockchain technology to the banking sector are as follows: Features that distinguish blockchain technology from the currently used bank payment system: While banks carry out their transactions through a central system, blockchain works in a decentralized network system. Unlike the banking system, Ledgers and banking transactions in blockchain are open to the public and data can be accessed in real time and transparently [33].

International money transfer transactions are made by using an intermediary and paying a transaction fee, and these transactions are slower. In blockchain technology, transactions can be made faster, transparently and with lower transaction fees. Various authentication methods are used to ensure security in banking transactions. In blockchain-based systems, verification processes can be carried out more securely.

Smart contracts are increasingly executed in applications fulfilling trustworthy and strong certifications in transactions between parties, whereas can be carried with immutable, transparent, traceable, and secure infrastructure. Recording transactions with blockchain's distributed ledger technology with blockchain decentralized consensus, allows accounting, bookkeeping and auditing in banking transactions to be carried out in a more functional, transparent, fast and secure way.

#### 3.6.2. Banking Quality Management System

(Bankacılıkta Kalite Yönetim Sistemi)

Banking sector should improve the quality, efficiency, effectiveness, productivity and resilience. Cucari et al [34] focused on the efficiency of processes, security and information network as types of banking applications. Customer portfolio is important in banking transactions in the banking sector. Credit operations are also based on customer reliability. Blockchain-based know your customer (KYC) system was developed to ensure that transactions that are reliable for both the customer and the institution. Blockchain strengthens identity verification and Know Your Customer (KYC) processes, for enhancing accurate customer data and supporting personal investment along with financial operations and customer relationship management (CRM) [35,36].

Banks should innovate through the implementation of blockchain. In this regard integration of new quality management systems along with new standards for blockchain technologies into banking sector. Integration of new generation technologies such as cloud data, artificial intelligence and IoT, blockchain technology into the banking sector is considered mandatory to increase efficiency, security and transparency [37].

CPA&AICPA [38] examined the impact of digital Blockchain Technology on auditing process of financial reports and services of quality assurance. Due to digitalization, utilization of digital tools is required for auditing process. Hashem et al. [39] recommend features (such as strategic, less time-consuming, continuity, review, scope expansion, consultancy) for more efficient audit processes of blockchain.

In the banking sector, integration of Accounting Information Systems (AIS) with Blockchain technology, features such as security, transparency and immutability come to the fore [40].

Al-Dmour et al. [41] revealed that there is a robust positive effect of Blockchain on AIS quality, significantly enhancing business performance.

The decentralized nature of Blockchain is used in cross-border payments and remittances by decreasing costs, increasing transaction speed, and transparency. As blockchain technology continues to develop, innovation and collaboration between industry stakeholders and regulators continue to transform the global financial system, making cross-border transactions more efficient, affordable, and

accessible for individuals and businesses around the world [42].

#### 4. QUALITY MANAGEMENT SYSTEM BASED BLOCKCHAIN TECHNOLOGY (KALİTE YÖNETİM SİSTEMİ TABANLI BLOK ZİNCİR TEKNOLOJİSİ)

There are eight quality management principles, which are kaizen (continuous improvement), customer focus, leadership, involvement of people, processes, approach, system approach to management, factual approach to decision-making, and mutual beneficial supplier relationship. As a matter of fact, kaizen should be taken into account within the quality management system. For this reason, continuous quality improvement should be concerned for all quality management based blockchain applications to improve performance along with effectiveness [43]. According to [43], there are limited applications for continuous quality improvement research. Apart from kaizen, other seven quality management principles should be broadly incorporated with blockchain systems.

Blockchain technology relies on cryptography which provides data compilation in blocks. Inevitably, validation should be performed considering kaizen technique via plan-do-check-act (PDCA cycle). After performing validation, block is formed and can be accessed by stakeholders. Afterwards, each block is incorporated with input and output hash values making blocks immutable for modification and having no variability. IoT technology is a system that has the ability to collect data by connecting devices to IoT based systems, and when integrated with blockchain technology, it enables the data to be stored in a database and shared with other relevant systems. At the same time, various options can be made within the quality management systems incorporated information obtained through IoT systems. Figure 3 shows the principal working logic of the intelligent quality management based blockchain system in terms of the PDCA cycle. There is a need for a quality management based blockchain system. In this regard, Blockchain technologies can solve mutability problems and enhance traceability considering kaizen.

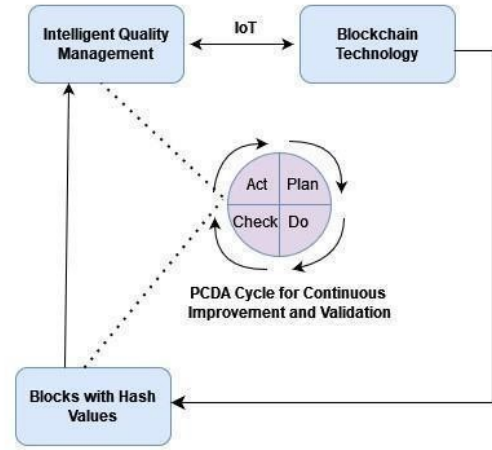
There are studies in the literature on the integration of blockchain technology with the ISO 9001:2015 standard. In [43], it is stated that the disruptions experienced in the operation and control of ISO 9001:2015-based Quality management system processes can be overcome by using blockchain

technology infrastructure with the considered scenarios.

ISO has released standards in the field of blockchain and distributed ledger technologies. Table 2 shows information on 12 ISO standards published since 2019 in this area.

Figure 3. Intelligent quality management system based blockchain

(Zeki kalite yönetimi sistemi tabanlı blok zincir)



#### 5. QUALITY CHALLENGES FOR BLOCKCHAIN APPLICATIONS (BLOK ZİNCİR UYGULAMALARI İÇİN KALİTE ZORLUKLARI)

##### 5.1. Quality Issues and Requirements (Kalite İşleri ve Gereksinimleri)

According to McCall's Software Quality Model 44, factors and related criteria are used to determine the quality of software. Blockchain structures consist of hardware and software systems and the quality criteria that should be considered for software and hardware systems are also valid for blockchain technologies. The factors of this model can be evaluated with establishing quality requirements for blockchain technology. These factors can be considered as follows: Efficiency, Reliability, Integrity, and Accuracy.

*Efficiency* factor is achieved by having appropriate and sufficient hardware and software for the system to operate as desired. Examples include having good storage capacity, high network speed, and using software with high processing capacity.

*Accuracy* factor can fulfill the required task completely when traceability and consistency criteria are ensured. In the blockchain system,

traceability of transactions is ensured with decentralized structures, while consistency of transactions is ensured by using smart contracts. In

this way, the accuracy of the transactions is ensured and the requirements of the reliability factor are met.

Table 2. Published ISO standards for blockchain and distributed ledger technologies

(Blok zincir ve dağıtık defter teknolojileri için yayınlanmış ISO standartları)

| Standards         | Type                                                                                                         | Published | Stage         |
|-------------------|--------------------------------------------------------------------------------------------------------------|-----------|---------------|
| ISO/TR 3242:2022  | Use cases                                                                                                    | 2022-10   | International |
| ISO/TR 6039:2023  | Identifiers of subjects and objects for the design of blockchain systems                                     | 2023-06   | International |
| ISO/TR 6277:2024  | Data flow models for blockchain and DLT use cases                                                            | 2024-02   | International |
| ISO 22739:2024    | Vocabulary                                                                                                   | 2024-01   | International |
| ISO/TR 23244:2020 | Privacy and personally identifiable information protection considerations                                    | 2020-05   | International |
| ISO/TR 23249:2022 | Overview of existing DLT systems for identity management                                                     | 2022-05   | International |
| ISO 23257:2022    | Reference architecture                                                                                       | 2022-02   | International |
| ISO/TS 23258:2021 | Taxonomy and ontology                                                                                        | 2021-11   | International |
| ISO/TR 23455:2019 | Overview of and interactions between smart contracts in blockchain and distributed ledger technology systems | 2019-09   | International |
| ISO/TR 23576:2020 | Security management of digital asset custodians                                                              | 2020-12   | International |
| ISO/TS 23635:2022 | Guidelines for governance                                                                                    | 2022-02   | International |
| ISO/TR 23644:2023 | Overview of trust anchors for DLT-based identity management                                                  | 2023-05   | International |

*Integrity* factor ensures that access to the system which is permitted to limited users and system control is performed by authorized users while unauthorized access is prevented in the blockchain system, possible fraud in transactions is not allowed. Blockchain is known as a transparent system in which users with authorized access can control.

In addition to McCall factors, the following factors can also be considered as blockchain system quality requirements.

*Installation cost:* Since blockchain is a costly system, a less costly system can be established by reducing costs.

*Network speed and capacity:* In order to carry out transactions quickly and accurately in the system, there must be a strong software and hardware infrastructure that can adapt to the network speed so that the users in the network can make fast transactions.

*Energy sustainability:* Blockchain mining requires high energy consumption. For this reason, systems that consume less energy are needed to be developed.

*Security:* Although blockchain has smart contracts and a decentralized structure, it is necessary to prevent attacks and fraud on the systems by ensuring the reliability of users and transaction information.

They are quantitative indicators that measure how the system works and how it achieves its goals. The system is examined and evaluations are made in terms of quality, performance and scalability. The performance metrics list is provided in Table 3.

It is difficult to measure the performance of a blockchain system accurately and reliably. Although the performance criteria for these are given above, priorities differ for each system and user. Performance tests of blockchain systems can be done using some special software. It is difficult to make a fair evaluation according to various criteria such as blockchain structures (public, private), network capacity, network speed, number of users, number of transactions made, hardware structure.

Awareness and usage rate of blockchain systems is increasing. Although the blockchain system is an advanced and complex structure for cryptocurrencies and integrated systems, the consumed energy is also quite high. There are concerns to reduce energy consumption to enhance ecological sustainability. Some blockchain systems are trying to build less energy consuming structures. In countries where usage of energy is widespread like India, there are initiatives in which the energy is consumed for mining of cryptocurrencies obtained from renewable energy sources.

## 5.2. Blockchain Performance Measurement Metrics (Blok Zincir Performans Ölçüm Metrikleri)

### 5.3. Blockchain Performance Measurement

**Methods** (Blok Zincir Performans Ölçüm Yöntemleri)

Blockchain performance evaluation methods are empirical and analytical methods. Empirical methods include benchmarking, testing, and checking

Table 3. Blockchain performance evaluation metrics

(Blok zincir performans değerlendirme metrikleri)

| Metrics      | Definition                                                                                                 |
|--------------|------------------------------------------------------------------------------------------------------------|
| Latency      | The time required for a transaction to be confirmed and recorded on the blockchain                         |
| Throughput   | The number of transactions performed per second                                                            |
| Scalability  | The capacity of the network, including the number of nodes it has and how many transactions it can process |
| Block Size   | Determines the number of transactions. The number of transactions increases with large size blocks.        |
| Resource     | Blockchain systems require high hardware capacity and consume a high amount of energy.                     |
| Security     | Ensuring the blockchain is protected against malicious attacks and fraud                                   |
| Storage      | It requires more capacity due to the size of the data stored in the blockchain system                      |
| Network Size | It refers to the total number of nodes actively participating in a blockchain network.                     |

Benchmarking compares the performance of a blockchain system with other systems. Testing performs experiments or simulations to test the functionality and reliability of a blockchain system. An audit examines and verifies the code, design, and architecture of a blockchain system for compliance with requirements.

Analytical methods include Stochastic Models, Queuing Models, Markov Chains, Emulation, Markov Decision Processes, Random Walks, Stochastic Petri Nets, Data Mining, Machine Learning.

Blockchain technology provides more transparent, secure, traceable and decentralized structure compared to traditional systems. Since records are not stored in a single center, it builds a trustful environment preventing possible data loss and security breaches and allowing simultaneous tracking of information and elimination of disruptions in processes. It designs a reliable environment in agreements between parties with smart contracts. The comparison of blockchain technology in food, automotive, healthcare, supply chain, security, energy and banking fields are performed with respect to various features presented in Table 4.

## 6. DISCUSSION AND CONCLUSION (TARTIŞMA

VE SONUÇ)

Table 4. Blockchain technology features considering various industries

(Çeşitli sektörleri göz önünde bulundurarak blok zincir teknolojisi özellikleri)

|                    | Food | Automotive | Healthcare | Supply Chain | Security | Energy | Banking |
|--------------------|------|------------|------------|--------------|----------|--------|---------|
| Traceability       | ✓    | ✓          | ✓          | ✓            | ✓        | ✓      | ✓       |
| Transparency       | ✓    | ✓          | ✓          | ✓            | ✓        | ✓      | ✓       |
| Immutability       | ✓    | ✓          | ✓          | ✓            | ✓        | ✓      | ✓       |
| Anonymousness      | ✗    | ✗          | ✗          | ✗            | ✗        | ✗      | ✗       |
| Dis-intermediation | ✓    | ✓          | ✓          | ✓            | ✓        | ✓      | ✓       |
| Security           | ✓    | ✓          | ✓          | ✓            | ✓        | ✓      | ✓       |
| Reliance           | ✓    | ✓          | ✓          | ✓            | ✓        | ✓      | ✓       |
| Smart contracts    | ✓    | ✓          | ✓          | ✓            | ✓        | ✓      | ✓       |
| Credibility        | ✓    | ✓          | ✓          | ✓            | ✓        | ✓      | ✓       |
| Scrutiny           | ✓    | ✓          | ✓          | ✓            | ✓        | ✓      | ✓       |
| Performance        | ✓    | ✓          | ✓          | ✓            | ✓        | ✓      | ✓       |
| Execution          | ✗    | ✗          | ✗          | ✗            | ✗        | ✗      | ✗       |
| Scalable           | ✗    | ✗          | ✗          | ✗            | ✗        | ✗      | ✗       |
| Cryptography       | ✓    | ✓          | ✓          | ✓            | ✓        | ✓      | ✓       |

In contrast to the benefits of blockchain technology, there are complexities along with difficulties of blockchain technology in terms of installation and use. High installation costs, lack of a common standard structure around the world and differences between countries pose obstacles to the widespread use of blockchain technology. There are complexities of blockchain technology.

Personal data are required for identification of patterns in delicate situations. Blockchain technologies are needed to be used for security and immutability considering storage of the highly sensitive data. The further directions of blockchain technologies will include a variety of implementations of decentralized blockchains along with super secured networks to enhance minimization of inherent vulnerability of centralized databases.

Artificial intelligence based blockchain technologies can contribute to tracing how algorithms work and how their input affects the output of machine learning. For further directions, new technological advancements are required with the integration of artificial intelligence and blockchain technologies considering cybersecurity challenges and having a double shield against cyberattacks by training machine learning algorithms to automate real time threat detection and to continuously learn about the behavior of attackers arisen in dynamic and uncertain environments.

For further directions, sustainability related researches should be considered within blockchain technology-based systems. Sectoral applications should consider three sustainability dimensions in terms of social, environmental, and economic aspects.

There are limitations in theory which need to be improved via empirical studies. As blockchain applications proliferate, data volume, volatilities, complexities, scales, diversities will be getting increased with the growth rate of blockchain technologies. rapid development of data applications have placed extremely high demands on the user amount, concurrency, and energy efficiency optimization of privacy protection service requests. Rapid development of blockchain applications will require analyzing data of blockchain technologies emerging methodological advancements in data science. As a matter of fact, there are still issues for standardization of blockchain technologies within

quality management involving especially considering new auditing developments and approaches via intelligent monitoring technologies, and dynamically exploring regulatory methods to improve quality assurance.

## REFERENCES (KAYNAKLAR)

- [1] S. Nakamoto, "Bitcoin: A Peer-to-Peer Electronic Cash System", [Online]. Available: <https://bitcoin.org/bitcoin.pdf/>. [Accessed: July. 10, 2024].
- [2] V. Buterin, "Ethereum white paper: a next generation smart contract & decentralized application platform," First Version, vol. 53, 2014.
- [3] S. Bistarelli, G. Mazzante, M. Micheletti, L. Mostarda, and F. Tiezzi, "Analysis of Ethereum Smart Contracts and Opcodes," *Advances in Intelligent Systems and Computing*, pp. 546–558, 2020. doi: 10.1007/978-3-030-15032-7\_46.
- [4] M. Pilkington, "Blockchain technology: principles and applications," *Research Handbook on Digital Transformations*, pp. 225-253, 2016. doi: 10.4337/9781784717766.00019.
- [5] T. Ahram, A. Sargolzaei, S. Sargolzaei, J. Daniels, and B. Amaba, "Blockchain technology innovations," 2017 IEEE Technology & Engineering Management Conference (TEMSCON), Santa Clara, CA, USA: IEEE, Jun. 2017, pp. 137–141. doi: 10.1109/TEMSCON.2017.7998367.
- [6] D. Macrinici, C. Cartoceanu, and S. Gao, "Smart contract applications within blockchain technology: A systematic mapping study," *Telematics and Informatics*, vol. 35, no. 8, pp. 2337–2354, Dec. 2018, doi: 10.1016/j.tele.2018.10.004.
- [7] H. Vranken, "Sustainability of bitcoin and blockchains," *Current Opinion in Environmental Sustainability*, vol. 28, pp. 1–9, Oct. 2017, doi: 10.1016/j.cosust.2017.04.011.
- [8] J. Xu, J. Lou, W. Lu, L. Wu, and C. Chen, "Ensuring construction material provenance using Internet of Things and blockchain: Learning from the food industry," *Journal of Industrial Information Integration*, vol. 33, p. 100455, Jun. 2023, doi: 10.1016/j.jii.2023.100455.
- [9] M. Torky and A. E. Hassanein, "Integrating blockchain and the internet of things in precision agriculture: Analysis, opportunities, and challenges," *Computers and Electronics in*

Agriculture, vol. 178, p. 105476, Nov. 2020, doi: 10.1016/j.compag.2020.105476.

[10] S. Musah, T. Medeni, and D. Soylu, "Assessment of Role of Innovative Technology through Blockchain Technology in Ghana's Cocoa Beans Food Supply Chains," Oct. 2019, pp. 1–12. doi: 10.1109/ISMSIT.2019.8932936.

[11] F. Tian, "An agri-food supply chain traceability system for China based on RFID & blockchain technology," 2016 13th International Conference on Service Systems and Service Management (ICSSSM), pp. 1–6, Jun. 2016, doi: 10.1109/ICSSSM.2016.7538424.

[12] S. Stranieri, F. Riccardi, M. P. M. Meuwissen, and C. Soregaroli, "Exploring the impact of blockchain on the performance of agri-food supply chains," Food Control, vol. 119, p. 107495, Jan. 2021, doi: 10.1016/j.foodcont.2020.107495.

[13] P. Pelé, J. Schulze, S. Piramuthu, and W. Zhou, "IoT and Blockchain Based Framework for Logistics in Food Supply Chains," Information Systems Frontiers, vol. 25, no. 5, pp. 1743–1756, Oct. 2023, doi: 10.1007/s10796-022-10343-9.

[14] F. Antonucci, S. Figorilli, C. Costa, F. Pallottino, L. Raso, and P. Menesatti, "A review on blockchain applications in the agri-food sector," Journal of the Science of Food and Agriculture, vol. 99, no. 14, pp. 6129–6138, Nov. 2019, doi: 10.1002/jsfa.9912.

[15] P. Burgess, F. Sunmola, and S. Wertheim-Heck, "Blockchain Enabled Quality Management in Short Food Supply Chains," Procedia Computer Science, vol. 200, pp. 904–913, 2022, doi: 10.1016/j.procs.2022.01.288.

[16] M. Fiore et al., "Blockchain for the Healthcare Supply Chain: A Systematic Literature Review," Applied Sciences, vol. 13, no. 2, Jan. 2023, doi: 10.3390/app13020686.

[17] S. Abdallah and N. Nizamuddin, "Blockchain-based solution for Pharma Supply Chain Industry," Computers & Industrial Engineering, vol. 177, p. 108997, Mar. 2023, doi: 10.1016/j.cie.2023.108997.

[18] A. R. C. Araújo, I. L. dos Santos, and A. da C. Reis, "A systematic review of the literature on the application of blockchain in the health supply chain," International Journal of Innovation, vol. 10, no. 4, Oct. 2022, doi: 10.5585/iji.v10i4.22060.

[19] F. Jamil, L. Hang, K. Kim, and D. Kim, "A Novel Medical Blockchain Model for Drug Supply Chain Integrity Management in a Smart Hospital," Electronics, vol. 8, no. 5, p. 505, May 2019, doi: 10.3390/electronics8050505.

[20] S. Yasmin and G. S. Devi, "Blockchain and Cloud-based Technology in Automotive Supply Chain," 2023 5th International Conference on Smart Systems and Inventive Technology (ICSSIT), Jan. 2023, pp. 771–775. doi: 10.1109/ICSSIT55814.2023.10060948.

[21] X. Xu, L. Tatge, X. Xu, and Y. Liu, "Blockchain applications in the supply chain management in German automotive industry," Production Planning & Control, vol. 0, no. 0, pp. 1–15, 2022, doi: 10.1080/09537287.2022.2044073.

[22] N. Ada et al., "Blockchain Technology for Enhancing Traceability and Efficiency in Automobile Supply Chain—A Case Study," Sustainability, vol. 13, no. 24, Jan. 2021, doi: 10.3390/su132413667.

[23] S. S. Kamble, A. Gunasekaran, N. Subramanian, A. Ghadge, A. Belhadi, and M. Venkatesh, "Blockchain technology's impact on supply chain integration and sustainable supply chain performance: evidence from the automotive industry," Annals of Operations Research, vol. 327, no. 1, pp. 575–600, Aug. 2023, doi: 10.1007/s10479-021-04129-6.

[24] IATF 16949: 2016 Quality Management System Requirements. [Online]. Available: <https://www.iatfglobaloversight.org/iatf-169492016/>. [Accessed: Aug. 10, 2024].

[25] ISO 9001:2015 Quality management systems — Requirements. [Online]. Available: <https://www.iso.org/standard/62085.html>. [Accessed: Aug. 10, 2024].

[26] S. Tönnissen and F. Teuteberg, "Analysing the impact of blockchain-technology for operations and supply chain management: An explanatory model drawn from multiple case studies," International Journal of Information Management, vol. 52, no. C, 2020, [Online]. Available: <https://ideas.repec.org/a/eee/ininma/v52y2020ics026840121930101x.html>. [Accessed: Aug. 30, 2024]

[27] B. Musigmann, H. Von Der Gracht, and E. Hartmann, "Blockchain Technology in Logistics and Supply Chain Management—A Bibliometric Literature Review From 2016 to January 2020," IEEE Transactions on Engineering Management,

vol. 67, no. 4, pp. 988–1007, Nov. 2020, doi: 10.1109/TEM.2020.2980733.

[28] S. T. Foster, “Towards an understanding of supply chain quality management,” *Journal of Operations Management*, vol. 26, no. 4, pp. 461–467, Jul. 2008, doi: 10.1016/j.jom.2007.06.003.

[29] S. Chen, R. Shi, Z. Ren, J. Yan, Y. Shi, and J. Zhang, “A Blockchain-Based Supply Chain Quality Management Framework,” 2017 IEEE 14th International Conference on e-Business Engineering (ICEBE), Nov. 2017, pp. 172–176. doi: 10.1109/ICEBE.2017.34.

[30] N. M. Kumar and P. K. Mallick, “Blockchain technology for security issues and challenges in IoT,” *Procedia Computer Science*, vol. 132, pp. 1815–1823, 2018, doi: 10.1016/j.procs.2018.05.140.

[31] ISO/IEC 27001:2022 Information security, cybersecurity and privacy protection — Information security management systems — Requirements. [Online]. Available: <https://www.iso.org/standard/27001>. [Accessed: June. 10, 2024].

[32] L. Cocco, A. Pinna, and M. Marchesi, “Banking on Blockchain: Costs Savings Thanks to the Blockchain Technology,” *Future Internet*, vol. 9, no. 3, p. 25, Jun. 2017, doi: 10.3390/fi9030025.

[33] Y. Guo and C. Liang, “Blockchain application and outlook in the banking industry,” *Financial Innovation*, vol. 2, no. 1, p. 24, Dec. 2016, doi: 10.1186/s40854-016-0034-9.

[34] N. Cucari, V. Lagasio, G. Lia, and C. Torriero, “The impact of blockchain in banking processes: the Interbank Spunta case study,” *Technology Analysis & Strategic Management*, vol. 34, no. 2, pp. 138–150, Feb. 2022, doi: 10.1080/09537325.2021.1891217.

[35] J. Frizzo-Barker, P. A. Chow-White, P. R. Adams, J. Mentanko, D. Ha, and S. Green, “Blockchain as a disruptive technology for business: A systematic review,” *International Journal of Information Management*, vol. 51, p. 102029, Apr. 2020, doi: 10.1016/j.ijinfomgt.2019.10.014.

[36] N. Rane, S. Choudhary, and J. Rane, “Metaverse for Enhancing Customer Loyalty: Effective Strategies to Improve Customer Relationship, Service, Engagement, Satisfaction, and Experience,” *SSRN Electronic Journal*, 2023, doi: 10.2139/ssrn.4624197.

[37] Y. Wang and A. Kogan, “Designing Privacy-Preserving Blockchain Based Accounting Information Systems,” *SSRN Electronic Journal*, 2017, doi: 10.2139/ssrn.2978281.

[38] CPA&AICPA, “Blockchain Technology and Its Potential Impact on the Audit and Assurance Profession.” 2017. [Online]. Available: <https://www.aicpa.org/content/dam/aicpa/interestareas/frc/assuranceadvisoryservices/downloadabledocuments/blockchaintechnology-and-its-potential-impact-on-the-audit-and-assurance-profession.pdf>. [Accessed: July. 20, 2024].

[39] R. Hashem, A. Abu-Musa, and E. Moubark, “The Impact of Blockchain Technology on Audit Process Quality: An Empirical Study on the Banking Sector,” *International Journal of Auditing and Accounting Studies*, vol. 5, pp. 87–118, Feb. 2023, doi: 10.47509/IJAAS.2023.v05i01.04.

[40] J. Schmitz and G. Leoni, “Accounting and Auditing at the Time of Blockchain Technology: A Research Agenda,” *Australian Accounting Review*, vol. 29, no. 2, pp. 331–342, Jun. 2019, doi: 10.1111/auar.12286.

[41] A. Al-Dmour, R. Al-Dmour, H. Al-Dmour, and A. Al-Adwan, “Blockchain applications and commercial bank performance: The mediating role of AIS quality,” *Journal of Open Innovation: Technology, Market, and Complexity*, vol. 10, no. 2, p. 100302, Jun. 2024, doi: 10.1016/j.joitmc.2024.100302.

[42] Dr. Paulo Leitão, “Cross-border Payments and Remittances on Blockchain: Exploring the use of blockchain for facilitating cross-border payments and remittances, reducing costs and improving transaction speed,” *Journal of AI-Assisted Scientific Discovery*, vol. 2, no. 1, pp. 1–8, Jun. 2024.

[43] R. Muruganandham, K. Venkatesh, S. R. Devadasan, and V. Harish, “TQM through the integration of blockchain with ISO 9001:2015 standard based quality management system,” *Total Quality Management & Business Excellence*, vol. 34, no. 3–4, pp. 291–311, Feb. 2023, doi: 10.1080/14783363.2022.2054318.

[44] J. A. McCall, P. K. Richards, and G. F. Walters, “Factors in Software Quality. Volume I. Concepts and Definitions of Software Quality;,” *Defense Technical Information Center, Fort Belvoir, VA*, Nov. 1977. doi: 10.21236/ADA049014.

[42] Dr. Paulo Leitão. (2024). Cross-border Payments and Remittances on Blockchain: Exploring the use of blockchain for facilitating cross-border payments and remittances, reducing costs and improving transaction speed. *Journal of AI-Assisted Scientific Discovery*, 2(1), 1–8.

[43] Muruganandham, R., Venkatesh, K., Devadasan, S. R., & Harish, V. (2023). TQM through the integration of blockchain with ISO 9001:2015 standard based quality management system. *Total Quality Management & Business Excellence*, 34(3–4), 291–311. <https://doi.org/10.1080/14783363.2022.2054318>

[44] McCall, J. A., Richards, P. K., & Walters, G. F. (1977). *Factors in Software Quality. Volume I. Concepts and Definitions of Software Quality*: Defense Technical Information Center. <https://doi.org/10.21236/ADA049014>



# MLP Temelli Yapay Zeka Modellerinde Çıktıları Etkileyen Giriř Veri Seti Niteliklerinin Model Mimarisine Göre Davranıřlarının Arařtırılması

Muhammer İLKUÇAR\*,<sup>a</sup>

<sup>a</sup>\*, Muğla Sıtkı Koçman Üniversitesi, Fethiye İřletme Fakültesi, Yönetim Biliřim Sistemleri Bölümü, Muğla, 48300, Türkiye

## MAKALE BİLGİSİ

Alınma: 01.11.2024  
Kabul: 23.06.2025

### Anahtar Kelimeler:

Açıklanabilir Yapay Zeka  
Sorumlu Yapay Zeka  
Derin Yapay Sinir Ağları  
Makine Öğrenmesi  
SHAP

### \*Sorumlu Yazar

e-posta:  
muhammerilkucar@mu.  
edu.tr

## ÖZET

Yapay zekanın yaygınlařması ile birlikte açıklanabilirlik, yorumlanabilirlik, řeffaflık gibi konular, özellikle saėlık, savunma sanayi, güvenlik, hukuk gibi alanlarda çok daha önemli hale gelmiřtir. Bu çalışmada; ileri beslemeli geri yayımlı çok katmanlı Yapay Sinir Aėı (MLP: Multi Layer Perceptron) yapay zeka modellerinde giriř veri seti nitelik deėerinin model çıkıřına olan etkilerinin model mimarisi ile olan iliřkisi arařtırılmıřtır. Model giriř veri özniteliklerinin model tahminine katkıları SHAP (SHapley Additive exPlanations) yöntemi ile ölçölmüřtür. MLP mimarisi deėiřtikçe giriř veri seti nitelik deėerlerinin model çıkıřına katkı oranları sıralaması da deėiřmektedir. Öznitelik etki sıralamasındaki deėiřimin çoėunlukla katkı düzeyleri birbirine göre yakın olan öznitelik deėerleri için geçerli olduėu, etki oranı diėer özniteliklerden biraz farklı olan özniteliklerin etki sıralamasının MLP mimarisi ile çok fazla deėiřmediėi gözlemlenmiřtir. Bu sonuçlara göre MLP model mimarisinin Açıklanabilir Yapay Zeka'da da belli bir oranda etkili olduėu, modelin doėruluk deėeri ile özniteliklerin önem oranları arasında anlamlı bir iliřki olmadıėı sonucuna varılabilir.

DOI: 10.59940/jismar.1577691

# Analysis of the Behavior of The Input Data Set Attributes Affecting the Outputs in MLP Based Artificial Intelligence Models According to the Model

## ARTICLE INFO

Received: 01.11.2024  
Accepted: 23.06.2025

### Keywords:

Explainable Artificial Intelligence  
Responsible AI  
Deep Neural Network  
Machine Learning  
SHAP

### \*Corresponding Authors

e-mail:  
muhammerilkucar@mu.  
edu.tr

## ABSTRACT

With the widespread use of artificial intelligence, explainability, interpretability, and transparency have become very important issues, especially in the health, defence, security, and law domains. In this study, the same datasets were used with different multilayer perceptron (MLP) architectures, and the effects of dataset attributes on MLP model output were analysed. The contributions of the model input data attributes to the model prediction were measured using the SHAP (SHapley Additive exPlanations) method. For the datasets, as the MLP architecture changed, the importance ranking levels of the input dataset attribute values also changed. It was observed that the change in the attribute influence ranking was mostly applicable and valid for attribute values whose contribution levels were relatively close to each other, and the influence ranking of the attributes whose influence ratio was slightly different from other attributes did not change significantly based on the MLP architecture. According to these results, it can be concluded that the model architecture also influences Explainable Artificial Intelligence results to a certain extent, and that there is no direct relationship between the model's accuracy and attribute importance ranking.

DOI: 10.59940/jismar.1577691

## 1. INTRODUCTION

Artificial intelligence has permeated all facets of life and is steadily progressing towards becoming an essential component of our lives. Artificial intelligence models have achieved significant success in various domains such as categorization, clustering, prediction, etc. As a result, they have been effectively utilized in diverse sectors such as education, engineering, health, transportation, finance, media, art, etc. Recently, different Artificial intelligence (AI) applications have become widespread in both business and individual daily life, such as Generative Adversarial Networks (GAN) and Large Language Models (LLM). The widespread use of artificial intelligence has prompted concerns about its outputs, including their explainability, interpretability, dependability, trustworthiness, and ethical.

Today, in addition to the impressive performance of artificial intelligence models, the significance of notions such as responsibility, explainability, and justice is increasing [1]. The field of Explainable Artificial Intelligence (XAI) has emerged due to the influence of input data set qualities on the output of machine learning (ML) models, and the need for transparency, traceability, and explainability of the obtained results [2]. XAI is essential in all domains that employ artificial intelligence, but it holds particular significance in sectors such as healthcare, defence, security, and law [2,3]. Explainable Artificial Intelligence (XAI) is a crucial technology that enhances the dependability and transparency of AI systems, fostering user trust in these systems. For AI to be adopted by people, it must be explainable and reliable [4].

According to Saeed and Omlin [3], XAI can be useful in fields including digital forensics, 5G, and the Internet of Things. However, there are certain general issues with XAI design, development, application, explainability, and interpretability. In order to guarantee the fairness, transparency, and accountability of machine learning systems, Barredo Arrieta et al. [5] contend that XAI is crucial for the creation and application of ML models. A study of the literature on interpretable machine learning in the field of artificial intelligence in healthcare was carried out by Tjoa and Guan [6]. Došilović, Brčić and Hlupić [7] examine the most recent developments in the application of XAI in supervised machine learning and contend that XAI should take precedence given the growing concern over AI and the issue of trust. XAI methods can be applied to many machine learning algorithms. Thanks to XAI, the interpretability of machine learning algorithms increases. Deep learning algorithms are more difficult to explain and comprehend than algorithms like

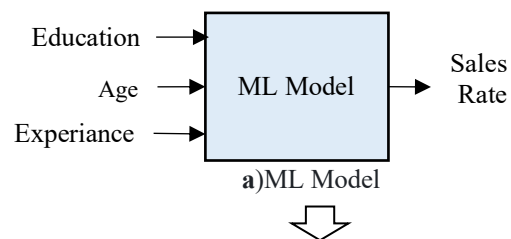
Decision Trees [8], Support Vector Machines [9,10], Logistic Regression [11], and Naive Bayes [12] classifiers. On the other hand, an artificial neural network have many parameters. The system is like a black box, especially in multilayer and deep artificial neural networks where the large number of layers and parameters makes them difficult to interpret.

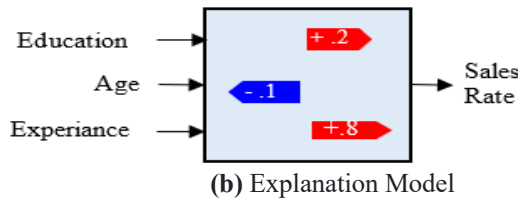
In this study, the effects of network architecture and input data attributes on XAI results in MLP were examined.

## 2. MATERIALS AND METHODS

### 2.1. SHapley Additive exPlanations (SHAP) Method

SHapley Additive exPlanations (SHAP) Method is a method based on game theory [13] that reveals the contribution of each input feature to the model's output in machine learning [14]. Game theory examines situations where there is more than one player and each player's actions affect the decisions of the other players. In game theory, methods exist to measure how much each player contributes to the outcome of a game. In the SHAP method, the input attributes of the ML model can be viewed as players, and their contribution to the model's prediction (the game's outcome) can be measured. In this way, the outputs of the model will be more explainable by measuring which attributes affect the output of an ML model and how. It can also be observed which attributes of the input data set are more important and which are less important for the ML Model with SHAP. A simple SHAP model is shown in Figure 1 [15]. The ML Model has three input features (age, education, and experiences) and one output (sales rate) (Figure 1.a). The sales rate will be predicted for age, education, and experience inputs with the trained ML model. As seen in Figure 1.b, the "experience" feature makes a positive contribution to the "sales rate" output with 0.8, while the "age" feature makes a negative contribution with 0.1. While education and experience feature positive contributions to the sales rate estimate, age features provide negative contributions.





**Figure 1.** Machine Learning model and explanation model (Makine öğrenmesi modeli ve açıklanabilir model).

### 3. DATASET

In the study two different datasets were used: Breast Cancer Wisconsin and Heart Disease [16]. Since these datasets are widely used in the literature, they were preferred for other researchers.

#### 3.1. Breast Cancer Wisconsin Dataset

The Wisconsin Breast Cancer dataset [16] consists of 699 samples, 9 features and one diagnostic feature (label) (Table 1). In each feature, the tumour diameters of X-ray images taken from the breast are expressed numerically. According to these features, whether the tumour is benign or malignant is given as a label in the diagnostic feature. A diagnosis value of 2 means a benign tumour, and 4 means a malignant tumour. Of the data in the dataset, 241 samples belong to malignant tumours, while the remaining 458 samples are benign tumour samples. There are missing data for 16 samples in the Features6 column in the dataset. Missing data was filled with 0. Various techniques can be employed to replace missing values [17,18,19]. Filling the missing data with one of these strategies, rather than using zero, will have a favourable impact on the model's performance. However, the aim of this study was not to improve the performance of the model. For this reason, missing data in the dataset were simply filled with 0.

**Table 1.** Breast cancer dataset, features and diagnosis (label) sample data (Meme kanseri veri seti, teşhis (tiket) ve örnek veriler).

| #   | Feature_0 | Feature_1 | Feature_2 | Feature_3 | Feature_4 | Feature_5 | Feature_6 | Feature_7 | Feature_8 | Diagnosis |
|-----|-----------|-----------|-----------|-----------|-----------|-----------|-----------|-----------|-----------|-----------|
| 0   | 5         | 4         | 4         | 5         | 7         | 10        | 3         | 2         | 1         | 2         |
| 1   | 3         | 1         | 1         | 1         | 2         | 2         | 3         | 1         | 1         | 2         |
| 2   | 6         | 8         | 8         | 1         | 3         | 4         | 3         | 7         | 1         | 2         |
| 3   | 4         | 1         | 1         | 3         | 2         | 1         | 3         | 1         | 1         | 2         |
| 4   | 8         | 10        | 10        | 8         | 7         | 10        | 9         | 7         | 1         | 4         |
| ... |           |           |           |           | ...       |           |           |           |           | ...       |

#### 3.2. Heart Disease Dataset

Table 2 displays the heart disease dataset [16], which consists of 13 features and a label indicating whether the person has heart disease or not. According to the findings, the diagnostic value was taken as 1 if the patient has heart disease and 0 if the patient does not have heart disease. The dataset consists of 1025 samples, all features are expressed numerically, and there are no missing data. In the dataset, 499 samples belong to patients with heart disease and 526 samples belong to patients without heart disease.

**Table 2.** Heart disease dataset, features and (label) diagnosis (Kalp hastalığı veri seti, nitelikleri ve teşhis/etiket).

| Features             | Explanation                                                                                    |
|----------------------|------------------------------------------------------------------------------------------------|
| Feature_0 (age)      | Age in years                                                                                   |
| Feature_1 (sex)      | Gender; 0:female; 1: male                                                                      |
| Feature_2 (cp)       | Chest pain type (1: atypical angina; 2: atypical angina; 3: non-anginal pain; 4: asymptomatic) |
| Feature_3 (trestbps) | Blood pressure at rest (mm Hg)-Tansion                                                         |
| Feature_4 (chol)     | Serum kolestorol (mg/dl)                                                                       |
| Feature_5 (fps)      | fasting blood sugar > 120 mg/dl (1: true; 0: false)                                            |
| Feature_6 (restech)  | Resting electrocardiographic results                                                           |
| Feature_7 (thalach)  | Maximum heart rate                                                                             |
| Feature_8 (exang)    | Exercise induced angina (1:Yes; 0: No)                                                         |
| Feature_9 (oldpeak)  | ST depression caused by exercise relative to rest                                              |
| Feature_10 (slope)   | Hill exercise Slope of segment (ST)                                                            |
| Feature_11 (ca)      | Number of large vessels colored by fluoroscopy (0-3)                                           |
| Feature_12 (thal)    | 3:normal; 6: fixed defect; 7: reversible defect                                                |
| Diagnostic           | 0: Not disease; 1: Disease                                                                     |

### 4. RESULTS AND DISCUSSION

In the study, MLP model training and SHAP tests with different architectures were performed for each data set. After training and testing Multi-Layers Artificial Neural Network (MLP) in different architectures with different data sets, the effects of the input data attributes in the model output and the importance levels of the input data attributes were observed by using the SHAP method. Thus, it was investigated whether the importance of the input data set attribute

values according to their effects on the output varies with the MLP model architecture. In the study, the hyperparameters were kept constant for all MLP architectures and the effect of these parameters on the output was stabilized.

XAI techniques such as SHAP, DeepSHAP, DeepLIFT, DeepLIFT, CXplain, and LIME (Local Interpretable Model-agnostic Explanations) [15, 20, 21] are used as XAI techniques. SHAP is a game theory [22] approach to explaining the output of any machine learning model. It combines optimal credential location with local explanations using classical SHapley values from game theory and its related extensions. This technique graphically and statistically illustrates the importance of the features that affect the decisions of the AI system, the statistical information that affects the decisions of the AI system, and the factors that affect the decisions of the system [23]. XAI techniques can show different suitability according to the artificial intelligence machine learning model. They state that LIME is more suitable for Artificial Neural Networks and Random Forest algorithms, while SHAP is more suitable for boosting-based algorithms. Since the SHAP method uses all attributes of the entire dataset [20], the SHAP method is preferred as the XAI technique for MLP and Deep Learning Machine Learning algorithm in this study in order to see the importance of all factors determining the result. The common parameters used in all MLP architectures in the study are as follows:

- *Training dataset* : 70% of the data
- *Test dataset* : 30 % of the data
- *Validation dataset* : 1% of training dataset
- *Epoch* : 500
- *Early Stopping* : true
- *Batch size* : 1
- *Optimizer* : adam
- *Hidden layerS activation function* : ReLu
- *Output layerS activation function* : Sigmoid
- *Loss function* : BinaryCrossentropy
- *Performance metric* : Accuracy

In order to the MLP algorithms to work better, the data was scaled in the range [0,1]. The model output is a binary classification (1: Patient, 0: Not Patient). Since the model output value is a continuous number between 0 and 1, it was converted to binary according to a certain threshold value. In this study, 0.5 was taken as the threshold. Accordingly, if the model output value was  $\geq 0.5$ , it was considered as 1; in the other case, it was considered as 0. 70% of the data was reserved for training and 30% for testing. While training, 1% of the training data was used as validation data. The ReLu was used as the activation function for all hidden layers of the model. Since the output of the model will be binary classification, the output layer activation function is Sigmoid, and Binary Cross

Entropy were used as the loss function. The model's optimization algorithm is *adam*, batch size 1, epoch 500, and no dropout used. With these hyperparameters, MLPs with different architectures were tested with two different data sets. In order to examine the effects of the input data set on the output with the MLP architecture, different architectures were selected. The architecture used 3, 4, 5, 6, and 7 layered architectures with high and close accuracy values. The number of nodes in the hidden layers was taken randomly between 5 and 50 as multiples of 5.

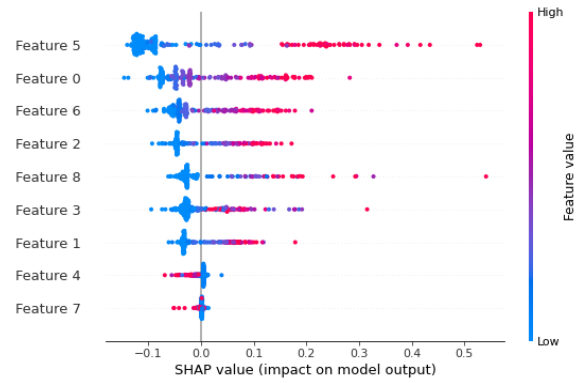
In the study, the accuracy score was used as the MLP model performance metric. Since the test data sets were balanced, the accuracy performance values reflect the success of the model. Since the study focused on XAI, the performance metrics such as precision, recall, f1 score, and r-square were not used. When the performance values of the MLP models in the study were compared with other studies conducted with the same data set in the literature [24-30], it was seen that the model was successful.

In Figure 2, for the Brest Cancer dataset, the model was trained separately using different MLP architectures, and the accuracy performance values and SHAP plots are shown. In the SHAP graphs, the effects of the dataset attributes of the model output (the importance of the attributes) can be seen graphically (Figure 2). For example, in Figure 2.a, the positive and negative effects of features on the outcome and the amount of these effects are shown in different colors. The color red indicates a positive effect, while blue indicates a negative effect. Also in the graph, the features are ranked according to their importance. For example, Feature5 is the most important feature while Feature7 is the least important feature. The accuracy values obtained for the different architectures and the importance ranks of the attributes (from important to unimportant) according to the SHAP graph results are given in Table 3. Upon examining Table 3 and Figure 2 show that the impact of the dataset attributes on the outcome also varies across the MLP architecture. For example, in Architecture 1, the importance ranking according to the effect of the attributes is 5-0-2-1-3-8-7-6-4, while in Architecture 4, the importance of the attributes ranking is 5-2-8-7-1-0-3-4-6. Here, Feature0 is ranked second from the beginning in Architecture 1, while in Architecture 4, it is ranked sixth from the beginning. Feature5 ranks first in all architectures, while the other features vary in importance (order) according to architecture. Since the importance rate of Feature5 is slightly different from the other features, it ranks first in all models. Since the importance ratios of other features are close to each other, it is seen that the ranking of importance ratios changes according to the model architecture. Examina the Table 3, it is seen that the ranking of the importance ratios of the features,

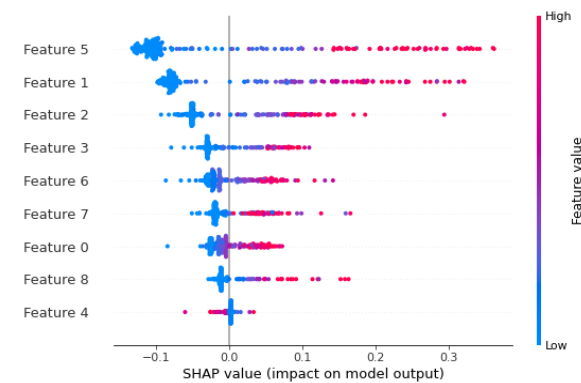
changes even if they have the same accuracy value in different architectures. For example; although both Architecture1 and Architecture8 have an accuracy of 98%, the importance ranking of the attributes is different. Thus, it can be concluded that there is no relationship between the model's accuracy value and the importance ratios of the attributes.

**Table 3.** Listing the effects of breast cancer input dataset features of the output according to architecture (Meme kanseri girdi veri kümesi özelliklerinin model çıktısına etkilerinin mimariye göre listesi).

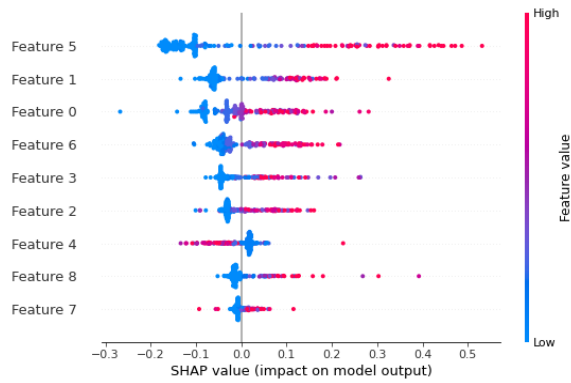
| MLP Architecture | Accuracy | The importance of the input data set attributes according to their SHAP values, listed in descending order |
|------------------|----------|------------------------------------------------------------------------------------------------------------|
| 9-5-10-5-1       | 0.98     | 5-0-2-1-3-8-7-6-4                                                                                          |
| 9-5-10-10-1      | 0.98     | 5-1-2-3-6-7-0-8-4                                                                                          |
| 9-5-5-1          | 0.95     | 5-0-6-2-8-3-1-4-7                                                                                          |
| 9-10-5-1         | 0.97     | 5-2-8-7-1-0-3-4-6                                                                                          |
| 9-10-15-1        | 0.96     | 5-1-0-6-3-2-4-8-7                                                                                          |
| 9-15-1           | 0.98     | 5-2-0-6-1-3-8-7-4                                                                                          |
| 9-45-1           | 0.98     | 5-0-2-6-3-1-8-7-4                                                                                          |
| 9-15-10-10-15-1  | 0.98     | 5-1-0-2-6-3-8-4-7                                                                                          |



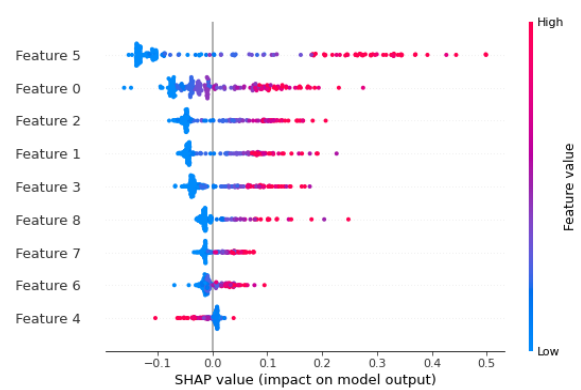
(c) MLP Architecture: 9-5-5-1



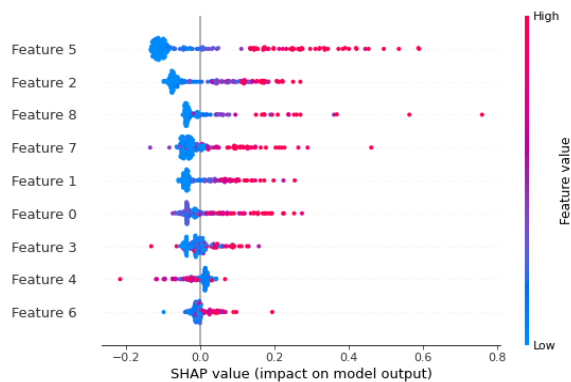
(d) MLP Architecture: 9-5-10-10-1



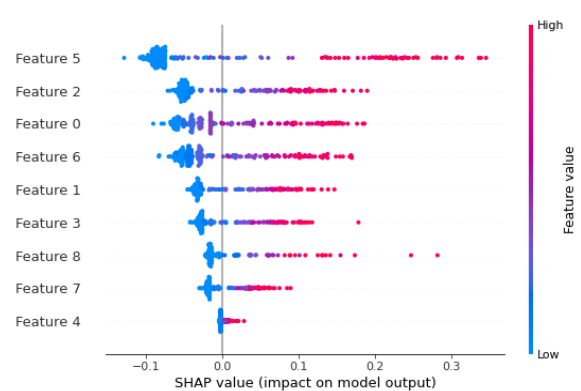
(a) MLP Architecture: 9-10-15-1



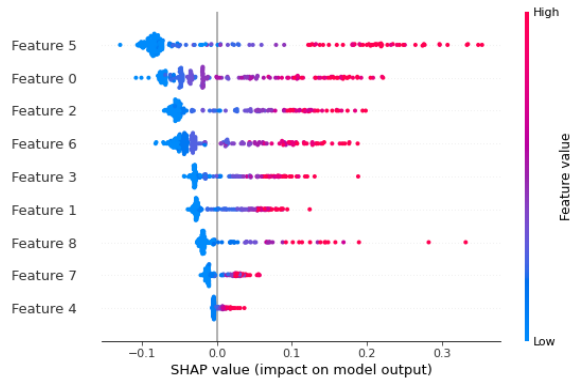
(e) MLP Architecture: 9-5-10-5-1



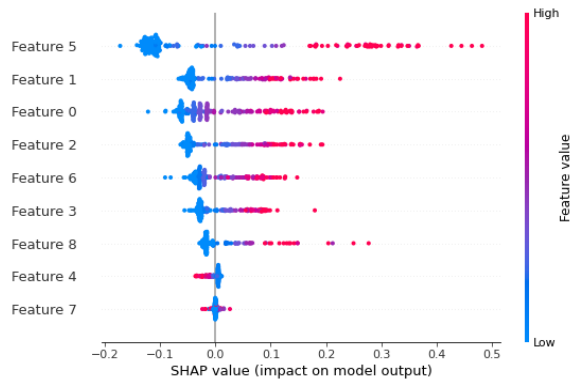
(b) MLP Architecture: 9-10-5-1



(f) MLP Architecture: 9-15-1



(g) MLP Architecture: 9-45-1



(h) MLP Architecture: 9-15-10-10-15-1

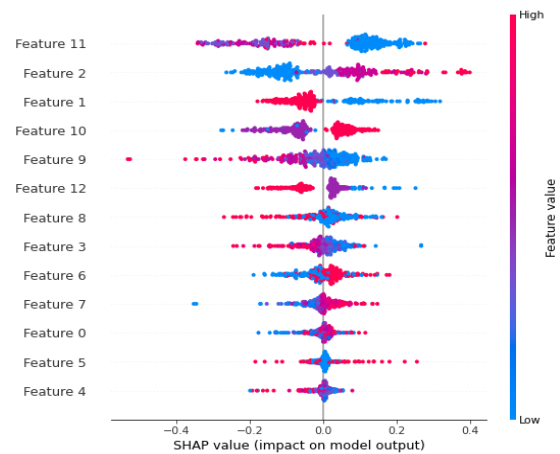
**Figure 2.** SHAP graphs of the breast cancer dataset obtained after MLP training for different architectures (Meme kanseri veri seti ile eğitilmiş, farklı mimarideki MLP'lerin, eğitim sonrasında elde edilen SHAP grafikleri).

The effects of the attributes of the input dataset on the output for different MLP architectures of the heart disease dataset are given in Figure 3. The model architecture was not taken according to any rule. Different architectures were created by using different numbers of hidden layers and different numbers of nodes (units) in each hidden layer. Table 4 shows the structure of the architectures, model performance, and the effects of the dataset attributes on the output according to the model, ranked in descending order of importance. Examine Figure 3 and, the Table 4 are analyzed, it is seen that the ranking of the impact values of the dataset attribute values on the model output changes. For example, in Architecture 1, Feature10 is ranked fourth from the beginning, while in Architecture 5 it is ranked seventh. The importance of some features remained the same in all architectures. For example, Feature1 has maintained its first rank in all architectures. Examining the Figure 3, it can be seen that these features are relatively more important than other features. Feature4 ranks last in

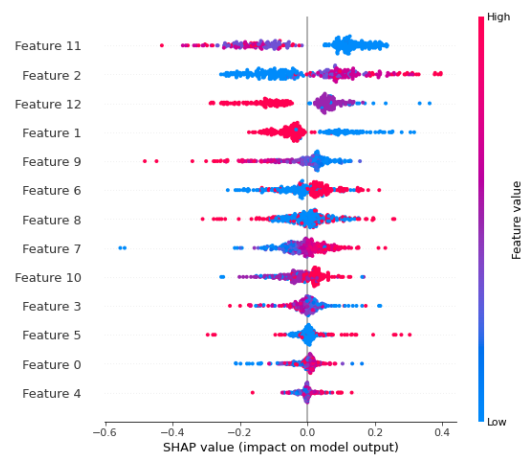
almost all models. This feature is the least important feature for all models.

**Table 4.** Listing the effects of the input dataset attributes on the output of the heart disease dataset according to different architectures (Kalp hastalığı veri kümesi özelliklerinin model çıktısına etkilerinin mimariye göre listesi)

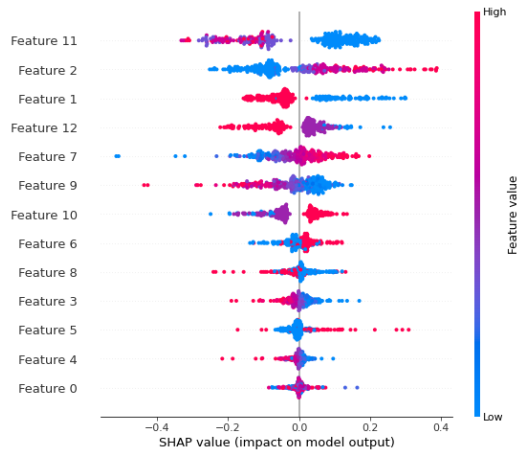
| MLP Architecture | Accuracy | The rating of input data set qualities' contribution to the output |
|------------------|----------|--------------------------------------------------------------------|
| 13-15-5-1        | 0.94     | 11-2-1-10-9-12-8-3-6-7-0-5-4                                       |
| 13-20-15-1       | 0.98     | 11-2-12-1-9-6-8-7-10-3-5-0-4                                       |
| 13-15-1          | 0.91     | 11-2-1-10-9-12-8-3-6-7-0-5-4                                       |
| 13-45-1          | 0.97     | 11-2-1-10-12-9-7-6-8-5-0-3-4                                       |
| 13-15-10-5-1     | 0.98     | 11-2-1-9-12-8-10-6-7-5-0-4-3                                       |
| 13-15-30-15-10-1 | 0.98     | 11-1-2-12-10-9-8-6-0-5-7-3-4                                       |



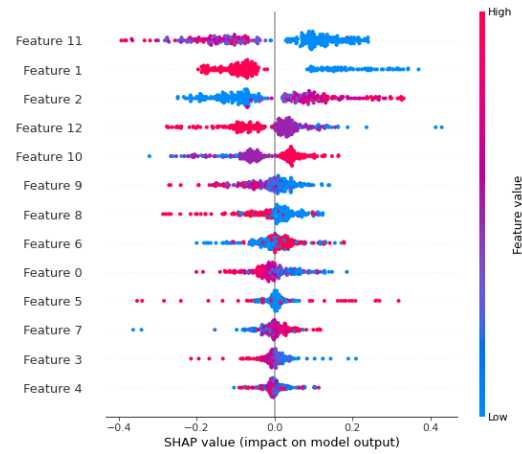
(a) MLP Architecture: 13-15-5-1



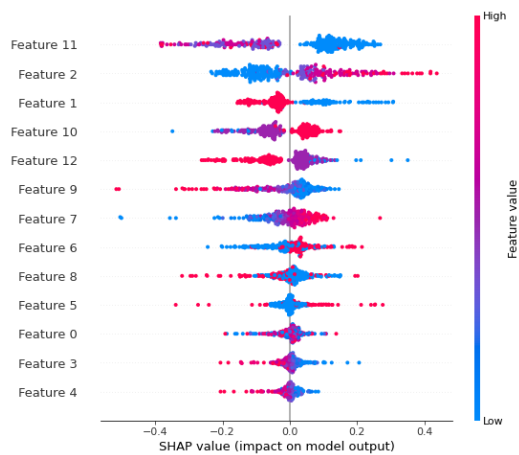
(b) MLP Architecture: 13-20-15-1



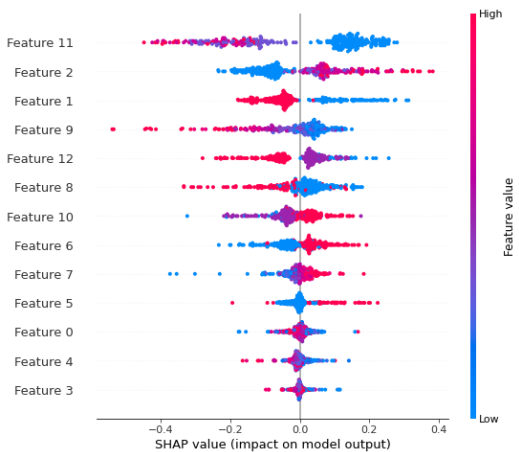
(c) MLP Architecture: 13-15-1



(f) MLP Architecture: 13-15-30-15-10-1



(d) MLP Architecture: 13-45-1



(e) MLP Architecture: 13-15-10-5-1

**Figure 3.** SHAP plots of the heart disease dataset obtained after MLP training for different architectures (Kalp hastalığı veri seti ile eğitilmiş, farklı mimarideki MLP'lerin, eğitim sonrasında elde edilen SHAP grafikleri).

## 5. CONCLUSION

In MLP systems, which contain too many parameters in artificial intelligence models and are seen as black-boxes, the explainability and interpretability of model outputs, the effects of the dataset on the output and the differences in the importance of dataset attributes according to MLP architecture were examined. Breast Cancer Wisconsin and Heart Disease datasets, which are frequently used in the literature, were used datasets. The reason for choosing this dataset is that issues such as explainability, interpretability, and transparency are much more important in the use of artificial intelligence in the health sector. Both datasets were trained and tested on MLP models with different architectures and the accuracy performance values of the models were measured. The SHAP method was used to track the impact and importance of the input dataset on the model output on the trained model. A test data set was used for the SHAP process. The graphs in Figure 2 and Figure 3 show the importance levels of the dataset attributes sequentially. The accuracy and importance of the dataset attributes obtained for both datasets are transferred to Table 3 and Table 4, respectively. When Figure 2 and Figure 3 and Table 2 and Table 3 are analyzed;

Input dataset features importance varies depending on the MLP architecture,

Different MLP architectures achieving similar accuracy can still produce different feature importance rankings,

MLP models with the same architecture but different performance can yield different feature importance rankings,

The impact of attributes with relatively high importance is not much affected by the MLP architecture,

The relative ranking of features with similar importance levels is more likely to change between different architectures,

No direct relationship between the model accuracy performance value and the ranking of attribute influence (importance) values in artificial intelligence MLP methods,

It was observed that the relative importance of input features could be measured using SHAP for MLP models. Accordingly, it can be concluded that the architecture of the MLP models is also an issue that needs to be considered in terms of explainability, interpretability and transparency of the input-output relationship of artificial intelligence models.

The study can be extended by using, in MLP, the relationships of other hyper-parameters such as learning algorithm, activation functions, loss function, etc. to the effect rates of data attribute.

## ACKNOWLEDGMENT

No potential conflict of interest was reported by the author(s).

## REFERENCES

- [1] R. A. S. Deliloğlu and A. Ç. Pehlivanlı, "Hybrid Explainable Artificial Intelligence Design and LIME Application", *European Journal of Science and Technology*, No. 27, pp. 228-236, November 202 Copyright © 2021 EJOSAT.
- [2] P. Gohel., Singh P., and Mohanty M., "Explainable AI: current status and future directions", *IEEE Access* (2021), 10.48550/arXiv.2107.07045, arXiv:2107.07045, DOI: <https://doi.org/10.48550/arXiv.2107.07045>.
- [3] W. Saeed and C. Omlin, "Explainable AI (XAI): A systematic meta-survey of current challenges and future opportunities, *Knowledge-Based Systems*", Volume 263, 110273, ISSN 0950-7051, 2023, <https://doi.org/10.1016/j.knosys.2023.110273>.
- [4] A. Chander, R. Srinivasan, S. Chelian, J. Wang, and K. Uchino, "Working with beliefs: AI transparency in the enterprise". In *IUI Workshops* (Vol. 1), 2018, March.
- [5] A. Barredo Arrieta, N. Díaz-Rodríguez, J. D. Ser, A. Bennetot, S. Tabik, A. Barbado, S. Garcia, S. Gil-Lopez, D. Molina, R. Benjamins, R. Chatil, F. Herrera, "Explainable Artificial Intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI", *Information Fusion*, Volume 58, Pages 82-115, ISSN 1566-2535, (2020), <https://doi.org/10.1016/j.inffus.2019.12.012>.
- [6] E. Tjoa and C. Guan, "A Survey on Explainable Artificial Intelligence (XAI): towards Medical XAI", *Journal of LaTeX Class Files*, Vol. 14, No. 8, August 2015.
- [7] F. K. Došilović, M. Brčić and H. Hlupić, "Explainable artificial intelligence: A survey", 2018 41st International Convention on Information and Communication Technology, Electronics and Microelectronics (MIPRO), Opatija, Croatia, 2018, pp. 0210-0215, doi: 10.23919/MIPRO.2018.8400040.
- [8] Y. Izza, A. Ignatiev, J. Marques-Silva, "On Explaining Decision Trees", arXiv:2010.11034 [cs.LG], <https://doi.org/10.48550/arXiv.2010.11034>.
- [9] M. A. Hearst, S. T. Dumais, E. Osuna, J. Platt and B. Scholkopf, "Support vector machines," in *IEEE Intelligent Systems and their Applications*, vol. 13, no. 4, pp. 18-28, July-Aug. 1998, doi: 10.1109/5254.708428.
- [10] N. Cristianini, E. Ricci, "Support Vector Machines", In: Kao, MY. (eds) *Encyclopedia of Algorithms*. Springer, Boston, MA, 2008, [https://doi.org/10.1007/978-0-387-30162-4\\_415](https://doi.org/10.1007/978-0-387-30162-4_415).
- [11] X. Zou, Y. Hu, Z. Tian and K. Shen, "Logistic Regression Model Optimization and Case Analysis," 2019 IEEE 7th International Conference on Computer Science and Network Technology (ICCSNT), Dalian, China, 2019, pp. 135-139, doi: 10.1109/ICCSNT47585.2019.8962457.
- [12] F. J. Yang, "An Implementation of Naive Bayes Classifier," 2018 International Conference on Computational Science and Computational Intelligence (CSCI), Las Vegas, NV, USA, 2018, pp. 301-306, doi: 10.1109/CSCI46756.2018.00065.
- [13] S. Norozpour and M. Safaei, "An Overview on Game Theory and Its Application", *IOP Conf. Series: Materials Science and Engineering* 993 (2020) 012114, ICMECE 2020, doi:10.1088/1757-899X/993/1/012114.

- [14] S. M. Lundberg and S. Lee, "A Unified Approach to Interpreting Model Predictions", 31st Conference on Neural Information Processing Systems (NIPS 2017), Long Beach, CA, USA,
- [15] SHAP documents, <https://shap.readthedocs.io/en/latest/index.html>, (Access Date: June 4 2024).
- [16] UCI, "Machine Learning Repository", <https://archive.ics.uci.edu/> (Access Date : June 05 2023).
- [17] M. Polatgil, "Investigation of The Effects of Data Scaling and Imputation of Missing Data Approaches on The Success of Machine Learning Methods", Duzce University Journal of Science and Technology, 11 78-88, 2023, DOI:10.29130/dubited.948564.
- [18] T. Emmanuel, T. Maupong, D. Mpoeleng, T. Semong, B. Mphago, O. Tabona, "A survey on missing data in machine learning", J Big Data 8, 140, (2021), <https://doi.org/10.1186/s40537-021-00516-9>.
- [19] S. Qinbao and M. J. Shepperd, "Missing Data Imputation Techniques", IJBIDM. 2. 261-291. 10.1504 / IJBIDM.2007. 015485.
- [20] C. Molna, "Interpretable Machine Learning A Guide for Making Black Box Models Explainable", Independently published (February 28, 2022), Page:168-188, ISBN-13 : 979-8411463330.
- [21] M. T. Ribeiro, S. Singh and C. Guestrin, "Why W. H. A. Wan Zunaidi, R. R. Saedudin, Z. Ali Shah, S. Kasim, C. Sen Seah, and M. Abdurrohman, "Performances Analysis of Heart Disease Dataset using Different Data Mining Classifications", Int. J. Adv. Sci. Eng. Inf. Technol., vol. 8, no. 6, pp. 2677–2682, Dec. 2018. Should I Trust You? Explaining the Predictions of Any Classifier", KDD 2016 San Francisco, CA, USA, Publication rights licensed to ACM. ISBN 978-1-4503-4232-2/16/08., 2016, DOI: <http://dx.doi.org/10.1145/2939672.2939778>.
- [22] H. W. Kuhn, "Classics in Game Theory", Princeton University Press, 1997, ISBN-13: 978-0691011929.
- [23] P. Nagaraj, V. Muneeswaran, A. Dharanidharan, K. Balanathanan, M. Arunkumar, C. Rajkumar, "A Prediction and Recommendation System for Diabetes Mellitus using XAI-based Lime Explainer", 2022 International Conference on Sustainable Computing and Data Communication Systems (ICSCDS), IEEE Xplore: 27 April 2022, DOI: 10.1109/ICSCDS53736.2022.9760847.
- [24] <https://archive.ics.uci.edu/dataset/17/breast+cancer+wisconsin+diagnostic>(Access Date: April 3 2025)
- [25] H. Doğan, A. B. Tatar, A. K. Tanyıldızı, B. Taşar, "Breast Cancer Diagnosis with Machine Learning Techniques", Bitlis Eren University Journal Of Science, ISSN: 2147-3129/e-ISSN: 2147-3188 Volume: 11 No: 2 Page: 594-603 Year: 2022 DOI: 10.17798/bitlisfen.1065685.
- [26] V. Asha, B. Saju, S. Mathew, A. M. V, Y. Swapna and S. P. Sreeja, "Breast Cancer classification using Neural networks," 2023 International Conference on Intelligent and Innovative Technologies in Computing, Electrical and Electronics (IITCEE), Bengaluru, India, 2023, pp. 900-905, doi: 10.1109/IITCEE57236.2023.10091020.
- [27] A. Çifci, M. İlkuçar, İ. Kırbaş, "Explainable Artificial Intelligence for Biomedical Applications/What Makes Survival of Heart Failure Patients? Prediction by the Iterative Learning Approach and Detailed Factor Analysis with the SHAP Algorithm". Publisher: River Publishers, 2023. Chapter:6, ISBN: 9788770228497 (Hardback), 9788770228848 (Ebook)
- [28] A.A. Ahmad, H. Polat, "Prediction of Heart Disease Based on Machine Learning Using Jellyfish Optimization Algorithm. Diagnostics", (Basel). 2023 Jul 17; 13(14):2392. doi: 10.3390/diagnostics13142392.
- [30] Y. Lin, "Prediction and Analysis of Heart Disease Using Machine Learning," 2021 IEEE International Conference on Robotics, Automation and Artificial Intelligence (RAAI), Hong Kong, Hong Kong, 2021, pp. 53-58, doi: 10.1109/RAAI52226.2021.9507928.



## RSO Kullanıcılarının RSO'ya İlişkin Algısı

Selçuk KIRAN<sup>a,\*</sup>, Sena SÜZÜK<sup>b</sup>

<sup>a,\*</sup> Marmara Üniversitesi İşletme Fakültesi, Yönetim Bilişim Sistemleri Bölümü, İstanbul, 348541, Türkiye

<sup>b</sup> Marmara Üniversitesi İşletme Fakültesi, Yönetim Bilişim Sistemleri Bölümü, İstanbul, 348541, Türkiye

### MAKALE BİLGİSİ

Alınma: 02.11.2024  
Kabul: 12.03.2025

**Anahtar Kelimeler:**  
RSO, Robotik Süreç  
Otomasyonu, Süreç  
Yönetimi, Otomasyon.

#### **\*Sorumlu Yazar**

e-posta:  
selcuk.kiran@marmara.e  
du.tr

### ÖZET

Tekrarlanan ve rutinleşen görevler için insan davranışını taklit edebilen, yazılım araçlarına dayalı bir otomasyon teknolojisi olan RSO teknolojisi son yıllarda, özellikle tekrar eden görevlerde hem bütçe tasarrufu sağlamakta hem de çalışanların yeteneklerini daha iyi gösterecekleri alanlar için vakitlerini ayırabilmelerini sağlamaktadır. Bu makalede de çoğunluğu yurt içinde yaşayan 221 RSO ile ilintili çalışanın (RSO kuran, RSO yazılımı programlayan, RSO ile iş süreçleri düzenlenen vb.) RSO teknolojisine bakışları, bir anket vasıtasıyla incelenmiştir. Çalışmanın sonucunda bazı konularda kadınların erkeklerden ve yurt dışında yaşayanların yurt içinde yaşayanlardan ayrışıkları görülmüştür. Diğer yandan RSO kullanışlığı ile ilgili ölçeklerin regresyon testi yardımıyla yapılan analizlerinde farklı sonuçlar bulunmuştur. Bunlara göre işlevsellik kullanıcı dostluğu getirmekte, RSO yazılımının avantajları projelerde başarı getirip ürünün benimsenmesini sağlamaktadır. Diğer yandan ankete katılanların çoğunun benzer markadaki bir ürünü kullandıkları ve bu ürünü kullananların diğerlerine göre daha mutlu oldukları görülmüştür..

DOI: 10.59940/jismar.1578001

## Perception of RPA Users on RPA

### ARTICLE INFO

Received: 02.11.2024  
Accepted: 12.03.2025

**Keywords:**  
RPA, Robotic Process  
Automation, Business  
Process, Automation.

#### **\*Corresponding Authors**

e-mail:  
selcuk.kiran@marmara.e  
du.tr

### ABSTRACT

In recent years, RPA technology, which is an automation technology based on software tools that can imitate human behavior for repetitive and routine tasks, has been providing budget savings, especially in repetitive tasks, and allowing employees to spare their time for areas where they can better demonstrate their talents. In this article, the opinions of 221 RPA-related employees (who install RPA, program RPA software, organize business processes with RPA, etc.), most of whom live in Turkey, on RPA technology were analyzed through a survey. As a result of the study, it was seen that women differ from men and those living abroad differ from those living in Turkey in some issues. On the other hand, different results were found in the analyzes of the scales related to RPA usability using regression testing. Accordingly, functionality brings user friendliness, and the advantages of RPA software bring success to projects and ensure product adoption. On the other hand, it was observed that most of the participants in the survey used a product of a similar brand and that those who used this product were happier than others.

DOI: 10.59940/jismar.1578001

### 1. GİRİŞ (INTRODUCTION)

Özellikle iş hayatındaki bazı süreçlere insanların yerine robotların dahil edilmesiyle ortaya çıkan

Robotik Süreç Otomasyonu (RSO – RPA Robotic Process Automation) teknolojisinin kısa sürede büyük gelişim göstermesi neticesinde firmalar iş süreçlerinde bu teknolojiyi kullanma eğilimine

girmişlerdir. Bu araştırmada önce RSO teknolojisinin hangi süreçlerden geçerek bugünlere geldiği incelenmektedir. Daha sonrasında RSO kullanıcılarına yapılan anketler vasıtasıyla hem Türkiye’deki hem de yurt dışındaki RSO kullanımı ve uzmanların bu teknolojiye bakışları incelenmeye çalışılmıştır.

Bunun için daha önce yapılmış olan çalışmalardan faydalanılarak üç bölümlü 19 sorudan oluşan bir anket hazırlanmıştır. Birinci bölümde katılımcının demografik bilgisi ve işiyle ilgili bilgiler, ikinci bölümde RSO kullanımı ve deneyimi ile ilgili bilgiler, son bölümde ise kullanıcının RSO’ları değerlendirmesi alınmaktadır. Eksik yanıtlar, analiz sürecinde veri kaybına neden olabilir ve bunları tamamlamak için kullanılan istatistiksel yöntemler (örneğin, ortalama ile doldurma, çoklu atama gibi) verinin doğasını değiştirebilir ve çalışma sonuçlarını etkileyebilir [1]. Diğer yandan eksik veriler, özellikle küçük örneklemelerde, bulguların güvenilirliğini ve geçerliliğini azaltabilir [2]. Bu sebeplerle katılımcıların tüm soruları yanıtlaması zorunlu tutulmuştur.

Bu anket vasıtasıyla RSO geliştiricileri, sorumluları ve ürün sahiplerinden oluşan 221 kişinin RSO teknolojisini benimseme nedenlerine, bu teknolojinin avantajlarına, kullanım kolaylığına, kullanışlılığına ve başarı kriterlerine bakışları ölçülmüştür. Bunun yapılması esnasında anket iki dilde (Türkçe ve İngilizce) hazırlanmış olup Google Forms vasıtasıyla yurt içindeki ve yurt dışındaki kullanıcılara sunulmuştur.

Doldurulan anketlerin üzerinde, SPSS yazılımı yardımıyla, Cronbach Alpha güvenilirlik testi, korelasyon testi, t-testleri ve regresyon testi uygulanmıştır. Bunların sonucunda bazı konularda kadınların erkeklerden ve yurt dışında yaşayanların yurt içinde yaşayanlardan ayrıştıkları görülmüştür. Diğer yandan RSO kullanışlılığı ile ilgili ölçeklerin regresyon testi yardımıyla yapılan analizlerinde ilginç sonuçlar bulunmuştur. Bunlara göre işlevsellik kullanıcı dostluğu getirmekte, RSO yazılımının avantajları projelerde başarı getirip ürünün benimsenmesini sağlamaktadır. Diğer yandan ankete katılanların çoğunun UIPath ürününü kullandıkları ve bu ürünü kullananların diğerlerine göre daha mutlu oldukları görülmüştür.

## 1. LİTERATÜR TARAMASI (LITERATURE SURVEY)

Dijital teknolojilerin ve iletişim ağının zamanla gelişmesi, pazara ve hammaddeye erişim kolaylığını, artan üretim hacmini ve ürün çeşitliliğini de beraberinde getirmiştir. Bunun sonucunda küresel

pazarda rekabet artmış, ülkeler pazarda rekabet avantajı elde edebilmek için maliyetlerini düşürmenin ve tüketici ihtiyaçlarını rakiplerine göre daha hızlı ve etkin bir şekilde karşılamının yollarını aramaya başlamışlardır. Bunun sonucunda sanayisi gelişmiş olan ülkeler, rekabet avantajı elde etmek için teknolojik olanaklarını geliştirmiş ve Almanya’nın öncülüğünde dördüncü sanayi devrimini başlatmışlardır [3].

Dördüncü sanayi devrimi kapsamında “Akıllı Üretim”, “Robot Sistemler”, “Nesnelerin İnterneti”, “Bulut Bilişim”, “Simülasyon”, “Artırılmış Gerçeklik” gibi bazı terimler popülerleşmişlerdir. Endüstri 4.0’ın unsurları olan bu teknolojiler birleşerek kendi kendini kontrol edebilen üretim sistemleri kurmayı amaçlamaktadır. Bu sistemler sayesinde firmalar esnekliklerini ve verimliliklerini artırabilir, maliyetlerini azaltabilir ve süreç hatalarını daha kolay gözlemleyip ortadan kaldırabilirler. Bunların yanı sıra çalışanların bilgisayarlarında özel olarak tanımlanmış süreçleri yönetmek için yazılım robotları kullanılmaktadır. Yazılım robotları endüstriyel robotların verimliliğini artırmanın yanı sıra çalışanların da verimliliğini artırmaktadır. Yazılım robotlarının iş dünyasında kullanılması fikrinin yaygınlaşmasıyla birlikte son beş yılda RSO teknolojisi öne çıkmış ve bu teknolojinin kullanımı şirketler arasında yaygınlaşmıştır. RSO teknolojisi, uygulandıkları süreçlerin otomatikleştirilmesine dayanmaktadır. Bu otomatikleşme sayesinde ilgili süreçleri çalışanların yerine yazılım robotları yürütebilmektedir.

İş akışı süreçlerindeki bazı adım ve eylemlerden oluşan bir dizi işlemi temsil eden ve bu eylemleri belirli bir uygulama için otomatikleştiren yazılıma makro denir [4]. Fare ve klavyeyle yapılan tüm hareketler makro sayesinde kaydedilir ve makro çalıştırıldığında, kaydedilen eylemler taklit edilir. Küçük de olsa bir problem çıkması (örneğin, işlemin çalışma zamanındaki gecikme), makronun işlemi çalıştıramamasına neden olabilir. Grafiksel kullanıcı arayüzünün (GUI) ortaya çıkmasından sonra, bilgisayarları kontrol etmek için görsel programlama dilleri ortaya çıkmıştır. Bu diller, pencere, menü, düğme gibi araçlarla insan gibi etkileşime girerek, kullanıcıların eylemleri kontrol etmelerini sağlar. Bu diller kullanılarak yazılan komut dosyaları otomatik zamanlayıcı vasıtasıyla kullanıcıdan bağımsız çalıştırılabilir. GUI kodlamaya örnek olarak VBScript dilini kullanır [5]. Makrolar ve GUI komut dosyaları en ilkel otomasyon araçları olarak kabul edilebilirler.

Ostdick'e göre RSO teknolojisinin üç farklı öncülü vardır [6], Bunlardan ilki olan ekran kazıma yazılımı 1990'larda ortaya çıkmıştır ve bir uygulama ekranında bulunan verilerin görüntülenip sonra kullanılmak üzere kaydedilmesini kapsar [7]. İkinci öncül olan iş akışı otomasyonu ise, iş akışına ilişkin görev, belge ve veri akışının, tanımlanmış iş kurallarına göre bağımsız olarak yürütüldüğü, 1990'larda ortaya çıkan bir yaklaşımdır. İş akışı otomasyonu, artan müşteri ve çalışan memnuniyeti, artan ölçeklenebilirlik ve üretim, tasarruf edilen maliyetler, daha az manuel işlem ihtiyacı ve insan hatası gibi birçok avantajı beraberinde getirir. Diğer yandan RSO'nun üçüncü öncülü olan, büyük miktarda veriyi karmaşık algoritmalar yardımıyla işleyerek örüntüleri tanıyan ve sonraki işlemlerde kullanabilen Yapay Zekâ [4], makineler tarafından yürütülen insan zekasının bir simülasyonu olarak da adlandırılabilir [8]. 2000'li yılların başında RSO, ekran kazıma, iş akışı otomasyonu ve yapay zekâ gibi teknolojilerin kullanıldığı ve bunların daha da geliştirildiği teknoloji olarak ortaya çıkmıştır [9].

Bahsedilen teknik yaklaşımların yanı sıra işletmeleri etkin bir şekilde yönetmek için bir takım yöntem ve disiplinler oluşturulmuştur. Bu disiplinlerden bazıları, İş Süreci Yeniden Yapılanması (BPR), İş Süreçleri Yönetimi (BPM) ve İş Süreçleri Otomasyonu'dur (BPA). Süreç bu disiplinlerin içindeki en küçük unsurdur ve belirli bir sonucu (ürün veya hizmet) üretmek için birbirleriyle etkileşime giren emek, ekipman, malzeme, yöntem ve çevresel unsurlarının toplamına denir [10]. Organizasyonel süreç, işi tamamlamak için gereken alt işlerden oluşturulan bir gruptur, bir başlangıcı ve sonu vardır. Diğer taraftan, bir veya daha fazla girdinin müşteri için değer yaratan bir çıktıya dönüştürüldüğü faaliyetler bütünü olarak da tanımlanabilir. Süreç yönetimi, süreçlerin sürekli ve düzenli olarak izlenmesini ve iyileştirilmesini sağlamak amacıyla yürütülen faaliyetler bütünüdür. Bir başka açıklamaya göreyse süreç yönetimi, süreçleri tasarlamak, sürdürmek ve müşteri ihtiyaçlarını daha iyi karşılamak için, sürekli değerlendirme, analiz ve iyileştirmeyi içeren bir döngüdür [11].

İş Süreci Yeniden Yapılanması (BPR), ilk kez 1993 yılında gündeme gelmiştir. BPR disiplinine göre, düzeltilmesi gereken süreçler doğrudan ortadan kaldırılarak gerekli süreçler yeniden planlanıp kurgulanır. Uygulamada BPR'nin önerdiği radikal değişikliklerin şirketlerde hayata geçirilmesi pek mümkün görünmemektedir. Bu disiplinin yerine daha bütünsel bakış açısına sahip, kademeli iyileştirmelere odaklanan bir disiplin aranmıştır [12]. Bu arayışın bir parçası olarak iş süreçleri yönetimi (BPM) devreye girebilir.

Süreç otomasyonuna geleneksel bir yaklaşım olan BPM, kuruluşlardaki iş akışlarını yönetmek için kullanılan süreç metodolojisidir. BPM'nin temel amacı, daha iyi getiri için iş çevikliğini ve operasyonel mükemmelliği teşvik ederek organizasyonel performansı artırmaktır. BPM, farklı modelleri, metrikleri ve değerlendirmeleri inceleyerek tüm gereksinimleri ve süreci tek bir çalışmada toplar ve daha iyi performans için gerekli iyileştirmeleri belirler [13]. BPM; veriler, kişiler ve sistemler arasında aktif koordinasyonu içerir ve operasyonel iş süreçleri için sağlam bir altyapı sunmayı amaçlar. BPM, bir yandan şirketteki tüm çalışan ve sistem hareketlerini belirlerken diğer yandan süreçte oluşan verilerin de depolanmasını sağlar.

Geleneksel süreç otomasyonu BPM kapsamında incelendiğinde İş Süreçleri Otomasyonu (BPA) olarak adlandırılır. BPA, mevcut süreçleri düzene sokup verimsizlikleri ortadan kaldırarak süreç iyileştirmelerini amaçlar. Daha iyi sistem entegrasyonu veya özel süreç yazılımı yoluyla otomatikleştirilebilecek tüm süreç adımlarını öne çıkaran, stratejik bir bilgi sistemi dönüşüm eylemidir [14]. Firmaların pazardaki verimliliğini ve rekabet gücünü artırmaya yönelik yeni çözüm arayışlarında ortaya çıkan BPM kavramının yardımıyla şirketler, iş süreçlerini kapsamlı bir şekilde iyileştirmeye odaklanırlar. Çalışanların süreçleri yürüttüğü sistemleri optimize etmek için İş Süreçleri Yönetim Yazılımı (BPMS) çözümlerinin kullanılmasıyla mevcut bilgi işlem sistemlerinde değişiklikler yapılması gerekmiştir. BPM kapsamında iş süreçlerinde ve Bilgi Teknolojisi (BT) sistemlerinde yaşanan köklü değişiklikler de şirketlere ek maliyetler getirmiştir. Bu nedenle şirketler, iş süreçlerini en az maliyet ve çabayla optimize etmelerine olanak tanıyan, kullandıkları BT sistemlerine hiçbir değişiklik yapmadan uyarlanabilecek yeni ürün arayışına girmişlerdir. RSO 2012 yılında bu gibi ihtiyaçlar kapsamında kullanılabilecek yeni bir teknoloji olarak ortaya çıkmıştır [4].

Aguirre ve Rodriguez, RSO'yu tekrarlanan ve rutinleşmiş görevler için insan davranışını taklit edebilen, yazılım araçlarına dayalı bir otomasyon teknolojisi olarak tanımlamıştır [15]. Diğer yandan Siderska, RSO'nun yaptıklarını, verilerin otomatik olarak yazılması, kopyalanması, yapıştırılması, çıkarılması, birleştirilmesi ve taşınması şeklinde örneklemiştir [16]. RSO, süreç içinde, "robotik otomasyon", "botlar", "yazılım robotları", "robot kontrolü", "süreç otomasyonu" gibi adlarla da isimlendirilmiştir [17]. Aguirre ve Rodriguez'e göre belirleyici sonuçlar içeren veri ve kural tabanlı iş süreçlerinin otomatikleştirilmesine RSO denir [15].

Bu tanımlara göre şirketlerin RSO kullanarak çalışanların iş yükünü azaltmayı amaçladıkları anlaşılmaktadır. Czarnecki ve Fettke, RSO'nun amacını, üretim süreçlerinde fiziksel robotların devreye girmesi gibi yazılım süreçlerinde de yazılım robotlarının insanların görevlerini devralması şeklinde özetlemektedir [18].

RSO, belirli kurallar ve yapılandırılmış veriler içeren süreçlerin otomatikleştirilmesi için kullanılabilir. Süreçlerin RSO ile otomatikleştirilebilmesi için basit etkinliklerden oluşmaları ve tekrarlanmaları gerekmektedir [19]. RSO, çalışanın eylemlerini kaydeder ve sonrasında bu eylemleri taklit eder. Syed vd.'nin, RSO alanında yazılmış olan 125 makaleyi inceleyen çalışmalarına bakıldığında, RSO teknolojisi, "kural tabanlı", "insan görevlerini yerine getirecek/yürütecek yazılımlar", "insan davranışlarının taklit edilmesi", "iş süreçlerinin teslimi" gibi konularda "yapılandırılmış veriler" ve "tekrarlanan görevler" ile ilişkilendirilmektedir [20].

RSO'yu diğer teknolojilerden ayıran en önemli özelliklerden biri RSO teknolojisinin, uygulamanın hem kullanıcı arayüzünü hem de Uygulama Programlama Arabirimini (API) kullanabilmesidir. API'si olmayan ancak erişilmesi gereken bir sistem varsa RSO, ilgili sisteme kullanıcı arayüzü üzerinden erişebilir yani kullanıcı arayüzü üzerinden erişilen sistemlerde herhangi bir değişiklik yapılmasına veya bir arayüz kurulmasına gerek yoktur [21][22]. Geleneksel otomasyon çözümleri "içten dışa" yaklaşımını takip eder, buna göre uygulama özelleştirilir veya sıfırdan yeniden geliştirilir. Geleneksel otomasyon çözümlerinin aksine RSO, "dışarıdan içeriye" yaklaşımını izler. Bu yaklaşıma göre bilgi sistemi değiştirilmez ve mevcut sistemin programlarına müdahale edilmez [22][23].

RSO teknolojisi programlama bilgisi olmayan sıradan geliştiriciler tarafından da kullanılabilir bir yapıya sahiptir. Sıradan geliştirici, yazılımı az kod yazarak veya hiç kod yazmadan geliştiren, kodlama konusunda tecrübesiz geliştiriciye verilen isimdir [24]. Bu görüşe göre RSO teknolojisinin bir parçası olarak kullanılan robotlar programlanmamakta, aksine eğitilmekte ve/veya yapılandırılmaktadır. Eğitim, sürük-le-bırak yöntemiyle kaydedilen kullanıcı etkileşimlerinin bir kombinasyonudur ve önceden oluşturulmuş modüllerin akış şeması benzeri bir süreç halinde yapılandırılması yoluyla gerçekleştirilir. Özellikle daha karmaşık otomasyonlar için yapılandırma, komut dosyası dillerinde ustalık gerektirir. Scheppeler ve Weber'in açıkça belirttikleri gibi, betik dillerine hakimiyet gerektiren durumlarda programlama bilgisinin gerekli olmadığı çok da doğru değildir [25].

Temel olarak RSO'lar, gözetimli (assisted), gözetimsiz (unassisted) ve hibrit (gözetimli ve gözetimsizin beraber kullanıldığı) olmak üzere üçe ayrılırlar. Gözetimli RSO aynı zamanda "Robotik Masaüstü Otomasyonu (RDA)" olarak da bilinir ve genellikle tekrarlayan ön büro süreçleri için kullanılmaktadır [26]. RDA, doğrudan kullanıcının masasüstünde çalışan yapılandırılmış bir yazılımdır ve operatör tarafından oluşturulan iş akışlarını veya görevleri yerine getirmede operatöre gerçek zamanlı olarak yardımcı olur [27][28]. Gözetimli RSO kullanımına örnek olarak, RSO'nun yardımıyla bir çağrı merkezi uzmanının, gelen çağrıya cevap vermeden önce müşterinin bilgilerini Kurumsal Kaynak Planlama (KKP), Müşteri İlişkileri Yönetimi (MİY) yazılımları gibi farklı uygulamalardan toplaması gösterilebilir.

Langmann ve Turi gözetimsiz RSO'ya doğrudan "Robotik Süreç Otomasyonu" olarak atıfta bulunurlar. Buna göre, gözetimsiz RSO işlemleri, genellikle bir sunucu üzerinden kullanıcı etkileşimi olmadan çalıştırır [28]. Bu tip RSO ile robotlar birçok farklı sistemle bağlantı kurarak görevleri kendi aralarında paylaştırabilmektedirler. Genellikle tekrarlayan arka ofis işlemleri için kullanılırlar. Gözetimsiz RSO, belirlenmiş bir zaman planına göre çalışır ve 7/24 kullanılabilir. Örneğin, gözetimsiz RSO, güne ait faturaları günün belirli bir saatinde toplayabilir ve işleyebilir. Daha sonra bot, kullanıcıya işlenmemiş faturaları gösteren otomatik bir rapor verir ve kullanıcı raporu inceleyerek müdahalesini gerektiren faturalar üzerinde çalışır [29].

RSO yazılımı genel olarak işlerin oluşturulduğu geliştirme ortamında, planlanan işleri yürüten "robot" veya "bot" olarak da adlandırılabilir sanal çalışanların oluşturulduğu, kontrol edilebilir bir kontrol panelinden oluşur [26]. Bir geliştirme ortamı, kullanıcıların robotlar tarafından gerçekleştirilen iş süreçlerini oluşturmasına, tasarlamasına ve otomatikleştirmesine olanak tanır. Robotlar, geliştirme ortamında otomatik çalıştırılmak üzere oluşturulan, çalışanların görevlerini ve iş süreçlerini yürüten sanal çalışanlardır [30].

RSO uygulamalarının başarılı bir şekilde hayata geçirilmesi için hangi süreçlerin RSO kullanımına uygun olduğunu bilmek önemlidir. RSO'nun uygulanmasına yönelik süreçlerin doğru şekilde seçilmesi için bazı seçim kriterleri önerilmiştir. Bu kriterlere göre firmalar, sıkı şekilde kurallara dayanan, yüksek hacimli, oldukça tekrarlı ve daha az karmaşık yani daha standart süreçlere sahip durumlarda RSO'yu seçmelidir [20][31][32]. Ayrıca Capgemini Consulting tarafından yapılan araştırmaya göre beş

dakikadan uzun ve otuz dakikadan kısa olan süreçler, RSO kullanımı için uygun adaylar olabilir [33]. Van der Aalst vd. yaptıkları çalışmada, frekansı çok yüksek fakat karmaşıklığı az olan süreçlerin otomasyonu için RSO'nun uygun olduğunu öte yandan, daha karmaşık, düşük frekanslı ve spesifik süreçlerin çalışanlar tarafından yürütülmesinin daha uygun olduğunu belirtmiştir [23]. RSO için seçilen süreçler tekrarlayan, kurallara dayalı, yüksek hacimli olmalı ve insan müdahalesi gerektirmemelidir [34]. Süreçlerin standardize edilmeleri, fazla istisna içermemeleri ve iyi belgelenmiş olmaları, kullanılan verilerin dijital, kaliteli ve yapılandırılmış olmaları RSO kullanımını kolaylaştırır.

RSO ile gelen avantajlardan ilki, robotların 7 gün 24 saat çalışabilmesidir. Buradan bir robotun çalışma süresinin normal bir mesaiden dört kat kadar daha fazla olduğu hesaplanabilir. Bazı kaynaklara göre bir robot 1,7 çalışanın işini devralmaktadır [31][35]. Şirketler bazı süreçlerde insanlar yerine robotları kullanarak iş yükünü yeniden dağıtmayı düşünebilir. Bu sayede şirketlerin büyümesine bağlı olarak ortaya çıkan gelecekteki iş yükü, yeni çalışanların yerine işlerinin bir kısmını robotlara devretmiş çalışanlar arasında dağıtılabilecektir [8][36]. Bir yazılım robotu, offshore tam zamanlı çalışanın (FTE – Tam Zamanlı Eşdeğer) fiyatının üçte biri kadar bir fiyata ya da karadaki bir FTE'nin fiyatının yalnızca beşte birine mal olabilir [37].

Otomotiv sektöründeki oniki projede yapılan mülakatlara göre, RSO sayesinde proje başına 0,02 ile 1,43 arası FTE tasarrufu sağlandığı gözlemlenmiştir. Buradaki düşük FTE tasarruflarının, işlem sayısının azlığı ve/veya kısa süreç süresi nedeniyle olduğu söylenebilir [38]. FTE ve işgücü maliyeti tasarruflarına ek olarak RSO, çalışanları zaman alıcı, yorucu, tekrarlayan, monoton görevlerden ve çok az zihinsel çaba gerektiren veya hiç gerektirmeyen işlerden de kurtarmış olur. Bu, çalışanların duygusal zekâ ve akıl yürütme gibi insani becerilerini kullanarak daha fazla değer katan ve yaratıcı düşünme, entelektüel muhakeme veya sosyal beceriler gerektiren görevlere odaklanmasına olanak tanır. Bu durum şirketlere olumlu yansyarak çalışan memnuniyetinin ve sadakatının artmasına olanak tanımaktadır. Genel olarak monoton görevlerin otomasyonu süreç verimliliğine %40'dan fazla katkıda bulunur [31][32][35][39][40]. RSO kullanılsa da entelektüel bilgi ve duygusal zekâ gerektiren karar verme durumlarında çalışanların sürecin devamı için müdahale etmesi gerekebilir [41].

Forrester, 2019'da temel iş alanlarında yönetici ve üzeri pozisyonlarda çalışan 100 kişiyle, RSO yatırımlarının çalışanların işlerini nasıl etkilediğine dair bir anket gerçekleştirmiştir [42]. Bu anketin sonuçlarına göre katılımcıların yüzde 66'sı RSO'nun

çalışanların mevcut işlerini yeniden yapılandırıldığını ve çalışanların insanlarla daha fazla etkileşim kurmasına olanak sağladığını belirtmişlerdir. Süreçler otomatikleştirildikten sonra çalışanlar karar vermeye ve müşterilerine daha fazla vakit ayırabilmektedirler. Diğer yandan katılımcıların yüzde 60'ı RSO'nun çalışanların daha anlamlı stratejik görevlere odaklanmasına yardımcı olduğunu açıkça belirtmiştir. Aynı araştırmada, katılımcıların yüzde 57'si, RSO'nun elle yapılan işlemlere göre hataları azalttığını ifade etmişlerdir. Süreçler ve robotlar doğru yapılandırılıp geliştirildikleri sürece robotların hata yapmadıkları gözlemlenmiştir. Bu nedenle süreçlerin robotlara aktarılması gibi unsurlar, olası hataların önlenmesi açısından önemlidir. Robotların süreçlerde yaptığı hataları tespit etmek, insanların yaptığı hataları tespit etmekten daha kolaydır. Çünkü robotun otomasyon sürecinde attığı her adım kayıt altına alınmaktadır. Böylece bu adımlar daha sonra gözden geçirilebilir hale gelir ve hata durumunda ilgili hata kolaylıkla bulunabilir [8][22][37][39][42].

Hata oranlarının azlığının yanı sıra robotlar, normal çalışanlardan çok daha hızlı çalışmaktadırlar. Süreçlerin hızlandırılması müşteri memnuniyetini artırır ve bu durum pazarda rekabet avantajı olarak kullanılabilir. Robot kullanımıyla elde edilen zaman tasarrufu, maliyet tasarrufuna da olumlu etki yapmaktadır. Artan hız, daha iyi tepki sürelerine sebep olur ve daha büyük hacimde görevlerin yerine getirilmesini sağlar. Robotların çalıştığı uygulamaların tepki süreleri, robotların hızına göre oldukça yavaş olabilmektedir. Bu gibi durumlarda robotların yürütme hızı, robotun birlikte çalıştığı uygulamanın hızına ve gecikmesine uyum sağlayacak şekilde azaltılmalıdır [22][35][36][39]. Robot kullanımıyla iş hacmi ve iş hızı arttıkça verimlilik de artmaktadır. Kullanılabilirlik ve verimlilik arttıkça şirketlerde operasyonel maliyetlerin ciddi oranda azaldıkları gözlemlenmektedir [22][39]. İş yüküne bağlı olarak kullanılan robot sayısının kolaylıkla azaltılabilir veya artırılabilir olması sayesinde kurumlarda esneklik de sağlamış olur. Kapasite darboğazları oluşturan yoğunluklar, robot sayısı bir süreliğine artırılarak aşılabılır [35][39]. RSO, mevcut BT sistemine değişiklik yapılmadan entegre edilebildiği için robotların analiz, planlama, lisanslama, entegrasyon ve test maliyetleri düşüktür. Ayrıca RSO yazılımındaki önceden tanımlanmış bileşenler ve kodlar birçok başka işlem için de kullanılabilir. Bu, gelecekte ihtiyaç duyulan robotların geliştirme sürelerini kısaltır ve maliyetlerini azaltır [31][35].

Yukarda belirtildiği üzere, robotların sisteme girişleri dahil her adımları sistemde kayıt altına alınmaktadır. İş akışları ve süreçler tam olarak belgelenir ve

istenildiği zaman kontrol edilebilir. Robotlar proselere uygun hareket ederek tüm görevleri yerine getirmektedir [35][39]. Süreçlerin ve iş akışlarının belgelenmesi, bir faaliyetin ne kadar sürede tamamlandığını veya ne zaman durdurulduğunu bilmek açısından son derece önemlidir. Gerçekleştirilen görevlerden toplanan verilere dayanarak, görevlerin zamanında tamamlanabilmesine ilişkin tahminler yapılabilir. Bu veriler analiz edildikten sonra süreçlerin iyileştirmeye açık alanları belirlenerek optimize edilir. Bu da süreçlerin yürütülmesindeki verimliliği artırır [36][37][39].

RSO kullanımıyla maliyetlerin azalması sayesinde RSO yatırımı hızlı bir şekilde geri döner [35][43]. İş denetim, risk ve uyumluluk, teknoloji, iş süreçleri, veri analizi ve finans gibi farklı alanlarda danışmanlık hizmetleri sunan Protiviti, 2019 yılında değişik bölge ve sektörlerde faaliyet gösteren, farklı büyüklükteki 450 şirketle bir anket gerçekleştirmiştir [44]. Bu anketin Tablo 1'deki sonuçlarına bakıldığında RSO açısından tüm sektörler arasında bir fikir birliğinin olduğu görülmektedir. Tüm sektörler için ortalama olarak RSO'nun getirdiği en önemli fayda verimliliğin artması, en az önemli fayda ise maliyetlerin azaltılması gibi görülmektedir. Avantajlar açısından sektörler arasında çok büyük fark olduğu değerlendirilmemektedir.

Tablo 1 - Farklı sektörlerde göre RSO kullanmanın en önemli avantajları [44]

(The most important advantages of using RPA according to different sectors)

|                                   | Finans | Teknoloji<br>Medya<br>Telekom | Sağ<br>lık | Enerji<br>Kamu | Üretim<br>Lojistik | Perak<br>ende |
|-----------------------------------|--------|-------------------------------|------------|----------------|--------------------|---------------|
| Artırılmış Üretkenlik             | 19%    | 19%                           | 22%        | 24%            | 23%                | 23%           |
| Daha İyi Kalite                   | 11%    | 21%                           | 16%        | 13%            | 18%                | 15%           |
| Daha Güçlü Rekabetçi Pazar Konumu | 18%    | 15%                           | 13%        | 16%            | 14%                | 15%           |
| Yüksek Müşteri Tatmini            | 12%    | 12%                           | 14%        | 10%            | 10%                | 12%           |
| Daha Yüksek Hız                   | 8%     | 10%                           | 9%         | 11%            | 14%                | 10%           |
| Artan Çalışan Tatmini             | 11%    | 5%                            | 6%         | 5%             | 8%                 | 8%            |
| Gelişmiş Uyumluluk                | 6%     | 6%                            | 5%         | 6%             | 4%                 | 6%            |
| Daha Az Hata                      | 6%     | 5%                            | 6%         | 6%             | 5%                 | 4%            |
| Yüksek Gelir Yaratma              | 5%     | 4%                            | 6%         | 5%             | 3%                 | 4%            |
| Azaltılmış Maliyetler             | 4%     | 3%                            | 3%         | 4%             | 1%                 | 3%            |

Everything Everywhere'den sonra Birleşik Krallık'taki en büyük ikinci mobil telekomünikasyon sağlayıcısı olan Telefónica O2, RSO projesinin bir parçası olarak 15 temel sürecini otomatize etmiştir. Kurulan 160 civarı robot sayesinde ayda 400.000 ila 500.000 işlem gerçekleştirilmektedir. Bunun sonucunda yüzlerce FTE kâra geçilmiş ve RSO

yatırım maliyetleri 12 ay içinde amorti edilmiştir. Şirketin üç yıllık yatırım getirisi %650 ile %800 arasında hesaplanmaktadır [45][46]. Diğer yandan Xchange isimli, birçok sektördeki müşterisine teknoloji ve satın alma hizmetleri sağlayan uluslararası bir şirkette yapılan RSO projesi sonucunda 14 temel süreç otomatikleştirilmiştir. Bunun sonucunda 27 robotla ayda 120.000 vaka işlenmekte ve süreç başına yüzde 30 oranında maliyet tasarrufu sağlanmaktadır [46][47].

Utility, ev ve işyerleri için elektrik ve gaz ile ilgili hizmetleri sağlayan Avrupa'nın en büyük enerji tedarikçilerinden biridir. Firma aynı zamanda kömür, gaz ve petrol ile çalışan enerji santrallerini işletmekte ve yönetmektedir. Şirketin RSO projesi sayesinde arka ofis süreçlerinin yüzde 35'i otomatikleştirilmiştir. Ayda yapılan 1 milyon işlem sayesinde RSO yatırım maliyetleri 12 ayda amorti edilmiş ve şirketin bir yıllık yatırım getiri oranı %200 olarak hesaplanmıştır. Bu süreçte iki kişi, 600 çalışanın görevlerini yerine getiren 300 robotu yönetmiştir. RSO'nun uygulanmasıyla birlikte süreçler paralel çalışmaya başlamış, hizmet kalitesi artmış ve hata oranı azalmıştır. Ayrıca süreçlerin ölçeklenebilirlikleri artmış, işlem süreleri de hızlanmıştır. RSO kullanımı sayesinde çalışanlar stratejik görevlere yönlendirilebilmiş ve FTE tasarrufu sağlanmıştır [46][48].

Siderska çalışmasında, CAWI (Bilgisayar Destekli Web Görüşmesi) anketini 110 Polonyalı hizmet şirketinin temsilcilerine uygulamıştır. Bu anket aracılığıyla, Covid-19 salgını sırasında RSO'nun benimsenmesini sağlayan faktörleri öğrenmeyi, RSO teknolojisinin algılanan kullanılabilirlik, kullanım kolaylığı, güvenlik ve işlevsellik açısından konumunu belirlemeyi ve RSO'nun benimsenmesini engelleyen faktörleri ortaya çıkarmayı amaçlamıştır. Bu araştırmanın anketi yalnızca bankacılık, finans ve sigorta, e-ticaret ve kamu hizmetleri sektörlerinde faaliyet gösteren Polonyalı şirketlere uygulanmıştır ve katılımcılardan RSO'nun işlevselliğini, kullanılabilirliğini ve kullanılabilirliğini derecelendirmeleri istenmiştir. Çıkan sonuca göre RSO'nun kullanılabilirliği, kullanım kolaylığı ve işlevselliğine göre daha yüksek ortalama değerlere sahip olarak çıkmıştır [49].

Şirketler, RSO'nun faydalarının yanı sıra RSO ile ilgili bazı zorluklarla da karşılaşabilmektedir. RSO teknolojisi kural tabanlı standartlaşmış işlemleri otomatikleştirmek için kullanılabilir ancak karar verme yeteneğine sahip değildir. Bu nedenle süreçlerde bazı kararların alınması gerektiğinde, çalışanların süreçlere müdahale etmesi gerekmektedir. RSO'nun uygulandığı süreçler standartlaştırılıp geliştirilmezlerse, RSO beklendiği

gibi çalışmayacak ve süreçler doğru şekilde yürütülmeyecektir. Piyasada bulunan RSO ürünlerinin çeşitliliği nedeniyle birçok şirketin farklı RSO yazılımlarının yeteneklerini ve özelliklerini değerlendirmesi zordur. Firmaların yanlış ürünü seçmesi durumunda, RSO'nun en önemli faydalarından biri olan maliyet tasarrufu sağlanamayabilir ve hatta tam tersine maliyetler artabilir. Bu nedenle doğru RSO ürününü ve tedarikçisini seçmek son derece önemlidir. Diğer yandan RSO kullanıldığında insan kaynaklarından tasarruf yapılabildiğinden çalışanlar, robotların işlerini devralmasından endişe etmektedirler. RSO'ların rutin işleri yüklenmeleriyle birlikte çalışanların eleştirel sorgulama veya yaratıcı yaklaşım geliştirme gibi yeni becerilere sahip olmaları gerekecektir. Bu noktada şirketler çalışanlarına yönelmeleri gereken yeni alanlar konusunda yol göstermelidirler [31][35][47].

RSO teknolojisinin, mevcut başvuru sistemlerine kayıt yapılması, form doldurulması, KKP sistemlerine giriş yapıp istenilen fonksiyonların yürütülmesi, sistem API'lerine bağlantı yapılması, yapılandırılmış verilerin ayıklanması, e-postaların ve eklerinin açılıp işlenmesi, rapor oluşturulması, internette veri kazınması, farklı kaynaklardan yapılandırılmış verilerin alınması, birleştirilmesi ve değerlendirilmesi, farklı BT sistemlerinden gelen istatistiksel verilerin bir araya getirilmesi, dosya ve klasörlerin depolanması, adlandırılması ve taşınması, uygulanan süreç kuralları çerçevesinde eğer-o halde kararlarının yürütülmesi, sosyal medya istatistiklerinin toplanması, veritabanlarının okunması ve yazılması gibi görevleri gerçekleştirmekte yaygın olarak kullanıldığı söylenebilir [34][36][50]. Diğer yandan RSO teknolojisi otomatik raporlama, fatura ve makbuzların kontrol edilmesi ve sisteme girişlerinin yapılması, müşteri hizmetleri kapsamında kural bazlı sorgulama yapılması, öğrenilen bilgilerin müşteriye iletilmesi, sahtekarlığın tespiti (örn.: kartla beş dakikada kaç kez ödeme yapılırsa dolandırıcılık ihtimali göz önünde bulundurulabilir), hesapların açılıp kapatılmasında, kredi ve kredi kartlarının hazırlanmasında da kullanılabilir [51][52]. Dijital verinin olmadığı, elle giriş yapılması gereken durumlarda OCR (Optik Karakter Tanıma) gibi teknolojilerden de faydalanılabilir. Bu sayede kâğıt faturalardaki veriler okunmakta ve dijital verilere dönüştürülmektedir. Veriler dönüştürüldükten sonra ilgili faturalar, robotlar tarafından sisteme girilebilir [50]. Diğer yandan yoğun emek ve zaman gerektiren geleneksel denetim prosedürleri için RSO'nun kullanılması, maliyetlerin, insan hatalarının ve işlem sürelerinin azaltılmasına, bunun sonucunda da denetim performansının artmasına yardımcı olur. RSO kullanımı yoluyla çalışanların iş yükünün

azaltılması, çalışanların mesleki muhakeme gerektiren prosedürlere (dolandırıcılık riskleri hakkında beyin fırtınası yapma, analitik prosedürlere ilişkin istisnaları analiz etme vb.) odaklanmasına olanak tanır [54].

RSO yazılımı, birçok sektörde, farklı departmanlarda kullanılmaktadır. 2018 yılında Samsung Electronic Vietnam'ın yöneticileri, departmanlarındaki durumu raporlar aracılığıyla takip edebilsinler diye bazı üretim, satın alma ve insan kaynakları süreçlerinde RSO hayata geçirilmiştir. Bu raporlar, çalışanlar tarafından oluşturulurken, proje ile birlikte robotlar tarafından oluşturulmaya başlanmıştır. Bu sistemde, satın alma departmanı, MES (Üretim Yönetimi Sistemi) sistemini kullanarak ürünlerin tahmini üretim tarihlerini ve üretimde kullanılacak malzeme gereksinimlerini belirlerler. Alıcılar gerekli malzeme bilgilerini Küresel Tedarik Zinciri Yönetimi (KTZY) sistemine girdikten sonra tedarikçiler gereksinimleri kontrol etmek ve hazırlamak için bu sisteme erişebilmektedir. Gereksinimlerde bir değişiklik olması durumunda alıcı malzeme verilerini güncellemeli ve daha önceki adımları tekrar etmelidir. Şirket bu satın alma sürecini RSO'ya uygularken günlük olarak büyük hacimli ve tekrarlayan görevleri olan insan kaynakları departmanı süreçlerinin de robotlara aktarılabilceği tespit edilmiştir [55]. Bir Pazar araştırması firması olan Dimensional Research'ün 2020 yılında yaptığı ankete göre şirketlerdeki RSO uygulamalarının yüzde 21'i BT departmanlarının, yüzde 79'u ise diğer departmanların süreçleri için kullanılmaktadır [56].

Şirketler COVID-19 salgını sırasında, uzaktan çalışma sistemine geçtikçe, iş akışlarını otomatikleştirme ihtiyacı artmış ve bu da RSO pazarının büyümesini hızlandırmıştır. Küresel RSO pazar büyüklüğünün değeri 2022 yılında 3,70 milyar ABD doları iken, satışların 2032 yılında 81,80 milyar ABD doları olacağı tahmin edilmektedir [57]. Pazarın hızla büyümesi ve şirketlerin günlük RSO desteğine duydukları ihtiyacın artması, pazardaki RSO ürün sağlayıcılarının sayısını artırmış ve bu da sağlayıcıların arasındaki rekabeti kızıştırmıştır. Piyasada bulunan RSO ürünlerinden bazıları, UiPath, Automation Anywhere, Blue Prism, Microsoft Power Automate, SAP iRPA (SAP Intelligent RPA), Nice, Kofax, WorkFusion, Kryon, Pegasystems, EdgeVerve ve IBM RPA şeklinde sıralanabilir.

### 3. YÖNTEM (METHODOLOGY)

Araştırma kapsamında bir anket uygulanarak çalışanların RSO teknolojisine yaklaşımları analiz edilmiştir. Anketi cevaplayacak kişiler, farklı sektörlerdeki RSO projelerinde çeşitli rollerde görev

yapmakta olan veya daha önce bu tür projelerde çalışmış olan kişilerden seçilmiştir. Araştırmacının en kolay ulaşabileceği bireyleri ve araştırmaya katılmaya istekli olan bireyleri seçmesi olarak tanımlanan, kolayda örnekleme metoduyla yapılan bu işlemde, LinkedIn platformu da kullanılarak, RSO ile ilişkili olarak çalışan, çoğunluğu analist ve geliştirici olan kişilerden anketi doldurmaları istenmiştir [62]. Anketler yaygın olarak kullanılan ve katılımcının özel bir yazılım yüklemesini gerektirmeyen Google Forms uygulaması kullanılarak Mayıs – Haziran 2022 tarihlerinde yapılmış, toplanan veriler SPSS Versiyon 28 programında analiz edilmiştir.

Hazırlanan anket üç bölümden oluşmaktadır. Birinci bölümde katılımcının demografik bilgisi ve işiyle ilgili bilgiler altı soruyla alınmaktadır. Anketin ikinci bölümünde RSO kullanımı ve deneyimi ile ilgili 6 ölçek yer almaktadır. Bu ölçekler değişik yayınlardan alınmıştır ve 5'li Likert ölçeği ile cevaplanmaları beklenmektedir. [34][44][49][50][58][59]. Anketin son bölümünde ise kullanıcının RSO'ları değerlendirmesi için değişik kaynaklardan alınan çoktan seçmeli veya çoklu seçmeli sorular yer almaktadır [56][58]. Anketler yurt dışındaki kullanıcılara İngilizce, yurt içindekilere de Türkçe olarak sunulmuştur. Katılımcıların tüm soruları yanıtlaması zorunlu tutulmuştur. Araştırma, Marmara Üniversitesi Sosyal Bilimler Araştırma Etik Kurulu'nun 2022-106 no.lu kararıyla etik yönden uygun bulunmuştur.

#### 4. BULGULAR (FINDINGS)

191'i Türkiye'de, 47'si yurt dışında yaşayan 238 katılımcıyla anket gerçekleştirilmiştir. Türkiye'de yaşayan bir çalışanın verisi, kişinin RSO teknolojisi ile çalışmaması nedeniyle analize dahil edilmemiştir. Diğer yandan box plot grafiğine bakılarak, çeyrekler arası açıklığın 1,5 katı üst ve alt çeyreklikten uzakta kalan 16 aykırı değer içeren satır silinmiştir. Böylece veri sayısı 221'e düşmüştür.

Toplanan veriler incelendiğinde katılımcıların yüzde 81'inin Türkiye'de, yüzde 19'unun ise yurt dışında ikamet ettikleri anlaşılmıştır. Diğer yandan katılımcıların %67,4'ü erkek, %32,6'sı kadındır. Ankete katılanların %44,3'ü teknoloji, bilgi teknolojileri ve yazılım danışmanlığı dallarında çalışırken, ikinci sırada finans ve bankacılık sektörü %15,8 ile, üçüncü sırada perakende ve satış alanı %8,2 ile, dördüncü sırada ise sigorta şirketleri %5,9 ile gelmektedir. Kişilerin RSO ile ilgili sorumlulukları, RSO geliştiricisi (%68,3), süreç sahibi veya anahtar kullanıcı (%14,5), RSO ürün sahibi veya ürün yöneticisi (%14,5) şeklindedir. Katılımcıların RSO

deneyimleri de ortalama olarak 6,45 yıl olarak ölçülmüştür.

Katılımcılar tarafından RSO ürünü olarak en fazla UiPath ürününün kullanıldığı dikkat çekmektedir. En çok kullanılan üründen en az kullanılan ürüne doğru doğru bir sıralama yapıldığında, sıralamanın UiPath (%83,3), Automation Anywhere (%14,5), Blue Prism (%13,1), Microsoft (%10,9) şeklinde olduğu görülmüştür. Kullanılan robot sayılarına bakıldığında her katılımcının ortalama 13,15 robot kullandığı görülmüştür. Katılımcılar RSO kullanmanın getirdiği avantajları, verimlilik ve kalite artışı (%94,1), daha yüksek hız (%90,5), daha az hata (%86), çalışma memnuniyetinin artması (%78,3), pazarda rekabet gücünün artması (%54,8) şeklinde sıralamışlardır. Araştırmada kullanılan altı ölçek için Cronbach Alpha güvenirlik analizi yapılmıştır. Cronbach Alpha analizinden elde edilen güvenirlik katsayısı 0,70 ve üzerinde ise ölçek güvenilir kabul edilmektedir [60]. Ölçekler ve bu ölçeklere ilişkin Cronbach Alpha güvenirlik analizi katsayıları Tablo 2'de verilmiştir.

Tablo 2 - Cronbach Alpha Analiz Sonuçları  
(Cronbach Alpha Analysis Results)

| Ölçek                                                           | Madde Sayısı | Cronbach Alpha |
|-----------------------------------------------------------------|--------------|----------------|
| RSO teknolojisini benimsemenin temel nedenleri                  | 7            | 0,752          |
| RSO teknolojisinin avantajları                                  | 4            | 0,757          |
| RSO teknolojisinin kullanım kolaylığı                           | 5            | 0,768          |
| RSO teknolojisinin kullanılabilirliği                           | 6            | 0,724          |
| RSO kullanımını popüler kılan faktörler                         | 5            | 0,566          |
| RSO projeleri için proje başarı kriterlerinin değerlendirilmesi | 3            | 0,712          |

Ölçekler için hesaplanan Cronbach Alpha güvenirlik katsayıları dikkate alındığında, beşinci ölçek olan "RSO kullanımını popüler kılan faktörler" dışındaki tüm ölçekler için hesaplanan değerlerin 0,70'in üzerinde olduğu tespit edilmiştir. Bu sebeple "RSO kullanımını popüler kılan faktörler" ölçeği çıkartılarak geriye kalan ölçeklerle analize devam edilmiştir. Bu aşamada t-testi vasıtasıyla, cinsiyet, yaşanan yer ve kullanılan RSO yazılımına göre, RSO benimsemek için sebepler, RSO'nun getirdiği faydalar, RSO'nun işlevselliği ve kullanım kolaylığı incelenmiştir. Bu hipotezlerden ispatlananlar Tablo 3'de verilmiştir.

Tablo 3 - t-test Sonuçları  
(t-test Results)

| Konu                                           | p     | Ortalamalar                     |
|------------------------------------------------|-------|---------------------------------|
| RSO'yu kullanışlı bulma                        | 0,018 | Kadın: 4,55<br>Erkek: 4,39      |
| RSO teknolojisini benimsemenin temel nedenleri | 0,015 | Yurtdışı: 4,25<br>Yurtiçi: 4,43 |
| RSO'nun desteklediği süreçlerin karmaşıklığı   | 0,000 | Yurtdışı: 2,42<br>Yurtiçi: 3,05 |

|                                                              |       |                                                         |
|--------------------------------------------------------------|-------|---------------------------------------------------------|
| RSO projelerini başarılı bulma                               | 0,000 | Yurtdışı: 4,45<br>Yurtiçi: 4,27                         |
| RSO teknolojisini benimsemenin temel nedenleri               | 0,039 | UIPath kullananlar: 4,42<br>UIPath kullanmayanlar: 4,26 |
| RSO'yu kullanışlı bulma                                      | 0,011 | UIPath kullananlar: 4,48<br>UIPath kullanmayanlar: 4,25 |
| RSO'yu kullanıcı dostu bulma                                 | 0,000 | UIPath kullananlar: 3,82<br>UIPath kullanmayanlar: 3,25 |
| RSO'nun proje başarı kriterlerine ilişkin derecelendirmeleri | 0,034 | UIPath kullananlar: 4,24<br>UIPath kullanmayanlar: 4,02 |

Bir sonraki aşamada ölçekler arasında Pearson korelasyon analizi kullanılarak ilişki var olup olmadığı, varsa bu ilişkilerin yönleri ve güçlü yönlerinin neler olduğu belirlenmeye çalışılmıştır. Değişkenlerin birbirleriyle olan ilişki derecelerini ölçen korelasyon katsayılarının değer aralıkları, 0,00-0,19 ise çok zayıf, 0,20 - 0,39 ise zayıf, 0,40 - 0,59 ise orta, 0,60 - 0,79 ise güçlü (yüksek) ve 0,80 - 1,00 ise çok güçlü (yüksek) şeklinde yorumlanmaktadır [61]. Değişkenler arasında güçlü ilişkinin korelasyon katsayılarına bakıldığında sadece “RSO Teknolojisini Benimsemenin Temel Nedenleri” ile “RSO Teknolojisinin Avantajları” değişkenleri arasında pozitif yönlü olarak olduğu görülmektedir. Diğerler ilişkilerin genelde orta dereceli oldukları dikkat çekmektedir. Korelasyon değerleri Tablo 4’de verilmiştir.

Tablo 4 – Pearson Korelasyon Analizi Sonuçları  
(Correlation Analysis Results)

|                            | RSO benimsenme nedenleri | RSO avantajları | RSO kullanım kolaylığı | RSO kullanışlılığı | RSO için başarı kriterleri |
|----------------------------|--------------------------|-----------------|------------------------|--------------------|----------------------------|
| RSO benimsenme nedenleri   | 1                        | 0,649           | 0,338                  | 0,412              | 0,496                      |
| RSO avantajları            | 0,649                    | 1               | 0,327                  | 0,365              | 0,467                      |
| RSO kullanım kolaylığı     | 0,338                    | 0,327           | 1                      | 0,515              | 0,460                      |
| RSO kullanışlılığı         | 0,412                    | 0,365           | 0,515                  | 1                  | 0,417                      |
| RSO için başarı kriterleri | 0,496                    | 0,467           | 0,460                  | 0,417              | 1                          |

Son olarak korelasyon katsayısı orta ve üzerinde (0,40 – 1,00) olan ölçek çiftleri için doğrusal regresyon analizleri yapılmıştır. Regresyon analizi, bir metrik bağımlı değişken ile bir veya daha fazla bağımsız değişken arasındaki ilişkiyi incelemek için kullanılan bir analiz yöntemidir [62]. Regresyon analizi için

gerekli olan, normal dağılım varsayımı, aykırı değer olmaması, otokorelasyonun olmaması ve varyansın homojenliği varsayımları [61] ölçekler tarafından sağlandığından regresyon analizine geçilmiştir.

Şekil 1’de görülen Doğrusal regresyon analizine göre “RSO teknolojisini benimsemenin temel nedenleri”nin % 42,1 oranında “RSO teknolojisinin avantajları” tarafından, “RSO teknolojisinin işlevselliği”nin ise % 26,6 oranında “RSO teknolojisinin kullanıcı dostu olması” tarafından, “RSO projeleri için proje başarı kriterlerinin değerlendirilmesi”nin % 24,6 oranında “RSO teknolojisinin benimsenmesinin temel nedenleri” tarafından, “RSO teknolojisinin avantajları”nın % 21,8 oranında “RSO projeleri için proje başarı kriterlerinin değerlendirilmesi” tarafından, “RSO teknolojisinin kullanılabilirliği”nin % 21,1 oranında “RSO projeleri için proje başarı kriterlerinin değerlendirilmesi” tarafından, “RSO teknolojisinin işlevselliği”nin %17,4 “RSO projeleri için proje başarı kriterlerinin değerlendirilmesi” tarafından ve “RSO Teknolojisinin İşlevselliği”nin % 17,0 oranında “RSO Teknolojisini Kullanmanın Temel Nedenleri” tarafından açıklandığı görülmüştür.

| Model Summary <sup>a</sup> |                   |          |                   |                            |               |
|----------------------------|-------------------|----------|-------------------|----------------------------|---------------|
| Model                      | R                 | R Square | Adjusted R Square | Std. Error of the Estimate | Durbin-Watson |
| 1                          | ,649 <sup>a</sup> | ,421     | ,418              | ,33071                     | 1,948         |

a. Predictors: (Constant), VortailederRPATechnologie\_Score

b. Dependent Variable: GründefürdieRPAAnnahme\_Score

| ANOVA <sup>a</sup> |                |     |             |         |                   |
|--------------------|----------------|-----|-------------|---------|-------------------|
| Model              | Sum of Squares | df  | Mean Square | F       | Sig.              |
| 1 Regression       | 17,408         | 1   | 17,408      | 159,170 | ,000 <sup>b</sup> |
| Residual           | 23,952         | 219 | ,109        |         |                   |
| Total              | 41,360         | 220 |             |         |                   |

a. Dependent Variable: GründefürdieRPAAnnahme\_Score

b. Predictors: (Constant), VortailederRPATechnologie\_Score

| Coefficients <sup>a</sup>       |                             |            |                           |        |      |
|---------------------------------|-----------------------------|------------|---------------------------|--------|------|
| Model                           | Unstandardized Coefficients |            | Standardized Coefficients | t      | Sig. |
|                                 | B                           | Std. Error | Beta                      |        |      |
| 1 (Constant)                    | 1,951                       | ,195       |                           | 9,986  | ,000 |
| VortailederRPATechnologie_Score | ,551                        | ,044       | ,649                      | 12,616 | ,000 |

a. Dependent Variable: GründefürdieRPAAnnahme\_Score

Şekil 1 - Doğrusal Regresyon Analizi Sonuçları

## 5. SONUÇ (CONCLUSION)

Bu çalışmanın amacı, çoğunluğu Türkiye’de yaşayan 221 katılımcının katıldığı bir anket aracılığıyla RSO teknolojisini kullanan veya kullanmış olan çalışanların bakış açısından RSO teknolojisini değerlendirmektir. Yapılan analizler sonucunda RSO’nun yararlılığını değerlendirmede kadın ve erkeklerin farklılaştığı, kadınların RSO’yu erkeklere göre daha kullanışlı buldukları tespit edilmiştir.

Türkiye’de yaşayan katılımcılar yurt dışında yaşayan katılımcılara göre RSO’yu benimseme nedenleri

konusunda daha olumlu görüş belirtmiş ve RSO'yu daha kolay benimsediklerini belirtmişlerdir. Bu arada Türkiye'de yaşayan katılımcılar, RSO'yu daha kolay benimseseler de yurt dışında yaşayan katılımcılara göre RSO projelerini daha az başarılı bulmuşlardır. Öte yandan yurt dışında yaşayan katılımcılar çalıştıkları firmalarda kullanılan yazılım robotlarının desteklediği süreçleri karmaşık olarak değerlendiren, Türkiye'de yaşayan katılımcılar süreçlerin karmaşıklık derecesi konusunda kararsız kalmışlardır.

Anket katılımcılarının kullandıkları RSO ürünleri incelendiğinde, katılımcıların çoğunlukla yıllık RSO tedarikçi değerlendirme raporlarında oldukça iyi puan alan UiPath, Automation Anywhere, Blue Prism ve Microsoft gibi popüler RSO ürünlerini kullandıkları tespit edilmiştir. Katılımcıların en sık kullandığı ürün UiPath olduğundan RSO ürün bazlı analizler için katılımcılar UiPath kullananlar ve kullanmayanlar olarak iki gruba ayrılmıştır. Analiz sonuçlarına göre UiPath kullanan katılımcılar, RSO'nun benimsenme nedenlerini UiPath kullanmayan katılımcılara göre daha olumlu değerlendirmişler ve RSO'yu daha kolay kabul etmişlerdir. Ayrıca UiPath kullanan katılımcılar ile UiPath kullanmayan katılımcıların RSO'nun kullanışlılığı ve RSO kullanım kolaylığına ilişkin değerlendirmeleri istatistiksel olarak farklılık göstermektedir. UiPath kullanan katılımcılar, UiPath kullanmayan katılımcılara göre RSO'yu daha kullanışlı ve kullanıcı dostu bulmaktadır. Ayrıca UiPath kullanan katılımcılar, RSO projelerine yönelik proje başarı kriterlerini UiPath kullanmayan katılımcılara göre daha iyi değerlendirmişlerdir. Bu sonuçlardan yola çıkarak, UiPath kullananların yaşadıkları RSO deneyimini daha olumlu olarak nitelendirdikleri söylenebilir.

Yapılan regresyon analizine göre “RSO teknolojisini benimsemenin temel nedenleri” değişkeni %42,1 oranında yerini “RSO teknolojisinin avantajları” değişkeni tarafından açıklanmaktadır. RSO teknolojisinin avantajlarının farkında olarak bu avantajları kullanmak amaçlı RSO kullananların bu teknolojiyi benimsemeleri doğal görünmektedir. Başka bir şekilde ifade etmek gerekirse RSO teknolojisinin, oluşturduğu beklentileri karşılamakta olduğu görülmektedir. Öte yandan “RSO teknolojisinin benimsenmesinin temel nedenleri” değişkeninin “RSO projeleri için proje başarı kriterlerinin değerlendirilmesi” değişkeni tarafından %24,6 oranında açıklanması RSO teknolojisinin benimsenmesinin bir diğer önemli sebebinin başarı kriterleri değerlendirmesine uyması olarak açıklanabilir. Diğer bir deyişle kullanıcılar başarı kriteri olarak gördükleri hedeflerin RSO çıktılarıyla uyumunu onaylamaktadır. Öte yandan “RSO

projeleri için proje başarı kriterlerinin değerlendirilmesi” değişkeni %21,8 oranında “RSO teknolojisinin avantajları” değişkeni tarafından açıklanması yukarıdaki sonuçlarla uyumlu görünmektedir. Bir başka regresyon analizine göre “RSO teknolojisinin kullanıcı dostu olması” değişkeni %26,6 oranında “RSO teknolojisinin işlevselliği” değişkeni tarafından açıklanmaktadır. RSO teknolojisinin işlevselliğinin, teknoloji kullanıcı dostu olmasına yol açması beklenen bir sonuçtur.

Diğer regresyon analizleri küçük oranda da olsa bazı ölçeklerin diğerleri aracılığıyla açıklanabildiğini göstermiştir. Çıkan sonuçlara göre, Blagoev'in 2021 yılında muhasebe ve finans alanında yaptığı çalışmanın sonuçlarına uygun olarak “RSO projeleri için proje başarı kriterlerinin değerlendirilmesi” değişkeni %21,1 oranında “RSO teknolojisinin kullanım kolaylığı” değişkeni tarafından açıklanmaktadır [63]. Diğer yandan “RSO projeleri için proje başarı kriterlerinin değerlendirilmesi” değişkeni %17,4 oranında “RSO teknolojisinin işlevselliği” değişkeni tarafından açıklanırken Blagoev'in çalışmasına ters bir sonuç vermiştir. Bu çalışmadaki katılımcıların, RSO projesinin kullanım kolaylığının proje memnuniyetine katkıda bulunduğunu ifade ederken Blagoev'in çalışmasının tersine işlevselliğin de memnuniyete olumlu etki yaptığının altını çizdiklerinin belirtmek gerekir. Öte yandan “RSO teknolojisini benimsemenin temel nedenleri” değişkeni “RSO teknolojisinin işlevselliği” değişkeni tarafından %17 oranında açıklanmaktadır.

Dey ve Das'ın 2019'da yaptıkları çalışmaya katılanların RSO teknolojisinin kullanışlılığını, RSO teknolojisinin işlevselliği ve kullanılabilirliğinden daha olumlu değerlendirdikleri görülmüştür [59]. Yapılan çalışmanın sonuçlarına göre RSO'nun kullanışlılığı, Siderska'nın bahsedilen ve Dey ve Das'ın çalışmalarıyla örtüşmekte olup işlevsellik ve kullanılabilirliğe göre daha yüksek bir ortalamaya sahiptir öte yandan Siderska'nın çalışmasından farklı olarak RSO'nun işlevselliği, RSO'nun kullanılabilirliğine göre daha yüksek bir ortalamaya sahiptir [49].

Ankete katılanların RSO olarak çoğunlukla, UiPath, Automation Anywhere, Blue Prism ve Microsoft gibi çok kullanılan ürünleri kullandıkları tespit edilmiştir. Kullanıcıların en sık kullandığı ürün olan UiPath'i kullananların RSO'yu daha kullanışlı ve kullanıcı dostu bulmanın yanı sıra daha kolay benimsedikleri görülmüştür. Bunda UiPath'in dünyada en çok kullanılan ürünlerden biri olmasından dolayı pratik olmasının etkili olduğu düşünülmektedir. Konuyla ilgili bir çalışma olmadığından ileride UiPath'in diğer RSO yazılımlarıyla karşılaştırıldığı bir araştırmanın

iyi olacağı düşünülmektedir. Diğer yandan UiPath kullanıcısının bu kadar fazla olmasının ve anketin katılımcılarının genel olarak BT ve diğer sektörlerde RSO alanında geliştirici ve/veya iş analisti olarak çalışmaları, çalışmanın sınırlılığı olarak görülebilir. Bunun önüne geçebilmek için ileride yapılacak çalışmalarda farklı RSO ürünleri kullanan ve farklı rollerde çalışan katılımcılara (anahtar kullanıcılar, süreç sahipleri vb.) daha fazla odaklanılmaya çalışılabilir.

RSO alanında gelecekte yapılacak çalışmalarda RSO teknolojisinin, yapay zekâ ürünleri (doğal dil işleme, chatbotlar, makine öğrenmesi vb.) ile birlikte kullanımı araştırılabilir. RSO teknolojisinin NLP, OCR, nesne tanıma ve duygu analizi gibi yapay zekanın çeşitli disiplinleriyle birlikte kullanılmasının RSO projelerine de büyük katkıda bulunacağı öngörülmektedir.

#### KAYNAKLAR (REFERENCES)

- [1] Little, R. J., & Rubin, D. B. (2019). Statistical analysis with missing data (Vol. 793). John Wiley & Sons.
- [2] Black, W. C., Babin, B. J., & Anderson, R. E. (2010). Multivariate data analysis: A global perspective. Pearson.
- [3] Pamuk, N. S., Soysal, M. (2018), “YENİ SANAYİ DEVRİMİ ENDÜSTRİ 4.0 ÜZERİNE BİR İNCELEME”, Verimlilik Dergisi, 1: 41–66.
- [4] Taulli, T. (2020), The Robotic Process Automation Handbook: A Guide to Implementing RPA Systems, (1.Baskı), New York: Apress.
- [5] SAP Blogs (2012), “Introduction to SAP Gui Scripting”. <https://community.sap.com/t5/additional-blogs-by-members/introduction-to-sap-gui-scripting/ba-p/12992466> adresinden alındı.
- [6] Ostlick, N. (2016, 26 Temmuz), “The Evolution of Robotic Process Automation (RPA): Past, Present, and Future”. <https://www.uipath.com/blog/rpa/the-evolution-of-rpa-past-present-and-future> adresinden alındı.
- [7] Ostlick, N. (2017, 15 Mart), “Everything You Need to Know About Screen Scraping Software”. <https://www.uipath.com/blog/rpa/screen-scraping-software-everything-you-need-to-know> adresinden alındı.
- [8] Doğuç, Ö. (2020), “Robot Process Automation (RPA) and Its Future”, Hacıoğlu, Ü. (Ed.), Handbook of Research on Strategic Fit and Design in Business Ecosystems, Pensilvanya: IGI Global: 469- 492. <https://doi.org/10.4018/978-1-7998-1125-1.ch021>
- [9] Çalışkan, L. S., Kıran, S. (2020), “İş Süreçlerinin Otomasyonunda RPA’nın Faydaları”, Yönetim Bilişim Sistemleri Dergisi, 6/1: 1- 13. <https://dergipark.org.tr/tr/pub/ybs/issue/54333/650397>
- [10] Ezech, M. O., Ogbu, A. D., & Heavens, A. (2023). The Role of Business Process Analysis and Re-engineering in Enhancing Energy Sector Efficiency.
- [11] Sebetci, O., Gunay, M. B., Sebetci, E. (2018), “İş Süreç Yönetimi (Bpm) ve İş Akış Yönetimi (Wfm) Kavramlarına Yaklaşım”, Ajit-e Online Academic Journal of Information Technology, 9/33: 115–126. doi: 10.5824/1309-1581.2018.3.007.x
- [12] Strömberg, K. (2018), Robotic Process Automation of Office Work: Benefits, Challenges and Capability Development, Yüksek Lisans Tezi, Aalto, Aalto University School of Business.
- [13] Mangu, S. (2020), “BUSINESS PROCESS MANAGEMENT:ROBOTIC PROCESS AUTOMATION APPROACH”, International Journal of Advanced Research in Engineering and Technology, 11/11: 831- 840. doi: 10.34218/IJARET.11.11.2020.078
- [14] Jovanović, S. Z., Đurić, J. S., Šibalića, T. V. (2018), “ROBOTIC PROCESS AUTOMATION: OVERVIEW AND OPPORTUNITIES”, International Journal ”Advanced Quality”, 46/3-4: 34-39.
- [15] Aguirre, S., & Rodriguez, A. (2017). Automation of a business process using robotic process automation (RPA): A case study. In Applied Computer Sciences in Engineering: 4th Workshop on Engineering Applications, WEA 2017, Cartagena, Colombia, September 27-29, 2017, Proceedings 4 (pp. 65-71). Springer International Publishing.
- [16] Siderska, J. (2020), “Robotic Process Automation — a driver of digital transformation?”, Engineering Management in Production and Services, 12/2: 21–31. doi:10.2478/emj-2020-0009
- [17] Häuser, M., Schmid, A. (2018), “Robotic Process Automation (RPA)”, Computer und Recht, 34/4: 266–276. doi:10.9785/cr-2018-340412
- [18] Czarnecki, C., Fettke, P. (2021), “Robotic Process Automation: Positioning, Structuring, and

Framing the Work”, Czarnecki, C. (Ed.), Fettke, P. (Ed.), Robotic Process Automation: Management, Technology, Applications, Berlin, Boston: De Gruyter Oldenbourg: 3-24.  
<https://doi.org/10.1515/9783110676693-001>

[19] Pratt, M. K., Gillis, A. S. (2022), “workflow automation”.  
<https://www.techtarget.com/searchcontentmanagement/definition/workflow-automation#:~:text=Workflow%20automation%20is%20an%20approach,accordance%20with%20defined%20business%20rules> adresinden alındı.

[20] Syed, R., Suriadi, S., Adams, M., Bandara, W., Leemans, S. J., Ouyang, C., ter Hofstede, A. H., van de Weerd, I., Wynn, M. T., Reijers, H. A. (2020), “Robotic Process Automation: Contemporary themes and challenges”, Computers in Industry, 115/1: 103162. doi: 10.1016/j.compind.2019.103162

[21] Allweyer, T. (2016), “Robotic Process Automation- Neue Perspektiven für die Prozessautomatisierung “. <https://www.kurze-prozesse.de/blog/wp-content/uploads/2016/11/Neue-Perspektiven-durch-Robotic-Process-Automation.pdf>

[22] Smeets, M., Erhard, R., Kaußler, T. (2019), Robotic Process Automation (RPA) in der Finanzwirtschaft: Technologie – Implementierung – Erfolgsfaktoren für Entscheider und Anwender (1.Baskı), Wiesbaden: Springer Gabler.

[23] van der Aalst, W. M. P., Bichler, M., Heinzl, A. (2018), “Robotic Process Automation”, Business Information Systems Engineering, 60/4: 269–272. doi: 10.1007/s12599-018-0542-4

[24] Oltrogge, M., Derr, E., Stransky, C., Acar, Y., Fahl, S., Rossow, C., Pellegrino, G., Bugiel, S., Backes, M. (2018), “The Rise of the Citizen Developer: Assessing the Security Impact of Online App Generators”, 2018 IEEE Symposium on Security and Privacy (SP), 634-647. doi: 10.1109/sp.2018.00005

[25] Scheppeler, B., Weber, C. (2020), “Robotic Process Automation”, Informatik Spektrum, 43/2: 152–156. doi:10.1007/s00287-020-01263-6

[26] ÖZDEM, H., BORA, M. P. (2021), “Türkiye’de Robotik Süreç Otomasyonu”, Bilgisayar Bilimleri ve Teknolojileri Dergisi, 1/1: 1-9. doi: 10.54047/bibtcd.1008340

[27] Sabharwal, A. (2018), Introduction To RPA, Independently published.

[28] Langmann, C., Turi, D. (2020), Robotic Process Automation (RPA) - Digitalisierung und Automatisierung von Prozessen: Voraussetzungen, Funktionsweise und Implementierung am Beispiel des Controllings und Rechnungswesens, (1. Baskı), Wiesbaden: Springer Gabler.

[29] Mullakara, N., Asokan, A. K. (2020), Robotic Process Automation Projects: Build real-world RPA solutions using UiPath and Automation Anywhere, Birmingham: Packt Publishing.

[30] Choi, D., R’bigui, H., Cho, C. (2021), “Candidate Digital Tasks Selection Methodology for Automation with Robotic Process Automation”, Sustainability, 13/16: 8980. doi: 10.3390/su13168980

[31] Santos, F., Pereira, R., Vasconcelos, J. B. (2019), “Toward robotic process automation implementation: an end-to-end perspective”, Business Process Management Journal, 26/2: 405–420. doi:10.1108/bpmj-12-2018-0380

[32] Richter, S., Heumüller, E. (2020), “Robotic Process Automation- Ein kurzer Einblick”, Augenstein, F. (Ed.), Heumüller, E (Ed.), Robotic Process Automation- Einführung, Vorgehen und Anwendungen, Simmozheim: Conmethos-Verlag: 1-28.

[33] Capgemini Consulting (2016), “Robotic Process Automation - Robots Conquer Business Processes in Back Offices”, <https://www.capgemini.com/consulting-de/wp-content/uploads/sites/32/2017/08/robotic-process-automation-study.pdf>

[34] Villar, A. S., Khan, N. (2021), “Robotic process automation in banking industry: a case study on Deutsche Bank”, Journal of Banking and Financial Technology, 5: 71-86. doi: 10.1007/s42786-021-00030-9

[35] Brettschneider, J. (2020), “Bewertung der Einsatzpotenziale und Risiken von Robotic Process Automation”, HMD, 57: 1097–1110. Doi: 10.1365/s40702-020-00621-y

[36] Koch, C., Fedtke, S. (2020), Robotic Process Automation, Ein Leitfaden für Führungskräfte zur erfolgreichen Einführung und Betrieb von Software-Robots im Unternehmen. Heidelberg: Springer Publishing.

[37] Kaya, C. T., Türkyılmaz, M., Birol, B. (2019), “RPA Teknolojilerinin Muhasebe Sistemleri

Üzerindeki Etkisi”, Muhasebe ve Finansman Dergisi, 82: 235–250. doi: 10.25095/mufad.536083

[38] Wewerka, J., Reichert, M. (2020), “Towards Quantifying the Effects of Robotic Process Automation”, 2020 IEEE 24th International Enterprise Distributed Object Computing Workshop (EDOCW), 11-19. doi: 10.1109/edocw49879.2020.00015

[39] Tripathi, A. M. (2018), Learning Robotic Process Automation: Create Software robots and automate business processes with the leading RPA tool – UiPath. Packt Publishing.

[40] Kedziora, D., Leivonen, A., Piotrowicz, W., Öörni, A. (2021), “Robotic Process Automation (RPA) Implementation Drivers: Evidence of Selected Nordic Companies”, Issues in Information Systems, 22/2: 21-40. doi: 10.48009/2\_iis\_2021\_21-40

[41] Hofmann, P., Samp, C., Urbach, N. (2019), “Robotic process automation”, Electronic Markets, 30/1: 99–106. doi: 10.1007/s12525-019-00365-8

[42] Forrester (2019), “The Impact Of RPA On Employee Experience”. [https://dfe.org.pl/wp-content/uploads/2019/04/Forrester\\_RPA-Impact\\_Employee-Engagement.pdf](https://dfe.org.pl/wp-content/uploads/2019/04/Forrester_RPA-Impact_Employee-Engagement.pdf)

[43] Kaelble, S., NICE RPA. (2018), Robotics Process Automation for Dummies, West Sussex: Wiley.

[44] Protiviti (2019), “Taking RPA to the Next Level”. [https://www.protiviti.com/sites/default/files/united\\_kingdom/insights/2019-global-rpa-survey-protiviti\\_global.pdf](https://www.protiviti.com/sites/default/files/united_kingdom/insights/2019-global-rpa-survey-protiviti_global.pdf)

[45] Lacity, M., Willcocks, L., Craig, A. (2015a), “Robotic Process Automation at Telefónica O2”, [https://eprints.lse.ac.uk/64516/1/OUWRPS\\_15\\_02\\_published.pdf](https://eprints.lse.ac.uk/64516/1/OUWRPS_15_02_published.pdf)

[46] Lacity, M., Willcocks, L., & Craig, A. (2015c), “The IT Function and Robotic Process Automation”, [https://eprints.lse.ac.uk/64519/1/OUWRPS\\_15\\_05\\_published.pdf](https://eprints.lse.ac.uk/64519/1/OUWRPS_15_05_published.pdf)

[47] Willcocks, L., Lacity, M., Craig, A. (2015), “Robotic Process Automation at Xchanging”, [http://eprints.lse.ac.uk/64518/1/OUWRPS\\_15\\_03\\_published.pdf](http://eprints.lse.ac.uk/64518/1/OUWRPS_15_03_published.pdf)

[48] Lacity, M., Willcocks, L., & Craig, A. (2015b), “Robotic Process Automation: Mature Capabilities in the Energy Sector”,

[http://eprints.lse.ac.uk/64520/1/OUWRPS\\_15\\_06\\_published.pdf](http://eprints.lse.ac.uk/64520/1/OUWRPS_15_06_published.pdf)

[49] Siderska, J. (2021), “The Adoption of Robotic Process Automation Technology to Ensure Business Processes during the COVID-19 Pandemic”, Sustainability, 13/14: 8020. doi: 10.3390/su13148020

[50] Czarnecki, C., Auth, G. (2018), “Prozessdigitalisierung durch Robotic Process Automation”, Barton, T. (Ed.), Müller, C. (Ed.), Seel, C. (Ed.), Digitalisierung in Unternehmen, Wiesbaden: Springer: 113- 131. [https://doi.org/10.1007/978-3-658-22773-9\\_7](https://doi.org/10.1007/978-3-658-22773-9_7)

[51] YETİZ, F., TURAN, Y., & CANPOLAT, B. (2021), “BANKACILIK SEKTÖRÜNDE ROBOTİK SÜREÇ OTOMASYONU ve VERİMLİLİK İLİŞKİSİ: BİR BANKA ÖRNEĞİ”, Verimlilik Dergisi, 2: 65–80. doi:10.51551/verimlilik.765336

[52] Williams, Z. (2019), A Guide to Robotic Process Automation For the Average Worker: RPA Use Cases, and How to Keep Your Job Safe from Bots [E-book]

[53] Sahu, S., Salwekar, S., Pandit, A., Patil, M. (2020), “Invoice Processing Using Robotic Process Automation”, International Journal of Scientific Research in Computer Science, Engineering and Information Technology, 6/2: 216–223. doi:10.32628/cseit2062106

[54] Huang, F., Vasarhelyi, M. A. (2019), “Applying robotic process automation (RPA) in auditing: A framework”, International Journal of Accounting Information Systems, 35: 1-11. doi: 10.1016/j.accinf.2019.100433

[55] Van Chuong, L., Hung, P. D., Diep, V. T. (2019), “Robotic Process Automation and Opportunities for Vietnamese Market”, Proceedings of the 2019 7th International Conference on Computer and Communications Management, 86-90. doi: 10.1145/3348445.3348458

[56] Dimensional Research (2020, Mayıs), “ROBOTIC PROCESS AUTOMATION DELIVERS BUSINESS VALUE QUICKLY Global RPA Survey Results—Adoption, Benefits, and Lessons Learned”, <https://content.microfocus.com/robotic-process-automation-tb/robotic-process-auto-1>

[57] Datahorizonresearch.com, 2024, <https://datahorizonresearch.com/robotic-process-automation-market-2192>

[58] Langmann, C. (2021), “Robotic Process Automation (RPA) - Study on Characteristics of Successful RPA Implementations”. Hochschule München University of Applied Sciences. doi: 10.13140/RG.2.2.28379.69921

[59] Dey, S., Das, A. (2019), “Robotic Process Automation: Assessment of the Technology for Transformation of Business Processes”, International Journal of Business Process Integration and Management, 9/3: 220- 230. doi: 10.1504/IJBPM.2019.100927

[60] KIRAN, S., & DEMİRÇEVİREN, F. G. (2021). The influence of culture on the motivation for the use of social media: A comparative study of Turkish and German high school students. Beykoz Akademi Dergisi, 9(1), 19-33.

[61] Demir, İ. (2020), SPSS İle İstatistik Rehberi (1.baskı), İstanbul: Efe Akademi Yayınları.

[62] Yazıcıoğlu, Y., Erdoğan, S. (2014), SPSS Uygulamalı Bilimsel Araştırma Yöntemleri, Detay Yayıncılık.

[63] Blagoev, B. (2021), “End-User Satisfaction As a Result of RPA - A Finance and Accounting Perspective”.  
[http://essay.utwente.nl/86762/1/Blagoev\\_BA\\_BMS.pdf](http://essay.utwente.nl/86762/1/Blagoev_BA_BMS.pdf)



## An Investigation of International Field Indexes in the Discipline of Management Information Systems: A Bibliometric Analysis

Fatma AKGÜN\*, Cevriye GENCER<sup>2</sup>, Hakan ÖZKÖSE

<sup>1</sup>Bartın University, Faculty of Economics and Administrative Sciences, Department of Management Information Systems, BARTIN, 74100, TURKEY <sup>2</sup>Gazi University, Informatics Institute, Department of Management Information Systems, ANKARA, 06500, TURKEY

<sup>3</sup>Gazi University, Faculty of Engineering, Department of Industrial Engineering, ANKARA, 06500, TURKEY

<sup>3</sup>Bartın University, Faculty of Economics and Administrative Sciences, Department of Management Information Systems, BARTIN, 74100, TURKEY

### ARTICLE INFO

Received: 20.12.2024  
Accepted: 17.04.2025

**Keywords:** Management information systems, big data, bibliometric analysis, text mining, scientific mapping

#### \*Corresponding Authors

e-mail:  
fakgun@bartin.edu.tr

### ABSTRACT

Management Information Systems (MIS) has emerged as a rapidly developing field in recent years, attracting increasing interest from researchers due to its integration of digital innovations and the diversity of interdisciplinary research areas. This study aims to contribute both to the literature and the MIS discipline by conducting a comprehensive examination of international subject indexes. The study focuses on international indexes designated by the Interuniversity Board (ÜAK), which play a significant role in academic promotion and incentive criteria, particularly ESCI and Scopus. The suitability of these indexes for the MIS field has been evaluated, and the identification of similar alternative indexes has been pursued. To analyze similarities between indexes, expert opinions and bibliometric analysis methods were employed. The relevance of categories within Scopus and ESCI to the MIS field was assessed based on expert evaluations, and the indexes covering journals within these categories were identified. As a result, a total of 73 indexes were identified. Relationships among these indexes were analyzed using the visual mapping technique in VosViewer software. Findings obtained from scenarios with different resolution values indicate that ESCI and Scopus exhibit strong connections with other international indexes, such as the Engineering Index and Inspec. Based on the results, the study emphasizes the necessity of expanding existing indexes and proposes alternative indexes for MIS researchers. It is anticipated that this expansion will enhance academic productivity and contribute to the further development of the MIS discipline.

DOI: 10.59940/jismar.1604854

## Yönetim Bilişim Sistemleri Disiplini Uluslararası Alan İndekslerinin İncelenmesi: Bibliyometrik Bir Analiz

### MAKALE BİLGİSİ

Alınma: 20.12.2024  
Kabul: 17.04.2025

#### Anahtar Kelimeler:

Yönetim bilişim sistemleri, büyük veri, bibliyometrik analiz, metin madenciliği, bilimsel haritalama

#### \*Sorumlu Yazar

e-posta:  
fakgun@bartin.edu.tr

### ÖZET

Yönetim Bilişim Sistemleri (YBS), dijitalleşmenin getirdiği yenilikleri bünyesinde barındırması ve disiplinler arası çalışma alanlarının çeşitliliği nedeniyle son yıllarda daha fazla araştırmacının ilgisini çeken, gelişmekte olan bir alandır. Bu çalışma hem literatüre hem de YBS disiplinine katkı sağlamak amacıyla uluslararası alan indeksleri üzerine kapsamlı bir inceleme gerçekleştirmektedir. Çalışmada, ÜAK tarafından belirlenen ve doçentlik ile akademik teşvik kriterlerinde önemli bir yer tutan uluslararası alan indeksleri kapsamında ESCI ve Scopus ele alınmıştır. Bu indekslerin YBS alanındaki uygunluğu değerlendirilmiş ve benzer nitelikteki alternatif indekslerin belirlenmesi hedeflenmiştir. Bu doğrultuda, indeksler arasındaki benzerlikleri analiz etmek için uzman görüşleri ve bibliyometrik analiz yöntemleri kullanılmıştır. Scopus ve ESCI indekslerinde yer alan kategorilerin YBS ile uyumluluğu uzman görüşleri doğrultusunda değerlendirilmiş, ardından tespit edilen kategorilerde yer alan dergilerin tarandığı diğer indeksler belirlenmiştir. Çalışma kapsamında toplamda 73 indeks tespit edilmiştir. İndeksler arasındaki ilişkiler, VosViewer yazılımı kullanılarak görsel haritalama tekniği ile analiz edilmiştir. Farklı çözünürlük değerleriyle oluşturulan senaryolardan elde edilen bulgular, ESCI ve Scopus indekslerinin Engineering Index ve Inspec gibi diğer uluslararası indekslerle güçlü bir ilişki içinde olduğunu ortaya koymuştur. Elde edilen sonuçlar doğrultusunda, mevcut indekslerin genişletilmesi gerektiği vurgulanmış ve YBS alanında çalışan akademisyenler için alternatif indeks önerileri sunulmuştur. Bu durumun, akademik üretkenliği artırarak YBS disiplininin gelişimine katkı sağlayacağı öngörülmektedir.

DOI: 10.59940/jismar.1604854

## 1. INTRODUCTION (GİRİŞ)

Data is recognized as the fundamental raw material of the information age, and its significance continues to grow daily. The processes of data collection, analysis, and interpretation play a crucial role in various fields, ranging from scientific research to business, from social sciences to natural sciences. Data is transformed into information for numerous purposes, including explaining phenomena, testing hypotheses, predicting future trends, developing new theories, understanding customer behavior, formulating marketing strategies, and enhancing operational efficiency. Although obtaining data is not particularly difficult, processing large volumes of raw data and converting them into meaningful information presents a significant challenge. This process becomes even more critical when the information derived from data is intended to support decision-making. In addition to transforming data into information, ensuring that this transformation occurs both accurately and efficiently is equally important. Properly processed and timely data can provide organizations with a competitive advantage, whereas decisions based on inaccurate or incomplete data may lead to significant losses. Therefore, it is essential that data processing procedures are both rapid and precise. To ensure the effectiveness of these processes, a systematic approach is required. Management Information Systems (MIS) are decision-support systems designed to assist decision-makers in efficiently planning, controlling, and monitoring the processes for which they are responsible. These systems facilitate data-driven decision-making, enabling organizations to optimize their strategic and operational processes effectively.

The concept of Management Information Systems (MIS) emerged in the 1950s and refers to the integrated systems that collect, store, and utilize information in organizations to support decision-making processes. MIS effectively utilizes the data it gathers to assist decision-makers when necessary. Rather than merely being a system that supports strategic decision-making, MIS can be defined as a discipline that integrates technology and the human factor into organizational processes. Given its interdisciplinary nature, various definitions of MIS exist in the literature. Schoderbeck et al. (1975) defined MIS as a human-machine system that supports decision-making processes in businesses and provides decision-makers with various types of information [1]. Stoner (1982) described it as a system that manages the processes of collecting and storing the accurate data necessary for organizational management and providing this data to decision-makers at the right time to support their actions [2]. Additionally, MIS

facilitates more effective and efficient decision-making by supplying data to various organizational functions, including planning, control, auditing, and execution. Culnan (1987) defined MIS as a discipline that interacts with different fields and provides data to facilitate decision-making at various managerial levels [3]. Long (1989) described it as a system that delivers the necessary information at the required time to enhance organizational efficiency and effectiveness [4]. Adeoti Adekeye (1997) defined MIS as a system that provides essential information for different processes such as decision-making and planning within organizations [5]. Lee (2001) described MIS as a system that enables the optimization of organizational processes by supplying the necessary data for more efficient operations [6]. Baskerville and Myers (2002) defined it as the development and effective use of information systems within organizations [7]. Bensghir (2002) characterized MIS as a dynamic and continuously evolving field that emerges from the convergence of multiple disciplines, including management science, computer science, information systems, and organizational behavior [8]. Laudon and Laudon (2004) defined MIS as a system that assists business decision-making processes through information technology systems [9]. Becta (2005) described it as a system that provides data to all necessary units within an organization while ensuring communication between these units [10]. O'Brien and George (2007) defined MIS as systems that support decision-makers through reports, graphics, and other documents during the decision-making process [11]. Gökçen (2011) described MIS as a system that enables organizations to use data and information more efficiently and effectively [12]. Pratap (2018) defined MIS as a multi-disciplinary system that integrates various fields and considers hardware, software, data, processes, and the human factor as a unified whole. Each of these components is critical for the functionality of MIS and constantly interacts with one another [13]. Overall, MIS requires the integration of information technology, management, business, and human factors within a holistic framework to optimize organizational decision-making and operational processes.

When examining the historical development of Management Information Systems (MIS), it is evident that, despite emphasizing different aspects over time, the primary focus has always been on enhancing the efficiency and effectiveness of organizational business processes. The multidisciplinary nature of MIS has also shaped its evolution throughout history. Table 1 presents the historical development and transformations that MIS has undergone over time. The evolution of MIS has closely followed advancements in computer technology, demonstrating a parallel progression.

Table 1. The roles of management information systems in the historical process  
(Yönetim bilişim sistemlerin tarihsel süreçte rolleri)

| Year           | Start                                              | Area of Use                                                                                                                                                                            |
|----------------|----------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| 1950           | Data Processing Systems                            | It was used to store large amounts of data and perform basic calculations.                                                                                                             |
| 1960           | Support for Decision Making                        | Data was analyzed to provide meaningful information for managers to use in decision-making processes.                                                                                  |
| 1970           | Database Management Systems                        | Data began to be used centrally in a shared repository.                                                                                                                                |
| 1980           | The proliferation of personal computers            | With the widespread use of personal computers, employees became able to access the information they needed instantly and make faster decisions..                                       |
| 1990           | Global Connections Through the Internet            | With the influence of the internet, the use of enterprise information systems has facilitated information sharing by consolidating all company data into a single platform.            |
| 2000 and After | The Age of Cloud Computing and Mobile Applications | With the development of mobile applications, Management Information Systems (MIS) have become accessible through smartphones and tablets as well.                                      |
| Nowadays       | The Impact of Artificial Intelligence and Big Data | With the integration of artificial intelligence and machine learning technologies, previously unseen relationships have been discovered, and more accurate predictions have been made. |

As can be seen from Table 1, the concept of Management Information Systems (MIS) has evolved alongside the advancements in computer technology, artificial intelligence, big data, and other technologies, shaping it into its current form. Being a discipline influenced by technological developments, it is expected to continue its development in future processes.

The primary objective of this study is to evaluate whether the international subject indices defined in the field of Management Information Systems (MIS) can be expanded. There is insufficient clarity regarding the criteria used by the Council of Higher Education (ÜAK) in determining the international subject indices and the selection process. In this study, analyses will be conducted based on the ESCI and Scopus indices, as determined by ÜAK. Initially, the suitability of these indices for the MIS field will be assessed through expert opinions. Based on expert evaluations, other international indices similar to these, found to be compatible with the MIS field, will be identified. The identified indices will be examined using various analytical methods, and the possibility of establishing new international subject indices will be investigated. In this regard, the relationship between only the ESCI and Scopus indices and the MIS field will not be evaluated, but the goal will also be to introduce new international indices to the field by expanding the scope of these existing indices.

## 2. LITERATURE REVIEW (LİTERATÜR TARAMASI)

MIS (Management Information Systems) is a complex structure that arises from the integration of various disciplines within the rapidly changing dynamics of technology and the business world. Consequently, it is difficult to find a comprehensive definition of MIS. Efforts to define MIS in the literature often emphasize different aspects of the system. This is because MIS is a system formed by the combination of fundamental components such as hardware, software, databases, networks, and human resources. The integration of these components enables MIS to meet the information needs of businesses, support decision-making processes, and enhance their competitive strength. Due to these

characteristics, the concept of MIS is a discipline that has been extensively studied in the literature.

Although the origins of Management Information Systems (MIS) date back to the 1950s, the development process in Turkey occurred later than the global average. The first academic studies and educational programs in this field were initiated in the 1990s. Most of the pioneering studies in MIS in Turkey have been carried out at Marmara University. The development of MIS education in Turkey began in 1991 with the establishment of the MIS department at Marmara University. While this field started relatively late in Turkey, it has developed rapidly. The process, which began with the first department opened at Marmara University, has led to the establishment of MIS programs at many universities today, providing opportunities to train qualified human resources in the field. Since the establishment of the first MIS departments in 1991, the number of departments offering MIS education in Turkey has steadily increased. Today, there are a total of 86 universities offering MIS programs, including 41 state universities and 45 private universities. This number continues to grow each year. While this has significantly contributed to the increase in the number of researchers and the diversification of studies in MIS, the field is still not at a sufficient level of development.

The fact that researchers in the MIS field come from various scientific disciplines further accentuates its interdisciplinary nature. Researchers from fields such as computer science, business, mathematics, statistics, and even sociology are conducting studies in MIS. On the one hand, this brings a wide range of research topics to the field, while on the other hand, it complicates the formation of a unified structure for the field. The relatively new status of MIS in Turkey also contributes to the interdisciplinary makeup of its researchers. This leads researchers to transfer their knowledge and skills from their original disciplines to the MIS field, which facilitates the integration of different disciplinary approaches and the development of new methods, but also contributes to the absence of a well-established methodology in the field. Numerous national and

international studies have been conducted to identify the contributions of MIS to science and society and to explore its general structure as an emerging academic discipline.

Barki et al. (1988) examined 792 articles published in the field of MIS between 1968 and 1988 to identify various metrics related to the field. As a result of the study, they identified key topics within the MIS discipline [14]. Seo and Han (1997) used subject analysis, citation analysis, and co-author analysis methods to uncover the general structure of the MIS field. They identified prominent journals and influential authors in the field [15]. Bensghir (2002), through content analysis, assessed the state of academic studies in MIS in Turkey, comparing them with global developments. The study found that academic research in Turkey was in line with global trends in the MIS field [8]. Cocosila et al. (2011) performed a bibliometric analysis of 452 papers presented at three MIS conferences between 1974 and 2008. The study revealed various metrics such as the most influential authors, frequently discussed topics, and the number of presentations over the years [16]. Mohanty (2014) conducted a bibliometric analysis of 596 articles published in the *MIS Quarterly* journal from 1995 to 2009, uncovering various metrics such as the most effective countries, authors, and article topics [17]. Yarlitaş (2015) analyzed completed graduate theses in the MIS field using content analysis, identifying related disciplines, research methods, and topics of focus. The study showed that, as expected, the field has strong relationships with disciplines such as management, information technology, and computer science, confirming the interdisciplinary nature of MIS. The study also found that both qualitative and quantitative research methods were used in the theses [18]. Lin et al. (2016) conducted a bibliometric analysis of 853 articles published between 1991 and 2014, identifying influential authors, topics, and the connections between publications [19]. Özköse (2017) analyzed articles from Scopus and WOS databases published in the MIS field between 1980 and 2015 using bibliometric methods.

The study identified influential authors, institutions, countries, the number of articles by year, the most cited publications, and journals [20]. Özköse and Gencer (2017) conducted a bibliometric analysis of 24 journals indexed in SCI-E and SSCI within the WOS database, revealing information about influential institutions, authors, and countries in the MIS field [21]. Beydoun et al. (2019) performed a bibliometric analysis of 855 articles published in *Information Systems Frontiers* between 1999 and 2018, identifying subject distributions in the journal's articles [22]. Çoşkun et al. (2019) used data and text mining methods to uncover the topics in MIS articles published in the WOS database from 2008 to 2019. The study revealed changing subject trends over the years [23]. Aytaç (2020) conducted a content analysis of graduate theses published between 2015 and 2020, identifying frequently discussed topics in the theses [24]. Jeyaraj and Zadeh (2020) used topic modeling to analyze 2962 articles published between 2003 and 2017 in five significant journals related to the MIS discipline, identifying 50 different subject headings [25]. Damar and Aydın (2021) analyzed 104 journals at the SCImago Q1 level from 2010 to 2021 using bibliometric methods. The study revealed various metrics, including frequently used topics and keywords, the most influential authors, and the field of study for Turkish researchers [26]. Cebeci (2021) analyzed 963 articles published between 2000 and 2020 in 107 journals in the Scopus database using multi-criteria decision-making methods. The study applied bibliometric analysis, trend analysis, text mining, and clustering methods, revealing that customer-focused topics and data mining were the most studied subjects related to MCDM methods [27]. Damar and Özdağoğlu (2022) conducted a bibliometric analysis of 1550 publications published in *MIS Quarterly* between 1980 and 2020. The study identified the disciplines and countries related to MIS and examined their productivity [28]. Özköse et al. (2023) conducted a bibliometric analysis of 25,304 articles published in the Scopus database between 2016 and 2021, identifying the most influential authors, institutions, countries, and journals in the field [29].

Table 2. Review of the MIS Literature  
(YBS Literatürünün incelenmesi)

|          | Year | Author/Authors          | Method Used                                                    |
|----------|------|-------------------------|----------------------------------------------------------------|
| National | 2002 | Bensghir[8]             | Content Analysis                                               |
|          | 2015 | Yarlitaş[18]            | Content Analysis                                               |
|          | 2017 | Özköse[20]              | Bibliometric Analysis                                          |
|          | 2017 | Özköse and Gencer[21]   | Bibliometric Analysis                                          |
|          | 2019 | Coşkun[23]              | Data and Text Mining                                           |
|          | 2019 | Ergüner Özkoç[40]       | Bibliometric Analysis                                          |
|          | 2020 | Aytaç[24]               | Content Analysis                                               |
|          | 2021 | Damar and Aydın[26]     | Bibliometric Analysis                                          |
|          | 2021 | Cebeci                  | Bibliometric Analysis, Trend Analysis, Text Mining, Clustering |
|          | 2022 | Damar and Özdağoğlu[27] | Bibliometric Analysis                                          |
|          | 2023 | Özköse et al.[29]       | Bibliometric Analysis                                          |
|          | 2023 | Sertçelik and Önder[39] | Apriori Algorithm                                              |
|          | 2024 | Güler and Zeren[43]     | Bibliometric Analysis                                          |

|               |      |                            |                                                                  |
|---------------|------|----------------------------|------------------------------------------------------------------|
| International | 2024 | Kırmızıyaka and Öztürk[44] | Bibliometric Analysis                                            |
|               | 2024 | Ünal Kestana[45]           | Bibliometric Analysis                                            |
|               | 2025 | Akgün et al.               | Bibliometric Analysis                                            |
|               | 1988 | Barki et al. [14]          | Data and Text Mining                                             |
|               | 1997 | Seo and Han[15]            | Subject Analysis, Citation Analysis, Author Co-Citation Analysis |
|               | 2011 | Cocosila et al.[16]        | Bibliometric Analysis                                            |
|               | 2014 | Mohanty[17]                | Bibliometric Analysis                                            |
|               | 2015 | Shiau[41]                  | Citation Analysis Method                                         |
|               | 2016 | Lin et al.[19]             | Bibliometric Analysis                                            |
|               | 2019 | Beydoun et al.[22]         | Bibliometric Analysis                                            |
|               | 2020 | Abedin et al.[42]          | Bibliometric Analysis                                            |
|               | 2020 | Jeyaraj and Zadeh[25]      | Topic Modeling Method                                            |
|               | 2024 | Aryawati[47]               | Bibliometric Analysis                                            |
|               | 2024 | Modina et al. [48]         | Bibliometric Analysis                                            |
|               | 2025 | Hidayat et al. [46]        | Bibliometric Analysis                                            |

In the literature of Management Information Systems (MIS), it is evident that this field is frequently studied and efforts are being made to establish a framework for it. However, due to both technological advancements and the interdisciplinary nature of MIS, its boundaries cannot be clearly defined. The topics studied and the methods used vary depending on the year and the expertise of the researchers. Research in MIS is carried out at both national and international levels, with a focus primarily on quantitative methods such as data and text mining. Data mining techniques help to uncover hidden information and relationships within large datasets, while text mining methods identify key terms and concepts in scientific publications, reports, and other documents, enabling the tracking of the field's development. Due to their nature, data and text mining methods facilitate the discovery of hidden information within patterns. These methods are often used to salvage a congested literature, providing a clearer image and an overall picture of the field. Therefore, these techniques are frequently employed to determine the boundaries of a field. When the literature on the subject is reviewed, it becomes clear that databases, journals, theses, articles, conferences, and papers related to MIS have been extensively studied. Various data and text mining methods such as content analysis, topic modeling, bibliometric analysis, citation analysis, clustering, classification, and apriori algorithms have been used in these studies.

As a result of the literature review, no studies specifically focusing on international indexing in the field of Management Information Systems (MIS) were found. One of the key challenges for academicians working in MIS is the uncertainty regarding which indexes the journals in this field are included in and whether these indexes are considered international. Although the Interuniversity Board (ÜAK) has made some regulations on this matter, existing studies do not provide sufficient clarity, leading to the limitation of the field to certain indexes. International indexes are typically defined as those outside the SSCI, SCI, SCI-E, and AHCI categories. However, this definition is not sufficiently clear. The interpretation of this definition varies across different universities according to their

academic promotion and incentive criteria, leading to inconsistencies in practice. Even though efforts were made to clarify certain indexes through regulations issued on June 15, 2023, uncertainty still remains regarding the criteria used to determine international indexes and how they differentiate from existing ones.

The main aim of this study is to evaluate the suitability of indexes outside the international area indexes determined by the Interuniversity Board (ÜAK) for the Management Information Systems (MIS) field, to determine whether these indexes can be considered international area indexes, and to analyze the relationship between the existing indexes and the MIS field. The study will seek answers to the following questions:

1. What are the categories within the field of Management Information Systems?
2. What are the other indexes that the journals in the Management Information Systems field are scanned in?

### 3. MATERIALS AND METHODS (MATERYAL VE YÖNTEM)

Science is a cumulative structure that progresses incrementally and continually develops. Each completed study forms the foundation for future research. Every new piece of knowledge can be tested and further developed by subsequent studies. Every academic work contributes to deepening and advancing knowledge. The benefit gained is not limited to academic growth alone but also facilitates the development and advancement of societies and improves their quality of life. Academic research, due to its collaborative nature with various scientific disciplines, fosters interdisciplinary approaches that accelerate scientific progress and provide solutions to more complex problems. This is because research is not confined to specific disciplines but progresses through interactions with different fields. As a result, science advances and offers new perspectives. Academic studies are communicated to the scientific community and society through scientific papers, conferences, and symposiums in various ways.

The progress, development, and dissemination of science are integral parts of scientific journals. Journals are included in various indexes when they meet specific criteria. Indexes are important tools that facilitate access to scientific works and categorize them in different ways. They provide researchers with the opportunity to access the information they need more easily. Given the variability in the quality and scope of different indexes, it is quite challenging to group them under one umbrella. The varying quality of journals and the diversity of the content they publish necessitate the classification of these indexes into specific categories. The classification of an index depends on various factors, which can change based on specific criteria. Therefore, researchers should be careful when choosing indexes and should determine which ones are most suitable for their field of study. For a scientific journal, being indexed in a reputable database increases its visibility, benefiting both the journal and the published works. Indexed journals reach a wider audience, resulting in more frequent citations. The increase in citations enhances the journal's prestige and makes it a more reliable source in the eyes of other researchers. There are many classifications for journal indexes, among which international area indexes will be discussed in this study. These indexes play a significant role in academic promotions, title changes, and academic incentives. Before June 15, 2023, there was ambiguity surrounding the concept of international area indexes, but after this date, the definition has been clarified for certain fields. Although ÜAK (Higher Education Council) previously published a list of 209 indexes, universities were making their own decisions and defining which indexes they would accept, leading to inconsistencies in practice. In the context of this study, the concept of international area indexes has been narrowed down to Scopus and ESCI. However, the distinctions and importance of these indexes have not been explained. This study aims to explore the differences between the selected indexes and determine if there are other indexes similar to these, based on the analysis conducted.

In summary, the aim of this study is to uncover certain dynamics of the field of Management Information Systems (MIS), which has gained significant popularity in recent years, particularly due to the innovations brought about by digitalization. The increasing diversity within the field and its wide appeal to numerous researchers have contributed to its growing prominence. However, due to the relatively new nature of the discipline, clear boundaries have not yet been established, and academic work is only just beginning to accelerate. This study will conduct an extensive review of the accepted indexes within the Management Information Systems field, aiming to expand the current list of indexes. By using bibliometric methods, the current status of these indexes will be mapped, and indexes with similar characteristics will be identified.

### 3.1. Index Selection (*İndeks Seçimi*)

In the study, the indexes to be addressed have been selected based on the international domain indexes recognized by the Yükseköğretim Kurulu (ÜAK) for academic promotion and appointment criteria. These indexes are Scopus and ESCI, as outlined by ÜAK.

Scopus, launched by Elsevier in 2004, is an index that provides content across various disciplines such as natural sciences, social sciences, medicine, and arts. Widely used in research and literature review, the Scopus database includes journals, conference papers, books, and patents. It offers comprehensive content, citation tracking, currency, filtering options, research performance measurement, and various integrated tools, making it a frequently used database by researchers. Scopus contains high-quality academic sources, which is one of the primary reasons for researchers' preference for the database. Currently, more than 25,000 journals are indexed in Scopus.

The Web of Science (WOS) is a database developed in 2015 that covers a wide range of disciplines, including natural sciences, social sciences, arts, and humanities. ESCI (Emerging Sources Citation Index) is a subset of the WOS database, designed to increase the visibility of journals that may not yet meet the criteria for WOS inclusion but still meet many standards. ESCI indexes journals that are close to meeting the required standards for inclusion in WOS but have not yet made the cut. Over time, as these journals meet the necessary standards, they can transition to the SCI, SSCI, or AHCI indexes. To be included in ESCI, a journal must meet specific criteria. Many academic journals in Turkey are indexed in ESCI.

### 3.2. Bibliometric Analysis (*Bibliyometrik Analiz*)

With the rapid advancement of digitalization, the importance of data and information has significantly increased. As societies transition into the information age, massive data accumulations have emerged, introducing the concept of big data into our lives. The exponential growth in data volume has made extracting meaningful insights from these vast datasets increasingly challenging. Data mining serves as a powerful analytical tool for uncovering hidden and meaningful patterns within large and complex datasets. One of the widely used data mining methods, bibliometric analysis, was first introduced by Pritchard in 1969. Bibliometric analysis employs statistical and mathematical techniques to examine scientific publications. Beyond merely providing publication-related metrics, bibliometric analysis offers insights into a given field, making scientific processes more comprehensible. This method allows for the assessment of scientific impact, the identification of research trends, and the examination of collaboration patterns. By analyzing bibliometric data, researchers can gain deeper insights into the dynamics of a discipline and uncover underlying patterns. Additionally, bibliometric analysis

helps outline the broader framework of a research field while identifying existing gaps, guiding future studies. Due to these attributes, bibliometric analysis plays a crucial role in advancing scientific research. Various methods are employed to derive relevant metrics in bibliometric analyses, ensuring a comprehensive understanding of the field.

#### Bibliometric Analysis Methods;

1. Citation Analysis: Examines how frequently a published work is cited by other studies, providing insights into its impact and influence.
2. Co-Citation Analysis: Investigates how often two or more studies are cited together, helping to identify relationships between research topics.
3. Co-Authorship Analysis: Analyzes collaborations between researchers, mapping networks and identifying patterns of academic cooperation.
4. Keyword Analysis: Evaluates the keywords used in publications to provide metrics related to research topics and emerging trends.
5. Scientific Mapping: Visually represents scientific fields and examines relationships between different disciplines.
6. H-Index: Measures the impact of researchers based on the number of their publications and citation frequency.

Each of these techniques provides a unique perspective on research outputs and academic interactions, playing a crucial role in evaluating scientific knowledge. Additionally, bibliometric methods are inherently objective as they rely on quantitative data and offer extensive coverage due to their ability to process large datasets. In this study, scientific mapping and clustering methods will be employed to analyze the relationships between academic works and identify key research trends.

### 3.3. Problem Identification (Problemin Belirlenmesi)

In academic literature, an important consideration is the index of the journal in which an article will be published. The quality of the target journal is crucial for authors, especially when the publication will be used for academic promotion and tenure evaluations. To meet these criteria, both the journal itself and the indexes in which it is listed hold significant importance for researchers. The starting point of this research is the international field indexes in the field of Management Information Systems (MIS). The Turkish Council of Higher Education (ÜAK) has designated two indexes for this field: Scopus and ESCI. The core research problem stems from the question: What criteria were used to select these indexes, and are there other indexes that meet these criteria? This question will guide the structure and direction of the study.

### 3.4. Determination of Research Constraints (Araştırma Kısıtlarının Belirlenmesi)

Within the scope of this study, certain limitations have been established to define the research framework and influence its direction. These limitations are as follows:

1. In certain sections of the study, expert evaluations will be sought. To ensure that experts can contribute effectively to the field of Management Information Systems (MIS), specific criteria have been established. The required qualifications for experts include: holding a Ph.D. in the field of MIS and having conducted research in this area; having supervised students in a thesis/non-thesis master's or doctoral program in MIS and having engaged in research within the field; and having taught courses in MIS for at least two years while also conducting research in this domain.
2. The selection of indices and journals to be evaluated in this study has been determined with consideration of the research domain. Although the indices recognized by the Interuniversity Board (ÜAK) were initially reviewed, the extensive number of indices and the presence of those unrelated to the field necessitated a more focused approach. Consequently, the study is based on the core indices acknowledged in academic appointment and promotion criteria. Furthermore, only journals within these indices that specifically publish in the field of Management Information Systems (MIS) have been included in the analysis.

### 3.5. Determination of Research Population (Araştırma Evreninin Belirlenmesi)

The research will focus on the internationally recognized indices for the MIS discipline, namely Scopus and ESCI. The categories and journals within these indices will be identified. By determining the categories on Scopus (334) and ESCI (542), a total of 876 categories have been obtained. Upon reviewing the journals within these categories, a total of 33,825 journals have been identified, with (25,837) from Scopus and (7,988) from ESCI.

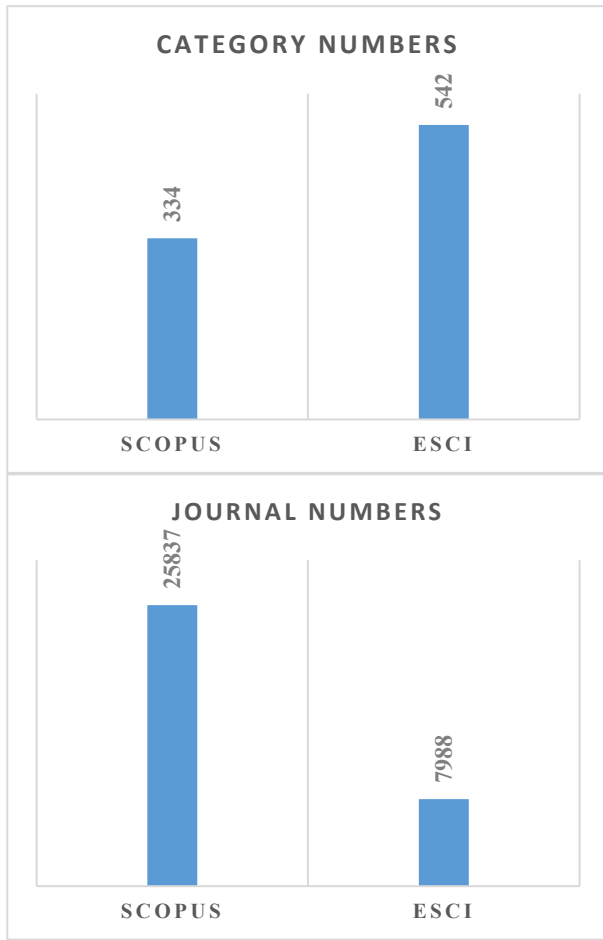


Figure 1. ESCI and Scopus Category and Journal numbers  
(ESCI ve Scopus kategori ve dergi sayıları)

Expert opinions will be used to include only the categories relevant to the Management Information Systems (MIS) discipline in the study, while other categories will be excluded. The Scopus and ESCI categories have been prepared using Microsoft Excel to facilitate the experts' evaluations. During the evaluation, the experts were asked to rate the suitability of the identified categories for the MIS discipline on a scale from 1 to 5. The data provided by the experts were then compiled using Microsoft Excel, and the geometric mean of the ratings was calculated. In the average, three equal value ranges were defined: 5, 5-4.5, and 4.5-4. The categories to be included in the research based on these threshold values are presented in Table 3 and Table 4.

Table 3. ESCI category and journal issues  
(ESCI kategori ve dergi sayıları)

| Code | Category                                                                                                                                                                                                         | Point | Journal Numbers |
|------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-------|-----------------|
| S1   | Computer Science, Information Systems                                                                                                                                                                            | 5     | 60              |
| S2   | Computer Science, Interdisciplinary Applications   Computer Science, Artificial Intelligence   Computer Science, Cybernetics   Computer Science, Information Systems                                             | 4,92  | 1               |
| S2   | Computer Science, Artificial Intelligence   Computer Science, Information Systems                                                                                                                                | 4,83  | 2               |
| S2   | Computer Science, Theory & Methods   Computer Science, Artificial Intelligence   Computer Science, Information Systems                                                                                           | 4,69  | 2               |
| S2   | Computer Science, Interdisciplinary Applications   Computer Science, Artificial Intelligence   Computer Science, Information Systems                                                                             | 4,69  | 1               |
| S2   | Computer Science, Artificial Intelligence                                                                                                                                                                        | 4,67  | 28              |
| S2   | Computer Science, Theory & Methods   Computer Science, Artificial Intelligence                                                                                                                                   | 4,67  | 1               |
| S2   | Computer Science, Artificial Intelligence   Computer Science, Cybernetics   Computer Science, Information Systems                                                                                                | 4,54  | 1               |
| S3   | Computer Science, Theory & Methods   Multidisciplinary Sciences                                                                                                                                                  | 4,46  | 1               |
| S3   | Multidisciplinary Sciences   Computer Science, Information Systems                                                                                                                                               | 4,456 | 1               |
| S3   | Computer Science, Interdisciplinary Applications   Computer Science, Theory & Methods   Engineering, Electrical & Electronic   Computer Science, Hardware & Architecture   Computer Science, Information Systems | 4,42  | 1               |
| S3   | Computer Science, Interdisciplinary Applications   Multidisciplinary Sciences   Computer Science, Information Systems                                                                                            | 4,32  | 1               |
| S3   | Computer Science, Interdisciplinary Applications                                                                                                                                                                 | 4,18  | 32              |
| S3   | Computer Science, Interdisciplinary Applications   Computer Science, Software Engineering   Computer Science, Information Systems                                                                                | 4,17  | 2               |
| S3   | Business                                                                                                                                                                                                         | 4,06  | 131             |
| S3   | Computer Science, Interdisciplinary Applications   Engineering, Multidisciplinary   Materials Science, Multidisciplinary                                                                                         | 4,04  | 1               |

S3

|                                                                                                                                                                           |      |            |
|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------|------------|
| Computer Science, Interdisciplinary Applications   Computer Science, Theory & Methods   Computer Science, Hardware & Architecture   Computer Science, Information Systems | 4,04 | 1          |
| <b>Total</b>                                                                                                                                                              |      | <b>267</b> |

Table 4. Scopus category and journal issues  
(Scopus kategori ve dergi sayıları)

| Code | Category                                    | Point | Journal Numbers |
|------|---------------------------------------------|-------|-----------------|
| 1404 | Management Information Systems              | 5     | 141             |
| 1709 | Human-Computer Interaction                  | 5     | 153             |
| 1710 | Information Systems                         | 5     | 413             |
| 1802 | Information Systems and Management          | 5     | 152             |
| 1405 | Management of Technology and Innovation     | 4,68  | 380             |
| 1702 | Artificial Intelligence                     | 4,58  | 298             |
| 2214 | Media Technology                            | 4,49  | 140             |
| 1705 | Computer Networks and Communications        | 4,34  | 425             |
| 1704 | Computer Graphics and Computer-Aided Design | 4,27  | 112             |
| 2614 | Theoretical Computer Science                | 4,25  | 139             |
| 1700 | Computer Science Applications               | 4,17  | 324             |
| 1706 | General Computer Science                    | 4,05  | 875             |
|      | <b>Total</b>                                |       | <b>3552</b>     |

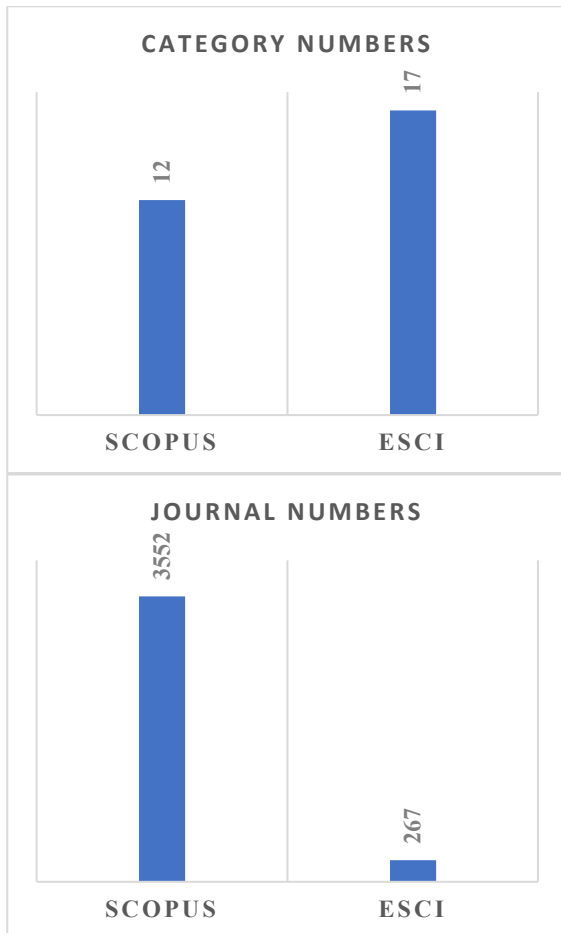


Figure 2. Number of categories and journals included in the study  
(Çalışmaya dahil edilen kategori ve dergi sayıları)

Based on expert opinions, 12 categories from the 334 Scopus categories and 17 categories from the 542 ESCI categories were included in the study, while the remaining categories were excluded. As a result of the expert evaluations, 3,552 journals from the 25,837 Scopus journals and 267 journals from the 7,998 ESCI journals were included in the research.

### 3.6. Research Methodology (Araştırma Metodolojisi)

In this study, expert opinions and bibliometric analysis methods will be utilized. Expert opinion involves a person with in-depth knowledge and experience on a particular subject analyzing and evaluating a specific issue or topic, and offering a suggestion or solution based on this evaluation. This opinion is generally based on expertise acquired through education, experience, and research. In short, expert opinion is a helpful method for gaining knowledge or making decisions about a topic. Bibliometric methods can be defined as mathematical techniques applied to derive statistical metrics of publications such as books, journals, and articles (Pritchard, 1969) [30]. Through bibliometric methods, the current structure and trends in the field will be identified. The research will proceed according to the following workflow diagram.

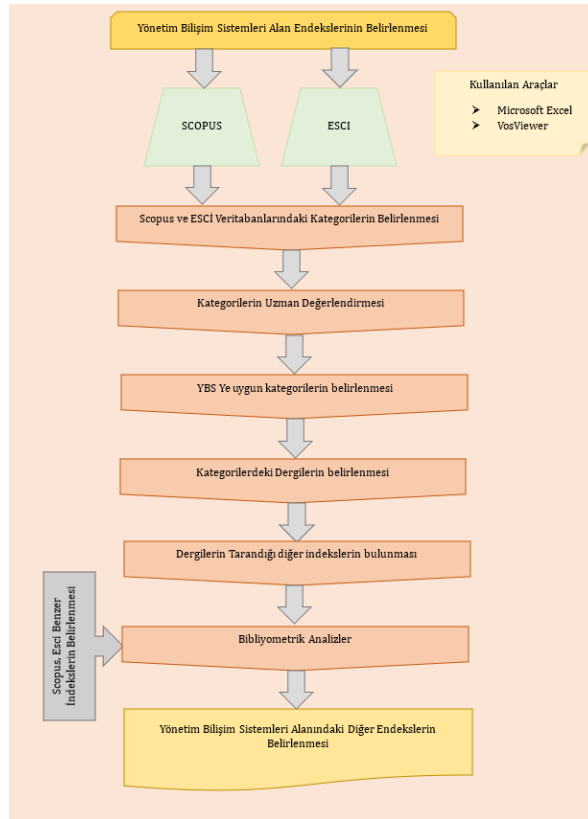


Figure 3. Research flow chart  
(Araştırma akış şeması)

### 3.7. Data Collection and Pre-processing (Veri Toplama ve Ön İşleme)

In this section of the study, certain processes will be applied to the data to enhance the quality of the analyses and increase the likelihood of obtaining meaningful

results from the analyses [37]. After defining the research population, the next step is data collection. At this stage, raw data is processed to prepare it for analysis [38]. The operations performed on the data prior to analysis are presented in Figure 4 below.

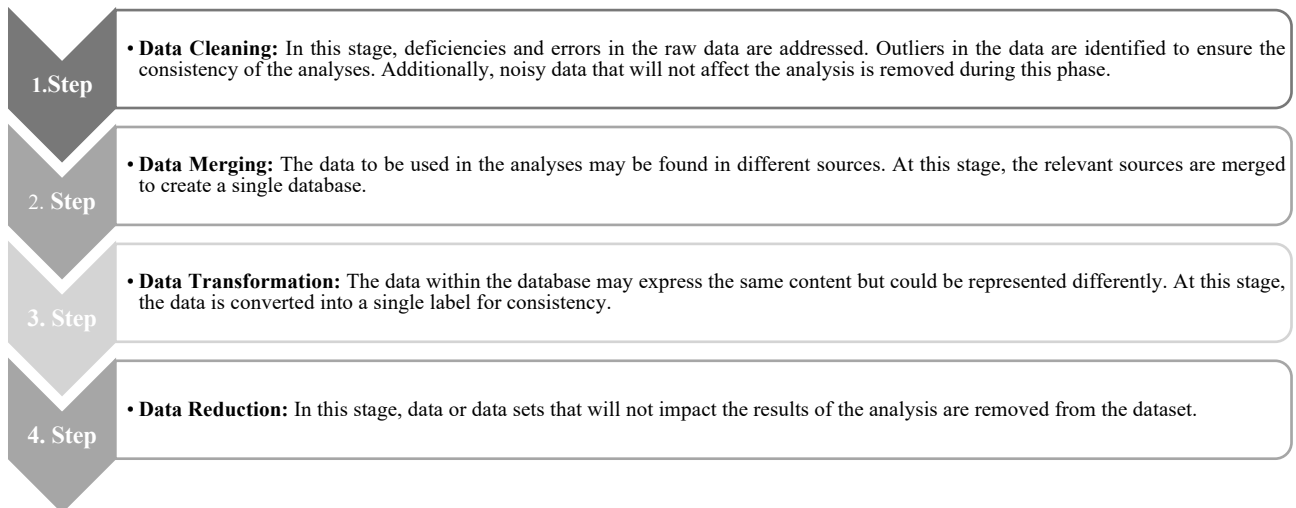


Figure 4. Data pre-cleaning stages  
(Veri ön temizleme aşamaları)

In the study, the 3,819 identified journals were consolidated using Microsoft Excel, and duplicate journals were detected. Each index was considered a separate database, and upon merging the two index datasets, it was found that many journals appeared in both indexes. The duplicate journals were removed from

the data file, leaving 1,471 unique journals. The index information for the identified journals was then scanned. From the scanned indexes, only those that were accepted as international field indexes were compiled using Microsoft Excel, resulting in a total of 73 international field indexes.

### 3.8. Importance of the Study and Contribution to the Field (*Çalışmanın Önemi ve Alana Katkısı*)

In this study, important indexes in the field of Management Information Systems (MIS) will be identified and the scope of the field will be expanded. A review of the literature reveals no existing studies on this topic, which highlights the significance of this research. It is not clearly known which criteria the indexes determined by YÖK (Higher Education Council) meet or how they are selected. Based on this, the presence of different indexes that meet similar criteria and characteristics as the selected ones will be investigated. As a result of the research, the identified indexes will be proposed, with the aim of expanding the scope of international indexes for MIS. This is expected to increase academic productivity and sustainability, as the limited number of indexes and the small number of journals in those indexes lead to long waiting times for researchers. This situation reduces academic productivity. If the number of indexes is increased, this will also lead to a rise in the number of journals, bringing diversity and more alternatives with it.

### 4. FINDINGS (*BULGULAR*)

In this study, scientific mapping was conducted using VosViewer. Through the scientific mapping method, the interactions and connections between the indexes will be examined. Different scenarios have been created to analyze and interpret the results more clearly. The number of occurrences used in creating these scenarios has been set as the threshold value. Table 5 presents the scenarios and relevant information related to those scenarios. For each scenario, an examination was carried out with three resolution values.

Table 5. Scenarios applied in the study  
(*Çalışmada uygulanan senaryolar*)

|    | Threshold Value | Number of Indexes | Resolution | Number of Clusters |
|----|-----------------|-------------------|------------|--------------------|
| S1 | 23              | 73                | 1          | 3                  |
|    |                 |                   | 1.15       | 5                  |
|    |                 |                   | 1.30       | 3                  |
| S2 | 53              | 51                | 1          | 2                  |
|    |                 |                   | 1.15       | 4                  |
|    |                 |                   | 1.30       | 5                  |
| S3 | 70              | 45                | 1          | 2                  |
|    |                 |                   | 1.15       | 3                  |
|    |                 |                   | 1.30       | 3                  |
| S4 | 100             | 36                | 1          | 2                  |
|    |                 |                   | 1.15       | 3                  |
|    |                 |                   | 1.30       | 4                  |
| S5 | 151             | 26                | 1          | 2                  |
|    |                 |                   | 1.15       | 2                  |
|    |                 |                   | 1.30       | 2                  |

Different scenarios were created to investigate how changes in the number of clusters and the associations between terms varied. Within the scenarios, clustering differences were examined using three different resolution values: 1, 1.15, and 1.30. The numbers in the scenarios were determined through trial and error to achieve the best possible results.

#### 4.1. Scenario 1 (*Senaryo 1*)

In this stage, the threshold value was set to 23 in order to display all index values in the dataset. All 73 indices in the dataset were included in the analysis. When the resolution value for the 73 indices was set to 1, the data was divided into 3 different clusters; when the resolution value was set to 1.15, the data was divided into 5 different clusters; and when the resolution value was set to 1.30, the data was divided into 5 different clusters. The cluster divisions are shown in Table 6.

Table 6. Clusters according to Scenario 1  
(*Senaryo 1'e göre kümeler*)

| Resolution =1 |                    | Resolution =1.15 |                    | Resolution =1.30 |                    |
|---------------|--------------------|------------------|--------------------|------------------|--------------------|
| Clusters      | Number of Elements | Clusters         | Number of Elements | Clusters         | Number of Elements |
| Red           | 27                 | Red              | 22                 | Red              | 21                 |
| Green         | 25                 | Green            | 19                 | Green            | 20                 |
| Blue          | 21                 | Blue             | 14                 | Blue             | 12                 |
|               |                    | Yellow           | 11                 | Yellow           | 11                 |
|               |                    | Purple           | 7                  | Purple           | 9                  |

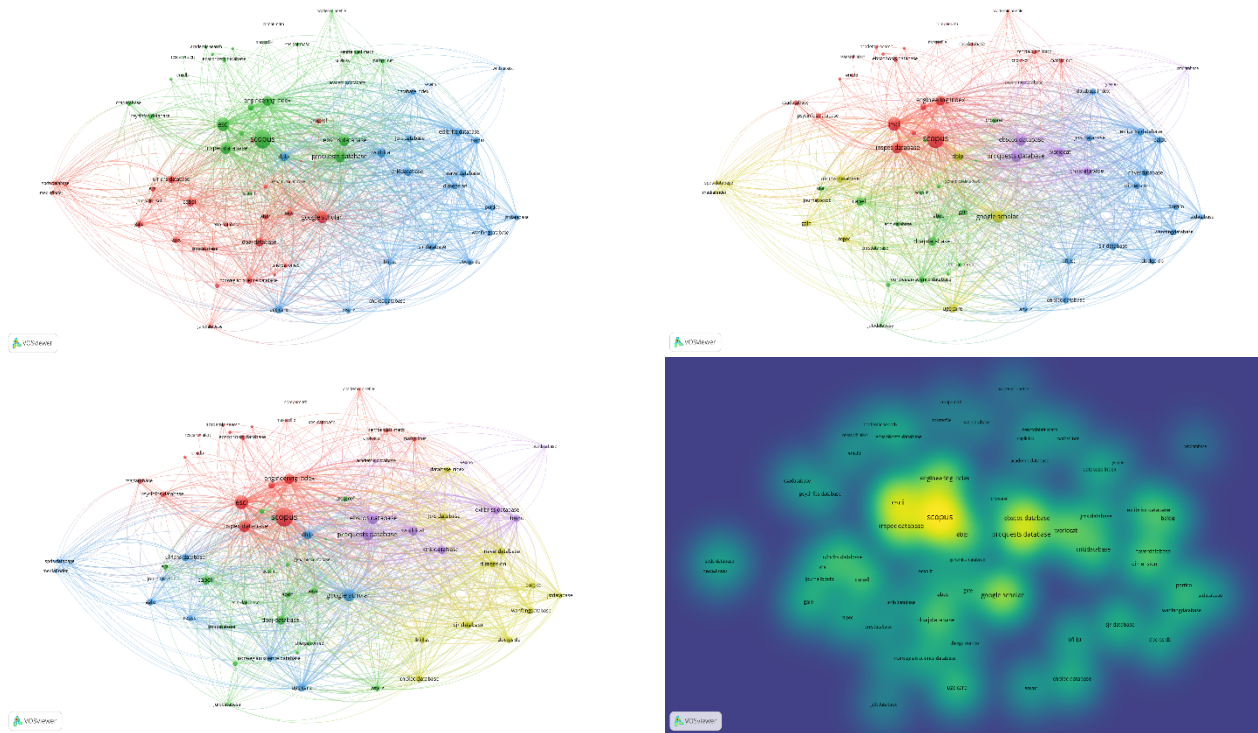


Figure 4. Index Clustering, Association and Density Analysis (N.T=23)  
(İndeks Kümeleme, Birliklilik ve Yoğunluk Analizi(T.S=23))

When Figure 4 is examined, an increase in the number of clusters is observed with the change in resolution values. The resolution value facilitates the convergence of clusters towards each other. As the resolution increases, the elements within the cluster become closer to one another. As the values within the cluster converge, values that are distant create another cluster, which leads to an increase in the number of clusters. When the results of the VosViewer analysis with a threshold value of 23 are examined;

When the resolution is set to 1, Scopus and ESCI are in the same cluster with 23 indexes. Upon examining their relationships, it is observed that they have strong associations with ProQuest, EBSCO, INSPEC, and Engineering Index. Although they share the same cluster with indexes like Abinform, PsycINFO, ACM Guide, EBSCOhost, CSA, Zentralblatt Math, MathSciNet, EconLit, CAS, Academic Search, ERIC, IBZ, Emerald, MasterFile, Premier, Academic OneFile, Research Alert, CNPLinker, and CompuMath, the relationships with these indexes are weaker. When the resolution is set to 1.15, Scopus and ESCI share the same cluster with 20 indexes. A strong relationship is observed with INSPEC and Engineering Index. Similar to the previous scenario, they are also in the same cluster with

Abinform, PsycINFO, ACM Guide, EBSCOhost, CSA, Zentralblatt Math, MathSciNet, CAS, Academic Search, ERIC, IBZ, Emerald, MasterFile, Premier, Academic OneFile, Research Alert, CNPLinker, and CompuMath, but the relationships with these indexes are weaker.

When the resolution is set to 1.30, Scopus and ESCI are in the same cluster with 19 indexes. As in the previous cases, there are strong relationships with INSPEC and Engineering Index. While they still share the same cluster with Abinform, PsycINFO, EBSCOhost, CSA, Zentralblatt Math, MathSciNet, CAS, Academic Search, ERIC, IBZ, Emerald, MasterFile, Premier, Academic OneFile, Research Alert, CNPLinker, and CompuMath, the relationships with these indexes remain weaker.

#### 4.2. Scenario 2 (Senaryo 2)

At this stage, the threshold value in the dataset is set to 53. A total of 51 indexes are included in the analysis. When the resolution value is set to 1, the data is divided into 2 different clusters. When the resolution value is set to 1.15, the data is divided into 4 different clusters. When the resolution value is set to 1.30, the data is divided into 5 different clusters. The cluster divisions are shown in Table 7.

Table 7. Clusters according to Scenario 2  
(Senaryo 2'ye göre kümeler)

| Resolution =1 |                    | Resolution =1.15 |          | Resolution =1.30   |               |
|---------------|--------------------|------------------|----------|--------------------|---------------|
| Clusters      | Number of Elements | Clusters         | Clusters | Number of Elements | Eleman Sayısı |
| Red           | 29                 | Red              | 18       | Red                | 17            |
| Green         | 22                 | Green            | 17       | Green              | 15            |
|               |                    | Blue             | 9        | Blue               | 9             |
|               |                    | Yellow           | 7        | Yellow             | 5             |
|               |                    |                  |          | Purple             | 5             |

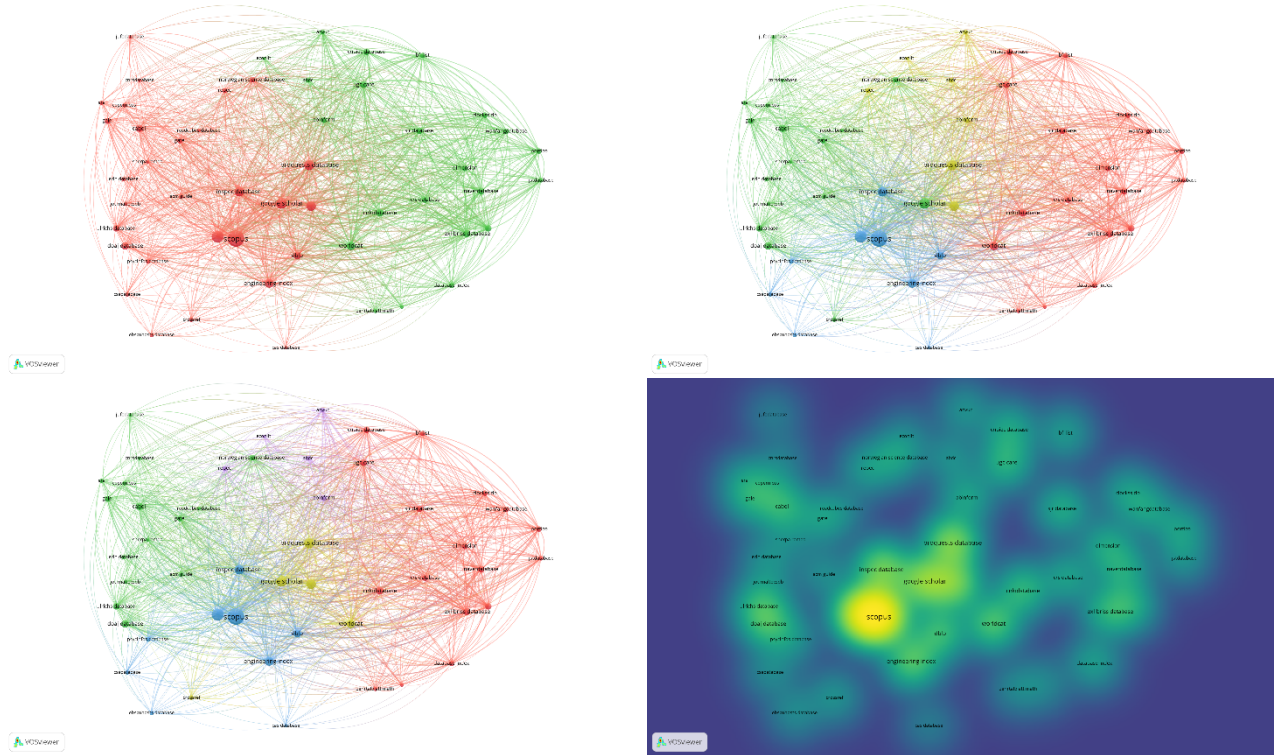


Figure 5. Index Clustering, Association and Density Analysis (N.T=53)  
(İndeks Kümeleme, Birliklilik ve Yoğunluk Analizi(T.S=53))

When the results of the VosViewer analysis with a threshold value of 53 are examined:

With a resolution value of 1, Scopus and ESCI are in the same cluster as 27 other indexes. When the relationship between them is analyzed, strong relationships are found with Google Scholar, ProQuest, EBSCO, INSPEC, Engineering Index, DBLP, and DOAJ. Although they share a cluster with Cabell, Ulrich, Gale, Norwegian Science, Gate, RePec, CrossRef, PsycINFO, ACM Guide, ERA, EBSCOhost, CSA, Copernicus, ERIH, JournalTOC, SHERPA Romeo, CAS, JUFO, ReadCube, and CNRS indexes, the relationships are weaker. With a resolution value of 1.15, Scopus and ESCI are in the same cluster as 7 other indexes. When the relationship between them is examined, strong connections are found with INSPEC, Engineering Index, and DBLP. While they share a cluster with PsycINFO, EBSCOhost, CSA, and CAS indexes, the relationships are weaker. With a

resolution value of 1.30, Scopus and ESCI are in the same cluster as 7 other indexes. The relationships between them and INSPEC, Engineering Index, and DBLP are still strong. However, the relationships with PsycINFO, EBSCOhost, CSA, and CAS indexes are weaker.

#### 4.3. Scenario 3 (Senaryo 3)

At this stage, the threshold value in the dataset has been set to 70. A total of 45 indexes are included in the analysis. When the resolution value is set to 1, the data is divided into 2 different clusters. When the resolution value is set to 1.15, the data is divided into 3 different clusters. When the resolution value is set to 1.30, the data is divided into 3 different clusters. The cluster divisions are shown in Table 8.

Tablo 8. Clusters according to Scenario 3  
(Senaryo 3'e göre kümeler)

| Resolution =1 |                    | Resolution =1.15 |                    | Resolution =1.30 |                    |
|---------------|--------------------|------------------|--------------------|------------------|--------------------|
| Clusters      | Number of Elements | Clusters         | Number of Elements | Clusters         | Number of Elements |
| Red           | 23                 | Red              | 21                 | Red              | 21                 |
| Green         | 22                 | Green            | 17                 | Green            | 15                 |
|               |                    | Blue             | 7                  | Blue             | 9                  |

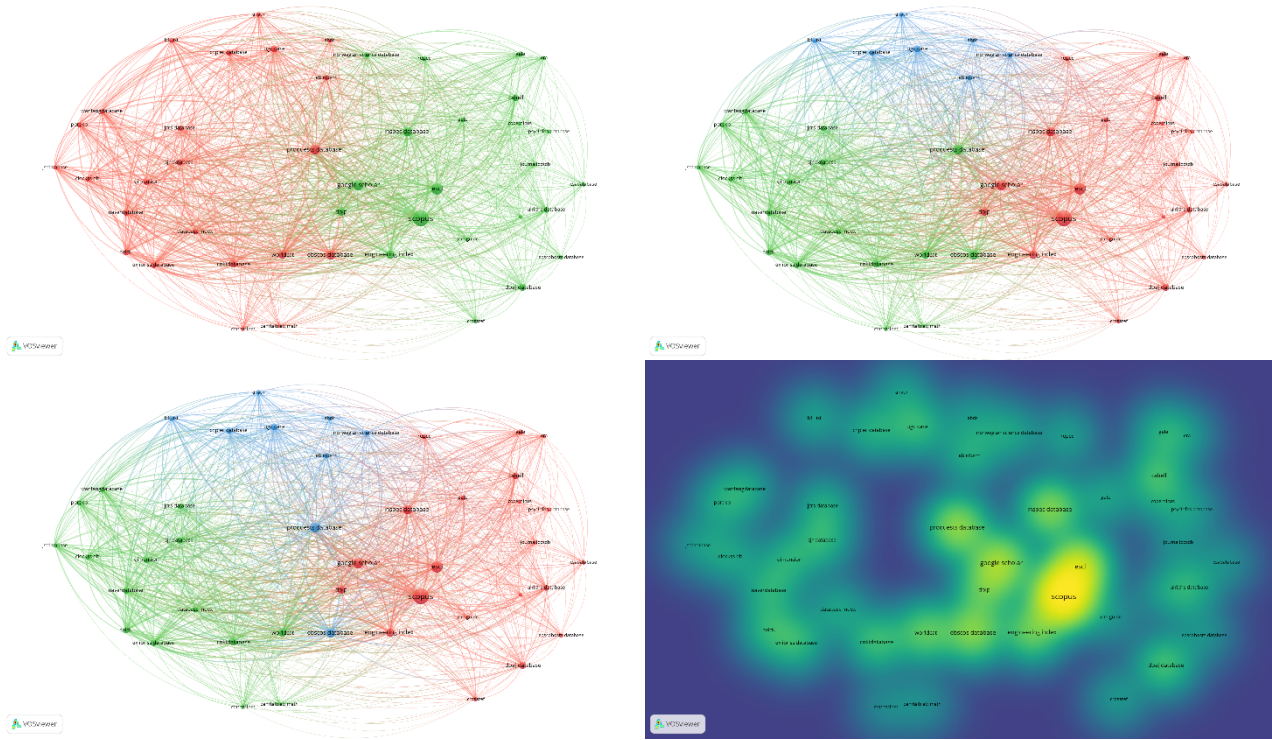


Figure 6. (Index Clustering, Association and Density Analysis (N.T=70))  
(İndeks Kümeleme, Birliklilik ve Yoğunluk Analizi(T.S=70))

When analyzing the results based on the threshold value of 70 in VosViewer:

With a resolution value of 1, Scopus and ESCI are placed in the same cluster with 20 indexes. Strong relationships are found with Google Scholar, Inspec, Engineering Index, DBLP, and DOAJ. Although Cabell, Ulrich, Gale, Norwegian Science, Gate, RePEc, Crossref, PsycInfo, ACM Guide, ERA, EBSCOhost, CSA, Copernicus, ERIH, and JournalTOC indexes are in the same cluster, the relationships between them are weaker. With a resolution value of 1.15, Scopus and ESCI are placed in the same cluster with 19 indexes. Strong relationships are found with Google Scholar, Inspec, Engineering Index, DBLP, and DOAJ. However, the relationships with Cabell, Ulrich, Gale, Gate, RePEc, Crossref, PsycInfo, ACM Guide, ERA, EBSCOhost, CSA, Copernicus, ERIH, and JournalTOC indexes are weaker. With a resolution value of 1.30,

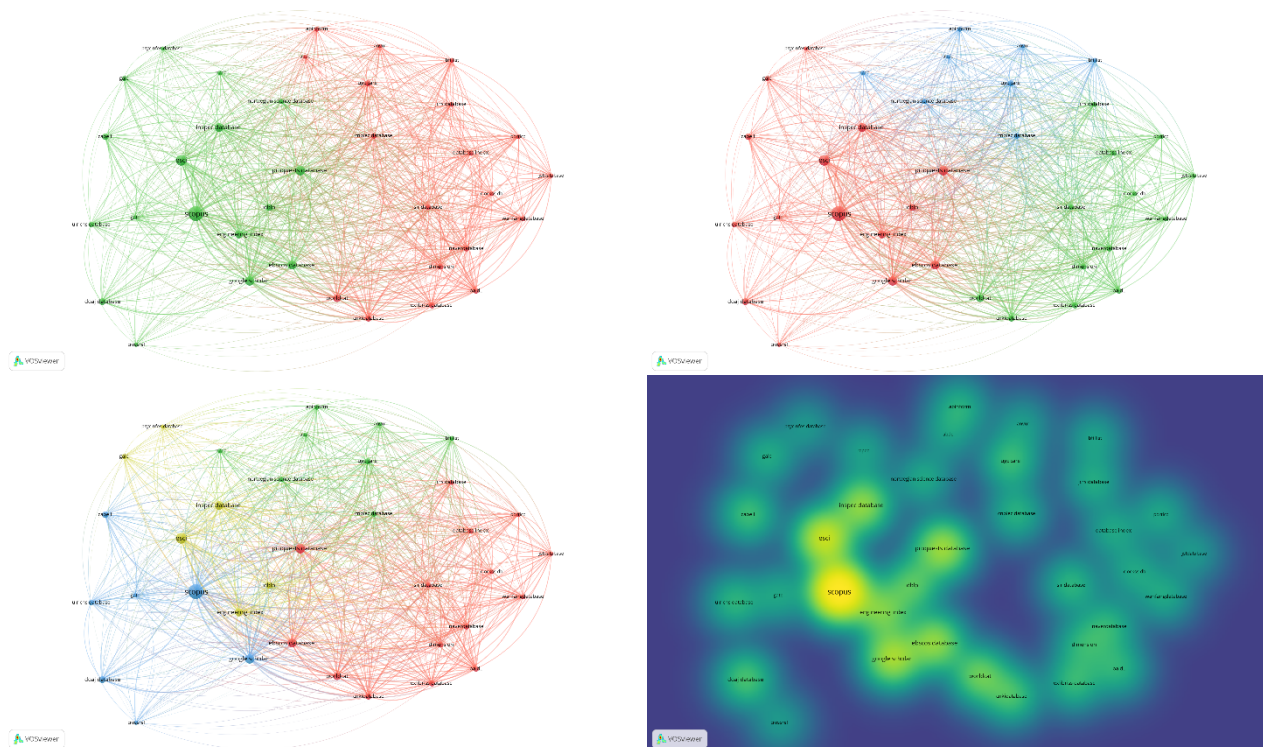
Scopus and ESCI are placed in the same cluster with 19 indexes. Strong relationships are again observed with Google Scholar, Inspec, Engineering Index, DBLP, and DOAJ. While Cabell, Ulrich, Gale, Gate, RePEc, Crossref, PsycInfo, ACM Guide, ERA, EBSCOhost, CSA, Copernicus, ERIH, and JournalTOC indexes are in the same cluster, their relationships remain weaker.

#### 4.4. Scenario 4 (Senaryo 4)

At this stage, the threshold value has been set to 100 in the dataset. There are 36 indexes in the analysis. When the resolution value is set to 1, the data is divided into 2 different clusters. With a resolution value of 1.15, the data is divided into 3 different clusters, and with a resolution value of 1.30, the data is divided into 4 different clusters. The cluster divisions are shown in Table 9.

Table 9. Clusters according to Scenario 4  
(Senaryo 4'e göre kümeler)

| Resolution =1 |                    | Resolution =1.15 |                    | Resolution =1.30 |                    |
|---------------|--------------------|------------------|--------------------|------------------|--------------------|
| Clusters      | Number of Elements | Clusters         | Number of Elements | Clusters         | Number of Elements |
| Red           | 19                 | Red              | 15                 | Red              | 15                 |
| Green         | 17                 | Green            | 13                 | Green            | 8                  |
|               |                    | Blue             | 8                  | Blue             | 7                  |
|               |                    |                  |                    | Yellow           | 6                  |

Figure 7. Index Clustering, Association and Density Analysis (N.T=100)  
(İndeks Kümeleme, Birliklilik ve Yoğunluk Analizi(T.S=100))

When the analysis results with a threshold value of 100 in VosViewer are examined:

With a resolution of 1, Scopus and ESCI are in the same cluster with 15 indexes. A strong relationship exists between Scopus and Google Scholar, ProQuest, EBSCO, Inspec, Engineering Index, DBLP, and DOAJ. While they are in the same cluster with Cabell, Ulrich, Gale, Norwegian Science, Gate, RePEc, Crossref, and PsycINFO indexes, the relationship between them is weaker. With a resolution of 1.15, Scopus and ESCI are in the same cluster with 13 indexes. A strong relationship exists between Scopus and Google Scholar, ProQuest, EBSCO, Inspec, Engineering Index, DBLP, and DOAJ. Although they are in the same cluster with Cabell, Ulrich, Gale, Gate, Crossref, and PsycINFO indexes, the relationship between them is weaker. With a resolution of 1.30, Scopus and ESCI are in different

clusters. Scopus is strongly related to Google Scholar and DOAJ. While it is in the same cluster with Cabell, Ulrich, Gate, and Crossref, the relationship is weaker. ESCI is strongly related to Inspec, Engineering Index, and DBLP. It shares the same cluster with Gale and PsycINFO indexes, although the relationship is weaker.

#### 4.5. Scenario 5 (Senaryo 5)

At this stage, the threshold value in the dataset has been set to 151. The dataset includes 26 indexes for analysis. When the resolution is set to 1, the data is divided into 2 different clusters. When the resolution is set to 1.15, the data is divided into 2 different clusters. When the resolution is set to 1.30, the data is divided into 2 different clusters. The cluster splits are shown in Table 10.

Table 10. Clusters according to Scenario 5  
(Senaryo 5'e göre kümeler)

| Resolution =1 |                    | Resolution =1.15 |                    | Resolution =1.30 |                    |
|---------------|--------------------|------------------|--------------------|------------------|--------------------|
| Clusters      | Number of Elements | Clusters         | Number of Elements | Clusters         | Number of Elements |
| Red           | 14                 | Red              | 13                 | Red              | 13                 |
| Green         | 12                 | Green            | 13                 | Green            | 13                 |

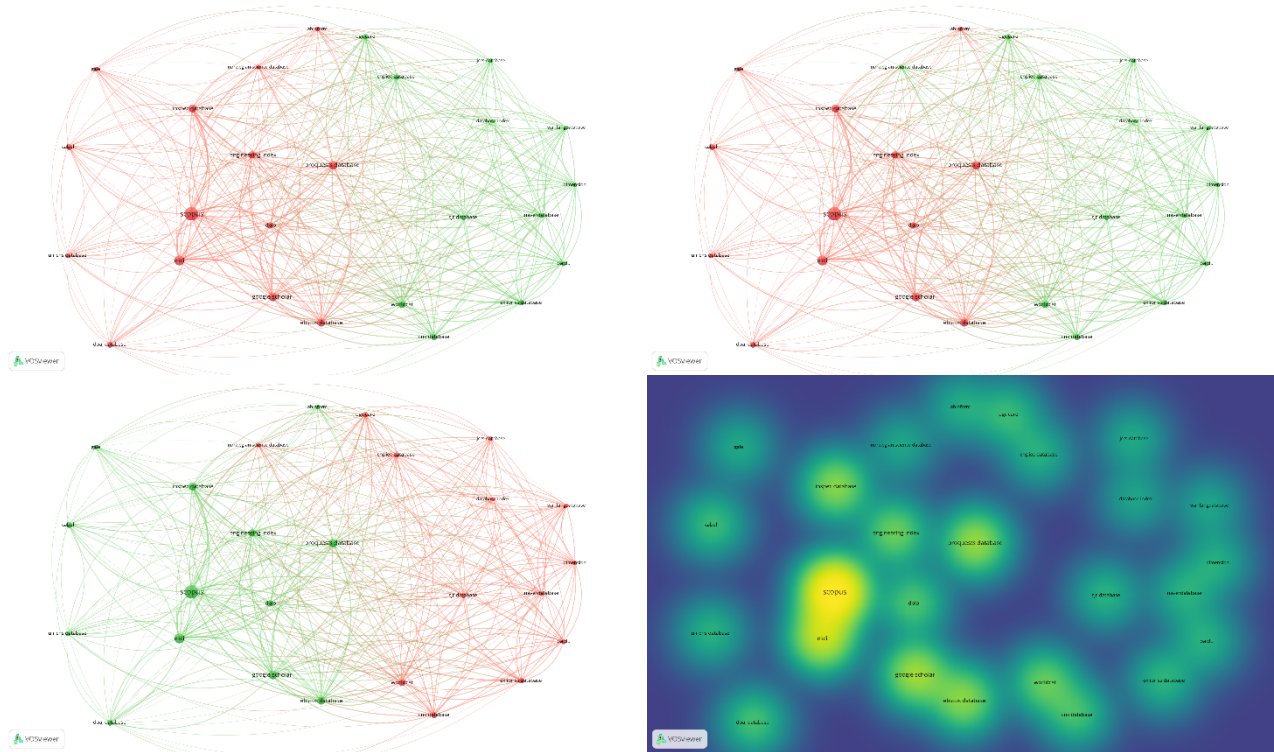


Figure 8. Index Clustering, Association and Density Analysis (N.T=151)  
(İndeks Kümeleme, Birliklilik ve Yoğunluk Analizi(T.S=151))

The results of the VosViewer analysis with a threshold value of 151 are as follows:

When the resolution is set to 1, Scopus and ESCI are placed in the same cluster with 12 indexes. Upon examining their relationships, strong connections are observed with Google Scholar, Proquest, EBSCO, Inspec, Engineering Index, DBLP, and DOAJ. Although they share a cluster with Cabell, Ulrichs, Abinform, Gale, and Norwegian Science indexes, the relationship with these indexes is weaker. When the resolution is set to 1.15, Scopus and ESCI are placed in the same cluster

with 11 indexes. A strong relationship is again observed with Google Scholar, Proquest, EBSCO, Inspec, Engineering Index, DBLP, and DOAJ, while a weaker relationship exists with Cabell, Ulrichs, Abinform, and Gale indexes. When the resolution is set to 1.30, Scopus and ESCI are still placed in the same cluster with 11 indexes. Strong relationships are found with Google Scholar, Proquest, EBSCO, Inspec, Engineering Index, DBLP, and DOAJ, whereas weaker relationships are observed with Cabell, Ulrichs, Abinform, and Gale indexes.

Table 11. Cluster and association analysis results  
(Kümeleme ve birliktelik analiz sonuçları)

| No | Terim               | T.S  | O=23-S=73 |        |        | O=53-S=51 |        |        | O=70-S=45 |        |        | O=100-S=36 |        |        | O=151-S=26 |        |        |
|----|---------------------|------|-----------|--------|--------|-----------|--------|--------|-----------|--------|--------|------------|--------|--------|------------|--------|--------|
|    |                     |      | R=1       | R=1.15 | R=1.30 | R=1       | R=1.15 | R=1.30 | R=1       | R=1.15 | R=1.30 | R=1        | R=1.15 | R=1.30 | R=1        | R=1.15 | R=1.30 |
| 4  | academic onefile    | 35   | Green     | Red    | Red    |           |        |        |           |        |        |            |        |        |            |        |        |
| 5  | academic search     | 50   | Green     | Red    | Red    |           |        |        |           |        |        |            |        |        |            |        |        |
| 13 | cas database        | 57   | Green     | Red    | Red    | Red       | Blue   | Blue   |           |        |        |            |        |        |            |        |        |
| 17 | cnplinker           | 28   | Green     | Red    | Red    |           |        |        |           |        |        |            |        |        |            |        |        |
| 19 | compumath           | 22   | Green     | Red    | Red    |           |        |        |           |        |        |            |        |        |            |        |        |
| 22 | csadatabase         | 90   | Green     | Red    | Red    | Red       | Blue   | Blue   | Green     | Red    | Red    |            |        |        |            |        |        |
| 27 | ebscohosts database | 94   | Green     | Red    | Red    | Red       | Blue   | Blue   | Green     | Red    | Red    |            |        |        |            |        |        |
| 30 | emerald             | 42   | Green     | Red    | Red    |           |        |        |           |        |        |            |        |        |            |        |        |
| 31 | engineering index   | 378  | Green     | Red    | Red    | Red       | Blue   | Blue   | Green     | Red    | Red    | Green      | Red    | Yellow | Red        | Red    | Green  |
| 33 | ericdb              | 45   | Green     | Red    | Red    |           |        |        |           |        |        |            |        |        |            |        |        |
| 35 | ESCI                | 681  | Green     | Red    | Red    | Red       | Blue   | Blue   | Green     | Red    | Red    | Green      | Red    | Yellow | Red        | Red    | Green  |
| 42 | ibzdatabase         | 43   | Green     | Red    | Red    |           |        |        |           |        |        |            |        |        |            |        |        |
| 44 | inspec database     | 464  | Green     | Red    | Red    | Red       | Blue   | Blue   | Green     | Red    | Red    | Green      | Red    | Yellow | Red        | Red    | Green  |
| 49 | masterfile          | 39   | Green     | Red    | Red    |           |        |        |           |        |        |            |        |        |            |        |        |
| 55 | premier database    | 38   | Green     | Red    | Red    |           |        |        |           |        |        |            |        |        |            |        |        |
| 57 | psycinfos database  | 105  | Green     | Red    | Red    | Red       | Blue   | Blue   | Green     | Red    | Red    | Green      | Red    | Yellow |            |        |        |
| 60 | research alert      | 29   | Green     | Red    | Red    |           |        |        |           |        |        |            |        |        |            |        |        |
| 62 | Scopus              | 1284 | Green     | Red    | Red    | Red       | Blue   | Blue   | Green     | Red    | Red    | Green      | Red    | Blue   | Red        | Red    | Green  |

#### 4.6. Interpretation of the Findings *(Bulguların Yorumlanması)*

According to Scenario 1, the threshold value was set to 23, and three different resolution settings were examined: 1, 1.15, and 1.30. When the resolution was set to 1, 3 different clusters were identified. When the resolution was set to 1.15, 5 different clusters were identified. When the resolution was set to 1.30, 5 different clusters were identified. For both resolution settings of 1.15 and 1.30, no significant differences were observed in the relationships between the Scopus and ESCI indexes and other indexes. The relationships remained consistent across these two resolutions.

According to Scenario 2, the threshold value was set to 53, and three different resolution settings were examined: 1, 1.15, and 1.30. When the resolution was set to 1, 2 different clusters were identified. When the resolution was set to 1.15, 4 different clusters were identified. When the resolution was set to 1.30, 5 different clusters were identified. For both resolution settings of 1.15 and 1.30, no significant differences were observed in either the clustering of Scopus and ESCI indexes or in the relationships between these indexes and the other indexes. The relationships remained consistent across these two resolutions.

According to Scenario 3, the threshold value was set to 70, and three different resolution settings were examined: 1, 1.15, and 1.30. When the resolution was set to 1, 2 different clusters were identified. When the resolution was set to 1.15, 3 different clusters were identified. When the resolution was set to 1.30, 3 different clusters were identified. Although changing the resolution caused slight variations in the number of elements within the clusters, there was no significant difference in the relationships between the indexes. The relationships between the indexes remained consistent regardless of the resolution setting.

According to Scenario 4, the threshold value was set to 100, and three different resolution settings were examined: 1, 1.15, and 1.30. When the resolution was set to 1, 2 different clusters were identified. When the resolution was set to 1.15, 3 different clusters were identified. When the resolution was set to 1.30, 4 different clusters were identified. No changes occurred when the resolution was set to 1 and 1.15. However, when the resolution was set to 1.30, the ESCI and Scopus indexes were separated and placed in different clusters, indicating a significant shift in their relationship compared to the lower resolution settings.

According to Scenario 5, the threshold value was set to 151, and three different resolution settings were examined: 1, 1.15, and 1.30. When the resolution was set to 1, 2 different clusters were identified. When the resolution was set to 1.15, 2 different clusters were identified. When the resolution was set to 1.30, 2 different clusters were identified. In Scenario 5, changing the resolution value did not affect the number

of clusters. Additionally, there was no difference in the clusters or the relationships between the indexes when the resolution was set to 1.15 and 1.30. The analysis results remained consistent across these resolution settings.

Due to the proximity of elements within the clusters and their distance from elements in other clusters, changes in the resolution value do not result in changes in the number of clusters. The similarities observed across different scenarios are due to the closeness of the cluster elements. The divergence observed in Scenario 4 is due to the increased distance between Scopus and ESCI at a resolution of 1.30, causing them to remain closer to other indices.

When evaluating the overall situation, as shown in Table 11, the indices related to ESCI and Scopus in the five different scenarios were the Engineering Index and IETinspec. The primary reasons for the similarity between these indices are as follows;

Table 12. Scopus Categories  
(Scopus Kategorileri)

|    |                                                                                                                                                                                                           |
|----|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| 1  | Management Information Systems                                                                                                                                                                            |
| 2  | Human-Computer Interaction                                                                                                                                                                                |
| 3  | Information Systems; Information Systems and Management                                                                                                                                                   |
| 4  | Management of Technology and Innovation                                                                                                                                                                   |
| 5  | Artificial Intelligence; Artificial Intelligence                                                                                                                                                          |
| 6  | Media Technology; Computer Networks and Communications                                                                                                                                                    |
| 7  | Computer Graphics and Computer-Aided Design                                                                                                                                                               |
| 8  | Theoretical Computer Science; Computer Science Applications                                                                                                                                               |
| 9  | General Computer Science)ve ESCI(Computer Science                                                                                                                                                         |
| 10 | Information Systems                                                                                                                                                                                       |
| 11 | Computer Science, Interdisciplinary Applications   Computer Science, Artificial Intelligence   Computer Science, Cybernetics   Computer Science, Information Systems                                      |
| 12 | Computer Science, Artificial Intelligence   Computer Science, Information Systems; Computer Science, Theory & Methods   Computer Science, Artificial Intelligence   Computer Science, Information Systems |
| 13 | Computer Science, Interdisciplinary Applications   Computer Science, Artificial Intelligence   Computer Science, Information Systems                                                                      |
| 14 | Computer Science, Artificial Intelligence; Computer Science, Theory & Methods   Computer Science, Artificial Intelligence                                                                                 |
| 15 | Computer Science, Artificial Intelligence   Computer Science, Cybernetics   Computer Science, Information Systems; Computer Science, Theory & Methods   Multidisciplinary Sciences                        |

|    |                                                                                                                                                                                                                  |
|----|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| 16 | Multidisciplinary Sciences   Computer Science, Information Systems                                                                                                                                               |
| 17 | Computer Science, Interdisciplinary Applications   Computer Science, Theory & Methods   Engineering, Electrical & Electronic   Computer Science, Hardware & Architecture   Computer Science, Information Systems |
| 18 | Computer Science, Interdisciplinary Applications   Multidisciplinary Sciences   Computer Science, Information Systems                                                                                            |
| 19 | Computer Science, Interdisciplinary Applications; Computer Science, Interdisciplinary Applications   Computer Science, Software Engineering   Computer Science, Information Systems; Business                    |
| 20 | Computer Science, Interdisciplinary Applications   Engineering, Multidisciplinary   Materials Science, Multidisciplinary                                                                                         |
| 21 | Computer Science, Interdisciplinary Applications   Computer Science, Theory & Methods   Computer Science, Hardware & Architecture   Computer Science, Information Systems                                        |

As a result of expert opinions, the similarity between these two indices was due to the fact that the categories obtained within Scopus, as shown in Table 12, are primarily related to engineering and technology fields.

When examining the scope of the IETinspec index, it covers topics in engineering and technology fields such as physics, electrical and electronic engineering, computer and information technology, mechanical and manufacturing engineering, robotics and automation, telecommunications and communication systems, energy and environmental technologies, nanotechnology, and interdisciplinary studies. Looking at the scope of the Engineering index, it includes research in the fields of engineering, technology, and applied sciences, covering topics such as mechanical engineering, civil and structural engineering, electrical and electronic engineering, computer and software engineering, materials engineering, chemical and process engineering, environmental engineering, industrial engineering, biomedical engineering, and energy and power systems. In addition, the Engineering index emphasizes interdisciplinary studies alongside traditional engineering topics.

The strong relationship between the Scopus and ESCI indices and the IETinspec and Engineering Index databases stems from the fact that the scope of these two indices is related to the fields of engineering and technology, which are aligned with the information systems (YBS) domain. Specifically, the strong relationship between Scopus and ESCI and the Engineering Index is due to the fact that the Engineering Index places importance on interdisciplinary studies, which further strengthens the connection.

## 5. CONCLUSION AND RECOMMENDATIONS (SONUÇ VE ÖNERİLER)

The importance of the Information and Business Systems (YBS) discipline, which has significantly increased with digitalization, is particularly notable. As YBS is closely related to many fields and includes technology, a key requirement of the modern age, its significance has grown in recent years. Especially with the COVID-19 pandemic, the importance of YBS in remote work processes has been emphasized once again. Recently, YBS has become a frequently studied field by researchers, primarily due to its multidisciplinary nature.

In this study, the current status of international field indexes in the area of Information and Business Systems (YBS) was analyzed, and the potential for expanding these indexes was evaluated. First, the existing international indexes used in the YBS field were identified, and the categories listed in these indexes were determined. The identified categories were scored on a scale of 1-5 based on expert opinions, and geometric averages were calculated. Based on the results obtained, categories with a geometric average above 4 were included in the study. In the next phase of the study, journals within the identified categories and the indexes that indexed these journals were analyzed. A total of 3,819 journals were identified. After cleaning the repeated data, a total of 1,473 journals were reached, and the indexes that scanned these journals were identified. As a result of the examination, a total of 73 different indexes were found. Bibliometric analysis methods were used to understand the relationships between these indexes and to group them. In the analyses performed with the help of the VosViewer software, different scenarios were applied to examine the clustering and connections between the indexes.

The analysis results revealed that, in addition to Scopus and ESCI indexes, especially the Engineering Index and Inspec indexes were grouped together and showed a strong relationship in the 15 different analysis results within the 5 different scenarios. It was determined that the scope information of these two indexes showed similarities with Scopus and ESCI, and it was also observed that the journal acceptance and publication criteria largely overlapped. Based on the findings, it was concluded that, in addition to ESCI and Scopus, which are considered international field indexes for the YBS discipline, the inclusion of the Engineering Index and Inspec indexes would be appropriate. It is anticipated that this expansion would contribute to the development of the discipline by increasing the international visibility of academic publications in the YBS field.

International field indexes are a challenging subject for researchers to understand and make decisions about, yet they hold significant importance in the academic publishing process. In addition to the difficulties in preparing a publication, selecting the journal and the index where the work will be published is also an

important decision-making process. This decision becomes even more critical, especially for research that will be evaluated under academic appointment and promotion criteria.

The starting point of this study is the examination of international field indexes defined by the Higher Education Council (YÖK). Before 2024, there was no clear distinction between these indexes, but in the post-2024 period, it became apparent that the ESCI and Scopus indexes were separated from other indexes. The study investigated the aspects in which these indexes differ and whether there are alternative indexes with similar characteristics. As a result of the analysis, it was revealed that there are other international indexes with similar qualities to ESCI and Scopus. When reviewing the criteria published by YÖK, it was found that similar international index regulations also exist in fundamental fields such as education sciences, natural sciences and mathematics, philology, fine arts, law, theology, architecture, planning and design, engineering, health sciences, agriculture, forestry and aquatic products, and sports sciences. In this context, the study has revealed the existence of new international indexes that can be considered in addition to ESCI and Scopus for the field of social and human sciences.

Based on the findings, it is anticipated that this study conducted in the field of social and human sciences could be applied to other fundamental fields in the future, thereby expanding the scope of international indexes. This would allow researchers to not only rely on ESCI and Scopus but also consider alternative indexes, enabling them to manage their publication processes more efficiently. Increasing the number of indexes would not only promote academic productivity but also contribute to the acceleration of the publication process.

## REFERENCES (KAYNAKLAR)

- [1]Schoderbeck, P. Kefalas,G. &Charles, C.(1975). Schoderbeck, Management Systems: Conceptual Considerations.Dallas, Texas.
- [2]Stoner, J. A. F. (1982). Management. Second Edition, New Jersey: Prentice-Hall, Englewood Cliffs.
- [3]Culnan, M. J. (1987). Mapping The Intellectual Structure Of MIS, 1980-1985: A Co-Citation Analysis. *Mis Quarterly*, 341-353.
- [4]Long, L. (1989). Management Information Systems (International Edition). NJ: Prentice Hall.
- [5]Adeoti-Adekeye, W. B. (1997). The Importance Of Management Information Systems. *Library Review*, 46(5), 318-327.
- [6]Lee, A. S. (2001). Editor's Comments: Research In Information Systems: What We Haven't Learned. *Mis Quarterly*, 25(4), V.
- [7]Baskerville, R. L., & Myers, M. D. (2002). Information Systems As A Reference Discipline. *MIS Quarterly*, 1-14.
- [8]BENSGHİR Kaya, T. (2002). Reflections on the Development of the Management Information Systems Discipline in Turkey. *Journal of Public Administration*, 35(1), 77-103.
- [9]Laudon, K. C., & Laudon, J. P. (2004). *Management information systems: Managing the digital firm*. Pearson Educación.
- [10]Becta (2005). School Management Information Systems And Value For Money. Coventry: Becta. [Online] Available: <http://www.egovmonitor.com/reports/rep12009.pdf>
- [11]O'Brien, J., & Marakas, G. (2007). Management Information Systems. McGraw Hill Irwin.
- [12]Gökçen H. (2011). Management Information Systems: Analysis and Design (Second Edition). Ankara: Afşar Printing.
- [13]Pratap, A. (2018). Management Information System: Role, Characteristics And Advantages. Retrieved July 11, 2019, <https://notesmatic.com/management-information-system-role-characteristics-and-advantages/>
- [14]Barki, H, Rivard, S, Talbot, J 1993 A Keyword Classification Scheme For IS Research Literature: An Update. *MIS Quarterly* 17 2 209–226 .
- [15]Seo, E. G., & Han, I. G. (1997). Bibliometric Analysis On MIS Research. *Asia Pacific Journal Of Information Systems*, 7(3), 145-165.
- [16]Cocosila, M., Serenko, A., & Turel, O. (2011). Exploring The Management Information Systems Discipline: A Scientometric Study Of ICIS, PACIS And ASAC. *Scientometrics*, 87(1), 1-16.
- [17]Mohanty, B. (2014). Management Information Systems Quarterly (MISQ): A Bibliometric Study. *Library Philosophy And Practice*, 0\_1.
- [18]Yarlıtaş, S. (2015). An Analysis of the Current State of the Management Information Systems Discipline in Turkey. *Journal Of Higher Education & Science*, 5(2).
- [19]Lin, A. J., Hsu, C. L., & Chiang, C. H. (2016). RETRACTED ARTICLE: Bibliometric study of Electronic Commerce Research in Information Systems & MIS Journals. *Scientometrics*, 109(3), 1455-1476.

- [20] Özköse, H. (2017). Bibliometric Analysis and Mapping of the Management Information Systems Field in Turkey and the World. Unpublished Doctoral Dissertation, Gazi University, Institute of Informatics, Department of Management Information Systems, Ankara.
- [21] Özköse, H., & Gencer, C. T. (2017). Bibliometric Analysis And Mapping Of Management Information Systems Field. *Gazi University Journal Of Science*, 30(4), 356-371.
- [22] Beydoun, G., Abedin, B., Merigó, J. M., & Vera, M. (2019). Twenty Years Of Information Systems Frontiers. *Information Systems Frontiers*, 21, 485-494.
- [23] Coşkun, E., Özdağoğlu, G., Damar, M., & Çallı, B. (2019). Scientometrics-Based Study Of Computer Science And Information Systems Research Community Macro Level Profiles.
- [24] Aytaç, A. M. (2020). *Trends in Academic Studies in the Field of Management Information Systems in Turkey (Master's Thesis, Ufuk University)*.
- [25] Jeyaraj, A., & Zadeh, A. (2020). Institutional Isomorphism In Organizational Cybersecurity: A Text Analytics Approach. *Journal Of Organizational Computing And Electronic Commerce*, 30(4), 361-380. <https://doi.org/10.1080/10919392.2020.1776033>
- [26] Damar, M., & Aydın, Ö. (2021). Turkey's Status in International Q1 Journals in the Field of Management Information Systems After 2010. *İzmir Journal of Economics*, 36(4), 811-842.
- [27] Cebeci, H. İ. (2021). Multi-Criteria Decision-Making Techniques in the Management Information Systems Literature: A Systematic Review. *Journal of Business Science*, 9(1), 111-147.
- [28] Damar, M., & Özdağoğlu, G. (2022). Forty Years Of Management Information Systems From The Window Of MIS Quarterly. *Acta Infologica*, 6(1), 99-125.
- [29] Özköse, H., Ozyurt, O., & Ayaz, A. (2023). Management Information Systems Research: A Topic Modeling Based Bibliometric Analysis. *Journal Of Computer Information Systems*, 63(5), 1166-1182.
- [30] Pritchard, A. (1969). *Statistical Bibliography; An Interim Bibliography*.
- [31] Ellegaard, O., & Wallin, J. A. (2015). The Bibliometric Analysis Of Scholarly Production: How Great Is The Impact?. *Scientometrics*, 105, 1809-1831.
- [32] García-Sánchez, P., Mora, A. M., Castillo, P. A., & Pérez, I. J. (2019). A Bibliometric Study Of The Research Area Of Videogames Using Dimensions. *Ai Database. Procedia Computer Science*, 162, 737-744.
- [33] Stremersch, S., Verniers, I., & Verhoef, P. C. (2007). "The Quest For Citations: Drivers Of Article Impact". *Journal Of Marketing*, 71(3), 171-193.
- [34] Rossetto, D.E., Bernardes, R.C., Borini, F.M. *Et Al*. Structure And Evolution Of Innovation Research In The Last 60 Years: Review And Future Trends In The Field Of Business Through The Citations And Co-Citations Analysis. *Scientometrics* 115, 1329-1363 (2018). <https://doi.org/10.1007/S11192-018-2709-7>
- [35] Donthu, N., Kumar, S., Mukherjee, D., Pandey, N. & Lim, W. M. (2021). How To Conduct A Bibliometric Analysis: An Overview And Guidelines. *Journal Of Business Research*, 133, 285-296.
- [36] Odenkirk, M. T., Horman, B. M., Dodds, J. N., Patisaul, H. B., & Baker, E. S. (2021). Combining Micropunch Histology And Multidimensional Lipidomic Measurements For In-Depth Tissue Mapping. *ACS Measurement Science Au*, 2(1), 67-75.
- [37] Gurcan, F., Cagiltay, N. E., & Cagiltay, K. (2021). "Mapping Human-Computer Interaction Research Themes And Trends From Its Existence To Today: A Topic Modeling-Based Review Of Past 60 Years". *International Journal Of Human Computer Interaction*, 37(3), 267-280.
- [38] Albayrak, M. (2008). *Detection of Epileptiform Activity in EEG Signals Using Data Mining Process (Doctoral Dissertation, Sakarya University, Turkey)*.
- [39] Sertçelik, Ş., & Önder, E. (2023). Examination of Academic Research Areas in the Scope of Management Information Systems: An Analysis Using the Apriori Algorithm. *Gümüşhane University Journal of Social Sciences*, 14(2), 680-690. <https://doi.org/10.36362/Gumus.1254125>
- [40] Özkoç, E. E. (2019, Nisan). Bibliometric Structure of Management Information Systems Studies in Turkey and Its Analysis Using Social Network Analysis. *YDÜ SOSBİLDİR*, 12(1),
- [41] Shiau, W. L., Chen, S. Y., & Tsai, Y. C. (2015). Management Information Systems Issues: Co-Citation Analysis Of Journal Articles. *International Journal Of Electronic Commerce Studies*, 6(1), 145-162.
- [42] Hajek, P., & Abedin, M. Z. (2020). A Profit Function-Maximizing Inventory Backorder Prediction System Using Big Data Analytics. *IEEE Access*, 8, 58982-58994.
- [43] Güler, G., & Zeren, D. (2024). Metaverse: Bibliometric Analysis of National Literature, *Akademik Hassasiyetler* 11(24), 599-623. <https://doi.org/10.58884/akademik-hassasiyetler.1430354>

[44] Kırmızııkaya, A.. & Öztürk, A. (2024). Bibliometric Analysis of Studies on Social Media Marketing Activities. *Gümüşhane University Journal of Social Sciences* 15(3), 830-844.

[45] Ünal Kestane, S. (2024). Bibliometric Analysis of Studies on the Concept of CRM (Customer Relationship Management) Using VosViewer. *Journal of the Institute of Social Sciences, Çukurova University*, 33(1), 311-329. <https://doi.org/10.35379/cusosbil.1396955>.

[46] Hidayat, EY , Hastuti, K ve Muda, AK (2025). Artificial Intelligence In Digital Image Processing: A Bibliometric Analysis, *Intelligent Systems With Applications*, 25(1).

[47] Aryawati, N. P. A., Triyuwono, I., Roekhudin, R., & Mardiaty, E. (2024). Bibliometric analysis on Scopus database related internal control in university: a mapping landscape. *Cogent Business & Management*, 11(1).

<https://doi.org/10.1080/23311975.2024.2422566>.

[48] Modina, M., Fedele, M. and Formisano, A.V. (2024), "Digital finance for SMEs and startups: a bibliometric analysis and future research direction", *Journal of Small Business and Enterprise Development*, Vol. ahead-of-print No. ahead-of-print. <https://doi.org/10.1108/JSBED-01-2024-0021>.



## COVID-19'un Yayılım Modelini Tahmin Etmek İçin Hibrit GA-ConvLSTM Modeli: řanghay İşbirlięi Örgütü İçin Bir Vaka Çalışması



Anıl UTKU\*,a

<sup>a,\*</sup> Munzur Üniversitesi, Mühendislik Fakültesi, Bilgisayar Mühendislięi Bölümü, TUNCELİ, 62000, TÜRKİYE

### MAKALE BİLGİSİ

Alınma: 20.12.2024  
Kabul: 29.01.2025

#### Anahtar Kelimeler:

Derin öğrenme,  
řangay İşbirlięi Örgütü,  
CNN,  
LSTM,  
Genetik algoritma,  
COVID-19

#### \*Sorumlu Yazar

e-posta:  
anilutku@munzur.edu.tr

### ÖZET

COVID-19 dünyanın hemen her yerine çok hızlı bir şekilde yayılarak birçok insanın ciddi semptomlar yaşamasına ve hayatını kaybetmesine neden olmuřtur. Bu çalışmada, saęlık sistemleri üzerindeki yükü hafifletmek ve salgının dağılımını tahmin etmek için planlar yapılabilmesi amacıyla derin öğrenme yöntemleriyle COVID-19'un yayılım örüntüsünü belirlemek amaçlandı. Bu amaçla CNN ve LSTM modelleri kullanılarak geliştirilen hibrit derin öğrenme modelinin hiper parametreleri genetik algoritma ile optimize edilerek daha başarılı bir tahmin performansı elde edilmiştir. GA-ConvLSTM modeli, SCO üyesi ülkelerde salgının yayılımını belirlemek için XGBoost, SVM, CNN, MLP, LSTM ve ConvLSTM ile test edilmiştir. Çalışmada, WHO tarafından sunulan 2020/01/03 ile 2022/05/31 tarihleri arasındaki günlük COVID-19 vaka ve ölüm verileri kullanılmıştır. Deneyler, GA-ConvLSTM'nin tüm ülkeler için vaka tahmininde 0,9'un üzerinde  $R^2$  değerine sahip olduğunu göstermiştir. Deneyler, GA-ConvLSTM'nin ölüm tahmininde ülkelerin çoęunluęu için 0,9'un üzerinde  $R^2$ 'ye sahip olduğunu göstermiştir. Ayrıca, COVID-19'un SCO ülkeleri arasındaki yayılım örüntüsü, 5 ve 14 günlük kuluçka dönemleri kullanılarak oluşturulan akor diyagramlarıyla belirlenmiştir.

DOI: 10.59940/jismar.1604942

## Hybrid GA-ConvLSTM Model for Predicting the Transmission Pattern of COVID-19: A Case Study for Shanghai Cooperation Organisation

### ARTICLE INFO

Received: 20.12.2024  
Accepted: 29.01.2025

#### Keywords:

Deep learning,  
řanghai Cooperation  
Organization,  
CNN,  
LSTM,  
Genetic algorithm,  
COVID-19

#### \*Corresponding Authors

e-mail:  
anilutku@munzur.edu.tr

### ABSTRACT

COVID-19 has spread very quickly to almost every part of the world, causing many people to experience severe symptoms and lose their lives. In this study, it is aimed to determine the transmission pattern of COVID-19 with deep learning methods so that plans can be made to alleviate the burden on healthcare systems and predict the distribution of the epidemic. For this purpose, the hyper-parameters of the hybrid deep learning model developed using CNN and LSTM models were optimized with genetic algorithm, and a more successful prediction performance was achieved. The GA-ConvLSTM model was tested with XGBoost, SVM, CNN, MLP, LSTM, and ConvLSTM to determine the spread of the epidemic in the member countries of SCO. The study used daily COVID-19 case and death data between 2020/01/03 and 2022/05/31 presented by WHO. Experiments showed that GA-ConvLSTM has over 0.9  $R^2$  in case prediction for all countries. Experiments showed that GA-ConvLSTM has above 0.9  $R^2$  for the majority of countries when it comes to death prediction. In addition, the transmission pattern of COVID-19 among the SCO countries was determined with the chord diagrams created using 5 and 14 days' incubation periods.

DOI: 10.59940/jismar.1604942

## 1. INTRODUCTION (GİRİŞ)

Coronavirus Disease 2019 (COVID-19) has high fever and respiratory symptoms that can feel like a cold, flu, or pneumonia. [1]. Symptoms of COVID-19 are high fever, shortness of breath, muscle, joint, and body aches, weakness, and severe cough [2]. The incubation period of coronavirus disease varies between 5-14 days [3]. These symptoms may be very mild or severe. The disease can be severe in those who pose a risk for COVID-19 and are most likely to be affected, especially the elderly, those with cancer or immune-suppressing diseases, and people with lung diseases [4]. It can be said that almost all of those who lost their lives due to COVID-19 had different underlying diseases.

COVID-19 has caused many people to be admitted to intensive care units or even die [5]. During this period, the health systems of most countries remained inadequate. Health personnel were insufficient in number, intensive care units were at capacity, and morgues were full [6]. For this reason, artificial intelligence emerges to make future predictions in epidemic management. Using artificial intelligence methods, it is essential to analyze epidemics such as COVID-19, predict epidemics' spread, and develop strategies to combat the epidemic. The inferences obtained can be used to ensure that countries' health systems can control the epidemic and prevent its spread.

Shanghai Cooperation Organization (SCO), whose main area of cooperation is security, is a regional cooperation organization [7]. SCO is a security-based cooperation organization founded in 1996 with the cooperation of China, Kazakhstan, Russia, Tajikistan and Kyrgyzstan [8]. The SCO, which Uzbekistan joined in 2001, India and Pakistan in 2017, and Iran in 2023, aims to ensure neighborly relations and trust among its member countries and to establish security and peace on a regional basis [7, 8].

In this study, we aimed to develop preventive strategies by determining the spread of pandemics. In this way, countries' health systems are optimized, the workload of healthcare professionals is lightened, and strategies are developed to prevent the spread of pandemics. The hyper-parameters of the developed ConvLSTM model were optimized with the GA to create the GA-ConvLSTM hybrid model. GA-ConvLSTM was tested with Convolutional Neural Network (CNN), Multilayer Perceptron (MLP), Extreme Gradient Boosting (XGBoost), ConvLSTM, Long-Short Term Memory (LSTM) and Support Vector Machine (SVM) using a dataset of daily COVID-19 deaths and cases provided by World

Health Organization (WHO). Additionally, experimental studies were conducted to determine the transmission pattern of COVID-19 among SCO member countries and chord graphs were created according to 5 and 14-day incubation times.

The innovations of this study are as follows:

- There is no study in the literature determining the distribution of COVID-19 in SCO member countries.
- GA-ConvLSTM was developed by optimizing the hyper-parameters of the ConvLSTM hybrid model using the GA.
- GA-ConvLSTM was compared in detail with SVM, CNN, XGBoost, SVM, MLP, and LSTM and the ConvLSTM hybrid model.
- To identify the transmission pattern of COVID-19 among SCO member countries, detailed analyzes were conducted for incubation periods.

## 2. RELATED WORKS (İLİŞKİLİ ÇALIŞMALAR)

Artificial intelligence is effectively used to predict the spread pattern of pandemics and diagnose diseases. Artificial intelligence techniques analyze big data and identify complex relationships and patterns between data, enabling planning and developing strategies for pandemics. This section examines studies in the literature on COVID-19 and deep learning.

Shahid et al. presented an evaluation of Autoregressive Integrated Moving Average (ARIMA), Bidirectional LSTM (Bi-LSTM), SVM, and LSTM for COVID-19 case prediction [10]. COVID-19 data from 10 countries were used as the dataset for approximately five months. Experiments showed that the Bi-LSTM model is more effective than the other models.

Abbasimehr and Paki presented a comparative analysis in which the parameters of CNN, multi-head attention, and LSTM were optimized with Bayes for predicting COVID-19 cases [11]. The applied models were compared with fuzzy fractals. Two different datasets consisting of 10 countries were used: long-term and short-term. Experiments showed that LSTM is more effective than the compared models in most countries.

Rauf et al. presented an evaluation of the Gated Recurrent Unit (GRU), Recurrent Neural Network (RNN), and LSTM to predict the spread of coronavirus [12]. Experiments showed that LSTM outperformed other models for the four compared countries.

Zhou et al. developed an application of Bi-LSTM, GRU, LSTM, and Dense-LSTM to determine COVID-19 deaths and cases [13]. Data from 12 countries until June 2022 were used as the dataset. Experiments showed that Dense-LSTM outperformed the other models.

Ukwuoma et al. presented an analysis of VGG-16, DenseNet, and InceptionV3 for coronavirus disease prediction from chest images [14]. A dataset consisting of approximately 21000 images was used in the study. Experiments showed that the DenseNet201 model outperformed the benchmark models with 96% and 98% accuracy in different test scenarios.

Al-Rashedi and Al-Hagery presented an application of LSTM, ARIMA, and CNN to determine the transmission pattern of COVID-19 [15]. Three different time intervals were used for predicting cases, recoveries, and deaths: long, medium, and short. Experiments showed that LSTM, in particular, and then CNN were quite successful.

Solayman et al. developed an application of ensemble models, k-Nearest Neighbor (kNN), CNN, Decision Tree (DT), LSTM, Logistic Regression (LR), SVM, and Random Forest (RF) for COVID-19 detection [16]. The dataset was preprocessed using synthetic oversampling. Experiments showed that the hybrid CNN-LSTM outperformed the benchmark models with over 96% accuracy.

### 3. MATERIAL AND METHOD (MATERİYAL VE METOT)

SCO is an important regional and political cooperation organization. The model presented in this study may be effective for SCO member countries to combat epidemics that burden healthcare systems, such as COVID-19. It can contribute to developing strategies for epidemic management and developing cooperation between member countries.

#### 3.1. Dataset (Veriseti)

This study aimed to develop a perspective for SCO member countries using the daily COVID-19 information presented by WHO. The dataset used includes the period between 2020/01/03 and 2022/05/31 when COVID-19 peaked. The dataset consists of the attributes date, region code determined by WHO, country name, country code, number of daily cases, total number of cases, total number of deaths, and daily deaths. For SCO member countries, the attributes of date and daily case and death numbers were selected according to country codes. Fig. 1 shows the daily COVID-19 cases for SCO member countries.

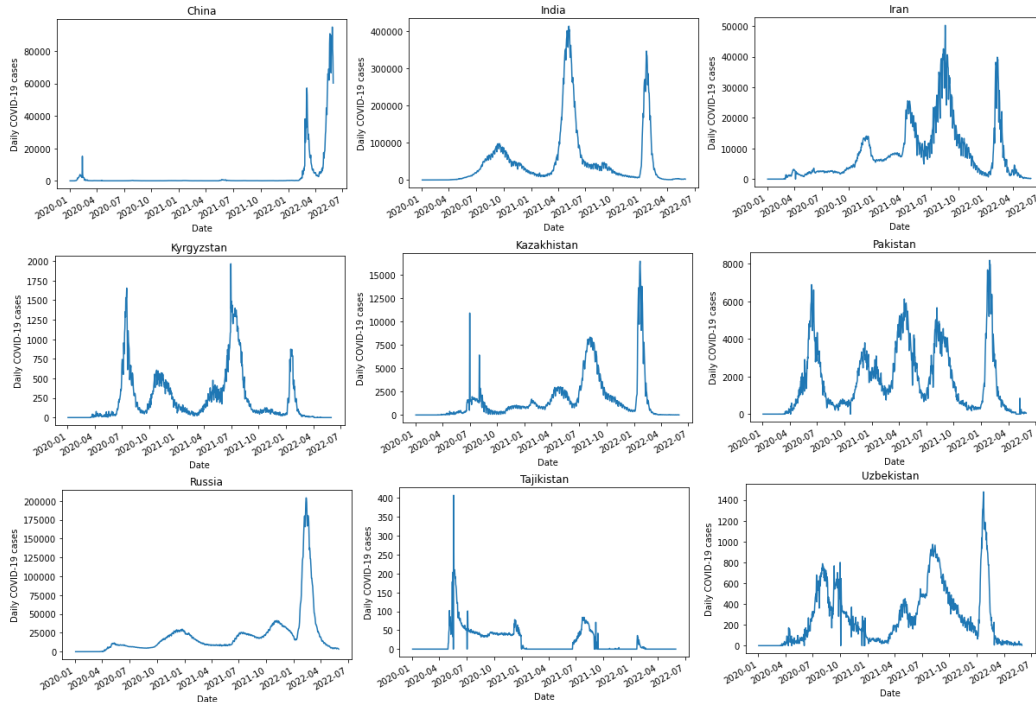


Figure 1. The daily COVID-19 cases for SCO member countries (*SCO üyesi ülkelerdeki günlük COVID-19 vaka sayıları*)

Fig. 2 shows the daily COVID-19 deaths for SCO member countries.

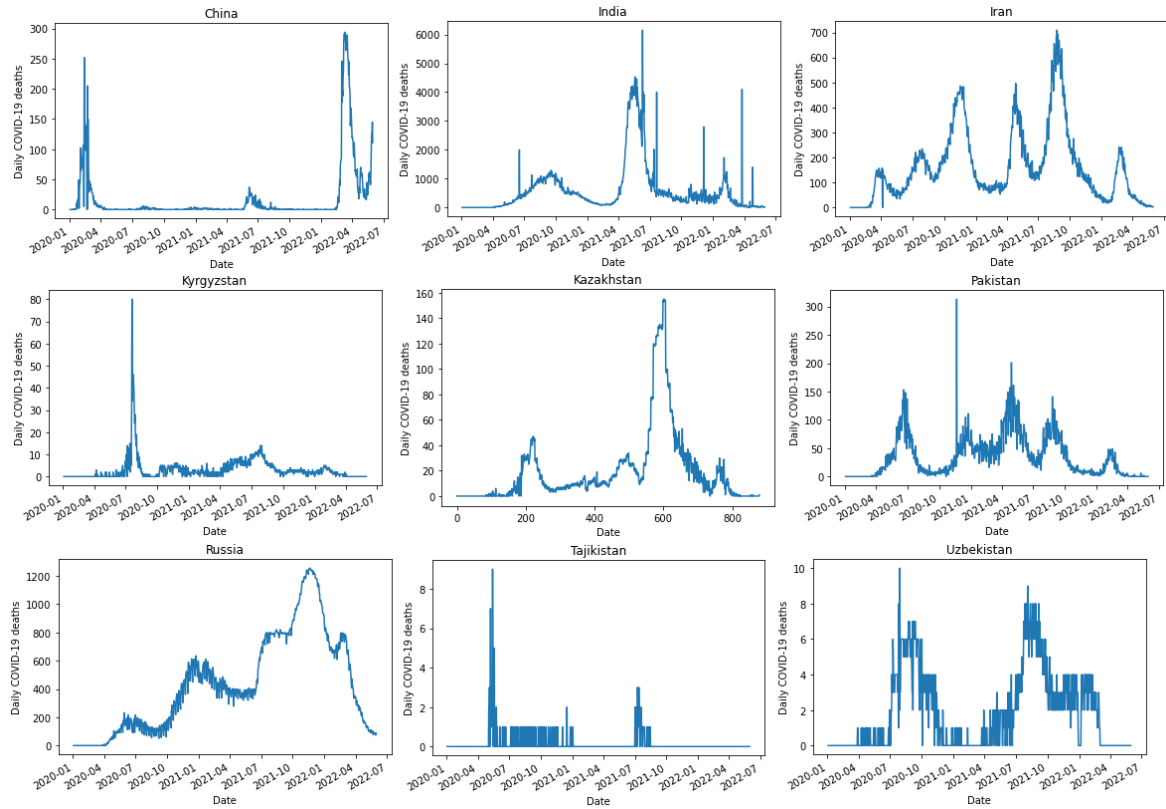


Figure 2. The daily COVID-19 deaths for SCO member countries (*SCO üyesi ülkelerdeki günlük COVID-19 ölüm sayıları*)

Fig. 3 indicates the total cases in the SCO member countries.

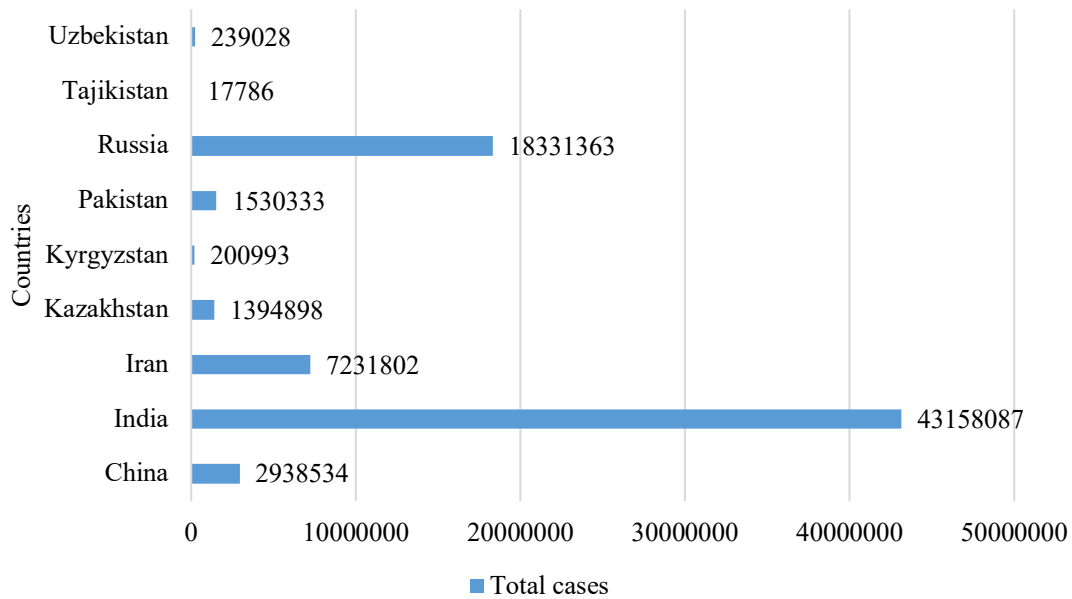


Figure 3. The total cases in the SCO member countries (*SCO üyesi ülkelerdeki toplam COVID-19 vaka sayıları*)

The total cases in the India is 43,158,087. After India, Russia, China and Iran have higher total cases. Fig. 4 shows the total deaths in the SCO member countries.

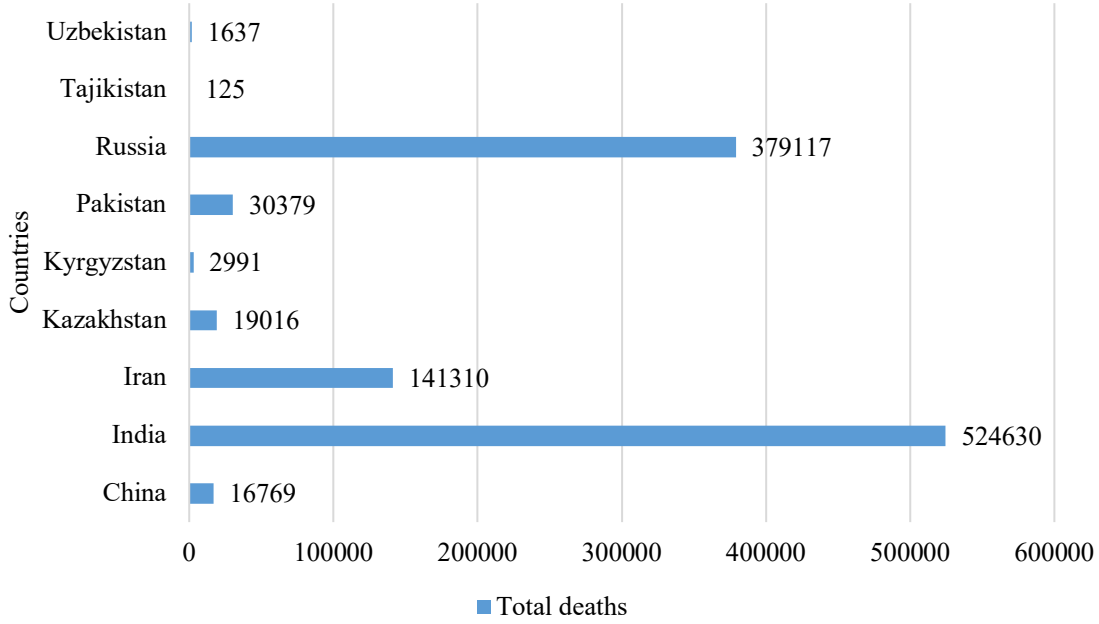


Figure 4. The cumulative number of deaths in the SCO member countries (*SCO üyesi ülkelerdeki toplam COVID-19 ölüm sayıları*)

The total deaths in India is 524,630. After India, Iran and Russia have maximum death counts.

Table 1 indicates the information of the deaths and cases in the SCO member countries.

Table 1. The information of the deaths and cases in the SCO member countries (*SCO üyesi ülkelerdeki ölüm ve vaka bilgileri*)

| Country    | Maximum case count | Maximum death count | Total cases | Total deaths |
|------------|--------------------|---------------------|-------------|--------------|
| China      | 94753              | 294                 | 2938534     | 16769        |
| India      | 414188             | 6148                | 43158087    | 524630       |
| Iran       | 50228              | 709                 | 7231802     | 141310       |
| Kazakhstan | 16442              | 155                 | 1394898     | 19016        |
| Kyrgyzstan | 1965               | 80                  | 200993      | 2991         |
| Pakistan   | 8183               | 313                 | 1530333     | 30379        |
| Russia     | 203949             | 1254                | 18331363    | 379117       |
| Tajikistan | 407                | 9                   | 17786       | 125          |
| Uzbekistan | 1478               | 10                  | 239028      | 1637         |

As seen in Table 1, India and Russia have the highest number of cases and deaths.

### 3.2. Data Pre-processing (*Veri Ön-işleme*)

The dataset used consists of daily COVID-19 time series data. For this reason, it is necessary to transform the dataset into a regression problem. The sliding window method was used for this purpose. This method ensures that the given input is set to the

specified window size, and the data point in the next time step is set as output [17]. For instance, if the

window size is 3,  $t_1$ ,  $t_2$  and,  $t_3$  will be the input, and  $t_4$  will be the output, as shown in Fig. 5.

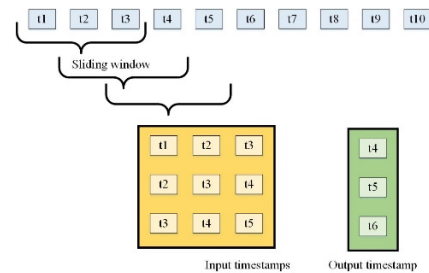


Figure 5. Transform the dataset into regression problem (*Verisetinin regresyon problemine dönüştürülmesi*)

Experimental studies showed that the lowest error rates were obtained when the sliding window size was 5. After transforming the data into the regression problem structure, it was scaled using MinMaxScaler. 33% of the dataset was used for testing and 67% for training. 10% of the train data was used for hyper-parameter optimization. For each compared model to achieve the most successful result, the hyper-parameters of the models were adjusted using grid search.

### 3.3. Prediction Models (Tahmin Modelleri)

XGBoost combines the prediction results of multiple decision tree predictors [18]. It creates new models by correcting errors in the models it creates until the training data is trained correctly [19]. Thanks to its ability to handle missing values, XGBoost can run without requiring data pre-processing. Additionally, it can work quickly on large data sets thanks to its parallel processing ability.

SVM determines a hyperplane that separates different classes in multidimensional space so that the margin between them is maximum [20]. Using the resulting hyperplane, data samples are classified according to their locations. SVM uses support vectors to maximize the margin between classes. Kernel functions are used to determine the decision boundary for the hyperplane [21].

MLP is a model inspired by the human brain's information-processing ability [22]. It has hidden layers consisting of interconnected neurons, except for input and output layers [23]. Hidden layers enable complex relationships to be learned in the dataset.

MLP updates the parameters used in the model by calculating the errors between the outputs obtained using backpropagation and the actual values [24].

CNN, generally used in image processing problems, consists of input and output layers, a pooling layer, a convolution layer, and a flatten layer [25]. The input layer receives the data and passes it to the convolution layer. The convolution layer extracts features from the input data using filters [26]. CNN determines patterns in the data by applying convolution along the temporal dimension in time series problems [27]. The flattened layer converts multidimensional feature maps into a 1D vector. Fully connected layers make predictions by learning representations of features [28]. The output layer also presents the prediction.

LSTM is a model that allows long-term dependencies in sequential data to be remembered [29]. LSTM uses gates and memory cells to control the flow of information, enabling information to be selectively forgotten or remembered. LSTM's forget gate decides which information to discard from the memory cell [30]. Thanks to memory cells, LSTM stores information from previous time steps. The outputs of LSTM cells are transferred to the next cell, allowing consecutive data to be processed [31].

### 3.4. Proposed Model (Önerilen Model)

The developed GA-ConvLSTM model enables the hyper-parameters of the ConvLSTM model to be optimized using GA. Proposed GA-ConvLSTM model architecture is seen in Fig. 6.

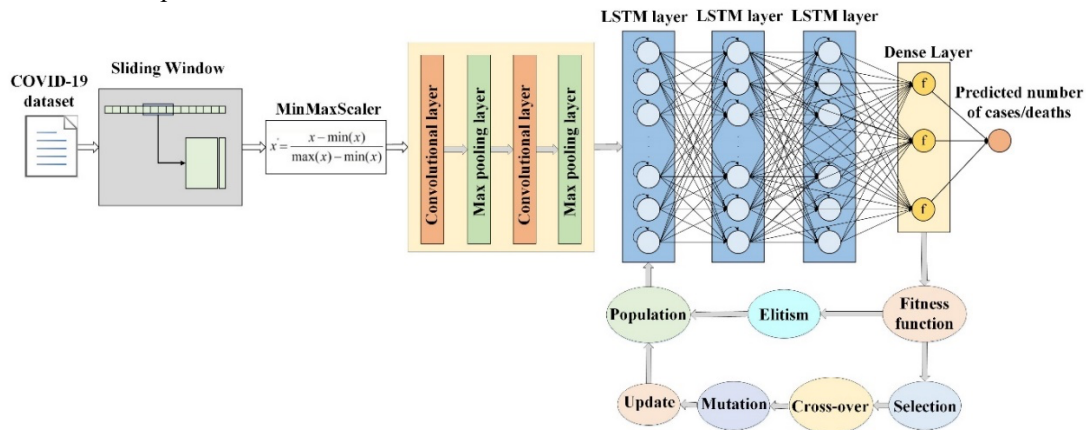


Figure 6. Proposed GA-ConvLSTM model architecture (Önerilen GA-ConvLSTM model mimarisi)

CNN is responsible for extracting features in time series data and learning the relationships and patterns between the data. LSTM enables increasing prediction performance by learning long- and short-term

dependencies between data. GA uses different hyper-parameters to find combinations with the lowest error rate. GA evolves the model to achieve lower error rates with each iteration using population-based search. GA speeds up the optimization process by

making fewer attempts than Grid Search. GA enables the evaluation of parameter combinations simultaneously by processing different hyper-parameters in parallel. Additionally, GA can determine parameter ranges in more detail than Grid Search. Grid Search works with specified fixed value ranges.

GA-ConvLSTM takes daily COVID-19 cases and the death numbers of SCO member countries as input. After the data is transformed into a regression problem structure using a sliding window, it is scaled using the MinMax scaler. The CNN component uses convolution layers to discover patterns in the data and perform feature extraction. The LSTM component models complex relationships in the data by processing feature maps sent from the CNN. In this way, LSTM learns the data's long and short-term dependencies. GA is used to optimize the hyper-parameters of CNN-LSTM. The hyper-parameter combinations are evaluated as an individual, and selection, crossover, and mutation are performed between individuals. The fitness function enables the determination of hyper-parameters with the lowest error value. When it comes to the GA-ConvLSTM model, the hyper-parameters optimized with GA play a crucial role. For the CNN component, the number of convolutional layers is set at 2, with an activation

function of ReLU, a pooling size of 2, and a number of filters at 64. The LSTM component, on the other hand, has 128 neurons, uses the Adam optimizer, has a dropout rate of 0.2, a batch size of 64, and runs for 50 epochs. As for GA itself, the crossover probability is set at 0.8, the population size at 50, the mutation probability at 0.1, and the number of generations at 100.

#### 4. EXPERIMENTAL RESULTS (DENEYSEL SONUÇLAR)

In this study, the transmission pattern of COVID-19 in SCO member countries was predicted, and chord diagrams were created for 5- and 14-day periods to determine the transmission of the epidemic among SCO member countries.

##### 4.1. Prediction of the Spread Pattern of COVID-19 in SCO Member Countries (SCO Üye Ülkelerinde COVID-19'un Yayılma Örüntüsünün Tahmini)

GA-ConvLSTM was comprehensively checked with CNN, XGBoost, SVM, LSTM, MLP, and CNN-LSTM according to Root Mean Squared Error (RMSE), R-Squared ( $R^2$ ), and Mean Absolute Error (MAE). Table 2 indicates the case prediction outcomes of RMSE in SCO member countries.

Table 2. The case prediction outcomes according to RMSE in SCO member countries (SCO üye ülkelerinde RMSE'ye göre vaka tahmin sonuçları)

| Country    | XGBoost  | SVM     | MLP     | CNN     | LSTM    | ConvLSTM | GA-ConvLSTM |
|------------|----------|---------|---------|---------|---------|----------|-------------|
| China      | 6477.20  | 3948.72 | 3542.48 | 3751.11 | 3433.66 | 3284.30  | 3018.18     |
| India      | 8470.77  | 8286.26 | 7673.31 | 8170.56 | 7515.76 | 7169.42  | 5086.96     |
| Iran       | 2847.86  | 2821.38 | 2180.75 | 2301.22 | 2113.95 | 1937.19  | 1814.97     |
| Kazakhstan | 1084.52  | 751.46  | 573.25  | 576.60  | 568.15  | 506.54   | 290.32      |
| Kyrgyzstan | 43.44    | 43.34   | 41.86   | 44.09   | 40.99   | 39.41    | 27.13       |
| Pakistan   | 387.46   | 383.17  | 379.19  | 382.98  | 365.04  | 342.89   | 240.18      |
| Russia     | 10245.13 | 4238.41 | 3654.05 | 3691.76 | 3602.80 | 3290.55  | 3074.55     |
| Tajikistan | 7.14     | 7.09    | 5.03    | 5.98    | 4.94    | 4.12     | 3.93        |
| Uzbekistan | 65.44    | 65.07   | 49.62   | 65.10   | 49.25   | 48.07    | 44.66       |

As seen in Table 2, GA-ConvLSTM is more successful than the other models according to the RMSE. After GA-ConvLSTM, ConvLSTM, LSTM, MLP, CNN, SVM, and XGBoost were successful, respectively. As seen in Table 2, RMSEs are low for Kyrgyzstan, Tajikistan, and Uzbekistan, where the

number of cases is low. However, RMSEs are also high for China, India, and Russia, where the number of cases is high. Table 3 indicates the case prediction outcomes of MAE in SCO member countries.

Table 3. The case prediction outcomes in terms of MAE in SCO member countries (SCO üye ülkelerinde MAE'ye göre vaka tahmin sonuçları)

| Country    | XGBoost | SVM     | MLP     | CNN     | LSTM    | ConvLSTM | GA-ConvLSTM |
|------------|---------|---------|---------|---------|---------|----------|-------------|
| China      | 2538.34 | 1705.76 | 1598.44 | 1547.50 | 1234.45 | 1146.30  | 1113.96     |
| India      | 4349.47 | 3771.15 | 3349.16 | 3607.85 | 3039.46 | 2984.93  | 2079.43     |
| Iran       | 1535.99 | 1561.28 | 1197.70 | 1229.09 | 1251.79 | 1088.78  | 1021.40     |
| Kazakhstan | 447.46  | 338.98  | 252.38  | 258.74  | 234.55  | 214.33   | 145.16      |
| Kyrgyzstan | 20.53   | 19.47   | 18.17   | 22.57   | 18.30   | 17.97    | 12.33       |
| Pakistan   | 221.48  | 218.10  | 217.99  | 225.87  | 204.98  | 194.49   | 139.23      |

|            |         |         |         |         |         |         |         |
|------------|---------|---------|---------|---------|---------|---------|---------|
| Russia     | 5130.96 | 2136.44 | 1751.11 | 2269.28 | 1858.14 | 1627.02 | 1501.16 |
| Tajikistan | 4.59    | 3.07    | 1.58    | 2.05    | 1.65    | 1.43    | 1.07    |
| Uzbekistan | 41.55   | 41.52   | 32.98   | 45.83   | 32.04   | 31.12   | 28.35   |

As seen in Table 3, GA-ConvLSTM is more successful than the other models, according to the MAE. After GA-ConvLSTM, ConvLSTM, LSTM, MLP, CNN, SVM, and XGBoost were successful, respectively. As seen in Table 3, MAE values are low for Kyrgyzstan, Tajikistan, and Uzbekistan, where the

number of cases is low. However, MAE values are also high for China, India, Iran, and Russia, where the number of cases is high. Fig. 7 and Table 4 show the case prediction results according to  $R^2$  in SCO member countries.

Table 4. The case prediction outcomes according to  $R^2$  in SCO member countries (*SCO üye ülkelerinde  $R^2$ 'ye göre vaka tahmin sonuçları*)

| Country    | XGBoost | SVM   | MLP   | CNN   | LSTM  | ConvLSTM | GA-ConvLSTM |
|------------|---------|-------|-------|-------|-------|----------|-------------|
| China      | 0.902   | 0.964 | 0.971 | 0.967 | 0.972 | 0.974    | 0.976       |
| India      | 0.986   | 0.986 | 0.988 | 0.986 | 0.988 | 0.989    | 0.994       |
| Iran       | 0.926   | 0.927 | 0.956 | 0.951 | 0.959 | 0.965    | 0.967       |
| Kazakhstan | 0.889   | 0.946 | 0.969 | 0.968 | 0.969 | 0.975    | 0.995       |
| Kyrgyzstan | 0.926   | 0.926 | 0.931 | 0.924 | 0.934 | 0.939    | 0.956       |
| Pakistan   | 0.957   | 0.958 | 0.959 | 0.958 | 0.962 | 0.966    | 0.976       |
| Russia     | 0.948   | 0.991 | 0.993 | 0.993 | 0.993 | 0.994    | 0.995       |
| Tajikistan | 0.605   | 0.611 | 0.804 | 0.724 | 0.810 | 0.869    | 0.965       |
| Uzbekistan | 0.961   | 0.961 | 0.977 | 0.962 | 0.978 | 0.979    | 0.982       |

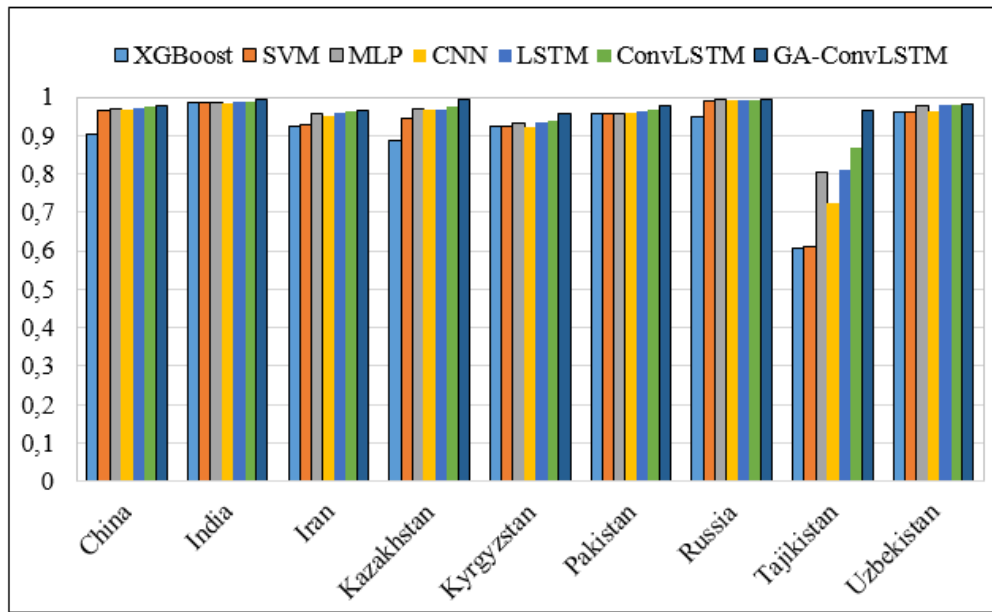


Figure 7. The case prediction results according to  $R^2$  in SCO member countries (*SCO üye ülkelerinde  $R^2$ 'ye göre vaka tahmin sonuçları*)

Table 4 and Fig. 7 show that all compared models have  $R^2$  above 0.9, except Tajikistan. For Tajikistan, GA-ConvLSTM has 0.965  $R^2$ , and ConvLSTM has

0.869  $R^2$ . Fig. 8 shows the prediction graphs of GA-ConvLSTM for predicting daily COVID-19 cases.

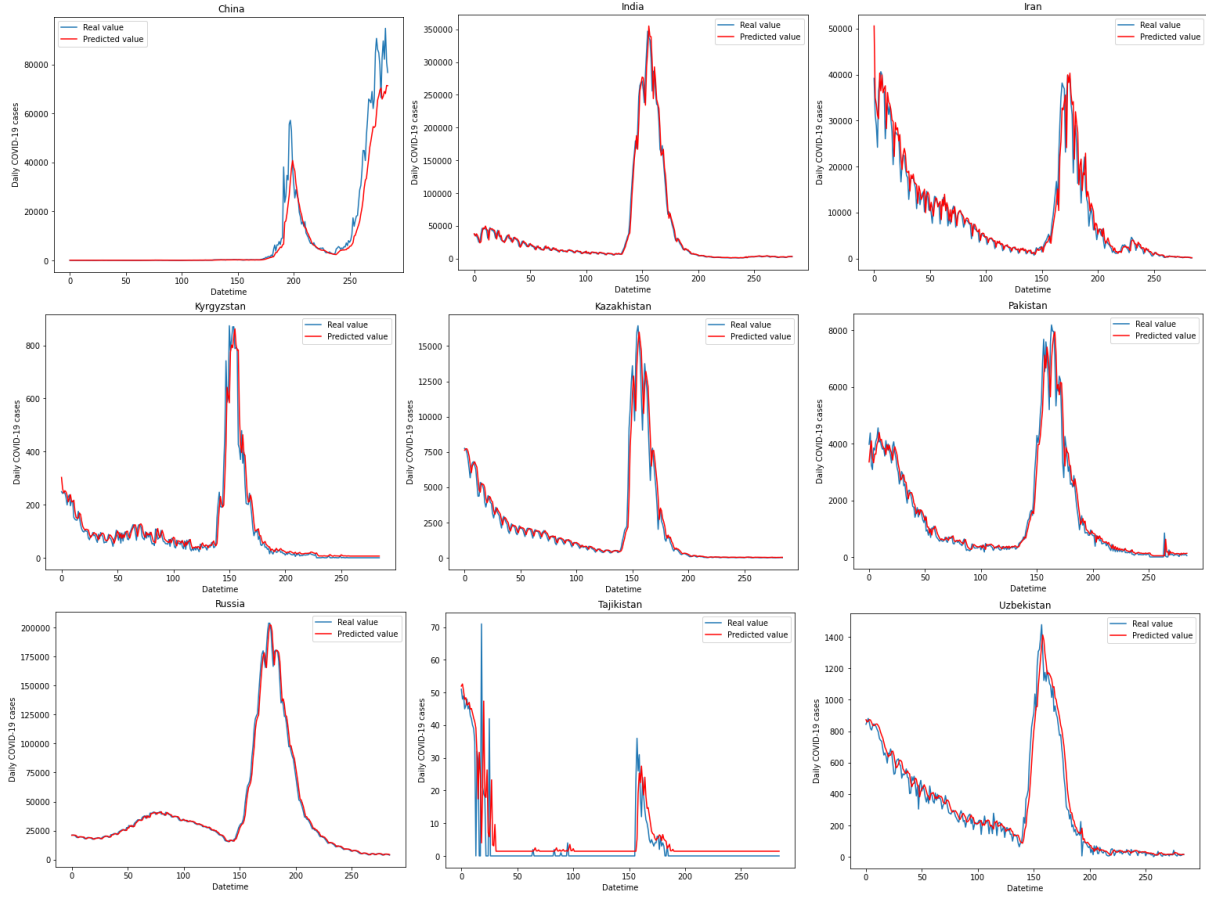


Figure 8. Prediction graphs of GA-ConvLSTM for predicting daily COVID-19 cases (*Günlük COVID-19 vakalarını tahmin etmek için GA-ConvLSTM tahmin grafikleri*)

As seen in Fig. 8, GA-ConvLSTM successfully modelled the changes in daily COVID-19 cases. Table

5 indicates the death prediction outcomes of RMSE in SCO member countries.

Table 5. The death prediction outcomes according to RMSE in SCO member countries (*SCO üye ülkelerinde RMSE'ye göre ölüm tahmini sonuçları*)

| Country    | XGBoost | SVM    | MLP    | CNN    | LSTM   | ConvLSTM | GA-ConvLSTM |
|------------|---------|--------|--------|--------|--------|----------|-------------|
| China      | 21.30   | 19.39  | 12.79  | 16.68  | 12.57  | 12.40    | 10.67       |
| India      | 350.54  | 313.29 | 312.02 | 314.36 | 311.91 | 311.14   | 242.15      |
| Iran       | 24.75   | 24.73  | 24.18  | 24.50  | 23.82  | 23.66    | 16.80       |
| Kazakhstan | 6.66    | 6.66   | 6.39   | 6.60   | 6.20   | 5.76     | 4.86        |
| Kyrgyzstan | 0.77    | 0.76   | 0.73   | 0.76   | 0.72   | 0.56     | 0.46        |
| Pakistan   | 9.10    | 9.01   | 8.91   | 8.97   | 8.85   | 8.79     | 6.52        |
| Russia     | 25.89   | 26.00  | 24.29  | 26.30  | 23.98  | 22.44    | 15.08       |
| Tajikistan | 0.06    | 0.03   | 0.02   | 0.04   | 0.03   | 0.02     | 0.01        |
| Uzbekistan | 0.82    | 0.81   | 0.78   | 0.78   | 0.77   | 0.75     | 0.70        |

As seen in Table 5, GA-ConvLSTM is more successful than the other models according to the RMSE. After GA-ConvLSTM, ConvLSTM, LSTM, MLP, CNN, SVM, and XGBoost were successful, respectively. As seen in Table 5, RMSEs are low for Kyrgyzstan, Kazakhstan, Pakistan, Tajikistan, and Uzbekistan, where the number of cases is low.

However, RMSEs are also high for China, India, Iran, and Russia, where the number of cases is high. Table 6 indicates The death prediction outcomes of MAE in SCO member countries.

Table 6. The deaths prediction outcomes according to MAE in SCO member countries (*SCO üye ülkelerinde MAE'ye göre ölüm tahmini sonuçları*)

| Country    | XGBoost | SVM    | MLP    | CNN    | LSTM   | ConvLSTM | GA-ConvLSTM |
|------------|---------|--------|--------|--------|--------|----------|-------------|
| China      | 9.17    | 9.00   | 5.46   | 7.04   | 5.15   | 5.09     | 4.46        |
| India      | 136.40  | 110.46 | 109.56 | 111.08 | 108.22 | 108.84   | 85.30       |
| Iran       | 15.94   | 15.89  | 14.62  | 14.70  | 14.96  | 14.68    | 10.28       |
| Kazakhstan | 3.85    | 3.84   | 3.65   | 3.78   | 3.69   | 3.11     | 2.41        |
| Kyrgyzstan | 0.62    | 0.56   | 0.54   | 0.56   | 0.56   | 0.56     | 0.53        |
| Pakistan   | 5.83    | 5.35   | 5.12   | 5.23   | 5.12   | 5.05     | 3.70        |
| Russia     | 19.52   | 19.76  | 17.80  | 19.48  | 17.64  | 15.50    | 12.81       |
| Tajikistan | 0.05    | 0.03   | 0.03   | 0.03   | 0.02   | 0.01     | 0.01        |
| Uzbekistan | 0.56    | 0.57   | 0.55   | 0.51   | 0.54   | 0.52     | 0.47        |

According to the MAE, GA-ConvLSTM is more successful than the other models. After GA-ConvLSTM, ConvLSTM, LSTM, MLP, CNN, SVM, and XGBoost were successful, respectively. As seen in Table 6, MAE values are low for Kyrgyzstan,

Tajikistan, and Uzbekistan, where the number of cases is low. However, MAE values are also high for India. Fig. 9 and Table 7 show the death prediction outcomes according to  $R^2$  in SCO member countries.

Table 7. The death prediction outcomes according to  $R^2$  in SCO member countries (*SCO üye ülkelerinde  $R^2$ 'ye göre ölüm tahmini sonuçları*)

| Country    | XGBoost | SVM   | MLP   | CNN   | LSTM  | ConvLSTM | GA-ConvLSTM |
|------------|---------|-------|-------|-------|-------|----------|-------------|
| China      | 0.917   | 0.931 | 0.970 | 0.949 | 0.971 | 0.972    | 0.982       |
| India      | 0.332   | 0.507 | 0.512 | 0.502 | 0.513 | 0.516    | 0.623       |
| Iran       | 0.977   | 0.978 | 0.978 | 0.978 | 0.979 | 0.979    | 0.994       |
| Kazakhstan | 0.963   | 0.963 | 0.966 | 0.964 | 0.969 | 0.972    | 0.975       |
| Kyrgyzstan | 0.719   | 0.730 | 0.747 | 0.730 | 0.755 | 0.760    | 0.765       |
| Pakistan   | 0.880   | 0.883 | 0.885 | 0.883 | 0.887 | 0.888    | 0.980       |
| Russia     | 0.994   | 0.994 | 0.995 | 0.994 | 0.995 | 0.995    | 0.997       |
| Tajikistan | 0.421   | 0.423 | 0.425 | 0.422 | 0.428 | 0.429    | 0.432       |
| Uzbekistan | 0.838   | 0.841 | 0.854 | 0.854 | 0.858 | 0.865    | 0.882       |

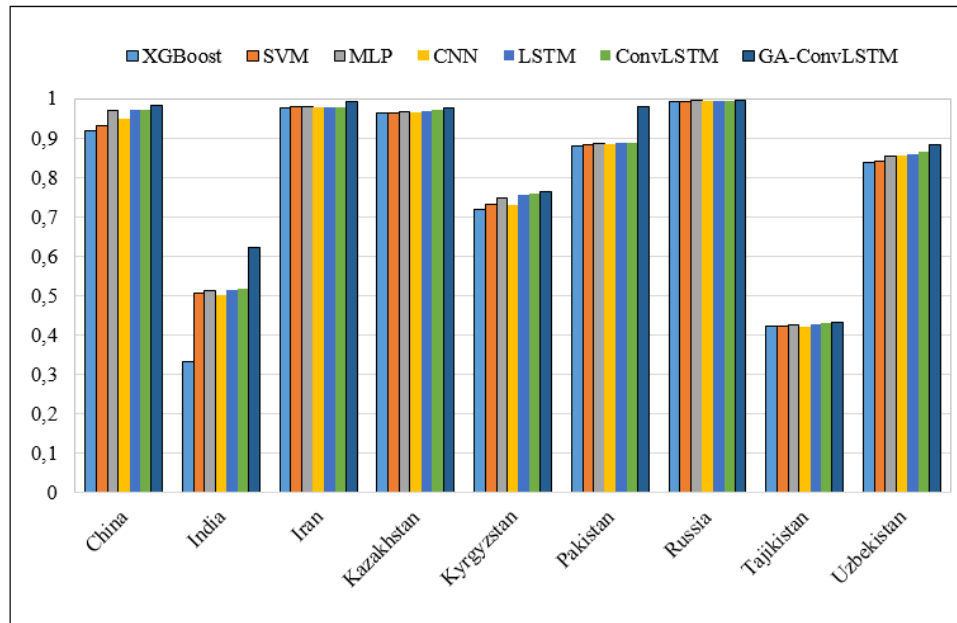
Figure 9. The death prediction results according to  $R^2$  in SCO member countries (*SCO üye ülkelerinde  $R^2$ 'ye göre ölüm tahmini sonuçları*)

Table 7 and Fig. 9 show that the  $R^2$  value is above 0.9 in all compared models for China, Iran, Kazakhstan,

and Russia. However, the  $R^2$  value of all models compared, especially for Tajikistan, is low. Because

the daily death numbers reported for Tajikistan in the dataset are close to 0. Fig. 10 shows the prediction

graphs of GA-ConvLSTM for predicting daily COVID-19 deaths.

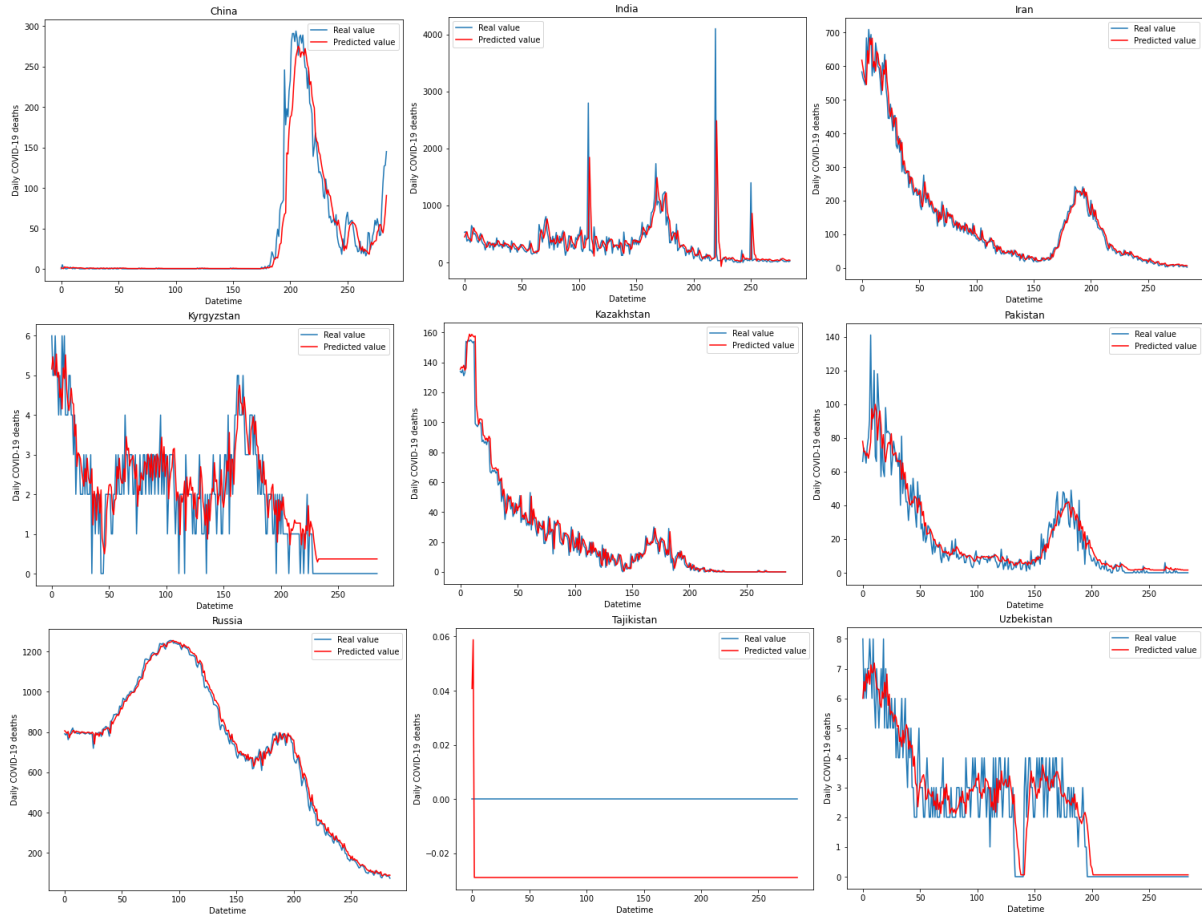


Figure 10. Prediction graphs of GA-ConvLSTM for predicting daily COVID-19 deaths (*Günlük COVID-19 ölümlerini tahmin etmek için GA-ConvLSTM tahmin grafikleri*)

As seen in Fig. 10, in predicting the daily COVID-19 deaths, GA-ConvLSTM has effectively modeled the changes in the daily deaths in countries other than Tajikistan and Uzbekistan. Daily death numbers for Tajikistan and Uzbekistan are reported as very low perhaps countries' failure to declare their death numbers prevented successful modeling of GA-ConvLSTM.

In this study, the GA-ConvLSTM model was created by optimizing the hyper-parameters of the ConvLSTM model with GA. GA-ConvLSTM model was compared with base ConvLSTM, LSTM, CNN, MLP, SVM, and XGBoost, where hyper-parameters were optimized with Grid Search. Experiments showed that the GA-ConvLSTM model was more successful than the other models.

ConvLSTM combines CNN and LSTM to analyze time series data, such as daily COVID-19 data. In the ConvLSTM hybrid model, CNN identifies local patterns in the data, while LSTM discovers time

dependencies and long-term relationships in time series data. Grid Search determined the hyper-parameters of base ConvLSTM. However, since Grid Search only tested the values in the specified parameter ranges, successful prediction performance needed to be improved. However, GA enabled the model to have a more successful prediction performance by testing a more extensive hyper-parameter range.

The effectiveness of ConvLSTM over LSTM can be interpreted by CNN's success in detecting local patterns. In this way, ConvLSTM extracts important features in time series data. The effectiveness of ConvLSTM over MLP and CNN can be interpreted by the time dependencies of MLP and CNN in the data and their limitations in capturing complex relationships in time series data. ConvLSTM is more effective than XGBoost and SVM because classical machine learning algorithms cannot model the dynamic structures and dependencies in time-dependent data.

#### 4.2. Determining the Intercountry Spread Pattern of COVID-19 in SCO Member Countries (SCO Üye Ülkelerinde COVID-19'un Ülkelerarası Yayılma Örüntüsünün Belirlenmesi)

It has been determined that COVID-19 symptoms emerge in most patients on the fifth day of the illness. It was observed that in some patients, the emergence of symptoms could take up to the 14th day of the disease. In this study, 5 and 14-day incubation times were applied to determine the distribution of COVID-

19 in SCO member countries. Moreover, the dates when the number of deaths and cases in each SCO member country peaked were determined, and the date ranges 5 and 14 days before and 5 and 14 days after these dates were determined. Chord diagrams were created according to these dates. The dates when the case counts in each SCO member country were highest, 5 and 14 days before and 5 and 14 days after these dates, are seen in the following table.

Table 8. Dates when each SCO member country has the maximum case counts and incubation periods (Her SCO üye ülkesinin maksimum vaka sayısına ve kuluçka süresine sahip olduğu tarihler)

| Country    | 14 days before        | 5 days before         | Peak date  | 5 days after          | 14 days after         |
|------------|-----------------------|-----------------------|------------|-----------------------|-----------------------|
| China      | 2022/05/14-2022/05/27 | 2022/05/23-2022/05/27 | 2022/05/28 | 2022/05/29-2022/06/02 | 2022/05/29-2022/06/11 |
| India      | 2021/04/24-2021/05/06 | 2021/05/02-2021/05/06 | 2021/05/07 | 2021/05/08-2021/05/12 | 2021/05/08-2021/05/21 |
| Iran       | 2021/08/04-2021/08/17 | 2021/08/13-2021/08/17 | 2021/08/18 | 2021/08/19-2021/08/23 | 2021/08/19-2021/09/02 |
| Kazakhstan | 2022/01/07-2022/01/20 | 2022/01/16-2022/01/20 | 2022/01/21 | 2022/01/22-2022/01/26 | 2022/01/22-2022/02/04 |
| Kyrgyzstan | 2021/06/16-2021/06/29 | 2021/06/25-2021/06/29 | 2021/06/30 | 2021/07/01-2021/07/05 | 2021/07/01-2021/07/14 |
| Pakistan   | 2022/01/15-2022/01/28 | 2022/01/24-2022/01/28 | 2022/01/29 | 2022/01/30-2022/02/03 | 2022/01/30-2022/02/12 |
| Russia     | 2022/01/28-2022/02/10 | 2022/02/06-2022/02/10 | 2022/02/11 | 2022/02/12-2022/02/16 | 2022/02/12-2022/02/25 |
| Tajikistan | 2020/05/05-2020/05/18 | 2020/05/14-2020/05/18 | 2020/05/19 | 2020/05/20-2020/05/24 | 2020/05/20-2020/06/02 |
| Uzbekistan | 2022/01/09-2022/01/22 | 2022/01/18-2022/01/22 | 2022/01/23 | 2022/01/24-2022/01/28 | 2022/01/24-2022/02/06 |

According to the dates presented in Table 8, the distributions of the transmission pattern of cases among SCO member countries were determined by drawing chord diagrams for the 5-day incubation time in Fig. 11 and the 14-day incubation time in Fig. 12.

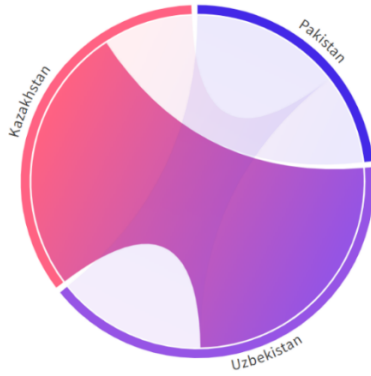


Figure 11. Spread pattern of COVID-19 cases among SCO member countries for 5-day incubation time (SCO üye ülkeleri arasında COVID-19 vakalarının 5 günlük kuluçka süresi boyunca yayılma örüntüsü)

Fig. 11 shows that the transmission pattern of cases in Pakistan, Kazakhstan and Uzbekistan showed a similar pattern. As seen in Fig. 11, Pakistan and Uzbekistan have 5 days in common, Pakistan and Kazakhstan have 3 days, and Kazakhstan and Uzbekistan have 9 days in common.

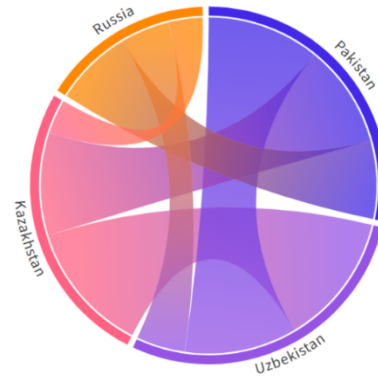


Figure 12. Spread pattern of COVID-19 cases among SCO member countries for 14-day incubation time (SCO üye ülkeleri arasında COVID-19 vakalarının 14 günlük kuluçka süresi boyunca yayılma örüntüsü)

Fig. 12 shows that the transmission pattern of cases in Pakistan, Kazakhstan, Russia, and Uzbekistan showed a similar pattern. As seen in Fig. 11, Pakistan and Uzbekistan have 23 days in common, Pakistan and Kazakhstan have 21 days in common, Kazakhstan and Uzbekistan have 27 days in common, Pakistan and Russia have 16 days in common, Russia and Uzbekistan have 10 days in common, and Kazakhstan and Russia have 7 days in common.

Table 9 shows the dates when the number of deaths in each SCO member country was highest, as well as the 5 and 14 days before and 5 and 14 days after these dates.

Table 9. Dates with the highest number of deaths and incubation periods of each SCO member country (*Her SCO üyesi ülkenin en fazla ölüm ve kuluçka döneminin olduğu tarihler*)

| Country    | 14 days before        | 5 days before         | Peak date  | 5 days after          | 14 days after         |
|------------|-----------------------|-----------------------|------------|-----------------------|-----------------------|
| China      | 2020/04/04-2020/04/17 | 2020/04/13-2020/04/17 | 2020/04/18 | 2020/04/19-2020/04/23 | 2020/04/19-2020/05/01 |
| India      | 2021/05/27-2021/06/09 | 2021/05/05-2021/06/09 | 2021/06/10 | 2021/06/11-2021/06/15 | 2021/06/11-2021/06/24 |
| Iran       | 2021/08/11-2021/08/24 | 2021/08/20-2021/08/24 | 2021/08/25 | 2021/08/26-2021/08/30 | 2021/08/26-2021/09/08 |
| Kazakhstan | 2021/08/12-2021/08/25 | 2021/08/21-2021/08/25 | 2021/08/26 | 2021/08/27-2021/08/31 | 2021/08/27-2021/09/09 |
| Kyrgyzstan | 2020/07/05-2020/07/18 | 2020/07/14-2020/07/18 | 2020/07/19 | 2020/07/20-2020/07/24 | 2020/07/20-2020/08/02 |
| Pakistan   | 2020/11/07-2020/11/20 | 2020/11/16-2020/11/20 | 2020/11/21 | 2020/11/22-2020/11/26 | 2020/11/22-2020/12/05 |
| Russia     | 2021/11/05-2021/11/18 | 2021/11/14-2021/11/18 | 2021/11/19 | 2021/11/20-2021/11/24 | 2021/11/20-2021/12/03 |
| Tajikistan | 2020/04/30-2020/05/13 | 2020/05/09-2020/05/13 | 2020/05/14 | 2020/05/15-2020/05/19 | 2020/05/15-2020/05/28 |
| Uzbekistan | 2020/07/15-2020/07/28 | 2020/07/24-2020/07/28 | 2020/07/29 | 2020/07/30-2020/08/03 | 2020/07/30-2020/08/12 |

According to the dates presented in Table 9, the distributions of the transmission pattern of deaths among SCO member countries were determined by drawing chord diagrams for the 14-day incubation period in Fig. 13 and the 14-day incubation time in Fig. 14.

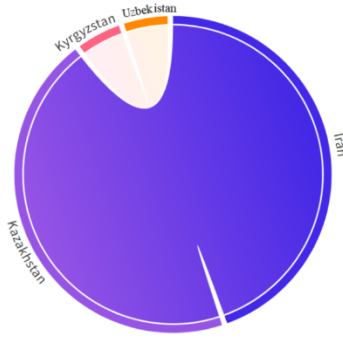


Figure 13. Spread pattern of COVID-19 deaths among SCO member countries for 5-day incubation time (*SCO üyesi ülkelerde COVID-19 ölümlerinin 5 günlük kuluçka süresi boyunca yayılma örüntüsü*)

Fig. 13 shows that for the 5-day incubation time, the transmission pattern of deaths in Iran and Kazakhstan, Kyrgyzstan and Uzbekistan showed a similar pattern. As seen in Fig. 13, Iran and Kazakhstan have 10 days in common and Kyrgyzstan and Uzbekistan have 1 days in common.

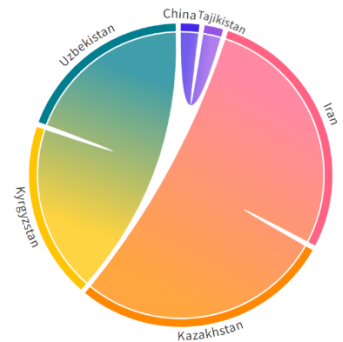


Figure 14. Spread pattern of COVID-19 deaths among SCO member countries for 14-day incubation time (*SCO üyesi ülkelerde COVID-19 ölümlerinin 14 günlük kuluçka süresi boyunca yayılma örüntüsü*)

Fig. 14 shows that for the 14-day incubation time, the transmission pattern of deaths in China, Iran, Kazakhstan, Kyrgyzstan, Tajikistan and Uzbekistan showed a similar pattern. As shown in Fig. 14, Tajikistan and China have 2 days in common, Iran and Kazakhstan have 28 days in common, Kyrgyzstan and Uzbekistan have 19 days in common.

## 5. CONCLUSIONS (SONUÇLAR)

COVID-19 has spread very quickly to almost every part of the world, causing many people to experience severe symptoms and lose their lives. Thanks to the intense and powerful scientific studies conducted in the past four years, many unknown issues regarding the disease have been clarified. With the introduction of vaccines that are effective in protecting against COVID-19 and antivirals that are effective in treatment and the increase in the immunity level of people who have the disease, the morbidity and mortality rates caused by the disease have decreased significantly. However, as COVID-19 spread worldwide, countries' health systems remained inadequate. Due to the complications caused by COVID-19, people were admitted to intensive care units and even lost their lives. Hospital intensive care units were complete, and healthcare personnel and medical equipment were insufficient.

In this study, the hybrid GA-ConvLSTM model was created to predict the daily deaths and cases due to COVID-19 in SCO member countries and to determine the distribution of the spread pattern of COVID-19 among SCO member countries. It aimed to obtain a more successful prediction performance by optimizing the hyper-parameters of the ConvLSTM model using GA. GA-ConvLSTM, XGBoost, SVM, MLP, CNN, LSTM, and ConvLSTM have been extensively tested using RMSE,  $R^2$ , and MAE. The study used daily COVID-19 case and death data provided by WHO from 2020/01/03 to 2022/05/31. Comparing models were evaluated. Experiments have shown that GA-ConvLSTM outperforms the compared models based on RMSE,  $R^2$ , and MAE.

For daily COVID-19 cases, GA-ConvLSTM had 0.9  $R^2$  for all SCO member countries. For daily COVID-19 deaths, GA-ConvLSTM had an  $R^2$  value above 0.9 for China, Iran, Kazakhstan, Pakistan, and Russia. However, it had a value of 0.623  $R^2$  for India, 0.765  $R^2$  for Kyrgyzstan, 0.432  $R^2$  for Tajikistan, and 0.882  $R^2$  for Uzbekistan. The relatively low  $R^2$  values obtained for India and Tajikistan can be expressed as the countries not reporting their actual mortality values. Specifically, the total number of deaths presented for Tajikistan was reported as 9. Unrealistic death numbers prevented the compared models from modeling the dataset successfully. However, the experimental results showed that GA-ConvLSTM can successfully model undulations in the daily deaths and cases of epidemic diseases such as COVID-19.

5 and 14-day incubation periods were used because WHO reported that symptoms caused by COVID-19 appear within the first 5 days but can extend up to the first 14 days in some patients. To determine the distribution of the spread of COVID-19 among SCO member countries, the dates with the highest number of COVID-19 deaths and cases were identified for SCO member countries. Moreover, chord graphs were created by determining 5 and 14 days before and 5 and 14 days after these peak dates.

The spread of cases in Pakistan, Kazakhstan, and Uzbekistan for the 5-day incubation time showed a similar pattern. Pakistan and Uzbekistan have 5 days in common, Pakistan and Kazakhstan have 3 days, and Kazakhstan and Uzbekistan have 9 days in common. For the 14-day incubation time, the transmission pattern of cases in Pakistan, Kazakhstan, Russia, and Uzbekistan showed a similar pattern. Pakistan and Uzbekistan have 23 days in common, Pakistan and Kazakhstan have 21 days, Kazakhstan and Uzbekistan have 27 days, Russia and Pakistan have 16 days, Russia and Uzbekistan have 10 days, and Kazakhstan and Russia have 7 days in common.

For the 5-day incubation time, the spread pattern of deaths in Iran and Kazakhstan, Kyrgyzstan, and Uzbekistan showed a similar pattern. Iran and Kazakhstan have 10 days in common, and Kyrgyzstan and Uzbekistan have 1 day in common. During the 14-day incubation period, the spread pattern of COVID-19 deaths in China, Iran, Kazakhstan, Kyrgyzstan, Tajikistan and Uzbekistan showed a similar pattern. China and Tajikistan have 2 days in common, Iran and Kazakhstan have 28 days in common, and Kyrgyzstan and Uzbekistan have 19 days in common.

## REFERENCES (KAYNAKLAR)

- [1] E. C. Abebe, T. A., Dejenie, M. Y., Shiferaw, and T. Malik, "The newly emerged COVID-19 disease: a systemic review" *Virology journal*, vol. 17, no. 1, 2020.
- [2] I. H. Elrobaa, and K. J. New, "COVID-19: pulmonary and extra pulmonary manifestations" *Frontiers in public health*, vol. 9, 2021.
- [3] C. Lai, R. Yu, M. Wang, W. Xian, X. Zhao, Q. Tang, and F. Wang, "Shorter incubation period is associated with severe disease progression in patients with COVID-19" *Virulence*, vol. 11, no. 1, pp. 1443-1452, 2020.
- [4] A. H. Al-Ani, R. E. Prentice, C. A. Rentsch, D. Johnson, Z. Ardalan, N. Heerasing, and B. Christensen, "Prevention, diagnosis and management of COVID-19 in the IBD patient" *Alimentary Pharmacology & Therapeutics*, vol. 52, no. 1, pp. 54-72, 2020.
- [5] A. Ayaz, A. Arshad, H. Malik, H. Ali, E. Hussain, and B. Jamil, "Risk factors for intensive care unit admission and mortality in hospitalized COVID-19 patients" *Acute and critical care*, vol. 35, no. 4, 2020.
- [6] S. Singh, G. C. Ambooken, R. Setlur, S. K. Paul, M. Kanitkar, S. S. Bhatia, and R. S. Kanwar, "Challenges faced in establishing a dedicated 250 bed COVID-19 intensive care unit in a temporary structure" *Trends in Anaesthesia and Critical Care*, vol. 36, pp. 9-16, 2021.
- [7] Y. Xue, and B. M. Makengo, "Twenty years of the Shanghai Cooperation Organization: Achievements, challenges and prospects" *Open Journal of Social Sciences*, vol. 9, no. 10, pp. 184-200, 2021.
- [8] I. Ahmad, "Shanghai Cooperation Organization: China, Russia, and Regionalism in Central Asia" *Initiatives of Regional Integration in Asia in Comparative Perspective: Concepts, Contents and Prospects*, pp. 119-135, 2018.
- [9] S. Azizi, "China's Belt and Road Initiative (BRI): The Role of the Shanghai Cooperation Organization (SCO) in Geopolitical Security and Economic Cooperation" *Open Journal of Political Science*, vol. 14, no. 1, pp. 111-129, 2024.
- [10] F. Shahid, A. Zameer, and M. Muneeb, "Predictions for COVID-19 with deep learning

- models of LSTM, GRU and Bi-LSTM” *Chaos, Solitons & Fractals*, vol. 140, 2020.
- [11] H. Abbasimehr, and R. Paki, “Prediction of COVID-19 confirmed cases combining deep learning methods and Bayesian optimization” *Chaos, Solitons & Fractals*, vol. 142, 2021.
- [12] H. T. Rauf, M. I. U. Lali, M. A. Khan, S. Kadry, H. Alolaiyan, A. Razaq, and R. Irfan, “Time series forecasting of COVID-19 transmission in Asia Pacific countries using deep neural networks” *Personal and Ubiquitous Computing*, pp. 1-18, 2023.
- [13] L. Zhou, C. Zhao, N. Liu, X. Yao, and Z. Cheng, “Improved LSTM-based deep learning model for COVID-19 prediction using optimized approach” *Engineering applications of artificial intelligence*, vol. 122, 2023.
- [14] C. C. Ukwuoma, D. Cai, M. B. B. Heyat, O. Bamisile, H. Adun, Z. Al-Huda, and M. A. Al-Antari, “Deep learning framework for rapid and accurate respiratory COVID-19 prediction using chest X-ray images” *Journal of King Saud University-Computer and Information Sciences*, vol. 35, no. 7, 2023.
- [15] A. Al-Rashedi, and M. A. Al-Hagery, “Deep learning algorithms for forecasting COVID-19 cases in Saudi Arabia” *Applied Sciences*, vol. 13, no. 3, 2023.
- [16] S. Solayman, S. A. Aumi, C. S. Mery, M. Mubassir, and R. Khan, “Automatic COVID-19 prediction using explainable machine learning techniques” *International Journal of Cognitive Computing in Engineering*, vol. 4, pp. 36-46, 2023.
- [17] M. Kim, and H. Kim, “A Dynamic Analysis Data Preprocessing Technique for Malicious Code Detection with TF-IDF and Sliding Windows” *Electronics*, vol. 13, no. 5, 2024.
- [18] M. A. Muslim, and Y. Dasril, “Company bankruptcy prediction framework based on the most influential features using XGBoost and stacking ensemble learning” *International Journal of Electrical and Computer Engineering (IJECE)*, vol. 11, no. 6, pp. 5549-5557, 2021.
- [19] A. Asselman, M. Khaldi, and S. Aammou, “Enhancing the prediction of student performance based on the machine learning XGBoost algorithm” *Interactive Learning Environments*, vol. 31, no. 6, pp. 3360-3379, 2023.
- [20] A. Rizwan, N. Iqbal, R. Ahmad, and D. H. Kim, “WR-SVM model based on the margin radius approach for solving the minimum enclosing ball problem in support vector machine classification” *Applied Sciences*, vol. 11, no. 10, 2021.
- [21] M. Aslani, S. Seipel, “Efficient and decision boundary aware instance selection for support vector machines” *Information Sciences*, vol. 577, pp. 579-598, 2021.
- [22] R. Sharma, M. Kim, and A. Gupta, “Motor imagery classification in brain-machine interface with machine learning algorithms: Classical approach to multi-layer perceptron model” *Biomedical Signal Processing and Control*, 71, 2022.
- [23] D. D. Oliveira, M. Rampinelli, G. Z. Tozatto, R. V. Andreão, and S. M. Müller, “Forecasting vehicular traffic flow using MLP and LSTM” *Neural Computing and applications*, vol. 33, pp. 17245-17256, 2021.
- [24] N. Calik, M. A. Belen, and P. Mahouti, “Deep learning base modified MLP model for precise scattering parameter prediction of capacitive feed antenna”. *International journal of numerical modelling: electronic networks, devices and fields*, vol. 33, no. 2, 2020.
- [25] R. Shyam, “Convolutional neural network and its architectures” *Journal of Computer Technology & Applications*, vol. 12, no. 2, pp. 6-14, 2021.
- [26] M. P. Akhter, Z. Jiangbin, I. R. Naqvi, M. Abdelmajeed, A. Mehmood, and M. T. Sadiq, “Document-level text classification using single-layer multisize filters convolutional neural network” *IEEE Access*, vol. 8, pp. 42689-42707, 2020.
- [27] H. V. Dudukcu, M. Taskiran, Z. G. C. Taskiran, and T. Yildirim, “Temporal Convolutional Networks with RNN approach for chaotic time series prediction” *Applied soft computing*, vol. 133, 2023.
- [28] K. Hermann, and A. Lampinen, “What shapes feature representations? exploring datasets, architectures, and training” *Advances in Neural Information Processing Systems*, vol. 33, pp. 9995-10006, 2020.
- [29] S. Patil, V. M. Mudaliar, P. Kamat, and S. Gite, “LSTM based Ensemble Network to enhance the learning of long-term dependencies in

- chatbot” *International Journal for Simulation and Multidisciplinary Design Optimization*, vol. 11, no. 25, 2020.
- [30] J. Duan, P. F. Zhang, R. Qiu, and Z. Huang, “Long short-term enhanced memory for sequential recommendation” *World Wide Web*, vol. 26, no. 2, pp. 561-583, 2023.
- [31] K. E. ArunKumar, D. V. Kalaga, C. M. S. Kumar, M. Kawaji, and T. M. Brenza, “Forecasting of COVID-19 using deep layer recurrent neural networks (RNNs) with gated recurrent units (GRUs) and long short-term memory (LSTM) cells” *Chaos, Solitons & Fractals*, vol. 146, 2021.



# Journal of Information Systems and Management Research

Bilişim Sistemleri ve Yönetim Araştırmaları Dergisi

<http://dergipark.gov.tr/jismar>

Derleme Makalesi / Review Article

## Kuantum Yazılım Test Teknikleri Ve Araçları Üzerine Bir İnceleme

Fatma Betül ÖZDEMİR<sup>a\*</sup>, Savaş ÖZTÜRK<sup>b</sup>

<sup>a</sup>Marmara Üniversitesi, Teknoloji Fakültesi, Bilgisayar Mühendisliği Bölümü, İSTANBUL, 34854, TÜRKİYE

<sup>b</sup>Marmara Üniversitesi, Teknoloji Fakültesi, Elektrik-Elektronik Mühendisliği Bölümü, İSTANBUL, 34854, TÜRKİYE

### MAKALE BİLGİSİ

Alınma: 10.02.2025  
Kabul: 27.04.2025

**Anahtar Kelimeler:**  
Kuantum Programı,  
Kuantum Yazılım Testi,  
Kuantum Yazılım  
Mühendisliği

**\*Sorumlu Yazar**

e-posta:  
fbetul23@marun.edu.tr

### ÖZET

Kuantum bilgisayarlarının gelişimiyle birlikte kuantum yazılımlarının üretilme ve test ihtiyacı doğmuştur. Geleneksel yazılım test yöntemleri, kuantum mekaniğinin doğasından kaynaklanan karmaşıklık nedeniyle kuantum yazılımlarını test etmede yetersiz kalmaktadır. Kuantum yazılım sistemlerinin doğruluk ve güvenilirliğini sağlamak için kuantum test teknikleri ve araçları, özel olarak geliştirilmiş çözümler sunmaktadır. Bu çalışmada kuantum bilgisayar teknolojisinin geleceği için temel niteliği taşıyan güncel çalışmalar anlatılmıştır. Başlıca QuanFuzz, QMutPy, Quito vb. araçlar ile metamorfik, kombinasyonel ve arama tabanlı test tekniklerinden bahsedilmiştir. Kuantum test mühendisliği alanında görülen standartlaşma problemi üzerinde değerlendirmeler yapılmıştır. Bu çalışma, kuantum test teknikleri ve araçları üzerinde çalışma yapmak isteyen araştırmacılara bir rehber niteliğindedir.

DOL: 10.59940/jismar.1636890

## A Review On Quantum Software Testing Techniques And Tools

### ARTICLE INFO

Received: 10.02.2025  
Accepted: 27.04.2025

**Keywords:**  
Quantum Program,  
Quantum Software  
Testing, Quantum  
Software Engineering

**\*Corresponding Authors**

e-mail:  
fbetul23@marun.edu.tr

### ABSTRACT

With the development of quantum computers, the need to produce and test quantum software has arisen. Traditional software testing methods are inadequate for testing quantum software due to the inherent complexity of quantum mechanics. Quantum testing techniques and tools provide specially developed solutions to ensure the accuracy and reliability of quantum software systems. In this study, current works that are fundamental for the future of quantum computing technology are described. The main tools such as QuanFuzz, QMutPy, Quito, etc. and metamorphic, combinatorial and search-based testing techniques are mentioned. The problem of standardisation in the field of quantum test engineering is evaluated. This study is a guide for researchers who want to work on quantum testing techniques and tools.

DOL: 10.59940/jismar.1636890

### GİRİŞ (INTRODUCTION)

Yazılım testleri, yazılım ürünlerinin belirlenen gereksinimleri karşılayıp karşılamadığını belirleyen, ürünün amaca uygunluğunu tespit eden,

kusurlarını bulan, hem statik hem de dinamik süreç aktivitelerini kapsar.

Gelişen teknoloji ile çeşitli donanımların yerini artık yazılımlar almaya başlamış, mobil teknolojiler

günlük hayatımızın vazgeçilmez parçası haline gelmiştir. Örneğin, elektrikli otomobillerde, mekanik unsurlar yazılımla kontrol edilmeye başlanmış, bu durum geliştirici ve testçileri yazılımlarda oluşacak hataların can kaybına yol açacak kazalara karşı savunmasız kalma riski ile karşı karşıya getirmiştir. Yazılımsal başarısızlıklar can, mal ve itibar kayıplarına yol açarlar. Yazılımların test edilmesi ve kalitelerinin artırılması her zamankinden daha fazla önem kazanmıştır [1]. Geleneksel hesaplama dan önemli farklılıklar içeren kuantum yazılımlar, kendine has yöntemlerle test edilmektedir.

Son yıllarda yapay zeka ve makine öğrenmesi alanındaki önemli gelişmeler işlem gücü ihtiyacını da artırmış, daha hızlı bilgisayarlara ihtiyaç her zamankinden fazla olmuştur. 2025 yılının henüz başında ortaya çıkan DeepSeek adlı büyük dil modeli, pekiştirmeli öğrenme tabanlı öğrenme modeli ile rakiplerine işlem gücünü az gerektirmesi avantajıyla üstünlük kurmaya çalışmıştır. Diğer yandan OpenAI ve Meta gibi rakipleri de Nvidia gibi çip üreticilerinin satış stratejisini yönlendirerek pazar payını koruma stratejisi gütmektedir. Hesaplama gücü gereksinimleri, daha hızlı ve özgün alternatiflere ve özellikle de kuantum bilgisayarlara olan ilgiyi artırmıştır. Klasik yazılım testlerinde hata ayıklama deterministik bir süreçtir; bir girdinin çıktısı her zaman aynıdır. Ancak kuantum yazılımları, kuantum mekaniğinin olasılıksal doğasına sahip olduğu için, aynı girdiye karşılık gelen çıktı her zaman değişebilir. Bu durum, geleneksel test metodolojilerinin kuantum sistemlerine uygulanmasını zorlaştırır. Örneğin, bir klasik program hata içermediğinde her zaman aynı çıktıyı verirken, bir kuantum programı hata içermese bile ölçüm varyansları nedeniyle farklı çıktılar üretebilir. Bu nedenle, kuantum yazılım testi de klasik yazılım testinden ayrılır.

Kuantum Bilgisayarları (Quantum Computer - QC)[4], süperpozisyon (üst üste binme) ve dolanıklık (Entanglement)[2] gibi kendine has özellikleri olan çok sayıda karmaşık sorunu çözmeye çalışır. Örneğin; süperpozisyon; dinamik ya da statik bir sistemde, farklı doğrultu ve şiddetlerde etkilemekte olan birden fazla kuvvetin sistem üzerindeki etkilerinin ayrı ayrı belirlenip sonuçlarının toplamının alınmasıdır. Örneğin iki teknenin oluşturduğu dalgaların birbirinin içinden geçmesi ve birbirini etkilemesi oldukça karmaşık fizik hesaplarıyla çözülebilmektedir. Kuantum bilgisayar, küçük ölçeklerde, fiziksel madde hem parçacıkların hem de dalgaların özelliklerini sergiler ve kuantum hesaplama, bu davranışı özel donanım kullanarak

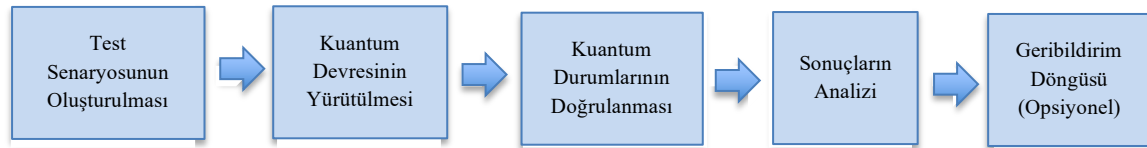
güçlendirir. Klasik fizik, bu kuantum aygıtlarının işleyişini açıklayamaz ve ölçeklenebilir bir kuantum bilgisayarı bazı hesaplamaları herhangi bir modern "klasik" bilgisayardan çok daha hızlı gerçekleştirme becerisini kuantum mekaniksel olgulardan alır[3]. Kuantum biti veya kısaca kübit, kuantum hesaplamasının temelidir. Klasik bir bit 0 ve 1'lerden oluşurken, kübitler ise  $|0\rangle$  ve  $|1\rangle$  biçiminde durumlardan oluşur. Bir kübitin genel durumu  $|0\rangle + b|1\rangle$ 'dir. Karmaşık sayılar ile ifade edilen eşitlik;  $|a|^2 + |b|^2 = 1$ 'i sağlayan genlik ifadesidir. Genlikler her baz durumunun olasılık oranını temsil eder. Kuantum aygıtlarında, kübitlerin bilgisi yalnızca ölçümle elde edilebilir. Bir kuantum sistemini ölçmek, karşılık gelen genliğe sahip olma olasılığı olan klasik değerlerin elde edilmesi ile mümkündür. Elde edilen durumlar, iki kübit dolanık olduğunda, tüm durum iki alakasız kübit olarak ayrılamaz [2]. Basitçe, iki nesnenin zaman ve uzaydan bağımsız olarak birbirinden haberdar olması, birbirini etkilemesi (kelebek etkisi) olarak ifade edilebilir. Birçok yazılım mühendisi kuantum devrelerini programlamak için gerekli bilgi ve tecrübeye sahip olmadığı için kuantum programlarının kalite güvencesine yönelik sistematik bir süreç henüz ortaya konmamıştır[4]. Bu amaçla, Kuantum Yazılım Mühendisliği (Quantum Software Engineering - QSE)[4], yazılım mühendislerinin kuantum yazılımları üretebilmeleri için daha yüksek bir soyutlama düzeyinde uygun yöntem ve araçlar sağlamayı amaçlamaktadır [4]. Kuantum algoritmalarındaki gelişmelere ve kaynak gereksinimlerinin optimizasyonuna rağmen, algoritmaların birçoğunun donanım gereksinimleri yakın vadeli kuantum bilgisayarlarının yeteneklerinin çok ötesindedir. hem kuantum hem de klasik kaynakları kullanarak geleneksel klasik bilgisayarların erişemediği özdeğer ve optimizasyon problemlerine varyasyonel çözümler bulmak için tasarlanmış hibrit kuantum-klasik algoritma olan varyasyonel kuantum özçözücü (Variational Quantum Eigensolver - VQE)[5]'ler geliştirilmiştir. VQE, erken kuantum bilgisayarlarının performansını keşfetmek ve algoritmaların belirli bir mimarinin güçlü yönlerinden yararlanmak üzere tasarlanmıştır[5].

Kuantum bilgisayarların, yazılımlarının ve testlerinin endüstrideki uygulamaları son kullanıcı düzeyinde karşılık bulmaya başlamıştır. IBM, kuantum yazılımlarını test etmek için IBM Quantum Experience platformunda kendi hata düzeltme ve test araçlarını geliştirmiştir. Örneğin, IBM Qiskit framework'ü, kuantum programlarının doğruluğunu test etmek için gürültüye duyarlı (noise-aware) test teknikleri kullanmaktadır[3]. Bu teknik, kuantum gürültüsünü minimize etmeye yönelik hata modellemeleri içermektedir. Buna karşılık, Google

Quantum AI ekibi, Sycamore kuantum bilgisayarı üzerinde kuantum hata oranlarını tespit etmek için metamorfik test yöntemlerini kullanmaktadır.

Kuantum yazılım test teknikleri ve araçları; yazılım geliştirmenin tüm aşamalarında (modelleme, analiz, test etme, hata ayıklama) kuantum yazılımının güvenilir mühendisliği için yöntemler ve yaklaşımlar geliştirmeye, programlardaki kuantum hatalarını düşük maliyetlerle keşfetmeye odaklanır. Klasik test yaklaşımları kuantum bilgisayarlarında çalışan karmaşık kuantum programları için ölçeklenebilen sistematik ve otomatik test yaklaşımları sunmamaktadır [4]. QC'nin uzmanlaşmış özellikleri (üst üste binme ve dolanıklık) nedeniyle test aşamasındaki zorlukları elimine edebilmek için önemli çalışmalar yapılmış ve çeşitli test araçları geliştirilmiştir. [6]. Bunların arasında bulanıklık test,

mutasyon analizi, arama tabanlı test, kombinasyonel test ile birlikte kuantum giriş çıkış testi (Quantum Input Output testing – Quito)[8] aracı yer almaktadır. Bir kuantum programını bir kara kutu olarak ele almak ve bir kuantum programı yürütmesinin sonunda durumları okumak bile olasılıksal çıktılarla sonuçlanır. Çıktıları değerlendirmek için, bir kuantum programını birden çok kez yürütmeye ve ilgili istatistiksel testleri uygulamaya başvurmak gerekir. Bu amaçla kuantum programlarının istatistiksel yaklaşımlarla yapılmış çalışmalar mevcuttur [7]. Ayrıca, birçok kuantum programının önceden tanımlanmış test araçları yoktur, bu da kuantum programlarının doğruluğunu belirlemede başka bir zorluktur. Bu amaçla, kuantum programlarını test etmek için metamorfik test [8] yaklaşımları, bu yönde değerli çalışmalardır.



Şekil 1. Kuantum Test Araçlarının Tipik İşleyişi (Typical Operation of Quantum Test Tools)

## YÖNTEMLER VE KAPSAMLAR (METHODS AND SCOPE)

### 1. Kapsayıcı Kriterler (Inclusive Criteria)

Test edilecek programın, test araçları veya yöntemleri ile tam kapsamlı ve uyumlu bir şekilde değerlendirilebilmesi için kriterler belirlenir. Kuantum programlarında fonksiyon ve doğruluk analizleri yapılırken, farklı hataları veya potansiyel problemleri tespit etmek için çeşitli yöntemler bir araya getirilir. Kapsayıcı kriterler, genellikle bir kuantum yazılımının farklı durumlarını (örneğin, süperpozisyon, dolanıklık, ölçüm süreçleri) kapsayacak şekilde bir test senaryosu sunmayı amaçlar.

Şekil 2[9]'de kuantum programlarının test edilmesine ilişkin akış diyagramı yer almaktadır. Bu diyagrama göre kuantum programı girdi olarak verildiğinde öncelikle yapısal analizi gerçekleştirilir. Yapısal analiz, kuantum programlarının alt rutin ve organizasyonlarını tespit etmek için gereklidir. Bu aşamada programın kuantum özelliklerini modellemek için Kuantum Yazılım Modelleme Dili (Quantum Software Modeling Language – Q-UML)[10] kullanılır. Q-UML kuantum yazılımını düzgün bir şekilde modellemesine olanak tanıyan Birleşik Modelleme Dili'nin (Unified Modeling Language - UML)[10]

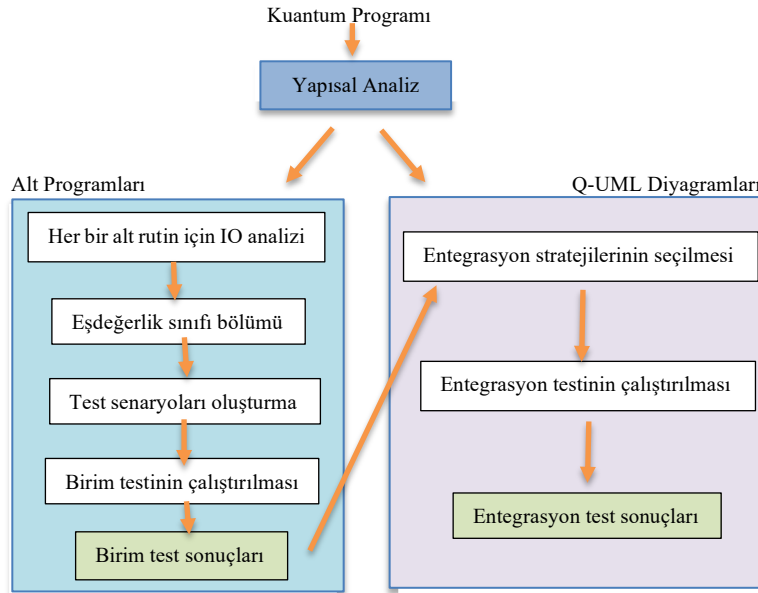
bir uzantısıdır. Her bir alt rutin üzerinde birim testi ve ardından tüm program üzerinde entegrasyon testi yapılır.

Alt rutin testleri için öncelikle girdi ve çıktı analizi yapılır. Kullanıcının atadığı değişkenler girdi olarak kabul edilirken program çalıştırdıktan sonra elde edilen kullanıcının ilgilendiği değişkenler çıktı olarak kabul edilir. Bu aşamanın ardından kuantum girdilerini bölümlendirmek için eşdeğerlilik sınıfı (Equivalence Classes) bölümleme testi uygulanır. Kübitlerin alt kümesi olarak tanımlanan değişkenlerinin taşıdığı değerler kuantum durumları olarak tanımlanır. Temel de üç kuantum durumu vardır bunlar süperpozisyon, karma ve klasik durumdur. Bölümleme kriteri kuantum değişkeninin durum türlerine bağlıdır. Bu bölümlemeye Klasik-Üst üste binme-Karma Bölme(Classical-Superposition-Mixed Partition- CSMP)[9] denir. Bununla birlikte karmaşık durumda çalışmayan kuantum durumları için klasik-süperpozisyon bölümü (Classical-Superposition Partition - CSP)[9] kriteri uygulanır. Süperpozisyon durumundan doğan dolanıklığında kapsanması gerekmektedir. Bunun için dolanıklık kapsamı (Entanglement Coverage - EntC)[9] ve süperpozisyon her bir kübit için üretilir. Öte yandan birden çok girdi değişkenine sahip olan kuantum programları için klasik kapsam kriterlerinden, Tüm Kombinasyon Kapsamı (All

Combination Coverage - ACoC)[9], Her Seçim Kapsamı (Each Choice Coverage - ECC)[9], Çiftler Halinde Kapsam (Pair-Wise Coverage - PWC)[9] ve Temel Seçim Kapsamı (Base Choice Coverage-BCC)[9] kuantum programlarına uyarlanabilir. Kapsam kriterleri uygulandıktan sonra test senaryoları oluşturulur. Test senaryosu temelde kuantum girdilerinin oluşturulması ve sonuçların tespit edilmesi olarak iki aşamadan oluşur. Klasik testlerden farkı beklenen çıktı kombinasyonları yerine İstatistik Tabanlı Algılama (Statistics Based Detection - SBD)[9], Kuantum Çalışma Zamanı İddiaları (Quantum Runtime Assertions - QRA)[9] gibi teknikler kullanılır. Klasik testlerdeki gibi girdi

durumları altında çıktıların elde edilmesi kuantum programları için yeterli doğruluk sunmaz. Hedef program üniter bir yapıdaysa ek birimsellik (Unitarity) kontrolü yapılır.

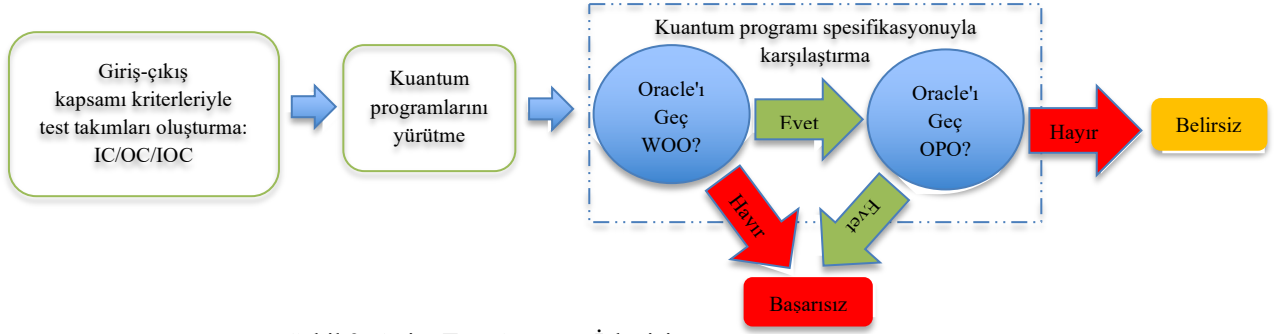
Alt programlar arasında oluşabilecek hataları elimine etmek için son olarak entegrasyon testine tabi tutulur. Bu aşamada Q-UML rehber olarak kullanılır. Klasik entegrasyon testlerinde olduğu gibi entegrasyon sırasını belirlemek için bağımlılık grafiği üzerinden topolojik olarak sıralama yapılır. Entegrasyon esnasında tanımlanmamış rutinler kullanılması gerekiyorsa test ikilisi kullanılır. Bu teknik kuantum programlarını Kahinlerle (Oracle) test etmek için uygundur [9].



Şekil 2. Çok Alt Rutinli Kuantum Programının Test Edilmesine İlişkin Genel Süreç Diyagramı (General Process Diagram for Testing a Quantum Program with Multiple Subroutines)

Programın test kalitesinin objektif bir ölçüsünü sağlamak için Quito tercih edilebilir. Bu araç özelleştirilebilir test senaryoları sunarak, programın hem kuantum (kapı yapılandırmaları) hem de klasik bileşenleri (algoritmalar) arasında tutarlılığı kontrol eder. Geniş bir test kapsama alanı sağlamak için metamorfik test stratejileriyle de uyumludur. Quito, genellikle kuantum devrelerinin beklenen özellikleri karşılayıp karşılamadığını doğrulamak için kullanılır. Dekoherans (Bağlı durumun kopması) etkiler, devre derinliği kısıtlamaları ve kuantum kapısı hataları gibi kuantum bilişimine özgü sorunları tespitinde de faydalanılmaktadır [2]. Şekil 3[2,11,12]'de Quito Test aracının işleyiş şeması verilmiştir. İlk olarak kapsayıcı kriterler ile test takımları oluşturulur. Bir test takımında her geçerli girdi ve çıktı için bir test vardır. Statik olarak oluşturulan test takımı girdi kapsamını (Input Coverage – IC) oluştururken çıktı kapsamı (Output Coverage – OC) ve girdi-çıkı kapsamı (Input -

Output Coverage – IOC) statik olarak elde edilemez[11]. Kapsayıcı kriterlerin kullanılması çok alt rutinli programlarda üstel karmaşıklığa sebep olması nedeniyle tercih edilmez. Girdi olarak verilen programın spesifikasyonları mevcutsa oluşturulan test paketleri iki test kahini kullanılarak değerlendirilir. Yanlış Çıktı Kahini (Wrong Output Oracle - WOO)[11] ve Çıktı Olasılık Kahini (Output Probability Oracle - OPO)[11]'dir. WOO, Test girdisi için döndürülen her çıktı değerinin geçerli olup olmadığı kontrol ederken, OPO ise her IO çifti için oluşturulan test takımları (Test Set - TS) arasında tek örnek Wilcoxon işaretli sıralama testi gerçekleştirir. Hipotez reddedilirse belirsiz, kabul edilirse başarısızlık bildirir.



Şekil 3. Quito Test Aracının İşleyişi (How the Quito Test Tool Works)

Çok alt rutinli kuantum programları için kuantum alt programlarının birleşimini destekleyerek alt rutinlerinin test gereksinimlerini belirlemeye ve yönetmeye yardımcı olması amacıyla çoklu alt rutin Q# programları için bir test aracı (QsharpTester)[9] geliştirilmiştir. Birincil programlama dili olarak Q# dili ile yazılan kuantum programlarının test edilmesini kolaylaştırır. Başlıca odak noktası, hibrit algoritmaların hem kuantum hem de klasik bileşenlerinin düzgün bir şekilde test edilmesini sağlamaktır. [9] yapılan çalışmada yedi orijinal doğru program ve bunların dört mutasyon türüne sahip 244 adet hatalı mutasyon içeren bir dizi Q# karşılaştırma programı kullanılmış ve deney sonucunda neredeyse tüm hatalı sonuçlar tespit edilmiştir. QSharpTester test çerçevesinin scaffold, silQ ve Qiskit gibi kuantum programlama dillerinde de uygulanabileceği önerilmiştir [9].

## 2. Test Teknikleri (Testing Techniques)

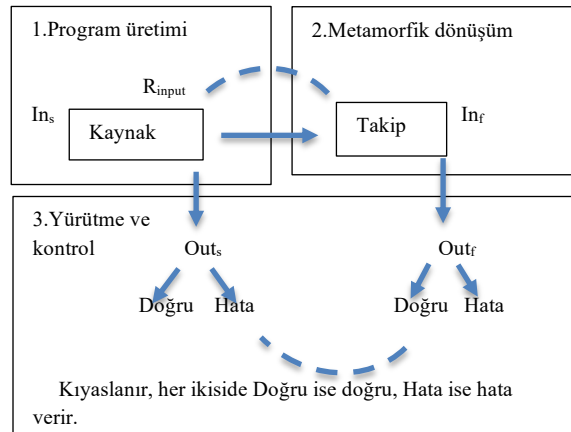
Kuantum yazılım test teknikleri yapı ve işleyiş bakımından çeşitlilik gösterir. Karmaşık kuantum davranışlarının doğruluğunu analiz etmek için metamorfik test tekniklerinden MorphQ kullanımı tercih edilirken, kuantum devrelerinin doğru çalışıp çalışmadığını denetlemek için mutasyon analizinden QMutPy ve Muskit test araçları veya kombinasyonel testlerden QuCAT test aracının kullanımı önerilir. QDiff gibi araçlarsa, farklı kuantum platformları arasında uyum kontrolü sağlayarak platformdan bağımsız testlerin gerçekleştirilmesine olanak

sağlar. Gerçekleştirilmek istenen test çalışmasının ihtiyaçları doğrultusunda birden fazla araç tercih edilebilir. Makalemizin bu bölümünde bahsi geçen teknik ve araçlara yer verilmiştir.

### 2.1. Metamorfik Testler (Metamorphic Tests)

Metamorfik testler, test vakaları arasındaki beklenen dönüşümlerin tutarlı ve mantıksal sonuçlar üretmesini sağlayarak kuantum algoritmalarının doğruluğunu test eder. Her bir çıktının doğruluğunu kontrol etmek yerine, girdilere karşılık gelen çıktılar arasındaki ilişkilerin beklediği gibi olup olmadığına odaklanır. Bu ilişkiler "metamorfik ilişkiler" olarak bilinir. Metamorfik test kavramını, kuantum programlarına uygulamak için MorphQ aracı kullanılmaktadır. On metamorfik ilişkiden oluşan MorphQ; devreyi değiştiren devre dönüşümleri, temsil dönüşümleri ve yürütme dönüşümlerinden oluşur[13].

Şekil 4[13]'de MorphQ ve üç ana adımına ilişkin bir genel bakış sunulmaktadır. İlk olarak, bir program üreticisi kaynak program olarak adlandırılan bir başlangıç kuantum programı ( $In_s$ ) oluşturur. Ardından, bir dizi metamorfik program dönüşümü uygulayarak, yaklaşım kaynak programla belirli bir ilişki içinde olan bir takip programı ( $In_r$ ) üretir. Son olarak, yaklaşım iki programı çalıştırır ve davranışlarının beklenen çıktı ( $Out_s$  ve  $Out_r$ ) ilişkisine uyup uymadığını kontrol eder. Ana döngü, sürekli olarak yeni kaynak ( $R_{input}$ ) çiftleri oluşturur ve kontrolünü sağlar [13].



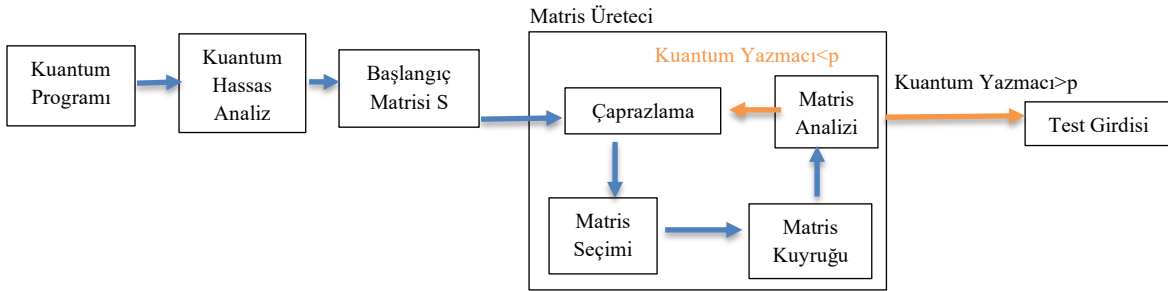
Şekil 4. MorphQ yaklaşımına genel bakış (Overview of the MorphQ approach)

## 2.2. Fuzz Testi (Fuzz Test)

Bulanık test (fuzz testing), klasik sistemlerde bir programa rastgele, beklenmeyen veya geçersiz veriler girerek güvenlik açıklarını ve hataları belirlemek için kullanılan bir yazılım test tekniğidir. Bu yöntem, yazılımın beklenmedik girdilere dayanabilmesini sağlamaya yardımcı olduğu için karmaşık girdi işleme özelliğine sahip programlar için kullanışlıdır. Kuantum sistemlerindeki belirsizlik, Kuantum yazılım testlerinde de olasılıksal sonuçların alınması gibi kısıtlamalara yol açar. Bu durum, fuzz testlerinde sağlıklı sonuçların üretilmesi için geleneksel 'beklenen çıktı' yöntemi yerine olasılıksal hata tespit algoritmalarının tercih edilmesini sağlar. QuanFuzz, kuantuma duyarlı bilgilerin tanımlanıp kuantum kayıtlarının durumunu değiştirerek kapsamı en üst düzeye çıkarmayı amaçlayan GreyBox bulanık test

aracını kullanır[14]. Başlangıç matrislerini mutasyona uğratarak kuantum duyarlı dalları seçmek için yüksek ağırlık değerli matrisleri üretir. Geleneksel test yöntemleriyle karşılaştırıldığında QuanFuzz, özellikle kuantum duyarlı dallarda kapsam %20 – 60 oranında artırmaktadır[14].

Şekil 5[14]'de QuanFuzz'un iş akışı gösterilmiştir. Başlangıçta girdi olarak verilen kuantum programının hassas analizi yapılır. Analiz sonucuna göre başlangıç matrisi oluşturulur. Matris sayısını artırmak için kuantum kapıları kullanılarak matrisler üretilir. Üretilen matrisler olasılık ağırlıklarına göre değerlendirilir. Seçilen matrisler, matris analizlerinin yapılması için sıraya alınır. Kuantum yazmacı olasılık ağırlığı eşik değeri  $p$ 'den büyük olanlar test girdisi olarak seçilir değilse çaprazlama aşamasına geri döndürülür [14].



Şekil 5. QuanFuzz genel iş akışı (QuanFuzz general workflow)

## 2.3. Arama Tabanlı Test (Search Based Testing)

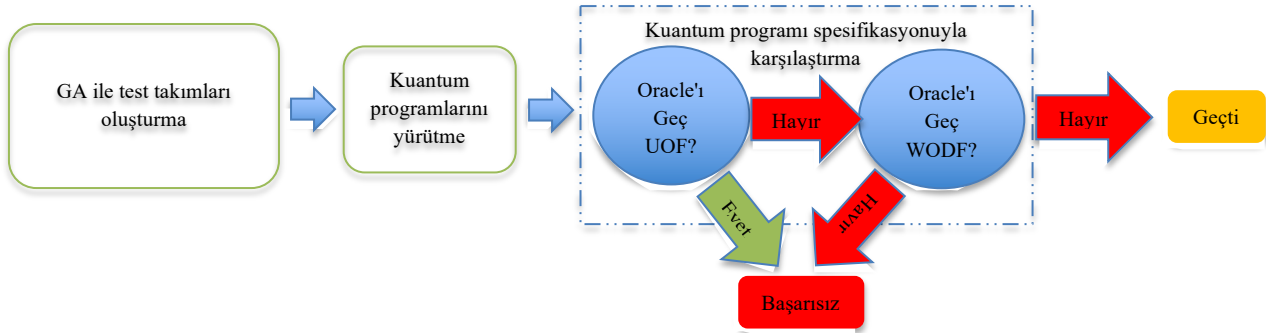
Yazılım test serilerinde genellikle, genetik uygulamalar ve arama tekniklerinden oluşan çeşitli senaryolar kullanılır. Arama tabanlı testler, yazılımın performans göstergeleri ve hataların bulunması için en iyi test durumlarını bulma işlemlerini içerir. Arama tabanlı test, klasik yazılımlarda olduğu gibi kuantum programlarında da mevcuttur. Testler, genellikle kuantum simülatörleri veya kuantum bilgisayarlarında çalıştırılır. Bu, test sürecini daha karmaşık ve maliyetli hale getirmesiyle beraber kuantum yazılımının doğası gereği olasılıksal çıktılarla sonuçlanmasına neden olur. Bu alanda Kuantum Arama Tabanlı Test (Quantum Search-Based Testing - QuSBT)[11] ve Çok Amaçlı Arama Tabanlı Yaklaşım (Multi-Objective Search-Based Approach - MutTG)[11] gibi araçlar kullanılmaktadır.

QuSBT bir genetik algoritma (Genetic Algorithm - GA)[11] kullanarak, belirli kuantum programları için olası hata durumlarını belirlemeye çalışır. Bu araç, programın öngörülen çalışmasına göre hata senaryoları oluşturup başarısız test

durumlarının bulunması için kullanılır. QuSBT'nin temel amacı, kuantum programlarının hata tespitini optimize etmektir. [15] yapılan çalışmada, QuSBT'yi beş kuantum programının on hatalı sürümü ile test edip, ortalama olarak test vakalarının en az %50'sinin başarısız olduğu test takımları üretmeyi başarmıştır. [11] yapılan çalışmada ise QuSBT 30 hatalı kuantum programı ile değerlendirilmiş, programların %87'sinde Rastgele Arama (Random Search - RS)'dan daha iyi performans göstermiştir. Şekil 6[11,15]'de rastgele aramanın baz alındığı QuSBT aracının işleyiş şeması sunulmuştur. İlk olarak, test edilen kuantum programının giriş bilgileri (giriş ve çıkış kubitlerinin listesi, toplam kubit sayısı ve Program Spesifikasyonu (Program Specification - PS)[11] sağlanır. Ardından sağlanan giriş bilgileri doğrultusunda GA yapılandırılmaları gerçekleştirilerek test takımları oluşturulur. Programın yürütme aşamasında ise programın giriş alanı üzerinde yapılan arama, tam sayı değişkenleri ile temsil edilir. Aramanın amacı, başarısız testlerin sayısını en üst düzeye çıkaran bir atama bulmaktır [15].

Test yürütme sonucunun doğruluğu PS karşılaştırmaları ile yapılır. İki hata türü açısından Beklenmeyen Çıkış Hatası (Unexpected Output Failure - UOF)[15] türünde bir hatanın meydana gelip gelmediğini, yani PS'ye göre meydana gelmemesi gereken bir çıktının üretilip üretilmediğini kontrol eder; durum buysa, sonuç başarısız olur ve Yanlış Çıkış Dağıtım Hatası (Wrong Output Distribution Failure - WODF)[15]

için değerlendirme yapılmaz. Aksi durumda ise PS durumunda beklenen dağılımın izlenip izlenilmediğini WODF'ye göre değerlendirir. Bu işlem verilen veya varsayılan önem düzeyini kullanarak Pearson'ın ki-kare testi ile bir uyum iyiliği testi gerçekleştirerek yapılır. İstatistiksel test önemli bir fark gösteriyorsa başarısız, aksi takdirde geçti olarak sonuç üretir [15].



Şekil 6. QuSBT Test Aracının İşleyişi (How the QuSBT Test Tool Works)

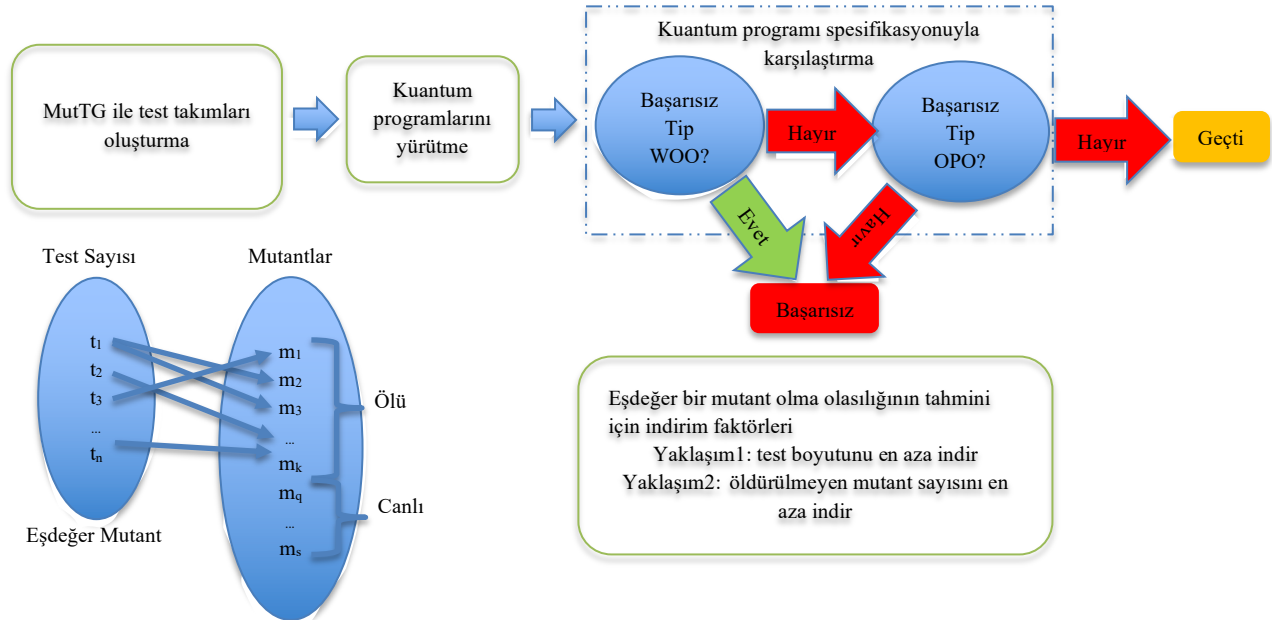
Bir diğeri ise çok amaçlı arama tabanlı yaklaşım (Multi-Objective Search-Based Approach – MutTG)[11]'dir. Kuantum programları için standart tabanlı bir test oluşturur. Şekil 7[11,16]'de MutTG aracının işleyiş şeması sunulmuştur. Kuantum programları için mutasyon analizi, bir kuantum programı (Quantum Program - QP) verildiğinde mutant ( $M$ ), devreye örneğin, yanlış bir kapı kullanımı gibi söz dizimsel bir hata sokularak elde edilen QP'nin biraz farklı bir versiyonu olarak üretilir. Bir testte ( $t$ ) $M$ 'nin yürütülmesi, orijinal program QP üzerinde  $t$ 'nin yürütülmesinden önemli ölçüde farklı sonuçlar üretirse, bir mutantın öldürüldüğü söylenir. Dolayısıyla, bir mutantın öldürülüp öldürülmediğini kontrol etmek için onu PS ile karşılaştırabiliriz. Özetle mutant yürütülmesi iki olası başarısızlıktan birine yol açarsa öldürülür.

Programın yürütülmesi aşamasında değişkenler tanımlanır. Arama sırasında, Muts'taki hangi mutantların öldürüldüğünü kontrol etmek için karşılık gelen testler yürütülür. Arama testlerinde öldürülen mutantlar (Killed Mutant – KM)[16] sözlüğüne, bunların dışındaki mutantlar (Not Killed Mutant – NKM)[16] sözlüğüne girdi olarak alınır. Bu işlemden sonra hedef fonksiyonlar değerlendirilir. İlk hedef, test takımı boyutunu en aza indirmektir. Bunun için öldürülmedi puanı (Not Killed Score - NKS)[16] isimli uygunluk fonksiyonu Denklem. [1,15] kullanılır.

$$= \sum_{M \in Muts} \frac{f_{notkilled}(v)}{notKilledScore(M, v, KM, NKM)} \quad [1]$$

Bu fonksiyon mutant  $M$ 'nin öldürülemez olma olasılığını belirten  $[0, 1]$ 'de tanımlanan bir endekstir. En iyi durum "0" en kötü durum "1" olarak derecelendirilir. Mutant  $M$ , mevcut bireyin bir testi (ilk durum) tarafından öldürülürse, "öldürülmemiş puanı" 0'dır. Mutant  $M$  mevcut bireyin bir testi tarafından öldürülmezse, ancak daha önce oluşturulan testlerden biri tarafından öldürülmüşse (ikinci durum), "öldürülmemiş puanı" 1'dir. Mutant  $M$  mevcut birey tarafından

öldürülmemiş, ancak daha önce oluşturulan testler tarafından öldürülebilir olduğu da bilinmiyorsa (üçüncü durum), "öldürülmemiş puanı" belirli bir indirim faktörünün en kötü değeri 1'den düşürülür. Bu faktör,  $M$ 'yi öldürmeyen testlerin olası girdilerin toplam sayısına oranıdır. Fonksiyon,  $M$  mutantını öldürmeyen ne kadar çok test denersek,  $M$ 'nin öldürülebilir olmayan eşdeğer bir mutant olma olasılığı o kadar artar. Bu indirim faktörünü kullanarak, aslında eşdeğer olan mutantları öldürmeye devam etmekten kaçınma amaçlanmıştır [16]. Test yürütme sonucunun doğruluğu PS karşılaştırmaları ile yapılır.

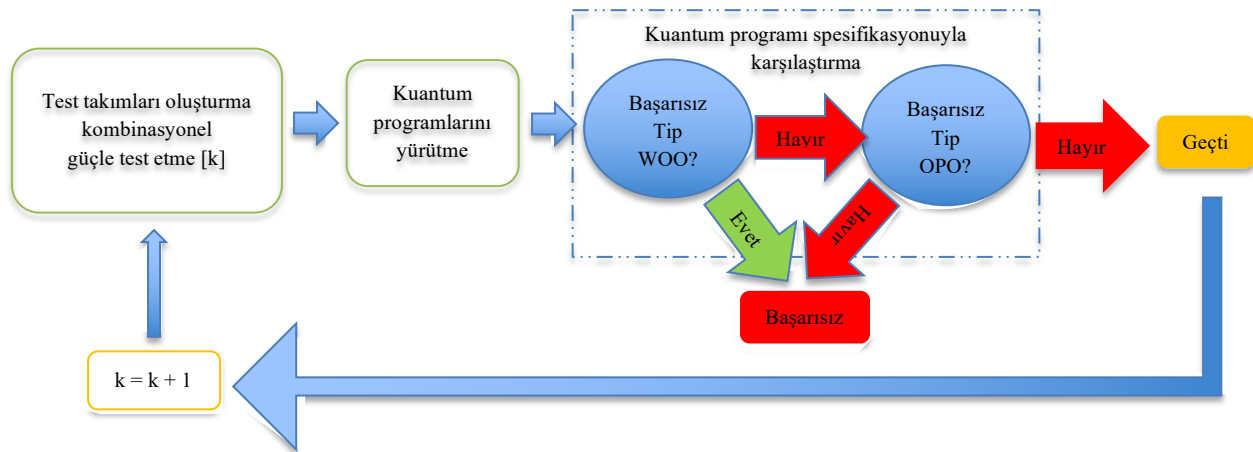


Şekil 7. MutTG Test Aracının İşleyişi (How the MutTG Test Tool Works)

#### 2.4. Kombinasyonel Test (Combinational Testing)

Kombinasyonel test, girdi kombinasyonlarını sistematik bir şekilde kullanarak kuantum programlarındaki hataları tespit eder. Kuantum programlarının doğruluğunu artırmayı hedefler[12]. Bu amaçla geliştirilmiş olan Quantum Kombinasyonel Test(Quantum Combinatorial Testing- QuCAT)[11], kuantum programlarının giriş ve çıkış yöntemlerini analiz eder. İki şekilde kullanılır; bunlardan ilki, kullanıcı sabit bir k değeri belirler ve kombinasyonel test takımı oluşturulur. Bir diğeri ise bir hata bulunana veya maksimum k değerine ulaşılan

kadar k'nın ilk değeri 2 den başlatılarak artırımlı olarak test takımları oluşturulmasıdır. Maksimum k değerine ulaşılan kadar PS geçti olarak sonuçlanır. Şekil 8 [11,17]'de rastgele aramanın baz alındığı QuCAT test aracının işleyiş şeması gösterilmiştir. Testlerin başarı durumu, diğer test araçlarında olduğu gibi PS karşılaştırmaları ile yapılır [11].

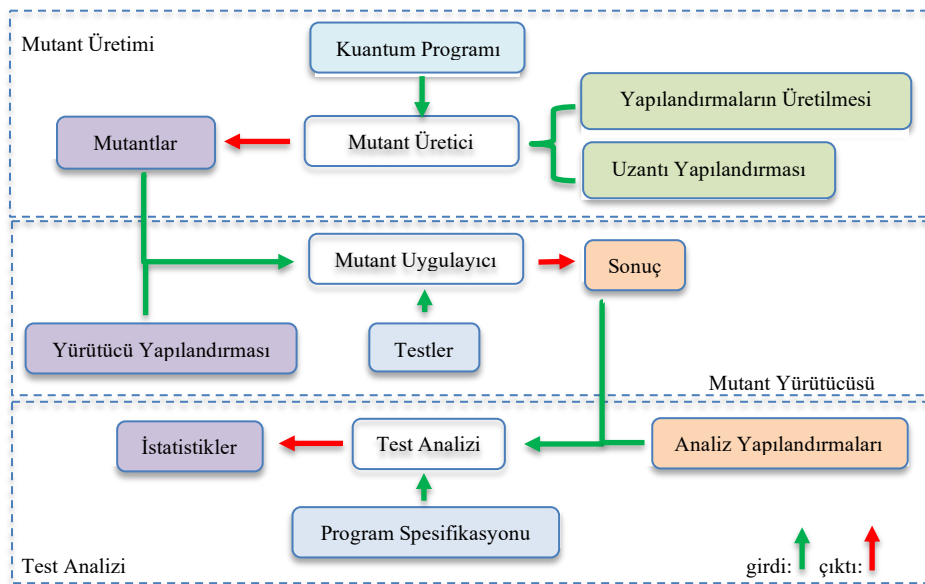


Şekil 8. QuCAT Test Aracının İşleyişi (How QuCAT Test Tool Works)

## 2.5. Kuantum Mutasyon Analizi (*Quantum Mutation Analysis*)

Kuantum programlama bağlamında, test teknikleriyle üretilen test vakalarının kalitesini değerlendirmek için hata depoları ve kıyaslama programları yeterli değildir. Bu duruma alternatif bir çözüm sunmak için mutasyon analizi kullanılmaktadır. Yazılım test süreçlerinde, özellikle de kuantum programları için test kalitesini değerlendirmek amacıyla kullanılan bir yöntemdir. Bir mutantı test vakasında yürüttüğümüzde, QP'nin program spesifikasyonuna (beklenen çıktı dahil) bağlı olarak önceden tanımlanmış kriterlere göre (örneğin, belirli bir girdi için yanlış bir çıktı gözlemlenmesi) başarısız olursa, bu mutantın öldürüldüğü anlamına gelir. Mutasyon skoru genellikle bir test takımının kalitesini değerlendirmek için kullanılır. Burada, test takımı tarafından öldürülen mutantların sayısını toplam mutant sayısından hesaplanır; eşdeğer mutantların sayısı biliniyorsa, hesaplama sırasında toplam mutant sayısından çıkarılır. Bu bağlamda iki önemli araç bulunmaktadır. İlki kuantum yazılım testi için bir mutasyon analizi aracı Muskit'tir[11]. Kuantum devrelerinde bulunan kapıları (Gates) hedef alarak çeşitli mutasyon operatörleri tanımlar. Muskit, iki ana bileşene sahiptir: Mutant Üretici(Mutant Producer) ve Mutant Yürütücü(Mutant Executor). Mutant üretici, belirli bir kuantum programı için mutantlar üretirken, mutant yürütücü, bu mutantlar üzerinde testleri çalıştırarak sonuçları değerlendirir.

Şekil 9[11,18]'da işleyiş diyagramı sunulmuştur. Muskit, kuantum devrelerinin mutasyona uğramış özelliklerine odaklanır. Bu amaçla, iki kavram tanımlar: kapı numarası ve konum. Kapı numarası, silme veya değiştirme operatörleri aracılığıyla mutasyona uğratmak istediğimiz bir kuantum devresindeki belirli bir kapıyı (örneğin, h, G1'dir) ifade eder. Konum ise kuantum devresinde mutasyona uğratmak istediğimiz yere, yeni bir kapı ekleyerek belirlenen yeri ifade eder. Üç operatör türüne göre mutasyon operatörü tanımlanır. Bunlar; Kapı Ekleme(Add Gate - AG)[18], Kapı Kaldırma(Remove Gate - RemG)[18] ve Kapı Değiştirme(Replace Gate - RepG)[18]. Giriş kubit sayısına göre kullanıcı bu operatörlerden bir veya daha fazlasını seçerek mutant üretimini kontrol eder. Mutant yöneticisi girdi olarak mutantlar ve test vakalarını alır, mutantlar üzerinde test vakalarını yürütür ve sonuç üretir. QP'lerin olasılıksal doğası göz önüne alındığında, her mutant belirtilen sayıda test vakasıyla yürütülür. Test Analizcisi(Test Analyzer), kullanıcının test değerlendirme kriterlerini uygular. Burada kullanıcı üç şeyi belirtir: ilki seçilen p değeri düzeyi (örneğin 0,05), diğeri girdi olarak kullanılan kubit kimlikleri ve son olarak ölçülecek kubit kimlikleridir. Bir program spesifikasyonu (örneğin, beklenen çıktılar ve her girdiye karşılık gelen olasılıklar) ve mutant yürütücü'den test sonuçları verildiğinde, test analizcisi bir mutantın bir test vakası tarafından öldürülüp öldürülmediğini söyleyebilir [18].



Şekil 9. Muskit Test Aracının İşleyişi (How Muskit Test Vehicle Works)

Bir diğer aracı ise QMutPy'dır. MutPy adlı açık kaynaklı mutasyon aracının kuantum programları için genişletilmiş bir versiyonudur[19]. QMutPy, kuantum ölçümleri ve kapıları temel alarak yeni mutasyon operatörleri oluşturur. IBM'in Kuantum Bilgi Yazılım Kit (Quantum Information Software Kit –Qiskit)[19] kütüphanesi kullanılarak yazılmış gerçek kuantum programları üzerinde testler gerçekleştirilmiştir. QMutPy'nin iş akışı MutPy benzer şekilde dört ana adımdan oluşur. Bir Python programı (P), onun test takımı (T) ve bir mutasyon operatörü kümesi (M) verildiğinde, QMutPy öncelikle P'nin kaynak kodunu ve test takımını yükler ardından T'yi orijinal (değiştirilmemiş) kaynak kodunda yürütür ve M'yi uygular. Son olarak P'nin tüm mutant sürümlerini üreterek T'yi her mutant sürümünde yürütür [19].

[20]'de, IBM'in Qiskit kütüphanesinde yazılmış, kuantum mutasyon operatörleri setinin 11 QP için toplam 696 mutant üretilmiş ve bunların 325'i (%46,7) programların test paketleri tarafından öldürülmüştür. Öldürülmeyen mutantlar ya test paketlerine kadar hayatta kalmış (307, %44,1) ya da test paketleri tarafından kullanılmamış (%0,3) veya zaman aşımına uğramıştır (62, %8,9). Elde edilen sonuçlar bu alanda QMutPy gibi araçların önemli bir rol oynadığını göstermektedir. Bunun yanı sıra Muskit test aracının şu anda eşdeğer mutantları tespit edemediği. Bunun başlıca nedeninin, mevcut bir gövdenin olmadığı belirtilmiştir[18].

## 2.6. Kuantum Platform Testi (*Quantum Platform Test*)

Kuantum mekaniğinin fiziksel özellikleri nedeniyle, kubitler ve kuantum kapıları klasik bitlerden ve kapı mantığından temelde farklıdır. Birçok kuantum yazılım yığını, alta yatan fiziksel ve matematiksel karmaşıklıkları soyutlayarak kuantum programlama için kullanıcı dostu üst düzey diller geliştirmiştir. Bunlar arasında CirQ, PyQuil ve Qiskit dilleri sayılabilir. Kuantum yazılım yığını (Quantum Software Stack – QSS)[18]; API'ler aracılığıyla verilen kuantum algoritmasını devre düzeyinde dönüştüren ve optimize eden bir derleyici ile ortaya çıkan kapıları klasik aygıtlarda simüle eden veya doğrudan kuantum donanımında yürüten bir arka uç yürütücü içerir. Programları derleyen ve optimize eden Qiskit Terra, yüksek performanslı gürültülü simülasyonları destekleyen Qiskit Aer, hata düzeltme, gürültü karakterizasyonu ve donanım doğrulaması için Qiskit Ignis, bir geliştiricinin bir kuantum algoritmasını veya uygulamayı ifade etmesine yardımcı olan Qiskit Aqua olmak üzere dört bileşenden oluşan Qiskit örnek verilebilir.

QS'i test etmeyi zorlaştıran başlıca etkenler mevcuttur. İlki çok platformlu program çevirmenidir(Multi-Platform Program Translator). Kuantum programları genellikle yeni programlama dillerinde yazılır veya mevcut dillerin üstündeki API'ler kullanılarak ifade edilir. Standartlaştırılmış ara gösterimlerin yokluğu, platformlar arası test için bir zorluk oluşturur. Bir diğer zorluk, kuantum program üretimidir. Platform testi çok sayıda program gerektirir. Ancak günümüzde yalnızca birkaç gerçek kuantum program örneği mevcuttur. Son olarak Çok Değişkenli İkili Dağılım Karşılaştırması(Multivariate Binary Distribution Comparison), kuantum ölçümlerinin olasılıksal doğası, çıktıların dağılımlarla temsil edilmesiyle karmaşıklık oluşturur. Bu zorlukla ilgili olarak kuantum hata ayıklamada Kolmogorov-Smirnov (KS) veya Çapraz Entropi (Cross Entropy) testleri kullanılarak yapılan ön çalışmalarla birlikte araştırmalar yürütülmüştür [21]. QSS testlerinde meydana gelen zorlukları aşmak adına çeşitli platform testleri geliştirilmiştir. Quantum Yazılım Yığınlarının Farklı Testleri (Differential Testing of Quantum Software Stacks - QDiff)[21] bunlardan biridir. QDiff, girdi olarak kuantum programı alır ve bir çift tanık programla (yani, aynı davranışı üretmesi beklenen mantıksal olarak eşdeğer program) olası hataları bildirir.

Bir başka zorluk, kuantum simülatörlerinden veya donanımdan gelen ölçümlerin yorumlanmasıdır. Kuantum ölçüm sonuçları olasılıksal değerler üretir. Bunun için içsel gürültüyü (örneğin, donanım kapısı hataları, okuma hataları ve dekoherans hataları) hesaba katmalı ve gözlemlenen sapmanın anlamlı bir kararsızlık olarak kabul edilebilecek kadar önemli olup olmadığı belirlenmelidir. Bu, her kuantum kapısının matematiksel olarak bir üniter matrislerle temsil edilebileceği anlayışına dayanır; bir kapı dizisi esasen üniter matrislerinin çarpımına karşılık gelir ve bu da bir üniter matris üretir. Bu nedenle, iki dizi aynı üniter matrisi üretirse, bir kapı dizisi diğerine anlamsal olarak eşdeğerdir. Aynı üniter matrisi üretmesi garantili farklı kapı dizileri üretmek için yedi kapı dönüşüm kuralından yararlanılır. Bu diziler esasen aynı davranışı üretmesi beklenen mantıksal olarak eşdeğer program varyantlarını tanımlar. Üretilen birden fazla mantıksal eşdeğer varyantlardan çalıştırılmaya değer olan devrelerin bir alt kümesi seçilir. Son olarak mantıksal eşdeğer devrelerin yürütmelerini karşılaştırmak için QDiff, güvenilir karşılaştırma için ne kadar ölçümün gerektiğini belirler. Ölçüm sayısını tahmin etmek için yakınlık testini uygular. Ardından gereken sayıda ölçümü gerçekleştirir. İki ölçüm kümesini karşılaştırmak için QDiff, dağıtım karşılaştırma

yöntemlerini kullanır. KS testine ve çapraz entropiye dayalı iki yöntem desteklenmektedir [21].

### 3. Kuantum Programları İçin Hata Tespiti (Error Detection for Quantum Programs)

Kuantum programlamadaki güncel araştırmalar esas olarak sorun analizi, dil tasarımı ve uygulamaya odaklanmaktadır. Kuantum yazılım geliştirme süreçlerini iyileştirmek, araştırmacılara gerçek hataları inceleme imkanı sağlar. Fakat, kuantum programlarındaki hataların ayıklanması konusu kuantum programlama paradigmasında çok az ilgi görmüştür. Kuantum programlamada tanıtılan üst üste binme, dolanıklık gibi belirli özellikler, kuantum programlarındaki hataları bulmayı zorlaştırmaktadır. Kuantum yazılımlarını hata ayıklamak ve test etmek için çeşitli yaklaşımlar önerilmiştir [22].

İlki Qbugs'tır. Qbugs, kuantum yazılım testinde ve hata ayıklamada kullanılmak üzere tasarlanmış bir altyapı sağlar. Farklı kuantum programlama dilleri (Q#, OpenQASM, Cirq, Quipper ve Scaffold) için desteklenir. Qbugs'ın bir prototipini oluşturmak için kuantum çerçeve depolarında bulunan kuantum algoritmalarının açık kaynaklı uygulamaları kullanılmıştır.

Bunlar; ProjectQ'nun çerçeve deposu, QiskitAqua'nın deposu ve O'Reilly'nin Kuantum Bilgisayarları Programlama" kitabının deposudur. Hangi hataların bildirildiğini otomatik olarak tespit edip gerekli düzeltmenin yapılması için Defects4J veritabanı ve BugsJS veritabanı kullanılmıştır [23].

Bir diğeri ise Bugs4Q'dur. Bu araç Qiskit programlarındaki yeniden üretilebilir hataları toplar ve kuantum yazılım testleri için test durumlarını indirmeyi ve çalıştırmayı destekler. Her gerçek hata ve karşılık gelen düzeltmeler erişime açıktır. Qiskit'in GitHub'daki mevcut hatalarının neredeyse tamamını toplar ve dört popüler Qiskit ögesinin (Terra, Aer, Ignis ve Aqua) gerçek zamanlı olarak günceller. Ayrıca, bu programlar orijinal olarak mevcut test durumları ve yeniden üretim desteği olan hatalar haricinde ayrı ayrı sıralanır ve filtrelenir. Bugs4Q, izole edilmiş hataların deneysel değerlendirmesi için mevcut hataları sınıflandırmak üzere hata türlerinin analizini içeren bir veritabanı sağlaması yönüyle kuantum programlarındaki gerçek hataların bir kataloğudur [22].

Tablo 1.'de Test araç ve tekniklerinin avantaj, dezavantaj ve performans değerlendirmeleri özetlenmiştir. (Advantages, disadvantages and performance evaluations of testing tools and techniques are summarized.)

| Kategori                 | Araç         | Avantajlar                                                                                             | Dezavantajlar                                                                                                                                              | Performans Değerlendirmesi                                                                                                                                 | Referanslar |
|--------------------------|--------------|--------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------|-------------|
| Kapsayıcı Kriter         | Quito        | Kuantum algoritmaları için özelleştirilmiş test çerçevesi sunar.                                       | Yalnızca Qiskit'te kodlanmış kuantum programlarının test edilmesini desteklemektedir.                                                                      | Sınırlı sayıda ölçüm yapılmıştır.                                                                                                                          | [2]         |
| Kapsayıcı Kriter         | QsharpTester | Simülatör desteği ile fiziksel kuantum sistemine ihtiyaç duymadan çalıştırma testi yapılabilir.        | Yürütme verimliliklerinin karşılaştırılabileceği yeterli sayıda çalışma henüz yapılmamıştır.                                                               | Sınırlı sayıda ölçüm yapılmıştır.                                                                                                                          | [9]         |
| Kuantum Mutasyon Analizi | Muskit       | -Kuantum kapılarına dayalı mutasyon testi sağlar.<br>- Özelleştirilebilir mutasyon operatörleri sunar. | Şimdilik eşdeğer mutantları tespit edemiyor.                                                                                                               | Sınırlı sayıda ölçüm yapılmıştır.                                                                                                                          | [18]        |
| Arama Tabanlı Test       | QuSBT        | Genetik algoritmalar kullanarak test senaryosu üretir. Bu da test çeşitliliğini artırır.               | Arama süresi asgari düzeyde olsa da, QuSBT'nin zamanının çoğunu simülatörü kullanarak QP'ler üzerinde test vakalarını yürütmeye harcadığını görmüştür[15]. | Sınırlı sayıda ölçüm yapılmıştır. [15] çalışmada, oluşturulan test takımında başarısız test vakalarının en az %50'sini bulmayı başardığı rapor edilmiştir. | [15]        |
| Kombinasyonel Test       | QuCAT        | Kuantum devrelerinin hata toleransını ölçen simülasyon tabanlı testler sunar.                          | Yüksek simülasyon maliyeti, büyük devrelerde performans kısıtlamalarını beraberinde getirir.                                                               | Sınırlı sayıda ölçüm yapılmıştır.                                                                                                                          | [17]        |
| Kuantum Platform Testi   | QDiff        | - Kuantum platformları arasında uyum kontrolü sağlar.                                                  | Yeterli sayıda çalışma henüz yapılmamıştır.                                                                                                                | Varyantları oluşturmada etkilidir, gereksiz kuantum donanım veya                                                                                           | [21]        |

|                          |          |                                                                                               |                                                                  |                                                                                                                   |      |
|--------------------------|----------|-----------------------------------------------------------------------------------------------|------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------|------|
|                          |          | - Farklı diller için kıyaslama yapar.                                                         |                                                                  | gürültülü simülasyon çağrılarını %66 oranında azaltır [21].                                                       |      |
| Kuantum Mutasyon Analizi | QMutPy   | - Mutasyon testleri için esnek ve genişletilebilir.<br>- Qiskit ile entegrasyon sağlanmıştır. | -Sadece belirli bir platform için optimize edilmiştir.           | Yüksek mutasyon skoru ile etkili sonuçlar sağlar, ancak platform bağımsız çalışmaması nedeniyle sınırlıdır.       | [19] |
| Fuzz Testi               | QuanFuzz | Hata tespiti için rastgele giriş yapar.                                                       | Rastgele testler bazen zayıf hata kapsamı oluşturabilir.         | Klasik test modellerine göre %20-60 daha fazla dal kapsamı elde etmektedir[14].                                   | [14] |
| Metamorfik Testler       | MorphQ   | Hata tespitinde kullanılır.                                                                   | Değişken özelliklerin gözlemi zordur.                            | Sınırlı sayıda ölçüm yapılmıştır.                                                                                 | [13] |
| Hata Tespiti             | Qbugs    | Hata tespitinde kullanılır.                                                                   | Karmaşık kuantum algoritmaları için sağladığı destek sınırlıdır. | Kuantum hesaplamalarında tekrarlanabilirliği kolaylaştırmak için henüz iyi tanımlanmış bir ölçüt bulunmamaktadır. | [23] |

Tablo 1., kuantum yazılım test araçlarının test senaryolarındaki performanslarını anlamak ve araçları seçerken avantajlarını ve sınırlılıklarını değerlendirmek için kullanılabilir. Her aracın kullanım alanları farklı olduğundan, ihtiyaçlara ve uygulama alanına göre seçim yapmak en iyi sonuçların elde edilmesi açısından katkı

sağlayacaktır. [24] Bu çalışmada, kuantum yazılım testindeki durumu kapsamlı bir şekilde görüntülemek için, seçilen yayınlar üzerinde bir SWOT analizi hazırlanmıştır. Tablo 2.[24], kuantum yazılım testinin mevcut durumunu ve gelişim potansiyelini anlamak açısından yararlıdır

Tablo 2. Kuantum Yazılım Test ve tekniklerinin SWOT Analizi (*SWOT Analysis of Quantum Software Testing and techniques.*)

| Kategori     | Açıklama                                                                                                                                                                                                                                                                                                                                                                                              |
|--------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Güçlü Yönler | - <b>Yüksek Hata Kapsamı:</b> Kuantum test yöntemleri, hataları tespit etme yeteneğine sahiptir.<br>- <b>Güvenilirlik:</b> Test yöntemleri, yazılımın doğru çalışmasını sağlamak için güvenilir bir çerçeveye sunar.<br>- <b>Platformlar Arası Uyum:</b> Araçlar, farklı kuantum donanım platformları arasındaki uyumluluğu test edebilir.                                                            |
| Zayıf Yönler | - <b>Yüksek Hesaplama Maliyeti:</b> Testler, büyük kuantum devrelerinde çalıştırıldığında yoğun hesaplama kaynakları gerektirir.<br>- <b>Olgunlaşmamış Teknoloji:</b> Çoğu test aracı gelişim aşamasındadır ve henüz tam anlamıyla olgunlaşmamıştır.<br>- <b>Yetersiz Standartlar:</b> Kuantum testleri için endüstri standartlarının olmaması uygulamayı zorlaştırır.                                |
| Fırsatlar    | - <b>Araştırma ve Geliştirme:</b> Kuantum test yöntemleri, akademik ve endüstriyel araştırmalara yatırım çekmektedir.<br>- <b>Gelişen Platformlar:</b> Kuantum donanım ve simülasyonlarındaki ilerlemeler, test süreçlerini iyileştirme fırsatları sunar.<br>- <b>Yeni Uygulama Alanları:</b> Kuantum yazılımlarının finans, kimya, lojistik gibi alanlarda artan uygulama potansiyeli bulunmaktadır. |
| Tehditler    | - <b>Klasik Bilgi İşlem ile Rekabet:</b> Kuantum yazılımın yavaş gelişimi, klasik sistemlerle rekabeti zorlaştırmaktadır.<br>- <b>Güvenlik Riskleri:</b> Henüz tespit edilemeyen hatalar, yazılım güvenliği için tehdit oluşturabilir.<br>- <b>Hızla Değişen Teknoloji:</b> Kuantum teknolojilerindeki hızlı değişim, mevcut test araçlarının hızla eskimesine yol açabilir.                          |

### Kuantum Test Mühendisliğinde Standartlaşma Problemi (*The Problem Of Standardization In Quantum Test Engineering*)

Kuantum test mühendisliği, klasik yazılım test mühendisliğine benzer prensipler içerse de, kuantum hesaplama ortamının getirdiği özel durumlar (belirsizlik, süperpozisyon ve doğrulama karmaşıklıkları vs.) nedeniyle klasik standartların doğrudan uygulanabilirliği sınırlıdır. Kuantum test mühendisliği henüz klasik yazılım test mühendisliği

kadar yerleşik bir standart setine sahip değildir. Ancak, kuantum yazılım ve sistemlerinin test edilmesi için önerilen yaklaşımlar ve gelişmekte olan yöntemler bulunmaktadır. Bu bağlamda, kuantum test mühendisliği standartları; büyük ölçüde akademik araştırmalara, uygulama odaklı önerilere ve mevcut klasik test standartlarının uyarlanmasına dayanmaktadır. Tablo 3.'de Klasik test mühendisliğinde yer alan bazı standartların kuantum test mühendisliğine ne ölçüde uyarlanabileceği değerlendirilmiştir.

Tablo 3. Klasik test mühendisliği standartlarının kuantum test mühendisliğine uygulanabilirliği (*Applicability of classical test engineering standards to quantum test engineering*)

| Standart      | Kuantum Testine Uygulanabilirlik                                                                                                                                             | Kaynaklar |
|---------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----------|
| ISO/IEC 29119 | Kısmen uygulanabilir, dinamik test süreçlerinden test durum spesifikasyonları kuantum spesifikasyonlarına göre uyarlanabilir.                                                | [25]      |
| IEEE 829      | Kısmen uygulanabilir, test tasarım ve vaka spesifikasyonlarının kuantum durumları (dolanıklık, süperpozisyon vs.) ve spesifikasyonları göz önüne alınarak revize edilebilir. | [26]      |
| TMMi          | Kısmen uygulanabilir, TMMi olgunluk seviyeleri, Kuantum test ortamının olgunluk seviyelerine göre tanımlanabilir.                                                            | [27]      |

Kuantum test mühendisliğinde, klasik test standartları çerçevesinde yeni model ve metodolojilerin geliştirme çalışmalarının yaygınlaştırılması, bu alanda uygulanabilirliğin artırılması için önem arz etmektedir. Quito, QuCAT, QSharpTester gibi test araçları; klasik yazılım test standartlarına uyumlu ve pratikte kullanılabilir olması sebebiyle standartlaşma potansiyeli taşımaktadır. Her bir test aracının odak noktası farklıdır. Tek bir standart setinin oluşturulabilmesi için kuantum yazılımının kullanılabilirliğinin artırılması, kuantum test mühendisliği alanında kullanılan araç ve metodolojilerin standartlaşmasına katkı sağlayacaktır.

#### ÖNERİLER VE TARTIŞMA (*SUGGESTIONS AND DISCUSSION*)

Kuantum programlarına yönelik test yöntem ve araçları, gelişmekte olan kuantum yazılım mühendisliği alanı için büyük önem taşımaktadır. Kuantum programlarının temel prensipleri (süperpozisyon, dolanıklık gibi) bu alanda klasik yazılım test yaklaşımlarını yetersiz kılmaktadır. Bu nedenle, yeni ve inovatif test yöntemleri gereklidir. Kuantum algoritmalarının hataya çok hassas olması nedeniyle, bu test araçları yazılımların doğruluğunu ve güvenilirliğini artırmada etkilidir. Özellikle metamorfik testler, karmaşık kuantum davranışlarının doğruluğunu analiz etmeye yardımcı olur. Kuantum programlarının doğruluğu hataların erken tespiti ile sağlanabilir. Kuantum test araçları, mutasyon analizi veya kombinasyonel testler ile kuantum devrelerinin doğru çalışıp çalışmadığı denetlenebilir. QDiff gibi araçlarsa, farklı kuantum platformları arasında uyum kontrolü sağlayarak, farklı kuantum işlemcilerindeki devre performansını kıyaslamaya yardımcı olur. Bu sayede platformdan bağımsız yazılım geliştirmeye olanak tanır. Quito aracı belirli kuantum test

senaryolarında etkin bir test ortamı sunarken, QMutPy platform bağımlılığı nedeniyle sınırlı kalmaktadır. Kuantum mekaniksel prensipleri temel alan test yöntemleri, oldukça karmaşıktır ve doğru şekilde uygulamak için ileri düzey bilgi gerektirir. Birçok kuantum test aracı, karmaşık devrelerde veya geniş veri kümelerinde çalıştırıldığında yoğun hesaplama gücü gerektirir. Bu durum, büyük ve karmaşık algoritmaların test edilmesini zorlaştırmaktadır. QSharpTester gibi araçlar belirli diller için optimize edilmiştir. Bu, genel bir çözüm sunmak yerine sadece belirli platformlar için kullanılabilirlik sağlamaktadır. Bu çalışma, kuantum yazılım test teknikleri arasında belirgin avantaj ve dezavantajlar olduğunu göstermektedir.

Kuantum mühendisliğinde resmi standartların olmamasının temel sebebi, kuantum yazılımının çoğunlukla deneysel ve akademik ortamda geliştirilmeye devam ediyor olmasıdır. Endüstriyel kullanımın yaygınlaşması; akademik ve sektörel işbirliğinin sağlanarak küresel standartların oluşturulması için kullanım alanına özel araç ve metodolojinin geliştirilmesi gerekmektedir.

Özetle, kuantum programlarının test edilmesi, hem teorik hem de pratik açıdan birçok zorluğu beraberinde getirir. Fakat, bu çalışmada bahsedilen araç ve yöntemler, güvenilir kuantum yazılım geliştirme açısından vazgeçilmezdir. Araştırmacılar, bu yöntemleri geliştirerek testlerin verimliliğini artırmayı hedeflemektedir. Test süreçleri iyileştirildiğinde, kuantum algoritmalarının günlük hayatta daha yaygın hale geleceği ve ortak standartların oluşturulacağı öngörülmektedir. Gelecek çalışmalar, kuantum bilgisayarlarına yönelik olarak geliştirilen yazılımları, makalemizde incelediğimiz araç ve yöntemlerle sınavarak deneysel bir karşılaştırma yapmayı amaçlamaktadır.

#### KAYNAKÇA (*Source*)

[1] ISTQB not-for-profit association [Internet]. 2024 [a.yer 17 Aralık 2024]. Managing the Test Team. Erişim adresi: <https://www.istqb.org/managing-the-test-team>

[2] Wang X, Arcaini P, Yue T, Ali S. Quito: a Coverage-Guided Test Generator for Quantum Programs. İçinde: 2021 36th IEEE/ACM International Conference on Automated Software Engineering (ASE) [Internet]. Melbourne, Australia:

IEEE; 2021 [a.yer 26 Aralık 2024]. s. 1237-41. Erişim adresi: <https://ieeexplore.ieee.org/document/9678798/>

[3] What Is Quantum Computing? | IBM [İnternet]. 2024 [a.yer 07 Şubat 2025]. Erişim adresi: <https://www.ibm.com/think/topics/quantum-computing>

[4] Ali S, Yue T. Quantum Software Testing: A Brief Introduction. İçinde: 2023 IEEE/ACM 45th International Conference on Software Engineering: Companion Proceedings (ICSE-Companion) [İnternet]. 2023 [a.yer 30 Kasım 2024]. s. 332-3. Erişim adresi: <https://ieeexplore.ieee.org/document/10172759>

[5] McClean JR, Romero J, Babbush R, Aspuru-Guzik A. The theory of variational hybrid quantum-classical algorithms. New J Phys. Şubat 2016;18(2):023023.

[6] Li G, Zhou L, Yu N, Ding Y, Ying M, Xie Y. Projection-based runtime assertions for testing and debugging Quantum programs. Proc ACM Program Lang. 13 Kasım 2020;4(OOPSLA):1-29.

[7] Huang Y, Martonosi M. Statistical Assertions for Validating Patterns and Finding Bugs in Quantum Programs [İnternet]. arXiv; 2019 [a.yer 30 Kasım 2024]. Erişim adresi: <http://arxiv.org/abs/1905.09721>

[8] Abreu R, Fernandes JP, Llana L, Tavares G. Metamorphic testing of oracle quantum programs: 3rd IEEE/ACM International Workshop on Quantum Software Engineering, Q-SE 2022. Proc - 3rd Int Workshop Quantum Softw Eng Q-SE 2022. 2022;16-23.

[9] Long P, Zhao J. Testing Quantum Programs with Multiple Subroutines [İnternet]. arXiv; 2023 [a.yer 30 Kasım 2024]. Erişim adresi: <http://arxiv.org/abs/2208.09206>

[10] Pérez-Delgado CA. A Quantum Software Modeling Language. İçinde: Serrano MA, Pérez-Castillo R, Piattini M, editörler. Quantum Software Engineering [İnternet]. Cham: Springer International Publishing; 2022 [a.yer 07 Şubat 2025]. s. 103-19. Erişim adresi: [https://link.springer.com/10.1007/978-3-031-05324-5\\_6](https://link.springer.com/10.1007/978-3-031-05324-5_6)

[11] Quantum Software Testing: A Brief Introduction [İnternet]. [a.yer 19 Nisan 2025]. Erişim adresi: <https://www.simula.no/research/quantum-software-testing-brief-introduction>

[12] Wang X, Arcaini P, Yue T, Ali S. Application

of Combinatorial Testing to Quantum Programs. İçinde: 2021 IEEE 21st International Conference on Software Quality, Reliability and Security (QRS) [İnternet]. Hainan, China: IEEE; 2021 [a.yer 30 Kasım 2024]. s. 179-88. Erişim adresi: <https://ieeexplore.ieee.org/document/9724888/>

[13] Paltenghi M, Pradel M. MorphQ: Metamorphic Testing of the Qiskit Quantum Computing Platform [İnternet]. arXiv; 2023 [a.yer 30 Kasım 2024]. Erişim adresi: <http://arxiv.org/abs/2206.01111>

[14] Wang J, Gao M, Jiang Y, Lou J, Gao Y, Zhang D, vd. QuanFuzz: Fuzz Testing of Quantum Program [İnternet]. arXiv; 2018 [a.yer 30 Kasım 2024]. Erişim adresi: <http://arxiv.org/abs/1810.10310>

[15] Wang X, Arcaini P, Yue T, Ali S. QuSBT: Search-Based Testing of Quantum Programs [İnternet]. arXiv; 2022 [a.yer 30 Kasım 2024]. Erişim adresi: <http://arxiv.org/abs/2204.08561>

[16] Wang X, Yu T, Arcaini P, Yue T, Ali S. Mutation-based test generation for quantum programs with multi-objective search. İçinde: Proceedings of the Genetic and Evolutionary Computation Conference [İnternet]. Boston Massachusetts: ACM; 2022 [a.yer 30 Kasım 2024]. s. 1345-53. Erişim adresi: <https://dl.acm.org/doi/10.1145/3512290.3528869>

[17] Wang X, Arcaini P, Yue T, Ali S. QuCAT: A Combinatorial Testing Tool for Quantum Software [İnternet]. arXiv; 2023 [a.yer 30 Kasım 2024]. Erişim adresi: <http://arxiv.org/abs/2309.00119>

[18] Mendiluze E, Ali S, Arcaini P, Yue T. Muskit: A Mutation Analysis Tool for Quantum Software Testing. İçinde: 2021 36th IEEE/ACM International Conference on Automated Software Engineering (ASE) [İnternet]. Melbourne, Australia: IEEE; 2021 [a.yer 30 Kasım 2024]. s. 1266-70. Erişim adresi: <https://ieeexplore.ieee.org/document/9678563/>

[19] Fortunato D, CAMPOS J, ABREU R. Mutation Testing of Quantum Programs: A Case Study With Qiskit. IEEE Trans Quantum Eng. 2022;3:1-17.

[20] Fortunato D, Campos J, Abreu R. QMutPy: a mutation testing tool for Quantum algorithms and applications in Qiskit. İçinde: Proceedings of the 31st ACM SIGSOFT International Symposium on Software Testing and Analysis [İnternet]. Virtual South Korea: ACM; 2022 [a.yer 13 Mart 2025]. s. 797-800. Erişim adresi: <https://dl.acm.org/doi/10.1145/3533767.3543296>

[21] Wang J, Zhang Q, Xu GH, Kim M. QDiff:

Differential Testing of Quantum Software Stacks. İçinde: 2021 36th IEEE/ACM International Conference on Automated Software Engineering (ASE) [Internet]. Melbourne, Australia: IEEE; 2021 [a.yer 30 Kasım 2024]. s. 692-704. Erişim adresi: <https://ieeexplore.ieee.org/document/9678792/>

[22] Zhao P, Zhao J, Miao Z, Lan S. Bugs4Q: A Benchmark of Real Bugs for Quantum Programs [Internet]. arXiv; 2021 [a.yer 30 Kasım 2024]. Erişim adresi: <http://arxiv.org/abs/2108.09744>

[23] Campos J, Souto A. QBugs: A Collection of Reproducible Bugs in Quantum Algorithms and a Supporting Infrastructure to Enable Controlled Quantum Software Testing and Debugging Experiments [Internet]. arXiv; 2021 [a.yer 30 Kasım 2024]. Erişim adresi: <http://arxiv.org/abs/2103.16968>

[24] García De La Barrera A, García-Rodríguez De

Guzmán I, Polo M, Piattini M. Quantum software testing: State of the art. J Softw Evol Process. Nisan 2023;35(4):e2419.

[25] ISO/IEC/IEEE 29119-3:2021(en), Software and systems engineering — Software testing — Part 3: Test documentation [Internet]. 2024 [a.yer 16 Aralık 2024]. Erişim adresi: <https://www.iso.org/obp/ui/en/#iso:std:iso-iec-ieee:29119:-3:ed-2:v1:en>

[26] IEEE Standard for Software Test Documentation. IEEE Std 829-1998. Aralık 1998;1-64.

[27] van Veenendaal E. Produced by the TMMi Foundation.



## Domates Yapraklarında Hastalık Tespiti İçin Transfer Öğrenme Kullanılması Ve Mobil Uygulamaya Entegre Edilmesi

Tahir Çağrı ÖZBEN<sup>a</sup>, Osman GÜLER<sup>\*b</sup>

<sup>a</sup> Çankırı Karatekin Üniversitesi, Mühendislik Fakültesi, Bilgisayar Mühendisliği Bölümü, ÇANKIRI, 18100, TÜRKİYE

<sup>b\*</sup> Çankırı Karatekin Üniversitesi, Mühendislik Fakültesi, Bilgisayar Mühendisliği Bölümü, ÇANKIRI, 18100, TÜRKİYE

### MAKALE BİLGİSİ

Alınma: 22.04.2025  
Kabul: 22.06.2025

#### Anahtar Kelimeler:

Yapay zeka, Transfer öğrenme, Mobil uygulama, Domates, Domates yaprak hastalığı.

#### \*Sorumlu Yazar

osmanguler@karatekin.edu.tr

### ÖZET

Bu çalışmada, tarımsal bitki sağlığı yönetiminde yapay zeka uygulamalarından faydalanılmaktadır. Özel olarak, domates bitkilerinde görülen dokuz farklı hastalık türü ve sağlıklı yaprakları tespit edebilmek için derin öğrenme tabanlı bir çözüm geliştirilmiştir. Toplam 16.011 görüntü içeren bir veri seti kullanılarak, MobileNetV1, MobileNetV2, MobileNetV3 Small ve MobileNetV3 Large mimarileri ile 20 eğitim devresi boyunca eğitimler gerçekleştirilmiştir. Eğitimler sonrasında en yüksek doğruluk oranına ulaşan model, Android platformunda çalışabilen bir mobil uygulamaya entegre edilmiştir. Uygulama, kullanıcıların domates yapraklarındaki hastalıkları gerçek zamanlı olarak tespit etmelerine olanak tanımaktadır. İlk denemeler, geliştirilen modelin yüksek doğrulukla hastalıkları tespit edebildiğini göstermiştir. Bu çalışma, mobil teknolojilerin tarımsal bitki hastalıklarını erken teşhis etmedeki potansiyelini vurgulamakta ve tarımsal verimliliği artırmaya yönelik yeni yollar sunmaktadır.

DOI: 10.59940/jismar.1681569

## Using Transfer Learning for Disease Detection in Tomato Leaves and Its Integration into a Mobile Application

### ARTICLE INFO

Received: 22.04.2025  
Accepted: 22.06.2025

#### Keywords:

Artificial intelligence, Transfer learning, Mobile application, Tomato, Tomato leaf disease.

#### \*Corresponding Authors

e-mail:  
osmanguler@karatekin.edu.tr

### ABSTRACT

In this study, artificial intelligence applications are utilized for agricultural plant health management. Specifically, a deep learning-based solution has been developed to detect nine different types of diseases observed in tomato plants, as well as healthy leaves. Using a dataset containing a total of 16,011 images, training was conducted over 20 epochs with MobileNetV1, MobileNetV2, MobileNetV3 Small, and MobileNetV3 Large architectures. Following the training, the model that achieved the highest accuracy was integrated into a mobile application compatible with the Android platform. The application enables users to detect diseases in tomato leaves in real time. Initial experiments demonstrated that the model could identify diseases with high accuracy. This study highlights the potential of mobile technologies in the early diagnosis of plant diseases and presents new avenues for improving agricultural productivity.

DOI: 10.59940/jismar.1681569

## 1. GİRİŞ (INTRODUCTION)

Tarım, küresel nüfusun artmasıyla birlikte gıda güvenliği sağlamada kritik bir rol oynamaktadır. Özellikle domates, dünya çapında geniş bir tüketim ağına sahip olup, önemli bir besin kaynağıdır. Ancak, hastalıkların erken teşhisi, verimliliği artırma ve ürün kayıplarını minimuma indirme açısından büyük zorluklar sunmaktadır. Türkiye, yılda ortalama 12,8 milyon tonluk domates üretmektedir. Türkiye, dünya domates üretiminde üçüncü sırada yer alırken, ihracat sıralamasında ise beşinci konumdadır [1]. Domates, toplam sebze tüketiminin yaklaşık %20'sini oluşturmaktadır ve yıllık kişi başı ortalama tüketim 20 kg civarındadır [2]. Dünya genelinde domates üretiminde öne çıkan ülkeler arasında Türkiye, Mısır ve Çin gibi ülkeler bulunmaktadır [3]. Yapılan araştırmalarda, özellikle Punjab ve Sindh bölgelerinde domates yapraklarında görülen hastalıkların, ürün kaybını %30 ile %40 arasında artırdığı belirtilmiştir [2]. Geleneksel tarım yöntemleri, hem zaman açısından verimsiz hem de hata yapmaya elverişli olabilmektedir. Bu nedenle, tarım sektöründe teknolojinin kullanımı giderek daha önemli hale gelmektedir.

Domates bitkileri, verimi ve kaliteyi önemli ölçüde etkileyebilecek çeşitli hastalıklara karşı hassastır. Bu hastalıkların erken ve doğru bir şekilde tanımlanması, etkili ürün yönetimi için çok önemlidir. Geleneksel hastalık tanımlama yöntemleri, zaman alıcı ve insan hatasına açık olabilen uzmanlar tarafından yapılan görsel incelemeye dayanır. Derin öğrenme alanındaki son gelişmeler, hastalık tanımlama sürecinin otomatikleştirilmesine olanak sağlayarak daha hızlı ve daha doğru bir alternatif sunmaktadır.

Yapay zekâ teknolojileri, tarımda verimliliği artırmanın yanı sıra hastalık teşhis ve kontrolünde önemli bir dönüşüm potansiyeline sahiptir. Özellikle derin öğrenme, görüntü işleme ve analizinde büyük başarılar göstererek, bitkisel hastalıkların erken teşhisinde önemli bir araç haline gelmiştir. Mobilite ve erişilebilirlik, bu teknolojilerin tarım sektörüne adaptasyonunu daha da değerli kılmaktadır.

Domates yaprak hastalıklarına mantarlar, bakteriler, virüsler ve çevresel faktörler neden olabilir. Bu çalışmada, domates bitkileri üzerinde yaygın olarak görülen dokuz farklı hastalık türü ve sağlıklı bitkileri ayırt edebilmek amacıyla MobileNet mimarileri kullanılarak eğitimler gerçekleştirilmiştir. Toplam 16,011 adet görüntü üzerinden eğitilen modeller, Android platformunda çalışan bir mobil uygulama aracılığıyla saha koşullarında test edilmiştir. Uygulamanın, kullanıcıların domates

yapraklarını tarayarak hastalıkları gerçek zamanlı olarak tespit etme yeteneği, tarımsal faaliyetlerde verimliliği artırma ve erken müdahale olanağı sağlama potansiyeli taşımaktadır. Günümüzdeki tarım uygulamaları, bitki hastalıklarının teşhisinde genellikle mikroskop gibi ek ekipmanlar ve uzmanların yardımıyla çiftlik çalışanlarının görsel incelemelerine dayanmaktadır [2]. Ancak, tarım uzmanları sahada sürekli olarak bulunmadığından kapsamlı izleme yapmak mümkün olmayabilir. Ayrıca, çiftçilerin çoğu bu tür teşhis süreçlerini yürütebilecek uzmanlığa sahip değildir [2]. Bu nedenle, domates yapraklarında görülen hastalıkların tespiti için derin öğrenme yöntemleri literatürde sıklıkla kullanılmaktadır [4], [5]. Derin öğrenmenin en önemli avantajlarından biri, özellik çıkarımını otomatik olarak gerçekleştirerek karmaşık ve üst düzey özelliklerin keşfedilmesine olanak tanımasıdır [6].

Bu çalışma, bu teknolojilerin pratik uygulamalarını gözler önüne sermekte ve yapay zekânın tarım sektöründe nasıl bir dönüşüm yaratabileceğine dair önemli bilgiler sunmaktadır. Ayrıca, bu girişim, yapay zekâ tabanlı çözümlerin tarım pratiğine entegrasyonunun önündeki teknik ve pratik engelleri aşmada önemli bir örnek olarak değerlendirilmektedir.

## 2. İLGİLİ ÇALIŞMALAR (RELATED WORK)

Makine öğreniminin bir alt kümesi olan derin öğrenme, verilerden özellikleri otomatik olarak öğrenme yeteneği nedeniyle görüntü tanıma görevlerinde yaygın olarak benimsenmiştir. Evrişimli Sinir Ağları (CNN'ler) görüntü sınıflandırma görevlerinde yüksek performans göstermeleri nedeniyle yaprak görüntülerinden bitki hastalıklarını tespit etmek için uygun hale getirir. Tarım alanında hastalık tespiti, verimlilik tespiti vb. farklı birçok uygulamada yapay zekâ yöntemleri tarım uzmanlarına destek sağlamaktadır.

Kılıçarslan ve Paçal, domates yapraklarında görülen hastalıkların tespiti için DenseNet, ResNet50 ve MobileNet mimarilerini kullanarak çeşitli deneyler gerçekleştirmişlerdir. Yapılan çalışmalar sonucunda, DenseNet modelinin en yüksek doğruluk oranına ulaştığı belirtilmiştir. Bu model ile 0.0269 hata oranı, 0.9900 doğruluk, 0.9880 kesinlik, 0.9892 F1-skor ve 0.9906 duyarlılık değerleri elde edilmiştir [2].

Literatürdeki benzer bir çalışmada, PlantVillage domates veri kümesi kullanılarak domates hastalıklarının sınıflandırılması için en uygun makine öğrenmesi (ML) ve derin öğrenme (DL) modelleri araştırılmıştır [7]. Bu çalışmada, yerel ikili desen (LBP) ve gri seviye eş oluşturma matrisi (GLCM)

yöntemleriyle 105 renk özelliği çıkarılmıştır. Test edilen modeller arasında ResNet34 ağı, %99,7 doğruluk, %99,6 kesinlik, %99,7 duyarlılık ve %99,7 F1-skoru ile en iyi sonuçları vermiştir.

Cengil ve Çıkar tarafından yapılan bir başka çalışmada, Taiwan veri kümesi kullanılarak domates yaprağı hastalıklarının tespiti için AlexNet, ResNet50 ve VGG16 gibi transfer öğrenme modelleri değerlendirilmiştir. Deneyler sonucunda, ResNet50 modelinin %96,10 doğruluk oranı ile en yüksek başarıyı sağladığı raporlanmıştır [8].

Bir diğer araştırmada, bitki hastalıklarının belirtilerini sağlıklı dokulardan ayırt etmek amacıyla bir algoritma önerilmiştir [9]. Bu algoritma, HSV renk uzayındaki H kanalı ve CIELAB renk uzayındaki a kanalının histogramlarını kullanmaktadır. Algoritma, 19 bitki türü ve 82 hastalık görüntüsü içeren geniş bir veri seti üzerinde test edilmiştir. Sonuçlara göre, a kanalı geçiş bölgesini hastalıklı doku olarak sınıflandırma eğilimindeyken, H kanalı aynı bölgeyi sağlıklı doku olarak değerlendirmiştir. Ayrıca, a kanalının geçiş bölgesinin %50'sinden fazlasını hastalıklı olarak tanımladığı rapor edilmiştir.

Derin öğrenme tekniklerinin kullanıldığı bir diğer çalışmada, Zaki ve arkadaşları [10] MobileNetV2 modelini kullanarak yaprak hastalıklarını tespit etmeyi amaçlamışlardır. Model, üç farklı domates hastalığını tanımak üzere ince ayarlanmış ve PlantVillage veri kümesinden alınan 4.671 görüntü üzerinde test edilmiştir. Sonuçlar, MobileNetV2'nin hastalıkları %90'ın üzerinde bir doğruluk oranıyla tespit edebildiğini ortaya koymuştur.

Diğer bir çalışmada ise, domatesler de dahil olmak üzere 14 mahsul türünde 26 hastalığı sınıflandırmak için derin öğrenme modeli kullanılmıştır [11]. Bu model, bitki hastalıklarının tanımlanmasında derin öğrenmenin etkili bir yöntem olduğunu ortaya koymuş ve yüksek doğruluk oranları elde etmiştir.

Sladojevic ve arkadaşları, domates yapraklarını etkileyen etmenleri de içeren bitki hastalıklarını tespit etmek ve sınıflandırmak amacıyla CNN tabanlı bir yaklaşım önermişlerdir [12]. Önerilen yaklaşım, 3.000 görüntüden oluşan bir veri setinde %96'nın üzerinde doğruluk oranı elde etmiştir.

Ferentinos, domatesler de dahil olmak üzere çeşitli bitkilerdeki hastalıkları tanımlamak için derin öğrenme modellerini kullanmıştır [13]. Bu çalışmada, farklı CNN mimarilerinin sınıflandırma başarıları karşılaştırılmış ve daha derin yapıya sahip ağların genellikle daha yüksek doğruluk oranlarına ulaştığı ifade edilmiştir.

Zhang ve arkadaşları, özellikle domates yaprağı hastalığının tanımlanması için tasarlanmış yeni bir CNN mimarisi önermişlerdir [14]. Önerilen model,

açık erişimli bir veri seti üzerinde yüksek doğruluk oranı ile başarılı sonuçlar elde etmiştir.

Karthik ve arkadaşları, domates bitkilerinde gerçek zamanlı hastalık tespiti için hafif bir CNN modeli önermişlerdir [15]. Model, mobil cihazlarda kullanım için optimize edilmiştir. Tarlada çalışan çiftçiler tarafından başarıyla kullanılmıştır.

Abbas ve arkadaşları, Koşullu Üretken Çelişkili Ağ (Conditional GAN – C-GAN) ile domates yapraklarının sentetik görüntülerini oluşturmuşlar ve derin öğrenmeye dayalı bir yaklaşım ile hastalık tahmini yapmışlardır [16]. Önerilen yaklaşım, açık erişimli PlantVillage veri kümesi üzerinde yapılan deneylerde sırasıyla 5 sınıfa, 7 sınıfa ve 10 sınıfa %99,51, %98,65 ve %97,11 doğruluk elde etmiştir.

Sanida ve arkadaşları, VGGNet ile çok nitelikli tanımlama görevi için iyileştirilmiş kategorik çapraz entropi kaybı fonksiyonu kullanan bir yaklaşım önermişlerdir [17]. Önerilen yaklaşım domates hastalık tespiti için %99,23 doğruluk değeri elde etmiştir.

Nag ve arkadaşları, AlexNet, ResNet-50, SqueezeNet-1.1, VGG19 ve DenseNet-121 gibi CNN mimarileri ile PlantVillage veri setinden domates yapraklarında akıllı hastalık tespiti için mobil uygulama tabanlı bir sistem önermişlerdir [18]. DenseNet-121 modeli %99,85 doğruluk oranı ile en iyi performansı göstermiştir.

### 3. MATERYAL VE METOD (MATERIALS AND METHODS)

#### 3.1. Veri Seti ve Görüntü Toplama (Dataset and Image Collection)

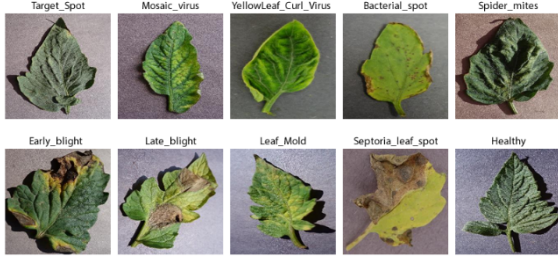
Bu çalışmada kullanılan veri seti, domates yapraklarına ait 16.011 adet görüntü içermektedir. Veri seti, sağlıklı yapraklar ve dokuz farklı hastalık türü olmak üzere toplam on sınıfa (Bacterial\_spot, Early\_blight, Late\_blight, Leaf\_Mold, Septoria\_leaf\_spot, Spider\_mites, Two-spotted\_spider\_mite, Target\_Spot, Yellow\_Leaf\_Curl\_Virus, Mosaic\_virus ve Healthy) ayrılmıştır. Görüntüler, tarla koşullarında farklı zamanlarda ve farklı açılardan çekilmiştir, böylece modelin gerçek dünya koşullarına adaptasyonu artırılmıştır.

Veri setindeki sınıf dağılımı homojen değildir. Sınıf dağılımı şu şekildedir: Tomato\_Yellow\_Leaf\_Curl\_Virus sınıfı 3.209, Tomato\_Bacterial\_spot 2.127, Tomato\_Late\_blight 1.909, Tomato\_Septoria\_leaf\_spot 1.771,

Tomato\_Spider\_mites 1.676, Tomato\_healthy 1.591, Tomato\_Target\_Spot 1.404, Tomato\_Early\_blight 1.000, Tomato\_Leaf\_Mold 952 ve Tomato\_Tomato\_mosaic\_virus 373 görüntü içermektedir.

Ayrıca, eğitim sırasında veri artırma (data augmentation) teknikleri uygulanmıştır. Görüntüler döndürme, yatay çevirme ve parlaklık ayarı gibi işlemlerle çeşitlendirilmiş; bu sayede modelin genelleme yeteneği artırılmış ve aşırı öğrenme (overfitting) riski azaltılmıştır.

Şekil 1'de veriseti sınıflarına ait örnek görseller verilmiştir.



Şekil 1. Sağlıklı ve hastalıklı domates yaprağı görüntüleri (Images of healthy and diseased tomato leaves.).

### 3.2. MobileNet Mimarileri (MobileNet Architectures)

MobileNet modelleri, özellikle mobil ve gömülü sistemlerde yüksek performanslı derin öğrenme modellerinin kullanılmasını sağlamak amacıyla geliştirilmiş hafif evrişimli sinir ağı (CNN) mimarileridir. Sınırlı hesaplama kaynaklarına sahip cihazlarda gerçek zamanlı uygulamalar için optimize edilmişlerdir.

Bu çalışmada kullanılan MobileNet mimarisinin üç farklı versiyonu bulunmaktadır: MobileNetV1, MobileNetV2 ve MobileNetV3Small. Her bir versiyon, daha az parametre kullanarak daha hızlı çıkarım yapabilme amacıyla farklı derinlik düzeyleri ve blok yapıları sunar. MobilNet mimarileri, mobil cihazlarda verimli çalışma açısından yaygın olarak tercih edilmektedir.

#### 3.2.1. MobileNetV1

MobileNetV1, Howard ve arkadaşları tarafından 2017 yılında tanıtılmıştır [19]. Bu modelde, standart evrişimli katmanlar yerine derinlik ayrılabilir evrişimler (depthwise separable convolutions) kullanılarak parametre sayısı ve hesaplama maliyeti önemli ölçüde azaltılmıştır. Derinlik ayrılabilir evrişimler, uzamsal filtreleme ve kanal birleştirme

işlemlerini ayrı ayrı gerçekleştirerek, standart evrişimlere kıyasla daha düşük hesaplama gereksinimi sağlar. Bu yapı sayesinde, MobileNetV1 düşük güç tüketimi ve yüksek işlem hızı sunarken, makul düzeyde doğruluk oranı da koruyabilmektedir. Ancak, derinlik ayrılabilir evrişimlerin doğruluk açısından bazı sınırlamalara sahip olduğu da literatürde rapor edilmiştir.

#### 3.2.2. MobileNetV2

MobileNetV2, Sandler ve arkadaşları tarafından 2018 yılında geliştirilmiştir [20]. Bu modelde, MobileNetV1'in eksikliklerini gidermek amacıyla doğrusal darboğaz (linear bottleneck) ve ters çevrilmiş artık bağlantılar (inverted residual connections) gibi yeni mimari öğeler eklenmiştir. Doğrusal darboğaz, düşük boyutlu uzayda gerçekleştirilen doğrusal dönüşümler aracılığıyla bilgi kaybını azaltmayı hedeflerken, ters çevrilmiş artık yapılar daha geniş ara katmanlar sayesinde modelin temsil kapasitesini artırmaktadır. Bu yenilikler, MobileNetV2'nin hem hesaplama verimliliği hem de doğruluk performansı açısından MobileNetV1'den daha iyi performans göstermesini sağlamıştır. Özellikle, düşük çözünürlüklü girdilerde daha iyi sonuçlar elde edilmiştir.

#### 3.2.3. MobileNetV3

MobileNetV3, Howard ve arkadaşları tarafından 2019 yılında tanıtılmıştır [21]. Bu model, MobileNetV2'nin temelini alarak, otomatik sinir ağı mimarisi arama ve donanım bilgili optimizasyonlar kullanılarak geliştirilmiştir. MobileNetV3, iki farklı versiyonla sunulmuştur: MobileNetV3-Small (MobileNetV3s) ve MobileNetV3-Large (MobileNetV3l). MobileNetV3s, daha küçük ve daha hızlı bir model olarak tasarlanırken, MobileNetV3l daha yüksek doğruluk oranları hedeflemiştir.

MobileNet serisi, mobil ve gömülü sistemler için hafif ve verimli derin öğrenme modelleri sunmaktadır. Her bir versiyon, önceki modelin eksikliklerini gidererek hem doğruluk hem de hesaplama verimliliği açısından önemli adımlar atmıştır. MobileNetV3, özellikle otomatik mimari arama ve donanım bilgili optimizasyonlar sayesinde, bu serinin en gelişmiş modeli olarak öne çıkmaktadır.

#### 3.2.4. Model Eğitimi Süreci (Training Process)

Model eğitimi sürecinde, sınıflandırma başarımını artırmak ve aşırı öğrenmeyi (overfitting) önlemek amacıyla çeşitli stratejiler uygulanmıştır.

Optimizasyon algoritması olarak Adam tercih edilmiş; öğrenme oranı (learning rate) 0.0001 olarak belirlenmiştir. Eğitim işlemi, 16 görüntüden oluşan mini-batch'lerle (batch size = 16) toplam 19 epoch boyunca gerçekleştirilmiş, ardından modelin son katmanları serbest bırakılarak 1 epoch süreyle ince ayar (fine-tuning) yapılmıştır.

Aşırı öğrenmenin önlenmesi için, eğitim veri kümesine yatay çevirme (horizontal flip), rastgele döndürme (random rotation), yakınlaştırma/uzaklaştırma (random zoom) ve kontrast değişikliği (random contrast) gibi veri artırma (data augmentation) teknikleri uygulanmıştır. Bu işlemler, modelin veri çeşitliliğine karşı daha dayanıklı hale gelmesini sağlamıştır. Ayrıca doğrulama kaybında belirli bir süre boyunca iyileşme gözlenmediğinde eğitimi sonlandıran erken durdurma (early stopping) yöntemi devreye alınarak aşırı öğrenme riski azaltılmıştır.

#### 4. DENEYSEL SONUÇLAR (EXPERIMENTAL RESULTS)

Bu çalışma kapsamında domates yaprak görüntülerini sınıflandırmak için keras kütüphanesinde yer alan MobileNetV1, MobileNetV2, MobileNetV3s ve MobileNetV3l modelleri kullanılmıştır. MobileNet mimarileri, hafifliği ve yüksek performansı nedeniyle mobil cihazlarda kullanım açısından uygundur. Modeller, Python programlama dili ile TensorFlow ve Keras kütüphaneleri kullanılarak geliştirilmiştir. Eğitim süreci 20 epoch boyunca gerçekleştirilmiştir. Veri seti %80 eğitim ve %20 doğrulama olacak şekilde ikiye ayrılmıştır. Ayrıca, modelin genelleme kabiliyetini artırmak amacıyla veri artırma teknikleri uygulanmıştır. Veri artırma sürecinde, eğitim veri kümelerinde yer alan görüntüler üzerinde çeşitli uzamsal dönüşümler uygulanarak orijinal görüntüye benzer ancak küçük farklılıklar içeren yeni görüntüler üretilir. Bu işlem, eğitim verilerinin çeşitliliğini ve miktarını artırarak derin öğrenme modellerinin daha geniş bir veri yelpazesinde eğitilmesini sağlar. Böylece, veri çeşitliliği artırılarak modelin daha geniş bir örneklem üzerinden eğitilmesi sağlanmıştır. [22].

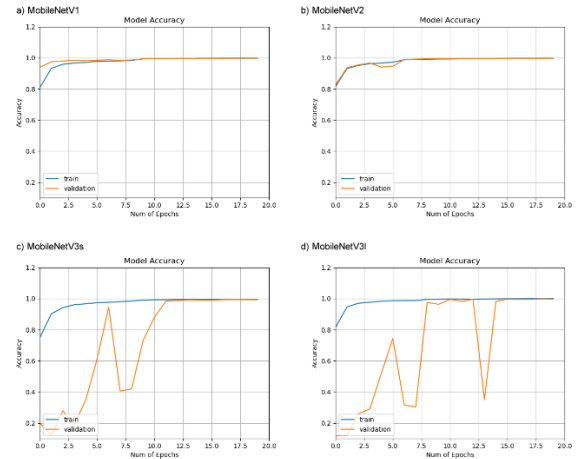
MobileNet modelleriyle 20 epoch süresince gerçekleştirilen eğitimlerin sonuçları Tablo 1'de, eğitim doğruluklarına ait grafikler ise Şekil 2'de sunulmaktadır.

**Tablo 1.** Deneysel sonuçlar (Experimental Results)

| Model       | Train Acc. | Val. Acc. | Train Loss | Val. Loss | Epoch Time    |
|-------------|------------|-----------|------------|-----------|---------------|
| MobileNetV1 | 0.9973     | 0.9978    | 0.8741     | 0.8732    | 144s<br>112ms |

|              |        |        |        |        |               |
|--------------|--------|--------|--------|--------|---------------|
| MobileNetV2  | 0.9973 | 0.9975 | 0.8751 | 0.8756 | 382s<br>298ms |
| MobileNetV3s | 0.9942 | 0.9928 | 0.8841 | 0.8862 | 141s<br>110ms |
| MobileNetV3l | 0.9991 | 0.9972 | 0.8711 | 0.8742 | 146s<br>114ms |

Tablo 1, incelendiği zaman, MobileNetV3l modeli en yüksek eğitim doğruluğu (%99.91) ve en düşük kayıp değerleriyle (0.8711) öne çıkarak en iyi performansı sergilerken, MobileNetV3s modeli daha düşük doğruluk (%99.28) ve daha yüksek kayıp değerleri (0.8862) göstererek daha sınırlı bir kapasiteye sahip olduğunu ortaya koymaktadır. MobileNetV1 ve MobileNetV2 ise benzer doğruluk ve kayıp değerleriyle dengeli bir performans sergilemektedir. Epoch süreleri açısından MobileNetV3s ve MobileNetV1 en hızlı modeller olurken, MobileNetV2 daha uzun eğitim süreleri gerektirmektedir.



Şekil 2. Doğruluk grafikleri (a) MobileNetV1, (b) MobileNetV2, (c) MobileNetV3s, (d) MobileNetV3l (Accuracy graphs: (a) MobileNetV1, (b) MobileNetV2, (c) MobileNetV3 Small, (d) MobileNetV3 Large).

Eğitimler sonucunda modellerin başarısını ölçmek için yapılan test sonuçları ve ölçülen performans metriklerine ait sonuçlar Tablo 2'de gösterilmiştir. Modellerin tahminlerini gösteren karmaşıklık matrisleri Şekil 3'de gösterilmiştir. Şekil 4'de ise eğitimden sonra yapılan test sonuçlarına ait görseller verilmiştir.

**Tablo 2.** Performans metrikleri sonuçları (Performance Metric Results)

| Model        | Test Accuracy | Classes        | Precision | Recall | F1 skor |
|--------------|---------------|----------------|-----------|--------|---------|
| MobileNet V1 | 0.91          | Bacterial Spot | 0.93      | 0.93   | 0.93    |
|              |               | Early Blight   | 0.86      | 0.76   | 0.80    |
|              |               | Late Blight    | 0.86      | 0.91   | 0.88    |
|              |               | Leaf Mold      | 0.89      | 0.87   | 0.88    |

| Model         | Test Accuracy | Classes            | Precision | Recall | F1 skor |
|---------------|---------------|--------------------|-----------|--------|---------|
|               |               | Septoria Leaf Spot | 0.85      | 0.89   | 0.87    |
|               |               | Spider Mite        | 0.84      | 0.91   | 0.87    |
|               |               | Target Spot        | 0.86      | 0.87   | 0.87    |
|               |               | Curl Virus         | 0.99      | 0.96   | 0.97    |
|               |               | Mosaic Virus       | 0.97      | 0.87   | 0.92    |
|               |               | Healthy            | 0.99      | 0.94   | 0.96    |
| MobileNet V2  | 0.91          | Bacterial Spot     | 1.0       | 0.79   | 0.88    |
|               |               | Early Blight       | 0.92      | 0.64   | 0.76    |
|               |               | Late Blight        | 0.99      | 0.91   | 0.95    |
|               |               | Leaf Mold          | 0.85      | 1.0    | 0.92    |
|               |               | Septoria Leaf Spot | 0.63      | 1.0    | 0.77    |
|               |               | Spider Mite        | 1.0       | 0.59   | 0.74    |
|               |               | Target Spot        | 0.75      | 0.93   | 0.83    |
|               |               | Curl Virus         | 1.0       | 0.96   | 0.98    |
|               |               | Mosaic Virus       | 0.73      | 1.0    | 0.84    |
|               |               | Healthy            | 0.96      | 0.99   | 0.98    |
| MobileNet V3s | 0.93          | Bacterial Spot     | 0.97      | 0.88   | 0.92    |
|               |               | Early Blight       | 0.98      | 0.70   | 0.82    |
|               |               | Late Blight        | 0.96      | 0.94   | 0.95    |
|               |               | Leaf Mold          | 0.89      | 0.93   | 0.91    |
|               |               | Septoria Leaf Spot | 0.85      | 0.94   | 0.89    |
|               |               | Spider Mite        | 0.90      | 0.92   | 0.91    |
|               |               | Target Spot        | 0.85      | 0.90   | 0.87    |
|               |               | Curl Virus         | 1.0       | 0.94   | 0.97    |
|               |               | Mosaic Virus       | 0.55      | 1.0    | 0.71    |
|               |               | Healthy            | 0.94      | 1.0    | 0.97    |
| MobileNet V3l | 0.99          | Bacterial Spot     | 0.99      | 0.99   | 0.99    |
|               |               | Early Blight       | 0.96      | 0.93   | 0.94    |
|               |               | Late Blight        | 0.97      | 0.99   | 0.98    |
|               |               | Leaf Mold          | 1.0       | 0.99   | 1.0     |
|               |               | Septoria Leaf Spot | 0.99      | 0.99   | 0.99    |
|               |               | Spider Mite        | 0.99      | 0.99   | 0.99    |
|               |               | Target Spot        | 0.97      | 1.0    | 0.99    |
|               |               | Curl Virus         | 1.0       | 0.99   | 1.0     |
|               |               | Mosaic Virus       | 1.0       | 0.98   | 0.99    |
|               |               | Healthy            | 1.0       | 0.99   | 1.0     |

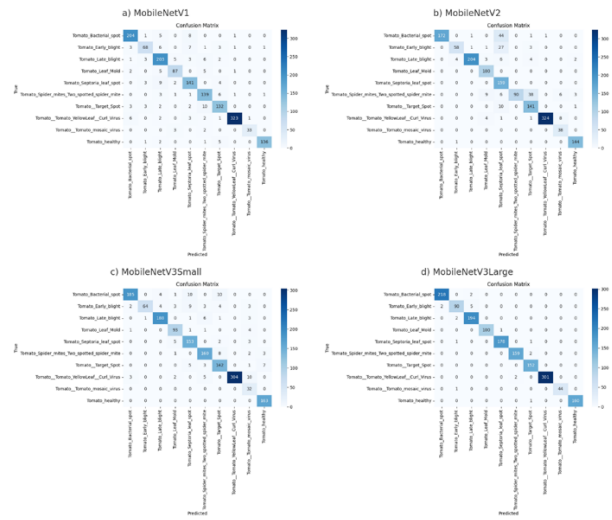
Tablo 2’de sunulan performans metrikleri incelendiğinde, dört farklı MobileNet mimarisinin test verisi üzerindeki sınıflandırma başarısı detaylı olarak değerlendirilmektedir. MobileNetV1 modeli, genel test doğruluğu %91 olup özellikle *Curl Virus* ve *Healthy* sınıflarında yüksek precision (%99) ve recall (%94-96) değerleriyle başarılı sonuçlar vermiştir. Ancak *Early Blight* sınıfında precision (%86) ve recall (%76) değerlerinin diğer sınıflara

göre daha düşük olması, bu sınıfın ayırt edilmesinde güçlük yaşandığını göstermektedir.

MobileNetV2 modeli de %91 doğruluk değeri ile benzer bir genel performans sergilemiştir. Ancak bazı sınıflarda precision ve recall değerleri arasında belirgin farklılıklar gözlenmiştir. Örneğin *Spider Mite* sınıfında %100 precision ve %59 recall değeri, modelin bu sınıfı doğru tanıma konusunda düşük duyarlılığa sahip olduğunu göstermektedir. Benzer şekilde *Septoria Leaf Spot* ve *Mosaic Virus* sınıflarında da düşük precision veya recall değerleri dikkat çekmektedir.

MobileNetV3Small modeli %93 test doğruluğu ile daha iyi bir genel başarı sağlamıştır. Bu model, *Late Blight*, *Curl Virus* ve *Healthy* sınıflarında yüksek precision ve recall değerleri ile öne çıkmaktadır. Ancak *Mosaic Virus* sınıfında precision (%55) düşüktür; bu durum modelin diğer sınıflardan bazı örnekleri bu sınıfa atadığını göstermektedir.

MobileNetV3Large modeli %99 test doğruluğu ile en yüksek genel başarıyı göstermiştir. Tüm sınıflarda precision, recall ve F1 skor değerleri %97’nin üzerindedir. Özellikle *Leaf Mold*, *Curl Virus* ve *Healthy* sınıflarında ölçülen metrikler %99 ve üzerindedir. Bu sonuçlar, modelin sınıflar arasında yüksek ayırım gücüne sahip olduğunu göstermektedir. Sonuç olarak, MobileNetV3Large modeli, sınıf bazlı metriklerde gösterdiği yüksek ve dengeli performans ile en başarılı mimari olmuştur.



Şekil 3. Karmaşıklık matrisleri (a) MobileNetV1, (b) MobileNetV2, (c) MobileNetV3s, (d) MobileNetV3l (Confusion matrices: (a) MobileNetV1, (b) MobileNetV2, (c) MobileNetV3 Small, (d) MobileNetV3 Large).

Şekil 3’te sunulan karmaşıklık matrisleri, her bir MobileNet mimarisinin test verisindeki sınıflandırma performansını görsel olarak ortaya koymaktadır. MobileNetV1 matrisinde, özellikle *Healthy*, *Curl*

Virus ve Mosaic Virus sınıflarında doğru sınıflandırma oranlarının yüksek olduğu görülmektedir. Bununla birlikte Early Blight ve Spider Mite sınıflarında yanlış sınıflandırmaların diğer modellere kıyasla daha fazla olduğu gözlemlenmiştir.

MobileNetV2 matrisinde sınıflar arası karışıklık daha belirgindir. Özellikle Spider Mite, Early Blight ve Septoria Leaf Spot sınıflarına ait örneklerin, bazı durumlarda benzer semptomlara sahip diğer sınıflarla karıştırıldığı görülmektedir. Bu da modelin bazı sınıflar arasında ayırım yapmakta zorlandığını göstermektedir.

MobileNetV3Small modeline ait matris, genel olarak daha dengeli bir dağılım sergilemektedir. Mosaic Virus sınıfında bazı hatalı sınıflandırmalar bulunsa da diğer sınıflarda başarı oranları daha yüksektir. Healthy, Late Blight ve Curl Virus sınıflarının yüksek doğrulukla sınıflandırıldığı gözlemlenmektedir.

MobileNetV3Large matrisinde ise sınıflar arası karışıklığın en az olduğu görülmektedir. Tüm sınıflar yüksek doğrulukla sınıflandırılmış, yanlış sınıflandırma oranı oldukça düşmüştür. Bu durum, modelin yüksek duyarlılık ve seçicilikle çalıştığını ve sınıflar arasındaki sınırları net olarak ayırt edebildiğini göstermektedir.

Genel olarak, karmaşıklık matrisleri, metrik tablosu ile paralel olarak MobileNetV3Large modelinin en yüksek sınıflandırma doğruluğunu sağladığını ve sınıflar arası ayırımı daha başarılı gerçekleştirdiğini göstermektedir.

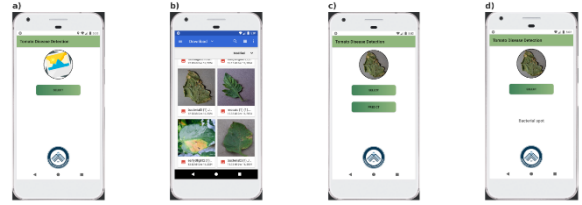


Şekil 4. Eğitimden sonra yapılan test sonuçları (Test results after training)

Sonuç olarak, eğitim süresi, doğruluk değerleri, test sonuçları, performans metrikleri, eğitim ve karmaşıklık grafikleri incelendiği zaman MobileNetV3l modeli en etkili model olarak öne çıkmaktadır. Mobil uygulamada MobileNetV3l mimarisi kullanılmıştır.

#### 4.1. Mobil Uygulama Entegrasyonu (Mobile Application Integration)

Eğitimden sonra, elde edilen model ağırlıkları, Android Studio ortamında geliştirilen bir mobil uygulamaya entegre edilmiştir. Uygulama, kullanıcıların akıllı telefonları aracılığıyla domates yapraklarını fotoğraflayıp, anında hastalık tespiti yapabilmelerini sağlar. Uygulama arayüzü, kullanıcı dostu ve etkileşimli özelliklerle donatılmıştır. Uygulama, çevrimdışı çalışabilme yeteneği sayesinde tarla gibi internet erişiminin sınırlı olduğu yerlerde de kullanılabilir.



Şekil 5. Uygulama tasarımı (a) Ana menü ekranı, (b) Fotoğraf seçim ekranı (c) Fotoğraf seçme-tahmin ekranı, (d) Tahmin sonuç ekranı (Application design: (a) Main menu screen, (b) Image selection screen, (c) Image selection and prediction screen, (d) Prediction result screen)

Şekil5’de gösterilen uygulama görüntülerinden uygulama ana menü ekranından select tuşu ile fotoğraf galerisine girilir. Galeri içerisinde daha önce fotoğrafını çekmiş olduğumuz domates yaprağını görüntüsünü seçeriz. Daha sonra seçilen görüntüden hastalık tahmini için predict butonuna basılarak tahmin sonucu görüntülenir, yeniden görüntü seçmek için select butonuna basılarak galeriye geri dönlür. Galeriden seçilen fotoğrafın hastalık tahmini yapılarak, sonuç ekranda gösterilir. Uygulama kullanıcı dostu olup işlemler çok basit şekilde yapılması için tasarlanmıştır.

#### 5. SONUÇ (CONCLUSION)

Derin öğrenme, domates yaprak hastalıklarının tanımlanmasını otomatikleştirmede büyük bir potansiyel göstermiştir ve geleneksel yöntemlere göre daha hızlı ve daha doğru bir alternatif sunmaktadır. Zorluklar devam ederken, derin öğrenme tekniklerindeki devam eden araştırmalar ve gelişmelerin bu sistemlerin etkinliğini ve erişilebilirliğini daha da iyileştirmesi ve nihayetinde

çiftçilere ve tarım endüstrisine bir bütün olarak fayda sağlaması muhtemeldir.

Bu çalışmada, MobileNet mimarileri kullanılarak domates yapraklarındaki hastalıkları tespit etmek için bir mobil uygulama geliştirilmiştir. Elde edilen sonuçlar, modelin yüksek doğruluk oranları ile hastalıkları etkili bir şekilde tespit edebildiğini göstermektedir. Mobil uygulama, tarım profesyonelleri ve çiftçiler için değerli bir araç olarak öne çıkmakta, erken teşhis ve müdahalede büyük bir potansiyel sunmaktadır. Bu çalışma, domates yapraklarındaki hastalıkları tespit etmek için geliştirilen bir yapay zekâ tabanlı mobil uygulamanın potansiyelini ortaya koymaktadır. Model, yüksek doğruluk, duyarlılık ve özgüllük oranları ile dikkat çekici sonuçlar elde etmiştir. Bununla birlikte, modelin performansı çeşitli faktörlere bağlı olarak değişkenlik gösterebilir. Örneğin, görüntü kalitesi, çekim koşulları ve yaprakların hasar durumu gibi faktörler, modelin doğruluğunu etkileyebilir. Bu faktörlerin, modelin genel performansına olan etkilerinin daha detaylı incelenmesi gerekmektedir.

Model, özellikle Leaf Mold, Curl Virus ve Healthy sınıflarında %99'un üzerinde başarı sergileyerek güçlü bir ayırt edicilik yeteneği ortaya koymuştur. Elde edilen bu sonuçlar, literatürdeki benzer çalışmalardaki performanslarla karşılaştırıldığında oldukça rekabetçidir. Örneğin, Zhang et al. (2020) tarafından kullanılan InceptionV3 modeli %96 doğruluk bildirirken; Fuentes et al. (2018) çalışmasında Faster R-CNN tabanlı sistemin doğruluk oranı %93 olarak rapor edilmiştir. MobileNet mimarilerinin daha düşük hesaplama gereksinimiyle bu denli yüksek başarı elde etmesi, bu çalışmayı önemli kılmaktadır.

Geliştirilen mobil uygulama, çevrimdışı çalışabilme özelliği sayesinde internet erişiminin kısıtlı olduğu kırsal bölgelerde de etkin bir şekilde kullanılabilir. Bu yönüyle, küçük ölçekli çiftçiler ve tarım danışmanları için pratik ve ulaşılabilir bir araç niteliği taşımaktadır. Bu teknolojinin ilerletilmesi, tarımsal üretkenliği artırmak ve gıda güvenliğini desteklemek için kritik öneme sahiptir.

#### 6. Gelecek Çalışmalar (Future Work)

Gelecek çalışmalarda, modelin farklı domates çeşitleri üzerindeki başarısı test edilerek genelleştirilebilirliği artırılmalıdır. Bu, modelin farklı ekolojik ve genetik çeşitlilik gösteren bitkilerde de doğru sonuçlar verip veremeyeceğini değerlendirmek açısından önemlidir. Ayrıca, modelin daha fazla hastalık türü üzerinde eğitilmesi ve farklı iklim koşullarında uygulanabilirliğinin test edilmesi de önerilmektedir. Bu tür gelişmeler, yapay zekâ tabanlı

tarım uygulamalarının daha geniş ölçekli entegrasyonuna katkı sağlayacaktır.

#### REFERANSLAR (REFERENCES)

- [1] TEPGE, "Tarım Ürünleri Piyasaları," Ankara, Haziran 2021.
- [2] Kılıçarslan, S., & Pacal, I. (2023). Domates Yapraklarında Hastalık Tespiti İçin Transfer Öğrenme Metotlarının Kullanılması. *Mühendislik Bilimleri ve Araştırmaları Dergisi*, 5(2), 215-222.
- [3] Zhao, S., Peng, Y., Liu, J., & Wu, S. (2021). Tomato leaf disease diagnosis based on improved convolution neural network by attention module. *Agriculture*, 11(7), 651.
- [4] Mansoor, S., Khan, S. H., Saeed, M., Bashir, A., Zafar, Y., Malik, K. A., & Markham, P. G. (1997). Evidence for the association of a bipartite geminivirus with tomato leaf curl disease in Pakistan. *Plant Disease*, 81(8), 958-958.
- [5] Raza, A., Shakeel, M. T., Umar, U. U. D., Tahir, M. N., Hassan, A. A., Katis, N. I., & Wang, X. (2020). First report of tomato chlorosis virus infecting tomato in Pakistan. *Plant Dis*, 104(2036), 10-1094.
- [6] Karaman, A., Karaboga, D., Pacal, I., Akay, B., Basturk, A., Nalbantoglu, U., ... & Sahin, O. (2023). Hyper-parameter optimization of deep learning architectures using artificial bee colony (ABC) algorithm for high performance real-time automatic colorectal cancer (CRC) polyp detection. *Applied Intelligence*, 53(12), 15603-15620.
- [7] Tan, L., Lu, J., & Jiang, H. (2021). Tomato leaf diseases classification based on leaf images: a comparison between classical machine learning and deep learning methods. *AgriEngineering*, 3(3), 542-558.
- [8] Cengil, E., & Çınar, A. (2022). Hybrid convolutional neural network based classification of bacterial, viral, and fungal diseases on tomato leaf images. *Concurrency and Computation: Practice and Experience*, 34(4), e6617.
- [9] Barbedo, J. G. A. (2016). A novel algorithm for semi-automatic segmentation of plant leaf disease symptoms using digital image processing. *Tropical Plant Pathology*, 41(4), 210-224.
- [10] Zaki, S. Z. M., Zulkifley, M. A., Stofa, M. M., Kamari, N. A. M., & Mohamed, N. A. (2020). Classification of tomato leaf diseases using MobileNet v2. *IAES International Journal of Artificial Intelligence*, 9(2), 290.
- [11] Mohanty, S. P., Hughes, D. P. ve Salathé, M. (2016). Görüntü tabanlı bitki hastalığı tespiti için derin öğrenmenin kullanılması. *Frontiers in Plant Science*, 7, 1419.

[12] Sladojevic, S., Arsenovic, M., Anderla, A., Culibrk, D. ve Stefanovic, D. (2016). Yaprak görüntü sınıflandırmasıyla bitki hastalıklarının derin sinir ağlarına dayalı tanınması. Hesaplamalı Zeka ve Sinirbilim, 2016.

[13] Ferentinos, K. P. (2018). Bitki hastalığı tespiti ve teşhisi için derin öğrenme modelleri. Tarımda Bilgisayarlar ve Elektronik, 145, 311-318.

[14] Zhang, S., Zhang, S., Zhang, C., Wang, X. ve Shi, Y. (2019). Küresel havuzlama genişlemiş evrişimli sinir ağı ile salatalık yaprağı hastalığı tanımlama. Tarımda Bilgisayarlar ve Elektronik, 162, 422-430.

[15] Karthik, R., Hariharan, M., Anand, S., Mathikshara, P., Johnson, A. ve Menaka, R. (2020). Domates yapraklarında hastalık tespiti için dikkat gömülü kalıntı CNN. Uygulamalı Yumuşak Hesaplama, 86, 105933.

[16] Abbas, A., Jain, S., Gour, M., & Vankudothu, S. (2021). Tomato plant disease detection using transfer learning with C-GAN synthetic images. Computers and Electronics in Agriculture, 187, 106279.

[17] Sanida, T., Sideris, A., Sanida, M. V., & Dasygenis, M. (2023). Tomato leaf disease identification via two-stage transfer learning approach. Smart Agricultural Technology, 5, 100275.

[18] Nag, A., Chanda, P. R., & Nandi, S. (2023). Mobile app-based tomato disease identification with fine-tuned convolutional neural networks. Computers and Electrical Engineering, 112, 108995.

[19] Howard, A. G., Zhu, M., Chen, B., Kalenichenko, D., Wang, W., Weyand, T., ... & Adam, H. (2017). Mobilenets: Efficient convolutional neural networks for mobile vision applications. arXiv preprint arXiv:1704.04861.

[20] Sandler, M., Howard, A., Zhu, M., Zhmoginov, A., & Chen, L. C. (2018). Mobilenetv2: Inverted residuals and linear bottlenecks. In Proceedings of the IEEE conference on computer vision and pattern recognition (pp. 4510-4520).

[21] Howard, A., Sandler, M., Chu, G., Chen, L. C., Chen, B., Tan, M., ... & Adam, H. (2019). Searching for mobilenetv3. In Proceedings of the IEEE/CVF International Conference on Computer Vision (pp. 1314-1324).

[22] Shorten, C., & Khoshgoftaar, T. M. (2019). A survey on image data augmentation for deep learning. Journal of big data, 6(1), 1-48.



## Can Law Deter in Cyberspace? Türkiye's Experience in the Context of the Turkish Penal Code

Mehmet Onur ÖZER<sup>a,\*</sup>, Mehmet KAHRAMAN<sup>b</sup>

<sup>a,\*</sup> University of Missouri, Truman School of Government and Public Affairs, Political Science, Columbia, 65211, UNITED STATES of AMERICA

<sup>b</sup> Hatay Mustafa Kemal University, Faculty of Economical and Administrative Sciences, Political Science and Public Administration, Hatay, 31060, TÜRKİYE

### ARTICLE INFO

Received: 09.05.2025  
Accepted: 18.06.2025

**Keywords:** Cyberspace, cyber security, cyber deterrence, Turkish Penal Code

#### \*Corresponding

#### Authors

e-mail:  
ozermehmetonur@gmail.com

### ABSTRACT

The widespread use of cyberspace has significantly transformed the concept of security, introducing complex and novel threats to both individuals and states. This article explores the idea of deterrence in cyberspace, particularly focusing on its legal dimensions within Türkiye's regulatory framework. It starts by tracing the historical shift in security from physical protection to cyber defence and discusses how the digital domain, now regarded as the fifth domain of warfare, presents unique challenges to traditional deterrence models. Drawing on theoretical frameworks, particularly those proposed by Libicki and Nye, the article examines the feasibility of cyber deterrence along with the challenges posed by attribution, asymmetry, and cost dynamics. It further investigates the role of legal deterrence, emphasizing that effective deterrence in cyberspace requires more than just severe penalties; it also depends on the certainty, promptness, and enforceability of legal consequences. The article reviews Türkiye's legal and institutional responses, from early reforms to the Penal Code to contemporary laws aligned with international conventions like the Budapest Convention. Despite Türkiye's significant progress in regulating cybercrime, practices such as the Postponement of the Announcement of the Verdict (HAGB) and effective remorse reductions pose key weaknesses that undermine the deterrent capacity of the legal system. The study concludes by asserting the importance of coherent national legislation, international cooperation, and the consistent application of legal norms to establish a strong deterrent framework in cyberspace. This article is derived from the doctoral dissertation titled "Digitalization and Cybersecurity Based on National and International Security Policies: A Legal and Administrative Assessment" defended in 2023 at Hatay Mustafa Kemal University, Institute of Social Sciences, Department of Political Science and Public Administration.

DOI: 10.59940/jismar.1695163

## Hukuk Siber Uzayda Caydırıcı Olabilir mi? Türk Ceza Kanunu Bağlamında Türkiye'nin Deneyimi

### MAKALE BİLGİSİ

Alınma: 09.05.2025  
Kabul: 18.06.2025

#### Anahtar Kelimeler:

Siber uzay, siber güvenlik, siber caydırıcılık, Türk Ceza Kanunu

### ÖZET

Siber uzayın yaygın kullanımı, güvenlik kavramını önemli ölçüde dönüştürerek hem bireyler hem de devletler için karmaşık ve yeni tehditleri beraberinde getirmiştir. Bu makale, özellikle Türkiye'de siber uzayın yasal olarak düzenlenmesi bağlamında caydırıcılık konusunu incelemektedir. Fiziksel korumadan siber savunmaya doğru güvenliğin tarihsel dönüşümünü ele alarak başlayan çalışma, dijital alanın artık savaşın beşinci boyutu olarak kabul edilmesiyle geleneksel caydırıcılık modelleri açısından oluşturduğu zorlukları tartışmaktadır. Libicki ve Nye tarafından önerilen teorik çerçevelerden hareketle, makalede siber caydırıcılığın uygulanabilirliği; atfedilebilirlik, asimetri ve maliyet dinamikleri gibi sorunlar çerçevesinde ele alınmaktadır. Ayrıca hukuki caydırıcılığın rolü incelenmekte ve siber uzayda etkili bir caydırıcılığın yalnızca ağır yaptırımlarla değil; aynı zamanda hukuki sonuçların kesinliği, zamanında uygulanabilirliği ve icra edilebilirliği ile mümkün olabileceği vurgulanmaktadır. Makale, Türkiye'nin erken dönem

**\*Sorumlu Yazar**

e-posta:  
ozermehmetonur@gmail  
.com

reformlardan başlayarak Ceza Kanunu ile Budapeşte Sözleşmesi gibi uluslararası düzenlemelere uyumlu çağdaş kanunlara uzanan hukuki ve kurumsal tepkilerini değerlendirmektedir. Türkiye, siber suçların düzenlenmesi konusunda önemli ilerlemeler kaydetmiş olsa da hükmün açıklanmasının geri bırakılması (HAGB) ve etkin pişmanlık gibi uygulamalar, hukuki sistemin caydırıcılık kapasitesini zayıflatan temel sorunlar arasında yer almaktadır. Çalışma, siber uzayda güçlü bir caydırıcılık çerçevesi oluşturmak için tutarlı ulusal mevzuat, uluslararası iş birliği ve hukuki normların istikrarlı şekilde uygulanmasının önemine dikkat çekerek son bulmaktadır. Bu makale 2023 yılında Hatay Mustafa Kemal Üniversitesi, Sosyal Bilimler Enstitüsü, Siyaset Bilimi ve Kamu Yönetimi Ana Bilim Dalında savunulan “Ulusal ve Uluslararası Güvenlik Politikaları Temelinde Dijitalleşme ve Siber Güvenlik: Hukuksal ve Yönetimsel Bir Değerlendirme” başlıklı doktora tezinden türetilmiştir.

DOI: 10.59940/jismar.1695163

## 1. INTRODUCTION (GİRİŞ)

The concept and perception of security have evolved throughout history according to the social, economic, and political conditions of the era in which humanity has lived. Before the emergence of the first cities, the perception of security was based primarily on protection from natural disasters and wild animals in the natural environment. However, with the rise of human communities and the formation of the first cities, driven by the tendency of people to live together, this perception evolved from a struggle against nature to one based on human interactions. The concept of security, like many other concepts in the social sciences, is one of the contested notions over which no consensus has been reached. The common point among studies on security is that it refers to freedom from threats to fundamental values, both at the individual and societal levels. However, where these studies diverge is the issue of what basis the analysis should rest on [1]. In different regions of the world, struggles among communities for various reasons evolved into inter-state conflicts with the emergence of the first states, and the notion of security began to be addressed on a much broader scale. The rise of nation-states brought the concept of security into sharp focus at national and international levels. Each era's political and economic conditions and consequences have led to semantic shifts in understanding security.

In this context, security has been perceived differently in various historical periods: as national security in the context of military threats between states following the First and Second World Wars, as strategic balances and nuclear deterrence during the Cold War, and as a fight against terrorism after the September 11, 2001 attacks [2]. The significance of technology in ensuring national and international security became even clearer during the two world wars of the twentieth century and the long-standing Cold War between the United States and the Soviet Union. These historical periods demonstrated that technological superiority is far more critical than traditional manpower [3].

With the widespread use of internet technologies in the 21st century, technology has become an indispensable part of daily life. Today, people can conduct banking operations, commerce, and shopping online. Many electronic devices we use at home such as computers, phones, and various digital appliances can also be controlled through internet technology. This shift has given rise to new and distinct security threats. In this regard, through their extensive use of technology, individuals and states generate security threats that states must address.

With the increasing use of information technologies in almost every field, a new virtual realm called "cyberspace", or the "cyber domain" has emerged. In recent years, ensuring this domain's security has become a priority for states and technology-developing private companies. As this artificial digital environment has permeated every aspect of life, its use has become unavoidable and indispensable at both the individual and state levels. As individuals integrate technology into every aspect of their lives and states, digitize many bureaucratic services, and adapt to technology across various domains, new security threats have emerged. The concept of cybersecurity has thus arisen in response to these growing threats. It represents an effort to keep pace with digitalization and ensure security within cyberspace which is a realm that remains relatively new and highly complex, especially for states.

The extensive reach of cyberspace into nearly all aspects of life, the increased use of computer and internet technologies, and the fact that information sharing occurs in digital environments have collectively introduced new security threats. These threats in cyberspace not only individually endanger people, particularly in terms of the security of personal data, but also pose risks to national security through the potential for cyberattacks targeting the operating systems of critical infrastructure within states. As a result, the concept of security has undergone a significant transformation. With cyberspace now recognized as the fifth domain of warfare alongside land, air, sea, and space, security has taken on a new dimension. States and international organizations are

now keenly aware of the importance of securing this artificial domain and are actively developing security policies to address its challenges.

This awareness has not only encouraged states to develop new security policies to address the challenges in cyberspace but also pushed them to develop new technologies to fight against the threats in this new realm. States have started building new offense and defense capabilities to create deterrence in cyberspace. However, even though states have been building these capabilities actively, the complexity of cyberspace threats has grown drastically. Creating a legal framework to regulate this new domain has also become necessary for states. Having defense and offense capabilities in cyberspace could create deterrence for states. Are those capabilities strong enough deterrents to prevent a cyber-attack before it happens? How about the law? Can it deter in cyberspace? This article will focus on creating legal frameworks to regulate cyberspace as a deterrence system by evaluating Türkiye's experience.

## 2. WHAT IS “DETERRENCE” IN CYBERSPACE? (SİBER UZAYDA CAYDIRICILIK NEDİR?)

The concept of deterrence became a significant reality, particularly during the Cold War era, with the emergence of nuclear deterrence as a balancing factor in the arms race between the United States (U.S.) and the Soviet Union (USSR). The concept of nuclear deterrence has not lost its significance in the new world order that emerged after the collapse of the USSR. However, it is fair to say that it has moved away from the kind of “balance of terror” seen during the Cuban Missile Crisis in 1962, in which Türkiye also played a key role. In recent years, many researchers have argued that cyber threats define the 21st century, and these threats are poised to replace nuclear weapons in terms of their strategic importance.

Deterrence theory refers to the idea that an adversary's potential attack can be prevented by convincing them that such an action would either have no chance of success or would result in unacceptable costs, especially when measured in terms of cost-benefit analysis. Therefore, the capacity to carry out a retaliatory response to a potential attack is critically important [4]. However, whether states can achieve a deterrent power in cyberspace through conventional capabilities remains debatable.

Martin Libicki (2009) argues that cyber deterrence fundamentally differs from nuclear deterrence, making it less effective as a policy tool. While nuclear deterrence relies on symmetry, where adversaries

understand each other's capabilities and consequences, cyber operations often lack this clarity. In nuclear strategy, mutual awareness enables rational cost-benefit analysis, and the threat of catastrophic retaliation discourages attacks. Cyberspace, by contrast, obscures attribution and scale, weakening the logic of predictable deterrence. Cyber deterrence, however, diverges at this point. In cyberspace, it may not be possible to identify in advance where the threat is coming from or determine the actors involved in a potential cyberattack. Cyberspace is an environment where states, non-state actors, and sometimes even individuals can effectively operate. Thus, a large-scale cyberattack may originate from a single state, a non-state actor, or a coalition of multiple actors, making attribution and response far more complex [5].

Based on the data obtained during the period of nuclear tension between the United States and the Soviet Union throughout the Cold War, it is possible to have cyber deterrence. In this context, Libicki poses three core and six supporting questions highlighting the differences between nuclear and cyber deterrence. The first core question he asks is: “*Do we know who did it?*” [5]. This question is crucial because in the concept of deterrence, particularly when it comes to retaliation, it is essential to know who launched the attack and against whom a response should be directed. From this perspective, it is often extremely difficult to identify the source of cyber-attacks in cyberspace, which stands out as one of the main factors that makes cyber deterrence problematic.

The other two core questions Libicki poses are: “Can the adversary's assets be held at risk?” and “Can this be repeated?” [5]. Deterrence becomes possible when a potential attack can be prevented through the threat of retaliation before the aggressor acts. Therefore, if a party planning an attack knows that a retaliatory response could put its assets at risk and that the defending side can repeat such retaliation, it may decide not to proceed.

If the damage expected in return outweighs the anticipated harm inflicted by the attack, the attacker will realize that the costs outweigh the benefits, making the attack irrational. In this sense, for cyber deterrence to be credible, the defender must be able to retaliate and repeat that retaliation if necessary.

The other six supporting questions posed by Libicki are as follows:

1. *If retaliation is not a deterrent, can it at least disarm the adversary?*
2. *Will third parties become involved in the conflict?*

3. *Does the retaliation send the right message to our side?*
4. *Do we have a limit to our response?*
5. *Can we avoid escalation?*
6. *Is it worth it for the attacker to respond or launch an attack?* [5]

These are the questions Libicki suggests should be asked in retaliation against a potential cyberattack. When evaluated, it becomes clear that these questions are not fundamentally different from those posed in conventional deterrence. However, what sets cyber deterrence apart is that the cost of conducting a cyberattack is generally much lower than that of a conventional attack, and its effects are quite different from those of a nuclear strike. These differences position the concept of cyber deterrence in a distinct category.

Like Libicki, Joseph Nye (2011) emphasizes that there are significant differences between cyber technology and nuclear technology [3]. Libicki highlights these differences by stating that the damage or disconnection of a cyber system can inflict massive economic harm, while to underscore the devastating effects of nuclear war, he notes that a large-scale nuclear conflict could return humanity to the Stone Age [5]. Unlike nuclear threats, cyber threats are not clearly identifiable. Therefore, deterrence in cyberspace is a highly complex phenomenon and not limited to retaliation alone. The views on deterrence that emerged during the Cold War when the nuclear arms race intensified were relatively simple, centering on the idea that deterrence depended on the ability to retaliate against a nuclear strike.

During the Cold War, retaliatory capacity was the core of the deterrence concept. However, later studies and theories concluded that deterrence, especially in the context of the use of power, is far more complex than originally conceived. Moreover, conventional military forces, clear policy declarations, changes in alert levels, and troop movements supported nuclear deterrence [3].

According to Joseph Nye, the view held by some researchers that deterrence does not work in cyberspace due to its nature is misguided and overly simplistic. Although cyber deterrence may lack the robustness of traditional deterrence, it persists, particularly when considering reciprocity and restraint in interstate relations. In the face of an attack with an uncertain origin, governments may suddenly find themselves caught in a web of interconnected relationships that produce unintended consequences. For instance, during the Cold War, there was a relatively straightforward military dependency

between the United States and the Soviet Union. In contrast, today, the United States, China, and other countries exist within complex, overlapping networks. Thus, a large-scale cyberattack that harms the U.S. economy could also cause significant losses for China. The reverse is equally possible. China could be negatively affected by disruptions to interconnected systems that damage its interests [3]. Nye strongly emphasizes the cost advantages of cyberattacks in cyberspace. In contrast to traditional domains of warfare, where achieving dominance and control through conventional military power is highly costly, cyberspace presents a cost-effective environment. This environment allows non-state actors and states with limited conventional capabilities to operate effectively. Nye also argues that, in cyberspace, superiority is more likely to be achieved through offense rather than defence [6, 3].

As highlighted earlier, the technological developments since the 2000s and the advancements in internet technologies have made not only individuals but also states and non-state actors as integral parts of cyberspace. The complexity of cyberspace makes it extremely difficult to detect cyberattacks. Additionally, determining the intent behind such attacks is another significant challenge. In this context, the principle of proportionality becomes increasingly complicated, especially when the source of an attack is unknown. It is often difficult (if not impossible) to determine the level of a cyberattack, whether it was carried out by a state actor, a non-state actor, or an individual [6]. As a result, it becomes very difficult to determine the appropriate nature of the response. If the response does not comply with the principle of proportionality or is directed at the wrong party, the consequences could escalate into an armed conflict.

The uncertainty surrounding attribution and capacity in cyberattacks makes the concept of deterrence in cyberspace both highly complex and sensitive. The logic of deterrence rests on the idea that a potential attacker refrains from acting due to fear of the likely consequences of a retaliatory response based on the perceived capabilities and capacities of the defender. However, the uncertainty and complexity of cyberspace raise serious questions about the feasibility of deterrence in this domain [5].

Cyber deterrence has emerged as states increasingly use cyberspace, particularly for managing critical infrastructures. The issue of securing critical infrastructure in cyberspace, especially given the potentially severe consequences of a possible cyberattack, has become a key national security concern. The fact that critical infrastructures such as

transportation, energy, communications, finance, industry, and health are managed through SCADA (Supervisory Control and Data Acquisition) systems makes these infrastructures vulnerable to potential cyberattacks.

While each state defines critical infrastructures differently, systems whose disruption by a potential attack could threaten national security, hinder vital societal functions, or bring economic activity to a standstill are usually the general description of critical infrastructure structures. The United States defines critical infrastructures as physical or virtual systems so vital that their incapacitation or destruction would have a debilitating impact on physical or economic security or public health [7].

Türkiye defines critical infrastructures based on a regulation issued by the Ministry of Transport, Maritime Affairs, and Communications. In 2013 [8] as “infrastructures that contain information or industrial control systems where the confidentiality, integrity, or availability of processed information, if compromised, could lead to loss of life, large-scale economic damage, national security vulnerabilities, or disruption of public order” [8]. Additionally, in Türkiye’s 2020–2023 National Cybersecurity Strategy and Action Plan, published by the Ministry of Transport and Infrastructure, the designated critical infrastructure sectors are listed as Electronic Communications, Energy, Finance, Transportation, Water Management, and Critical Public Services [9].

The cyber-attacks in Estonia in 2007 and the "Stuxnet Attack" in Iran in 2010 increased the importance of deterrence capabilities for states in cyberspace. Cyberattacks in Estonia were a cornerstone for deterrence studies in cyberspace. Following these large-scale cyberattacks against Estonia in 2007, interest in cyber deterrence and related studies increased significantly both at the level of academic research and in the form of state-level measures and responses. Although the perpetrators and exact origin of the attacks were never definitively identified, it was widely claimed that Russia was behind them [10]. The attacks, which lasted for three weeks, rendered the websites and systems of the presidency, parliament, ministries, political parties, major newspapers, banks, and companies managing inter-institutional communication inoperable, creating a full-blown digital crisis in the country [11]. This massive cyberattack on Estonia caused widespread disruption of services and brought inter-agency communication to a near halt. Just one year later, similar cyberattacks were launched against Georgia, reportedly again by Russia, and produced comparable effects [12]

In 2010, a cyberattack allegedly carried out by the United States with assistance from Israel targeted Iran’s Natanz nuclear facility near Isfahan. The attack reportedly disrupted Iran’s uranium enrichment operations by destroying almost 1,000 centrifuges and even caused the reactors to function dangerously uncontrolled. This attack used a virus called Stuxnet, which has since entered the literature as the "Stuxnet Attack" [13, 14].

## 2.1. Deterrence in Cyberspace in the Context of Laws and Regulations *(Hukuk ve Mevzuat Bağlamında Siber Uzayda Caydırıcılık)*

The complex nature of cyberspace not only makes it difficult for states to maintain deterrent capabilities in this artificial domain in terms of national security but also complicates the establishment of legal deterrence through the regulation of cyberspace to prevent potential criminal activities. In legal discourse, experts discuss deterrence as an extension of criminal law, and domestic and international literature offers various theories on this issue.

In general, deterrence in the fight against crime is evaluated by the nature of the punishment imposed for a given offense. At this point, legal deterrence is often understood as the deterrent effect of punishment. While the idea that effective deterrence comes from harsh penalties is widespread, legal deterrence should not be viewed solely in terms of punishment severity. It must also be examined in connection with the overall structure and functioning of a country's criminal justice system [15]. Therefore, legal deterrence aimed at preventing crimes in cyberspace through its legal regulation depends on the precise definition of offenses and penalties in law and, more importantly, on their enforceability.

For the criminal justice system to have a preventive and deterrent effect against offenses, certain principles must be in place regarding the enforceability of punishments for clearly defined crimes in law. These principles are certainty, swiftness, and severity of punishment. The principle of certainty means that everyone is judged equally before the law and that if an individual commits a crime, the corresponding punishment will inevitably be applied sooner or later. The principle of swiftness refers to the prompt apprehension of offenders after a crime has occurred, followed by timely investigation, prosecution, and adjudication, leading to the finalization of the sentence. The severity of punishment means that the penalty imposed must be proportionate to the offense committed. If these three principles are absent, punishments lose their deterrent effect [16].

## 2.2. Deterrence in Cyberspace in International Legal Context *(Uluslararası Hukuk Bağlamında Siber Uzayda Caydırıcılık)*

Cyberspace has become a fundamental domain for most states. In terms of competition, nations seek to deter adversaries from malicious cyber activities through threats of retaliation or denial of benefits. However, creating a deterrence system in cyberspace requires unique legal and practical solutions. Unlike the conventional understanding of deterrence issues such as military attacks, cyber incidents often create a complexity that differentiates the line between crime, espionage, and armed aggression, challenging states to respond within the bounds of international law. Therefore, the international community has gradually recognized that existing international legal norms should be extended to cyberspace.

The international community hasn't been able to create a globally accepted treaty that regulates cyberspace. However, existing international law provides a normative framework for states that cyber deterrence strategies must operate. While no single treaty is dedicated solely to cyber operations by states, various existing legal frameworks regulate state actions in cyberspace. These include the UN Charter's rules on the prohibition of force and the right to self-defence, core international law principles such as sovereignty and non-intervention, and specific agreements like the Budapest Convention on cybercrime. Additionally, non-binding instruments and expert manuals have helped establish norms and guide state behaviour in the digital domain.

### 2.2.1. International Legal Frameworks Relevant to Cyber Deterrence *(Siber Caydırıcılıkla İlgili Uluslararası Hukuki Çerçeveseler)*

Several international legal norms and instruments are relevant to cyber deterrence. Legally binding instruments such as the UN Charter establish core rules prohibiting using force that also applies to cyber operations. At the same time, customary international law addresses areas not covered explicitly, including sovereignty and state responsibility. Additionally, non-binding initiatives like the UN Group of Governmental Experts (GGE) norms and expert analyses like the Tallinn Manual offer additional guidance on appropriate conduct. Collectively, these frameworks create a structured environment that informs how states design cyber deterrence strategies by clarifying unacceptable behaviours and legitimate responses in cyberspace.

Although no single treaty comprehensively regulates state behaviour in cyberspace, several binding and non-binding legal instruments have emerged to shape

expectations, responsibilities, and consequences surrounding cyber operations. Together, these frameworks establish the normative foundation upon which states design and implement cyber deterrence strategies.

This framework's core is the UN Charter [17], which applies fully to cyberspace. Article 2(4) prohibits the use of force against the territorial integrity or political independence of any state. In contrast, Article 51 affirms the inherent right of self-defence in the event of an "armed attack." These provisions form the legal bedrock for deterrence by punishment in cyberspace: a sufficiently severe cyber operation—causing death, injury, or significant physical destruction could be interpreted as a use of force, thus justifying a forcible response [17]. However, the Charter offers limited utility for deterring most cyber activities below the armed conflict threshold, creating persistent challenges in addressing so-called "grey zone" operations [18].

Customary international law fills some of these regulatory gaps. The principles of sovereignty, non-intervention, and state responsibility apply to cyberspace and help delineate acceptable conduct. Sovereignty protects a state's control over its cyber infrastructure, while non-intervention prohibits coercive interference in domestic affairs, such as election manipulation or fomenting unrest [19]. The doctrine of state responsibility, codified in the Articles on State Responsibility, enables using proportionate countermeasures in response to internationally wrongful cyber acts [20]. These principles support deterrence by norms, emphasizing that breaches of international obligations, if attributable to a state, can prompt diplomatic, legal, or cyber retaliatory measures. However, the difficulty of attribution remains a critical weakness in operationalizing these norms as effective deterrents [21].

International Humanitarian Law (IHL) becomes applicable in armed conflict. Instruments such as the Geneva Conventions and the Hague Regulations impose obligations to respect the principles of distinction, proportionality, humanity, and necessity, even in cyber warfare [22]. IHL constrains cyber operations targeting civilian infrastructure and reinforces the notion that cyberspace is not exempt from wartime legal constraints. While IHL does little to deter peacetime activities, it plays a vital role in preventing escalatory cyber actions during conflicts by classifying certain cyberattacks as potential war crimes.

On the criminal enforcement side, the Budapest Convention on Cybercrime (2001) strengthens

deterrence through legal accountability. By mandating the criminalization of specific cyber offenses and facilitating international cooperation in cyber investigations, the Convention contributes to deterrence by law enforcement, particularly against non-state actors and proxy groups [23]. Nevertheless, the Convention's normative reach is limited by the absence of key cyber powers such as Russia and China, who reject what they perceive as Western-centric legal standards [24].

A range of non-binding but influential instruments also shape state behaviour in cyberspace. The UN Group of Governmental Experts (GGE) and Open-Ended Working Group (OEWG) processes have produced voluntary norms for responsible state conduct, including commitments to refrain from targeting critical infrastructure or emergency response teams during peacetime [24, 25]. These norms serve as a basis for deterrence through shared expectations, enabling collective condemnation and sanctions in response to violations.

The Tallinn Manual 2.0, an expert commentary that analyzes how existing international law applies to cyber operations in peacetime and wartime, provides further interpretive guidance. Although not a binding legal instrument, it has significantly influenced state practice and legal doctrine [18]. The Manual helps states articulate "red lines" by elaborating on when cyber operations might constitute uses of force or armed attacks, thus supporting more credible deterrence postures grounded in legal reasoning.

Lastly, regional and multistakeholder initiatives, including NATO's cyber policy, the EU's Cyber Diplomacy Toolbox, and global norms like the Paris Call for Trust and Security in Cyberspace, bolster deterrence by signalling collective responses and enhancing resilience [27, 28]. NATO's declaration that a major cyberattack could trigger Article 5 collective defence obligations adds weight to deterrence by alliance commitments. Meanwhile, EU-led sanctions and private-sector engagement increase the cost of cyber aggression through diplomatic, economic, and reputational consequences.

### 3. TÜRKİYE'S EXPERIENCE IN REGULATING CYBERSPACE (SİBER UZAYIN DÜZENLENMESİNDE TÜRKİYE'NİN DENEYİMİ)

Since the early 1990s, Türkiye has undertaken numerous legal and administrative measures to address potential cyber security threats. Although various actors implemented many of these regulations independently, they still represent necessary steps toward ensuring cyberspace security. While the legal

regulation of this field through various laws, regulations, circulars, and communiqués has sometimes created inconsistencies, the legal measures introduced remain highly significant in establishing deterrence against potential threats that may arise in cyberspace.

No single law in Türkiye comprehensively regulates crimes committed in cyberspace. Instead, incorporating relevant provisions into existing laws has addressed offenses in the field of information technologies [29]. The first legal regulation regarding cyber-related crimes was introduced on June 6, 1991, through the "Law No. 3756 on the Amendment of Certain Articles of the Turkish Penal Code." By the early 2000s, with the increasing use of cyberspace, Türkiye began to take more concrete and serious steps toward ensuring cybersecurity and establishing deterrence in cyberspace. In this context, a far more comprehensive regulation than the 1991 amendment was enacted in 2004 when the concept of cybercrime was legally defined. Under the heading "Crimes in the Field of Information Technology" Chapter Ten of the Turkish Penal Code No. 5237 included significant legal provisions, particularly focused on offenses committed in the cyber domain.

In Türkiye, the most comprehensive legal regulation of the Internet was enacted in 2007 through Law No. 5651 on the Regulation of Publications on the Internet and Combating Crimes Committed Through Such Publications [30]. Another significant legal regulation to ensure cyberspace security was the Electronic Communications Law No. 5809, adopted in 2008 [31]. This law intends to prevent unfair competition in the electronic communications sector and to ensure that services in this field are delivered actively and effectively. It was an important step toward safeguarding individuals' freedom and security of communication in cyberspace, especially in protecting fundamental rights and freedoms.

In addition to the Electronic Communications Law, another critical piece of legislation for ensuring cybersecurity was Law No. 6698 on the Protection of Personal Data, which was submitted to parliament in the same year [32]. Although it took a long time, it was enacted and published on April 7, 2016. Today, with the widespread use of e-government applications and the storage of personal data in digital environments across nearly all public institutions, not to mention digital storage on shopping websites and social media platforms, this law serves as an essential deterrent against malicious actors, particularly in terms of protecting the privacy of personal life.

### 3.1. Turkish Penal Code and Deterrence in Cyberspace (*Türk Ceza Kanunu ve Siber Uzayda Caydırıcılık*)

Although several laws regulate cyberspace in Türkiye, the Turkish Penal Code is a key legal framework that creates deterrence. As highlighted earlier, Türkiye is one of the countries that recognized the importance of cyberspace and its security at an early stage and took legislative action accordingly. Aware of the issue's significance as early as the 1990s, Türkiye introduced its first regulation on crimes committed in cyberspace on June 6, 1991, through Law No. 3756 on the Amendment of Certain Articles of the Turkish Penal Code No. 765. Article 20 of this law, titled "Crimes in the Field of Information Technology," made it a criminal offense to unlawfully obtain, use, transmit, or reproduce programs, data, or other elements from an automated data processing system, especially if done with the intent to harm others. The law also set forth the provisions for penalties related to such offenses [33].

In 2004, a far more comprehensive regulation than the 1991 amendment was introduced when the concept of cybercrime was formally defined by law. Under the title "Crimes in the Field of Information Technology" in Chapter Ten of the Turkish Penal Code No. 5237, provisions were made concerning unauthorized access to information systems, obstruction or disruption of systems, deletion or alteration of data, and the misuse of bank and credit cards. These offenses are independently regulated under Articles 243, 244, and 245 of the Turkish Penal Code [34]. Additionally, Türkiye became a party to the Council of Europe Convention on Cybercrime, signed in Budapest in 2001 (commonly referred to as the Budapest Convention) after being ratified by the Grand National Assembly of Türkiye (TBMM) in 2012. Following its ratification, Türkiye amended the Turkish Penal Code to align with the provisions of this international agreement.

Article 20 of Law No. 3756 added a new "Crimes in the Field of Information Technology" section to the Turkish Penal Code No. 765 as "Chapter Eleven" to follow Article 525. Articles 21, 22, 23, and 24 of the same law also introduced Articles 525a, 525b, 525c, and 525d, which were appended to the Penal Code under the same new chapter.

These articles are particularly significant as they represent the first legislative amendments made in Türkiye to ensure cyberspace security. Accordingly, the following provisions are set forth as they appear in the law:

Article 525a:

"Any person who unlawfully obtains programs, data, or any other elements from a system that processes information automatically shall be sentenced to imprisonment from one to three years and a heavy fine ranging from one million to fifteen million Turkish lira. The same penalty shall also apply to any person who uses, transmits, or reproduces a program, data, or any other element in a system that processes information automatically, intending to cause harm to another." [35].

Article 525b:

"Any person who, with the intent to cause harm to another or to obtain benefit for themselves or others, partially or completely destroys, alters, deletes, obstructs the operation of, or causes the incorrect functioning of a system that processes information automatically, or its data or any other element, shall be sentenced to imprisonment from two to six years and a heavy fine ranging from five million to fifty million Turkish lira. Any person who unlawfully obtains a benefit for themselves or others by using a system that processes information automatically shall be sentenced to imprisonment from one to five years and a heavy fine ranging from two million to twenty million Turkish lira." [35].

Article 525c:

"Any person who, for the purpose of creating a forged document to be used as legal evidence, inputs data or other elements into a system that processes information automatically or alters existing data or elements shall be sentenced to imprisonment from one to three years. Those knowingly using the forged or altered data shall be imprisoned for six months to two years." [35].

The amendments to the Turkish Penal Code (TCK) by Law No. 3756 gained further significance in 1993 when Türkiye was introduced to the Internet through initiatives led by Middle East Technical University (METU). In 1991, when the law was enacted, individual computer use in Türkiye was still very limited. Therefore, adopting a legal regulation when internet technology had not yet begun to be widely used aimed at creating deterrence against crimes in cyberspace should be considered a noteworthy development.

Another critical point to emphasize is that this regulation holds great significance within the scope of the principle of legality. This principle was first formulated by German criminal law scholar Anselm von Feuerbach as "*nullum crimen, nulla poena sine lege*", translated as "*no crime, no punishment without law*." In Türkiye, the principle of legality in crimes

and punishments is guaranteed under Article 13 of the 1982 Constitution, which states:

*"Fundamental rights and freedoms may be restricted only by law and solely for the reasons set forth in the relevant articles of the Constitution, without infringing upon their essence. These restrictions shall not violate the letter and spirit of the Constitution, the requirements of the democratic order of society, or the principles of the secular Republic, and shall comply with the principle of proportionality."* [36]

This provision represents the constitutional embodiment of the principle of legality. Similarly, Article 38 of the Constitution, titled "Principles Relating to Offenses and Penalties," further reinforces this principle by stating, *"No one shall be punished for any act that was not defined as a crime by law at the time it was committed; nor shall anyone be subjected to a heavier penalty than the one prescribed by law at the time the offense was committed. The provisions of the above paragraph shall also apply to statutes of limitations for offenses and penalties, as well as to the legal consequences of criminal convictions. Criminal penalties and security measures in lieu of penalties may only be imposed by law."* [36]

Therefore, the 1991 amendment to the Turkish Penal Code essentially paved the way for acts committed in cyberspace to be legally recognized as crimes, thereby enabling the initiation of investigation and prosecution processes related to such actions.

Instead of terms such as "computer" or "information technology," the law used the phrase "a system that processes information automatically." Considering that the widespread use of computer and software technologies had not yet begun at the time, this definition was intended to encompass all technological devices, from the simplest data processing systems to the most advanced computers of the period. Through this regulation, the law established a legal basis for acts committed using or through such devices, assigning them a material and legal meaning. At the time, this regulation was enacted when computer use in Türkiye was still relatively new, and it is true that threats in the context of cybersecurity were quite limited. Therefore, no specific definition was provided regarding the nature of the offense in the regulation. However, with the advent of the internet and its integration into daily life, the concept of crimes committed in cyberspace began to take on real meaning. Thus, this regulation marked an important step for Türkiye in establishing a legal basis for such crimes.

In 2004, Law No. 5252 on the Enforcement and Implementation of the Turkish Penal Code repealed the Turkish Penal Code No. 765, which was replaced by Penal Code No. 5237. The new Penal Code addressed crimes committed in cyberspace much more comprehensively than the 1991 regulation. Under the title "*Crimes in the Field of Information Technology*" Chapter Ten of the Turkish Penal Code No. 5237 introduced provisions related to unauthorized access to information systems, obstruction or disruption of systems, deletion or alteration of data, and the misuse of bank and credit cards [37]. These offenses are independently regulated under Articles 243, 244, and 245 of the Penal Code [34].

Article 243 of the Turkish Penal Code (TCK) regulates the offense of unauthorized access to information systems. According to this article, any person who unlawfully accesses all or part of an information system is subject to up to one year of imprisonment or a judicial fine. The second paragraph of the same article states that if this act is committed against systems that are available for use in exchange for payment, the penalty shall be reduced by half. Finally, the third paragraph stipulates that if, because of this act, the data contained in the system is deleted or altered, the offender shall be sentenced to imprisonment from six months to two years [34].

Article 244 of the Turkish Penal Code (TCK) addresses the crimes of obstructing a system, disrupting its functioning, and deleting or altering data. Compared to Article 243, this article provides a more detailed regulation of the offense of interfering with an information system. Article 244 of the Turkish Penal Code states: (1) Any person who obstructs or disrupts the operation of an information system shall be imprisoned for one to five years. (2) Any person who corrupts, deletes, alters, renders inaccessible, inserts data into, or transfers existing data from an information system shall be imprisoned from six months to three years. (3) If the act is committed against an information system belonging to a bank, credit institution, or a public institution or organization, the penalty shall be increased by half. (4) If these acts are committed in such a way as to benefit oneself or another unjustly, and if the act does not constitute another offense, the offender shall be punished with imprisonment from six months to two years and a judicial fine of up to five thousand days." [34]. This provision regulates the offenses of obstructing, disrupting, deleting, or altering a system or its data. The purpose behind defining this offense is to ensure compliance with the "data interference" provision in Article 4 and the "system interference" provision in Article 5 of the Budapest Convention.

Article 245 of the Turkish Penal Code (TCK) regulates the offense of misuse of bank or credit cards. Under this article, acts involving the misuse of bank and credit cards are defined as a distinct crime category, aiming to prevent financial harm to banks or their customers and the unlawful acquisition of benefits through such means. According to the article, any person who uses a bank or credit card belonging to someone else without the consent of the cardholder or the person authorized to possess the card, thereby obtaining a benefit, shall be punished with imprisonment from three to six years and a judicial fine of up to five thousand days. Furthermore, anyone who produces, sells, transfers, purchases, or accepts counterfeit bank or credit cards using fake bank accounts shall be imprisoned for three to seven years and a judicial fine of up to ten thousand days. The third paragraph of this article states that if a counterfeit bank or credit card is used to obtain a benefit, and if this act does not constitute another offense that requires a more severe penalty, the offender shall be sentenced to imprisonment from four to eight years and a judicial fine of up to five thousand days [34].

With the amendment introduced in 2016, Article 245/A, titled "Prohibited Devices and Programs," was added to the section on Crimes in the Field of Information Technology in the Turkish Penal Code. This article established a significant regulation regarding the use and production of devices and software employed to commission offenses regulated under this section. According to the article: *"If a device, computer program, password, or other security code is manufactured or created exclusively for the purpose of committing the offenses outlined in this section or other crimes that can be committed using information systems as tools, any person who manufactures, imports, dispatches, transports, stores, accepts, sells, offers for sale, purchases, distributes to others, or possesses such items shall be punished with imprisonment from one to three years and a judicial fine of up to five thousand days."* [34]. With this regulation, lawmakers introduced criminal sanctions for the hardware and software tools used or produced for the commission of cybercrimes.

In the Turkish Penal Code (TCK), the regulation of cybercrimes within the scope of substantive criminal law is not limited to the section titled "Crimes in the Field of Information Technology." Although not specifically designed for cybercrimes, the Turkish Penal Code addresses offenses committed using or through information technologies in various other contexts. In particular, the entry into force of the Budapest Convention and the obligations arising from this international treaty prompted harmonization

efforts in domestic law. Recognizing that traditional crimes can also be committed through information technologies, the TCK incorporates relevant provisions across several articles.

### 3.1.1. Articles of the Turkish Penal Code Associated with or Potentially Applicable to Crimes Committed Using Information Technologies or Through These Technologies (*Türk Ceza Kanunu'nda Bilişim Teknolojileri Kullanarak veya Bu Teknolojiler Aracılığıyla İşlenen Suçlarla İlişkilendirilen ya da İlişkilendirilebilecek Maddeler*)

In addition to Articles 243, 244, and 245, which specifically address cybercrimes under the category of information technology offenses in the Turkish Penal Code (TCK), numerous other articles are associated with or potentially applicable to crimes committed using or through information technologies. Many of these provisions were introduced or amended after the signing of the Budapest Convention as part of Türkiye's efforts to harmonize its domestic legislation with the Convention's requirements. The relevant articles can be listed as follows:

**Article 123/A – Persistent Stalking** (*Added: 12/5/2022 – Law No. 7406, Article 8*): In the first paragraph of the article, the following provision is introduced: *"Anyone who persistently follows a person physically or attempts to make contact using communication tools, information systems, or third parties in a way that causes serious discomfort to that person or makes them fear for their own or a relative's safety shall be sentenced to imprisonment from six months to two years."* [34]

This regulation considers persistent stalking not only in its physical form but also when carried out through communication tools and information systems. Although it does not provide a detailed definition of stalking via information systems, it encompasses acts of persistent harassment conducted in cyberspace that cause discomfort or create concerns for personal security.

This article was added to the TCK in 2022, reflecting the growing recognition that such behaviours increasingly occur in cyberspace and thus must be addressed through appropriate legal measures.

**Article 124 – Obstruction of Communication:** This article does not include specific provisions or references to information technologies or cyberspace. Nor does it clarify how or through which means the offense may be committed. Nevertheless, the article defines the offense of obstructing communication as

follows: (1) Anyone who unlawfully obstructs communication between individuals shall be imprisoned for six months to two years or a judicial fine. (2) Anyone who unlawfully obstructs communication between public institutions shall be imprisoned for one to five years. (3) If the unlawful obstruction concerns any form of media or broadcasting outlet, the penalty stated in the second paragraph shall be applied [34].

Although the article does not explicitly mention cyber or digital means, its broad wording allows for interpretation that may include acts committed via information systems, especially as cyber-based disruptions to communication become more prevalent.

**Article 132 – Violation of the Confidentiality of Communication:** This article sets out the criminal sanctions to be applied in cases where the confidentiality of communication between individuals is violated, including the recording of communication content, the unlawful disclosure of such content, and the unlawful disclosure of communications involving the person themselves. Given the widespread use of smartphones and the current level of internet technology, internet-based messaging, and video call applications have become the primary means of communication between individuals. Especially during the COVID-19 pandemic, when curfews restricted people from leaving their homes, internet technology became the most important communication tool, leading to the decline of traditional communication methods. In this context, it becomes evident that the relevant article of the Turkish Penal Code is directly related to the offense of violating the confidentiality of communication as it may occur in cyberspace.

**Article 133 – Listening to and Recording Conversations Between Individuals:** Although, as in other articles, this provision does not explicitly address the use of information technologies or the commission of such acts through digital means, cyberspace is the primary medium where such offenses are committed or can be committed today. In an era where internet technology is heavily used to communicate, listening to, recording, and disclosing private conversations between individuals increasingly occurs via internet-based platforms, making this a cyber-enabled offense in practice. Smartphones, now used by nearly everyone and always carried, serve as communication devices and tools for audio and video recording. Furthermore, as discussed in the section on cybersecurity threats, unauthorized access to networks for the purpose of illegal surveillance and recording has become a highly

probable occurrence. In this context, it can be argued that the relevant article of the Turkish Penal Code is directly related to cybersecurity threats.

**Article 134 – Violation of Privacy:** This article of the Turkish Penal Code regulates the violation of an individual's right to privacy, specifically when such a violation is committed through the recording of images or audio and the unlawful disclosure of these recordings. The provision refers to privacy infringement by recording visual or auditory content. Still, it does not address whether this is done using information technologies or specify the platforms through which such acts occur. However, the lack of a specific reference to digital means does not prevent the application of this article to offenses committed in cyberspace. There is no legal barrier to interpreting and applying this provision to privacy violations that occur via digital or cyber platforms.

**Article 135 – Unlawful Recording of Personal Data:** This article of the Turkish Penal Code regulates the offense of unlawfully recording personal data. According to Article 135: (1) Any person who unlawfully records personal data shall be imprisoned for one to three years. (2) If the personal data concerns individuals' political, philosophical, or religious views; racial origins; or their unlawful moral tendencies, sexual lives, health conditions, or trade union affiliations, the penalty under the first paragraph shall be increased by half." [34]

The article does not distinguish whether the offense is committed through information technologies or by other means. However, considering that virtually all types of data, from government institutions to individual users, are now stored in digital environments, the primary medium through which this offense is likely to occur today is cyberspace, particularly through computers and digital systems. Therefore, while the article does not explicitly prescribe a penalty for committing this crime in cyberspace, there is no legal barrier to applying it to cases involving personal data theft in the digital realm.

**Article 136 – Unlawful Transfer or Acquisition of Data:** As with many other articles that can be associated with the security of cyberspace, this article does not address the use of digital devices or the commission of the offense through information technologies, nor does it provide any specific explanation regarding this issue. The article regulates the unlawful acquisition and transfer of personal data as follows: "(1) Any person who unlawfully transfers, disseminates, or acquires personal data shall be sentenced to imprisonment from two to four years. (2) If the subject of the offense involves statements and

*recordings made in accordance with the fifth and sixth paragraphs of Article 236 of the Code of Criminal Procedure, the penalty shall be doubled.” [34]*

As can be seen, the article does not specify the tools used in committing the offense or the environment in which it is carried out. The second paragraph, which was added by an amendment in 2019, refers to statements and recordings taken from child victims during the investigation phase of the offense defined under Article 103 of the Penal Code ("Sexual Abuse of Children"), in accordance with paragraphs five and six of Article 236 of the Code of Criminal Procedure (CMK). Therefore, the "statements and recordings" referred to in the second paragraph of Article 136 pertain specifically to those obtained during investigations related to child sexual abuse cases.

**Article 142/2-e – Aggravated Theft:** The theft offense, addressed in Chapter Ten of the Turkish Penal Code under "Crimes Against Property," is examined under two categories: theft and aggravated (qualified) theft. According to subparagraph (e) of paragraph 2 in Article 142, if the offense is committed using information systems, the perpetrator shall be sentenced to imprisonment from five to ten years [34]. The Cyber Crimes Department of the Turkish National Police defines aggravated theft as a cybercrime involving unauthorized data acquisition from a system or during data transmission between systems through malicious software. Examples include the theft of in-game characters in online games and the unauthorized transfer of money from one bank account to another [37]. The explicit inclusion in the TCK of the offense of aggravated theft committed using information systems can be considered a preventive legal measure aimed at addressing such crimes committed in cyberspace.

**Articles 213–218 – Offenses Against Public Peace:** The offenses listed under the section "Crimes Against Public Peace" in the Turkish Penal Code do not explicitly address acts committed in cyberspace or through the use of information technologies. However, if these offenses are carried out via digital technologies, there is no legal obstacle to applying the relevant articles in such cases.

The relevant articles are:

- Article 213: Threat intended to cause fear and panic among the public
- Article 214: Incitement to commit a crime
- Article 215: Praising an offense or offender
- Article 216: Incitement to hatred and hostility or public denigration
- Article 217: Incitement to disobey the law
- Article 217/A: Public dissemination of misleading information

Although the offenses described above are traditionally understood as conventional crimes, each can easily be committed in digital environments. Moreover, social media platforms, now used by nearly everyone, are among the primary digital spaces where such crimes can occur. On these platforms, where each individual can act like a personal media outlet, a single post can reach massive audiences within minutes.

Therefore, the crimes addressed in Chapter Five of the Turkish Penal Code can be committed easily through such channels. In this context, clearly defined penalties for these offenses in the law can be seen as a deterrent factor. However, cyberspace's complex, vast, and borderless nature makes it increasingly difficult to identify and apprehend perpetrators of such crimes.

The borderless nature of cyberspace allows these offenses to be committed from beyond national jurisdictions. As a result, although relevant provisions exist in the Turkish Penal Code, they sometimes fail to function effectively as deterrents. For this reason, international cooperation is critical in combating cross-border cyber offenses and establishing effective deterrence mechanisms in cyberspace.

**Article 226 – Obscenity:** The offense of obscenity, addressed in the section "Crimes Against Public Morality" of the Turkish Penal Code, is particularly significant in cyberspace due to the ease with which this offense can be committed online. Although the article does not explicitly address the commission of the offense using or through information technologies, it does define as a criminal act the sale, rental, distribution, publication via press and media, or facilitation of the distribution of obscene images, texts, or expressions, and prescribes a prison sentence of six months to five years, depending on the method of commission.

When considered within the scope of Article 9 of the Budapest Convention, which deals with Offenses Related to Child Pornography [38], the importance of Article 226 increases. The provision criminalizes the display, reading, distribution, or provision of obscene materials in places accessible to children or directly to children. More importantly, it foresees a prison sentence of five to ten years and a judicial fine of up to five thousand days for individuals who use children, child-like representations, or persons made to look like children in the production of such materials.

Given the current capabilities of AI, deepfake, and animation technologies, the criminalization of obscene images featuring persons made to appear as

children in digital environments is a crucial measure for preventing such crimes in cyberspace. Although the explicit criminalization of virtual child pornography remains a significant gap in the law, there is no legal barrier to prosecuting such acts under the existing provisions of this regulation.

**Article 228 – Providing a Place and Opportunity for Gambling:** This article prescribes imprisonment from one to three years and a judicial fine of no less than two hundred days for individuals who provide a place or opportunity for gambling. The increasing use of information and internet technologies creates gambling environments in virtual spaces. has become much easier.

In 2017, an amendment was made to this article specifying that if the offense is committed through the use of information systems, the penalty shall be three to five years of imprisonment and a judicial fine of one thousand to ten thousand days. However, as emphasized earlier, the borderless nature of cyberspace makes it very difficult to apprehend individuals committing this offense.

In cases where online gambling is facilitated through websites hosted on servers located abroad, apprehension and prosecution of the offenders are impossible without international cooperation. In such instances, the most that authorities can do is restrict access to the website. Thus, even though this offense is addressed in the legislation, there is a clear need for stronger international collaboration to combat it effectively.

**Articles 209–301 – Offenses Against the Symbols of State Sovereignty and the Dignity of State Institutions:** This section of the Turkish Penal Code addresses offenses such as insulting the President (Article 299), denigrating the symbols of state sovereignty (Article 300), and insulting the Turkish Nation, the State of the Republic of Türkiye, or its institutions and organs (Article 301). Although these articles do not include specific provisions for cases where such offenses are committed using information technologies, there is no legal obstacle to applying these provisions when such acts occur in cyberspace. Given that these offenses can easily be committed via digital platforms, particularly on social media, these laws can also be enforced in response to online conduct.

**Articles 326–339 – Offenses Against State Secrets and Espionage:** This section of the Turkish Penal Code addresses offenses such as the acquisition, destruction, forgery, and disclosure of information and documents that relate to the security of the state

or its domestic and foreign political interests, and which are required to be kept confidential. Given that today, most information is stored digitally and a significant portion of communication and data exchange between institutions occurs over internet-connected networks, the unauthorized acquisition of such information in virtual environments is highly likely. Although the relevant articles do not specifically reference the commission of these offenses in cyberspace or through information technologies, it is clear in the current information age that a large portion of sensitive data is stored electronically and can potentially be accessed via cyberattacks on these systems.

It is also useful to examine judicial practices and court rulings in assessing the effectiveness of cybercrime provisions in the Turkish Penal Code. However, implementing measures such as the Postponement of the Announcement of the Verdict (HAGB) and the Effective Remorse Reduction has been viewed as a major weakness in preventing cyber offenses.

The Postponement of the Announcement of the Verdict (HAGB) is regulated under Article 231 of the Criminal Procedure Code (CMK). According to paragraph 5 of the article: *“If the sentence imposed upon the defendant because of the trial for the charged crime is imprisonment of two years or less or a judicial fine, the court may decide to postpone the announcement of the verdict. The provisions regarding reconciliation remain reserved. The postponement of the announcement of the verdict means that the judgment will not have legal consequences for the defendant.”* [39]

Paragraph 6 outlines the conditions for applying HAGB: *“To decide on the postponement of the announcement of the verdict: a) The defendant must not have been previously convicted of an intentional crime; b) The court must be convinced, based on the defendant's character, behaviour in court, and other personal qualities, that they are unlikely to re-offend; c) The harm caused to the victim or public due to the offense must be fully compensated by restitution, restoration, or reparation. The defendant's consent is required for the decision.”* [39]

In this context, HAGB may be applied to offenses such as:

- Unauthorized Access to Information Systems [34]
- Disruption or Destruction of Systems or Data [34]
- Misuse of Bank or Credit Cards [34]

According to Turkish Penal Code Article 245(5):

*“For the acts listed in the first paragraph, the provisions on effective remorse for crimes against property shall apply.” Thus, under Article 245(1): “Any person who, by any means, obtains or retains another person’s bank or credit card and uses it or has it used without the consent of the cardholder or authorized party, to benefit themselves or another, shall be punished with imprisonment from three to six years and a judicial fine of up to five thousand days.” [34]*

If the offender fulfils the conditions set forth in Turkish Penal Code, Article 168 (Effective Remorse), a reduction in the sentence may apply. If the victim's damages are compensated during the investigation phase, the sentence may be reduced by up to two-thirds. If compensation occurs during the prosecution phase (i.e. after the case has been filed), the sentence may be reduced by up to one-half [34].

Although the cybercrime provisions of the TCK may be seen as a legal deterrent against offenses in cyberspace, HAGB and effective remorse reductions weaken this deterrent effect. These practices undermine the enforceability of penalties and diminish the dissuasive power of the legal framework regarding cyber offenses.

#### 4. CONCLUSION and SUGESTIONS (SONUÇ ve ÖNERİLER)

As cyberspace becomes increasingly central to modern life and national security, the concept of deterrence, traditionally rooted in kinetic warfare, must be reinterpreted for the digital age. Cyberspace's characteristics, including anonymity, low cost of entry, and difficulty of attribution, challenge the classical assumptions of deterrence theory. While traditional deterrence relies heavily on the threat of retaliation and visible capabilities, cyber deterrence must also incorporate legal, normative, and institutional mechanisms.

This article has shown that legal frameworks play a critical role in cyber deterrence, particularly when retaliatory action is unfeasible or ineffective. Türkiye's legislative evolution demonstrates an early recognition of this reality, with successive reforms aimed at addressing cyber threats through criminal law. From the initial amendments in 1991 to the more structured provisions of the Turkish Penal Code (TCK) and the country's accession to the Budapest Convention, Türkiye has laid a legal foundation to define, punish, and thus deter cyber offenses.

However, the existence of legal norms alone does not ensure deterrence. The efficacy of deterrent laws

depends on their consistent enforcement, the severity and proportionality of penalties, and the elimination of loopholes that undermine punishment, such as the overuse of HAGB and adequate remorse provisions. These practices, though well-intended, often reduce the dissuasive power of the law in the cyber realm, where certainty and swiftness of justice are critical.

Türkiye's experience highlights the need for integrated deterrence strategies that combine legal frameworks with technological capabilities and international collaboration in the face of rapidly evolving cyber threats. Cybersecurity cannot rely solely on reactive measures; it must be supported by proactive legal systems that deter malicious actors before they strike. As the digital domain continues to expand, the challenge for all states will be to ensure that their laws are robust on paper and effective in practice.

To enhance the deterrent effect of legal frameworks, particular attention must be paid to:

*Strengthening Enforcement Mechanisms:* Beyond merely having laws, the consistent and timely application of these laws is paramount. This requires well-resourced law enforcement agencies, judiciaries with specialized knowledge in cybercrime, and efficient judicial processes to ensure the swiftness of punishment.

*Re-evaluating Sentencing Practices:* Practices like the Postponement of the Announcement of the Verdict (HAGB) and effective remorse reductions, while aiming for rehabilitation, inadvertently diminish the perceived certainty and severity of punishment for cyber offenses. A critical review of these mechanisms is necessary to ensure they do not undermine the deterrent impact of the law, especially for crimes that can have far-reaching national security and economic consequences.

*Fostering International Cooperation:* Given the borderless nature of cyberspace, no single nation can effectively combat cyber threats in isolation. Türkiye's experience, particularly with online gambling and other cross-border offenses, highlights the indispensable need for robust international agreements, intelligence sharing, and collaborative law enforcement efforts to identify, apprehend, and prosecute perpetrators operating beyond national jurisdictions. Harmonization of legal standards, as seen with the Budapest Convention, remains crucial, though efforts must continue to bring key global players into consensus.

*Developing Dynamic Legal Frameworks:* The rapid evolution of technology, including advancements in AI and deepfake technologies, means that legal frameworks must be agile and adaptable. Laws should

be periodically reviewed and updated to address emerging cyber threats, ensuring that new forms of malicious activity are clearly defined and subject to appropriate legal sanctions.

*Integrating Legal Deterrence with Broader Cybersecurity Strategies:* Legal measures are just one pillar of a comprehensive cybersecurity strategy. They must be seamlessly integrated with technological defense capabilities, offensive measures for deterrence by punishment, and public awareness campaigns to foster a culture of cybersecurity. The goal is to create a layered defense where legal consequences, technological resilience, and international partnerships collectively raise the cost and risk for malicious actors, thereby creating a more formidable deterrent in the digital realm.

Ultimately, the effectiveness of cyber deterrence will not only depend on the strength of a nation's digital defenses or its capacity for retaliation but increasingly on its ability to build and enforce a credible legal infrastructure that resonates both domestically and internationally, fostering accountability and predictability in the inherently unpredictable domain of cyberspace.

## REFERENCES (KAYNAKLAR)

- [1] Baylis, J. (2008). Uluslararası İlişkilerde Güvenlik Kavramı. *Uluslararası İlişkiler*, 5(18): 69-85.
- [2] Booth, K. (2007). *Theory of World Security*. University Press. <https://doi.org/10.1017/CBO9780511840210>
- [3] Nye, J. S. (2011a). Nuclear Lessons for Cyber Security? *Strategic Studies Quarterly*, 5(4), 18-38.
- [4] Mazarr, M. J. (2018). Understanding Deterrence. *Rand Corporation Perspective*. [https://www.rand.org/content/dam/rand/pubs/perspectives/PE200/PE295/RAND\\_PE295.pdf](https://www.rand.org/content/dam/rand/pubs/perspectives/PE200/PE295/RAND_PE295.pdf) [Accessed: April 23, 2025]
- [5] Libicki, M. C. (2009). *Cyberdeterrence and cyberwar*. The Rand Corporation.
- [6] Nye, J. S. (2010). Cyber Power. *Belfer Center for Science and International Affairs*. <https://apps.dtic.mil/sti/pdfs/ADA522626.pdf> [Accessed: April 20, 2025]
- [7] CISA (2022). *Federal Information Security Modernization Act*. <https://www.cisa.gov/topics/cyber-threats-and-advisories/federal-information-security-modernization-act> [Accessed: April 25, 2025]
- [8] Resmi Gazete (2013). 2818 sayılı Siber Olaylara Müdahale Ekiplerinin Kuruluş, Görev ve Çalışmalarına Dair Usul ve Esaslar Hakkında Tebliğ. <https://www.resmigazete.gov.tr/eskiler/2013/11/20131111-6.htm> [Accessed: April 15, 2024]
- [9] UAB (2020). 2020-2023 Ulusal Siber Güvenlik Stratejisi ve Eylem Planı. <https://hgm.uab.gov.tr/uploads/pages/siber-guvenlik/ulusal-siber-guvenlik-stratejisi-ep-2020-2023.pdf> [Accessed: July 16, 2023]
- [10] Boulos, S. (2017). Cyberspace: Risks and Benefits for Society, Security and Development, Ramírez, J. M., & García-Segura, L. A. (Eds.). *The Tallinn Manual and Jus as bellu: Some Critical Notes*, 231-242, Springer International Publishing. <https://doi.org/10.1007/978-3-319-54975-0>
- [11] The Guardian (2007). "Russia accused of unleashing cyberwar to disable Estonia". <https://www.theguardian.com/world/2007/may/17/to-pstories3.russia> [Accessed: March 25, 2025]
- [12] Healy, J. & Jordan, K.T. (2014). *Nato's Cyber Capabilities: Yesterday, Today, and Tomorrow*, Atlantic Council. [https://www.atlanticcouncil.org/wp-content/uploads/2014/08/NATOs\\_Cyber\\_Capabilities.pdf](https://www.atlanticcouncil.org/wp-content/uploads/2014/08/NATOs_Cyber_Capabilities.pdf) [Accessed: March 20, 2025]
- [13] Singer, P. W. & Friedman, A. (2014). *Cybersecurity and cyberwar: What everyone needs to know*. Oxford University Press.
- [14] Holloway, M. (2015). *Stuxnet Worm Attack on Iranian Nuclear Facilities*. Stanford University. <http://large.stanford.edu/courses/2015/ph241/holloway1/> [Accessed: April 18, 2025]
- [15] Orhan, U. (2021). Cezaları Ağırlaştırmak Caydırıcılığı Artırır mı? (Does Aggravating Penalties Increase Deterrence?). *Türkiye Barolar Birliği Dergisi*, 34(56), 65-85. <http://tbbdergisi.barobirlik.org.tr/m2021-156-1997> [Accessed: January 3, 2023]
- [16] Dolu, O., Büker, H.& Uludağ, Ş. (2012). *Türk Ceza Adalet Sisteminin Caydırıcılık Kapasitesine İlişkin Eleştirel Bir Değerlendirme*, Ankara Üniversitesi Hukuk Fakültesi Dergisi, 61(1), 69-106. [https://doi.org/10.1501/Hukfak\\_0000001651](https://doi.org/10.1501/Hukfak_0000001651)

- [17] United Nations Charter (1945). <https://www.un.org/en/about-us/un-charter/full-text> [Accessed: March 18, 2025]
- [18] Schmitt, M. N., Vihul, L., & NATO Cooperative Cyber Defence Centre of Excellence (Ed.). (2017). *Tallinn manual 2.0 on the international law applicable to cyber operations* (First published 2017). Cambridge University Press.
- [19] Tikk, E., Kaska, K. & Vihul, L. (2010). International Cyber Incidents Legal Considerations. NATO Cooperative Cyber Defence Centre of Excellence [https://ccdcoc.org/uploads/2018/10/legalconsiderations\\_0.pdf](https://ccdcoc.org/uploads/2018/10/legalconsiderations_0.pdf) [Accessed: March 15, 2025]
- [20] Moynihan, H. (2020). The vital role of international law in the framework for responsible state behaviour in cyberspace. *Journal of Cyber Policy*, 6(3), 394-410. <https://doi.org/10.1080/23738871.2020.1832550> [Accessed: April 18, 2025]
- [21] Nye, J.S. (2017). Deterrence and Dissuasion in Cyberspace. *International Security*, Vol. 41, No. 3 (Winter 2016/17). [https://www.belfercenter.org/sites/default/files/pantheon\\_files/files/publication/isec\\_a\\_00266.pdf](https://www.belfercenter.org/sites/default/files/pantheon_files/files/publication/isec_a_00266.pdf) [Accessed: April 4, 2025]
- [22] International Committee of the Red Cross (ICRC). (2015). *International humanitarian law and the challenges of contemporary armed conflicts*. <https://www.icrc.org/en/publication/international-humanitarian-law-and-challenges-contemporary-armed-conflicts-building> [Accessed: April 18, 2025]
- [23] Council of Europe. (2001). *Convention on Cybercrime* (ETS No. 185). <https://www.coe.int/en/web/conventions/full-list/-/conventions/treaty/185> [Accessed: April 23, 2025]
- [24] Carr, E.H. (2001). *The Twenty Years' Crisis: An Introduction to the Study of International Relations*, New York: Palgrave.
- [25] United Nations. (2015). *Report of the Group of Governmental Experts on Developments in the Field of Information and Telecommunications in the Context of International Security* (A/70/174). <https://digitallibrary.un.org/record/799853?ln=en&v=pdf> [Accessed: April 29, 2025]
- [26] United Nations. (2021). *Report of the Open-Ended Working Group on Developments in the Field of Information and Telecommunications in the Context of International Security* (A/75/816). <https://documents.un.org/doc/undoc/gen/n21/068/72/pdf/n2106872.pdf> [Accessed: April 29, 2025]
- [27] NATO. (2019). *Cyber defence*. [https://www.nato.int/cps/en/natohq/topics\\_78170.htm](https://www.nato.int/cps/en/natohq/topics_78170.htm) [Accessed: April 29, 2025]
- [28] European Commission. (2020). *Cyber Diplomacy Toolbox*. <https://www.consilium.europa.eu/en/policies/cyber-diplomacy/> [Accessed: April 29, 2025]
- [29] BTK (2017). *Siber Güvenlik Kurulu*. <https://www.btk.gov.tr/siber-guvenlik-kurulu> [Accessed: June 13, 2023]
- [30] 5651 Kanun (2004). <https://www.mevzuat.gov.tr/mevzuatmetin/1.5.5651.pdf> [Accessed: June 13, 2023]
- [31] Resmi Gazete (2008). *5809 sayılı Elektronik Haberleşme Kanunu*. <https://www.resmigazete.gov.tr/eskiler/2008/11/20081110M1-3.htm> [Accessed: May 12, 2023]
- [32] Resmi Gazete (2016). *6698 Sayılı Kişisel Verilerin Korunması Kanunu*. <https://www.resmigazete.gov.tr/eskiler/2016/04/20160407-8.pdf> [Accessed: April 29, 2023]
- [33] Resmi Gazete (1991). *3756 sayılı 765 sayılı Türk Ceza Kanunu'nun Bazı Maddelerinin Değiştirilmesine Dair Kanun*. <https://www.resmigazete.gov.tr/arsiv/20901.pdf> [Accessed: Feb 2, 2023]
- [34] 5237 Sayılı Türk Ceza Kanunu (2004). <https://www.mevzuat.gov.tr/MevzuatMetin/1.5.5237.pdf> [Accessed: July 12, 2023]
- [35] *3756 sayılı 765 sayılı Türk Ceza Kanunu'nun Bazı Maddelerinin Değiştirilmesine Dair Kanun*. <https://www.resmigazete.gov.tr/arsiv/20901.pdf> [Accessed: Feb 2, 2023]
- [36] Türkiye Cumhuriyeti Anayasası (1982) <https://www.mevzuat.gov.tr/mevzuat?MevzuatNo=2709&MevzuatTur=1&MevzuatTertip=5> [Accessed: Feb 2, 2023]
- [37] EGM (2019). *Siber Suç Nedir?*. <https://www.egm.gov.tr/siber/sibersucnedir> [Accessed: April 2, 2023]
- [38] Council of Europe Budapest Convention (2004). *The Budapest Convention (ETS No.185) and its Protocols*.

<https://www.coe.int/en/web/cybercrime/the-budapest-convention> [Accessed: June 21, 2023]

[39] 5271 Sayılı Ceza Muhakemesi Kanunu (2004).  
<https://www.mevzuat.gov.tr/mevzuatmetin/1.5.5271.pdf> [Accessed: June 27, 2023]