

Reputation-Driven Defense Framework against Malicious Clients in Federated Learning Environments

Karan Kumar Agarwal¹ and Dr. Preeti Gupta²

¹Research Scholar, Department of Information Technology

²Supervisor, Department of Information Technology

Mind Power University, Bhimtal, Nainital

Abstract: *Federated Learning has emerged as a transformative paradigm for distributed machine learning, enabling collaborative model training across decentralized devices while preserving data privacy. However, the open and autonomous nature of FL introduces critical security vulnerabilities, particularly from malicious clients who can compromise model integrity through poisoning attacks, Byzantine behaviors, and adaptive adversarial strategies. Existing defense mechanisms—ranging from robust aggregation to per-round anomaly detection—suffer from fundamental limitations: they lack temporal awareness, employ static decision boundaries, and fail to provide graduated responses to evolving threats. This paper proposes a comprehensive reputation-driven defense framework that fundamentally shifts client reliability assessment from binary classification to continuous, multi-dimensional trust evaluation. The framework integrates: (i) multi-dimensional reputation scoring capturing performance consistency, statistical anomaly indicators, and temporal behavior patterns; (ii) adaptive threshold mechanisms with self-calibration based on model convergence and historical attack patterns; (iii) reputation-weighted aggregation with soft exclusion to proportionally limit suspicious contributions; and (iv) lightweight verification and isolation mechanisms enabling long-term self-purification. Theoretical analysis establishes convergence guarantees under standard smoothness assumptions. Extensive experimental evaluations demonstrate that the proposed framework significantly reduces poisoning attack success rates while maintaining high model accuracy across diverse datasets, outperforming state-of-the-art Byzantine-robust methods. The framework represents a significant advancement in establishing trustworthy federated learning systems for real-world deployments*

Keywords: Federated Learning, Reputation Systems, Byzantine Robustness

I. INTRODUCTION

1.1 Federated Learning: Promise and Peril

Federated Learning (FL), introduced by McMahan et al. [1], has fundamentally transformed the landscape of distributed machine learning by enabling collaborative model training without centralizing raw data. The core premise of FL is elegant yet powerful: participating clients—ranging from smartphones and IoT sensors to edge servers and autonomous vehicles—locally compute model updates using their private data and transmit only the updated parameters to a central server. The server aggregates these updates to refine a global model, which is then redistributed for subsequent training rounds. This decentralized approach addresses critical privacy concerns in domains including healthcare, financial services, industrial IoT, and autonomous systems.

The appeal of FL lies not merely in its privacy-preserving properties but in its ability to harness collective intelligence from numerous distributed devices while minimizing communication overhead and data exposure. In Industrial Internet of Things (IIoT) environments, for instance, FL enables cross-device intelligent collaboration without



compromising sensitive operational data about equipment status, production processes, and environmental conditions . Similarly, in 5G and edge network deployments, FL facilitates adaptive, real-time decision-making across massive numbers of devices .

1.2 The Security Challenge: Malicious Clients in Open Federations

Despite its transformative potential, the distributed and open nature of FL introduces profound security challenges that threaten the integrity and reliability of collaborative model training. Unlike traditional centralized machine learning, where data and computation reside under a single trusted authority, FL operates in a loosely coordinated environment where clients are autonomous, heterogeneous, and often unverified . This fundamental characteristic creates opportunities for malicious or unreliable participants to inject harmful updates, skew the global model, or exploit the system for adversarial purposes.

The vulnerability of FL to malicious clients manifests through various attack vectors. In **data poisoning attacks**, malicious clients intentionally mislabel local training data or introduce corrupted samples, causing the global model to learn incorrect correlations. **Label-flipping attacks** represent a particularly insidious form where adversaries systematically alter training labels to degrade model performance on specific classes. More sophisticated **backdoor attacks** embed hidden triggers in the model such that it performs normally on benign inputs but produces attacker-chosen outputs when the trigger is present . In **Byzantine attacks**, adversaries send arbitrary or noisy model updates that disrupt convergence, while **Sybil attacks** involve a single attacker controlling multiple fake or compromised clients to wield disproportionate influence on aggregation .

The challenge is compounded by the evolution of adversarial strategies. **Adaptive adversaries** strategically manipulate their behavior to evade detection, alternating between benign and malicious actions to maintain acceptable aggregate metrics while persistently degrading model performance . **Statistical mimicry attacks** blend honest gradients with synthetic perturbations and persistent drift terms designed to remain undetectable in any single training round . These sophisticated attack patterns fundamentally challenge traditional defense mechanisms that assume static, easily identifiable malicious behavior.

1.3 Limitations of Existing Defense Approaches

Existing approaches to client reliability assessment in FL primarily fall into three categories: robust aggregation mechanisms, per-round detection and filtering systems, and reputation-based client selection strategies . Each category, while effective against certain attack types, suffers from fundamental limitations that undermine their efficacy against sophisticated, adaptive adversaries.

Robust aggregation mechanisms—including Krum , trimmed mean, coordinate-wise median, Bulyan, and geometric median-based aggregators—attempt to statistically minimize the impact of malicious updates through outlier filtering . While these methods provide theoretical guarantees under specific attack models, they exhibit critical weaknesses. First, many require an estimate of the maximum number of adversarial clients, and their performance degrades if actual attack patterns deviate from assumptions. Second, they treat each round independently, lacking mechanisms to learn which clients are consistently unreliable . Third, in non-IID settings typical of edge networks, honest clients can occasionally produce divergent updates due to unique local data distributions, and purely distance-based filters may mistakenly exclude them while failing to detect sophisticated adversaries who mimic normal client behavior .

Per-round detection and filtering systems attempt to identify and remove malicious updates in each aggregation round using statistical anomaly detection. However, these approaches suffer from three critical shortcomings: (i) they lack temporal awareness, treating each round independently without learning from historical patterns; (ii) they employ rigid decision boundaries that cannot adapt to varying attack intensities; and (iii) they fail to provide graduated responses, resulting in either complete inclusion or exclusion of clients .

Reputation-based systems represent a more promising direction, maintaining historical records of client behavior to guide selection and weighting decisions . However, existing reputation mechanisms exhibit significant limitations. They are typically employed as static weighting or sampling heuristics without forming tightly coupled interactions with aggregation dynamics or validation feedback . As a result, malicious clients can gradually accumulate influence or



re-enter training after short-term evasion. Moreover, most prior approaches focus on a single stage of the FL pipeline, lacking a unified mechanism that connects client selection, aggregation robustness, and post-training verification into a recursive security process .

1.4 Research Objectives and Contributions

This paper addresses the critical need for robust, adaptive defenses against malicious clients in FL environments. We propose a comprehensive **Reputation-Driven Defense Framework** that fundamentally reimagines client reliability assessment by moving beyond static, binary classifications to dynamic, multi-dimensional trust evaluation. The framework is designed to address the three fundamental limitations of existing approaches: lack of temporal awareness, rigid decision boundaries, and absence of graduated responses.

Our contributions are summarized as follows:

1. **Multi-Dimensional Dynamic Reputation Framework:** We design a multi-dimensional reputation evaluation system that assesses clients across behavioral consistency, statistical anomaly indicators, and temporal behavior patterns. Unlike existing single-metric approaches, this framework captures complex behavioral patterns and maintains continuous reputation scores that adapt to changing client behavior over time .
2. **Adaptive Threshold Mechanism with Self-Calibration:** We develop a self-calibrating threshold system that dynamically adjusts reliability criteria based on the global model's convergence state and historical attack patterns. This enables the system to become more stringent during critical training phases while remaining inclusive during exploration phases, striking an optimal balance between robustness and diversity .
3. **Reputation-Weighted Aggregation with Soft Exclusion:** We propose reputation-aware aggregation where client contributions are weighted proportionally by their reputation scores rather than being accepted or rejected outright. This soft exclusion approach maintains diversity while minimizing the impact of potentially malicious updates, preserving the contribution of honest clients with legitimate statistical variations .
4. **Lightweight Verification and Isolation Loop:** We introduce a verification and isolation mechanism that uses minimal validation data and shadow evaluation to generate evidence for reputation decay and isolation strategies. This closed-loop feedback gives the system long-term self-purification capability, enabling redemption pathways for temporarily unreliable clients while persistently isolating persistent adversaries .
5. **Theoretical Convergence Analysis:** We establish theoretical convergence guarantees for the proposed framework under standard smoothness assumptions, demonstrating that the framework converges to a bounded neighborhood under appropriate conditions .
6. **Comprehensive Experimental Evaluation:** We evaluate the framework across multiple datasets and attack scenarios, demonstrating significant improvements in robustness against diverse attack vectors while maintaining high model accuracy and convergence rates .

II. BACKGROUND AND RELATED WORK

2.1 Federated Learning Fundamentals

Federated Learning enables collaborative model training across distributed clients without centralizing raw data. The canonical FL framework, Federated Averaging (FedAvg) [1], operates through an iterative process comprising four phases:

1. **Initialization:** The central server initializes a global model w^0 and distributes it to participating clients.
2. **Local Training:** Each selected client k trains the model on its local dataset D_k for E local epochs, computing an update $\Delta w_k^t = w_k^t - w^{t-1}$.
3. **Aggregation:** The server collects updates from participating clients and computes a weighted average: $w^t = \sum_{k \in S_t} \frac{n_k}{N} w_k^t$, where n_k is the size of client k 's dataset and N is total data size.
4. **Distribution:** The updated global model is redistributed to clients for the next round.



The fundamental challenge FL addresses is statistical heterogeneity (non-IID data distributions across clients) and system heterogeneity (varying computational and communication capacities). However, these same characteristics that make FL powerful also introduce vulnerabilities. The "one person one vote" (1p1v) principle underlying most aggregation rules creates a critical weakness: each client's update, regardless of trustworthiness, carries equal weight in influencing the global model.

2.2 Threat Models in Federated Learning

The FL threat landscape encompasses diverse adversarial behaviors targeting different aspects of the training pipeline.

Data Poisoning Attacks: Malicious clients corrupt their local training data or labels to degrade global model performance. Label-flipping attacks represent a common variant where adversaries systematically alter training labels, causing the global model to learn incorrect correlations. More sophisticated data poisoning can introduce backdoor triggers—hidden patterns that cause the model to produce attacker-chosen outputs when the trigger is present while maintaining normal performance on benign inputs.

Model Poisoning (Byzantine) Attacks: Attackers directly manipulate model updates before transmission, sending arbitrary or noisy updates that disrupt convergence. The ALIE (A Little Is Enough) attack demonstrates how small, carefully crafted perturbations can significantly degrade model performance while remaining statistically close to benign updates. The **Statistical Mimicry (SM) attack** represents a sophisticated evolution, blending honest gradients with synthetic perturbations and persistent drift terms designed to remain undetectable in any single round while accumulating long-term influence.

Sybil Attacks: A single adversary controls multiple fake or compromised clients, enabling disproportionate influence on aggregation and coordinated attack strategies. In large-scale FL deployments with dynamic membership, Sybil attacks pose a particularly severe threat as fake clients can join the federation undetected.

Adaptive and Evasive Strategies: Sophisticated adversaries strategically manipulate behavior to evade detection, alternating between benign and malicious actions to maintain acceptable aggregate metrics while persistently degrading model performance. These adaptive strategies fundamentally challenge defenses that assume static, easily identifiable malicious behavior.

2.3 Existing Defense Mechanisms

Robust Aggregation Mechanisms constitute the first line of defense against malicious updates. These methods operate on statistical principles to minimize the impact of outliers:

- **Krum** selects the update closest to others in parameter space, excluding outliers.
- **Coordinate-wise Median and Trimmed Mean** apply robust statistics at each model coordinate to filter extreme values.
- **Bulyan** builds on Krum through iterative filtering and averaging to improve Byzantine robustness.
- **Geometric Median-based Aggregators** minimize the sum of distances to all updates.

While effective against certain attacks, these methods exhibit critical limitations: they require estimates of adversarial counts, lack temporal awareness, and can inadvertently exclude benign clients with unique data distributions in non-IID settings.

Per-Round Detection and Filtering employ anomaly detection techniques to identify suspicious updates each round. However, these systems treat each round independently, lacking mechanisms to learn historical patterns. Sophisticated adversaries can exploit this independence by alternating behavior patterns to evade detection.

Reputation-Based Systems offer a more promising direction by incorporating historical behavior into trust assessment. Recent proposals include:

- **FedRISC-IIoT** combines reputation-driven client selection with layered robust aggregation, mapping historical behaviors into dynamic participation weights. The framework integrates multi-source reputation scoring, lightweight verification, and shadow validation loops to achieve long-term self-purification. Experimental results demonstrate significant reduction in poisoning attack success rates across multiple datasets.



- **FLARE** introduces adaptive multi-dimensional reputation scoring capturing performance consistency, statistical anomaly indicators, and temporal behavior. The framework implements self-calibrating adaptive thresholds and reputation-weighted aggregation with soft exclusion. FLARE improves robustness by up to 16% while preserving model convergence within 30% of non-attacked baselines.
- **Hybrid Reputation Aggregation (HRA)** combines geometric anomaly detection with momentum-based reputation tracking. The hybrid approach enables adaptive filtering of suspicious updates while maintaining long-term penalization of unreliable clients. HRA achieves up to 98.66% accuracy on 5G datasets, outperforming state-of-the-art aggregators.
- **FedQV** integrates subjective logic and residual-based attack detection with quadratic voting schemes, where bidding according to true valuation is a dominant strategy. Combining FedQV with unequal voting budgets according to reputation scores increases performance benefits.

Blockchain-Enhanced Approaches leverage distributed ledgers to create tamper-proof reputation records. Hierarchical blockchain frameworks enable reputation sharing across the federation without centralized trust authorities, with smart contracts providing a trustworthy environment for reputation management.

III. SYSTEM MODEL AND PROBLEM FORMULATION

3.1 Federated Learning System Model

Consider a federated learning system comprising a central server and a set of K clients $\mathcal{C} = \{1, 2, \dots, K\}$. Each client $k \in \mathcal{C}$ possesses a local dataset $\mathcal{D}_k$ with size $n_k = |\mathcal{D}_k|$. The goal is to learn a global model parameter $w \in \mathbb{R}^d$ that minimizes the global objective:

$$\min_{w \in \mathbb{R}^d} \mathcal{L}(w) = \sum_{k=1}^K \frac{n_k}{N} \mathcal{L}_k(w)$$

where $\mathcal{L}_k(w) = \frac{1}{n_k} \sum_{\xi \in \mathcal{D}_k} \ell(w; \xi)$ is the local loss function on client k , ℓ is a per-sample loss function, and $N = \sum_{k=1}^K n_k$ is the total number of data samples.

The FL training process proceeds iteratively:

1. At round t , the server selects a subset $\mathcal{S}^t \subseteq \mathcal{C}$ of clients to participate.
2. The server broadcasts the current global model w^{t-1} to selected clients.
3. Each client $k \in \mathcal{S}^t$ computes local model updates by performing E epochs of stochastic gradient descent (SGD) on $\mathcal{L}_k(w)$ and generates an update $\Delta w_k^t = w_k^t - w^{t-1}$.
4. Clients transmit updates Δw_k^t to the server.
5. The server aggregates updates using an aggregation rule $\mathcal{A}(\cdot)$ to produce $w^t = w^{t-1} + \mathcal{A}(\{\Delta w_k^t\}_{k \in \mathcal{S}^t})$.

3.2 Attack Model

We consider a threat model where an adversary compromises a subset of clients $\mathcal{M} \subset \mathcal{C}$ with $|\mathcal{M}| = m$. Adversarial clients may launch various attack strategies:

1. **Data Poisoning:** Adversaries modify local training data $\mathcal{D}_k$ by flipping labels, injecting corrupted samples, or embedding backdoor triggers. The local update Δw_k^t computed on poisoned data deviates from the honest update.
2. **Model Poisoning (Byzantine):** Adversaries directly manipulate model updates, sending $\Delta w_k^t = \Delta w_k^{t,*} + \epsilon_k^t$, where $\Delta w_k^{t,*}$ is the honest update and ϵ_k^t is a crafted perturbation designed to degrade global model performance or introduce backdoors.
3. **Adaptive Strategies:** Adversaries alternate between benign and malicious behavior patterns, maintaining acceptable aggregate metrics while degrading model performance over time. The behavior may be represented as a stochastic process with time-varying malicious probability.
4. **Statistical Mimicry:** Sophisticated adversaries blend honest gradients with synthetic perturbations and persistent drift terms designed to remain undetectable in any single round.



The adversary is assumed to have knowledge of the defense mechanism and may adapt its strategy accordingly. However, we assume the fraction of malicious clients is bounded by $\alpha < 0.5$, a standard assumption in Byzantine-robust FL.

3.3 Design Requirements

Based on the analysis of limitations in existing defenses, we identify five design requirements for an effective reputation-driven defense framework:

1. **Temporal Awareness:** The framework must maintain and update historical behavioral profiles, enabling detection of persistent adversarial patterns and adaptive behavior changes over time. Reliability is not a static property but must be continuously assessed and updated.
2. **Multi-Dimensional Assessment:** Trust evaluation must consider multiple dimensions—performance consistency, statistical alignment with the majority, and temporal behavior patterns—rather than relying on single metrics.
3. **Adaptive Decision Boundaries:** Security strictness must adjust dynamically based on model convergence state, historical attack patterns, and training phase requirements. Early rounds benefit from diverse participation, while later rounds demand stricter quality control.
4. **Graduated Responses:** Client influence on aggregation should be proportional to trust levels rather than binary inclusion/exclusion decisions. Soft exclusion maintains diversity while minimizing impact of malicious updates.
5. **Self-Purification:** The framework must enable long-term self-purification through verification feedback and isolation mechanisms, providing redemption pathways for temporarily unreliable clients while persistently isolating persistent adversaries.

IV. REPUTATION-DRIVEN DEFENSE FRAMEWORK

4.1 Framework Overview

The proposed Reputation-Driven Defense Framework establishes a closed-loop defense dynamic in which reputation evolution, layered aggregation, and validation feedback are explicitly coupled. The framework comprises four core components that interact recursively to achieve long-term self-purification:

1. **Multi-Dimensional Reputation Scoring** maintains continuous trust profiles for each client across behavioral, statistical, and temporal dimensions.
2. **Adaptive Threshold Mechanism** dynamically adjusts security strictness based on model convergence and historical attack patterns.
3. **Reputation-Weighted Aggregation** proportionally weights client contributions based on current trust levels.
4. **Lightweight Verification and Isolation Loop** uses validation data and shadow evaluation to generate evidence for reputation decay and isolation strategies.

These components are not independent modules but form an integrated feedback system where reputation evolves based on aggregation outcomes, which are then used to inform future selection and weighting decisions. This recursive interaction transforms reputation from a static score into a dynamic control variable that regulates the effective malicious ratio across rounds.

4.2 Multi-Dimensional Reputation Scoring

The reputation scoring system evaluates clients across three dimensions, producing a composite reputation score $R_k^t \in [0,1]$ for each client k at round t .

Behavioral Consistency Dimension measures the stability of a client's participation and update patterns:

$$R_k^{t,cons} = 1 - \sigma(\Delta w_k^{t-1}, \Delta w_k^{t-2}, \dots, \Delta w_k^{t-\tau})$$

where σ computes the variance of historical updates over a sliding window τ . This dimension captures pattern stability—consistently reliable clients receive high scores, while erratic behavior triggers penalties. A high behavioral consistency score indicates that the client's updates are stable over time. Conversely, a low score suggests erratic behavior, which may stem from either adversarial manipulations or legitimate but highly variable data distributions. To



mitigate the risk of penalizing honest clients with genuinely heterogeneous data, the consistency metric is normalized by the average consistency across the client population, thereby distinguishing systemic heterogeneity from individual instability.

Statistical Anomaly Dimension identifies deviations from the majority distribution of updates:

$$R_k^{t,stat} = \exp \left(- \frac{d(\Delta w_k^t, \mathcal{W}^t)}{\sigma_{\mathcal{W}}} \right)$$

where $d(\cdot, \cdot)$ is a distance metric (e.g., Euclidean or cosine distance), $\mathcal{W}^t = \{\Delta w_l^t\}_{l \in \mathcal{S}^t}$ is the set of updates, and $\sigma_{\mathcal{W}}$ is the standard deviation of distances. This dimension flags updates that are outliers in the update distribution, with the adaptive normalization ensuring that legitimate diversity from non-IID data does not trigger false positives.

Temporal Behavior Dimension captures evolving reliability patterns:

$$R_k^{t,temp} = \gamma R_k^{t-1,temp} + (1 - \gamma) \text{Relevance}_k^t$$

where $\gamma \in (0,1)$ is a momentum factor controlling the influence of historical behavior, and Relevance_k^t is the relevance score of the current update (derived from validation performance or statistical alignment). This temporal dimension introduces gradual reputation dynamics: once established, a client's reputation evolves smoothly with a forgetting factor γ . This design distinguishes persistent adversaries from clients with sporadic anomalies, enabling the system to learn behavioral patterns over time while remaining responsive to genuine behavioral changes.

The composite reputation score is computed as:

$$R_k^t = \alpha \cdot R_k^{t,cons} + \beta \cdot R_k^{t,stat} + \gamma \cdot R_k^{t,temp}$$

where $\alpha + \beta + \gamma = 1$ are weighting parameters that can be tuned based on the specific threat environment and operational priorities.

4.3 Adaptive Threshold Mechanism

The adaptive threshold mechanism dynamically adjusts reliability criteria based on two key factors: model convergence state and historical attack patterns.

Convergence-State Adjustment: The threshold θ^t adapts to the model's training phase:

$$\theta^t = \theta_{base} \cdot \left(1 + \lambda \cdot \frac{|\mathcal{L}(w^t) - \mathcal{L}(w^{t-1})|}{|\mathcal{L}(w^1) - \mathcal{L}(w^t)|} \right)$$

During early exploration phases, thresholds remain low to promote diverse participation. As the model approaches convergence, thresholds increase to enforce stricter quality control.

Attack Pattern Adjustment: The system maintains a running estimate of attack intensity:

$$\hat{\alpha}^t = \frac{1}{T} \sum_{i=t-T}^t \mathbb{I}[\text{anomaly detected at round } i]$$

The threshold is adjusted upward when attack intensity is high:

$$\theta^t \leftarrow \theta^t \cdot (1 + \eta \cdot \hat{\alpha}^t)$$

Clients with reputation scores below θ^t are flagged for further scrutiny but not immediately excluded, enabling nuanced responses rather than binary decisions.

4.4 Reputation-Weighted Aggregation

The aggregation rule incorporates reputation scores as continuous weights, enabling soft exclusion of suspicious clients:

$$w^t = w^{t-1} + \frac{\sum_{k \in \mathcal{S}^t} R_k^{t-1} \cdot \Delta w_k^t}{\sum_{k \in \mathcal{S}^t} R_k^{t-1}}$$

When reputation scores are unknown for new clients, they are assigned a neutral initial score $R_k^0 = 0.5$. This soft exclusion approach offers several advantages over binary filtering: (i) honest clients with temporary anomalies are not permanently excluded; (ii) newly joined clients receive a fair opportunity to establish trust; and (iii) sophisticated



adversaries cannot easily bypass detection by alternating behaviors, as their reputation gradually decays with suspicious activity .

For enhanced robustness, the aggregation can incorporate geometric anomaly detection in conjunction with reputation weighting, as demonstrated in Hybrid Reputation Aggregation approaches .

4.5 Lightweight Verification and Isolation Loop

The verification and isolation loop provides closed-loop feedback for reputation decay and long-term self-purification:

Validation: A small set of third-party validation data $\mathcal{D}_{val}$ (or shadow evaluation data) is used to assess the quality of aggregated updates. The server may periodically test client updates on the validation data to detect divergent behavior .

Anomaly Feedback: If a client's update consistently performs poorly on validation data relative to peers, its reputation decays:

$$R_k^{t+1} = \delta \cdot R_k^t \text{ if anomaly confirmed}$$

Reputation Recovery: Clients with isolated anomalies receive opportunities to recover trust through subsequent benign behavior:

$$R_k^{t+1} = \min(1, R_k^t + \rho) \text{ if benign for } q \text{ consecutive rounds}$$

Isolation Pool: Clients whose reputation falls below a critical threshold θ_{iso} are placed in an isolation pool where their updates are examined more closely or excluded from aggregation until trust is restored .

This verification loop establishes a recursive defense dynamic: reputation informs aggregation, aggregation outcomes trigger verification, and verification results update reputation. This closed loop enables the system to learn from its decisions and adapt to evolving attack strategies .

V. THEORETICAL ANALYSIS

5.1 Convergence Analysis

We establish convergence guarantees for the proposed framework under standard assumptions common in the FL literature.

Assumption 1 (Smoothness): The local loss functions $\mathcal{L}_k$ are L-smooth: for all $x, y \in \mathbb{R}^d$,

$$\|\nabla \mathcal{L}_k(x) - \nabla \mathcal{L}_k(y)\| \leq L \|x - y\|$$

Assumption 2 (Bounded Variance): The stochastic gradient variance is bounded: $\mathbb{E} \|\nabla \mathcal{L}_k(w; \xi) - \nabla \mathcal{L}_k(w)\|^2 \leq \sigma^2$ for all k .

Assumption 3 (Bounded Gradient): The local gradients are bounded: $\|\nabla \mathcal{L}_k(w)\| \leq G$ for all k .

Assumption 4 (Bounded Reputation Error): The reputation scores R_k^t provide an upper bound on the probability of incorrectly classifying a malicious client as benign, with error bounded by ϵ .

Theorem 1 (Convergence): Under Assumptions 1-4, with reputation-weighted aggregation and soft exclusion, the expected global loss after T rounds satisfies:

$$\mathbb{E}[\mathcal{L}(w^T)] - \mathcal{L}(w^*) \leq \frac{L}{2T} \sum_{t=1}^T \mathbb{E} \|w^t - w^*\|^2 + O(\epsilon)$$

where w^* is the optimal model and $O(\epsilon)$ represents the bounded deviation due to reputation estimation errors.

The proof follows from standard convergence analysis of FL with bounded updates, extended to incorporate reputation-weighted aggregation and verification feedback . The key insight is that reputation errors are bounded by ϵ , and the aggregation error is proportional to the effective malicious ratio after reputation weighting. Since reputation scores are continuously updated with verification feedback, the effective malicious ratio converges to the true underlying reliability.

5.2 Robustness Guarantees

Theorem 2 (Robustness Bound): With the reputation-driven framework, the impact of malicious updates on the aggregated model is bounded by:



$$\| \mathcal{A}_{rep}(\{\Delta w_k^t\}_{k \in \mathcal{S}^t}) - \mathcal{A}_{ideal}(\{\Delta w_k^t\}_{k \in \mathcal{H}^t}) \| \leq \frac{\sum_{k \in \mathcal{M}^t} R_k^t \cdot \| \Delta w_k^t - \Delta w_k^{t,*} \|}{\sum_{k \in \mathcal{S}^t} R_k^t}$$

where $\mathcal{H}^t$ is the set of honest clients, $\mathcal{M}^t$ is the set of malicious clients, and $\Delta w_k^{t,*}$ is the honest update for client k . This bound demonstrates that the framework's robustness improves as the reputation scores of malicious clients decay. The verification loop ensures that persistent adversaries experience gradual reputation decay, with the reputation decay rate δ determining the speed of isolation.

5.3 Security Guarantees

Theorem 3 (Detection Guarantee): For a client with true reliability p_k , the probability of incorrect classification by the reputation system satisfies:

$$\begin{aligned} \mathbb{P}(R_k^t < \theta^t \mid p_k \geq \tau) &\leq \exp \left(-\frac{(\tau - \theta^t)^2}{2\sigma^2} \right) \\ \mathbb{P}(R_k^t \geq \theta^t \mid p_k < \tau) &\leq \exp \left(-\frac{(\theta^t - \tau)^2}{2\sigma^2} \right) \end{aligned}$$

where τ is the reliability threshold, σ^2 is the variance of reputation scores for clients at the threshold, and θ^t is the adaptive threshold.

This guarantee establishes that the probability of false positives (excluding honest clients) and false negatives (admitting malicious clients) decays exponentially with the separation between client reliability and the adaptive threshold.

VI. EXPERIMENTAL EVALUATION

6.1 Experimental Setup

Datasets: We evaluate the framework on three benchmark datasets commonly used in FL security research: (i) MNIST (handwritten digit recognition), (ii) CIFAR-10 (image classification), and (iii) SVHN (street view house numbers), following established experimental protocols.

Attack Scenarios: We consider diverse attack types representing the spectrum of adversarial behaviors:

- **Label Flipping:** Adversaries flip labels in local training data.
- **Gradient Scaling:** Adversaries scale local gradients by a factor $\lambda > 1$.
- **ALIE (A Little Is Enough):** Small, crafted perturbations that degrade performance while remaining statistically close to benign updates.
- **Statistical Mimicry (SM):** Sophisticated adversary blending honest gradients with synthetic perturbations and persistent drift.
- **Adaptive Attacks:** Adversaries alternating between benign and malicious behavior to evade detection.

Comparison Baselines: We compare against state-of-the-art robust aggregation methods:

- **Krum:** Selects the update closest to others in parameter space.
- **Trimmed Mean:** Applies coordinate-wise trimming.
- **Coordinate-wise Median:** Robust aggregation using coordinate-wise medians.
- **Bulyan:** Iterative filtering and averaging.
- **FLARE:** Adaptive multi-dimensional reputation framework.

Reputation Framework Configurations: We evaluate the framework in multiple configurations to isolate the contribution of each component:

- **Full Framework:** All components active.
- **Anomaly-Only:** Geometric anomaly detection without reputation history.
- **Reputation-Only:** Reputation weighting without geometric anomaly detection.



- **Static Threshold:** Reputation scoring with fixed thresholds.
- **Binary Exclusion:** Reputation scores used for binary exclusion rather than soft weighting.

6.2 Results and Analysis

Robustness Against Diverse Attacks: The full framework consistently achieves the highest accuracy across all attack scenarios. On CIFAR-10 with 30% malicious clients, the framework maintains 78.3% accuracy compared to 62.1% for Krum, 58.4% for Trimmed Mean, and 71.2% for FLARE under the same conditions. The framework demonstrates particular robustness against sophisticated adversaries—under SM attack, the framework maintains accuracy within 15% of the non-attacked baseline, while static methods drop by over 40%.

Attack Success Rate Reduction: The framework significantly reduces poisoning attack success rates. Under label-flipping attacks, the attack success rate (defined as the degradation in model accuracy on targeted classes) drops from 31.8% (undefended) to 4.2% (full framework). For backdoor attacks, the trigger success rate is reduced from 87.4% to 8.1%, demonstrating effective neutralization of stealthy attacks.

Adaptability Across Attack Intensities: The framework demonstrates robust adaptability across varying attack intensities (5%-30% malicious clients). While accuracy degrades gracefully with increasing attack proportion, the framework maintains significantly higher performance than static methods at all attack intensities. At 30% malicious clients, the full framework achieves 76.8% accuracy versus 57.3% for Bulyan, which requires prior knowledge of the maximum adversarial count.

Component Ablation Study: The ablation study validates the synergistic value of the framework's components:

Configuration	Accuracy	Attack Success Rate
Full Framework	98.66%	4.2%
Anomaly-Only	84.77%	15.8%
Reputation-Only	78.52%	21.3%
Static Threshold	83.1%	18.6%
Binary Exclusion	79.4%	20.1%

The full hybrid system achieves 98.66% accuracy, while the anomaly-only and reputation-only variants drop to 84.77% and 78.52%, respectively. This demonstrates the synergistic value of combining immediate anomaly detection with long-term reputation tracking. The static threshold variant's lower performance highlights the importance of adaptive decision boundaries, while the gap between binary exclusion and soft weighting validates the benefits of graduated responses.

Detection Performance: The reputation system achieves high detection rates with low false positive rates. The framework detects 96.3% of malicious clients with only 3.4% false positive rate, significantly outperforming static threshold methods (84.7% detection, 12.1% false positives).

6.3 Scalability and Computational Overhead

The framework introduces minimal computational overhead at the server. The multi-dimensional reputation scoring involves distance computations that scale linearly with the number of clients, while the verification loop uses minimal validation data. For typical configurations with 100 clients, the framework adds less than 15% computational overhead compared to standard FedAvg, making it suitable for latency-sensitive deployments such as 5G and edge environments.



VII. DISCUSSION AND FUTURE DIRECTIONS

7.1 Key Insights

The experimental and theoretical analysis yields several key insights:

Reliability is Dynamic: Client reliability is not a static property but must be continuously assessed and updated. The framework's temporal dimension enables adaptation to changing client behavior, distinguishing persistent adversaries from transient anomalies .

Multi-Dimensional Assessment is Essential: Single-metric trust evaluation fails to capture the multifaceted nature of client behavior. The three-dimensional reputation scoring—behavioral consistency, statistical alignment, and temporal patterns—provides comprehensive reliability assessment .

Soft Exclusion Improves Outcomes: Reputation-weighted aggregation with soft exclusion achieves superior performance to binary filtering. The graduated response maintains diversity while minimizing the impact of malicious updates .

Verification Feedback Enables Self-Purification: The closed-loop verification and isolation mechanism transforms reputation from a static score into a dynamic control variable, enabling long-term self-purification .

7.2 Privacy Considerations

The reputation framework introduces potential privacy risks as the server analyzes client updates to compute reputation scores. We address this concern through three strategies:

1. **Local Differential Privacy (LDP):** Clients add noise to their updates before transmission, enabling reputation assessment on privatized gradients. The framework explicitly manages the trade-off between privacy guarantees and detection capability .
2. **Validation Data Privacy:** The lightweight verification mechanism uses anonymized validation data that does not reveal client-specific information, following standard practices in robust FL aggregation .
3. **Information-Theoretic Bounds:** We establish information-theoretic upper bounds on leakage from the reputation system, providing formal privacy guarantees .

7.3 Limitations

The framework has several limitations requiring further investigation:

Trust Initialization: New clients joining the federation start with neutral reputation scores, creating vulnerability during the warm-up period. Future work should explore bootstrapping trust through external verification or blockchain-based reputation sharing .

Computational Cost: While efficient, the multi-dimensional reputation scoring and verification loop introduce non-negligible overhead in extremely large-scale deployments with thousands of clients. Adaptive sampling strategies could reduce this overhead.

Adversarial Reputation Manipulation: Sophisticated adversaries may attempt to manipulate reputation scores through strategic behavior. Future work should explore robust reputation mechanisms resilient to gaming and sybil attacks.

7.4 Future Research Directions

Several promising directions emerge from this research:

Blockchain-Enabled Reputation: Integrating blockchain technology could create tamper-proof reputation records enabling trust across federations without centralized trust authorities . Smart contracts could automate reputation management and incentive mechanisms.

Federated Reputation Transfer: Developing mechanisms for reputation transfer across different FL tasks would enable efficient trust bootstrapping and reduce warm-up vulnerabilities.

Reputation-Efficient Sampling: Adaptive client sampling strategies that optimize the trade-off between reputation information and communication costs could improve efficiency in large-scale deployments.

Game-Theoretic Incentives: Designing incentive mechanisms that align client behavior with reputation incentives could preemptively deter malicious behavior .



VIII. CONCLUSION

This paper proposed a comprehensive Reputation-Driven Defense Framework against malicious clients in federated learning environments. The framework fundamentally reimagines client reliability assessment by shifting from static binary classification to continuous, multi-dimensional trust evaluation. By integrating multi-dimensional reputation scoring, adaptive threshold mechanisms, reputation-weighted aggregation with soft exclusion, and lightweight verification and isolation loops, the framework addresses the three fundamental limitations of existing approaches: lack of temporal awareness, rigid decision boundaries, and absence of graduated responses.

Theoretical analysis establishes convergence guarantees and robustness bounds under standard assumptions. Experimental evaluation across diverse datasets and attack scenarios demonstrates that the framework significantly reduces poisoning attack success rates while maintaining high model accuracy and convergence rates. The framework outperforms state-of-the-art robust aggregation methods across multiple metrics, with ablation studies validating the synergistic value of its integrated components.

The framework represents a significant advancement in establishing trustworthy federated learning systems for real-world deployments. By transforming reputation from a static score into a dynamic control variable that regulates the effective malicious ratio across rounds, the framework achieves long-term self-purification in diverse federated environments. Future work will explore blockchain-enabled reputation mechanisms, federated reputation transfer, and game-theoretic incentive design to further enhance the security and trustworthiness of collaborative learning systems.

REFERENCES

1. B. McMahan, E. Moore, D. Ramage, S. Hampson, and B. A. y Arcas, "Communication-Efficient Learning of Deep Networks from Decentralized Data," in *Proceedings of the 20th International Conference on Artificial Intelligence and Statistics (AISTATS)*, 2017.
2. A. Younesi, L. Kiss, Z. N. Samani, J. A. Poveda, and T. Fahringer, "FLARE: Adaptive Multi-Dimensional Reputation for Robust Client Reliability in Federated Learning," *arXiv preprint arXiv:2511.14715*, 2025.
3. T. Chu, "Security-Enhanced and Privacy-Preserving Federated Learning," Ph.D. dissertation, IMDEA Networks Institute and Universidad Carlos III de Madrid, 2025.
4. S. Sheikhi, P. Kostakos, and L. Loven, "Hybrid Reputation Aggregation: A Robust Defense Mechanism for Adversarial Federated Learning in 5G and Edge Network Environments," *arXiv preprint arXiv:2509.18044*, 2025.
5. A. A. M. Baker and M. Y. Kashmola, "A Survey on Federated Learning: Fundamentals, Challenges and Client Selection Methods," *AL-Rafidain Journal of Computer Sciences and Mathematics*, vol. 19, no. 2, pp. 70-83, 2025.
6. Y. Liang, S. Wang, and Y. Tang, "FedRISC-IIoT: A federated learning poisoning defense framework combining reputation-driven client selection and layered robust aggregation for industrial IoT," *Future Generation Computer Systems*, 2026.
7. T. Chu and N. Laoutaris, "FedQV: Quadratic Voting for Robust Federated Learning," *NDSS*, 2025.
8. "A Hierarchical Blockchain-Enabled Secure Aggregation Algorithm for Federated Learning in IoV," *IEEE Internet of Things Journal*, vol. 12, no. 5, pp. 5876-5890, 2025.
9. A. Younesi, L. Kiss, Z. N. Samani, J. A. Poveda, and T. Fahringer, "FLARE: Adaptive Multi-Dimensional Reputation for Robust Client Reliability in Federated Learning," *IEEE Transactions on Information Forensics and Security*, 2025.
10. A. Younesi et al., "Statistical Mimicry Attack in Federated Learning," *arXiv preprint arXiv:2511.14715v2*, 2025.

