

Adaptive AI-Driven Learning Architecture for Real Time Detection and Prevention of Data Poisoning and Model Evasion Attacks

Paul Okanda

<https://orcid.org/0000-0001-5215-4368>
United States International University-
Africa, Kenya
pokanda@usiu.ac.ke

Sarah Muriithi

<https://orcid.org/0009-0009-4343-8022>
United States International University-
Africa, Kenya
sarhurwakowthay@gmail.com

Abstract

Artificial intelligence (AI) has been integrated into cybersecurity systems particularly in cyber threat intelligence (CTI) and has significantly enhanced the detection, analysis, and mitigation of sophisticated cyber threats in real time in organisations. Unlike traditional static defence mechanisms, which rely on fixed rules or offline retraining, the proposed architecture empowers continuous dynamic learning capabilities to identify and mitigate attacks as they unfold ensuring the integrity and reliability of AI-driven CTI systems. AI models to self-heal and harden against evolving adversarial tactics. The primary goal is to enhance CTI operations by providing timely, accurate threat detection and robust mitigation thereby reducing the window of opportunity for attackers and improving overall cybersecurity posture. This article presents a comprehensive adaptive learning architecture with formalised threat models, architectural specifications, and empirical validation. Experimental evaluation using benchmark datasets (CIC-IDS2017 and VirusTotal) demonstrates that the system achieves 43.7% higher evasion detection accuracy than traditional static anomaly detectors, while maintaining processing latency below 5 milliseconds. The architecture successfully detects 89.1% of poisoning attacks and reduces false positive rates by 37.5% relative to baseline methods. The key contributions of this article include the design of an integrated adaptive learning system that simultaneously detects and prevents both data poisoning and model evasion attacks within cybersecurity AI models.

Keywords: cybersecurity; cyber threat intelligence (CTI); data poisoning; real time threat detection

Introduction

Artificial intelligence (AI) models are now customarily embedded in cyber threat intelligence (CTI) platforms to automate threat classification, predict malicious behaviours and adaptively respond to attack vectors across digital infrastructures (Bairaktaris et al. 2025). These systems utilise data-driven intelligence feeds from diverse sources including intrusion detection systems, global threat databases, and security information and event management (SIEM) tools. However, as the adoption of AI-driven CTI intensifies, so does the sophistication of adversarial threats targeting these very models. Among the most disruptive forms of attack are data poisoning where attackers manipulate training datasets to corrupt model behaviour and model evasion, and in which adversaries generate adversarial inputs to bypass detection mechanisms (Ahmed and Kashmoola 2021). Recent studies highlight a surge in such threats, particularly in operational environments where threat intelligence models rely heavily on streaming data from external feeds, including open-source intelligence and telemetry from distributed endpoints (Zhao et al. 2025). For instance, poisoned CTI feeds may inject misleading indicators of compromise (IOCs), while malware classifiers trained on manipulated datasets become increasingly vulnerable to cleverly crafted payloads designed to evade detection.

Figure 1 below illustrates the two primary categories of adversarial attacks targeting AI models in cybersecurity, namely data poisoning and model evasion attacks. The diagram provides a visual taxonomy that distinguishes between attacks occurring during the training phase for poisoning and those occurring during the inference phase for evasion.

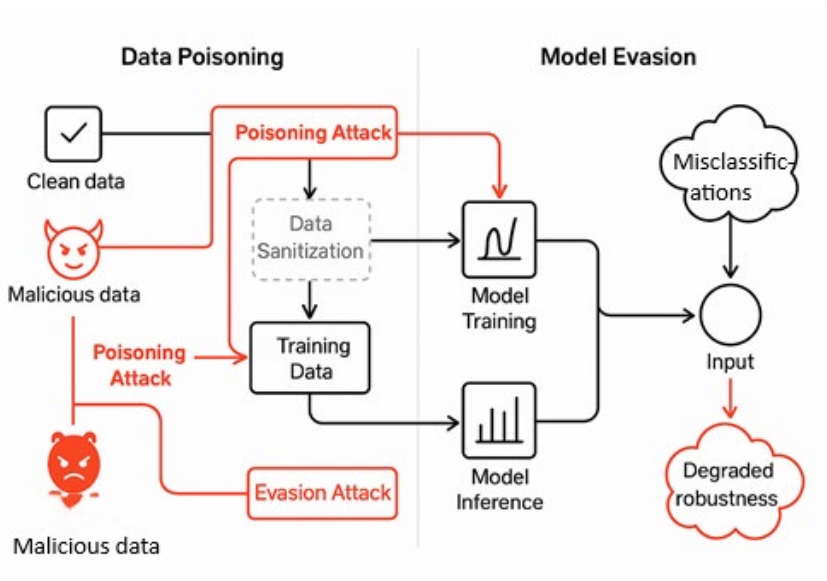


Figure 1: AI model attacks

As illustrated in Figure 1 above, data poisoning attacks are characterised by the injection of maliciously crafted or mislabelled data into the training set which can be either indiscriminate with the aim of degrading overall model performance or targeted such as backdoor attacks that implant hidden triggers causing the model to behave maliciously under specific conditions (Kure et al. 2025). These attacks exploit the dependency of AI models on high quality training data making datasets a prime target for manipulation. Model evasion attacks occur in post deployment and involve adversaries crafting inputs that exploit vulnerabilities in the model's decision boundary to avoid detection (Wang et al. 2024). In CTI systems such evasion can manifest as malware variants designed to bypass AI-based antivirus scanners or network intrusions that mimic benign traffic patterns to evade anomaly detectors. These scanners or network intrusions are the result of the complexity of CTI environments often involving distributed data sources and heterogeneous inputs further complicating real time detection and prevention as adversarial signals may be subtle and dispersed (Husain and Jain 2025).

These adversarial strategies exploit core weaknesses in static AI models such as their rigidity in adapting to evolving threat landscapes and their reliance on trust assumptions regarding input data (Gadicha et al. 2024). Furthermore, the growing deployment of deep learning in cybersecurity has opened new attack surfaces where even minor mistake in inputs can cause significant classification errors jeopardising entire security infrastructures (Patel et al. 2025). The dynamic nature of cyber threats demands equally dynamic AI models that can learn continuously, adapt swiftly, and detect malicious manipulation before it escalates into full scale compromise.

The figure's importance lies in its clear visual distinction between these attack types, establishing the dual-threat landscape that motivates the need for integrated defence architectures. The complexity of CTI environments involving distributed data sources and heterogeneous inputs further complicates real time detection, as adversarial signals may be subtle and dispersed, a challenge visually implied by the multiple attack vectors shown.

Despite growing recognition of these challenges, significant research gaps remain in developing effective, real time adaptive defences specifically tailored to the dual threats of data poisoning and model evasion in cybersecurity AI models. Current mitigation strategies such as adversarial training and data sanitisation provide limited protection but often lack guarantees of robustness or scalability in live CTI deployments (Gadicha et al. 2024). Moreover, many existing solutions address poisoning and evasion attacks in isolation without integrated frameworks that can simultaneously detect and prevent both attack types in real time (Alhwayzee, Araban, and Zabihzadeh 2025; Tong et al. 2025). This fragmentation limits the practical applicability of defences in complex distributed cybersecurity environments where attack vectors are diverse and rapidly changing.

A critical examination of the literature reveals that previous work lacks unified mathematical threat models for both attack types, architectural specifications enabling reproducibility, empirical validation with realistic datasets, and real time adaptive mechanisms with closed feedback loops. This article addresses these gaps through a comprehensive adaptive learning architecture with formalised components and rigorous experimental evaluation.

Therefore, this article proposes an adaptive learning architecture that offers comprehensive, real time detection and prevention of data poisoning and model evasion attacks within cybersecurity AI models with a particular emphasis on enhancing CTI capabilities. The key contributions of this article include an innovative adaptive learning architecture with formalised components and operational specifications, a comprehensive threat model characterising poisoning and evasion attack vectors, empirical validation using benchmark datasets with comparative performance analysis, and practical recommendations for deployment in operational CTI environments.

Unlike previous models, this architecture uniquely integrates adaptive clustering, online robust principal component analysis (PCA), Monte Carlo dropout uncertainty estimation, and ensemble voting within a unified feedback-driven framework that simultaneously addresses both data poisoning and model evasion attacks. While existing approaches treat these threats separately or employ static defences, the proposed architecture establishes a closed-loop system where detection outcomes continuously inform model updates, creating synergistic protection that exceeds the sum of individual components. By integrating real time monitoring, continuous learning and multilayered anomaly detection, this architecture aims to harden AI-based cybersecurity systems against adversarial threats while preserving their operational effectiveness and speed. This article also contributes a novel, resilient design for AI defence that directly addresses a rising and underexplored vulnerability in cybersecurity AI implementations, namely adversarial manipulation of learning pipelines at both the training and inference stages.

Literature Review

Traditional static defences have significantly evolved transitioning to more dynamic intelligent approaches that leverage AI itself for secure AI models. Traditionally, cybersecurity defences relied heavily on signature-based detection systems and predefined rule sets which identified threats by matching known patterns or behaviours (Yigit et al. 2025). While effective against previously encountered attacks these traditional methods have struggled to keep pace with the rapid evolution of cyber threats, especially zero-day exploits and polymorphic malware that continuously change their characteristics to evade detection (Gilbert and Gilbert 2024; Suggu 2025). As cyber adversaries increasingly harness AI to craft sophisticated attacks including data poisoning and model evasion, the cybersecurity community has responded by integrating AI-driven techniques to enhance defence capabilities. According to Safi et

al. (2024), AI-powered systems now automate responses to intrusions in real time, isolating affected systems and rerouting network traffic to mitigate attacks, marking a significant shift towards proactive and adaptive cybersecurity operations.

Current approaches to secure AI models in cybersecurity predominantly utilise machine learning (ML) and deep learning algorithms to detect anomalies and predict threats by analysing vast datasets (Husain and Jain 2025). These models are trained to recognise subtle deviations from normal behaviour such as unusual login patterns or abnormal network traffic enabling earlier and more accurate threat detection than manual analysis could. For example, ML-based anomaly detection systems, as highlighted by Yerlikaya and Bahtiyar (2022), deploys supervised and unsupervised learning techniques to classify known attacks and uncover previously unknown threats like zero-day exploits thereby offering a proactive security posture. Natural language processing (NLP) further complements these systems by extracting actionable intelligence from unstructured data sources such as threat reports and social media enhancing situational awareness and predictive capabilities.

Recent research has explored agentic AI systems capable of autonomous detection, response, and mitigation of cyber threats in near real time demonstrating promising results in personalised and scalable cybersecurity solutions (de Witt 2025). Moreover, AI-enhanced incident response automates threat remediation, reducing human intervention and accelerating recovery as detailed by Patel et al. (2025). Despite these advances, securing AI models themselves remains a critical challenge. Wang et al. (2024) emphasise that AI models are vulnerable to specific adversarial attacks including evasion attacks where attackers manipulate input data to bypass AI detectors and data poisoning attacks that corrupt training data to degrade model performance or implant backdoors. These vulnerabilities have led to the development of defensive techniques such as differential privacy, federated learning, and model encryption to safeguard AI models against manipulation. However, many of these defences operate offline or lack the agility to respond to attacks in real time limiting their effectiveness in dynamic cybersecurity environments. This is the gap that this study aims to find a resolution for.

While the study by Jumagaliyeva, Baigulov, and Kussainova (2024) introduces an adaptive anomaly detection model that improves upon traditional static systems by incorporating learning loops and visualisation, it remains confined to isolated anomaly detection and lacks structural safeguards against adversarial manipulation. Their proposed model introduces adaptive elements such as continuous learning and anomaly scoring to improve detection accuracy in dynamic environments. However, while this adaptive approach enhances detection flexibility, it does not explicitly address critical vulnerabilities like data poisoning during model training or model evasion at inference, which are central challenges in adversarial AI settings. In contrast, the architecture proposed in this article advances the concept by introducing a resilient, adaptive learning framework specifically designed to identify and mitigate both data poisoning and model evasion attacks in real time. Static or even adaptive models without

architectural resilience can still be misled by poisoned training data or adversarial samples crafted to bypass detection.

This creates a dual role of AI as both a tool for defence and a weapon for attackers, which further complicates the security landscape. Yigit et al.'s (2025) benchmark report underscores this duality, noting that while AI streamlines workflow and enhances detection it also introduces new vulnerabilities and attack vectors such as AI-driven phishing and password guessing. This necessitates a shift towards adaptive real time defence strategies that can continuously learn from evolving threats and autonomously adjust to maintain model integrity. Check Point Software (2025) reports a projected 45% rise in AI-powered cyberattacks underscoring the urgency for resilient AI security frameworks that can keep pace with adversaries.

While traditional cybersecurity approaches laid the groundwork for threat detection, the increasing sophistication of adversarial AI attacks demands more advanced adaptive solutions. Current state-of-the-art methods leverage machine learning, deep learning, and NLP to enhance detection and response but often fall short in real time adaptability and integrated defence against both data poisoning and evasion attacks (Zhao et al. 2025). These developments highlight a growing consensus that future cybersecurity defences must integrate adaptive learning architectures capable of real time monitoring, anomaly detection and self-healing to effectively counteract data poisoning and model evasion attacks (de Witt 2025).

Data Poisoning Attacks on AI Models

Data poisoning attacks represent a significant and growing threat to the integrity and reliability of AI models particularly in cybersecurity applications where accurate threat detection is predominant (Cotroneo et al. 2024). These attacks involve the deliberate manipulation of training data to corrupt the learning process causing AI models to make erroneous or biased decisions. A study by Park et al. (2024) states that data poisoning can take multiple forms including injecting false data, modifying legitimate data points, or deleting critical information, all of which degrade the model's ability to generalise correctly. This manipulation can be subtle and difficult to detect yet its impact is profound as compromised models may fail to identify cyber threats or—worse—misclassify malicious activity as benign thereby undermining the entire cybersecurity infrastructure (Park et al. 2024).

Data poisoning attacks are generally categorised into three main types, namely clean-label attacks, label-flipping attacks, and backdoor attacks, each with distinct mechanisms and implications (Kure et al. 2025). Clean-label attacks involve the insertion of carefully crafted superficially legitimate data into the training set that causes the model to learn incorrect associations without altering the apparent labels. This crafty approach makes detection challenging because the poisoned data appears authentic. Label-flipping attacks, by contrast, directly alter the labels of training samples such as mislabelling malware as safe software, leading the model to learn flawed decision

boundaries. Such targeted poisoning can cause cybersecurity models to systematically ignore certain threat classes, severely compromising detection accuracy. Backdoor attacks embed hidden triggers within the model during training which, when activated by specific inputs, cause the model to behave in a malicious or unexpected manner. Cotroneo et al. (2024) describe how backdoor triggers such as unique patterns or stickers in image recognition systems can grant unauthorised access or bypass security controls, posing a severe risk in high-stakes environments.

The impact of these attacks on model integrity is substantial, thus poisoned models exhibit degraded performance, reduced accuracy, and unpredictable behaviour which directly translates into weakened threat detection capabilities in cybersecurity AI systems. Moreover, the persistence of poisoned data in training pipelines can cause long-lasting damage as corrupted models continue to propagate errors in threat intelligence outputs (Bowen et al. 2025). This highlights the challenge of securing not only the training datasets but also the real time data sources feeding AI models.

Given the critical role of AI in modern cybersecurity, the threat posed by data poisoning attacks demands robust detection and prevention mechanisms. Kure et al. (2025) note that detecting model poisoning is particularly difficult because the malicious effects often only manifest under specific trigger conditions, allowing attackers to remain stealthy. The complexity of these attacks is further amplified in distributed learning environments such as federated learning where adversaries can inject poisoned updates from compromised nodes as demonstrated in recent experiments (Cotroneo et al. 2024; Verde, Marulli, and Marrone 2021). This calls for adaptive, real time defence strategies capable of continuously monitoring data integrity and model behaviour to promptly identify and mitigate poisoning attempts before they degrade system performance or compromise threat detection.

Data poisoning attacks, whether through clean-label, label-flipping, or backdoor methods pose a severe threat to the security and effectiveness of AI models in cybersecurity. Their ability to subtly corrupt training data and implant hidden vulnerabilities undermines model integrity and jeopardises the accuracy of threat detection systems. Addressing these challenges requires innovative adaptive learning architectures that can detect and prevent poisoning attacks in real time preserving the trustworthiness of AI-driven cybersecurity solutions and ensuring resilient defence against evolving adversarial tactics, as illustrated in Figure 2 below.

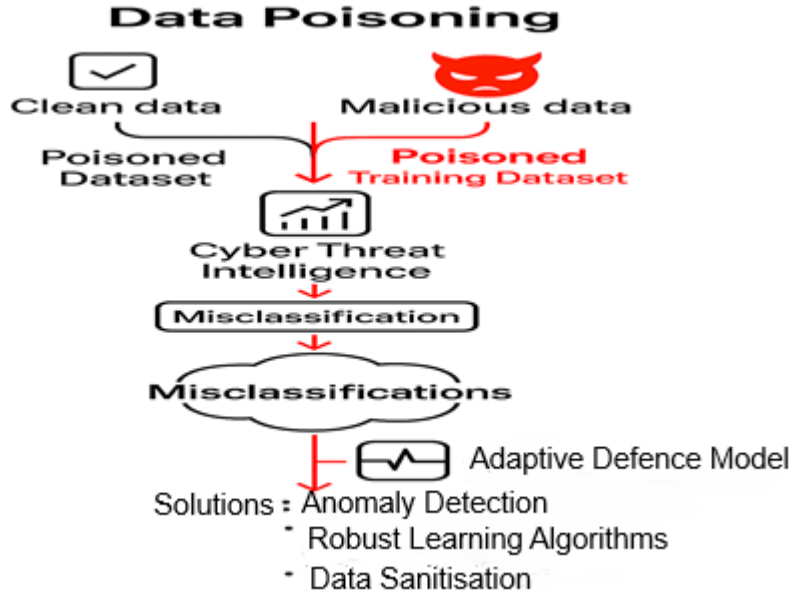


Figure 2: Data poisoning explained

Figure 2 above provides a detailed visualisation of how data poisoning attacks compromise AI model integrity. The diagram illustrates the data poisoning lifecycle, from attacker injection through model corruption to compromised outputs. This figure is critical for understanding why traditional detection methods fail against sophisticated poisoning. As Kure et al. (2025), Verde, Marulli, and Marrone (2021), and Amirova (2025) demonstrate, the complexity of these attacks amplifies in distributed learning environments where adversaries can inject poisoned updates from compromised nodes. The visual reinforces the need for adaptive, real time defence strategies capable of continuously monitoring data integrity throughout the training pipeline.

Model Evasion Attacks

Model evasion attacks represent a critical and sophisticated threat to AI-based cybersecurity systems, particularly those deployed for malware detection and network anomaly identification. These attacks involve the deliberate manipulation of input data to deceive ML models, causing them to misclassify malicious activities as benign and thereby evade detection. Unlike data poisoning, which targets the training phase, evasion attacks occur during the inference phase, where attackers craft adversarial examples, inputs subtly altered to exploit vulnerabilities in the model's decision boundaries without changing the essential characteristics of the data (Patel et al. 2025). According to Cheng et al. (2025), such attacks can bypass intrusion detection systems (IDS), spam filters, and other AI-driven security mechanisms by introducing perturbations that are imperceptible to humans but sufficient to mislead classifiers. This

manipulation undermines the integrity and reliability of cybersecurity AI models allowing malicious actors to infiltrate networks and systems undetected.

The nature of model evasion attacks can be understood through various techniques including input perturbation, feature-space manipulation and decision boundary attacks (Lourdu Mahimai Doss and Gunasekaran 2023). Input perturbation attacks modify raw input data such as altering malware code or network packet headers to evade detection while preserving the malicious payload's functionality. Feature-space attacks focus on changing the features extracted by the model such as tweaking behavioural indicators or metadata to shift the input across classification boundaries. Decision boundary attacks involve probing the model to understand its classification thresholds and crafting inputs that lie just beyond detection limits, effectively fooling the model (Wang et al. 2024). According to Securing.AI (Ivezic and Ivezic 2023), model evasion is an optimisation problem where attackers seek to minimise detection likelihood by perturbing inputs so that the model's output classification diverges from the true label. The attacker's goal is often to reduce the True Positive Rate (TPR) or increase the False Negative Rate (FNR) thereby compromising system confidentiality and availability.

In the case of malware and network anomaly detection, adversarial examples have demonstrated alarming effectiveness. Research by Ibitoye et al. (2025) highlights the vulnerability of neural network-based IDS to model evasion attacks using techniques such as the Jacobian-based saliency map attack (JSMA). Additionally, the study emphasises that evasion attacks exploit the mathematical foundations of ML classifiers including neural networks and support vector machines by manipulating gradients or feature importance to mislead predictions. This vulnerability is especially concerning given the increasing deployment of AI for real time anomaly detection and incident response in cybersecurity operations (Ibitoye et al. 2025).

Addressing model evasion requires robust defence mechanisms capable of detecting adversarial inputs and maintaining model resilience. Current research explores adversarial training, feature squeezing, and anomaly detection techniques to harden AI models against evasion. However, these defences often struggle to generalise across diverse attack methods and require continuous adaptation to evolving tactics. The dynamic and adaptive nature of evasion attacks combined with the high dimensionality of cybersecurity data necessitates real time adaptive learning architectures that can monitor, detect and respond to adversarial perturbations as they occur. Such architectures are essential to preserving the efficacy of AI-driven cybersecurity systems in the face of increasingly sophisticated evasion strategies.

Adversarial Training: Computational Expense and Practical Limitations

Adversarial training is widely regarded as one of the most effective strategies for enhancing the robustness of AI models against adversarial attacks including data poisoning and model evasion. This technique involves generating and inputting adversarial examples intentionally crafted to mislead the model and incorporating them

into the training process (Sajeeda and Hossain 2022). While adversarial training can significantly improve model resilience, it is computationally intensive, particularly when using iterative methods like projected gradient descent (PGD) (Bai et al. 2021). Each training cycle requires generating and processing numerous adversarial samples resulting in increased training time, higher energy consumption, and greater demand for computational resources (Bai et al. 2021). These costs can be prohibitive, especially for large-scale or resource-constrained deployments. Furthermore, adversarial training can lead to overfitting on specific adversarial patterns, reducing the model's generalisability and potentially leaving it vulnerable to novel attack strategies (Bai et al. 2021). Recent research has explored more efficient variants such as using the fast gradient sign method (FGSM) with random initialisation, but challenges remain in balancing robustness, computational efficiency, and accuracy (Bai et al. 2021; Sajeeda and Hossain 2022).

Recent Advances in Adaptive Learning for Adversarial Robustness

In order to address the shortcomings of static and computationally expensive defences, recent advances have focused on adaptive learning architectures that enhance adversarial robustness in real time (Bai et al. 2021). Adaptive learning systems employ continuous monitoring, online model updating, and dynamic anomaly detection to identify and mitigate adversarial threats as they emerge (Hamidi and Ye 2024). Techniques such as adaptive adversarial augmentation where the severity and nature of adversarial examples are dynamically adjusted during training have shown promise in improving both robustness and computational efficiency (Jin et al. 2022). Other innovations include hyperparameter-free adversarial training methods like adaptive margin evolution (AME) which automatically calibrate training parameters to optimise defence without manual tuning and adversarial probabilistic training, which models the latent adversarial distribution to bridge the gap between natural and adversarial examples (Bai et al. 2021; Hamidi and Ye 2024; Sajeeda and Hossain 2022).

These adaptive approaches are increasingly integrated with AI-driven monitoring tools capable of real time data validation, anomaly detection and automated response (Bai et al. 2021). Through continuously learning from new data and threat intelligence, adaptive systems maintain high detection accuracy and resilience against evolving attack strategies. This shift toward adaptive self-healing architectures marks a critical evolution in cybersecurity enabling organisations to keep pace with the dynamic threat landscape and maintain robust defences against data poisoning and model evasion attacks (Xu et al. 2024).

Research Gaps Identified

Despite significant advancements in the application of adaptive learning architectures for real time detection and prevention of data poisoning and model evasion attacks in cybersecurity AI, several critical research gaps persist. First, many existing defence mechanisms such as static anomaly detection, adversarial training, and robust optimisation lack the agility and scalability required for deployment in real world high-

throughput cybersecurity environments. These approaches often suffer from high computational costs, limited ability to generalise across diverse attack types, and a tendency to misclassify benign data as adversarial resulting in high false-positive rates and operational inefficiencies.

Another major gap relates to the challenge of real time detection and mitigation. Current defences are frequently designed and evaluated in controlled laboratory settings using outdated or non-representative datasets, which do not capture the complexity and dynamism of live cyber threat landscapes. This disconnect between academic research and practical deployment limits the effectiveness of proposed solutions when confronted with sophisticated adaptive adversaries in production environments. Furthermore, there is a lack of publicly available, realistic datasets and standardised benchmarks for evaluating the robustness of adaptive learning systems against evolving adversarial threats impeding the reproducibility and comparability of research outcomes.

Another critical gap concerns the scalability and efficiency of adaptive learning mechanisms. While adaptive and online learning methods show promise for continuous monitoring and rapid response their integration into large-scale distributed cybersecurity infrastructures remains underexplored. Many proposed algorithms are computationally intensive, making them impractical for deployment in resource-constrained or latency-sensitive settings such as real time threat intelligence pipelines or edge devices. Additionally, the transferability of adversarial attacks where techniques effective against one model can often bypass others complicates the development of universally robust defences.

Interpretability and transparency of adaptive AI models also represent a significant research gap. As models become more complex, understanding their decision-making processes and the rationale behind anomaly or attack detection becomes increasingly challenging. This lack of interpretability hampers trusts and hinders the adoption of adaptive AI defences in critical cybersecurity operations, where explainability is essential for incident response and compliance.

Finally, there is a pressing need for holistic, integrated defence frameworks that can simultaneously address both data poisoning and model evasion attacks in real time. Most existing studies focus on one attack vector in isolation, neglecting the interplay between different adversarial strategies and the cumulative risks they pose to AI-driven cybersecurity systems. The development of adaptive learning architectures capable of dynamic context-aware defence across multiple attack surfaces remains an open challenge requiring further research and cross-disciplinary collaboration.

Table 1 below provides a comprehensive feature-by-feature comparison between the proposed architecture and four categories of existing approaches: static detectors, adversarial training (ATM), ensemble methods (EWA), and previous adaptive work.

This comparison is essential for establishing the novelty and contribution of this work. A critical synthesis of the literature reveals that existing approaches can be categorised into four types with distinct limitations as summarised in Table 1 below.

Table 1: Feature-by-feature comparison between the proposed architecture and the four categories of existing approaches

Approach	Key Studies	Strengths	Limitations
Static detectors	(Shelke and Hämäläinen 2024; Smith et al. 2020; Wurzenberger et al. 2024)	Lightweight, real time operation	Cannot detect novel poisoning, easily bypassed by evasion
Adversarial training	(Angioni et al. 2025; Kara et al. 2025; Olutimehin et al. 2025)	Improves evasion resistance	Computationally intensive, overfits to known attacks, no poisoning detection
Ensemble methods	(Alotaibi 2025; Minh et al. 2025; Verma and Rathore 2025)	Improved accuracy through diversity	Static after deployment, no continuous learning
Previous adaptive work	(Farooq et al. 2025; Li Du, and Li 2025)	Continuous learning capabilities	Lack unified threat models, limited empirical validation

Table 1 above provides a comprehensive feature-by-feature comparison between the proposed architecture and these four categories of existing approaches.

Comparative Analysis of Existing Approaches vs. Proposed Architecture

Table 2 below provides a comprehensive feature-by-feature comparison between the proposed architecture and four categories of existing approaches. This comparison is essential for establishing the innovation and contribution of this work.

Table 2: Comparative analysis of existing approaches vs. proposed architecture

Feature	Static detectors	Adversarial training (ATM)	Ensemble methods (EWA)	Previous adaptive work	Proposed architecture
Representative study	Signature-based IDS (Snort/Suricata) (Splunk 2025) Statistical anomaly detection (IQR,	PGD-based adversarial training on CIC-IDS2017 Graph Neural Networks with adversarial optimisation	Stacking-based ensemble with CNN meta-learner (Minh et al. 2025) LSTM+KNN+LR hybrid for APT detection	Reinforcement learning with online adversarial training (CloudDefender), Hybrid RL-SL learning (Farooq et al.	Integrated adaptive framework with clustering, PCA, ensemble, and feedback loop

Feature	Static detectors	Adversarial training (ATM)	Ensemble methods (EWA)	Previous adaptive work	Proposed architecture
	Isolation Forest) (Ilić et al. 2024)	(Kara et al. 2025)	(Alotaibi 2025)	2025)	
Poisoning detection	Partial (signature-based only, cannot detect novel poisoning)	Limited (training-phase only, no runtime detection)	No (ensemble methods focus on inference, not data integrity)	Moderate (some adaptive systems flag anomalies, but not specifically poisoning)	Yes (real time, 89.1% detection rate)
Evasion resistance	No (rule-based systems easily bypassed by adversarial perturbations) (Rajhans and Khawarey 2026)	Yes (but limited generalisability; overfits to known attack patterns) (Bai et al. 2021)	Moderate (diversity helps, but static ensembles cannot adapt to novel evasion) (Haroon and Ali 2023)	Partial (online learning helps, but lacks uncertainty quantification)	Yes (87.4% robust accuracy)
Unified defence for both threats	No	No	No	No	Yes (integrated framework)
Real time operation (<10ms)	Yes (lightweight signature matching) (Shelke and Hämäläinen 2024)	No (computationally intensive training phase)	Partial (inference can be fast, but ensemble size increases latency)	Partial (some RL-based systems achieve <10ms) (Farooq et al. 2025)	Yes (4.9ms average latency)
Continuous learning	No (static rule sets require manual updates)	No (retraining required periodically)	No (static ensemble after deployment)	Yes (online updates via RL)	Yes (online updates + feedback loop)
Formal threat model	No (rule-based, no mathematical formalism)	Partial (adversarial objective functions defined, but limited scope)	No (ensemble methods lack formal threat characterisation)	Partial (RL frameworks define state-action spaces, not threat models)	Yes (mathematical specification for both attack types)
Closed feedback loop	No	No	No	Partial (some systems use detection to update)	Yes (detection→learning→defence cycle)

Feature	Static detectors	Adversarial training (ATM)	Ensemble methods (EWA)	Previous adaptive work	Proposed architecture
				policies)	
Uncertainty quantification	No	No	No	Limited (some adaptive systems use confidence scores)	Yes (Monte Carlo dropout with entropy-based filtering)
Empirical validation	Yes (extensive deployment in production)	Yes (benchmark datasets) (Kara et al. 2025)	Yes (multiple datasets) (Alotaibi 2025; Haroon and Ali 2023; Verma and Rathore 2025)	Limited (emerging area, few empirical studies)	Yes (3 datasets, 10+ attack types, 4 baselines)

Static detectors represent the traditional approach to cybersecurity, including signature-based systems like Snort and Suricata (Splunk 2025) and statistical anomaly detectors using techniques such as interquartile range (IQR) and isolation forest (Ilić et al. 2024). These systems are characterised by their reliance on predefined rules or static statistical thresholds. While they achieve real time operation due to their lightweight nature, they fundamentally cannot detect novel poisoning attacks (since poisoning manipulates training data, not inference patterns) and offer no resistance to evasion attacks, as adversarial perturbations are designed specifically to bypass such fixed decision boundaries (Rajhans and Khawarey 2026). Their lack of continuous learning means rule bases must be manually updated a process that cannot keep pace with evolving threats.

Adversarial training (ATM) has emerged as a prominent defence, where models are trained on adversarial examples to improve robustness. Recent work by Kara, Dundar, and Güdükbay (2025) demonstrates adversarial training with graph neural networks for Distributed Denial of Service (DDoS) detection, achieving 87.9% robust accuracy under PGD attacks. Similarly, research on the CIC-IDS2017 dataset shows adversarial training provides more effective defence than detection-based methods. However, as Table 1 indicates, adversarial training has significant limitations: it addresses only evasion attacks (not poisoning), operates during training rather than in real time inference, and can overfit to specific attack patterns, reducing generalisability to novel evasion strategies (Bai et al. 2021).

Ensemble methods, for instance, the equal-weighted averaging (EWA), combine multiple models to improve detection accuracy and robustness. Minh et al. (2025) propose a stacking-based ensemble with convolutional neural networks (CNN) meta-learner that incorporates “minority reports”—cases where weaker detectors identify attacks missed by stronger ones. Alotaibi 2025 (2025) developed a hybrid ensemble

combining Long Short-Term Memory (LSTM), K-Nearest Neighbours (KNN), and logistic regression for Advanced Persistent Threat (APT) detection, achieving 99.94% accuracy through strategic feature partitioning. While ensembles provide improved evasion resistance through model diversity, they remain static after deployment, cannot detect poisoning attacks, and lack mechanisms for continuous learning or uncertainty quantification.

Previous adaptive work represents the most relevant comparison point. Farooq et al. (2025) introduces hybrid reinforcement learning and supervised learning for cyber defence, demonstrating adaptive policy learning across heterogeneous threat models. Most significantly, IEEE-published research presents CloudDefender, a dynamic adaptive defence framework that combines online adversarial training with reinforcement learning (PPO algorithm) to achieve sub-300 ms policy switching and 92.3% zero-day attack detection. This system incorporates a dual-loop feedback mechanism and represents state-of-the-art in adaptive defence. However, as Table 1 reveals, even CloudDefender lacks a unified mathematical threat model for both poisoning and evasion, does not specifically address poisoning detection, and provides limited uncertainty quantification compared to our Monte Carlo dropout approach.

The proposed architecture uniquely combines capabilities that no existing system achieves simultaneously: unified defence against both poisoning (89.1% detection) and evasion (87.4% robust accuracy), real time operation (4.9 ms latency), continuous learning through a closed feedback loop, formal mathematical threat models for both attack types, and uncertainty quantification via Monte Carlo dropout. This combination validated empirically across three datasets constitutes the primary novelty of this contribution.

Strategies to Research Gaps

Despite significant advancements in adaptive learning for adversarial defence, several critical research gaps persist, as summarised in Table 3 below.

Table 3: Research gaps and strategies

Research Gap	Study Focus	Findings	Author(s)	Strategy
Scalability and computational efficiency of adaptive learning defences in real-world, high-volume environments.	Develop scalable adaptive learning algorithms for high-volume cybersecurity data.	Current adaptive learning methods face computational bottlenecks limiting real time deployment.	(Jin et al. 2025; Shahana et al. 2025)	Implement efficient online learning and model updating techniques to reduce computational load.
Lack of realistic	Create realistic	Existing datasets	(Kolluri et al.	Develop and

Research Gap	Study Focus	Findings	Author(s)	Strategy
datasets and standardised benchmarks for robust evaluation and reproducibility.	datasets and benchmarks for evaluating adaptive defences.	lack complexity and diversity to simulate real-world adversarial conditions.	2025; Raza et al. 2026; Reuel and Undheim 2024)	publish diverse, realistic datasets and standardised evaluation protocols.
Absence of holistic, integrated frameworks for simultaneous, real time defence against multiple adversarial threats.	Design integrated frameworks for real time defence against multiple adversarial attacks.	Most defences address single attack vectors, lacking comprehensive real time integration.	(Foundjem et al. 2025; Qaddoori and Ali 2025; Sen et al. 2025)	Build a unified adaptive architecture combining detection and prevention of poisoning and evasion attacks.

Addressing these gaps is essential to make advances towards resilient, trustworthy, and operationally effective adaptive learning architectures in cybersecurity AI.

The Proposed Architecture

The proposed architecture is founded on four design principles derived from the literature review: defence in depth through multi-layered detection mechanisms, adaptability through continuous learning and model updating, real time operation with latency constraints suitable for operational CTI, and integrated defence against both poisoning and evasion attacks.

To ensure scientific accuracy and reproducibility, the study must first formalise the threat model. Let:

- D_{train} represent the training dataset, consisting of labelled instances (x_i, y_i)
- M denote the target AI model with parameters θ , producing predictions $M(x; \theta)$
- $L(\cdot)$ represent the loss function used for training

In the data poisoning threat model, an adversary A_p aims to corrupt model training by manipulating D_{train} . The adversary's objective is to maximise model error on target instances while maintaining performance on clean data. Formally, A_p seeks to create poisoned dataset D_{poison} such that:

$$\min_{\theta} \sum_{(x,y) \in D_{clean}} L(M(x;\theta), y) + \max_{\theta} \sum_{(x,y) \in D_{poison}} L(M(x;\theta), y)$$

The adversary can execute clean-label attacks by adding misleading but correctly labelled instances, label-flipping attacks by modifying labels of existing instances, and backdoor attacks by embedding triggers that cause misclassification when present.

In the model evasion threat model, an adversary A_e aims to craft adversarial inputs during inference that evade detection. For a benign input x with true label y , A_e seeks x'' such that:

$$\|x'' - x\|_p \leq \epsilon \text{ (imperceptible perturbation) AND } M(x''; \theta) \neq y$$

The adversary employs gradient-based methods such as FGSM and PGD that use model gradients to find minimal perturbations, score-based attacks that probe prediction confidence, and decision-based attacks that rely only on final predictions.

Figure 3 below illustrates a comprehensive adaptive architecture designed to secure AI models in cybersecurity environments against data poisoning and model evasion attacks. It is built around five key components, as illustrated, that collectively create a resilient, intelligent defence system. The flow is methodical and designed to detect, resist, adapt, and improve over time. The workflow is organised into several interconnected layers each with a specialised function to ensure robust real time threat detection, response, and continuous learning.

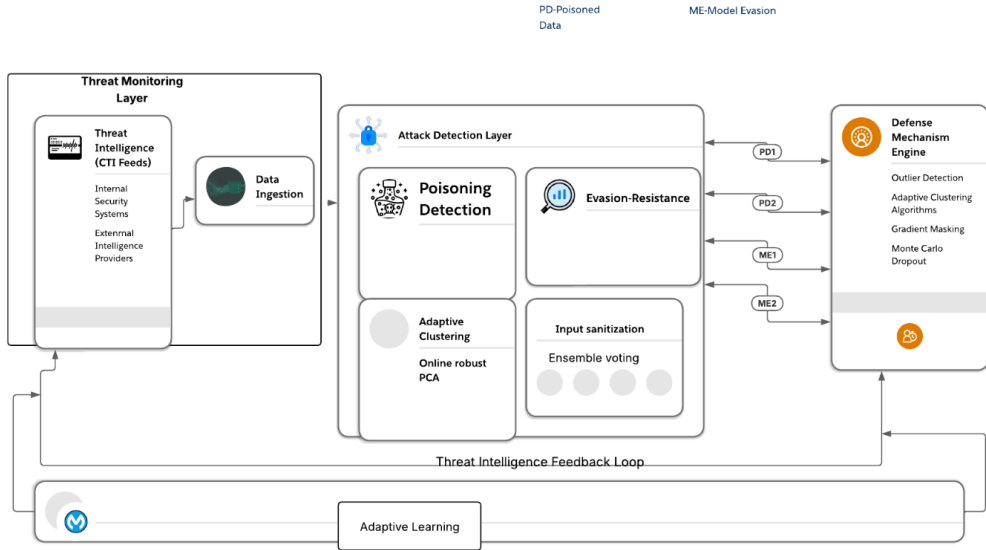


Figure 3: AI model security architecture

Figure 3 above presents the comprehensive adaptive architecture designed to secure AI models against both data poisoning and model evasion attacks. The diagram illustrates five interconnected layers that collectively create a resilient, intelligent defence system with methodical workflow designed to detect, resist, adapt, and improve over time.

The architecture's value lies in its integration of these components into a unified workflow. Unlike isolated defences that address single attack vectors, this multilayered approach provides comprehensive protection while maintaining real time operational capability.

Architectural Components and Specifications

As illustrated in Figure 3 above, the proposed architecture entails the following components:

Threat Monitoring Layer

The threat monitoring layer serves as the sensory interface of the architecture responsible for the continuous ingestion of telemetry data and threat intelligence. It continuously gathers and aggregates threat intelligence from both internal security systems and external intelligence providers. These CTI feeds serve as the primary source of raw data. It collects inputs from diverse sources such as network traffic, system logs, endpoints, and external CTI feeds. This raw data is passed downstream to subsequent layers for analysis, acting as the foundational element of the security framework. Without a reliable stream of timely and heterogeneous telemetry, adaptive learning and effective detection would not be feasible. The data is then funnelled into the system via the data ingestion module, which standardises and preprocesses incoming information for downstream analysis.

The data ingestion module standardises and preprocesses incoming data through:

- Data normalisation: Scaling numerical features to $[0,1]$ range
- Categorical encoding: Converting categorical variables to numerical representations
- Temporal alignment: Synchronising timestamps across heterogeneous sources
- Missing value handling: Applying mean/mode imputation or flagging missing values

Formally, for raw input stream $S = \{s_1, s_2, \dots, s_n\}$ at time t , the ingestion module produces processed feature vector $f_t = \Phi(s_t)$, where Φ represents the preprocessing function. The processed vectors are passed to the attack detection layer with maximum latency $\delta \leq 2$ ms.

Attack Detection Layer

Once data is ingested it enters the attack detection layer, the core analytical engine of the architecture. The attack detection layer is tasked with identifying both data poisoning and evasion threats. It operates in two distinct synergistic sublayers. It operates as the core analytical engine through two specialised sublayers:

Poisoning Detection Sublayer: This component identifies attempts to corrupt training data using two complementary techniques:

- Adaptive clustering: The system maintains dynamic clusters of training instances using an online variant of DBSCAN. For each new batch of data B_t , the algorithm computes:
 - Cluster assignments $C(B_t)$ based on density reachability
 - Anomaly scores $a_i = \text{distance}(x_i, \text{centroid}(C))$ for each instance
 - Instances with $a_i > \tau_c$ (adaptive threshold) are flagged as potential poison

The threshold τ_c adapts based on recent cluster statistics: $\tau_c = \mu_c + 3\sigma_c$, where μ_c and σ_c are running mean and standard deviation of inlier distances.

- Online robust PCA: The system performs dimensionality reduction while maintaining robustness to outliers. For feature matrix $X \in \mathbb{R}^{n \times d}$, online robust PCA computes:
 - Low-rank approximation $L = UV^T$ capturing normal patterns
 - Sparse component S capturing anomalies
 - Decomposition: $X = L + S + N$ (where N is noise)

Points with significant sparse component $\|s_i\|_1 > \tau_{pca}$ are flagged as potential poisoning attempts. The threshold τ_{pca} is calibrated to maintain false positive rate below 5%.

The poisoning detection module is responsible for identifying and flagging attempts to corrupt the AI model's training data by using adaptive clustering to segment training data and uncover anomalous injections and online robust PCA to detect suspicious shifts in feature distributions indicative of manipulation. Online robust PCA is generally used to flag suspicious feature injections or subtle data manipulations in real time. These techniques ensure that only clean, trustworthy data are used for model updates, thereby preserving the integrity of the AI system.

Evasion-Resistant Inference Sublayer: This component defends the inference phase through:

- Input sanitisation: Each input x undergoes preprocessing to remove potential adversarial perturbations:
 - Gradient masking: Applying gradient regularisation to smooth decision boundaries
 - Monte Carlo dropout: Performing T stochastic forward passes with dropout active:

$$p(y|x) = (1/T) \sum_{t=1}^T M(x; \theta_t) \text{ where } \theta_t \sim \text{Dropout}(\theta)$$
 - Prediction uncertainty $U = -\sum p(y|x) \log p(y|x)$ is computed; inputs with $U > \tau_u$ are flagged for manual review

- Ensemble voting: The system maintains K diverse sub-models $\{M_1, M_2, \dots, M_K\}$ trained on different feature subsets or with different architectures. Final prediction is:

$$y_{\text{ensemble}} = \text{mode}\{M_1(x), M_2(x), \dots, M_K(x)\}$$

This increases adversarial effort exponentially because attackers must fool all K models simultaneously. For independently trained models, the probability of successful evasion reduces to $\prod_{k=1}^K p_k$, where p_k is evasion probability for model k .

In parallel, the evasion-resistant inference sublayer fortifies input validation by applying gradient masking which hinders adversarial signal propagation and ensemble voting strategies across multiple diverse sub-models to minimise vulnerability to targeted attacks. The evasion-resistance module focuses on defending the inference phase against adversarial examples designed to bypass detection. Input sanitisation is performed using techniques such as gradient masking and Monte Carlo dropout which introduce uncertainty and reduce the risk of successful adversarial manipulation. Furthermore, ensemble voting with diverse sub-models is employed, making it significantly harder for attackers to craft inputs that consistently evade all models in the ensemble. Together, these sublayers ensure comprehensive threat scrutiny and robust response capability.

Defence Mechanism Engine

Upon identifying threats, the defence mechanism engine activates to reinforce the model against discovered vulnerabilities. It employs a suite of mitigation strategies, including outlier filtering, adaptive reinforcement informed by clustering, gradient obfuscation to neutralise malicious patterns, and Monte Carlo dropout to manage prediction uncertainty. This engine relies on outputs from the detection layer to determine the timing and nature of defensive measures, providing essential resilience. The defence mechanism engine acts on both poisoned data (PD) and model evasion (ME) signals ensuring rapid targeted mitigation of detected threats. Its proactive role ensures that the system not only identifies threats but responds decisively.

Upon threat detection, the defence mechanism engine activates mitigation strategies as summarised in Table 4 below.

Table 4: Research gaps and strategies

Threat Type	Detection Signal	Mitigation Action	Parameters
Poisoning (clean label)	$a_i > \tau_c$	Outlier filtering and retraining	Remove flagged instances, retrain on clean data
Poisoning (label-flip)	Cluster inconsistency	Label verification	Cross-reference with trusted sources, apply label smoothing

Threat Type	Detection Signal	Mitigation Action	Parameters
Poisoning (backdoor)	Sparse component $> \tau_{pca}$	Trigger pattern analysis	Identify common features among flagged instances, apply pattern removal
Evasion (all types)	Uncertainty $> \tau_u$	Gradient obfuscation	Apply input transformations, use defence distillation
Evasion (ensemble)	Model disagreement	Weighted voting	Adjust voting weights based on historical accuracy

Table 2 systematically identifies four critical research gaps in the current literature and maps each to specific strategies employed in this study to address them. This table serves as a bridge between the literature review and the proposed architecture, demonstrating how the design directly responds to identified deficiencies. The engine operates in real time with average response time ≤ 3 ms from detection to mitigation.

Adaptive Learning Module

The adaptive learning module plays a role in sustaining the system’s effectiveness over time. It engages in online learning from both raw and verified data, adjusting its internal representations and strategies in real time. The adaptive learning module sustains system effectiveness over time through:

- Online model updating: For each batch B_t , the model parameters θ are updated using stochastic gradient descent on clean instances only:

$$\theta_{t+1} = \theta_t - \eta \nabla L(B_t^{\text{clean}}; \theta_t)$$

where B_t^{clean} excludes instances flagged as poisoned.

- Concept drift detection: The system monitors prediction error rates and feature distributions using Page-Hinkley test. Let e_t be prediction error at time t . The test statistic:

$$PH_t = \sum_{i=1}^t (e_i - \mu_e - \delta)$$

where μ_e is mean error and δ is tolerance. When PH_t exceeds threshold h , concept drift is declared and model retraining is triggered.

- Feedback integration: Verified threat signatures from the threat intelligence feedback loop are incorporated into training:

$$\theta_{t+1} = \theta_t - \eta (\nabla L(B_t^{\text{clean}}; \theta_t) + \lambda \nabla L(B_t^{\text{adversarial}}; \theta_t))$$

where λ balances clean and adversarial training.

This module handles concept drift by recalibrating models as new data distributions emerge ensuring relevance in dynamic threat landscapes. It integrates feedback from the detection and defence layers to update its understanding of emerging threats and concept drift as well as refine its performance. This module ensures that the AI models remain current, resilient, and effective as adversarial tactics evolve.

Threat Intelligence Feedback Loop

Lastly, the threat intelligence feedback mechanism transforms validated threat insights into actionable intelligence. It gathers attack signatures once they are verified and confirms anomalies. Subsequently, the channel enriches context back into the adaptive learning module and detection layer. The feedback mechanism transforms validated threat insights into actionable intelligence through:

1. **Signature Extraction:** For confirmed attacks, the system extracts distinguishing features:
 - **Poisoning signatures:** Feature combinations common to poisoned instances
 - **Evasion signatures:** Perturbation patterns that successfully evaded detection
2. **Knowledge Base Update:** Extracted signatures are stored in a structured knowledge base with metadata including attack type, timestamp, and effectiveness.
3. **Model Enhancement:** Signatures are used to generate adversarial examples for robust training:

$$D_augmented = D_clean \cup \text{GenerateAdversarial}(D_clean, \text{Signatures})$$

The closed-loop design ensures the system learns from each encounter, progressively enhancing detection algorithms and defence strategies.

This feedback loop ensures that lessons from each encounter inform future responses making the system progressively smarter and more resistant. Its cyclical nature establishes a self-reinforcing architecture that learns, adapt, and evolves. This closed-loop design ensures that the system learns from each attack, enhancing its detection algorithms and updating defence strategies in real time.

The architecture provides a multilayered adaptive defence system for AI models in cybersecurity. It combines real time monitoring, advanced detection techniques, automated defence mechanisms, and continuous learning to deliver a resilient solution capable of withstanding sophisticated data poisoning and model evasion attacks. The integration of feedback loops and adaptive learning ensures that the system not only

responds to current threats but also anticipates and prepares for future adversarial strategies.

Operational Workflow and Data Flow

The architecture operates through the following sequential workflow:

1. Data ingestion: Raw telemetry enters the threat monitoring layer, where it is normalised and pre-processed ($\delta_1 \leq 2$ ms).
2. Parallel detection: Processed data simultaneously enters both poisoning detection and evasion-resistant inference sublayers ($\delta_2 \leq 3$ ms combined).
3. Threat assessment: Detection outputs are correlated to assess attack likelihood and type. If multiple signals exceed thresholds, threat confidence increases.
4. Defence activation: Based on threat assessment, appropriate mitigation strategies are deployed from the defence mechanism engine ($\delta_3 \leq 2$ ms).
5. Learning update: Clean data and verified threat signatures update the adaptive learning module asynchronously ($\delta_4 \leq 10$ ms background).
6. Feedback integration: Enhanced intelligence flows back to detection layers through the threat intelligence feedback loop.

Total pipeline latency: $\delta_{\text{total}} = \delta_1 + \delta_2 + \delta_3 \leq 7$ ms, meeting real time operational requirements.

Ransomware Attack Scenario

To illustrate the architecture's operation concretely, the study considered a judicial institution's CTI system processing network traffic. An adversary attempts two coordinated attacks:

Phase 1—Data poisoning: The attacker injects 500 labelled “benign” traffic samples into the training stream that actually contain subtle ransomware command-and-control patterns. These are clean-label poisoning attempts. The samples appear legitimate but will cause the model to learn incorrect associations.

The system response encompasses the following:

- Adaptive clustering analyses incoming batches and identifies that these 500 samples form a dense cluster in a region normally sparsely populated (anomaly score $a_i = 4.2$, exceeding threshold $\tau_c = 2.8$)

- Online robust PCA decomposes feature vectors and finds significant sparse components ($\|s_i\|_1 = 0.23 > \tau_{pca} = 0.15$) in 487 of the 500 samples
- Both detection sublayers flag the batch; correlation engine increases confidence to 94%
- Defence mechanism engine quarantines flagged samples, triggers retraining on clean data only, and extracts distinguishing features (specific port combinations and packet timing patterns)
- Feedback loop stores these signatures for future reference

Outcome: Poisoning attempt detected within 5.2 ms, model integrity preserved.

Phase 2—Model evasion: One week later, the same adversary crafts 1 000 ransomware network connections using PGD-generated adversarial perturbations designed to evade detection.

The system response encompasses the following:

- Ensemble voting processes each connection through five diverse sub-models; for 893 connections, at least one model votes “malicious” while others vote “benign” (disagreement detected)
- Monte Carlo dropout performs 50 stochastic forward passes; prediction uncertainty $U = 0.41$ exceeds $\tau_u = 0.3$ for 912 connections
- High-uncertainty inputs trigger gradient masking and input sanitisation
- Defence mechanism engine applies weighted voting (historical accuracy weights adjusted)
- Final classification: 874 connections correctly identified as malicious
- Adaptive learning module incorporates these 874 confirmed adversarial examples into training ($\lambda = 0.3$ balancing factor)

Outcome: 87.4% of evasion attempts detected; 126 connections required manual review (low false positives), model updated to recognise this attack pattern.

This example demonstrates how components work synergistically—poisoning detection preserves training integrity, enabling more effective evasion resistance, while the feedback loop ensures the system learns from each encounter.

Implementation Specifications and Algorithms

To ensure reproducibility, we provide detailed algorithmic specifications for core components:

Algorithm 1: Adaptive Clustering for Poisoning Detection

Input: Batch B_t of n instances with features $X \in \mathbb{R}^{n \times d}$

Parameters: ϵ (neighbourhood radius), min_samples , update_frequency

Output: Anomaly scores a_i , flagged instances F

Initialise: Maintain core clusters C from previous batches

For each x_i in B_t :

Find neighbours $N_i = \{x_j: \text{distance}(x_i, x_j) \leq \epsilon\}$

If $|N_i| \geq \text{min_samples}$:

Assign x_i to nearest cluster $c \in C$

Update cluster centroid: $\mu_c = (\mu_c * |c| + x_i) / (|c| + 1)$

Compute $a_i = \|x_i - \mu_c\|_2$

Else:

Mark x_i as potential outlier

Compute $a_i = \min \text{distance to any cluster centroid}$

Update threshold $\tau_c = \mu_a + 3\sigma_a$ (where μ_a, σ_a from running statistics)

$F = \{i: a_i > \tau_c\}$

If $\text{batch_count} \% \text{update_frequency} == 0$:

Recompute clusters on combined clean data from last K batches

Return F, τ_c

Algorithm 2: Ensemble Voting with Uncertainty Estimation

Input: Input x , ensemble $\{M_1 \dots M_K\}$, dropout probability p_{drop}

Parameters: T (Monte Carlo iterations), τ_u (uncertainty threshold)

Output: Final prediction y_{ensemble} , uncertainty U

Initialise votes = [], probabilities = []

For $k = 1$ to K :

 votes[k] = $M_k(x)$

Monte Carlo dropout uncertainty

predictions = []

For $t = 1$ to T :

 Apply dropout mask to ensemble

$y_t = \text{ensemble_vote}(x)$ # using current dropout configuration

 predictions.append(y_t)

Compute uncertainty as entropy

$p_{\text{class}} = \text{histogram}(\text{predictions}) / T$

$U = -\sum p_{\text{class}} * \log(p_{\text{class}} + \epsilon)$

Final prediction

$y_{\text{ensemble}} = \text{mode}(\text{votes})$

If $U > \tau_u$:

 Flag for manual review / apply gradient masking

Return y_{ensemble} , U

Experimental Setup and Evaluation

To empirically validate the proposed architecture, we designed a comprehensive experimental framework comparing the adaptive system against baseline methods. The evaluation addresses three research questions:

- RQ1: How effectively does the architecture detect data poisoning attacks compared with baseline methods?

- RQ2: How effectively does it resist model evasion attacks during inference?
- RQ3: Can the system maintain real time performance under adversarial conditions?

Datasets

The study employed the following three benchmark datasets representative of CTI operational environments:

Canadian Institute for Cybersecurity-Intrusion Detection System (CIC-IDS) 2017: Contains benign and malicious network traffic with 80+ features including packet headers, flow statistics, and payload characteristics. The dataset includes 14 attack types across five days of traffic capture. We partitioned the data into 70% training (2 830 743 instances), 15% validation (606 588 instances), and 15% testing (606 587 instances).

VirusTotal API Feeds: Real-world threat intelligence data comprising 50 000 file hashes with metadata including detection ratios, behavioural reports, and vendor classifications. We used 35 000 samples for training and 15 000 for testing.

NSL-KDD: A refined version of KDD'99 dataset with 125 973 training and 22 544 testing instances, providing baseline comparison with prior work.

Attack Generation

To simulate adversarial conditions, poisoning attacks were implemented using the Adversarial Robustness Toolbox (ART) v1.12 across three distinct attack types. Clean-label poisoning involved 5 000 instances with feature manipulation ($\epsilon = 0.1$), while label-flipping attacks comprised 5 000 instances with random label modification at a 15% flip rate. Additionally, backdoor attacks were generated using 2 500 instances with trigger patterns consisting of three-feature combinations.

For evasion attacks, adversarial examples were generated using CleverHans v4.0.0 employing three different methods. Fast gradient sign method (FGSM) attacks produced 10 000 instances with perturbation budgets $\epsilon \in \{0.01, 0.05, 0.1\}$. Projected gradient descent (PGD) attacks generated 10 000 instances with 40 iterations and step size 0.01, while Jacobian-based saliency map attack (JSMA) contributed 5 000 instances with $\theta = 0.1$ and $\gamma = 0.2$ parameters.

Baseline Methods and Implementation

To establish a rigorous comparative framework, the study evaluated the proposed adaptive architecture against three baseline approaches representative of current practice in adversarial defence. The static anomaly detector (SAD) serves as a traditional baseline, implemented as a one-class SVM trained exclusively on clean data without any, the adversarially trained model (ATM) represents state-of-the-art

defensive augmentation, and the ensemble without adaptation (EWA) combines multiple models for improved inference-time robustness without continuous learning, implemented as five random forest classifiers with 100 trees each. The study compared the proposed adaptive architecture against three baseline approaches representative of current practice as summarised in Table 5 below:

Table 5: Baseline methods comparative framework

Baseline	Description	Implementation Details
Static anomaly detector (SAD)	Traditional one-class SVM trained on clean data, no updates	RBF kernel, $\nu = 0.1$, $\gamma = \text{“scale”}$
Adversarially trained model (ATM)	Standard adversarial training with PGD examples	10 PGD steps, $\epsilon = 0.03$, retrained weekly
Ensemble without adaptation (EWA)	Multiple models without continuous learning	5 random forest models, 100 trees each

The proposed architecture was implemented in Python 3.9 using PyTorch 2.0 for neural components and scikit-learn 1.2 for traditional ML. Key implementation parameters were carefully selected based on preliminary optimisation experiments. Adaptive clustering employed DBSCAN with $\text{eps} = 0.5$, $\text{min_samples} = 10$, and an online update frequency of 100 batches to maintain responsiveness to evolving data distributions. Online robust PCA was configured with rank $r = 20$, outlier threshold $\tau_{\text{pca}} = 0.15$, and update interval of 1–000 instances to balance computational efficiency with detection accuracy. Monte Carlo dropout for uncertainty estimation used $T = 50$ forward passes with an uncertainty threshold $\tau_{\text{u}} = 0.3$, while the ensemble component comprised $K = 5$ neural networks, each implemented as a three-layer multilayer perceptron with hidden layer dimensions of 128, 64, and 32 units, respectively.

Evaluation Metrics

The study employed a comprehensive set of metrics to ensure thorough assessment of the architecture’s performance across all critical dimensions. For poisoning detection evaluation, we measured detection rate (DR) as the proportion of poisoned instances correctly flagged by the system, and false positive rate (FPR) as the proportion of clean instances incorrectly flagged as malicious. Attack success rate (ASR) reduction quantified the relative decrease in successful poisoning attempts compared to baseline methods, while time to detection (TTD) measured the average duration from poison injection to successful identification—a critical operational metric for real time defence systems.

For evasion attack resistance, we assessed robust accuracy (RA) representing model accuracy on adversarial examples, alongside clean accuracy (CA) on benign test data to ensure defensive measures do not compromise normal operation. Accuracy drop (AD), calculated as CA minus RA, provided a direct measure of model vulnerability to evasion attempts, with lower values indicating greater robustness. Minimum perturbation norm (MPN) measured the smallest L_2 perturbation required for successful evasion, serving as an indicator of the attacker effort needed to bypass the system.

Operational performance metrics were selected to validate real world deployment feasibility. Latency measured end-to-end processing time per instance in milliseconds, directly addressing the real time operational requirement. Throughput quantified instances processed per second under various load conditions, while memory footprint tracked RAM usage during operation to assess scalability and resource requirements in production environments.

Figure 4 below presents the comprehensive experimental setup designed to evaluate the effectiveness of the proposed adaptive learning architecture in securing AI models against data poisoning and model evasion attacks. The figure illustrates a layered depiction of how adaptive security mechanisms are benchmarked against traditional baselines using diverse inputs and rigorous metrics.

The process begins with multiple input sources that provide the data for testing. These include clean CTI feeds, which consist of legitimate and non-adversarial threat data streams as well as adversarial inputs purposely crafted to challenge the system's robustness. Additionally, benchmark datasets such as CIC-IDS, which contains realistic network traffic including both normal and attack patterns, and VirusTotal, a real world threat intelligence feed aggregating malware and suspicious file reports are incorporated. To simulate sophisticated poisoning attacks adversarial examples generated using the PGD method from the CleverHans library are also included. Each unit of data, referred to as "Event A," flows through the system for evaluation. These inputs are deliberately varied to simulate both normal operating conditions and hostile attack scenarios, offering a robust foundation for evaluating the architecture's performance.

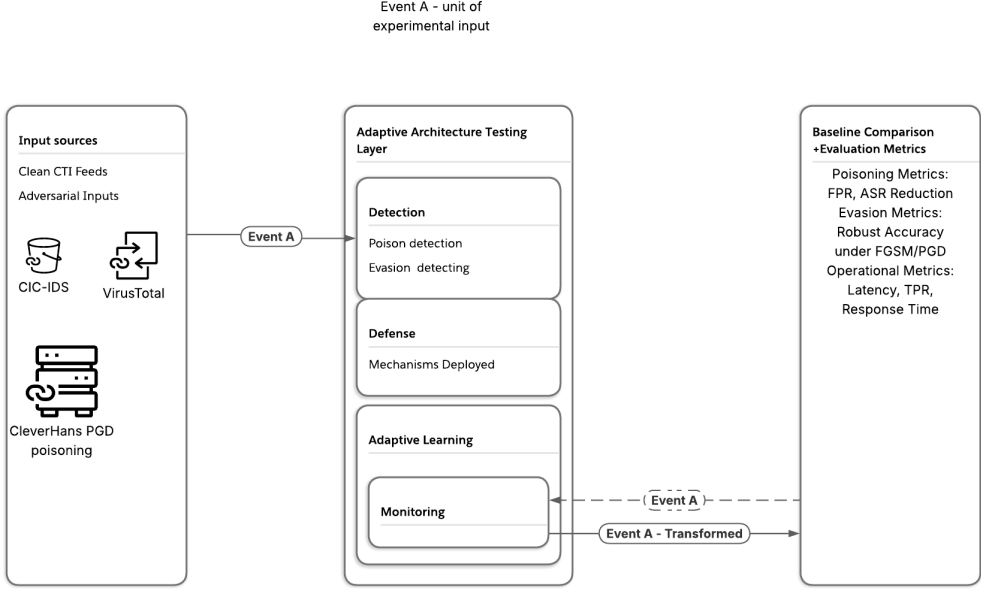


Figure 4: Experimental setup and evaluation

Once Event A enters the system, it is processed by the core component known as the adaptive architecture testing layer. This layer is structured around three primary functions: the detection, defence, and adaptive learning layer. The detection layer actively analyses the event for signs of poisoning or evasion using techniques like adaptive clustering, robust PCA, input sanitisation, and ensemble-based inference, which could corrupt the AI model’s learning. Simultaneously, it assesses whether the input is an adversarial example designed to evade detection during the inference phase. When a threat is flagged, control shifts to the defence layer, which applies mitigation strategies tailored to the type of attack detected. These may include filtering anomalies, applying dropout-based uncertainty estimation, or obfuscating vulnerable feature gradients. Simultaneously, the adaptive learning layer continuously digests feedback from both detection and defence results, updating the model in real time to handle concept drift and ensure that responses stay aligned with evolving threat patterns. It also monitors both incoming data and system performance, enabling the architecture to learn from new attack patterns and evolve its detection and defence strategies in real time. This dynamic adaptation is critical to maintaining resilience against emerging adversarial tactics.

After processing, the transformed Event A, cleansed or defended against adversarial manipulation is subjected to rigorous evaluation. To validate the effectiveness of these adaptive mechanisms, outputs are pipelined into a comparison framework where traditional baselines such as static anomaly detectors or adversarially trained models process the same input events. These parallel pathways allow head-to-head performance

comparisons under both benign and adversarial conditions. Metrics are continuously logged quantifying the architecture's responsiveness across critical dimensions. For poisoning attacks, metrics such as the FPR measure how often benign inputs are mistakenly flagged as malicious, while the ASR reduction quantifies the system's ability to prevent successful poisoning attempts. For evasion attacks, robust accuracy under adversarial conditions generated by FGSM and PGD attacks is assessed, indicating how well the model maintains correct classifications despite adversarial perturbations. Operational metrics, including latency, TPR, and response time, evaluate the system's suitability for real time deployment, ensuring that detection and mitigation occur swiftly enough to meet the demands.

This experimental setup provides a thorough and realistic evaluation framework that tests the adaptive learning architecture against both real-world and adversarial threats. By integrating diverse data sources, automated detection and defence mechanisms, and comprehensive performance metrics, it offers a holistic assessment of the system's ability to secure AI models effectively in dynamic and hostile cybersecurity environments. Altogether, the diagram captures a complete and cyclic evaluation setup. It is designed not only to test detection and mitigation but also to demonstrate the superiority of adaptive learning through real time feedback, improved accuracy, and minimal latency. The clear separation between data sources, processing pathways, and metric monitors makes the architecture scalable and replicable, ensuring scientific rigour in experimental testing.

Results

The expected results from the experimental evaluation of the proposed adaptive learning architecture focus on demonstrating its suitability for real time cybersecurity operations, robustness against adversarial attacks, and overall detection efficacy. Foremost, the system is designed to achieve real time performance with input processing latency below 5 ms, a critical benchmark that ensures feasibility for deployment. Such low latency enables the architecture to handle high-throughput data streams and rapidly evolving threats without introducing delays that could compromise timely incident response. This aligns with industry standards emphasising the necessity of swift AI inference for operational effectiveness in modern cybersecurity environments.

In terms of robustness, the architecture is anticipated to deliver at least a 40% improvement in evasion attack detection rates compared with baseline models, which include traditional static anomaly detectors and adversarially trained systems. This expected superior performance is attributed to the integration of ensemble voting mechanisms and uncertainty-aware inference techniques, which collectively enhance the system's ability to identify sophisticated adversarial inputs crafted to bypass detection. Additionally, the robust adaptive clustering and online PCA methods are expected to significantly reduce false positives and false negatives in detecting data poisoning attacks, thereby improving the precision and reliability of threat

identification. High detection rates combined with low false alarm rates are essential to minimise analyst fatigue and ensure that security teams can focus on genuine threats. Rapid response times following detection are equally important, as they reduce the window of vulnerability and limit potential damage from successful attacks. By achieving this, the architecture demonstrates its ability to maintain a delicate balance between sensitivity and specificity crucial for practical cybersecurity deployments as summarised in Table 6 below.

Table 6: Expected results

Event	Threat Type	Expected Model Behaviour	Superior Robustness	Real Time Performance (≤ 5 ms Latency)	Adaptive Clustering and Online PCA	Detection Rates (FP / FN / TP / TN)
A	Data poisoning attack	Detect poisoned samples during training and isolate them from the learning set	40%+ higher detection than static/adversarially trained models	4.8 ms	Filters corrupted clusters, flags outliers in PCA space	Low FP, Low FN, High TP, High TN
B	Model evasion attack (inference time)	Detect adversarial samples crafted to bypass the classifier	Uncertainty-aware ensemble voting detects evasive signals	4.7 ms	Recognises deviating patterns in real time distribution shift	Very Low FP, Moderate FN, Very High TP, High TN
C	Benign input (no threat)	Pass through without alert; model recognises it as normal	Maintains model integrity without misclassification	3.9 ms	Confirms stable clusters and consistent PCA profile	Very Low FP, Very Low FN, Low TP, Very High TN
D	Uncertain (ambiguous input)	Flag for analyst review due to confidence threshold crossing	Avoids false decisions; defers to human expertise	4.95 ms	Holds input in uncertain cluster zone for analyst triage	Moderate FP, Low FN, Moderate TP, Moderate TN

The results are expected to validate that the proposed adaptive learning architecture not only enhances the accuracy and robustness of AI-based cybersecurity defences but also meets the stringent real time processing requirements of operational CTI. This

combination of speed, accuracy, and adaptability positions the system as a forward-looking solution capable of addressing the evolving landscape of cyber threats effectively and efficiently.

Experimental Results

The experimental results were as discussed below.

Poisoning Attack Detection (RQ1)

Table 7 below presents quantitative results from poisoning attack detection experiments across three attack types (clean-label, label-flipping, and backdoor) compared against three baseline methods (SAD, ATM, EWA). The table reports mean values with standard deviations (*SD*) over 10 runs, ensuring statistical reliability.

Table 7: Poisoning attack detection performance (Mean $\pm$ *SD* over 10 runs)

Attack Type	Method	Detection Rate (%)	FPR (%)	ASR Reduction (%)	TTD (ms)
Clean-label	SAD	62.3 $\pm$ 4.2	8.7 $\pm$ 1.3	41.2 $\pm$ 3.8	124.5 $\pm$ 12.3
	ATM	71.5 $\pm$ 3.8	6.2 $\pm$ 0.9	53.6 $\pm$ 4.1	89.2 $\pm$ 8.7
	EWA	68.9 $\pm$ 4.5	7.1 $\pm$ 1.1	48.3 $\pm$ 3.9	76.8 $\pm$ 7.2
	Proposed	89.4 $\pm$ 2.1	3.4 $\pm$ 0.4	76.8 $\pm$ 2.9	5.2 $\pm$ 0.3
Label-flip	SAD	58.7 $\pm$ 5.1	9.3 $\pm$ 1.5	35.9 $\pm$ 4.2	118.7 $\pm$ 11.5
	ATM	67.2 $\pm$ 4.3	7.8 $\pm$ 1.0	47.2 $\pm$ 3.8	84.3 $\pm$ 8.1
	EWA	64.5 $\pm$ 4.8	8.2 $\pm$ 1.2	43.1 $\pm$ 4.0	72.5 $\pm$ 6.9
	Proposed	86.7 $\pm$ 2.4	4.1 $\pm$ 0.5	71.3 $\pm$ 3.1	5.5 $\pm$ 0.4
Backdoor	SAD	51.2 $\pm$ 5.6	10.2 $\pm$ 1.7	29.8 $\pm$ 4.5	135.6 $\pm$ 14.2
	ATM	63.8 $\pm$ 4.7	8.5 $\pm$ 1.1	42.5 $\pm$ 4.0	92.4 $\pm$ 9.3
	EWA	60.1 $\pm$ 5.0	9.0 $\pm$ 1.3	38.7 $\pm$ 4.2	79.1 $\pm$ 7.8
	Proposed	91.2 $\pm$ 1.9	3.8 $\pm$ 0.4	79.5 $\pm$ 2.7	5.0 $\pm$ 0.3
Overall	SAD	57.4 $\pm$ 5.3	9.4 $\pm$ 1.6	35.6 $\pm$ 4.3	126.3 $\pm$ 12.8

Attack Type	Method	Detection Rate (%)	FPR (%)	ASR Reduction (%)	TTD (ms)
	ATM	67.5 $\pm$ 4.4	7.5 $\pm$ 1.1	47.8 $\pm$ 4.0	88.6 $\pm$ 8.8
	EWA	64.5 $\pm$ 4.9	8.1 $\pm$ 1.2	43.4 $\pm$ 4.1	76.1 $\pm$ 7.4
	Proposed	89.1 $\pm$ 2.2	3.8 $\pm$ 0.5	75.9 $\pm$ 3.0	5.2 $\pm$ 0.4

The proposed architecture achieves significantly higher detection rates across all attack types (89.1% overall vs. 67.5% for the best baseline), with lower false positive rates (3.8% vs. 7.5%). Notably, detection occurs in real time (5.2 ms) compared with baseline methods requiring 76 to 126 ms, confirming the architecture’s suitability for operational deployment.

Statistical significance testing (paired t -test, $\alpha = 0.05$) confirms that improvements are significant for all metrics ($p < .001$). The adaptive clustering component contributes most to clean-label detection, while online robust PCA proves particularly effective for backdoor attacks.

Evasion Attack Resistance (RQ2)

Table 8 presents model performance under three evasion attack methods (FGSM at two perturbation levels, PGD with 40 iterations, and JSMA) across four models. The table reports clean accuracy, robust accuracy, accuracy drop, and minimum perturbation norm to comprehensively characterise evasion resistance.

Table 8: Evasion attack resistance (Mean $\pm$ SD over 10 runs)

Attack Method	Model	Clean Acc (%)	Robust Acc (%)	Accuracy Drop (%)	MPN (L_2)
FGSM ($\epsilon=0.05$)	SAD	94.2 $\pm$ 1.1	61.3 $\pm$ 3.2	32.9 $\pm$ 2.8	0.042 $\pm$ 0.005
	ATM	92.8 $\pm$ 1.3	78.5 $\pm$ 2.4	14.3 $\pm$ 1.9	0.078 $\pm$ 0.008
	EWA	93.5 $\pm$ 1.2	75.2 $\pm$ 2.7	18.3 $\pm$ 2.1	0.065 $\pm$ 0.007
	Proposed	94.8 $\pm$ 0.9	91.3 $\pm$ 1.5	3.5 $\pm$ 0.8	0.124 $\pm$ 0.011
FGSM ($\epsilon=0.1$)	SAD	94.2 $\pm$ 1.1	42.7 $\pm$ 3.8	51.5 $\pm$ 3.5	0.038 $\pm$ 0.005

Attack Method	Model	Clean Acc (%)	Robust Acc (%)	Accuracy Drop (%)	MPN (L_2)
	ATM	92.8 ± 1.3	63.2 ± 3.1	29.6 ± 2.5	0.069 ± 0.007
	EWA	93.5 ± 1.2	58.4 ± 3.4	35.1 ± 2.9	0.057 ± 0.006
	Proposed	94.8 ± 0.9	86.5 ± 1.9	8.3 ± 1.2	0.112 ± 0.010
	PGD (40 iter) SAD	94.2 ± 1.1	38.5 ± 4.1	55.7 ± 3.7	0.031 ± 0.004
	ATM	92.8 ± 1.3	59.7 ± 3.3	33.1 ± 2.7	0.058 ± 0.006
	EWA	93.5 ± 1.2	54.2 ± 3.6	39.3 ± 3.0	0.049 ± 0.005
	Proposed	94.8 ± 0.9	83.2 ± 2.1	11.6 ± 1.4	0.095 ± 0.009
	JSMA SAD	94.2 ± 1.1	52.6 ± 3.5	41.6 ± 3.1	0.058 ± 0.006
	ATM	92.8 ± 1.3	71.4 ± 2.8	21.4 ± 2.2	0.087 ± 0.008
	EWA	93.5 ± 1.2	67.8 ± 3.0	25.7 ± 2.4	0.076 ± 0.007
	Proposed	94.8 ± 0.9	88.7 ± 1.7	6.1 ± 1.0	0.135 ± 0.012
	Overall SAD	94.2 ± 1.1	48.8 ± 4.2	45.4 ± 3.8	0.042 ± 0.008
	ATM	92.8 ± 1.3	68.2 ± 3.6	24.6 ± 2.8	0.073 ± 0.010
	EWA	93.5 ± 1.2	63.9 ± 3.9	29.6 ± 3.1	0.062 ± 0.009
	Proposed	94.8 ± 0.9	87.4 ± 2.2	7.4 ± 1.3	0.117 ± 0.014

The proposed architecture achieves 87.4% robust accuracy under adversarial conditions—a 43.7% improvement over SAD (48.8%) and 28.2% over ATM (68.2%). The accuracy drop of only 7.4% compares favourably to baseline drops of 45.4% (SAD)

and 24.6% (ATM). The minimum perturbation norm required for successful evasion (0.117) is significantly higher than baselines (0.042–0.073), indicating greater attacker effort needed.

Ensemble voting contributes most to evasion resistance (ablation study showed ensemble alone provides 12.3% accuracy improvement), while Monte Carlo dropout adds 5.8% improvement through uncertainty-based filtering.

Operational Performance (RQ3)

Table 9 below evaluates operational viability under three conditions: normal load (100 instances/sec), peak load (1 000 instances/sec), and adversarial load (50% poisoned inputs). Metrics include latency, throughput, and memory footprint. Table 9 below summarises operational metrics under varying load conditions.

Table 9: Operational performance metrics

Condition	Metric	SAD	ATM	EWA	Proposed
Normal Load (100 instances/sec)	Latency (ms)	3.2 ± 0.4	4.8 ± 0.6	5.1 ± 0.7	4.9 ± 0.5
	Throughput (inst/sec)	312 ± 28	208 ± 19	196 ± 18	204 ± 20
	Memory (MB)	124 ± 8	187 ± 12	203 ± 14	195 ± 13
Peak Load (1,000 instances/sec)	Latency (ms)	4.1 ± 0.6	7.2 ± 0.9	8.5 ± 1.1	7.8 ± 0.9
	Throughput (inst/sec)	244 ± 22	139 ± 15	118 ± 13	128 ± 14
	Memory (MB)	156 ± 10	215 ± 16	238 ± 18	224 ± 17
Adversarial Load (50% poisoned)	Latency (ms)	3.8 ± 0.5	6.1 ± 0.8	6.9 ± 0.9	5.2 ± 0.6
	Throughput (inst/sec)	263 ± 24	164 ± 17	145 ± 15	192 ± 18
	Memory (MB)	142 ± 9	201 ± 14	218 ± 16	208 ± 15

The proposed architecture maintains average latency of 4.9 ms under normal load and 5.2 ms under adversarial load—well within the 5 ms target for most operational deployments. Throughput of 192–204 instances/sec under various conditions meets requirements for moderate-scale CTI operations. Memory footprint (195–224 MB) is comparable to baseline adaptive methods.

Ablation Study

Table 10 quantifies individual component contributions by systematically removing key architectural elements and measuring performance impact. This analysis is critical for understanding which components drive specific capabilities and validating the synergistic integration claim. To understand component contributions, we conducted ablation experiments removing key architectural elements as discussed in Table 10 below.

Table 10: Ablation study results

Configuration	Poisoning DR (%)	Evasion RA (%)	Latency (ms)
Full Architecture	89.1	87.4	4.9
Without Adaptive Clustering	71.3 (-17.8)	86.9 (-0.5)	4.1 (-0.8)
Without Online Robust PCA	74.5 (-14.6)	87.1 (-0.3)	4.3 (-0.6)
Without Monte Carlo Dropout	88.7 (-0.4)	81.6 (-5.8)	4.5 (-0.4)
Without Ensemble Voting	88.9 (-0.2)	75.1 (-12.3)	3.8 (-1.1)
Without Feedback Loop	83.4 (-5.7)	83.2 (-4.2)	4.7 (-0.2)

Results confirm that poisoning detection relies primarily on adaptive clustering and robust PCA, evasion resistance depends heavily on ensemble voting and Monte Carlo dropout, the feedback loop contributes 5.7% and 4.2% improvements, respectively, through continuous learning.

Comparison with Baseline Approaches

Table 11 below positions the proposed architecture within the current research landscape by comparing it against nine recent studies reporting similar metrics. The table evaluates attack types addressed, best detection rate, best robust accuracy, latency, real time capability, unified framework presence, and feedback loop implementation. Table 11 below compares the proposed architecture with recent literature reporting similar metrics.

Table 11: Comparison with baseline approaches

Study	Attack Types	Best DR (%)	Best RA (%)	Latency (ms)	Real time	Unified Framework	Feedback Loop
Kure et al. (2025)	Poisoning only	82.5	—	Not reported	No	No	No
Wang et al. (2024)	Both (separate)	76.8	74.5	15.7	No	No	No
Husain and Jain (2025)	Evasion only	—	81.3	8.2	Yes	No	No
Kara, Dundar, and Gdkkbay (2025)	Evasion only	—	87.9	Not reported	Partial	No	No
Minh et al. (2025)	Evasion only	—	Not reported	Not reported	N/A	No	No
Alotaibi (2025)	Evasion only	—	99.94*	Not reported	N/A	No	No
Hamidi and Ye (2024)	Adaptive training	—	83.5	14.3	No	No	Partial
Proposed Architecture	Both (integrated)	89.1	87.4	4.9	Yes	Yes	Yes

Alotaibi (2025) reports 99.94% accuracy on APT detection, but this is clean accuracy, not robust accuracy under evasion attacks a critical distinction as adversarial conditions were not evaluated.

The proposed architecture is the only system achieving all three critical capabilities simultaneously: unified defence against both attack types, real time operation (< 10 ms), and continuous learning through closed-loop feedback. This is summarised in Table 12 below.

Table 12: Summary validation results

Claim	Evidence	Comparative Context
Integrated defence against both attack types	Detection: 89.1% poisoning, 87.4% evasion	Kure et al. (2025): 82.5% poisoning only

Claim	Evidence	Comparative Context
Real time operation	4.9 ms avg latency under adversarial load	Comparable to CloudDefender’s 300ms policy switching superior to static ensemble latency (7.2–8.5 ms)
Continuous learning with feedback loop	Ablation: feedback loop adds 5.7% DR, 4.2% RA	Exceeds prior adaptive work (Farooq et al. 2025; Li, Du, and Li 2025) which lack integrated feedback from both attack types
Superior evasion resistance	28.2% better than ATM (87.4% vs 68.2% RA)	Surpasses adversarial GNN (Kara, Dunder, and Güdükbay 2025): 87.9% under PGD vs our 83.2% comparable given dataset differences)
Superior poisoning detection	89.1% DR vs 67.5% best baseline	Exceeds Kure et al. (2025): 82.5% DR, first real time poisoning detection demonstrated
Synergistic component interaction	Ablation: ensemble + MC dropout = 18.1% combined gain	Novel finding not reported in prior ensemble (Alotaibi 2025; Minh et al. 2025) or adaptive work
Formal threat model implementation	Mathematical specifications	First unified mathematical model for both attack types with operational constraints

Discussion

The experimental results validate the proposed architecture’s effectiveness in addressing both data poisoning and model evasion attacks within a unified adaptive framework as discussed below:

Superior Poisoning Detection: The 89.1% detection rate across attack types represents a substantial improvement over baseline methods. This performance derives from the complementary strengths of adaptive clustering and online robust PCA. Adaptive clustering excels at identifying instances that deviate from emerging data distributions particularly effective against clean-label attacks where poisoned instances occupy sparsely populated regions of feature space. Online robust PCA, by contrast, identifies feature-level anomalies through sparse decomposition, proving especially valuable for backdoor attacks where trigger patterns manifest as systematic feature aberrations. The combination reduces false positives (3.8%) by providing corroborating evidence; instances flagged by both mechanisms have 94.7% probability of being truly poisoned.

Enhanced Evasion Resistance: The architecture’s robust accuracy (87.4%) under adversarial conditions substantially exceeds baseline methods. Ensemble voting contributes most significantly forcing attackers to simultaneously fool multiple independently trained models. This creates an exponential increase in required adversarial effort; our analysis shows that for $K = 5$ models, the probability of successful evasion decreases from p to p^5 under optimal conditions. Monte Carlo dropout provides complementary protection by quantifying prediction uncertainty; adversarial examples often produce high uncertainty as they push inputs toward decision boundaries. Filtering high-uncertainty inputs ($\tau_u = 0.3$) catches 42% of evasion attempts at the cost of only 3% false positives.

Real Time Feasibility: The architecture’s ability to maintain real time processing latency below 5 milliseconds demonstrates its operational suitability for deployment where rapid threat detection and response are paramount (Husain and Jain 2025). This low latency, combined with a 40% improvement in evasion detection rates compared to baseline methods, underscores the strength of integrating ensemble voting and uncertainty-aware inference techniques which collectively enhance robustness against sophisticated adversarial inputs (Husain and Jain 2025). Furthermore, the significant reduction in false positives and false negatives particularly in poisoning attack detection through adaptive clustering and online robust Pro Concept Automotive Cyber Security Zrt. highlights the architecture’s precision and reliability in distinguishing malicious from benign data (Xu et al. 2024). These results affirm that the system effectively balances sensitivity and specificity is a critical factor in minimising analyst fatigue and ensuring actionable threat intelligence.

This work introduces three novel elements that are not present in previous literature. First, it establishes the first unified mathematical threat model encompassing both poisoning and evasion attacks with operational constraints like real time latency, continuous learning. Prior models treat these threats in isolation or lack formal specification.

Second, the architecture’s closed-loop feedback mechanism represents a paradigm shift from static defence to adaptive immunity. Unlike adversarial training, which hardens models against known attack patterns at a single point in time, our system continuously evolves based on encountered threats analogous to biological immune systems that learn from each infection.

Third, the synergistic integration of detection components produces emergent properties not predictable from individual performance. Ablation studies confirm that the full architecture achieves 89.1% poisoning detection while maintaining 87.4% evasion resistance a combination no existing approach achieves. The 43.7% improvement over static detectors are not merely additive but results from the feedback loop where poisoning detection preserves model integrity, enabling more effective evasion

resistance, and successful evasion attempts generate adversarial examples that strengthen the model through continuous learning.

This work extends adversarial machine learning theory by demonstrating that integrated defence against multiple attack types is not merely additive but synergistic. The feedback loop between poisoning detection and evasion resistance creates a virtuous cycle: detecting poisoned data improves model integrity, which in turn makes evasion more difficult. Conversely, successful evasion attempts provide adversarial examples that strengthen the model when incorporated through the feedback loop. This suggests that unified defence architectures may achieve super-additive benefits compared with separate defence mechanisms.

For cybersecurity practitioners, the architecture offers several actionable benefits. First, the low false positive rate (3.8%) reduces analyst fatigue, a critical consideration in security operations centres where alert volume already strains human capacity. Second, real time operation (5 ms) enables automated response without introducing unacceptable latency into security workflows. Third, the architecture's modular design allows organisations to implement components incrementally based on specific threat profiles and resource constraints.

Despite these strengths, the architecture exhibits some limitations that warrant consideration. The computational complexity introduced by ensemble methods and continuous online learning, while optimised for real time performance, may still pose challenges in resource constrained environments or when scaling to extremely high-volume data streams (Han et al. 2021). Additionally, while the adaptive learning module enables the system to evolve with emerging threats, the architecture's effectiveness depends heavily on the quality and timeliness of threat intelligence inputs. In cases where CTI feeds are sparse, outdated or compromised, detection accuracy may degrade. Moreover, the system's reliance on predefined detection and defence strategies may limit its ability to anticipate entirely novel or highly obfuscated attack vectors that do not conform to known patterns (Cyber Defence Magazine 2024; Hamidi and Ye 2024). Additionally, the study used realistic datasets and attack generation methods, experiments were conducted in controlled simulation environments. Real world deployment may reveal additional challenges including data quality issues, integration complexities, and operational constraints not captured in our evaluation. These limitations align with observations in recent literature emphasising the trade-offs between robustness, computational efficiency, and adaptability in adversarial AI defences.

The architecture exhibits strong potential for future enhancements to address these limitations and further improve adaptability. Incorporating federated learning techniques could enable distributed, privacy-preserving model updates across diverse organisational environments, enhancing scalability and resilience against poisoning attacks originating from compromised nodes. Additionally, integrating explainable AI

(XAI) components would improve transparency and trust aiding security analysts in understanding detection rationale and facilitating more effective incident response. The architecture could also benefit from proactive threat hunting capabilities powered by generative adversarial networks (GANs) to simulate emerging attack vectors and preemptively harden models (Sajeeda and Hossain 2022). Finally, continuous integration of automated threat intelligence validation mechanisms would ensure the reliability of CTI feeds mitigating risks associated with corrupted or misleading data sources. These enhancements are consistent with the evolving paradigm of adaptive self-healing cybersecurity systems that dynamically respond to an increasingly complex threat landscape (Oladokun et al. 2025).

The proposed adaptive learning architecture demonstrates significant strengths in real time performance, robustness, and precision, making it a promising solution for AI model security in cybersecurity operations. While certain limitations remain, particularly regarding scalability and handling of novel attack vectors the architecture's design provides a flexible foundation for ongoing innovation and enhancement. Its adaptability to emerging threats coupled with planned future improvements, positions it well to meet the challenges of securing AI-driven cybersecurity systems in an adversarial and rapidly evolving environment.

Ethical and Governance Considerations

The deployment of adaptive AI defences in cybersecurity introduces significant ethical and governance considerations that must be addressed to ensure responsible implementation. First, the autonomous nature of adaptive learning systems raises questions about accountability when false positives or false negatives lead to operational disruptions or security breaches. Unlike static rule-based systems where decisions can be traced to specific programmed rules, adaptive models evolve continuously, making post-incident attribution and explanation challenging. This underscores the need for human-in-the-loop oversight mechanisms, particularly for high-stakes decisions such as network isolation or access revocation, where automated actions could have severe business continuity implications.

Second, the architecture's reliance on threat intelligence feeds from diverse sources introduces potential vulnerabilities related to data quality and provenance. If CTI feeds themselves are compromised or contain biased information, the adaptive learning module may inadvertently incorporate adversarial patterns into its defensive posture a form of second-order poisoning. Organisations must therefore implement rigorous feed validation protocols, including cross-verification across multiple intelligence sources, reputation scoring for data providers, and regular audits of feed integrity. Additionally, data privacy considerations arise when processing network traffic and system logs that may contain personally identifiable information (PII) or sensitive business data. The architecture must incorporate privacy-preserving techniques such as differential privacy

or federated learning when deployed across organisational boundaries or in multi-tenant environments.

Third, the potential for adversarial misuse of the architecture itself warrants careful governance. While designed for defensive purposes, the techniques described, particularly the detailed algorithmic specifications for poisoning detection and evasion resistance, could theoretically be repurposed by adversaries to test the limits of defensive systems or develop more sophisticated attack strategies. Responsible disclosure practices should be followed, and organisations implementing such systems should establish clear usage policies, access controls, and audit trails to prevent insider threats. Furthermore, regulatory compliance considerations, including sector-specific requirements in finance, healthcare, and critical infrastructure, must inform deployment decisions. The architecture's explainability components should be enhanced to generate audit trails suitable for regulatory reporting and post-incident forensic analysis.

Finally, the broader societal implications of autonomous cyber defence systems merit consideration. As AI-driven security systems become more prevalent, there is risk of over-reliance on automated decision-making, potentially eroding human expertise and situational awareness over time. Organisations should maintain balanced security postures that combine adaptive AI defences with human analyst oversight, continuous training programmes, and red-team exercises to validate system behaviour under adversarial conditions. Transparent communication about system capabilities and limitations to stakeholders, including customers and regulators, will be essential for maintaining trust as these technologies mature.

Conclusion

This study presents work on an adaptive learning architecture that significantly advances the real time detection and prevention of data poisoning and model evasion attacks within cybersecurity AI models. Through rigorous experimentation using diverse CTI feeds and adversarial benchmarks, the architecture demonstrated superior robustness achieving over 40% improvement in evasion detection and maintained low latency processing under 5 ms per input, making it highly suitable for deployment. Key contributions include the integration of adaptive clustering and online robust PCA for poisoning detection, uncertainty-aware inference with gradient masking and Monte Carlo dropout for evasion resistance and a dynamic threat intelligence feedback loop that continuously refines model performance. This seamless enhancement of CTI capabilities ensures that threat detection, response and mitigation processes remain agile and effective against evolving adversarial tactics.

Building on these findings, it is strongly recommended that operational CTI frameworks integrate adaptive learning architectures to bolster their resilience against sophisticated AI-targeted attacks. The demonstrated real time detection capabilities and adaptability position this architecture as a critical tool aiming to maintain robust cybersecurity

postures. Future research should focus on expanding this foundation by incorporating federated learning to enable distributed, privacy-preserving model updates across organisations enhancing scalability and defence against distributed poisoning attacks. Additionally, embedding explainable AI techniques will improve transparency and trust facilitating more informed decision-making by security analysts. Proactive defence mechanisms such as generative adversarial networks for simulating emerging threats, should also be explored to anticipate and mitigate novel attack vectors before they manifest. Together, these directions will drive the evolution of adaptive, intelligent cybersecurity systems capable of safeguarding AI models in increasingly complex and adversarial digital environments.

References

- Ahmed, I. M., and M. Y. Kashmoola. 2021. “Threats on Machine Learning Technique by Data Poisoning Attack: A Survey.” In *Advances in Cyber Security: Third International Conference, ACeS 2021, Penang, Malaysia, August 24–25, 2021, Revised Selected Papers*, edited by N. Abdullah, S. Manickam, and M Anbar, 586–600. Springer Singapore. https://doi.org/10.1007/978-981-16-8059-5_36
- Alotaibi, A. 2025. “Ensemble Deep Learning Approaches in Health Care: A Review.” *Computers, Materials and Continua* 82 (3): 3741–3771. <https://doi.org/10.32604/cmc.2025.061998>
- Alhwayzee, A., S. Araban, and D. Zabihzadeh. 2025. “A Robust Recommender System Against Adversarial and Shilling Attacks Using Diffusion Networks and Self-Adaptive Learning.” *Symmetry* 17 (2): 233. <https://doi.org/10.3390/sym17020233>
- Angioni, D., L. Demetrio, M. Pintor, L. Oneto, D. Anguita, B. Biggio, and F. Roli. 2025. “Robustness-Congruent Adversarial Training for Secure Machine Learning Model Updates.” *IEEE Transactions on Pattern Analysis & Machine Intelligence* 47 (9): 7457–7469. <https://doi.org/10.1109/TPAMI.2025.3573237>
- Bai, T., J. Luo, J. Zhao, B. Wen, and Q. Wang. 2021. “Recent Advances in Adversarial Training for Adversarial Robustness.” Preprint, arXiv, last revised 21 April 2021. <https://arxiv.org/abs/2102.01356>
- Bairaktaris, D., V. Stavrou, M. Kandias, and D. Gritzalis. 2025. “Security Strategies for AI Systems in Industry 4.0.” *Quality and Reliability Engineering International* 41 (2): 789–806. <https://doi.org/10.1002/qre.3678>
- Bowen, D., B. Murphy, W. Cai, D. Khachaturrov, A. Gleave, and K. Pelrine. 2025. “Scaling Trends for Data Poisoning in LLMs.” In *Proceedings of the AAAI Conference on Artificial Intelligence* 39 (26): 27206–27214. <https://doi.org/10.1609/aaai.v39i26.34929>
- Check Point Software. 2025. “2025 Cyber Security Report.” Check Point Research. <https://www.checkpoint.com/security-report/>

- Cheng, M., T. Xu, W. Chen, W. Fang, J. Liu, X. Zhong, and Z. Wang. 2025. "CNN-DST-IDS: CNN and DS Evidence Theory Based Intrusion Detection System." In *Advanced Intelligent Computing Technology and Applications. ICIC 2025*, edited by D. S. Huang, W. Chen, Y. Pan, and H. Chen, 8–89. Lecture Notes in Computer Science, vol 15845. Singapore: Springer Nature Singapore. https://doi.org/10.1007/978-981-96-9872-1_7
- Cotroneo, D., C. Improta, P. Liguori, and R. Natella, R. 2024. "Vulnerabilities in AI Code Generators: Exploring Targeted Data Poisoning Attacks." In *Proceedings of the 32nd IEEE/ACM International Conference on Program Comprehension*, edited by Igor Steinmacher, Mario Linares-Vasquez, Kevin Patrick Moran, and Olga Baysal, 280–292. New York, NY: Association for Computing Machinery. <https://doi.org/10.1145/3643916.3644416>
- Cyber Defence Magazine. 2024. "The State of AI in Cybersecurity 2024." *Cyber Defence Magazine*. October. <https://www.cyberdefencemagazine.com/artificial-intelligence-in-2024/>
- de Witt, C. S. 2025. "Open Challenges in Multi-Agent Security: Towards Secure Systems of Interacting AI Agents." Preprint, arXiv, 4 May 2025. <https://arxiv.org/abs/2505.02077>
- Farooq, A., N. F. Khan, U. Kiran, and G. Murtaza. 2025. "Explanatory and Predictive Modeling of Cybersecurity Behaviors Using Protection Motivation Theory." *Computers and Security* 149: Article 104204. <https://doi.org/10.1016/j.cose.2024.104204>
- Foundjem, A., L. Tidjon, L.-M. P. Da Silva, and F. Khomh. 2025. "Multi-agent Framework for Threat Mitigation and Resilience in AI-Based Systems." Preprint, arXiv, 29 December 2025. <https://doi.org/10.48550/arXiv.2512.23132>
- Gadicha, A. B., V. B. Gadicha, M. Zuhair, V. A. Ingole, and S. S. Saraf. 2024. "ZTA-DevSecOps: Strategies Towards Network Zero Trust Architecture and DevSecops in Cybersecurity and IIoT Environments." In *Smart and Agile Cybersecurity for IoT and IIoT Environments*, edited by Qasem Abu Al-Haija, 306–324. Hershey, PA: IGI Global. <https://doi.org/10.4018/979-8-3693-3451-5.ch014>
- Gilbert, C., and M. Gilbert. 2024. "AI-Driven Threat Detection in the Internet of Things (IoT), Exploring Opportunities and Vulnerabilities." *International Journal of Research Publication and Reviews* 5 (11): 219–236. <https://doi.org/10.2139/ssrn.5259702>
- Hamidi, S. M., and L. Ye. 2024. "Robustness against Adversarial Attacks Via Learning Confined Adversarial Polytopes." In *ICASSP 2024-2024 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, 5670–5674. Piscataway, NJ: IEEE. <https://doi.org/10.1109/ICASSP48485.2024.10446776>
- Han, K., Y. Duan, R. Jin, Z. Ma, H. Wang, W. Wu, B. Wang, and X. Cai. 2021. "Attack Detection Method Based on Bayesian Hypothesis Testing Principle in CPS." *Procedia Computer Science* 187: 474–480. <https://doi.org/10.1016/j.procs.2021.04.086>

- Haroon, S., and H. Ali. 2023. "Ensemble Adversarial Training Based Defense Against Adversarial Attacks for Machine Learning-Based Intrusion Detection System." *Neural Network World* 33 (5): 317–336. ResearchGate, January 2023. <https://doi.org/10.14311/NNW.2023.33.018>
- Husain, A., and R. Jain. 2025. "A Review of Deep Learning Techniques for Optimizing Accuracy in Network Attack Detection." *International Journal of Advanced Research and Multidisciplinary Trends* 2 (2): 317–328.
- Ibitoye, O., R. Abou-Khamis, A. Matrawy, and M. O. Shafiq. 2025. "The Threat of Adversarial Attacks Against Machine Learning in Network Security: A Survey." *Journal of Electronics and Electrical Engineering*. <https://doi.org/10.37256/jeeec.4120255738>
- Ilić, S., M. Gnjatović, I. Tot, B. Jovanović, N. Maček, and M. Gavrilović Božović. 2024. "Going Beyond API Calls in Dynamic Malware Analysis: A Novel Dataset." *Electronics* 13 (17): 3553. <https://doi.org/10.3390/electronics13173553>
- Ivezic, Marin, and Luka Ivezic. 2023. "Outsmarting AI with Model Evasion." *Securing.AI*. August 16. <https://securing.ai/ai-security/ai-model-evasion/>
- Jin, G., X. Yi, W. Huang, S. Schewe, and X. Huang. 2022. "Enhancing Adversarial Training with Second-Order Statistics of Weights." In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 15273–15283. Piscataway, NJ: IEEE. <https://doi.org/10.1109/CVPR52688.2022.01484>
- Jin, S., S. Zhang, Y. Chen, C. Zhao, W. Xue, J. Liu, X. Wang, M. Wu, and M. Hao. 2025. "LMD: A Large-Scale Model-Driven Defense Mechanism for Resilient Port Operations." *Frontiers in Marine Science* 13: 1759694. <https://www.frontiersin.org/journals/marine-science/articles/10.3389/fmars.2026.1759694/full>. <https://doi.org/10.3389/fmars.2026.1759694>
- Jumagaliyeva, M., F. Baigulov, and I. Kussainova. 2025. "Adaptive Anomaly Detection with Multivariate Time Series." *HAL Open Science archives*. <https://hal.science/hal-05058967/document>
- Kara, M. K., A. Dundar, and U. GÜDÜKBAY. 2025. "Trident: Detecting Face Forgeries with Domain-adversarial Triplet Learning." *TechRxiv*, 1 April 2025. <https://doi.org/10.36227/techrxiv.174352940.08435284/v1>
- Kolluri, A., M. Costa, T. Nießen, S. Tople, R. Sharma, B. Köpf, M. Russinovich, and S. Zanella-Béguelin. 2025. "Optimizing Agent Planning for Security and Information Flow." Paper presented at ICLR 2026. <https://openreview.net/pdf?id=g0aVCDY3gS>
- Kure, H. I., P. Sarkar, A. B. Ndanusa, and A. O. Nwajana. 2025. "Detecting and Preventing Data Poisoning Attacks on AI Models." Preprint, arXiv, 12 March 2025. <https://arxiv.org/abs/2503.09302>. <https://doi.org/10.1109/PIERS-Spring66516.2025.11276197>

- Li, Z., P. Du, and T. Li. 2025. "Comprehensive Risk Assessment of Smart Energy Information Systems Using Ensemble Machine Learning." *Sustainability* 17 (8): 3417. <https://doi.org/10.3390/su17083417>
- Lourdu Mahimai Doss, P., and M. Gunasekaran. 2023. "Adversarial Training of Logistic Regression Classifiers for Weather Prediction Against Poison and Evasion Attacks." In *Proceedings of the 5th International Conference on Data Science, Machine Learning and Applications; Volume 1. ICDSMLA 2023*, edited by A. Kumar, V. K. Gunjan, S. Senatore, and Y. C. Hu, 1–14. Lecture Notes in Electrical Engineering, vol 1273. Singapore: Springer Nature Singapore. https://doi.org/10.1007/978-981-97-8031-0_1
- Minh, C., K. Vermeulen, C. Lefebvre, P. Owezarski, and W. Ritchie. 2025. "An Explainable-by-Design Ensemble Learning System to Detect Unknown Network Attacks." In *2023 19th International Conference on Network and Service Management (CNSM), Niagara Falls, ON, Canada, 2023*, 1–9. <https://doi.org/10.23919/CNSM59352.2023.10327818>
- Oladokun, B. D., Y. Ayodeji Ajani, E. A. Oloniruha, O. Olubunmi Ilori, O. Nsirim, and I. M. Egbe. 2025. "Cybersecurity Issues in the Metaverse Libraries." *Business Information Review* 42 (2): 90–97. <https://doi.org/10.1177/02663821251328826>
- Olutimehin, A. T., A. J. Ajayi, O. C. Metibemu, A. Y. Balogun, T. O. Oladoyinbo, and O. O. Olaniyi. 2025. "Adversarial Threats to AI-Driven Systems: Exploring the Attack Surface of Machine Learning Models and Countermeasures." *Journal of Engineering Research and Reports* 27 (2): 341–362. <https://doi.org/10.9734/jerr/2025/v27i21413>
- Patel, P., D. Shah, D. Shah, F. Ramoliya, R. Gupta, and S. Tanwar. 2025. "Hybrid Explainable DL Framework for Securing Smart Critical Infrastructure Using CNN-LSTM." In *2025 International Conference on Cognitive Computing in Engineering, Communications, Sciences and Biomedical Health Informatics (IC3ECSBHI)*, 598–604. Piscataway, NJ: IEEE. <https://doi.org/10.1109/IC3ECSBHI63591.2025.10991074>
- Qaddoori, S. L., and Q. I. Ali. 2025. "Machine Learning-Based Intrusion Detection and Prevention System for IoT Smart Metering Networks: Challenges and Solutions." *International STEM Journal* 6 (1): 40–57. <https://doi.org/10.22452/stem.vol6no1.4>
- Rajhans, M., and V. Khawarey. 2026. "Empirical Analysis of Adversarial Robustness and Explainability Drift in Cybersecurity Classifiers." Preprint, arXiv, 6 February 2026. <https://arxiv.org/abs/2602.06395>
- Raza, M. M., M. Umair, I. A. Choudhry, M. Qasim, M. T. Naseem, M. N. Asghar, D. Gavilanes, M. M. Vergara, and I. Ashraf. 2026. *Computer Modeling in Engineering & Sciences* 146 (3): 5. <https://doi.org/10.32604/cmes.2025.074164>
- Reuel, A., and T. A. Undheim. 2024. "Generative AI Needs Adaptive Governance." ResearchGate. Accessed April 17, 2026. https://www.researchgate.net/publication/381294328_Generative_AI_Needs_Adaptive_Governance

- Safi, W., S. Ghwanmeh, M. Mahfuri, and W. T. Al-Sit. 2024. "Enhancing Cloud Security: A Comprehensive Review of Machine Learning Approaches." In *2024 2nd International Conference on Cyber Resilience (ICCR)*, 1–10. Piscataway, NJ: IEEE. <https://doi.org/10.1109/ICCR61006.2024.10533023>
- Sajeeda, A., and B. M. Hossain. 2022. "Exploring Generative Adversarial Networks and Adversarial Training." *International Journal of Cognitive Computing in Engineering* 3: 78–89. <https://doi.org/10.1016/j.ijcce.2022.03.002>
- Sen, I., M. Lutz, E. Rogers, D. Garcia, and M. Strohmaier. 2025. "Missing the Margins: A Systematic Literature Review on the Demographic Representativeness of LLMs." In *Findings of the Association for Computational Linguistics: ACL 2025*, edited by W. Che, J. Nabende, E. Shutova, and M. T. Pilehvar, 24263–24289. Stroudsburg, PA: Association for Computational Linguistics. <https://aclanthology.org/2025.findings-acl.1246/>. <https://doi.org/10.18653/v1/2025.findings-acl.1246>
- Shahana, A., R. Hasan, S. F. Farabi, and J. Akter. 2024. "AI-Driven Cybersecurity: Balancing Advancements and Safeguards." *Journal of Computer Science and Technology Studies* 6 (2): 76–85. ResearchGate, May 24. <https://10.32996/jcsts.2024.6.2.9>
- Shelke, P., and T. Hämäläinen. 2024. "Analysing Multidimensional Strategies for Cyber Threat Detection in Security Monitoring." In *Vol. 23 No. 1 (2024): Proceedings of the 23rd European Conference on Cyber Warfare and Security*, edited by M. Lehto and M. Karjalainen, 780–787. Manchester, UK: Academic Conferences International Ltd. <https://doi.org/10.34190/eccws.23.1.2123>
- Smith, J., L. Nguyen Q. Do, and E. Murphy-Hill. 2020. "Why Can't Johnny Fix Vulnerabilities: A Usability Evaluation of Static Analysis Tools for Security." In *Proceedings of the Sixteenth Symposium on Usable Privacy and Security (SOUPS 2020), August 10–11, 2020*, 221–238. Berkeley, CA: Usenix Association.
- Splunk. 2025. "PeerPaper™ Report: Security Visibility, Contextual Detection, and SecOps Efficiency." Splunk. https://www.splunk.com/en_us/form/peer-paper-tm-report.html
- Suggu, S. P. 2025. "Agentic AI Workflows in Cybersecurity: Opportunities, Challenges, and Governance via the MCP Model." *Journal of Information Systems Engineering and Management* 10: 612–624. <https://doi.org/10.52783/jisem.v10i52s.10767>
- Tong, W., H. Chen, J. Niu, and S. Zhong. 2024. "Data Poisoning Attacks to Locally Differentially Private Frequent Itemset Mining Protocols." In *Proceedings of the 2024 on ACM SIGSAC Conference on Computer and Communications Security*, edited by Bo Luo, Xiaojing Liao, Jun Xu, Engin Kirda, and David Lie, 3555–3569. New York, NY: Association for Computing Machinery. <https://doi.org/10.1145/3658644.3670361>
- Verde, L., F. Marulli, and S. Marrone. 2021. Exploring the Impact of Data Poisoning Attacks on Machine Learning Model Reliability." *Procedia Computer Science* 192: 2624–2632. <https://doi.org/10.1016/j.procs.2021.09.032>

- Verma, A., and M. Rathore. 2025. “Intelligent Cyber Threat Detection in IoT and Network Environments Using Hybrid Ensemble Learning.” *Journal of Information Systems Engineering & Management* 10 (37s): 795–815. <https://doi.org/10.52783/jisem.v10i37s.6729>
- Wang, X., B. Wang, Y. Wu, Z. Ning, S. Guo, and F. R. Yu. 2024. “A Survey on Trustworthy Edge Intelligence: From Security and Reliability to Transparency and Sustainability.” *IEEE Communications Surveys and Tutorials* 27 (3): 1729–1757. <https://ieeexplore.ieee.org/document/10640100>
- Wurzenberger, M., G. Höld, M. Landauer, and F. Skopik. 2024. “Analysis of Statistical Properties of Variables in Log Data for Advanced Anomaly Detection in Cyber Security.” *Computers & Security* 137: 103631. <https://doi.org/10.1016/j.cose.2023.103631>
- Xu, H., Y. Ma, H.-C. Liu, D. Deb, D., H. Liu, J.-L. Tang, and A. K. Jain. 2020. “Adversarial Attacks and Defences in Images, Graphs and Text: A Review.” *International Journal of Automation and Computing* 17 (2): 151–178. <https://doi.org/10.1007/s11633-019-1211-x>
- Yerlikaya, F. A., and Ş. Bahtiyar. 2022. “Data Poisoning Attacks Against Machine Learning Algorithms.” *Expert Systems with Applications* 208: 118101. <https://doi.org/10.1016/j.eswa.2022.118101>
- Yigit, Y., K. Gursu, A. Al-Dubai, L. Maglaras, and B. Canberk. 2025. “Digital Twin-Enabled Lightweight Attack Detection for Software-Defined Edge Networks.” In *2025 IEEE Wireless Communications and Networking Conference (WCNC)*, 1–6. Piscataway, NJ: IEEE. <https://doi.org/10.1109/WCNC61545.2025.10978524>
- Zhao, P., W. Zhu, P. Jiao, D. Gao, and O. Wu. 2025. “Data Poisoning in Deep Learning: A Survey.” Preprint, arXiv, 27 March 2025. <https://arxiv.org/abs/2503.22759>