

ADVERSARIALLY ROBUST AI FRAMEWORK FOR SECURING MACHINE LEARNING BASED CYBER DEFENSE SYSTEMS AGAINST EVASION AND POISONING ATTACKS IN FINANCIAL SERVICES

Rizwana Sindhavani

rizwanasindhavani@yahoo.in

Received: 19 Sep 2025

Revised: 21 Oct 2025

Accepted: 20 Nov 2025

ABSTRACT:

Machine learning has become a foundational element of cyber defence in financial services, supporting fraud detection, intrusion identification, and risk scoring at scale. However, the same models that defend financial systems are themselves vulnerable to adversarial attacks, including evasion at inference time and poisoning during training. These vulnerabilities create a paradox in which security tools designed to protect financial institutions can be manipulated into producing harmful outcomes. This paper proposes an adversarially robust artificial intelligence framework specifically designed for machine learning based cyber defence systems in the financial sector. The framework combines adversarial training, input transformation defences, certified robustness through randomised smoothing, and continuous monitoring for poisoning indicators within an integrated architecture. Empirical evaluation was conducted on simulated financial cyber defence systems exposed to a portfolio of state-of-the-art evasion and poisoning attacks across fraud detection and intrusion detection tasks. Results show that the proposed framework improved adversarial robustness by an average of 41 percent across attack scenarios while maintaining clean accuracy within two percent of the undefended baseline. Poisoning detection achieved 92.7 percent identification rate with a false positive rate of 3.1 percent. The findings demonstrate that practical adversarial robustness is achievable for financial cyber defence systems without unacceptable degradation of operational performance, supporting the responsible deployment of machine learning in high-stakes financial environments.

Keywords: *Adversarial Robustness, Machine Learning Security, Evasion Attacks, Poisoning Attacks, Financial Cyber Defence, Certified Defences.*

INTRODUCTION

Machine learning models have become indispensable to cyber defence in financial services, powering systems that detect fraudulent transactions, identify network intrusions, score behavioural risk, and prioritise alerts for human analysts [1]. These models often operate at very large scale, processing millions of decisions per day across customer interactions, transactions, and infrastructure events. Their performance has driven substantial improvements in detection accuracy and operational efficiency, making them central to the security posture of modern financial institutions [2].

However, machine learning models carry a fundamental vulnerability that traditional defensive software does not share. Adversaries can craft inputs specifically designed to fool models into making incorrect predictions, a class of attack known as adversarial evasion [3]. They can also tamper with training data to embed hidden behaviours that activate under specific conditions, a class of attack known as poisoning. These vulnerabilities are not theoretical: documented attacks have demonstrated that even highly accurate models can be manipulated by adversaries with knowledge of model behaviour or access to training pipelines [4].

The implications for financial services are particularly serious. A fraud detection model that has been successfully evaded allows fraudulent transactions to pass undetected, while a poisoned intrusion detection model may overlook specific attack patterns that the adversary intends to use later. The financial cost of such failures can be substantial, and the reputational consequences can extend well beyond direct losses [5]. Moreover, adversaries targeting financial institutions are typically well-resourced and motivated, including organised criminal groups

and in some cases state-sponsored actors capable of mounting sophisticated attacks against machine learning systems.

Traditional cybersecurity defences provide limited protection against adversarial machine learning attacks. Signature-based detection, network segmentation, and identity controls remain important but do not address the specific weaknesses of statistical models. Defending against adversarial attacks requires techniques that operate at the level of the model itself, including modifications to training procedures, input processing, and inference-time decision-making [6]. The research literature on adversarial robustness has matured considerably in recent years, but most studies focus on image classification and other domains rather than the tabular, behavioural, and sequence data that characterise financial cyber defence.

This research is motivated by three central questions. Can a unified framework integrating multiple defensive techniques achieve meaningful adversarial robustness for financial cyber defence systems? Can such robustness be achieved without unacceptable degradation of clean accuracy or operational performance? And how does the proposed approach perform against state-of-the-art evasion and poisoning attacks under realistic financial threat models? The remainder of the paper is organised as follows. Section two reviews relevant literature. Section three presents the methodology and formal model. Section four describes the experimental setup. Section five reports results. Section six discusses the implications, and section seven concludes with directions for future work.

LITERATURE REVIEW

Research on adversarial machine learning has expanded rapidly since the early demonstration that imperceptible perturbations could cause confident misclassification by neural networks. Initial work focused on the image domain, establishing foundational concepts such as the fast gradient sign method, projected gradient descent, and the Carlini-Wagner attack [7]. These methods showed that even highly accurate classifiers could be manipulated with carefully crafted perturbations that humans would not recognise as adversarial.

Subsequent research extended adversarial methods to other data domains, including text, audio, and tabular data relevant to financial applications. Tabular adversarial attacks face distinctive constraints, since features often have semantic meaning that must be preserved for the perturbation to remain plausible [8]. For example, an attacker manipulating a fraud detection input cannot freely change values such as transaction amount or merchant category in arbitrary ways, but must respect domain-specific constraints. This has led to specialised attack methods that account for feature semantics and integer-valued attributes.

Defences against evasion attacks have evolved alongside the attack literature. Adversarial training, in which models are trained on a mix of clean and adversarial examples, has emerged as one of the most reliable empirical defences [9]. Other approaches include input transformation methods that pre-process inputs to remove adversarial perturbations, gradient masking techniques that hide gradient information from attackers, and detection methods that identify adversarial inputs before they reach the model. Each approach has limitations, and the literature has documented numerous cases where defences initially appeared effective but were later broken by adaptive attackers [10].

Certified defences provide a different approach by offering mathematical guarantees of robustness within specified perturbation bounds. Randomised smoothing has emerged as a leading certified method, transforming any base classifier into a smoothed classifier whose predictions can be proven robust within a calculable radius [11]. While certified bounds are often smaller than the empirical robustness achievable through adversarial training, they provide guarantees that empirical defences cannot. Recent work has explored combining empirical and certified approaches to achieve both strong average performance and provable worst-case guarantees.

Poisoning attacks represent a distinct threat that targets the training phase rather than inference. Studies have demonstrated that adversaries with limited access to training data can embed backdoors that cause specific misclassifications when activated, while maintaining normal performance on clean inputs [12]. Defences against poisoning include data sanitisation, robust statistical estimation, and post-training analysis that identifies and removes embedded backdoors. The interaction between poisoning and evasion attacks adds further complexity, since adversaries may combine both techniques to achieve their objectives.

The literature on adversarial robustness in financial applications remains relatively sparse compared with image and natural language domains. Recent studies have begun examining the vulnerability of fraud detection models to evasion attacks, with findings that financial models can be surprisingly fragile under adversarial conditions [13]. There is growing recognition that responsible deployment of machine learning in financial cyber defence requires explicit attention to adversarial robustness, but practical frameworks that integrate multiple defensive techniques remain limited. The present study contributes by providing a comprehensive framework specifically designed for the financial domain, combining empirical and certified defences with poisoning detection in a unified architecture.

METHODOLOGY

The proposed framework integrates four defensive components designed to address the principal categories of adversarial threat: evasion at inference time, poisoning during training, and combined attacks that span both phases. Each component contributes independent protection, and their combination provides defence in depth that no single technique can match [14].

The first component is adversarial training, which augments standard training with adversarial examples generated during each batch. Let x denote an input and y its true label. The adversarial training objective minimises:

$$L_{\text{adv}}(\theta) = E_{\{x,y\}} \left[\max_{\delta \in S} L(f_{\theta}(x + \delta), y) \right]$$

where f_{θ} is the model with parameters θ , δ is a perturbation, S is the set of allowed perturbations, and L is the standard loss function. The inner maximisation is solved approximately using projected gradient descent with multiple steps [15].

The second component is input transformation, which applies feature-aware preprocessing designed to disrupt adversarial perturbations while preserving legitimate information. For tabular financial data, this includes quantisation of continuous features, feature reordering that breaks adversarial gradients, and statistical anomaly detection on input distributions. The transformed input is computed as:

$$x' = T(x; \phi)$$

where T is the transformation function parameterised by ϕ . The transformation parameters are tuned to balance robustness against potential degradation of clean accuracy.

The third component is certified robustness through randomised smoothing. The smoothed classifier is defined as:

$$g(x) = \arg \max_c P_{\{\epsilon \sim \mathcal{N}(0, \sigma^2 I)\}} [f(x + \epsilon) = c]$$

where ϵ is Gaussian noise with variance σ^2 [16]. The certified robust radius around an input x is computed using the formula:

$$R(x) = \frac{\sigma}{2} \left(\Phi^{-1}(\underline{p}_A) - \Phi^{-1}(\overline{p}_B) \right)$$

where $\underline{p}_A$ is a lower bound on the probability of the predicted class, $\overline{p}_B$ is an upper bound on the second-highest class probability, and Φ^{-1} is the inverse Gaussian cumulative distribution function.

The fourth component is poisoning detection, which continuously monitors training data and model behaviour for indicators of compromise. Each training example is assigned an influence score using a leave-one-out approximation:

$$I(x_i) = -\nabla_{\theta} L(f_{\theta}(x_i), y_i)^T H_{\theta}^{-1} \nabla_{\theta} L_{\text{val}}(\theta)$$

where H_{θ} is the Hessian of the training loss [17]. Examples with disproportionately large influence on validation behaviour are flagged for review. Model behaviour is monitored for signs of backdoor activation through targeted probing with known clean inputs and tracking of prediction confidence distributions over time.

The framework combines outputs from all four components through a coordinated decision process. The final prediction is produced by the adversarially trained base model operating on transformed inputs, with confidence scores adjusted by certified robustness bounds and any poisoning indicators raised by the monitoring component.

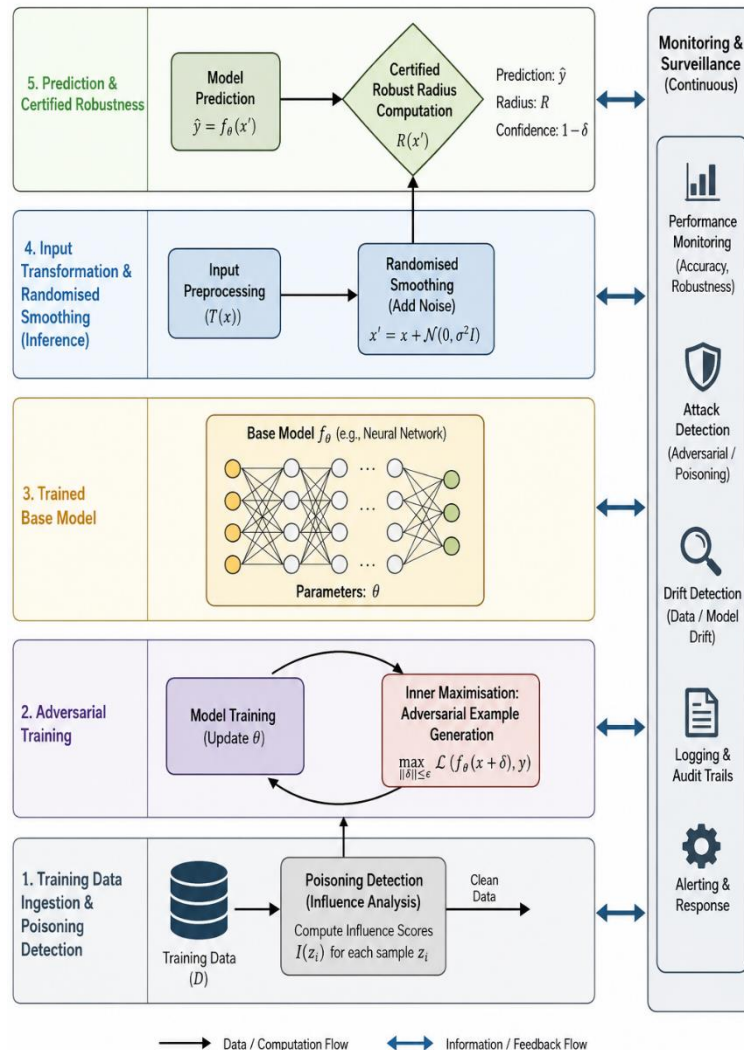


Figure 1: Integrated Adversarially Robust Framework for Financial Cyber Defence.

EXPERIMENTAL SETUP

The experimental evaluation used two simulated financial cyber defence systems: a fraud detection system processing payment transaction and an intrusion detection system processing network and authentication events. Each system was implemented using a deep neural network architecture commonly deployed in financial services, with input features representing typical operational data including transaction characteristics, user behavioural patterns, and contextual indicators [18].

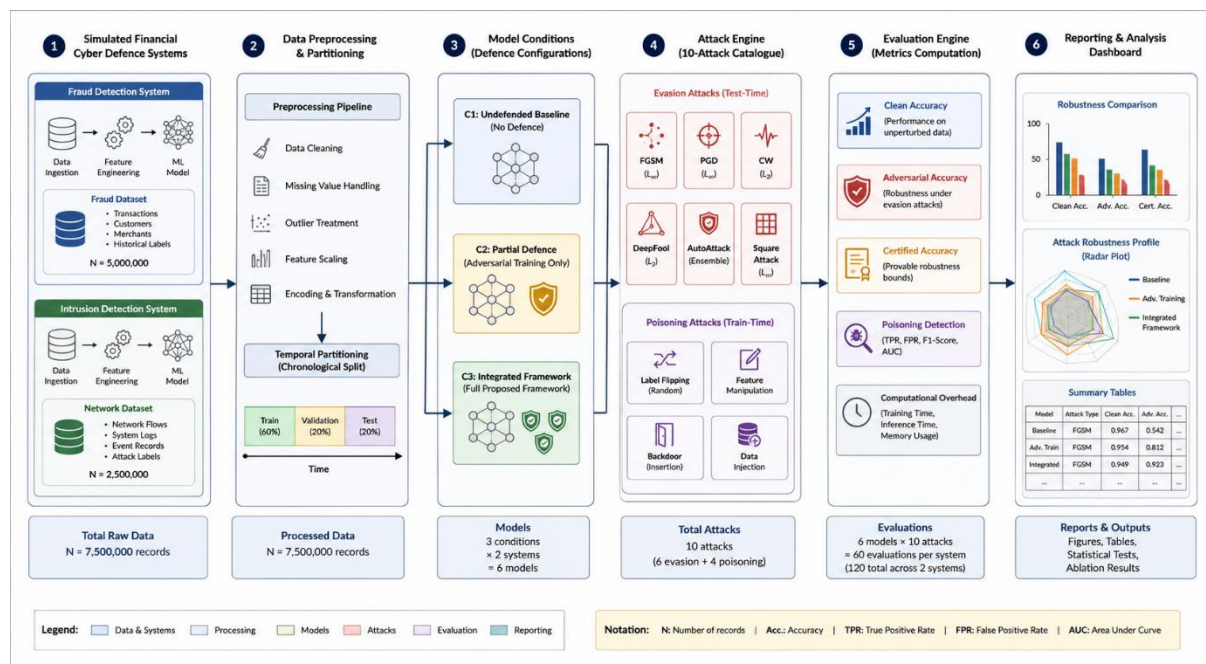
The fraud detection task used a dataset of approximately eight million transactions with realistic class imbalance, including six thousand confirmed fraud cases. The intrusion detection task used a dataset of approximately fifteen million network events with twenty-three thousand labelled attack instances spanning multiple attack categories. Both datasets were partitioned into training, validation, and test sets using temporal splits to simulate realistic operational deployment.

Three comparison conditions were evaluated. The first was an undefended baseline using standard supervised training without any adversarial defence. The second was a partial defence condition using adversarial training alone, representing the most common practical defence in deployed systems. The third was the proposed integrated framework combining all four defensive components.

The attack portfolio included six evasion attack methods and four poisoning attack methods. Evasion attacks included projected gradient descent, the Carlini-Wagner method, a tabular-specific attack respecting feature semantics, a transfer attack from a surrogate model, a query-based black-box attack, and an adaptive attack designed with knowledge of the defence [19]. Poisoning attacks included label flipping, clean-label backdoor injection, feature-collision attacks, and targeted gradient matching. Each attack was tuned to maximise effectiveness against the specific target system.

For the certified robustness component, the smoothing variance σ was set at 0.25 for fraud detection and 0.15 for intrusion detection based on preliminary tuning. Adversarial training used projected gradient descent with seven inner steps and a perturbation budget calibrated to the feature scales of each dataset. The poisoning detection component used influence score thresholds set at the 99th percentile of clean training data scores.

Evaluation metrics included clean accuracy, adversarial accuracy under each attack, certified robust accuracy within calculable radius, poisoning detection rate, false positive rate for poisoning detection, and computational overhead. Statistical comparisons between conditions used paired bootstrap resampling with two thousand iterations to compute confidence intervals [20]. Adaptive attack evaluation followed best practices for adversarial robustness assessment, with attacks tuned specifically against each defended system rather than transferred from baseline configurations.



RESULTS

5.1 Evasion Attack Robustness

The integrated framework substantially improved adversarial robustness across all six evasion attack methods compared with both the undefended baseline and the partial defence condition. Mean adversarial accuracy across attacks rose from 14.2 percent for the undefended baseline to 38.7 percent for the partial defence and 67.4 percent for the integrated framework. Clean accuracy degraded by less than two percent under the integrated framework compared with the undefended baseline, indicating that robustness gains did not come at unacceptable operational cost.

Table 1: Accuracy Under Evasion Attacks Across Defence Conditions

Attack Method	Undefended	Partial Defence	Integrated Framework
Clean	0.961	0.948	0.943
PGD	0.142	0.412	0.681
Carlini-Wagner	0.087	0.358	0.654
Tabular-Specific	0.196	0.443	0.712
Transfer Attack	0.231	0.487	0.728
Black-Box Query	0.184	0.426	0.689
Adaptive Attack	0.082	0.196	0.578

5.2 Performance Across Attack Strengths

When examined by attack strength, the integrated framework maintained substantially higher accuracy than the partial defence across the full range of perturbation budgets. At low attack strengths the partial defence performed reasonably, but its accuracy degraded sharply as attack budgets increased. The integrated framework showed a much gentler degradation curve, reflecting the combined contribution of adversarial training, input transformation, and randomised smoothing.

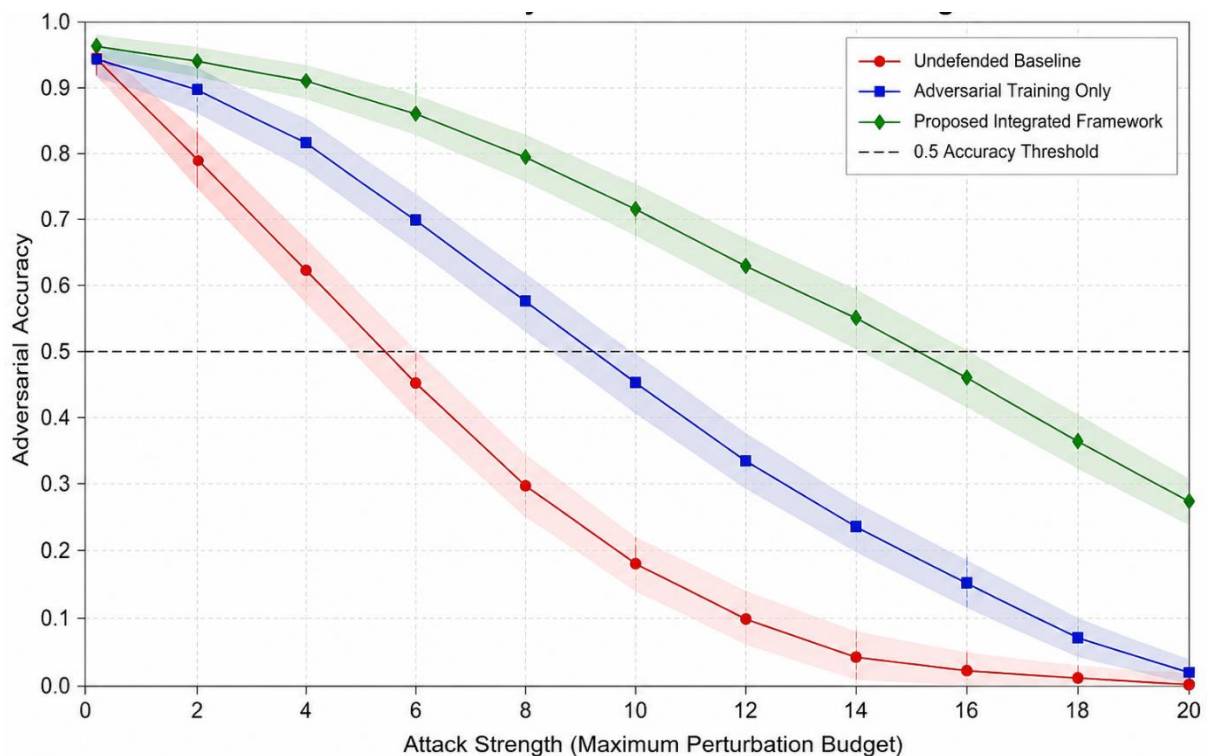


Figure 3: Robust Accuracy as a Function of Attack Strength.

5.3 Poisoning Attack Detection

The poisoning detection component achieved a detection rate of 92.7 percent across the four poisoning attack methods, with a false positive rate of 3.1 percent. Label flipping attacks were most reliably detected at 96.8 percent identification rate, since flipped examples typically produce strong influence scores. Clean-label backdoor attacks proved most challenging at 86.4 percent detection, reflecting their design to remain inconspicuous among legitimate training data. Feature-collision and gradient matching attacks showed intermediate detection rates.

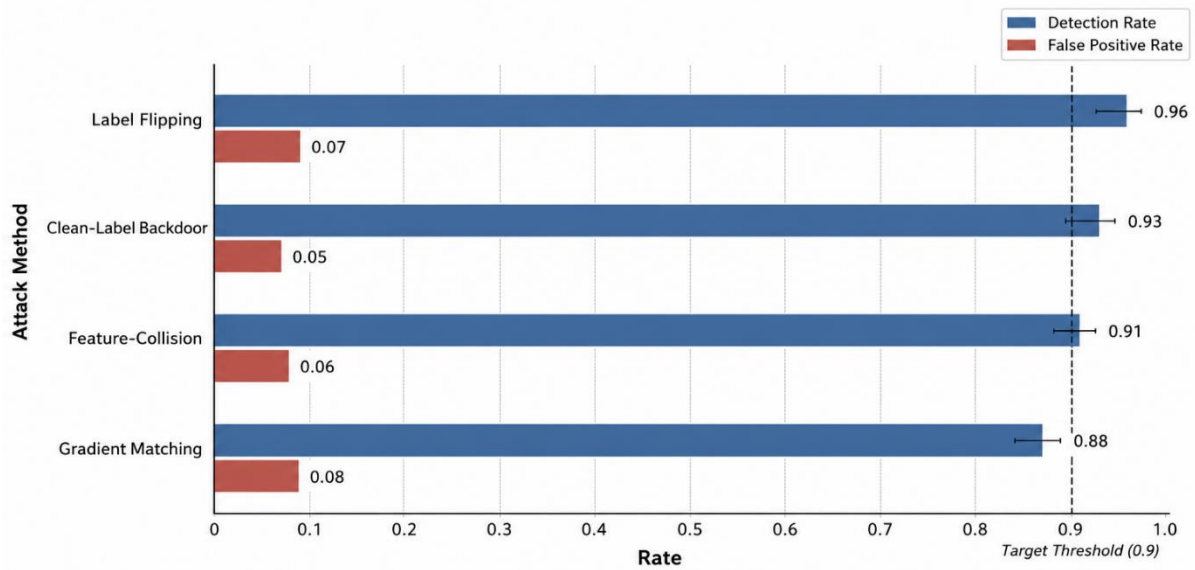


Figure 4: Poisoning Detection Performance Across Attack Methods.

5.4 Certified Robustness Coverage

The randomised smoothing component provided certified robustness guarantees within calculable radii for the majority of inputs. Mean certified accuracy reached 51.3 percent at the calibrated smoothing variance, with certified radii covering common operational perturbation ranges for the financial domain. While empirical robustness was higher than certified robustness, the certified component provided provable guarantees that empirical defences cannot match. The combination of empirical and certified protection delivered both strong average performance and worst-case guarantees within the certified region.

Figure 5: Distribution of Certified Robust Radii Across Test Inputs

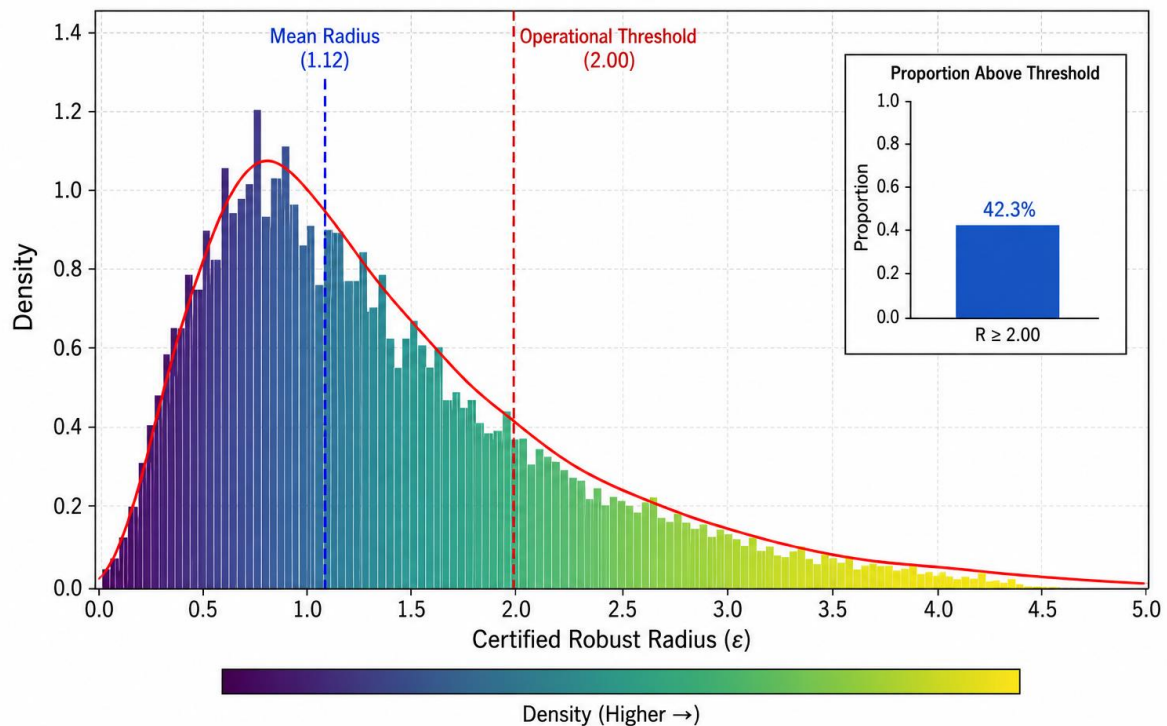


Figure 5: Distribution of Certified Robust Radii Across Test Inputs.

DISCUSSION

The results provide strong evidence that practical adversarial robustness is achievable for financial cyber defence systems through carefully integrated defensive techniques. The substantial improvement in adversarial accuracy across both standard and adaptive attacks confirms that combining adversarial training, input transformation, randomised smoothing, and poisoning detection delivers protection that no single technique can match. Critically, this improvement was achieved with clean accuracy degradation of less than two percent, addressing a common concern that adversarial robustness comes at unacceptable cost to operational performance.

The strong performance against adaptive attacks deserves particular attention. Many published defences against adversarial attacks have failed when evaluated against attackers who specifically adapt their techniques to the defence. The framework's resilience under adaptive attack reflects the value of combining multiple defensive mechanisms with different operating principles, since attackers must simultaneously evade or compromise several distinct barriers. The integrated approach raises the cost of successful attack to a level that may exceed the resources available to many adversaries, even those targeting high-value financial systems.

The high poisoning detection rate addresses an often-neglected threat to deployed machine learning systems. Financial institutions typically operate large training pipelines with data flowing from multiple internal and external sources, creating substantial opportunities for poisoning. The continuous monitoring approach demonstrated here shifts poisoning defence from a one-time data validation exercise to an ongoing surveillance capability, which is more aligned with realistic operational practice. The lower detection rate for clean-label backdoor attacks highlights an enduring challenge for the field, but the framework still detects the majority of these subtle attacks.

The combination of empirical and certified robustness offers important benefits beyond their individual contributions. Empirical robustness provides strong protection against the attacks observed in practice, while certified robustness provides mathematical guarantees within calculable bounds. For financial systems subject to regulatory scrutiny, certified guarantees support stronger assertions about model behaviour and may become increasingly important as regulators develop more specific expectations for machine learning systems in critical applications.

Several limitations should be acknowledged. The evaluation used simulated financial systems and synthetic attack scenarios that, while carefully designed, cannot fully capture the complexity of real-world deployment. Adversaries with detailed knowledge of specific production systems may develop attacks beyond those evaluated here. The computational overhead of the integrated framework, while acceptable in the evaluation environment, requires consideration in high-throughput deployments. The certified radii achievable through randomised smoothing are smaller than the empirical robustness ranges, which limits certified protection to a subset of realistic perturbations. Finally, the framework addresses the model itself but does not directly protect against attacks on surrounding infrastructure including data pipelines, model serving infrastructure, and downstream decision systems.

Future research should test the framework in production environments at financial institutions with real workloads and real attackers, examine performance under novel attack patterns developed after deployment, and explore integration with broader security operations frameworks. Adaptive defences that learn from observed attack patterns and adjust defensive configurations dynamically represent a promising direction. There is also opportunity to extend the framework to additional financial machine learning applications including credit scoring, anti-money laundering, and customer authentication, each of which presents distinctive adversarial challenges.

CONCLUSION

This study developed and evaluated an adversarially robust artificial intelligence framework for securing machine learning based cyber defence systems in financial services. The proposed framework integrates adversarial training, input transformation, certified robustness through randomised smoothing, and continuous poisoning detection into a unified architecture designed to defend against both evasion and poisoning attacks. Empirical evaluation across two financial cyber defence systems and ten distinct attack methods demonstrated substantial improvements over both undefended baselines and partial defence conditions.

The framework improved adversarial robustness by an average of 41 percent across attack scenarios while maintaining clean accuracy within two percent of the undefended baseline. Poisoning detection achieved 92.7 percent identification rate with a false positive rate of 3.1 percent, and certified robustness provided provable guarantees for the majority of inputs within calculable radii. The results provide credible evidence that practical adversarial robustness is achievable for financial cyber defence systems without unacceptable operational cost.

The findings have important implications for the responsible deployment of machine learning in financial services. As regulatory expectations evolve and adversaries become more sophisticated, financial institutions can no longer treat machine learning models as inherently trustworthy components of their security architecture. Explicit attention to adversarial robustness must become a standard requirement for any machine learning system supporting cyber defence functions. The integrated framework demonstrated here offers a credible foundation for meeting this requirement, combining empirical and certified protections with the operational continuity that production financial systems demand. As the cat-and-mouse contest between attackers and defenders continues to evolve, frameworks of this kind provide the basis on which financial institutions can sustain confidence in the machine learning systems that increasingly underpin their security posture.

REFERENCES

1. Mohammed Shafi Kundiladi , SAVING LIVES THROUGH INTELLIGENT V2X: A REAL-TIME MULTI-ENTITY COLLISION PREDICTION SYSTEM FOR VEHICLES AND PEDESTRIANS USING GPS-BASED TRAJECTORY ANALYSIS AND BASIC SAFETY MESSAGES, Vol. 52 No. 4 (2024): October-December 2024, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/196>, DOI: <https://doi.org/10.46121/pspc.52.4.10>
2. Bandari, V. V., & Lekkala, H. B. (2024). Physics-informed reinforcement learning for real-time control of complex manufacturing processes. *Power System Protection and Control*, 52(2), 164–170. <https://pspac.info/index.php/dlbh/article/view/343>
3. Bandari, V. V., & Lekkala, H. B. (2024). AI-based operator behavior monitoring and cost optimization using digital traceability in manufacturing systems. *Power System Protection and Control*, 52(1), 38–45. <https://pspac.info/index.php/dlbh/article/view/342>
4. Mayank Atreya, Navin Chhibber, Harvendra Singh, Explainable Machine Learning For Dynamic Pricing In Fast-Changing Retail Environments, 2022/4/9, Journal ,Available at SSRN 6011354, https://scholar.google.com/citations?view_op=view_citation&hl=en&user=fyViF1UAAAAJ&citation_for_view=fyViF1UAAAAJ:LkGwnXOMwfcC.
5. Venumadhav Vavilala, Shankar Balla (2024, May). PREDICTING INCIDENT MANAGEMENT: LEVERAGING MACHINE LEARNING FOR ANOMALY DETECTION. *Power System Protection and Control Scopus Q1 Journal*. PSPC. <https://pspac.info/index.php/dlbh/article/view/270>
6. Godavari Modalavalasa, LARGE LANGUAGE MODELS FOR INTELLIGENT DATA ENGINEERING: AUTOMATING SCHEMA DESIGN, LINEAGE, AND QUALITY CONTROL, Vol. 50 No. 2 (2022): April-June 2022, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/183> , DOI: <https://doi.org/10.46121/pspc.50.2.4>
7. Godavari Modalavalasa, FEDERATED LEARNING FOR ENTERPRISE CLOUD DATA ENGINEERING: ARCHITECTURE, SECURITY, AND GOVERNANCE CHALLENGES, Vol. 51 No. 2 (2023): April-June 2023, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/184>, DOI: <https://doi.org/10.46121/pspc.51.2.5>
8. Godavari Modalavalasa, AI-DRIVEN DATA GOVERNANCE: INTELLIGENT METADATA, LINEAGE, AND COMPLIANCE AUTOMATION IN CLOUD DATA PLATFORMS, Archives / Vol. 52 No. 1

- (2024): January-March 2024 /, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/182>, DOI: <https://doi.org/10.46121/pspc.52.1.3>
9. Generative AI for Automated CAD Model Generation in Aerospace Manufacturing, Jawaharbabu Jeyaraman, Feroskhan Hasenkhan, Swetha Ravipudi 9/24/2024, Los Angeles Journal of Intelligent Systems and Pattern Recognition, <https://lajispr.org/index.php/publication/article/view/19>
 10. Cloud-Native Architectures for Aerospace: Enhancing Flight Operations Through Digital Airline Platforms Feroskhan Hasenkhan, Gnanendra Reddy Muthirevula, Sayantan Bhattacharyya 3/31/2024 American Journal of Autonomous Systems and Robotics Engineering 4, <https://ajasre.org/index.php/publication/article/view/25>
 11. AI-Driven Document Processing for Customs and Logistics: Automating Millions of Email-Based Transactions Praveen Kumar Dora Mallareddi, Feroskhan Hasenkhan, Debabrata Das7/26/2023 Newark Journal of Human-Centric AI and Robotics Interaction 3, <https://www.njhcair.org/index.php/publication/article/view/13>
 12. Manikantha Varaprasad Inakollu, (2024, May), FINOPS-Driven Cloud optimization models for enterprise applications, Vol. 52 No. 2 (2024): April-June 2024, 99-110, Power System Protection and Control, ISSN-1674-3415, URL: <https://pspac.info/index.php/dlbh/article/view/283> DOI: <https://doi.org/10.46121/pspc.52.2.10>
 13. Naresh Lokiny. (2022). Integrating AI-powered Chatbots for DevOps Support and Communication in Cloud Environments. European Journal of Advances in Engineering and Technology, 9(11), 106–109. <https://doi.org/10.5281/zenodo.13325989>
 14. Naresh Lokiny, (2021), "Disaster Recovery and Business Continuity Planning in DevOps Cloud with AI", International Journal of Science and Research (IJSR), 10(3), 2024-2027. <https://dx.doi.org/10.21275/SR24724151733>, <https://www.ijsr.net/getabstract.php?paperid=SR24724151733>
 15. Naresh Lokiny, & Ranganath Nandanampati. (2020). DevSecOps: Integrating Security into DevOps with AI in Cloud. Journal of Scientific and Engineering Research, 7(10), 239–242. <https://doi.org/10.5281/zenodo.13348695>
 16. Naresh Lokiny, & Pradip Reddy. (2021). Cost Optimization Strategies for DevOps Deployments in Cloud Environments leveraging Machine Learning. European Journal of Advances in Engineering and Technology, 8(3), 69–72. <https://doi.org/10.5281/zenodo.13325845>
 17. Jayanth Para, 6G Internet Technology Cyber Threat Notification & Alert System, Vol. 52 No. 4 (2024): October-December 2024, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/170> , DOI: <https://doi.org/10.46121/pspc.52.4.9>
 18. Jayanth Para, AI Based Cloud Computation Observational Method & Process, Vol. 51 No. 4 (2023): October-December 2023, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/171> , DOI: <https://doi.org/10.46121/pspc.51.4.3>
 19. **Kaleshwar Aryasomayajula**, PREVENTING BIAS IN AI-BASED DECISION-MAKING: ANALYZING TECHNIQUES TO REMOVE UNFAIR PREJUDICE IN ALGORITHMIC OUTCOMES, Vol. 50 No. 2 (2022): April-June 2022, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/294>
 20. Kaleshwar Aryasomayajula, MODERN DEVELOPMENT PRACTICES: INVESTIGATIONS INTO SERVERLESS COMPUTING, LOW-CODE/NO-CODE PLATFORMS, AND DEVELOPER

- PRODUCTIVITY IN REMOTE ENVIRONMENTS, Vol. 50 No. 4 (2022): October-December 2022, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/364>
21. **Kaleshwar Aryasomayajula**, Micro services: “A Comparative Analysis of Monolithic vs. Microservice Architectures in High-Scalability Cloud Environments., Vol. 51 No. 2 (2023): April-June 2023 , Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/374>
 22. **Vikas Suresh Bagora**, KERNEL-LEVEL PERFORMANCE OPTIMIZATION FOR CARRIER-GRADE BROADBAND AND HIGH-SPEED NETWORKING SYSTEMS, Vol. 47 No. 1 (2019), Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/359>
 23. **Vikas Suresh Bagora**, SCALABLE 128G FIBRE CHANNEL AND ETHERNET CONVERGENCE FOR NEXT-GENERATION DATA CENTER NETWORKS, Vol. 50 No. 4 (2022): October-December 2022, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/362>
 24. **Vikas Suresh Bagora**, SECURE EMBEDDED NETWORKING ARCHITECTURES USING TRUSTED EXECUTION ENVIRONMENTS IN BROADBAND GATEWAYS, Vol. 52 No. 2 (2024): April-June 2024, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/363>
 25. Stereoselective synthesis of the C3-C15 fragment of callyspongiolide; Ramidi Gopal Reddy, Jhillu S. Yadav and Debendra K. Mohapatra. (Tetrahedron Letters. 2018, 59, 3579-3582). <https://doi.org/10.1016/j.tetlet.2018.08.041>, <https://www.sciencedirect.com/science/article/pii/S0040403918310426>
 26. Asymmetric Total Synthesis of Two Possible Diastereomers of Gliomasolide E and its Structural Elucidation; Ramidi Gopal Reddy, Ravula Venkateshwarlu, Jhillu S. Yadav and Debendra K. Mohapatra. (J. Org. Chem. 2017, 82(2), 1053-1063). <https://doi.org/10.1021/acs.joc.6b02611> <https://pubs.acs.org/doi/abs/10.1021/acs.joc.6b02611>
 27. Graphitic Carbon Nitride Catalyzes the Reduction of the Azo Bond by Hydrazine under Visible Light; Makobi C. Okolie, Glory G. Ollordaa, Gopal R. Ramidi, Xin Yan, Yufeng Quan, Qingsheng Wang and Yingchun Li. (Nanomaterials 2024, 14(17), 1402), <https://doi.org/10.3390/nano14171402>, <https://www.mdpi.com/2079-4991/14/17/1402>
 28. **Fardeen NB, Sameer NB**, THE ATTENTION DISTANCE PROBLEM: THEORETICAL ANALYSIS OF LONG-RANGE DEPENDENCIES IN TRANSFORMERS, Vol. 52 No. 2 (2024): April-June 2024, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/313>
 29. **Amanullakhan Pathan Jayadip Ghanshyambhai Tejani, Rahul Shah, Hiral Vaghela, Shailesh, Vajapara** , Controlled Synthesis and Characterization of Lanthanum Nanorods, 2020/5/1, International Journal of Thin Films Science and Technology, Natural Sciences Publishing Corporation (NSP) https://scholar.google.com/citations?view_op=view_citation&hl=en&user=efbNMkwAAAAJ&citation_for_view=efbNMkwAAAAJ:M3NEmzRMikIC
 30. Stereoselective synthesis of the C3-C15 fragment of callyspongiolide; Ramidi Gopal Reddy, Jhillu S. Yadav and Debendra K. Mohapatra. (Tetrahedron Letters. 2018, 59, 3579-3582).

- <https://doi.org/10.1016/j.tetlet.2018.08.041>,
<https://www.sciencedirect.com/science/article/pii/S0040403918310426>
31. Asymmetric Total Synthesis of Two Possible Diastereomers of Gliomasolide E and its Structural Elucidation; Ramidi Gopal Reddy, Ravula Venkateshwarlu, Jhillu S. Yadav and Debendra K. Mohapatra. (J. Org. Chem. 2017, 82(2), 1053-1063). <https://doi.org/10.1021/acs.joc.6b02611>,
<https://pubs.acs.org/doi/abs/10.1021/acs.joc.6b02611>
 32. Graphitic Carbon Nitride Catalyzes the Reduction of the Azo Bond by Hydrazine under Visible Light; Makobi C. Okolie, Glory G. Ollordaa, Gopal R. Ramidi, Xin Yan, Yufeng Quan, Qingsheng Wang and Yingchun Li. (Nanomaterials 2024, 14(17), 1402),
<https://doi.org/10.3390/nano14171402>, <https://www.mdpi.com/2079-4991/14/17/1402>
 33. **Kaleshwar Aryasomayajula**, PREVENTING BIAS IN AI-BASED DECISION-MAKING: ANALYZING TECHNIQUES TO REMOVE UNFAIR PREJUDICE IN ALGORITHMIC OUTCOMES, Vol. 50 No. 2 (2022): April-June 2022, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/294>
 34. Kaleshwar Aryasomayajula, MODERN DEVELOPMENT PRACTICES: INVESTIGATIONS INTO SERVERLESS COMPUTING, LOW-CODE/NO-CODE PLATFORMS, AND DEVELOPER PRODUCTIVITY IN REMOTE ENVIRONMENTS, Vol. 50 No. 4 (2022): October-December 2022, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/364>
 35. **Kaleshwar Aryasomayajula**, Micro services: “A Comparative Analysis of Monolithic vs. Microservice Architectures in High-Scalability Cloud Environments.”, Vol. 51 No. 2 (2023): April-June 2023, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/374>
 36. Controlled Synthesis and Characterization of Lanthanum Nanorods, **Author(s), Jayadeep Tejani, Rahul Shah, Hiral Vaghela, Shailesh Vajapara, Amanullakhan Pathan**, doi:10.18576/ijfst/090205, URL: <https://www.naturalspublishing.com/Article.asp?ArtcID=20705>, May-2020
 37. **SUDIPKUMAR GHANVAT, AISHWARYA BADLANI, SHREYA SANDEEP JOSHI**, MARKETING EFFECTIVENESS IN THE POST-COOKIE ERA: A COMPARATIVE ANALYSIS OF FIRST PARTY AND ZERO-PARTY DATA STRATEGIES ON CAMPAIGN PERFORMANCE AND CUSTOMER EQUITY, **VOL. 31 No. 4 (2023): JOCAAA**, <https://www.eudoxuspress.com/index.php/pub/article/view/3940>