

Trustworthy and Adversarial-Resilient AI for Financial Crime Defense: A Unified Framework for Fraud Detection, Anti-Money Laundering, Credit Risk, and Secure Payments

ROHIT JADAV¹, RAJ KUMAR MISHRA², PANKAJ RAI³
^{1,2,3}Cybernetic Research Labs

Abstract- Financial institutions face an increasingly interconnected set of threats spanning transaction fraud, money laundering, adversarial manipulation of scoring models, payment-system exploitation, and network-level intrusions against the infrastructure that hosts these services. Point solutions built for a single threat category leave institutions exposed at the seams between systems, since a fraud model, an anti-money-laundering (AML) pipeline, a credit engine, and a network intrusion detector are rarely designed, evaluated, or governed together. This paper proposes a Unified Financial Crime Defense (UFCD) framework that integrates transformer-based fraud detection with conformal risk control, quantum-classical graph learning for AML entity resolution, smart-contract-verified payment security, calibrated and fair credit risk scoring, adversarial stress testing and certified robustness, homomorphically secured federated fraud analytics, and transformer-based intrusion detection into a single layered architecture with shared governance, explainability, and audit controls. The framework is organized into six layers: data and telemetry ingestion, representation and graph construction, domain-specific detection models, adversarial and privacy hardening, decision intelligence and explainability, and institutional governance. We synthesize findings from the fraud, AML, credit-risk, and network-security literatures to justify each layer, present a comparative assessment of detection paradigms, and discuss deployment challenges around data heterogeneity, adversarial adaptation, regulatory alignment, and cross-institutional trust. The paper contributes a reference architecture that treats robustness, privacy, and explainability as first-class design requirements rather than post-hoc additions, and it outlines validation directions for future empirical work.

Keywords— fraud detection, anti-money laundering, adversarial robustness, federated learning, graph neural networks, credit risk, intrusion detection, financial cybersecurity, explainable ai.

I. INTRODUCTION

Digital financial systems now process fraud, laundering, and credit decisions at a scale and velocity that manual review cannot match, and machine learning has become the default control layer across banking, payments, and capital markets [18]. This shift has produced deep, largely separate literatures: deep learning survey work on transaction fraud [5], graph-based approaches to money laundering detection [46], federated and privacy-preserving fraud analytics across institutions [2], and network intrusion detection for the infrastructure that underpins financial services [1]. Each of these threads has matured independently, but real financial-crime incidents rarely respect these boundaries: a laundering network may exploit a payment rail whose fraud model has been adversarially probed, while the same infrastructure is simultaneously targeted by network-level intrusions against transaction-processing endpoints.

Three gaps motivate this paper. First, fraud, AML, credit-risk, and network-security models are typically trained, evaluated, and governed in isolation, which leaves seams that adaptive adversaries can exploit [37]. Second, robustness and privacy are frequently treated as optional extensions rather than integral design requirements, even though adversarial perturbations measurably degrade discrimination and calibration in production credit and fraud models [4], and centralized data pooling across institutions is often precluded by regulation [48]. Third, explainability is unevenly distributed across the pipeline: interpretable outputs are common in credit scoring but far less consistent in graph-based AML

and intrusion-detection systems [13]. The central question addressed here is: how can fraud detection, AML, credit risk, payment security, and network intrusion detection be organized as a single, governable AI architecture without sacrificing the specialized performance each domain requires?

This paper makes three contributions. It proposes the Unified Financial Crime Defense (UFCD) framework, a six-layer architecture spanning data ingestion through institutional governance. It maps recent domain-specific systems — including transformer-based fraud scoring with conformal risk control, quantum-classical AML graph learning, smart-contract-secured payments, Bayesian credit calibration, adversarial stress testing, homomorphically secured federated training, and transformer-based intrusion detection — onto the layers where each contributes the most value. Finally, it identifies governance, robustness, and interoperability requirements that determine whether such an integrated system can be deployed responsibly.

II. BACKGROUND AND RELATED WORK

A. Machine Learning for Transaction Fraud

Classical fraud detection relied on rule-based systems and shallow classifiers, but comparative studies consistently find that machine learning approaches outperform static rules as fraud patterns evolve [16]. Deep learning surveys document a shift toward hybrid architectures that combine convolutional, recurrent, and attention-based components with ensemble classifiers [3], while explainable-AI additions such as SHAP and LIME are increasingly used to justify fraud decisions to auditors and regulators [44]. Bibliometric analysis of this literature identifies adversarial robustness, multimodal learning, and federated training as the fastest-growing themes, suggesting that isolated accuracy-only fraud models are already considered insufficient by the research community itself [38].

B. Graph Learning for Fraud and Money Laundering

Because fraud and laundering are inherently relational — shared devices, beneficial ownership chains, and coordinated transaction bursts — graph

neural networks (GNNs) have become a dominant paradigm. Foundational work applied graph convolutional networks to Bitcoin transaction graphs for anti-money-laundering forensics [46], and self-supervised node-embedding methods have since improved laundering detection without requiring exhaustive labeled data [20]. Surveys of the field report that GNN-based detectors consistently outperform tabular baselines on relational fraud tasks, particularly when combined with temporal attention over evolving transaction graphs [36], and quantum graph neural networks have recently been explored as a way to represent higher-order entity relationships more efficiently in AML entity resolution [9]. Case studies using leaked financial records such as the FinCEN Files have used heterogeneous graph isomorphism networks to recover laundering typologies that transaction-level models miss [39]. Federated variants of graph learning extend this relational approach across institutions without centralizing raw transaction data [32].

C. Credit Risk, Uncertainty, and Fairness

Credit-risk scoring shares much of its methodological toolkit with fraud detection but carries additional regulatory obligations around fairness and explainability. Antifraud scoring systems built for banking computer security incident response teams (CISIRTs) illustrate how gradient-boosted and calibrated models are already embedded in institutional response workflows [31]. Conformal prediction provides a distribution-free way to attach calibrated uncertainty to individual decisions [37], building on earlier work establishing the theoretical guarantees of conformal inference [45]; this matters directly for credit decisions, where a point score without a confidence interval understates the model's actual reliability under distribution shift.

D. Adversarial Robustness and Privacy

Financial ML models are attractive adversarial targets because a successful evasion translates directly into monetary gain. Recent work applying gradient-based attacks to tabular credit-scoring and fraud models shows measurable degradation in discrimination and calibration, with direct economic and governance consequences [4]. Privacy risk compounds this: membership-inference attacks

demonstrate that trained models can leak information about the individual records used to train them [37], which is a material concern when institutions pool transaction data for collaborative fraud modeling. Federated learning addresses the data-centralization problem by keeping raw data on-premise and exchanging only model updates [22]; foundational federated-averaging work [49] has since been extended with differential privacy and secure aggregation for credit-card fraud detection specifically [50], and with defenses against poisoning attacks in adversarial federated settings [43].

E. Payment Security and Network Intrusion Detection

Beyond the analytical layer, the infrastructure carrying financial traffic is itself a target. Synthetic transaction simulators such as PaySim have been widely used to benchmark fraud detection under controlled mobile-money conditions [21], while GAN-based deepfake and synthetic-identity fraud in online payments represents an emerging attack vector that traditional classifiers were not designed to catch [15]. At the network layer, intrusion detection research has converged on transformer and autoencoder architectures capable of generalizing to previously unseen attack types [12], evaluated against established benchmark datasets including KDD-derived sets [41], UNSW-NB15 [24], Bot-IoT [17], CICIDS2017 [28], and IoT botnet traffic captured in N-BaIoT [23]. Federated deployment of these detectors at the network edge reduces reliance on centralized monitoring while preserving detection latency [26], and explainable variants make alerts auditable for security operations teams rather than opaque black-box flags [13].

Across these five sub-fields, a consistent pattern emerges: each has independently rediscovered the need for robustness, privacy, and explainability, but almost no published work integrates fraud, AML, credit, payment security, and intrusion detection into a single governed architecture. The framework below addresses that integration gap directly.

III. PROPOSED FRAMEWORK: UNIFIED FINANCIAL CRIME DEFENSE (UFCD)

UFCD is organized into six layers, summarized in Table 1. The first layer, data and telemetry ingestion, collects transaction logs, device and session metadata, KYC/beneficial-ownership records, network traffic, and payment-gateway events into a common governed data plane. The second layer, representation and graph construction, builds transaction and entity graphs alongside tabular feature sets, since fraud and AML detection increasingly depend on relational structure rather than record-level features alone [19].

The third layer, domain-specific detection, hosts specialized models mapped to the sub-problems identified in Section II. A hybrid transformer with gated token mixing and conformal risk control provides calibrated, uncertainty-aware fraud scoring with statistically guaranteed error bounds rather than a single confidence number [25]. A hybrid quantum-classical graph learning pipeline performs entity resolution and evidence-subgraph discovery for AML transaction screening, extending classical GCN-based laundering detection with quantum-enhanced relational representations [26]. A smart-contract-verified payment layer, informed by decentralized-identifier (DID) assisted hybrid fraud mining, secures payment execution itself rather than only the post-hoc detection of fraudulent payments [27]. A calibrated credit-intelligence module applies shift-robust, fair risk scoring using Bayesian uncertainty combined with gradient boosting, addressing both the distribution-shift and fairness gaps identified in Section II-C [28].

The fourth layer, adversarial and privacy hardening, is applied horizontally across every model in Layer 3 rather than being domain-specific. Adversarial stress testing and certified robustness checks probe each detector for evasion vulnerabilities before deployment [29], while homomorphically secured federated aggregation with poisoning-resilient training allows multiple institutions to jointly improve fraud and AML models without exposing raw transaction data or accepting a single point of trust for model updates [30]. This layer also

incorporates network-facing protection: a transformer-based intrusion detector with zero-day generalization and explainable attribution monitors the infrastructure hosting the payment and scoring services themselves, closing the gap between application-layer fraud defenses and network-layer security controls [31]. Relational anomaly detection extends this hardening to graph-structured fraud signals directly, using self-supervised graph neural methods with adversarial robustness and explainable reasoning to catch coordinated attack patterns that single-transaction models miss [32].

The fifth layer, decision intelligence and explainability, converts model outputs from all four

preceding layers into ranked alerts, calibrated risk scores, and human-readable justifications, reconciling conflicting signals — for example, a transaction flagged as low-risk by the fraud transformer but high-risk by the AML graph model — before it reaches a human analyst. The sixth layer, institutional governance, implements audit trails, model-risk management, and regulatory reporting aligned with recognized risk-management and AI-management standards [25] and with financial-crime reporting obligations defined by national and international authorities [7].

Table 1. UFCD Framework Layers and Representative Components.

Layer	Function	Representative Components
Data & Telemetry Ingestion	Collects transaction, KYC, and network data	Payment logs, device/session data, network flow records
Representation & Graph Construction	Builds tabular and relational views	Entity graphs, transaction graphs, feature pipelines
Domain-Specific Detection	Scores fraud, AML, credit, and payment risk	Transformer fraud scoring, quantum-classical AML graphs, credit calibration, smart-contract payment security
Adversarial & Privacy Hardening	Stress-tests and protects all models	Adversarial testing, certified robustness, homomorphic federated training, transformer IDS, graph anomaly detection
Decision Intelligence & Explainability	Reconciles and explains alerts	Ranked alerts, calibrated scores, human-readable rationale
Institutional Governance	Provides oversight and reporting	Audit trails, model-risk controls, regulatory reporting

IV. COMPARATIVE DISCUSSION

Table 2 summarizes the trade-offs among the major detection paradigms surveyed in Section II, positioned against the design requirements each UFCD layer imposes. Tabular gradient-boosted and transformer models offer the fastest inference and the most mature explainability tooling, but they treat transactions independently and therefore miss coordinated, multi-account laundering activity that only becomes visible in graph structure [19]. Graph neural approaches recover this relational signal and

have been shown to outperform tabular baselines on relational fraud and laundering benchmarks [36], but they are more expensive to train, harder to explain to non-technical auditors, and more sensitive to graph-construction choices such as edge definition and time-window aggregation [46].

Federated approaches solve the cross-institutional data-sharing problem central to both AML and fraud detection [2], but they introduce new attack surfaces: without secure aggregation and poisoning defenses, a single malicious participant can degrade the shared

model [43], and membership-inference risk persists even in federated settings unless differential privacy or homomorphic protections are layered on top [37]. This is precisely why UFCD treats privacy hardening as a horizontal layer rather than a per-model afterthought: a fraud detector that is individually accurate but trained on a poisoned federated update is a governance failure, not merely a modeling one.

Conformal and Bayesian uncertainty methods add relatively little computational overhead while providing calibrated confidence bounds that materially improve both credit-decision defensibility

and fraud-alert triage [37], which is why UFCD places conformal risk control directly inside the fraud-scoring component rather than as a separate post-processing step. Finally, extending intrusion detection into the same governance layer as fraud and AML models — rather than leaving it with a separate security team and a separate risk register — directly addresses the observation from Section I that adaptive adversaries do not respect organizational boundaries between application-layer fraud and network-layer intrusion [1].

Table 2. Comparative Assessment of Detection Paradigms Within UFCD.

Paradigm	Strength	Limitation
Tabular / Transformer Scoring	Fast inference, mature explainability	Misses multi-account relational patterns
Graph Neural Networks	Captures relational fraud and AML structure	Higher training cost, harder to audit
Federated Learning	Enables cross-institution training without data pooling	Vulnerable to poisoning without hardening
Conformal / Bayesian Uncertainty	Calibrated, defensible risk scores	Requires careful calibration data
Transformer-Based IDS	Generalizes to zero-day network attacks	Needs continuous retraining against drift

V. IMPLEMENTATION CHALLENGES AND LIMITATIONS

Several practical barriers stand between the UFCD architecture and production deployment. Data heterogeneity across institutions — differing KYC formats, transaction schemas, and labeling conventions — complicates the shared representation layer, and synthetic benchmarks such as PaySim, while useful for controlled evaluation, do not fully capture real institutional data distributions [21]. Regulatory alignment is non-trivial: AML reporting obligations [7] and evolving AI-specific risk-management expectations [25] must be satisfied simultaneously, and the governance layer must be able to demonstrate, not merely assert, that automated decisions are traceable to an accountable control.

Adversarial adaptation is an ongoing arms race rather than a one-time hardening exercise; certified robustness guarantees hold only against the threat models they were designed for, and attackers targeting financial ML have both the incentive and the resources to search for new evasion strategies [4]. Finally, this paper's evaluation is analytical rather than empirical: the framework was assessed against completeness, interoperability, privacy-by-design, and auditability criteria drawn from the reviewed literature, but end-to-end validation on real, cross-institutional data — with the access and privacy constraints that implies — remains future work.

VI. CONCLUSION

This paper proposed the Unified Financial Crime Defense (UFCD) framework as a way to close the integration gap between transaction fraud detection,

anti-money-laundering analytics, credit-risk scoring, payment security, and network intrusion detection. By treating adversarial robustness, privacy-preserving federated training, and explainability as horizontal layers applied across every domain-specific model rather than isolated features of individual systems, UFCD reframes financial-crime defense as a governance and integration problem as much as a modeling one. The framework maps recent transformer-based, graph-based, quantum-enhanced, and federated techniques onto a coherent architecture and identifies the regulatory, data-heterogeneity, and adversarial-adaptation challenges that must be addressed before such integration can be validated empirically across real institutions. Future work should prioritize cross-institutional pilot deployments, standardized relational fraud-and-AML benchmarks, and longitudinal evaluation of certified robustness guarantees under live adversarial pressure.

REFERENCES

- [1] S. A. Abdulkareem, C. H. Foh, M. Shojafar, F. Carrez, and K. Moessner, "Network intrusion detection: An IoT and non-IoT-related survey," *IEEE Access*, vol. 12, pp. 147167–147191, 2024. [Crossref](#)
- [2] A. Alhasawi et al., "A federated approach to scalable and trustworthy financial fraud detection," *Security and Privacy*, Wiley, 2025.
- [3] A. N. Angelopoulos and S. Bates, "A gentle introduction to conformal prediction and distribution-free uncertainty quantification," arXiv:2107.07511, 2021.
- [4] S. Baviskar, "Adversarial robustness in financial machine learning: Defenses, economic impact, and governance evidence," arXiv:2512.15780, 2025.
- [5] R. H. Chowdhury, "Advancing fraud detection through deep learning: A comprehensive review," *World J. Adv. Eng. Technol. Sci.*, vol. 12, no. 2, pp. 606–613, 2024.
- [6] R. R. Devi, J. E. Raja, and Y. B. Chin, "Reinforcement learning with graph neural network (RL-GNN) fusion for real-time financial fraud detection: A context-aware community mining approach," *Scientific Reports*, 2025. [Crossref](#)
- [7] Financial Action Task Force, "Money Laundering National Risk Assessment Guidance," FATF, 2024.
- [8] Financial Crimes Enforcement Network, "What We Do," FinCEN, 2023. [Online]. Available: <https://www.fincen.gov/what-we-do>
- [9] N. Innan et al., "Financial fraud detection using quantum graph neural networks," *Quantum Machine Intelligence*, vol. 6, no. 1, p. 7, 2024. [Crossref](#)
- [10] International Organization for Standardization, "ISO/IEC 42001:2023 Artificial intelligence management system," ISO/IEC, 2023.
- [11] S. Islam, G. Raj Gupta, A. Chakraborty, S. Singh, A. Soni, and C. Patle, "Detecting fraudulent transactions for different patterns in financial networks using layer weighted GCN," *Human-Centric Intelligent Systems*, vol. 5, no. 2, pp. 181–195, 2025. [Crossref](#)
- [12] R. Kalakoti, S. Nömm, and H. Bahsi, "Explainable transformer-based intrusion detection in Internet of Medical Things (IoMT) networks," in Proc. ICMLA, 2024, pp. 1164–1169.
- [13] R. Kalakoti, S. Nömm, and H. Bahsi, "Federated learning of explainable AI (FedXAI) for deep learning-based intrusion detection in IoT networks," *Computer Networks*, 2025.
- [14] R. Kalakoti, R. Vaarandi, H. Bahsi, and S. Nömm, "Evaluating explainable AI for deep learning-based network intrusion detection system alert classification," arXiv:2506.07882, 2025.
- [15] Z. Ke, S. Zhou, Y. Zhou, C. H. Chang, and R. Zhang, "Detection of AI deepfake and fraud in online payments using GAN-based models," in Proc. ICAACE, 2025, pp. 1786–1790.
- [16] A. Khanum, K. Chaitra, B. Singh, and C. Gomathi, "Fraud detection in financial transactions: A machine learning approach vs. rule-based systems," in Proc. IITCEE, 2024.
- [17] N. Koroniotis, N. Moustafa, E. Sitnikova, and B. Turnbull, "Towards the development of realistic botnet dataset in the Internet of Things for network forensic analytics: Bot-IoT dataset," *Future Generation Computer Systems*, vol. 100, pp. 779–796, 2019. [Crossref](#)

- [18] G. Kou and Y. Lu, "FinTech: A literature review of emerging financial technologies and applications," *Financial Innovation*, vol. 11, no. 1, p. 1, 2025. [Crossref](#)
- [19] E. Kurshan and H. Shen, "Graph computing for financial crime and fraud detection: Trends, challenges and outlook," *International Journal of Semantic Computing*, vol. 14, no. 4, pp. 565–589, 2020.
- [20] W. W. Lo, G. K. Kulatilleke, M. Sarhan, S. Layeghy, and M. Portmann, "Inspection-L: Self-supervised GNN node embeddings for money laundering detection in Bitcoin," *Applied Intelligence*, vol. 53, pp. 19406–19417, 2023. [Crossref](#)
- [21] E. Lopez-Rojas, A. Elmir, and S. Axelsson, "PaySim: A financial mobile money simulator for fraud detection," in Proc. European Modelling Symposium (EMS), 2016.
- [22] H. B. McMahan, E. Moore, D. Ramage, S. Hampson, and B. A. y Arcas, "Communication-efficient learning of deep networks from decentralized data," in Proc. AISTATS, 2017, pp. 1273–1282.
- [23] Y. Meidan, M. Bohadana, Y. Mathov, Y. Mirsky, A. Shabtai, D. Breitenbacher, and Y. Elovici, "N-BaIoT: Network-based detection of IoT botnet attacks using deep autoencoders," *IEEE Pervasive Computing*, vol. 17, no. 3, pp. 12–22, 2018. [Crossref](#)
- [24] N. Moustafa and J. Slay, "UNSW-NB15: A comprehensive data set for network intrusion detection systems," in Proc. MilCIS, 2015. [Crossref](#)
- [25] S. Nayak and A. R. Bushara, "Uncertainty-aware fraud detection using hybrid transformer with gated token mixing and conformal risk control," *IEEE Access*, vol. 14, pp. 77557–77573, 2026.
- [26] S. Nayak and R. Kumar, "Quantum-enhanced AML: Hybrid quantum–classical graph learning with entity resolution and evidence subgraph discovery for transaction screening," *IEEE Access*, vol. 14, pp. 74837–74850, 2026.
- [27] S. Nayak, "SecurePayChain-AI: Smart-contract-verified secure payments with DID-assisted hybrid fraud mining," in Proc. 5th Asia Conf. Algorithms, Computing and Machine Learning (CACML), Guangzhou, China, 2026, pp. 1–8.
- [28] S. Nayak, "Calibrated credit intelligence: Shift-robust and fair risk scoring with Bayesian uncertainty and gradient boosting," in Proc. IEEE Int. Research Conf. Smart Computing and Systems Engineering (SCSE), Kelaniya, Sri Lanka, 2026, pp. 1–6.
- [29] S. Nayak, "BHF-Guard: Breaking and hardening financial ML with adversarial stress tests and certified robustness checks," in Proc. 14th Int. Symp. Digital Forensics and Security (ISDFS), Boston, MA, USA, 2026, pp. 1–7.
- [30] S. Nayak, "TrustFed-He: Trustworthy federated fraud analytics for banks with homomorphic secure aggregation and poisoning-resilient training," in Proc. 14th Int. Symp. Digital Forensics and Security (ISDFS), Boston, MA, USA, 2026, pp. 1–7.
- [31] S. Nayak, "ThreatFormer-IDS: Robust transformer intrusion detection with zero-day generalization and explainable attribution," in Proc. 4th Cognitive Models and Artificial Intelligence Conf. (AICCONF), Prague, Czech Republic, 2026, pp. 1–8.
- [32] S. Nayak, "FraudGNN: Self-supervised graph neural anomaly detection for real-time financial fraud with adversarial robustness and explainable reasoning," in Proc. IEEE 5th Int. Conf. AI in Cybersecurity (ICAIC), Houston, TX, USA, 2026, pp. 1–7.
- [33] National Institute of Standards and Technology, "Artificial Intelligence Risk Management Framework (AI RMF 1.0)," NIST AI 100-1, 2023.
- [34] J. Pope et al., "Intrusion detection at the IoT edge using federated learning," in *Security and Privacy in Smart Environments*, LNCS vol. 14800, Springer, 2024, pp. 98–119.
- [35] Z. Rouhollahi, "Towards artificial intelligence enabled financial crime detection," arXiv:2105.10866, 2021.
- [36] I. Sharafaldin, A. H. Lashkari, and A. A. Ghorbani, "Toward generating a new intrusion detection dataset and intrusion traffic characterization," in Proc. ICISSP, 2018, pp. 108–116.
- [37] R. Shokri, M. Stronati, C. Song, and V. Shmatikov, "Membership inference attacks against machine learning models," in Proc.

- IEEE Symp. Security and Privacy (S&P), 2017, pp. 3–18. [Crossref](#)
- [38] N. Shone, T. N. Ngoc, V. D. Phai, and Q. Shi, "A deep learning approach to network intrusion detection," *IEEE Trans. Emerging Topics in Computational Intelligence*, vol. 2, no. 1, pp. 41–50, 2018. [Crossref](#)
- [39] M. Srokosz, A. Bobyk, B. Ksiezopolski, and M. Wydra, "Machine-learning-based scoring system for antifraud CISIRTS in banking environment," *Electronics*, vol. 12, no. 10, 2023.
- [40] T. Suzumura et al., "Towards federated graph learning for collaborative financial crimes detection," arXiv:1909.12946, 2019.
- [41] M. Tavallae, E. Bagheri, W. Lu, and A. A. Ghorbani, "A detailed analysis of the KDD CUP 99 data set," in Proc. IEEE CISDA, 2009, pp. 1–6. [Crossref](#)
- [42] UK Finance Limited, "Annual Fraud Report 2025," London, 2025.
- [43] A. Vaswani et al., "Attention is all you need," in Proc. NeurIPS, 2017.
- [44] I. Vorobyev, A. Kireev, and E. Fedulova, "Graph neural networks for financial fraud detection: A survey," *ACM Computing Surveys*, vol. 56, no. 8, pp. 1–35, 2024.
- [45] V. Vovk, A. Gammerman, and G. Shafer, *Algorithmic Learning in a Random World*. Springer, 2005.
- [46] M. Weber et al., "Anti-money laundering in Bitcoin: Experimenting with graph convolutional networks for financial forensics," in Proc. KDD Workshop on Anomaly Detection in Finance, 2019.
- [47] F. Wójcik, "An analysis of novel money laundering data using heterogeneous graph isomorphism networks: FinCEN files case study," *Econometrics*, vol. 28, no. 2, pp. 32–49, 2024.
- [48] Q. Yang, Y. Liu, T. Chen, and Y. Tong, "Federated machine learning: Concept and applications," *ACM Trans. Intelligent Systems and Technology*, vol. 10, no. 2, pp. 1–19, 2019. [Crossref](#)
- [49] W. Yang, Y. Zhang, K. Ye, L. Li, and C. Z. Xu, "FFD: A federated learning based method for credit card fraud detection," in Proc. IEEE Int. Conf. Big Data, 2019.
- [50] Q. Yu, Z. Ke, G. Xiong, Y. Cheng, and X. Guo, "Identifying money laundering risks in digital asset transactions based on AI algorithms," in Proc. EIECC, 2024, pp. 1081–1085.
- [51] H. Zhang, J. Yu, Y. Wang, J. Qi, and J. Cao, "Robust federated learning against poisoning attacks in industrial IoT," *Future Generation Computer Systems*, vol. 152, pp. 563–578, 2024.