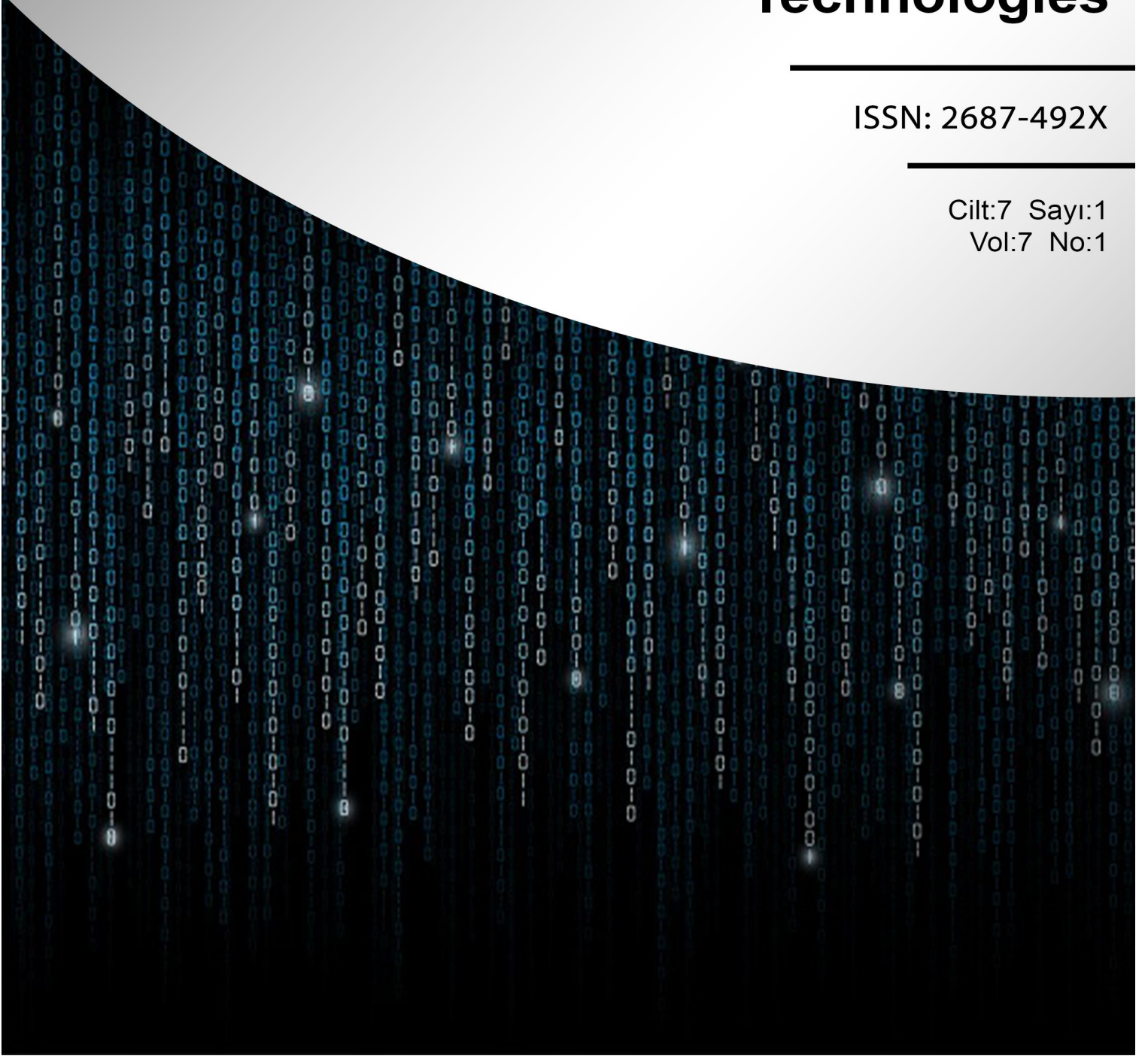


Bilgi ve İletişim Teknolojileri Dergisi

Journal of Information and Communication Technologies

ISSN: 2687-492X

Cilt:7 Sayı:1
Vol:7 No:1





BİLGİ VE İLETİŞİM TEKNOLOJİLERİ DERGİSİ

JOURNAL OF INFORMATION AND COMMUNICATION TECHNOLOGIES

ULUSLARARASI HAKEMLİ DERGİ / INTERNATIONAL REFEREED JOURNAL

Volume/Cilt: 7, Issue/Sayı: 1, 2025

Editor-in-Chief

Prof. Dr. Fatma Gizem KARAOĞLAN YILMAZ,
Bartın University

Associate Editor

Asst. Prof. Dr. Şeyma ÇAĞLAR ÖZHAN, Bartın
University

Editorial Board

Prof. Dr. Hafize KESER, Ankara University, Turkey
Prof. Dr. Hüseyin UZUNBOYLU, Near East
University, Turkish Republic of Northern Cyprus
Prof. Emeritus, James Lee MOSELEY, Wayne State
University, United States
Prof. Dr. Jesús García LABORDA, Alcalá University,
Spain
Prof. Dr. Piet KOMMERS, Twente University,
Netherlands
Prof. Dr. Ramazan YILMAZ, Bartın University, Turkey

Secretariat

Foreign Language and Pre-Review Specialists

Res. Asst. Rumeysa ERDOĞAN, Bartın University,
Turkey
Res. Asst. Fatma Rumeysa UÇAR, Bartın University,
Turkey

Publishing Preparation

Res. Asst. Rumeysa ERDOĞAN, Bartın University,
Turkey
Res. Asst. Fatma Rumeysa UÇAR, Bartın University,
Turkey

Technical Assistants

Res. Asst. Rumeysa ERDOĞAN, Bartın University,
Turkey
Res. Asst. Fatma Rumeysa UÇAR, Bartın University,
Turkey

Contact

Journal of Information and Communication
Technologies
e-mail: bilgiveiletisimdergisi@gmail.com

Editör

Prof. Dr. Fatma Gizem KARAOĞLAN YILMAZ,
Bartın Üniversitesi

Yardımcı Editör

Dr. Öğr. Üyesi Şeyma ÇAĞLAR ÖZHAN,
Bartın Üniversitesi

Editörler Kurulu (Yayın Kurulu)

Prof. Dr. Hafize KESER, Ankara Üniversitesi, Türkiye
Prof. Dr. Hüseyin UZUNBOYLU, Yakın Doğu
Üniversitesi, Kuzey Kıbrıs Türk Cumhuriyeti
Prof. Emeritus, James Lee MOSELEY, Wayne State
Üniversitesi, Birleşik Devletler
Prof. Dr. Jesús García LABORDA, Alcalá Üniversitesi,
İspanya
Prof. Dr. Piet KOMMERS, Twente Üniversitesi,
Hollanda
Prof. Dr. Ramazan YILMAZ, Bartın Üniversitesi,
Türkiye

Sekreteryä

Yabancı Dil ve Ön Hazırlık Sorumluları

Arş. Gör. Rumeysa ERDOĞAN, Bartın Üniversitesi,
Türkiye
Arş. Gör. Fatma Rumeysa UÇAR, Bartın Üniversitesi,
Türkiye

Yayıma Hazırlık

Arş. Gör. Rumeysa ERDOĞAN, Bartın Üniversitesi,
Türkiye
Arş. Gör. Fatma Rumeysa UÇAR, Bartın Üniversitesi,
Türkiye

Teknik Sorumlular

Arş. Gör. Rumeysa ERDOĞAN, Bartın Üniversitesi,
Türkiye
Arş. Gör. Fatma Rumeysa UÇAR, Bartın Üniversitesi,
Türkiye

İletişim

Bilgi ve İletişim Teknolojileri Dergisi
e-posta: bilgiveiletisimdergisi@gmail.com

Journal of Information and Communication
Technologies; is an **online, open access, free
international peer-reviewed** journal published in
Turkish or English.

Bilgi ve İletişim Teknolojileri Dergisi; araştırma ve
derleme çalışmalarını Türkçe veya İngilizce olarak
çevrimiçi yayımlanan, **açık erişime sahip, ücretsiz,
uluslararası hakemli** bir dergidir.

Index List / Dizin Listesi

Google Scholar, Index Copernicus, Asos Index, CiteFactor, J-Gate, ESJI Index, Directory of Research Journal
Indexing, Academic Resource Index, ROAD, Türk Eğitim İndeksi, Rootindexing, Journals Directory, Journal Factor,
International Servicesfor Impact Factor and Indexing (ISIFI), The Scientific Literature Database, Akademik
Dokümanlar Dizini (Index of Academic Documents [IAD])

BİLİM KURULU / EDITORIAL BOARD

- Prof. Dr. Apisak Bobby PUIPAT**, Thammasat Üniversitesi, Tayland
Prof. Dr. Cindy WALKER, Duquesne Üniversitesi, Pittsburgh, Birleşik Devletler
Prof. Dr. Ertuğrul USTA, Necmettin Erbakan Üniversitesi, Türkiye
Prof. Dr. Gary N. MCLEAN, Minnesota Üniversitesi, Minnesota, Birleşik Devletler
Prof. Dr. Hafize KESER, Ankara Üniversitesi, Türkiye
Prof. Dr. Halil YURDUGÜL, Hacettepe Üniversitesi, Türkiye
Prof. Dr. Huda AYYASH-ABDO, Lebanese American Üniversitesi, Lübnan
Prof. Dr. Hüseyin UZUNBOYLU, Yakın Doğu Üniversitesi, Kuzey Kıbrıs Türk Cumhuriyeti
Prof. Dr. Jesús García LABORDA, Alcalá Üniversitesi, İspanya
Prof. Dr. Lotte Rahbek SCHOU, Aarhus Üniversitesi, Danimarka
Prof. Dr. Michael K. THOMAS, Illinois Üniversitesi, Chicago, Birleşik Devletler
Prof. Dr. Michele BIASUTTI, Padova Üniversitesi, İtalya
Prof. Dr. Piet KOMMERS, Twente Üniversitesi, Hollanda
Prof. Dr. Rita Alexandra CAINÇO DIAS CADIMA, Polytechnic of Leiria, Portekiz
Prof. Dr. Rolf GOLLOB, Zürih Üniversitesi, İsviçre
Prof. Dr. Rosalina Abdul SALAM, Science Üniversitesi, Malezya
Prof. Dr. Saouma BOUJAOUDE, Beirut American Üniversitesi, Lübnan
Prof. Dr. Todd Alan PRICE, National Louis Üniversitesi, Illinois, Birleşik Devletler
Prof. Dr. Vinayagum CHINAPAH, Stockholm Üniversitesi, İsveç
Prof. Dr. Vladimir A. FOMICHOV, National Research Üniversitesi, Rusya
Doç. Dr. Agah Tuğrul KORUCU, Necmettin Erbakan Üniversitesi, Türkiye
Doç. Dr. Ctibor HATÁR, Constantine the Philosopher Üniversitesi, Slovakya
Doç. Dr. Fezile ÖZDAMLİ, Yakın Doğu Üniversitesi, Kuzey Kıbrıs Türk Cumhuriyeti
Doç. Dr. Hüseyin BİÇEN, Yakın Doğu Üniversitesi, Kuzey Kıbrıs Türk Cumhuriyeti
Doç. Dr. Ramazan YILMAZ, Bartın Üniversitesi, Türkiye
Doç. Dr. Tuğba ÖZTÜRK, Ankara Üniversitesi, Türkiye
Dr. Öğr. Üyesi Ahmet Berk ÜSTÜN, Bartın Üniversitesi, Türkiye
Dr. Öğr. Üyesi Barış SEZER, Hacettepe Üniversitesi, Türkiye
Dr. Öğr. Üyesi Hilal KAYA, Ankara Yıldırım Beyazıt Üniversitesi, Türkiye
Dr. Öğr. Üyesi Seyfullah GÖKOĞLU, Bartın Üniversitesi, Türkiye
Dr. Agnaldo ARROIO, São Paulo Üniversitesi, Brezilya
Dr. Ayşe Begüm ASLAN, Wayne State Üniversitesi, ABD
Dr. Chryssa THEMELIS, Lancaster Üniversitesi, İngiltere
Dr. Nurbiha A. SHUKOR, Malezya Teknoloji Üniversitesi, Malezya
Dr. Vina ADRIANY, Universitas Pendidikan Indonesia, Endonezya

CONTENT / İÇİNDEKİLER

Koray DEMİROK-Fatih EREN-Yusuf Samet AKKAYA-Dilek ÇAYIRLI

Generative AI Performance Benchmarking: The Impact of ChatGPT and DeepSeek on User Interactions
(Research Article)

Üretken Yapay Zeka Performans Kıyaslaması: ChatGPT ve DeepSeek'in Kullanıcı Etkileşimleri Üzerindeki Etkisi
(Araştırma Makalesi)

Atıf: Demirok, K., Eren, F., Akkaya Y. S., Çayırılı, D., (2025). Üretken yapay zeka performans kıyaslaması: ChatGPT ve DeepSeek'in kullanıcı etkileşimleri üzerindeki etkisi. *Bilgi ve İletişim Teknolojileri Dergisi* 7(1), 1-39. <https://doi.org/10.53694/bited.1710862> 1-39

Cite: Eren, F., Akkaya Y. S., Çayırılı, D., (2025). Generative AI performance benchmarking: The impact of ChatGPT and DeepSeek on user interactions. *Journal of Information and Communication Technologies*, 7(1), 1-39. <https://doi.org/10.53694/bited.1710862>

Özgür ATAMAN- Fatma Gizem KARAOĞLAN YILMAZ

Applications of Artificial Intelligence in Education: A Systematic Review of Postgraduate Theses
(Research Article)

Eğitim Alanında Yapay Zekâ Uygulamaları: Lisansüstü Tezlerin Sistematik İncelemesi
(Araştırma Makalesi)

Atıf: Ataman, Ö., Karaoğlan Yılmaz F. G., (2025). Oyunlaştırılmış e-öğrenme destekli öğretimin öğrencilerin bilgisayar programlama öz yeterlik algıları, programlamaya yönelik tutumları ve derse katılımlarına etkisi. *Bilgi ve İletişim Teknolojileri Dergisi*, 7(1), 40-62. <https://doi.org/10.53694/bited.1603352> 40-62

Cite: Ataman, Ö., Karaoğlan Yılmaz F. G., (2025). The effect of gamified e-learning supported teaching on students' computer programming self-efficacy perceptions, attitudes towards programming and classroom engagement. *Journal of Information and Communication Technologies*, 7(1), 40-62. <https://doi.org/10.53694/bited.1603352>

Şafak KAMAN

Investigation of Classroom Teachers' Artificial Intelligence Literacy Levels According to Various Variables
(Research Article)

Sınıf Öğretmenlerinin Yapay Zekâ Okuryazarlık Düzeylerinin Çeşitli Değişkenlere Göre İncelenmesi
(Araştırma Makalesi)

Atıf: Kaman, Ş. (2025). Sınıf öğretmenlerinin yapay zekâ okuryazarlık düzeylerinin çeşitli değişkenlere göre incelenmesi. *Bilgi ve İletişim Teknolojileri Dergisi*, 7(1), 63-77. <https://doi.org/10.53694/bited.1628589> 63-77

Cite: Kaman, Ş. (2025). Investigation of classroom teachers' artificial intelligence literacy levels according to various variables. *Journal of Information and Communication Technologies*, 7(1), 63-77. <https://doi.org/10.53694/bited.1628589>

Üretken Yapay Zeka Performans Kıyaslaması: ChatGPT ve DeepSeek'in Kullanıcı Etkileşimleri Üzerindeki Etkisi

Koray DEMİROK*¹ Fatih EREN² Yusuf Samet AKKAYA³ Dilek ÇAYIRLI⁴

Anahtar Sözcükler

Üretken Yapay Zeka
ChatGPT
DeepSeek
Kullanıcı Etkileşimi
Performans
Kıyaslaması
Makale Hakkında
Gönderim Tarihi
31 Mayıs 2025
Kabul Tarihi
27 Haziran 2025
Yayın Tarihi
30 Haziran 2025

Öz

Bu çalışma, üretken yapay zekâ modelleri olan ChatGPT ve DeepSeek'in kullanıcı etkileşimleri üzerindeki performansını ve etkisini karşılaştırmalı bir analizle değerlendirmeyi amaçlamaktadır. Araştırmada, Kaggle.com'dan temin edilen ve Temmuz 2023 ile Şubat 2025 arasını kapsayan, 10.000'den fazla kullanıcı etkileşim verisi içeren sentetik bir veri seti kullanılmıştır. Makine öğrenmesi algoritmaları (regresyon, karar ağaçları, kümeleme) ve istatistiksel testler (t-test, ANOVA) aracılığıyla yanıt doğruluğu, oturum süresi, kullanıcı memnuniyeti gibi değişkenler analiz edilmiştir. Bulgular, DeepSeek'in özellikle teknik görevlerde (hata ayıklama, kod optimizasyonu) daha yüksek kullanıcı derecelendirmeleri aldığını, ChatGPT'nin ise içerik üretimi ve eğitim gibi alanlarda öne çıktığını göstermiştir. DeepSeek kullanıcıları genel olarak daha yüksek deneyim puanları ve derecelendirmeler bildirirken, daha düşük terk etme oranlarına sahip olduğu gözlemlenmiştir. Çalışma, her iki platformun güçlü yönlerini, sınırlılıklarını ve kullanıcı etkileşim dinamiklerini ortaya koyarak, platformların kullanım bağlamına göre performans farklılıkları sergilediğini ve veri gizliliği ile etik konuların önemli olduğunu vurgulamaktadır.

Makale Türü

Araştırma Makalesi

Generative AI Performance Benchmarking: The Impact of ChatGPT and DeepSeek on User Interactions

Keywords

Generative Artificial Intelligence
ChatGPT
DeepSeek
User Interaction
Performance
Comparison
Article Info
Received
May 31, 2025
Accepted
June 27, 2025
Published
June 30, 2025

Abstract

This study aims to evaluate and compare the performance and impact of generative AI models ChatGPT and DeepSeek on user interactions through a comparative analysis. The research utilizes a synthetic dataset obtained from Kaggle.com, encompassing over 10,000 user interaction records collected between July 2023 and February 2025. Variables such as response accuracy, session duration, and user satisfaction were analyzed using machine learning algorithms (regression, decision trees, clustering) and statistical tests (t-test, ANOVA). The findings indicate that DeepSeek received higher user ratings particularly in technical tasks (e.g., debugging, code optimization), while ChatGPT stood out in areas such as content generation and education. DeepSeek users generally reported higher experience scores and ratings, along with lower abandonment rates. The study highlights the strengths, limitations, and user interaction dynamics of both platforms, emphasizing that performance varies depending on the context of use and underlining the importance of data privacy and ethical considerations.

Article Type

Research Paper

Atıf: Demirok, K., Eren, F., Akkaya Y. S., Çayırli, D., (2025). Üretken yapay zeka performans kıyaslaması: ChatGPT ve DeepSeek'in kullanıcı etkileşimleri üzerindeki etkisi *Bilgi ve İletişim Teknolojileri Dergisi*, 7(1), 1-39. <https://doi.org/10.53694/bited.1710862>

Cite: Demirok, K., Eren, F., Akkaya Y. S., Çayırli, D., (2025). Generative AI performance benchmarking: The impact of ChatGPT and DeepSeek on user interactions. *Journal of Information and Communication Technologies*, 7(1), 1-39. <https://doi.org/10.53694/bited.1710862>

*Sorumlu Yazar/Corresponding Author demirkoray74@gmail.com

¹ M.Sc. Student, Bartın University, Science Faculty, Bartın/Turkey, demirkoray74@gmail.com, <https://orcid.org/0009-0003-2796-845X>

² M.Sc. Student, Bartın University, Science Faculty, Bartın/Turkey, fatiheren.15@outlook.com, <https://orcid.org/0009-0009-2259-0420>

³ M.Sc. Student, Bartın University, Science Faculty, Bartın/Turkey, yusufsametakkaya@gmail.com, <https://orcid.org/0009-0006-3927-9100>

⁴ M.Sc. Student, Bartın University, Science Faculty, Bartın/Turkey, dicayirli@gmail.com, <https://orcid.org/0000-0003-0621-2074>

Extended Abstract

Introduction

Artificial Intelligence (AI) has progressively evolved from rule-based systems to advanced generative models capable of mimicking human-like reasoning and creativity. Generative Artificial Intelligence (GenAI), as a subdomain of AI, focuses on producing novel content such as text, images, and code by learning patterns from large datasets (Al-Busaidi et al., 2024). The current study examines the comparative performance of two prominent GenAI platforms—OpenAI's ChatGPT and DeepSeek—through the lens of user interaction metrics, leveraging synthetic datasets derived from realistic user behavior.

As GenAI systems become increasingly integrated into education, software development, and digital content creation, understanding their performance and interaction dynamics is critical (Rivas & Liang, 2023). ChatGPT has garnered widespread use due to its fluency in natural language understanding and generation, while DeepSeek, albeit less popular, is recognized for its exceptional accuracy in programming and logic-intensive tasks (Guo et al., 2024; Manik, 2025). Despite their common foundations in transformer architecture, the models differ in architectural specialization, data access policies, and reinforcement learning strategies (Du et al., 2022; Ouyang et al., 2022).

Method

This research adopts a comparative quantitative methodology, utilizing a synthetic dataset titled "DeepSeek vs ChatGPT: AI Platform Comparison" retrieved from Kaggle.com. The dataset simulates over 10,000 user interactions from July 2023 to February 2025 and includes diverse variables such as session duration, user satisfaction, response accuracy, platform retention rate, and device type. Data preprocessing involved encoding categorical variables (e.g., Query_Type, Device_Type), handling missing values, and normalization using StandardScaler. Outlier detection employed IQR-based trimming.

Machine learning techniques such as linear regression, random forest, and gradient boosting were implemented to predict user satisfaction and session behaviors. Statistical analyses, including independent samples t-tests and one-way ANOVA, evaluated significant differences between the two platforms in terms of user ratings and experience scores.

Findings

User Satisfaction and Experience

Descriptive and inferential analyses indicated notable differences in the quality of user interactions between the two platforms. ChatGPT sessions were generally associated with moderate satisfaction scores, commonly ranging between 3 and 4 stars. These moderate evaluations were frequently observed in use cases such as content creation and educational assistance. In contrast, DeepSeek consistently achieved higher user ratings, predominantly between 4 and 5 stars, especially in sessions focused on debugging, algorithm development, and code optimization.

The average user experience score for DeepSeek was approximately 2.03, compared to 1.23 for ChatGPT. A strong positive correlation ($r \approx 0.75$) was observed between user ratings and experience scores, indicating that higher interaction satisfaction typically coincided with a more positive overall experience. This correlation was evident across both platforms, with DeepSeek exhibiting a more pronounced relationship between the two metrics.

Device-Type Variations

Interaction quality varied depending on the type of device used. Sessions conducted on desktop and laptop computers yielded the highest levels of both satisfaction and experience scores for both platforms. In contrast, interactions via tablets and smart speakers displayed greater variability and lower average experience values. These results suggest that device type plays a considerable role in shaping user perceptions and experiences during interactions with generative AI platforms.

Inferential Statistics

Statistical tests confirmed that the observed differences between ChatGPT and DeepSeek were not due to chance. Independent samples t-tests revealed statistically significant differences ($p < 0.001$) in both user satisfaction and experience metrics. Analysis of variance (ANOVA) further supported these findings, with F-statistics exceeding 4000 for satisfaction and 20,000 for experience scores. These results affirm that the type of AI platform used has a significant effect on user outcomes.

Machine Learning Modeling

Regression Models

Among the various regression approaches tested, the Gradient Boosting Regression model achieved the best performance, with a Root Mean Squared Error (RMSE) of approximately 0.3825 and an R^2 value of 0.2954. Despite this relatively better performance, all models yielded R^2 values below 0.30, indicating that a large portion of the variance in user satisfaction could not be explained by the available features. This suggests the influence of other unmeasured or subjective factors that affect user perceptions.

Classification Models

In attempts to classify user satisfaction into categories such as 'Excellent', 'Fair', and 'Good', the Logistic Regression model demonstrated the highest overall accuracy, around 57.5%. The model performed particularly well in predicting the most frequently occurring class but had lower accuracy in identifying the less common categories. Support Vector Machines, by comparison, showed weaker performance, especially for underrepresented satisfaction classes, highlighting the potential need for future work to address class imbalance in the dataset.

Churn and Loyalty Modeling

A Random Forest classifier was employed to predict user churn based on interaction data, achieving perfect accuracy on the test set. This model successfully identified users who were likely to discontinue usage of the

platform. Additionally, regression models were used to predict session duration and user return frequency, with RMSE values of approximately 14.57 seconds and 3.16 respectively. These results indicate a reasonable level of predictive performance, though improvements are still possible.

Response Accuracy and Correction Need

Analysis showed that DeepSeek generally provided more accurate responses and required fewer corrections compared to ChatGPT. A classification model was used to predict whether a given response would require correction. While the model achieved high accuracy in identifying correct responses, its ability to detect errors was limited. The macro-average F1 score was approximately 0.51, suggesting that performance was moderate and could be enhanced through further optimization and data balancing strategies.

Discussion and Conclusion

This study affirms that platform-specific architectural and training differences manifest distinctly in user interactions. ChatGPT's broad adaptability (Bahrini et al., 2023; Hariri, 2023) makes it effective in educational and creative scenarios, while DeepSeek's MoE-based architecture (Du et al., 2021) provides efficiency and precision in logic-based tasks, especially programming. DeepSeek's superior performance in debugging, its use of large-scale token windows (128K), and high efficiency in energy consumption (IBM, 2025b) reflect an optimization strategy tailored to technical use-cases. However, its closed training datasets and limited transparency in RLHF protocols (Sapkota et al., 2025; CTOL Digital, 2025) raise ethical concerns around reproducibility and bias. In contrast, ChatGPT, while offering broader linguistic fluency, is challenged by hallucination problems (Alkaissi & McFarlane, 2023) and ethical constraints related to misinformation, bias, and data privacy (Rozado, 2023; Zirpoli, 2023). Both models face scrutiny regarding user data handling under GDPR and KVKK regulations, as reflected in the temporary ban of ChatGPT in Italy (Reuters, 2024).

This study provides a comprehensive comparative analysis of ChatGPT and DeepSeek based on synthetic user interaction data. While both platforms excel in distinct domains, DeepSeek demonstrates a performance edge in technical accuracy and user satisfaction. These findings suggest that user context and task type should guide the deployment of GenAI systems. Future research should integrate real-world datasets, address classification imbalance, and further explore ethical implications such as transparency, data security, and hallucination mitigation.

The implications of this research extend to developers, AI ethicists, and organizational decision-makers seeking to optimize GenAI adoption. As generative models become increasingly embedded in daily workflows, ongoing assessment of their sociotechnical impact remains essential.

1.Giriş

“Yapay zekâ (YZ), insan davranışındaki zekâ ile ilişkilendirdiğimiz özellikleri sergileyen sistemler geliştirme teori ve pratiği ile ilgili Bilim ve Mühendislik alanıdır” (Tecuci, 2012).

Yapay zekâ teknolojileri, insan zekasını taklit etme yetenekleriyle karakterize edilir. Eren (2024), bu tür sistemlerin öğrenme, mantıksal çıkarım yapma ve karar verme becerilerine sahip olduğunu belirtmektedir. Yapay zekâ ise, algılama, doğal dil işleme, problem çözme, planlama, öğrenme, uyum sağlama ve çevreyle etkileşim kurma gibi insan zekasına özgü davranışları sergileyen sistemlerin geliştirilmesine odaklanan bir Bilim ve Mühendislik disiplindir (Tecuci, 2012). Yapay zekâ'nın temel bilimsel amacı, insanlarda, hayvanlarda ve yapay zekalarda akıllı davranışı mümkün kılan prensipleri anlamaktır. Bu bilimsel amaç, akıllı ajanlar geliştirmek, bilgiyi biçimlendirmek ve muhakemeyi mekanikleştirmek, bilgisayarlarla çalışmayı insanlarla çalışmak kadar kolaylaştırmak ve insan ve otomatik muhakemenin tamamlayıcılığından yararlanan insan-makine sistemleri geliştirmek gibi çeşitli mühendislik hedeflerini doğrudan destekler (Tecuci, 2012).

Üretken Yapay Zekâ (GenAI), mevcut veri kümelerinden örüntüler çıkararak metin, görsel, müzik, yazılım kodu ve video gibi özgün içerikler üretebilen özel bir yapay zekâ biçimini tanımlar (Al-Busaidi vd., 2024). (Bell ve Bell, 2023) ile (Dwivedi vd., 2023), geleneksel yapay zekâ yaklaşımlarının çoğunlukla veri analizi ve sınıflandırma gibi belirli görevlerle sınırlı kaldığını belirtmektedir. Buna karşın üretken yapay zekâ, insan tarafından oluşturulmuş içeriklere oldukça yakın ve yaratıcı nitelikte çıktılar üretebilme yeteneğiyle öne çıkmaktadır. Yapay zekâ, çok katmanlı ve öğrenme tabanlı bilgisayar programlarının geliştirilmesi ve kullanılmasıyla insanların yapabileceği her şeyi yapabilen bir teknolojidir. Yapay zekâ aynı zamanda, insan davranışını ve zihinsel yeteneklerini taklit eden algoritmalar kullanarak, öğrenmek ve karar vermek için insanların yapamayacağı şeyleri yapabilen bir bilgisayar sistemi olarak da tanımlanabilir (Şentürk, 2023).

Generatif yapay zekâ modellerinin hızla gelişmesi, insanların teknolojiyle etkileşim biçimlerinde köklü bir değişimi beraberinde getirmiştir (Rivas ve Liang 2023). Bu dönüşümün öncüsü olarak ChatGPT ve DeepSeek gibi modellerin kilit bir rol oynadığını vurgulamaktadır. Doğal dil işleme ve makine öğrenmesi algoritmalarından yararlanan bu yapay zekâ sistemleri, çeşitli sektörleri dönüştürmekte ve kullanıcı beklentilerini yeniden şekillendirmektedir (Rivas ve Liang, 2023). OpenAI tarafından geliştirilen ChatGPT, doğal dili anlama konusundaki yetkinliği, bağlama duyarlılığı ve insan benzeri yanıtlar üretebilme kapasitesiyle, çok çeşitli kullanım alanlarına hitap eden esnek ve etkili bir yapay zekâ aracı olarak dikkat çekmektedir (Bahrini vd., 2023; Hariri, 2023). ChatGPT; müşteri hizmetleri, sağlık ve eğitim gibi birçok alanda uygulama bulmuştur (Hariri, 2023; Ray, 2023). DeepSeek ise belki daha az bilinen bir model olmakla birlikte, generatif yapay zekâ alanında önemli bir gelişmeyi temsil etmekte olup, kendine özgü yetenekleri ve performans özellikleriyle karşılaştırmalı bir analiz yapılmasını gerektirmektedir (Wang vd., 2024). Bu teknolojiler gelişimini sürdürürken, güçlü yönlerinin, sınırlılıklarının ve kullanıcı etkileşimleri üzerindeki etkilerinin anlaşılması araştırmacılar, geliştiriciler ve son kullanıcılar açısından büyük önem taşımaktadır. Bu modellerin performansını incelemek, yalnızca anlamlı ve bağlama uygun metinler üretme yeteneklerini değil, aynı zamanda kullanıcılarla anlamlı ve verimli diyaloglar kurabilme kapasitelerini de dikkate alan çok yönlü bir yaklaşımı gerektirir (Nazir & Wang, 2023).

1.2.Üretken Yapay Zekâ ile Kullanıcı Etkileşimi Dinamikleri

Bireylerin ChatGPT gibi üretken yapay zekâ araçlarıyla etkileşime girmeyi neden seçtiklerini anlamak, performanslarını ve etkilerini değerlendirmek için çok önemlidir (Skjuve vd., 2024). Kullanıcı etkileşimi, algılanan değer, kullanım kolaylığı ve etkileşimden elde edilen keyif dahil olmak üzere çeşitli faktörlerden etkilenir (Al-Abdullatif ve Alsubaie, 2024). Yapay zekanın eğitim ortamlarına entegrasyonu hem fırsatlar hem de zorluklar sunar ve pedagojik olarak güçlü ve etik açıdan sağlam yapay zekâ entegreli öğrenme ortamlarını teşvik etmek için kullanıcı deneyimlerinin eleştirel bir şekilde incelenmesini gerektirir (Mogavi vd., 2023; Qawqzeh, 2024). Yapay zekâ ile kullanıcı etkileşimleri, özellikle yapay zekâ tarafından üretilen içeriklerin sıklıkla görüldüğü sosyal medya bağlamlarında giderek daha yaygın hale geliyor (Chein vd., 2024). Yapay zekanın eğitimi devrim niteliğinde değiştirme potansiyeli önemlidir; kişiselleştirilmiş öğrenme deneyimleri sunar ve idari görevleri otomatikleştirir (Adıgüzel vd., 2023). Yapay zekâ sistemleri daha karmaşık hale geldikçe, sanal asistanlar olarak hizmet verebilir, öğrencilere anında geri bildirim ve destek sağlayabilir ve eğitimcilerin daha karmaşık öğretim etkinliklerine odaklanmasını sağlayabilir (Sharples, 2023). Üniversitelerde yapay zekâ sohbet robotlarının konuşlandırılması, öğrenci katılımını artırmanın ve öğrenme başarılarını geliştirmenin olası bir yolu olarak önemli ilgi görmektedir (Fırat, 2023). Öğrencilerin yapay zekâ ile etkileşim kurma biçimi, özellikle işbirlikçi öğrenme senaryolarında, kritik öneme sahiptir.

1.3.ChatGPT

ChatGPT'nin Tanımı: ChatGPT, OpenAI tarafından geliştirilen güçlü bir GPT modeline dayanan ve özellikle optimize edilmiş konuşma yeteneklerine sahip bir sohbet modelidir (Zhao vd, 2023). OpenAI tarafından geliştirilen üretken bir yapay zekâ dil modelidir. Transformer mimarisini kullanarak tutarlı ve bağlamsal olarak uygun, insan benzeri yanıtlar üreten bir OpenAI üretken yapay zekâ dil modelidir. Gelişmiş bir doğal dil işleme modelidir (Markov vd, 2023).

ChatGPT ve GPT modellerinin temel prensibi, dünya bilgisini dil modellemesi yoluyla yalnızca-kod çözücülü (decoder-only) Transformer modeline sıkıştırmaktır. Bu sayede dünya bilgisinin semantiğini geri kazanabilir (veya ezberleyebilir) ve genel amaçlı bir görev çözücü olarak işlev görebilir. Başarısının iki temel noktası şunlardır: (I) Bir sonraki kelimeyi doğru bir şekilde tahmin edebilen yalnızca-kod çözücülü Transformer dil modellerinin eğitilmesi ve (II) dil modellerinin boyutunun ölçeklendirilmesi. ChatGPT, InstructGPT'ye benzer bir şekilde eğitilmiştir. Ancak ChatGPT'nin eğitiminde, insan tarafından oluşturulan konuşmalar (kullanıcı ve AI rollerini oynayan) InstructGPT veri setiyle diyalog formatında birleştirilmiştir (Zhao vd, 2023). GPT modelleri, geniş miktarda metin verisi üzerinde derin öğrenme tekniği olan transformer'ları kullanarak eğitilmiştir. ChatGPT'nin farklı versiyonları bulunmaktadır. Kasım 2022'de genel erişime açılan ChatGPT-3.5, geniş kullanıcı kitlesiyle buluşarak dil modelleme alanında önemli bir adım olmuştur. Bunu takiben, Mart 2023'te sunulan GPT-4 modeli; artırılmış doğruluk, daha güvenli yanıtlar ve kullanıcı niyetine daha duyarlı tepkiler verme kabiliyetiyle öne çıkmıştır. Özellikle akademik yazım, proje üretimi ve içerik geliştirme gibi alanlarda tercih edilen bu modellerin eğitim verilerindeki zamansal sınırlama ise ChatGPT-3.5 için Eylül 2021 itibarıyla belirlenmiştir. Bu tarih, modelin bilgi tabanının güncellik sınırını oluşturmaktadır, bir araştırma sürecinde bu sınır göz önünde bulundurulmalıdır. Bu, ChatGPT'nin bu tarihten sonraki olayları, keşifleri veya görüşleri entegre edemeyeceği anlamına gelir. ChatGPT4o'nun en son veriye erişim tarihi konusunda çelişkili bilgiler bulunmaktadır, bazı kaynaklar Ekim 2023 derken, ChatGPT4o'nun kendisi Eylül 2021'i belirtmektedir (Spennemann, 2025).

1.3.1.ChatGPT'nin Temel Özellikleri:

İnsanlarla mükemmel düzeyde iletişim kurma kapasitesine sahip olup, geniş bir bilgi birikimiyle donatılmıştır. Matematiksel problemler üzerinde akıl yürütebilir, çok turlu diyaloglarda bağlamı doğru bir şekilde takip edebilir ve insan değerleriyle uyum sağlayarak güvenli kullanım sunar. Ayrıca eklenti mekanizmasını desteklemesi, mevcut araçlar ve uygulamalarla entegrasyonunu mümkün kılarak kapasitesini daha da genişletir. Doğal dil konuşmaları üretme konusunda güçlü bir temele sahip olan sistem, doğal sesli yanıtlar oluşturabilir. Dil çevirisi, görüntü tanıma ve tahmine dayalı modelleme gibi çeşitli yapay zekâ uygulamalarında etkin biçimde kullanılabilir. Bununla birlikte, metin kaynaklarından veri çıkarma ve özetleme, insan benzeri bağlamsal yanıtlar sunma, metin sınıflandırma ve duygu analizi gibi doğal dil görevlerini başarıyla yerine getirme yeteneğine sahiptir. E-posta ya da programlama kodu gibi temel bağlamsal metinlerin taslağını hazırlayabilmesi de önemli bir avantajdır. Tüm bu özelliklerinin ötesinde, GPT-4 ile birlikte daha yüksek olgusal doğruluk ve kullanıcı niyetine daha isabetli yanıtlar verme gibi gelişmiş yetenekler de sunmaktadır (OpenAI, 2023).

Yapay zekâ teknolojileri sağlık alanında biyolojik verilerden bilgi çıkarma, tıbbi tavsiye danışmanlığı, ruh sağlığı analizi ve raporların sadeleştirilmesi gibi görevleri yerine getirebilmekte; hatta USMLE sınavlarında uzman düzeyinde başarı gösteren Med-PaLM modelleri gibi özel sistemler geliştirilmiştir (Mesko, 2023). Eğitimde ise öğretim materyali üretme, sunum hazırlama, araştırma fikirleri ve metodolojileri önerme, istatistiksel analiz ve veri yorumlama, yazma becerilerini geliştirme ve akademik yazım desteği sağlama gibi çok yönlü katkılar sunmaktadır (Zheldibayeva, 2024). İş dünyasında operasyon ve yönetim süreçlerinden tedarik zinciri yönetimine, iş analitiğinden pazarlama ve e-ticaret uygulamalarına kadar geniş bir yelpazede kullanılmaktadır. Bilimsel araştırmalarda ise araştırma fikri geliştirme, yöntem belirleme, literatür taraması yapma ve hipotez oluşturma gibi çalışmalarda etkin rol oynamaktadır (Acemoglu vd., 2023). Hukuk alanında dava taslağı hazırlama, yasal araştırmaları hızlandırma, sözleşme analizi ve belge inceleme süreçlerinde kullanılabilirken; finans alanında da çeşitli uygulamalara entegre edilebilmektedir. Sanat ve kültür alanında müzik, sanat ve tasarım, yaratıcı yazarlık, oyun geliştirme ve sanal gerçeklik gibi yaratıcı süreçlere katkı sağlamaktadır. Programlama konusunda ise kod üretimi, hata ayıklama, akademik projelerde destek ve mülakatlara hazırlık gibi işlevleriyle yaygın şekilde kullanılmaktadır. Müşteri hizmetlerinde geçmiş konuşmaları da dikkate alarak müşteri etkileşimlerini otomatikleştirme, soruları yanıtlama ve destek sağlama görevlerini üstlenirken; veri bilimi alanında kişisel bir veri bilimcisi gibi işlev görebilmektedir. Ayrıca akademik yazımda makale üretimine katkı sağlayabilmekte ve hakemlik süreçlerinde makale değerlendirme çalışmalarında da kullanılabilir (Liu vd., 2024).

1.3.2.ChatGPT ile İlişkili Etik Endişeler:

ChatGPT'nin kullanımında bazı önemli sınırlılıklar ve riskler bulunmaktadır. Bunların başında, zaman zaman gerçek dışı veya mantıksız çıktılar üretmesiyle ortaya çıkan "hallucination" sorunu gelmektedir; özellikle akademik referanslar gibi bilgi temelli içeriklerde, uydurma veriler sunabilme eğilimi dikkat çekicidir (Alkaissi & McFarlane, 2023). Ayrıca, eğitim verilerindeki önyargıları yansıtma riski taşıyan sistem, din, kültür, cinsiyet, ırk ve sosyoekonomik durum gibi çeşitli konularda istemeden taraflı ifadeler üretebilir; bazı analizler, ChatGPT'nin siyasi yönelim testlerinde özgürlükçü, ilerici ve sol eğilimli görüşlere yatkınlık sergilediğini göstermektedir (Rozado, 2023). Fikri mülkiyet ve telif hakları da bir başka önemli kaygı alanıdır; zira üretken yapay zekâ modelleri, telif hakkına tabi içeriklerden öğrenmekte ve bu da potansiyel olarak hukuki sorumluluk doğurabilecek

durumlar yaratmaktadır (Zirpoli, 2023). Kötüye kullanım potansiyeli de göz ardı edilemez; bu teknolojiyle kötü amaçlı yazılım üretimi, toksik dil kullanımı, yanlış bilgi yayma ve akademik etik ihlalleri gibi olumsuz senaryolar mümkün hale gelebilmektedir (Weidinger vd., 2021). Sistemlerin arkasındaki algoritmik yapıların karmaşıklığı nedeniyle şeffaflık sınırlıdır; bu da tahminlerin ya da önerilerin nasıl üretildiğini anlamayı güçleştirmektedir. Dahası, yapay zekâ tarafından oluşturulan içerikler konusunda sorumluluğun kime ait olduğu hâlâ net değildir; bu nedenle ChatGPT'nin bilimsel makalelerde yazar olarak kabul edilmemesi gerektiği belirtilmektedir. Son olarak, bu tür sistemlere aşırı güven duymak, kullanıcıların eleştirel düşünme ve doğrulama becerilerini zayıflatma riskini beraberinde getirmektedir (Stokel-Walker, 2023).

ChatGPT ve benzeri yapay zekâ tabanlı sohbet sistemleri, veri gizliliği ve güvenliği açısından çeşitli riskler barındırmaktadır. Kullanıcıların hassas veya kişisel bilgilerini ticari sunuculara yüklemesi, özellikle sağlık verileri gibi yüksek derecede özel içeriklerde ciddi gizlilik ihlali potansiyeli taşımaktadır (Borgesius, 2018). Bu sistemlerin işleyebilmesi için büyük miktarda veriye ihtiyaç duyulması ise, onları siber saldırılara karşı daha savunmasız hale getirmekte ve veri güvenliği risklerini artırmaktadır (ICO, 2023). Ayrıca, ChatGPT gibi modellerin kişisel verileri işlemesi, özellikle eğitimlerinde doğruluğu tartışmalı ya da eksik veri setleri kullanılmışsa, hatalı veya yanıltıcı sonuçlar üretme olasılığını beraberinde getirmektedir (Weidinger vd., 2021). Avrupa Birliği Genel Veri Koruma Tüzüğü - General Data Protection Regulation (GDPR) ve Türkiye'deki Kişisel Verilerin Korunması Kanunu (KVKK) gibi düzenlemeler bağlamında bu tür sistemlerin veri işleme uygulamalarının ne ölçüde uyumlu olduğu tartışmalı olup, İtalya'nın ChatGPT'yi gizlilik gerekçesiyle geçici olarak yasaklaması bu alandaki kaygıların altını çizmektedir (Reuters, 2024). Ayrıca, kullanıcıların sohbet geçmişi ve girdilerinin nasıl saklandığı, kimlerle ve ne amaçla paylaşıldığı konusundaki şeffaflık eksikliği önemli bir endişe kaynağıdır; nitekim Google'ın Gemini modeli için yayımladığı gizlilik bildiriminde, kullanıcıların kişisel bilgilerini yapay zekâ ile paylaşmamaları yönündeki uyarı, bu endişeleri açıkça yansıtmaktadır (Search Engine Journal, 2024).

1.4.DeepSeek

DeepSeek, mimari olarak OpenAI'nin GPT serisi ya da Meta'nın LLaMA modellerine benzer şekilde transformatör tabanlı büyük dil modelleri sınıfına ait olsa da, özellikle kod üretimi ve mantıksal çıkarım gerektiren görevlerde optimize edilmiş yapısıyla öne çıkmaktadır (DeepSeek-AI, 2024). GPT-4 gibi modeller, geniş kapsamlı doğal dil anlama ve üretimi süreçlerinde geneli bir paradigma benimseyerek çok çeşitli görevlerde etkin performans sergilemeyi amaçlamaktadır. Buna karşın, DeepSeek daha dar kapsamlı ve belirli görev odaklı senaryolarda, doğruluk ve bağlamsal bütünlük açısından daha yüksek başarı elde etmeye yönelik, özelleştirilmiş ve hedef odaklı bir mühendislik yaklaşımını benimsemektedir. Kod üretiminde, özellikle Python, Java ve C++ gibi dillerde yüksek başarı gösteren DeepSeek, CodeLlama ya da DeepMind'ın AlphaCode modeliyle aynı ligde değerlendirilmekte, ancak daha istikrarlı bağlam takibi ve hata ayıklama hassasiyetiyle öne çıkmaktadır (Guo vd., 2024). Modelin açık ağırlıklı olması, araştırma ve geliştirici topluluğu için erişim imkânı sunsa da, tam anlamıyla açık kaynak olmaması nedeniyle yeniden eğitim, veri incelemesi veya modifikasyon gibi işlemler sınırlı düzeyde gerçekleştirilebilmektedir. Performans açısından bakıldığında, DeepSeek'in tımdengelimli akıl yürütme, çok adımlı mantıksal çözümleme ve programatik problem çözme gibi görevlerde GPT-3.5'i geçtiği; ancak GPT-4 ya da Claude 3 Opus gibi modellerle karşılaştırıldığında bağlam genişliği ve genel dil üretim kalitesinde hâlen sınırlı kaldığı görülmektedir. Bununla birlikte, kaynak verimliliği, inference (çıkı üretme) hızı ve geliştirici dostu

arayüzü sayesinde dar amaçlı yüksek doğruluk gerektiren görevlerde tercih edilebilir bir çözüm sunmaktadır (Manik, 2025).

DeepSeek ekosistemi, farklı görev ve kullanım senaryolarına uyarlanmış çeşitli modeller içermektedir. Örneğin, DeepSeek-R1 modeli, özellikle matematiksel akıl yürütme ve programlama gibi bilişsel yoğunluk gerektiren görevler doğrultusunda optimize edilmiş olmakla birlikte, aynı zamanda genel amaçlı bir dil modeli olarak da işlev görmektedir (Guo vd., 2024). Buna karşılık, DeepSeek-V3 modeli, çoklu disiplinleri kapsayan geniş ölçekli bir veri kümesi üzerinde eğitilmiş olup, genel amaçlı kullanım senaryolarında yüksek performans sergileyen kapsamlı bir yapay zeka çözümü olarak öne çıkmaktadır. Kodlama alanına odaklanan DeepSeek Coder V2, 338 dili destekleyen ve gelişmiş kodlama yeteneklerine sahip bir kod dil modeli olup, açık kaynaklı bir Karışık Uzmanlar (Mixture-of-Experts, MoE) mimarisine sahiptir. Bunun yanında, DeepSeek VL, gerçek dünya görme ve dil görevlerinde başarılı olabilmek üzere tasarlanmış açık kaynaklı bir Görsel-Dil (Vision-Language, VL) modelidir (DeepSeek-AI, 2024).

Ayrıca, DeepSeek LLM modelleri de dikkat çekmektedir. DeepSeek LLM, 67 milyar parametreye sahip olup tescilli ağırlıklarla eğitilmiş ve 2023'e kadar olan Kitaplar+Wiki verileri üzerinde geliştirilmiştir. Daha büyük ölçekteki DeepSeek LLM V2 modeli ise 236 milyar parametreye ve tescilli bir Karışık Uzmanlar (MoE) mimarisine sahiptir (Guo vd., 2024). Son olarak, DeepSeek-V3 modeli 671 milyar parametre kapasitesine ulaşmakta ve FP8 eğitimi gibi yüksek verimli öğrenme tekniklerini kullanan tescilli bir Karışık Uzmanlar mimarisiyle tasarlanmıştır. Bu çeşitlilik, DeepSeek ekosisteminin hem genel amaçlı hem de uzmanlaşmış görevlerde güçlü performans sunmasına olanak tanımaktadır (DeepSeek-AI, 2024).

1.4.1. DeepSeek'in Temel Özellikleri

MoE (Mixture-of-Experts) mimarisi, büyük dil modellerinin hesaplama verimliliğini artırmak amacıyla geliştirilen ve yalnızca belirli parametre alt kümelerini etkinleştirerek çalışan dinamik bir hesaplama yöntemi olarak tanımlanmaktadır. Bu mimari, her bir katmanda çok sayıda “uzman” (expert) bileşen barındırmakta olup, model her bir girdi için yalnızca sınırlı sayıda uzmanın çıktısını hesaba katarak işlemi gerçekleştirmektedir. Bu sayede, hem hesaplama maliyetleri azaltılmakta hem de modelin ölçeklenebilirliği önemli ölçüde artırılmaktadır (Du vd., 2021). Bu yaklaşımın başlıca avantajlarından biri kaynak verimliliğidir: yüz milyarlarca parametrelili modellerin eğitimi ve çalıştırılması oldukça maliyetli iken, MoE yalnızca gerekli uzmanları aktif ederek hesaplama yükünü önemli ölçüde azaltmaktadır. Ayrıca, ölçeklenebilirlik açısından parametre sayısı artırılarak daha yüksek kapasiteli modeller tasarlanabilirken, çıkarım (inference) sırasında maliyet sabit tutulabilmektedir (Artetxe vd., 2021).

MoE mimarisinin bir diğer avantajı da uyarlanabilirlik özelliğidir; çünkü her uzman, farklı görevler veya bağlamlar için eğitilebilmekte ve bu sayede model çoklu görevlerde daha esnek bir şekilde çalışabilmektedir. Örneğin, DeepSeek-R1 modeli, MoE mimarisi temel alınarak geliştirilmiş olup, toplamda 671 milyar parametre içermesine rağmen, çıkarım aşamasında yalnızca 37 milyar parametreyi etkinleştirmektedir. Bu yaklaşım, modelin yüksek parametre kapasitesine sahip olmasına rağmen hesaplama verimliliğinin korunmasını sağlamaktadır. Bu durum hem yüksek kapasiteyi hem de maliyet avantajını birlikte sunmakta olup, muhtemelen 16 veya 32 uzmandan oluşan bir top-k yönlendirme mekanizması (örneğin $k=2$) kullanılarak gerçekleştirilmektedir (DeepSeek R1, t.y.).

Pekiştirmeli öğrenme (Reinforcement Learning), modelin çıktılarını optimize etmek amacıyla bir “ödül sinyali” kullanılarak gerçekleştirilen bir öğrenme yöntemidir (Du vd., 2021). Büyük dil modellerinde (LLM’lerde), özellikle insan geri bildirimle pekiştirmeli öğrenme (Reinforcement Learning from Human Feedback, RLHF) formu öne çıkmaktadır. Bu yaklaşım, muhakeme ve karar verme becerilerinin geliştirilmesi, daha insani, bağlamsal ve faydalı yanıtlar üretilmesi ve zararlı ya da ilgisiz içeriklerin azaltılması gibi önemli kullanım alanlarıyla dikkat çekmektedir (Ouyang vd., 2022). DeepSeek-R1 modeli, yalnızca gözetimli öğrenme yöntemleriyle değil, aynı zamanda pekiştirmeli öğrenme süreçleriyle de geliştirilmiştir. Böylece, model yalnızca verileri taklit etmekle kalmayıp, amaca yönelik mantıksal çıkarımlarda bulunmayı da öğrenebilmektedir (Gupta, 2025). Örneğin, uzun adımlı çıkarım süreçlerinde ortaya çıkan hataları analiz ederek kendi yanıtlarını düzeltebilme kapasitesine sahiptir. Bu sürecin teknik olarak uygulanmasında sıklıkla gözetimli ince ayar (supervised fine-tuning), ödül modeli eğitimi (reward model training) ve Proximal Policy Optimization (PPO) gibi algoritmalarla gerçekleştirilen RLHF aşamaları yer almaktadır (Neptune.ai, 2025).

DeepSeek-V3 modeli, büyük olasılıkla GPT mimarisi temel alınarak geliştirilmiş olup, rotary embedding, katman normalizasyonu ve dikkat mekanizması optimizasyonları gibi güncel teknikleri bünyesinde barındırmaktadır (DeepSeek-V3, t.y.). Bu temel yapının üzerine inşa edilen DeepSeek-R1 ise, MoE (Mixture-of-Experts) mimarisi ile entegre edilmiş ve pekiştirmeli öğrenme (RL) yöntemleriyle desteklenmiş, performans ve verimlilik açısından iyileştirilmiş özel bir varyant olarak dikkat çekmektedir (Gupta, 2025). Böylece, hem model kapasitesi artırılmış hem de öğrenme süreçlerinde adaptasyon yeteneği güçlendirilmiştir. DeepSeek-R1 modeli, kod anlama, tümdengelimli muhakeme ve yüksek doğruluk gerektiren görevlerde geliştirilmiş olup, çok büyük bir parametre havuzuna sahip olmasına rağmen çıkarım sırasında yalnızca küçük bir kısmını kullanarak verimliliği artırmaktadır (Du vd., 2021). Ayrıca, DeepSeek-Coder, kod üretimi, analiz ve hata ayıklama gibi görevler için özelleştirilmiş bir alt model olarak konumlanmakta ve OpenAI Codex, CodeLlama veya AlphaCode gibi modellerle rekabet edebilecek düzeyde performans sunmayı hedeflemektedir (DeepSeek-Coder, t.y.).

DeepSeek-R1’in eğitimi, sentetik ve gerçek dünya verilerinden derlenen toplamda 9.8 trilyonluk geniş kapsamlı ve karma bir veri seti kullanılarak gerçekleştirilmiştir (DeepSeek-R1, t.y.). Söz konusu veri seti tescilli ve gizli bilgi içermesi sebebiyle, yalnızca yetkilendirilmiş kişiler tarafından erişilebilmekte ve sıkı güvenlik protokolleriyle korunmaktadır. Bu yaklaşım, modelin hem çeşitlilik hem de kalite açısından zengin bir eğitim materyaliyle desteklenmesini sağlamaktadır. Eğitim sürecini daha verimli hâle getirmek amacıyla, DeepSeek Gelişmiş Yeniden Parametrelendirme Optimizasyonu (GRPO) yöntemini kullanarak Pekiştirmeli Öğrenme (RL) için gerekli olan denetimli ince ayar (SFT) gibi ön adımları ortadan kaldırmış ve böylece maliyetleri önemli ölçüde düşürmüştür (Gupta, 2025). Şirketin başarısının temel faktörlerinden biri, titiz veri etiketleme uygulamalarıdır; liderlik kadrosu bu etiketleme süreçlerine doğrudan katılarak kaliteyi sürekli olarak denetlemektedir. Öte yandan, DeepSeek-R1 modeli, MIT Lisansı kapsamında açık ağırlık (open-weights) politikasıyla yayınlanmış olup, bu sayede ticari ve akademik araştırma amaçlı kullanımlar için nöral ağ parametrelerine sınırsız erişim imkânı sağlamaktadır (DeepSeek-R1, t.y.). Bu yaklaşım, geniş kullanıcı kitlesinin model üzerinde özgürce çalışabilmesine ve yenilikçi uygulamalar geliştirebilmesine olanak tanımaktadır. Ayrıca, DeepSeek’in bazı modelleri (özellikle V2 ve V3 sürümleri), FP8 eğitimi gibi yüksek verimli eğitim tekniklerini kullanarak daha hızlı ve verimli bir öğrenme süreci sağlamaktadır (DeepSeek-V3, t.y.).

DeepSeek-R1, zorlu problem çözme görevlerinde dikkat çekici muhakeme yetenekleri sergileyerek güçlü bir performans göstermektedir. Özellikle DeepSeek-Coder-V2, kod üretimi ve kod analizi alanlarında yüksek doğruluk oranları sergilemekte olup, bu başarısını GPT-4 Turbo ile karşılaştırılabilir düzeyde gerçekleştirmektedir (BytePlus, 2025). Model, karmaşık programlama görevlerinde etkin performans göstererek yazılım geliştirme süreçlerinde önemli bir araç olarak öne çıkmaktadır. 338 programlama dilini desteklemesi ve 128K'ya kadar uzayan bağlam uzunluğu, modelin esneklik ve derinlik açısından güçlü bir kapasite sunduğunu göstermektedir (Dataloop, t.y.). Matematiksel muhakeme alanında, DeepSeek-R1'in MATH veri seti üzerinde gerçekleştirdiği 30 zorlu matematik problemi, ChatGPT ve Gemini'yi geride bırakarak üstün bir performans ortaya koymuştur (Shirani, 2025). Ayrıca, bilimsel hesaplama görevlerinde de yüksek başarı elde etmiştir, özellikle nümerik integrasyon gibi alanlarda başarılı sonuçlar elde etmiştir. Bağlamsal anlama ve mantıksal çıkarım yetenekleri açısından ileri düzeyde olan DeepSeek, bazı modellerinde 128K token'a kadar uzayabilen uzun bağlam pencereleri sayesinde karmaşık ve kapsamlı metinleri etkin şekilde işleyebilmektedir (Dataloop, t.y.). Ayrıca, eğitim sürecinde sağladığı maliyet etkinliği de önemli bir rekabet avantajı sunmaktadır; DeepSeek-R1 modeli, yalnızca 5.5 milyon dolarlık bir bütçe ile geliştirilmiş olup, performans açısından birçok muadilini geride bırakmayı başarmıştır (IBM, 2025a). Karışık Uzmanlar (MoE) mimarisi sayesinde, yoğun eşdeğerlerine kıyasla enerji tüketimini %58 oranında azaltarak enerji verimliliği konusunda da önemli bir adım atmıştır (IBM, 2025b). DeepSeek V3'ün teknik özellikleri arasında, 7/8 aktivasyon parametresi ve toplamda 67 milyar parametre ile oldukça güçlü bir model yapısına sahiptir. Son olarak, DeepSeek modelleri, genellikle yapılandırılmış ve tutarlı referanslar sunabilme kapasitesiyle öne çıkmakta; bu özellik, modelin bilgi doğruluğunu artırmakta ve elde edilen çıktılara yönelik güvenilirliği önemli ölçüde yükseltmektedir (DeepSeek-V3, t.y.).

1.4.2. DeepSeek ile İlişkili Etik ve Veri Koruma Endişeleri

DeepSeek, ağırlıklarını ve bazı mimari detaylarını açıkça paylaşmasına rağmen (açık ağırlık), eğitim veri seti kompozisyonu ve Reinforcement Learning from Human Feedback (RLHF) gibi kritik eğitim detaylarını tescilli tutmaktadır (Sapkota vd., 2025). Bu durum, modelin potansiyel önyargılarını ve etik kaynak kullanımını incelemeyi zorlaştırmakta, aynı zamanda modelin tam anlamıyla yeniden üretilebilirliğini sınırlamaktadır. Açık kaynak belirsizliği ise, DeepSeek-R1'in MIT lisansı altında sunulmasına rağmen, eğitim verileri ve metodolojilerinin tam olarak paylaşılmaması nedeniyle, aslında tam anlamıyla "açık kaynak" tanımına uymadığını göstermektedir. Bu durum, "açık ağırlık" ve "açık kaynak" kavramları arasındaki karışıklığı ortaya koymaktadır (CTOL Digital, 2025). Önyargılar açısından, eğitim verilerinin tescilli olması, DeepSeek modellerinde olası önyargı kaynaklarını belirlemeyi ve bu önyargıları azaltmayı zorlaştırmaktadır (Witness AI, 2025). Veri gizliliği konusunda ise, DeepSeek'in veri işleme uygulamaları ve kullanıcı verilerinin nasıl saklandığına dair yeterli bilgi bulunmamaktadır; bu durum, kullanıcıların hassas verilerini bu tür platformlara yüklerken gizlilik endişeleri yaratmaktadır (Security Magazine, 2025). Bunun yanı sıra, diğer dil modellerinde olduğu gibi, DeepSeek'te de zaman zaman yanlış veya uydurma bilgi üretme (hallucination) riski mevcuttur; model, bazı durumlarda gerçeklikten sapmış veya mantıksal tutarsızlık içeren çıktılar üretebilmekte ve referans gösterme süreçlerinde hatalar yapabilmektedir (Vectara, 2025). Kötüye kullanım potansiyeli ise, DeepSeek'in kod üretme yeteneği sayesinde kötü amaçlı yazılım geliştirme gibi etik dışı kullanım endişelerini gündeme getirmektedir (SecurityWeek, 2025). Özetlemek gerekirse, DeepSeek; gelişmiş mantıksal muhakeme, matematiksel işlem ve kodlama becerileriyle dikkat çeken, potansiyeli yüksek bir dil modeli olarak öne çıkmaktadır. Açık ağırlık stratejisi

sayesinde geniş kullanıcı kitlesine erişim ve kapsamlı ince ayar imkânları sağlarken, eğitim süreçlerinin tam anlamıyla şeffaf olmaması bazı etik kaygılar ve yeniden üretilebilirlik sorunlarını gündeme getirmektedir. Kullanıcıların DeepSeek'i kullanırken bu potansiyel risklerin farkında olmaları önemlidir.

Bu çalışmanın temel amacı, üretken yapay zekâ modelleri olan ChatGPT ve DeepSeek'in kullanıcı etkileşimleri üzerindeki performansını ve etkisini karşılaştırmalı bir analizle değerlendirmektir. Kaggle.com'dan elde edilen ve 10.000'den fazla sentetik kullanıcı etkileşim verisini içeren bir veri seti kullanılarak, bu iki platformun yanıt doğruluğu, oturum süresi, kullanıcı memnuniyeti, farklı görevlerdeki (örneğin içerik üretimi, teknik görevler) performansları ve kullanıcı bağlılığı gibi çeşitli metrikler açısından farklılıkları incelenmiştir. Makine öğrenmesi algoritmaları ve istatistiksel testler aracılığıyla, platformların güçlü yönleri, sınırlılıkları ve kullanım bağlamına göre performans farklılıkları ortaya konulmaya çalışılmış, aynı zamanda veri gizliliği ve etik konuların önemi vurgulanmıştır.

Bu araştırma sonucunda aşağıdaki sorulara cevap verilmesi hedeflenmektedir:

1. Kullanıcı memnuniyeti ve deneyim puanları açısından ChatGPT ve DeepSeek platformları arasında istatistiksel olarak anlamlı bir fark var mıdır?
2. Cihaz türü (mobil, masaüstü, tablet, akıllı hoparlör) platformlar kullanılırken nasıl farklılaşmaktadır?
3. Makine öğrenmesi ile geliştirilen regresyon ve sınıflandırma modelleri kullanıcı memnuniyetini ve bağlılık davranışlarını ne ölçüde doğru tahmin edebilmektedir?
4. Üretken yapay zekâ platformlarının kullanıcı sadakati ve terk oranları üzerindeki etkileri nelerdir ve bu etkiler hangi kullanıcı segmentlerinde daha belirgindir?
5. Yanıt doğruluğu ile kullanıcı tarafından bildirilen deneyim ve memnuniyet skorları arasında nedensel bir ilişki kurulabilir mi?
6. Üretken yapay zekâ platformlarının kullanıcı etkileşimi performansları zamansal olarak nasıl değişim göstermektedir? Mevsimsellik veya eğilim (trend) etkisi var mıdır?
7. Farklı coğrafi bölgelerdeki kullanıcıların ChatGPT ve DeepSeek'e yönelik memnuniyet düzeyleri arasında anlamlı farklar bulunmakta mıdır?
8. Üretken yapay zekâ platformlarında (ChatGPT ve DeepSeek) kullanıcıların tercih ettiği diller ile platform seçimi arasında anlamlı bir ilişki var mıdır ve bu ilişki kullanıcı yoğunluğunun dağılımı açısından ne tür farklılıklar ortaya koymaktadır?

2.Yöntem

Bu çalışma, karşılaştırmalı bir analiz yöntemiyle, ChatGPT ve DeepSeek'in kullanıcı etkileşimleri üzerindeki performansını değerlendirmeyi amaçlayan nicel bir araştırmadır. Araştırmada, Kaggle.com'dan temin edilen bir veri seti kullanılacak ve bu veriler üzerinde makine öğrenmesi algoritmaları, özellikle regresyon, karar ağaçları ve kümeleme teknikleriyle analizler yapılacaktır. Katılımcılar, farklı coğrafi bölgelerden, çeşitli cihaz türleri (mobil, masaüstü, tablet, akıllı hoparlör) aracılığıyla ChatGPT ve DeepSeek platformlarıyla etkileşime gireceklerdir. Kullanıcı etkileşimleri, yanıt doğruluğu, oturum süresi, kullanıcı memnuniyeti (1-5 puan aralığında), sorgu türü ve cihaz türü gibi değişkenler üzerinden ölçülüp kaydedilecektir. Verilerin görselleştirilmesi, Google Looker Studio aracılığıyla yapılacak, böylece her iki yapay zeka platformunun performans karşılaştırılması görsel olarak

sunulacaktır. Elde edilen veriler, deskriptif istatistikler ile analiz edilecek ve platformlar arasındaki performans farkları, t-test veya ANOVA gibi karşılaştırmalı istatistiksel testler kullanılarak değerlendirilecektir. Çalışma, her iki platformun kullanıcı etkileşimlerine olan etkilerini anlamak için kapsamlı bir karşılaştırmalı değerlendirme sunacak ve araştırma, nicel analiz yöntemleriyle veri toplama, analiz etme ve görselleştirme aşamalarını içeren bir deneysel araştırma türüne girmektedir.

2.1. Veri Toplama Araçları

Bu çalışmada kullanılan veri seti, çevrimiçi veri platformu Kaggle.com üzerinden erişilen “DeepSeek vs ChatGPT: AI Platform Comparison” başlıklı kaynaktan temin edilmiştir.

Söz konusu veri seti, önde gelen yapay zeka modelleri olan ChatGPT (GPT-4-turbo) ve DeepSeek (DeepSeek-Chat 1.5) arasındaki performansı ve kullanıcı davranışlarını karşılaştırmak amacıyla Python programlama dili (özellikle “faker” kütüphanesi kullanılarak) ve SQL aracılığıyla yapay olarak oluşturulmuştur. Gerçekçi bir yapay zeka etkileşimi temsili sunmayı hedefleyen veri seti, Temmuz 2023 ile Şubat 2025 arasındaki yaklaşık 1.5 yıllık bir dönemi kapsamakta olup, 10.000'den fazla satır içermektedir.

Veri seti, analiz kapsamında kritik öneme sahip çeşitli kullanıcı etkileşim metriklerini, platform performans göstergelerini ve yapay zeka yanıt karakteristiklerini barındırmaktadır. Veri setinin temel özellikleri arasında şunlar yer almaktadır:

- **Zaman Serisi Uygunluğu:** Eğilim (trend) ve mevsimsellik analizleri yapılmasına olanak tanıyan ayrıntılı tarih ve saat sütunları içermektedir.
- **Karşılaştırmalı YZ Analizi:** Kullanıcı etkileşimi, platforma elde tutma oranları ve yanıt kalitesi gibi kıyaslama metriklerini karşılaştırma imkanı sunmaktadır.
- **Kullanıcı Davranışı İlgörüleri:** Oturum süreleri, kullanıcı girdilerinin metinsel karmaşıklığı ve kullanıcılar tarafından verilen derecelendirmeler gibi kullanıcı davranışlarına dair bilgiler sağlamaktadır.
- **Teknik Performans Metrikleri:** Yapay zeka yanıtlarının doğruluğu ve işlem/yanıt hızı gibi teknik performans göstergelerini içermektedir.
- **Veri Ön İşleme İmkani:** Veri temizleme ve ön işleme adımları için pratik sunmak amacıyla bilerek eklenmiş eksik (null) değerler de mevcuttur.

Bu veri seti, yapay zeka modellerinin kıyaslanması, platform performansının değerlendirilmesi, zaman serisi analizleri ve kullanıcı davranışı araştırmaları gibi konular için uygun bir yapıya sahiptir. Veri seti, araştırma, proje ve analiz amaçlı kullanımlar için MIT Lisansı ile lisanslanmıştır. Çalışmamızın kapsamına uygunluğu ve geniş veri içeriği nedeniyle bu veri seti tercih edilmiştir.

2.2. Veri Analizi

Bu çalışmada, "Üretken Yapay Zeka Performans Kıyaslaması: ChatGPT ve DeepSeek'in Kullanıcı Etkileşimleri Üzerindeki Etkisi" başlığı altında, iki farklı yapay zeka platformunun (ChatGPT ve DeepSeek) kullanıcı etkileşim verileri analiz edilmiştir. Araştırmanın bu aşamasındaki temel amaç, platformlar arası performans farklılıklarını ve kullanıcı etkileşim örüntülerini ortaya koymak için veriyi hazırlamak ve temel metrikler üzerinden tanımlayıcı analizler gerçekleştirmektir. Analizde kullanılan veri seti, ChatGPT ve DeepSeek platformlarından elde edilen kullanıcı etkileşimlerini içermektedir. Veri seti; Date, Month_Num, Weekday, AI_Platform, AI_Model_Version, Active_Users, New_Users, Churned_Users, Daily_Churn_Rate, Retention_Rate, User_ID, Query_Type, Input_Text, Input_Text_Length, Response_Tokens, Topic_Category, User_Rating, User_Experience_Score, Session_Duration_sec, Device_Type, Language, Response_Accuracy, Response_Speed_sec, Response_Time_Category, Correction_Needed, User_Return_Frequency, Customer_Support_Interactions ve Region gibi çeşitli değişkenleri kapsamaktadır. Bu değişkenler, kullanıcı demografisi, etkileşim yoğunluğu, platform performansı ve kullanıcı memnuniyeti gibi farklı boyutları temsil etmektedir. Analizlerin güvenilirliğini ve tutarlılığını sağlamak amacıyla ham veri seti üzerinde kapsamlı bir ön işleme süreci uygulanmıştır. Bu süreçte veri setindeki eksik değerler tespit edilmiş ve değişkenin tipine ve eksiklik oranına bağlı olarak uygun stratejiler kullanılarak yönetilmiştir. AI_Platform (örneğin, ChatGPT=0, DeepSeek=1 şeklinde), Weekday, Query_Type, Topic_Category, Device_Type, Language, Response_Time_Category ve Region gibi nominal veya ordinal yapıda olan kategorik değişkenler, analizlerde kullanılabilir sayısal formata dönüştürülmüştür. Bu dönüşüm için genellikle "Label Encoding" veya "One-Hot Encoding" teknikleri tercih edilmiştir. Özellikle, iki platformun doğrudan karşılaştırılabilmesi için AI_Platform sütunu sayısal olarak kodlanmıştır. Date sütunu, zaman içindeki eğilimlerin analiz edilebilmesi için datetime objelerine dönüştürülmüş; gerekirse analizlerde kullanılmak üzere Yıl, Ay, Haftanın Günü gibi ek zaman özellikleri türetilmiştir. Session_Duration_sec gibi süre bildiren metrikler, doğrudan sayısal değerler olarak kullanılmıştır. Session_Duration_sec, Input_Text_Length, Response_Tokens gibi sürekli sayısal değişkenlerdeki aykırı değerler, istatistiksel yöntemler (örneğin, IQR – Çeyrekler Açıklığı yöntemi) ve görsel araçlar (örneğin, kutu grafikleri) kullanılarak tespit edilmiştir. Tespit edilen aykırı değerlerin analiz sonuçları üzerindeki potansiyel etkisini azaltmak amacıyla, veri setinin doğasına uygun olarak kırpma (winsorizing) veya logaritmik dönüşüm gibi teknikler uygulanmıştır. Farklı ölçeklerdeki sayısal değişkenlerin analizler sırasında eşit ağırlıkta değerlendirilmesi ve bazı görselleştirme tekniklerinin performansını artırmak amacıyla, MinMaxScaler veya StandardScaler gibi ölçeklendirme yöntemleri kullanılarak özellikler belirli bir aralığa (örneğin, 0-1) veya standart normal dağılıma (ortalama=0, standart sapma=1) dönüştürülmüştür. Veri ön işleme aşamasının tamamlanmasının ardından, ChatGPT ve DeepSeek platformlarının performansını ve kullanıcı etkileşimlerini karşılaştırmak amacıyla tanımlayıcı istatistiksel analizler yapılması planlanmıştır. Bu kapsamda, her iki platform için temel performans göstergeleri (örneğin, ortalama User_Rating, ortalama Session_Duration_sec, Daily_Churn_Rate, Response_Speed_sec vb.) hesaplanmış ve karşılaştırılmıştır. Görselleştirmelerde, zaman serisi grafikleri, çubuk grafikler, pasta grafikler, dağılım grafikleri ve ısı haritaları gibi çeşitli grafik türlerinden faydalanılarak iki platform arasındaki temel farklılıklar, benzerlikler ve zaman içindeki eğilimler vurgulanacaktır.

3. Bulgular

3.1. Kullanıcı Memnuniyeti ve Deneyimi Karşılaştırmasına İlişkin Bulgular

ChatGPT ve DeepSeek platformları arasındaki kullanıcı memnuniyeti ve deneyimi karşılaştırması analizi, kullanıcı derecelendirmeleri ve deneyim puanlarına ilişkin değerli içgörüler sunmuştur. Bulguların ayrıntılı bir dökümü aşağıda sunulmaktadır.

ChatGPT oturumlarının analizinde, en yaygın derecelendirmelerin 3, 4 ve 5 olduğu gözlemlenmiştir. Özellikle 5 puanlık yüksek derecelendirmelerin, "İyi Uygulamalar", "İçerik Üretimi", "Profesyonel Yazarlık" ve "Eğitim" gibi başlıklarla ilişkilendirildiği tespit edilmiştir. Buna karşın, 3 puanlık düşük derecelendirmeler ise genellikle "İçerik Üretimi", "Eğitim" ve "Profesyonel Yazarlık" gibi konularla bağlantılı olmakla birlikte, daha kısa oturum sürelerinde yoğunlaşmaktadır. DeepSeek oturumlarında ise derecelendirmelerin ağırlıklı olarak 4 ile 5 arasında yoğunlaştığı görülmektedir. "Hata Ayıklama", "Kod Optimizasyonu", "Uygulama Yardımı" ve "Algoritma Açıklaması" gibi konularda tutarlı bir şekilde yüksek derecelendirmeler (4-5 puan) elde edilmiştir. Bununla birlikte, deneyim puanlarının 2.0 ve üzeri olduğu durumlarda derecelendirmelerde daha fazla değişkenlik gözlemlenmektedir.

ChatGPT oturumlarına ilişkin deneyim puanlarının genellikle 0.5 ile 2.2 arasında değiştiği görülmektedir. Bu bağlamda, daha yüksek deneyim puanlarının (2.0 ve üzeri) genellikle daha yüksek kullanıcı derecelendirmeleri (4-5 puan) ile pozitif bir ilişki sergilediği dikkat çekmektedir. Öte yandan, yaklaşık 0.5 ile 1.0 arasındaki daha düşük deneyim puanlarının ise sıklıkla 3 puan gibi daha düşük kullanıcı derecelendirmeleriyle örtüştüğü tespit edilmiştir. DeepSeek oturumlarında ise deneyim puanlarının genellikle daha yüksek seviyelerde (2.0 ve üzeri) olduğu gözlemlenmektedir. Bu yüksek deneyim puanları, kullanıcı derecelendirmeleriyle güçlü bir pozitif korelasyon göstermekte, özellikle 2.0'ın üzerindeki puanlarda bu ilişki daha da belirginleşmektedir.

Kullanıcı derecelendirmeleri ile kullanıcı deneyimi puanları arasında dikkate değer bir korelasyon tespit edilmiştir. Özellikle daha yüksek deneyim puanlarının, genellikle daha yüksek kullanıcı derecelendirmeleri ile pozitif bir ilişki içerisinde olduğu gözlemlenmiştir. Bu iki değişken arasındaki ilişkinin gücünü göstermek amacıyla yapılan analizde, Pearson korelasyon katsayısı yaklaşık olarak $r \approx 0.75$ olarak hesaplanmıştır. Bu sonuç, deneyim puanları ile kullanıcı derecelendirmeleri arasında güçlü bir pozitif korelasyon olduğunu ortaya koymaktadır.

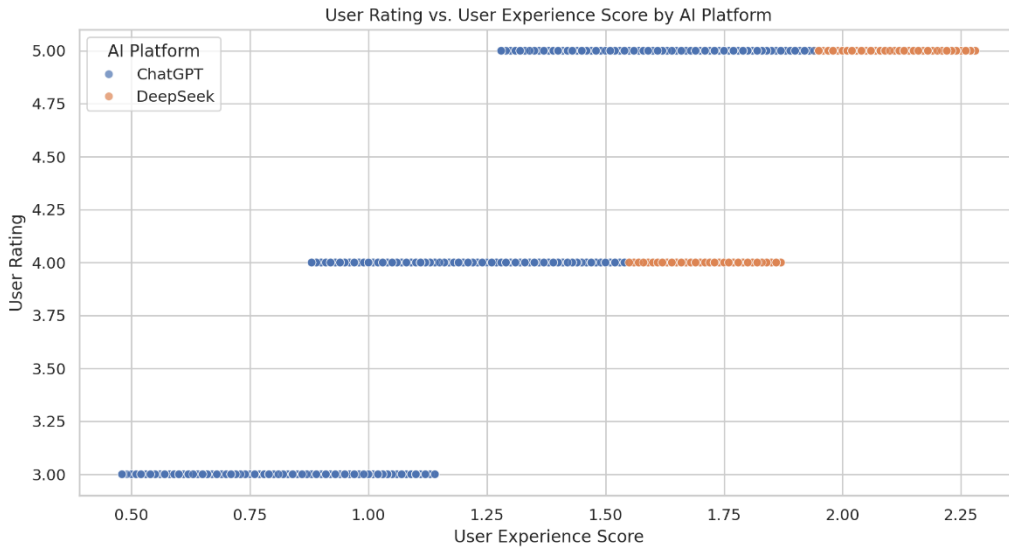
Profesyonel yazarlık oturumlarında, kullanıcıların ortalama deneyim puanlarının nispeten daha düşük olduğu (yaklaşık 1.0-1.5 aralığında) belirlenmiştir. Bu oturumlarda kullanıcı derecelendirmelerinin çoğunlukla 3 ile 4 arasında yoğunlaştığı gözlemlenmiştir. İçerik üretimi oturumlarında ise deneyim puanlarının orta ile yüksek düzeyde (yaklaşık 1.5-2.0 aralığında) olduğu ve kullanıcı derecelendirmelerinin genellikle 4 veya 5 olma eğiliminde bulunduğu dikkat çekmektedir. Hata ayıklama ve uygulama yardımı oturumlarında deneyim puanlarının sıklıkla 2.0'ı aştığı ve kullanıcı derecelendirmelerinin ağırlıklı olarak 4 ile 5 arasında olduğu görülmektedir. Bu durum, bu alanlarda kullanıcı deneyimlerinin pozitif bir algı yarattığını ortaya koymaktadır.

Mobil cihazlarda yapılan oturumlarda, deneyim puanlarında biraz daha fazla değişkenlik gözlemlenmektedir. Buna rağmen, kullanıcı derecelendirmeleri genellikle yüksek düzeylerde (3-5 puan) seyretmektedir. Dizüstü ve masaüstü bilgisayar oturumlarında ise ortalama deneyim puanlarının daha yüksek olduğu (yaklaşık 2.0 ve üzeri)

ve kullanıcı derecelendirmelerinin tutarlı bir şekilde yüksek düzeylerde (4-5 puan) gerçekleştiği tespit edilmiştir. Tablet ve akıllı hoparlör oturumlarında ise ortalama deneyim puanlarının nispeten daha düşük olduğu dikkat çekmektedir. Buna karşın, kullanıcı derecelendirmeleri yine 3-5 arasında yoğunlaşmakla birlikte, daha fazla dalgalanma ve değişkenlik göstermektedir.

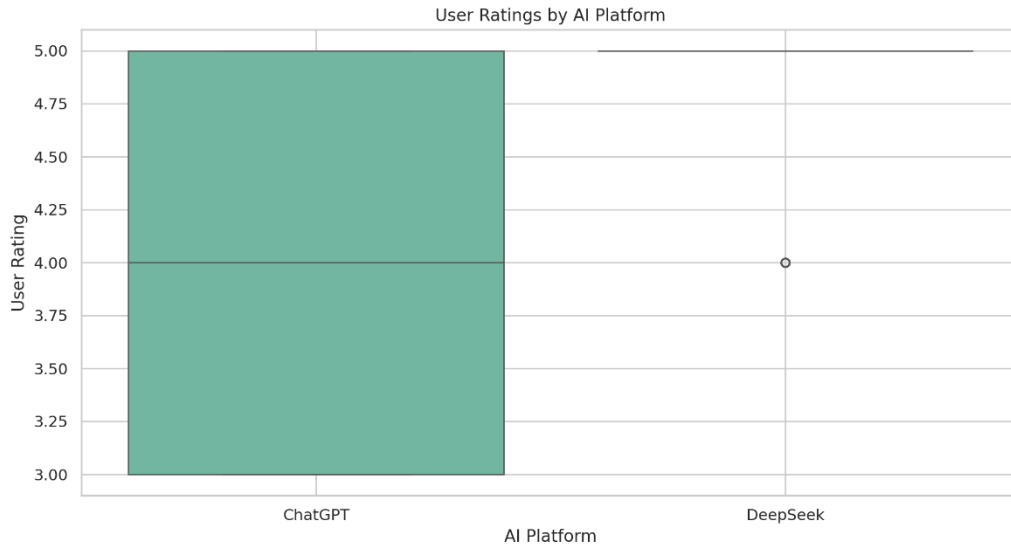
Genel olarak, daha yüksek kullanıcı deneyimi puanları, daha yüksek kullanıcı derecelendirmeleri ile korelasyon göstermektedir. DeepSeek kullanıcıları, ChatGPT kullanıcılarına kıyasla daha yüksek deneyim puanları ve derecelendirmeler bildirme eğilimindedir. "Hata Ayıklama", "Uygulama Yardımı" ve "Kod Optimizasyonu" gibi teknik konular, deneyim puanları ve kullanıcı derecelendirmeleri arasında güçlü pozitif ilişkiler sergilemektedir. Cihaz türü, deneyim puanlarını kısmen etkilemekte olup, dizüstü/masaüstü bilgisayarlar daha iyi kullanıcı deneyimleriyle ilişkilendirilmektedir.

Şekil 1'de sunulan saçılım grafiği, kullanıcı oyları ile kullanıcı deneyimi skorları arasındaki ilişkiyi platform bazında göstermektedir. Grafikte yer alan her bir nokta, analiz kapsamındaki bir oturumu temsil etmekte olup, noktalar kullanılan yapay zeka platformuna (ChatGPT veya DeepSeek) göre farklı renklendirilmiştir.



Şekil 1. Kullanıcı Oyları ve Kullanıcı Deneyimi Skoru İlişkisi

Şekil 2 ise, her bir yapay zeka platformu için kullanıcı oylarının dağılımını özetleyen bir kutu grafiğidir. Bu grafik, platformlara göre kullanıcı oylarının medyanını, çeyrekler açıklığını ve genel değişkenliğini vurgulamaktadır.



Şekil 2. Kullanıcı Oyları Dağılımı

Elde edilen bulgular incelendiğinde:

Saçılım grafiği (Şekil 1), kullanıcı deneyimi skorları ile kullanıcı oyları arasında pozitif bir korelasyon olduğunu, özellikle DeepSeek platformu için bu ilişkinin daha belirgin olduğunu göstermektedir. DeepSeek'in hem kullanıcı deneyimi skorlarının hem de kullanıcı oylarının genel olarak daha yüksek değerlerde yoğunlaştığı gözlenmiştir. Kutu grafiği (Şekil 2) bulgularını destekler nitelikte, DeepSeek platformunun ChatGPT'ye kıyasla genel olarak daha yüksek kullanıcı oyları aldığı ve bu oylardaki değişkenliğin (yayılımın) daha az olduğu tespit edilmiştir.

3.2. ChatGPT ve DeepSeek Platformlarının Kullanıcı Oyu ve Deneyimi Skorları Açısından İstatistiksel Karşılaştırılması

Bu bölümde, farklı yapay zeka platformlarının (ChatGPT ve DeepSeek) Kullanıcı Derecelendirmeleri ve Kullanıcı Deneyimi Puanları ortalamaları arasındaki farkların istatistiksel anlamlılığını değerlendirmek amacıyla Bağımsız Örneklem T-Testi ve Tek Yönlü Varyans Analizi (ANOVA) uygulanmıştır.

Her iki değişken için platformlar arası ortalamaların karşılaştırılması amacıyla yapılan bağımsız örneklem T-testi sonuçları Tablo 1'de sunulmuştur.

Tablo 1. T-testi Sonuçlarına Göre Kullanıcı Derecelendirmesi ve Deneyimi

Ölçüm Değişkeni	t(sd)	p-değeri
Kullanıcı Derecelendirmesi	-65.37 (9998)	< .001
Kullanıcı Deneyimi Puanı	-142.12 (9998)	< .001

Not. Negatif t-istatistik değerleri, ikinci grubun ortalamasının daha yüksek olduğunu göstermektedir. $p < .001$ değeri istatistiksel olarak yüksek derecede anlamlı sonuçlara işaret eder.

T-testi sonuçları, hem Kullanıcı Derecelendirmeleri hem de Kullanıcı Deneyimi Puanları için istatistiksel olarak anlamlı düzeyde düşük p-değerleri ($< 0,001$) vermiştir. Bu durum, ChatGPT ve DeepSeek platformlarının ilgili

ortalama değerleri arasındaki gözlemlenen farkların istatistiksel olarak anlamlı olduğunu göstermektedir. Ortalama Kullanıcı Derecelendirmesi ChatGPT için yaklaşık 4,00 iken, DeepSeek için yaklaşık 4,80 olarak hesaplanmıştır. Benzer şekilde, ortalama Kullanıcı Deneyimi Puanı ChatGPT için yaklaşık 1,23 iken, DeepSeek için yaklaşık 2,03'tür. Bu bulgular, kullanıcıların DeepSeek platformunu ChatGPT'ye kıyasla genel olarak daha yüksek derecelendirdiğini ve DeepSeek'in daha olumlu bir kullanıcı deneyimi sunduğunu işaret etmektedir.

Platformun (ChatGPT vs. DeepSeek) Kullanıcı Derecelendirmeleri ve Kullanıcı Deneyimi Puanları üzerindeki etkisini incelemek için yapılan ANOVA sonuçları Tablo 2'de gösterilmiştir.

Tablo 2. ANOVA Sonuçlarına Göre Kullanıcı Derecelendirmesi ve Deneyimi

Ölçüm Değişkeni	F-İstatistik	p-değeri
Kullanıcı Derecelendirmesi	4273.78	< .001
Kullanıcı Deneyimi Puanı	20199.47	< .001

Not. F-İstatistik değerleri ve p-değerleri, gruplar arasında anlamlı farklar olduğunu göstermektedir. $p < .001$ değeri, istatistiksel olarak yüksek derecede anlamlı farklara işaret eder.

ANOVA sonuçları, hem Kullanıcı Derecelendirmeleri ($F(1, 9998) = 4273,78, p < 0,001$) hem de Kullanıcı Deneyimi Puanları ($F(1, 9998) = 20199,47, p < 0,001$) açısından platformlar arasında istatistiksel olarak anlamlı farklılıklar olduğunu doğrulamaktadır. Yüksek F-istatistik değerleri, yapay zeka platformunun söz konusu değişkenler üzerinde güçlü bir etkiye sahip olduğunu göstermektedir. Elde edilen p-değerlerinin belirlenen anlamlılık düzeyi olan 0,05'ten düşük olması nedeniyle, sıfır hipotezi reddedilmekte ve ChatGPT ile DeepSeek arasında kullanıcı memnuniyeti düzeylerinde önemli farklılıklar olduğu sonucuna varılmaktadır. Gözlemlenen farklılıkların rastgele şanstın kaynaklanma olasılığı çok düşüktür.

3.3. Regresyon Modeli Sonuçları

Bu bölümde, Kullanıcı Derecelendirmelerindeki değişkenliği açıklamak amacıyla uygulanan farklı regresyon modellerinin performans sonuçları sunulmaktadır. Doğrusal Regresyon, Rastgele Orman ve Gradyan Artırma modelleri test edilmiş ve modellerin başarıları Ortalama Karesel Hata (OKH) ve R^2 Skoru metrikleri üzerinden değerlendirilmiştir. Modellere ait performans sonuçları Tablo 3'te gösterilmiştir.

Tablo 3. Modellere Ait Performans Sonuçları

Model	Ortalama Karesel Hata (OKH)	R^2 Skoru
Doğrusal Regresyon	0.4089	0.2467
Rastgele Orman	0.4090	0.2466
Gradyan Artırma	0.3825	0.2954

Not. Ortalama Karesel Hata (OKH), modelin tahminlerinin gerçek değerlere olan ortalama karesel uzaklığını temsil eder. R^2 Skoru, modelin veri üzerindeki açıklayıcılığını ölçer.

Elde edilen regresyon sonuçları incelendiğinde, test edilen üç model arasında en iyi performansı Gradyan Artırma modelinin sergilediği görülmektedir. Gradyan Artırma modeli, en düşük Ortalama Karesel Hataya (0.3825) ve en

yüksek R^2 skoruna (0.2954) sahip olup, bu durum modelin Kullanıcı Derecelendirmelerindeki varyansın diğer modellere kıyasla daha büyük bir oranını açıkladığını düşündürmektedir.

Bununla birlikte, elde edilen R^2 skorlarının (hepsi 0.3'ün altında olması) tüm modeller için nispeten düşük olduğu dikkat çekmektedir. Bu durum, modellerin Kullanıcı Derecelendirmelerindeki toplam varyansın büyük bir kısmını açıklayamadığını göstermektedir. Düşük R^2 değerleri, mevcut özellik setinde yakalanamayan, kullanıcı derecelendirmelerini etkileyen başka önemli faktörlerin olabileceğine işaret etmektedir.

3.4. Sınıflandırma Modelleri Sonuçları

Bu çalışmada, sınıflandırma görevi için Lojistik Regresyon, Rastgele Orman Sınıflandırıcısı (Random Forest Classifier) ve Destek Vektör Makinesi (Support Vector Machine - SVM) olmak üzere üç farklı sınıflandırma modeli uygulanmıştır. Modellerin performansları, Doğruluk (Accuracy), Kesinlik (Precision), Duyarlılık (Recall) ve F1 Skoru metrikleri kullanılarak değerlendirilmiştir. Her bir model için genel doğruluk değeri ve her bir sınıf ('Excellent', 'Fair', 'Good') bazında elde edilen detaylı performans metrikleri aşağıda sunulmuştur.

Yapılan sınıflandırma analizleri sonucunda, lojistik regresyon modeli %57.51 doğruluk oranı ile en yüksek performansı göstermiştir. Rastgele orman (random forest) algoritması %57.25 doğruluk oranı ile lojistik regresyona oldukça yakın bir performans sergilemiştir. Buna karşılık, destek vektör makineleri (SVM) %52.88 doğruluk oranı ile diğer iki modele kıyasla daha düşük bir başarı elde etmiştir. Bu sonuçlar, kullanılan veri seti ve özellikler doğrultusunda lojistik regresyon modelinin sınıflandırma görevinde diğer modellere göre daha etkili olabileceğini göstermektedir.

Tablo 4. Sınıflandırma Modellerinin Performans Metrikleri

Sınıf	Model	Kesinlik	Duyarlılık	F1 Skoru	Destek
Mükemmel	Lojistik Regresyon	0.6497	0.8418	0.7334	1018
	Rastgele Orman	0.6934	0.7908	0.7389	1018
	SVM	0.5288	1.0000	0.6918	1018
Orta	Lojistik Regresyon	0.3386	0.1448	0.2028	297
	Rastgele Orman	0.2672	0.2088	0.2344	297
	SVM	0.0000	0.0000	0.0000	297
İyi	Lojistik Regresyon	0.4322	0.3393	0.3802	610
	Rastgele Orman	0.4417	0.3852	0.4116	610
	SVM	0.0000	0.0000	0.0000	610
Makro Ort.	Lojistik Regresyon	0.4735	0.4420	0.4388	1925
	Rastgele Orman	0.4674	0.4616	0.4616	1925
	SVM	0.1763	0.3333	0.2306	1925
Ağırlık. Ort.	Lojistik Regresyon	0.5328	0.5751	0.5396	1925

Rastgele Orman	0.5479	0.5725	0.5573	1925
SVM	0.2797	0.5288	0.3659	1925

Not. Kesinlik (Precision), Duyarlılık (Recall), ve F1 Skoru, sınıflandırma modellerinin doğruluk ölçütlerini temsil eder. Destek, her sınıftaki örnek sayısını göstermektedir.

Elde edilen sınıflandırma sonuçları incelendiğinde Lojistik Regresyon Modelinin genel olarak makul bir performans sergilediği görülmüştür. Bu model özellikle 'Excellent' (Mükemmel) kategorisini tahmin etmede daha başarılı iken, 'Fair' (Orta) ve 'Good' (İyi) kategorilerinde performansı daha düşük kalmıştır.

Rastgele Orman Modeli, Lojistik Regresyon modeline benzer bir performans göstermiştir. Genel doğruluk değeri hafifçe daha düşük olmasına rağmen, Kesinlik ve Duyarlılık gibi diğer performans metrikleri karşılaştırılabilir düzeydedir. SVM modelinin ise diğer modellere kıyasla genel doğruluğunun daha düşük olduğu ve özellikle 'Fair' ve 'Good' kategorilerini etkili bir şekilde tahmin edemediği tespit edilmiştir. Bu sonuçlar, kullanılan modellerin özellikle yaygın olan 'Excellent' sınıfını tahmin etmede bir miktar başarı gösterdiğini, ancak daha az temsil edilen sınıflarda (Fair, Good) sınıflandırma performansının belirgin şekilde düştüğünü ortaya koymaktadır.

3.5.Kullanıcı Etkileşimi ve Sadakati Karşılaştırması

Bu bölümde, kullanıcı etkileşimi ve bağlılığına ilişkin farklı yönleri modellemek amacıyla geliştirilen makine öğrenmesi modellerinin sonuçları sunulmaktadır. Analizler hem sınıflandırma hem de regresyon görevlerini içermektedir.

Günlük Terk Oranı (Daily_Churn_Rate) değişkeni, belirlenen bir eşik değeri (0.1) kullanılarak ikili bir sınıflandırma problemine dönüştürülmüştür. Bu eşik değere göre, günlük terk oranı 0.1'i aşan kullanıcılar 'terk etmiş' (1), aşmayanlar ise 'terk etmeyen' (0) olarak sınıflandırılmıştır. Bu ikili hedef değişkeni tahmin etmek üzere bir Rastgele Orman Sınıflandırıcısı (Random Forest Classifier) eğitilmiş ve modelin değerlendirme sonuçları aşağıda sunulmuştur. Modelin performansına ilişkin detaylı sınıflandırma metrikleri — kesinlik (precision), duyarlılık (recall), F1 skoru ve destek (support) — Tablo 5'te sunulmuştur. Bu metrikler, modelin sınıflandırma başarısını daha kapsamlı bir şekilde değerlendirmek amacıyla raporlanmıştır. Rastgele Orman Sınıflandırıcısının genel doğruluk değeri ise %100 olarak elde edilmiştir.

Tablo 5. Model Performansına Göre Sınıf 0 ve Ortalama Değerler

Sınıf / Ortalama	Kesinlik	Duyarlılık	F1 Skoru	Destek
0	1.000	1.000	1.000	1925
Makro Ortalama	1.000	1.000	1.000	1925
Ağırlıklı Ortalama	1.000	1.000	1.000	1925

Not. Tüm metriklerin 1.000 olması, sınıflandırma modelinin hata yapmadan tüm örnekleri doğru tahmin ettiğini göstermektedir.

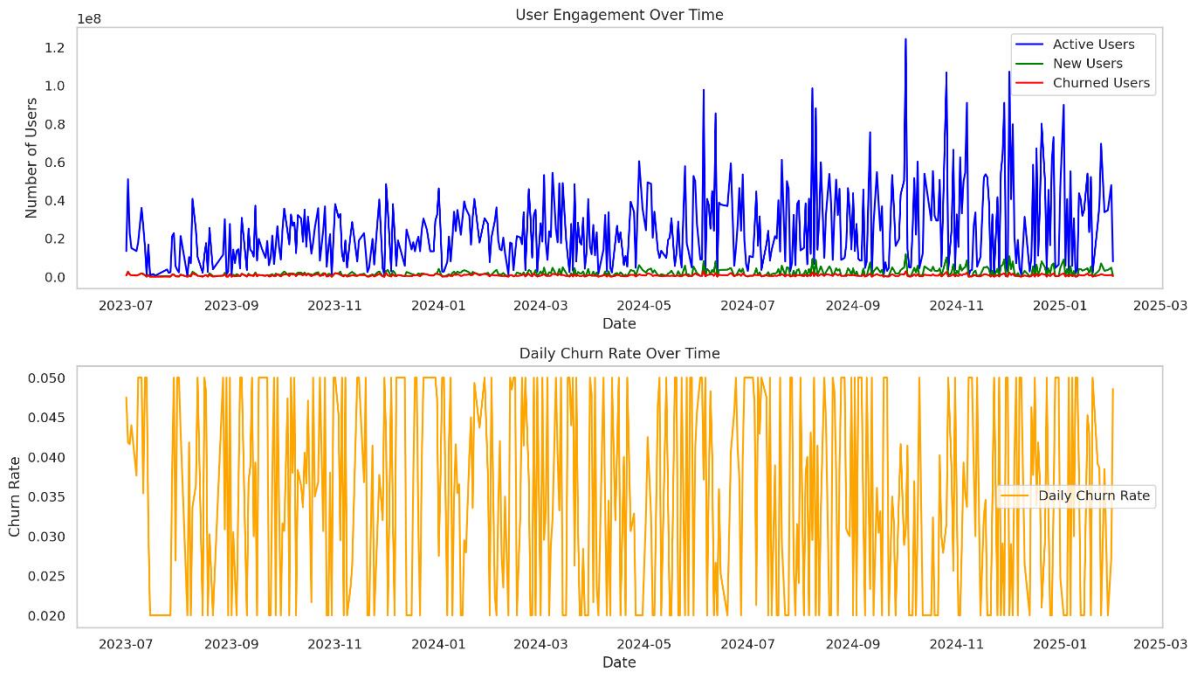
Elde edilen sonuçlara göre, Rastgele Orman sınıflandırma modeli test setinde %100 doğruluk elde etmiştir, bu da modelin tahminlerinin mükemmel olduğunu göstermektedir. Sunulan verilere göre, '0' sınıfı (terk etmeyen) için

Kesinlik, Duyarlılık ve F1 Skoru değerlerinin tamamı 1.0'dır. Bu durum, modelin test setindeki 'terk etmeyen' kullanıcıları hatasız bir şekilde belirleyebildiğini işaret etmektedir.

Sürekli değişkenler olan Oturum Süresi (Session_Duration_sec) ve Kullanıcı Dönüş Sıklığı (User_Return_Frequency) gibi kullanıcı davranışlarını tahmin etmek amacıyla bir Rastgele Orman Regresyoncusu (Random Forest Regressor) eğitilmiştir. Regresyon modellerinin performansını değerlendirmek için Ortalama Karesel Hatanın Karekökü (Root Mean Squared Error - RMSE) metriği kullanılmıştır.

Regresyon modellerine ait KOKH (RMSE) sonuçları aşağıda sunulmuştur:

- **Oturum Süresi KOKH (RMSE):** Yaklaşık 14.57 saniye
- **Kullanıcı Dönüş Sıklığı KOKH (RMSE):** Yaklaşık 3.16



Şekil 3. Oturum Süresi ve Kullanıcı Dönüş Sıklığı

Elde edilen KOKH (RMSE) değerleri, oturum süresi tahmini için yaklaşık 14.57 saniye ve kullanıcı dönüş sıklığı tahmini için yaklaşık 3.16 olarak hesaplanmıştır. KOKH değeri, tahmin edilen değerlerin gerçek değerlerden ortalama sapmasını gösterir; dolayısıyla daha düşük KOKH değerleri daha iyi performansı işaret eder. Bu sonuçlar, Rastgele Orman Regresyon modellerinin oturum süresi ve kullanıcı dönüş sıklığını tahmin etmede makul derecede etkili olduğunu düşündürmektedir.

3.6.Yanıt Kalitesi ve Performansı Karşılaştırması

Yanıt doğruluğunu tahmin etmek amacıyla uygulanan regresyon modelinin performansını değerlendirmek için Kök Ortalama Karesel Hata (RMSE) metriği kullanılmıştır. Yanıt doğruluğu tahmini için elde edilen KOKH (RMSE) değeri yaklaşık **0.064**'tür. Bu düşük ortalama sapma değeri, modelin yanıt doğruluğu skorlarını tahmin etmede iyi bir performans sergilediğini düşündürmektedir.

Yanıtın düzeltmeye ihtiyacı olup olmadığını (0 = Hayır, 1 = Evet) sınıflandırmak amacıyla uygulanan modelin performansını gösteren sınıflandırma raporu aşağıdadır:

Tablo 6. Sınıflandırma Performans Metrikleri (Sınıf 0 ve 1)

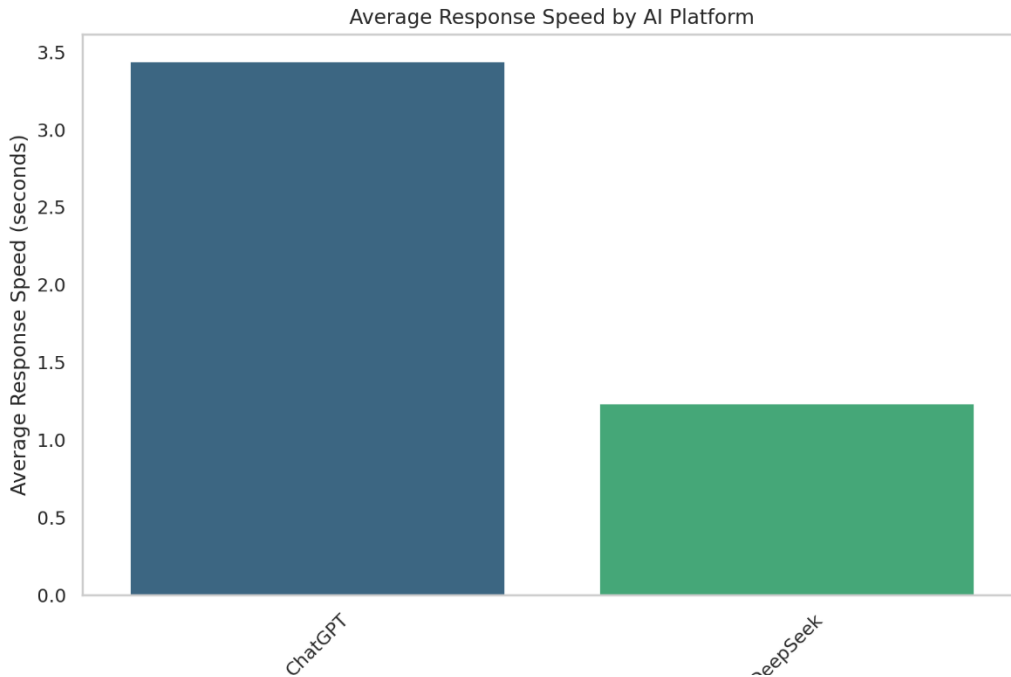
Sınıf / Ortalama	Kesinlik	Duyarlılık	F1 Skoru	Destek
0	0.863	0.913	0.888	1655
1	0.177	0.115	0.139	270
Makro Ortalama	0.520	0.514	0.513	1925
Ağırlıklı Ortalama	0.767	0.801	0.783	1925

Not. Makro ortalamalar sınıflar arasında eşit ağırlıkla hesaplanırken, ağırlıklı ortalamalar sınıf destek (örnek sayısı) oranında ağırlıklandırılarak elde edilir.

Sınıflandırma modeli, düzeltme gerektirmeyen yanıtlar (Sınıf 0) için yaklaşık **%86.34** Kesinlik ve yaklaşık **%91.30** Duyarlılık sergilemiştir. Düzeltme gerektiren yanıtlar (Sınıf 1) için ise Kesinlik yaklaşık **%17.71**, Duyarlılık ise yaklaşık **%11.48** olarak bulunmuştur. Bu metrikler, modelin düzeltme gerektirmeyen durumları başarılı bir şekilde belirleyebildiğini, ancak düzeltme gerektiren durumları belirlemede zorlandığını göstermektedir. Modelin genel doğruluğu ise yaklaşık **%80.10**'dur.

Yanıt kalitesi ve performansına ilişkin analizler, farklı yapay zeka platformlarının etkinliği hakkında değerli içgörüler sağlamıştır. Yanıt doğruluğu ve düzeltme gereksinimi değerlendirmelerinin ardından, çeşitli platformlar arasındaki ortalama yanıt hızı görselleştirilmiştir.

Aşağıdaki çubuk grafik, her bir yapay zekâ platformu için ortalama yanıt hızını (saniye olarak) göstermektedir:

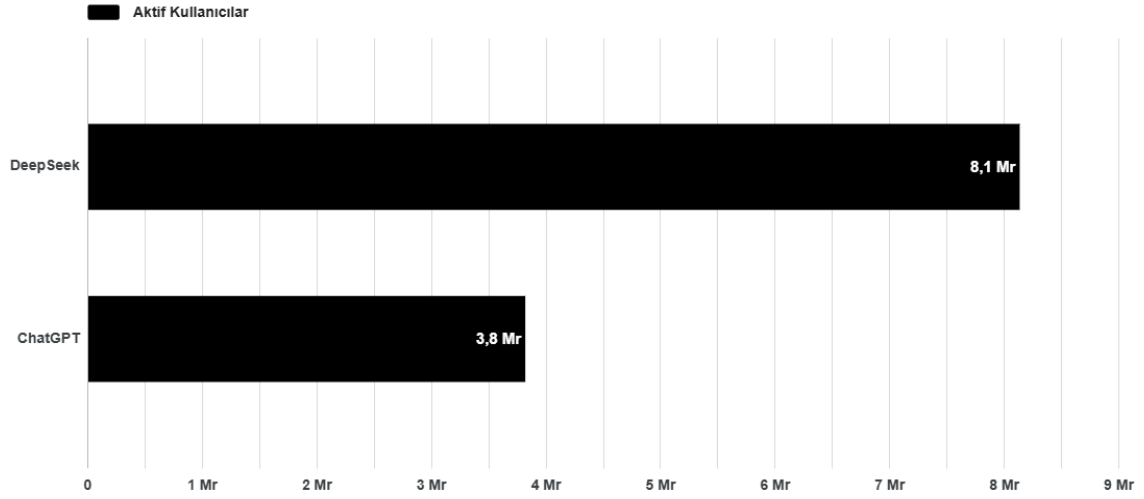


Şekil 4. ChatGPT ve Deepseek İçin Ortalama Yanıt Hızı (Saniye Olarak)

Bu grafik (Şekil 4), her bir yapay zekâ platformunun kullanıcı sorgularına ne kadar hızlı yanıt verdiğini karşılaştırmamıza olanak tanımaktadır. Bu bilgi, kullanıcı memnuniyeti için kritik öneme sahip olan yanıt süresi açısından hangi platformların daha iyi performans gösterdiğini belirlemeye yardımcı olabilir.

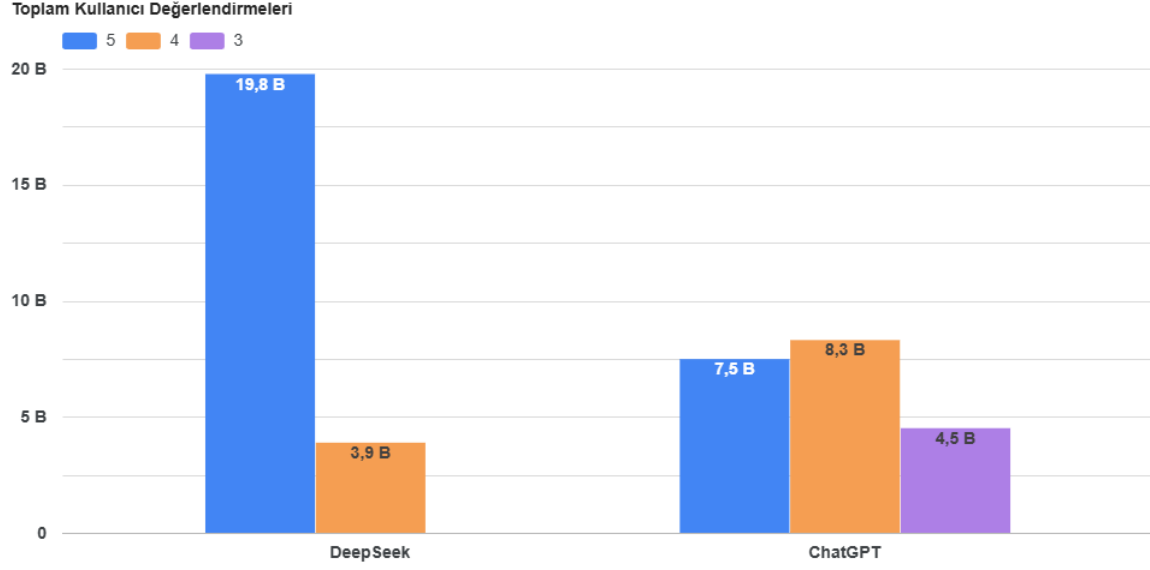
3.7. Aktif Kullanıcı Sayısı ve Toplam Kullanıcı Değerlendirmeleri

Şekil 5, DeepSeek ve ChatGPT platformlarının aktif kullanıcı sayılarını karşılaştırmaktadır. Elde edilen verilere göre, DeepSeek platformu 8.1 milyon aktif kullanıcıya sahipken, ChatGPT'nin aktif kullanıcı sayısı 3.8 milyon olarak belirlenmiştir. Bu sonuç, DeepSeek'in ChatGPT'ye kıyasla daha geniş bir aktif kullanıcı kitlesine ulaştığını göstermektedir.



Şekil 5. Aktif Kullanıcı Sayısı

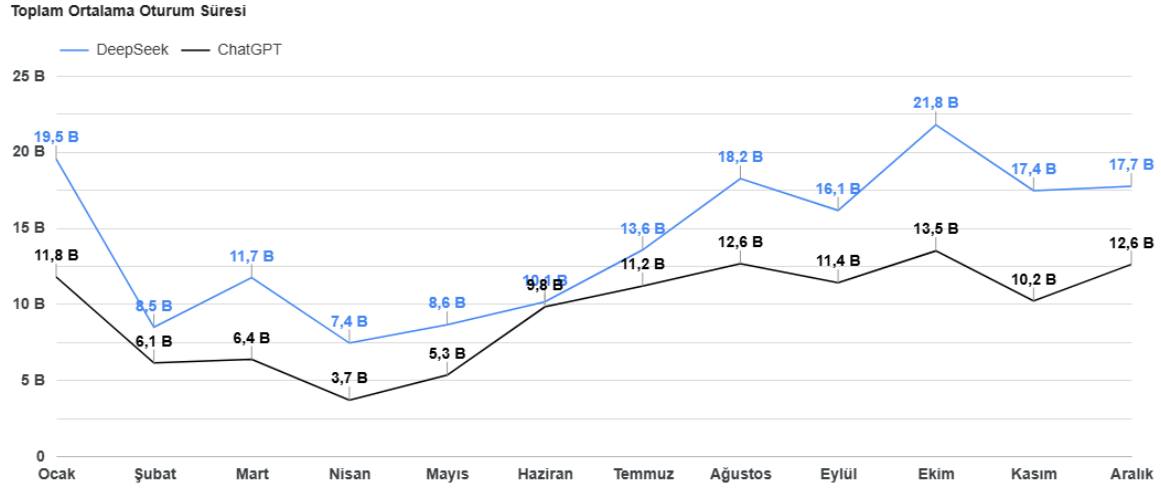
Şekil 6, DeepSeek ve ChatGPT platformlarının toplam kullanıcı değerlendirmelerini karşılaştırmaktadır. Elde edilen verilere göre, DeepSeek platformu 5 üzerinden ortalama 19.8 milyar, 4 üzerinden 3.9 milyar değerlendirme almıştır. ChatGPT ise 5 üzerinden 7.5 milyar, 4 üzerinden 8.3 milyar ve 3 üzerinden 4.5 milyar değerlendirme almıştır. Bu sonuçlar, her iki platformun da yüksek kullanıcı etkileşimi ve geri bildirimi aldığını göstermektedir. DeepSeek'in 5 yıldız üzerinden aldığı yüksek değerlendirme sayısı dikkat çekicidir.



Şekil 6. Toplam Kullanıcı Değerlendirmeleri Karşılaştırması

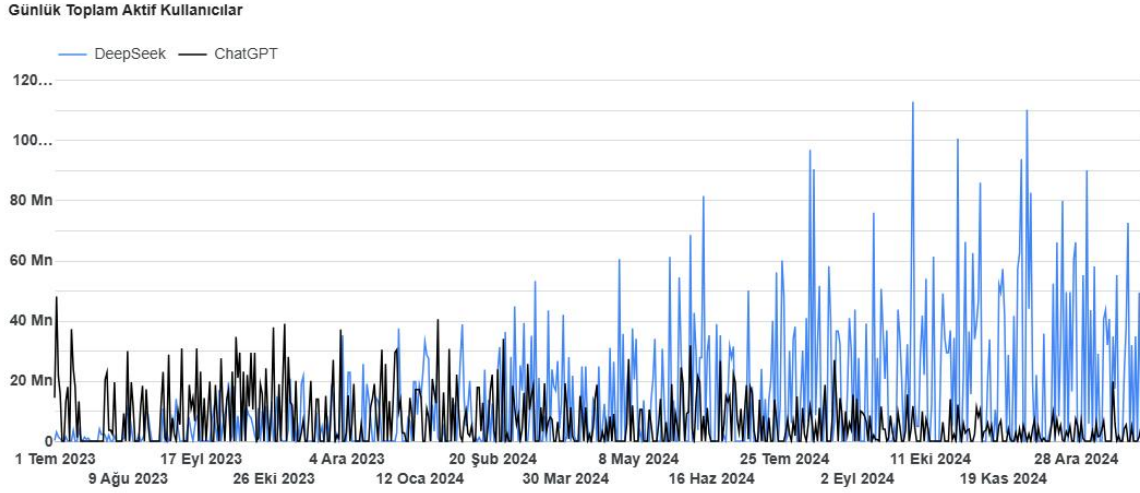
3.8. Toplam Ortalama Oturum Süresi ve Günlük Toplam Aktif Kullanıcılar

Şekil 7, DeepSeek ve ChatGPT platformlarının aylık bazda toplam ortalama oturum sürelerini milyar cinsinden göstermektedir. Yıl boyunca her iki platformda da oturum sürelerinde dalgalanmalar gözlemlenmektedir.



Şekil 7. Toplam Ortalama Oturum Süresi Karşılaştırması

DeepSeek'in oturum süreleri genellikle ChatGPT'den daha yüksek seyretmekle birlikte, bazı aylarda (örneğin Şubat, Nisan) ChatGPT'nin oturum sürelerinin DeepSeek'i geçtiği görülmektedir. Özellikle Ekim ayında DeepSeek'in oturum süresinde belirgin bir artış yaşanırken, Kasım ayında düşüş gözlemlenmiştir. ChatGPT'nin oturum süreleri ise daha istikrarlı bir seyir izlemekle birlikte, Aralık ayında belirgin bir artış göstermiştir. Bu bulgular, kullanıcıların platformlarla etkileşim sürelerinin zaman içinde değişiklik gösterebildiğini ve platform özelliklerinin bu değişikliklerde etkili olabileceğini düşündürmektedir.

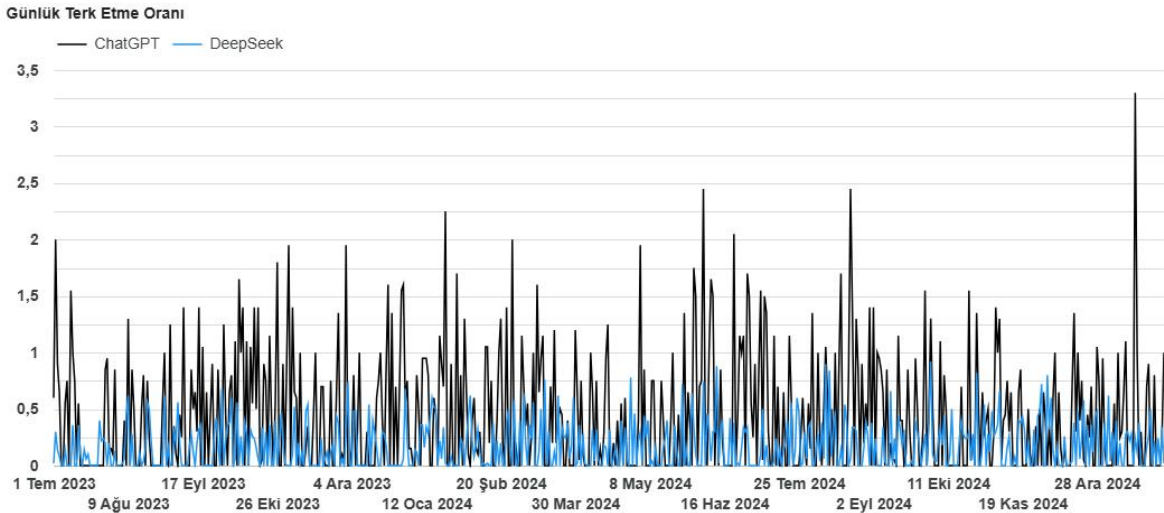


Şekil 8. Günlük Toplam Aktif Kullanıcı Sayısı Karşılaştırması

Şekil 8, DeepSeek ve ChatGPT platformlarının günlük toplam aktif kullanıcı sayılarını milyon cinsinden göstermektedir. Grafik incelendiğinde, DeepSeek'in günlük aktif kullanıcı sayısında ChatGPT'ye kıyasla genel olarak daha yüksek ve daha dalgalı bir seyir izlendiği görülmektedir. Özellikle 2024 yılının ortalarından itibaren DeepSeek'in aktif kullanıcı sayısında belirgin pikler yaşanmıştır. ChatGPT'nin günlük aktif kullanıcı sayısı ise daha stabil bir görünüm sergilemekte ve DeepSeek'in zirve yaptığı dönemlerde bile daha düşük seviyelerde kalmaktadır. Bu durum, DeepSeek'in zaman zaman önemli ölçüde daha fazla kullanıcı etkileşimi yarattığını, ancak bu etkileşimin ChatGPT'ye göre daha değişken olabileceğini düşündürmektedir.

3.9.Günlük Terk Etme Oranı ve Ortalama Oturum Süresi

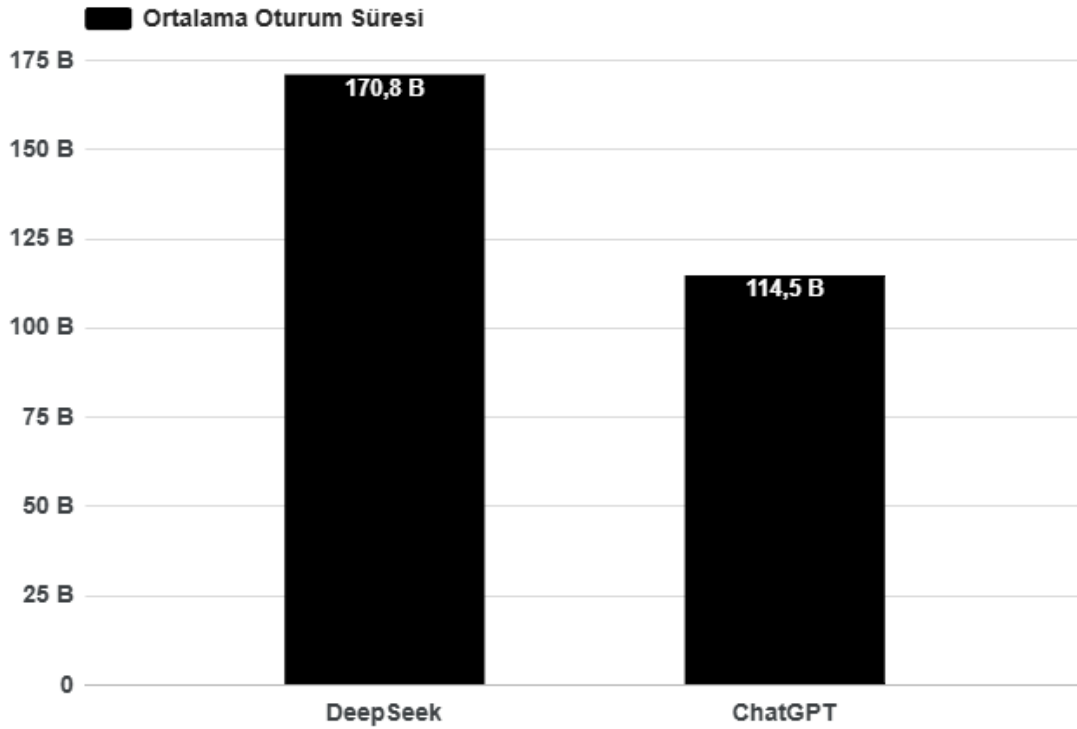
Şekil 9, DeepSeek ve ChatGPT platformlarının günlük terk etme oranlarını göstermektedir. Grafik incelendiğinde, her iki platformda da zaman zaman belirgin yükselişler olsa da, genel olarak DeepSeek'in terk etme oranının ChatGPT'ye göre daha düşük seyrettiği görülmektedir.



Şekil 9. Günlük Terk Etme Oranı Karşılaştırması

ChatGPT'nin terk etme oranında özellikle 2023'ün sonlarına doğru ve 2024'ün ilk aylarında daha yüksek değerler gözlemlenirken, DeepSeek'te bu dönemde daha stabil ve düşük bir oran söz konusudur. 2024 yılı boyunca her iki platformda da terk etme oranlarında inişli çıkışlı bir seyir izlenmekle birlikte, DeepSeek'in genel ortalamasının ChatGPT'den daha iyi olduğu söylenebilir. Bu bulgu, kullanıcıların DeepSeek ile etkileşimlerinden ChatGPT'ye göre daha fazla memnun kaldıkları veya platformun kullanıcı beklentilerini daha iyi karşıladığı şeklinde yorumlanabilir.

Şekil 10, DeepSeek ve ChatGPT platformlarının ortalama oturum sürelerini milyar cinsinden karşılaştırmaktadır.

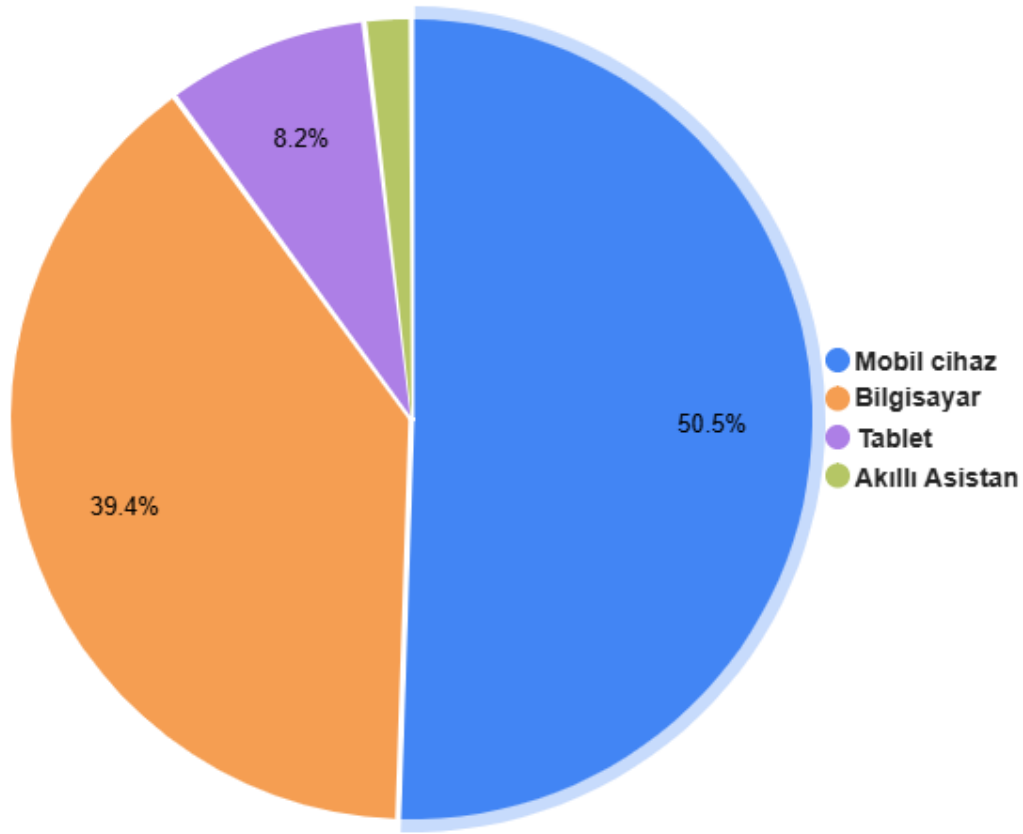


Şekil 10. Ortalama Oturum Süresi Karşılaştırması

Elde edilen verilere göre, DeepSeek platformunun ortalama oturum süresi 170.8 milyar olarak belirlenirken, ChatGPT'nin ortalama oturum süresi 114.5 milyardır. Bu sonuç, kullanıcıların DeepSeek ile etkileşimlerinin ChatGPT'ye kıyasla önemli ölçüde daha uzun sürdüğünü göstermektedir. Bu durum, DeepSeek'in kullanıcıları daha fazla etkileşimde tutabilen veya daha karmaşık sorgulara yanıt verebilen özelliklere sahip olabileceği şeklinde yorumlanabilir.

3.10. Kullanıcıların Cihaz Tercihleri ve Dil Bazında Kullanıcı Dağılımı

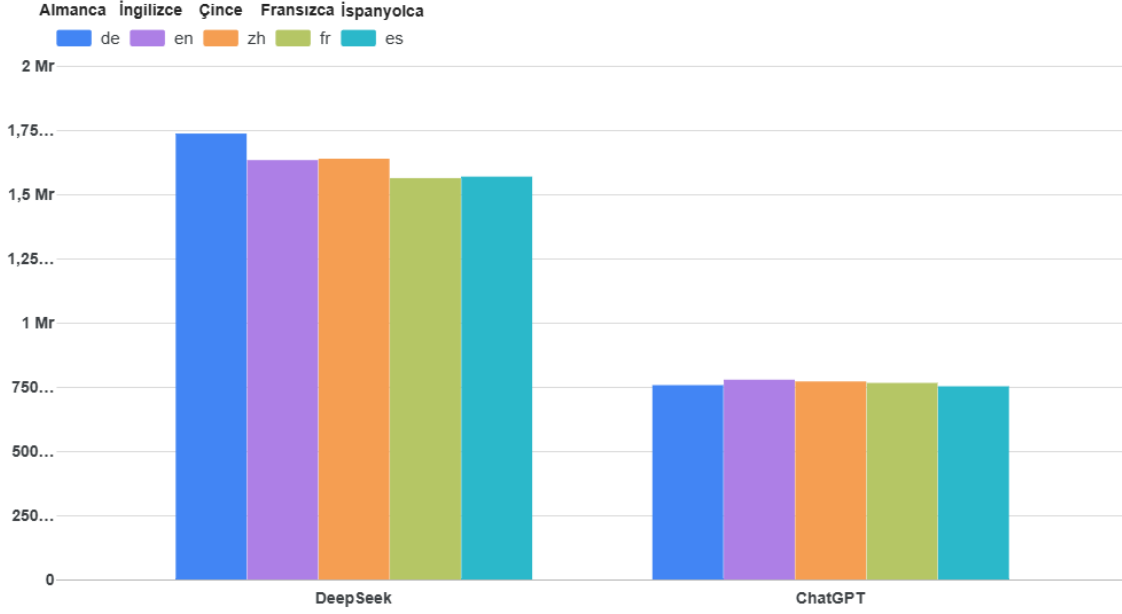
Şekil 11, kullanıcıların platforma erişim için kullandıkları cihazların dağılımını göstermektedir.



Şekil 11. Kullanıcıların Cihaz Tercihleri

Elde edilen verilere göre, kullanıcıların büyük çoğunluğu (%50.5) platforma mobil cihazlar aracılığıyla erişmektedir. Bilgisayar (%39.4) ise en sık kullanılan ikinci erişim yöntemi olarak belirlenmiştir. Tablet (%8.2) ve akıllı asistanlar (%1.9) üzerinden erişim oranları ise daha düşüktür. Bu bulgu, platforma erişimde mobil cihazların önemli bir rol oynadığını ve mobil uyumluluğun kullanıcı deneyimi açısından kritik olduğunu göstermektedir.

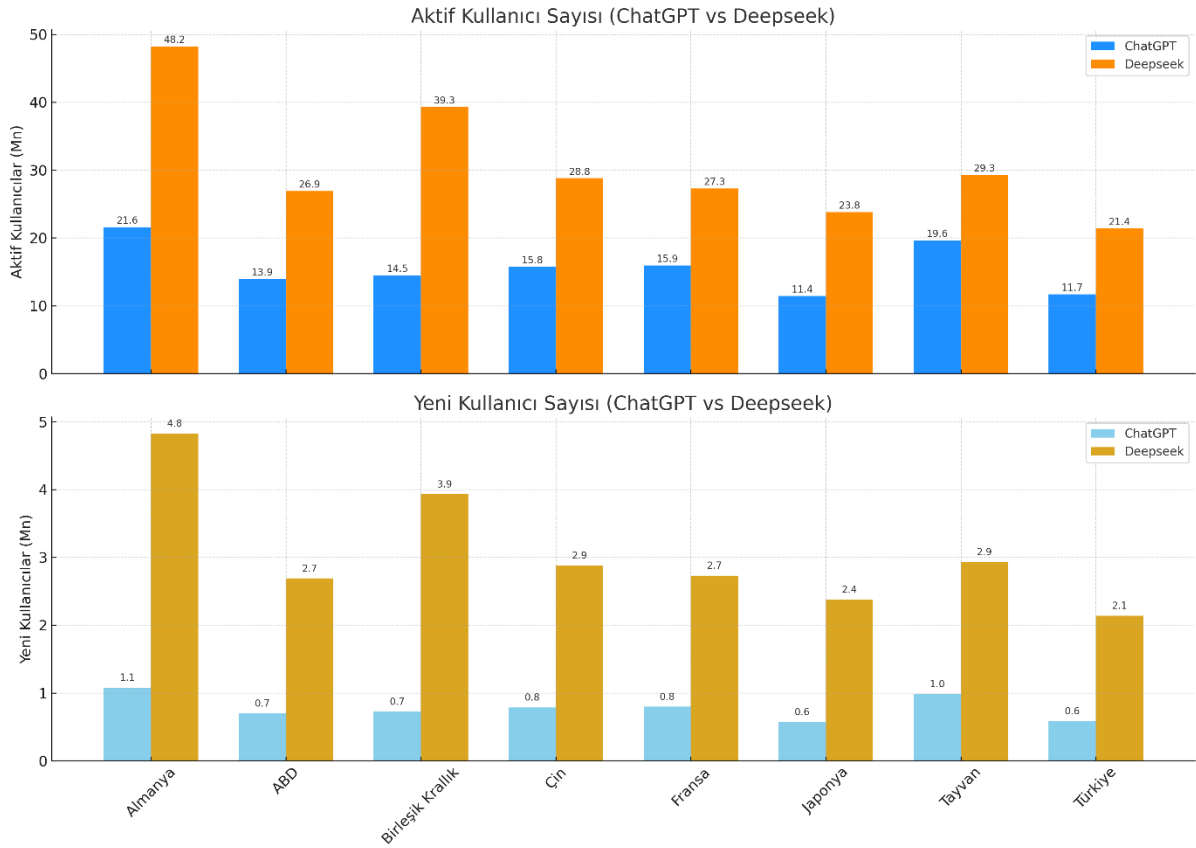
Şekil 12, DeepSeek ve ChatGPT platformlarının farklı dillerdeki kullanıcı dağılımlarını milyon cinsinden göstermektedir.



Şekil 12. Dil Bazında Kullanıcı Dağılımı

DeepSeek platformunda en fazla kullanıcının İngilizce (en) konuştuğu (yaklaşık 1.75 milyon) görülmektedir. Bunu sırasıyla Almanca (de), Çince (zh) ve İspanyolca (es) takip etmektedir. Fransızca (fr) konuşan kullanıcı sayısı ise diğer dillere göre daha düşüktür. ChatGPT platformunda ise İngilizce (en) ve Çince (zh) konuşan kullanıcı sayıları birbirine yakın (yaklaşık 0.75 milyon) ve diğer dillere göre daha yüksektir. Almanca (de), Fransızca (fr) ve İspanyolca (es) konuşan kullanıcı sayıları ise benzer seviyelerde ve İngilizce ile Çinceye göre daha düşüktür. Bu bulgular, DeepSeek'in İngilizce konuşan kullanıcılar arasında daha popüler olduğunu gösterirken, ChatGPT'nin İngilizce ve Çince konuşan kullanıcılar arasında daha dengeli bir dağılıma sahip olduğunu düşündürmektedir. Her iki platformda da diğer Avrupa dillerini konuşan kullanıcı sayısının İngilizce ve Çinceye göre daha az olması dikkat çekicidir.

3.11.Ayrıntılı İncelenen 8 Ülkenin Aktif Kullanıcı Sayısı ve Yeni Kullanıcı Sayısı Karşılaştırılması (Platform Bazında)



Şekil 13. Ayrıntılı İncelenen 8 Ülkenin Aktif Kullanıcı Sayısı ve Yeni Kullanıcı Sayısı Karşılaştırılması (Platform Bazında)

Şekil 13'ün üst panelinde yer alan bar grafik, incelenen sekiz ülkedeki (Almanya, ABD, Birleşik Krallık, Çin, Fransa, Japonya, Tayvan, Türkiye) ChatGPT ve DeepSeek'in milyon cinsinden aktif kullanıcı sayılarını karşılaştırmaktadır. Genel olarak, DeepSeek'in çoğu ülkede ChatGPT'ye kıyasla önemli ölçüde daha yüksek aktif kullanıcı sayısına sahip olduğu görülmektedir. Özellikle Almanya'da DeepSeek'in 48.7 milyon aktif kullanıcısı varken, ChatGPT'nin aktif kullanıcı sayısı 21.6 milyonda kalmıştır. Benzer şekilde, Çin (DeepSeek: 38.8 Mn, ChatGPT: 15.8 Mn), Fransa (DeepSeek: 27.5 Mn, ChatGPT: 15.9 Mn) ve Tayvan (DeepSeek: 29.3 Mn, ChatGPT: 21.6 Mn) gibi ülkelerde DeepSeek belirgin bir üstünlük sağlamıştır. ABD'de ise DeepSeek (26.9 Mn) aktif kullanıcı sayısında ChatGPT'yi (13.9 Mn) geride bırakmasına rağmen, aradaki fark diğer bazı ülkelere göre daha düşüktür. Türkiye'de ise DeepSeek'in (21.4 Mn) ChatGPT'ye (12.7 Mn) kıyasla daha fazla aktif kullanıcıya sahip olduğu tespit edilmiştir. Bu veriler, DeepSeek'in küresel çapta, özellikle belirli pazarlarda, mevcut kullanıcı tabanı açısından ChatGPT'ye göre daha geniş bir erişime sahip olduğunu göstermektedir.

Şekil 13'ün alt panelindeki bar grafik ise aynı ülkelerdeki yeni kullanıcı sayılarını milyon cinsinden kıyaslamaktadır. Bu paneldeki bulgular, aktif kullanıcı sayısındaki eğilime paralel olarak, DeepSeek'in yeni kullanıcı kazanımında da ChatGPT'ye kıyasla dikkate değer bir başarı elde ettiğini ortaya koymaktadır. Almanya'da DeepSeek 4.8 milyon yeni kullanıcı kazanırken, ChatGPT'nin yeni kullanıcı sayısı yalnızca 1.3 milyondur. ABD'de DeepSeek 2.7 milyon yeni kullanıcıya ulaşırken, ChatGPT 0.7 milyonda kalmıştır. Birleşik Krallık (DeepSeek: 3.9 Mn, ChatGPT: 0.7 Mn), Çin (DeepSeek: 2.9 Mn, ChatGPT: 0.9 Mn), Fransa (DeepSeek:

2.7 Mn, ChatGPT: 0.8 Mn), Japonya (DeepSeek: 2.4 Mn, ChatGPT: 0.6 Mn), Tayvan (DeepSeek: 2.9 Mn, ChatGPT: 1.0 Mn) ve Türkiye (DeepSeek: 2.1 Mn, ChatGPT: 0.8 Mn) verileri de DeepSeek'in yeni kullanıcı kazanma oranlarının ChatGPT'ye göre belirgin şekilde yüksek olduğunu teyit etmektedir. Bu durum, DeepSeek'in pazarlama stratejilerinin, ürün özelliklerinin veya yerel pazar uyumluluğunun yeni kullanıcıları çekme konusunda daha etkili olduğunu düşündürmektedir.

Her iki grafiğin birleşik analizi, DeepSeek'in incelenen pazarlarda hem mevcut kullanıcı tabanını genişletme hem de yeni kullanıcıları çekme konusunda ChatGPT'ye karşı genel bir üstünlük sağladığını göstermektedir. Bu bulgu, üretken yapay zekâ pazarındaki rekabetin dinamiklerini ve DeepSeek gibi platformların belirli bölgelerde önemli bir kullanıcı tabanı oluşturma potansiyelini ortaya koymaktadır.

4. Tartışma ve Sonuç

Bu çalışmada, ChatGPT ve DeepSeek üretken yapay zekâ modellerinin kullanıcı etkileşimleri üzerindeki etkileri çok boyutlu olarak analiz edilmiştir. Elde edilen bulgular, her iki platformun da güçlü yönleri ve geliştirilmesi gereken alanları olduğunu ortaya koymakta, literatürdeki benzer çalışmalarla da paralellikler ve farklılıklar göstermektedir.

4.1. Bulguların Yorumlanması ve Alanyazınla İlişkisi

Bulgulara göre, DeepSeek'in özellikle teknik görevlerde (örneğin hata ayıklama, kod optimizasyonu gibi) kullanıcılar tarafından daha yüksek derecelendirmelerle değerlendirilmesi, modelin uzmanlaşmış işlevsellikleriyle uyumludur. Bu durum, DeepSeek'in mimari olarak transformatör tabanlı kodlama görevleri için optimize edilmesinin (Wang vd., 2024) ve özellikle kod üretimi ile mantıksal çıkarım gerektiren görevlerde öne çıkmasının (DeepSeek-AI, 2024) doğal bir sonucu olarak değerlendirilebilir. Nitekim, Guo vd. (2024), DeepSeek'in Python, Java ve C++ gibi dillerde yüksek başarı gösterdiğini ve CodeLlama ya da DeepMind'in AlphaCode modeliyle aynı ligde değerlendirildiğini, hatta daha istikrarlı bağlam takibi ve hata ayıklama hassasiyetiyle öne çıktığını belirtmektedir. Bu uzmanlaşma, kullanıcıların belirli ve karmaşık teknik problemlere çözüm ararken DeepSeek'i daha tatmin edici bulmasını açıklayabilir ve Manik (2025) tarafından da işaret edildiği gibi, dar amaçlı yüksek doğruluk gerektiren görevlerde tercih edilmesine yol açabilir. DeepSeek'in bu teknik üstünlüğü, BytePlus (2025) tarafından da DeepSeek-Coder-V2 modelinin kod üretimi ve analizi alanlarındaki GPT-4 Turbo ile karşılaştırılabilir yüksek doğruluk oranlarıyla desteklenmektedir.

ChatGPT ise özellikle içerik üretimi, profesyonel yazarlık ve eğitim senaryolarında daha geniş bir kullanım alanına sahiptir. Bu, Hariri (2023) ve Ray (2023) tarafından da ifade edilen, ChatGPT'nin çok yönlü doğasına dair literatür bulgularını desteklemektedir. ChatGPT'nin doğal dili anlama konusundaki yetkinliği ve insan benzeri yanıtlar üretebilme kapasitesi (Bahrini vd., 2023), onu eğitim materyali üretme, sunum hazırlama ve yazma becerilerini geliştirme gibi alanlarda (Zheldibayeva, 2024) popüler kılmaktadır. Ancak, bu geniş kullanım alanı, bazen teknik derinlik gerektiren konularda DeepSeek kadar odaklanmış bir performans sunamamasına neden olabilmektedir.

ChatGPT kullanıcılarının daha kısa oturum süresine rağmen daha yüksek memnuniyet bildirdiği bazı oturumlar, yapay zekâ ile olan etkileşimin sadece teknik doğrulukla değil, aynı zamanda iletişim tarzı, dil doğallığı ve bağlamsal uyumla da şekillendiğini ortaya koymaktadır. Bu durum, Nazir ve Wang (2023) tarafından da vurgulandığı üzere, üretken yapay zekâ sistemlerinin “anamlı diyalog” kurma kapasitesinin kullanıcı deneyiminde belirleyici olduğunu göstermektedir. Skjuve vd. (2024) de kullanıcıların ChatGPT gibi araçları kullanma motivasyonlarının çeşitliliğine işaret ederek, algılanan değer ve kullanım kolaylığının önemini vurgular.

4.2.Kullanıcı Etkileşimi ve Cihaz Türleri Üzerindeki Etkiler

Kullanıcıların platformlara erişim için kullandıkları cihaz türleri de deneyim puanları üzerinde anlamlı etkiler yaratmıştır. Masaüstü/dizüstü bilgisayarlarda gerçekleşen etkileşimlerde kullanıcı deneyimi puanlarının daha yüksek olması, bu cihazların hem işlem gücü hem de ekran boyutu açısından daha uygun bir kullanıcı ortamı sunmasından kaynaklanabilir. Bu bulgu, Qawqzeh (2024) ve Sharples (2023) tarafından belirtilen “pedagojik etkileşim ortamının fiziksel yapısının öğrenme verimliliğine etkisi” kavramıyla örtüşmektedir; daha geniş ekranlar ve klavye kullanımı, özellikle karmaşık görevlerde ve içerik üretiminde daha verimli bir etkileşim sağlayabilir. Mobil cihazlar üzerinden yapılan etkileşimlerde ise deneyim puanlarının daha değişkenlik göstermesi, arayüz farklılıkları, ekran boyutu kısıtlamaları ve etkileşim süresi ile açıklanabilir. Bu durum, Adıgüzel vd. (2023) tarafından belirtilen yapay zekanın eğitimde kişiselleştirilmiş deneyimler sunma potansiyelinin, kullanılan cihazın ergonomisi ve işlevselliği ile yakından ilişkili olduğunu düşündürmektedir.

4.3.Regresyon ve Sınıflandırma Modelleri Işığında Derinlemesine Değerlendirme

Çalışmada uygulanan regresyon analizleri, kullanıcı derecelendirmelerinin düşük R^2 değerleri ile modellenemediğini ortaya koymuştur. Bu durum, kullanıcı memnuniyetini etkileyen faktörlerin yalnızca oturum süresi, cihaz türü veya yanıt doğruluğu gibi gözlemlenebilir değişkenlerle sınırlı olmadığını, aksine bireysel kullanıcı özellikleri, duygusal bağlam, önceki etkileşim geçmişi gibi daha derin psikometrik faktörlerin de önemli olduğunu düşündürmektedir. Al-Abdullatif ve Alsubaie (2024), kullanıcı deneyiminin çok boyutlu olduğunu ve kişisel beklenti düzeylerinin, yapay zekâ okuryazarlığının ve algılanan değerın memnuniyet üzerinde belirleyici rol oynadığını vurgulamaktadır. Bu bulgu, yapay zekâ etkileşimlerinde insan faktörünün karmaşıklığını ve standart metriklerle ölçülmesinin zorluğunu bir kez daha teyit etmektedir.

Sınıflandırma modelleri özelinde, “Excellent” olarak etiketlenen kullanıcı geri bildirimlerinin görece başarıyla tahmin edilebilmesine karşın, “Fair” ve “Good” sınıflarının ayırt edilememesi, kullanıcı değerlendirme ölçeklerinin sübjektifliğini ve potansiyel dengesiz veri dağılımını ortaya koymaktadır. Kullanıcıların “orta” ve “iyi” gibi derecelendirmeleri farklı yorumlayabileceği ve bu durumun model performansını etkileyebileceği düşünülmektedir. Bu bağlamda, kullanıcıların değerlendirme yaparken hangi kriterleri temel aldıkları ve bu kriterlerin platformdan platforma nasıl değiştiği daha derinlemesine incelenmelidir.

4.4.Kullanıcı Sadakati, Terk Etme Oranı ve Platform Performansı

Çalışmanın dikkat çeken sonuçlarından biri, DeepSeek'in genel olarak daha düşük terk etme oranlarına sahip olmasıdır. Bu durum, kullanıcıların DeepSeek ile kurdukları etkileşimin sürdürülebilirliğini ve platformun belirli kullanıcı ihtiyaçlarını karşılama başarısını göstermektedir. Mogavi vd. (2023), yapay zekâ destekli dijital araçların

sürekli kullanımının, sadece işlevsel başarı değil, aynı zamanda psikolojik tatmin ile de ilişkili olduğunu belirtmektedir. Bu bağlamda, DeepSeek'in daha teknik, hedef odaklı ve kodlama görevlerine dayalı yapısı, profesyonel kullanıcılar için daha cazip ve tatmin edici bir seçenek olarak değerlendirilebilir ve bu da daha düşük terk oranlarını açıklayabilir. Ayrıca, DeepSeek'in aktif kullanıcı sayısının ve ortalama oturum süresinin ChatGPT'ye kıyasla daha yüksek olması, bu platformun kullanıcıları daha uzun süre etkileşimde tutabildiğini veya daha karmaşık ve zaman alıcı görevler için tercih edildiğini göstermektedir.

Öte yandan ChatGPT'nin özellikle yıl sonlarına doğru kullanıcı bağlılığında artış göstermesi, eğitim-öğretim dönemlerinin etkisiyle ve genel amaçlı kullanımının yaygınlığıyla açıklanabilir. Hariri (2023), ChatGPT'nin akademik takvimle uyumlu kullanım örüntülerine sahip olduğunu ifade etmiş ve bu modeli öğrenci destek aracı olarak tanımlamıştır. Fırat (2023) da üniversitelerde yapay zekâ sohbet robotlarının öğrenci katılımını artırma potansiyeline dikkat çekmektedir. Bu durum, ChatGPT'nin geniş bir kullanıcı kitlesine hitap etme ve çeşitli ihtiyaçlara yanıt verme esnekliğinden kaynaklanıyor olabilir.

4.5. Etik ve Veri Koruma Açısından Değerlendirme

Her iki modelin de kullanıcı verilerini işlerken önemli etik sorumluluklar taşıdığı görülmektedir. ChatGPT için belirtilen en önemli sorunlardan biri, "hallucination" adı verilen yanlış veya uydurma bilgi üretme eğilimidir; bu durum Alkaissi & McFarlane (2023) tarafından da vurgulanmıştır. Bu sorun özellikle akademik ya da tıbbi içeriklerde ciddi sonuçlar doğurabilir (Bahrini vd., 2023; Spennemann, 2025). Ayrıca, eğitim verilerindeki önyargıları yansıtmaya riski (Rozado, 2023) ve fikri mülkiyet hakları (Zirpoli, 2023; Al-Busaidi vd., 2024) ChatGPT ile ilişkili önemli etik kaygılardır. DeepSeek'in eğitim verilerinin tescilli olması (Sapkota vd., 2025), modelin potansiyel önyargılarını (Witness AI, 2025) ve eğitim seti kalitesini değerlendirmeyi güçleştirmektedir. "Açık ağırlık" politikasının (DeepSeek-R1, t.y.) tam anlamıyla "açık kaynak" anlamına gelmediği (CTOL Digital, 2025) ve bu durumun yeniden üretilebilirlik ve şeffaflık açısından sınırlamalar getirdiği belirtilmelidir. DeepSeek'in de, diğer dil modelleri gibi, zaman zaman yanlış bilgi üretme (hallucination) riski (Vectara, 2025) ve kod üretme yeteneği nedeniyle kötü amaçlı yazılım geliştirme gibi kötüye kullanım potansiyeli (SecurityWeek, 2025) bulunmaktadır. Kullanıcı verilerinin güvenliğine ilişkin olarak, DeepSeek ve ChatGPT platformlarının şeffaflık düzeylerinin farklılık gösterdiği görülmektedir. Özellikle Avrupa Birliği Genel Veri Koruma Tüzüğü (GDPR) ve Türkiye'deki Kişisel Verilerin Korunması Kanunu (KVKK) gibi düzenlemelere uyum düzeylerinin platformlar bazında açıklığa kavuşturulması gerekmektedir, nitekim İtalya'nın OpenAI'ye gizlilik kuralları ihlali nedeniyle ceza vermesi (Reuters, 2024) bu konunun hassasiyetini göstermektedir. Kullanıcıların hassas bilgilerini bu platformlara yüklemesi, özellikle sağlık verileri gibi özel içeriklerde ciddi gizlilik ihlali potansiyeli taşımaktadır (Borgesius, 2018). Google'ın Gemini modeli için yayımladığı gizlilik uyarısı (Search Engine Journal, 2024) ve ICO (2023) tarafından belirtilen yapay zekâ ve veri koruma rehberliği, bu alandaki genel endişeleri yansıtmaktadır. Kullanıcı sohbet geçmişinin nasıl saklandığı, hangi amaçlarla analiz edildiği ve üçüncü taraflarla paylaşılıp paylaşılmadığı gibi konular, etik sorumluluklar açısından kritik önemdedir.

4.6. Araştırmanın Sınırlılıkları

Bu çalışmanın bazı metodolojik sınırlılıkları bulunmaktadır. İlk olarak, kullanılan veri seti sentetik olarak oluşturulmuş olup, kullanıcı davranışlarını tamamen gerçekçi biçimde yansıtmayabilir. Bu nedenle, çalışmanın dış

geçerliliği (external validity) sınırlı olabilir. İkinci olarak, zaman içinde yapay zekâ modellerinde gerçekleşen güncellemeler (örneğin, GPT-4'ten GPT-4o'ya geçiş gibi), model performansını etkilemekte ve elde edilen sonuçların geçerlilik süresini sınırlamaktadır. Çalışmada kullanılan modellerin versiyonları (GPT-4-turbo ve DeepSeek-Chat 1.5) en son çıkan modeller olmayabilir ve bu da güncel karşılaştırmalar açısından bir kısıtlılık oluşturabilir. Ayrıca, çalışmada kullanıcıların demografik özelliklerine, kültürel geçmişlerine veya yapay zekâ okuryazarlık düzeylerine ilişkin detaylı veri bulunmaması, analizlerin daha özelleştirilmiş ve derinlemesine çıkarımlar sunmasını zorlaştırmıştır. Bu tür faktörlerin kullanıcı deneyimi ve memnuniyeti üzerinde önemli etkileri olabileceği (Al-Abdullatif ve Alsubaie, 2024) bilinmektedir. Son olarak, "hallucination" veya önyargı gibi etik sorunların tespiti ve analizi için kullanılan metriklerin bu çalışmada sınırlı olması, bu konulardaki değerlendirmelerin derinliğini kısıtlamıştır.

4.7.Sonuçlar ve Öneriler

Bu çalışmadan elde edilen bulgular, üretken yapay zekâ platformları olan ChatGPT ve DeepSeek'in performanslarının ve kullanıcı etkileşimlerinin; kullanım bağlamı, görevin türü ve erişim sağlanan cihazın özelliklerine bağlı olarak anlamlı farklılıklar gösterdiğini kapsamlı bir şekilde ortaya koymaktadır. Her iki platform da kendi içinde belirli alanlarda üstünlükler sergilemekle birlikte, etik ve veri gizliliği gibi konular tüm yapay zekâ sistemlerinin geliştirilmesinde temel öncelikler arasında yer almaktadır.

Genel bulgular doğrultusunda, ChatGPT'nin içerik üretimi, akademik yazım desteği, eğitim uygulamaları ve sosyal etkileşim gibi alanlarda geniş kullanım alanına sahip olduğu ve yüksek kullanıcı memnuniyeti sağladığı görülmüştür. Platformun güçlü doğal dil işleme yetenekleri ve esnek yapısı, farklı kullanıcı profilleri için cazip bir seçenek haline gelmesine katkı sağlamaktadır. Öte yandan, DeepSeek'in özellikle teknik ve uzmanlık gerektiren görevlerde —örneğin kodlama, hata ayıklama, algoritma açıklamaları— optimize edilmiş mimarisi sayesinde daha yüksek doğruluk oranları sunduğu ve bu bağlamda kullanıcı sadakatini artırdığı tespit edilmiştir. DeepSeek, aktif kullanıcı sayısı ve ortalama oturum süresi gibi metriklerde de dikkat çekici üstünlükler göstermiştir.

Kullanıcı deneyimini etkileyen önemli değişkenlerden biri olan cihaz türü de çalışmada öne çıkan bir başka bulgudur. Masaüstü ve dizüstü bilgisayar kullanıcılarının, daha gelişmiş arayüz etkileşimi ve işlem kapasitesi nedeniyle daha pozitif deneyimler yaşadığı rapor edilmiştir. Buna karşılık, mobil cihazlarda kullanıcı deneyimi puanlarının daha yüksek oranda değişkenlik gösterdiği gözlemlenmiştir.

Çalışmada yürütülen regresyon analizleri, kullanıcı memnuniyetini etkileyen çok sayıda değişken bulunduğunu ve mevcut bağımsız değişkenlerin bu varyansın yalnızca belirli bir kısmını açıklayabildiğini ortaya koymuştur. Bu durum, kullanıcı memnuniyetinin öznel doğasını ve kişisel beklentiler gibi ölçülmesi zor faktörlerin etkisini vurgulamaktadır.

Etik hususlar ve veri gizliliği hem ChatGPT hem de DeepSeek için kritik endişe alanları olarak öne çıkmaktadır. Model şeffaflığı, önyargıların azaltılması, "hallucination" riski ve zararlı içerik üretimi gibi konular, yapay zekâ teknolojilerinin güvenilirliğini ve toplumsal kabulünü doğrudan etkilemektedir. Bu bağlamda platformların etik kurallara uyumu ve yasal düzenlemelere (örneğin GDPR, KVKK) uygunluğu, kullanıcı güveni açısından büyük önem taşımaktadır.

Bu bulgular doğrultusunda geliştirilen öneriler ise üç temel paydaş grubu etrafında şekillenmiştir. İlk olarak, platform geliştiricileri için; eğitim veri setlerinin içeriği, önyargılar ve bu önyargıların azaltılmasına yönelik stratejiler konusunda daha şeffaf olunması gerekmektedir. Kullanıcı verilerinin işlenmesi, saklanması ve korunmasına ilişkin politikalar açık biçimde sunulmalı ve veri koruma yasalarına tam uyum sağlanmalıdır. Ayrıca, farklı cihazlar için optimize edilmiş arayüzler geliştirilerek kullanıcı geri bildirimleri sistematik şekilde değerlendirilmelidir. Sürekli Ar-Ge çalışmalarıyla etik risklerin (örneğin hallucination, kötüye kullanım) önlenmesine yönelik mekanizmalar güçlendirilmelidir. DeepSeek örneğinde görüldüğü üzere, genel amaçlı ve uzmanlaşmış modeller arasında bir denge kurulması da önemli bir strateji olacaktır.

İkinci olarak, eğitim kurumları ve profesyonel kullanıcılar için; yapay zekâ okuryazarlığını artırmak amacıyla kullanıcıların farklı araçların kapasite ve sınırlılıklarını anlayabilecekleri rehberler hazırlanmalı, eleştirel düşünme ve doğrulama becerileri teşvik edilmelidir. Kullanım bağlamına uygun araç seçiminde farkındalık oluşturulmalı; örneğin yaratıcı içerik üretimi için ChatGPT, teknik görevler için ise DeepSeek tercih edilmelidir. Ayrıca, akademik ve profesyonel alanlarda yapay zekânın etik kullanımına ilişkin ilkeler açık biçimde tanımlanmalı ve bu ilkelere uyum teşvik edilmelidir.

Son olarak, araştırmacılar için; demografik faktörlerin etkisini göz önünde bulunduran daha kapsamlı ve karşılaştırmalı çalışmalar yapılması önerilmektedir. Kullanıcıların üretken yapay zekâ ile olan uzun vadeli etkileşimlerinin; öğrenme, problem çözme ve yaratıcılık üzerindeki etkileri derinlemesine araştırılmalıdır. Önyargıların tespiti, hallucination probleminin çözümü ve yapay zekâ güvenliği konularında disiplinlerarası çalışmalar teşvik edilmeli; sentetik veri setlerinin ötesine geçilerek gerçek dünya senaryolarına dayalı risk analizleri yapılmalıdır.

Teşekkür ve Bilgilendirme / Acknowledgements

Makale, kısmen veya tamamen yayınlanmaması şartıyla, bir proje veya tez olarak devam eden bir bildiri olarak sunuluyorsa, bu bölümde belirtilmelidir. Makale bir araştırma kurumu veya bir fon tarafından destekleniyorsa, vakfın adı, proje numarası ve tamamlanma tarihi burada belirtilmelidir. İstenirse, makale bağlamında bir kişi veya vakıf için takdirler burada belirtilmelidir. / If the article is submitted as a proceeding, on the condition of not being partially or fully published, a project or dissertation, it should be stated in this section. If the article is supported by a research institution or a fund, the name of the foundation, project number and completion date should be stated here. If desired, appreciations to a person or a foundation within the context of the article should be stated here.

Yayın Etiği Bildirimi / Research Ethics

Yazar araştırmanın etik dışı bir sorunu olmadığını, araştırma ve yayın etiği konusunu gözlemlediğini beyan etmektedir. / The author declares that the research has no unethical problem and observes the research and publication ethics.

Araştırmacıların Katkı Oranı / Contribution Rate of Researchers

Çalışmanın her aşamasına tüm yazarlar eşit derecede katkı sunmuştur. / All authors contributed equally to every stage of the study.

Çıkar Çatışması / Conflict of Interest

Çalışmada herhangi bir çıkar çatışması bulunmamaktadır. / The study has no conflict of interest.

Fon Bilgileri / Funding

Bu çalışmada herhangi bir fon kullanılmamıştır. / There is no funding for this study.

Etik Kurul Onayı / The Ethical Committee Approval

Etik kurul kararı: Bu araştırmada, tüm araştırmacılara açık, uluslararası veri tabanında yer alan veriler kullanıldığından etik kurul kararı gerektirmemektedir. / The Ethical Committee Approval: This research does not require an ethics committee decision, since data in an international database open to all researchers are used.

Kaynakça / References

- Acemoglu, D., Autor, D., & Johnson, S. (2023). Can we have pro-worker ai. <https://shapingwork.mit.edu/wp-content/uploads/2023/09/Pro-Worker-AI-Policy-Memo.pdf>
- Adiguzel, T., Kaya, M. H., & Cansu, F. K. (2023). Revolutionizing education with AI: Exploring the transformative potential of ChatGPT. *Contemporary Educational Technology*, 15(3).
- Al-Abdullatif, A. M., & Alsubaie, M. A. (2024). ChatGPT in learning: Assessing students' use intentions through the lens of perceived value and the influence of AI literacy. *Behavioral Sciences*, 14(9), 845.
- Al-Busaidi, A. S., Raman, R., Hughes, L., Albashrawi, M. A., Malik, T., Dwivedi, Y. K., ... & Walton, P. (2024). Redefining boundaries in innovation and knowledge domains: Investigating the impact of generative artificial intelligence on copyright and intellectual property rights. *Journal of Innovation & Knowledge*, 9(4), 100630.
- Alkaissi, H., & McFarlane, S. I. (2023). Artificial hallucinations in ChatGPT: implications in scientific writing. *Cureus*, 15(2).
- Artetxe, M., Bhosale, S., Goyal, N., Mihaylov, T., Ott, M., Shleifer, S., & Stoyanov, V. (2021). Efficient large scale language modeling with mixtures of experts. <https://arxiv.org/pdf/2112.10684>
- Bahrini, A., Khamoshifar, M., Abbasimehr, H., Riggs, R. J., Esmaeili, M., Majdabadkohne, R. M., & Pasehvar, M. (2023). ChatGPT: Applications, opportunities, and threats. In *2023 Systems and Information Engineering Design Symposium (SIEDS)*, 274-279.
- Bell, R., & Bell, H. (2023). Entrepreneurship education in the era of generative artificial intelligence. *Entrepreneurship Education*, 6(3), 229–244. <https://doi.org/10.1007/s41959-023-00099-x>
- Borgesius, F. (2018). Discrimination, artificial intelligence, and algorithmic decision-making. *Council of Europe, Directorate General of Democracy*.
- BytePlus. (2025). DeepSeek-Coder-V2: Revolutionizing AI-powered code generation. <https://www.byteplus.com/en/topic/375561?title=deepseek-coder-v2-revolutionizing-ai-powered-code-generation>
- Chein, J., Martinez, S., & Barone, A. (2024). Can human intelligence safeguard against artificial intelligence? Exploring individual differences in the discernment of human from AI texts. *Research Square*, rs-3.
- CTOL Digital. (2025). Is DeepSeek Truly Open Source or Just Following Industry Norms? <https://www.ctol.digital/news/is-deepseek-truly-open-source-or-following-industry-norms/>
- Dataloop. (t.y.). DeepSeek Coder V2 Instruct · Models. https://dataloop.ai/library/model/deepseek-ai_deepseek-coder-v2-instruct/
- DeepSeek R1: DeepSeek-V3. (t.y.). <https://deepseeksr1.com/v3/>
- DeepSeek-AI. (2024). DeepSeek-V3 Technical Report. <https://arxiv.org/pdf/2412.19437>
- DeepSeek-Coder. (t.y.). DeepSeek Coder. <https://github.com/deepseek-ai/DeepSeek-Coder>
- DeepSeek-R1. (t.y.). DeepSeek-R1 Training Overview. <https://github.com/deepseek-ai/DeepSeek-R1>

- DeepSeek-V3. (t.y.). DeepSeek-V3. <https://github.com/deepseek-ai/DeepSeek-V3>
- Du, N., Huang, Y., Dai, A. M., Tong, S., Lepikhin, D., Xu, Y., ... & Cui, C. (2022, June). Glam: Efficient scaling of language models with mixture-of-experts. <https://arxiv.org/pdf/2112.06905>
- Dwivedi, Y. K., Kshetri, N., Hughes, L., Slade, E. L., Jeyaraj, A., Kar, A. K., & Wright, R. (2023). Opinion Paper:“So what if ChatGPT wrote it?” Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research, practice and policy. *International Journal of Information Management*.
- Eren, F., Dülek, L. N., Uraz, Ö. A., Kuşcu, B., & Sakallı, M. (2024). Yapay zeka endeksi kapsamında ülke bazlı yapay zeka politika stratejilerinin etkinliği: Performans ve verimlilik değerlendirmesi incelenmesi. *Bilgi ve İletişim Teknolojileri Dergisi*, 6(2), 113-148.
- Firat, M. (2023). What ChatGPT means for universities: Perceptions of scholars and students. *Journal of Applied Learning and Teaching*, 6(1), 57-63.
- Guo, D., Zhu, Q., Yang, D., Xie, Z., Dong, K., Zhang, W., ... & Liang, W. (2024). DeepSeek-Coder: When the large language model meets programming -- The rise of code intelligence. <https://arxiv.org/pdf/2401.14196>
- Gupta, M. (2025). RLHF (OpenAI) vs Simple RL (DeepSeek). Medium. <https://medium.com/data-science-in-your-pocket/rlhf-openai-vs-simple-rl-deepseek-93054cb509a4>
- Hariri, W. (2023). Unlocking the potential of ChatGPT: A comprehensive exploration of its applications, advantages, limitations, and future directions in natural language processing. <https://arxiv.org/abs/2304.02017>
- IBM. (2025a). DeepSeek: sorting through the hype. <https://www.ibm.com/think/topics/deepseek>
- IBM. (2025b). DeepSeek's reasoning AI shows power of small models, efficiently. <https://www.ibm.com/think/news/deepseek-rl-ai>
- Information Commissioner's Office (ICO). (2023). Guidance on AI and data protection. ICO.
- Liu, H., Zhou, Y., Li, M., Yuan, C., & Tan, C. (2024). Literature meets data: A synergistic approach to hypothesis generation. <https://arxiv.org/pdf/2410.17309>
- Manik, M. M. H. (2025). ChatGPT vs. DeepSeek: A comparative study on AI-Based code generation. <https://arxiv.org/pdf/2502.18467>
- Markov, T., Zhang, C., Agarwal, S., Eloundou, T., Lee, T., Adler, S., Jiang, A., & Weng, L. (2023). New and improved content moderation tooling. <https://openai.com/blog/new-and-improved-content-moderation-tooling/>
- Mesko, B. (2017). The role of artificial intelligence in precision medicine. *Expert Review of Precision Medicine and Drug Development*, 2(5), 239-241.
- Mogavi, R. H., Deng, C., Kim, J. J., Zhou, P., Kwon, Y. D., Metwally, A. H. S., ... & Hui, P. (2024). ChatGPT in education: A blessing or a curse? A qualitative study exploring early adopters' utilization and perceptions. *Computers in Human Behavior: Artificial Humans*, 2(1).

- Nazir, A., & Wang, Z. (2023). A comprehensive survey of ChatGPT: Advancements, applications, prospects, and challenges. *Meta-radiology*.
- Neptune.ai (2025). Reinforcement learning from human feedback for LLMs. <https://neptune.ai/blog/reinforcement-learning-from-human-feedback-for-llms>
- OpenAI. (2023). GPT-4 Technical Report. <https://arxiv.org/pdf/2303.08774>
- Ouyang, L., Wu, J., Jiang, X., Almeida, D., Wainwright, C., Mishkin, P., ... & Lowe, R. (2022). Training language models to follow instructions with human feedback. <https://arxiv.org/pdf/2203.02155>
- Qawqzeh, Y. (2024). Exploring the influence of student interaction with ChatGPT on critical thinking, problem solving, and creativity. *International Journal of Information and Education Technology*, 14(4), 596-601.
- Ray, P. P. (2023). ChatGPT: A comprehensive review on background, applications, key challenges, bias, ethics, limitations and future scope. *Internet of Things and Cyber-Physical Systems*, 3, 121-154.
- Reuters. (2024). Italy fines OpenAI 15 million euros over privacy rules breach. Reuters. <https://www.reuters.com/technology/italy-fines-openai-15-million-euros-over-privacy-rules-breach-2024-12-20/>
- Rivas, P., & Zhao, L. (2023). Marketing with chatgpt: Navigating the ethical terrain of gpt-based chatbot technology. *AI*, 4(2), 375-384.
- Rozado, D. (2023). The political biases of ChatGPT. *Social Sciences*, 12(3), 148.
- Sapkota, R., Raza, S., & Karkee, M. (2025). Comprehensive analysis of transparency and accessibility of ChatGPT, DeepSeek and other Sota large language models. <https://arxiv.org/pdf/2502.18505>
- Search Engine Journal. (2024). Google gemini warning: Don't share confidential information. <https://www.searchenginejournal.com/google-gemini-privacy-warning/507818/>
- Security Magazine. (2025). Dangers of DeepSeek's privacy policy: Data risks in the age of AI. <https://www.securitymagazine.com/articles/101374-dangers-of-deepseeks-privacy-policy-data-risks-in-the-age-of-ai>
- SecurityWeek. (2025). DeepSeek's malware-generation capabilities put to test. <https://www.securityweek.com/deepseeks-malware-generation-capabilities-put-to-test/>
- Sharples, M. (2023). Towards social generative AI for education: theory, practices and ethics. *Learning: Research and Practice*, 9(2), 159-167.
- Shirani, M. (2025). Comparing the performance of ChatGPT 4o, DeepSeek R1, and Gemini 2 Pro in answering fixed prosthodontics questions over time. *The Journal of Prosthetic Dentistry*.
- Skjuve, M., Brandtzaeg, P. B., & Følstad, A. (2024). Why do people use ChatGPT? Exploring user motivations for generative conversational AI. *First Monday*.
- Spennemann, D. H. (2025). The origins and veracity of references 'Cited' by generative artificial intelligence applications: Implications for the quality of responses. *Publications*, 13(1), 12.
- Stokel-Walker, C. (2023). ChatGPT listed as author on research papers: many scientists disapprove.

- Şentürk, Ö. (2023). İç denetim faaliyetlerinde yapay zekadan beklentiler: chatgpt uygulaması örneği. *TIDE AcademIA Research*, 4(2), 51-82.
- Tecuci, G. (2012). Artificial intelligence. *Wiley Interdisciplinary Reviews: Computational Statistics*, 4(2), 168–180. <https://doi.org/10.1002/wics.200>
- Vectara, (2025). DeepSeek-R1 hallucinates more than DeepSeek-V3. <https://www.vectara.com/blog/deepseek-r1-hallucinates-more-than-deepseek-v3>
- Wang, K., Ruan, Q., Zhang, X., Fu, C., & Duan, B. (2024). Pre-service teachers' GenAI anxiety, technology self-efficacy, and TPACK: Their structural relations with behavioral intention to design GenAI-assisted teaching. *Behavioral sciences*, 14(5), 373.
- Weidinger, L., Mellor, J., Rauh, M., Griffin, C., Uesato, J., Huang, P. S., ... & Gabriel, I. (2021). Ethical and social risks of harm from language models. <https://arxiv.org/abs/2112.04359>
- Witness AI. (2025). Open R1 vs DeepSeek: security implications of open vs. closed data. <https://witness.ai/blog-open-r1-vs-deepseek-security-implications-of-open-vs-closed-data/>
- Zhao, W. X., Zhou, K., Li, J., Tang, T., Wang, X., Hou, Y., ... & Wen, J. R. (2023). A survey of large language models. <https://arxiv.org/abs/2303.18223>
- Zheldibayeva, R. (2024). The impact of AI and peer feedback on research writing skills: a study using the CGScholar platform among Kazakhstani scholars. <https://arxiv.org/pdf/2503.05820>
- Zirpoli, C. T. (2023). Generative artificial intelligence and copyright law.

Oyunlaştırılmış E-Öğrenme Destekli Öğretimin Öğrencilerin Bilgisayar Programlama Öz Yeterlik Algıları, Programlamaya Yönelik Tutumları ve Derse Katılımlarına Etkisi

Özgür ATAMAN*  ¹ Fatma Gizem KARAOĞLAN YILMAZ  ²

Anahtar Sözcükler

Oyunlaştırma
Eba
Programlama
Öz-yeterlik
Tutum

Makale Hakkında

Gönderim Tarihi

17 Aralık 2024

Kabul Tarihi

09 Ocak 2025

Yayın Tarihi

30 Haziran 2025

Makale Türü

Araştırma Makalesi

Öz

Bu araştırmanın amacı, oyunlaştırılmış e-öğrenme destekli öğretimin öğrencilerin bilgisayar programlamaya yönelik öz yeterlik algıları, tutumları ve derse katılım üzerindeki etkilerini incelemektir. Karma araştırma yöntemi kullanılan çalışmada, “Bilgisayar Programlama Öz Yeterlik Ölçeği,” “Programlamaya Yönelik Tutum Ölçeği,” “Derse Katılım Envanteri” ve yüz yüze görüşme formları veri toplama aracı olarak kullanılmıştır. Araştırmaya 6 hafta süresince 60 lise öğrencisi katılmıştır. Deney grubundaki öğrencilere, Milli Eğitim Bakanlığı’nın Eğitim Bilişim Ağı (EBA) platformu üzerinden oyunlaştırılmış bir e-öğrenme ortamı sunulmuştur. Bu süreçte öğrencilere düzenli olarak ders materyalleri, videolar ve testler sağlanmış; oyunlaştırma kapsamında katılımları puanlanarak rozetler verilmiş ve haftalık liderlik tahtası oluşturulmuştur. Araştırma sonucunda, öğrencilerin derse katılım ve programlamaya yönelik tutumlarında istatistiksel olarak anlamlı bir fark gözlenmemiştir. Ancak, programlama öz yeterlik algılarında anlamlı bir artış olduğu tespit edilmiştir. Çalışma, programlama öğreniminde karşılaşılan zorluklara yönelik olarak oyunlaştırılmış e-öğrenme desteğinin etkisini deneysel olarak göstermektedir. Ayrıca, bu araştırmanın EBA sistemi içerisinde farklı derslerin oyunlaştırılmasına yönelik bir model sunabileceği düşünülmektedir.

The Effect of Gamified E-Learning Supported Teaching on Students’ Computer Programming Self-Efficacy Perceptions, Attitud Towards Progammig and Classroom Engagement

Keywords

Gamification
Eba
Programming
Self-efficacy
Attitude

Article Info

Received

December 17, 2024

Accepted

January 09, 2025

Published

June 30, 2025

Article Type

Research Paper

Abstract

The aim of this study is to examine the effects of gamified e-learning-supported instruction on students' self-efficacy perceptions, attitudes toward computer programming, and classroom engagement. In the study using mixed research method, “Computer Programming Self-Efficacy Scale,” “Attitude Scale Towards Programming,” “Classroom Engagement Inventory” and face-to-face interview forms were used as data collection tools. 60 high school students participated in the study for 6 weeks. A gamified e-learning environment was provided to the students in the experimental group via the Education Informatics Network (EBA) platform of the Ministry of National Education. During this process, students were regularly provided with course materials, videos and tests; their participation was scored within the scope of gamification and badges were given and a weekly leader board was created. As a result of the study, no statistically significant difference was observed in students' classroom engagement and attitudes towards programming. However, a significant increase was found in their programming self-efficacy perceptions. The study experimentally demonstrates the effect of gamified e-learning support on the difficulties encountered in programming learning. Additionally, it is thought that this research may provide a model for gamification of different courses within the EBA system.

Atf: Ataman, Ö., Karaoğlan Yılmaz F. G., (2025). Oyunlaştırılmış e-öğrenme destekli öğretimin öğrencilerin bilgisayar programlama öz yeterlik algıları, programlamaya yönelik tutumları ve derse katılımlarına etkisi. *Bilgi ve İletişim Teknolojileri Dergisi*, 7(1), 40-62. <https://doi.org/10.53694/bited.1603352>

Cite: Ataman, Ö., Karaoğlan Yılmaz F. G., (2025). The effect of gamified e-learning supported teaching on students’ computer programming self-efficacy perceptions, attitud towards progammig and classroom engagement. *Journal of Information and Communication Technologies*, 7(1), 40-62. <https://doi.org/10.53694/bited.1603352>

*Sorumlu Yazar/Corresponding Author zgrataman@gmail.com.

¹ Prof. Dr., Bartın University, Science Faculty, Bartın/Turkey, gkaraoglanyilmaz@gmail.com,  <https://orcid.org/0000-0003-4963-8083>

² M.Sc. Student, Bartın University, Science Faculty, Bartın/Turkey, zgrataman@gmail.com,  <https://orcid.org/0000-0002-3847-8427>

Extended Abstract

Introduction

In learning programming, due to the abstractness and complexity of programming, it becomes a difficult obstacle for beginner students and causes students to lose interest in programming and not to put in effort (Ersoy et al., 2011). One way to reduce these problems may be to evaluate e-learning environments. The advantages of e-learning are equal opportunity, elimination of time limitations in accessing information and reducing costs (Özgür, 2011). These advantages it adds support both distance education and blended education (Özgür, 2011).

One way to solve the problems in programming education may be gamification. It is possible to come across many different definitions of gamification. Gamification, which is generally defined as the inclusion of game structuring elements and mechanics in environments where games do not exist, can be adapted to all subjects in theory (Deterding et al., 2011; Muntean, 2011). When looking at the literature, although there are serious differences between gamification and game-based learning (Codish & Ravid, 2014), it is noteworthy that this concept is confused at some points. Kim et al. (2009) defined game-based learning as participants achieving educational goals within the game. Although gamification is not the context itself, it is a driving force that enables difficult processes to continue more motivated and determinedly.

In this study, the relationship between classroom engagement, programming self-efficacy and programming attitudes of the student group that was given gamified e-learning support and the student group that was not given gamified e-learning support was examined. The study was carried out on the EBA platform and the use of gamification of this e-learning environment on a course basis in terms of gamification, which is a current approach, was also discussed. It is thought that the study will contribute to the literature with the effect of gamified e-learning support on problems in programming learning.

Method

This study employed a mixed-methods research approach, integrating both quantitative and qualitative phases. The quantitative aspect utilized a "Pretest-Posttest Control Group Quasi-Experimental Design," which is one of the methods under quantitative research. While the experimental and control groups for the quasi-experimental study were assigned objectively, the participants within these groups were not randomly selected. Since the research was conducted during the formal education period, reassigning students to different groups could have disrupted the ongoing educational process. Additionally, grouping students through EBA (Education Informatics Network) could have caused confusion in the gamification process. Consequently, two randomly assigned classes were designated as the control and experimental groups.

Measurements were taken on the dependent variables of both groups before and after the study. After the experimental study, measurements were made on the dependent variables of the groups again and analysis was performed to determine whether there was a statistically significant difference. The study lasted six weeks and involved a total of 60 students from two 9th-grade classes of a Science High School in Zonguldak during the 2023-2024 academic year. The study group was selected based on several criteria, including the availability of a computer laboratory at the school, the ease of access for the researcher, the teacher's qualifications (a master's degree and ongoing academic studies in information technologies), and the voluntary participation of students.

The study was carried out during the second semester of the 2023-2024 academic year within the scope of the "Information Technologies and Software" course for 9th-grade students. The study group consisted of one class of 30 students serving as the experimental group and another class of 30 students as the control group. Coding instruction was delivered using Python, a text-based programming language.

Quantitative data were collected using the "Programming Self-Efficacy Scale for Secondary School Students" developed by Özkara and Yelken (2020), the "Programming Attitude Scale" designed by Altay and Kışla (2018), the "Classroom Participation Inventory" originally developed by Wang et al. (2014) and translated into Turkish by Sever (2014), and qualitative data were gathered using a face-to-face interview form developed by the researcher.

The gamification components used in the study included points, leaderboards, levels, and badges. Learners earned badges based on the learning points they accumulated through the activities provided. The types of badges in the gamification setup were as follows: activity badges, level badges, weekly ranking badges, and overall ranking badges. Within the framework of the gamification design, students' participation in activities was scored, and their progress was monitored weekly using a tabulation program. Weekly, students were awarded entrance badges, level badges, and, where applicable, ranking badges based on their scores. At the end of the study, overall badge tables were published.

Quantitative data from the study were analyzed using the SPSS statistical software. Analysis of Covariance (ANCOVA) was applied to determine whether the effects of gamified e-learning on classroom engagement, programming self-efficacy, and attitudes toward programming were statistically significant.

Qualitative data were analyzed using content analysis. The findings were expressed in terms of percentages and frequencies and were interpreted accordingly. During the analysis process, the researcher and a subject matter expert independently derived themes from the interview texts covering the interview questions. They then convened to reach a consensus and finalized the themes for the interviews. Content analysis was conducted based on the themes derived from the responses to the student opinion determination form.

Findings

Quantitative Research Findings

In this study, while a significant difference was found between the computer programming self-efficacy perceptions of the student group that was given gamified e-learning supported instruction and the student group that was not given, no significant difference was found in terms of their attitudes towards programming and classroom engagement.

Qualitative Research Findings

The fourth sub-problem of the research is what are the opinions of the students who were given gamified e-learning supported education about the process? The opinions of the students about the gamified e-learning process they went through were obtained with an open-ended opinion determination form. The information obtained with the opinion determination form was examined with the content analysis method. The opinions obtained as a result of the content analysis were evaluated as frequency (f) and percentage (%). The evaluations that stood out from the participant opinions were stated with percentage expressions.

In this study, 46.7% of participants provided "positive" feedback on gamification. 46.7% of the participants thought that the scoring in the gamification setup in this activity should be changed. The rate of participants who stated that the “badges” attracted their attention the most in gamification was again 46.7%. 50% of those who chose badges stated that the badges attracted attention. 80% of the participants stated that the badges were encouraging. 93.3% of the participants stated that the leaderboard was motivating. 73.3% of the participants stated that it increased their classroom engagement. 80% of the participants stated that gamified e-learning supported teaching had a positive effect on the subjects they had difficulty with. 86.72% of the participants stated that other courses should also be gamified. 50% of those who thought that other courses should be gamified preferred science courses. Finally, 93.3% of the participants stated that EBA should be gamified on a course basis, as in this activity.

Discussion and Conclusion

Evaluation of the Programming Attitude Scale Results

In this study, no statistically significant difference was found between the group receiving gamified e-learning-supported instruction and the group not using this method in terms of attitudes toward programming. Similar findings have also been reported in the literature. For example, Başer (2013) found that computer engineering students exhibited positive attitudes, whereas computer teacher candidates demonstrated neutral attitudes in his study.

In the gamified group, it is suggested that the competition created by elements such as earning points, obtaining badges, and appearing on the leaderboard may lead to negative attitudes among some students. Yıldırım and Demir (2016), along with Hebebe and Usta (2018), noted that competition in gamification activities could result in certain negative outcomes. Similarly, Yıldız (2018) argued that the level of competition and difficulty might generate negative perceptions among students. Furthermore, Ibanez, Di-Serio ve Delgado-Kloos (2014) found that student participation in the system decreased when competition ended.

Additionally, the abstract nature of programming concepts may pose challenges for participants without prior experience in programming, potentially leading to negative attitudes. According to Başer (2013), perceiving programming courses as difficult may cause learners to develop unfavorable attitudes toward the subject, ultimately resulting in failure. Moreover, Yıldız (2018) emphasized that virtual rewards in gamification do not effectively foster intrinsic motivation.

Evaluation of the Classroom Engagement Scale Results

No statistically significant difference was found between the group receiving gamified e-learning-supported instruction and the group not using this method in terms of class participation. A review of the literature reveals similar findings in other studies. For instance, Ronimus et al. (2014) reported comparable results in their research.

Various factors influence students' attitudes, one of which is the perception of programming as a challenging subject, which can have negative effects on students. For instance, Yıldız (2018) conducted a study examining the impact of gamified block-based algorithmic thinking activities on variables such as attitude toward programming and participation. In the experimental groups, gamified activities involving competition, rewards, and challenges were conducted weekly for three weeks. As a result of the study, a significant difference was found between the

experimental and control groups, although no significant difference was observed between the reward and challenge weeks in the experimental group. Furthermore, motivational elements such as rewards and incentives have been found to occasionally fail in producing the desired effects on students.

Hakulinen et al. (2015) discovered that while some types of rewards served their intended purpose, others were ineffective for certain students. Similarly, Hanus and Fox (2015) stated that rewards in some studies did not produce meaningful effects and, in some cases, even reduced intrinsic motivation. Erümit (2016) also reported that the use of rewards can sometimes decrease intrinsic motivation and lead to negative outcomes.

Evaluation of the Self-Efficacy Scale Results

The study found statistically significant differences in programming self-efficacy between the class receiving gamified e-learning-supported instruction and the class not using this method. Similar findings have been reported in the literature. For instance, Kaya (2024) identified significant differences in self-efficacy perceptions related to coding through the use of gamification in block-based robotics and coding education. Likewise, Özyurt and Özyurt (2015), in their study involving 325 students, determined that participants had a moderate level of programming self-efficacy. Serim (2019) reported that coding education structured with a gamification approach positively impacted students' self-efficacy perceptions in a statistically significant manner. Kasalak (2017), in a study involving middle school students, observed that robotic coding activities positively influenced self-efficacy perceptions related to block-based programming. Similarly, Yağcı (2016) found significant differences in self-efficacy perceptions among computer programming and information technology teacher candidates in his study.

Evaluation of the Face-to-Face Interview Data

Overall, face-to-face interviews with students revealed that they often forgot to read the weekly announcement component within the activity. The EBA software includes a real-time evaluation feature, which addresses such challenges. Similarly, it is believed that providing teachers with the ability to create gamification scenarios tailored to specific lessons or topics—allowing students to earn points, badges, and participate in weekly group leaderboards—could effectively mitigate these issues.

Conclusion

Gamification is an effective tool with a positive impact on enhancing self-efficacy perceptions related to programming. It is suggested that gamification can provide effective solutions to self-efficacy challenges encountered in programming education, which inherently involves abstract concepts. However, it has been observed that the effects of gamification are not always consistent and may vary depending on factors such as different learning environments, topics, cognitive levels, and instructors. Additionally, findings indicate that varying gamification components may produce different outcomes depending on students' interests.

This study was conducted with a group of students who have academic scores above the general average at an educational institution that admits students through exams. This situation strengthens the possibility that gamification may have different effects on various student groups. Therefore, it is suggested that practitioners should design gamification frameworks that are suitable for the specific student group.

In conclusion, it is understood that the careful design of the gamification framework, tailored to the characteristics of the instructor, students, institution, time, and course, as well as its diversification when necessary throughout

the process, is of critical importance. Additionally, the Educational Information Network (EBA), due to its ease of use, speed, accessibility, and the variety of educational materials it houses, is regarded as one of the fundamental tools for e-learning in the Turkish National Education System. EBA, which includes elements of gamification within itself, offers teachers the opportunity to implement gamification at the course level, which could yield significant educational outcomes. The results of this study are expected to make a significant contribution to uncovering the potential and limitations of gamification in education and to provide valuable insights for future research in this field.

Giriş

Günümüzde bilişim teknolojilerinde yaşanan hızlı gelişmeler, teknolojik ürün ve yazılımların hayatın hemen her alanında kullanılmasına olanak tanımaktadır. Bu ilerlemeler, programlama ve robotik kodlama gibi derslerin popülerliğini artırmış ve eğitim sistemlerinde bu alanlara daha fazla önem verilmesine neden olmuştur. Milli Eğitim Bakanlığı (MEB) da bu değişimlere paralel olarak bilişim, programlama ve robotik kodlama gibi alanlara yönelik çalışmalarını artırmış, eğitim materyalleri ve altyapılarında yenilikler yapmıştır. Bu yenilikler, ders içeriklerinin teknolojideki değişimle uyum sağlamasıyla birlikte, programlama eğitimine verilen önemin daha belirgin hale gelmesini sağlamıştır.

Son yıllarda Türk eğitim sistemine yönelik yenilikler incelendiğinde, MEB tarafından 2011 yılında başlatılan FATİH (Fırsatları Arttırma ve Teknolojiyi İyileştirme Hareketi) Projesi ve 2012 yılında "Eğitim Daima" sloganıyla hayata geçirilen EBA (Eğitim Bilişim Ağı) gibi uygulamalar dikkat çekmektedir. Bu projeler sayesinde okullarda hızlı internet altyapısı sağlanmış ve EBA ile bir e-öğrenme ortamı oluşturulmuştur. EBA, K12 seviyesindeki tüm öğrenci gruplarına hitap eden bir öğrenme yönetim sistemi olarak kullanılmaktadır (Keleş, 2022).

Programlama ve yazılım eğitimine verilen önem, yalnızca okulların teknolojik altyapısının iyileştirilmesiyle sınırlı kalmamış, ders içeriklerinin güncellenmesiyle daha da belirgin hale gelmiştir. Bilişim teknolojileri eğitimi, Türkiye'de ilk olarak 1998 yılında Milli Eğitim Bakanlığı tarafından "Seçmeli Bilişim Teknolojileri" dersi olarak ilköğretim programlarına dahil edilmiştir (Uzgun ve Aykaç, 2016). Bu dersin temel amacı, öğrencilere bilgisayar okuryazarlığı kazandırmaktır. Daha sonra, 4+4+4 eğitim sistemiyle birlikte "Bilişim Teknolojileri ve Yazılım" dersi olarak yeniden adlandırılmış ve ilköğretim beşinci sınıftan itibaren programlama eğitimi verilmeye başlanmıştır. 2024 yılında ise "Türkiye Yüzyılı Maarif Modeli" kapsamında matematik derslerinde "algoritma ve bilişim odaklı" bir müfredatın uygulanacağı açıklanmıştır (MEB, 2024). Bu yenilikler, bilişim ve programlamaya verilen önemin giderek arttığını göstermektedir.

Hem bireylerin hem de eğitim sistemlerinin giderek daha fazla önem verdiği programlama eğitiminde çeşitli zorluklarla karşılaşmaktadır. Yapılan çalışmalar, bu zorlukları programlamaya özgü sorunlar, öğretmen yaklaşımlarından kaynaklanan sorunlar ve yeni teknolojik imkanların yeterince kullanılmaması olarak gruplandırmıştır. Demir (2015), bu sorunları; programlamanın soyut yapısı, programlamanın algılanışı, programlama kaygısı ve öğretim yöntemlerinin yetersizliği olarak ele almıştır. Literatürde, karşılaşılan bu sorunlara yönelik çözüm önerileri de tartışılmaktadır.

Programlama eğitimi sırasında, programlamanın soyutluk ve karmaşıklık özellikleri nedeniyle çeşitli zorluklarla karşılaşmaktadır. Gomes ve Mendes (2007), bu zorlukları öğretim yöntemleri, öğrencilerin becerileri ve davranışları, programlamanın soyut yapısı ve psikolojik etkenler gibi başlıklar altında incelemiştir. Programlamanın karmaşık sözdizimi ve soyut yapısı, başlangıç seviyesindeki öğrenciler için zorlu bir engelle dönüşmekte, bu durum öğrencilerin programlamadan soğumalarına ve çaba sarf etmemelerine yol açmaktadır (Ersoy ve diğerleri, 2011). Soyut programlama kavramlarının öğrenciler için anlaşılması güç olması, öz yeterliliklerini ve motivasyonlarını düşürmektedir (Atabaş, 2018). Başer (2013), programlama derslerinin öğrenenler tarafından zor olarak algılanmasının, olumsuz tutumlara ve başarısızlıklara yol açtığını ifade

etmektedir. Bu nedenle, soyut programlama kavramlarının somutlaştırılması ve eğlenceli bir şekilde anlatılması önerilmektedir (Başer, 2013; Ersoy ve diğerleri, 2011; Özyurt ve Özyurt, 2016).

Byrne ve Lyons (2001), programlama eğitimindeki zorlukların temel nedeninin geleneksel öğretim yöntemleri olduğunu belirtmiştir. Geleneksel yöntemlerde, öğrencilerin kitaplarda yer alan kuralları ezberlemesi beklenmekte ve bu durum öğrenme sürecini olumsuz etkilemektedir. Connolly, Murphy ve Moore (2007) ise programlama kaygısının öğrencilerin yeteneklerini yanlış değerlendirmelerine neden olduğunu ve bunun, öğrencilerin derse aktif katılımını engellediğini ifade etmiştir.

Özyurt ve Özyurt (2016), programlama derslerindeki başarısızlık nedenlerinden birinin yetersiz pratik olduğunu vurgulamıştır. Bu sorunun çözümü için derslerde bol örnek yapılması ve geçmiş konuları içeren küçük proje ödevleri verilmesi önerilmektedir. Ayrıca, her öğrencinin öğrenme hızı farklı olduğundan, ders içeriklerinin öğrencilerin erişebileceği bir platformda paylaşılması, bireysel öğrenme ihtiyaçlarını karşılamada etkili olacaktır.

Sorunların Çözümüne Yönelik Öneriler

Programlama eğitiminde karşılaşılan sorunları çözmeye e-öğrenme ortamları önemli bir potansiyele sahiptir. E-öğrenme; fırsat eşitliği, bilgiye erişimde zaman sınırlamalarını ortadan kaldırması ve maliyetleri düşürmesi gibi avantajlar sunmaktadır (Özgür, 2011). Bu avantajlar hem uzaktan eğitimi desteklemekte hem de harmanlanmış eğitim modellerine katkı sağlamaktadır. E-öğrenme ortamlarının oyunlaştırma teknikleri ile desteklenmesi ise, öğrencilerin motivasyonunu artırarak öğrenme sürecini daha etkili hale getirebilir.

Oyunlaştırma, zorlu süreçleri daha motive edici hale getirerek eğitimde başarıyı artırma potansiyeline sahiptir (Deterding ve diğerleri, 2011; Muntean, 2011). Eğitim bağlamında oyunlaştırma, öğrenenlerde ilgi ve motivasyon oluşturmak amacıyla oyun tasarım ilkelerinin kullanıldığı bir yöntemdir (Kocadere ve Çağlar, 2018). Bu teknik, e-öğrenme ortamlarının sınırlılıklarını gidermede etkili bir araç olarak öne çıkmaktadır.

Çalışmanın Amacı ve Alt Problemler

Bu çalışmanın amacı, 9. sınıf bilişim teknolojileri yazılım dersinde oyunlaştırılmış e-öğrenme destekli öğretimin, öğrencilerin bilgisayar programlama öz yeterlik algıları, programlamaya yönelik tutumları ve derse katılımları üzerindeki etkilerini incelemektir. Araştırma kapsamında şu alt problemler ele alınmıştır:

1. Oyunlaştırılmış e-öğrenme destekli öğretim verilen öğrenci grubu ile verilmeyen grup arasında programlamaya yönelik tutumlar açısından anlamlı bir fark var mıdır?
2. Oyunlaştırılmış e-öğrenme destekli öğretim verilen öğrenci grubu ile verilmeyen grup arasında derse katılım açısından anlamlı bir fark var mıdır?
3. Oyunlaştırılmış e-öğrenme destekli öğretim verilen öğrenci grubu ile verilmeyen grup arasında programlama öz yeterlik algıları açısından anlamlı bir fark var mıdır?
4. Oyunlaştırılmış e-öğrenme destekli öğretim verilen öğrencilerin süreçle ilgili görüşleri nelerdir?

Yöntem

Araştırma Modeli

Bu araştırmada, nicel ve nitel aşamaların birlikte kullanıldığı karma araştırma yöntemi kullanılmıştır. Çalışmada nicel araştırma yöntemlerinden “Öntest-Sontest Kontrol Gruplu Yarı Deneysel Desen” kullanılmıştır. Bu araştırmada yarı deneysel çalışma için deney ve kontrol grupları yansız belirlenirken, gruplardaki denekler ise yansız belirlenmemiştir. Araştırma örgün eğitimin devam ettiği süreçte öğrenciler üzerinde gerçekleştirildiği için öğrencilerin gruplarının değişmesi devam etmekte olan eğitim öğretimi aksatabilir ve EBA (eğitim bilişim ağı) üzerinden gruplama ise oyunlaştırmada karışıklığa sebep olabilir. Çalışmada rastgele atanan iki sınıftan biri kontrol diğeri ise deney grubu olarak çalışmaya dâhil edilmiştir. Çalışma öncesinde her iki grubun da bağımlı değişkenler üzerinden ölçümleri alınmıştır. Deney grubu ile yapılan deneysel çalışma sonrasında tekrar grupların bağımlı değişkenler üzerinden ölçümleri alınmış ve istatistiksel olarak anlamlı farklılık olup olmadığına dair analiz yapılmıştır. Çalışmanın deneysel deseni Tablo 1’de gösterilmektedir.

Tablo 1. Ön Test – Son Test Kontrol Gruplu Desenin Sembolik Gösterimi

Grup		Ön test	İşlem	Son test
D (Deney)	M	Q _{1,2,3}	X	Q _{1,2,3,4}
K (kontrol)	M	Q _{1,2,3}		Q _{1,2,3}

D: Oyunlaştırılmış e-öğrenme ortamıyla desteklenen grup

K: Sadece MEB öğretim programını kullanan grup

X: Uygulama

Q₁: Programlama öz yeterlik ölçeği

Q₂: Programlama tutum ölçeği

Q₃: Derse katılım envanteri

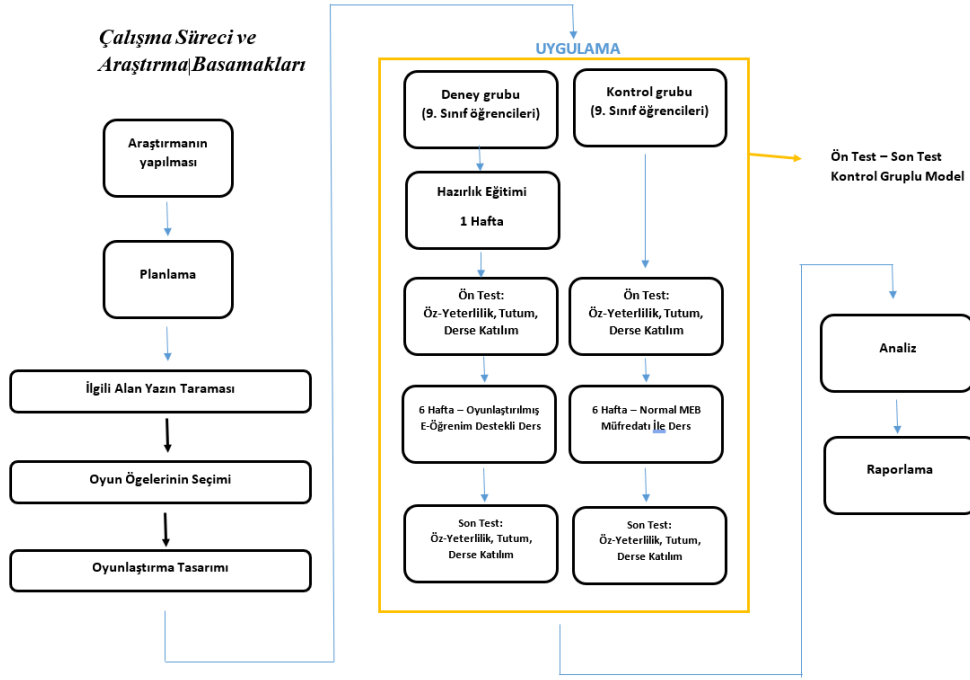
Q₄: Yarı yapılandırılmış görüşme formu

M: Eşleştirilmiş grupların uygulama gruplarına yansız atanması

Çalışma Grubu

Çalışma grubu, 2023-2024 eğitim öğretim yılında, Zonguldak ilinde bir Fen Lisesinde iki şubede öğrenim gören gönüllü 60 öğrencidir. Çalışma grubu çalışmanın uygulandığı okulda bilgisayar sınıfı bulunması, araştırmacı tarafından kolay ulaşılabilir olması ve öğrencilerin gönüllü olması sebebiyle tercih edilmiştir.

Çalışmanın süresi altı hafta olarak belirlenmiştir. Yapılan araştırmanın işlem basamakları Şekil 2’te gösterilmektedir.



Şekil 1. Çalışma Süreci ve Araştırma İşlem Basamakları

Çalışma 9. Sınıflarda okutulan Bilişim Teknolojileri ve Yazılım dersinde yürütülmüştür. Kodlama eğitiminde metin tabanlı programlama dili olan Python kullanılmıştır.

Çalışma grubunda 30 öğrenciden oluşan şube deney grubu, 30 öğrenciden oluşan başka bir şube ise kontrol grubu olarak belirlenmiştir.

Veri Toplama Araçları

Araştırmada nicel verileri toplamak için Özkara ve Yelken (2020) tarafından geliştirilen ortaöğretim öğrencilerine yönelik programlama öz yeterlik ölçeği, Altay ve Kışla (2018) tarafından geliştirilen programlama tutum ölçeği, Orjinali “Classroom Engagement Inventory” ismiyle Wang ve diğerleri (2014) tarafından İngilizce olarak geliştirilip Sever (2014) tarafından Türkçeye uyarlanan derse katılım envanteri ve nitel verileri almak için ise araştırmacı tarafından geliştirilen yüz yüze görüşme formu kullanılmıştır.

Oyunlaştırma Sürecinin Kurgusu ve Çalışmada Kullanılan Rozetler

Çalışmada kullanılan oyunlaştırma bileşenleri aşağıda verilmiştir:

- 1- Puan
- 2- Lider Tahtası
- 3- Seviye
- 4- Rozet

Öğrenenler, paylaşılan etkinliklerden topladıkları öğrenme puanları doğrultusunda araştırmacı tarafından geliştirilen dijital rozetleri kazanabilirler. Bu rozetler hakkında bilgi aşağıdaki listede verilmiştir.

Rozet Türleri ve Özellikleri

Oyunlaştırma kurgusunda, öğrencilerin katılımını ve motivasyonunu artırmayı amaçlayan çeşitli rozetler tasarlanmıştır. Araştırmacı tarafından tasarlanan rozetler aşağıdaki şekilde sınıflandırılmıştır:

1. **Etkinlik Rozetleri:** Giriş rozeti, ilk haftanın ilk etkinliğine özel olarak tasarlanmıştır ve bu tür etkinliklerde başlangıç motivasyonunu desteklemeyi hedeflemektedir.
2. **Seviye Rozetleri:** Katılımcıların ilerleme durumlarını göstermek amacıyla 1'den 10'a kadar farklı seviyelerde tasarlanmış rozetler sunulmuştur.
3. **Haftalık Derece Rozetleri:** Haftanın performansına göre dereceye giren katılımcılar için üç farklı rozet tasarlanmıştır: *Zirve (1. sırada olanlar)*, *Zirve Ortağı (2. sırada olanlar)* ve *Azimli (3. sırada olanlar)*.
4. **Genel Derece Rozetleri:** Tüm etkinlikler sonunda genel sıralamaya giren katılımcılara özel olarak üç farklı rozet verilmiştir: *Efsane (1. sırada olanlar)*, *Bilge (2. sırada olanlar)* ve *Usta (3. sırada olanlar)*.

Oyunlaştırma Sürecinin Takibinde Puanlama ve Rozet Tabloları

Yapılan oyunlaştırma kurgusu çerçevesinde öğrencilerin etkinliklere katılımı puanlanmıştır. Çalışma haftalık olarak Şekil 2'de gösterildiği gibi tablolama programı aracılığıyla takip edilmiştir.

A	B	C	D	E	F	G	H	I	J	K	L	M	N
			HAFTA 1	HAFTALIK TOPLAM PUAN		1-1 DUYURU	1-2 DOKÜMAN	1-3 VİDEO					14-TEST
						Katıl 20puan	Gördüyse 20puan	İzlemişse 20puan				Testten puanı 100üzerinden	Dönüştürüldü 40puana
Sıra	Şube	No	İsim										
1				92		20	20	20				80	32
2				58			20	20				45	18
3				68		20	20	20				20	8
4				0									0
5				44			20	20				10	4
6				92		20	20	20				80	32
7				52			20	20				30	12
8				92		20	20	20				80	32
9				0									0
10				72			20	20				80	32
11				50			20					75	30
12				92		20	20	20				80	32
13				74		20	20	20				35	14

Şekil 2. Haftalık Puan Takip Tablosu

Haftalık Tahta ve Rozet Kazananların Takibi

Haftalık olarak öğrencilerin aldıkları puana göre giriş rozeti, seviye rozeti ve varsa derecelerine dair derece rozetleri ilan edilmiştir. Tüm haftalar sonucunda rozet kazanan ilk 10 Şekil 3'te gösterildiği gibi ilan edilmiştir.

Oyunlaştırılmış E-Öğrenme Destekli Öğretimin Öğrencilerin Bilgisayar Programlama Öz Yeterlik Algıları,
Programlamaya Yönelik Tutumları ve Derse Katılımlarına Etkisi

ÖĞRENCİ ROZET LİSTESİ - İLK 10 – TÜM HAFTALAR SONUC													
S	Toplam Puan	Öğrenci Adı	Giriş Rozeti	Seviye Rozeti	Haftalık Derece Rozetleri			Beğeni Rozeti	Okuma Rozeti	İzleme Rozeti	Son Derece Rozetleri		
1	2	3	4	5	6	7	8	9	10	11	12	13	14
			Merhaba Dünya	1-10	1 ZİRVE	2 ZİRVE ORTAĞI	3 AZİMLİ	ÇALIŞKAN	OKUR	SEYİRCİ	1 EFSANE	2 BİLGE	3 USTA
1	568	M.O.											
2	554	S.D.											
3	512	Y.A.											
4	513	N.C.											
5	500	K.A.											
6	484	U.K.											
7	477	E.B.	-										
8	478	A.K.											
9	466	B.K.											
10	468	B.Ç.											

Şekil 3. Oyunlaştırma Rozet Kazanan İlk 10

Verilerin Analizi

Araştırmanın nicel veriler ise SPSS istatistik programı ile irdelenmiştir. Oyunlaştırılmış e-öğrenmenin derse katılımlarına, programlama öz yeterlik etkisine ve programlamaya karşı tutumlarına etkilerinin anlamlı olup olmadığına bakmak için Kovaryans Varyans Analizi (ANCOVA) uygulanmıştır.

Öğrenci görüşlerini belirleme formuna verilen yanıtların analiz edilmesinde içerik analizinden yararlanılmış olup bulgular yüzde ve frekans şeklinde ifade edilip yorumlanmıştır. Analiz sürecinde görüşme sorularını kapsayacak şekilde araştırmacı ve bir alan uzmanı, birbirinden bağımsız olarak görüşme metinlerinden temalar çıkartılmıştır, daha sonra bir araya gelerek aralarında fikir birliğine varmış ve görüşme temalarına son şeklini vermişlerdir. İçerik analizi öğrenci görüşlerini belirleme formuna verilen yanıtlardan çıkarılan temalara dayalı olarak gerçekleştirilmiştir.

Bulgular

Nicel Bulgular

Bu çalışmada oyunlaştırılmış e-öğrenme destekli öğretim verilen öğrenci grubuyla verilmeyen öğrenci grubu arasında bilgisayar programlama öz yeterlik algıları arasında anlamlı fark olduğu tespit edilirken, programlamaya yönelik tutumları ve derse katılımları açısından anlamlı fark tespit edilememiştir. Tüm ölçeklerin Kolmogorov-Smirnov testi sonucundan anlamlılık değerleri $P>0.5$ olarak tespit edilmesi nedeniyle verilerin normal dağılım gösterdiği sonucuna ulaşılmıştır. Öğrencilerin test ortamlarının birbirine denk olması nedeniyle dış geçerlik

varsayımı sağlandığı kabul edilmiştir. Çalışma öncesi ve sonrasında uygulanan ölçeklere dair yapılan analiz sonuçları Tablo 2’de gösterilmiştir.

Tablo 2. Programlama Öz Yeterlik, Tutum ve Derse Katılım Ölçeği Analiz Sonuçları

Ölçek ve Alt Boyutlar	Gruplar	$\bar{X}$		Ss	
		Ön Test	Son Test	Ön Test	Son Test
Programlama Öz Yeterlik	Deney	61.90	66.46	8.21	8.64
	Kontrol	62.96	62.10	10.27	9.49
Programlama Tutum	Deney	46.90	45.70	7.35	10.08
	Kontrol	45.50	41.63	11.17	9.55
Derse Katılım	Deney	77.03	79.53	12.58	15.54
	Kontrol	75.00	76.50	14.44	12.30

Programlama öz yeterlik ölçeği kovaryans sonuçları [$F=(1,57)=4.961$; $p=0.30<0.05$] şeklinde, programlama tutum ölçek kovaryans sonuçları [$F=(1,57)=2.191$; $p=0.144<0.05$] şeklinde ve derse katılım ölçeği kovaryans sonuçları da [$F=(1,57)=4.291$; $p=0.591<0.05$] şeklinde olmuştur. Bu analizlere dair veriler Tablo 3’te verilmiştir.

Tablo 3. Programlama Öz yeterlik, Tutum ve Derse Katılım Ölçeklerine Dair Kovaryans Analiz Sonuçları.

Programlama öz yeterlik ölçeği kovaryans analiz sonuçları					
Varyansın Kaynağı	Kareler Toplamı	Serbestlik Derecesi (sd)	Kareler Ortalaması	F	Anlamlılık Düzeyi (P)
Ön test	827.028	1	827.028	11.913	.001
Grup	34.376	1	34.376	4.961	.030
Hata	3957.138	57	69.423		
Toplam	253011.000	60			
Programlama tutum ölçeği kovaryans analiz sonuçları					
Varyansın Kaynağı	Kareler Toplamı	Serbestlik Derecesi (sd)	Kareler Ortalaması	F	Anlamlılık Düzeyi (P)
Ön test	624.217	1	624.217	7.149	.010
Grup	191.276	1	191.276	2.191	.144
Hata	4977.050	57	87.317		
Toplam	120256.000	60			
Derse katılım ölçeği kovaryans analiz sonuçları					
Varyansın Kaynağı	Kareler Toplamı	Serbestlik Derecesi (sd)	Kareler Ortalaması	F	Anlamlılık Düzeyi (P)
Ön test	3882.134	1	3882.134	29.427	.000
Grup	38.436	1	38.436	.291	.591
Hata	7519.633	57	131.923		
Toplam	377655.000	60			

Nitel Bulgular

Araştırmamızın dördüncü alt problemi, oyunlaştırılmış e-öğrenme destekli öğretim verilen öğrencilerin süreçle ilgili görüşlerinin neler olduğudur? Öğrencilerin geçirdikleri oyunlaştırılmış e-öğrenme sürecine ilişkin görüşleri açık uçlu görüş belirleme formu ile alınmıştır. Görüş belirleme formu ile elde edilen bilgiler içerik analizi yöntemiyle incelenmiştir. İçerik analizi sonucu elde edilen görüşlerin frekans(f) ve yüzde (%) olarak değerlendirilmiş. Katılımcı görüşlerinden öne çıkan değerlendirmeler Tablo 4’teki gibidir.

Oyunlaştırılmış E-Öğrenme Destekli Öğretimin Öğrencilerin Bilgisayar Programlama Öz Yeterlik Algıları,
Programlamaya Yönelik Tutumları ve Derse Katılımlarına Etkisi

Tablo 4. Yüz Yüze Görüşme Bulguları

Yüz yüze Görüşme Sonuçları	Oran (%)
Oyunlaştırılmış e-öğrenme destekli öğretim sürecinin olumlu terimlerle ifade edenler	46.7
Bu çalışmadaki oyunlaştırma senaryosundaki puanlamanın değişmesi gerektiğini düşünenler	46.7
Oyunlaştırmada en çok “rozet” bileşeninin dikkatini çektiğini ifade edenler.	46.7
“Rozet” için “ilgi çekici” ifadesini kullananlar.	50
“Rozet kazanmak teşvik edici miydi?” sorusuna “Evet” belirenler.	80
Liderlik tahtasının motive edici olduğunu düşünenler.	93.3
Oyunlaştırmanın derse katılımlarına etki ettiğini düşünenler.	73.3
Zorlanılan konularda oyunlaştırılmış e-öğrenme ile katkı sağlandığını düşünenler	80
Diğer derslerin oyunlaştırılmış e-öğrencim destekli olmasını tercih	86.7
Hangi dersi oyunlaştırılarak işlenmesini isterdiniz soruna “fen” derslerini yazanlar.	50
EBA’nın bu etkinlikte olduğu gibi ders bazında oyunlaştırılarak güncellenmesinin olumlu olacağını düşünenler	93.3

Tartışma ve Sonuç

Kullanılan ölçeklerden elde edilen veriler doğrultusunda oyunlaştırılmış e-öğrenme desteği verilen deney grubu ile standart müfredatın geleneksel yöntemlerle işlendiği kontrol grubu arasında programlama öz yeterlikleri açısından istatistiksel olarak anlamlı farklar bulunurken, programlamaya yönelik tutumlar ve derse katılım açısından anlamlı bir fark tespit edilmemiştir.

Tartışma

Programlama Tutum Ölçeği Sonuçlarının Değerlendirilmesi

Yapılan çalışmada, kullanılan ölçeklerden elde edilen verilere göre, oyunlaştırılmış e-öğrenme destekli öğretim yapılan grup ile bu yöntem uygulanmayan grup arasında, programlamaya karşı tutum açısından istatistiksel olarak anlamlı bir fark bulunamamıştır. Literatürde yapılan benzer çalışmalarda da bu yönde sonuçlar elde edilmiştir. Örneğin, Başer (2013), bilgisayar mühendisliği öğrencileri ve bilgisayar öğretmeni adaylarıyla yaptığı çalışmada, mühendis adaylarının tutumlarının olumlu, öğretmen adaylarının tutumlarının ise nötr olduğunu tespit etmiştir.

Oyunlaştırma yapılan grupta, puan, rozet kazanma ve lider tahtasına girme gibi unsurların yarattığı rekabetin, bazı öğrencilerde olumsuz tutumlara neden olduğu düşünülmektedir. Literatüre bakıldığında araştırmacıların bazıları çalışmalarında ödül unsurunu öne çıkartırken (Hakulinen ve diğerleri, 2013), bazıları rekabet ögesine önem vermekte (Hakulinen ve diğerleri, 2015) ve bazıları da ödül, rekabet, zorlaştırma gibi öğeleri öne çıkarmaktadır. Yıldırım ve Demir (2016) ile Hebece ve Usta (2018) oyunlaştırma aktivitelerinde ortaya çıkan rekabetin bazı olumsuz sonuçlarının olabileceğini belirtmişlerdir. Yıldız’a (2018) göre, rekabet ve zorluk seviyesi, öğrencilerde olumsuz algılar oluşturabilmektedir. Benzer şekilde, Ibanez, Di-Serio ve Delgado-Kloos (2014), çalışmalarında rekabet sona erdiğinde öğrencilerin sisteme katılımlarında azalma olduğunu tespit etmişlerdir.

Bununla beraber, programlamanın soyut kavramlar içermesi, programlamayla doğrudan ilişkisi olmayan katılımcıların zorlanmasına ve bu zorlanmanın tutumlarını olumsuz yönde etkilemesine yol açmış olabilir.

Nitekim Başer'e (2013) göre, programlama dersinin öğrenenler tarafından zor olarak algılanması, bu derse yönelik olumsuz tutum sergilemelerine ve dolayısıyla başarısız olmalarına yol açabilmektedir. Yıldız'a (2018) göre oyunlaştırmadaki sanal ödül ise gerçek manada içsel teşvik oluşturmamaktadır.

Alanyazın incelendiğinde, oyunlaştırmının tutum üzerinde anlamlı sonuçlar ürettiği çalışmaların mevcut olduğu görülmektedir. Örneğin, Polat (2014) tarafından yapılan bir çalışmada, dil öğreniminde olumlu akademik sonuçlar elde edilememesine rağmen katılımcılarda olumlu tutumlar gözlemlenmiştir. Yıldırım (2016), karma desenli çalışmada oyunlaştırmının öğrencilerin tutumları üzerinde istatistiksel olarak anlamlı ve pozitif bir etki ortaya koyduğunu ifade etmiştir.

Gürer ve Tokumacı (2020), gerçekleştirdikleri bir araştırmada üniversite öğrencilerinin programlamaya karşı tutumlarını incelemiş ve bu tutumlarda duyuşsal ve davranışsal boyutlarda yüksek düzeyde sonuçlar elde ederken, bilişsel boyutta ise orta düzeyde sonuçlara ulaşmıştır. Özdemir (2024), öğretmen adayları üzerinde yaptığı çalışmada, programlama öz yeterlik ve tutum ölçeğini kullanmış; analizler sonucunda, katılımcıların programlamaya karşı orta düzeyde bir tutum sergilediklerini belirtmiştir. Ayrıca, katılımcılar oyunlaştırma sayesinde dersin tekdüzeliğinden kurtulduğunu, rekabet hissettiklerini, kalıcı öğrenme sağladıklarını ve derse aktif katılım gösterdiklerini ifade etmişlerdir.

Son olarak, Özyurt ve Özyurt (2015), katılımcıların programlamaya karşı tutumlarını, programlama öz yeterliklerini ve bu iki unsur arasındaki ilişkiyi inceledikleri çalışmada, öğrencilerin programlamaya karşı olumlu bir tutum sergilediklerini tespit etmiştir.

Derse Katılım Ölçeği Sonuçlarının Değerlendirilmesi

Kullanılan ölçeklerden elde edilen veriler doğrultusunda, oyunlaştırılmış e-öğrenme destekli öğretim yapılan grup ile bu yöntemin uygulanmadığı grup arasında derse katılım ölçeği analizi sonuçlarına göre istatistiksel olarak anlamlı bir farklılık bulunmamıştır. Alanyazın incelendiğinde, bu yönde benzer sonuçlar üreten çalışmalar da mevcuttur. Örneğin, Ronimus ve diğerleri (2014) yaptıkları bir çalışmada benzer bulgulara ulaşmışlardır.

Öğrencilerin tutumlarını etkileyen çeşitli faktörler bulunmaktadır. Bunlardan biri, programlamanın zorluk algısının öğrenciler üzerinde olumsuz etkiler yaratmasıdır. Nitekim, Yıldız'a (2018) da oyunlaştırılmış blok tabanlı algoritmik düşünme etkinliklerinin programlamaya karşı tutum ve katılım gibi değişkenlere etkisini incelediği bir çalışma yapmıştır. Deney gruplarda 3 hafta boyunca her hafta oyunlaştırılmış rekabet, ödül ve zorluk etkinlikleri yapmıştır. Çalışma sonucunda deney ve kontrol gruplarında anlamlı fark çıkarken, deney grubunda, ödül ve zorluk haftalarına arasında anlamlı fark ortaya çıkmamıştır. Ayrıca, ödül ve teşvik gibi motivasyon unsurlarının bazı durumlarda öğrenciler üzerinde beklenen etkiyi yaratmadığı görülmüştür. Hakulinen ve diğerleri (2015) yaptığı bir çalışmada, ödül çeşitlerinden bazılarının amaca hizmet ettiği, bazılarının ise kimi öğrenciler üzerinde etkisiz kaldığı tespit edilmiştir. Hanus ve Fox (2015), bazı çalışmalarda ödüllerin kayda değer seviyede bir etki yaratmadığını, hatta bazı durumlarda ödülün içsel motivasyonu azaltabileceğini ifade etmişlerdir. Benzer şekilde, Erümit (2016) ödül kullanımının bazı durumlarda içsel motivasyonu düşürdüğünü ve olumsuz etkiler oluşturduğunu belirtmiştir.

Öğrenci katılımlarını etkileyen bir diğer etmenin de e-öğrenme ortamına erişimdeki sınırlı süre kısıtı olduğu düşünülmektedir. Deneysel çalışma sonrasında yapılan yüz yüze görüşmelerde, öğrencilere "Uygulamada nelerin değişmesini isterdiniz?" sorusu yöneltmiş ve öğrenciler, oyunlaştırma içeriklerine sınırlı süre sonra

ulaşılamamasını bir sorun olarak işaret etmişlerdir. Bu bulgu, Aban ve Fontanil (2015) ile Albayrak'ın (2012) çalışmalarında da desteklenmiştir. Araştırmalar, zaman kısıtlamasının öğrenciler için rahatsızlık verici olduğunu ve bu durumun katılımcıların içsel motivasyonunu azaltarak etkinliğe ilgilerini düşürdüğünü göstermektedir. Ayrıca, zaman kısıtlaması süre bitmeden gelişi güzel işlem yapmalarına, heyecanlanmalarına ve başarısızlık hissi yaşamalarına neden olabilmektedir (Hakulinen ve diğerleri, 2015). Dolayısıyla, oyunlaştırma ve e-öğrenme sürecindeki zaman kısıtlamasının yanı sıra, e-öğrenme ortamının tasarım özelliklerinin de öğrencilerin deneyimlerini ve tutumlarını etkileyebileceği düşünülmektedir. Yılmaz ve Keser'in (2016) çalışması, öğrencilerin görsel ve işitsel açıdan zenginleştirilmiş e-öğrenme ortamlarını tercih ettiğini göstermekte ve bu tür tasarımların, katılımcıların motivasyon ve ilgisini artırmada önemli bir rol oynayabileceğini vurgulamaktadır.

Ayrıca tüm öğrenenler üzerinde aynı oyunlaştırma bileşenlerinin benzer etkiler yaratmadığı düşünülmektedir. Bazı çalışmalar, öğrenci katılımını etkileyen farklı oyunlaştırma bileşenlerinin farklı sonuçlar ortaya çıkardığını göstermektedir. Örneğin, Ibanez ve diğerleri (2014), oyunlaştırma unsurlarının bilgisayar bilimleri öğrencilerinin öğrenme sürecine katılımı üzerindeki etkilerini inceledikleri çalışmalarında, bu unsurların bireyler arasında farklı etkiler yarattığını gözlemlemişlerdir. Çalışmada, rozeti motive edici bulan öğrenciler ile liderlik tahtasının etkisini daha fazla hisseden öğrencilerin farklılık gösterdiği, dolayısıyla farklı oyunlaştırma bileşenlerinin öğrenciler üzerinde farklı düzeylerde motive edici güç oluşturduğu tespit edilmiştir. Bu çeşitliliği destekleyen bir diğer nokta ise araştırmacıların bazılarının ödül unsurunu (Hakulinen ve diğerleri, 2013), bazılarının rekabet ögesini (Hakulinen ve diğerleri, 2015) ve bazılarının ise ödül, rekabet ve zorlaştırma gibi unsurları ön plana çıkarmasıdır (Ronimus ve diğerleri., 2014).

Bununla birlikte, literatürde oyunlaştırmanın derse katılım üzerinde anlamlı sonuçlar ürettiği çalışmalar da bulunmaktadır. Tunga'nın (2016) yaptığı bir çalışmada, kontrol grubunda klasik müfredat yüz yüze işlenirken, deney grubunda oyunlaştırılmış e-öğrenme ortamları kullanılmış ve deney grubu öğrencilerinin derse katılım oranlarının kontrol grubundan daha yüksek olduğu tespit edilmiştir. Çalışmaya katılan öğrenciler, oyunlaştırma bileşenleri içeren e-öğrenme ortamlarına daha istekli olduklarını ifade etmişlerdir.

Moccozet ve diğerleri (2013) gerçekleştirdiği bir çalışmada, oyunlaştırmanın derse katılımı artırdığı tespit edilmiştir. Barata ve diğerleri (2013), çevrimiçi olarak iki yıl süren oyunlaştırma çalışmalarında, öğrenci katılımını artırmak için deneyim puanları, lider tahtası ve seviye kademeleri gibi oyunlaştırma bileşenleri kullanmışlardır. Bu çalışmada, mühendislik öğrencilerinin katılımının olumlu yönde etkilendiği ve sistemde geçirdikleri sürenin arttığı tespit edilmiştir. Meşe ve Dursun'un (2018), üniversite öğrencileriyle 13 hafta boyunca yürüttükleri çalışmada, oyunlaştırmanın çevrimiçi ortamlara katılımı olumlu yönde etkilediği gözlemlenmiştir. Sarı ve Altun'un (2016) lise öğrencileri ile yaptığı bir başka çalışmada, oyunlaştırmanın derse katılımı artırdığı tespit edilmiştir.

Hamzah ve diğerleri (2015), e-öğrenme uygulamalarında oyun öğeleri içermesinin öğrenci katılımını artırdığını belirtmişlerdir. Landers ve Armstrong (2015) ise oyunlaştırmanın standart yöntemlere göre daha etkili bir yöntem olduğunu, öğrenenlerin motivasyonunu ve derse katılımını geliştirdiğini ifade etmişlerdir. Oyunlaştırma ve oyunlar, doğası gereği öğrenme ortamını daha eğlenceli hale getirerek öğrenenlerin içsel motivasyonlarına katkıda bulunur ve derse olan katılımı artırır (Erümit, 2016).

Bu çalışmaların birçoğu, derse katılım, motivasyon ve katılımcı memnuniyeti açısından olumlu sonuçlar ortaya koymuştur (Dominguez ve diğerleri, 2013; Polat, 2014; Tunga, 2016; Meşe, 2016; Erümit, 2016). Aguilar ve diğerleri (2015) tasarım temelli yaptıkları bir çalışmada, yüz yüze derslere oyun elemanları eklenmiş ve bu düzenlemenin öğrencilerin çalışmalarını artırdığı gözlemlenmiştir.

Öz yeterlik Ölçeği Sonuçlarının Değerlendirilmesi

Çalışmada, oyunlaştırılmış e-öğrenme destekli öğretim yapılan sınıf ile bu yöntemin uygulanmadığı sınıf arasında programlama öz yeterlikleri açısından istatistiksel olarak anlamlı farklar olduğu tespit edilmiştir. Alanyazında benzer sonuçlar mevcuttur. Kaya (2024), blok tabanlı robotik ve kodlama eğitimi üzerine yaptığı çalışmada, oyunlaştırma yöntemiyle kodlamaya ilişkin öz yeterlik algılarında anlamlı farklar bulmuştur. Özyurt ve Özyurt (2015) tarafından yapılan 325 öğrenciyle yapılan çalışmada, katılımcıların programlama öz yeterliklerinin orta düzeyde olduğu belirlenmiştir. Serim (2019), oyunlaştırma yaklaşımıyla yapılandırılan kodlama eğitiminde öğrencilerin öz yeterlik algılarının istatistiksel olarak olumlu yönde etkilendiğini ifade etmiştir. Kasalak (2017) ortaokul öğrencileri üzerinde yaptığı araştırmada, robotik kodlama etkinliğinin blok temelli programlamaya yönelik öz yeterlik algılarını olumlu yönde etkilediğini gözlemlemiştir. Yağcı (2016) ise bilgisayar programcılığı ve bilişim teknolojileri öğretmen adayları üzerinde yaptığı çalışmada, bilişim teknolojileri öğretmen adaylarının öz yeterlik algılarındaki anlamlı farkı tespit etmiştir.

Öğrencilerin öz yeterlik algılarını olumlu yönde etkileyen bir unsur rekabet olabilir. Örneğin Bicen ve Kocakoyun (2018) gerçekleştirdikleri bir çalışmada, oyunlaştırmadaki iş birliği ve rekabetin öğrenenleri pozitif yönde etkilediğini bildirmiştir. Ayrıca, öğrenenlerin geçmiş deneyimleri de etkili olabilmektedir. Mazman ve Altun (2013), bilgisayar öğretmenliği öğrencilerinden oluşan 64 katılımcı üzerinde yaptıkları çalışmada, öncesinde programlama dersi almamış olan katılımcılarda öz yeterlik algılarındaki artışın daha belirgin olduğunu gözlemlemiştir.

Alanyazında elde edilen olumlu sonuçların yanı sıra, oyunlaştırmının her zaman öz yeterlik algıları üzerinde kayda değer etki oluşturmadığını ortaya koyan araştırmalar da bulunmaktadır. Örneğin, Ortiz ve diğerleri (2017), mühendislik öğrencilerinin temel programlama kursunda oyunlaştırmının öz yeterlik algıları üzerinde önemli bir etkisi olmadığını belirtmişlerdir. Yine Eray (2022) tarafından yapılan bir çalışmada 8. sınıf öğrencilerine yönelik matematik dersinde uygulanan oyunlaştırma tabanlı etkinliklerin de öz yeterlik algıları üzerinde kayda değer bir etki ortaya koymadığı görülmüştür.

Yüz yüze Görüşme Verilerinin Değerlendirilmesi

Genel olarak değerlendirildiğinde, deney grubundaki öğrencilerin büyük çoğunluğu (%46,7), oyunlaştırılmış e-öğrenme destekli öğretim süreci hakkında olumlu görüş bildirmiştir. Öğrenciler, bu durumu oyunlaştırılmış e-öğrenme desteğinin derse eğlence katması ve dersin daha kolay ve kalıcı bir şekilde anlaşılmasına yardımcı olmasıyla açıklamaktadır. Araştırma sonuçlarına göre, öğrencilerin yaklaşık yarısı (%46,7) oyunlaştırmada kullanılan puanlama sisteminin değiştirilmesi gerektiğini belirtmiştir. Bu durumun sebebi olarak, duyuru/katıl puanının gözden kaçması veya unutulması ifade edilmiştir. Öğrenciler, bu puanların video, test gibi alanlara aktarılmasını önermiştir.

Ayrıca, Öğrenciler, etkinlik sırasında haftalık duyuruları okumayı unuttuklarını belirtmiştir. Mevcut oyunlaştırma senaryosunda, öğrenciler kendilerine verilen süre içinde görevlerini tamamlamakta, araştırmacı süre sonunda

raporları olarak değerlendirme yapmakta ve puanlama sonrasında oyunlaştırma bileşenlerine göre kazananları belirlemektedir. Ancak, öğrencilerin kazandıkları puanları görebilmeleri için araştırmacının değerlendirmesini tamamlamasını ve sonuçları ilan etmesini beklemeleri gerekmektedir. Bu durumun, değerlendirme sürecinin anlık ve senkronize bir şekilde yazılım tarafından gerçekleştirilmesi durumunda ortadan kalkacağı düşünülmektedir. Böyle bir uygulama, aynı zamanda öğrencilere hızlı geri bildirim sağlanmasını da mümkün kılacaktır. Bahsedilen anlık değerlendirme özelliğinin, EBA yazılımının mevcut oyunlaştırma kurgusunda yer aldığı görülmektedir. Benzer şekilde, öğretmenlerin ders ya da konu bazında oluşturacakları oyunlaştırma senaryolarında; puanlama, rozet kazanımı ve haftalık liderlik tahtası belirleme gibi özelliklerin sunulmasının, bu tür sorunlara etkili çözümler sağlayabileceği değerlendirilmektedir.

Öğrencilere, oyunlaştırma sürecinde en çok ilgilerini çeken öge hakkında sorulduğunda, yarısı (%46,7) "rozet" cevabını vermiştir. Literatüre bakıldığında ise Özgür, Çuhadar ve Akgün (2018) tarafından yapılan bir çalışmada, 2008-2017 yılları arasındaki ulusal ve uluslararası çalışmaların taraması sonucunda öğrencilerin en çok puan bileşenini öne çıkardıkları tespit edilmiştir. Öğrencilerin rozetleri seçme sebepleri arasında yarısı (%50,0), rozetlerin ilgi çekici olduğunu belirtmiştir. Araştırma sürecinde, puana ve etkinlik katılım durumuna göre verilen rozetlerin öğrencileri rozet kazanmaya teşvik edip etmediği sorusuna ise öğrencilerin büyük çoğunluğu (%80,0) "Evet" yanıtını vermiştir. Literatürde, Çelik, Yıldırım, Kaban ve Yıldırım (2017) tarafından yapılan bir çalışmada dijital rozetlerin öğrencilerin hem motivasyonlarını hem de derse katılımlarını artırdığı tespit edilmiştir.

Öğrencilerin haftalık liderlik tahtasının motivasyon ilişkisine dair cevaplarına bakıldığında, büyük çoğunluğunun (%93,3) liderlik tahtasının onları motive ettiğini belirttiği gözlemlenmiştir. Bu veriler, Bayrak (2023) tarafından yapılan çalışmanın bulgularıyla örtüşmektedir. Bu çalışmada da benzer şekilde oyunlaştırılmış e-öğrenme desteğinin etkileri incelenmiş ve deney grubu üzerinde motivasyon açısından istatistiksel olarak anlamlı bir fark elde edilmiştir. Araştırmada öğrencilerin dersi oyunlaştırarak işlemelerinin derse katılımlarına etkisi büyük çoğunlukla (%73,3) olumlu bulunmuştur. Tunga (2016) ile Sarı ve Altun (2016) tarafından yapılan çalışmalar da benzer sonuçlar verdiği gözlemlenmiştir.

Öğrenciler, zorlandıkları konulara oyunlaştırılmış e-öğrenme destekli öğrenim sürecinin yüksek oranda (%80,0) olumlu katkı sağladığını belirtmişlerdir. Katılımcılar, oyunlaştırmanın dersi daha iyi anlamalarına yardımcı olduğu ve çalışma isteklerini artırdığı görüşündedir. Araştırmaya katılan öğrencilerin büyük çoğunluğu (%86,7) diğer derslerin de oyunlaştırılmış e-öğrenme destekli olmasını tercih ettiklerini ifade etmişlerdir. Bu öğrencilerin yarısı (%50,0) fen dersini tercih etmiştir. Öğrenciler, genel olarak (%93,3) EBA'nın ders bazında oyunlaştırılmasının olumlu olduğunu ve bu şekilde sistemin güncellenmesini önermişlerdir. Aksoy (2022) tarafından EBA'da kullanılan oyunlaştırma bileşenleri üzerine yapılan bir çalışmada, bu sonucun desteklendiği gözlemlenmiştir.

Sonuç

Programlamaya yönelik öz yeterlik algılarını geliştirmede oyunlaştırmanın, olumlu bir etkiye sahip etkili bir araç olduğu görülmüştür. Doğası gereği soyut kavramlar içeren programlama öğretiminde karşılaşılan öz yeterlik sorunlarına, oyunlaştırma yoluyla etkili çözümler sunulabileceği düşünülmektedir. Bununla birlikte, oyunlaştırmanın tutum ve derse katılım gibi etkilerinin öğrenme ortamı, konu ve bilişsel düzey gibi faktörlere bağlı olarak tutarsızlık gösterebildiği ve öğrenci ilgisine göre değişken sonuçlar üretebildiği literatürde vurgulanmıştır.

Sonuç olarak, öğretim sürecinde oyunlaştırma kurgusunun, öğretmen, öğrenci, eğitim ortamı ve dersin özellikleri dikkate alınarak titizlikle tasarlanmasının ve gerektiğinde süreç içerisinde çeşitlendirilmesinin önemi ortaya konulmuştur. Ayrıca, oyunlaştırma uygulamalarında değerlendirme süreçlerinin anlık yapılması gerektiği önerilmektedir. Türk Milli Eğitim sisteminde temel bir e-öğrenme aracı olarak değerlendirilen EBA (Eğitim Bilişim Ağı), kullanım kolaylığı, erişilebilirlik, hız ve farklı eğitim materyallerini barındırması gibi özellikleriyle öne çıkmaktadır. EBA'nın, öğretmenlere ders bazında oyunlaştırma imkânı sunması ve anlık değerlendirme yapılabilmesi, rozet ve madalya gibi oyunlaştırma bileşenlerini eş zamanlı olarak belirleyerek öğrencilerin ödev yapma alışkanlıklarını geliştirebileceği, ders takibini kolaylaştırabileceği ve motivasyonlarını artırabileceği ifade edilmektedir.

Teşekkür ve Bilgilendirme / Acknowledgements

Bu makale ikinci yazarın danışmanlığında yapılan birinci yazarın lisansüstü çalışmasının özetidir. / This article is a summary of the first author's postgraduate study conducted under the supervision of the second author..

Yayın Etiği Bildirimi / Research Ethics

Yazarlar araştırmanın etik dışı bir sorunu olmadığını, araştırma ve yayın etiği konusunu gözlemlediklerini beyan etmelidir. / The authors should declare that the research does not have an unethical problem and that they observe the topic of research and publication ethics.

Araştırmacıların Katkı Oranı / Contribution Rate of Researchers

Çalışma yazarların iş birliğiyle tasarlanmış ve yürütülmüştür. / The study was designed and conducted with the authors' collaboration.

Çıkar Çatışması / Conflict of Interest

Yazarlar çalışmanın herhangi bir çıkar çatışması olmadığını beyan etmiştir. / The authors declared that the study has no conflict of interest.

Fon Bilgileri / Funding

Yazarlar bu çalışma için herhangi bir fon olmadığını beyan etmiştir. / The authors declared that they had no funds for this study

Etik Kurul Onayı / The Ethical Committee Approval

Bu çalışma için etik kurul onayı, sorumlu yazar tarafından 25 Ocak 2024 tarihinde Bartın Üniversitesi Sosyal Bilimler Etik Kurulu'na yapılan başvuru sonucunda, 2024-SBB-0090 numarası ile alınmıştır. / Ethical approval for this study was obtained with the application submitted by the responsible author to the Bartın University Social Sciences Ethics Committee on January 25, 2024, under the approval number 2024-SBB-0090.

Kaynakça / References

- Aban, A. K. A., & Fontanil L. L. (2015). Maintaining students' involvement in a math lecture using countdown timers. *Asia Pacific Journal of Multidisciplinary Research*, 3(5), 1-9.
- Aguilar, S. J., Holman, C. ve Fishman, B. J. (2015). Game-inspired design: Empirical evidence in support of gameful learning environments. *Games and Culture*, 13(1): 1-27.
- Albayrak, D. (2012). *Sosyal ağların eğitim ve öğretim amaçlı kullanımı (Yayımlanmamış doktora tezi)*. Orta Doğu Teknik Üniversitesi Fen Bilimleri Enstitüsü, Ankara
- Aksoy, N.(2022). *Eğitim bilişim ağında (EBA) kullanılan oyunlaştırma bileşenlerine yönelik öğrenci öğretmen ve veli görüşleri*. Yüksek Lisans Tezi.Necmettin Erbakan Üniversitesi. Eğitim Bilimleri Enstitüsü. Konya.
- Altay, G. ve Kışla, T. (2018). Programlamaya yönelik tutum ölçeği ve psikometrik özellikleri. *Ege Eğitim Dergisi*, 19(2): 559-574.
- Barata, G., Gama, S., Jorge, J. ve Gonçalves, D. (2013). Engaging engineering students with gamification. *In Games and virtual worlds for serious applications (VSGAMES), 2013 5th international conference on*, 1-8.
- Başer, M. (2013). Developing attitude scale toward computer programming. *International Journal of Social Science*, 6(6), 199-215.
- Bayrak, F. (2023). *Fen bilimleri dersi eğitiminin oyunlaştırılmış e-öğrenme ortamlarıyla desteklenmesinin motivasyon ve başarıya etkisi*. Yüksek Lisans Tezi. Hacettepe Üniversitesi Eğitim Bilimleri Enstitüsü, Matematik ve Fen Bilimleri Eğitimi Ana Bilim Dalı, Ankara, 113 s.
- Bicen, H. & Kocakoyun, S. (2018). Perceptions of students for gamification approach: Bir Mobil Uygulamanın Geliştirilmesi, *Annual Congress*, 126-132
- Codish, D. ve Ravid, G. (2014). Personality Based Gamification – Educational Gamification for Extroverts and Introverts. *Proceedings of the 9th Chais Conference for the Study of Innovation and Learning Technologies: Learning in the Technological Era*, Israel, 36-44.
- Çelik, E., Yıldırım, S., Kaban, A. & Yıldırım, G. (2017). Değerlendirme Sürecinde Dijital Rozetlerin Kullanımı ve Öğrenen Görüşleri, *Mehmet Akif Ersoy Üniversitesi Sosyal Bilimler Enstitüsü Dergisi*, 9 (22), 173-193. DOI: 10.20875/makusobed.310614
- Deterding, S., Dixon, D., Khaled, R. ve Nacke, L. (2011). *From Game Design Elements to Gamefulness: Defining "Gamification"*. MindTrek'11, Tampere, Finland, 9-15.
- Domínguez, A., Saenz-de-Navarrete, J., De-Marcos, L., Fernández-Sanz, L., Pagés, C., & Martínez-Herráiz, J. (2013). Gamifying learning experiences: Practical implications and outcomes. *Computers & Education*, 63, 380-392.

- Eray, F. (2022). *Ortaokul 8. sınıf öğrencileri üzerinde yürütülen oyunlaştırma tabanlı etkinliklerin öğrencilerin motivasyon, öz yeterlik ve matematik kaygılarına etkisi* (Doktora Tezi, Aydın Adnan Menderes Üniversitesi, Fen Bilimleri Enstitüsü).
- Ersoy, H., Madran, R. O. ve Gülbahar, Y. (2011). A model proposed for teaching programming languages: Robotic programming, In *XIII. Academic Information Conference, Malatya, Türkiye*, 731-736.
- Erümit, S. F. (2016). *Oyunlaştırma yaklaşımlarının eğitimde kullanımı: Tasarım tabanlı bir araştırma*. Doktora Tezi (yayımlanmamış), Atatürk Üniversitesi Bilgisayar ve Öğretim Teknolojileri Eğitimi Ana Bilim Dalı, Erzurum, 221 s.
- Gürer, M.D. & Tokumacı, S. (2020). Mühendislik fakültesi öğrencilerinin programlamaya yönelik tutumları, *Cumhuriyet International Journal of Education*, 9(4), 1064-1082.
- Hamzah, W. M., Ali, N. H., Saman, M., Mohd, Y., Yusoff, M. H., & Yacob, A. (2015). Influence of Gamification on Students' Motivation in using E-Learning Applications Based on the Motivational Design Model. *International Journal of Emerging Technologies in Learning*, 10(2).
- Hakulinen, L., Auvinen, T. & Korhonen, A. (2015). The effect of achievement badges on students' behavior: An empirical study in a university-level computer science course. *International Journal of Emerging Technologies in Learning* 10(1), 18-29.
- Hakulinen, L., Auvinen, T., & Korhonen, A., (2013). Empirical Study on the Effect of Achievement Badges in TRAKLA2 Online Learning Environment. *Learning and Teaching in Computing and Engineering (LaTiCE)*, (1) 47-54, Macau
- Hanus, M. D., & Fox, J. (2015). Assessing the effects of gamification in the classroom: A longitudinal study on intrinsic motivation, social comparison, satisfaction, effort, and academic performance. *Computers & Education*, 80, 152-161.
- Hebebcı, M. T., & Usta, E. (2018). Eğitim Ortamlarında Dijital Rozet Kullanımına İlişkin Öğretmen Görüşleri. *Turkish Journal of Computer and Mathematics Education*, 9(2), 192-210.
- Ibáñez, M. B., Di-Serio, A., & Delgado-Kloos, C. (2014). Gamification for engaging computer science students in learning activities: A case study. *IEEE Transactions on Learning Technologies*, 7(3), 291-301
- Kasalak, İ. (2017). *Robotik kodlama etkinliklerinin ortaokul öğrencilerinin kodlamaya ilişkin özyeterlik algılarına etkisi ve etkinliklere ilişkin öğrenci yaşantıları*. (Yayımlanmamış Yüksek Lisans Tezi) Hacettepe Üniversitesi Bilgisayar ve Öğretim Teknolojileri Anabilim Dalı, 103 s.
- Kaya, D.A. (2024). *Oyunlaştırma yöntemiyle tasarlanan blok tabanlı robotik ve kodlama eğitiminde ortaokul öğrencilerinin bilgisayarca düşünme becerileri ve kodlamaya ilişkin öz-yeterlik algılarının incelenmesi*. (Yayımlanmamış Yüksek Lisans Tezi) Van Yüzüncü Yıl Üniversitesi.
- Kim, B., Park, H. ve Baek, Y. (2009). Not just fun, but serious strategies: Using metacognitive strategies in game-based learning. *Computers & Education*, 52(4): 800- 810.
- Landers, R. N., & Armstrong, M. B. (2015). Enhancing instructional outcomes with gamification: An empirical test of the Technology-Enhanced Training Effectiveness Model. *Computers in Human Behavior*, 1-9.

- Mazman, S.G. ve Altun, A. (2013). Programlama dersinin böte bölümü öğrencilerinin programlamaya ilişkin öz yeterlik algıları üzerine etkisi. *Journal of Instructional Technologies & Teacher Education JITTE*, 2(3): 24-29.
- Meşe, C. (2016). *Harmanlanmış öğrenme ortamlarında oyunlaştırma bileşenlerinin etkililiği*. (Yayınlanmamış doktora tezi). Anadolu Üniversitesi, Eskişehir
- Meşe C. ve Dursun Ö. (2018). Oyunlaştırma bileşenlerinin duygu, ilgi ve çevrimiçi katılıma etkisi. *Eğitim ve Bilim*, 43: 67-95.
- Moccozet, L., Tardy, C., Opprecht, W. ve Léonard, M. (2013). Gamificationbased assessment of group work. In Interactive Collaborative Learning (ICL), *2013 International Conference*, pp 171-179.
- Muntean, C.I. (2011). Raising engagement in e-learning through gamification. *The 6th International Conference on Virtual Learning ICVL 2012*, Romania, 323-329.
- Ortiz Rojas, M. E., Chiluiza, K., & Valcke, M. (2017). Gamification in computer programming: Effects on learning, engagement, self-efficacy and intrinsic motivation. In 11th European Conference on Game-Based Learning (ECGBL) (507-514).
- Özdemir, M. (2024). İlköğretim matematik öğretmen adaylarının oyunlaştırma ile programlama eğitimine ilişkin görüşlerinin incelenmesi. (Yayınlanmamış Yüksek Lisans Tezi). Kırıkkale Üniversitesi.
- Özgür, H. (2011). *Syracuse modeli ile e-öğrenme ortamı için tasarlanmış bir dersin öğrencilerin başarısına etkisi: "Trakya Üniversitesi Eğitim Fakültesi Örneği"*. Doktora Tezi, Trakya Üniversitesi Fen Bilimleri Enstitüsü, Bilgisayar Mühendisliği Ana Bilim Dalı, Edirne, 298 s.
- Özgür, H., Çuhadar, C. & Akgün, F. (2018). Eğitimde Oyunlaştırma Araştırmalarında Güncel Eğilimler. *Kastamonu Eğitim Dergisi*, 26 (5), 1479-1488. DOI: 10.24106/kefdergi.380982
- Özkara, C. E. ve Yelken, T. (2020). Ortaöğretim öğrencilerine yönelik programlama öz yeterlilik ölçeğinin geliştirilmesi: geçerlilik ve güvenilirlik çalışması. *Eğitim Teknolojisi Kuram ve Uygulama*, 10(2): 345-365.
- Özyurt, Ö. & Özyurt, H. (2015). Bilgisayar Programcılığı Öğrencilerinin Programlamaya Karşı Tutum ve programlama Öz-Yeterliklerinin Belirlenmesine Yönelik Bir Çalışma. *Eğitimde Kuram ve Uygulama*, 11(1), 51-67.
- Polat, Y. (2014). Bir vaka incelemesi: *Oyunlaştırma yöntemi ve İngilizce öğrencilerinin motivasyonu üzerine etkisi*. Yüksek Lisans Tezi (yayınlanmamış), Çığ Üniversitesi, Adana, 79 s.
- Ronimus, M., Kujala J., Tolvanen, A., Lyytinen H., (2014). Children's engagement during digital game-based learning of reading: The effects of time, rewards, and challenge, *Computers & Education*, 71, 237–246.
- Sarı, A. ve Altun T. (2016). Oyunlaştırma yöntemi ile işlenen bilgisayar derslerinin etkililiğine yönelik öğrenci görüşlerinin incelenmesi, 1. *Turkish Journal Of Computer And Mathematics Education*, 7: 553-577.

- Serim, E. (2019). *Oyunlaştırma yöntemiyle tasarlanan kodlama eğitimi ile öğrencilerin hesaplamalı düşünme becerileri ve kodlamaya ilişkin öz-yeterlik algılarının incelenmesi*. (Yayımlanmamış Yüksek Lisans Tezi) Fen Bilimleri Enstitüsü, Bilgisayar ve Öğretim Teknolojileri Eğitimi Anabilim Dalı, Balıkesir Üniversitesi.
- Sever, M. (2014). Derse katılım envanterinin Türk kültürüne uyarlanması. *Eğitim VE Bilim*, 39(176): 171-182.
- Tunga, Y. (2016). *E-öğrenme ortamlarında oyunlaştırma kullanımının öğrenenlerin akademik başarısına ve derse katılım durumuna etkisinin incelenmesi*. (Yayımlanmamış Yüksek Lisans Tezi), Ege Üniversitesi, İzmir, 91.
- Wang, Z., Bergin, C. ve Bergin, D. A. (2014). Measuring engagement in fourthtotwelfth grade classrooms: The classroom engagement inventory. *School Psychology Quarterly*, 29(4): 517-535.
- Yağcı, M. (2016). Bilişim teknolojileri (BT) öğretmen adaylarının ve bilgisayar programcılığı (BP) öğrencilerinin programlamaya karşı tutumlarının programlama öz yeterlik algılarına etkisi. *International Journal of Human Sciences*, 13(1): 1418-1432.
- Yıldırım, İ. (2016). *Oyunlaştırma temelli “öğretim ilke ve yöntemleri” dersi öğretim programının geliştirilmesi, uygulanması ve değerlendirilmesi*. Doktora Tezi (yayımlanmamış), Gaziantep Üniversitesi, Eğitim Bilimleri Enstitüsü, Eğitim Bilimleri Ana Bilişim Dalı, Gaziantep, 133 s.
- Yıldırım, İ. ve Demir, S. (2016). Oyunlaştırma temelli “öğretim ilke ve yöntemleri” dersi öğretim programı hakkında öğrenci görüşleri. *Uluslararası Eğitim Programları ve Öğretim Çalışmaları Dergisi*, 2(6), 85-102
- Yıldız, M. (2018) *Oyunlaştırılmış blok tabanlı algoritmik düşünme etkinliklerinin öğrencilerin programlamaya yönelik tutum, katılım ve becerilerine etkisi*. Yüksek Lisans Tezi. Ankara: Ankara Üniversitesi. Eğitim Bilimleri Enstitüsü.
- Yılmaz, F. G. ve Keser, H. (2016), The impact of reflective thinking activities in e-learning: A critical review of the emprical research. *Computer and Education*, 95: 163-173.



Sınıf Öğretmenlerinin Yapay Zekâ Okuryazarlık Düzeylerinin Çeşitli Değişkenlere Göre İncelenmesi

Şafak KAMAN*  ¹

Anahtar Sözcükler

Sınıf öğretmenleri
Teknoloji
Yapay zekâ

Makale Hakkında

Gönderim Tarihi

28 Ocak 2025

Kabul Tarihi

1 Haziran 2025

Yayın Tarihi

30 Haziran 2025

Makale Türü

Araştırma Makalesi

Öz

Bu araştırma, sınıf öğretmenlerinin yapay zekâ okuryazarlık düzeylerini belirlemeyi amaçlamakta olup nicel araştırma yöntemlerinden ilişkisel tarama modeli kullanılarak gerçekleştirilmiştir. Çalışma grubunu, 2024-2025 eğitim-öğretim yılında Kastamonu ve Kırşehir illerinde Millî Eğitim Bakanlığı'na bağlı devlet okullarında görev yapan toplam 255 sınıf öğretmeni oluşturmaktadır. Veri toplama aracı olarak araştırmacı tarafından geliştirilen kişisel bilgi formu ile "Yapay Zekâ Okuryazarlık Ölçeği" kullanılmıştır. Elde edilen veriler SPSS 22 paket programı aracılığıyla analiz edilmiştir. Araştırma bulguları, sınıf öğretmenlerinin yapay zekâ okuryazarlık düzeylerinin genel olarak orta düzeyin üzerinde olduğunu ve bazı demografik değişkenler açısından anlamlı bir farklılık göstermediğini ortaya koymuştur. Cinsiyet değişkenine ilişkin analizlerde, kadın ve erkek öğretmenler arasında istatistiksel olarak anlamlı bir fark bulunmamıştır. Benzer şekilde, mesleki kıdem değişkeni ile yapay zekâ okuryazarlık düzeyi arasında da anlamlı bir ilişki tespit edilmemiştir. Ayrıca, öğretmenlerin eğitim durumlarına göre yapay zekâ okuryazarlık düzeylerinde anlamlı bir farklılık görülmemiştir. Bu bulgular, sınıf öğretmenlerinin yapay zekâ okuryazarlığı konusunda genel bir farkındalık ve bilgi düzeyine sahip olduklarını, ancak bu düzeyin cinsiyet, kıdem ve eğitim durumu gibi bireysel özelliklerden bağımsız olarak geliştiğini göstermektedir.

Investigation of Classroom Teachers' Artificial Intelligence Literacy Levels According to Various Variables

Keywords

Classroom teachers
Technology
Artificial
intelligence

Article Info

Received

January 28, 2025

Accepted

June 1, 2025

Published

June 30, 2025

Article Type

Research Paper

Abstract

This study aims to determine the artificial intelligence literacy levels of classroom teachers and was conducted using the relational survey model, one of the quantitative research methods. The study group consists of a total of 255 classroom teachers working in public schools affiliated to the Ministry of National Education in Kastamonu and Kırşehir provinces in the 2024-2025 academic year. The personal information form developed by the researcher and the "Artificial Intelligence Literacy Scale" were used as data collection tools. The data obtained were analysed using SPSS 22 package programme. The findings of the study revealed that the artificial intelligence literacy levels of classroom teachers were generally above the middle level and did not show a significant difference in terms of some demographic variables. In the analyses related to gender variable, no statistically significant difference was found between male and female teachers. Similarly, no significant relationship was found between the professional seniority variable and artificial intelligence literacy level. In addition, there was no significant difference in the artificial intelligence literacy levels of teachers according to their educational background. These findings show that classroom teachers have a general level of awareness and knowledge about artificial intelligence literacy, but this level develops independently of individual characteristics such as gender, seniority and educational status.

Atıf/Cite: Kaman, Ş. (2025). Sınıf öğretmenlerinin yapay zekâ okuryazarlık düzeylerinin çeşitli değişkenlere göre incelenmesi. *Bilgi ve İletişim Teknolojileri Dergisi*, 7(1), 63-77. <https://doi.org/10.53694/bited.1628589>

Atıf/Cite: Kaman, Ş. (2025). Investigation of classroom teachers' artificial intelligence literacy levels according to various variables. *Journal of Information and Communication Technologies*, 7(1), 63-77. <https://doi.org/10.53694/bited.1628589>

*Sorumlu Yazar/Corresponding Author safakkaman@gmail.com

¹ Dr. Lecturer, Kastamonu University, Faculty of Education, Kastamonu, Turkey, safakkaman@gmail.com, <https://orcid.org/0000-0002-8378-6963>

Extended Abstract

Introduction

In the 21st century, the rapid evolution of technology has brought 21st-century skills into sharp focus. These essential skills require individuals to adapt to ongoing changes and to critically analyze and effectively utilize information. Additionally, the integration of artificial intelligence (AI) applications in education has become indispensable due to technological progress and globalization, with support from international organizations and national governments (Eker & Halıcı Gürbüz, 2024). Consequently, education systems need to be redesigned to equip individuals with these vital competencies. These 21st-century skills include not only technological proficiency but also critical thinking, problem-solving, collaboration, creativity, and communication.

A review of the literature reveals various studies related to AI applications. These studies cover topics such as the future of AI (Çetin & Akkaş, 2021; İşler & Kılıç, 2021), educators' perspectives (Dülger & Gümüşelli, 2023), different AI applications (Arslan, 2020), comprehensive reviews of prior research (Akdeniz & Özdiñ, 2021), assessments of AI literacy levels among teacher candidates (Erdoğan & Çakır, 2024; Karaca & Çiçek, 2024), and attitudes toward AI (Banaz & Maden, 2024; Sarıkaya & Kavan, 2024). However, it is notable that there is a lack of research focusing specifically on in-service classroom teachers. Classroom teachers, who are the first educators children encounter and play a critical role in teaching foundational skills such as reading and writing, hold an equally important position in education. Therefore, determining their AI literacy levels is crucial. Higher levels of AI literacy among classroom teachers will enhance their ability to effectively utilize AI in educational settings and to guide their students accordingly. Being literate in AI also helps teachers raise awareness early on by educating students about the positive and negative aspects of AI applications.

Against this background, the main objective of this study is to identify the artificial intelligence literacy levels of classroom teachers. In line with this goal, the study sought answers to the following questions:

1. What are the artificial intelligence literacy levels of classroom teachers?
2. Do classroom teachers' AI literacy scores differ significantly by gender?
3. Do classroom teachers' AI literacy scores differ significantly based on years of professional experience?
4. Do classroom teachers' AI literacy scores differ significantly according to their educational qualifications?

Method

This research aims to determine the AI literacy levels of classroom teachers, focusing on describing the current situation. For this purpose, a relational survey design was employed. The study sample comprised 255 classroom teachers working in the provinces of Kırşehir and Kastamonu during the 2023-2024 academic year. To measure AI literacy, data were collected via two instruments: a personal information form and the Artificial Intelligence Literacy Scale. The personal information form included three questions regarding gender, length of service, and educational background. The Artificial Intelligence Literacy Scale, developed by Çelebi et al. (2023), contains 12 items measured on a seven-point Likert scale ranging from strongly disagree (1) to strongly agree (7). The scale demonstrated a Cronbach's alpha reliability coefficient of 0.85. The data gathered were analyzed using SPSS 22 software.

Findings, Discussion, and Conclusion

The study examined classroom teachers' AI literacy levels in relation to several variables and aimed to describe the results. It was found that the average score on the AI Literacy Scale was 5.28 (approximately 75%), indicating that the AI literacy level of classroom teachers is above average. When gender differences were analyzed, no significant difference was found between male and female teachers, suggesting comparable AI literacy levels across genders. Similarly, the length of professional experience had no significant effect on AI literacy. Additionally, the educational background of teachers, whether undergraduate or graduate, did not significantly influence their AI literacy levels.

Previous research on AI applications involving teachers and teacher candidates covers topics such as attitudes (Aksakal, Emre & Özbek, 2024; Karacan & Çiçek, 2024), self-efficacy (Demir & Beyazhançer, 2024), perceptions (Eker & Halıcı Gürbüz, 2024; Seyrek, Yıldız, Emeksiz, Şahin & Türkmen, 2024), awareness (İçöz & İçöz, 2024), and literacy (Banaz & Demirel, 2024; Erdoğan & Çakır, 2024). Nonetheless, there remains a scarcity of studies specifically addressing AI literacy (Ng et al., 2021). This study attempts to situate its findings within this limited body of research.

The discovery that classroom teachers possess a 75% AI literacy level indicates they have a literacy level above the average. Banaz and Demirel's (2024) study on Turkish teacher candidates, which investigated AI literacy across various variables, supports this result by finding high AI literacy among pre-service teachers. Other researchers have also reported elevated AI awareness among teacher candidates (Çam, Çelik, Turan Güntepe & Durukan, 2024; İçöz & İçöz, 2024). Furthermore, Kose et al. (2023) observed that teachers have developed positive attitudes toward AI applications. The growing interest of teachers and teacher candidates in technology today may explain their heightened engagement with AI.

Moreover, this study's analysis showed that gender, professional experience, and educational status do not have a significant impact on classroom teachers' AI literacy. Supporting evidence can be found in a study by Eker and Halıcı Gürbüz (2024), which examined secondary school mathematics teachers' perceptions of their AI competence and found no significant differences based on gender, age, or education. Other studies by Uyak et al. (2023) and Gerlich (2023) also support these findings.

In summary, classroom teachers generally demonstrate a positive level of AI literacy. This can be attributed to their interest in technology and their adaptability to digital advancements. However, since demographic and professional factors do not significantly influence AI literacy, it appears that this competence depends more on individual curiosity, personal interest, and self-directed learning efforts.

Giriş

21. yüzyılda hızla gelişen teknoloji, bireylerin değişime uyum sağlamalarını ve bilgiyi analiz ederek kullanmalarını gerektiren temel beceriler olan 21. yüzyıl becerilerini öne çıkarmıştır. Özellikle bilgi toplumuna dönüşüm sürecinde, eğitim sistemlerinin bu becerileri kazandıracak şekilde yeniden yapılandırılması giderek daha büyük bir önem taşımaktadır. Bu noktada, yapay zekâ teknolojilerinin farklı alanlarda artan etkisi ve kullanım alanının genişlemesiyle birlikte, eğitimde de bir dönüşüm kaçınılmaz hâle gelmiştir. Uluslararası kuruluşlar ve ülkeler tarafından desteklenen yapay zekâ uygulamaları, teknoloji ve küreselleşmenin etkisiyle eğitim alanında bir zorunluluk haline gelmiştir (Eker & Halıcı Gürbüz, 2024). Bu bağlamda, eğitim sistemleri bireylerin bu becerileri edinmelerine olanak sağlayacak şekilde yeniden yapılandırılmalıdır. 21. yüzyıl becerileri teknolojiyi kullanma yetkinliklerini kapsadığı gibi eleştirel düşünme, problem çözme, iş birliği, yaratıcılık ve iletişim gibi becerileri de kapsamaktadır.

Yapay zekanın tanımı tam olarak yapılamasa da (Banaz & Demirel, 2024) alanyazında araştırmacılar tarafından yapılmış tanımlar bulunmaktadır (Akyürek, 2013; Coppin, 2004; Haenlein & Kaplan, 2019; Nabiyev, 2012; Pena, 2021). Yapılan tanımlar günün koşullarına göre seçilmiş ifadeler ile yapılmıştır. Yapay zekâ en sade haliyle insan yaşamını kolaylaştıran ve belirli yönergelerle insanların işlerinde onlara destek olmak amacıyla geliştirilen bir teknolojidir (Gansser & Reich, 2021). Yapay zekâ, doğal mekanizmaların gerçekleştirebildiği mantıklı veya mantıksız herhangi bir bilişsel etkinliği, robotlar ya da yazılımlar aracılığıyla yapay bir sisteme nasıl devredebileceğimizi ve bu süreçte ne kadar başarılı olunabileceğini inceleyen bir bilim dalıdır (Say, 2018). Yapay zekâ sayesinde verimlilik artacak, insanlara yeni fırsatlar sunacak, insan kaynaklı hatalar azalacak ve karmaşık sorunlar çözüme kavuşacaktır (Hartwig, 2025). Yapay zekanın insan hayatından uzak duramayacağı ve insan hayatına etki edeceği kaçınılmaz bir gerçektir.

Yapay zekâ, günümüzde telefon uygulamaları, akıllı ev aletleri ve otonom araçlar gibi pek çok alanda farkında olunmaksızın yaşamın bir parçası hâline gelmiştir. Nesnelerin interneti ile birbirine bağlı cihazların yapay zekâ ile entegre edilmesi, yaşamı kolaylaştıran akıllı sistemlerin ve robot teknolojilerinin gelişimini hızlandırmaktadır. Eğitim, tıp, mühendislik, endüstri ve psikoloji gibi çok sayıda disiplinde yapay zekâ, ihtiyaçlara yönelik etkili çözümler sunarak kritik bir rol üstlenmektedir (Eker & Halıcı Gürbüz, 2024). Tarım alanında yapay zekâ, sulama süreçlerinin yönetimi ve ürün verimliliğinin artırılması gibi uygulamalarla öne çıkarken; sağlık sektöründe hastalıkların tanı ve tedavisinde, finans alanında ise piyasa eğilimlerinin analiz edilmesi ve yatırım stratejilerinin geliştirilmesinde etkin biçimde kullanılmaktadır. Ayrıca, eğlence sektöründe içerik öneri sistemleri, mühendislikte otonom sistemlerin tasarımı, eğitimde bireyselleştirilmiş öğrenme deneyimlerinin desteklenmesi; iletişim ve psikoloji gibi alanlarda ise duygusal süreçlerin analiz edilmesi gibi çok çeşitli uygulama alanlarına sahiptir (Elçiçek, 2024). Özellikle eğitim alanında yapay zekâ teknolojileri her düzeyde yaygınlaşmakta ve tüm paydaşlar tarafından ilgiyle karşılanmaktadır. Yapay zekânın günlük yaşama büyük bir hızla entegre olması, artık kaçınılmaz bir gerçek olarak kabul edilmektedir. Bu teknolojilerin eğitim ortamlarına aktarılması, eğitimin dijital dönüşüm sürecinde önemli bir adım olarak değerlendirilmektedir. Eğitimde yapay zekâ destekli uygulamaların yaygınlaşması, öğretim süreçlerinin daha verimli, bireyselleştirilmiş ve etkileşimli hâle gelmesine katkı sağlamaktadır.

Yapay zekâ, eğitimde bireysel öğrenme süreçlerini destekleme, etkileşimli içerikler oluşturma ve sunma, ayrıca eğitim süreçlerini hızlı bir şekilde değerlendirerek öğretmenlere zaman kazandırma gibi birçok avantaj

sunmaktadır (Banaz & Demirel, 2024). Yapay zekânın eğitimdeki kullanım alanlarının oldukça çeşitli olduğunu belirten Korucu ve Biçer (2022, s. 50), yapay zekâ uygulamalarının eğitimde kullanılabileceği alanları şu şekilde sıralamaktadır: a) Öğrenci başarılarının tahmin edilmesi ve mevcut başarı düzeylerinin belirlenmesi, b) bireyselleştirilmiş ders içeriklerinin önerilmesi ve oluşturulması, c) anlık geri bildirimlerin sağlanması ve d) teknik destek sunulması. Buna ek olarak biyometrik verilerden; yüz, ses, konuşma ve parmak izi gibi verilerin tanınması, doğal dil işleme teknolojileri sayesinde farklı dillerde öğrenme fırsatlarının sunulması ve optik karakter tanıma (OCR) özellikleri aracılığıyla sınav değerlendirmelerinin yapılabilmesi de yapay zekâ teknolojilerinin eğitimdeki kullanım alanları arasında yer almaktadır. Bu doğrultuda, yapay zekâ araçlarının doğru ve etkili biçimde kullanılması, eğitim ortamlarının niteliğini artırmada önemli bir rol oynamaktadır (Banaz & Demirel, 2024). Eğitim ortamında kullanılan teknolojilere her geçen gün yenileri eklenmektedir. Özellikle son yıllarda, eğitimde yaygın bir şekilde kullanılan ve üzerinde sıkça tartışmalar yapılan yapay zekâ teknolojisi, eğitim süreçlerinde hem olumlu hem de olumsuz birçok durumu beraberinde getirmektedir (Kutlucan & Seferoğlu, 2024). Bu bağlamda öğretmenlerin teknoloji okuryazarlığının yanında yapay zekâ okuryazarlık becerileri önem kazanmaktadır. Yapay zekâ okuryazarlık becerisi gelişmiş bir öğretmen sınıfında bulunan öğrencilerine bu teknolojik yeniliği en verimli şekilde öğretecek ve kullanabilecektir.

Alanyazın incelendiğinde yapay zekâ uygulamalarına yönelik farklı çalışmaların olduğu görülmektedir. Bu çalışmalar; yapay zekanın geleceği (Çetin & Akkaş, 2021; İşler & Kılıç, 2021), eğitimcilerin görüşleri (Dülger & Gümüştelli, 2023), yapay zekâ uygulamaları (Arslan, 2020), yapılan çalışmaların incelenmesi (Akdeniz & Özding, 2021), öğretmen adaylarının yapay zekâ okuryazarlık düzeylerinin belirlenmesi (Erdoğan & Çakır, 2024; Karaca & Çiçek, 2024), yapay zekaya yönelik tutumdur (Banaz & Maden, 2024; Sarıkaya & Kavan, 2024). Yapılan bu çalışmalar incelendiğinde sınıf öğretmenleri ile ilgili yapılmış bir çalışmanın olmaması dikkat çekicidir. Oysa ki eğitim öğretim denilince ilk akla gelen sınıf öğretmenleri çocuklar içinde bir o kadar önemlidir. Onlara ilk okuma yazma öğreterek hayatlarında izler bırakan sınıf öğretmenlerinin yapay zekâ okuryazarlık düzeylerinin belirlenmesinin önemli olduğu düşünülmektedir. Sınıf öğretmenlerinin yapay zekâ okuryazarlık düzeyleri ne kadar yüksek olursa eğitim ortamlarında verimli kullanma ve öğrencilerine bu konuda rehberlik etmeleri mümkün olacaktır. Öğrencilerine yapay zekâ uygulamalarının olumlu ve olumsuz yanlarını öğreterek küçük yaşlarda bir farkındalık oluşturmaları açısından yapay zekâ okuryazarı olmak önemlidir. Ayrıca sınıf öğretmenlerinin yapay zekâ okuryazarlık düzeylerinin belirlenmesinde cinsiyet, mesleki kıdem ve mezuniyet durumu gibi bireysel değişkenlerin incelenmesi, mevcut durumun çok yönlü değerlendirilmesine olanak tanımaktadır. Cinsiyet değişkeni, teknolojik bilgiye erişim ve kullanım açısından toplumsal cinsiyet rollerinin etkisini anlamaya yardımcı olarak farklılıkların tespiti, cinsiyet temelli eşitlik odaklı eğitim politikalarının oluşturulmasına katkı sunabilir. Mesleki kıdem, öğretmenlerin teknolojiye yaklaşımı ve öğrenme alışkanlıkları üzerinde etkili bir değişkendir. Daha kıdemli öğretmenlerin teknolojik yeniliklere daha temkinli yaklaşabileceği, genç öğretmenlerin ise bu alanlara daha hızlı adapte olabileceği göz önünde bulundurulduğunda, hizmet süresine göre farklı ihtiyaçlara sahip öğretmen grupları için özel destekleyici eğitimlerin tasarlanması önem arz etmektedir. Mezuniyet durumu da öğretmenlerin teknolojik gelişmeleri takip etme düzeyleriyle ilişkili olabilir. Lisansüstü eğitim almış öğretmenlerin yapay zekâ okuryazarlığına dair daha yüksek bir farkındalığa sahip olmaları beklenebilir.

Tüm bu açıklamalar doğrultusunda, bu araştırmanın temel amacı “sınıf öğretmenlerinin yapay zekâ okuryazarlık düzeylerinin farklı değişkenler açısından incelenmesi” olarak belirlenmiştir. Bu kapsamda, aşağıdaki sorulara yanıt aranmıştır:

1. Sınıf öğretmenlerinin yapay zekâ okuryazarlık düzeyleri nedir?
2. Sınıf öğretmenlerinin yapay zekâ okuryazarlık puanları cinsiyete göre anlamlı farklılık göstermekte midir?
3. Sınıf öğretmenlerinin yapay zekâ okuryazarlık puanları mesleki kıdem yılına göre anlamlı farklılık göstermekte midir?
4. Sınıf öğretmenlerinin yapay zekâ okuryazarlık puanları eğitim durumuna göre anlamlı farklılık göstermekte midir?

Yöntem

Araştırmanın Deseni

Bu çalışma, sınıf öğretmenlerinin yapay zekâ okuryazarlık düzeylerini belirlemeye yönelik olup, mevcut durumu açıklamayı amaçlamaktadır. Bu doğrultuda, araştırmada genel tarama yöntemi kullanılmıştır. İlişkisel tarama yöntemleri, iki veya daha fazla değişken arasındaki ilişkileri incelemeye odaklanan araştırma yaklaşımlarıdır (Karasar, 2020). Araştırmada herhangi bir müdahale veya uygulama yapılmamıştır. Araştırmanın amaçlarına uygun olarak betimsel tarama modeli tercih edilmiştir.

Çalışma Grubu

Sınıf öğretmenlerinin yapay zekâ okuryazarlık düzeylerinin belirlenmeye çalışıldığı bu araştırmanın örneklem grubunu 2024-2025 eğitim-öğretim yılında Kastamonu ve Kırşehir illerinde görev yapan sınıf öğretmenleri arasından 255 sınıf öğretmeni oluşturmuştur. Örneklem grubuna ilişkin betimsel veriler aşağıdaki tabloda verilmiştir:

Tablo 1. Örneklem Grubuna İlişkin Betimsel Veriler

Değişken	Grup	f	%
Cinsiyet	Kadın	136	53
	Erkek	119	47
	Toplam	255	100
Mesleki kıdem	1 ve 10 yıl	28	11
	10 ve 20 yıl	105	41
	20 ve üzeri yıl Toplam	122	48
	Toplam	255	100
Eğitim Durumu	Lisans	197	77
	Lisansüstü	58	23
	Toplam	255	100

Tablo 1 incelendiğinde örneklem grubunu oluşturan öğretmenlerin %53'ü kadın öğretmenler oluştururken %47'sini erkek öğretmenler oluşturmuştur. Öğretmenlerin %11'i 1 ve 10 yıl arasında hizmet kıdemine, %41'i 10 ve 20 arasında hizmet kıdemine ve %48'i 20 ve üzeri hizmet kıdemine sahip öğretmenler oluşturduğu görülmektedir. Ayrıca öğretmenlerin %77'si lisans mezunu %23'ü lisansüstü mezundur. Araştırma toplamda 255 öğretmen ile yürütülmüştür.

Veri Toplama Araçları

Sınıf öğretmenlerinin yapay zekâ okuryazarlık düzeylerini belirlemek için kişisel bilgi formu ve yapay zekâ okuryazarlık ölçeği kullanılmıştır. Kişisel bilgi formunda sınıf öğretmenlerinin; cinsiyet, hizmet süreleri ve eğitim

durumu olmak üzere üç adet soru bulunmaktadır. Çelebi, Yılmaz, Demir ve Karakuş (2023) tarafından geliştirilen “Yapay Zekâ Okuryazarlık Ölçeği” 12 maddeden oluşmaktadır. Kesinlikle katılmıyorum (1) ve kesinlikle katılıyorum (7) aralığında olmak üzere yedili likert şeklindedir. Ölçeğin güvenirliğini belirlemek amacıyla Cronbach’s Alpha iç tutarlılık katsayısı hesaplanmıştır. Dört alt boyut ve toplam 12 maddeden oluşan Yapay Zekâ Okuryazarlığı Ölçeği’ne yönelik yapılan doğrulayıcı faktör analizi (DFA) sonucunda; $\chi^2/df = 1.82$, RMSEA = 0.04, RMR = 0.03, NFI = 0.95, CFI = 0.98, GFI = 0.96 ve AGFI = 0.94 değerleri elde edilmiştir. Bu uyum indeksleri, ölçek modelinin iyi düzeyde uyum gösterdiğini ortaya koymaktadır. Güvenirlik analizinde ise ölçeğin alt boyutlarına ilişkin Cronbach’s Alpha katsayıları sırasıyla 0.72, 0.74, 0.76 ve 0.72 olarak bulunmuştur. Ölçeğin tamamı için hesaplanan Cronbach’s Alpha iç tutarlılık katsayısı ise 0.85’tir. Bu sonuçlar, ölçeğin genel olarak yüksek düzeyde güvenilir olduğunu göstermektedir.

Veri Toplama Süreci

Sınıf öğretmenlerinin yapay zekâ okuryazarlık düzeylerini belirlemeyi amaçlayan bu araştırmanın verileri 2024-2025 eğitim öğretim yılında Milli Eğitim Bakanlığı’na bağlı devlet okullarında görev yapan sınıf öğretmenlerinden toplanmıştır. Veriler araştırmacı tarafından toplanmış ve gönüllülük esasına uygun şekilde yürütülmüştür. Veri toplama süreci öncesinde belirlenmiş olan ilkokullar ziyaret edilerek okul idaresine çalışma hakkında gerekli açıklamalar yapılmıştır. Gönüllü olan sınıf öğretmenlerine ölçek basılı olarak verilmiş ve öğretmenlerin yapay zekâ okuryazarlık düzeylerini değerlendirmeleri istenmiştir.

Veri Analizi

Verilerin normal dağılıma uygunluğunu belirlemek amacıyla Kolmogorov-Smirnov ve çarpıklık (Skewness) ile basıklık (Kurtosis) değerleri incelenmiştir. Yapay Zekâ Okuryazarlığı Ölçeği genelinde, Kolmogorov-Smirnov testi anlamlı bulunmuştur ($p < 0.05$). Skewness = -,010 ve Kurtosis = -,438 değeri olarak tespit edilmiştir. Cinsiyet değişkenine göre incelendiğinde, erkek öğretmenler için Kolmogorov-Smirnov testi $p < 0.05$ düzeyinde anlamlı bulunmuş; Skewness = -,324 ve Kurtosis = 1,195 değeri olarak hesaplanmıştır. Kadın öğretmenlerde de benzer şekilde Kolmogorov-Smirnov testi anlamlıdır ($p < 0.05$) ve Skewness ile Kurtosis değerleri sırasıyla -,324 ve 1,195’tir. Mesleki kıdem değişkenine göre yapılan değerlendirmede; 1-10 yıl deneyime sahip öğretmenlerde Kolmogorov-Smirnov testi $p < 0.05$ düzeyinde anlamlı olup, Skewness değeri 1,325 ve Kurtosis değeri 1,537’dir. 10-20 yıl arası kıdeme sahip öğretmenlerde ise Kolmogorov-Smirnov testi anlamlı değildir ($p > 0.05$); Skewness = -,900 ve Kurtosis 1,756’dır. 20 yıl ve üzeri deneyime sahip öğretmenlerde ise test sonucu anlamlıdır ($p < 0.05$) ve Skewness ile Kurtosis değerleri sırasıyla -,569 ve -,494’tür. Eğitim durumu değişkenine göre lisans mezunu öğretmenlerde Kolmogorov-Smirnov testi $p < 0.05$ düzeyinde anlamlıdır; Skewness = ,376 ve Kurtosis = -,411 değeri olarak bulunmuştur. Lisansüstü mezunlarda ise Kolmogorov-Smirnov testi anlamlı değildir ($p > 0.05$); Skewness = -,299 ve Kurtosis = -,531’dir. George ve Mallery (2010) göre Skewness ve Kurtosis değerlerinin -2 ile +2 aralığında olması normal dağılım için kabul edilebilir sınırlar içinde değerlendirilmektedir. Bu doğrultuda elde edilen analiz sonuçları, sınıf öğretmenlerinin Yapay Zekâ Okuryazarlığı Ölçeğinden aldıkları puanların ve değişkenlere göre dağılımın normal olduğu belirlenmiştir. Verilerin analizinde cinsiyet değişkenleri arasındaki ilişkiyi belirlemek amacıyla bağımsız ölçümlerde t-testi uygulanmıştır. Sınıf öğretmenlerinin ölçekten aldığı puanları mesleki kıdem değişkeni arasındaki ilişkiyi ortaya koymak için ise One Way ANOVA kullanılmıştır. Bunlara ek olarak ise verilerin betimlenmesi kapsamında frekans, yüzde ve ortalama puanlardan yararlanılmıştır.

Bulgular

Sınıf öğretmenlerinin yapay zekâ okuryazarlık düzeylerini belirlemek amacı ile 255 sınıf öğretmenine kişisel bilgi formu ve “Yapay Zekâ Okuryazarlık Ölçeği” uygulanmıştır. Elde edilen bulgular alt maddeler halinde aşağıda sunulmuştur.

“Sınıf öğretmenlerinin yapay zekâ okuryazarlık düzeyleri nasıldır?” alt amacına ilişkin sınıf öğretmenlerine uygulanan ölçekten elde edilen sonuçlar Tablo 2’te verilmiştir:

Tablo 2. Yapay Zekâ Okuryazarlık Ölçeği Maddelerine İlişkin Betimsel Verileri

İfadeler	Ortalama	Σ
Akıllı cihazlar ile akıllı olmayan cihazları birbirinden ayırt edebilirim	6,58	,80
Yapay zekâ teknolojisinin kötü amaçlı kullanılmaması için her zaman dikkatliyimdir.	6,03	1,15
Yapay zekâ uygulamalarını veya ürünlerini kullanırken her zaman etik ilkelere uyarım.	6,00	,94
Yapay zekâ uygulamalarını veya ürünlerini kullanırken gizlilik ve bilgi güvenliği konularına asla dikkat etmem.	5,61	1,65
İş verimliliğimi artırmak için yapay zekâ uygulamalarını veya ürünlerini kullanabilirim.	5,40	1,16
Yapay zekâ tarafından sunulan çeşitli çözümler arasından uygun olanını seçebilirim.	5,29	1,49
Bir yapay zekâ uygulamasını veya ürününü bir süre kullandıktan sonra kapasitesini ve sınırlarını değerlendirebilirim.	5,21	1,36
Belirli bir görev için çeşitli yapay zekâ uygulamaları veya ürünleri arasından en uygun olanını seçebilirim.	5,20	1,53
Kullandığım uygulama ve ürünlerde kullanılan yapay zekâ teknolojisini tanımlayabilirim.	4,70	1,68
Yeni bir yapay zekâ uygulamasını veya ürününü kullanmayı öğrenmek benim için genellikle zordur.	4,65	1,72
Günlük işlerimde bana yardımcı olması için yapay zekâ uygulamalarını veya ürünlerini ustalıkla kullanabilirim.	4,58	1,65
Yapay zekâ teknolojisinin bana nasıl yardımcı olacağını bilmiyorum.	4,12	2,13
Toplam	5,28	,63

Tablo 2 incelendiğinde, sınıf öğretmenlerinin Yapay Zekâ Okuryazarlık Ölçeğinden elde ettikleri puan ortalamasının 5,28 olduğu görülmektedir. Bu değer, ölçeğin maksimum puanı dikkate alındığında yaklaşık %75’lik bir orana karşılık gelmektedir. Bu durum, sınıf öğretmenlerinin yapay zekâ okuryazarlık düzeylerinin orta düzeyin üzerinde olduğunu göstermektedir. Ölçek maddelerine ait ortalamalar ayrı ayrı değerlendirildiğinde, “Akıllı cihazlar ile akıllı olmayan cihazları birbirinden ayırt edebilirim.” maddesinin 6,58 ortalama ile en yüksek değere sahip olduğu belirlenmiştir. Bu bulgu, öğretmenlerin akıllı cihazları tanıma ve diğer teknolojik araçlar içerisinde ayırt edebilme konusunda oldukça yeterli olduklarını göstermektedir. İkinci en yüksek ortalamanın “Yapay zekâ teknolojisinin kötü amaçlı kullanılmaması için her zaman dikkatliyimdir.” maddesinde elde edildiği görülmektedir (6,03). Bu, öğretmenlerin yapay zekâ teknolojilerini kullanırken etik sorumluluklarının bilincinde olduklarını ve olası kötüye kullanımlara karşı duyarlı davrandıklarını ortaya koymaktadır. Benzer şekilde, “Yapay zekâ uygulamalarını veya ürünlerini kullanırken her zaman etik ilkelere uyarım.” maddesine verilen yanıtların ortalaması 6,00 olarak hesaplanmıştır. Bu sonuç, sınıf öğretmenlerinin yapay zekâ teknolojilerini etik ilkeler doğrultusunda kullandıklarını göstermektedir.

Tablo sınıf öğretmenlerin ölçeğe vermiş oldukları maddelerden en az ortalamanın “Yapay zekâ teknolojisinin bana nasıl yardımcı olacağını bilmiyorum.” (4,12), “Günlük işlerimde bana yardımcı olması için yapay zekâ uygulamalarını veya ürünlerini ustalıkla kullanabilirim.” (4,58) ve “Yeni bir yapay zekâ uygulamasını veya ürününü kullanmayı öğrenmek benim için genellikle zordur.” (4,65) olduğu görülmektedir. Bu bulgu, sınıf öğretmenlerinin güncel yapay zekâ uygulamalarını kullanmakta zorlandıklarını ve günlük işlerinde kullanmakta zorlandıklarını ortaya koymaktadır.

“Sınıf öğretmenlerinin yapay zekâ okuryazarlık puanları cinsiyete göre anlamlı farklılık göstermekte midir?” alt amacı kapsamında bağımsız ölçümlerde t-testi yapılmıştır ve elde edilen bulgular Tablo 3’te verilmiştir:

Tablo 3. Sınıf Öğretmenlerinin Yapay Zekâ Okuryazarlık Ölçeği Puanlarının Cinsiyete Göre T-Testi Sonuçları

Cinsiyet	N	$\bar{x}$	S	sd	t	p
Erkek	136	5,34	,62	253	1,598	,11
Kız	119	5,21	,65			

“Sınıf öğretmenlerinin yapay zekâ okuryazarlık puanları mesleki kıdem yılına göre anlamlı farklılık göstermekte midir?” alt amacı kapsamında Tek Faktörlü ANOVA yapılmıştır ve elde edilen bulgular Tablo 4’te verilmiştir:

Tablo 4. Sınıf Öğretmenlerinin Yapay Zekâ Okuryazarlık Düzeyleri Mesleki Kıdem Yılı Puanlarına İlişkin Betimsel İstatistikler

Yıllar	N	$\bar{x}$	SS
1 ve 10 yıl	28	5,36	,47
10 ve 20 yıl	105	5,29	,63
20 ve üzeri yıl	122	5,25	,67
Toplam	255	5,28	,63

Analiz sonuçları, sınıf öğretmenlerinin yapay zekâ okuryazarlık düzeylerinin mesleki kıdem yıllarına göre anlamlı bir farklılık ortaya koymadığını göstermektedir, $F(254) = ,318$, $p.01$. Bir başka ifade ile sınıf öğretmenlerinin mesleki kıdemleri yapay zekâ okuryazarlık düzeyleri üzerinde benzer nitelikte olduğunu göstermektedir.

Tablo 5. Sınıf Öğretmenlerinin Yapay Zekâ Okuryazarlık Düzeyleri Mesleki Kıdem Yılı Puanlarına Göre ANOVA Sonuçları

Varyansın Kaynağı	Kareler Toplamı	sd	Kareler Ortalaması	F	p
Gruplar arası	,262	2	,131	,318	,72
Gruplar içi	103,738	252	,412		
Toplam	104,000	254			

Yapılan tek yönlü ANOVA analizi, sınıf öğretmenlerinin yapay zekâ okuryazarlık düzeyleri açısından mesleki kıdem yıllarına göre anlamlı bir farklılık göstermediğini ortaya koymuştur, $F(2, 252) = 0.318$, $p = 0.72$.

“Sınıf öğretmenlerinin yapay zekâ okuryazarlık puanları eğitim durumuna göre anlamlı farklılık göstermekte midir?” alt amacı kapsamında bağımsız ölçümlerde t-testi yapılmıştır ve elde edilen bulgular Tablo 6’da verilmiştir:

Tablo 6. Sınıf Öğretmenlerinin Yapay Zekâ Okuryazarlık Ölçeği Puanlarının Eğitim Durumuna Göre T-Testi Sonuçları

Eğitim	N	$\bar{x}$	S	sd	t	p
Lisans	197	5,32	,58	253	1,99	,93
Lisansüstü	58	5,13	,78			

Tablo 6’deki verilere göre, sınıf öğretmenlerinin yapay zekâ okuryazarlık düzeyleri ölçeği puanları eğitim durumuna göre anlamlı bir farklılık göstermemiştir, $t(253) = 1,99$, $p<.05$. Bu bulgu, lisans ($\bar{x}=5,32$) mezunu öğretmenlerin lisansüstü ($\bar{x}=5,21$) öğretmenleri arasında yapay zekâ okuryazarlık düzeylerinin anlamlı düzeyde farklılaşmadığını ortaya koymuştur.

Tartışma ve Sonuç

Araştırmada sınıf öğretmenlerinin yapay zekâ okuryazarlık düzeyleri bazı değişkenler açısından incelenmiş ve sonuçlar betimlenmeye çalışılmıştır. Araştırmanın sonucunda yapay zekâ okuryazarlık ölçeği ortalama puanı ($X=5,28$) ile %75'e denk geldiği bulunmuştur. Bu bulgu, sınıf öğretmenlerinin yapay zekâ okuryazarlıklarının orta düzeyin üzerinde olduğunu göstermektedir. Cinsiyet değişkenine ilişkin puanlarda, erkek ve kadın öğretmenler arasında anlamlı bir fark bulunmamıştır. Elde edilen bu bulgu erkek ve kadın öğretmenlerin yapay zekâ okuryazarlık düzeylerinin benzer olduğu şeklinde yorumlanabilir. Araştırmanın bir diğer bulgusunda ise sınıf öğretmenlerinin meslekte geçirmiş oldukları sürelerin yapay zekâ okuryazarlık düzeyinde etkisinin olmadığı görülmüştür. Ayrıca sınıf öğretmenlerinin eğitim durumları da incelenmiş ve lisans-lisansüstü eğitim almalarının yapay zekâ okuryazarlık düzeylerine anlamlı derecede etki etmediği tespit edilen bir diğer bulgudur.

Konuya ilişkin yapılan çalışmalar incelendiğinde, yapay zekâ uygulamalarına yönelik öğretmenler ve öğretmen adayları ile gerçekleşen çalışmalarda; tutum (Aksakal, Emre & Özbek, 2024; Karacan & Çiçek, 2024), öz-yeterlik (Demir & Beyazhançer, 2024), algı (Eker & Halıcı Gürbüz, 2024; Seyrek, Yıldız, Emeksiz, Şahin & Türkmen, 2024), farkındalık (İçöz & İçöz, 2024), okuryazarlık (Banaz & Demirel, 2024; Erdoğan & Çakır, 2024) gibi kavramlar üzerinde durulduğu görülmektedir. Alanyazında yapay zekâ okuryazarlığına yönelik çalışmaların yetersizliği (Ng, Leung, Chu & Qiao, 2021) görülmektedir. Bu bağlamda araştırmanın bulguları ilgili literatür ile tartışılırken bahsedilen çalışmalar ile ele alınmaya çalışılmıştır.

Araştırmada sınıf öğretmenlerinin yapay zekâ okuryazarlık düzeylerinin %75 seviyelerinde olduğu bulunmuştur. Bu bulgu sınıf öğretmenlerinin yapay zekâ okuryazarlık seviyelerinin ortalamanın üzerinde olduğu şeklinde yorumlanabilir. Banaz ve Demirel (2024)'de Türkçe öğretmen adaylarının yapay zekâ okuryazarlıklarının farklı değişkenlere göre inceledikleri çalışmalarında öğretmen adaylarının yapay zekâ okuryazarlık düzeylerinin yüksek olduğunu bularak araştırma bulgusunu desteklemişlerdir. Öğretmen adaylarının yapay zekâ farkındalık düzeylerinin yüksek olduğu da farklı araştırmacılar tarafından bulunmuştur (Çam, Çelik, Turan Güntepe & Durukan, 2024; İçöz & İçöz, 2024). Dahası farklı bir araştırmada öğretmenlerin yapay zekâ uygulamalarına yönelik olumlu tutum geliştirdikleri de tespit edilmiştir (Köse, Radıf, Uyar, Baysal & Demirci, 2023). Araştırma sonuçlarının olumlu olmasında günümüz dünyasında öğretmen ve öğretmen adaylarının teknoloji ile ilgili olmaları yapay zekâ uygulamalarına karşı da ilginin yüksek olmasını sağlamış olabilir.

Araştırmanın diğer bulgularında ise sınıf öğretmenlerinin yapay zekâ okuryazarlık düzeyleri cinsiyet, mesleki kıdem ve eğitim değişkenlerine göre incelenmiş ve elde edilen bulgular bu değişkenlerin sınıf öğretmenlerinin yapay zekâ okuryazarlık düzeylerine anlamlı düzeyde etki etmediği bulunmuştur. Araştırma bulgularını destekleyen çalışmalar da vardır. Eker ve Halıcı Gürbüz (2024) tarafından yapılmış olan ortaokul matematik öğretmenlerinin yapay zekâ kullanımına yönelik yeterlik algılarını inceledikleri çalışmalarında cinsiyet, yaş ve eğitim durumu yapay zekâ algıları üzerinde anlamlı bir farklılık oluşturmamıştır. Araştırmanın bulgularını destekleyen farklı çalışmalarda vardır (Gerlich, 2023; Uyak ve diğerleri, 2023).

Araştırmada elde edilen bulgular, sınıf öğretmenlerinin yapay zekâ okuryazarlığı açısından genel olarak olumlu bir düzeyde olduklarını göstermektedir. Bu durum, öğretmenlerin teknolojiye olan ilgileri ve dijital dünyadaki gelişmelere adaptasyon becerileriyle ilişkilendirilebilir. Ancak araştırmada cinsiyet, mesleki kıdem ve eğitim durumunun yapay zekâ okuryazarlığı üzerinde belirleyici bir etki oluşturmadığı sonucuna ulaşılmıştır. Bu durum,

yapay zekâ okuryazarlığının bireylerin kişisel ilgi, merak ve bireysel öğrenme çabalarına daha fazla bağlı bir yetkinlik olduğunu düşündürmektedir.

Bu araştırmanın bulgularından hareketle, öğretmenlerin yapay zekâ okuryazarlık düzeylerini daha da geliştirebilmek adına şu önerilerde bulunulabilir:

1. Öğretmenlerin mesleki gelişim süreçlerine yapay zekâ temelli eğitimlerin dâhil edilmesi, bu konudaki farkındalık ve yetkinlik düzeylerini artırabilir.
2. Sınıf öğretmenlerinin yapay zekâ teknolojilerini kullanma becerilerini geliştirmek için uygulamalı çalıştaylar düzenlenebilir.
3. Okul ortamında yapay zekâ araçlarının entegrasyonunu kolaylaştıracak rehberlik hizmetleri sağlanabilir.
4. Farklı branşlardaki öğretmenlerin yapay zekâ okuryazarlık düzeylerini inceleyen daha kapsamlı çalışmalar yapılabilir.

Araştırma sonuçlarının genel eğitim politikalarına ve uygulamalarına katkı sağlayabileceği düşünülmektedir. Bu bağlamda, öğretmenlerin teknoloji ile etkileşimlerinin artırılması ve dijital dönüşüm süreçlerine daha aktif katılımlarının teşvik edilmesi önem taşımaktadır.

Sonuç olarak, bu çalışma, sınıf öğretmenlerinin yapay zekâ okuryazarlığı konusunda ortalamanın üzerinde bir seviyede olduklarını ortaya koymakla birlikte, bu düzeyin daha ileri seviyelere taşınabilmesi için uygulamaya dönük adımların gerekliliğini vurgulamaktadır. Gelecekte yapılacak araştırmalarda, bireysel ve kurumsal düzeydeki farklı faktörlerin yapay zekâ okuryazarlığı üzerindeki etkileri daha ayrıntılı bir şekilde ele alınabilir.

Teşekkür ve Bilgilendirme / Acknowledgements

Bu çalışma “3-5 Ekim 2024 tarihinde 17. Uluslararası Bilgisayar ve Öğretim Teknolojileri Sempozyumu’nda” sözlü bildiri olarak sunulmuştur. / This study was presented as an oral presentation at the 17th International Symposium on Computer and Instructional Technologies on 3-5 October 2024.

Yayın Etiği Bildirimi / Research Ethics

Araştırmada, çalışma grubuna uygulanan veri toplama araçları hakkında bilgilendirme yapılmış ve verilerin kullanım amacı açıklanmıştır. Katılımcılardan kişisel bilgiler talep edilmemiştir. Ayrıca, araştırmacının iletişim bilgileri ve ulaşılabilecek e-posta adresi paylaşılmıştır. / In the study, information was provided about the data collection tools applied to the study group, and the purpose of using the data was explained. Personal information was not requested from the participants. Additionally, the researcher's contact information and accessible email address were shared.

Araştırmacıların Katkı Oranı / Contribution Rate of Researchers

Yazar bu çalışmayı tek başına uygulamış ve raporlaştırmıştır. / The author carried out and reported this study alone.

Çıkar Çatışması / Conflict of Interest

Bu araştırmada herhangi bir çıkar çatışması yoktur. / There is no conflict of interest in this research.

Fon Bilgileri / Funding

Araştırmayı destekleyen herhangi bir kuruluş yoktur. / There is no organisation supporting the research.

Etik Kurul Onayı / The Ethical Committee Approval

Araştırmanın yürütülebilmesi için Kastamonu Üniversitesi Sosyal ve Beşeri Bilimler Bilimsel Araştırmalar ve Yayın Etiği Kurulundan alınan 07.01.2025 tarihli ve 1/33 belge sayılı etik kurul izni doğrultusunda gerçekleştirilmiştir.

Kaynakça / References

- Akdeniz, M., & Özdiñ, F. (2021). Eğitimde yapay zekâ konusunda Türkiye adresli çalışmaların incelenmesi. [Analyzing Turkish studies on artificial intelligence in education] *Van Yüzüncü Yıl Üniversitesi Eğitim Fakültesi Dergisi*, 18(1), 912-932. <https://doi.org/10.33711/yyuefd.938734>
- Aksakal, Ş., Emre, İ., & Özbek, M. (2024). Sınıf öğretmenlerinin yapay zekaya ilişkin tutumlarının belirlenmesi. [Determination of classroom teachers' attitudes towards artificial intelligence] *Eğitimde Yeni Yaklaşımlar Dergisi*, 7(1), 1-13.
- Akyürek, H. A. (2013). *Yapay zekâ teknikleri kullanarak akıllı iş gücü yönetimi [Intelligent workforce management using artificial intelligence techniques]*, (Yüksek lisans tezi/ Master thesis), Mevlana Üniversitesi.
- Arslan, K. (2020). Eğitimde yapay zekâ ve uygulamaları [Artificial intelligence and applications in education]. *Batı Anadolu Eğitim Bilimleri Dergisi*, 11(1), 71-88.
- Banaz, E., & Demirel, O. (2024). Türkçe öğretmen adaylarının yapay zekâ okuryazarlıklarının farklı değişkenlere göre incelenmesi [Investigation of Turkish pre-service teachers' artificial intelligence literacy according to different variables]. *Dokuz Eylül Üniversitesi Buca Eğitim Fakültesi Dergisi*, (60), 1516-1529. <https://doi.org/10.53444/deubefd.1461048>
- Banaz, E., & Maden, S. (2024). Türkçe öğretmen adaylarının yapay zekâ tutumlarının farklı değişkenler açısından incelenmesi [Investigation of Turkish pre-service teachers' artificial intelligence attitudes in terms of different variables]. *Trakya Eğitim Dergisi*, 14(2), 1173-1180. <https://doi.org/10.24315/tred.1430419>
- Coppin, B. (2004) *Artificial Intelligence Illuminated*, Jones and Bartlet.
- Çam, M. B., Çelik, N. C., Turan Güntepe, E., & Durukan, Ü. G. (2021). Öğretmen adaylarının yapay zekâ teknolojileri ile ilgili farkındalıklarının belirlenmesi [Determination of prospective teachers' awareness of artificial intelligence technologies]. *Mustafa Kemal Üniversitesi Sosyal Bilimler Enstitüsü Dergisi*, 18(48), 263- 285.
- Çelebi, C., Yılmaz, F., Demir, U., & Karakuş, F. (2023). Artificial intelligence literacy: An adaptation study. *Instructional Technology and Lifelong Learning*, 4(2), 291-306. <https://doi.org/10.52911/ital.1401740>
- Çetin, M., & Aktaş, A. (2021). Yapay zekâ ve eğitimde gelecek senaryoları [Artificial intelligence and future scenarios in education]. *OPUS International Journal of Society Researches*, 18(Eğitim Bilimleri Özel Sayısı), 4225-4268. <https://doi.org/10.26466/opus.911444>
- Demir, B., & Beyazhançer, R. (2024). İlköğretim matematik öğretmeni adaylarının yapay zekâ öz-yeterliklerinin bazı değişkenler açısından incelenmesi [Investigation of artificial intelligence self-efficacy of prospective elementary mathematics teachers in terms of some variables]. *International Journal of Social and Humanities Sciences Research (JSHSR)*, 11(113), 2393-2398. <https://doi.org/10.5281/zenodo.14279357>
- Dülger, E. D., & Gümüseli, A. İ. (2023). Okul müdürleri ve öğretmenlerin eğitimde yapay zekâ kullanılmasına ilişkin görüşleri [Opinions of school principals and teachers on the use of artificial intelligence in

- education]. *ISPEC International Journal of Social Sciences & Humanities*, 7(1), 133-153. <https://doi.org/10.5281/zenodo.7766578>
- Eker, C. ve Halıcı Gürbüz, S. (2024). Matematik öğretmenlerinin matematik dersinde yapay zekâ kullanımına yönelik yeterlilik algıları [Mathematics teachers' perceptions of competence towards the use of artificial intelligence in mathematics lessons]. *Sosyal, Beşeri ve İdari Bilimler Dergisi*, 7(7): 513-528. <https://doi.org/10.26677/TR1010.2024.1425>
- Elçiçek, M. (2024). Öğrencilerin yapay zekâ okuryazarlığı üzerine bir inceleme [A study on students' artificial intelligence literacy]. *Bilgi ve İletişim Teknolojileri Dergisi*, 6(1), 24-35. <https://doi.org/10.53694/bited.1460106>
- Erdoğan, F., & Çakır, O. (2024). Öğretmen adaylarının yapay zekâ okuryazarlıklarının ve yapay zekâyâ ilişkin algılarının belirlenmesi [Determination of pre-service teachers' artificial intelligence literacy and their perceptions of artificial intelligence]. *Uluslararası Türk Kültür Coğrafyasında Sosyal Bilimler Dergisi*, 9(2), 63-95. <https://doi.org/10.55107/turksosbilder.1594635>
- Gansser, O. A., & Reich, C. S. (2021). A new acceptance model for artificial intelligence with extensions to UTAUT2: An empirical study in three segments of application. *Technology in Society*, 65, 101535. <https://doi.org/10.1016/j.techsoc.2021.101535>
- George, D., & Mallery, M. (2010). *SPSS for Windows Step by Step: A Simple Guide and Reference, 17.0 update* (10a ed.) Boston: Pearson
- Gerlich, M. (2023). Perceptions and acceptance of artificial intelligence: A multi-Dimensional study. *Social Sciences*, 12(9), 502. <https://doi.org/10.3390/socsci12090502>
- Haenlein, M., & Kaplan, A. (2019). A brief history of artificial intelligence: On the past, present, and future of artificial intelligence. *California Management Review*, 61(4), 5-14. <https://doi.org/10.1177/0008125619864925>
- Hartwig, B. (2025). *Benefits of artificial intelligence*. Hackr.io. Retrieved January 18, 2025, from <https://hackr.io/blog/benefits-of-artificial-intelligence>
- İçöz, S., & İçöz, E. (2024). Türkçe öğretmen adaylarının yapay zekâ uygulamalarına yönelik farkındalık düzeylerinin incelenmesi [Investigation of Turkish pre-service teachers' awareness levels towards artificial intelligence applications]. *Ulusal Eğitim Dergisi*, 4(3), 987-1001. <https://doi.org/10.5281/zenodo.10909458>
- İşler, B., & Kılıç, M. (2021). Eğitimde yapay zekâ kullanımı ve gelişimi [Use and development of artificial intelligence in education]. *Yeni Medya Elektronik Dergisi*, 5(1), 1-11. https://doi.org/10.17932/IAU.EJNM.25480200.2021/ejnm_v5i1001
- Karacan, S. B., & Çiçek, Ş. (2024). İlahiyat fakültesi öğrencilerinin yapay zekâ okuryazarlık düzeyleri ile yapay zekâyâ yönelik tutumları [Artificial intelligence literacy levels of theology faculty students and their attitudes towards artificial intelligence]. *Din Bilimleri Akademik Araştırma Dergisi*, 24(3), 259-292. <https://doi.org/10.33415/daad.1577561>

- Karasar, N. (2020). *Bilimsel araştırma yöntemleri [Scientific research methods]*. Ankara: Nobel Yayınları.
- Korucu, A. T. & Biçer, H. (2022). Eğitimde yapay zekânın rolleri ve eğitsel yapay zekâ uygulamaları [The roles of artificial intelligence in education and educational artificial intelligence applications]. Nabiye, V. & Erümit, A.K. (Ed.), Eğitimde yapay zekâ, kuramdan uygulamaya (38-56) içinde [*Artificial intelligence in education, from theory to practice (38-56) in*]. Pegem Akademi.
- Köse, B., Radıf, H., Uyar, B., Baysal, İ., & Demirci, N. (2023). Öğretmen görüşlerine göre eğitimde yapay zekanın önemi [The importance of artificial intelligence in education according to teachers' views]. *Journal of Social, Humanities and Administrative Sciences*, 9(71):4203-4209. <https://doi.org/10.29228/JOSHAS.74125>
- Kutlucan, E. & Seferoğlu S. S. (2024). Eğitimde yapay zekâ kullanımı: ChatGPT'nin KEFE ve PEST analizi [Use of artificial intelligence in education: KEFE and PEST analysis of ChatGPT]. *TEBD*, 22(2), 1059-1083. <https://doi.org/10.37217/tebd.1368821>
- Nabiye, V. V. (2012). *Yapay zekâ: İnsan-bilgisayar etkileşimi [Artificial intelligence: Human-computer interaction]*. Seçkin Yayıncılık.
- Ng, D. T. K., Leung, J. K. L., Chu, K. W. S., & Qiao, M. S. (2021). AI Literacy: Definition, teaching, evaluation and ethical issues. *Wiley Online Library*, 58(1), 504–509. <https://doi.org/10.1002/pr2.487>
- Pena, A. (2021). Why introducing kids to machine learning? Erişim adresi: <https://medium.com/code-explorers-worldwide/how-to-introduce-kids-to-machine-learning-career-explorations-26d46f6feb12>
- Sarıkaya, B., & Kavan, N. (2024). Türkçe öğretmen adaylarının yapay zekâyâ yönelik tutumlarının incelenmesi [Investigation of Turkish teacher candidates' attitudes towards artificial intelligence]. *Elektronik Eğitim Bilimleri Dergisi*, 13(26), 191-203. <https://doi.org/10.55605/ejedus.1550010>
- Say, C. (2018). *50 Soruda yapay zekâ [50 questions about artificial intelligence]*. İstanbul: Bilim ve Gelecek Kitaplığı.
- Seyrek, M., Yıldız, S., Emeksiz, H., Şahin, A., & Türkmen, M. T. (2024). Öğretmenlerin eğitimde yapay zekâ kullanımına yönelik algıları [Teachers' perceptions towards the use of artificial intelligence in education]. *International Journal of Social and Humanities Sciences Research (JSHSR)*, 11(106), 845-856. <https://doi.org/10.5281/zenodo.11113077>
- Uyak, S., Güngör Uyak, S., Ürey, D., Keskin, Ö., Aymaz, A., & Aydın, İ. (2023). Okul öncesi eğitim kurumlarında yapay zekâ uygulamaları: yönetici ve öğretmen görüşleri [Artificial intelligence applications in preschool education institutions: administrator and teacher views]. *International Social Mentality and Researcher Thinkers Journal*, 9 (75): 4625- 4636. <https://doi.org/10.29228/smryj.72414>