

FLamma: Adaptive Gamma-Based Game Theoretic Framework for Fair Federated Learning

Simin Javaherian¹, Bryce Turney¹, Xu Yuan², Li Chen¹, Nian-Feng Tzeng¹

Abstract—Federated Learning (FL) has gained prominence as a decentralized machine learning paradigm, allowing clients to collaboratively train a global model while preserving data privacy. Despite its potential, FL faces significant challenges in Internet of Things (IoT) heterogeneous environments, where varying client resources and capabilities can undermine overall system performance. Existing approaches primarily focus on maximizing global model accuracy, often at the expense of fairness among clients and system efficiency, particularly in non-IID (non-Independent and Identically Distributed) settings. In this paper, we introduce *FLamma*, a novel Federated Learning framework based on an adaptive gamma-based Stackelberg game, designed to address these limitations. *FLamma* defines fairness as *reducing both the variance of client accuracies and the deviation of client contributions from the mean over time*, ensuring equitable participation and outcomes. Our framework allows the server, acting as the leader, to adaptively adjust a decay factor, while clients, acting as followers, optimally select their numbers of local epochs to maximize their utility. Over time, the server incrementally balances client influence, initially rewarding higher-contributing clients and gradually leveling their impact, driving the system toward a Stackelberg Equilibrium. This strategic balancing not only improves fairness but also enhances collaboration across a diverse set of clients in heterogeneous or resource-variable environments. Extensive experiments on both IID and non-IID datasets show that *FLamma* significantly improves fairness—according to our defined metrics—without compromising model performance or convergence speed, outperforming traditional FL baselines in IoT-relevant scenarios.

Index Terms—Federated Learning, Internet of Things, Game Theory, Fairness, Data Heterogeneity

I. INTRODUCTION

Federated learning (FL) [1] enables collaborative model training across large-scale Internet of Things (IoT) devices, where vast numbers of distributed sensors, edge devices, and embedded systems continuously generate data. By allowing model training directly on resource-constrained IoT nodes, FL helps preserve data privacy and reduces communication overhead. Although numerous solutions have been proposed to address notoriously well-known challenges from different perspectives [2]–[7], these efforts often rely on the assumption that IoT nodes are inherently motivated to participate actively, which is far from practice. To establish an effective FL system, it is crucial to attract and retain a large number of

participating clients [8]. However, this is particularly challenging in heterogeneous IoT environments where devices have varying capabilities, resources, and data distributions, which may affect their willingness or abilities to contribute effectively. Without appropriate incentives, high-performing clients may lose motivation or become discouraged, resulting in their reduced participation and, thus, compromised global model performance in the FL system. Therefore, designing effective incentive mechanisms is essential to ensure the active and beneficial participation of the entire set of IoT devices involved in FL.

This motivates the use of game theory [9] with incentive mechanisms as a powerful tool for FL in heterogeneous IoT environments. By leveraging game theory, we can tailor incentives that align individual client objectives with the collective goals of the learning network. This involves formulating a payoff structure that rewards IoT devices based on multiple factors, including their accuracy and consistency in contributing to the global model. For example, devices providing high-quality data or demonstrating consistent participation can receive higher rewards, fostering a cooperative and equitable environment. Additionally, by accounting for client-specific limitations, such as computational power or network stability, incentive schemes can be customized to encourage low-resource IoT nodes to contribute effectively. Furthermore, these game-theoretic strategies can dynamically adjust to varying network conditions and client capabilities, ensuring a robust and adaptive FL system that optimizes both individual and collective outcomes. In this way, game theory offers opportunities to not only enhance the efficiency of FL but also ensure its scalability and fairness, making it viable for diverse FL deployments.

However, previous investigations into game theory and incentive mechanisms for FL have notable limitations. Many focus on short-term incentives—such as fixed or utility-proportional payments for resource usage or participation—without adapting to the evolving and variable contributions of clients over time [10], [11]. These approaches overlook the dynamic interplay between client behavior, resource availability, and participation frequency, which is especially pronounced in IoT environments. Existing game-theoretic models [12], [13] and non-game-theoretic methods [14] generally do not address long-term fairness or explicitly account for temporal participation patterns, such as shifts in which clients dominate the training process over multiple rounds. As a result, some clients may disproportionately influence the early stages of training, while others remain underrepresented, leading to persistent disparities in both contribution and per-

¹University of Louisiana at Lafayette, Email: {simin.javaherian1, bryce.turney1, li.chen, nianfeng.tzeng}@louisiana.edu. ²University of Delaware, Email: xyuan@udel.edu

Copyright (c) 2026 IEEE. Personal use of this material is permitted. However, permission to use this material for any other purposes must be obtained from the IEEE by sending a request to pubs-permissions@ieee.org.

This work was supported in part by NSF under Grants OIA-2019511 and OIA-2327452.

formance. Addressing these gaps is essential for sustaining balanced engagement and maintaining high-quality global models in heterogeneous IoT-based FL systems.

Our approach, *FLamma* (Federated Learning framework based on an adaptive gamma-based Stackelberg game), addresses these shortcomings by explicitly modeling the co-dependence between server policies and client strategies. Unlike prior reweighting schemes [15], our decay factor γ is *jointly optimized* with client decisions on local epochs in a leader–follower framework. This allows the server to reward high-contributing IoT devices in the short term while gradually normalizing their influence to prevent long-term dominance. By adapting γ to evolving contributions and participation patterns, *FLamma* ensures a fairer and more stable training dynamic in heterogeneous IoT-based FL environments.

We design a novel utility function with dynamic allocation coefficients that modulate each IoT device influence based on both their historical and current engagement. The resulting client–server interaction is formalized as a Stackelberg game, for which we derive a Stackelberg equilibrium that characterizes optimal strategies for both parties. By integrating the decay factor into client selection and aggregation weights, the framework adapts to temporal variations in participation, improving fairness while maintaining both a theoretically supported competitive convergence rate and strong overall model performance. In this way, *FLamma* addresses both immediate disparities and evolving imbalances, ensuring that initial high contributions are recognized without allowing persistent overrepresentation, resulting in a stable and equitable training dynamic.

Through experimental validation, we demonstrate the effectiveness of our approach in addressing these challenges and enhancing FL performance in IoT environments. *FLamma* is evaluated against FedAvg [1], FedProx [16], q-FFL [14], and incentivization-based methods [12], [17]. We conduct experiments on four benchmark datasets—CIFAR100, CIFAR10, FMNIST, and MNIST—under both IID and Non-IID settings using ResNet-18 and LeNet-5 models. In the high-skew data distribution setting, *FLamma* improves accuracy over FedAvg by 19.8% on MNIST, 26.8% on FMNIST, 88.4% on CIFAR10, and 78.3% on CIFAR100. Compared with the closest counterpart (IAFL) [17], it achieves gains of 2.2% on MNIST, 1.9% on FMNIST, performs on par with CIFAR10, and outperforms by 12.2% on CIFAR100. With respect to fairness across clients at test time, *FLamma* reduces accuracy variance relative to FedAvg by 91.6%, 89.8%, 78.8%, and 75.1% on MNIST, FMNIST, CIFAR10, and CIFAR100, respectively. During the learning process, CIFAR10 also achieves the lowest steady-state contribution deviation, with reductions of 83% compared to FedAvg, and 72% compared to the closest competing method (IAFL), indicating substantially fairer participation over the rounds.

We summarize our contributions as follows:

- We model the interaction between the server and its clients using a Stackelberg game-theoretic framework, where the server dynamically adjusts a decay factor, and clients optimize their local epochs to maximize utility. This approach achieves Stackelberg Equilibrium, promot-

ing fairness, balancing participation rates, and reducing update variance and communication imbalance across clients.

- The decay factor adapts to temporal variations in client engagement, preventing over-contribution by certain clients and ensuring balanced participation throughout the training process. This fosters a more fair and collaborative learning environment, particularly suited for dynamic IoT settings.
- We provide experimental validation results showing that our approach achieves a fairer accuracy distribution across clients and reduces contribution deviation from the mean, while maintaining comparable or improved global model performance. Our approach outperforms traditional FL methods in fairness, participation balance, and overall system performance.
- We provide a convergence analysis of our proposed approach, demonstrating that it maintains competitive convergence rates, and ensures timely model updates without compromising accuracy or fairness.

II. BACKGROUND AND RELATED WORKS

In FL, many studies aim to improve collaborative training under heterogeneous data and resource conditions, focusing on three interconnected themes: incentive mechanisms, client selection, and fairness in aggregation.

Incentive Mechanisms. Extensive research has explored incentive mechanisms for FL, including game-theoretic, contract-based, and reputation-driven approaches. Stackelberg models have been used to allocate rewards based on model quality or contribution [12], [18], while non-cooperative and coalition games address device selection, cost efficiency, and bias–variance trade-offs [19]–[21]. Other approaches incorporate hierarchical incentives, training-time rewards, or coalition stability mechanisms [17], [22]–[24]. More recent frameworks such as IncFL [25], hierarchical Stackelberg–reinforcement learning [26], and FedTune-SGM [27] extend incentives to model adoption and personalized fine-tuning. Despite these advances, existing incentive mechanisms primarily focus on participation or reward allocation rather than dynamically regulating each client’s influence during aggregation. In contrast, *FLamma* integrates fairness directly into the training process via a Stackelberg formulation, where the server adjusts a per-round decay factor to continuously balance client contributions and prevent long-term overrepresentation.

Contract-theoretic and reputation-based mechanisms have been widely used in FL to incentivize participation when client resources or data quality are private or uncertain [28]–[30]. These approaches design contracts, pricing, and payment schemes to balance cost, latency, and incentive compatibility [31], [32]. Reputation systems further encourage reliable participation by leveraging historical behavior and dynamic trust modeling [33]–[35]. Reinforcement learning methods have also been applied to optimize incentive policies and resource allocation [36], [37]. In contrast, *FLamma* regulates client influence directly through a game-theoretic decay mechanism that dynamically adjusts contribution impact during

training, improving fairness under resource variability and participation uncertainty without relying on contract design, reputation tracking, or learned incentive policies.

Client Selection and Fairness. Client selection aims to balance performance, efficiency, and fairness in heterogeneous IoT environments. Prior work has explored fairness-aware selection and aggregation strategies. Donahue *et al.* [38] studied egalitarian and proportional fairness, while FairFedCS [39] and AdaFL [40] adjust participation probabilities over time. FedFair³ [41], Eiffel [42], FairDPFL-SCS [43], PraFFL [44], and F³ [45] improve fairness through adaptive selection, personalization, or global penalty terms. Classical aggregation methods such as Agnostic Federated Learning (AFL) [46], q-FFL [14], and AdaFed [15] promote fairness by reweighting client losses but rely on static or loss-driven weighting and do not regulate long-term influence. In contrast, FLamma dynamically controls client contributions using a game-theoretic decay factor, enabling fairness over time rather than only correcting instantaneous disparities. Unlike personalization approaches that train client-specific models, FLamma maintains a single global model and enforces fairness directly at the aggregation level. Through its Stackelberg design, the server adapts a per-round decay factor to limit dominance and balance participation throughout training.

Motivation: Ensuring fairness while maintaining efficiency remains a key challenge in FL, particularly in heterogeneous and resource-constrained IoT environments where clients differ in resources, data quality, and availability. Existing approaches either allocate influence proportional to contributions [12], which risks persistent dominance by high-resource clients, or enforce fairness objectives through loss-based weighting [14], which may limit efficiency or fail to regulate long-term influence. This motivates the need for a dynamic and adaptive mechanism that balances fairness and performance. To address this, we propose FLamma, which introduces a decay factor to regulate client influence over time. Unlike prior methods [12], [47], FLamma initially allows high-contributing clients to accelerate convergence but gradually attenuates their influence, preventing long-term dominance and promoting balanced participation. The decay factor dynamically adjusts client contributions based on alignment and training dynamics, while clients independently select their local epochs to maximize utility. This leader-follower design enables resource-aware participation, reduces contribution imbalance, and improves fairness while maintaining stable and efficient training across heterogeneous IoT devices.

III. PRELIMINARIES

A. Federated Learning

The primary objective in FL is to minimize the global loss across all participating clients. This global loss function, denoted as $F(w)$, aggregates the local loss functions $F_i(w)$ computed on the individual datasets of each client. The objective function is expressed as: $\arg \min_w F(w) = \sum_{i=1}^N p_i F_i(w)$, where $F_i(w)$ represents the local loss function for client i , and p_i is the weight associated with each client, calculated as $p_i = \frac{|D_i|}{|D|}$, with $|D_i|$ being the number of data points owned by

client i , and $|D|$ being the total number of data points across all clients.

B. Game Theory and Stackelberg Games

Game theory [9] models interactions among rational agents whose payoffs depend on joint strategic choices. Each player selects a strategy to maximize a utility of the form $U_i = R_i - C_i$, where R_i represents the obtained benefit and C_i the associated cost.

A Stackelberg game [48] is a hierarchical leader-follower model in which:

- **Leader:** moves first by committing to a strategy.
- **Followers:** observe the leader's action and respond with their best-response strategies.

A Stackelberg equilibrium is reached when the leader chooses a strategy that maximizes its utility while accounting for the followers' optimal responses. In our FL setting, the server acts as the leader and clients as followers, with the formal definitions of their utilities and decision variables introduced later in Sec. IV.

IV. DESIGN OF FLamma

This section presents our proposed game-theoretic FL framework.

A. Overview

In conventional FL, uncontrolled participation can cause over-represented or misaligned clients to dominate the aggregation. FLamma replaces the implicit uniform weight with an adaptive term, $\underbrace{\gamma_t \cdot \omega_i^t \cdot \tau_i}_{\text{FLamma weight}}$, where ω_i^t measures alignment with

the global model, τ_i denotes local effort, and $\gamma_t \in [0, 1]$ is a server-controlled decay factor. By adjusting γ_t over rounds, the server progressively normalizes contributions, preventing dominance and improving fairness. Fig. 1 provides an overview of the proposed workflow.

B. Game Strategy

FLamma models the interaction between server and clients as a Stackelberg game [49]: the server (leader) selects the decay factor γ_t , and each client i (follower) chooses its number of local epochs τ_i to maximize its own utility. Both utilities follow the standard benefit-cost structure commonly used in FL incentive modeling.

Server Utility. Following the observation in [50] that heterogeneous updates bias convergence, FLamma incorporates an alignment proxy $\omega_i^t = 1 - \frac{\|w_i^t - w^t\|}{\|w^t\| + \epsilon}$, which increases when a client's update aligns well with the global direction. The server's per-round utility is $U_{\text{server}}(\gamma_t, \tau) = \eta \sum_{i=1}^N s_i^t p_i \gamma_t \tau_i \omega_i^t - \lambda_t \gamma_t^2$, where λ_t controls regularization and gradually reduces the optimal γ_t , promoting fairness over time.

Client Utility. Each client i solves

$$U_i(\gamma_t, \tau_i) = \gamma_t \omega_i^t \tau_i - c_i \tau_i^2, \quad (1)$$

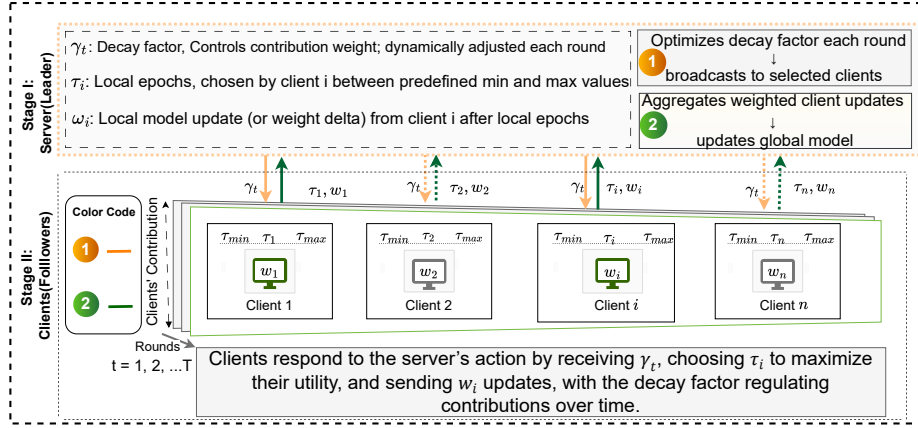


Fig. 1. Overview of *FLamma*, an adaptive gamma-based game theoretic framework for fair federated learning, where γ is adjusted over time to modulate the influence of every client's contribution along with training, avoiding overrepresentation or underrepresentation of clients.

where the first term reflects the contribution benefit, and the second captures the computation cost associated with local training. The cost parameter c_i models system heterogeneity through device-specific computation, energy, and communication constraints, enabling adaptive participation under heterogeneous device capabilities and availability.

Definition 1. Individual Rationality (IR). In our FL game, the IR constraint ensures that each client's utility in round t is non-negative: $U_i(\gamma_t, \tau_i) \geq 0$, $\forall i$. This condition guarantees that participating yields no loss for any client, so each has an incentive to follow its utility-maximizing strategy as defined in Eq. (1) rather than opting out of training.

C. Stackelberg Equilibrium

In the Stackelberg game, the server commits to the decay factor γ_t , and the clients respond optimally by selecting the number of local iterations τ_i^* . Differentiating the client utility $U_i(\gamma_t, \tau_i) = \gamma_t \omega_i^t \tau_i - c_i \tau_i^2$ with respect to τ_i gives $\frac{\partial U_i}{\partial \tau_i} = \gamma_t \omega_i^t - 2c_i \tau_i$. Setting this to zero yields the unconstrained maximizer $\tau_i = \frac{\gamma_t \omega_i^t}{2c_i}$.

Lemma 1. The client utility function $U_i(\gamma_t, \tau_i)$ is concave with respect to τ_i .

Proof. The second derivative is $\frac{\partial^2 U_i}{\partial \tau_i^2} = -2c_i < 0$. Thus, U_i is strictly concave in τ_i , ensuring a unique maximizer.

Lemma 2. Client utility is maximized at the optimal number of local epochs τ_i^* satisfying $\frac{\partial U_i}{\partial \tau_i} = 0$.

Proof. Since $U_i(\gamma_t, \tau_i)$ is strictly concave in τ_i by Lemma 1, its unique stationary point obtained from $\frac{\partial U_i}{\partial \tau_i} = 0$ is a global maximizer, which defines τ_i^* .

Definition 2 (Best Response). For $\tau_i \in [\tau_{\min}, \tau_{\max}]$, the best response of client i is $B_i(\gamma_t) = \arg \max_{\tau_i \in [\tau_{\min}, \tau_{\max}]} U_i(\gamma_t, \tau_i)$, which results the optimal local epochs $\tau_i^* = B_i(\gamma_t) = \min \left\{ \tau_{\max}, \max \left\{ \tau_{\min}, \frac{\gamma_t \omega_i^t}{2c_i} \right\} \right\}$.

Substituting τ_i^* into the server utility yields $U_{\text{server}}^{(t)}(\gamma_t) = \eta \sum_{i=1}^N s_i^t p_i \gamma_t \tau_i^* \omega_i^t - \lambda_t \gamma_t^2$, $0 \leq \gamma_t \leq 1$. Let $A_t \triangleq \eta \sum_{i=1}^N s_i^t p_i \tau_i^* \omega_i^t \geq 0$, so the utility becomes: $U_{\text{server}}^{(t)}(\gamma_t) = A_t \gamma_t - \lambda_t \gamma_t^2$. Differentiating and enforcing the constraint gives $\gamma_t^* = \min \left\{ 1, \max \left\{ 0, \frac{A_t}{2\lambda_t} \right\} \right\}$.

Lemma 3. The server utility is concave with respect to γ_t .

Proof. From the updated formulation, $U_{\text{server}}^{(t)}(\gamma_t) = A_t \gamma_t - \lambda_t \gamma_t^2$: Since the second derivative is with respect to γ_t , $A_t \geq 0$ is constant, and we have: $\frac{\partial^2 U_{\text{server}}^{(t)}}{\partial \gamma_t^2} = -2\lambda_t$. Since $\lambda_t > 0$ by definition, this derivative is negative, implying that $U_{\text{server}}^{(t)}$ is concave in γ_t for all t . Therefore, γ_t^* maximizes the utility function.

Definition 3 (Nash Equilibrium). A strategy profile $\tau^* = (\tau_1^*, \dots, \tau_N^*)$ is a Nash equilibrium if no client i can increase its utility by unilaterally deviating from τ_i^* . Formally, $U_i(\gamma_t, \tau_i^*) \geq U_i(\gamma_t, \tau_i) \quad \forall \tau_i, \forall i$.

Lemma 4. The client sub-game has at least one Nash equilibrium.

Proof. Each client's utility function is given by Eq. (1), where τ_i denotes the number of local epochs. From Lemma 1, $U_i(\gamma_t, \tau_i)$ is concave in τ_i , and the strategy set $\tau_i \in [\tau_{\min}, \tau_{\max}]$ is compact and convex. The best-response mapping $B_i(\gamma_t) = \tau_i^*$ is therefore continuous. By Rosen's Fixed Point Theorem [51], a continuous best-response mapping on a compact, convex strategy space admits at least one fixed point. Hence there exists a the strategy vector $\tau^* = (\tau_1^*, \dots, \tau_N^*)$ such that $\tau_i^* = B_i(\gamma_t)$ for all i . At this fixed point, no client can increase its utility by unilaterally choosing a different τ_i , so τ^* is a Nash equilibrium of the client sub-game. Therefore, the client sub-game admits at least one Nash equilibrium.

Given this context, we summarize Algorithm 1 as follows:

- **Step 1.** The server samples a subset of clients and broadcasts the current global model and decay factor (Lines 3–4).
- **Step 2.** Each selected client computes its best-response epochs, performs local training, and sends its update to the server (Lines 5–11).
- **Step 3.** The server evaluates client contribution proxies and updates the decay factor γ_t using the equilibrium rule (Lines 12–14).
- **Step 4.** The server aggregates updates using contribution-aware *FLamma* weighting (Line 15).
- **Step 5.** The final global model is obtained after T rounds (Line 16).

Algorithm 1: FLamma: Adaptive γ -based Game Theoretic Framework for Fair Federated Learning

Input : Initial global model w_0 , rounds T , base learning rate η , client weights $\{p_i\}$, participation size K , epoch bounds $\tau_{\min}, \tau_{\max}$, client costs $\{c_i\}$, stability constant ε , regularizer λ_t

Output: Final global model w_T

```

1 Initialize  $\gamma_0 \in (0, 1]$ ;
2 for  $t = 0, 1, \dots, T - 1$  do
3   Server (leader): sample a subset  $S_t$  of  $K$  clients;
4   broadcast  $(w_t, \gamma_t, \omega_{i,0}^{t-1})$  to all  $i \in S_t$ ;
5   foreach  $i \in S_t$  do // in parallel
6     Client  $i$  (follower): set  $w_{i,0}^t \leftarrow w_t$ ;
7     compute best-response epochs
8      $\tau_i^t \leftarrow \min\{\tau_{\max}, \max\{\tau_{\min}, \frac{\gamma_t \omega_{i,0}^{t-1}}{2c_i}\}\}$ ;
9     for  $e = 0, 1, \dots, \tau_i^t - 1$  do
10      sample mini-batch  $\xi_{i,e}^t$ ;
11       $w_{i,e+1}^t \leftarrow w_{i,e}^t - \eta \nabla F_i(w_{i,e}^t; \xi_{i,e}^t)$ ;
12      send  $\Delta w_i^t \leftarrow w_{i,\tau_i^t}^t - w_t$  to server;
13   Server: compute alignment proxy
14    $\omega_i^t \leftarrow 1 - \frac{\|w_i^t - w_t\|}{\|w_t\| + \varepsilon} = 1 - \frac{\|\Delta w_i^t\|}{\|w_t\| + \varepsilon}, \quad \forall i \in S_t$ ;
15   compute  $A_t \leftarrow \eta \sum_{i \in S_t} p_i \tau_i^t \omega_i^t$ ,  $\gamma_t^* \leftarrow \min\{1, \max\{0, \frac{A_t}{2\lambda_t}\}\}$ ;
16   set  $\gamma_{t+1} \leftarrow \gamma_t^*$ ;
17   Aggregation (FLamma-weighted):
18    $w_{t+1} \leftarrow w_t + \sum_{i \in S_t} p_i \gamma_t \omega_i^t \Delta w_i^t$ .
19 return  $w_T$ ;

```

Lemma 5. The equilibrium of the clients' sub-game always satisfies the IR constraint.

Proof. Consider an equilibrium τ^* , and let i be a client such that: $U_i(\gamma_t, \tau_i^*) < 0$. Now, substitute τ_i^* with 0: $U_i(\gamma_t, 0) = 0 > U_i(\gamma_t, \tau_i^*)$. This contradicts the definition of Nash equilibrium. Therefore, we have: $U_i(\gamma_t, \tau_i^*) \geq 0$ for all i .

How does FLamma respect fairness? As training progresses, the server adaptively adjusts γ_t to reduce the influence of consistently high-contributing clients and to strengthen the impact of those with lower contributions. This dynamic scaling balances client updates, prevents dominance, and mitigates overfitting to data from stronger clients, leading to a more uniform distribution of accuracy and contribution. Empirically, this effect is reflected in FLamma's reduction in accuracy variance (Table III) and contribution deviation (Fig. 2). By continuously regulating contribution weights, FLamma ensures that all clients—including resource-constrained or intermittently available IoT devices—retain meaningful influence throughout training, promoting fairness in both participation and model outcomes.

Computational and Communication Overhead. FLamma introduces minimal additional overhead beyond standard FL.

Computation. Each client computes its optimal local epochs τ_i^* via the closed-form best response $\tau_i^* =$

TABLE I
SENSITIVITY OF FLAMMA TO THE INITIAL DECAY FACTOR γ_0 ON CIFAR10 UNDER HIGH HETEROGENEITY ($\alpha = 0.5$).

γ_0	Acc (%) (at round 100)	Acc (%)	Variance	Deviation
1.00	55.32 $\pm$ 0.71	59.82 $\pm$ 0.64	274.52 $\pm$ 4.85	15.11 $\pm$ 0.25
0.99	56.90 $\pm$ 0.65	60.21 $\pm$ 0.57	265.10 $\pm$ 2.41	10.50 $\pm$ 0.13
0.95	55.01 $\pm$ 0.62	58.03 $\pm$ 0.59	263.65 $\pm$ 4.47	10.12 $\pm$ 0.10
0.90	54.76 $\pm$ 0.72	57.74 $\pm$ 0.63	260.83 $\pm$ 5.32	9.87 $\pm$ 0.14
0.85	52.34 $\pm$ 0.78	54.21 $\pm$ 0.68	259.83 $\pm$ 3.22	9.79 $\pm$ 0.18

TABLE II
SENSITIVITY OF FLAMMA TO THE REGULARIZER λ ON CIFAR10 UNDER HIGH HETEROGENEITY ($\alpha = 0.5$), WITH $\gamma_0 = 0.99$.

λ	Acc (%) (at round 100)	Acc (%)	Variance	Deviation
0.05	57.05 $\pm$ 0.66	60.10 $\pm$ 0.58	275.20 $\pm$ 3.60	13.44 $\pm$ 0.21
0.10	56.90 $\pm$ 0.65	60.21 $\pm$ 0.57	265.10 $\pm$ 2.41	10.50 $\pm$ 0.13
0.50	55.10 $\pm$ 0.60	58.80 $\pm$ 0.56	260.90 $\pm$ 3.40	10.22 $\pm$ 0.11

$\min\{\tau_{\max}, \max\{\tau_{\min}, \frac{\gamma_t \omega_i^t}{2c_i}\}\}$, which requires only scalar arithmetic. The server computes γ_t^* in closed form. Both operations are $O(1)$ and negligible compared to local SGD updates and model aggregation, making them feasible even for resource-constrained IoT devices.

Communication. FLamma broadcasts a single scalar decay factor γ_t together with the global model, and clients compute their optimal local epochs τ_i^* locally without additional feedback. Thus, FLamma does not introduce extra communication rounds, synchronization steps, or handshaking beyond standard FL protocols such as FedAvg. The additional payload per round is one scalar ($O(1)$), while model transmission requires $O(d)$ parameters. For realistic models where $d \gg 1$, this overhead is negligible and does not increase overall communication cost, bandwidth usage, or training latency, even in bandwidth-constrained or high-latency IoT environments.

Contribution Refresh. FLamma refreshes contribution estimates every 10 rounds to avoid staleness, incurring one additional local epoch per refresh. The relative overhead is $\frac{T/10}{\tau}$, which is negligible in practice (e.g., $\approx 1\%$ when $T = 100$, $\tau = 1000$). Importantly, FLamma remains adaptive because the decay factor γ_t and client epochs τ_i^t are recomputed every round, mitigating potential staleness effects. The refresh interval is tunable and can be reduced or made adaptive in highly dynamic environments.

V. CONVERGENCE ANALYSIS

We analyze FLamma under the standard FedAvg setting [52]. Let $\gamma_t \in [0, 1]$ denote the server's decay factor at round t and define the *effective server step* $\alpha_t \triangleq \eta \gamma_t$, with $\alpha_{\max} \triangleq \sup_{1 \leq t \leq T} \alpha_t$. For clarity of analysis, we consider the normalized case $\omega_i^t = 1$. Since $\gamma_t \in [0, 1]$ and $\tau_i^t \in [\tau_{\min}, \tau_{\max}]$, both the effective step size α_t and accumulated local updates are uniformly bounded. Consequently, the global update magnitude remains controlled, and convergence guarantees hold despite dynamic fairness adjustments.

Let $w_{i,e}^t$ denote the local model of client i after completing the e -th local epoch at round t , with $w_{i,0}^t = w_t$. In each round t , the server samples a set S_t of K clients. Each client $i \in S_t$

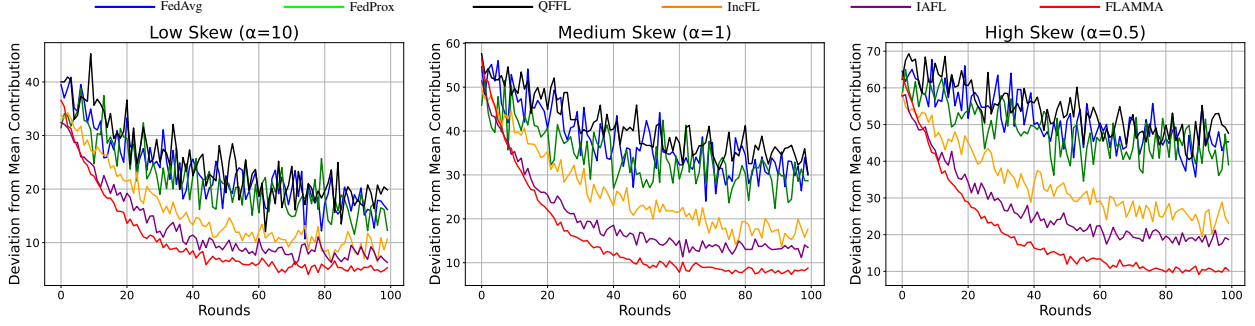


Fig. 2. Comparison of client contribution deviation from the mean contribution over 100 rounds on the CIFAR10 dataset.

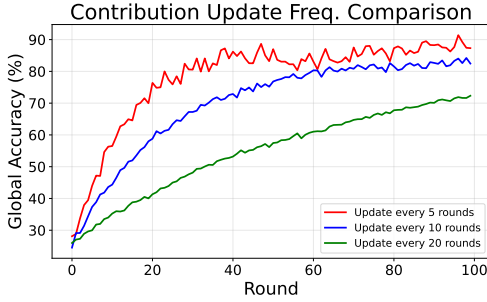


Fig. 3. Effect of contribution-refresh frequency on global accuracy over 100 rounds. Refreshing every 10 rounds strikes the best trade-off between rapid progress and stable performance; refreshing every 5 rounds accelerates early gains but shows larger round-to-round fluctuations, while every 20 rounds slows learning and yields lower final accuracy.

performs τ_i^t local epochs and returns its local model w_i^t . The server update is $w_{t+1} = w_t - \alpha_t \sum_{i \in S_t} p_i g_i^t$, $g_i^t \triangleq \frac{w_t - w_i^t}{\eta}$, where η is the client SGD step size. Equivalently, if $w_{i,0}^t = w_t$ and $w_{i,e+1}^t = w_{i,e}^t - \eta h_i(w_{i,e}^t)$, then $g_i^t = \sum_{e=0}^{\tau_i^t-1} h_i(w_{i,e}^t)$, where $h_i(\cdot)$ is the stochastic gradient in Assumption 2.

Assumption 1 (Strong convexity and smoothness). For all clients i and all w, w' (where $\langle \cdot, \cdot \rangle$ denotes the standard inner product between two vectors):

- (a) $F_i(w') \geq F_i(w) + \langle \nabla F_i(w), w' - w \rangle + \frac{\rho}{2} \|w' - w\|^2$,
- (b) $\|\nabla F_i(w) - \nabla F_i(w')\| \leq \beta \|w - w'\|$.

These imply that $F(w) = \sum_i p_i F_i(w)$ is ρ -strongly convex and β -smooth.

Assumption 2 (Stochastic gradients). On client i at round t , let $h_i(w_i^t)$ be an unbiased stochastic gradient of F_i at w_i^t , with $\mathbb{E}[h_i(w_i^t)] = \nabla F_i(w_i^t)$, $\mathbb{E}[\|h_i(w_i^t) - \nabla F_i(w_i^t)\|^2] \leq \sigma_i^2$, $\mathbb{E}[\|h_i(w_i^t)\|^2] \leq G^2$.

Assumption 3 (Bounded update variance with effective step). Let the aggregation weights satisfy $\sum_{i \in S_t} p_i = 1$. Under bounded stochastic gradients and bounded local epochs, the variance of the global update satisfies $\mathbb{E}\|w_{t+1} - w_t\|^2 \leq 4\alpha_t^2(\tau_{\max})^2 G^2 \mathbb{E}_{S_t}[\sum_{i \in S_t} p_i^2]$. Moreover, since $\sum_{i \in S_t} p_i^2 \leq \sum_{i=1}^N p_i^2$ for any subset S_t , $\mathbb{E}\|w_{t+1} - w_t\|^2 \leq 4\alpha_t^2(\tau_{\max})^2 G^2 \sum_{i=1}^N p_i^2$.

Assumption 4 (Gradient heterogeneity). The variance of client gradients around the global gradient is uniformly bounded:

$\mathbb{E}_{S_t}[\sum_{i \in S_t} p_i \|\nabla F_i(w_t) - \nabla F(w_t)\|^2] \leq \zeta^2$, for some constant $\zeta^2 > 0$ independent of round t .

Assumption 5 (Bounded local-global divergence). Given Assumption 2 and bounded gradients, the divergence between the local models and the global model at round t satisfies $\mathbb{E}_{S_t}[\sum_{i \in S_t} p_i \|w_t - w_i^t\|^2] \leq 4\eta^2(\tau_{\max} - 1)^2 G^2$, where $\tau_{\max}$ is the maximum local epochs, and G bounds the stochastic gradient norms. The expectation is over client sampling, mini-batch sampling, and any randomness in local updates. Because τ_i^t is constrained by the client best-response to lie within $[\tau_{\min}, \tau_{\max}]$, the local-global divergence remains uniformly bounded despite dynamic epoch selection.

Theorem (Convergence of FLamma). Let $\kappa = \beta/\rho$, $\xi = \max\{8\kappa, \tau_{\max}\}$, choose $\eta = \frac{2}{\rho\xi}$, and define the effective step $\alpha_t = \eta\gamma_t$ with $\alpha_{\max} = \sup_t \alpha_t$. Let $C = 4(\tau_{\max})^2 G^2 \sum_{i=1}^N p_i^2$. Then

$$\mathbb{E}[F(w_T)] - F^* \leq \frac{\kappa}{T} \left(\frac{2(B_\alpha + \alpha_{\max}^2 C)}{\rho} + \frac{\rho\xi}{2} M \right), \quad (2)$$

where $M = \mathbb{E}\|w_1 - w^*\|^2$ and $B_\alpha = \frac{\alpha_{\max}^2}{\rho} \left(\sum_{i=1}^N p_i^2 \sigma_i^2 + \zeta^2 + 6\beta \right) + 4\eta^2(\tau_{\max} - 1)^2 G^2$. The last term bounds local-global model divergence (Assumption 5), while C bounds global update magnitude (Assumption 3).

Proof. Following FedAvg [52], by Assumption 1 (strong convexity and smoothness),

$$\begin{aligned} \mathbb{E}\|w_{t+1} - w^*\|^2 &\leq (1 - \alpha_t \rho) \mathbb{E}\|w_t - w^*\|^2 + \alpha_{\max}^2 \zeta^2 \\ &\quad + \alpha_{\max}^2 \sum_{i=1}^N p_i^2 \sigma_i^2 + 6\beta \alpha_{\max}^2 + 4\eta^2(\tau_{\max} - 1)^2 G^2 \end{aligned} \quad (3)$$

Applying Assumption 4 and $\alpha_t \leq \alpha_{\max}$ to the heterogeneity term gives $\alpha_t^2 \mathbb{E}_{S_t}[\sum_{i \in S_t} p_i \|\nabla F_i(w_t) - \nabla F(w_t)\|^2] \leq \alpha_{\max}^2 \zeta^2$. Since F is β -smooth, $F(w_{t+1}) \leq F(w_t) + \langle \nabla F(w_t), w_{t+1} - w_t \rangle + \frac{\beta}{2} \|w_{t+1} - w_t\|^2$. Taking expectations and summing over $t = 1, \dots, T$,

$$\begin{aligned} \sum_{t=1}^T \mathbb{E}[F(w_{t+1})] &\leq \sum_{t=1}^T \mathbb{E}[F(w_t)] + \sum_{t=1}^T \mathbb{E}[\langle \nabla F(w_t), w_{t+1} - w_t \rangle] \\ &\quad + \frac{\beta}{2} \sum_{t=1}^T \mathbb{E}\|w_{t+1} - w_t\|^2. \end{aligned} \quad (4)$$

The inner-product term produces the standard descent contribution under strong convexity and smoothness, and can

be bounded using the distance recursion in Eq. (3) together with the smoothness inequality. The quadratic term is upper-bounded by Assumption 3:

$$\frac{\beta}{2} \sum_{t=1}^T \mathbb{E} \|w_{t+1} - w_t\|^2 \leq \frac{\beta}{2} T \alpha_{\max}^2 C, \quad C = 4(\tau_{\max})^2 G^2 \sum_{i=1}^N p_i^2.$$

Combining these with Eq. (3), using $\eta = \frac{2}{\rho\xi}$ and the standard relations $\mathbb{E} \|w_t - w^*\|^2 \leq \frac{2}{\rho} (\mathbb{E}[F(w_t)] - F^*)$ and $F(w_T) - F^* \leq \frac{\beta}{2} \mathbb{E} \|w_T - w^*\|^2$, and telescoping over $t = 1, \dots, T$ (as in [52]), we obtain

$$\mathbb{E}[F(w_T)] - F^* \leq \frac{\kappa}{T} \left(\frac{2(B_\alpha + \alpha_{\max}^2 C)}{\rho} + \frac{\rho\xi}{2} M \right),$$

This completes the proof. This bound holds uniformly since ζ^2 does not depend on the round index t .

In our method, the server scales updates using the effective step $\alpha_t = \eta\gamma_t$.

In the worst case where $\gamma_t = 1$, we have $\alpha_{\max} = \eta$, and the bound reduces to standard FedAvg. As training progresses and γ_t decreases, the effective step α_t becomes smaller, reducing variance and drift terms that scale with α_t^2 and improving empirical stability. Since γ_t and τ_i^t are bounded, the update magnitude remains controlled and convergence guarantees remain valid under dynamic fairness adjustments. Although γ_t and τ_i^t are dynamically generated by the game, they remain explicitly bounded by design. Therefore, the effective step size sequence $\{\alpha_t\}$ is time-varying but uniformly bounded, which satisfies the standard conditions for convergence in stochastic optimization with variable step sizes. Since the fairness mechanism only scales aggregated updates and does not modify the underlying gradient distributions, the heterogeneity bound ζ^2 remains valid across all rounds, ensuring that the convergence guarantees remain robust under dynamic fairness adjustments.

VI. EXPERIMENTS AND RESULTS

This section provides the specifics of our evaluation and performance results, along with a discussion of the proposed method.

A. Setup

Datasets. We evaluate FLamma using convolutional neural network (CNN) models, specifically ResNet-18 [53] and LeNet-5 [54], on four widely used benchmark datasets: MNIST [55], FashionMNIST [56], CIFAR10, and CIFAR100 [57]. Each dataset is tested under both independent and identically distributed (IID) and non-IID settings, with non-IID simulated at three skew levels: *Low* ($\alpha = 10$), *Medium* ($\alpha = 1$), and *High* ($\alpha = 0.5$). Here, α is the Dirichlet concentration parameter controlling the degree of label distribution skew among clients, where smaller values indicate greater heterogeneity. Training is conducted for a maximum of 100 global rounds for MNIST and FMNIST, and 150 rounds for CIFAR10 and CIFAR100.

Hyperparameter Sensitivity. We evaluate FLamma under different initial decay factors γ_0 and regularizers λ under severe heterogeneity ($\alpha = 0.5$). Tables I and II show a clear

fairness–performance trade-off. The highlighted rows in the tables are the FLamma setting values. Setting $\gamma_0 = 1.00$ starts similarly to standard FL, while smaller values improve fairness with gradual accuracy reduction, without instability. Since $\gamma_t \in [0, 1]$, updates are scaled but not eliminated, ensuring continued convergence. Because γ_t is updated via a bounded closed-form Stackelberg solution, its adjustment remains smooth and adaptive, preventing instability. Sensitivity to λ shows that smaller values accelerate convergence but increase imbalance, while larger values improve fairness at modest accuracy cost. We select $\gamma_0 = 0.99$ and $\lambda = 0.1$ as they provide the best balance between accuracy, fairness, and stability. We use a standard learning rate $\eta = 0.001$, and since the effective step size $\alpha_t = \eta\gamma_t$ remains bounded, stability is preserved across standard learning-rate ranges.

FL Baselines. We evaluate FLamma against representative FL baselines, including FedAvg [1], FedProx [16], q-FFL [14], the Incentive method (referred to in this paper as IncFL) [12], and IAFL [17]. IncFL builds upon earlier incentive mechanisms such as [36] and FiFL [58], extending them to model client contributions more explicitly. FedAvg and FedProx serve as standard FL baselines, q-FFL provides a fairness-oriented comparison, and IncFL/IAFL represent incentive-based methods. These baselines capture the main FL challenges relevant to assessing FLamma in IoT-oriented heterogeneous systems.

- **FedAvg:** This baseline serves as the simplest baseline with random client selection without any specialized mechanism to address client heterogeneity.
- **FedProx:** FedProx enhances the FedAvg algorithm by incorporating a proximal term into the objective function. This adjustment helps address the challenges posed by non-IID data distributions, commonly encountered in FL.
- **q-FFL:** Improves fairness by weighting clients according to their empirical losses. In contrast, FLamma uses a game-theoretic decay factor γ that dynamically equalizes client contributions over time.
- **IncFL:** This approach evaluates each client's marginal contribution to the global model, prioritizing clients whose updates improve accuracy.
- **IAFL:** An incentive-aware method that provides differentiated training-time model rewards at each FL iteration, proportional to each client's contribution.

By contrasting conventional, fairness-focused, and recent incentivization baselines, our evaluation provides a comprehensive and relevant comparison aligned with modern developments in FL for heterogeneous IoT systems.

B. Evaluation Metrics

Our evaluation focuses on three key performance metrics: Testing accuracy, accuracy variance, and contribution deviation from the mean, which together assess the effectiveness, fairness, and balance of the FL system in heterogeneous IoT environments.

Testing Accuracy. Testing accuracy measures the overall performance of the global model after aggregating local models, expressed as the percentage of correct predictions.

TABLE III
ACCURACY AND ACCURACY VARIANCE (MEAN $\pm$ STD) ACROSS SKEW LEVELS (IID, LOW $\alpha=10$, MEDIUM $\alpha=1$, HIGH $\alpha=0.5$).

Dataset	Setting	FedAvg		FedProx		q-FFL		IncFL		IAFL		FLamma	
		Acc(%)	Var	Acc(%)	Var	Acc(%)	Var	Acc(%)	Var	Acc(%)	Var	Acc(%)	Var
MNIST	IID	98.1 $\pm$ 0.24	5.7 $\pm$ 0.9	98.2 $\pm$ 0.36	5.5 $\pm$ 0.9	98.3 $\pm$ 0.27	5.2 $\pm$ 0.7	94.2 $\pm$ 0.64	5.9 $\pm$ 0.8	98.5 $\pm$ 0.22	2.1 $\pm$ 0.6	98.7$\pm$0.16	1.2$\pm$0.5
	Low ($\alpha=10$)	98.0 $\pm$ 0.26	14.0 $\pm$ 0.5	98.1 $\pm$ 0.28	12.1 $\pm$ 0.6	98.0 $\pm$ 0.57	13.5 $\pm$ 0.4	97.2 $\pm$ 0.34	13.5 $\pm$ 0.5	98.9$\pm$0.18	10.1 $\pm$ 0.3	98.8 $\pm$ 0.19	2.5$\pm$0.3
	Medium ($\alpha=1$)	90.4 $\pm$ 0.58	60.1 $\pm$ 1.2	93.3 $\pm$ 0.31	28.0 $\pm$ 0.9	93.2 $\pm$ 0.33	30.0 $\pm$ 1.0	92.1 $\pm$ 0.40	34.0 $\pm$ 1.1	96.9 $\pm$ 0.28	12.0 $\pm$ 0.6	98.0$\pm$0.22	8.0$\pm$0.5
	High ($\alpha=0.5$)	80.3 $\pm$ 0.74	143.0 $\pm$ 1.8	88.1 $\pm$ 0.44	63.1 $\pm$ 1.2	86.4 $\pm$ 0.48	68.0 $\pm$ 1.3	85.2 $\pm$ 0.53	72.0 $\pm$ 1.4	94.1 $\pm$ 0.36	20.0 $\pm$ 0.7	96.2$\pm$0.30	12.0$\pm$0.6
FMNIST	IID	98.3 $\pm$ 0.15	6.7 $\pm$ 0.7	97.9 $\pm$ 0.18	6.1 $\pm$ 0.6	98.1 $\pm$ 0.22	5.9 $\pm$ 0.6	84.4 $\pm$ 0.54	3.84 $\pm$ 0.7	98.8$\pm$0.21	1.8 $\pm$ 0.5	98.5 $\pm$ 0.19	1.4$\pm$0.4
	Low ($\alpha=10$)	97.1 $\pm$ 0.22	18.0 $\pm$ 0.9	97.6 $\pm$ 0.21	16.0 $\pm$ 0.7	97.2 $\pm$ 0.23	12.5 $\pm$ 0.8	96.8 $\pm$ 0.25	10.0 $\pm$ 0.6	98.5 $\pm$ 0.18	4.0$\pm$0.4	98.6$\pm$0.18	4.2 $\pm$ 0.4
	Medium ($\alpha=1$)	85.2 $\pm$ 0.65	119.0 $\pm$ 2.2	94.1 $\pm$ 0.27	70.4 $\pm$ 1.4	93.4 $\pm$ 0.39	88.5 $\pm$ 0.42	62.1 $\pm$ 1.5	96.1 $\pm$ 0.26	20.2 $\pm$ 0.8	97.0$\pm$0.24	14.1$\pm$0.6	97.0$\pm$0.24
	High ($\alpha=0.5$)	75.1 $\pm$ 0.82	239.0 $\pm$ 3.0	82.4 $\pm$ 0.36	142.0 $\pm$ 1.5	91.4 $\pm$ 0.52	71.1 $\pm$ 2.1	80.1 $\pm$ 0.58	127.5 $\pm$ 1.9	93.4 $\pm$ 0.33	32.9 $\pm$ 0.9	95.2$\pm$0.30	24.3$\pm$0.8
CIFAR10	IID	75.3 $\pm$ 0.25	43.2 $\pm$ 1.3	75.1 $\pm$ 0.18	36.18 $\pm$ 1.1	75.6 $\pm$ 0.21	27.12 $\pm$ 1.2	57.1 $\pm$ 0.98	9.04 $\pm$ 0.9	80.1 $\pm$ 0.40	8.10 $\pm$ 0.9	84.4$\pm$0.32	7.5$\pm$0.8
	Low ($\alpha=10$)	62.4 $\pm$ 0.52	521.0 $\pm$ 6.1	66.1 $\pm$ 0.49	422.0 $\pm$ 5.4	68.2 $\pm$ 0.45	365.0 $\pm$ 4.9	70.3 $\pm$ 0.56	247.0 $\pm$ 3.2	74.1 $\pm$ 0.41	156.0 $\pm$ 2.1	78.5$\pm$0.48	122.5$\pm$2.0
	Medium ($\alpha=1$)	40.1 $\pm$ 0.73	921.0 $\pm$ 7.2	45.2 $\pm$ 0.68	788.0 $\pm$ 6.5	48.5 $\pm$ 0.61	607.0 $\pm$ 5.6	52.0 $\pm$ 0.70	428.0 $\pm$ 3.9	60.2 $\pm$ 0.55	261.0 $\pm$ 2.6	65.3$\pm$0.62	202.4$\pm$2.5
	High ($\alpha=0.5$)	31.9 $\pm$ 0.81	1253.0 $\pm$ 8.1	38.4 $\pm$ 0.74	1045.0 $\pm$ 7.4	42.1 $\pm$ 0.70	823.0 $\pm$ 6.6	46.2 $\pm$ 0.77	600.0 $\pm$ 4.6	60.7$\pm$0.63	322.0 $\pm$ 2.9	60.1 $\pm$ 0.69	265.0$\pm$2.4
CIFAR100	IID	61.2 $\pm$ 0.91	58.0 $\pm$ 1.4	61.1 $\pm$ 0.27	52.0 $\pm$ 1.3	61.5 $\pm$ 0.26	45.0 $\pm$ 1.2	55.4 $\pm$ 0.52	18.0 $\pm$ 0.9	65.1 $\pm$ 0.36	9.8$\pm$0.7	68.4$\pm$0.33	10.1 $\pm$ 0.7
	Low ($\alpha=10$)	50.1 $\pm$ 0.55	688.2 $\pm$ 6.8	53.2 $\pm$ 0.51	591.3 $\pm$ 6.0	54.6 $\pm$ 0.48	524.9 $\pm$ 5.4	56.1 $\pm$ 0.57	389.9 $\pm$ 4.1	58.9 $\pm$ 0.45	260.1 $\pm$ 2.7	61.8$\pm$0.49	228.7$\pm$2.4
	Medium ($\alpha=1$)	35.4 $\pm$ 0.78	1082.4 $\pm$ 7.6	39.1 $\pm$ 0.70	925.2 $\pm$ 5.8	41.3 $\pm$ 0.66	767.1 $\pm$ 6.0	44.2 $\pm$ 0.71	543.2 $\pm$ 4.7	47.9 $\pm$ 0.58	365.5 $\pm$ 3.5	50.6$\pm$0.62	304.5$\pm$2.9
	High ($\alpha=0.5$)	25.3 $\pm$ 0.90	1451.1 $\pm$ 8.5	29.5 $\pm$ 0.84	1280.5 $\pm$ 7.8	32.4 $\pm$ 0.79	982.6 $\pm$ 6.8	35.1 $\pm$ 0.85	720.5 $\pm$ 5.2	40.2 $\pm$ 0.69	423.9 $\pm$ 3.2	45.1$\pm$0.74	361.4$\pm$2.8

Accuracy Variance. Accuracy variance measures the inconsistency in performance across different clients. A lower accuracy variance indicates that all clients achieve similar accuracy levels, reflecting fairness in the learning process.

Contribution Deviation from Mean. This metric quantifies the fairness of client participation by measuring the standard deviation of individual client contributions from the mean contribution across all clients. A lower value indicates that client contributions are more evenly balanced, preventing over-representation of certain clients and under-utilization of others. The metric is data-independent, making the results comparable across different datasets and training scenarios, and Fig. 2 illustrates the results for CIFAR10 as a representation.

Details of Experimental Setup. Each experiment was run 3 times to ensure correctness and accuracy. The experiments were executed on a server equipped with an NVIDIA GeForce RTX 3080Ti GPU, Intel(R) Core(TM) i9-109000X CPU, and 64G of RAM. We run both the server and clients on the same machine in a simulated environment where the clients train their local models for a number of local iterations on their own data. Once completing their local iterations, they communicate their updated weights to the server. This configuration is validated by the fact that our evaluated performance metrics are independent of the physical separation between the server and its clients, reflecting typical IoT-based FL deployments.

Clients' Contributions. Clients train for ten local iterations, after which the server computes their contributions based on the Euclidean distance between each client's weights and the global model. Clients whose weights are closer to the global model receive higher contribution values (between 0 and 1), while those farther away receive lower values. Contributions are updated every ten global rounds and guide client selection throughout training. As selected clients continue contributing, a decay factor gradually reduces their contribution values over time. We use LeNet-5 for MNIST and FMNIST, and ResNet-18 for CIFAR10 and CIFAR100.

Decay Factor Adjustment. The decay factor plays a pivotal role in balancing client contributions throughout the training process. Initially set closer to 1, it allows for greater exploration by ensuring all clients to have substantial impacts on

the global model updates during the early stages of training. This is crucial for capturing general patterns from diverse data distributions. As training progresses, the decay factor is gradually reduced, diminishing the influence of each initial high contributor and shifting focus toward low-contributing clients.

C. Performance Results and Discussion

Performance under IID Data Split. As shown in Table III, FLamma consistently outperforms the closest baselines, IncFL and IAFL, under IID settings. On CIFAR10, FLamma improves accuracy from 80.1% (IAFL) to 84.4% and further reduces variance from 8.10 to 7.5. Similar trends appear in MNIST and FMNIST, where FLamma matches or slightly surpasses IAFL in accuracy while achieving the lowest variance across all methods. On CIFAR100, FLamma improves accuracy from 65.1% (IAFL) to 68.4% while maintaining comparable fairness. These results demonstrate that even when data is uniformly distributed, FLamma provides both competitive accuracy and stronger fairness by minimizing accuracy disparities across clients.

Performance under Non-IID Data Split. Under heterogeneous data splits, FLamma shows even larger gains. On CIFAR10 (High skew), it improves accuracy from 46.2% (IncFL) to 60.1% (with IAFL achieving a comparable 60.7%), while reducing variance from 600.0 to 265.0. In FMNIST, FLamma boosts accuracy from 93.4% (IAFL) to 95.2% and reduces variance from 32.9 to 24.3. Similarly, on MNIST, FLamma raises accuracy from 94.1% (IAFL) to 96.2% while lowering variance from 20.0 to 12.0. In CIFAR100, FLamma improves accuracy from 40.2% (IAFL) to 45.1% and reduces variance from 423.9 to 361.4. These results confirm FLamma's robustness under severe heterogeneity, consistently enhancing fairness while maintaining strong accuracy across all datasets.

VII. LIMITATION AND FUTURE WORK

Despite improving fairness and efficiency, FLamma has several limitations. First, it relies on estimating client contributions, which is challenging in heterogeneous IoT environments. FLamma uses Euclidean distance between local

and global models as an efficient proxy for contribution quality. While computationally lightweight, this metric may be less reliable under severe non-IID settings, where divergence may reflect data heterogeneity rather than poor contribution. More robust alternatives, such as gradient similarity or loss-based metrics, could improve reliability. Second, no single decay factor is optimal across all tasks, requiring task-specific tuning. Future work may explore adaptive or learning-based approaches, such as reinforcement learning, to dynamically optimize hyperparameters under diverse IoT conditions. Third, FLamma currently uses a fixed refresh interval for contribution estimates. Although the framework remains adaptive through per-round decay updates and client best responses, the refresh interval is a tunable hyperparameter that must be specified by the user and is not automatically adjusted. Future work may investigate fully adaptive refresh strategies based on system dynamics, such as variance- or participation-triggered updates, to improve responsiveness in highly dynamic environments.

VIII. CONCLUSION

In this paper, we proposed an FL framework inspired by Stackelberg game theory, which dynamically assigns local training epochs and regulates client contributions through a decay factor. The key innovation of our method is the dynamic allocation of both local epochs and the decay factor to incentivize clients to maximize their contributions while preventing any device from dominating the updates, ensuring balanced participation across all clients and promoting fairness without undermining the overall learning process. To further address fairness, FLamma dynamically assigns weights to clients based on their proximity to the global model. By gradually balancing the contributions of clients over time, it promotes equity and enhances collaboration in heterogeneous IoT environments. Our experimental results demonstrate that FLamma significantly reduces the variance in accuracy among clients compared to baseline methods and highlight FLamma's ability to harmonize client participation and maintain stable learning despite variation in IoT device capabilities. Overall, FLamma offers a scalable and practical solution to fairness and efficiency challenges in FL, supporting deployment across diverse real-world IoT settings.

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)*, April 2017.
- [2] M. S. Islam, S. Javaherian, F. Xu, X. Yuan, L. Chen, and N.-F. Tzeng, "Fedclust: Tackling data heterogeneity in federated learning through weight-driven client clustering," in *Proceedings of the 53rd International Conference on Parallel Processing*, 2024, pp. 474–483.
- [3] T. Li, S. Hu, A. Beirami, and V. Smith, "Ditto: Fair and robust federated learning through personalization," in *Proceedings of International Conference on Machine Learning*. PMLR, 2021, pp. 6357–6368.
- [4] F. Lai, X. Zhu, H. V. Madhyastha, and M. Chowdhury, "Oort: Efficient federated learning via guided participant selection," in *Proceedings of 15th USENIX Symposium on Operating Systems Design and Implementation (OSDI 21)*, 2021, pp. 19–35.
- [5] Y. H. Ezzeldin, S. Yan, C. He, E. Ferrara, and A. S. Avestimehr, "Fairfed: Enabling group fairness in federated learning," in *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 37, no. 6, 2023, pp. 7494–7502.
- [6] A. Ghosh, J. Chung, D. Yin, and K. Ramchandran, "An efficient framework for clustered federated learning," in *Proceedings of Advances in Neural Information Processing Systems*, vol. 33, pp. 19 586–19 597, 2020.
- [7] Y. J. Cho, J. Wang, and G. Joshi, "Client selection in federated learning: Convergence analysis and power-of-choice selection strategies," *arXiv preprint arXiv:2010.01243*, 2020.
- [8] A. Blum, N. Haghtalab, R. L. Phillips, and H. Shao, "One for one, or all for all: Equilibria and optimality of collaboration in federated learning," in *Proceedings of International Conference on Machine Learning*. PMLR, 2021, pp. 1005–1014.
- [9] M. Maschler, S. Zamir, and E. Solan, *Game theory*. Cambridge University Press, 2020.
- [10] L. U. Khan, S. R. Pandey, N. H. Tran, W. Saad, Z. Han, M. N. Nguyen, and C. S. Hong, "Federated learning for edge networks: Resource optimization and incentive mechanism," *IEEE Communications Magazine*, vol. 58, no. 10, pp. 88–93, 2020.
- [11] Y. Sarikaya and O. Ercetin, "Motivating workers in federated learning: A stackelberg game perspective," *IEEE Networking Letters*, vol. 2, no. 1, pp. 23–27, 2019.
- [12] H. Wu, X. Tang, Y.-J. A. Zhang, and L. Gao, "Incentive mechanism for federated learning with random client selection," *IEEE Transactions on Network Science and Engineering*, vol. 11, pp. 1922–1933, March-April 2024.
- [13] A. Murhekar, Z. Yuan, B. Ray Chaudhury, B. Li, and R. Mehta, "Incentives in federated learning: Equilibria, dynamics, and mechanisms for welfare maximization," in *Proceedings of Advances in Neural Information Processing Systems*, vol. 36, 2024.
- [14] T. Li, M. Sanjabi, A. Beirami, and V. Smith, "Fair resource allocation in federated learning," *arXiv preprint arXiv:1905.10497*, 2019.
- [15] S. M. Hamidi and E.-H. Yang, "Adafed: Fair federated learning via adaptive common descent direction," *arXiv preprint arXiv:2401.04993*, 2024.
- [16] T. Li, A. K. Sahu, M. Zaheer, M. Sanjabi, A. Talwalkar, and V. Smith, "Federated optimization in heterogeneous networks," in *Proceedings of Machine learning and systems*, vol. 2, pp. 429–450, 2020.
- [17] Z. Wu, M. M. Ghosh, R. Raskar, and B. K. H. Low, "Incentive-aware federated learning with training-time model rewards," in *Proceedings of 12th International Conference on Learning Representations*, 2024.
- [18] W. Zhang, Q. Wang, H. Zhao, W. Xia, and H. Zhu, "Incentivizing quality contributions in federated learning: A stackelberg game approach," in *Proceedings of 2024 IEEE 99th Vehicular Technology Conference (VTC2024-Spring)*. IEEE, 2024, pp. 1–5.
- [19] K. Khawam, H. Taleb, S. Lahoud, H. Fawaz, D. Quadri, and S. Martin, "Non-cooperative edge server selection game for federated learning in iot," in *Proceedings of NOMS 2024-2024 IEEE Network Operations and Management Symposium*. IEEE, 2024, pp. 1–6.
- [20] K. Donahue and J. Kleinberg, "Model-sharing games: Analyzing federated learning under voluntary participation," in *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 35, no. 6, 2021, pp. 5303–5311.
- [21] N. Zhang, Q. Ma, W. Mao, and X. Chen, "Coalitional fl: Coalition formation and selection in federated learning with heterogeneous data," *IEEE Transactions on Mobile Computing*, vol. 23, pp. 10 494–10 508, November 2024.
- [22] G. Hu, J. Han, J. Lu, J. Yu, S. Qiu, H. Peng, D. Zhu, and T. Li, "Game-theoretic design of quality-aware incentive mechanisms for hierarchical federated learning," *IEEE Internet of Things Journal*, 2024.
- [23] S. Arisdakessian, O. A. Wahab, A. Mourad, and H. Otrok, "Coalitional federated learning: Improving communication and training on non-iid data with selfish clients," *IEEE Transactions on Services Computing*, vol. 16, no. 4, pp. 2462–2476, 2023.
- [24] C. Hasan, "Incentive mechanism design for federated learning: Hedonic game approach," *arXiv preprint arXiv:2101.09673*, 2021.
- [25] Y. J. Cho, D. Jhunjunwala, T. Li, V. Smith, and G. Joshi, "To federate or not to federate: incentivizing client participation in federated learning," in *Proceedings of Workshop on Federated Learning: Recent Advances and New Challenges (in Conjunction with NeurIPS 2022)*, 2022.
- [26] J. Cao, Y. Gao, and J. Huang, "A service-oriented adaptive hierarchical incentive mechanism for federated learning," *arXiv preprint arXiv:2509.10512*, 2025.
- [27] N. Singh and M. Adhikari, "Fedtone-sgm: A stackelberg-driven personalized federated learning strategy for edge networks," *IEEE Transactions on Parallel and Distributed Systems*, 2025.
- [28] Y. Liu, M. Tian, Y. Chen, Z. Xiong, C. Leung, and C. Miao, "A contract theory based incentive mechanism for federated learning," in *Federated and Transfer Learning*. Springer, 2022, pp. 117–137.

- [29] N. Ding, Z. Fang, and J. Huang, "Optimal contract design for efficient federated learning with multi-dimensional private information," *IEEE Journal on Selected Areas in Communications*, vol. 39, no. 1, pp. 186–200, 2020.
- [30] Y. Zhan and J. Zhang, "An incentive mechanism design for efficient edge learning by deep reinforcement learning approach," in *Proceedings of 2020 IEEE Conference on Computer Communications*. IEEE, 2020, pp. 2489–2498.
- [31] N. Ding, L. Gao, and J. Huang, "Incentive mechanism design for federated learning with dynamic network pricing," *IEEE Transactions on Mobile Computing*, 2025.
- [32] L. Yu, Z. Chang, and Z. Zhao, "Contract-based incentive mechanism for federated learning in edge computing system," in *Proceedings of 2024 IEEE Wireless Communications and Networking Conference (WCNC)*. IEEE, 2024, pp. 1–6.
- [33] J. Kang, Z. Xiong, D. Niyato, S. Xie, and J. Zhang, "Incentive mechanism for reliable federated learning: A joint optimization approach to combining reputation and contract theory," *IEEE Internet of Things Journal*, vol. 6, no. 6, pp. 10700–10714, 2019.
- [34] M. Tang and V. W. Wong, "An incentive mechanism for cross-silo federated learning: A public goods perspective," in *Proceedings of 2021 IEEE Conference on Computer Communications*. IEEE, 2021, pp. 1–10.
- [35] M. M. Rashid, Y. Xiang, M. P. Uddin, J. Tang, K. Sood, and L. Gao, "Trustworthy and fair federated learning via reputation-based consensus and adaptive incentives," *IEEE Transactions on Information Forensics and Security*, vol. 20, pp. 2868–2882, 2025.
- [36] Y. Zhan, P. Li, Z. Qu, D. Zeng, and S. Guo, "A learning-based incentive mechanism for federated learning," *IEEE Internet of Things Journal*, vol. 7, no. 7, pp. 6360–6368, 2020.
- [37] S. Shen, C. Liu, and T. J. Lim, "Federated learning with heterogeneous client expectations: A game theory approach," *IEEE Transactions on Knowledge and Data Engineering*, 2024.
- [38] K. Donahue and J. Kleinberg, "Fairness in model-sharing games," in *Proceedings of the ACM Web Conference 2023*, 2023, pp. 3775–3783.
- [39] Y. Shi, Z. Liu, Z. Shi, and H. Yu, "Fairness-aware client selection for federated learning," in *Proceedings of 2023 IEEE International Conference on Multimedia and Expo (ICME)*. IEEE, 2023, pp. 324–329.
- [40] Q. Li, X. Li, L. Zhou, and X. Yan, "Adafl: Adaptive client selection and dynamic contribution evaluation for efficient federated learning," in *Proceedings of 2024 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*. IEEE, 2024, pp. 6645–6649.
- [41] S. Javaherian, S. Panta, S. Williams, M. S. Islam, and L. Chen, "Fedfair³: Unlocking threefold fairness in federated learning," in *Proceedings of 2024 IEEE International Conference on Communications*, 2024, pp. 3622–3627.
- [42] A. Sultana, M. M. Haque, L. Chen, F. Xu, and X. Yuan, "Eiffel: Efficient and fair scheduling in adaptive federated learning," *IEEE Transactions on Parallel and Distributed Systems*, vol. 33, no. 12, pp. 4282–4294, 2022.
- [43] F. Sabah, Y. Chen, Z. Yang, A. Raheem, M. Azam, N. Ahmad, and R. Sarwar, "Fairdpfl-scs: Fair dynamic personalized federated learning with strategic client selection for improved accuracy and fairness," *Information Fusion*, vol. 115, p. 102756, 2025.
- [44] R. Ye, W.-B. Kou, and M. Tang, "Prafl: A preference-aware scheme in fair federated learning," in *Proceedings of the 31st ACM SIGKDD Conference on Knowledge Discovery and Data Mining V. 1*, 2025, pp. 1797–1808.
- [45] J. Pei, "F3: Fair federated learning framework with adaptive regularization," *Knowledge-Based Systems*, vol. 316, p. 113392, 2025.
- [46] M. Mohri, G. Sivek, and A. T. Suresh, "Agnostic federated learning," in *International conference on machine learning*. PMLR, 2019, pp. 4615–4625.
- [47] F. E. Dorner, N. Konstantinov, G. Pashaliev, and M. Vechev, "Incentivizing honesty among competitors in collaborative learning and optimization," in *Proceedings of Advances in Neural Information Processing Systems*, vol. 36, 2024.
- [48] R. Amir and I. Grilo, "Stackelberg versus cournot equilibrium," *Games and Economic Behavior*, vol. 26, no. 1, pp. 1–21, 1999.
- [49] D. Fudenberg and D. K. Levine, *The theory of learning in games*. MIT press, 1998, vol. 2.
- [50] J. Wang, Q. Liu, H. Liang, G. Joshi, and H. V. Poor, "A novel framework for the analysis and design of heterogeneous federated learning," *IEEE Transactions on Signal Processing*, vol. 69, pp. 5234–5249, 2021.
- [51] J. B. Rosen, "Existence and uniqueness of equilibrium points for concave n-person games," *Econometrica: Journal of the Econometric Society*, pp. 520–534, 1965.
- [52] X. Li, K. Huang, W. Yang, S. Wang, and Z. Zhang, "On the convergence of fedavg on non-iid data," *arXiv preprint arXiv:1907.02189*, 2019.
- [53] K. He, X. Zhang, S. Ren, and J. Sun, "Deep residual learning for image recognition," in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2016, pp. 770–778.
- [54] Y. LeCun, B. Boser, J. S. Denker, D. Henderson, R. E. Howard, W. Hubbard, and L. D. Jackel, "Backpropagation applied to handwritten zip code recognition," *Neural Computation*, vol. 1, no. 4, pp. 541–551, 1989.
- [55] L. Deng, "The mnist database of handwritten digit images for machine learning research [best of the web]," *IEEE Signal Processing Magazine*, vol. 29, no. 6, pp. 141–142, 2012.
- [56] H. Xiao, K. Rasul, and R. Vollgraf, "Fashion-mnist: a novel image dataset for benchmarking machine learning algorithms," *arXiv preprint arXiv:1708.07747*, 2017.
- [57] A. Krizhevsky, G. Hinton *et al.*, "Learning multiple layers of features from tiny images," *cs.utoronto.ca*, 2009.
- [58] L. Gao, L. Li, Y. Chen, W. Zheng, C. Xu, and M. Xu, "Fifi: A fair incentive mechanism for federated learning," in *Proceedings of the 50th International Conference on Parallel Processing*, 2021, pp. 1–10.



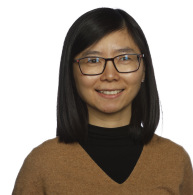
Simin Javaherian (S'23) received her M.S. degree in Computer Science from the Center for Advanced Computer Studies (CACS), at the University of Louisiana at Lafayette. She is currently pursuing a Ph.D. degree in Computer Science from the Center for Advanced Computer Studies (CACS), at the University of Louisiana at Lafayette. Her current areas of interest include federated learning, game theory and parallel computing.



Bryce Turney (S'22) received his B.S. degree in Electrical Engineering from the Department of Electrical and Computer Engineering at the University of Louisiana at Lafayette. He is currently pursuing a Ph.D. degree in Computer Engineering from the Center for Advanced Computer Studies (CACS), at the University of Louisiana at Lafayette. His current areas of interest include machine learning, parallel computing, embedded systems, and autonomous systems.



Xu Yuan (M'16-SM'22) received the Ph.D. degree from the Bradley Department of Electrical and Computer Engineering, Virginia Tech, Blacksburg, in 2016. He is an Associate Professor in the Department of Computer and Information Sciences at the University of Delaware, Newark. His research focuses on artificial intelligence (AI), cybersecurity, networking, and cyber-physical systems. He received the Best Paper Award and the Distinguished Paper Award from DSN 2023.



Li Chen (M'18-SM'24) received her Ph.D. degree from the Department of Electrical and Computer Engineering, University of Toronto, in 2018. She is currently an associate professor with the School of Computing and Informatics, University of Louisiana at Lafayette, USA. Her research interests include big data analytics systems, machine learning systems, cloud computing, datacenter networking, and resource allocation.



Nian-Feng Tzeng (M'86-SM'92-F'10-LF'22) received his Ph.D. degree in Computer Science from the University of Illinois, Urbana-Champaign in 1986. He has been with the School of Computing and Informatics, the University of Louisiana at Lafayette, since 1987. His current research interest includes high-performance computer systems, computer networks, and parallel and distributed processing. He received the Outstanding Paper Award of IEEE ICDCS 1990 and the University Foundation Distinguished Professor Award in 1997.