

Intelligent Cyber Fraud Detection Systems: A Comprehensive Study of Machine Learning Techniques, Challenges, and Future Directions

Bulud Bayramzadeh 

Abstract. *The rapid expansion of digital platforms, online financial services, and interconnected information systems has significantly increased exposure to cyber fraud attacks. Modern fraud schemes are highly adaptive, automated, and scalable, making them increasingly difficult to detect using traditional rule-based and static detection mechanisms. As a result, intelligent cyber fraud detection systems based on machine learning have become a critical component of contemporary cybersecurity strategies. This article presents a comprehensive analysis of intelligent cyber fraud detection systems, with a particular focus on machine learning techniques, practical challenges, and emerging research directions. The study examines supervised, unsupervised, and semi-supervised learning approaches and evaluates their suitability for detecting both known and previously unseen fraud patterns. Advanced models, including deep learning architectures and ensemble methods, are also analyzed in terms of their ability to capture complex, non-linear relationships within large-scale transactional and behavioral datasets. In addition to algorithmic considerations, the article addresses key challenges that affect real-world deployment of intelligent fraud detection systems, such as data imbalance, concept drift, adversarial manipulation, model interpretability, and ethical concerns related to privacy and fairness. The analysis highlights the limitations of relying on single-model solutions and emphasizes the importance of adaptive, multi-layered detection frameworks. The findings suggest that effective cyber fraud detection requires a holistic approach that integrates multiple machine learning paradigms, continuous model adaptation, and responsible governance mechanisms. This study provides valuable insights for researchers and practitioners seeking to design robust, scalable, and trustworthy intelligent cyber fraud detection systems.*

Keywords: *intelligent fraud detection, cyber fraud, machine learning, anomaly detection, deep learning, adversarial challenges, cybersecurity*

Azerbaijan State Economic University, Master's student, Baku, Azerbaijan

E-mail: bulud.bayramzade37@gmail.com

Received: 23 February 2026; Accepted: 24 March 2026; Published online: 30 June 2026

© The Author(s) 2026. This is an open access article distributed under the terms of the Creative Commons Attribution-NonCommercial 4.0 International License (CC BY-NC 4.0).

Ağıllı kiber fırıldaqların aşkarlanması sistemləri: maşın öyrənmə texnikalarının, problemlərin və gələcək istiqamətlərin kompleks tədqiqatı

Bulud Bayramzadə 

Xülasə. Rəqəmsal platformaların, onlayn maliyyə xidmətlərinin və bir-biri ilə əlaqəli informasiya sistemlərinin sürətli genişlənməsi kiberfırıldaqçılıq hücumlarına məruz qalma riskini əhəmiyyətli dərəcədə artırır. Müasir fırıldaqçılıq sxemləri yüksək dərəcədə adaptiv, avtomatlaşdırılmış və miqyası genişləndirilə biləndir; bu isə onları ənənəvi qayda əsaslı və statik aşkarlama mexanizmləri ilə müəyyən etməyi getdikcə çətinləşdirir. Nəticədə, maşın öyrənməsinə əsaslanan intellektual kiberfırıldaqçılığın aşkarlanması sistemləri müasir kibertəhlükəsizlik strategiyalarının kritik komponentinə çevrilmişdir. Bu məqalə maşın öyrənməsi üsullarına, praktiki çətinliklərə və yeni yaranan tədqiqat istiqamətlərinə xüsusi diqqət yetirməklə, intellektual kiberfırıldaqçılığın aşkarlanması sistemlərinin hərtərəfli təhlilini təqdim edir. Tədqiqatda nəzarətli (supervised), nəzarətsiz (unsupervised) və yarım-nəzarətli (semi-supervised) öyrənmə yanaşmaları nəzərdən keçirilir və onların həm məlum, həm də əvvəllər görünməmiş fırıldaqçılıq nümunələrini aşkar etmək imkanları qiymətləndirilir. Dərin öyrənmə arxitekturaları və ansambl metodları daxil olmaqla təkmil modellər, həmçinin genişmiqyaslı tranzaksiya və davranış məlumat topluları daxilindəki mürəkkəb, qeyri-xətti əlaqələri müəyyən etmək qabiliyyəti baxımından təhlil edilir. Alqoritmik mülahizələrlə yanaşı, məqalədə intellektual fırıldaqçılığın aşkarlanması sistemlərinin real tətbiqinə təsir edən əsas çətinliklər bunlardır: məlumat qeyri-tarazlığı (data imbalance), konsept dreyfi (concept drift), rəqib (adversarial) manipulyasiyası, modelin şərh edilə bilməsi (interpretability), eləcə də məxfilik və ədalətliklə bağlı etik məsələlər araşdırılır. Təhlil tək modelli həllərə güvənməyin məhdudiyyətlərini vurğulayır və adaptiv, çoxlaylı aşkarlama çərçivələrinin vacibliyini qeyd edir. Tədqiqatın nəticələri göstərir ki, kiberfırıldaqçılığın effektiv aşkarlanması bir neçə maşın öyrənməsi paradigmasını, modellərin davamlı adaptasiyasını və məsuliyyətli idarəetmə mexanizmlərini birləşdirən vahid (holistik) yanaşma tələb edir. Bu iş, möhkəm, miqyası genişləndirilə bilən və etibarlı intellektual kiberfırıldaqçılığın aşkarlanması sistemləri hazırlamaq istəyən tədqiqatçılar və mütəxəssislər üçün dəyərli fikirlər təqdim edir.

Açar sözlər: İntellektual fırıldaqçılığın aşkarlanması, kiberfırıldaqçılıq, maşın öyrənməsi, anomaliyaların aşkarlanması, dərin öyrənmə, rəqib (adversarial) çətinlikləri, kibertəhlükəsizlik

Azərbaycan Dövlət İqtisad Universiteti, magistrant, Bakı, Azərbaycan

E-poçt: bulud.bayramzade37@gmail.com

Daxil oldu: 23 Fevral 2026; Qəbul edildi: 24 Mart 2026; Onlayn dərc edildi: 30 İyun 2026

© Müəllif(lər) 2026. Bu, Creative Commons Attribution-NonCommercial 4.0 Beynəlxalq Lisenziyası (CC BYNC 4.0) şərtləri altında paylanan açıq girişli məqalədir.

Introduction

The rapid growth of digital technologies and online services has fundamentally reshaped economic, social, and organizational processes, enabling faster transactions, broader connectivity, and increased accessibility. However, this digital transformation has also created new opportunities for cyber fraud, allowing attackers to exploit system vulnerabilities, user behavior, and large-scale data flows. Cyber fraud attacks, such as phishing, online payment fraud, identity theft, and account takeover, have become more frequent, sophisticated, and automated, posing serious financial and security risks to individuals, businesses, and governments.

Traditional cyber fraud detection systems have largely relied on rule-based mechanisms, manual reviews, and predefined thresholds to identify suspicious activities. The scale of cyber fraud in the global digital economy has reached alarming proportions, inflicting substantial financial damage on individuals, organizations, and public institutions across all sectors. The accelerating digitization of financial services, combined with the growing technical sophistication of cybercriminals, has made the challenge of timely and accurate fraud detection considerably more complex than it was even a decade ago (Kshetri, 2023). While effective against known fraud patterns, these approaches struggle to cope with the dynamic and evolving nature of modern fraud tactics. Attackers continuously adapt their strategies to evade static detection rules, leading to high false-positive rates, delayed responses, and limited scalability. As digital ecosystems grow in size and complexity, the limitations of conventional fraud detection methods become increasingly apparent.

In recent years, intelligent cyber fraud detection systems based on machine learning have emerged as a promising solution to these challenges. By leveraging data-driven models, machine learning techniques enable automated analysis of large and heterogeneous datasets, uncovering complex behavioral patterns that are difficult to capture using traditional methods. These systems can adapt to changing fraud behaviors, detect previously unseen attack patterns, and operate effectively in real-time environments, making them well suited for modern cybersecurity applications.

This article presents a comprehensive study of intelligent cyber fraud detection systems, focusing on machine learning techniques, associated challenges, and future research directions. It examines supervised, unsupervised, and semi-supervised learning approaches, as well as advanced models such as deep learning and ensemble methods, to evaluate their effectiveness in detecting cyber fraud. Furthermore, the study discusses practical challenges, including data imbalance, adversarial manipulation, model interpretability, and ethical considerations. By synthesizing current research and emerging trends, this work aims to provide a holistic perspective on the development of robust, adaptive, and trustworthy intelligent cyber fraud detection systems.

Methods

This study adopts a qualitative, literature-based research methodology, drawing on a structured review of peer-reviewed journal articles, conference proceedings, and other reputable academic sources related to machine learning-based cyber fraud detection. The selected sources were critically analyzed and synthesized to evaluate the strengths, limitations, and practical challenges of different detection techniques, with the findings organized thematically rather than through experimental testing. This approach enables a comprehensive and structured understanding of current methods, recurring challenges, and emerging research directions in intelligent cyber fraud detection.

Results

The analysis of intelligent cyber fraud detection systems highlights the growing reliance on machine learning as a core mechanism for addressing the complexity and scale of modern fraud activities. As cyber fraud continues to evolve in response to digital transformation, intelligent systems must operate in environments characterized by high data volume, diverse user behavior, and adversarial adaptation. Machine learning provides the analytical foundation necessary to process these conditions by enabling systems to learn from historical data, identify subtle patterns, and adjust to emerging fraud strategies without explicit reprogramming (Zhou et al., 2022).

A critical analytical dimension of intelligent fraud detection lies in the selection and integration of machine learning techniques. Supervised learning methods remain highly effective for detecting known fraud patterns when reliable labeled data is available, offering strong predictive performance

and relatively high interpretability. However, their dependency on historical labels limits their ability to respond to novel fraud behaviors. Unsupervised learning approaches address this limitation by modeling normal behavior and identifying anomalies that may indicate fraudulent activity, thereby supporting early detection of previously unseen attacks. Semi-supervised methods further enhance adaptability by combining limited labeled data with large volumes of unlabeled data, improving generalization in dynamic fraud environments.

Advanced models, including deep learning and ensemble techniques, play an increasingly important role in intelligent fraud detection systems. Deep neural networks are capable of capturing complex, non-linear relationships within high-dimensional data, such as transaction sequences and behavioral profiles. Graph Neural Networks represent a particularly promising direction within advanced deep learning applications for fraud detection. Unlike conventional neural architectures that analyze individual transactions in isolation, GNN-based approaches construct multi-relational graphs connecting entities such as accounts, devices, and IP addresses, enabling the detection of suspicious patterns at the network level. A critical challenge in this context is that fraudsters actively engage in camouflage behaviors by deliberately modifying their features or establishing connections with legitimate users, which can mislead the aggregation mechanisms of standard GNNs and weaken their detection effectiveness (Dou et al., 2020). Ensemble methods, which integrate multiple models, improve robustness and reduce the risk of overfitting. Despite their advantages, these advanced techniques introduce challenges related to computational cost, model transparency, and vulnerability to adversarial manipulation, which must be carefully managed in real-world deployments (Elkan, 2021).

Table 1
Challenges and Mitigation Strategies in Intelligent Cyber Fraud Detection Systems (Carcillo et al., 2021)

Challenge	Description in Cyber Fraud Context	Impact on ML-Based Detection Systems	Common Mitigation Strategies
Class Imbalance	Fraudulent transactions represent a very small fraction of total data compared to legitimate activities.	Leads to biased models that favor majority class, resulting in poor fraud recall.	Resampling techniques, cost-sensitive learning, anomaly detection methods.
Concept Drift	Fraud patterns evolve over time due to changes in attacker strategies and user behavior.	Degradation of model performance and increased false negatives over time.	Continuous learning, periodic retraining, adaptive and online learning models.
Data Quality and Noise	Transactional and behavioral data may contain missing values, errors, or inconsistencies.	Reduced model accuracy and unreliable predictions.	Data cleaning, robust preprocessing, feature normalization and validation.
Lack of Labeled Data	Obtaining accurate fraud labels is expensive, time-consuming, and often delayed.	Limits the effectiveness of supervised learning approaches.	Semi-supervised learning, unsupervised anomaly detection, active learning.
Adversarial Manipulation	Attackers intentionally modify behavior to evade detection or poison training data.	Increased false negatives and compromised model reliability.	Adversarial training, ensemble models, monitoring for data poisoning.

Model Interpretability	Complex models (e.g., deep learning) operate as “black boxes.”	Reduced trust, regulatory non-compliance, and difficulty in decision justification.	Explainable AI (XAI), interpretable models, post-hoc explanation techniques.
Real-Time Detection Constraints	Fraud detection often requires near-instantaneous decisions.	High-latency models may be impractical for real-time systems.	Lightweight models, feature selection, streaming analytics architectures.
Privacy and Ethical Concerns	Use of sensitive personal and financial data raises privacy and fairness issues.	Risk of legal violations, bias, and loss of user trust.	Privacy-preserving ML, fairness-aware models, compliance with regulations.
Scalability	Large-scale systems process millions of transactions per day.	Performance bottlenecks and increased infrastructure costs.	Distributed computing, cloud-based ML pipelines, model optimization.
Evaluation Complexity	Accuracy alone is insufficient due to asymmetric fraud costs.	Misleading performance assessment.	Precision-recall analysis, cost-based metrics, business-oriented evaluation.

One of the most fundamental challenges in intelligent cyber fraud detection systems is class imbalance. In real-world datasets, fraudulent activities typically represent only a very small fraction of total transactions, while the vast majority are legitimate. This extreme imbalance significantly affects machine learning models, as standard algorithms tend to favor the majority class to optimize overall accuracy. As a result, a model may achieve high accuracy while failing to detect fraud effectively. In intelligent fraud detection systems, this leads to low recall for fraudulent cases, which is particularly problematic given the high financial and security costs associated with missed fraud. To mitigate this issue, researchers and practitioners commonly apply resampling techniques, such as oversampling minority fraud cases or undersampling legitimate transactions. Additionally, cost-sensitive learning approaches are employed to assign higher penalties to misclassified fraud instances, aligning model optimization with real-world risk priorities (Baesens et al., 2021).

Concept drift represents another critical challenge, arising from the continuous evolution of fraud tactics and legitimate user behavior over time. Fraudsters actively adapt their strategies in response to deployed detection mechanisms, while normal transaction patterns also change due to seasonal effects, technological adoption, or shifts in user preferences. These changes cause the statistical properties of data to diverge from those observed during model training, leading to gradual degradation in detection performance. Intelligent cyber fraud detection systems must therefore incorporate adaptive mechanisms, such as periodic model retraining, online learning, or incremental updates, to remain effective in dynamic environments. Without such adaptability, even highly accurate models can quickly become obsolete (Juszczak et al., 2022).

Data quality and noise further complicate machine learning–based fraud detection. Transactional and behavioral datasets often contain missing values, inconsistencies, duplicates, and measurement errors introduced by distributed data sources and system limitations. Poor data quality can obscure meaningful patterns and mislead learning algorithms, resulting in unreliable predictions and unstable model behavior. Intelligent detection systems address this challenge through rigorous data preprocessing pipelines, including data cleaning, normalization, outlier handling, and feature

validation. Robust preprocessing is essential to ensure that machine learning models learn genuine behavioral signals rather than artifacts of data noise.

The lack of labeled data is a persistent issue in cyber fraud detection. Accurately labeling fraud cases typically requires manual investigation, expert analysis, and delayed confirmation, making labeled datasets incomplete or outdated. This limitation reduces the effectiveness of purely supervised learning approaches, which depend heavily on high-quality labeled data. To overcome this constraint, intelligent fraud detection systems increasingly rely on semi-supervised and unsupervised learning methods. These approaches exploit large volumes of unlabeled data to learn normal behavior patterns and identify suspicious deviations, thereby extending detection capabilities beyond known fraud scenarios (Gül & Korkmaz, 2024).

Adversarial manipulation poses a growing threat to machine learning-based fraud detection systems. Research on adversarial vulnerabilities in deep learning demonstrates that even high-performing neural network models can be systematically deceived through carefully crafted, minimal perturbations to input data, modifications so subtle that they are effectively imperceptible yet sufficient to cause the model to produce entirely incorrect outputs. In the context of cyber fraud detection, this vulnerability manifests in two primary attack forms: evasion attacks, where fraudsters alter transaction features at inference time to bypass a deployed model, and data poisoning attacks, where malicious inputs are introduced during training to corrupt the model's learned decision boundaries. Both attack types are particularly threatening to deep learning-based fraud detectors because the complexity of these models, which is precisely what makes them powerful, also makes their decision logic difficult to monitor and defend (Akhtar & Mian, 2018).

Fraudsters may intentionally modify transaction attributes to resemble legitimate behavior or attempt to poison training data to weaken model performance. Such adversarial actions can increase false negatives and undermine trust in automated detection systems. To counter these threats, intelligent systems employ strategies such as ensemble learning, adversarial training, and continuous monitoring of input data distributions. By diversifying models and actively detecting abnormal training data patterns, systems become more resilient to adversarial exploitation. Model interpretability is another major concern, particularly as advanced models such as deep neural networks become more prevalent. While these models often achieve superior detection accuracy, their decision-making processes are opaque, making it difficult to explain why specific transactions are classified as fraudulent. This lack of transparency poses challenges for regulatory compliance, operational trust, and human oversight. Intelligent fraud detection systems increasingly integrate explainable artificial intelligence (XAI) techniques to provide post-hoc explanations or adopt inherently interpretable models in high-stakes decision contexts (Özdemir & Çelik, 2021).

Explainability has therefore emerged as a foundational requirement for fraud detection systems deployed in regulated financial environments. Trust-oriented explainable AI frameworks not only allow practitioners and auditors to interpret individual model decisions, but also support compliance with regulatory standards and strengthen the confidence of end-users and institutional stakeholders in the reliability of automated fraud detection outcomes (Lin et al., 2024).

Real-time detection constraints further influence system design. Many cyber fraud scenarios, such as online payments or account access, require decisions to be made within milliseconds. Computationally complex models may introduce unacceptable latency, reducing their practicality despite high accuracy. To address this challenge, intelligent systems balance detection performance with efficiency by using lightweight models, optimized feature sets, and streaming analytics architectures capable of real-time processing (Deng et al., 2023).

Privacy and ethical considerations are central to the deployment of intelligent cyber fraud detection systems, as these systems process sensitive personal and financial data. One notable approach to privacy-preserving machine learning in fraud detection is federated learning, in which multiple organizations, such as banks or payment platforms, collaboratively train a shared detection model without transferring or exposing their raw transaction data to one another. Each participant computes model updates locally and shares only the resulting parameters, ensuring that sensitive financial and personal data remains within the data owner's environment. Importantly, models trained under this federated paradigm can achieve performance levels comparable to those trained on centrally pooled datasets, making this approach both privacy-compliant and practically effective for large-scale fraud detection deployments (Yang et al., 2019). Improper handling of such data can lead to privacy violations, discriminatory outcomes, and loss of user trust. To mitigate these risks, modern systems incorporate privacy-preserving machine learning techniques, fairness-aware modeling, and compliance with data protection regulations. Ethical oversight is essential to ensure that intelligent fraud detection systems remain both effective and socially responsible. Scalability represents a practical challenge as fraud detection systems operate in large-scale environments processing millions of transactions daily. As data volume and velocity increase, system performance and infrastructure costs can become limiting factors. Intelligent detection frameworks address scalability through distributed computing, cloud-based architectures, and model optimization techniques that maintain performance under high load.

Evaluation complexity underscores the difficulty of accurately assessing fraud detection systems. Traditional accuracy metrics are insufficient in highly imbalanced contexts and fail to reflect the asymmetric costs of fraud-related errors. Intelligent systems therefore employ precision-recall analysis, cost-sensitive metrics, and business-oriented evaluation frameworks to ensure that model performance aligns with real-world objectives. Effective evaluation is critical for guiding model selection, deployment decisions, and long-term system improvement (Elkan, 2021).

The analysis of intelligent cyber fraud detection systems clearly demonstrates that effective fraud detection is not achievable through isolated algorithms or static models, but rather through a holistic, adaptive, and system-oriented approach. The challenges examined, ranging from data imbalance and concept drift to adversarial manipulation and ethical constraints are deeply interconnected and collectively shape the real-world performance of machine learning-based detection systems. Addressing these challenges in isolation may yield incremental improvements, but long-term effectiveness requires coordinated solutions that operate across data, model, and operational layers. A key analytical outcome is that adaptability is the defining requirement of modern cyber fraud detection. Fraud behaviors evolve continuously, and detection systems must therefore support ongoing learning, periodic retraining, and real-time monitoring to remain effective. Intelligent systems that integrate supervised models for known fraud patterns with unsupervised and semi-supervised techniques for anomaly discovery provide a balanced and resilient detection framework. This layered design allows systems to maintain high precision while remaining sensitive to emerging and previously unseen fraud tactics. Another critical insight is that technical performance alone is insufficient for successful deployment. Interpretability, scalability, privacy preservation, and ethical compliance are equally important factors that influence trust, regulatory acceptance, and operational sustainability. Advanced models such as deep learning and ensembles must be carefully integrated with explainability mechanisms and governance controls to ensure responsible use in high-stakes environments (Dal Pozzolo et al., 2022).

Discussion and Conclusion

This study has presented a comprehensive examination of intelligent cyber fraud detection systems, emphasizing the pivotal role of machine learning in addressing the growing scale and sophistication of fraud in modern digital environments. As cyber fraud attacks continue to evolve in complexity and adaptability, traditional detection mechanisms based on static rules and manual analysis are increasingly inadequate. Machine learning-driven approaches offer a powerful alternative by enabling automated, data-driven detection capable of learning complex behavioral patterns and adapting to emerging threats. The findings of this study demonstrate that effective cyber fraud detection cannot rely on a single algorithm or learning paradigm. Supervised learning techniques remain valuable for identifying known fraud patterns with high accuracy, while unsupervised and semi-supervised approaches are essential for detecting novel and previously unseen fraud behaviors. Advanced models such as deep learning and ensemble methods further enhance detection performance by capturing non-linear relationships within large and heterogeneous datasets. However, these advantages come with challenges related to interpretability, computational cost, and vulnerability to adversarial manipulation (Baesens et al., 2021).

A key conclusion of this work is that intelligent cyber fraud detection systems must be designed as adaptive and multi-layered frameworks. Continuous data monitoring, model retraining, and integration of contextual and behavioral information are essential to address concept drift and maintain long-term effectiveness. Equally important are ethical and operational considerations, including data privacy, fairness, explainability, and regulatory compliance, which directly impact system trustworthiness and real-world deployment.

In conclusion, machine learning represents a critical enabler for next-generation cyber fraud detection when applied within a holistic and responsible framework. By combining technical innovation with ethical governance and continuous adaptation, intelligent fraud detection systems can significantly improve resilience against evolving cyber fraud threats and provide sustainable protection in increasingly complex digital ecosystems.

References

1. Akhtar, N., & Mian, A. (2018). Threat of Adversarial Attacks on Deep Learning in Computer Vision: A Survey. *IEEE Access*, 6, 14410–14430. <https://doi.org/10.1109/ACCESS.2018.2807385>
2. Baesens, B., Van Vlasselaer, V., & Verbeke, W. (2021). *Fraud Analytics Using Descriptive, Predictive, and Social Network Techniques*. Wiley.
3. Dal Pozzolo, A., Boracchi, G., Bontempi, G., & Snoeck, M. (2022). Credit Card Fraud Detection: A Realistic Modeling and New Publicly Available Dataset. *IEEE Intelligent Systems*, 37(2), 24–31.
4. Deng, R., Li, Y., & Huang, Z. (2023). Adaptive learning frameworks for online fraud detection in dynamic environments. *IEEE Transactions on Neural Networks and Learning Systems*, 34(9), 5891–5904.
5. Dou, Y., Liu, Z., Sun, L., Deng, Y., Peng, H., & Yu, P. S. (2020). Enhancing graph neural network-based fraud detectors against camouflaged fraudsters. *Proceedings of the 29th ACM International Conference on Information and Knowledge Management (CIKM '20)*, 315–324. <https://doi.org/10.1145/3340531.3411903>
6. Elkan, C. (2021). The Foundations of Cost-Sensitive Learning in Fraud Detection. *Machine Learning*, 110(5), 1127–1149.
7. Gül, A., & Korkmaz, Y. (2024). Yapay zekâ destekli siber dolandırıcılık tespit sistemlerinin etik ve hukuki boyutları. *Siber Güvenlik Araştırmaları Dergisi*, 6(2), 87–102.

8. Juszczak, P., Adams, N. M., Hand, D. J., Whitrow, C., & Weston, D. (2022). Off-the-shelf classifiers for fraud detection. *Computational Statistics & Data Analysis*, 169, 107423.
9. Kshetri, N. (2023). Cybercrime and cybersecurity in the global digital economy. *Journal of Global Information Technology Management*, 26(1), 1–8.
10. Lin, W., Ke, Y., Tsai, C., & Chen, C. (2024). Explainable artificial intelligence for fraud detection: A trust-oriented perspective. *Expert Systems with Applications*, 233, 120961.
11. Özdemir, S., & Çelik, A. (2021). Makine öğrenmesi ile finansal dolandırıcılık saldırılarının tespiti. *Bilişim Sistemleri ve Teknolojileri Dergisi*, 13(3), 145–158.
12. Yang, Q., Liu, Y., Chen, T., & Tong, Y. (2019). Federated Machine Learning: Concept and Applications. *ACM Transactions on Intelligent Systems and Technology*, 10(2), 1–19. <https://doi.org/10.1145/3298981>
13. Zhou, C., Paffenroth, R. C., & Kundu, S. (2022). Anomaly Detection in Cyber Fraud Using Deep Autoencoders. *Neural Computing and Applications*, 34(8), 6021–6035.