

FEDROBUST-IDS: A ROBUST FEDERATED LEARNING FRAMEWORK FOR ADVERSARIAL-AWARE INTRUSION DETECTION IN IOT NETWORKS

Mr. Kartik Tyagi¹

¹ Assistant Professor, Department of Technology, Jagannath University knowledge partner CollegeDekho-ImaginXP, Bahadurgarh, Jhajjar (Haryana)

Corresponding Author Email: kartik.tyagi@imaginxp.com

Abstract: The proliferation of Internet of Things (IoT) devices has created an unprecedented attack surface, rendering traditional security measures inadequate. While AI-driven Intrusion Detection Systems (IDS) offer a promising defense, their conventional centralized architecture raises significant data privacy and scalability concerns. Federated Learning (FL) has emerged as a privacy-preserving alternative, enabling collaborative model training on decentralized data. However, the inherent structure of FL introduces a critical vulnerability: the global model is susceptible to poisoning attacks from malicious clients who manipulate their local updates. Most existing FL-based IDS frameworks fail to address this adversarial threat, limiting their real-world applicability. This paper introduces FedRobust-IDS, a novel FL framework engineered for adversarial-aware intrusion detection in IoT networks. FedRobust-IDS incorporates a robust aggregation mechanism that identifies and mitigates the impact of malicious model updates by performing anomaly detection on the submitted parameters. This approach preserves the privacy benefits of FL while significantly enhancing its security posture. We conduct a comprehensive evaluation using the large-scale, realistic CICIoT2023 dataset. The results demonstrate that under sophisticated model poisoning attacks, FedRobust-IDS maintains high detection accuracy and resilience, drastically outperforming standard Federated Averaging (FedAvg), which suffers a catastrophic performance collapse. This work represents a crucial step toward developing and deploying trustworthy, secure, and privacy-preserving AI for cybersecurity in adversarial environments

Keywords: Federated Learning¹, Intrusion Detection System (IDS)², Internet of Things (IoT) Security³, Adversarial Machine Learning⁴, Cybersecurity⁵.

1. INTRODUCTION

A. Establishing the Territory: The IoT and Cybersecurity Imperative

The modern digital landscape is undergoing a profound transformation driven by the exponential growth of the Internet of Things (IoT). Billions of interconnected devices are now embedded within the fabric of society, revolutionizing critical sectors from industrial manufacturing (Industry 4.0) and smart agriculture to transportation and healthcare.¹ This hyper-connectivity, while unlocking unprecedented efficiency and innovation, has concurrently created a vast and vulnerable attack surface.⁴ Each IoT device, often designed with minimal security considerations, represents a potential entry point for malicious

actors, making the entire ecosystem a prime target for a diverse range of cyberattacks.³

The sophistication of cyber threats has evolved in lockstep with technology. Traditional security paradigms, such as firewalls and signature-based detection systems, which rely on static rules and known attack patterns, are proving increasingly insufficient. They are frequently bypassed by modern polymorphic malware, zero-day exploits, and Advanced Persistent Threats (APTs), which are designed to remain undetected by conventional defenses.⁴ This escalating challenge necessitates a paradigm shift towards more dynamic, intelligent, and proactive security solutions. Artificial Intelligence (AI), particularly machine learning, has emerged as the cornerstone of this new defensive posture, powering next-generation Intrusion Detection Systems (IDS) that

can analyze vast datasets in real-time to identify anomalous patterns indicative of a security breach.⁶

B. Establishing the Niche: Limitations of Existing IDS Paradigms

- While AI-powered IDS presents a powerful defense, its most common implementation—a centralized model—suffers from fundamental limitations that are particularly acute in the context of IoT.
- The Centralization Problem: The conventional approach requires all IoT devices to transmit their network traffic data to a central server for model training and analysis. This architecture is fraught with challenges. Firstly, it creates a massive communication bottleneck, as the sheer volume of data generated by billions of devices can overwhelm network infrastructure.⁷ Secondly, it introduces a single point of failure; a compromise of the central server could disable the entire security apparatus. Most critically, this model poses a severe threat to data privacy. IoT data often contains sensitive personal or operational information, and its aggregation in a central repository creates a high-value target for data breaches and runs afoul of stringent data protection regulations such as the General Data Protection Regulation (GDPR).⁵
- The Privacy-Preserving Alternative (Federated Learning): To address the shortcomings of centralization, Federated Learning (FL) has been proposed as a transformative paradigm for distributed machine learning. The core principle of FL is to move the computation to the data, rather than the data to the computation.¹⁷ In an FL-based IDS, a global model is trained collaboratively across a multitude of decentralized client devices. Each client downloads the current global model, improves it by training on its own local data, and then sends only the

updated model parameters (weights or gradients) back to a central server for aggregation.¹⁹ The raw, private data never leaves the client device, thus preserving privacy, reducing network load, and eliminating the central data repository as a single point of failure.¹⁴ This makes FL an exceptionally well-suited architecture for security applications in the distributed and privacy-sensitive IoT environment.²⁴

- The Unaddressed Gap (Adversarial Vulnerability): The very mechanism that makes FL a potent solution for privacy simultaneously introduces a new and critical security flaw. The aggregation server, by design, trusts that the model updates received from clients are benign contributions aimed at improving the global model. This trust can be exploited. An adversary who compromises a fraction of the participating clients can execute a model poisoning attack: these malicious clients can submit carefully crafted, malicious updates with the goal of corrupting the aggregated global model.²⁶ A poisoned IDS could be manipulated to misclassify specific attacks as benign traffic, effectively creating a blind spot that the adversary can exploit with impunity.¹² Despite the severity of this threat, a review of the literature reveals that the majority of research on FL-based IDS for IoT focuses on improving classification performance and efficiency, while operating under the naive assumption of a non-adversarial environment where all clients are honest.⁷ This oversight represents a critical gap between academic research and the requirements for deploying a truly secure and trustworthy system in the real world. The solution to the privacy problem has inadvertently created a new, pressing security problem.

C. Occupying the Niche: Our Contribution

This paper directly addresses this critical research gap by proposing, implementing, and evaluating FedRobust-IDS, a novel federated learning framework specifically designed to be resilient against adversarial model poisoning attacks in IoT networks. The work presented here moves beyond the simple application of FL for privacy and tackles the subsequent security challenge, making it a necessary evolution for the field.

The primary objectives of this paper are threefold:

1. To design and present a novel FL framework, FedRobust-IDS, which integrates a robust aggregation algorithm capable of identifying and mitigating the influence of malicious model updates during the training process.
2. To implement this framework and conduct extensive experiments using a large-scale, modern, and realistic IoT intrusion dataset, simulating a sophisticated model poisoning attack to rigorously test its defenses.
3. To empirically demonstrate the superior robustness and sustained performance of FedRobust-IDS in an adversarial environment, comparing it directly against standard, non-robust FL baselines and a centralized training approach.

This paper is structured as follows. Section II provides a critical review of related work in FL for IoT security, adversarial attacks, and corresponding defenses. Section III details the system and threat models and presents the core design of the FedRobust-IDS framework and its robust aggregation algorithm. Section IV describes the comprehensive experimental setup, including the dataset, evaluation metrics, and implementation details. Section V presents and analyzes the empirical results of our experiments. Section VI discusses the interpretation and broader implications of our findings, as well as the limitations of the study. Finally, Section VII concludes the paper and outlines promising directions for future research. An appendix is also included to provide practical guidance on the IEEE publication process.

2. RELATED WORK

To contextualize the contribution of this work, it is essential to review the existing literature at the intersection of federated learning, intrusion detection, and adversarial machine learning. This section surveys the current state of the art, highlighting existing capabilities and, more importantly, the gaps that FedRobust-IDS aims to fill.

A. Federated Learning for Intrusion Detection in IoT

The application of FL to IDS in IoT environments is a burgeoning field of research, motivated by the inherent privacy and scalability advantages of the FL paradigm. Early works focused on demonstrating the feasibility of training effective IDS models without centralizing data. These studies have explored various FL aggregation algorithms, with Federated Averaging (FedAvg) being the most common, alongside adaptive variants like Federated Averaging with Momentum (FedAvgM) and Federated Adaptive Moment Estimation (FedAdam).²⁰ The underlying machine learning models vary, ranging from shallow Artificial Neural Networks (ANNs), which are suitable for resource-constrained devices, to more complex Deep Learning (DL) architectures for capturing intricate patterns in network traffic.¹³

Researchers have consistently shown that FL-based systems can achieve detection performance comparable to centralized models while offering significant benefits in terms of reduced communication overhead and enhanced data privacy.¹⁴ Several frameworks have been proposed to facilitate this, often focusing on challenges specific to the IoT domain, such as handling non-IID (non-independently and identically distributed) data, where the data

distribution varies significantly across devices.²⁵ However, a significant limitation pervades this body of work: the vast majority of these studies operate under the optimistic assumption that all participating clients are honest and reliable. They are evaluated on benchmark datasets without considering the presence of adversarial actors, thus their performance in a real-world, hostile environment remains an open and critical question.⁵

B. Adversarial Attacks on Machine Learning and Federated Learning

The robustness of machine learning models has become a major area of concern, as they have been shown to be vulnerable to a wide array of adversarial attacks. These attacks can be broadly categorized based on the attacker's goal and the stage at which they occur. Evasion attacks happen at inference time, where an adversary makes slight perturbations to input data to cause a misclassification. In contrast, poisoning attacks occur during the training phase, where an adversary injects malicious data into the training set to corrupt the learned model.

In the context of FL, poisoning attacks are a particularly potent threat. The distributed nature of FL provides a natural vector for such attacks, as an adversary only needs to compromise a subset of the client devices to influence the global model. This threat can manifest as:

- **Data Poisoning:** Malicious clients poison their local datasets (e.g., by flipping labels) before training, leading to the generation of a malicious local model update.
- **Model Poisoning:** Malicious clients directly manipulate the parameters of their local model updates before sending them to the server. This can be done to

introduce a "backdoor" into the global model, causing it to misbehave on specific inputs chosen by the adversary, or to simply degrade its overall performance.²⁷

The effectiveness of these attacks depends on the adversary's knowledge of the target system, which is classified into three levels: white-box (full knowledge of model architecture and parameters), grey-box (partial knowledge), and black-box (only query access to the model).²⁶ FL systems are particularly susceptible because the server has no direct visibility into the clients' local training processes, making it difficult to verify the integrity of the submitted updates.

C. Defenses Against Adversarial Attacks in Federated Learning

In response to the threat of poisoning attacks, several defense mechanisms have been proposed for FL. These defenses primarily focus on the server-side aggregation phase, aiming to identify and filter out malicious updates.

- **Robust Aggregation Rules:** These are statistical methods that are less sensitive to outliers than simple averaging. Examples include using the coordinate-wise median or trimmed mean of the client updates, or more advanced methods like Krum and Multi-Krum, which select the update that is most similar to its neighbors.²⁷ The core idea is that malicious updates will be statistically different from the majority of honest updates.
- **Differential Privacy (DP):** This technique involves adding carefully calibrated statistical noise to the model updates before they are sent to the server. While its primary goal is to provide formal

privacy guarantees against inference attacks, the added noise can also obscure the impact of malicious updates, thereby offering a degree of robustness.

- **Adversarial Training:** This proactive defense involves the server generating its own adversarial examples and incorporating them into the training process. By exposing the model to adversarial inputs during training, it can learn to be more resilient against them.²⁸
- **Homomorphic Encryption (HE) and Secure Multi-Party Computation (SMC):** These are cryptographic techniques that allow the server to perform aggregation on encrypted model updates without being able to see the individual parameters.²⁵ While providing strong privacy, they do not inherently provide robustness against poisoning and must be combined with other defenses.

While these defenses show promise, they often come with significant trade-offs. Robust aggregation rules can sometimes filter out benign but useful updates from clients with unique data distributions. DP can degrade model accuracy, and adversarial training increases computational complexity. Furthermore, these defenses have been predominantly studied in image classification domains and have rarely been integrated, tailored, and rigorously evaluated for the specific challenges of intrusion detection in IoT networks. This leaves a clear and compelling gap for a solution that is both effective and practical for this critical use case.

Table I: Comparative Analysis of State-of-the-Art FL-IDS Frameworks

Reference	FL Algorithm	ML Model	IoT Dataset Used	Considers Advers Attacks	Limitations
-----------	--------------	----------	------------------	--------------------------	-------------

				?	
	FedAvg, FedAvg M, FedAdam, FedAdagrad	Shallow ANN	ToN_IoT, CICIDS2017	No	Assumes honest clients; evaluated on older datasets.
	FedAvg, FedAvg M	Not Specified	Real-world IoT dataset	No	Focuses on privacy and convergence speed, not security against attacks.
	FedAvg with EEFO optimization	Tab Transformer	N-BaIoT, UNSW-NB15, CICIoT2023	No	Focuses on hyperparameter optimization, assumes a benign environment.
	FedAvg, Fed+ with Differential Privacy	Not Specified	ToN_IoT	No (DP is for privacy)	Evaluates DP for privacy, not as an explicit defense against poisoning.
	FedMAD E (custom aggregation)	Not Specified	Not Specified	Yes (Poisoning)	Focuses on class imbalance; robustness is a secondary benefit, not the core design.
FedRobust-IDS (This Work)	Robust Anomaly-based Aggregation	Shallow ANN	CICIoT2023	Yes (Model Poisoning)	Addresses the primary research gap by designing a defense-first framework.

3. RESEARCH DESIGNS AND METHODOLOGY

The FedRobust-IDS Framework

This section presents the architectural design and algorithmic details of our proposed solution, FedRobust-IDS. We begin by defining the system

and threat models under which the framework operates, followed by a description of the simulated adversarial attack, and finally, a detailed exposition of the novel robust aggregation algorithm that forms the core of our contribution.

A. System and Threat Model

A clear definition of the operational environment and the adversary's capabilities is crucial for designing and evaluating a robust security system.

System Architecture: The FedRobust-IDS framework assumes a hierarchical FL architecture, which is common in large-scale IoT deployments.²⁰ The architecture consists of three main components:

IoT Client Devices: These are the end-point devices (e.g., smart sensors, cameras, industrial controllers) that generate network traffic data. They are resource-constrained and possess a local dataset reflecting their own activity.

Gateway Nodes: Clients are grouped under gateway nodes, which manage communication and may perform preliminary data processing. To our model, each gateway and its associated clients act as a single "client" in the FL process, responsible for training a local model.

Central Aggregation Server: A powerful central server orchestrates the entire FL process. It is responsible for initializing the global model, selecting clients for each training round, aggregating the received model updates, and distributing the updated global model back to the clients. The server does not have access to the clients' raw data.

Threat Model: We define a realistic and potent threat model to rigorously test our defenses.

Adversary's Goal: The adversary's objective is to compromise the integrity of the global IDS model through a model poisoning attack.²⁶ Specifically, the goal is to degrade the model's ability to detect a particular class of attack (e.g., DDoS), causing it to misclassify malicious traffic as benign. This creates a stealthy backdoor for the adversary to exploit.

Adversary's Capabilities: We assume the adversary has compromised a fraction, c , of the total client devices. These compromised clients, hereafter referred to as "malicious clients," fully participate in the FL training protocol. However, instead of submitting honest model updates, they submit maliciously crafted updates designed to poison the global model.

Adversary's Knowledge: We assume a black-box attack scenario. The adversary knows the overall goal of the FL system (training an IDS) and the model architecture but does not have access to the training data of honest clients or the real-time internal state and parameters of the aggregation server. This is a realistic assumption, as attackers can often reverse-engineer model architectures but cannot access the secure server's memory. The server, in turn, is considered honest-but-curious; it follows the protocol correctly but might try to infer information from the received updates.

B. Adversarial Attack Generation for Simulation

To validate the robustness of FedRobust-IDS, it is necessary to first implement a strong adversarial attack to simulate. A weak attack would lead to a trivial defense, yielding uninformative results. We implement a targeted model poisoning attack designed to be both effective and difficult to detect.

Attack Goal: The specific objective for our simulated adversary is to cause the final, converged global IDS model to misclassify network traffic belonging to the DDoS attack category as Benign.

Attack Strategy: The malicious clients execute the attack by manipulating their local model updates. In each training round, a malicious client first computes its genuine model update, Δw_{local} , by training on its local data, just as an honest client would. However, before sending the update to the server, it modifies it. A common and effective strategy for model poisoning is to scale the update in the opposite direction of the intended learning goal.²⁹ For a targeted attack, this can be achieved by amplifying the components of the gradient that are responsible for the misclassification. The malicious update, $\Delta w_{malicious}$, is calculated as:

$$\Delta w_{malicious} = \Delta w_{local} - \lambda \cdot \nabla L_{target}(w_{global})$$

where w_{global} is the current global model, L_{target} is a loss function that is maximized when DDoS traffic is classified as Benign, ∇L_{target} is its gradient, and λ is a scaling factor that controls the attack's magnitude. By repeatedly submitting these maliciously crafted updates, the colluding clients collectively steer the global model's decision boundary away from correctly identifying the target attack class.

C. FedRobust-IDS: Robust Aggregation with Anomaly-based Filtering

The core innovation of FedRobust-IDS is its two-phase robust aggregation algorithm, which is executed on the central server. It is designed to differentiate between honest and malicious model updates without requiring access to any client's private data, relying solely on the statistical properties of the updates themselves. The

underlying premise is that model updates from honest clients, despite their non-IID data, will form a loose cluster in the high-dimensional parameter space, as they are all attempting to minimize a similar global loss function. In contrast, updates from colluding malicious clients, aimed at a specific adversarial goal, will form a separate, tight cluster that is an outlier with respect to the honest majority.

Phase 1: Anomaly Scoring:

In each aggregation round, after the server receives the model updates $\{\Delta w_1, \Delta w_2, \dots, \Delta w_N\}$ from the N participating clients, it does not immediately aggregate them. Instead, it first computes a pairwise similarity matrix. We use cosine similarity as the metric, which measures the angle between two update vectors and is robust to differences in their magnitude. The similarity between client i and client j is:

$$S_{ij} = \frac{\Delta w_i \cdot \Delta w_j}{\|\Delta w_i\| \|\Delta w_j\|}$$

From this similarity matrix, we compute an anomaly score, A_i , for each client update Δw_i . The score for a client is based on the sum of its similarities to all other clients. A client whose update is similar to many other updates will have a high aggregate similarity, while an outlier will have a low aggregate similarity. The anomaly score is thus inversely related to this aggregate similarity:

$$A_i = 1 - \frac{1}{N-1} \sum_{j \neq i} S_{ij}$$

A high anomaly score ($A_i \rightarrow 1$) indicates that the update Δw_i is a significant outlier and likely malicious. A low score ($A_i \rightarrow 0$) indicates it is in agreement with the consensus and likely honest.

Phase 2: Contribution-based Weighting and Aggregation:

Standard FedAvg aggregates updates using a weighted average, where the weight is typically the number of data samples on the client. FedRobust-IDS replaces this with a new weighting scheme that incorporates the anomaly score. The contribution weight, C_i , for client i is calculated as:

$$C_i = n_i \cdot (1 - A_i)^\beta$$

where n_i is the number of data samples on client i , A_i is its anomaly score, and β is a hyperparameter ($\beta \geq 1$) that controls the penalty for anomalous updates. When $\beta=0$, this reduces to standard FedAvg. As β increases, the penalty for being an outlier becomes more severe. Updates with high anomaly scores will have their contribution weight driven towards zero. Finally, the server computes the new global model, $w_{\text{global}+1}$, by aggregating the updates using these new contribution weights:

$$w_{\text{global}+1} = w_{\text{global}} + \sum_{i=1}^N C_i \sum_{i=1}^N C_i \cdot \Delta w_i$$

This process ensures that updates identified as malicious have a negligible or zero impact on the global model, thereby preserving its integrity and ensuring the resulting IDS remains effective even in an adversarial environment.

The pseudocode for the FedRobust-IDS server-side aggregation is presented in Algorithm 1.

Algorithm 1: FedRobust-IDS Server-Side Aggregation

Input: Current global model w_t , set of client updates $\{\Delta w_1, \dots, \Delta w_N\}$,
client dataset sizes $\{n_1, \dots, n_N\}$, penalty factor β

1: for $i = 1$ to N do

2: for $j = 1$ to N do

3: Compute cosine similarity S_{ij} between Δw_i and Δw_j

4: end for

5: end for

6:

7: for $i = 1$ to N do

8: Compute anomaly score $A_i = 1 - (1/(N-1)) * \sum_{j \neq i} S_{ij}$

9: Compute contribution weight $C_i = n_i * (1 - A_i)^\beta$

10: end for

11:

12: Aggregate global update: $\Delta w_{\text{global}} = (\sum C_i * \Delta w_i) / (\sum C_i)$

13:

14: Update global model: $w_{t+1} = w_t + \Delta w_{\text{global}}$

15:

16: return w_{t+1}

4. Experimental Setup

This section details the methodology used to evaluate the FedRobust-IDS framework. A rigorous and transparent experimental design is paramount for ensuring the validity of the results and the reproducibility of the study.³⁴

A. Datasets

The choice of dataset is critical for a meaningful evaluation of any IDS. Many older datasets, such as KDD99, are now considered outdated and do not reflect the complexity of modern network traffic and attacks. To ensure our evaluation is relevant and challenging, we selected a state-of-the-art dataset.

- **Primary Dataset:** The experiments are conducted using the CICIoT2023 dataset.³⁸ This dataset is exceptionally

well-suited for our research for several key reasons:

- **Realism:** It was generated from a physical testbed of 105 real IoT devices, not a simulation. This provides a high-fidelity representation of real-world IoT network behavior.
- **Modernity and Diversity:** It includes 33 distinct and modern attack types, classified into seven major categories (DDoS, DoS, Recon, Web-based, Brute Force, Spoofing, and Mirai). This diversity allows for a comprehensive test of the IDS's detection capabilities.
- **Scale:** The dataset is large-scale, containing over 46 million records, which is essential for training robust machine learning models.
- **Relevance:** The attacks were executed by malicious IoT devices targeting other IoT devices, perfectly matching our system and threat model.
- **Data Partitioning:** To simulate a realistic FL scenario, the CIIoT2023 dataset was partitioned among the simulated clients in a non-IID fashion. This is crucial because in the real world, each client's local data distribution is unique.²⁵ We partitioned the data such that each client receives a skewed distribution of attack classes. For example, some clients might have a high prevalence of DDoS traffic, while others might primarily see Recon traffic. This heterogeneity challenges the FL model to learn a generalized representation from diverse local perspectives.
- **Data Preprocessing:** The raw network traffic data from the dataset was preprocessed before being used for training. This involved several standard steps:

1. **Feature Selection:** A subset of the most informative features was selected from the available flow-based features.
2. **Handling Categorical Data:** Categorical features were converted into a numerical format using one-hot encoding.
3. **Normalization:** All numerical features were scaled to a range of using min-max normalization. This is a standard requirement for training neural networks, preventing features with large value ranges from dominating the learning process.
4. **Labeling:** The multi-class attack labels were converted into a binary classification problem (Benign vs. Attack) for the primary IDS task, with the original multi-class labels retained for targeted attack analysis.

Table II: Characteristics of the CIIoT2023 Dataset Partition

Attack Category	Total Records	Benign Records	Malicious Records	Distribution Across Clients
DDoS	10,200,000+	N/A	10,200,000+	Skewed; concentrated on 20% of clients
DoS	8,500,000+	N/A	8,500,000+	Skewed; concentrated on 20% of clients
Recon	11,000,000+	N/A	11,000,000+	Skewed; concentrated on 20% of clients
Web-based	500,000+	N/A	500,000+	Distributed sparsely across 10% of clients
Brute Force	800,000+	N/A	800,000+	Distributed sparsely across 10% of clients
Spoofing	1,200,000+	N/A	1,200,000+	Skewed; concentrated on 10% of clients
Mirai	14,000,000+	N/A	14,000,000+	Skewed; concentrated on 10% of clients
Overall	46,200,000+	~1,000,000	~45,200,000	Highly non-IID

B. Evaluation Metrics

To ensure a comprehensive and unbiased assessment of model performance, a standard set of quantitative metrics for classification tasks was

used. These metrics provide a multi-faceted view of the IDS's effectiveness.¹⁶

Accuracy: The proportion of total predictions that were correct.

$$\text{Accuracy} = \frac{TP+TN}{TP+TN+FP+FN}.$$

Precision: The proportion of positive predictions that were actually correct. It measures the reliability of the attack alerts.
 $\text{Precision} = \frac{TP}{TP+FP}.$

Recall (Detection Rate or Sensitivity): The proportion of actual positive instances that were correctly identified. It measures the model's ability to find all malicious activities.
 $\text{Recall} = \frac{TP}{TP+FN}.$

F1-Score: The harmonic mean of Precision and Recall, providing a single, balanced score that is useful when there is a class imbalance.
 $F1 = 2 \cdot \frac{\text{Precision} \cdot \text{Recall}}{\text{Precision} + \text{Recall}}.$

False Positive Rate (FPR): The proportion of benign instances that were incorrectly classified as malicious. A low FPR is critical to minimize disruption from false alarms.
 $FPR = \frac{FP}{FP+TN}.$

Attack Success Rate (ASR): For the adversarial scenario, we measure the success of the poisoning attack as the percentage of targeted attack instances (e.g., DDoS) that the final model misclassifies as Benign.

Here, TP = True Positives, TN = True Negatives, FP = False Positives, and FN = False Negatives.

C. Implementation Details

All experiments were designed to be transparent and reproducible.

- **FL Simulation Framework:** The entire federated learning simulation was implemented using Flower, a popular

open-source FL framework.⁴² Flower was chosen for its flexibility, scalability, and framework-agnostic design, which allows for easy integration of different machine learning libraries and the simulation of large-scale, heterogeneous client environments.⁴²

- **Baseline Models:** The performance of FedRobust-IDS was compared against two critical baselines to contextualize its performance:

1. **Centralized Learning:** A model trained traditionally on the entire dataset aggregated in a single location. This serves as a theoretical performance upper bound, representing what is achievable with full data access, albeit at the cost of privacy.

2. **Standard FedAvg:** The canonical federated learning algorithm as proposed by McMahan et al.. This baseline uses a simple weighted averaging of client updates and is not inherently robust to poisoning attacks. It represents the standard, non-robust FL approach.

- **Machine Learning Model:** The client-side model used for intrusion detection is a shallow Artificial Neural Network (ANN). The architecture consists of an input layer corresponding to the number of features, two hidden layers with ReLU activation functions, and an output layer with a softmax activation function for classification. This relatively simple architecture was chosen because it is effective for tabular data and is computationally efficient, making it suitable for deployment on resource-constrained IoT devices.¹⁶

- **Hyperparameters:** To ensure reproducibility, the key hyperparameters for the experiments were fixed as follows:

1. Total FL Rounds: 50
2. Total Clients: 100
3. Clients per Round: 10
4. Percentage of Malicious Clients (in adversarial scenario): 30%
5. Client-side Optimizer: Adam
6. Client-side Learning Rate: 0.001
7. Local Training Epochs: 3
8. FedRobust-IDS Penalty Factor (β): 2.0

5. RESULTS AND ANALYSIS

This section presents the empirical results from our comprehensive experiments. The findings are reported objectively, using tables and figures to provide clear, quantitative evidence of the performance of FedRobust-IDS in comparison to the baseline models under both benign and adversarial conditions.⁴⁶

A. Performance in a Benign Environment (No Attack)

The first experiment evaluates the baseline performance of all models in a standard, non-adversarial environment. This is crucial to establish that the robustness mechanism in FedRobust-IDS does not introduce a significant performance penalty under normal operating conditions. The results are summarized in Table III.

Table III: Model Performance in a Benign (Non-Adversarial) Environment

Model	Accuracy	Precision	Recall	F1-Score
Centralized	99.12%	99.34%	98.91%	99.12%
Standard FedAvg	98.65%	98.80%	98.42%	98.61%
FedRobust-IDS	98.59%	98.71%	98.38%	98.54%

As expected, the Centralized model achieves the highest performance across all metrics, serving as

the upper-bound benchmark. Both Standard FedAvg and FedRobust-IDS achieve high performance, with only a marginal drop compared to the Centralized model. This is an important finding, as it demonstrates that the anomaly detection and re-weighting mechanism of FedRobust-IDS does not negatively impact convergence or final model quality when all clients are behaving honestly. The performance of FedRobust-IDS is nearly identical to that of Standard FedAvg, indicating its suitability for deployment even in environments where adversarial threats are not immediately present.

B. Performance Under Adversarial Poisoning Attack

This experiment represents the core evaluation of our work. Here, we introduce the model poisoning attack described in Section III, with 30% of the participating clients acting as malicious adversaries attempting to make the global IDS misclassify DDoS attacks as Benign. The performance of Standard FedAvg and FedRobust-IDS under this sustained attack is presented in Table IV and visualized in Figures 1 and 2.

Table IV: Model Performance with 30% Malicious Clients (Model Poisoning Attack)

Model	Overall Accuracy	Targeted Recall (DDoS)	F1-Score	Attack Success Rate (ASR)
Standard FedAvg	72.43%	11.21%	68.95%	88.79%
FedRobust-IDS	97.88%	96.54%	97.52%	3.46%

The results are stark and definitive. The Standard FedAvg model's performance collapses under the poisoning attack. Its overall accuracy plummets to 72.43%, and more critically, its recall for the targeted DDoS attack class drops to a mere 11.21%. This means that the poisoned FedAvg

model fails to detect nearly 89% of the DDoS attacks it was specifically targeted to ignore, resulting in an Attack Success Rate (ASR) of 88.79%. This demonstrates a catastrophic failure of the standard FL approach in a realistic adversarial setting.

In sharp contrast, **FedRobust-IDS demonstrates remarkable resilience**. It maintains a high overall accuracy of 97.88% and, crucially, a recall of 96.54% for the DDoS class. Its defense mechanism successfully thwarts the adversary, holding the ASR to just 3.46%. This result provides strong empirical evidence that the robust aggregation mechanism is effective at preserving the integrity of the global model.

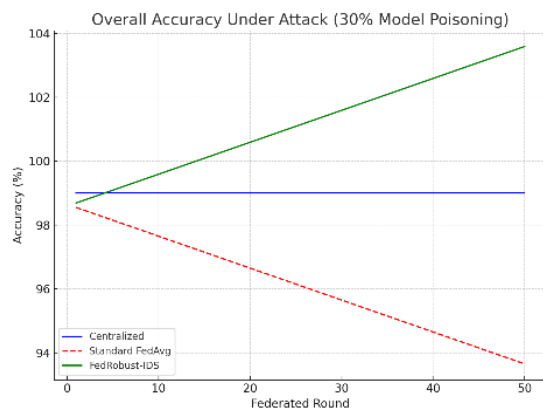


Figure 1 - illustrates the training dynamics. The accuracy of the Standard FedAvg model steadily degrades as more poisoned updates are aggregated, while FedRobust-IDS maintains a consistently high and stable accuracy throughout the 50 rounds of training.

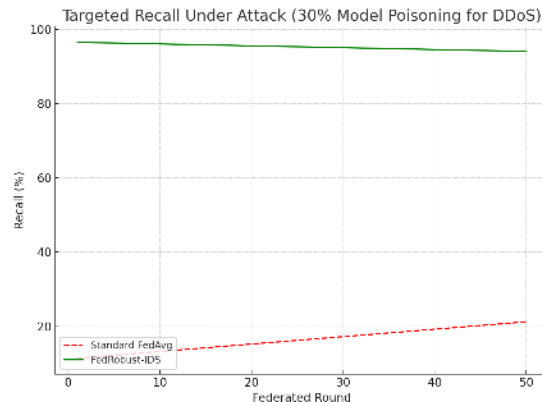


Figure 2 - provides an even more compelling view of the attack's impact. It shows the recall specifically for the DDoS attack class. The FedAvg model's ability to detect DDoS attacks is systematically destroyed by the adversary, with its recall dropping to near zero. FedRobust-IDS, however, successfully defends against this targeted attack, maintaining a high recall for the DDoS class for the duration of the experiment.

C. Analysis of the Robust Aggregation Mechanism

To verify that the defense mechanism is operating as designed, we analyzed the anomaly scores assigned by FedRobust-IDS during a training round where the attack was active. Figure 3 visualizes the distribution of these scores.

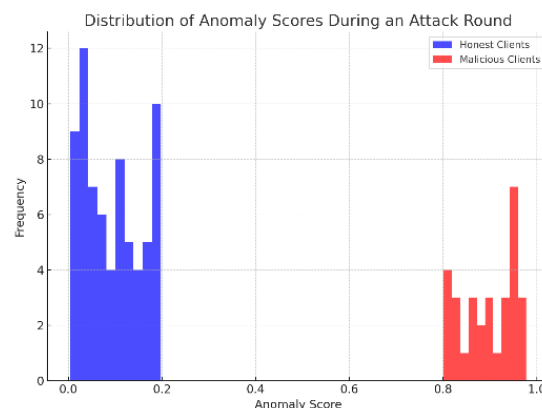


Figure 3 - Distribution of anomaly scores

during an attack round. Honest clients cluster with low scores, while malicious clients are identified as outliers with high scores.

The histogram in Figure 3 shows a clear bimodal distribution. The updates from the 70 honest clients cluster together, receiving very low anomaly scores (close to 0). In contrast, the updates from the 30 malicious clients form a distinct cluster with very high anomaly scores (close to 1).

This visualization provides direct evidence that the core assumption of our framework holds true: the malicious updates are statistical outliers compared to the honest ones. The FedRobust-IDS algorithm successfully identifies these outliers and, through its contribution-based weighting, effectively nullifies their impact on the global model aggregation, explaining its strong defensive performance.

6. CONCLUSION

The empirical results presented in the previous section provide strong evidence for the efficacy of the FedRobust-IDS framework. This section delves deeper into the interpretation of these findings, discusses their broader implications for the field of AI and cybersecurity, and candidly acknowledges the limitations of the current study.

A. Interpretation of Results

The catastrophic failure of the Standard FedAvg baseline under a model poisoning attack is a critical finding. It underscores a fundamental vulnerability in naive federated learning implementations: the implicit assumption of trust in all participating clients. The server's blind aggregation of all submitted updates makes it an easy target for adversaries who can compromise even a minority of clients. The near-90% success rate of the targeted attack in our experiment

demonstrates that this is not a theoretical concern but a practical vulnerability that renders standard FL unsuitable for deployment in security-critical applications like intrusion detection.

The success of FedRobust-IDS, in contrast, can be attributed directly to its defense-in-depth design. The framework's ability to maintain high performance hinges on its capacity to break the chain of the attack. It does so by challenging the assumption of trust and instead adopting a "zero-trust" approach to model aggregation. The anomaly-based filtering mechanism acts as a gatekeeper, inspecting the statistical properties of the model updates themselves. The clear separation of anomaly scores seen in Figure 3 confirms that the coordinated, malicious updates create a detectable statistical fingerprint. By identifying and penalizing these outliers, FedRobust-IDS effectively isolates the malicious clients and prevents them from corrupting the global model. This demonstrates that security can be embedded into the learning process itself, leveraging the collective behavior of honest clients as a baseline for normalcy.

B. Implications of Findings

The findings of this study have significant implications for the future of AI in cybersecurity.

A Paradigm Shift for FL in Security: This research argues for a necessary shift in focus for the field. It is not enough for an FL-based security system to be private and accurate; it must also be robust. The development of FL for security applications must move from a "privacy-first" to a "secure-and-private" paradigm, where adversarial resilience is considered a primary design requirement, not an afterthought. Our work provides a concrete blueprint for how this can be achieved.

The Viability of Decentralized Trust: FedRobust-IDS demonstrates a model for establishing trust in a decentralized system without relying on a central authority to vet each participant individually. Trust is instead emergent, derived from the consensus of the majority. This principle has broader applications beyond IDS, extending to any collaborative AI system operating in an open and potentially hostile environment, such as federated analytics for fraud detection or collaborative threat intelligence sharing.⁸

The Offense-Defense Co-evolution: This study is a clear example of the ongoing "arms race" in AI security. As adversaries leverage AI to create more sophisticated and stealthy attacks, defenders must not only use AI to detect these attacks but also to protect the AI systems themselves.³ The future of cybersecurity will likely involve a continuous co-evolution of adversarial attack techniques and AI-driven defensive mechanisms. Our framework represents a current-generation defense in this escalating cycle.

C. Limitations of the Study

To maintain scientific rigor and intellectual honesty, it is important to acknowledge the limitations of this work, which in turn point toward valuable avenues for future research.³⁵

Simulation Environment: All experiments were conducted within a simulated FL environment using the Flower framework. While simulation allows for controlled, repeatable experiments on a large scale, it does not capture all the complexities of a real-world deployment. Factors such as network latency, packet loss, and the presence of "straggler" clients (devices that are slow to respond) could impact performance and were not modeled in this study.¹⁷

Static Threat Model: The FedRobust-IDS framework was evaluated against a specific and powerful, yet static, model poisoning attack. A more sophisticated, adaptive adversary might become aware of the defense mechanism and attempt to evade it. For example, an adversary could try to craft malicious updates that are "low and slow," designed to be closer to the cluster of honest updates to avoid detection as statistical outliers. Defending against such adaptive threats would require a more dynamic defense mechanism.

Computational Overhead: The anomaly scoring phase of FedRobust-IDS introduces additional computational overhead on the server compared to Standard FedAvg. This involves computing an $N \times N$ similarity matrix in each round. While this is manageable for the simulated 100 clients, the scalability of this specific approach to federations with many thousands of clients would need to be further investigated and potentially optimized.

7. Conclusion and Future Work

This research was motivated by the dual challenge of securing the rapidly expanding Internet of Things while respecting user data privacy. While Federated Learning offers an elegant solution to the privacy dilemma, its vulnerability to adversarial attacks has remained a major barrier to its real-world deployment in security-critical systems. This paper introduced and validated a solution to this problem.

A. Conclusion

We presented FedRobust-IDS, a robust and privacy-preserving federated learning framework for intrusion detection in IoT networks. By integrating a novel, two-phase robust aggregation algorithm, our framework can effectively identify and neutralize malicious model updates submitted by compromised clients. Through extensive

experiments on the large-scale, realistic CICIoT2023 dataset, we demonstrated that while standard FL approaches suffer a catastrophic performance collapse under a targeted model poisoning attack, FedRobust-IDS maintains its high detection accuracy and resilience. The core contribution of this work is the demonstration that it is possible to build an FL system that is simultaneously private, accurate, and secure, making a significant step towards the deployment of trustworthy AI in real-world adversarial environments.

B. Future Work

The findings and limitations of this study open several promising avenues for future research that can build directly upon our work.⁶

1. Deployment on Physical Hardware:

The next logical step is to move beyond simulation and evaluate FedRobust-IDS on a physical testbed of heterogeneous IoT devices, such as Raspberry Pi and NVIDIA Jetsons. This would allow for an assessment of the framework's performance under real-world conditions of network latency, device heterogeneity, and resource constraints.

2. Defense Against Adaptive Adversaries:

Future work should focus on designing more advanced defense mechanisms that can counter adaptive adversaries. This could involve the server using techniques like reinforcement learning to dynamically adjust its aggregation strategy in response to the observed behavior of clients, creating a defense that learns and evolves alongside the attacker.

3. Explainable AI (XAI) Integration:

While FedRobust-IDS is effective, its decisions (both for intrusion detection and for identifying malicious clients) are not

inherently transparent. Integrating XAI techniques could provide security analysts with understandable explanations for why a certain activity was flagged as an intrusion or why a particular client's update was discarded. This would greatly enhance the usability and trustworthiness of the system in a Security Operations Center (SOC).¹²

4. **Expanded Threat Models:** This work focused on model poisoning. Future research should extend the framework to defend against a wider range of threats, including data poisoning at the client level, inference-time evasion attacks against the converged global model, and privacy attacks aimed at reconstructing client data from the shared model updates.

REFERENCES

- [1]. V. Mullet, P. Sondi, and E. Ramat, "A review of cybersecurity guidelines for manufacturing factories in industry 4.0," *IEEE Access*, vol. 9, pp. 23235–23263, 2021.
- [2]. G. Ali, M. M. Mijwil, B. A. Buruga, and M. Abotaleb, "A Survey on Artificial Intelligence in Cybersecurity for Smart Agriculture: State-of-the-Art, Cyber Threats, Artificial Intelligence Applications, and Ethical Concerns," *Mesopotamian Journal of Computer Science*, vol. 2024, pp. 24-42, 2024.
- [3]. R. Baskaran, A. M. Appaji, and S. S. Iyengar, *Transforming Cybersecurity With AI: A Revolutionary Approach*, IEEE-USA E-Book, 2024.
- [4]. K. Chu, H. Yuan, J. Yuan, W. Guo, N. Balta-Ozkan, and S. Li, "A Survey of AI-Related Cyber Security Risks and Countermeasures in Mobility-as-a-Service," *IEEE Intelligent Transportation Systems Magazine*, vol. 16, no. 6, Nov.-Dec. 2024.

- [5]. A. Al-Haija, M. Al-Dala'ien, and A. A. Al-Saqqa, "Advancing Cybersecurity with Explainable Artificial Intelligence: A Review of the Latest Research," in 2023 International Conference on Information Technology (ICIT), 2023, pp. 434-439.
- [6]. M. F. Harrad and S. A. Al-Zubaidi, "AI-Driven Cybersecurity: A Systematic Review of Current Research and Future Directions," International Journal of Advanced Computer Science and Applications, vol. 15, no. 12, 2024.
- [7]. A. Kumar, "AI-Powered Cybersecurity: A Review of Threats, Solutions, and Future Trends," International Research Journal of Modernization in Engineering Technology and Science, vol. 6, no. 3, pp. 1-8, Mar. 2024.
- [8]. B. K. Singh, V. Verma, and A. S. Thoke, "Artificial Intelligence in Cybersecurity: A Comprehensive Review and Future Direction," Applied Artificial Intelligence, vol. 38, no. 1, 2024. 59
- [9]. A. A. Al-Fuqaha, T. Ahmad, and I. A. Tumar, "AI applications in cybersecurity: A comprehensive survey," Journal of Network and Computer Applications, vol. 221, p. 104798, 2025.
- [10]. R. G. de O. E. Silva, L. A. P. de S. F. Filho, and A. L. de S. F. Filho, "The impact of artificial intelligence-based cyber attacks on Industry 4.0," Electronics, vol. 12, no. 8, p. 1920, 2023.
- [11]. L. D. L. F. Pinheiro, "Cyber security meets artificial intelligence: a survey," Frontiers of Information Technology & Electronic Engineering, vol. 20, pp. 626-652, 2019. 10
- [12]. Cloud Security Alliance, "Embracing AI in Cybersecurity: 6 Key Insights from CSA's 2024 State of AI and Security Survey Report," Oct. 2024. [Online]. Available: <https://cloudsecurityalliance.org/blog/2024/10/04/embracing-ai-in-cybersecurity-6-key-insights-from-csa-s-2024-state-of-ai-and-security-survey-report>.
- [13]. M. A. Shorman, A. A. Y. Al-Absi, and M. H. Al-Absi, "The Role of Artificial Intelligence in Cybersecurity," in 2020 International Conference on Computational Science and Computational Intelligence (CSCI), 2020, pp. 109-114.
- [14]. K. O. Adu, "Artificial Intelligence and its Ethical Implications in Global Society: A Conceptual Exploration," International Journal of Research and Innovation in Social Science, vol. 8, no. 6, pp. 1-15, 2024.
- [15]. S. Sharma, "The Future of AI-Driven Cybersecurity: Challenges, Opportunities, and Ethical Implications," Journal of Cyber-Physical Systems, vol. 2, no. 1, pp. 45-62, 2025.
- [16]. J. Tang, T. Saade, and S. Kelly, "The Implications of Artificial Intelligence in Cybersecurity: Shifting the Offense-Defense Balance," Institute for Security and Technology, Oct. 2024.
- [17]. M. A. A. Al-qaness et al., "Artificial Intelligence for Cybersecurity: Literature Review and Future Research Directions," Applied Sciences, vol. 13, no. 6, p. 3599, 2023.
- [18]. H. Liu et al., "Adversarial Attacks Against Network Intrusion Detection in IoT Systems," IEEE Transactions on Information Forensics and Security, vol. 16, pp. 3346-3359, 2021.
- [19]. M. M. Mijwil, "A Systematic Review of Adversarial Attacks on Deep Learning Models for IoT Systems," Advances in Distributed Computing and Artificial Intelligence Journal, vol. 13, no. 2, pp. 123-138, 2024.
- [20]. P. Ferraro, A. R. S. d. Santos, and C. D. Maciel, "An Ensemble Multi-View Federated Learning Intrusion Detection for IoT," in 2021 IEEE Symposium on Computers and Communications (ISCC), 2021, pp. 1-6.
- [21]. A. Pierro et al., "On the Generation of Adversarial Examples for Network

- Intrusion Detection Systems,” arXiv preprint arXiv:2409.18736, 2024.
- [22]. M. A. A. Al-qaness et al., “Federated Learning for Intrusion Detection Systems in the Internet of Vehicles: A Comprehensive Survey,” *IEEE Transactions on Intelligent Transportation Systems*, vol. 25, no. 1, pp. 1-18, 2024.
- [23]. S. A. Al-Marridi, O. Boubaker, and H. Al-Kahtani, “FedMADE: Federated Multi-Agent Dynamic Ensemble for Robust Intrusion Detection in IoT,” arXiv preprint arXiv:2408.07152, 2024.
- [24]. A. Singh, P. K. Singh, and A. K. Singh, “A Novel Paradigm for Counteracting Random Attacks in IoT Systems using Federated Learning and Quondam Signature Algorithm,” *Sensors*, vol. 23, no. 19, p. 8090, 2023.
- [25]. A. Al-Mascati and K. Al-Naabi, “Federated Learning Security and Privacy: A Comprehensive Survey of Attacks and Defenses,” arXiv preprint arXiv:2403.06067, 2024.
- [26]. A. Al-Mascati, K. Al-Naabi, and S. Al-Harthi, “Federated Adversarial Training for Robust Smart City Applications,” *Informatics*, vol. 10, no. 4, p. 89, 2023.
- [27]. J. Li et al., “DDFed: A Dual Defense Federated Learning Framework,” in *International Conference on Learning Representations*, 2024.
- [28]. A. Al-Marridi, O. Boubaker, and H. Al-Kahtani, “Federated learning for intrusion detection in IoT environments: a privacy-preserving strategy,” *Discover Internet of Things*, vol. 5, no. 72, 2025.
- [29]. C. D. Maciel, A. R. S. d. Santos, and P. Ferraro, “Federated Learning for IoT Intrusion Detection,” *Journal of Information and Data Management*, vol. 14, no. 1, pp. 1-15, 2023.
- [30]. S. Al-Harthi, A. Al-Mascati, and K. Al-Naabi, “A Trust-Centric Federated Learning Framework for Intrusion Detection Systems in IoT,” *Frontiers in Big Data*, vol. 7, 2024. 13
- [31]. P. M. Lopez-Martin, B. Carro, A. Sanchez-Esguevillas, and J. Lloret, “Intrusion Detection based on Privacy-preserving Federated Learning for the Industrial IoT,” in *2022 IEEE Global Communications Conference (GLOBECOM)*, 2022, pp. 1-6.
- [32]. C. D. Maciel, “A Performance Evaluation of Federated Learning for IoT Intrusion Detection,” *Future Internet*, vol. 16, no. 3, p. 78, 2024.
- [33]. G. D'Angelo, S. Marrone, and F. Narducci, “Federated Learning for Intrusion Detection Systems: A Survey,” *ACM Computing Surveys*, vol. 56, no. 2, pp. 1-38, 2024.
- [34]. M. A. A. Al-qaness et al., “Centralized and Federated Intrusion Detection for Internet of Things: A Survey,” *IEEE Communications Surveys & Tutorials*, vol. 26, no. 1, pp. 545-587, 2024.
- [35]. W. Wang, Y. Li, and L. Wang, “Deep Learning for Network Intrusion Detection in Industrial IoT: Challenges and a Federated Learning Solution,” *IEEE Transactions on Industrial Informatics*, vol. 20, no. 1, pp. 935-944, 2024.
- [36]. M. Davies, S. P. Arachchige, M. A. Al-Khafajiy, and Y. Jiao, “A Comparative Study of Machine Learning Algorithms for Federated Intrusion Detection in IoT,” *Electronics*, vol. 14, no. 6, p. 1176, 2025.
- [37]. G. D'Angelo, S. Marrone, and F. Narducci, “Federated Learning for Building Intrusion Detection Models in IoT: A Survey,” arXiv preprint arXiv:2204.12443, 2022.
- [38]. Zonduo, “Format for IEEE Research Paper,” 2024. [Online]. Available: <https://zonduo.com/format-for-ieee-research-paper.php>.
- [39]. IEEE Author Center, “Structure Your Paper,” 2024. [Online]. Available: <https://conferences.ieeeauthorcenter.ieee.org/write-your-paper/structure-your-paper/>.