



ISSN: 1309-1581

AJIT-e

*Academic Journal of
Information Technology*

Volume 16 • Issue 3 • Summer 2025

187 – 206

Teorik Makale
Theoretical Article

DOI: 10.5824/ajite.2025.03.001.x

Algorithmic Nostalgia: How Machine Learning Curates Our Emotional Connections to the Past

Bahtiyar Ahu Alpaslan, Ecemnur Delibalta

207 –231

Araştırma Makalesi
Research Article

DOI: 10.5824/ajite.2025.03.002.x

Advanced Mobile Money Fraud Detection Using CNN-BiLSTM and Optimized SGD with Momentum

Niamatu Yussif, Kate Takyi, Rose-Mary Owusuaa Mensah Gyening, Samuelson Israel Boadu-Acheampong

232 –250

Araştırma Makalesi
Research Article

DOI: 10.5824/ajite.2025.03.003.x

EmotionUnet: Konuşma Duygu Tanıma için U-Net Tabanlı Özgün Derin Öğrenme Modeli

Yasin Görmez

251 – 271

Araştırma Makalesi
Research Article

DOI: 10.5824/ajite.2025.03.004.x

Türkiye'nin Bilgi Bilim Alanında Yapay Zeka Araştırmaları: Web of Science İndeksli Yayınların Bibliyometrik Analizi (1994-2024)

Okan Koç

Supported by

ABA

Akademik Bilişim Araştırmaları
Derneği

ISSN: 1309-1581

AJIT-e

*Academic Journal of
Information Technology*

Volume • 16
Cilt

Issue • 3
Sayı

Summer • 2025
Yaz

Owner - Editor-in-Chief

Sahibi - Bař Editör

Prof. Dr. Özhan TINGÖY

Marmara Üniversitesi, İletişim Fakültesi, Gazetecilik Bölümü, Bilişim Ana Bilim Dalı, İstanbul, Turkey

Editörler

Editors

Prof. Dr. ŞEVKİ IŞIKLI

*Marmara Üniversitesi, İletişim Fakültesi, Gazetecilik
Bölümü, Bilişim Ana Bilim Dalı
İstanbul, Turkey*

Dr. Öğr. Üyesi Alaattin ASLAN

*Marmara Üniversitesi, İletişim Fakültesi, Gazetecilik
Bölümü, Bilişim Ana Bilim Dalı
İstanbul, Turkey*

Field Editors

Alan Editörleri

FEN BİLİMLERİ SCIENCE

Prof. Dr. NAZMİ EKREN

*Marmara Üniversitesi
Teknoloji Fakültesi
Elektrik-Elektronik Mühendisliği
Devreler ve Sistemler Anabilim
Dalı
İstanbul, Turkey*

Prof. Dr. Faik Nüzhet OKTAR

*Marmara Üniversitesi
Mühendislik Fakültesi
Biyomühendislik Bölümü
Biyomühendislik Anabilim Dalı
İstanbul, Turkey*

Doç. Dr. Nilüfer YURTAY

*Sakarya Üniversitesi
Bilgisayar ve Bilişim Bilimleri
Fakültesi
Bilgisayar Mühendisliği
Bölümü
Bilgisayar Mühendisliği
Anabilim Dalı
Sakarya, Turkey*

Doç. Dr. Rıdvan ŞAHİN

*Gümüşhane Üniversitesi
Mühendislik ve Doğa Bilimleri
Fakültesi
Matematik Mühendisliği Bölümü
Uygulamalı Mekanik Anabilim
Dalı
Gümüşhane, Turkey*

Doç. Dr. Oğuzhan GÜNDÜZ

*Marmara Üniversitesi
Teknoloji Fakültesi
Metalurji ve Malzeme
Mühendisliği Bölümü
Seramik Anabilim Dalı
İstanbul, Turkey*

Doç. Dr. Yusuf ALİYEV

*Azerbaijan State Pedagogical
University
Azerbaijancan, Turkey*

SOSYAL BİLİMLER SOCIAL SCIENCES

Doç. Dr. İhsan KARLI

*Kocaeli Üniversitesi
İletişim Fakültesi
Gazetecilik Bölümü
Genel Gazetecilik Ana Bilim Dalı
Kocaeli, Turkey*

Dr. Öğr. Üyesi Ali ÖZCAN

*Gümüşhane Üniversitesi
İletişim Fakültesi
Gazetecilik Bölümü
Bilişim Enformasyon
Teknolojileri Ana Bilim Dalı
Gümüşhane, Turkey*

Dr. Öğr. Üyesi Alihan Limoncuoğlu

*İstanbul Aydın Üniversitesi
İktisadi ve İdari Bilimler
Fakültesi
Siyaset Bilimi ve Uluslararası
İlişkiler Bölümü
Siyaset Bilimi ve Uluslararası
İlişkiler Pr.
İstanbul, Turkey*

Dr. Öğr. Üyesi Münevver SOYAK

*Marmara Üniversitesi
Sosyal Bilimler Meslek
Yüksekokulu
Dış Ticaret Bölümü
Dış Ticaret Pr.
İstanbul, Turkey*

Dr. Öğr. Üyesi Yusuf BUDAK

*Kocaeli Üniversitesi
İletişim Fakültesi
Gazetecilik Bölümü
Bilişim (Bilgisayar Teknikleri ve
İletişim) Ana Bilim Dalı
Kocaeli, Turkey*

Foreign Language Editor <i>Yabancı Dil Editörü</i>		
Doç. Dr. Gulshan AGABAY <i>Azerbaijan State Pedagogical University Azerbaijancan, Turkey</i>	Doç. Dr. Süheyla Nil MUSTAFA <i>Marmara Üniversitesi İletişim Fakültesi Gazetecilik Bölümü Yayınçılık Yönetimi Anabilim Dalı İstanbul, Turkey</i>	Doç. Dr. Serkan BAYRAKÇI <i>Marmara Üniversitesi İletişim Fakültesi Gazetecilik Bölümü Bilişim Anabilim Dalı İstanbul, Turkey</i>

Redaktör <i>Redactor</i>	
Fırat Mert DEMİR (M.A.) <i>Marmara Üniversitesi Doktora Öğrencisi İstanbul, Turkey</i>	Hilal GÖKBAYRAK (M.A.) <i>Marmara Üniversitesi Doktora Öğrencisi İstanbul, Turkey</i>

Editorial Secretariat <i>Editöryal Sekreteryası</i>
Mustafa ÇOKYAŞAR (M.A.) <i>Marmara Üniversitesi editor@ajit-e.org İstanbul, Turkey</i>

Editorial Board <i>Yayın Kurulu</i>		
Prof. Dr. Rauf Nurettin NİŞEL <i>Piri Reis Üniversitesi Mühendislik Fakültesi Endüstri Mühendisliği Bölümü Endüstri Mühendisliği Pr. İstanbul, Turkey</i>	Prof. Dr. Halil İbrahim GÜRCAN <i>Anadolu Üniversitesi/İletişim Bilimleri Fakültesi Basın ve Yayın Bölümü Basın Yayın Tekniği Ana Bilim Dalı Eskişehir, Turkey</i>	Prof. Dr. Murat ÖZGEN <i>İstanbul Üniversitesi İletişim Fakültesi Gazetecilik Bölümü Genel Gazetecilik Ana Bilim Dalı İstanbul, Turkey</i>
Prof. Dr. Oya KALIPSIZ <i>Yıldız Teknik Üniversitesi Elektrik-Elektronik Fakültesi Bilgisayar Mühendisliği Bölümü Bilgisayar Yazılımı Ana Bilim Dalı İstanbul, Turkey</i>	Prof. Dr. Özhan TINGÖY <i>Marmara Üniversitesi İletişim Fakültesi Gazetecilik Bölümü Bilişim Ana Bilim Dalı İstanbul, Turkey</i>	Prof. Dr. Derman KÜÇÜKALTAN <i>İzmir Kavram Meslek Yüksekokulu Otel Lokanta ve İkram Hizmetleri Bölümü Aşçılık Pr. İzmir, Turkey</i>
Prof. Dr. Yavuz AKPINAR <i>Boğaziçi Üniversitesi Eğitim Fakültesi Bilgisayar ve Öğretim Teknolojileri Eğitimi Bölümü Bilgisayar ve Öğretim Teknolojileri Eğitimi Ana Bilim Dalı İstanbul, Turkey</i>	Prof. Dr. Süleyman ÖZDEMİR <i>İstanbul Esenyurt Üniversitesi İstanbul, Turkey</i>	Prof. Dr. Ahmet KALENDER <i>Selçuk Üniversitesi İletişim Fakültesi Halkla İlişkiler ve Tanıtım Bölümü Halkla İlişkiler Ana Bilim Dalı Konya, Turkey</i>
Prof. Dr. Özgür ÇENGEL <i>İstanbul Arel Üniversitesi İktisadi ve İdari Bilimler Fakültesi İşletme Bölümü İşletme Pr. İstanbul, Turkey</i>	Doç. Dr. İhsan KARLI <i>Kocaeli Üniversitesi İletişim Fakültesi Gazetecilik Bölümü Genel Gazetecilik Ana Bilim Dalı Kocaeli, Turkey</i>	Prof. Dr. ŞEVKİ IŞIKLI <i>Marmara Üniversitesi İletişim Fakültesi Gazetecilik Bölümü Bilişim Ana Bilim Dalı İstanbul, Turkey</i>

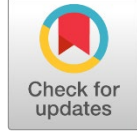
Doç. Dr. Fatime Neşe KAPLAN İLHAN Marmara Üniversitesi İletişim Fakültesi Radyo, Televizyon ve Sinema Bölümü Sinema Anabilim Dalı İstanbul, Turkey	Dr. Öğr. Üyesi Ali Barış KAPLAN İstanbul Üniversitesi İletişim Fakültesi Radyo, Televizyon ve Sinema Bölümü Radyo-Televizyon Anabilim Dalı İstanbul, Turkey	Dr. Öğr. Üyesi Yusuf BUDAK Kocaeli Üniversitesi İletişim Fakültesi Gazetecilik Bölümü Bilişim (Bilgisayar Teknikleri ve İletişim) Ana Bilim Dalı Kocaeli, Turkey
Dr. Öğr. Üyesi Ali ÖZCAN Gümüşhane Üniversitesi İletişim Fakültesi Gazetecilik Bölümü Bilişim Enformasyon Teknolojileri Ana Bilim Dalı Gümüşhane, Turkey	Dr. Öğr. Üyesi Ahmet ÖZTÜRK Manisa Celâl Bayar Üniversitesi Gördes Meslek Yüksekokulu Pazarlama ve Dış Ticaret Bölümü Halkla İlişkiler ve Tanıtım Pr. Manisa, Turkey	

International Board of Overseers Uluslararası Danışma Kurulu		
Prof. Amjad Hadjikhani Uppsala University, Department of Business Studies Sweden	Prof. Angappa Gunasekaran California State University, School of Business and Public Administration USA	Prof. David Benyon Edinburgh Napier University, School of Computing Scotland
Prof. David Gunkel Northern Illinois University, Department of Communication USA	Prof. Ian Ruthven University of Strathclyde, Computer and Information Sciences Scotland	Prof. Maria Manuela Cruz da Cunha Escola Superior de Tecnologia – IPCA Portugal
Prof. Sayed Abdul Muneem Pasha Jamia Millia Islamia, Department of Political Science India	Prof. Thomas Bauer University of Münster, Islamic and Arab Studies Germany	Prof. Umit Sezer Bititci Heriot-Watt University, Edinburgh Business School Scotland
Assoc. Prof. Anvarjon Ahmedov Ahatjonovich Universiti Malaysia Pahang Faculty of Industrial Sciences & Technology Malaysia	Assoc. Prof. Awang Dharmawan Universitas Negeri Surabaya Indonesia	Assoc. Prof. Berk Kucukaltan University of Bradford United Kingdom
Assoc. Prof. Lala Jabbarova Baku State University Azerbaijan	Assoc. Prof. Romel Aceron Batangas State University Philippines	Assoc. Prof. Vafa Isgandarova Baku State University Azerbaijan
Ayşe Goker, (Ph.D.) AmbieSense (Co-founder, Director) United Kingdom	Charalambos Tsekeris, (Ph.D.) National Centre for Social Research Greece	David Fernández Quijada, (Ph.D.) European Broadcasting Union Switzerland
Ismet Anitsal, (Ph.D.) Missouri State University USA	Sidharth Verma, (Ph.D.) - India	Tim Marsh (Ph.D.) Griffith University, Griffith Film School Australia

Dergide yayınlanan makalelerde belirtilen görüşler ve fikirler sadece yazar(lar)ın görüşüdür. Yayınlanan içeriklerle ilgili bütün sorumluluklar yazar(lar)a aittir. Yayınlanan eserlerde yer alan tüm içerik kaynak gösterilmeden kullanılamaz.



The opinions and ideas stated in the articles published in the journal are only the opinion of the author (s). All responsibilities regarding the published content belong to the author (s). The published contents in the articles cannot be used without being cited.



Tüm makaleler DOI ve Crossmark ile kayıt altına alınmaktadır.



All articles are registered with DOI and Crossmark.



AJIT-e has an Open Access policy and is licensed under the [Creative Commons Attribution-Same License Share 4.0 International License](#). Access to published articles is free.



Tarandığımız İndeksler – Indexes



© 2010 – 2025

AJIT-e – Academic Journal of Information Technology

Address: Kazım Özalp Sk. No: 15 Kat: 2 34740 Şaşkınbakkal / Suadiye / KADIKÖY / ISTANBUL / TURKEY

Tel: +90 216 355 56 19

Faks: +90 216 368 43 30

Email: editor@ajit-e.org

Supported by

ABA

Akademik Bilişim Araştırmaları
Derneği

www.ajit-e.org



www.abilar.org

Yeni iletişim ortamları hız ve yayın süreçleri açısından yazılı basına göre çok daha avantajlı olduğundan, akademik yayıncılığın geleceği, Internet gibi yeni iletişim ortamları etrafında şekillenmeye başlamıştır. Makaleler dergilerin basılı versiyonlarından önce yayınlanabilmektedir. AJIT-e de iletişim ve bilişim alanına ilgi duyan araştırmalar için bir kaynak ve yayın ortamı sağlamak amacıyla 2010 yılında yayın hayatına başlamıştır.

AJIT-e, uluslararası hakemli bir dergidir. Türkçe ve İngilizce, iki dilde yılda dört sayı yayınlanır. AJIT-e yayın alanları arasında başlıca şu konular yer alır:

Yeni Medya ve İletişim Bilimleri, Teknoloji, Adli Bilişim, Belge ve Kayıt Yönetimi, Bilgi Güvenliği, Bilgi Yönetimi, Bilişim Etiği, Bilişim Hukuku, Dağıtık Bilişim Sistemleri, E-Öğrenme, E-Dönüşüm, E-Devlet, E-Pazarlama, E-Reklam, E-Scm, E-Yayıncılık, E-Yayın, E-Yönetim, Tıp Bilişimi, Karar Destek Sistemleri, Sayısal Eğlence ve Oyun, Sayısal Hak Yönetimi, Sosyal Ağlar, Tedarik Zinciri Yönetimi, Telekomünikasyon, Veri Madenciliği, Veritabanları, Yapay Zekâ, Yönetim Bilişim Sistemleri



As new communication environments are much more advantageous than print media in terms of speed and broadcast processes, the future of academic publishing has begun to take shape around new communication environments such as the Internet. Articles can be published long before the printed versions of journal. AJIT-e started publication in 2010 to provide a resource and publication environment for research interested in the field of communication and informatics.

AJIT-e is an international refereed journal. It is published four times a year in both languages, in Turkish and English. AJIT-e publication areas include the following topics:

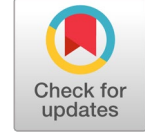
New Media and Communication Sciences, Technology, Computer Forensics, Document and Records Management, Information Security, Information Management, Information Ethics, Distributed Information Systems, E-Learning, E-Transformation, E-Government, E-Marketing, E-Advertising, E-Scm, E-Publishing, E-Management, Medical Informatics, Decision Support Systems, Digital Entertainment and Gaming, Digital Rights Management, Social Networks, Supply Chain Management, Telecommunications, Data Mining, Databases, Artificial Intelligence, Management information systems

Prof. Dr. Özhan TINGÖY
Editor-in-Chief

Contents

İçindekiler

187 – 206 Teorik Makale Theoretical Article	DOI: 10.5824/ajite.2025.03.001.x Algorithmic Nostalgia: How Machine Learning Curates Our Emotional Connections to the Past <i>Bahtiyar Ahu Alpaslan, Ecemnur Delibalta</i>
207 – 231 Araştırma Makalesi Research Article	DOI: 10.5824/ajite.2025.03.002.x Advanced Mobile Money Fraud Detection Using CNN-BiLSTM and Optimized SGD with Momentum <i>Niamatu Yussif, Kate Takyi, Rose-Mary Owusuaa Mensah Gyening, Samuelson Israel Boadu-Acheampong</i>
232 – 250 Araştırma Makalesi Research Article	DOI: 10.5824/ajite.2025.03.003.x EmotionUnet: Konuşma Duygu Tanıma için U-Net Tabanlı Özgün Derin Öğrenme Modeli <i>Yasin Görmez</i>
251 – 271 Araştırma Makalesi Research Article	DOI: 10.5824/ajite.2025.03.004.x Türkiye'nin Bilgi Bilim Alanında Yapay Zeka Araştırmaları: Web of Science İndeksli Yayınların Bibliyometrik Analizi (1994-2024) <i>Okan Koç</i>



Algorithmic Nostalgia: How Machine Learning Curates Our Emotional Connections to the Past

Bahtiyar Ahu Alpaslan, Yeditepe University, Faculty of Communication, Journalism, Asst. Prof.,
aerdogdu@yeditepe.edu.tr, 0000-0002-3296-3050

Ecemnur Delibalta, Yeditepe University, Faculty of Communication, Media and Communication
Management, M.A., ecemnurdelibalta@gmail.com, 0009-0002-7506-0422

ABSTRACT

In an era increasingly shaped by algorithmically curated media environments, nostalgia has become a powerful emotional tool within digital platforms. This study investigates how algorithmic recommendation systems on Instagram and Twitter facilitate the emergence of nostalgia-themed content and influence users' emotional engagement. The analysis explores how brand-generated throwback posts, strategically crafted to evoke collective memory and sentiment, intersect with user interaction patterns and algorithmic visibility mechanisms.

Using a qualitative digital ethnography approach, this research is based on non-intrusive observation of selected public brand accounts on Instagram and Twitter. Content using hashtags such as #TBT and #ThrowbackThursday was examined to understand how digital nostalgia is amplified through algorithmic circulation. The findings reveal that algorithm-driven promotion of nostalgic content not only enhances users' emotional connectivity and sense of community but also generates concerns about data visibility and user autonomy. Some users were observed to later remove their engagements, suggesting a potential discomfort with the visibility or implications of their digital interactions.

The study highlights that digital nostalgia—while fostering emotional continuity—can also reinforce the complexities of personalization and data ethics in platform economies. By focusing on the interplay between emotional experience and algorithmic mediation, the study contributes to a deeper understanding of how affective engagement and concerns over data privacy co-evolve in contemporary digital culture.

Keywords : Algorithmic Recommendations, Digital Nostalgia, Machine Learning, Media, Communication

Algoritmik Nostalji: Makine Öğreniminin Geçmişle Duygusal Bağlantılarımızı Şekillendirme Yöntemleri

ÖZ

Algoritmik olarak şekillenen medya ortamlarının giderek etkisini artırdığı günümüzde, nostalji dijital platformlar içinde güçlü bir duygusal araç hâline gelmiştir. Bu çalışma, Instagram ve Twitter'daki algoritmik öneri sistemlerinin nostalji temalı içeriklerin oluşumunu nasıl kolaylaştırdığını ve kullanıcıların duygusal katılımını nasıl etkilediğini incelemektedir. Marka hesapları tarafından paylaşılan geçmişe dönük içeriklerin, kolektif belleği harekete geçirecek biçimde nasıl kurgulandığı ve algoritmalar aracılığıyla görünürlük kazandığı analiz edilmiştir.



Çalışma, nitel dijital etnografi yöntemiyle yürütülmüş; Instagram ve Twitter’da seçilen halka açık marka hesapları üzerinde müdahalesiz gözlem yapılmıştır. Özellikle #TBT ve #ThrowbackThursday etiketleriyle paylaşılan içerikler incelenerek dijital nostaljinin algoritmalar yoluyla nasıl yeniden üretildiği değerlendirilmiştir. Bulgular, algoritmaların nostalji temalı içerikleri öne çıkararak kullanıcıların duygusal bağlarını ve topluluk hissini güçlendirdiğini, ancak aynı zamanda veri görünürlüğü ve kullanıcı özerkliği konularında endişeleri de artırdığını ortaya koymaktadır. Bazı kullanıcıların içeriklerle kurdukları etkileşimleri sonradan kaldırdıkları gözlemlenmiş, bu durum dijital etkileşimlerinin yarattığı görünürlük ya da sonuçlara dair bir rahatsızlık hissine işaret etmektedir.

Sonuç olarak, dijital nostaljinin duygusal sürekliliği teşvik ettiği kadar, kişiselleştirme ve veri etiği meselelerini de pekiştirdiği görülmektedir. Bu çalışma, duygusal deneyimle algoritmik arabuluculuk arasındaki ilişkiyi ele alarak, dijital kültürde duygusal bağlılık ve veri mahremiyetine dair endişelerin birlikte nasıl şekillendiğine dair derinlemesine bir anlayış sunmayı amaçlamaktadır.

Anahtar Kelimeler : Algoritmik Öneriler, Dijital Nostalji, Makine Öğrenmesi, Medya, İletişim

INTRODUCTION

In today’s algorithmically mediated media landscape, nostalgia has become a strategic and emotional device embedded within digital content flows. Rather than existing merely as a sentimental longing for the past, digital nostalgia now plays an active role in how users engage with social platforms, reinforcing emotional ties and personal narratives through curated experiences. Content that recalls past decades, cultural moments, or personal memories is not only widely consumed but is often algorithmically elevated due to its high engagement potential.

This study focuses on the intersection of nostalgic effect, algorithmic visibility, and user interaction within Instagram and Twitter. Both platforms have established mechanisms such as memory prompts or throwback hashtags that routinely bring past content back into circulation. Through machine-learning-driven personalization, these platforms do not simply surface content from the past but actively frame and filter what kind of memories are worth re-experiencing.

Building on Sedikides and Wildschut’s (2018) conceptualization of nostalgia as a self-regulatory psychological resource, and Van Dijck’s (2013) critique of datafication, this study approaches digital nostalgia as both a personalized emotional affordance and a socio-technical strategy. It explores how recommender systems contribute to memory-making in digital contexts, and how emotional resonance and data privacy anxieties co-evolve within these processes.

By observing public brand accounts on Instagram and Twitter, this research investigates how algorithmically circulated nostalgic content not only elicits emotional responses but also reveals users’ shifting awareness of data control and content visibility. In

doing so, it contributes to broader discussions on affective infrastructures, platform logics, and the ethics of emotional engagement in digital media.

1. CONCEPTUAL FRAMEWORK

This study draws on existing theories of nostalgia, affect, datafication, and algorithmic culture to explain how digital platforms shape emotional engagement through curated memory experiences. In particular, it situates digital nostalgia as both a psychological process and a socio-technical construct, where algorithms do not merely reflect user preferences but actively shape memory consumption through personalization.

From a psychological perspective, nostalgia is understood as a self-regulatory emotional resource that fosters continuity, belonging, and emotional resilience (Sedikides & Wildschut, 2018). This form of affective recall becomes especially significant when it is mediated by algorithmic systems that personalize content to resonate with users' emotional histories. Digital platforms such as Instagram and Twitter operationalize nostalgia through repeated exposure to past-themed content, often under hashtags like #TBT or "Throwback Thursday," which serve as structured prompts for emotional recall.

Sociologically, nostalgia is not only a personal sentiment but also a shared cultural phenomenon. As Davis (1979) notes, nostalgia often functions collectively, reactivating shared identities through symbols and narratives. In algorithmically managed platforms, these collective moments are encoded and re-surfaced through engagement metrics and recommender systems. Van Dijck's (2013) framework of datafication helps to explain how user memories are transformed into data points—engagement histories, visual preferences, and interaction patterns—that then feed back into algorithmic logics.

Furthermore, recent studies on platform governance and recommender systems show that nostalgic content is not neutral. Leong et al. (2019) and Milligan (2020) argue that nostalgia cues are strategically employed by platforms to boost emotional retention and extend user attention. This strategic use of nostalgia raises ethical concerns regarding data agency, transparency, and the emotional manipulation of users.

Thus, this study conceptualizes digital nostalgia as an algorithmically modulated emotional experience, shaped by user data, curated by platform infrastructures, and circulated through collective engagement. The framework integrates perspectives from communication studies, digital media theory, and emotional psychology to better understand how digital nostalgia both sustains emotional connection and invites critical reflection on algorithmic influence and data ethics.

2. METHODOLOGY

This study adopted a qualitative digital ethnography approach to explore how algorithmic recommendation systems on social media platforms contribute to the formation and circulation of nostalgia-oriented content. This method was chosen for its ability to deeply contextualize and interpret user-platform interactions within algorithmically shaped digital environments (Pink et al., 2016).

The research focused on Instagram and Twitter, two high-traffic platforms that actively use algorithmic recommendation engines and regularly surface memory-based prompts. A purposive sampling strategy was applied to identify brand-generated content that reflected nostalgic themes, particularly posts using hashtags such as #TBT and #ThrowbackThursday.

In addition to content sampling, two public brand accounts were selected for close observation. These accounts were chosen based on their engagement with retro-themed content, either by posting historically framed brand imagery or by participating in memory-based interaction patterns.

Data were collected through non-intrusive observation of publicly accessible brand accounts. During this process, nostalgic content, interaction metrics, and algorithmic prompts were monitored and documented. The observed data categories included:

- Captions and titles referencing specific timeframes
- User engagement behaviors (likes, shares, comments, mentions)
- Platform-generated algorithmic signals (e.g., resurfaced content or “Suggested for you” sections)

These observations highlighted how nostalgia-driven content is shaped both by brand strategy and algorithmic visibility. Rather than emerging purely from user activity, nostalgic narratives were found to be amplified through the platforms’ curation systems, reinforcing emotional and temporal continuity within user experiences.

Data analysis followed an interpretive communication research framework, using iterative thematic coding to identify three core analytical dimensions:

- Emotional Resonance

Nostalgic posts frequently used time-based references and cultural imagery to elicit emotional responses, such as longing, joy, or reflection.

- Interactive Participation

Users actively engaged with nostalgic content by liking, commenting, or sharing, especially during recurring digital rituals such as #TBT. These actions contributed to the co-construction of communal digital memory.

- Privacy Awareness and Content Removal

Although no direct interviews or comments were recorded, it was observed that some users removed their prior interactions such as comments or tags on nostalgic brand content. This behavior, while not explicitly explained, may reflect concerns around data visibility, personal exposure, or platform surveillance.

To ensure analytical reliability, two independent reviewers conducted parallel coding of the observational data. Through iterative discussion and cross-checking, consensus was reached on all thematic classifications, enhancing the credibility of the findings.

3. WHAT IS DIGITAL NOSTALGIA?

Where a new generation is concerned, the term "digital nostalgia" quite literally describes the resurfacing of memories, images, or events from the past through technology and digital media. Digital platforms are able to reconnect people more emotionally to their past because they are sources of content that remind users of their past. It may be some old photographs, old catchy songs, or anything else that comes under the definition of cultural elements associated with any definite era-in that respect, the nostalgic trend is a tool to facilitate remembering individually and collectively.

Boym (2001, p. 16) comments, While defining nostalgia as that sense of an individual's longing for the past, in the digital age, this gets reconfigured, especially through social media and content platforms. Grainge (2002, p. 54) refers to "the capacity of digital media to refigure nostalgia, to mobilize user feeling about the past. Digital nostalgia is a way for people to attach themselves to the past by means of access and the sharing of certain periods on those very platforms.

While digital nostalgia enables people to reattach themselves to their past, the very fact that this contact is made through digital platforms also reveals the other side of nostalgia-that one related to the effects on consumer behavior. Niemeyer (2014, p. 83-94) discusses "digital nostalgia in practice within marketing activities and mentions that it is efficient in enhancing emotional and social bonds".

4. WHAT IS ALGORITHMIC PERSONALIZATION?

It basically refers to the process through which digital platforms and applications use user information to deliver tailored content or services that match the needs, preferences, and

behaviors of every individual user. In the context of technologies such as big data, artificial intelligence, and machine learning, this topic has become increasingly relevant. Algorithmic personalization refers to a technique to personalize streams of content or recommendation systems by considering the interests, habits, and characteristics of users (Smith, 2019, p. 45 & Chen, 2020, p. 112)

These systems monitor the interaction made by users to predict the user preferences in order to view, read, or to buy something. For instance, social network sites rank what the user has engaged with in anticipation of what he or she might find most relevant to their interests. More precisely, the process is algorithmic, and it personalizes user experiences of content according to Gillespie (2014, p. 67).

The key goal of algorithmic personalization is increasing users' satisfaction and exposure to digital platforms. Recommendation algorithms of different content to individual users fuel other major digital platforms, like Instagram, Twitter, YouTube or Netflix. Quite a great deal of work is done with regard to how these recommendation systems truly work (Van Dijck, 2013, p. 112, Pariser, 2011, p. 45). Precisely, as the Filter Bubble theory explained by Eli Pariser states, algorithmic personalization will expose the users to specific content only that may foster particular biases and narrow-mindedness in acquiring information (ibid, 2011, p. 45).

4.1. Algorithmic Recommender Systems and User Experience

The algorithmic recommendation systems relate to various sophisticated analytics applied in the creation of personalized digital experiences where the user is targeted with customized content. These systems compare users' past behavior and interaction and identify which content may be of interest to the user. Netflix and Spotify classify and suggest content to their users according to the history of viewed or listened items. Such recommendations not only increase the time spent by users on the platform but also enhance users' satisfaction and loyalty (Davidson, 2010, p. 213, Gómez-Urbe & Hunt, 2016, p. 47).

Especially, content recommendations regarding the past drive the nostalgic feelings in the user. This is where algorithms trigger positive emotional responses by recommending content related to popular songs during one's childhood, old series, or even any key event. While doing so, it personalizes them with memories of the past rather well, which increases the dwell time spent on the platform even more. It learns about users' preference through algorithmic recommendations but also shows responses to their emotional needs, as expressed in Kaplan & Haenlein, 2019, p. 17, Chaney, Stewart & Engelhardt, 2018, p. 245).

Nostalgic Content and Algorithmic Recommendations Social networks, digital media, and video streaming services come up with special days for content creation as a way of

retaining users on the platforms. The algorithms on those platforms study the tendencies of users that have content that provides a retro flair and then promote content related to music, movies, or cultural events that were really popular during the user's childhood years. For example, the social media platforms develop the nostalgic-themed posts, like #TBT (The "Throwback Thursday" (TBT) trend gained popularity on **Instagram** around **2011**, where users began sharing nostalgic content using the hashtag #TBT. As the trend evolved, **social media algorithms actively promoted throwback content**, reinforcing nostalgia-driven engagement loops (Leong et al., 2019; Milligan, 2020). This phenomenon exemplifies how **digital nostalgia is systematically engineered**, shaping user engagement through **AI-driven personalization** (van Dijck, 2013).) and contents from specific historic periods that gain the attention of the users so as to enhance user interest in the nostalgic contents. The content improves the user experience on sites through the attainment of emotional satisfaction, as indicated by Rosenbaum & Wong, (2015, p. 2346), Fairclough (2018, p. 30).

To better understand how algorithmic recommender systems operate and their role in curating nostalgic content, it is essential to explore the mechanisms behind different recommendation techniques used in digital platforms.

Algorithmic recommender systems operate through various techniques that detect user behavior to provide personalized content. These techniques include:

Content-Based Filtering: This approach suggests content based on a user's past interactions. It relies on analyzing similarities between consumed and recommended content (Ricci et al., 2015).

Collaborative Filtering: This method predicts user preferences by analyzing the behavior of similar users. It is widely used in streaming services like Netflix and Spotify to enhance engagement through nostalgia-driven recommendations (Gómez-Urbe & Hunt, 2016).

Hybrid Recommendation Systems: Combining both content-based and collaborative filtering, these models aim to improve recommendation accuracy (Davidson et al., 2010).

Deep Learning-Based Models: Advanced AI-driven algorithms that refine recommendations by continuously learning from user interactions (Zhang & Chen, 2020).

The way algorithms deliver nostalgic content is based on analyzing user interaction data and behavioral patterns over time.

Recommendation systems use multiple data points to determine which nostalgic content should be shown to a user. These data points include:

- **User Interaction Data:** Algorithms track watch history, likes, and engagement to recommend past content (Milligan, 2020).
- **Demographic Insights:** Platforms shows user age and generational preferences to provide relevant nostalgic material (Sedikides & Wildschut, 2018).
- **Temporal Trends:** Some platforms prioritize nostalgic content during key events, such as anniversaries or seasonal periods (Grainge, 2002).

To illustrate the recommendation process, the following flowchart outlines how digital platforms generate personalized nostalgic content suggestions.

Step 1: User engagement data is collected (watch history, likes, shares).

Step 2: The algorithm processes data and categorizes user preferences.

Step 3: Personalized nostalgic content is recommended based on analysis.

Step 4: User interactions update the recommendation system for future accuracy.

However, the video streaming sites have also managed to develop various ways of offering nostalgic content to their users. For instance, Netflix tries to revive the viewer's interests in the past to include their childhood cartoons and television series by categorizing them under "Classics." The mere highlighting of such stuff results in users automatically spending more time on the platform and creating some emotional bond with it. More specifically, on Instagram and other social networks, such engagers make this content even catchier with their instant sharings that call others to take part either virtually or face-to-face. More precisely, #TBT calls its consumers to engage in an act of mutual participation through refreshing memories of times gone by, thus being closely and emotionally connected with users sheltered under a common hashtag.

In this respect, nostalgic content represents one of the most efficient means by which platforms can satiate users' emotional needs and increase loyalty towards them, according to Sedikides, Wildschut, & Baden (2004, p. 308) and Belk (1990, p. 670). By the content, users are able to feel connected with their past while digital platforms can create a more holistic user experience.

4.2. Algorithmic Personalization and Privacy Concerns

In the present information age, the chance to receive and the range of information is developing more and more. It interests them in how information is treated and then again

represented to the user in personalized forms. While algorithmic personalization definitely offers a lot of advantages from the point of view of digital platforms and applications, it has brought on varied concerns about data privacy on the part of users. This regular monitoring and analysis of users' behaviors raise an ethical debate at the very core of privacy violation of individuals. This is a problem that many works focused their attention on: the security of data collected in the process of personalization and the threat of its sharing with third parties. Among the many works are the works of Zuboff (2019, p. 45), Bucher (2018, p. 89).

In this respect, algorithmic personalization came to the fore as a system that, while improving the user experience in the digital world, was also threatening the privacy of users.

4.3. Data Privacy and Trust

While the potential of nostalgic content recommendations in building emotive connections on digital platforms, it raises a number of issues regarding data privacy and trust. In reality, personalized recommendations on digital platforms are based on various data regarding past user interactions, analyses of past content consumption, and interests. However, users' interest in such personalized content recommendations also raise questions about the extent and manner of their data usage (Zuboff, 2019, p. 78, Acquisti, 2004, p. 45). They are scared that, as consumers in such retro-memoir content delivery, their private information might be violated during processing.

Accordingly, following the famous hierarchy of needs by Maslow, basic needs for survival and participation in society, which are claimed by social scientists today, embrace psychological safety, one of the important needs to self-actualization, as creating a successful performance in all spheres of life, including personal, educational, and professional life.

Psychological safety describes the conditions under which people feel able to take the interpersonal risks of expressing their opinions, concerns, and ideas, and providing feedback-all without fear of negative consequences or needing to soften difficult truths. It is the culture that fosters the generation of ideas from all individuals without the fear of judgment or confrontation. In this kind of environment, for example, people can express their opinions against ideas, including those of the management, where they feel things are not working right and should be improved. It also encourages individuals to confess to their errors, express vulnerability, and address authority candidly with regard (McKinsey, November 2024).

Users' data privacy concerns against the tide of nostalgic content usually fall under three major points:

- 1.Data Processing of Personal Information and Ethical Issues: While processing data in order to provide personalized content, algorithmic recommendation systems have to monitor users' behaviour constantly. Data gathered in this process gives not just the history of individuals' viewing but also their emotional tendencies. Particular

recommendations based on nostalgic content demonstrate users based on their emotional tendencies, which again raises ethical issues with regard to users' privacy. That raises concerns that personal data cannot only be used for commercial purposes but also for emotional manipulation (Nissenbaum, 2010, p. 78, Tufekci, 2014, p. 56).

2. Lack of Data Privacy and Algorithmic Transparency: Most users do not have enough information about how algorithms work and what data is collected. This lack of transparency undermines user trust, especially in user data-driven services such as social media platforms. Without transparency about the decision-making processes of algorithms, users have difficulty understanding the potential risks to their personal data. This increases data privacy concerns (Pasquale, 2015, p. 112).

3. Data Security and Third Party Sharing: The digital platforms can sell or share the collected user data with third parties. Due to the uprising of popular nostalgic content its related data requirements are growing and so do the concerned questions on the security of user data, access to this data by third party companies opens up possibilities of misuse of the credentials of users and increasing risks of data security (Acquisti 2004, p. 35).

While such nostalgic content might come at the price of personalization of the user's experience, the privacy and security of the data on which such a process relies surely raises very important ethical issues. More precisely, uncertainty over the purposes for which and by whom users' personal data will be stored and used decreases trust in digital platforms. In this context, concerns about the privacy of data obtained through nostalgic content should be discussed more in the field of digital media ethics and legal regulations to protect users' rights should be considered (Zuboff, 2019, p. 102, Nissenbaum, 2010, p. 79).

4.4. Techniques of Emotional Bonding

Storytelling is among the most ancient strategies that humans have ever developed for the purpose of communication. Beyond the development of organizational culture, stories can enable people to connect with each other, create a unique perspective, and rally together on common goals (Fisher 1984, p. 8, Guber, 2011, p. 43). Digital platforms offer the user the opportunity to tell their story to their audience with every single piece of content created. While opening the doors to a new world of communication for the following audiences in this digital world that is going to be created with this content, it starts to create an emotional bond with the audience who follows the systematic content sharing over time. The bond that then emerged allowed the platforms to invent new influencers with a developing audience of their own, and for influencers to act as a bridge between them and the brands themselves, creating an organic system through which the recommendation mechanisms result in sales. Matching user data obtained from here with the right creators and followers is the algorithmic

mechanism which digital platforms are working on most intensively. The identification of users' past posts, interests in nostalgic contents, and creating an opportunity for them to make new stories is one of the newest and trendiest ways of creating emotional bonding with the users via digital platforms. The most obvious example of the user decisions taken with the use of the emotional brain can be found within the developing circle based on reliability.

In this situation , it is also useful to mention nostalgia and personalization. This, combined with personalized content, makes the longing for the past with nostalgic content even more effective. Algorithms used on digital platforms offer customized nostalgic content based on the individual's previous preferences and past interests. It is this content that evokes the past memories of the user hereby creating personalized experiences, and thus stronger emotional bonds, as suggested by Holbrook 1993, p. 252; Sedikides et al., 2008, p. 306. Algorithms can shows users' digital history and flash up, for them, the most pop music, photos, or content from a specific era; this makes individuals' ties to the past more tangible, deepening their experiences (Tufekci, 2014, p. 33).

"Nostalgic feeling has the ability to revive people's memories and fill them with positive feelings. More precisely, finding the content in an emotional response within an individual's memories has made users more interested in social media platforms and spending more time on interacting with these platforms as stated by Davis (1979, p. 422)". With this context, the algorithm-provided nostalgic content keeps the sense of being "special" and leads to emotional satisfaction in the digital environment.

4.5. Interaction and Engagement

Nostalgia-themed content functions as a powerful driver of user engagement on social media platforms. As Boym (2001) emphasizes, nostalgic narratives allow individuals to re-establish emotional connections with the past, satisfying a desire for continuity, identity, and emotional grounding. Social media platforms, particularly Instagram and Twitter, harness this affective appeal through algorithmic mechanisms that promote emotionally resonant content, thereby sustaining user attention and interaction.

Research suggests that nostalgic posts generate higher engagement rates—measured by likes, comments, and shares—than non-nostalgic content (Bolton et al., 2013, p. 249). This engagement is amplified when users participate in platform-specific rituals such as Instagram's "Throwback Thursday" (#TBT) or nostalgic hashtag chains on Twitter. These practices encourage users to revisit and share their own past, while also engaging with others' memories through responses, retweets, and quote tweets, transforming personal nostalgia into a social and interactive experience.

In both platforms, nostalgia operates as a feedback loop. On Instagram, the visual nature of the platform facilitates emotional immersion through filtered aesthetics and dated imagery. On Twitter, textual memory fragments often accompanied by hashtags or cultural

references allow for real-time interaction and discursive participation. In both cases, engagement is not only emotional but algorithmically rewarded—visible in increased visibility, trending tags, and prolonged screen time.

This interplay between user interaction and algorithmic amplification illustrates how engagement is not merely reactive, but shaped through strategic platform design. Nostalgic content thereby becomes a vector for emotional expression, community-building, and data generation—demonstrating how affective and algorithmic logics are deeply intertwined in digital participation.

4.6. Data Privacy and Ethical Debates

Today, it is pivotal that technological advancements and the problems and issues arising from those very opportunities have come under debate. In this realm, one of the major issues is related to the security of personal information regarding every individual on the face of this earth. Websites, internet search engines, and social media applications have been used widely in modern digital life and periodically capture their users' personal data. Such situations may bring great risks with respect to data security. These services also enable users to be better informed on matters of personal data security, one's 'digital footprint', principles of ethics in the digital world, the advantages and disadvantages of using digital platforms, and of their contribution to making the digital ecosystem more secure and sustainable.

“Although personal data is a concept that has existed throughout human history, it has gained more importance today, especially with the developments in the fields of law and technology (Köksalan, 2022, p. 15)”. The first official document on the concept of personal data was published by the Organization for Economic Cooperation and Development (OECD) on 23 September 1980 (OECD, 1980, p. 12). It gives the guiding principle on how personal data breaches can be avoided, including how security of information flow could be regulated.

Examples are personally identifiable information such as email addresses, social media, and online shopping activities; clinical information captured and retained within health systems; and information on financial transactions. While Article 6 of Law No. 6698 on the Protection of Personal Data defines personal data as "data relating to race, ethnic origin, political opinion, philosophical belief, religion, sect or other beliefs, dress, membership of associations, foundations or trade unions, health, sexual life, criminal convictions and security measures, and biometric and genetic data" (Official Gazette, 2016). Therefore, this definition has led to the legal construction that any information which identifies a person will be considered as personal data.

Personal data security is significantly vulnerable, especially in the digital context. Security breaches in data are allowed when some malicious people or hackers get

unauthorized access to systems and reach personal data. For example, the outcome of a cyber-attack on any database may result in theft, disclosure, or alteration of users' personal information. Commonly, these types of breaches have brought into danger very important information like financial information, health records, or social security.

Another severe threat to data security in a digital environment is malware. These malicious types of software can infiltrate users' computers and access their personal data for malicious processing. This type of software operates very often in forms that are not even noticed and realized by users. It creates dangers regarding the tracking, stealing, and manipulation of data.

Another type of threat is phishing attacks. With this attack method, "malicious people aim to obtain personal information by misleading users with fake websites, emails or messages (Scott, 2019, p. 38, Tchakounté et al., 2020, p. 95)". For example, users may be directed to a fake website and tricked into entering their bank account details. Such attacks can lead to the disclosure of personal data and exposure to fraud.

Finally, the other type of threat for personal data security is social engineering. In this attack method, "there would be plans to access personal data by manipulating individuals. Methods involving phone scams, on-site information requests, or misdirection capture people's information (Akça, 2016, p. 59)". The attackers would try to do this through psychological manipulations, distorted tales, and ways of gaining the trust of an individual.

With the rise in such threats, awareness of personal data security becomes an urgent need so that precautionary measures can be taken for security in the digital environment. The use of user data is also of particular importance in the context of data privacy and ethical debates.

Algorithms indicate users' personal data to create nostalgic content and emotional connection. In this process, algorithms collect and betray a wide range of data sets such as the user's browsing history, likes, shares and even location information. This data is used to provide customized content recommendations to the user, aiming to create a deeper connection with the user (Zuboff, 2015, p. 197). Yet, processing of such personal data by algorithms for making recommendations is susceptible to cross the boundaries of data privacy. More importantly, offering retro/nostalgic content specifically entails collecting and storing data on the user's pasts, perceived by many as an act against digital privacy per se (van Dijck 2014, p. 75). Adding to this, the fact that the users are not informed about what data is being collected and for what purpose enhances these ethical concerns (Tufekci, 2014, p. 153).

Algorithms demonstrate users' personal data to create nostalgic content and emotional connection. In this process, algorithms collect and indicate a wide range of data sets such as the user's browsing history, likes, shares and even location information. This data is used to provide customized content recommendations to the user, aiming to create a deeper

connection with the user (Zuboff, 2015, p. 197). Yet, processing of such personal data by algorithms for making recommendations is susceptible to cross the boundaries of data privacy. More importantly, offering retro/nostalgic content specifically entails collecting and storing data on the user's pasts, perceived by many as an act against digital privacy per se (van Dijck 2014, p. 75). Adding to this, the fact that the users are not informed about what data is being collected and for what purpose enhances these ethical concerns (Tufekci, 2014, p. 153). User trust is at the forefront of these ethical concerns.

Whether or not users are aware of how their data is being processed while being connected to nostalgic content has a significant negative impact on user trust. "Social media and digital platforms may not inform users about this process when collecting personal data through algorithms. This can undermine users' trust in data privacy and undermine their trust in the platforms" (Acquisti, Brandimarte, & Loewenstein, 2015, p. 510). It follows from studies that in the case of users perceiving a violation of their privacy, their engagement and loyalty for the digital platforms goes down. "The lack of privacy in the digital environment makes the user feel less "safe" and more reluctant to disclose data (Nissenbaum, 2010, p. 77)". In this respect, "ethical issues concerning the protection of privacy influence users' confidence in the platform (Solove, 2008, p. 121)".

5. DISCUSSION AND PARTICIPATION

Nostalgic content has proved a most effective instrument in the creation of user engagement on social media platforms. Boym (2001) says nostalgic actions enable the user to re-activate affectively dense moments in the face of a requirement for individual and collective memory (p. 49). Algorithms on social media platform websites have actively incorporated the affective proclivity by algorithmically showcasing content in a way designed to bring about emotional resonance and thus increase user engagement.

Experiments verify the perception that nostalgically involving material will more willingly be passed on, endorsed, annotated and hence algorithmically appealing (Bolton et al., 2013, p. 249). The "Throwback Thursday" custom on Instagram is the ideal exemplar of the dynamic: it turns individual reminiscence into a formalised weekly interactivity ritual. With the reanimation of archived content and the user's inclusion in a group construct of time, the platform enhances interaction levels as well as rebrand itself as a digital repository for reminiscences.

This repeated interaction enables lasting emotional linkage among platform users and platforms. Referencing Zhang et al. (2019), emotional content reinforces content creation through the validation of the presence of users and digital reminiscence (p. 4). Such algorithmic growth hence aids platform development and individual expression.

Most critically, nostalgic engagement is crafted and extended as opposed to organic or neutral. Platforms utilize affect triggers in the quest towards extending user activity, platform utilization, and user data production upon which the recommendations are channeled. Such recursive logic undermines the co-construction of user engagement by algorithmic constructs and affect practices and in the process raises critical questions regarding autonomy, manipulation, and emotional labor operative in mundane digital practices.

5.1. Case Studies

To gain deeper insight into how nostalgia-themed content is surfaced and engaged with on digital platforms, a focused case study approach was applied to selected examples from the observed data. The purpose was to illustrate how algorithms and user interactions together shape nostalgic experiences online. Below are two representative cases, one from each platform.

Instagram and Twitter: Hashtag-Driven Nostalgia Loops

In order to consider how content by nostalgia circulates in and into attention on algorithmic social media sites, a comparative case research design was conducted. Two model cases were discovered in both Twitter and Instagram, based on a clear invoking on the broadly felt “Throwback Thursday” (#TBT) custom. The cases both involve two different companies beverages (Pepsi; reach date: 06.07.2025) and logistics (Maersk; reach date: 06.07.2025) turning the attention towards an industry-crossover consideration on the manner in which nostalgic imagery is employed in a conscious way towards audience engagement.

The data was gathered through a manual search for publicly released platform feeds and hashtag archives from 2021 through 2023. All publicly released posts from brand accounts officially authorized on the platform were included in accordance with ethics rules for gathering digital data. The lone selective filter was overt discussion of nostalgia in the content, adoption of the #TBT hashtag, and clear indicators of algorithmic amplification (e.g., presence in “Explore” modules, large engagement numbers, or indicators of virality).

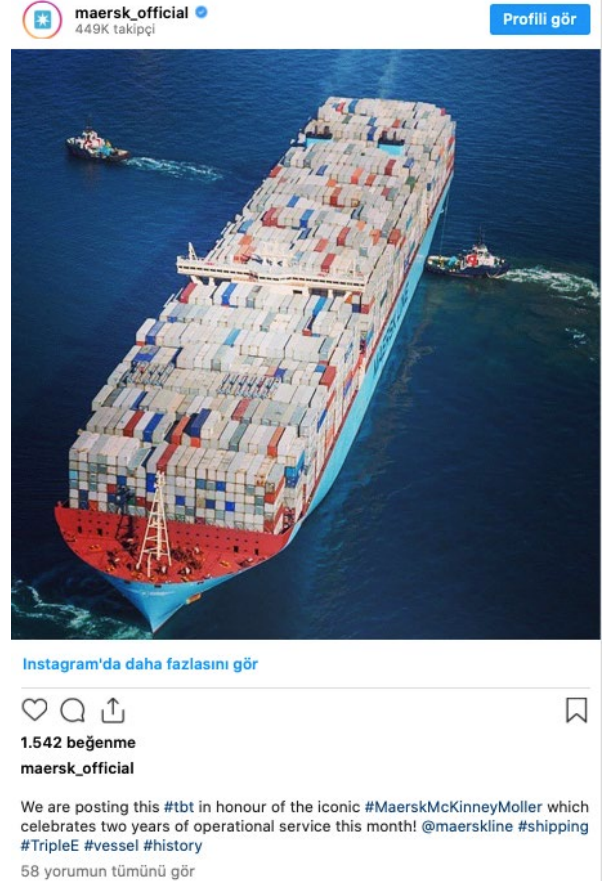
The analysis was undertaken through close reading of the visual and written content in each platform, platform-specific engagement indicators (shares, likes, comments), and the contextual placement of the content within the broader trends in hashtags. From this vantage

point, the research aims at revealing the manner in which algorithmic logics and user practices both construct and solidify digital loops of nostalgia on the various platforms.

Figure 1: Pepsi Official Twitter Account



Figure 2: Maersk Official Instagram Account



5.2. Findings

The analysis of the user engagement and tweet analysis on Twitter and Instagram shows the multifaceted interaction among participatory dynamics, algorithmic mediation, and emotional engagement. On social media websites, the activity of nostalgia occurs in more multifaceted manners than simple reminiscence but as a dynamic and active process where the user responds affectively, socially engages, and critically negotiates the curating on content filtered by the platform. The results discussed in the next sections greatly elaborate the related dynamics.

Emotional Engagement

Nostalgic content will necessarily yield robust affect reactions by re-activating both individual memories as well as public-cultural histories in users. Affective resonance will continually act on users to like, write about the content, or re-share the content by re-infusing

in the affectual flow of the platform the affect of the past. The affective dimension is especially dominant in the posts re-iteratively capturing the eye on the past through retro packaging stands, aged photograph filter compositions, or universal temporal referents like “Throwback Thursday.”

Participatory Interaction Participants engage in the nostalgia content through dozens of various interaction manners commenting, tagging acquaintances, making the same hashtags, and re-sharing on the stories or the feeds. Such activities help constitute the digital group and the reciprocal recognition in the process making the nostalgia from the solitary sentiment the common social practice. Memory caters the public performance through the participatory dynamics co-authored by the users and stimulated through the engagement measures by the platforms.

Algorithmic Knowledge and Ethic Issues

The increasing degree of user consciousness towards the practices of platform data and algorithmic presentations of content adds a crucial layer to the nostalgic effect. While there is empathetic response on the emotional level towards user curation of content, there is more discrimination on the level on which the algorithms themselves generate memory-driven content. That note brings forth questions on the level of manipulation, surveillance, and commodification of affect. Particularly the repetition in the popular nostalgic hashtags such as #TBT betrays the layer in which emotional memory is transmitted but systematized and commodified by algorithmic means.

CONCLUSIONS AND RECOMMENDATIONS

This research has covered digital nostalgia from the perspective of a sociotechnical phenomenon in the domain of algorithmic suggestion systems and emotional interaction between consumers and memory-driven content. The findings from the research indicate digital nostalgia on the web transcends a sentiment staring back and a data-driven and mediated process by the platform's algorithmic logic. More significantly, recommender systems by themselves are busy creating memory-driven content groupings based on consumers' prior interactions and hence hold platform attachment and emotional attachment.

While customized nostalgic material according to the specific users' digital footprints allows communion with the past through algorithms, the personalized treatment itself creates essential ethical unease. When increasingly many recognize how much material is collected from them, operated on by them, and exploited in order to generate the content world within which they participate, fear about privacy, observation, and unreadability by the algorithm intensifies. The ostensibly smooth interweaving between affective reminiscence and data-driven personalized material creates the two-fold effect the users are both critically disturbed and emotively engaged.

This tension reveals there is rising need for added algorithmic transparency in the area of how content invested with memory is aggregated and distributed. Operators should offer clear explanations regarding the procedures by which the personal data are being utilized in distribution of content in manners making the user an active participant in digital culture and not a passive recipient. More algorithmic transparency can facilitate confidence building, prevent misinformation, and alleviate ethical risk by way of affect manipulation.

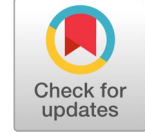
In addition, there needs to be more research into how digital nostalgia converges with wider psychological and cultural dynamics especially within the contexts of identification construction, consumerism, and digital well-being. There is also a pressing requirement for the establishment of regulatory and ethical guidelines protecting user data in a way that does not compromise the cultural and affective roles performed by nostalgic content. A lasting digital media environment will bring together affective engagement and lasting ethical responsibility. Platforms that recognize the affective aspect of algorithmic curation and making commitments towards data openness and user agency will more likely gain credibility, lasting usage, and positive digital interactions.

REFERENCES

- Acquisti, A. (2004). *Privacy in electronic commerce and the economics of immediate gratification*. Proceedings of the 5th ACM Conference on Electronic Commerce, 21–29. <https://doi.org/10.1145/988772.988773>
- Acquisti, A., Brandimarte, L., & Loewenstein, G. (2015). Privacy and human behavior in the age of information. *Science*, 347(6221), 509–514. <https://doi.org/10.1126/science.aaa1465>
- Akça, F. (2016). Siber güvenlikte insan faktörü ve sosyal mühendislik. *Journal of Security Studies*, 2(1), 58–65.
- Belk, R. W. (1990). The role of possessions in constructing and maintaining a sense of past. *Advances in Consumer Research*, 17, 669–676.
- Bolton, R. N., Parasuraman, A., Hoefnagels, A., Migchels, N., Kabadayi, S., Gruber, T., ... & Solnet, D. (2013). Understanding Generation Y and their use of social media: A review and research agenda. *Journal of Service Management*, 24(3), 245–267. <https://doi.org/10.1108/09564231311326987>
- Boym, S. (2001). *The future of nostalgia*. Basic Books.
- Bucher, T. (2018). *If... then: Algorithmic power and politics*. Oxford University Press.
- Chaney, A. J., Stewart, B. M., & Engelhardt, B. E. (2018). *How algorithmic confounding in recommendation systems increases homogeneity and decreases utility*. Proceedings of the 27th International Conference on World Wide Web, 239–248.
- Chen, X. (2020). *Personalization and algorithmic curation in the digital age*. Oxford University Press.
- Davidson, J., Liebald, B., Liu, J., Nandy, P., Van Vleet, T., Gargi, U., ... & Sampath, D. (2010). *The YouTube video recommendation system*. Proceedings of the Fourth ACM Conference on Recommender Systems, 293–296. <https://doi.org/10.1145/1864708.1864770>
- Davis, F. (1979). *Yearning for yesterday: A sociology of nostalgia*. Free Press.
- Fairclough, K. (2018). Classics in the streaming age: The formation and performance of cultural memory on Netflix. *International Journal of Cultural Studies*, 21(1), 25–41.
- Fisher, W. R. (1984). Narration as a human communication paradigm: The case of public moral argument. *Communication Monographs*, 51(1), 1–22.
- Gillespie, T. (2014). The politics of “platforms.” In T. Gillespie, P. J. Boczkowski, & K. A. Foot (Eds.), *Media technologies: Essays on communication, materiality, and society* (pp. 67–94). MIT Press.
- Gómez-Urbe, C. A., & Hunt, N. (2016). The Netflix recommender system: Algorithms, business value, and innovation. *ACM Transactions on Management Information Systems*, 6(4), Article 13. <https://doi.org/10.1145/2843948>
- Grainge, P. (2002). *Monochrome memories: Nostalgia and style in retro America*. Praeger.
- Guber, P. (2011). *Tell to win: Connect, persuade, and triumph with the hidden power of story*. Crown Business.

- Holbrook, M. B. (1993). Nostalgia and consumption preferences: Some emerging patterns of consumer tastes. *Journal of Consumer Research*, 20(2), 245–256.
- Kaplan, A. M., & Haenlein, M. (2019). Siri, Siri, in my hand: Who's the fairest in the land? On the interpretations, illustrations, and implications of artificial intelligence. *Business Horizons*, 62(1), 15–25. <https://doi.org/10.1016/j.bushor.2018.08.004>
- Köksalan, H. (2022). Dijital ortamda kişisel verilerin korunması ve etik sorunlar. *Academic Perspectives Journal*, 4(3), 15–28.
- Leong, C. W., Beigman Klebanov, B., Hamill, C., Stemle, E., Ubale, R., & Chen, X. (2020). A report on the 2020 VUA and TOEFL metaphor detection shared task. *Proceedings of the Second Workshop on Figurative Language Processing*, 18–29. <https://doi.org/10.18653/v1/2020.figlang-1.3>
- Maersk [@maerskline]. (2015, July 9). *We are posting this #tbt in honour of the iconic #MaerskMcKinneyMoller which celebrates two years of operational service this month!* [Instagram post]. Instagram. <https://www.instagram.com/p/5egnLkqexv/>
- McKinsey & Company. (2024, November). *What is psychological safety?* <https://www.mckinsey.com/featured-insights/mckinsey-explainers/what-is-psychological-safety>
- Milligan, C. (2020). The affective mechanics of nostalgic media. *New Media & Society*, 22(3), 387–405.
- Niemeyer, K. (Ed.). (2014). *Media and nostalgia: Yearning for the past, present and future*. Palgrave Macmillan.
- Nissenbaum, H. (2010). *Privacy in context: Technology, policy, and the integrity of social life*. Stanford University Press.
- OECD. (1980). *Guidelines on the protection of privacy and transborder flows of personal data*. OECD Publishing.
- Official Gazette. (2016). *Law No. 6698 on the Protection of Personal Data*. Türkiye Cumhuriyeti Resmî Gazete.
- Pariser, E. (2011). *The filter bubble: What the internet is hiding from you*. Penguin Press.
- Pasquale, F. (2015). *The black box society: The secret algorithms that control money and information*. Harvard University Press.
- Pepsi [@pepsi]. (2013, December 12). Showing some love on #ThrowbackThursday. <3 #TBT [Tweet]. X. Retrieved July 06, 2025, from <https://x.com/pepsi/status/411217908476682240>
- Ricci, F., Rokach, L., & Shapira, B. (Eds.). (2015). *Recommender systems handbook (2nd ed.)*. Springer. <https://doi.org/10.1007/978-1-4899-7637-6>

- Rosenbaum, M. S., & Wong, I. A. (2015). If you remember, I'll remember: Nostalgia, relationship maintenance, and Chinese consumers. *Journal of Business Research*, 68(11), 2346–2352.
- Scott, J. (2019). *Cybersecurity and data protection: Trends and challenges*. Cambridge Security Press.
- Sedikides, C., Wildschut, T., & Baden, D. (2004). Nostalgia: Past, present, and future. *Current Directions in Psychological Science*, 17(5), 304–307. <https://doi.org/10.1111/j.1467-8721.2008.00595.x>
- Smith, J. (2019). *Understanding algorithmic personalization: A critical approach*. Cambridge University Press.
- Solove, D. J. (2008). *Understanding privacy*. Harvard University Press.
- Tchakounté, F., Bellomo, D., & Gildas, L. (2020). Digital threat analysis in modern systems. *Cybersecurity Advances*, 7(3), 91–97.
- Tufekci, Z. (2014). Engineering the public: Big data, surveillance, and computational politics. *First Monday*, 19(7). <https://doi.org/10.5210/fm.v19i7.4901>
- Van Dijck, J. (2013). *The culture of connectivity: A critical history of social media*. Oxford University Press.
- Van Dijck, J. (2014). Datafication, dataism and dataveillance: Big data between scientific paradigm and ideology. *Surveillance & Society*, 12(2), 197–208. <https://doi.org/10.24908/ss.v12i2.4776>
- Zhang, C., Ko, M., & Carpenter, B. (2019). Exploring nostalgic engagement in social media: How memories and algorithms drive participation. *International Journal of Social Media Studies*, 6(2), 1–18.
- Zhang, S., & Chen, J. (2020). Explainable recommendation: A survey and new perspectives. *Foundations and Trends® in Information Retrieval*, 14(1), 1–101. <https://doi.org/10.1561/15000000066>
- Zuboff, S. (2015). Big other: Surveillance capitalism and the prospects of an information civilization. *Journal of Information Technology*, 30(1), 75–89. <https://doi.org/10.1057/jit.2015.5>
- Zuboff, S. (2019). *The age of surveillance capitalism: The fight for a human future at the new frontier of power*. PublicAffairs.



Advanced Mobile Money Fraud Detection Using CNN-BiLSTM and Optimized SGD with Momentum

Niamatu Yussif, Kwame Nkrumah University of Science and Technology, Department of Computer Science, yussifniamatu@gmail.com, 0009-0005-4637-8740

Kate Takyi, Kwame Nkrumah University of Science and Technology, Department of Computer Science, Ph.D., takyikate@knust.edu.gh, 0000-0002-2016-6169

Rose-Mary Owusuaa Mensah Gyening, Kwame Nkrumah University of Science and Technology, Department of Computer Science, Ph.D., rmo.mensah@knust.edu.gh, 0000-0002-8087-5207

Samuelson Israel Boadu-Acheampong, Kwame Nkrumah University of Science and Technology, Department of Computer Science, samuelsonacheampong@gmail.com, 0009-0003-9287-5989

ABSTRACT

The accelerated adoption of mobile money systems has significantly increased fraudulent activity, compromising their security and trustworthiness. This research presents an enhanced method for detecting mobile money fraud by modifying a CNN-BiLSTM model with momentum using Stochastic Gradient Descent (SGD). We computed salient features from transaction data using a pre-processed hybrid CNN-BiLSTM model and trained the model to identify trends in the data that included geographical and temporal aspects. The model performed remarkably using industry-standard testing approaches: an F1 score of 0.9928, precision of 0.9927, accuracy of 0.9928, and recall of 0.9929. The proposed model can identify dishonesty and has a low false positive rate. According to the study, the model improves feature selection and incorporates various optimization techniques, making it more flexible and suitable for different mobile money systems.

Keywords

: Fraud Detection, Machine Learning, Neural Networks, Stochastic Gradient Descent

CNN-BiLSTM ve Momentum ile Optimize Edilmiş SGD Kullanılarak Gelişmiş Mobil Para Dolandırıcılığı Tespiti

ÖZ

Mobil para sistemlerinin hızla benimsenmesi, dolandırıcılık faaliyetlerinde önemli bir artışa yol açarak güvenliklerini ve itibarlarını tehlikeye atmıştır. Bu araştırma, Stokastik Gradyan İnişi (SGD) kullanarak momentumla bir CNN-BiLSTM modelini değiştirerek mobil para dolandırıcılığını tespit etmek için gelişmiş bir yöntem sunmaktadır. Önceden işlenmiş bir hibrit CNN-BiLSTM modeli kullanarak işlem verilerinden belirgin özellikleri hesapladık ve modeli coğrafi ve zamansal yönleri içeren verilerdeki eğilimleri belirlemek üzere eğittik. Model, endüstri standardı test yaklaşımlarını kullanarak dikkat çekici bir performans gösterdi: 0.9928'lik bir F1 puanı, 0.9927'lik bir hassasiyet, 0.9928'lik bir doğruluk ve 0.9929'lık bir geri çağırma. Önerilen model, sahtekârlığı tespit edebilir ve düşük bir yanlış pozitif oranına



sahiptir. Çalışmaya göre, model özellik seçimini iyileştirir ve çeşitli optimizasyon tekniklerini birleştirerek onu daha esnek ve farklı mobil para sistemleri için uygun hale getirir.

Anahtar Kelimeler : Sahtecilik Tespiti, Makine Öğrenmesi, Sinir Ağları, Stokastik Gradyan İnişi

INTRODUCTION

Given the growing number of people using mobile banking, it becomes vital to identify and stop fraud. Mobile money, which provides convenience and access to banking services through mobile devices, has revolutionized financial transactions, especially in developing nations. The quick expansion of mobile financial services has, however, also resulted in a sharp rise in fraudulent activity, which puts users and service providers at great risk. Banks find it more difficult to spot fraud with more mobile money companies. Unauthorized transactions, identity theft, and social engineering attempts are common forms of fraud in mobile money networks. Due to the high amount of transactions, changing fraud trends, and the unbalanced structure of fraud datasets, where fraudulent cases are far less common than genuine ones, these threats are challenging to identify. Therefore, it is both technically and practically necessary to design reliable and effective fraud detection models. Complex fraud activities are frequently too difficult to detect with traditional rule-based systems. The challenge of traditional fraud detection systems in adjusting to different setups raises the probability of thievery. The study must use fresh approaches if fraud detection is to overcome these difficulties (El Kafhali & Tayebi, 2024: 1-25). This study presents a unique and more exact approach based on robust deep learning algorithms and momentum-optimized stochastic gradient descent (SGD) to detect mobile money fraud (Lokanan, 2023: 1-24). Combining the most potent SGD model with momentum, Convolutional Neural Network (CNN), and the Bidirectional Long Short-Term Memory (BiLSTM) model improves network performance and creates a fraud detection tool (Agarwal et al., 2021: 2552-2560). The two most successful machine learning methods are deep learning and efficient stochastic gradient descent (SGD with momentum). This study presents a unique method using a momentum-optimal SGD and CNN-BiLSTM architecture to identify mobile money offenders. Researchers inability to successfully fight mobile money theft is their neglect of modern approaches such as CNN-BiLSTM and upgraded SGD + Momentum algorithms (Igweni, 2023). CNNs and BiLSTM provide location information and study temporal correlations, thereby improving fraud detection's speed, accuracy, and data management. Using state-of-the-art deep learning methods, including CNN-BiLSTM structures and momentum SGD implementation, our study aims to enhance the training process and help identify mobile money scams. In machine learning, the SGD optimizer is one of the most straightforward and widely utilized optimization algorithms, especially for applications like financial fraud detection (Yang et al., 2023: 1-26). Our approach involves gathering a dataset and its documentation, standardizing the number of columns, and maintaining the category data to ensure every attribute is

regularly scaled (Botchey et al., 2021: 45-56). In this work, CNN-BiLSTM deep learning model together with the Decision Tree Classifier is used to find the most essential characteristics for fraud detection.

1. RELATED WORKS

With an emphasis on the application of Convolutional Neural Networks (CNN), Bidirectional Long Short-Term Memory (BiLSTM), and the Stochastic Gradient Descent (SGD) optimizer improved with momentum, this literature review examines recent developments in mobile money fraud detection. The capacity of these algorithms to identify temporal and spatial trends in transaction data has drawn a lot of attention, making them ideal for identifying complex fraudulent activities. This review outlines the fundamental approaches, points out gaps in the literature, and suggests fresh lines of inquiry for future research with the goal of enhancing detection accuracy and practical applicability through a methodical analysis of previous works.

The paper aimed to make a system that can spot mobile money scams using deep-learning tools. The RNN-based system did much better than the ANN and LSTM systems separately by a significant amount. Bayesian methods were used to fine-tune the RNN model (El Kafhali & Tayebi, 2024: 1-25).

(Botchey et al., 2021: 45-56) examines the issue of identifying fraudulent activity in mobile money transactions, focusing on developing nations with increased financial inclusion. The authors use synthetic minority oversampling and adaptive synthetic sampling, among other sample techniques, to raise the models' performance.

Discussed in the paper is a fresh approach (Zhang et al., 2020: 25210-25220) presented for better spotting fraud for customers who do not make regular transactions. The NM (Naive Bayes-based Multi-behavior) model, which integrates transaction status, individual user conduct, and group behaviour, significantly enhances detection accuracy. Despite occasional recall shortcomings, the NM model outperforms competing models in accuracy and F1 score. The model reduces the disturbance rate, preventing the incorrect classification of legitimate transactions as fraudulent. Further research aims to improve the model's internet-based updates.

Long short-term memory (LSTM) is employed in this study to create a deep learning model for finding bad financial habits. This model has constantly shown excellent results in forecast tasks over time. The model often does a better job than well-known machine learning methods at finding cases of financial wrongdoing in large datasets. (Alghofaili & Rassam, 2020: 498-516).

According to (Almazroi & Ayub, 2023: 137188-137203), The RXT-J model uses advanced machine learning algorithms to identify online payment fraud, overcoming

challenges like imbalanced data and extensive models. It uses SMOTE, Ensemble Autoencoder, and ResNet models for feature recovery and attribute extraction.

This project aims to build a network architecture to detect fraudulent online transactions, especially those involving Zimbabwean institutions. It will do this by studying expenditure trends hidden Markov model and deep neural network anomaly detection (Mbunge et al., 2015: 2319-7064).

Mobile money transfer (MMT) services, controlled by carriers, are growing globally. To combat money laundering, MMT systems should identify fake chains. A new method, Predictive Security Analysis at Runtime (PSA@R), outperforms conventional detection methods with 99.81% precision and 90.18% recall (Zhdanova et al., 2014: 11-20).

Researchers found that when there isn't enough training data, understanding what people are saying may help mobile e-commerce systems make better predictions. This includes help from experts, blacklists, and past dishonest behaviour. Data mining methods raise memory and accuracy rates by taking risk variables out of transaction data (Sun et al., 2021: 1-17).

This study will use the Kaggle IEEE-CIS dataset to test new deep-learning methods for finding fraudulent scams. Using deep neural networks and attention methods to find wrong classifications. It shows a CNN-Bi-LSTM-ATTENTION model about 95% of the time (Agarwal et al., 2021: 2552-2560).

The detection of mobile money fraud has been the subject of numerous studies, few have integrated CNN-BiLSTM architectures with momentum-optimized SGD. Although traditional approaches continue to capturing the complete intricacy of transaction patterns. Further study of hybrid deep learning models that can successfully handle the magnitude and dynamic nature of mobile money fraud is obviously needed.

2. METHODOLOGY

The aim of this research is to enhance the detection of mobile money fraud with the implementation of a momentum-based approach and CNN-BiLSTM. Data collection and project success assessment are two critical components. The integration of CNN-BiLSTM with momentum and SGD efficiency techniques can facilitate detection. The study uses an existing dataset to develop a deep learning system for fraud detection. This dataset provided a realistic basis for challenges involving fraud detection by simulating mobile money transfers. Following data conversion and standardization, we used resampling to balance the classes in order to address the significant disparity between fraud and non-fraud cases. The data was divided into training and testing sets. Feature importance was assessed using a Decision Tree to guide model design. We combined CNN and BiLSTM to create a hybrid deep learning

model that can identify both short-term and long-term trends. A model's F1 score, accuracy, recall, and ROC curve are frequently assessed by users to determine its ability to detect fraud with minimal false positives. Figure 1 shows the conceptual framework of the designed for the proposed work.

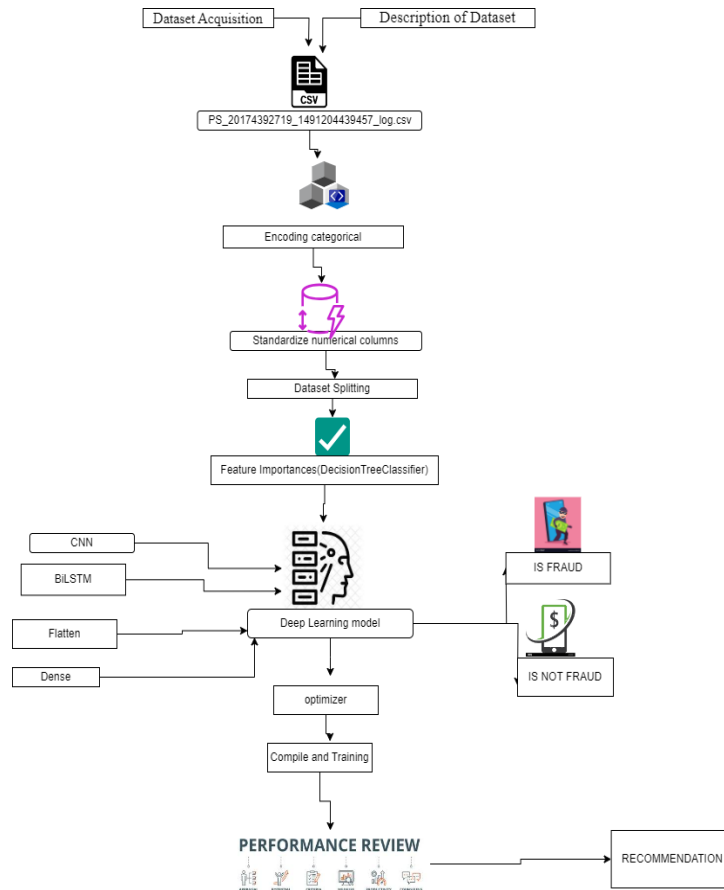


Figure 1: The proposed Conceptual Framework

2.1. Data Collection

To identify mobile money fraud, the CNN-BiLSTM model dataset was created using the Paysim simulator found on the Kaggle dataset named Synthetic Financial Datasets for Fraud Detection (Kaggle, 2023). The dataset includes columns with names such as sorts, quantities, old balance, new balance, fraud, etc. The dataset has 6,362,620 rows and 11 columns, with 6354407 labelled as not fraudulent transaction (0) and 8213 labelled as a fraudulent transaction (1). This dataset was used in the detection of mobile money fraud using the CNN-BiLSTM algorithm to improve security in mobile money transactions. It simulates mobile money transactions in a unique way and incorporates fraudulent activity. This dataset made it possible to create and assess sophisticated fraud detection models like CNN-BiLSTM by offering a realistic depiction of financial transactions. The extensive structure of the dataset, which included a variety of transaction kinds and difficult features, allowed for in-depth testing

and analysis of the suggested model. Utilizing this information, the study sought to develop fraud detection and security techniques for mobile banking services. The dataset were chosen because they offer a special combination of simulated fraudulent activities and real-world financial transactions, making them the perfect setting for testing and training sophisticated fraud detection models like CNN-BiLSTM.

2.2. Encoding Categorical

The Momentum SGD optimization and CNN-BiLSTM scam detection for mobile money were enhanced through categorical column encoding, requiring sci-kit-learn LabelEncoders to convert category data into numbers. The "type" section displays transaction types like cash-in, cash-out, DEBIT, PAYMENT, or TRANSFER. "nameDest" and "nameOrig" indicate trade sender and receiver. The 'fit_transform' function assigned a LabelEncoder object to each category column, with a different number assigned to each group. Machine learning techniques were employed to speed up data processing.

2.3. Standardize Numerical

Algorithms like CNN-BiLSTM and Momentum performed better against mobile money frauds when using consistent number fields. For this, the study turned to the Scikit-learn StandardScaler. By setting the mean to 0 and the standard deviation to 1, it is possible to scale numerical data. When there was no bias, and each attribute had an equal probability of progressing the model, the fraud detection approach performed better.

2.4. Dataset Balancing (Resample)

The research identifies mobile money theft using CNN-BiLSTM and Momentum-optimized SGD. Due to a discrepancy in the datasets' classes, the SGD was adjusted. When looking for solutions to societal issues, minorities are given top priority. The majority group is more likely to be a member of both the legitimate and dishonest categories in the sample. Consistently distributed phonetic samples may be obtained by resampling using Scikit-learn. The algorithm can more easily detect unusual activity with an even collection.

2.5. K-Fold Cross Validation

Applying the K-Fold cross validation method with five folds brings up the most accurate model evaluation. Stratified K-Fold was used for the class distribution of all folds to be preserved. Each fold is used for both training and testing, but in the process, a different set is achieved. Training data was used to train a CNN-BiLSTM model, and the model was then improved using a modified SGD with the momentum method. The results of each fold were collected to be later compared to the other folds and were used to measure the extent of the

developed model's new capabilities. This same procedure was repeated for the next five original folds until the whole dataset was covered.

2.6. Training and Testing Set

CNN-BiLSTM is a model used to detect mobile money fraud by analyzing attributes and behaviours. It is easily understandable and verifiable. The model uses Scikit-learn to train and test the model. This method expands the applicability and stability of the model, enabling the identification of scammers using mobile money.

2.7. Feature Importances (Decision Tree Classifier)

Decision Tree Classifier was used for feature importance assessment to improve CNN-BiLSTM-based mobile money fraud detection. This approach identified the dataset's fraud detection strengths—first, a Decision Tree Classifier assessed feature significance. X indicates features, and y fraud/non-fraud was used to train the classifier. The relevance of each feature was determined by the trained decision tree classifier using the 'feature importances_' option. Lower-ranked characteristics related to fraud detection were judged to be the most important. Significant ratings were shown in the end, and feature importance was arranged for every characteristic. This made clear the standards the model used to identify mobile money fraud. Investigation of feature values allowed them to identify dishonesty. The pseudo-code for adopted the decision tree approach is depicted in Figure 2 (Hambali Moshood, n.d.)

Algorithm 1 – DECISION TREE

```

DecTree (Sample Sm, Features Ft)
Steps:
1. Ifstopping_condition (Sm, Ft) = true then
    a. Leaf = createNode()
    b. leafLabel = classify(s)
    c. return leaf
2. root = createNode()
3. root.test_condition = findBestSpilt(Sm, Ft)
4. Q = {q | q a possible outcome of root.test_condition}
5. For each value q ∈ Q:
    a. Smy = {s | root.test condition(s) = q};
    b. Child = TreeGrowth (Smy, Ft);
    c. Add derived child as descent of the root and
    label the edge {root — child} as q
6. return root
    
```

Figure 2: Algorithm for Decision Tree classifier

2.8. Deep Learning Model

The deep learning model for Advanced Mobile Money Fraud Detection was built on a sequential architecture. It started with a Conv1D layer that localized features from transaction sequences, then a MaxPooling1D layer reduced the dimensionality. The application of BatchNormalization was meant to make training more stable and faster. A Bidirectional wrapper around the BiLSTM layer helped in getting the inputs from both directions along the time axis. Dropout was used here to eliminate overfitting. The network also had a Flatten layer that transformed the features into a vector format, then it was followed by two fully connected Dense layers that were used for capturing complex patterns. It was trained with an optimized SGD optimizer with momentum to accelerate convergence.

2.9. CNN Layer

To capture important local temporal patterns in the sequential mobile money transaction data, the model added a Conv1D layer with 64 filters and a kernel size of 3. This layer, using the ReLU activation, re-mapped the raw input signals such as transaction time gaps and amounts into feature maps that were more informative and highlighted potential critical indicators of fraud. The input shape was specified to correspond to the preprocessed training data that allowed the same data format to be passed to the following layers. This convolutional method enabled the system to directly learn complex fraud-related behaviours without the need for manual feature engineering, which is very important for the fast-changing and evolving nature of mobile money fraud detection cases.

2.10. MaxPooling Layer

The research utilised a MaxPooling1D layer of size 2 for the purpose of lowering the dimensionality of the convolutional feature maps. This pooling operation targeted the most prominent signals related to fraud while eliminating the less relevant noise, thus enabling better computation efficiency. By aggregating adjacent features, it contributed to the steadiness of detection accuracy across different transaction patterns. This phase therefore was crucial in feature extraction compression, avoiding overfitting, and allowing the following BiLSTM layer to effectively read the most important time-series fraud hints hidden in the transaction sequences.

2.11. Batch Normalization

Following the pooling, the investigation utilised Batch Normalisation to standardize the feature maps, thus helping to solve the issue of internal covariate shift. This normalisation operation not only made training more stable and quicker, but also ensured that each feature across the batches had the same mean and variance, which is very important for financial data that is susceptible to scale changes. In our highly sophisticated mobile money fraud detection

pipeline, this layer played a significant role in improved convergence when we combined it with Optimised SGD with Momentum, thus allowing the model to handle various fraud patterns, different types of transactions, and the changes in transaction behaviours that occur in different seasons effectively.

2.12. BiLSTM Layer

Researchers conducted experiments with a Bidirectional wrapper around an LSTM layer of 64 units and return sequences = True, which allowed the model to understand transaction sequences from both forward and backwards directions. This dual-context approach enable the model to not only understand simple dependencies but also more complicated ones, such as regurently fraud and transaction reversals. By jointly considering the past and future context, the BiLSTM boosted the model's capability to spot tricky fraud patterns that are usually overlooked by unidirectional models. This move was essential to getting the model more temporally aware, which made it fit the nature of the continously changing scam cases in the mobile money sector perfectly.

2.13. Dropout

In order to avoid overfitting, we also used a Dropout layer which has a dropout rate of 0.2. This layer, throughout the training, randomly switched off 20% of the neurons, thus reducing the co-adaptation of features and helping the model to be more robust when faced with fraud cases that it has not seen before. Dropout, in the case of mobile money fraud detection, ensured that the network learned the fraud representations which were spread out rather than actually memorizing the specific transactions, therefore, it was able to maintain high performance even if the fraud strategies changed.

2.14. Flatten Layer

After going through the sequential and convolutional layers, we used a Flatten layer to change the multi-dimensional feature maps into one vector. This operation not only facilitated the dense layers by collapsing the spatial and temporal fraud cues in to a single format but also made it easier for the model to move from extracting temporal features to high-level fraud risk scoring. Flattening was a key step in connecting sequential learning with fully connected classification.

2.15. Dense Layer

To conclude, this research also employed two Dense layers: initially with 64 neurons and ReLU activation to extract more fraud features and secondly with a single neuron and sigmoid activation for the final fraud probability output. The dense layers carried out intricate nonlinear operations to distinguish fraud cases from non-fraud ones. The Optimised SGD with Momentum facilitated efficient and stable convergence, hence the model was able to attain accurate fraud detection by recognizing inconspicuous but vital transaction patterns.

2.16. Optimizer

The focus of the study is to enhance mobile money fraud detection using a momentum-based stochastic gradient descent (SGD) model using a CNN-BiLSTM optimizer. The optimizer accelerated and stabilized the training process, adjusting model weights and learning rate of 0.01. The momentum coefficient of 0.9 and learning rate of 0.01 improved the consistency and fluidity of parameter updates. This strategy allowed the model to understand complex patterns associated with fraudulent mobile money transactions quickly. Researchers may modify the momentum and learning rate to optimize the model's fraud detection capabilities.

2.17. Compiling and Training

CNN-BiLSTM models using the TensorFlow backend and Keras framework were trained to enhance the detection of mobile money fraud. We used the SGD with momentum optimizer to improve the model by measuring the difference between the predicted and real class labels. The performance of the model was assessed using accuracy measures. The fit approach used data from all batch sizes of 64 to train the algorithm. The train, test and validation data obtained is used to evaluate the performance of the performance. Researchers wanted to make mobile money transfers more secure; therefore, they worked to make the model better at detecting fraud.

2.18. Performance Metrics

This study examines key performance indicators in two stages: training and testing. The evaluation metrics Accuracy, Precision, recall and F1-score, False Positive Rate (FPR) and False Negative Rate (FNR) are mostly considered in measuring the performance of a model. The values from True Positive (TP), False Positive (FP), True Negative (TN), False Negative (FN) and N (where N is the sum of (TP + TN + FP + FN)) are the metrics used to calculate these measures. The different metrics are as follows:

- a) F1 score: $2 * \frac{\text{Prec} * \text{Rec}}{\text{Prec} + \text{Rec}}$ (Chatterjee & Namin, 2019: 227-32)
- b) Recall: $\frac{TP}{TP + FN}$ (Chatterjee & Namin, 2019: 227-32)
- c) Precision: $\frac{TP}{TP + FP}$ (Gibson et al., 2020: 187914-187932)
- d) Accuracy: $\frac{(TP + TN)}{N}$ (Gibson et al., 2020: 187914-187932)
- e) False Positive Rate: $\frac{(FP)}{(FP + TN)}$
- f) False Negative Rate: $\frac{(FN)}{(FN + TP)}$

These performance metrics are considered in the study and used to evaluate the effectiveness proposed model with Batch Normalisation for detecting mobile money fraud, and any necessary adjustments are made to increase the model's performance.

3. RESULTS AND ANALYSIS

The paper explores a method for detecting mobile money fraud using CNN-BiLSTM and Momentum-optimized SGD. It emphasizes the model's usefulness in detecting fraud and its impact on developing safe digital financial ecosystems. The dataset, which includes transaction details, account names, balances, and binary fraud flags, is analysed, highlighting the need for dataset-balancing strategies. The findings will guide future efforts in fraud detection.

3.1.Results of Balance the Dataset (Resample)

The "isFraud" field of the dataset shows a way to do resampling that is meant to help find more mobile money scams. The class difference in the dataset indicates that almost all the trades are real, while only a tiny number are fake. After the resampling method, there were no clear differences between fake activities and those that were not. This showed that the distribution was fair. The fair distribution of the dataset makes it easier for the model to find mobile money transfer networks because it makes the dataset more representative and less biased. Results obtained from the data balancing process are shown in Table 1.

Table 1: The results obtained from the balancing of the dataset

CLASS	FROM	TO
IS NOT FRAUD (0)	6354407	6354407
IS FRAUD (1)	8213	6354407

3.2. Results of Feature Importance Selection (Decision Tree Classifier)

The "newbalanceOrig" which is a core component of the Decision Tree approach is crucial and beneficial for identifying mobile money fraud. Some attributes, such as "type" and "step," are more significant than others. The system achieved a memory score 1.0 due to its accurate identification of all fraudulent transactions. It may be worth contemplating for legal and illicit enterprises due to its durability and favourable F1 rating. The system became perplexed while attempting to sort through the 207 authentic offers. Figure 3 shows the results of feature improvements while Figure 4 shows the ROC curve for the featured result. Figure 5 shows the bar graph of the evaluation Metric score for the feature improvements.

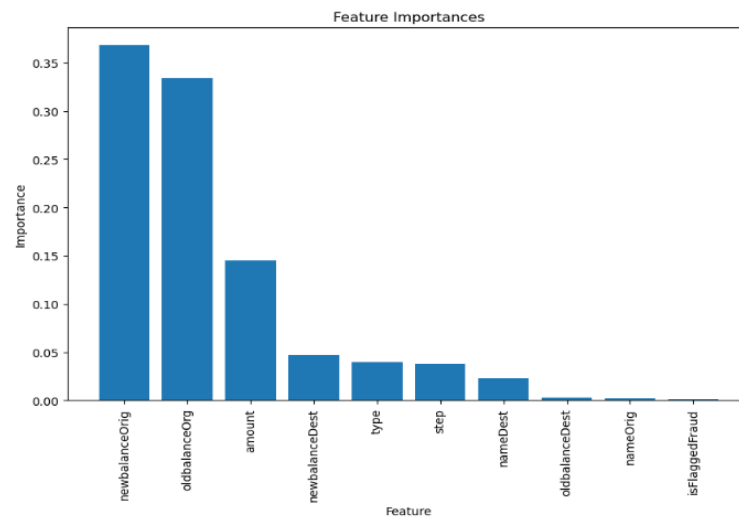


Figure 3: Bar graph representation of feature improvements against its importance

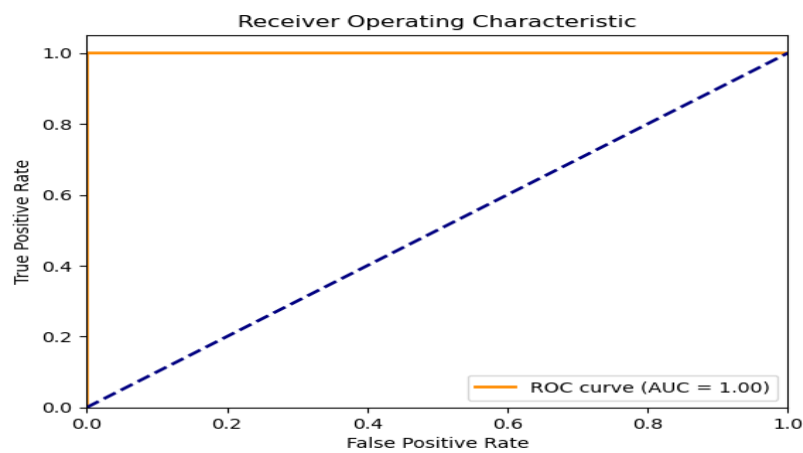


Figure 4: ROC curve of feature improvements

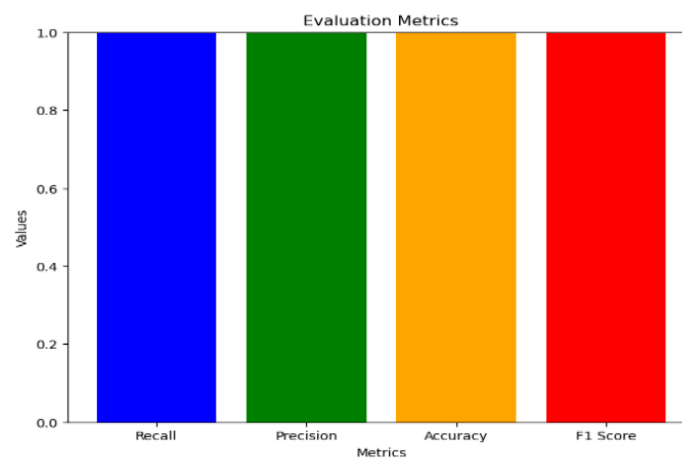


Figure 5: Evaluation Metric Score

3.3. Test Results with Variations of Epoch and Batch Size

Optimal SGD with Momentum is used in this work to test how well the CNN-BiLSTM model can find mobile money plans. For the training, 5-fold cross validation was chosen among the the range of 5 to 10 to reduce the computational costs, in addition to its quick iteration. Furthermore 5-fold cross validation gives a little higher variance with respect to performance estimation which is acceptable in the case of improved speed exchange, which is the expected goal of the training.

Fold 1: Loss was gradually getting smaller by each subsequent batch, eventually getting to 0.0464 in the 10th epoch. The peak of accuracy was 99.50%, while the validation accuracy went up to 99.55%, hence a clear sign of the better learning stability from the increment of the training.

Fold 2: The initial loss was 0.6500, but the error descended to 0.0535 by epoch 10. The accuracy became 99.41% at best with the validation accuracy using 99.55%, leading to the fact that it was a learning and generalisation process that was efficient.

Fold 3: The journey started with a loss of 0.6302 and got down to 0.0455 by the 10th epoch. The validation accuracy hit the highest bar of 99.8%, which suggests that the model's performance has grown significantly with the increase of the batch.

Fold 4: By having obtained 0.6186 as initial loss, the model improved to 0.0456. The accuracy grew to 99.55%, and the validation accuracy maintaind stability over 99.00%, hence the process of learning dynamics was resilient across all batch sizes.

Fold 5: The last part of the loss climbed down from 0.6562 to 0.0624, and the final accracy was 99.19%. The figures for validation accuracy were consistently above 99.00%, even at the highest of 99.55, which showed the strength of the model in detecting fraud. Table 2, Table 3, Table 4, Table 5 and Table 6 shows the results with variations of epoch and batch size for fold 1, fold 2, fold 3, fold 4 and fold 5 respectively. Figure 6 and Figure 7 shows the training and validation loss of the epoch and batch size and accuracies obtained resepectively.



Figure 6: The training and validation loss with respect to epoch and batch size

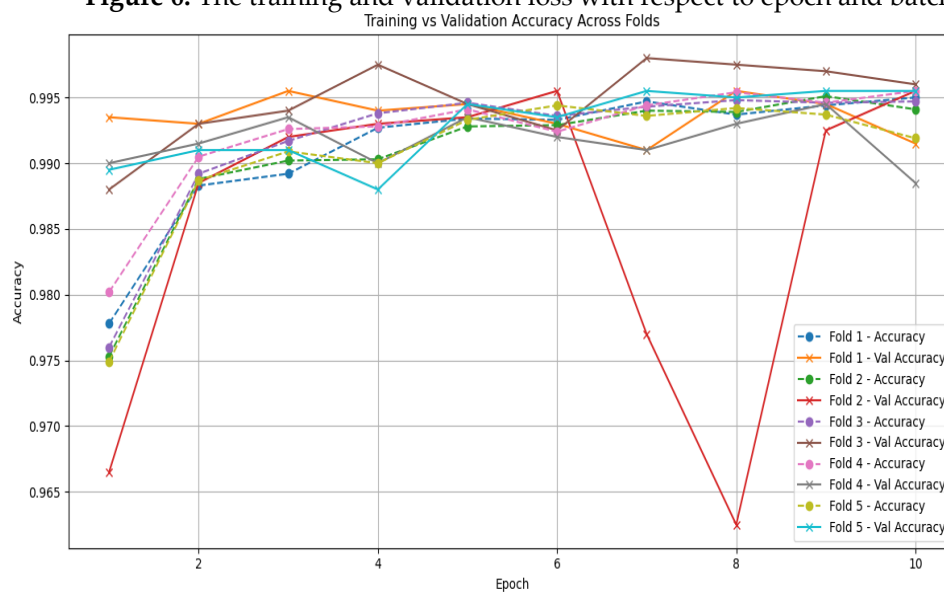


Table 2: The results with variations of epoch and batch size

No of Batches	Loss	Accuracy	Val_Loss	Val_Accuracy
1	0.7713	0.9778	0.1777	0.9935
2	0.1359	0.9883	0.0902	0.9930
3	0.1015	0.9892	0.0701	0.9955
4	0.0746	0.9927	0.0626	0.9940
5	0.0666	0.9934	0.0675	0.9945

6	0.0635	0.9933	0.0590	0.9930
7	0.0536	0.9947	0.0746	0.9910
8	0.0582	0.9937	0.0550	0.9945
9	0.0541	0.9944	0.0474	0.9945
10	0.0464	0.9950	0.0550	0.9915

Table 3: The results with variations of epoch and batch size for fold 2

No of Batches	Loss	Accuracy	Val_Loss	Val_Accuracy
1	0.6500	0.9753	0.2376	0.9665
2	0.1298	0.9888	0.1019	0.9885
3	0.0914	0.9902	0.0812	0.9920
4	0.0865	0.9903	0.0768	0.9930
5	0.0706	0.9928	0.0719	0.9935
6	0.0662	0.9929	0.0574	0.9955
7	0.0582	0.9940	0.1309	0.9770
8	0.0546	0.9939	0.1675	0.9625
9	0.0521	0.9951	0.0554	0.9925
10	0.0535	0.9941	0.0514	0.9955

Table 4: The results with variations of epoch and batch size for fold 3

No of Batches	Loss	Accuracy	Val_Loss	Val_Accuracy
1	0.6302	0.9760	0.1519	0.9880
2	0.1212	0.9892	0.0913	0.9930
3	0.0854	0.9917	0.0582	0.9940
4	0.0677	0.9938	0.0466	0.9975
5	0.0550	0.9946	0.0482	0.9945
6	0.0570	0.9936	0.0557	0.9925
7	0.0537	0.9943	0.0369	0.9980
8	0.0497	0.9948	0.0373	0.9975
9	0.0463	0.9946	0.0373	0.9970
10	0.0455	0.9947	0.0380	0.9960

Table 5: The results with variations of epoch and batch size for fold 4

No of Batches	Loss	Accuracy	Val_Loss	Val_Accuracy
1	0.6186	0.9802	0.1642	0.9900
2	0.1154	0.9905	0.0942	0.9915
3	0.0797	0.9926	0.0804	0.9935
4	0.0721	0.9928	0.0971	0.9900
5	0.0635	0.9941	0.0601	0.9935
6	0.0675	0.9924	0.0678	0.9920
7	0.0546	0.9944	0.0686	0.9910
8	0.0486	0.9954	0.0552	0.9930
9	0.0500	0.9946	0.0548	0.9945
10	0.0456	0.9955	0.0832	0.9885

Table 6: The results with variations of epoch and batch size for fold 5

No of Batches	Loss	Accuracy	Val_Loss	Val_Accuracy
1	0.6562	0.9749	0.1707	0.9895
2	0.1267	0.9887	0.0898	0.9910
3	0.0864	0.9909	0.0941	0.9910
4	0.0979	0.9900	0.0844	0.9880
5	0.0719	0.9933	0.0682	0.9945
6	0.0590	0.9944	0.0673	0.9935
7	0.0603	0.9936	0.0466	0.9955
8	0.0556	0.9942	0.0537	0.9950
9	0.0590	0.9937	0.0359	0.9955
10	0.0624	0.9919	0.0494	0.9955

3.4. Results Of The K-Folds Validation

The model's performance was consistently strong and impressive in all of the K-Folds validation outcomes. The prototype from the first fold was marked by an accuracy of 99.15% and a loss of 0.0560. The second fold had further ascended in its performance as the associated loss was down to 0.0514 with an improved accuracy rate of 99.55%. Following this pattern, the third fold was better than the previous one since the loss was 0.0360, and the accuracy was the highest of 99.60%. An insignificant drop was noticed in the fourth fold, leading to a 0.0832 loss and 98.85% accuracy. The last fold, like the second one, obtained the same level of accuracy at 99.55% with 0.0495

loss. Table 7, shows the loss and accuracy of each fold. Figure 8 shows the loss and accuracy concerning each fold.

Table 7: The loss and accuracy with respect to each fold

No of Folds	Loss	Accuracy
1	0.0560	0.9915
2	0.0514	0.9955
3	0.0360	0.9960
4	0.0832	0.9885
5	0.0495	0.9955

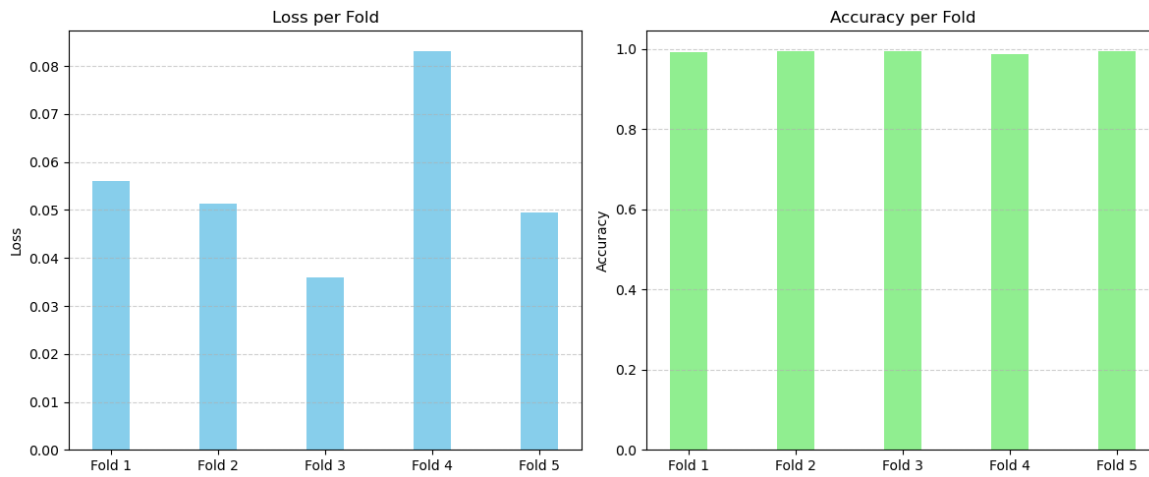


Figure 8: The loss and accuracy with respect to each fold

3.5. Results of the CNN + Bi-LSTM

For the detection of high-tech mobile money fraud, the CNN-BiLSTM model set a benchmark for the evaluation metrics. It had recorded an average recall score of 0.9929, a precision score of 0.9927, and an accuracy of 0.9928. The average F1 score also conferred with it with the highest value of 0.9928, thus an excellent balance of precision and recall was achieved. The false positive rate was very low at 0.0073, and the false negative rate was 0.0071. The model made no mistakes in classifying 1,261,877 fraud cases (TP) and 1,261,550 non-fraud cases (TN). Moreover, it was able to detect 1,261,877 fraud cases (TP) and 1,261,550 situations of non-fraud (TN) with 9,300 false positives and 9,036 false negatives, respectively. Table 8 summarizes the results obtained from the CNN + BiLSTM model. Figure 9 shows the ROC curve of the CNN + BiLSTM model.

Table 8: Summary of the proposed CNN + Bi-LSTM model

MODEL EVALUATION	MODEL RESULT
Recall Score	0.9929
Precision Score	0.9927
Accuracy Score	0.9928
F1 Score	0.9928
FPR	0.0073
TN	1,261,550
FP	9,300
FN	9,036
TP	1,261,877

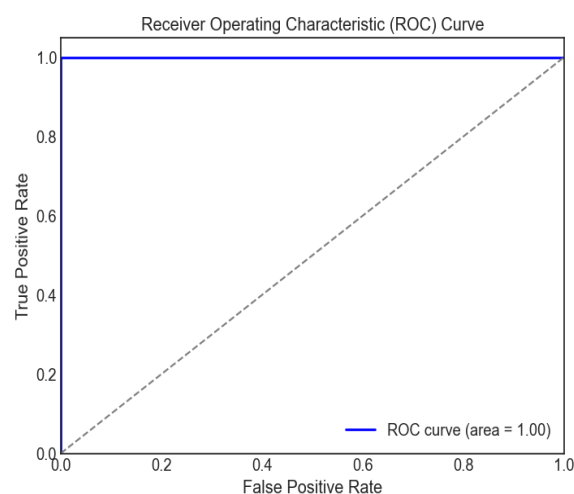


Figure 9: ROC curve of the CNN + BiLSTM model with Area under the Curve (AUC)

3.6. Model Prediction

The model's performance was evaluated on the dataset for validation, testing and training.. The research includes a wide variety of model success metrics for the model. Table 9 shows the classification report. Figure 10 shows the Confusion Matrix of the CNN + BiLSTM model while Figure 11 depicts the bar graph of the CNN + BiLSTM, with the value for each evaluation metric. Figure 12 shows the bar graph of the classification report. The CNN-

BiLSTM model is a great way to spot mobile money scams. It has a 99% success rate in class 1 and a 99% success rate in class 0. Because deals are legal or unfair, it is easy to distinguish

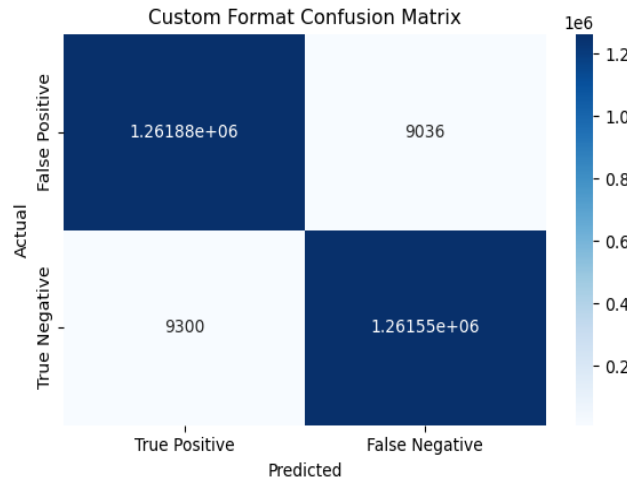


Figure 10: Confusion Matrix of the proposed CNN + BiLSTM model

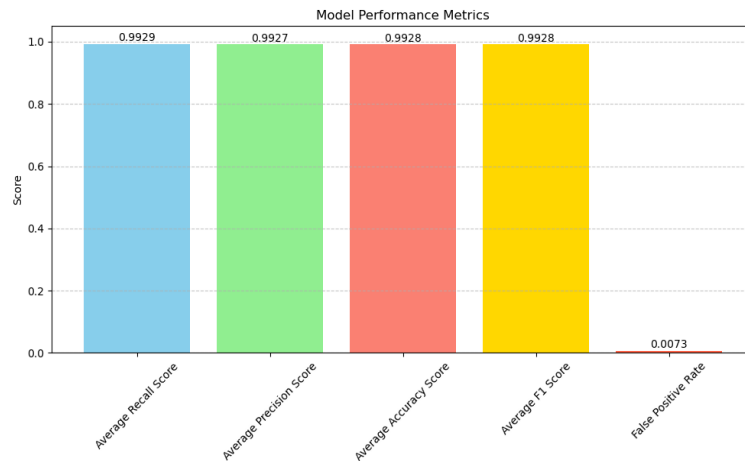


Figure 11: Bar graph showing the evaluation metric score for the proposed CNN + BiLSTM model

Table 9: The classification report

CLASS	PRECISION	RECALL	F1-SCORE	ACCURACY
IS NOT FRAUD (0)	0.99	0.98	0.99	0.99
IS FRAUD (1)	0.98	0.99	0.99	0.99

between real and fake ones. Customers trust financial companies and their services more because this technology can easily find mobile money theft.

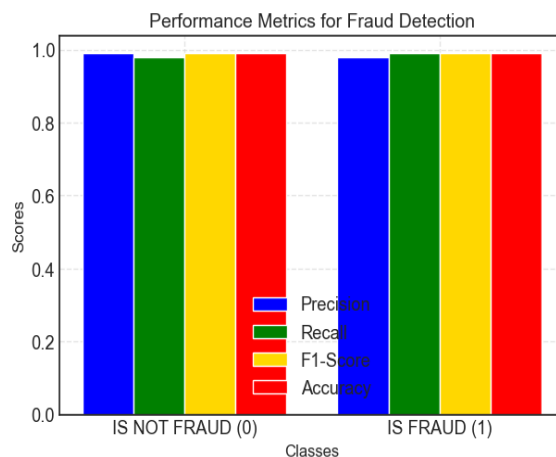


Figure 12: The classification report

The CNN-BiLSTM model effectively detects mobile money fraud with a 99% accuracy rate, distinguishing between authorized and fraudulent transactions. Its recall rate of 98% demonstrates its effectiveness in preventing false negatives. This system is particularly beneficial when prompt and accurate detection of fraudulent behaviour is crucial, as it boosts customer trust in financial institutions.

3.7.Model Comparison

Predicting fraud in mobile money transactions with machine learning: how sampling methods affect a sample that isn't fair (Botchey et al., 2021: 45-56). Revised: Detection of financial fraud in the healthcare sector via the use of machine learning and deep learning techniques has been retracted (Networks, 2023: 1-1). Methods for Detecting Fraudulent Digital Payments in the Digital Era and Various Industries (Chang et al., 2022: 1-31). Table 10 shows the Metric performance of the applied model with their respective datasets.

Table 10: Metric performance of the proposed applied model compared to other detection algorithms with their respective datasets.

MODEL	ACCURACY	PRECISION	RECALL	F1-SCORE
RF (Networks , 2023: 1-1)	98%	97%	97%	97%
DT (Chang et al., 2022: 1-31)	98%	92%	83%	100%

ADASYN (Botchey et al., 2021: 45-56)	98%	0.84%	99%	1.68%
CNN+ BiLSTM (proposed model)	99%	99%	99%	99%

Using stochastic gradient descent with momentum for optimization, the suggested CNN-BiLSTM model showed good classification performance, attaining 99% accuracy, 99% precision, 99% recall, and 99% F1-score. These outcomes outperform related studies using the same dataset and assumption produced by other machine learning models. This model employs machine learning methods to identify fraudulent online payment transactions (Almazroi & Ayub, 2023: 137188-137203). Using machine learning techniques to predict fraudulent mobile money transactions (Lokanan, 2023: 1-24). Detecting fraud with the Kolmogorov-Arnold Network (Gislain & Yurievic, 2025: 1-14). Table 11 shows the Metric performance of the applied model.

Table 11: Metric performance of the proposed applied model compared to other detection algorithms with the proposed dataset.

MODEL	ACCURACY	PRECISION	RECALL	F1- SCORE
RF (Lokanan, 2023: 1-24)	0.89	0.96	0.80	0.87
RXT-J (Almazroi & Ayub, 2023: 137188- 137203)	96%	96%	98%	97%
KAN (Gislain & Yurievic, 2025: 1-14)	97%	97%	97%	97%
CNN+ BiLSTM (proposed model)	99%	99%	99%	99%

With a 99% success rate, the CNN-BiLSTM model is better than most modern ways of finding mobile money schemes. Using convolutional neural networks and bidirectional long short-term memory (BiLSTM) layers, this model is better than all the others. These include Decision Tree (DT), ADASYN, RXT-J, Kolmogorov-Arnold Network (KAN) , and Random

Forest (RF). There are almost no fake wins or rejections produced by deep learning systems. These results show they might make banking institutions safer against mobile money theft.

For sophisticated mobile money fraud detection, the CNN-BiLSTM and Optimised SGD with Momentum model produced exceptional results for a number of reasons. Subtle irregularities in sequential data were captured by CNN layers, which successfully retrieved local transactional patterns. BiLSTM layers improved pattern identification for both fraudulent and genuine acts by learning intricate temporal relationships in both past and future contexts. By smoothing gradient updates, the Optimised SGD with Momentum sped up convergence and prevented local minima, resulting in accurate and steady learning. When combined, these elements yielded a balanced precision of 0.9927, which decreased false alarms, and a high recall of 0.9929, which is essential for identifying fraud efforts. Strong overall performance is reflected in the outstanding F1 score of 0.9928. This architecture successfully handled imbalanced data issues, resulting in a low false positive rate (0.0073) and balanced classification accuracy for fraud and non-fraud classes, which makes it ideal for real-world mobile money fraud detection.

3.8. Limitations

It's critical to recognize the PaySim dataset's limitations even though it offers a useful approximation of mobile money transactions. It might not fully represent the intricacy, volatility, and variety of actual financial behavior because it is a simulated dataset. Because they are predicated on set assumptions, the fraudulent patterns produced might not accurately represent the changing tactics employed by actual fraudsters. Furthermore, simulated data is typically more organized and tidy than real data, which could cause model performance to be exaggerated. These elements imply that even though the outcomes are encouraging, additional verification utilizing actual mobile money data is required to validate the model's resilience and capacity for generalization.

4. CONCLUSION AND RECOMMENDATION

The study demonstrates the effectiveness of CNN-BiLSTM and optimized SGD + Momentum in detecting mobile money fraud. The model has a recall score of 0.9929, a Precision score of 0.9927, and an overall accuracy of 0.9928 in distinguishing between legal and fraudulent transactions. It correctly recognized and rejected occurrences, indicating that advanced deep learning techniques like CNN-BiLSTM architecture and improved SGD + Momentum optimization approaches can mitigate mobile money fraud.

4.1. Conclusion

This study uses momentum-optimized stochastic gradient descent (SGD) and CNN-BiLSTM to detect mobile money theft. The CNN-BiLSTM model achieves high generalisability

and convergence by identifying patterns in time and space and optimizing its momentum and learning rate components. The model outperforms older models like ADASYN, Decision Tree (DT), Random Forest (RF), Kolmogorov-Arnold Network (KAN) and RXT-J, with a precision of 0.9927 and a success rate of 0.9928. The model's effectiveness is evaluated using various methods.

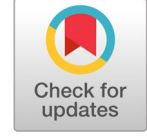
4.2. Recommendation and Future Works

A new method for detecting mobile money fraud has been developed using CNN-BiLSTM and Optimised SGD plus momentum models. This method aims to enhance real-time fraud detection, improve system transparency, and improve feature engineering. Collaboration between police, IT firms, and institutions can efficiently tackle mobile money fraud. Future research will use innovative methodologies like CNN-BiLSTM and federated learning to enhance detection further in addition to Real-time transaction monitoring systems can also help identify and prevent fraudulent behaviour early.

REFERENCES

- Agarwal, A., Iqbal M., Mitra, B., Kumar, V., & Lal, N.(2021). "Hybrid CNN-BILSTM-attention based identification and prevention system for banking transactions," ... *Essent. OILS J.* ..., vol. 8, no. 5, pp. 2552–2560, [Online]. Available: <http://www.nveo.org/index.php/journal/article/view/809>
- Alghofaili, Y., Albattah, A., & Rassam, M. A. (2020). "A Financial Fraud Detection Model Based on LSTM Deep Learning Technique," *J. Appl. Secur. Res.*, vol. 15, no. 4, pp. 498–516, <https://doi.org/10.1080/19361610.2020.1815491>.
- Almazroi, A. A., & Ayub, N. (2023) "Online Payment Fraud Detection Model Using Machine Learning Techniques," *IEEE Access*, vol. 11, no. November, pp. 137188–137203, <https://doi.org/10.1109/ACCESS.2023.3339226>.
- Botchey, F. E., Qin, Z., Hughes-Lartey, K., & Ampomah, K. E.(2021). "Predicting Fraud in Mobile Money Transactions using Machine Learning: The Effects of Sampling Techniques on the Imbalanced Dataset," *Inform.*, vol. 45, no. 7, pp. 45–56, <https://doi.org/10.31449/inf.v45i7.3179>.
- Chatterjee, M., & Namin, A. S. (2019). "Detecting Phishing Websites through Deep Reinforcement Learning." *Proceedings - International Computer Software and Applications Conference* 2(1): 227–32.
- Chang, V., Doan, L. M. T., Di Stefano, A., Sun, Z., & Fortino, G. (2022). "Digital payment fraud detection methods in digital ages and Industry 4.0," *Comput. Electr. Eng.*, vol. 100, pp. 1–31, <https://doi.org/10.1016/j.compeleceng.2022.107734>.
- El Kafhali, S., Tayebi, M., & Sulimani, H.(2024). "An Optimized Deep Learning Approach for Detecting Fraudulent Transactions," *Inf.*, vol. 15, no. 4, pp. 1–25, <https://doi.org/10.3390/info15040227>.
- Gibson, S., Issac, B., Zhang, L., & Jacob, S. M. (2020) "Detecting spam email with machine learning optimized with bio-inspired metaheuristic algorithms," *IEEE Access*, vol. 8, pp. 187914–187932, <https://doi.org/10.1109/ACCESS.2020.3030751>.
- Gislain, Z. N. T., & Yurievich, K. E. (2025). Fraud detection using Kolmogorov-Arnold Network, pp. 1-14.
- Hambali Moshood, A., "Comparative Analysis of Decision Tree Algorithms for Predicting Undergraduate Students' Performance in Computer Programming".
- Igwesi, C. (2023). "Enhancing Authentication and Fraud Detection in Financial Technology," no. December, 2023.

- Kaggle. (n.d.). Kaggle datasets: Paysim1. Retrieved [4th August, 2023], from <https://www.kaggle.com/datasets/ealaxi/paysim1>
- Lokanan, M. E. (2023). "Predicting mobile money transaction fraud using machine learning algorithms," *Appl. AI Lett.*, vol. 4, no. 2, pp. 1–24, <https://doi.org/10.1002/ail2.85>.
- Mbunge, E., Makuyana, R., Chirara, N., & Chingosho, A. (2015). "Fraud Detection in E-Transactions using Deep Neural Networks-A Case of Financial Institutions in Zimbabwe," *Int. J. Sci. Res.*, vol. 6, no. 9, pp. 2319–7064, <https://doi.org/10.21275/ART20176804>.
- Networks, C. (2023). "Retracted: Financial Fraud Detection in Healthcare Using Machine Learning and Deep Learning Techniques," *Secur. Commun. Networks*, vol. 2023, pp. 1–1, <https://doi.org/10.1155/2023/9758612>.
- Sun, Q., Tang, T., Chai, H., Wu, J., & Chen, Y. (2021). "Boosting fraud detection in mobile payment with prior knowledge," *Appl. Sci.*, vol. 11, no. 10, pp. 1–17, <https://doi.org/10.3390/app11104347>.
- Yang, X., Zhang, C., Sun, Y., Pang, K., Jing, L., Wa, S., & Lv, C. (2023) "FinChain-BERT: A High-Accuracy Automatic Fraud Detection Model Based on NLP Methods for Financial Scenarios," *Inf.*, vol. 14, no. 9, pp. 1–26, <https://doi.org/10.3390/info14090499>.
- Zhang, Z., Chen, L., Liu, Q., & Wang, P. (2020). "A Fraud Detection Method for Low-Frequency Transaction," *IEEE Access*, vol. 8, pp. 25210–25220, , <https://doi.org/10.1109/ACCESS.2020.2970614>.
- Zhdanova, M., Repp, J., Rieke, R., Gaber, C., & Hemery, B. (2014). "No smurfs: Revealing fraud chains in mobile money transfers," *Proc. - 9th Int. Conf. Availability, Reliab. Secur. ARES 2014*, pp. 11–20, <https://doi.org/10.1109/ARES.2014.10>.



EmotionUnet: Konuşma Duygu Tanıma için U-Net Tabanlı Özgün Derin Öğrenme Modeli

Yasin Görmez, Sivas Cumhuriyet Üniversitesi, Yönetim Bilişim Sistemleri, Doktor Öğretim Üyesi, yasingormez@cumhuriyet.edu.tr, 0000-0001-8276-2030, 0000-0001-8276-2030

ÖZ

Konuşma, insanlar arasındaki iletişimin en temel ve etkili yolu olarak değerlendirilmektedir. İnsanlar konuşma yolu ile duygu, düşünce ve bilgilerini paylaşmakta, ilişkilerini güçlendirmekte ve toplumsal bağlarını pekiştirmektedir. Konuşma sırasında karşıdaki kişinin duygu durumunun anlaşılması, empati kurarak daha etkili ve anlamlı bir iletişim sağlamak için önemlidir. Günümüzde telefon gibi araçlarla yapılan uzaktan konuşmalarda ifade edilen duygu tonlarının anlaşılması için konuşma duygu tanıma yöntemlerinden sıklıkla faydalanılmaktadır. Konuşma duygu tanıma müşteri hizmetleri, sağlık, eğitim, eğlence ve akıllı sistemler gibi birçok alanda kullanılmaktadır. Konuşma duygu tanımada sinyal işleme, istatistiksel analiz ve biyometrik teknikler gibi yöntemler kullanılırken, son zamanlarda derin öğrenme yöntemleri de yaygınlaşmıştır. Bu çalışmada konuşma duygu tanıma için evrimsel sinir ağları kullanılarak U-Net tabanlı özgün derin öğrenme modeli önerilmiştir. Önerilen modelin hiper-parametre optimizasyonları için Bayesian optimizasyon yönteminden faydalanılmıştır. Önerilen model Türkçe, İngilizce, Arapça ve Bangla dillerinden dört farklı veri ile analiz edilmiştir. Önerilen model ile farklı veri setlerinde %56,55 ile %99,71 arasında doğruluk hesaplanmıştır.

Anahtar Kelimeler

: Konuşma Duygu Tanıma, Derin Öğrenme, Evrimsel Sinir Ağları, U-Net, Makine Öğrenmesi

EmotionUnet: A Novel Deep Learning Model Based on U-Net for Speech Emotion Recognition

ABSTRACT

Speech is considered to be the most basic and effective way of communication between people. Through speaking, people share their feelings, thoughts and information, strengthen their relationships and reinforce their social bonds. It is important to understand the emotional state of the other person during the conversation in order to provide a more effective and meaningful communication by empathizing. Today, speech emotion recognition methods are frequently used to understand the emotional tones expressed in remote conversations via tools such as telephones. Speech emotion recognition is used in many fields such as customer service, healthcare, education, entertainment, and intelligent systems. While signal processing, statistical analysis and biometric techniques are used in speech emotion recognition, deep learning methods have recently become widespread. In this study, a novel U-Net based deep learning model for speech emotion recognition using convolutional neural networks is



proposed. Bayesian optimization method is used for hyper-parameter optimization of the proposed model. The proposed model is analyzed with four different datasets from Turkish, English, Arabic and Bangla languages. The accuracy of the proposed model is calculated between 56,55% and 99,71% on different datasets.

Keywords : Speech Emotion Recognition, Deep Learning, Convolutional Neural Network, U-Net, Machine Learning

EXTENDED ABSTRACT

INTRODUCTION

Effective communication relies not only on the exchange of thoughts, feelings, and information, but also on the ability to strengthen relationships and reinforce social bonds. A key component of this process is understanding the emotional state of the other person, as it facilitates empathy and leads to more meaningful interactions. In face-to-face communication, cues such as gestures and facial expressions make it easier to discern emotional states. However, in remote communication—such as phone conversations—this becomes more challenging. In such cases, emotional understanding is often derived from vocal elements such as tone, accent, speech rate, volume, pauses, and word choice. Speech Emotion Recognition (SER) methods have been developed to automatically detect emotional states in spoken language, offering a technological solution to this challenge.

In this study, a novel deep learning model for SER is developed using the U-Net architecture. The model incorporates both Convolutional Neural Network (CNN) and Fully Connected Layer (FCL) structures. To enhance model performance, hyper-parameters are optimized using the Bayesian optimization approach. The proposed model is evaluated on datasets from four languages: Turkish, English, Arabic, and Bangla. The primary objective of this study is to assess the performance of the U-Net-based deep learning model in the context of SER. A comprehensive review of the existing literature reveals that U-Net-based architectures have not previously been applied to SER tasks. Therefore, the model developed in this study is considered to be a novel contribution to the field.

MATERIALS AND METHODS

In this study, a novel U-Net-based deep learning model incorporating CNN and FCL layers was developed. The model was trained on the TurEV-DB, BanglaSER, TESS, and ANAD datasets, and its performance was evaluated on each. Key performance metrics, including accuracy, precision, recall, and area under the curve, were employed to assess the model's effectiveness. Spectrogram representations of the audio data, used as inputs, were generated using Short-Time Fourier Transform. Additionally, the model's hyperparameters were optimized through a Bayesian optimization approach.

RESULTS AND DISCUSSION

In the initial phase of the study, each dataset was divided into three subsets: training, validation, and testing. Specifically, 25% of the data was randomly selected for the testing set, 10% for the validation set, and the remaining 65% for the training set. The training set was used to fit the model, the validation set to monitor model performance during hyper-parameter optimization and to trigger callback functions during training, and the test set to evaluate the final performance of the trained models. Once the datasets were partitioned, the hyper-parameters of the proposed model were optimized individually for each dataset. This optimization was performed using the Bayesian approach, implemented through the scikit-optimize library in Python. The results of the hyper-parameter optimization are presented in Table 2.

After the hyper-parameter tuning phase, the model was trained separately on each dataset using the optimized hyper-parameters. During training, two callback functions were incorporated for learning rate reduction and early stopping. The learning rate reduction function halved the learning rate if the validation loss did not improve for two consecutive iterations, while the early stopping function halted training if the validation loss remained unchanged for six consecutive iterations, preventing overfitting. The performance metrics, calculated after training, are displayed in Table 3. These results indicate that the model achieved the highest accuracy on the ANAD dataset and the lowest accuracy on the BanglaSER dataset, with noticeably lower accuracy rates observed on the BanglaSER dataset compared to the others.

To further evaluate model performance, it is critical to analyze results on a class-by-class basis. For this purpose, precision-recall curves were plotted for each dataset, as shown in Figures 2–5. For the TurEV-DB dataset, the best performance was observed in the "sad" class, while the "calm" class yielded the weakest performance. The "angry" and "happy" classes showed similar patterns. In the ANAD dataset, all classes exhibited near-perfect curves. For the BanglaSER dataset, the "neutral" and "angry" classes outperformed the others. In the TESS dataset, nearly all classes showed near-perfect curves except for the "surprised" class, which had a slightly lower performance.

When compared to the existing literature, the proposed model demonstrates superior results on the TESS and ANAD datasets, achieving almost flawless performance. For the BanglaSER dataset, the model's results are close to the literature but still show an improvement over previously reported finding.

CONCLUSIONS

In this study, a novel U-Net-based deep learning model for Speech Emotion Recognition is proposed, utilizing CNN and FCL layers. The model was trained on datasets from four different languages, and its performance was evaluated based on key metrics.

Hyper-parameter tuning was conducted using Bayesian optimization, which offers significant advantages over traditional methods such as grid search for deep learning applications. The experimental results reveal that the proposed model achieved accuracies ranging from 56.55% to 99.71%. When compared to the existing literature, the model underperforms on the TurEV-DB dataset, while slightly outperforming the literature on the BanglaSER dataset. For the ANAD and TESS datasets, the model demonstrates superior performance compared to previously reported results.

GİRİŞ

Konuşma insanlar arasındaki iletişimin en temel ve etkili yoludur. Konuşma sayesinde insanlar; duygu, düşünce ve bilgilerini paylaşabilmekte, ilişkilerini güçlendirebilmekte ve toplumsal bağları pekiştirebilmektedir. Konuşma esnasında karşıdaki kişinin duygu durumunun bilinmesi, empati kurarak daha etkili ve anlamlı bir iletişim sağlanması için büyük önem arz etmektedir. Bu ön bilgi sayesinde insanlar rahatlıkla iş birliği yaparak, birlikte sorunlara çözümler bulabilmektedir. Bu nedenle insanlar konuşma esnasında karşıdaki insanın duygu durumunu anlamak için bir çaba sarf etmektedir. Yüz yüze konuşma esnasında karşıdaki kişinin duygu durumunu anlamak, jest ve mimik gibi unsurlar sayesinde daha rahat olmakta ancak sohbet esnasında kişilerin birbirini görmediği durumlarda duygu durumlarının anlaşılması daha zor olmaktadır. Günümüzde sıklıkla telefon vb. araçlar sayesinde yapılan uzaktan konuşmalarda duygu durumunu; ses tonu, vurgular, hız, ses seviyesi, duraklamalar ve kelime seçimi gibi unsurların analiziyle anlaşılabilir. Konuşma esnasındaki kişilerin duygu durumlarının otomatik bir şekilde elde edilmesi için ise konuşma duygu tanıma (Speech Emotion Recognition- SER) yöntemlerinden faydalanılmaktadır.

SER, konuşma sırasında ifade edilen duygusal tonların ve ruh hallerinin otomatik olarak belirlendiği süreci ifade etmektedir. SER; müşteri hizmetleri, sağlık, eğitim, eğlence ve akıllı sistemler gibi birçok alanda farklı amaçlar ile kullanılmaktadır (Madanian et al., 2023; Zhu et al., 2017). SER kullanım alanlarına acil çağrı merkezlerinde tehlikeli veya yaşamı tehdit eden durumları tanımlamak, araç içi interaktif sesli yanıt sisteminde yorgun sürücülerini tanımlamak, zihinsel sağlık teşhislerini desteklemek ve çevrimiçi eğitim sistemlerinde öğrenci tepki durumlarını anlamak örnek olarak verilebilmektedir (Ahmad et al., 2016; Kerkeni et al., 2020; D. Mishra & rawat, 2015; Zhou et al., 2016).

Geliştirilen SER uygulamalarında konuşulan dil büyük önem arz etmektedir çünkü her dilin kendine özgü ses tonları, ritimleri, vurgu düzenleri ve duygu ifadeleri bulunmaktadır. Bu unsurlar duyguların nasıl ifade edildiğini ve algılandığını büyük ölçüde etkilemekte bu nedenle farklı dillerdeki vurgu, ritim ve tonlama farklılıkları duygusal durumların

tanınmasında kritik rol oynamaktadır. Bu nedenle bir dilde başarılı sonuçlar elde eden bir modelin başka bir dilde istenilen başarıyı elde etmemesi olası bir durumdur (Ververidis & Kotropoulos, 2006). Bahsedilen sebeplerden ötürü tasarlanan bir uygulamanın farklı dillerde de analizinin yapılması, uygulamanın genelleme yapabilme kabiliyetini ölçmek için önemlidir.

SER uygulamalarında sinyal işleme, istatistiksel analiz ve biyometrik teknikler gibi yöntemler kullanılabilir. Bu yöntemlerin arasında fourier dönüşümü, wavelet dönüşümü, cepstral analiz, elektromiyografi, elektrodermal aktivite ve kural tabanlı sistemler gibi yaklaşımlar bulunmaktadır (Wagner et al., 2005). Bu yöntemlerin yanı sıra makine öğrenmesi ve derin öğrenme yöntemleri de SER için sıklıkla kullanılmaktadır. Son zamanlarda birçok problemde olduğu gibi SER uygulamalarında da derin öğrenme yaklaşımlarının model performanslarında iyileşmeler yaptığı görülmektedir (Satt et al., 2017).

Bu çalışmada SER için özgün bir derin öğrenme modeli U-Net mimarisi kullanılarak geliştirilmiştir. Geliştirilen modelde evrimsel sinir ağları (Convolutional Neural Network-CNN) ve tam bağlı katman (Fully Connected Layer- FLC) yapılarından faydalanılmıştır. Derin öğrenme modellerinde başarı oranını etkileyen en önemli unsurlardan biri ise kullanılan hiper-parametrelerdir. Çalışmada önerilen modelin hiper-parametreleri Bayesian optimizasyon yaklaşımı ile optimize edilmiştir. Önerilen model Türkçe, İngilizce, Arapça ve Bangla dili olmak üzere dört farklı dilden veri seti kullanılarak analiz edilmiştir. Bu çalışmanın amacı U-Net tabanlı geliştirilen derin öğrenme modelinin SER üzerindeki performansının ölçülmesidir. Çalışma kapsamında yapılan literatür taramaları doğrultusunda daha önce U-Net tabanlı mimarinin SER için kullanılmadığı görülmüştür. Bu nedenle çalışma kapsamında geliştirilmiş olan model özgün olarak değerlendirilmektedir. Çalışmanın devamında sırası ile literatür taramasından bahsedilmiş, yöntemler anlatılmış, analiz sonuçlarına değinilmiş ve tartışma kısmı ile çalışma sonlandırılmıştır.

1. LİTERATÜR TARAMASI

Literatürde hem geleneksel makine öğrenmesi hem de derin öğrenme yöntemleri kullanılarak SER çalışmaları yapılmıştır. Lin ve Wei yapmış oldukları çalışmada farklı öznitelik teknikleri kullanarak eğittikleri saklı Markov modelleri, destek vektör makinaları ve k en yakın komşu modellerinde sırasıyla %99,5, %88,9 ve %84,5 doğruluk elde etmişlerdir (Lin & Wei, 2005). Jha ve diğerleri RAVDESS veri setinden altı farklı yöntemle öznitelik çıkararak eğittikleri naive Bayes, rastgele orman, k en yakın komşu, destek vektör makinaları ve yapay sinir ağı modellerinden en başarılı yöntemlerin destek vektör makinaları ve yapay sinir ağları olduğu sonucuna varmışlardır (Jha et al., 2022). Krishna ve diğerleri destek vektör makinaları ve yapay sinir ağı sınıflandırma yöntemleri ile geliştirmiş oldukları modellerde %86,50 doğruluğa ulaşmışlardır (Krishna et al., 2022). Liu ve diğerleri CASIA veri seti üzerine korelasyon analizi ve Fisher yöntemi ile öznitelik seçimi uygulamış ve aşırı öğrenme

makinaları yöntemi kullanarak eğittikleri modelde %89,6 doğruluk elde etmişlerdir (Liu et al., 2018). Ghai ve diğerleri mel-frekansı kepsral katsayıları ile öznitelik çıkararak eğittikleri destek vektör makinaları, rastgele karar ormanı ve gradyan artırma yöntemleriyle %81,05 doğruluğa ulaşmışlardır (Ghai et al., 2017).

Harrar ve arkadaşları önermiş oldukları CNN tabanlı derin öğrenme modelini EMO-DB veri seti ile eğiterek %96,97 doğruluğa ulaşmışlardır (Harár et al., 2017). Issa ve arkadaşları geliştirmiş oldukları CNN tabanlı derin öğrenme modelini kullanarak RAVDESS ve EMO-DB veri setlerinde sırasıyla 71,61% ve 95,71% doğruluğa ulaşmışlardır (Issa et al., 2020). Anvarjon ve arkadaşları duygu konuşma tanıma için önerdikleri hafif etkili CNN modeliyle IEMOCAP veri setinde %77,01, EMO-DB veri setinde %92,02 doğruluk elde etmişlerdir (Anvarjon et al., 2020). Mustaqeem ve arkadaşları kısa zamanlı Fourier dönüşümünden (Short-time Fourier transform- STFT) elde edilen spektrogram değerlerini girdi olarak kullanan CNN ve uzun kısa süreli bellek (Long Short-Term Memory - LSTM) tabanlı derin öğrenme modelinde IEMOCAP, EMO-DB ve RAVDESS veri setlerinde sırasıyla 74%, 86% ve 81% F1 skoru elde etmişlerdir (Mustaqeem et al., 2020). Mustaqeem ve Kwon SER için önermiş oldukları tek boyutlu genişletilmiş CNN tabanlı çoklu öğrenme mimarisinde IEMOCAP ve EMO-DB veri setlerinde sırasıyla %73 ve %90 doğruluk elde etmişlerdir (Mustaqeem & Kwon, 2021). Nediyanath ve diğerleri çok kafalı dikkat ağları ve log mel filtre bankası enerjileri kullanarak geliştirmiş oldukları modelle IEMOCAP veri setinde %76,4 doğruluğa ulaşmışlardır (Nediyanath et al., 2020). Singh ve arkadaşları CNN ve destek vektör makinalarını birlikte kullanarak önerdikleri modelde farklı spektrogram yöntemlerini deneyerek EMO-DB veri setinde %79,86, RAVDESS veri setinde %52,54 doğruluğa ulaşmışlardır (Singh et al., 2023). Liu ve arkadaşları dikkat tabanlı LSTM ve CNN kullanarak geliştirmiş oldukları modelde IEMOCAP ve IMPROV veri setlerinde sırasıyla %70,27 ve %60,90 ortalama duyarlılık skoru elde etmişlerdir (Liu et al., 2023). Khan ve diğerleri önerdikleri çok modlu SER modeli sayesinde IEMOCAP ve MELD veri setlerinde doğruluğu %4,5 iyileştirmişlerdir (Khan et al., 2024). Mishra ve arkadaşları farklı öznitelik çıkarma tekniklerini kullanarak elde ettikleri girdiler ile derin sinir ağı eğiterek %87,48 doğruluğa ulaşmışlardır (S. P. Mishra et al., 2024). Mary Little Flower ve diğerleri veri çoğaltma teknikleri sayesinde CNN tabanlı modelin doğruluğunu AESDD, CAFE, EmoDB, IEMOCAP ve MESD veri setlerinde sırasıyla %99,2, %99,5, %97,5, %92,4, ve %96,9 olarak hesaplamışlardır (Mary Little Flower et al., 2024).

2. MATERYAL VE METOT

Çalışmada CNN ve FLC katmanları kullanılarak U-Net tabanlı özgün derin öğrenme modeli geliştirilmiştir. Yapılan literatür taraması doğrultusunda konuşma duygu tanıma alanında U-Net tabanlı bir mimari yapının kullanılmamış olması ise çalışmamızda önerilen modelin özgün olarak değerlendirilmesini sağlamaktadır. Geliştirilen model dört farklı dilden

veri setleri kullanılarak eğitilmiş ve modelin bu veri setlerindeki performans skorları hesaplanmıştır. Girdi olarak kullanılan ses verilerinin spektrogram gösterimleri STFT kullanılarak elde edilmiş ve modelin hiper-parametreleri Bayesian optimizasyon yaklaşımı ile optimize edilmiştir. Girdi olarak STFT spektrogram verileri kullanmak zamana bağlı frekans bilgisini modele sunarken, sesin içeriğini daha iyi temsil eden bir özellik uzayı oluşturmayı sağlamıştır. Hiper-parametre optimizasyonu aşamasında Bayesian yaklaşımının kullanılması ise daha az deneme sayısı ile modeller için en uygun hiper-parametrelerin tespit edilmesini sağlamıştır.

2.1. VERİ SETİ

Çalışmada Türkçe dilinden TurEV-DB (Canpolat et al., 2020), Arapça dilinden ANAD (Altamimi & Alayba, 2023), Bangla dilinden BanglaSER (Das et al., 2022) ve İngilizce dilinden TESS (Pichora-Fuller & Dupuis, 2020) veri setleri kullanılmıştır. Bu veri setleri, her biri için sunulan referanslarda yer alan bağlantılardan indirilmiştir. ANAD veri setinde 741 kızgın, 505 mutlu ve 137 şaşırılmış olmak üzere toplam 1383 örnek bulunmaktadır. BanglaSER veri setinde 306 kızgın, 306 mutlu, 306 üzgün, 306 şaşırılmış ve 243 nötr olmak üzere toplam 1467 örnek bulunmaktadır. TurEV-DB veri setinde 487 kızgın, 408 sakin, 357 mutlu ve 483 üzgün olmak üzere toplam 1735 örnek bulunmaktadır. TESS veri setinde ise korku, şaşırılmış, üzgün, sinirli, iğrenmiş, mutlu ve nötr sınıflarının her birinden 400 olmak üzere 2800 örnek bulunmaktadır. Veri setleri hakkındaki detaylı bilgilere ilgili makalelerden ulaşılabilmektedir.

2.2. KISA ZAMANLI FOURIER DÖNÜŞÜMÜ

STFT, sinyali kısa zaman aralıklarına bölerek her bir aralığın Fourier dönüşümünü hesaplayan zaman-frekans analiz yöntemidir. STFT sinyalin zaman içerisindeki frekans içeriğini incelemeye olanak tanıdığı için özellikle konuşma ve müzik sinyalleri gibi zamanla değişen sinyalleri analiz etmek için sıklıkla kullanılmaktadır. Her bir kısa zaman aralığı, genellikle Hamming veya Hanning gibi bir pencere fonksiyonu ile çarpılmakta ve bu pencere fonksiyonu sinyalin belirli bir kısmını izole etmektedir. STFT hem zamansal hem de frekans bilgilerini koruyarak sinyalin daha kapsamlı bir analizini sağlamaktadır. Zaman-frekans çözünürlüğü, pencere uzunluğuna bağlı olarak değişmektedir. Daha kısa pencereler daha iyi zaman çözünürlüğü sağlarken, daha uzun pencereler daha iyi frekans çözünürlüğü sağlamaktadır (Allen & Rabiner, 1977).

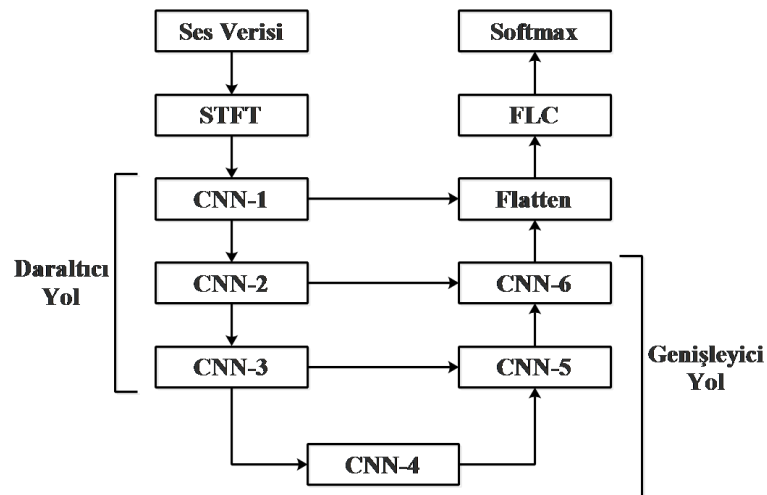
2.3. BAYESIAN OPTIMIZATION

Bayesian optimizasyon, özellikle hesaplamalı olarak pahalı ve zaman alıcı fonksiyonları optimize etmek için kullanılmaktadır. Bayesian optimizasyon yaklaşımı, bir fonksiyonun belirli bir noktadaki değerini tahmin etmek için Bayes teoremini kullanılmaktadır. Bu yöntemde genellikle fonksiyonun bilinmeyen kısımlarını modellemek

için Gaussian süreçlerinden faydalanılmaktadır. Bayesian optimizasyon, özellikle hiper-parametre optimizasyonu gibi alanlarda yaygın olarak kullanılmakta ve derin yapay sinir ağları gibi hiper-parametre sayısı çok ve hiper-parametre uzayı geniş olan modellerde tercih edilmektedir. Bayesian optimizasyon, daha önce kendisine sunulan örneklerden öğrenerek yeni denemelerin nerede yapılması gerektiğini belirleyerek optimize edilecek fonksiyonun minimum veya maksimum değerine daha hızlı ulaşmayı amaçlamaktadır. Öngörücü model ve kazanım fonksiyonu olmak üzere iki temel bileşene dayanmaktadır. Kazanım fonksiyonu, yeni denemelerin yapılacağı noktaları seçerken keşif ve sömürü arasında denge kurmaktadır. Bu denge, yeni bölgeleri keşfetmek ile mevcut bilgiye dayanarak en iyi bilinen çözümü sömürmek arasında bir seçim yapılmasını sağlamaktadır (Snoek et al., 2012).

2.4. U-NET TABANLI DERİN ÖĞRENME MODELİ

Çalışmada altı adet CNN modülü ve bir adet FLC katmanı içeren, EmotionUnet ismi verilen U-Net tabanlı derin öğrenme modeli geliştirilmiştir. ProteinUnet (Kotowski et al., 2021) modeline benzer bir şekilde, Python dilinde var olan keras (*Keras: Deep Learning for Humans*, 2024) kütüphanesi ile geliştirilmiş olan bu modelin mimari yapısı Şekil 1’de gösterilmektedir.



Şekil 1: Emotionunet Modeli Mimari Diyagramı

Şekil 1’de de görüleceği üzere model, CNN katmanlarını daraltıcı yol (contracting path) ve genişleyici yol (expansive path) olmak üzere iki ana katmana ayırmıştır. Daraltıcı yol birbirine seri şekilde bağlı üç CNN modülünden oluşurken genişleyici yol hem birbirine seri şekilde bağlı hem de daraltıcı yolda kendisine karşılık gelen CNN modülüne bağlı üç CNN modülünden oluşmaktadır. Her bir CNN modülü sırasıyla tek boyutlu evrişim ve Batch Normalizasyonu katmanlarından oluşmaktadır. Tek boyutlu evrişim katmanı için filtre sayıları her bir modül için ayrı ayrı optimize edilmiştir. Daraltıcı yoldaki CNN modülleri için çekirdek genişliği 7, genişleyici yoldaki CNN modülleri için ise çekirdek genişliği 5 olarak

ayarlanmıştır. Bu katmandaki aktivasyon fonksiyonu relu ve padding parametresi same olarak ayarlanırken diğer parametreler ise varsayılan olarak bırakılmıştır. CNN katmanlarından sonra ise modele aktivasyon fonksiyonu relu olan bir FLC katmanı eklenmiş daha sonra model softmax katmanı ile sonlandırılmıştır. FLC katmanındaki nöron sayısı da optimize edilmiş, kalan parametreler ise varsayılan ayarda bırakılmıştır. Modelde optimizasyon algoritması Adam ve kayıp fonksiyonu categorical_crossentropy olarak ayarlanmıştır. Modelin öğrenme adımı, döngü sayısı ve parça sayısı değerleri ise optimize edilmiştir. Çalışmada konuşma duygu tanıma için önerilmiş olan U-Net mimarisinin çeşitli faydalarının olduğu düşünülmektedir. Bunlardan ilki iniş-çıkış yaklaşımıyla hem geniş kapsamlı (global) hem de ince ayrıntılı (local) desenleri yakalayabilen yapısı sayesinde ses verisinde farklı zaman aralıklarında oluşan kısa ya da uzun süreli olayları tanıyabilmesidir. Özellikle görsel verilerinde başarılı sonuçlar veren U-Net mimarisi, spektrogram gibi görsel özellikleri kullanan ses verileri ile yüksek uyuma sahiptir. Daraltıcı yol bölümünde özellikler sıkıştırılırken, genişleyici yol bölümünde özellikler genişletilmektedir. Bu duruma en olarak daraltıcı yol ve genişleyici yol bölümlerinde geçiş bağlantısının olması ise veri kaybını minimize etmektedir.

3. DENEY SONUÇLARI

Çalışmanın ilk aşamasında her bir veri seti eğitim, test ve validasyon olmak üzere üçe bölünmüştür. Bu kapsamda örneklerin rastgele %25'i seçilerek test, %10'u seçilerek validasyon ve kalanlar ile ise eğitim veri seti oluşturulmuştur. Her bir veri seti için eğitim, test ve validasyon veri setindeki örnek sayıları tablo 1'de gösterilmektedir.

Tablo 1: Her Bir Veri Setinden Elde Edilen Eğitim, Test ve Validasyon Veri Setleri İçin Örnek Sayıları

Veri Seti Adı	Eğitim Örnek Sayısı	Test Örnek Sayısı	Validasyon Örnek Sayısı
TurEV-DB	1129	433	173
ANAD	900	345	138
BanglaSER	955	366	146
TESS	1431	549	219

Oluşturulan veri setlerinden, eğitim veri seti modeli eğitmek, validasyon veri seti hiper-parametre optimizasyonu aşamasında model performansı hesaplamak ve model eğitimi esnasında geri çağırma (callback) fonksiyonlarını tetiklemek, test veri seti ise eğitilmiş modellerin performansını hesaplamak için kullanılmıştır. Eğitim, test ve validasyon veri setleri oluşturulduktan sonra önerilen modelin hiper-parametreleri her bir veri seti için ayrı ayrı optimize edilmiştir. Hiper-parametre optimizasyonu için kullanılan Bayesian optimizasyon yöntemi Python dilinde var olan scikit-optimize (skopt) (*Scikit-Optimize*, 2025) kütüphanesi kullanılarak geliştirilmiştir. Yöntem için bu kütüphanede var olan gp_minimize

fonksiyonu $acq_func='EI'$, $n_calls=50$ parametre ayarları ile kullanılmıştır. Bu yöntem çizge aramadan farklı olarak, her bir hiper-parametre için bir aralık ya da kategorik bir uzay alarak, belirlenen uzaydan en uygun değerleri Gauss sürecini kullanarak tespit etmeye çalışmaktadır. Optimize edilen hiper-parametreler için hiper-parametre uzayı, hiper-parametre adı, hiper-parametre türü ve her bir veri seti için bulunan optimum değer tablo 2’de gösterilmektedir.

Tablo 2: Optimize edilen hiper-parametreler için hiper-parametre adı, uzayı ve optimum değerler

Hiper-parametre Adı	Hiper-parametre Türü	Hiper-parametre uzayı	TurEV-DB değeri	optimum	ANAD optimum değeri	BanglaSE R optimum değeri	TESS optimum değeri
öğrenme adımı	Gerçek Sayı	En düşük = 10^{-4} , En yüksek = 10^{-1}	0,000232		0,0001208	0,0001	0,0011830
döngü sayısı	Tam Sayı	En düşük = 5, En yüksek = 100	100		52	85	29
parça sayısı	Kategorik	[1, 2, 4, 8, 16, 32, 64, 128]	16		4	8	16
filtre sayısı (CNN-1)	Tam Sayı	En düşük = 16, En yüksek = 256	256		128	68	34
filtre sayısı (CNN-2)	Tam Sayı	En düşük = 16, En yüksek = 256	106		39	16	242
filtre sayısı (CNN-3)	Tam Sayı	En düşük = 16, En yüksek = 256	256		57	256	194
filtre sayısı (CNN-4)	Tam Sayı	En düşük = 16, En yüksek = 256	196		93	256	124

filtre sayısı (CNN-5)	Tam Sayı	En düşük = 16, En yüksek = 256	203	256	116
filtre sayısı (CNN-6)	Tam Sayı	En düşük = 16, En yüksek = 256	169	256	186
nöron sayısı (FLC)	Tam Sayı	En düşük = 100, En yüksek = 1500	596	794	864

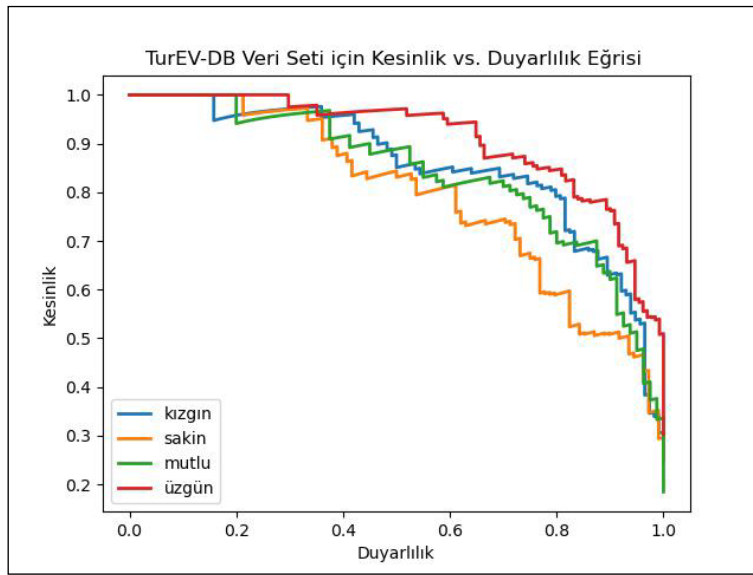
Tablo 2’de yer alan öğrenme adımı (learning rate), her adımda model ağırlarının değiştirilme katsayısını; döngü sayısı (epoch), modelinin eğitim verisinin tamamını kaç kez geçeceğini; parça sayısı (batch size), modelinin eğitim sırasında aynı anda işlediği veri örneklerinin sayısını; filtre sayısı, ilgili evrişim katmanındaki pencere sayısını; nöron sayısı ise FLC katmanındaki birim sayısını ifade etmektedir. Filtre sayısı hiper-parametresi her bir CNN modülü için ayrı ayrı optimize edildiği için tabloda altı farklı filtre sayısı verilmiştir. Filtre sayısı hiper-parametresinde hangi modül için kullanıldığı parantez içerisinde verilmiştir.

Hiper-parametre optimizasyonu aşamasından sonra, belirlenen hiper-parametreler kullanılarak önerilen model her bir veri seti için eğitim kümeleri kullanılarak ayrı ayrı eğitilmiştir. Eğitim esnasında modele öğrenme hızı azalması ve erken durdurma için iki adet geri çağırma fonksiyonu eklenmiştir. Öğrenme hızı azaltma fonksiyonu, validasyon veri kümesi üzerinde art arda iki kez kayıp değeri değişmez ise öğrenme oranının iki ile bölünmesini sağlarken erken durdurma fonksiyonu, validasyon veri kümesi üzerinde art arda altı kez kayıp değeri değişmez ise maksimum döngü sayısına ulaşmadan model eğitiminin bitirilmesini sağlamaktadır. Model eğitimi tamamlandıktan sonra modellerin performanslarını gözlemlemek için test verisi üzerinde doğruluk (accuracy), kesinlik (precision), duyarlılık (recall) ve roc eğrisi altında kalan alan (auc), değerleri hesaplanmıştır. Tablo 3’te her bir veri seti için hesaplanmış olan değerler gösterilmektedir.

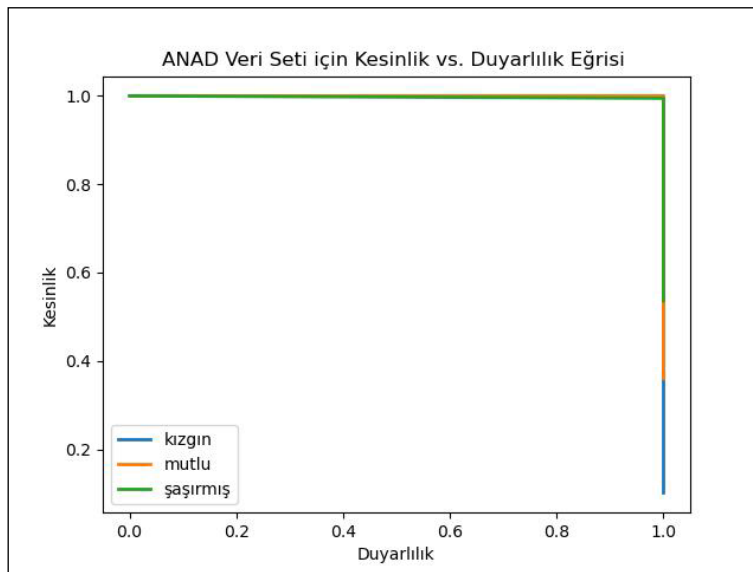
Tablo 3: Optimum Hiper-Parametreler Kullanılarak Eğitilmiş Modellerin Performans Ölçüt Skorları

Veri Seti Adı	Doğruluk	Kesinlik	Duyarlılık	AUC
TurEV-DB	%77,82	%78,01	%76,95	%84,97
ANAD	%99,71	%99,82	%99,04	%99,68
BanglaSER	%59,55	%59,08	%61,17	%74,98
TESS	%98,90	%98,80	%98,77	%99,34

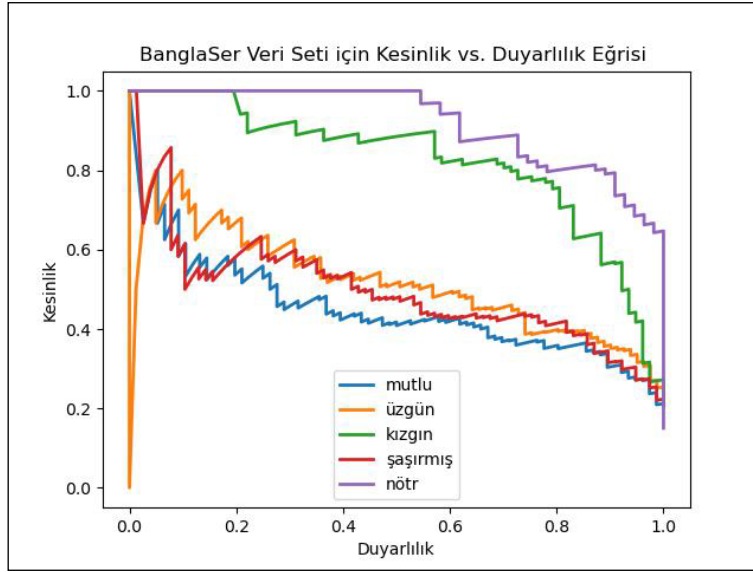
Tablo 3'te yer alan sonuçlar incelendiğinde en yüksek doğruluğun ANAD veri setinde, en düşük doğruluğun ise BanglaSER veri setinde elde edildiği görülmüştür. Özellikle BanglaSER veri setinde doğruluk oranları diğer veri setlerine göre düşük olarak gözlemlenmiştir. Model performanslarının daha iyi değerlendirilebilmesi için sınıf bazında performans skorlarına da bakılması büyük önem arz etmektedir. Bu bağlamda her bir veri seti için sınıf bazında kesinlik-duyarlılık grafikleri çizilmiştir. Şekil 2-5 arasında çizilen bu grafikler gösterilmektedir.



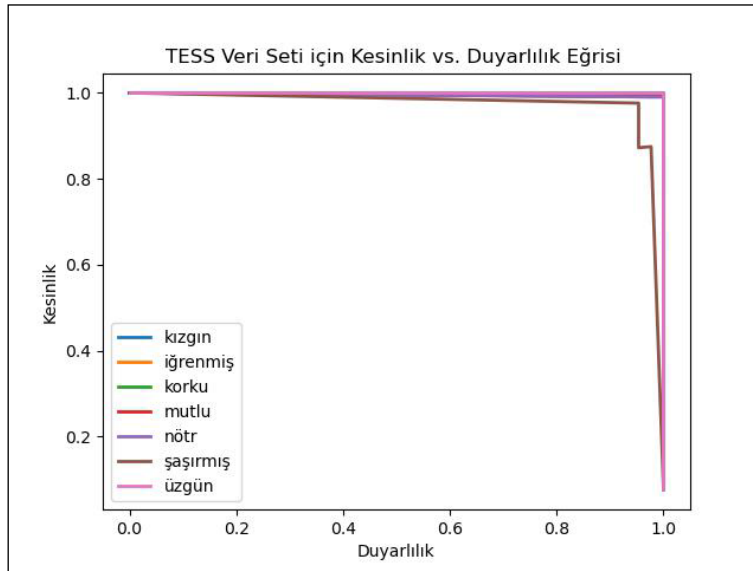
Şekil 2: TurEV-DB Veri Seti için Kesinlik vs Duyarlılık Eğrisi



Şekil 3: ANAD Veri Seti için Kesinlik vs Duyarlılık Eğrisi



Şekil 4: BanglaSER Veri Seti için Kesinlik vs Duyarlılık Eğrisi



Şekil 5: TESS Veri Seti İçin Kesinlik vs Duyarlılık Eğrisi

Şekil 2’de yer alan grafik incelendiğinde, TurEV-DB veri seti için en iyi eğrinin “üzgün”, en kötü eğrinin ise “sakin” sınıf ile elde edildiği gözlemlenmiştir. “Kızgın” ve “mutlu” sınıflarına ait eğrilerin ise birbirine yakın bir çizge takip ettiği görülmektedir. Şekil 3’te yer alan grafik incelendiğinde ANAD veri set için tüm sınıfların neredeyse mükemmel bir eğriye sahip olduğu görülmektedir. Şekil 4’te yer alan grafik incelendiğinde BanglaSER veri seti için “nötr” ve “kızgın” veri sınıflarına ait eğrilerin diğer sınıflara göre daha iyi bir performansa sahip olduğu görülmektedir. Şekil 5’te yer alan grafik incelendiğinde TESS veri seti için “şaşırmış” sınıfı hariç tüm sınıfların neredeyse mükemmel bir eğriye sahip olduğu, “şaşırmış” sınıfının ise diğerlerine göre biraz daha kötü bir eğriye sahip olduğu görülmektedir.

Elde edilen sonuçlar literatür ile karşılaştırıldığında TurEV-DB veri seti için literatürden daha kötü sonuç elde edilmiş, BanglaSER veri seti için literatür ile benzer fakat biraz daha iyi sonuçlar elde edilmiş, ANAD ve TESS veri seti için ise literatürden daha iyi sonuçlar elde edilmiştir. Bu bağlamda, çalışmamızda elde edilen sonuçlar BanglaSER veri seti için Aziz ve diğerleri tarafından veri çoğaltma yapılmadan (çalışmamızdakine benzer şekilde) elde edilen sonuçlar ile (Aziz et al., 2023), TESS veri seti için Sankara Pandiammal ve diğerleri tarafından yapılan çalışma ile (Sankara Pandiammal et al., 2024), ANAD veri seti için ise Albasan tarafından yapılan ve İsmail ve diğerleri tarafından yapılan çalışmalar (Alsabhan, 2023; Ismaiel et al., 2024) ile karşılaştırılmıştır. Model karşılaştırma sonuçları tablo 4'te gösterilmektedir.

Tablo 4: Önerilen Model ile Literatür Sonuçlarının Karşılaştırılması

Model Adı	Veri Seti Adı	Doğruluk
(Aziz et al., 2023)	BanglaSER	59,00%
(Sankara Pandiammal et al., 2024)	TESS	91,00%
(Alsabhan, 2023)	ANAD	96,72%
(Ismaiel et al., 2024)	ANAD	97,40%
Önerilen Model	BanglaSER	59,55%
Önerilen Model	TESS	98,90%
Önerilen Model	ANAD	99,71%

Tablo 4'te yer alan sonuçlar incelendiğinde önerilen modelin TESS ve ANAD veri setleri için en iyi sonuçları elde ettiği ve neredeyse hiç hata yapmadığı gözlemlenmiştir. BanglaSER veri seti için ise literatüre yakın ama literatürden daha iyi bir sonuç elde edildiği görülmektedir.

SONUÇLAR VE TARTIŞMA

Bu çalışmada, SER için U-Net tabanlı özgün bir derin öğrenme modeli önerilmiş ve modelin performansı Türkçe (TurEV-DB), İngilizce (TESS), Arapça (ANAD) ve Bangla (BanglaSER) dillerinden dört farklı veri seti kullanılarak değerlendirilmiştir. Önerilen model, CNN ve FLC yapılarını bir araya getiren U-Net mimarisini temel almakta olup, hiper-parametre optimizasyonu için Bayesian optimizasyon yöntemi kullanılmıştır. Deney sonuçları, modelin veri setlerine bağlı olarak %56,55 ile %99,71 arasında doğruluk oranları elde ettiğini göstermektedir. Deney sonuçları kapsamında elde edilen bu sonuçlar, modelin farklı dillerdeki genelleme kabiliyetini ve U-Net mimarisinin SER alanındaki potansiyelini ortaya koymaktadır.

Modelin performansı, veri setleri arasında belirgin farklılıklar göstermiştir. Özellikle ANAD veri setinde %99,71 doğruluk oranıyla neredeyse mükemmel bir performans elde edilmişken, TESS veri setinde %98,90 doğruluk oranıyla yüksek bir başarı sağlanmıştır. Bu

sonuçlar, literatürdeki diğer çalışmalara kıyasla önemli bir gelişme sunmaktadır. Önerilen modelin ANAD veri setindeki üstün performansı, U-Net mimarisinin spektrogram verilerindeki zaman-frekans özelliklerini etkin bir şekilde yakalayabildiğini ve Bayesian optimizasyonun hiper-parametre ayarlamasında sağladığı avantajları göstermektedir. Benzer şekilde.

Bununla birlikte, BanglaSER veri setinde elde edilen %59,55 doğruluk oranı, diğer veri setlerine kıyasla daha düşük bir performans sergilemiştir. BanglaSER veri setindeki düşük performansın, veri setinde var olan ses özelliklerinin benzerliği gibi faktörlerden kaynaklandığı ön görülmektedir. Özellikle Şekil 4'te gösterilen kesinlik-duyarlılık eğrileri, "nötr" ve "kızgın" sınıfların diğer sınıflara göre daha iyi performans sergilediğini ortaya koymaktadır. Bu durum, Bangla dilindeki duygu ifadelerinin tonlama ve ritim farklılıklarının model tarafından ayırt edilmesinde zorluklar yaratabileceğini düşündürmektedir. TurEV-DB veri setinde ise %77,82 doğruluk oranı elde edilmiş olup, bu sonuç literatürdeki diğer çalışmalara kıyasla daha düşük bir performans göstermiştir. Şekil 2'de sunulan kesinlik-duyarlılık eğrileri, "üzgün" sınıfının en iyi performansı sergilediğini, ancak "sakin" sınıfının daha düşük bir performans gösterdiğini ortaya koymaktadır. Bu durum, Türkçe dilindeki sakın duygu ifadelerinin spektrogram özelliklerinde diğer duygulara kıyasla daha az belirgin olabileceğini işaret etmektedir.

Önerilen U-Net tabanlı modelin başarısının ardındaki temel faktörler, mimarinin daraltıcı ve genişleyici yollarının spektrogram verilerindeki hem yerel hem de küresel desenleri etkin bir şekilde yakalayabilmesi ve Bayesian optimizasyonun hiper-parametre seçiminde sağladığı verimliliğidir. Özellikle, daraltıcı yolun özellik sıkıştırma kapasitesi ve genişleyici yolun özellik genişletme yeteneği, ses verilerindeki kısa ve uzun süreli olayların tanınmasında önemli bir avantaj sağlamaktadır. Bununla birlikte, modelin bazı veri setlerinde (özellikle BanglaSER ve TurEV-DB) daha düşük performans göstermesi, dil-spesifik özelliklerin ve veri seti boyutlarının modelin genelleme kabiliyetini etkilediğini göstermektedir.

Gelecek çalışmalar için önerilen modelin performansını artırmak amacıyla birkaç strateji önerilmektedir. İlk olarak, ses verileri gibi zaman serisi verilerinde sıklıkla başarılı sonuçlar veren uzun kısa vadeli bellek (LSTM) ağları ve Transformer tabanlı katmanlar modele entegre edilerek özellikle BanglaSER ve TurEV-DB veri setlerindeki performansın iyileştirilmesi hedeflenmektedir. İkinci olarak, veri çoğaltma teknikleri kullanılarak veri setlerinin boyutunun artırılmasının, modelin genelleme kabiliyetini güçlendirebileceği ön görülmektedir. Üçüncü olarak, farklı dillerden daha fazla veri seti ile modelin dil bağımsızlığı test edilerek, U-Net tabanlı mimarinin evrensel bir SER çözümü olarak potansiyelinin değerlendirilmesi amaçlanmaktadır.

TEŞEKKÜR

Bu çalışmada yer alan tüm nümerik hesaplamalar TÜBİTAK ULAKBİM, Yüksek Başarım ve Grid Hesaplama Merkezi'nde (TRUBA kaynaklarında) gerçekleştirilmiştir.

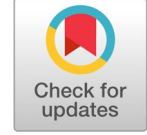
KAYNAKÇA

- Ahmad, J., Muhammad, K., Kwon, S., Baik, S. W., & Rho, S. (2016). Dempster-Shafer Fusion Based Gender Recognition for Speech Analysis Applications. *2016 International Conference on Platform Technology and Service (PlatCon)*, 1–4. <https://doi.org/10.1109/PlatCon.2016.7456788>
- Allen, J. B., & Rabiner, L. R. (1977). A unified approach to short-time Fourier analysis and synthesis. *Proceedings of the IEEE*, 65(11), 1558–1564. <https://doi.org/10.1109/PROC.1977.10770>
- Alsabhan, W. (2023). Human–Computer Interaction with a Real-Time Speech Emotion Recognition with Ensembling Techniques 1D Convolution Neural Network and Attention. *Sensors*, 23(3), Article 3. <https://doi.org/10.3390/s23031386>
- Altamimi, M., & Alayba, A. M. (2023). ANAD: Arabic news article dataset. *Data in Brief*, 50, 109460. <https://doi.org/10.1016/j.dib.2023.109460>
- Anvarjon, T., Mustaqeem, & Kwon, S. (2020). Deep-Net: A Lightweight CNN-Based Speech Emotion Recognition System Using Deep Frequency Features. *Sensors*, 20(18), Article 18. <https://doi.org/10.3390/s20185212>
- Aziz, S., Arif, N. H., Ahbab, S., Ahmed, S., Ahmed, T., & Kabir, Md. H. (2023). Improved Speech Emotion Recognition in Bengali Language using Deep Learning. *2023 26th International Conference on Computer and Information Technology (ICCIT)*, 1–6. <https://doi.org/10.1109/ICCIT60459.2023.10441053>
- Canpolat, S. F., Ormanoğlu, Z., & Zeyrek, D. (2020). Turkish Emotion Voice Database (TurEV-DB). In D. Beermann, L. Besacier, S. Sakti, & C. Soria (Eds.), *Proceedings of the 1st Joint Workshop on Spoken Language Technologies for Under-resourced languages (SLTU) and Collaboration and Computing for Under-Resourced Languages (CCURL)* (pp. 368–375). European Language Resources association. <https://aclanthology.org/2020.sltu-1.52>
- Das, R. K., Islam, N., Ahmed, Md. R., Islam, S., Shatabda, S., & Islam, A. K. M. M. (2022). BanglaSER: A speech emotion recognition dataset for the Bangla language. *Data in Brief*, 42, 108091. <https://doi.org/10.1016/j.dib.2022.108091>
- Ghai, M., Lal, S., Duggal, S., & Manik, S. (2017). Emotion recognition on speech signals using machine learning. *2017 International Conference on Big Data Analytics and Computational Intelligence (ICBDAC)*, 34–39. <https://doi.org/10.1109/ICBDACI.2017.8070805>

- Harár, P., Burget, R., & Dutta, M. K. (2017). Speech emotion recognition with deep learning. *2017 4th International Conference on Signal Processing and Integrated Networks (SPIN)*, 137–140. <https://doi.org/10.1109/SPIN.2017.8049931>
- Ismail, W., Alhalangy, A., Mohamed, A. O. Y., & Musa, A. I. A. (2024). Deep Learning, Ensemble and Supervised Machine Learning for Arabic Speech Emotion Recognition. *Engineering, Technology & Applied Science Research*, 14(2), Article 2. <https://doi.org/10.48084/etasr.7134>
- Issa, D., Fatih Demirci, M., & Yazici, A. (2020). Speech emotion recognition with deep convolutional neural networks. *Biomedical Signal Processing and Control*, 59, 101894. <https://doi.org/10.1016/j.bspc.2020.101894>
- Jha, T., Kavya, R., Christopher, J., & Arunachalam, V. (2022). Machine learning techniques for speech emotion recognition using paralinguistic acoustic features. *International Journal of Speech Technology*, 25(3), 707–725. <https://doi.org/10.1007/s10772-022-09985-6>
- Keras: Deep Learning for humans. (2024, July 20). <https://keras.io/>
- Kerkeni, L., Serrestou, Y., Mbarki, M., Raoof, K., Ali Mahjoub, M., & Cleder, C. (2020). Automatic Speech Emotion Recognition Using Machine Learning. In A. Cano (Ed.), *Social Media and Machine Learning*. IntechOpen. <https://doi.org/10.5772/intechopen.84856>
- Khan, M., Gueaieb, W., El Saddik, A., & Kwon, S. (2024). MSER: Multimodal speech emotion recognition using cross-attention with deep fusion. *Expert Systems with Applications*, 245, 122946. <https://doi.org/10.1016/j.eswa.2023.122946>
- Kotowski, K., Smolarczyk, T., Roterman-Konieczna, I., & Stapor, K. (2021). ProteinUnet—An efficient alternative to SPIDER3-single for sequence-based prediction of protein secondary structures. *Journal of Computational Chemistry*, 42(1), 50–59. <https://doi.org/10.1002/jcc.26432>
- Krishna, K. V., Sainath, N., & Posonia, A. M. (2022). Speech Emotion Recognition using Machine Learning. *2022 6th International Conference on Computing Methodologies and Communication (ICCMC)*, 1014–1018. <https://doi.org/10.1109/ICCMC53470.2022.9753976>
- Lin, Y.-L., & Wei, G. (2005). Speech emotion recognition based on HMM and SVM. *2005 International Conference on Machine Learning and Cybernetics*, 8, 4898–4901 Vol. 8. <https://doi.org/10.1109/ICMLC.2005.1527805>
- Liu, Z.-T., Han, M.-T., Wu, B.-H., & Rehman, A. (2023). Speech emotion recognition based on convolutional neural network with attention-based bidirectional long short-term

- memory network and multi-task learning. *Applied Acoustics*, 202, 109178. <https://doi.org/10.1016/j.apacoust.2022.109178>
- Liu, Z.-T., Wu, M., Cao, W.-H., Mao, J.-W., Xu, J.-P., & Tan, G.-Z. (2018). Speech emotion recognition based on feature selection and extreme learning machine decision tree. *Neurocomputing*, 273, 271–280. <https://doi.org/10.1016/j.neucom.2017.07.050>
- Madanian, S., Chen, T., Adeleye, O., Templeton, J. M., Poellabauer, C., Parry, D., & Schneider, S. L. (2023). Speech emotion recognition using machine learning—A systematic review. *Intelligent Systems with Applications*, 20, 200266. <https://doi.org/10.1016/j.iswa.2023.200266>
- Mary Little Flower, T., Jaya, T., & Christopher Ezhil Singh, S. (2024). Data augmentation using a 1D-CNN model with MFCC/MFMC features for speech emotion recognition. *Automatika*, 65(4), 1325–1338. <https://doi.org/10.1080/00051144.2024.2371249>
- Mishra, D., & rawat, A. (2015). Emotion Recognition through Speech Using Neural Network. *International Journal of Advanced Research in Computer Science and Software Engineering (IJARCSSE)*, 5.
- Mishra, S. P., Warule, P., & Deb, S. (2024). Speech emotion recognition using MFCC-based entropy feature. *Signal, Image and Video Processing*, 18(1), 153–161. <https://doi.org/10.1007/s11760-023-02716-7>
- Mustaqeem, & Kwon, S. (2021). MLT-DNet: Speech emotion recognition using 1D dilated CNN based on multi-learning trick approach. *Expert Systems with Applications*, 167, 114177. <https://doi.org/10.1016/j.eswa.2020.114177>
- Mustaqeem, Sajjad, M., & Kwon, S. (2020). Clustering-Based Speech Emotion Recognition by Incorporating Learned Features and Deep BiLSTM. *IEEE Access*, 8, 79861–79875. IEEE Access. <https://doi.org/10.1109/ACCESS.2020.2990405>
- Nediyanchath, A., Paramasivam, P., & Yenigalla, P. (2020). Multi-Head Attention for Speech Emotion Recognition with Auxiliary Learning of Gender Recognition. *ICASSP 2020 - 2020 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, 7179–7183. <https://doi.org/10.1109/ICASSP40776.2020.9054073>
- Pichora-Fuller, M. K., & Dupuis, K. (2020). *Toronto emotional speech set (TESS)* [Dataset]. Borealis. <https://doi.org/10.5683/SP2/E8H2MF>
- Sankara Pandiammal, K., Karishma, S., Harine Sakthe, K., Manimaran, V., Kalaiselvi, S., & Anitha, V. (2024). Emotion Recognition from Speech – an LSTM approach with the TESS Dataset. *2024 5th International Conference on Innovative Trends in Information Technology (ICITIIT)*, 1–6. <https://doi.org/10.1109/ICITIIT61487.2024.10580351>

- Satt, A., Rozenberg, S., & Hoory, R. (2017). Efficient Emotion Recognition from Speech Using Deep Learning on Spectrograms. *Interspeech 2017*, 1089–1093. <https://doi.org/10.21437/Interspeech.2017-200>
- scikit-optimize: Sequential model-based optimization toolbox*. (Version 0.10.2). (2025). [Python; MacOS, Microsoft :: Windows, POSIX, Unix]. <https://scikit-optimize.readthedocs.io/en/latest/contents.html>
- Singh, P., Sahidullah, M., & Saha, G. (2023). Modulation spectral features for speech emotion recognition using deep neural networks. *Speech Communication*, 146, 53–69. <https://doi.org/10.1016/j.specom.2022.11.005>
- Snoek, J., Larochelle, H., & Adams, R. P. (2012). Practical Bayesian Optimization of Machine Learning Algorithms. *Advances in Neural Information Processing Systems*, 25. <https://proceedings.neurips.cc/paper/2012/hash/05311655a15b75fab86956663e1819cd-Abstract.html>
- Ververidis, D., & Kotropoulos, C. (2006). Emotional speech recognition: Resources, features, and methods. *Speech Communication*, 48(9), 1162–1181. <https://doi.org/10.1016/j.specom.2006.04.003>
- Wagner, J., Kim, J., & Andre, E. (2005). From Physiological Signals to Emotions: Implementing and Comparing Selected Methods for Feature Extraction and Classification. *2005 IEEE International Conference on Multimedia and Expo*, 940–943. <https://doi.org/10.1109/ICME.2005.1521579>
- Zhou, X., Guo, J., & Bie, R. (2016). Deep Learning Based Affective Model for Speech Emotion Recognition. *2016 Intl IEEE Conferences on Ubiquitous Intelligence & Computing, Advanced and Trusted Computing, Scalable Computing and Communications, Cloud and Big Data Computing, Internet of People, and Smart World Congress (UIC/ATC/ScalCom/CBDCCom/IoP/SmartWorld)*, 841–846. <https://doi.org/10.1109/UIC-ATC-ScalCom-CBDCCom-IoP-SmartWorld.2016.0133>
- Zhu, L., Chen, L., Zhao, D., Zhou, J., & Zhang, W. (2017). Emotion Recognition from Chinese Speech for Smart Affective Services Using a Combination of SVM and DBN. *Sensors*, 17(7), Article 7. <https://doi.org/10.3390/s17071694>



Türkiye'nin Bilgi Bilim Alanında Yapay Zeka Araştırmaları: Web of Science İndeksli Yayınların Bibliyometrik Analizi (1994-2024)

Okan Koç, Balıkesir Üniversitesi, Tıbbi Hizmetler ve Teknikler Bölümü, Doç. Dr.,
okan.koc@balikesir.edu.tr, 0000-0002-5356-5940

ÖZ

Bu çalışmada, Türkiye adresli ve Web of Science (WoS) veri tabanında indekslenen, bilgi bilimi alanında yapay zeka konulu 74 makale bibliyometrik yöntemlerle analiz edilmiştir. 1994–2024 yılları arasında kapsayan bu analiz kapsamında, Türkiye’de bilgi bilimi disiplini yapay zeka araştırmalarının zaman içerisindeki gelişimini değerlendirmek amacıyla; yazar iş birliği ağları, en çok atıf alan çalışmalar, yayınların dergi dağılımları, tematik odak noktaları ve araştırma eğilimleri gibi çeşitli boyutlar incelenmiştir.

Çalışmanın sonuçları, 2000’li yıllardan itibaren Türkiye merkezli bilgi bilimi yayınlarının artış gösterdiğini, özellikle son on yılda uluslararası iş birliklerinin ve atıf performansının önemli ölçüde arttığını göstermektedir. Bununla birlikte, analizler, en çok tercih edilen dergileri, sık kullanılan anahtar kelimeleri ve ilgili alandaki etkili yazarları belirlemiştir. Alandaki araştırma eğilimlerini ve akademik üretimin dinamiklerini daha iyi anlamak için bibliyometrik yöntemler kullanılmıştır.

Bu çalışmanın, bilgi bilimi alanında Türkiye’nin uluslararası bilimsel görünürlüğünü ve katkılarını anlamamıza yardımcı olacağı beklenmektedir. Araştırmacıların ve politika yapıcıların bu çalışmanın bulgularını kullanarak stratejik planlama yapabileceği düşünülmektedir.

Anahtar Kelimeler : Bilgi ve Belge Yönetimi, Bilgi Bilim, Yapay Zeka, Bibliyometrik Analiz

Artificial Intelligence Research in the Field of Information Science in Turkey: A Bibliometric Analysis of Web of Science Indexed Publications (1994–2024)

ABSTRACT

In this study, 74 articles addressing artificial intelligence within the field of information science, indexed in the Web of Science (WoS) database and affiliated with Turkey, were analyzed using bibliometric methods. Covering the period between 1994 and 2024, the analysis examines various dimensions such as the evolution of artificial intelligence research within Turkey’s information science discipline over time, author collaboration networks, the most highly cited works, journal distributions, thematic focuses, and research trends.

The results of the study reveal that, starting from the 2000s, information science publications originating from Turkey have shown a notable increase, with a significant rise in international collaborations and citation performance particularly over the past decade. Furthermore, the



analyses identified the most preferred journals, frequently used keywords, and influential authors within the field. Bibliometric methods were utilized to gain a deeper understanding of research trends and the dynamics of academic production in the domain.

It is anticipated that this study will contribute to a better understanding of Turkey's international scientific visibility and contributions in the field of information science. It is also considered that researchers and policymakers may utilize the findings of this study for strategic planning purposes.

Keywords : Information and Document Management, Information Science, Artificial Intelligence, Bibliometric Analysis

GİRİŞ

Yapay zeka (YZ) bilimi, son yıllarda dünya genelinde hızla gelişen bir araştırma alanı hâline gelmiştir. YZ teknolojileri; bilgi yönetimi, bilgi analitiği, bilgi keşfi ve bilgiye erişim gibi birçok uygulama alanında köklü değişimlere yol açmıştır (Russell & Norvig, 2016; LeCun, Bengio & Hinton, 2015). Türkiye'de de yapay zeka alanında yürütülen araştırmaların ve bilimsel çalışmaların önemli ölçüde arttığı gözlemlenmektedir. Bu bağlamda, Web of Science (WoS) veri tabanında indekslenen Türkiye adresli araştırmaların incelenmesi; bilimsel üretim süreçlerinin, uluslararası iş birliği düzeylerinin ve araştırma eğilimlerinin anlaşılması açısından önemli bilgiler sunma potansiyeline sahiptir.

Bibliyometrik analizler, bilimsel literatürün gelişimini anlamak, araştırma eğilimlerini ortaya koymak ve akademik iş birliklerini belirlemek amacıyla yaygın olarak kullanılan yöntemlerdir (Leydesdorff & Bornmann, 2011; Aghaei et al., 2013; Velmurugan, 2013). Bu yöntemler, belirli bir alandaki literatürün yapısal ve tematik dinamiklerini kavramak ve bu doğrultuda etkili araştırma politikaları geliştirmek açısından büyük önem taşımaktadır (Uzun & Öztürk, 2022). Bilgi bilimi alanında Türkiye'nin yapay zeka temelli çalışmalarını incelemek ise, ulusal bilim politikalarının yönlendirilmesine katkı sağlarken aynı zamanda ülkenin uluslararası bilimsel görünürlüğünü artırma açısından da önemli bir fırsat sunmaktadır.

Bu çalışma, 1994–2024 yılları arasında Web of Science (WoS) veri tabanında indekslenen ve Türkiye adresli bilgi bilimi alanındaki yapay zeka temalı 74 makaleyi bibliyometrik yöntemlerle incelemektedir. Araştırma kapsamında, Türkiye'de bilgi bilimi disiplini içinde yapay zeka çalışmalarının gelişimi; yazar iş birliği ağları, tematik odaklar, araştırma eğilimleri ve atıf performansları gibi çeşitli boyutlarda analiz edilmiştir. Çalışmanın temel amacı, Türkiye'nin bu alandaki bilimsel üretimini ve bu üretimin uluslararası düzeydeki görünürlüğünü ortaya koymaktır. Ayrıca, en çok tercih edilen dergiler, öne çıkan yazarlar ve sık kullanılan anahtar kelimelere ilişkin bulgular, bilgi bilimi alanındaki yapay zeka araştırmalarının temel dinamiklerini anlamada önemli katkılar sunmaktadır. Elde edilen bulguların, Türkiye'de yeni bilim politikalarının geliştirilmesine ve akademik iş birliklerinin

güçlendirilmesine katkı sağlayacağı öngörülmektedir. Ülkelerin bilimsel başarılarını artırmaları ve küresel ölçekte rekabet edebilmeleri açısından bilim politikaları stratejik bir önem taşımaktadır. Özellikle 2000'li yıllardan itibaren Türkiye'de, uluslararası atıf performansları ve araştırma iş birliklerine odaklanan çalışmaların sayısında belirgin bir artış gözlemlenmektedir. Bu eğilim, Türkiye'nin akademik alanda daha etkin bir konum benimsediğini ve bilimsel üretim süreçlerine daha fazla odaklandığını ortaya koymaktadır.

1. LİTERATÜRE BAKIŞ

Bilimsel üretim süreçlerinin nicel yönlerini analiz etmeye yönelik bibliyometrik çalışmalar, özellikle bilgi bilimi ve yapay zeka gibi hızla gelişen alanlarda, akademik etkinin değerlendirilmesi ve araştırma eğilimlerinin izlenmesi açısından önemli bir araç olarak kullanılmaktadır. Bu bölümde, bibliyometrik analiz yöntemlerinin literature dayalı olarak geçmiş olduğu süreç kronolojik bir çerçevede ele alınmaktadır.

Bibliyometrik yöntemlerin olanakları ve sınırlılıklarına ilişkin erken değerlendirmelerden biri Wallin (2005) tarafından yapılmıştır. Wallin, bibliyometrik analizlerin bilimsel kaliteyi ön plana çıkarmaktan ziyade, görünürlük ve etkiyi ölçtüğünü belirterek bu yöntemlerin sınırlılığına dikket çekmiştir. Wallin'in çalışması, bilimsel değerlendirme süreçlerinde bibliyometrik analizlerin nasıl daha etkili kullanılabileceğine dair önemli bir uyarı niteliğindedir.

Ding, Yan, Frazho ve Caverlee (2009), PageRank algoritması temelinde yazar eş-atıf ağlarını incelemiş ve bilimsel etki ölçümünde tek başına atıf sayısının yeterli olmayacağını, ağ içindeki merkeziyetin de önemli olduğunu ortaya koymuştur. İlgili çalışma, klasik atıf analizlerine alternatif olarak aynı zamanda daha nitelikli ve derinlemesine bağlantıları ortaya çıkaracak yöntemlerin bibliyometrik analizlerde uygulanabilirliğini göstermiştir.

Chang, Huang ve Lin (2015) ise bilgi ve belge yönetimi alanındaki araştırma konularının zaman içindeki evrimini anahtar kelime, bibliyografik eşleşme ve eş-atıf analizleri ile incelemiştir. Çalışma, "bilgi arama ve erişim" temalarındaki genel değişimin seyrini ortaya koymakla birlikte, "bibliyometrik analizin" önemini de ortaya koymuştur.

Aynı yıl yayımlanan bir diğer kapsamlı analizde Ellegaard ve Wallin (2015), bibliyometrik yöntemlerin bilimsel üretim üzerindeki etkisini incelemiş ve bu yöntemlerin farklı disiplinlerde yaygınlaşmasını, çok yazarlı, çok disiplinli yayınları öne çıkarmada etkili bir yöntem olarak kullanılması gerektiğini belirtmiştir. Ayrıca, disiplinler arası çalışmaların yazar ve atıf düzeyi açısından daha yüksek etki yarattığı vurgulanmıştır.

Liu, Yin, Liu ve Dunford (2015), yenilik sistemleri araştırmalarının bibliyometrik bir görselleştirmesini sunarak, alanın entelektüel yapısını ve evrimsel dönüşümünü CiteSpace aracılığıyla ortaya koymuşlardır. İlgili çalışma, bibliyometrik araçların sosyal bilimlerdeki uygulamalarına örnek teşkil etmekte, alana özgü iş birliği ağlarının nasıl ortaya çıkarılacağını

vurgulamakta ve araştırma alanlarının tematik kümelenmesini görsel olarak ortaya koymaktadır.

Tahamtan, Safipour Afshar ve Ahamdzadeh (2016), bibliyometrik analizler ile atıf alma oranlarını etkileyen faktörlerin tespit edilmesi, makale yapısı, dergi prestiji ve yazar geçmiş gibi etmenlerin etki düzeyinin ölçülmesinin bilimsel üretim açısından önemini ortaya koymuşlardır. İlgili çalışma, niceliksel göstergelerin ötesine geçerek atıf dinamiklerine dair daha karmaşık ilişkiler analiz edilmesinin önemini vurgulamaktadır.

Van Eck ve Waltman (2010) tarafından geliştirilen VOSviewer yazılımı, bilimsel yayınlara ilişkin haritalama sürecinde görselleştirmenin önemini ortaya koymuştur. Bu araç, özellikle anahtar kelime ağları, eş-atıf yapıları ve tematik kümelenmeleri analiz etmek için önemli bir metodolojik yenilik sunmuştur. Bibliyometrik analizlerin daha erişilebilir ve yorumlanabilir olmasına katkı sağlayan bu araç, literatür haritalama tekniklerinde yeni bir dönemi başlatmıştır.

Son olarak Donthu, Kumar, Mukherjee, Pandey ve Lim (2021), bibliyometrik analiz yöntemlerine yönelik kapsamlı bir rehber sunmuş ve performans analizi, bilimsel haritalama, atıf ve eş-atıf analizleri gibi teknikleri uygulamalı örneklerle detaylandırmıştır. Bu çalışma, araştırmacılara alan taraması ve stratejik yayın politikaları geliştirme açısından yön gösterici niteliktedir. YZ teknolojileri; bilgi yönetimi, bilgi analitiği, bilgi keşfi ve bilgiye erişim gibi birçok uygulama alanında bibliyometrik analizlerin etkinliğini artırarak araştırmacılara daha kapsamlı, hızlı ve stratejik değerlendirmeler yapma olanağı sunmaktadır.

2. AMAÇ VE YÖNTEM

2.1. Amaç

Bu çalışmanın amacı, 1994–2024 yılları arasında Web of Science (WoS) veri tabanında indekslenen ve Türkiye adresli bilgi bilimi alanındaki yapay zeka temalı akademik yayınları bibliyometrik yöntemlerle inceleyerek, bu alandaki bilimsel üretim eğilimlerini, iş birliği ağlarını, tematik odakları ve atıf performanslarını ortaya koymaktır. Elde edilen bulgular doğrultusunda, Türkiye'nin bilgi bilimi disiplini içinde yapay zeka alanındaki gelişiminin değerlendirilmesi ve bu gelişimin ulusal bilim politikaları ile uluslararası bilimsel görünürlük açısından nasıl bir potansiyel taşıdığına ışık tutulması hedeflenmektedir. Ayrıca, öne çıkan yazarlar, tercih edilen dergiler ve sık kullanılan anahtar kelimeler gibi göstergeler aracılığıyla, alandaki temel dinamiklerin belirlenmesi amaçlanmaktadır.

2.2. Yöntem

Bu çalışmada, bibliyometrik analizi yöntemi kullanılmıştır. Bibliyometrik analiz, bilimsel literatürün niceliksel ve niteliksel yönlerini incelemek, araştırma eğilimlerini

belirlemek ve bilimsel işbirliği yapılarını haritalamak için yaygın olarak kullanılan bir yöntemdir (Donthu, Kumar, Mukherjee, Pandey & Lim, 2021).

Araştırmanın temel aşamaları aşağıdaki gibi düzenlenmiştir;

- 1. Veri Toplama:** Araştırmanın veri seti, "bilgi bilimi" ve "yapay zeka" terimlerini içeren WoS veritabanından toplanmıştır. Veriler, 1994–2024 yılları arasında yayınlanan 74 makaleyi içerir.
- 2. Veri Analizi:** Toplanan veriler bibliyometrik analize tabi tutulmuş olup, analiz için R Studio'dan Bibliometrix sistemi kullanılmıştır. Bibliometrix aracılığıyla en etkili yayınlar ve en çok atıf alan çalışmalar belirlenmiş, çalışmalara yapılan atıf sayıları incelenmiş, kelime analizi noktasında tematik haritalar oluşturulmuştur.

3. BULGULAR

3.1. Genel Bilgi

Bu bölümde, bibliyometrik analiz doğrultusunda elde edilen bulgulara yer verilmiştir.



Görsel 1: Genel Değerlendirme

Bu çalışma, genel olarak 1994–2024 yılları arasındaki otuz yıllık bir dönemi kapsamaktadır. Bu sayede, bilgi bilimi alanında yapay zeka temelli araştırmaların tarihsel gelişimini değerlendirmek mümkün hâle gelmiştir. Analiz kapsamında 36 farklı kaynak kullanılmış olup, bu durum çalışmanın kapsamlı bir literatür taramasıyla desteklendiğini göstermektedir. İncelenen toplam 74 makale, bilgi bilimi literatüründeki yapay zeka odaklı akademik üretimin önemli bir örneklemini oluşturmaktadır.

Görsel 1’de yer alan bulgulara göre, yıllık yayın artış oranı %8,64 olarak hesaplanmıştır. Bu oran, yapay zeka konusunun bilgi bilimi disiplini içerisinde özellikle son yıllarda artan bir ilgiyle ele alındığını göstermektedir. Çalışmalara toplamda 194 yazar katkı sağlamış; bu durum, alandaki akademik üretimin hem yazar çeşitliliği hem de iş birliğine dayalı bir yapı sergilediğini ortaya koymuştur. İncelenen yayınların 12’si tek yazarlı olarak gerçekleştirilmiştir. Ayrıca, makalelerin %27’sinde uluslararası iş birliği tespit edilmiştir. Bu

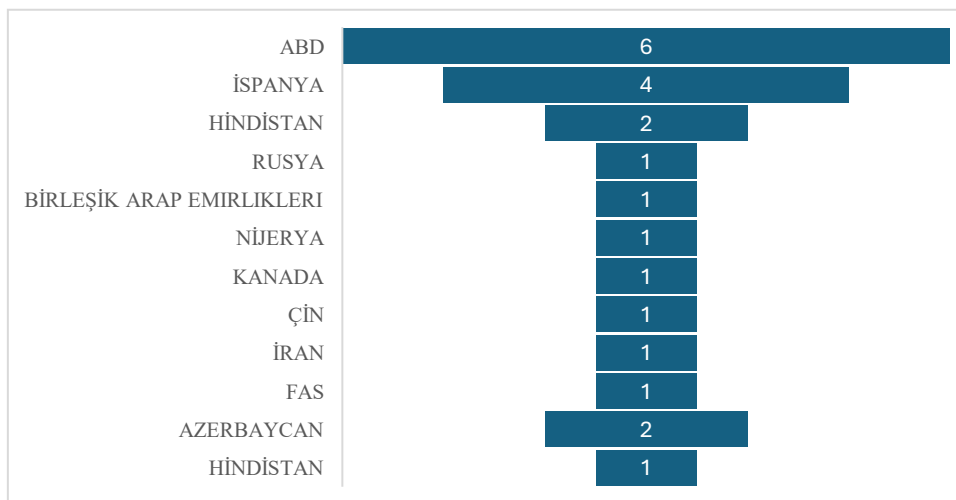
oran, Türkiye'nin bilgi bilimi alanında küresel ölçekte güçlü bağlantılara sahip olduğunu ve uluslararası akademik projelere açık bir konumda bulunduğunu göstermektedir.

Görsel 1'de yer alan bibliyometrik verilere göre, her bir belgeye ortalama 2,82 ortak yazar düşmektedir. Bu durum, bilgi bilimi alanında yapay zeka temelli bilimsel üretimin büyük ölçüde ekip çalışmasına dayandığını ve iş birliği düzeyinin yüksek olduğunu göstermektedir. Analizde kullanılan toplam 306 anahtar kelime, yapay zeka ve bilgi bilimi alanındaki çalışmaların geniş bir tematik çeşitlilik barındırdığını ortaya koymaktadır.

Ayrıca, analiz edilen 74 makalede toplam 3440 atıf referansının yer alması, çalışmanın literatürle güçlü bir bağ kurduğunu ve kapsamlı bir akademik zemine dayandığını göstermektedir. Belgelerin ortalama yaşı 6,18 yıl olarak hesaplanmıştır. Bu veri, analiz edilen çalışmaların büyük ölçüde güncel olduğunu ve mevcut araştırma eğilimlerini yansıttığını göstermektedir. Öte yandan, her bir yayının ortalama 28,04 atıf almış olması, bu çalışmaların akademik etki düzeyinin yüksek olduğunu ve alanda dikkate değer bir görünürlüğe sahip olduğunu ortaya koymaktadır.

Ayrıca, analiz edilen 74 makalede toplam 3440 atıf referansının yer alması, çalışmanın literatürle güçlü bir bağ kurduğunu ve kapsamlı bir akademik zemine dayandığını göstermektedir. Belgelerin ortalama yaşı 6,18 yıl olarak hesaplanmıştır. Bu veri, analiz edilen çalışmaların büyük ölçüde güncel olduğunu ve mevcut araştırma eğilimlerini yansıttığını göstermektedir. Öte yandan, her bir yayının ortalama 28,04 atıf almış olması, bu çalışmaların akademik etki düzeyinin yüksek olduğunu ve alanda dikkate değer bir görünürlüğe sahip olduğunu ortaya koymaktadır.

3.2. Yayınların Ülkelere Göre Dağılımı



Grafik 1: Yayınların ülkelere göre dağılımı

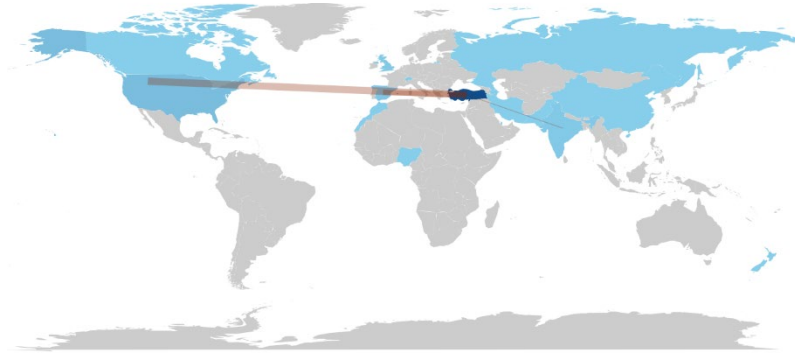
Bayesian optimizasyon, özellikle hesaplamalı olarak pahalı ve zaman alıcı fonksiyonları optimize etmek için kullanılmaktadır. Bayesian optimizasyon yaklaşımı, bir Grafik 1'de yer alan bulgulara göre, Türkiye adresli yapay zeka ve bilgi bilimi temalı yayınlarda iş birliği yapılan ülkeler arasında ilk sırada Amerika Birleşik Devletleri (ABD) yer almaktadır.

ABD (6 yayın): En fazla yayına sahip ülke ABD'dir. Bu, Amerika Birleşik Devletleri'nin bilgi bilimi ve yapay zeka alanlarında güçlü akademik altyapısına ve liderliğine işaret etmektedir. Amerika Birleşik Devletleri'nin önde gelen üniversiteleri ve araştırma merkezlerinin bu alanda önemli bir rol oynadığı düşünülmektedir.

İspanya (4 yayın): İspanya, Avrupa'daki bilgi bilimi araştırmalarında ikinci sırada yer almaktadır. Araştırma ve iş birliği fonlarının Avrupa Birliği tarafından sağlanmasının bu yaygınlığı artırdığı düşünülmektedir.

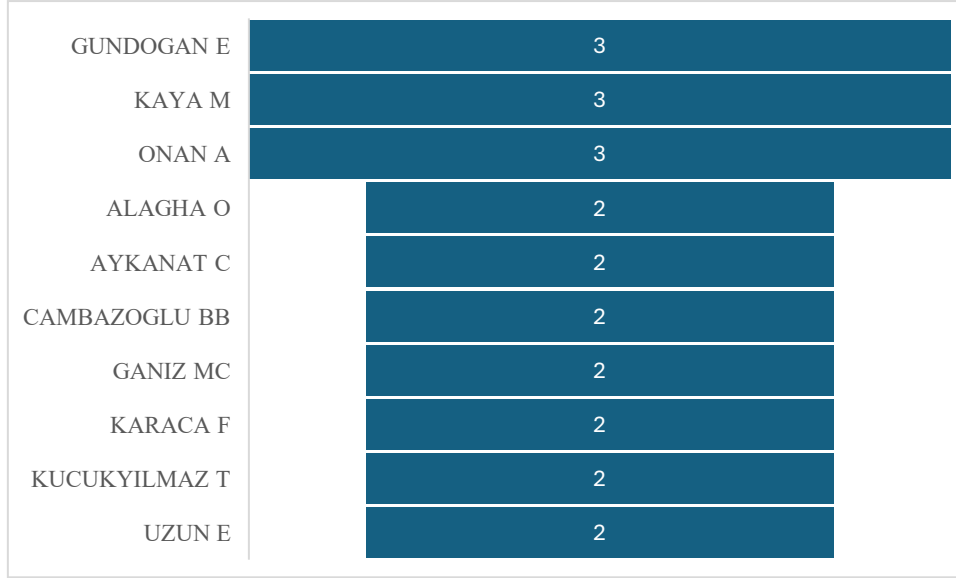
Hindistan (2 yayın): Hindistan, gelişmekte olan ülkelere bilgi bilimi ve yapay zeka alanlarında önemli katkılar sağlamıştır. Daha fazla teknoloji yatırımı ve geniş akademik ağların bu sonuca katkıda bulunduğu düşünülmektedir.

Diğer Ülkeler (1 yayın): Rusya, Birleşik Arap Emirlikleri, Nijerya, Kanada, Çin, İran, Fas ve Azerbaycan birer yayından oluşmaktadır. Bu ülkeler, bilgi bilimi ve yapay zekaya çok az katkıda bulunmalarına rağmen bu alanlarda önemli bir yere sahiptir.



Görsel 2: En Çok Birlikte Çalışılan Ülkeler

3.3. Yayınların Yazarlara Göre Dağılımı



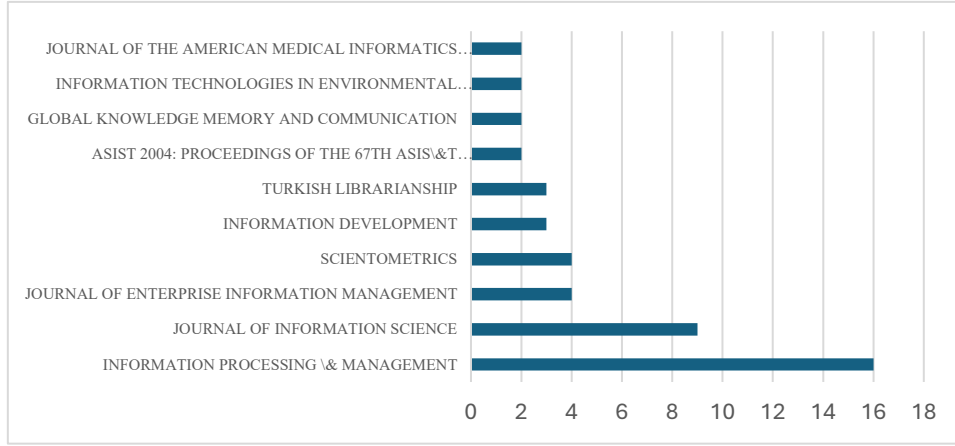
Grafik 2: Yayınların Yazarlara Göre Dağılımı

Grafik 2 incelendiğinde, bilgi bilimi alanında yapay zekâ konulu yayınlarda en fazla sayıda yayına sahip yazarlar olarak Gündoğan E., Kaya M. ve Onan A. öne çıkmaktadır. Bu bulgu, söz konusu araştırmacıların alanda aktif olarak çalıştıklarını ve yüksek düzeyde bilimsel üretkenliğe sahip olduklarını göstermektedir. İlgili yazarların araştırma profilleri, yapay zekâ teknolojilerinin bilgi yönetimi, metin madenciliği ve bilgi erişimi gibi alt alanlarına odaklanarak bilgi bilimi literatürüne önemli katkılar sunduklarını ortaya koymaktadır. Bu durum aynı zamanda Türkiye'de bu alanda uzmanlaşmış akademik odakların oluştuğunu da göstermesi açısından dikkat çekicidir.

Diğer Yazarlar (2 Yayın): Alagha O., Aykanat C., Cambazoğlu B.B., Ganiz M.C., Karaca F., Küçükylmaz T. ve Uzun E. Bu yazarlar da bilgi bilimi alanında gerçekleştirdikleri başarılı çalışmalarla bilimsel üretime önemli katkılarda bulunmuşlardır.

Grafik 2'ye göre, alandaki yayınların önemli bir kısmı ilk on yazarın katkılarıyla şekillenmiştir; bu durum, belirli araştırmacıların bilgi bilimi ve yapay zekâ kesişimindeki çalışmalarda öne çıktığını ve literatüre yön veren bir konumda bulunduklarını göstermektedir.

3.4. Yayınların Dergilere Göre Dağılımı



Grafik 3: Yayınların Dergilere Göre Dağılımı

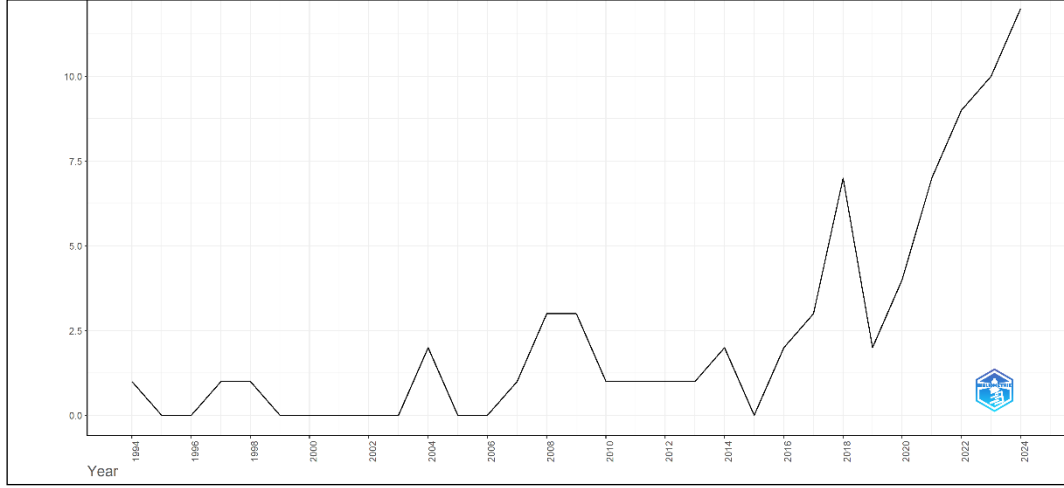
Grafik 3'e göre, incelenen 74 yayından 36'sının yalnızca 10 dergide yayımlandığı görülmektedir. Bu durum, söz konusu dergilerde yayımlanan çalışmaların toplam yayınların yaklaşık %48,6'sını oluşturduğunu göstermektedir.

Information Processing & Management: Analiz edilen yayınlar arasında 16 yayın ile en fazla tercih edilen dergi olarak öne çıkmaktadır. Bu durum, derginin bilgi yönetimi, bilgi işleme ve ilgili alanlarda önemli bir akademik etkiye sahip olduğunu göstermektedir. Yüksek yayın sayısı, bu derginin Türkiye'de alanında lider bir konumda olduğunu ve araştırmacıların çalışmalarını yayımlamak için sıklıkla bu dergiyi tercih ettiklerini ortaya koymaktadır.

Journal of Information Science: 8 yayın ile en çok yayın yapılan ikinci dergidir. İlgili dergi, bilgi bilimi alanındaki araştırmaları yayımlamak için etkili bir platform olarak öne çıkmaktadır. Araştırmacıların bu dergiyi tercih etme nedeninin, derginin alandaki geniş kapsama alanı ve sahip olduğu akademik prestij olduğu değerlendirilmektedir.

Bu iki dergide yayımlanan çalışmaların oranı toplamda %32,4 olarak hesaplanmıştır. Bu durum, ilgili alanda çalışan araştırmacıların büyük ölçüde bu iki dergiyi tercih ettiklerini ve söz konusu dergilerin, hedef kitle açısından yayın süreçlerinde kritik bir rol oynadığını göstermektedir. Bu tercihin, dergilerin akademik etki düzeyi ve alan içi görünürlüğüyle ilişkili olduğu düşünülmektedir.

3.5. Yayınların Yıllara Göre Dağılımı



Grafik 4: Yayınların Yıllara Göre Dağılımı

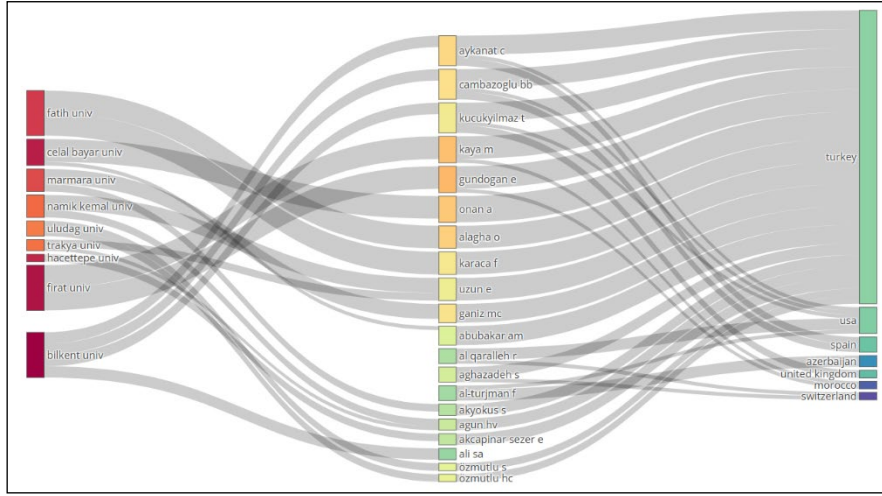
Grafik 4 değerlendirildiğinde bilimsel yayınların yıllara göre değişiminde üç temel dönem dikkat çekmektedir:

1990'lardan 2010'lara Kadar Düşük Üretim: 1990'ların başından 2010'lara kadar olan dönemde, bilimsel çalışmalara ayrılan kaynakların sınırlı olması veya odaklanılan araştırma alanlarının kısıtlı kalması nedeniyle yayın sayılarının oldukça düşük seviyelerde (genellikle 1–2 yayın) seyrettiği görülmektedir.

2010'dan Sonra Artış Başlangıcı: Konuya ilişkin yayın sayılarının 2010 yılından itibaren artmaya başladığı, özellikle 2016 yılından sonra belirgin bir yükseliş eğilimi gösterdiği tespit edilmiştir. Yayın sayısındaki bu artışın; yeni projelerin finanse edilmesi, teknolojik altyapının gelişmesi ve akademik iş birliklerinin artması gibi faktörlerden kaynaklandığı düşünülmektedir.

2018 ve Sonrası Hızlı Büyüme: 2018 yılından itibaren dikkat çekici bir artış ivmesi gözlemlenmekte olup, bu eğilimin 2024 yılına kadar hızlanarak devam ettiği anlaşılmaktadır. 2024 yılında en fazla makalenin yayımlandığı belirlenmiştir. Bu durum, bilimsel çalışmaların daha düzenli bir şekilde desteklendiğini ve ilgili alanlara yönelik ilgi ile kaynak tahsisinin arttığını göstermektedir.

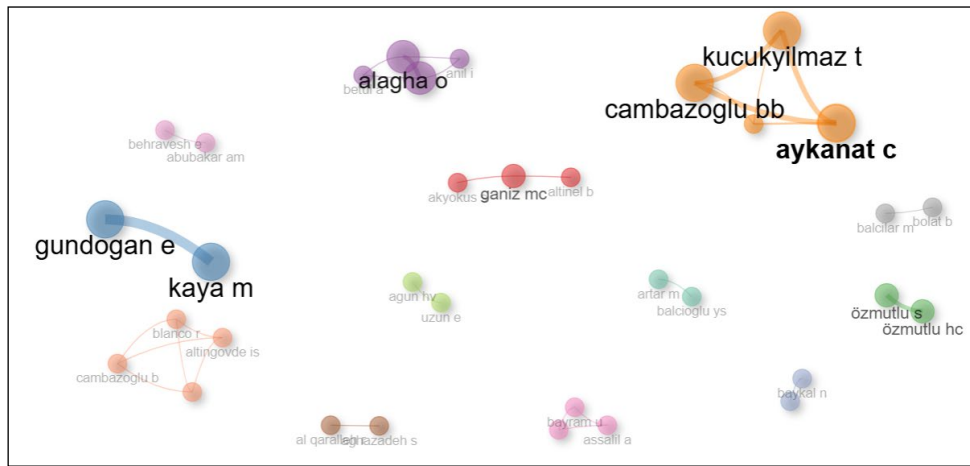
3.6. Yayınların Kurum, Yazar ve Ülkelere Göre Dağılımı



Grafik 5: Yayınların Kurum, Yazar ve Ülkelere Göre Dağılımı Veren Sankey Diyagramı

Yazarların büyük bir kısmının Fatih Üniversitesi, Celal Bayar Üniversitesi, Marmara Üniversitesi, Namık Kemal Üniversitesi, Hacettepe Üniversitesi ve Bilkent Üniversitesi gibi Türkiye'nin önde gelen üniversitelerinde görev yaptığı görülmektedir. Akyanat C., Kaya M., Gündoğan E., Karaca F. ve Al-Turjman F. gibi araştırmacıların ise farklı kurumlar ve ülkelerle iş birliği içinde çalıştıkları tespit edilmiştir. Yayınların, özellikle ABD, İspanya, Birleşik Krallık, İsviçre ve Azerbaycan gibi ülkelerle olan ilişkiler çerçevesinde uluslararası iş birliği kapsamında gerçekleştirildiği anlaşılmaktadır. İlgili diyagram, Türk üniversiteleri ve araştırmacılarının yapay zekâ alanındaki hem yerel hem de uluslararası bilimsel çalışmalarda aktif bir şekilde yer aldıklarını ortaya koymaktadır.

3.7. Yazarlar Arası İşbirliği



Grafik 6: Yazarlar Arası İş Birliği

Güçlü İş birliği Grupları:

Küçükıylmaz T, Cambazoğlu BB ve Aykanat C: Bu grup, iş birliği ağında dikkat çekici bir rol oynamaktadır. Üç yazar arasında güçlü bir bağ ve iş birliği olduğu görülmektedir.

Gundogan E ve Kaya M: Bu yazarların da birlikte çalıştığı görülmektedir. İki yazar arasındaki güçlü bağlantı, akademik ortaklığın etkin ve sürekli olduğunu göstermektedir.

Alagha O ve Betul A: Bu grup daha küçük bir araştırma grubunu temsil etmekte olup, iş birliği ağında daha az bağlantıya sahiptir.

Bağımsız Gruplar:

Görselde yer alan bazı küçük araştırma grupları (Özmutlu S. ve Özmutlu H.C.; Balcılar M. ve Bolat B.) yalnızca kendi aralarında çalışmaktadır. Bu durum, söz konusu çalışmaların aynı kurumda yürütüldüğünü veya belirli bir proje etrafında şekillendiğini düşündürmektedir.

Geniş İşbirliği Ağları Eksikliği:

Grafik 6, genel olarak küçük gruplar ve sınırlı bağlantılardan oluşan bir yapıyı ortaya koymaktadır. Çok sayıda yazarın yer aldığı geniş bir iş birliği ağına rastlanmamaktadır. Bu durum, yazarların çoğunlukla bireysel ya da küçük gruplar hâlinde çalıştığını veya belirli tematik alanlara yoğunlaştığını göstermektedir.

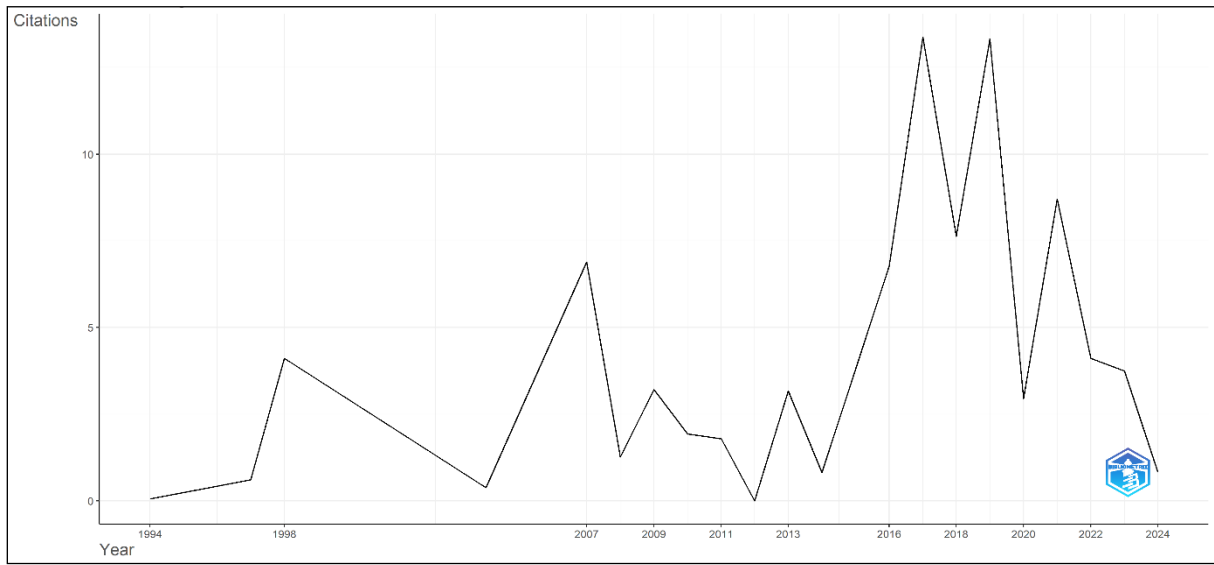
3.7. Yıllara Göre Atıf Analizi**Tablo 1: Yıllara Göre Atıf Analizi**

Yıl	Makaleye Düşen Ortalama Atıf	İlgili Yılda Yayımlanan Makale	Yıllık Ortalama Atıf	Bazda Atıf Yıl	Alınabilecek
1994	2,00	1	0,06	31	
1997	17,00	1	0,61	28	
1998	111,00	1	4,11	27	
2004	8,00	2	0,38	21	
2007	124,00	1	6,89	18	
2008	21,67	3	1,27	17	
2009	51,33	3	3,21	16	
2010	29,00	1	1,93	15	
2011	25,00	1	1,79	14	
2012	0,00	1	0,00	13	

Yüksek Atıf Alınan Yıllar: Tablo 1'de görüldüğü üzere, en fazla atıf alan makale 2007 yılında yayımlanmış olup, makale başına ortalama 124 atıf almıştır. Bu durum, ilgili

çalışmanın bilgi bilimi ve yapay zekâ alanında önemli bir etki yarattığını göstermektedir. Benzer şekilde, 1998 yılında yayımlanan makale 111 atıf/makale ortalamasına, 2009 yılında yayımlanan makale ise 51,33 atıf/makale ortalamasına sahiptir. Bu veriler, söz konusu yıllarda yayımlanan çalışmaların alanda yüksek düzeyde akademik görünürlüğe ve etkiye sahip olduğunu ortaya koymaktadır. Özellikle 2007 yılı, 6,89 atıf ortalaması ile diğer yıllara kıyasla açık ara öne çıkmaktadır.

Çalışmada ayrıca, en çok atıf alan ilk 10 makale Tablo 1'de detaylı olarak listelenmiştir. Bu durum, az sayıda çalışmanın bile bilgi bilimi ve yapay zekâ alanındaki akademik etkiyi önemli ölçüde şekillendirebildiğini göstermektedir.

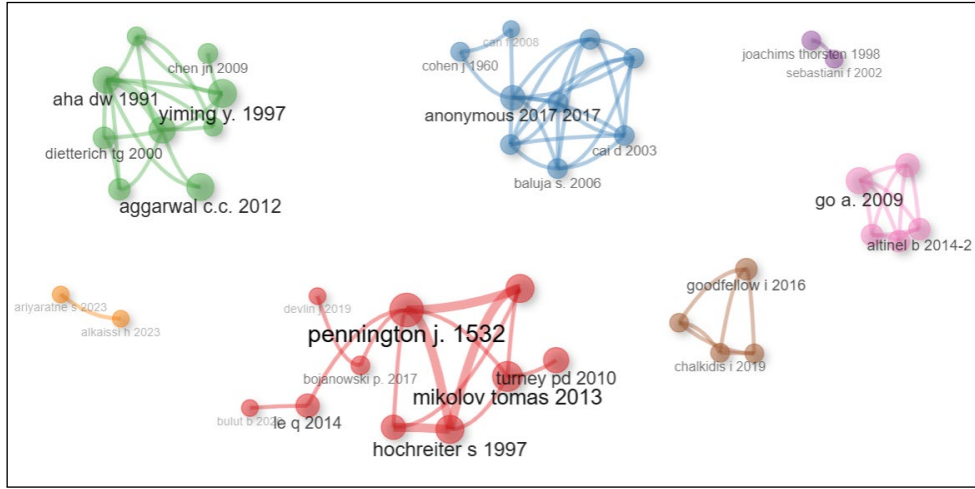


Grafik 7: Yıllara Göre Ortalama Atıf

Atıf Sayılarında Dalgalanma:

Yıllara göre makalelere yapılan ortalama atıf sayısını gösteren grafikte, 1990'lı yıllardan 2010'lara kadar genel olarak düşük bir atıf eğilimi dikkat çekmektedir. Özellikle 2007 ve 2018 yıllarında belirgin bir artış yaşanmış olsa da bu artışlar kısa süreli olup uzun vadeli bir yükseliş trendi oluşturmamıştır.

2007 yılı civarında ortalama atıf sayısında dikkat çekici bir artış yaşanmış ve bu dönemde yayımlanan çalışmaların daha fazla ilgi gördüğü söylenebilir. Benzer şekilde, 2018 yılı ortalama atıf sayısının en yüksek olduğu dönemdir. Bu durum, özellikle bu yılda yayımlanan çalışmanın geniş bir etki alanına sahip olduğunu göstermektedir. Ancak 2020 sonrası dönemde ortalama atıf sayısında belirgin bir düşüş gözlemlenmiştir. Bu düşüşün nedeni, güncel çalışmaların henüz yeterli sayıda atıf almamış olması olabilir (Grafik 7).

**Grafik 8: Yayınlar Arası Atıf İlişkisi**

Grafik 8'de, akademik yayınlar arasındaki atıf ilişkileri ve kümelenmeler görselleştirilmektedir. Her nokta veya düğüm bir yayını temsil ederken, düğümler arasındaki bağlantılar bu yayınlar arasındaki atıf ilişkilerini göstermektedir. Renkler ise, yayınların belirli tematik gruplara ayrıldığını ve kendi içinde daha yoğun atıf ilişkileri bulunan kümelerin oluştuğunu ortaya koymaktadır.

Yeşil Küme (Aha DW 1991, Yiming Y 1997, Dietterich TG 2000): Bu küme, veri madenciliği ve makine öğrenimi ile ilgili kapsamlı çalışmaları içermektedir.

Kırmızı Küme (Pennington J 1532, Mikolov Tomas 2013, Hochreiter S 1997): Bu grup, derin öğrenme ve doğal dil işleme (NLP) ile ilgili araştırmalara odaklanmıştır.

Mavi Küme (Anonymous 2017, Cohen J 1960, Baluja S 2006): Bu grup istatistiksel öğrenmeye odaklanmıştır.

Pembe Küme (Go A 2009, Altinel B 2014): Bu grup duygu analizi ve sosyal medya üzerine çalışmalar yapmıştır.

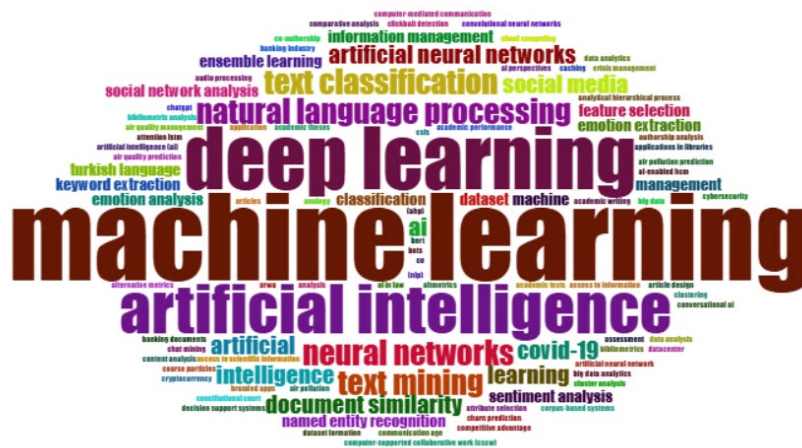
Kahverengi Küme (Goodfellow I 2016, Chalkidis I 2019): Derin öğrenme alanında ve diğer alanlarda benzer çalışmalar yapmıştır.

Bilimsel alandaki disiplinler arası ilişkileri ve belirli araştırmaların etkisini anlamada, yayınlar arası atıf ilişkileri güçlü bir analiz aracı olarak kullanılmaktadır. İlgili görsel, Hochreiter (1997), Mikolov (2013) ve Pennington (2013) gibi çalışmaların bilgi bilimi alanında öne çıkan etkili yayınlar olduğunu ortaya koymaktadır.

3.7. Anahtar Kelime Analizi

Araştırmaya dâhil edilen makalelerde kullanılan 306 anahtar kelime analiz edilmiştir. Anahtar kelime bulutu, bir çalışmanın temel odak noktalarını görsel olarak sunan ve bu odakları hızlı biçimde kavramaya yardımcı olan etkili bir görselleştirme aracıdır. Bu tür

görseller, yalnızca öne çıkan terimleri değil, aynı zamanda kelimeler arasındaki ilişkileri ve birlikte kullanım sıklıklarını da ortaya koyarak, alanın yapısal dinamiklerine dair daha derinlemesine bilgiler sunmaktadır. Bu bağlamda, anahtar kelimeler arasındaki ilişkilerin analiz edilmesi, araştırma eğilimlerinin, kavramsal kümelerin ve tematik yönelimlerin anlaşılmasında kritik bir rol oynamaktadır. Anahtar kelime bulutunun, veri kaynaklarının kapsamını ve disiplinler arası etkileşimleri ortaya koymak açısından önemli bir araç olduğu düşünülmektedir.



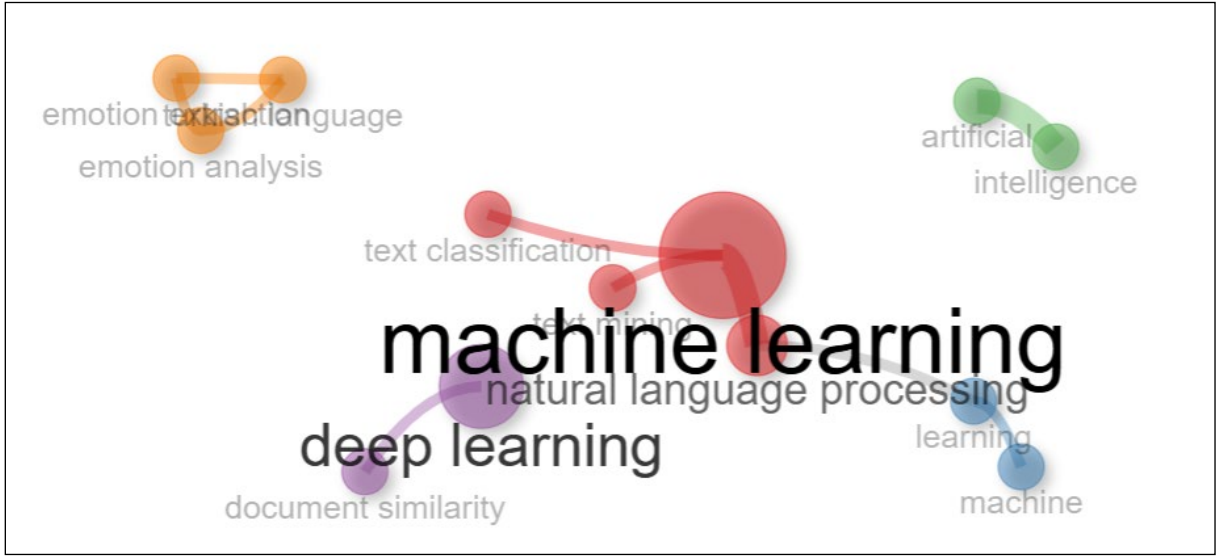
Görsel 3: Anahtar Kelime Bulutu

Görsel 3'e göre, en sık kullanılan anahtar kelimeler arasında doğal dil işleme, derin öğrenme, makine öğrenimi ve yapay zekâ öne çıkmakta olup, bu terimlerin çalışmalarda temel konular olarak ele alındığı görülmektedir.

Machine Learning (Makine Öğrenimi): Çalışmalarda en çok tercih edilen anahtar kelimedir ve araştırmaların büyük ölçüde bu kavram etrafında şekillendiğini göstermektedir.

Deep Learning (Derin Öğrenme): Makine öğrenimi ile yakından ilişkili olan derin öğrenme kavramı da sıklıkla kullanılan anahtar kelimelerden biridir.

Artificial Intelligence (Yapay Zekâ): Yapay zekâ ifadesi de anahtar kelimeler arasında öne çıkmakta olup, çalışmalarda yaygın bir şekilde tercih edildiği görülmektedir.



Grafik 9: Birlikte En Çok Tercih Edilen Anahtar Kelimeler

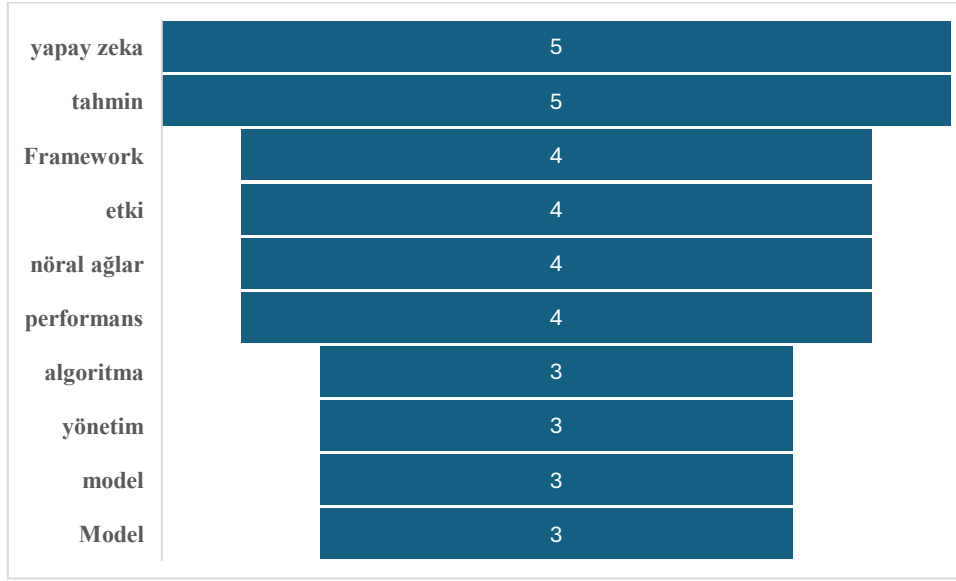
Grafik 9'da, anahtar kelimeler ve kavramlar arasındaki ilişkileri gösteren bir bağlantı haritası sunulmaktadır. Haritada her bir düğüm bir anahtar kelimeyi temsil ederken, düğümler arasındaki çizgiler bu kelimeler arasındaki kavramsal ve bağlamsal bağlantıları göstermektedir. Ayrıca, düğüm boyutları kelimenin literatürdeki kullanım sıklığını ve göreceli önemini yansıtmaktadır.

Görselleştirmeye göre, "Machine Learning" (Makine Öğrenimi) en merkezi ve en büyük düğüm olarak öne çıkmaktadır. Bu durum, analiz edilen çalışmaların büyük ölçüde makine öğrenimi konusuna odaklandığını ve bu kavramın bilgi bilimi alanında temel bir yapı taşı olarak kullanıldığını göstermektedir.

"Deep Learning" (Derin Öğrenme) kavramının, "Machine Learning" ile doğrudan ve güçlü bir bağlantı içinde olduğu görülmektedir. Bu ilişki, derin öğrenmenin makine öğreniminin bir alt dalı olarak değerlendirildiğini ve çoğu çalışmada bu şekilde ele alındığını düşündürmektedir.

Benzer şekilde, "Natural Language Processing" (Doğal Dil İşleme) terimi de hem "Machine Learning" hem de "Deep Learning" ile doğrudan bağlantılıdır. Bu durum, doğal dil işlemenin bu iki teknolojiye yoğun şekilde yararlandığını ve yapay zekâ uygulamalarında sıklıkla birlikte kullanıldıklarını göstermektedir.

Tüm bu ilişkiler, bilgi bilimi alanında yapay zekâ temelli çalışmaların, birbirine sıkı bağlı kavramsal yapılar üzerinden şekillendiğini ve belirli teknoloji kümeleri etrafında yoğunlaştığını düşündürmektedir.



Grafik 10: Başlıkta En Sık Geçen Kelimeler

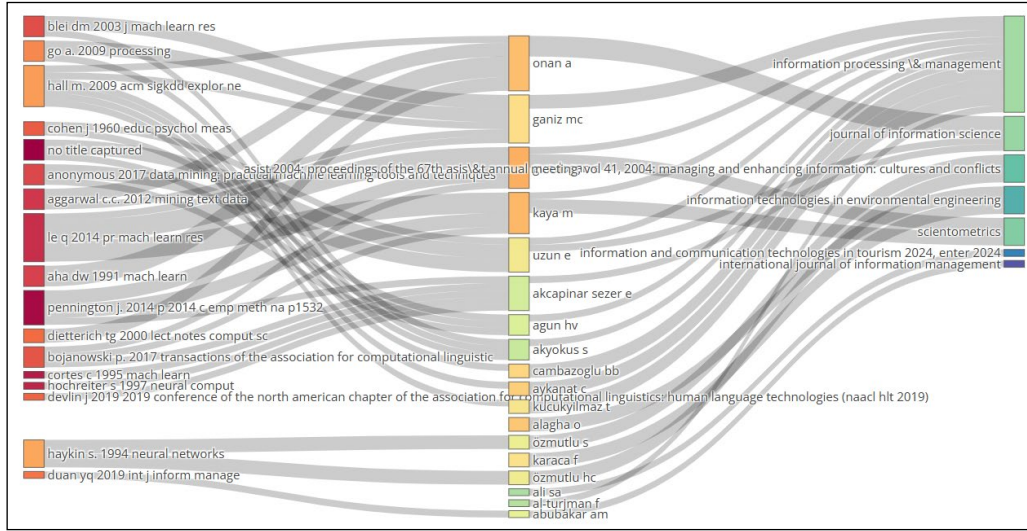
Grafik 10'da, akademik çalışmaların başlıklarında en sık kullanılan kelimeler ve bu kelimelerin tekrar sayıları gösterilmektedir. Elde edilen veriler, araştırmaların hangi konulara daha fazla önem verdiğini anlamak açısından önemlidir.

"Yapay zeka" (5 kez): Başlıklarda en yaygın kullanılan kelime "yapay zekâ"dır. Bu durum, çalışmaların yapay zekâyı birincil odak noktası olarak ele aldığını göstermektedir.

"Tahmin" (5 kez): "Tahmin" kelimesi, başlıklarda sık tercih edilen bir diğer terim olarak öne çıkmaktadır. Bu durum, çalışmaların büyük ölçüde modelleme odaklı olduğunu ve bu yaklaşımın doğrudan yapay zekâ uygulamalarıyla ilişkili olduğunu düşündürmektedir.

3.8. Kaynakça Analizi

Akademik çalışmalar arasındaki ilişkiyi belirlemek ve hangi kaynaklara hangi yazarların atıfta bulunduğunu tespit etmek için kaynakça analizi kullanılmıştır. Kaynakça analizi, çalışmanın hangi dergilerde yayınlandığını da gösterebilir. Aşağıdaki görsel, akademik yayınlardaki etkileşimlerin bir haritasını içermektedir.

**Grafik 11: Kaynakça Analizi**

Anahtar Kaynaklar: Grafik 11'e bakıldığında, bazı kaynakların diğerlerine göre daha fazla bağlantı içerdiği görülmektedir. Örneğin, birçok yazarın Pennington J. (2014) ve Dietterich T.G. (2000) gibi çalışmalara sıklıkla atıfta bulunduğu tespit edilmiştir. İlgili literatürde bu kaynakların oldukça önemli referanslar olarak kabul edildiği düşünülmektedir.

Yazarların Yayınlarla İlişkisi: Bazı yazarların belirli yayınlarla güçlü bağlantılar kurduğu gözlemlenmektedir. Örneğin, Onan A.'nın Pennington J. (2014) gibi çalışmalara sık sık atıfta bulunduğu görülmektedir.

Dergi Dağılımı: Grafik aynı zamanda atıf yapılan çalışmaların yayımlandığı dergileri de göstermektedir. Atıfta bulunulan kaynaklar arasında Journal of Information Science dergisinin önemli bir konumda yer aldığı anlaşılmaktadır. Ayrıca, Information Processing & Management ve Scientometrics dergileri de sıklıkla başvurulanan diğer yayın organları olarak öne çıkmaktadır. Bu durum, söz konusu dergilerin alandaki araştırmalarda etkili ve yönlendirici bir rol oynadığını göstermektedir.

ANALİZ VE NİHAİ DEĞERLENDİRME

Bu çalışma, 1994'ten 2024'e kadar Türkiye'de yapay zekâ ile ilgili akademik yayınların bilimsel gelişimini ve akademik etkisini incelemektedir. Elde edilen bulgular, Türkiye'nin özellikle son yıllarda bilgi bilimi ve yapay zekâ alanlarında önemli bir üretim artışı yaşadığını ortaya koymaktadır. Uluslararası iş birliğinde gözlenen artış ve akademik etkisi yüksek yayınların varlığı, Türkiye'nin bu alandaki küresel bilimsel görünürliğini güçlendirmektedir.

Yayınların dergi dağılımı, atıf performansı, yazarlar arası iş birliği ve anahtar kelime analizi gibi parametreler, bilgi bilimi kapsamında yapay zekâ konusuna duyulan ilginin

giderek arttığını ve bu alanda gelişmekte olan bir akademik altyapının oluştuğunu göstermektedir. Konuya ilişkin bibliyometrik analiz sonuçları, yürütülen çalışmaların bilimsel etkisinin ne düzeyde olduğunu ve bu etkinin zamanla nasıl şekillendiğini ortaya koymaktadır.

Bu doğrultuda, konuyla ilgili tavsiyeleri aşağıdaki gibi sıralamak mümkündür;

- Türkiye'deki akademisyenlerin, bilgi bilimi ve yapay zekâ alanındaki uluslararası iş birliklerini artırmak ve bilimsel görünürlüklerini güçlendirmek amacıyla daha fazla uluslararası akademik etkinlikte yer almaları gerekmektedir.
- Ulusal çalışmaların genişletilmesi ve uluslararası kongre, sempozyum ve toplantılara katılımın teşvik edilmesi gerekmektedir.
- Bilgi bilim ve yapay zeka alanında yüksek kaliteli dergilere yönelik yayın sayısını artırmak için yüksek kaliteli araştırma fonları ve teşvik programları oluşturulması gerekmektedir.
- Anahtar kelime analizinde elde edilen bulgular, araştırmaların büyük ölçüde teknik kavramlar—özellikle derin öğrenme, duygu analizi ve doğal dil işleme (NLP)—etrafında yoğunlaştığını göstermektedir. Ancak, yapay zekâ etiği ve toplumsal etkiler gibi daha bütüncül ve eleştirel perspektiflerin literatürde yeterince yer bulamadığı dikkat çekmektedir. Yapay zekâ teknolojilerinin etik boyutları, sosyal sonuçları ve politik etkileri gibi konuların, bilgi bilimi alanında daha fazla çalışılması gerektiği düşünülmektedir.
- Yapay zeka ve bilgi bilimi alanındaki genç araştırmacılar için proje bazlı burslar ve eğitim programları oluşturulmalıdır.
- Bilgi bilim ve yapay zekâ araştırmalarının dünya çapında tanınmasını sağlamak için açık erişim yayıncılığa daha fazla önem verilmelidir. Bilimsel etkinliklerin (çalıştaylar, paneller vb.) sayısındaki artışın bir sonucu olarak, bu etkinliklerin bulguları uluslararası platformlarda paylaşılmalıdır.
- Yapay zekâ ve bilgi bilimi alanındaki ulusal ve uluslararası eğilimleri göz önünde bulundurarak, bilim politikaları güncellenmelidir. Bilgi bilimi alanında multidisipliner yaklaşımlar benimsenmeli ve diğer alanlarla entegre edilmelidir.

REFERENCES

- Aghaei Chadegani, A., Salehi, H., Yunus, M. M., Farhadi, H., Fooladi, M., Muniandy, B., & Farhadi, M. (2013). A comparison between two main academic literature collections: Web of Science and Scopus databases. *Asian Social Science*, 9(5), 18-26.
- Chang, Y.-W., Huang, M.-H., & Lin, C.-W. (2015). Evolution of research subjects in library and information science based on keyword, bibliographical coupling, and co-citation analyses. *Scientometrics*, 105(3), 2071–2087. <https://doi.org/10.1007/s11192-015-1762-8>
- Ding, Y., Yan, E., Frazho, A., & Caverlee, J. (2009). PageRank for ranking authors in co-citation networks. *Journal of the American Society for Information Science and Technology*, 60(11), 2229–2243. <https://doi.org/10.1002/asi.21171>
- 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. <https://doi.org/10.1016/j.jbusres.2021.04.070>
- Ellegaard, O., & Wallin, J. A. (2015). The bibliometric analysis of scholarly production: How great is the impact? *Scientometrics*, 105(3), 1809–1831. <https://doi.org/10.1007/s11192-015-1645-z>
- LeCun, Y., Bengio, Y., & Hinton, G. (2015). Deep learning. *nature*, 521(7553), 436-444.
- Leydesdorff, L., & Bornmann, L. (2011). Mapping (USPTO) patent data on US inventors: Patent analysis as a tool for measuring science and technology interaction. *Journal of the American Society for Information Science and Technology*, 62(4), 660-669. <https://doi.org/10.1002/asi.21434>.
- Liu, Z., Yin, Y., Liu, W., & Dunford, M. (2015). Visualizing the intellectual structure and evolution of innovation systems research: A bibliometric analysis. *Scientometrics*, 103(1), 135–158. <https://doi.org/10.1007/s11192-014-1517-y>
- Russell, S., & Norvig, P. (2016). *Artificial Intelligence: A Modern Approach* (3rd ed.). Pearson Education.
- Tahamtan, I., Safipour Afshar, A., & Ahamdzadeh, K. (2016). Factors affecting number of citations: A comprehensive review of the literature. *Scientometrics*, 107(3), 1195–1225. <https://doi.org/10.1007/s11192-016-1889-2>
- Uzun, K., & Öztürk, O. (2022). Bibliometric Analysis of Organizational Ecology Theory (OET): To Review Past for Directing the Future of the Field. *Ege Academic Review*, 22(2), 195-212.
- Van Eck, N. J., & Waltman, L. (2010). Software survey: VOSviewer, a computer program for bibliometric mapping. *Scientometrics*, 84(2), 523–538. <https://doi.org/10.1007/s11192-009-0146-3>

- Velmurugan, C. (2013). Bibliometric analysis with special reference to Authorship Pattern and Collaborative Research Output of Annals of Library and Information Studies for the Year 2007–2012. *International Journal of Digital Library Services*, 3(3), 13-21.
- Wallin, J. A. (2005). Bibliometric methods: Pitfalls and possibilities. *Basic & Clinical Pharmacology & Toxicology*, 97(5), 261–275. https://doi.org/10.1111/j.1742-7843.2005.pto_139.x

ABA

Akademik Biliřim Arařtırmaları Derneęi

Suadiye Mah. Kazım Özalp Sok. No:15 Kat:2

řařkınbakkal Kadıköy/İSTANBUL

Tel: 0216 355 56 19 • Fax: 0216 368 43 30

www.abilar.org