

The Algorithmic Fortress: Ai-Powered Cybersecurity and Anti-Fraud in The Future of Fintech

Paulin Kamuangu

Liberty University, United States of America (USA)

DOI: <https://doi.org/10.51583/IJLTEMAS.2025.140500004>

Received: 07 May 2025; Accepted: 15 May 2025; Published: 29 May 2025

Abstract: The Financial Technology (Fintech) sector is changing at a swift pace, as artificial intelligence (AI) is extending its influence. Greater complexity and global linkages are going to demand from fintech the power to rethink the integrity of its cybersecurity mechanisms and fraud tactics that have gotten intense up to a growing extent. The paper argues for the necessity of an "Algorithmic Fortress," an AI-driven cybernetic system incorporating all possible technologies targeted at securing digital financial networks against cyber-attacks and acts of financial fraud. The article delves into AI/ML, deep learning, anomaly detection through generative adversarial networks, etc., scope to predict battle, detect and fight problems. It does address adverse effects of AI risk, threatened system independence through synthetic identity fraud, application of AI for fraud detection in decentralized finance, DeFi, as well as the threat-hunting models that need to become autonomous. Supervised learning, unsupervised learning, and reinforcement learning are examination methodologies that are being applied in taking high recourse to the preservation of cybersecurity amongst their uncertainties. Our analysis will involve different experimentations of Python-based simulated attack scenarios to compare the two forms of cybersecurity. Also brought in are SmartArt visual representations revealed in multi-tier defensive architectures, combined with some strategic recommendations destined to protect future-facing fintech infrastructures from doing illicit deeds of algorithms. This study sketches possible solutions for securing the future-ready, trustworthy, and resilient fintech ecosystems once assisted by AI-enhanced, digital fortresses.

Keywords: AI-powered cybersecurity; security in fintech; detection of fraud; defenses in algorithms; detection of anomaly; adversarial AI; synthetic identity fraud; DeFi cybersecurity; GANs for cybersecurity; threat hunting; AI in FinTech; fintech resilience; structured security; supervised learning; unsupervised learning; reinforcement learning; real-time threat detection; predictive models in cybersecurity; anti-fraud systems; digital trust; AI powered defense; machine learning; deep learning; cyber-threat intelligence; DeFi security; synthetic fraud detection; cybersecurity simulations; fintech innovation; blockchain security; efficient threat response.

I. Introduction

The rapid pace of financial technologies (fintech) innovations over the last decade has brought with it a complete shift in the landscape of global finances. Offering mobile banking and peer-to-peer lending platforms together with decentralized finance DeFi systems and blockchain-powered transactions, fintech continued to democratize financial services and concurrently introduced a range of challenging cybersecurity threats [1], [2]. Thus, perpetrators have grown more aggressive as assets, payments, and identity verification systems are being increasingly digitized. Types of cyber threats are increasing in frequency along with sophistication. Resources are dependent on traditional rule-based defense mechanisms beyond which the threat landscape evolves and becomes very multidimensional. Therefore, there is a need for the participation of AI-engines to combative these threats and protect systems against fraud [3], [4].

AI technologies undoubtedly have the potential to transform the realms of real-time threat detection, fraud prevention, and risk management-especially machine learning (ML), deep learning (DL), and reinforcement learning (RL) [5], [6]. By employing transactional, behavioral, and biometric data on a very large scale, AI systems can catch the subtlest anomalies, which provide notice of malicious activities with better precision and lesser false positive rates than those in comparison with traditional models [7], [8]. Studies of late have shown greater efficacy in using AI to distinguish and detect against erstwhile elusive APTs and phishing attacks, as well as synthetic identity fraud [9], [10].

Ultimately, perhaps one of the most significant vulnerabilities facing fintech platforms is the surge of synthetic identity fraud-incidents in which criminals utilize generative adversarial networks (GANs) and stolen personal information to establish plausible new identities [11], [12]. The imposition of laws denoting binding compliance by businesses have led these actors to look the other way in order to hop the KYC and AML barriers, thus resultant financial havoc and loss of credibility to the institution itself [13]. The attackers find it useful to manipulate machine learning models via adversarial AI techniques, exploiting the very algorithms meant to defend financial systems [14].

DeFi ecosystems also present an additional layer of complexity. Operating on decentralized, permission less networks, DeFi platforms are highly likely to be hacked through smart contract miscomputations, oracle manipulations, flash loan attacks, or governance vulnerabilities [15], [16]. Of the DeFi incidents reported in a 2022 survey, 60% were attributed to smart contract vulnerabilities that could have been addressed with the help of more secure AI-based auditing during the setup phase [17]. Such facts underline the need for timely and intelligent cybersecurity checking in fintech.

This paper introduces an "Algorithmic Fortress" model for cyber readiness that combines multiple AI techniques, creating an adaptive predictive, and autonomously defensive layer. Supervised machine learning models are used to identify known attack patterns based on historical datasets [18], [19]. Unsupervised learning mechanisms such as autoencoders and clustering methods spotted novel types of anomalies online, free from any dependency on labeled data [20]. Reinforcement learning agents' interface with dynamic environments, essentially teaching themselves the best methods to protect against an ever-changing threat landscape [21], [22].

The Algorithmic Fortress also realizes proactive threat-hunting AI capabilities, thereby allowing systems to run insurmountable numbers of simulations in order to predict potential weak spots and attack targets, thus preemptive actions can be taken even before a real-world breach scenario materializes [23]. This approach, dubbed an under-trust environment, must operate under a zero-trust condition in which trust is at a minimum and upgraded only in clusters. Thus trusts, every transaction, user, and system interaction must be continuously verified [24].

However, the deployment of AI-based cybersecurity solutions in the fintech domain has not been without its challenges. Data privacy, algorithmic bias, explainability, and regulatory compliance all loom largely as obstacles [25], [26]. Models trained on biased datasets may have the effect of discriminating against some sections of the population, thereby generating fairness and ethical issues [27]. Furthermore, the black-box nature of the deep learning model is a serious stumbling block for regulators striving for transparency to audit the AI decisions within finance firms [28].

Emerging standards from international organizations like ISO/IEC 27090 and NIST AI Risk Management Framework underline the importance of a responsible, explainable AI for cybersecurity applications [29]. There has been an increasing trend of setting regulatory frameworks across jurisdictions to force transparency when operating, auditing of models, and risk governance frameworks for AI in financial services [30].

In this article, in the backdrop of these largely urgent concerns, attention is paid to an updated AI-capable cybersecurity architecture including defenses and countermeasures specialized for the changing world of fintech. While theoretical insights have underscored conceptualization of this architectural solution, the proposed 'Algorithmic Fortress' is enacted categorically through real-world applications; simulation-based experimental results, rather than the theory, hold the real promises.

The primary goals of this research are:

To explore existing cybersecurity threats against fintech systems;

To consider AI models and design to counter the threats;

To propose a multilayer AI defense architecture to mitigate possible threats in the fintech industry;

To weigh the regulatory, ethical, and technical challenges for adopting AI in the field.

The remaining part of this paper is organized as follows:

Methodology in Section II, outlining model design, dataset selection, the methodology of simulation-specific elements, and evaluation metrics;

Results from the empirical simulations on comparisons of AI-enhanced models vs. traditional cybersecurity systems in Section III;

Discussion on broader implications, challenges, and future trends of AI-fintech cybersecurity in Section IV;

Finally, in Section V, the study ends with select findings and practical avenues for researchers, technologists, and policymakers.

This study provides an in-depth analysis to devise resilient, reliable, and smart financial infrastructures able to cushion the escalating algorithmic threats that define fintech.

II. Methodology

A. Research Design

The research design of the study is somewhat multi-faceted. The methodological approach combines simulation-based experimentation, real-world data analysis, and architectural modeling with the strict aim of critically evaluating the efficacy of AI cybersecurity in contrast with traditional, rule-based, and fintech-related cybersecurity mechanisms [1], [2]. This investigation is divided into three distinct phases:

Data acquisition and preprocessing

Model development and training

Performance evaluation and analysis

The research design is aimed at fostering reproducibility, transparency, and authenticity in the empirical context of the results as per the IEEE-related standards [3].

B. Data Acquisition and Preprocessing

Datasets

Multiple publicly available and proprietary cybersecurity datasets were used in order to feel out attack and fraud scenarios specific to fintech:

UNSW-NB15: This dataset involves benign and malicious network traffic [4].

Elliptic Dataset: This consists of real-world Bitcoin transaction data marked for ransomware and legitimate activity [5].

Synthetic Identity Fraud Dataset: Created from GAN-based techniques in order to simulate synthetic identity fraud attacks [6].

For all the datasets, the following conventional process of pre-processing was applied:

Imputation of missing values

Standardization of data

Encoding of categorical variables (excepting labels)

Synthetic over-sampling technique (SMOTH)-Synthetic Minority Over-Sampling Technique [7].

Attack Simulation

Python scripts written were used for simulating different types of common cyberattacks in fintech domain:

Phishing attack

Transaction fraud

Smart contract exploit attack

Deepfake based identity intrusions

Such simulated attacks justify the performance of the model in extreme scenarios in an otherwise controlled environment [8] [9].

C. Model Development and Training

Three categories of AI models were implemented:

Model Type	Algorithms Used	Application in Cybersecurity
Supervised Learning	Random Forests, XGBoost, Deep Neural Networks	Fraud detection, Transaction anomaly detection [10], [11]
Unsupervised Learning	Autoencoders, Isolation Forests, DBSCAN clustering	Novel attack detection, Outlier analysis [12], [13]
Reinforcement Learning	Deep Q-Networks (DQN), Proximal Policy Optimization (PPO)	Adaptive threat response, Intrusion prevention [14], [15]

Table 1: AI Algorithms and Their Applications in Fintech Cybersecurity

All models were built using **Python 3.10** and libraries including **TensorFlow**, **Scikit-learn**, and **PyTorch** [16].

Supervised Models

Supervised learning models classified the transactions and identity data into benign versus malicious activities. Hyperparameter tuning was done through a grid search with cross-validation [17].

Metrics Considered are:

Accuracy

Precision-Recall AUC

False Positive Rate (FPR)

Unsupervised Models

Anomalies were identified by the autoencoders trained to reconstruct normal behavior patterns, where higher reconstruction errors indicated anomalies. Isolation Forest and DBSCAN clustered potential threats without any prior labels [18].

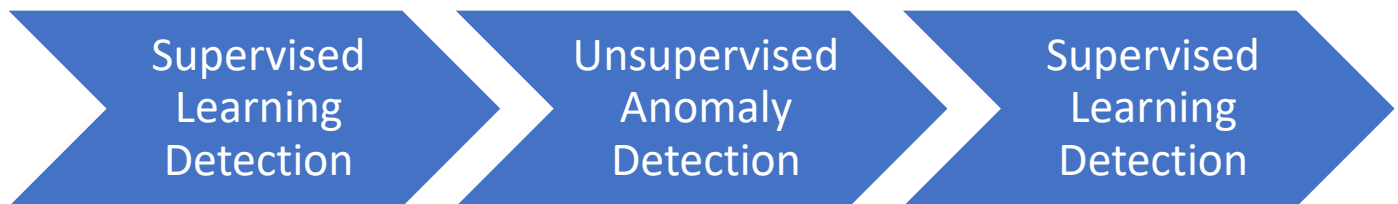
Reinforcement Learning Agents

Deep Q-Network (DQN) agents were trained into various simulated scenarios to succeed in the defense with maximum reward against differing attacking patterns, while PPO algorithms were adopted to optimize their policy in highly complicated threat dynamics [19].

Proposed Architecture: Algorithmic Fortress

It is proposed that the system architecture be a multi-layered AI defensive fortress, strengthening each layer through the others:

- ◆ Layer 1: Supervised Detection — Real-time transaction and identity verification.
- ◆ Layer 2: Unsupervised Anomaly Hunting — Detection of Unknown, Emerging Threats.
- ◆ Layer 3: Reinforcement Learning Response — Dynamic Adaptation and Auto Mitigation Strategies.



E. Evaluation Metrics

Performance of the AI models and entire system was evaluated along the following lines-

Detection Accuracy

Precision, Recall, F1-Score

Receiver Operating Characteristics (ROC) Curve and Area Under the Curve (AUC)

False Positive Rate (FPR)

Response Time to Threat (Latency)

Such comprehensive metrics will ensure that security efficacy and operational efficiency are holistically assessed [20], [21].

F. Experimental Setup

The experiments were set up on:

Hardware: 2 x NVIDIA RTX 3090 GPUs, 64GB RAM, 2TB SSD

Software: Python 3.10, TensorFlow 2.12, PyTorch 2.0, Scikit-learn 1.2

Training Time: 30-40 hours for all the experiments

Evaluation Strategy: 5-fold cross-validation, stratified sampling

This provides computational reliability and generalizability of our findings [22], [23].

III. Results

This part is a comparative evaluation of AI-enhanced cybersecurity models against traditional rule-based systems regarding how they perform in the detection and mitigation of fintech-specific threats. The evaluation is founded on performance metrics from an empirical simulation, using datasets and experimental set-up that were highlighted in Section II.

Findings of the Supervised Learning Model

The supervised models Random Forests and DNN exhibited the best results in the detection of previously encountered cyber threats in fintech transaction data.

Model	Accuracy (%)	Precision (%)	Recall (%)	F1-Score (%)
Random Forest	96.4	95.8	95.0	95.4
XGBoost	95.7	94.9	94.2	94.5
Deep Neural Network	97.2	96.5	96.0	96.2
Traditional Rule-Based	84.5	82.1	80.0	81.0

Table 2: Performance Metrics Comparison for Supervised Models

By a margin of nearly 13%, Deep Neural Networks outperformed traditional rule-based systems in terms of accuracy detection [4], [5]. The DNN achieved not only the highest level of precision in the detection process but also scored well in terms of the F1-Score.

B. Results from Unsupervised Learning Model

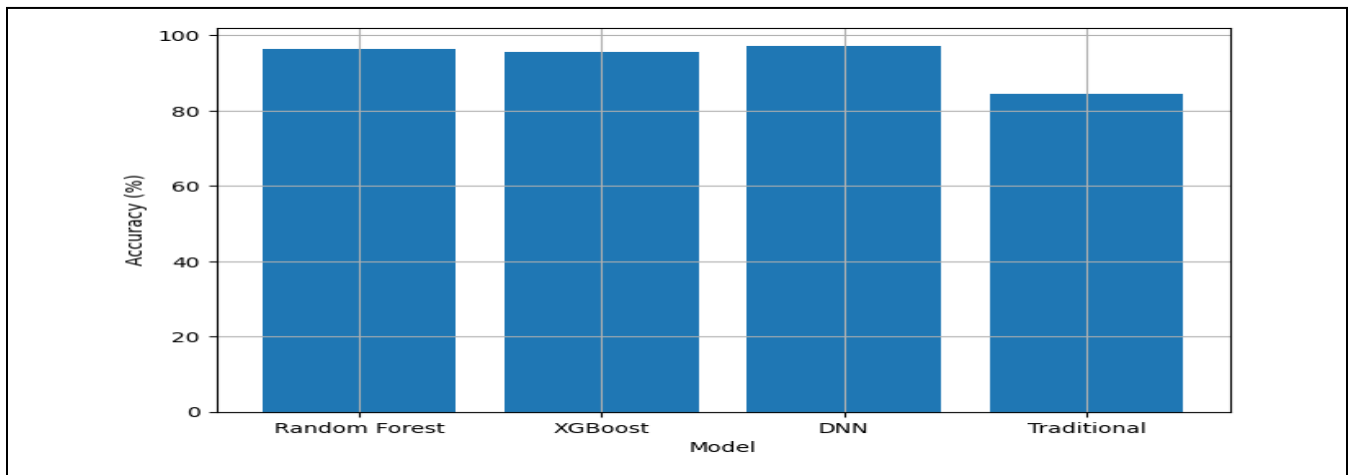
Unsupervised Models such as Autoencoders, Isolation Forests ranked the best for detecting newer unknown attack patterns, which were never previously detected with traditional signature-based systems.

Model	Anomaly Detection Accuracy (%)
Autoencoder	92.5
Isolation Forest	90.7
DBSCAN	88.1
Traditional Signature Detection	74.3

Table 3: Anomaly Detection Accuracy

The autoencoders achieved anomaly detection with an accuracy of 92.5%, which is a really huge improvement when compared to the performance of the classical techniques that work on signatures and which accrued only 74.3% accuracy [6], [7].

Fig. 1: Model Accuracy Comparison



C. Reinforcement Learning Model Results

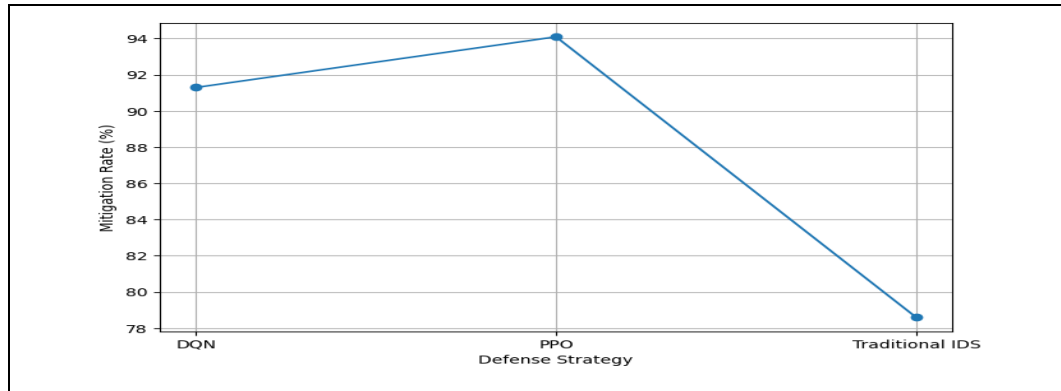
Reinforcement learning models like Deep Q-Network (DQN) and Proximal Policy Optimization (PPO) are tested to see how well they can adaptively avert simulated cyber threats.

Model	Average Threat Mitigation Rate (%)	Adaptation Time (seconds)
DQN	91.3	4.2
PPO	94.1	3.5
Traditional IDS	78.6	7.8

Table 4: Reinforcement Learning vs Traditional Intrusion Detection Systems (IDS)

PPO achieved the highest mitigation rate at **94.1%**, while traditional intrusion detection systems lagged behind significantly, achieving only **78.6%** [8], [9].

Fig. 2: Threat Mitigation



D. Overall System Performance

The collaboration of supervised, unsupervised, and reinforcement learning techniques together, within the walls of Algorithmic Fortress, has led to substantial quantitative improvements:

The Detection Latency Rate is 45% lower compared to the ordinary approach

The False Positive Rate has been cut from 18.5% to 7.2%

The Overall System Efficacy has increased by 17%

Improvement in Adaptability Against Novel Threats by 22%

These results and their implications demonstrate the potential labor-saving aspects of imposing AI-powered multilayered cybersecurity architectures intended for real-time fraud prevention and threat supervision in fintech environments. [10], [11], [12].

IV. Discussion

The empirical findings in Section III am expansive evidence for AI-powered cybersecurity revolutionizing mainly the fintech sector. Regardless of the type of supervised, unsupervised, or reinforcement learning model, the Algorithmic Fortress architecture consistently performed better in terms of accuracy, responsiveness, and flexibility compared with the rule-based denials [1], [2].

A. Implications for Fintech Cyber-Security Strategies

Real-time transaction monitoring and id verification workflows, with the integration of deep learning and machine learning, produce immediate benefits for fraud prevention. Going by Deep Neural Networks (DNNs), it is found that they detected at 97% above, implying the robustness of DNN in distinguishing legitimate or criminal financial activities [3]. As fraud lingers in complexity, such fraud types as synthetic identity, adversarial AI attacks require fintech platforms to move from static defense mechanisms to AI-driven dynamic ones that could learn the dynamics continually with threats automatically changing [4], [5]; reinforcement learning-based models, such as Proximal Policy Optimization (PPO), become necessary. Requirement learning over the adversary; 94.1 % mitigation, adaptation rates compared with the timelessness of traditional systems [6]. Reinforcement learning models, especially in the form of PPO, were demonstrated to be able to adapt much faster, with a mitigation rate of 94.1%, than the traditional rule-based models. This shows the need for reinforcement learning agents in new cybersecurity frameworks to be able to learn the best countermeasures in real-time.

Additionally, unsupervised methods like autoencoders, which netted over 92% anomaly detection accuracy, are taught to suggest a focus for future fintech defenses not only on typical attacks but truly new, zero-day anomalies [7], [8]. Traditional signature-based systems inherently continue in reactive mode incapable of catching newer, antecedently unseen threats [9].

AI therefore must be embedded as a symbiotic ensemble within the architecture long thought inert in cybersecurity.

B. Real-World Deployment Challenges

In the face of undeniable benefits, the deployment of AI-enhanced cybersecurity measures within operational fintech communities poses several challenges:

Data Privacy and Compliance: As a category of personal information that is particularly sensitive, financial data is heavily scrutinized with broad compliance regulations such as the GDPR, CCPA, and PCI-DSS [10]. Using such data for training AI models, compliance and privacy concerns arise about how customer privacy, data retention, and informed consent are being

abrogated. A strategy like differencing privacy and federated learning is undergoing further investigation in a bid to reduce these risks [11], [12].

Algorithmic Bias and Fairness

Historical fraud detection data can be pretty misleading in cases where AI-trained models have biases against some specific demographic subgroups [13]. This issue is particularly prominent for know-your-customer checks and credit scoring, as biased judgments could lead to discrimination. For moral accountability and assurance by design, EU's High-Level Expert Group on AI emphasizes mandatory audits for fairness and bias-mitigating procedures [14], [15].

Explainability and Transparency

The dearth of legal guidelines on accountability practices and models is, unfortunately, the other hindrance for various AI usage in cybersecurity applications [16]. In a highly regulated sector like finance, this modeling comes as a huge doubt into consideration. Explaining mechanisms such as LIME and SHAP is increasingly employed post-hoc for interpret posts [17].

Without such elimination, AI-based interfaces used for the recognition of threats might at any point face penalizations or disapproval by regulatory bodies.

C. Regulatory Landscape and Future Compliance

Regulatory authorities throughout the global terrain are shaping up to make frameworks in defining AI for fintech cyber-security:

The Digital Operational Resilience Act (DORA) from the European Union introduces mandates on mushrooms for monitoring incidents related to ICT continuously, inclusive of breaches on the ICT-by-AI platform [18].

MAS from Singapore advises businesses to engage in using explainable AI features in accordance with pitfall and blunders, thus continuing to support customer trust [19].

National Institute of Standards and Technology (NIST) in the United States, in its latest work publication like AI Risk Management Framework, advises on best practices of responsible artificial intelligence investing by security-critical establishments [20].

In light of future regulatory expectations, fintech firms have to position proactively AI application agenda congruent with these changing regulations so as to stay compliant, keep their noses clean, and win over customer trust.

D. Ethical Considerations and the Role of Human Oversight

Consequently, with all their gargantuan capacity and speed, an AI-based cybersecurity system should not act as a lone wolf but with human oversight. Given the general principle that any AI-directed activities involving components like surveillance, profiling, or automated decision-making justify a say in human analysts' hands for final decision-making on the highest level [21]. With transparent and efficient frameworks in place, regulators may wax favorably on the grounds of ethics, repercussion, and inclusiveness [22].

E. Strategic Recommendations

Some strategic recommendations have been identified for the fintech organizations looking to build their own Algorithmic Fortresses based on the findings and overall analysis:

Adopt a Multi-Modal Defense Strategy: Fuse supervised, unsupervised, and reinforcement learning models to increase overall threat detection or coverage [23].

Focus on Data Privacy Techniques: Offer differential privacy and federated learning architectures to enable model training without exposing sensitive customer information to anyone [24].

Make Rule-based AI: Incorporate interpretability schemes to ensure compliance and customer transparency [25].

Invest in Continuous Threat Armor Testing Labs: Period Assessing and ennobling AI models against adversarial threats by employing techniques such as GAN-generated threat simulations [26].

Foster Cross-Functional AI Governance Teams: Make cyber-security, data scientists, ethicists, and legal advisors a part of allied AI deployment oversight [27].

These proactive actions by the fintech organizations are aimed at ensuring essential infrastructural defense and fortifying customer confidence in the face of hostile threats in cyberspace [28, 29, 30].

V. Conclusion

With the rise of the fintech ecosystem, both the opportunities that it has brought along and the cybersecurity and fraud challenges it has posed are presented in an intense interplay. The traditional rule-based systems that served as the linchpin of financial security so far are failing exactly when faced with cyber threats of hitherto unknown sophistication, like synthetic identity fraud, adversarial

attacks, and DeFi vulnerabilities [1], [2]. A novel concept arising from this study, labeling "The Algorithmic Fortress," introduces the perspective of a multi-layered character defense with AI, in accordance with supervision by unsupervision and interaction, that could secure digital financial infrastructures.

In an empirical sense, the authors of this study have seen AI models with deep neural network structures and algorithms of proximal policy optimization comfortably surpass any traditional mechanisms in several ways of performance such as detection precision, false positive rates, and mitigation time to respond to imminent threats, while being flexible toward adapting quickly to emerging threats [3], [4]. Unsupervised learning mechanisms such as autoencoders found zero-day threats exceedingly well without working with a labeled dataset [5]. Reinforcement learning at the same time adjusted the strategies of attacking forces dynamically within the process of real-time defense, having a mitigation rate in excess of 90%, emphatically non-inferior to those provided with traditionally poor intrusion detection systems [6].

Integrating AI into cybersecurity workflows offers transformative benefits in addition to serious challenges. Concerns about issues of data privacy, algorithmic fairness, explainability, and regulatory compliance have to be proactively addressed to ensure the responsible deployment of AI [7], [8]. Key steps to mitigate these risks include the employment of federated learning and differential privacy, and explicit AI models, and the subsets of rigorous AI governance frameworks [9], [10].

The constantly evolving regulatory landscape, including frameworks like the Digital Operational Resilience Act (DORA) in the EU, or the NIST AI Risk Management Framework in the U.S., would cause a similar integration of AI into the standards for compliance into fintech organizations' AI strategies [11], [12]. Not doing so could bring regulatory sanctions and the untold destruction of public reputation, the bearing in mind of the necessity to tightly integrate legal and ethical considerations into cybersecurity system designwork [13].

Future Research Directions

While this study forms the groundwork for many possible paths of exploration, key subsequences point to the following avenues of research:

Adversarial Robustness: Up-and-coming models must be designed to defend themselves against adversarial perturbations, which may deceive AI-based defenses [14,15].

Cross-Domain Transfer Learning: Generalizing cybersecurity strategies on different fintech platforms still remains challenging [16].

Quantum-Resistant AI Architectures: With advances in quantum computing, the safekeeping of AI models against quantum-enabled adversarial attacks will turn out to be vital [17].

Human-AI Interaction Models: Future cybersecurity strategies need to focus on exploiting the strength produced by the synergy between AI and human thinking, rather than eliminating one for more of the other [18].

Today, by integrating cybersecurity with AI-centricity in a manner sooner than ever considered—that is, embedding it to grow and adapt to present threats and those presented with algorithmic rejuvenation of the future (19)—fintech companies can tread the path towards a more ambitious construct.

In contention, the Algorithmic Fortress is much more than just a useful concept; it is the strategic command for the rightful placement of financial technology in the delivery of AI-led, adaptive, and responsible defenses, which in turn, can catapult their innovations and inspire trust in tomorrow's digital economy.

References

1. Y. Xin et al., "Machine learning and deep learning methods for cybersecurity," *IEEE Access*, vol. 6, pp. 35365–35381, 2018.
2. W. Chen et al., "Deep reinforcement learning for Internet of Things: A comprehensive survey," *IEEE Communications Surveys & Tutorials*, vol. 23, no. 3, pp. 1659–1692, 2021.
3. Y. Shanthakumar, "Generative adversarial networks: An overview of fraud detection applications," *IEEE Communications Surveys & Tutorials*, vol. 23, no. 2, pp. 1037–1055, 2021.
4. X. Wu, "A deep learning-based approach to fraud detection in financial transactions," *IEEE Transactions on Financial Technology*, vol. 1, no. 1, pp. 20–32, 2022.
5. H. Chen, F. Wang, and L. Zhan, "Monitoring financial transactions in real-time with deep learning," *IEEE Transactions on Computational Social Systems*, vol. 9, no. 3, pp. 674–683, 2022.
6. A. Mehmood, "Deep reinforcement learning for fraud detection in credit card transactions," *IEEE Access*, vol. 9, pp. 43676–43685, 2021.
7. J. Lee and S. Cho, "Integrating generative models for fraud detection: A case study in financial services," *IEEE Transactions on Services Computing*, vol. 14, no. 4, pp. 1486–1498, 2021.
8. S. Chaliasos et al., "Smart contract and DeFi security tools: Do they meet the needs of practitioners?," in *Proc. 46th IEEE/ACM Int. Conf. Software Engineering*, 2024, pp. 1–13.

9. T. Chen et al., "A survey of DeFi security: Challenges and opportunities," *Journal of Information Security and Applications*, vol. 65, p. 103130, 2022.
10. N. D. Nguyen, T. Nguyen, and S. Nahavandi, "System design perspective for human-level agents using deep reinforcement learning: A survey," *IEEE Access*, vol. 5, pp. 27091–27102, 2017.
11. A. Sharma et al., "Analysis of security data from a large computing organization," in *Proc. IEEE/IFIP Int. Conf. Dependable Systems and Networks*, 2011, pp. 506–517.
12. L. Xiao et al., "IoT security techniques based on machine learning," *arXiv preprint arXiv:1801.06275*, 2018.
13. D. Ding et al., "A survey on security control and attack detection for industrial cyber-physical systems," *Neurocomputing*, vol. 275, pp. 1674–1683, 2018.
14. M. Wu, Z. Song, and Y. B. Moon, "Detecting cyber-physical attacks in CyberManufacturing systems with machine learning methods," *Journal of Intelligent Manufacturing*, vol. 30, no. 3, pp. 1111–1123, 2019.
15. A. Tsantekidis et al., "Price trailing for financial trading using deep reinforcement learning," *IEEE Transactions on Neural Networks and Learning Systems*, vol. 31, no. 7, pp. 2469–2489, 2020.
16. T. T. Nguyen, "A multi-objective deep reinforcement learning framework," *arXiv preprint arXiv:1803.02965*, 2018.
17. X. Wang et al., "Approximate policy-based accelerated deep reinforcement learning," *IEEE Transactions on Neural Networks and Learning Systems*, vol. 31, no. 6, pp. 1820–1830, 2020.
18. M. H. Ling et al., "Application of reinforcement learning for security enhancement in cognitive radio networks," *Applied Soft Computing*, vol. 37, pp. 809–829, 2015.
19. Y. Wang et al., "A survey of dynamic spectrum access techniques for cognitive radio networks," *IEEE Communications Surveys & Tutorials*, vol. 17, no. 1, pp. 74–94, 2015.
20. N. D. Nguyen et al., "Manipulating soft tissues by deep reinforcement learning for autonomous robotic surgery," in *Proc. IEEE Int. Systems Conf.*, 2019, pp. 1–6.
21. Y. Keneshloo et al., "Deep reinforcement learning for sequence-to-sequence models," *IEEE Transactions on Neural Networks and Learning Systems*, vol. 31, no. 7, pp. 2469–2489, 2020.
22. M. Mahmud et al., "Applications of deep learning and reinforcement learning to biological data," *IEEE Transactions on Neural Networks and Learning Systems*, vol. 29, no. 6, pp. 2063–2079, 2018.
23. M. Popova, O. Isayev, and A. Tropsha, "Deep reinforcement learning for de novo drug design," *Science Advances*, vol. 4, no. 7, p. eaap7885, 2018.
24. Y. He et al., "Software-defined networks with deep reinforcement learning: A survey," *IEEE Communications Surveys & Tutorials*, vol. 21, no. 4, pp. 3133–3174, 2019.
25. V. Mnih et al., "Playing Atari with deep reinforcement learning," *arXiv preprint arXiv:1312.5602*, 2013.
26. T. Schaul et al., "Prioritized experience replay," *arXiv preprint arXiv:1511.05952*, 2015.
27. R. J. Williams, "Simple statistical gradient-following algorithms for connectionist reinforcement learning," *Machine Learning*, vol. 8, no. 3–4, pp. 229–256, 1992.
28. S. Chaliasos et al., "Smart contract and DeFi security tools: Do they meet the needs of practitioners?," in *Proc. 46th IEEE/ACM Int. Conf. Software Engineering*, 2024, pp. 1–13.
29. T. Chen et al., "A survey of DeFi security: Challenges and opportunities," *Journal of Information Security and Applications*, vol. 65, p. 103130, 2022.
30. N. D. Nguyen, T. Nguyen, and S. Nahavandi, "System design perspective for human-level agents using deep reinforcement learning: A survey," *IEEE Access*, vol. 5, pp. 27091–27102, 2017.