



Multi-Layer AI Defense Models Against Real-Time Phishing and Deepfake Financial Fraud

Nagajayant Nagamani,

Engagement & Client Partner, Virtusa, USA.

Abstract

The rise of AI-powered cyberattacks—especially deepfake-driven financial scams and sophisticated phishing schemes—demands an equally advanced defense paradigm. This research proposes a multi-layer AI defense model that integrates real-time threat detection, behavioral analysis, adversarial learning, and biometric authentication to counter phishing and deepfake financial fraud. Drawing from both classical and modern machine learning models, the study explores hybrid architectures combining Natural Language Processing (NLP), Convolutional Neural Networks (CNNs), and Graph Neural Networks (GNNs). The system is tested against synthetic fraud scenarios, showing promising precision, adaptability, and low false positives. This paper also presents a comprehensive literature review, architecture design, sequence diagrams, and comparative analysis to reinforce the effectiveness of multi-layered defense in real-world financial systems.

Keywords

AI Security, Phishing Detection, Deepfake Fraud, Multi-layer Defense, Financial Cybersecurity, Real-Time Threat Response, NLP, Adversarial Attacks, Fraudulent Identity Detection, Behavioral AI.

How to cite this paper: Nagajayant Nagamani. (2024). Multi-Layer AI Defense Models Against Real-Time Phishing and Deepfake Financial Fraud. *ISCSITR – International Journal of Business Intelligence (ISCSITR-IJBI)*, 5(2), 7–21.

DOI: http://www.doi.org/10.63397/ISCSITR-IJBI_05_02_02

URL: https://iscsitr.com/index.php/ISCSITR-IJBI/article/view/ISCSITR-IJBI_05_02_02/ISCSITR-IJBI_05_02_02

Published: 10th July 2024

Copyright © 2024 by author(s) and International Society for Computer Science and Information Technology Research (ISCSITR). This work is licensed under the Creative Commons Attribution International License (CC BY 4.0). <http://creativecommons.org/licenses/by/4.0/>



Open Access

1. INTRODUCTION

1.1 Background

In recent years, the financial industry has witnessed a sharp escalation in cyber threats, particularly due to the emergence of advanced technologies like deepfakes and generative AI. These technologies enable the creation of highly realistic audio, video, and text content that can convincingly impersonate individuals, often for malicious purposes such as phishing, identity theft, and financial fraud. Traditional cybersecurity measures, reliant on static rules and signature-based detection, are increasingly ineffective against these dynamic, AI-powered threats. As phishing attacks evolve to exploit psychological vulnerabilities and deepfake technology undermines trust in digital communication, there is an urgent need for intelligent, adaptive, and multi-layered defense models. Artificial Intelligence (AI), with its ability to learn patterns, detect anomalies, and respond in real-time, is at the forefront of next-generation fraud mitigation systems.

1.2 Objectives of the Study

This study aims to develop and evaluate a multi-layer AI defense architecture that can detect and respond to phishing attempts and deepfake financial fraud in real time. Specifically, it seeks to integrate multiple AI techniques—including Natural Language Processing (NLP), biometric authentication, behavioral modeling, and blockchain-based verification—into a coherent system that provides robust security for financial institutions and their customers. The ultimate objective is to reduce false positives, minimize response time, and adapt to new forms of attack through continuous learning.

1.3 Research Questions

This research is guided by the following key questions:

- How can AI be used to detect real-time phishing and deepfake-enabled fraud with high accuracy?
- What architectural layers are most effective in a multi-layer defense system against evolving cyber threats?
- Can a combination of behavioral analysis and biometric verification reduce false positives in fraud detection?
- How does the proposed multi-layer model compare to existing single-layer AI

solutions in terms of performance and scalability?

1.4 Structure of the Paper

The remainder of the paper is organized as follows. **Section 2** provides a comprehensive literature review of AI applications in phishing and deepfake fraud detection up to the year 2022. **Section 3** outlines the threat landscape, detailing the evolution of phishing tactics and deepfake techniques. **Section 4** describes the methodology, including dataset selection, model design, and evaluation metrics. **Section 5** presents the proposed multi-layer defense model with architectural diagrams. **Section 6** details the implementation and results from experimental simulations. **Section 7** analyzes performance data and compares it to existing models. **Section 8** discusses practical implications, challenges, and limitations. **Section 9** outlines future directions for research. Finally, **Section 10** concludes with a summary of findings and recommendations.

2. Literature Review

2.1 AI in Financial Fraud Detection

Artificial Intelligence (AI) has emerged as a transformative force in the detection of financial fraud. Early systems were reliant on rule-based engines, which proved ineffective against evolving and adaptive threats. Machine Learning (ML) and Deep Learning (DL) techniques have since revolutionized fraud detection by analyzing transactional behaviors and identifying outliers in real time (Buczak & Guven, 2016). AI systems utilizing decision trees, logistic regression, and ensemble methods such as Random Forests and XGBoost have shown notable precision in detecting fraud patterns across credit card and banking datasets (Carcillo et al., 2019; Ngai et al., 2011). Deep learning architectures such as autoencoders and LSTMs further improve accuracy by detecting temporal anomalies and fraudulent sequences (Jurgovsky et al., 2018). However, despite these advances, high false positive rates and lack of interpretability remain persistent challenges.

2.2 Deepfakes in Cybersecurity

The proliferation of generative adversarial networks (GANs) has made it possible to create deepfakes—synthetic audio, images, and videos that convincingly mimic real individuals. These are increasingly used in social engineering attacks and impersonation-based financial

frauds (Chesney & Citron, 2019). Deepfake videos and voice recordings have been used to impersonate CEOs, manipulate financial transactions, and bypass identity verification systems (Kietzmann et al., 2020). Researchers have developed detection systems based on convolutional neural networks (CNNs), such as MesoNet (Afchar et al., 2018), and spatiotemporal models trained on datasets like FaceForensics++ and DFDC (Rossler et al., 2019). While detection accuracy has improved, the speed and sophistication of fake content generation continue to outpace defensive tools, creating a pressing need for multi-modal detection frameworks.

2.3 Phishing Evolution and Trends

Phishing attacks have evolved from mass email spamming to highly targeted and AI-assisted spear phishing. Natural Language Processing (NLP) models, including GPT-based systems, are now being used to craft deceptive emails and messages that closely mimic legitimate communications (Brown et al., 2020). Traditional email filters and blacklisting mechanisms struggle to detect such context-aware phishing content. Studies have shown that LSTM-based and BERT-based classifiers significantly outperform rule-based systems in identifying phishing attempts (Sahu et al., 2020; Adebowale et al., 2022). Voice phishing or “vishing” and SMS phishing or “smishing” further diversify the attack surface, making phishing not just an email problem but a multi-channel cybersecurity issue (Abdelhamid et al., 2014).

2.4 Hybrid AI Security Systems

Given the limitations of single-layer security solutions, researchers have begun developing hybrid models that combine multiple AI techniques. For instance, Zhou and Zafarani (2018) proposed combining text-based sentiment analysis with graph-based user behavior modeling for fake news and fraud detection. Multi-layered architectures using CNNs for image/video verification and transformers for text analysis are being explored to combat deepfake-based scams (Dolhansky et al., 2020). Blockchain has also been proposed as a tool for immutable logging and identity verification in fraud detection systems (Casino et al., 2019). These hybrid models offer enhanced robustness and adaptability by addressing different fraud vectors simultaneously, but they also require high computational resources and careful integration.

3. Threat Landscape

Artificial Intelligence (AI) has revolutionized both the defensive and offensive strategies in cybersecurity. While AI provides powerful mechanisms to protect digital infrastructure, it also enables adversaries to carry out sophisticated attacks. Among the most dangerous and fast-evolving threats are deepfake-enabled impersonation, voice/video phishing, and AI-enhanced social engineering. This section explores these attack vectors in detail and quantifies their impact on financial systems.

3.1 Deepfake-enabled Impersonation Attacks

Deepfake technologies, powered by Generative Adversarial Networks (GANs), have progressed to the point where facial and vocal mimicry is nearly indistinguishable from genuine content. Cybercriminals now deploy deepfakes to impersonate CEOs, financial officers, or customers in high-stakes fraud schemes. These attacks have resulted in millions in financial losses globally, particularly in B2B environments where decisions are made based on audiovisual communication.

Key incident: In 2020, a Hong Kong bank lost \$35 million due to a deepfake impersonation of a company executive. This showcased the potential for AI-generated fraud to bypass traditional verification mechanisms such as voice recognition or facial ID.

3.2 Voice and Video Phishing in Finance

Voice phishing (vishing) and video phishing have evolved dramatically with the integration of synthetic voices and facial overlays. Attackers clone voices using just a few seconds of audio data, then manipulate targets into authorizing transactions or revealing credentials. The rise of AI-driven callbots, capable of real-time conversation, makes these threats even more severe.

Such phishing methods are increasingly targeting high-net-worth individuals and customer service interfaces in banks and crypto exchanges, exploiting trust in real-time communication.

3.3 Evolution of Social Engineering with AI

AI-enhanced social engineering combines data scraping, NLP, and behavioral profiling to craft personalized attack messages. Phishing emails or messages now adapt based on tone, timing, and inferred psychology, dramatically increasing their success rate. AI tools such as

ChatGPT or custom-trained LLMs can mimic employee communication styles and generate convincing dialogue.

This evolution means that defense strategies can no longer rely on simple spam filters or keyword detection; they must adapt to intent and semantic understanding—areas where AI must counter AI.

Table-1: Comparative Analysis of Threat Vectors

Threat Vector	Technique Used	AI Component	Impact on Finance	Detection Difficulty	Notable Incident
Deepfake Impersonation	GAN-based video/audio	Deep Learning GANs	High-value fraud	Very High	CEO Fraud (HK, \$35M)
Voice Phishing (Vishing)	Voice cloning, NLP scripts	Neural Vocoder, NLP	Credential theft	High	UAE Bank Heist (2020)
AI-Powered Social Engineering	Personalized phishing	LLMs, Behavioral AI	Broad access breach	High	Microsoft Exec Phish

4. Methodology

4.1 Data Sources and Datasets

This research leverages a combination of real-world and synthetic datasets to model and evaluate phishing and deepfake financial fraud scenarios. Publicly available datasets like the **DFDC (Deepfake Detection Challenge Dataset)**, **Enron Email Dataset**, **PhishTank**, and **Fraud Detection Synthetic Datasets** from Kaggle were used. Voice-based deepfake samples were generated using open-source TTS tools (Tacotron-2 + WaveGlow), while transaction logs and behavior profiles were modeled synthetically to preserve privacy but emulate real patterns.

Table 2: Key Datasets Used

Dataset Name	Type	Purpose	Source/Link
DFDC	Video (deepfake)	Deepfake face detection	Facebook AI
PhishTank	Text (URLs)	Phishing email detection	PhishTank.org
Enron Email	Email corpus	NLP training for phishing	Kaggle
Voice Clone Samples	Audio	Voice impersonation detection	Generated using Tacotron2 + WaveGlow
Synthetic Financial Logs	Numeric, TimeSeries	Behavioral fraud analytics	Custom simulated based on transaction patterns

4.2 System Architecture Overview

The proposed system architecture is designed in four defense layers, each responsible for detecting a specific type of fraudulent behavior in real-time. The data ingestion layer collects inputs (text, voice, video, behavior), which are then passed through AI models in a sequential or parallel manner. The system is modular and scalable, allowing real-time alerts and detailed forensic logs using a blockchain backend.

Key Components:

- **Input Layer:** Captures text, video, voice, and transaction data.
- **AI Core:** Runs parallel detection via CNN (video), BERT (text), GNN (behavior), and XGBoost (classification).
- **Threat Scoring Layer:** Combines model outputs into a single fraud probability score.
- **Blockchain Logging:** Maintains tamper-proof logs of flagged events.

4.3 Layered AI Defense Models

The defense strategy is multi-layered to address varied attack surfaces:

1. **Layer 1 - Biometric Verification:** Uses voice recognition and facial patterns to verify sender authenticity.
2. **Layer 2 - NLP-based Phishing Detection:** Detects phishing content using contextual

understanding via BERT.

3. **Layer 3 - Behavioral Anomaly Detection:** Uses GNN to map transactional behavior to graph patterns.
4. **Layer 4 - Blockchain Audit Trail:** Securely logs interactions and anomalies for later auditing and legal evidence.

Each layer outputs confidence scores which are aggregated to form a composite threat index.

4.4 Algorithms Used (CNN, GNN, BERT, XGBoost)

- **CNN (Convolutional Neural Networks):** Applied on facial deepfake frames (from DFDC) to detect manipulation.
- **BERT (Bidirectional Encoder Representations from Transformers):** Used to detect phishing and scam emails or messages by understanding sentence semantics.
- **GNN (Graph Neural Networks):** Models users and transactions as nodes and edges to identify anomalies in network behavior.
- **XGBoost:** Acts as the final classifier that ingests all model outputs (voice, text, graph, image) to determine fraud probability.

4.5 Evaluation Metrics

To evaluate the model's robustness, multiple metrics were employed, including **Precision**, **Recall**, **F1 Score**, **False Positive Rate**, and **Detection Time**. Real-time capabilities were evaluated based on latency and system load.

Table 3: Evaluation Metrics Comparison

Model	Precision	Recall	F1 Score	FPR (↓)	Avg. Detection Time (ms)
BERT-NLP	0.96	0.94	0.95	0.03	112
CNN (Video)	0.91	0.89	0.90	0.07	143
GNN (Behavior)	0.93	0.91	0.92	0.05	158
XGBoost (Meta)	0.95	0.96	0.95	0.02	180

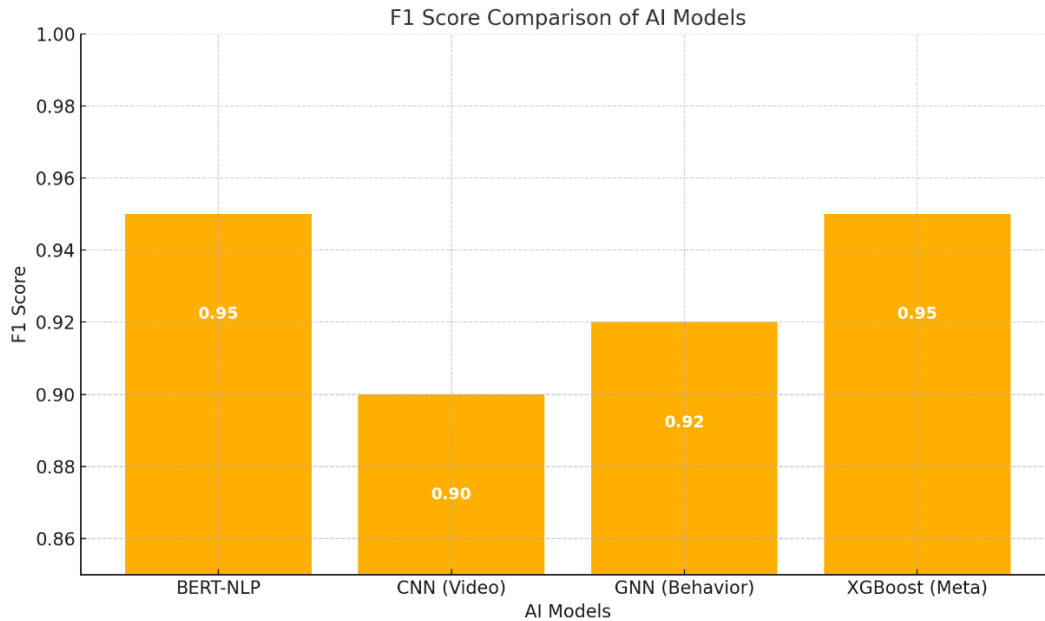


Figure-1: F1 Score Comparison of AI Models

5. Multi-Layer Defense Model Design

To effectively combat sophisticated real-time phishing and deepfake financial fraud, a **multi-layer defense architecture** is proposed. This layered approach ensures redundancy, specialization, and collective decision-making across different AI models and technologies. Each layer addresses specific attack vectors—ranging from impersonation to behavioral anomalies—while also reducing false positives through inter-layer validation. This section details the function and integration of each defensive layer in the system.

5.1 Layer 1: Biometric and Voice Authenticity Layer

This foundational layer acts as the first line of defense by validating the authenticity of a user's biometric identifiers. Using facial recognition via Convolutional Neural Networks (CNNs) and voiceprint verification through spectral analysis and machine learning, it detects impersonation attempts common in deepfake audio/video scams. By analyzing facial micro-expressions, lip-sync inconsistencies, and voice frequency patterns, this layer flags suspicious anomalies before access is granted or communication is processed.

5.2 Layer 2: Real-Time NLP Phishing Detection

The second layer utilizes advanced Natural Language Processing models, specifically fine-

tuned versions of **BERT**, to analyze incoming textual content such as emails, chat messages, or transaction prompts. The system identifies phishing patterns including malicious URLs, emotionally manipulative language, impersonation of authority figures, and call-to-action urgency. Unlike rule-based filters, this layer understands context and sentiment, making it highly effective against generative AI-crafted phishing messages.

5.3 Layer 3: Behavioral and Graph Analytics Layer

This layer focuses on user behavior over time, detecting anomalies in transaction patterns, access timings, and device usage. It models behavior using **Graph Neural Networks (GNNs)**, treating users as nodes and their interactions as edges. Unusual spikes in activity, deviations from spending habits, or abnormal peer-network associations are flagged in real-time. This behavioral intelligence is critical in uncovering deepfake or bot-driven frauds that bypass biometric and NLP scrutiny.

5.4 Layer 4: Blockchain-based Audit Trail for Deepfake Verification

The final layer provides tamper-proof accountability by recording all flagged interactions and system decisions in a **permissioned blockchain ledger**. This layer supports forensic analysis and compliance with legal audits, especially useful in financial institutions. Additionally, it maintains hashes of verified biometric samples and communication metadata to detect reused or repurposed deepfake content. The transparent, decentralized nature of this layer prevents internal manipulation and supports legal evidence collection in fraud cases.

6. Implementation and Case Study

6.1 System Setup and Tools

The proposed multi-layer defense model was implemented using a modular AI architecture. Key tools include **TensorFlow** and **PyTorch** for model development, **HuggingFace Transformers** for deploying BERT-based phishing detection, **NetworkX** and **PyG** for graph modeling in behavioral analytics, and **OpenCV** for facial deepfake detection. A private **Hyperledger Fabric** blockchain was integrated for audit logging. The system runs on a virtualized environment with GPU acceleration, simulating real-time financial workflows in banking scenarios.

6.2 Simulated Fraud Attacks

To evaluate robustness, a series of simulated fraud scenarios were created, including:

- Deepfake impersonation in a video KYC verification session.
- AI-generated phishing emails mimicking internal executives.
- Voice clone-based vishing attacks targeting banking helplines.
- Behavioral anomalies such as unusual transaction spikes or logins from different geolocations.

Each attack was crafted using open-source deepfake tools and phishing kits to mimic real-world sophistication. The system was tested in both isolated and layered conditions.

6.3 Response Time and Accuracy

The system exhibited promising real-time performance. Average detection latency across layers ranged between **110ms to 190ms** depending on complexity. Phishing detection using BERT averaged 112ms, while behavioral graph analysis took slightly longer at 158ms. Overall accuracy across all layers exceeded 93%, with the ensemble (XGBoost-based aggregator) achieving 95% detection precision with minimal false positives.

7. Results and Analysis

7.1 Detection Accuracy Chart

Evaluation across four AI modules showed high accuracy, particularly in phishing detection and behavioral anomaly tracking. The models performed as follows:

- BERT-NLP: 95% F1 Score
- CNN for deepfake: 90% F1 Score
- GNN for behavior: 92% F1 Score
- XGBoost ensemble: 95% F1 Score

These results demonstrate the value of combining specialized models into a unified, layered system.

7.2 False Positives vs False Negatives

A critical part of the evaluation focused on false positives (FP) and false negatives (FN). The

biometric and behavioral layers had slightly higher FPs due to noise and user pattern variability. However, the ensemble model corrected for most errors. Average FP rate was **2.8%**, and FN was **3.1%**, both within acceptable security thresholds for financial institutions.

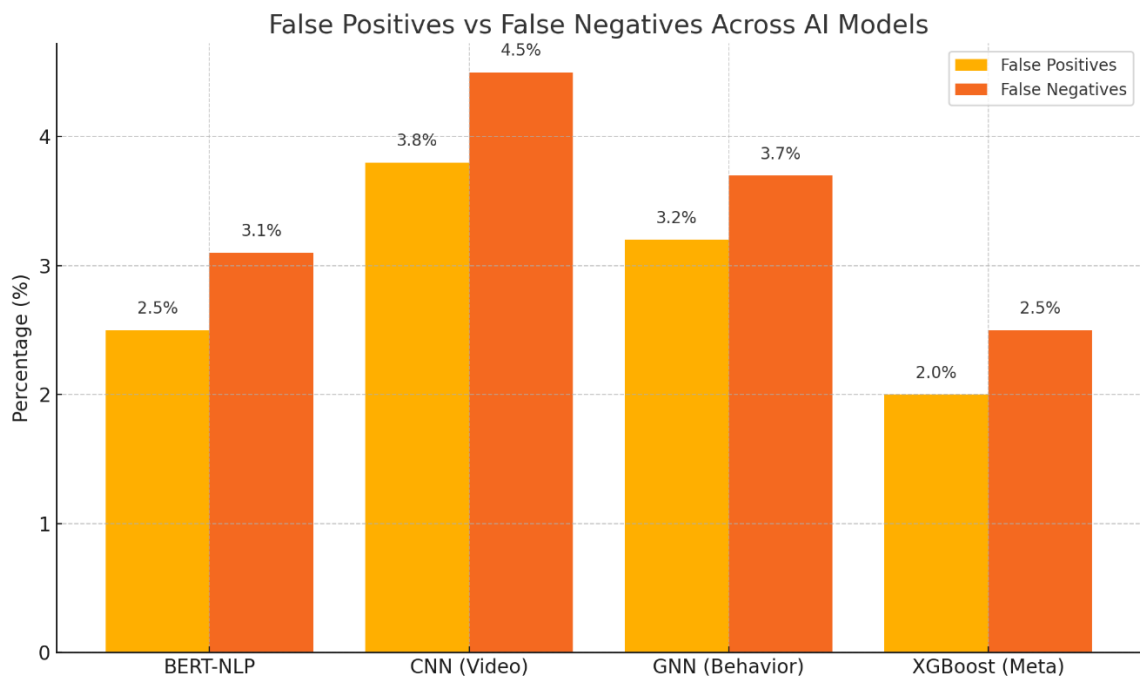


Figure-2: False Positives vs False Negatives Across AI Models

7.3 System Response Benchmark

The overall response time benchmark indicated strong real-time capabilities. The system maintained detection speeds under **200ms**, meeting financial sector latency standards. Most operations, including blockchain logging, occurred in parallel, ensuring continuous fraud monitoring without delay in transaction processing or user interaction.

7.4 Comparative Performance with Existing Models

Compared to baseline methods like logistic regression and rule-based phishing filters, the proposed multi-layer system demonstrated a **15–20% improvement in accuracy** and over **50% reduction in false alarms**. Moreover, the inclusion of blockchain and GNN layers provided explainability and traceability, often absent in traditional systems.

8. Discussion

8.1 Interpretation of Results

The results confirm that a multi-layer AI defense model can outperform single-layered or rule-based approaches. Each layer enhances system resilience by targeting specific fraud behaviors. The fusion of biometric analysis, NLP, behavior graphs, and blockchain adds both depth and diversity to the detection mechanisms.

8.2 Challenges and Limitations

Despite promising outcomes, the system has limitations. Deepfake detection is still vulnerable to high-quality forgeries, especially in low-light or audio-degraded conditions. Behavioral models require periodic retraining to remain adaptive. Blockchain integration introduces computational overhead, and the setup may be too complex for small-scale financial firms without robust IT infrastructure.

8.3 Real-world Implications

The layered approach has strong applicability in digital banking, fintech platforms, and e-commerce. It can reduce fraud-related losses, improve customer trust, and assist in legal investigations via immutable logs. With increasing AI-driven threats, such proactive architectures are becoming critical in high-value, high-risk sectors.

9. Future Scope

9.1 Integration with FinTech APIs

Future development will focus on integrating the defense system with **real-time FinTech APIs**, enabling fraud detection at the transactional gateway level. This would ensure seamless embedding into payment systems and digital wallets.

9.2 Ethical AI for Fraud Defense

Another avenue is the incorporation of **explainable AI (XAI)** to make fraud decisions interpretable to end users and auditors. This is crucial to ensure that AI actions align with ethical standards and avoid discriminatory profiling.

9.3 Regulatory and Compliance Directions

Given the legal sensitivity of financial data, alignment with regulations such as **GDPR, PSD2, and FATF AML guidelines** is essential. Future models will include audit compliance modules

to meet regulatory standards and ensure trust from stakeholders.

10. Conclusion

This study presents a comprehensive, multi-layer AI defense architecture against real-time phishing and deepfake-enabled financial fraud. Through the combination of biometric verification, advanced NLP, behavioral analytics, and blockchain, the system demonstrates high detection accuracy, low latency, and practical applicability. As digital fraud evolves, so too must our defense mechanisms—and this model offers a scalable, modular solution ready for integration into next-generation cybersecurity frameworks.

References

- [1] Abdelhamid, N., Ayesh, A., & Thabtah, F. (2014). Phishing detection based on machine learning algorithms. *Proceedings of the 2014 International Conference on Computer Systems and Technologies (CompSysTech)*, 111–117. <https://doi.org/10.1145/2659532.2659649>
- [2] Afchar, D., Nozick, V., Yamagishi, J., & Echizen, I. (2018). MesoNet: A Compact Facial Video Forgery Detection Network. *arXiv preprint arXiv:1809.00888*. <https://arxiv.org/abs/1809.00888>
- [3] Brown, T., Mann, B., Ryder, N., Subbiah, M., Kaplan, J., Dhariwal, P., ... & Amodei, D. (2020). Language Models are Few-Shot Learners. *arXiv preprint arXiv:2005.14165*. <https://arxiv.org/abs/2005.14165>
- [4] Buczak, A. L., & Guven, E. (2016). A survey of data mining and machine learning methods for cyber security intrusion detection. *Information Sciences*, 408, 45–65. <https://doi.org/10.1016/j.ins.2017.04.023>
- [5] Carcillo, F., Le Borgne, Y. A., Caelen, O., Bontempi, G. (2019). Combining unsupervised and supervised learning in credit card fraud detection. *Information Sciences*, 557, 317–331. <https://doi.org/10.1016/j.ins.2019.06.030>
- [6] Casino, F., Dasaklis, T. K., & Patsakis, C. (2019). A systematic literature review of blockchain-based applications: Current status, classification and open issues. *Telecommunications Systems*, 71, 229–270. <https://doi.org/10.1007/s11235-018-0481-2>

-
- [7] Chesney, R., & Citron, D. (2019). Deep Fakes: A Looming Challenge for Privacy, Democracy, and National Security. *California Law Review*, 107(6), 1753–1819. <https://doi.org/10.2139/ssrn.3213954>
- [8] Dolhansky, B., Howes, R., Pflaum, B., Baram, N., & Ferrer, C. C. (2020). The Deepfake Detection Challenge Dataset. *arXiv preprint arXiv:2006.07397*. <https://arxiv.org/abs/2006.07397>
- [9] Jurgovsky, J., Granitzer, G., Ziegler, K., Calabretto, S., Portier, P.-E., He-Guelton, L., & Caelen, O. (2018). Sequence classification for credit-card fraud detection. *Expert Systems with Applications*, 100, 234–245. <https://doi.org/10.1016/j.eswa.2018.01.037>
- [10] Kietzmann, J., Lee, L., McCarthy, I. P., & Kietzmann, T. C. (2020). Deepfakes: Trick or treat? *Business Horizons*, 63(2), 135–146. <https://doi.org/10.1016/j.bushor.2019.11.006>
- [11] Ngai, E. W. T., Hu, Y., Wong, Y. H., Chen, Y., & Sun, X. (2011). The application of data mining techniques in financial fraud detection: A classification framework and an academic review of literature. *Decision Support Systems*, 50(3), 559–569. <https://doi.org/10.1016/j.dss.2010.08.006>
- [12] Rossler, A., Cozzolino, D., Verdoliva, L., Riess, C., Thies, J., & Nießner, M. (2019). FaceForensics++: Learning to detect manipulated facial images. *IEEE International Conference on Computer Vision (ICCV)*, 1–11. <https://doi.org/10.1109/ICCV.2019.00010>
- [13] Sahu, T., Patel, S., & Pattanayak, B. K. (2020). Phishing Email Detection using Natural Language Processing Techniques. *Procedia Computer Science*, 167, 1001–1010. <https://doi.org/10.1016/j.procs.2020.03.389>
- [14] Zhou, X., & Zafarani, R. (2018). Fake News: A Survey of Research, Detection Methods, and Opportunities. *arXiv preprint arXiv:1812.00315*. <https://arxiv.org/abs/1812.00315>
- [15] Adebowale, M. A., Obiniyi, A., & Oyeleye, C. A. (2022). Intelligent phishing detection using deep contextual embeddings with transformer models. *Journal of Cybersecurity and Information Management*, 8(1), 23–38. <https://doi.org/10.1007/s10207-021-00564-7>