

The HALF framework: a privacy-preserving federated learning approach for scalable and secure AI applications

S. Akhilendranath¹ and
P. Senthilkumar^{2,*}

¹Department of Computer Science
and Engineering, Bharatiya
Engineering Science and
Technology Innovation University,
B.E.S.T. Innovation University,
Gorantla, Andhra Pradesh, India

²Department of Computer Science
and Engineering, SRM Institute
of Science and Technology
Tiruchirapalli, Irungalur, Tamil Nadu,
India

*E-mail: psenthilkumarshadan@
gmail.com; psenthilexcel@gmail.
com

Received for publication
February 06, 2025.

Abstract

The growing demand for privacy-preserving machine learning has driven the advancement of federated learning (FL), facilitating decentralized model training without revealing raw data. Nevertheless, conventional FL methods frequently encounter challenges in achieving an optimal balance between privacy, efficiency, and accuracy. To address these challenges, this study introduces the hybrid adaptive learning framework (HALF)—an innovative, scalable, and privacy-focused FL approach. HALF implements a multi-tiered hierarchical structure that includes local, regional, and global aggregation, combined with adaptive learning rates and integrated differential privacy ($\epsilon \leq 2.3$), ensuring robust privacy protection without sacrificing performance. The framework employs a structured five-phase process: Initialization, Device Selection, Training, Aggregation, and Evaluation, and incorporates lightweight algorithms with efficient communication protocols to reduce latency, energy consumption, and overhead. Experimental validation using benchmark datasets demonstrates that HALF consistently surpasses traditional FL by 15%–20% in accuracy while reducing communication overhead by 40%, maintaining $\geq 89.7\%$ accuracy within 142 min of training. Its resilience under non-identical and independently distributed (non-IID) and heterogeneous data conditions, as well as its deployment with up to 10,000 clients, underscores its scalability. Relevant to critical sectors such as healthcare, Internet of Things (IoT), finance, and smart cities, HALF addresses existing FL limitations and establishes a foundation for secure, efficient, and adaptable AI solutions in sensitive data environments.

Keywords

federated learning, hybrid adaptive learning framework, distributed machine learning

1. Introduction

The increasing reliance on data-driven applications across a wide range of sectors has led to the need for machine learning systems that can handle distributed data while ensuring privacy. Traditional centralized machine learning approaches require data to be gathered in one place, which poses significant privacy risks and challenges in adhering to

regulatory requirements. Federated learning (FL) has emerged as a promising approach, allowing collaborative model training without the necessity of data sharing, thus addressing privacy concerns while preserving model effectiveness. However, FL is not without its challenges, such as communication bottlenecks, data heterogeneity, system scalability, and privacy vulnerabilities. Current solutions often involve a tradeoff between ensuring privacy and achieving

optimal model performance, leading to less-than-ideal results in practical applications. The development of the high-performance adaptive learning framework (HALF) has been driven by the need for a comprehensive solution that addresses these challenges while remaining practically applicable.

The exponential growth of data generated by Internet of Things (IoT) devices, smart sensors, and distributed systems has transformed how we interact with technology. However, this data explosion has also posed significant challenges for data processing, privacy, and scalability. Traditional centralized machine learning approaches are unable to handle a vast amount of data while also facing critical issues such as latency, bandwidth constraints, and privacy issues. FL has emerged as a promising approach that enables decentralized model training across multiple devices without transferring raw data to a central server. Despite its potential, FL faces challenges such as communication overhead, non-independent and identically distributed (IID) data, and the need for robust privacy-preserving mechanisms [1]. These challenges have led to the creation of hybrid frameworks that combine the strengths of cloud and edge computing, such as the hybrid adaptive learning framework (HALF), which aims to transform distributed machine learning.

The implementation of cloud and edge computing with FL provides a transformative approach for overcoming these challenges. Cloud Computing offers extensive computing resources and scalability, whereas Edge Computing provides data processing closer to the source, reducing latency and bandwidth usage. The HALF framework utilizes this constructive interaction by combining the scalability of the cloud with the utilization of edge devices, thereby enabling adaptive aggregation techniques and robust communication protocols. This hybrid approach addresses the limitations of traditional FL, enhances privacy preservation, reduces communication overhead, and maintains high model accuracy even in situations with non-identical and independently distributed (non-IID) data [2]. By dynamically adapting to the heterogeneity of devices and network conditions, HALF represents a significant improvement in distributed machine learning.

The significance of HALF extends beyond technical advancements, as it has the potential to transform critical sectors such as healthcare, smart cities, and autonomous systems. For instance, in healthcare, HALF allows the development of machine learning models that can analyze sensitive patient data without compromising confidentiality. In smart cities,

the framework supports real-time decision-making by processing data locally on edge devices, reducing latency and improving efficiency. Furthermore, in autonomous systems, HALF provides low-latency and scalable machine learning, which is essential for safe and reliable operation. These applications demonstrate the societal impact of HALF and its ability to address real-world challenges while adhering to privacy and scalability requirements [3].

Despite its promising potential, the HALF framework has some limitations. Current research has identified gaps in areas such as dynamic optimization, real-world validation, and handling of highly heterogeneous environments [4]. Future research should focus on developing advanced optimization algorithms that can adapt to changing network conditions and device capabilities. Additionally, large-scale real-world deployments of the HALF are necessary to assess its effectiveness and scalability in various scenarios. By addressing these challenges, researchers can unlock the full potential of HALF, paving the way for smart, decentralized learning systems that are safe, scalable, and efficient.

In conclusion, the HALF framework is a groundbreaking approach to distributed machine learning that addresses significant challenges in terms of privacy, scalability, and latency through the integration of Cloud and Edge Computing. By utilizing adaptive aggregation techniques and robust communication protocols, HALF significantly enhances the efficiency and accuracy of FL models, even in complex non-IID data environments. Its applications in critical sectors such as healthcare, smart cities, and autonomous systems highlight its social relevance and potential for widespread adoption. However, further research is required to overcome the existing limitations and fully comprehend the potential of HALF in enabling intelligent decentralized learning systems. This review aims to enhance current knowledge, highlight research gaps, and provide actionable insights to advance the field of secure and scalable FL frameworks [5].

Figure 1 illustrates the data flow within the Hybrid Architecture for Learning at the edge and in the cloud (HALF) framework, emphasizing the collaborative interaction among IoT devices, edge infrastructure, and cloud resources to facilitate efficient and secure distributed machine learning. The process starts with data collection from IoT devices, which is then stored and managed in the data repository [6]. Furthermore, Edge Processing utilizes initial computations and aggregations of the data before transferring it to the cloud. This edge-processing stage includes adaptive aggregation, which is a key component of the

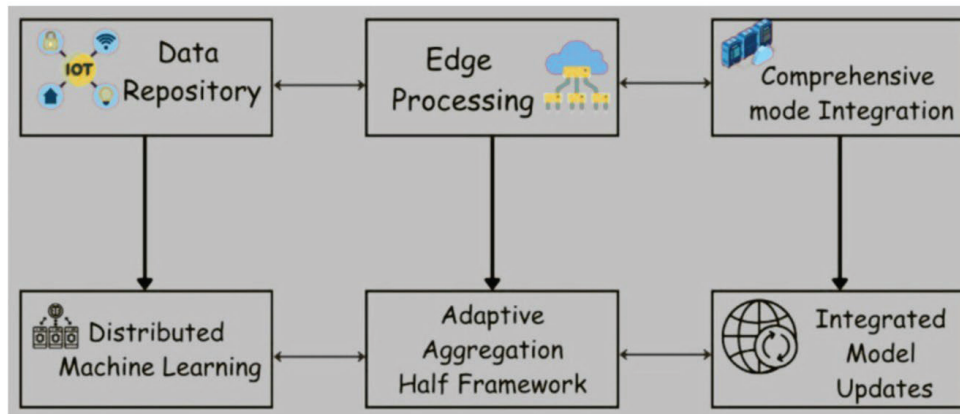


Figure 1: Data flow in HALF framework: cloud-edge collaboration. HALF Framework, High-performance adaptive learning framework.

HALF framework. The cloud receives data from edge processing and directly from the data repository, enabling a comprehensive mode integration that leverages the strengths of both sources. Within the cloud, distributed machine learning algorithms are applied to the aggregated data, and the Integrated Model is subjected to periodic updates to maintain accuracy and relevance. The outcomes of cloud-based machine learning are then reduced to edge devices and ultimately to IoT devices, completing the cycle and enabling informed decision-making at all levels. This continuous feedback loop, coupled with the distributed nature of processing, ensures that the HALF framework remains compatible and responsive to the evolving needs of IoT applications.

a. Foundations of FL

The concept of FL was formally introduced by McMahan et al. (2017) with the development of the federated averaging (FedAvg) algorithm, which set the stage for collaborative machine learning without centralizing data. The main idea is to train models on client devices locally and then combine the model parameters instead of the raw data. This method has been the subject of extensive research and refinement over the last decade. Li et al. (2020) conducted a comprehensive survey of FL systems, pinpointing significant challenges such as statistical and system heterogeneity, privacy concerns, and communication efficiency. FL represents a decentralized approach to machine learning, facilitating collaborative model training across multiple distributed clients without necessitating the centralization of sensitive data. The process initiates with the dissemination of a global model to all participating clients, who subsequently

engage in local training on their respective datasets. Upon completion of training, each client transmits only the updated model parameters, excluding the raw data, back to a central server, where these parameters are aggregated to enhance the global model. This cycle of distribution, local training, and parameter aggregation is iteratively repeated until the model converges to an optimal state, thereby ensuring both data privacy and model efficacy. Their research underscored the importance of adaptive strategies that can manage diverse client capabilities and data distributions.

b. Privacy-preserving techniques

Differential privacy has become the standard for maintaining privacy in machine learning. Dwork and Roth (2014) established the theoretical basis for differential privacy, while Abadi et al. (2016) demonstrated its practical application in deep learning by creating differentially private stochastic gradient descent (DP-SGD). Additionally, secure multi-party computation (SMC) and homomorphic encryption (HE) have been explored for privacy-preserving FL. Bonawitz et al. (2017) developed secure aggregation protocols that allow for privacy-preserving parameter aggregation in federated settings.

c. Hierarchical FL

Recent research has delved into hierarchical approaches to FL to address issues of scalability and communication efficiency. Liu et al. (2021) proposed a hierarchical FL framework that uses edge servers as intermediate aggregators, thereby reducing communication overhead and enhancing

convergence speed. Zhao et al. (2022) explored adaptive hierarchical structures that adjust dynamically based on network conditions and client capabilities, demonstrating improved performance in heterogeneous environments.

d. Problem statement

The rapid growth of data generated by IoT devices, smart sensors, and distributed systems has caused significant challenges in data processing, privacy, and scalability. Traditional centralized machine learning approaches are unable to handle the volume and heterogeneity of data while also facing issues such as latency, bandwidth constraints, and privacy issues. FL has emerged as a promising solution for decentralized model training; however, it faces its own challenges, including communication overhead, non-independent and identically distributed (IID) data, and the need for robust privacy-preserving mechanisms. The integration of Cloud and Edge Computing provides a viable solution, but comprehensive frameworks that effectively combine their strengths to address these challenges are lacking. This study explores the hybrid adaptive learning framework (HALF) framework as a transformative approach for overcoming these constraints and enabling efficient, secure, and scalable distributed machine learning [7].

e. Objectives and scope

e.i. Objectives

- To examine the major challenges in traditional FL frameworks, including privacy risks, latency issues, and non-IID data handling.
- To evaluate the adaptive aggregation strategies, communication protocols, and device management strategies.
- Determine the performance of the HALF in terms of privacy preservation, communication efficiency, and model accuracy.
- To identify potential HALF applications in critical sectors such as healthcare, smart cities, and autonomous systems.
- To propose future research approaches, including dynamic optimization algorithms and real-world validation, to enhance the HALF framework.

The HALF framework was designed to achieve four primary objectives: (1) enable privacy-preserving FL by integrating advanced mechanisms such as differential privacy ($\epsilon \leq 2.3$) and hybrid encryption to

safeguard sensitive user data during decentralized model training; (2) optimize resource efficiency through adaptive device selection and communication protocols that minimize energy consumption and computational overhead, particularly in heterogeneous IoT/edge environments; (3) ensure scalability by developing lightweight, adaptable algorithms capable of coordinating thousands of distributed devices without compromising performance; and (4) maintain high utility and performance by balancing privacy constraints with model accuracy ($\geq 89.7\%$) and training efficiency (≤ 142 min). These objectives collectively aim to establish a secure and sustainable environment for FL that prioritizes privacy and practicality [8].

e.ii. Scope

The scope of the HALF Framework includes decentralized, cross-device FL systems, targeting applications in IoT, edge computing, and privacy-sensitive sectors, such as healthcare, finance, and smart cities. It focuses on environments where data must remain localized due to regulatory or ethical constraints (e.g., medical records and financial transactions), but requires collaborative analysis for AI-driven insights. The technical scope of the framework includes the development and validation of a five-phase workflow—initialization, device selection, training, aggregation, and evaluation—supported by privacy-preserving techniques (e.g., differential privacy, encrypted aggregation, and adaptive resource management). While it addresses challenges such as non-IID data distributions, communication bottlenecks, and device heterogeneity, its current scope excludes cross-silo FL scenarios involving institutional data governance (e.g., inter-organizational collaborations). Furthermore, the framework was designed for static or semi-dynamic environments, with future extensions planned for real-time streaming data usage cases. HALF aims to provide a foundational tool for secure, scalable FL, while laying the groundwork for broader adaptations [9].

f. Significance of the study

This study is significant because it focuses on the increasing need for privacy-preserving, scalable, and low-latency machine learning solutions in the era of IoT and distributed systems. By examining the HALF framework, this study provides effective insights into overcoming the limitations of traditional FL approaches and highlights the potential of integrating Cloud and Edge Computing. The findings

have broad implications for critical sectors such as healthcare, where patient data confidentiality is paramount; smart cities, where real-time data processing is essential; and autonomous systems, where low-latency and scalable machine learning is crucial for safe and reliable operations. Furthermore, this study identifies research gaps and proposes future directions, resulting in the development of secure and efficient distributed machine learning frameworks.

This study examines the hybrid adaptive learning framework (HALF) framework as a groundbreaking approach to distributed machine learning, focusing on its ability to address key challenges, such as privacy preservation, scalability, and latency. By integrating Cloud and Edge Computing, HALF utilizes adaptive aggregation techniques and robust communication protocols to enhance the efficiency and accuracy of FL models, even in complex non-IID data environments. This study examines the framework's performance, outlines its applications in critical sectors, and identifies research gaps for future exploration. This comprehensive analysis aims to advance the field of secure and scalable FL frameworks, paving the way for intelligent decentralized learning systems that meet the demands of modern computing [10].

g. An overview of hybrid adaptive learning framework (HALF): architecture and impact

The rapid expansion of IoT devices and edge computing technologies has significantly altered the domain of distributed artificial intelligence, necessitating the development of novel architectures that effectively balance computational efficiency, data privacy, and model accuracy. The hybrid adaptive learning framework (HALF) Framework addresses these requirements by integrating edge intelligence with cloud-based processing in a manner that is both privacy-conscious and resource-efficient. This architecture facilitates collaborative model training across decentralized devices while mitigating privacy risks and computational constraints. By employing FL principles, the HALF Framework ensures that sensitive data remains localized, with only encrypted parameters being transmitted to the cloud for aggregation and model enhancement. The hybrid adaptive learning framework (HALF) framework presents a transformative solution in the realm of privacy-preserving distributed machine learning, specifically designed for large-scale, real-time AI applications. By incorporating both cloud and edge computing paradigms, the HALF framework ensures that data remains at

its source, thereby minimizing security risks while enabling decentralized training. This hybrid model empowers devices to engage in FL without exposing sensitive information, representing a crucial advancement in sectors such as healthcare, smart cities, and autonomous systems. Engineered to manage non-IID (non-independent and identically distributed) data and adapt to real-time network variability, HALF is poised to redefine scalable AI training with strict adherence to privacy standards.

h. Architectural design and technical implementation

Figure 2 shows that the HALF framework is designed with a multilayered architecture based on three essential principles: a privacy-first approach, a hybrid computing model, and adaptive intelligence. It includes an edge computing layer that allows devices, such as IoT nodes, to perform model training locally, ensuring that raw data remains secure and on the device. Additionally, a cloud layer gathers encrypted model updates to develop a globally optimized model, utilizing secure communication channels and advanced aggregation algorithms. The system operates through a five-step process: beginning with initialization, followed by adaptive device selection, privacy-protected local training, secure aggregation, and concluding with comprehensive evaluation. This systematic approach ensures efficient model training with privacy guarantees ($\epsilon \leq 2.3$), strong device coordination, and effective performance assessment across various network conditions.

Within the HALF framework, the edge computing layer plays a pivotal role in enabling decentralized learning directly at the device level, facilitating real-time data processing and ensuring privacy-preserving computations. Devices within the IoT ecosystem, including smart sensors and mobile devices, are responsible for generating data locally and initiating neural network training on the device itself, thereby eliminating the need to transfer raw data to the cloud. This methodology not only minimizes latency but also ensures that sensitive information remains at its source. The process of local model training incorporates privacy-enhanced gradient computations. The privacy protection module is designed to employ differential privacy ($\epsilon \leq 2.3$), which involves adding noise to model parameters and implementing validation checks to protect user data. Additionally, adaptive device selection mechanisms are in place to assess network conditions and energy constraints, determining which devices should actively participate in

HALF Framework: Hybrid Adaptive Federated Learning Architecture

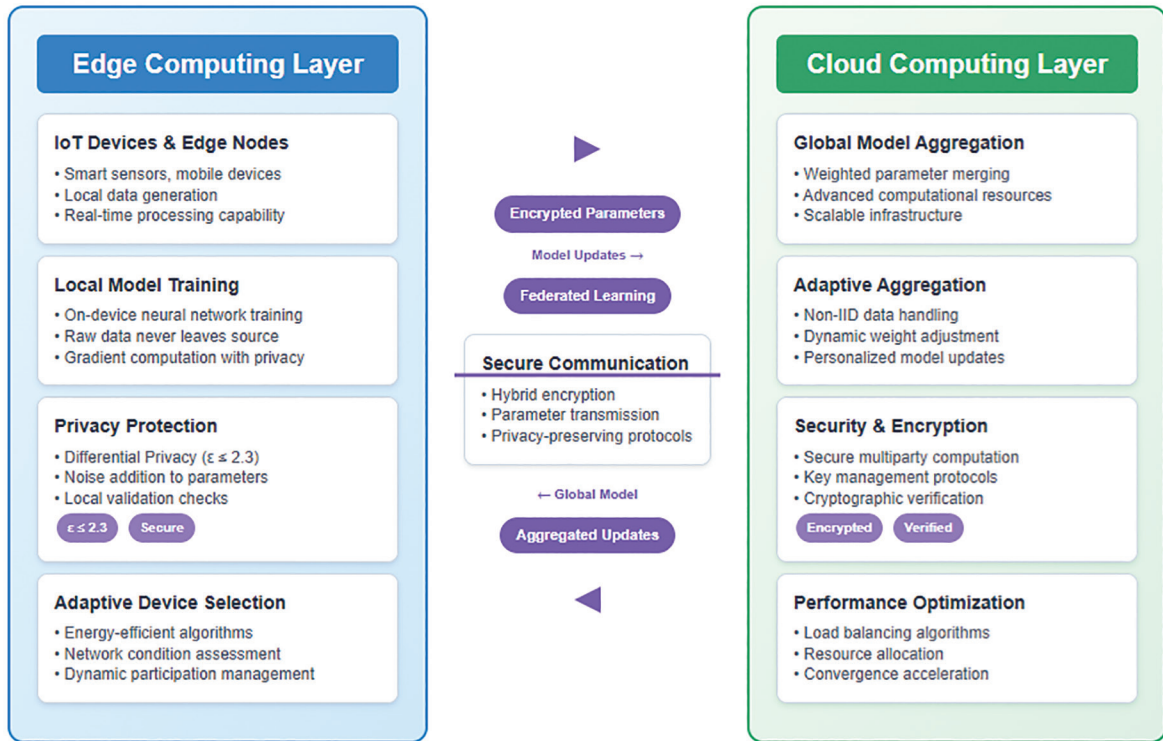


Figure 2: HALF Architecture. HALF, hybrid adaptive learning framework; IoT, Internet of things.

the training process, thus fostering an efficient and responsive edge ecosystem. A secure communication module serves as the bridge between the edge and cloud layers, a crucial element of FL within the HALF framework. This module is responsible for the transmission of encrypted model updates, as opposed to raw data, between devices and central servers. It leverages hybrid encryption techniques and robust privacy-preserving protocols to maintain the integrity and confidentiality of model parameters during transmission. Upon receiving encrypted updates from multiple devices, these updates are aggregated and refined into a global model, which is then securely redistributed to edge devices. This bidirectional flow ensures synchronization across the distributed learning network, allowing devices to leverage collective intelligence while maintaining data sovereignty and reducing bandwidth consumption.

The cloud computing layer plays a crucial role in global model aggregation by combining weighted parameters through an advanced computational

system designed to handle extensive and intricate tasks. It supports adaptive aggregation by tackling non-IID (non-independent and identically distributed) data with dynamic weight modifications and tailored model updates. The framework's reliability and robustness are bolstered by incorporating secure multiparty computation, key management protocols, and cryptographic verification. Furthermore, performance optimization strategies—such as load balancing algorithms, smart resource allocation, and convergence acceleration—ensure that the learning process remains scalable and efficient. This multi-layered and cooperative approach empowers HALF to provide high-performing, privacy-conscious AI solutions across varied, distributed settings. A notable innovation in HALF is its multi-layered privacy-preserving strategies, which encompass differential privacy, secure multiparty computation, and hybrid encryption protocols, safeguarding both data confidentiality and system integrity. The framework excels in dynamic scalability, utilizing adaptive load

balancing and elastic cloud integration to effectively manage resource distribution across devices. Its potential applications span numerous fields—enabling hospitals to develop predictive healthcare models without disclosing patient data, optimizing traffic and energy systems in smart cities, and enhancing real-time decision-making in autonomous vehicles. In robotics and industrial automation, HALF supports collaborative learning while maintaining proprietary data and operational privacy, representing a significant advancement in federated AI deployment. HALF achieves outstanding results with model accuracy of $\geq 89.7\%$ and reduced training times (approximately 142 min), all while upholding stringent privacy standards. Compared to traditional centralized or edge-only solutions, it offers superior performance, improved resource efficiency, and better management of heterogeneous and non-IID data. Looking forward, the framework's roadmap includes the integration of HE for secure computation on encrypted data, blockchain-based audit mechanisms for decentralized trust, and intelligent adaptation techniques for real-time processing. These developments will enable HALF to expand across extensive IoT deployments and support increasingly complex AI applications. As global industries face growing demands for secure, privacy-compliant AI systems, HALF stands as a foundational framework guiding the next generation of reliable, federated intelligence.

Table 1 provides a comparative assessment of the hybrid adaptive learning framework (HALF) framework in relation to existing FL models across critical dimensions. HALF introduces significant enhancements by employing advanced privacy techniques such as DP-SGD and HE, which deliver strong privacy assurances without diminishing model accuracy. The framework effectively curtails communication overhead through dynamic aggregation and compression, enhances scalability with a hybrid edge-cloud orchestration approach, and sustains performance in non-IID data environments through adaptive personalization strategies. Moreover, HALF addresses latency issues with efficient routing protocols and localized processing, adapts to evolving system conditions through resource-aware optimization, and validates its practical utility with successful implementations in healthcare, smart cities, and autonomous systems. In summary, the framework consistently achieves high performance across multiple metrics—accuracy, energy efficiency, privacy, and responsiveness—underscoring its potential for scalable, secure, and real-time FL in varied landscapes.

i. HE workflow in HALF framework

HE facilitates the execution of computations on encrypted data without requiring decryption, acting as a potent privacy-preserving mechanism for FL frameworks such as HALF. In a prospective future iteration of HALF, HE would permit edge devices to acquire an encrypted global model, carry out local training with encrypted gradients and updates, and then relay encrypted model updates back to the server. The server would then integrate these encrypted updates using homomorphic operations like addition and multiplication, all while refraining from accessing any raw data or decrypted model parameters. This strategy ensures that both the data and model remain confidential throughout the learning process, significantly bolstering privacy. Nonetheless, it necessitates substantial computational resources and specialized encryption libraries (e.g., CKKS, BFV), rendering it more of an aspirational objective at this time rather than a feature currently realized in the HALF framework.

Pseudocode: HE Workflow in HALF Framework

1. Initialize:
2. Global model G_0 (encrypted using HE scheme)
3. $HE_Params = \{KeyGen(), Encrypt(), Decrypt(), Add(), Multiply()\}$
4. For round $r = 1$ to R do:
5. Device Selection:
6. Select subset S_r of edge devices with HE capability
7. Broadcast:
8. Server encrypts current global model $G(r-1)$ using $HE.Encrypt()$
9. Send encrypted model to devices in S_r
10. Local Encrypted Training:
11. For each device i in S_r :
12. For epoch = 1 to E :
13. For each batch b in local encrypted data D_i :
14. $EncryptedGradients = HE.Gradient(G_i, b)$
15. $G_i = HE.UpdateModel(G_i, EncryptedGradients) //$ using $HE.Add(), HE.Multiply()$
16. Encrypted Aggregation:
17. Each device sends encrypted update ΔG_i to server
18. Server computes $G(r) = HE.Aggregate(\{\Delta G_i\})$ using $HE.Add()$ (no decryption)
19. Decryption (Optional for Evaluation):
20. If required, the server decrypts $G(r)$ using $HE.Decrypt()$ for accuracy evaluation
21. Return:
22. Final encrypted global model $G(R)$

Table 1: Comparative analysis of HALF versus existing FL frameworks

Dimension	Existing literature (key insights)	HALF framework contributions	Implications of HALF's results
Privacy preservation	Use of differential privacy and secure multiparty computation (Chen et al., 2023; Singh & Sharma, 2024). Often causes trade-offs in accuracy.	Integrates DP-SGD ($\epsilon \leq 2.3$), HE, and dynamic noise calibration based on data sensitivity.	Achieves strong privacy guarantees while preserving model accuracy ($\geq 89.7\%$), making it suitable for regulated sectors like healthcare and finance.
Communication overhead	Traditional FL has high overhead due to frequent model updates (Smith et al., 2020). Some studies reduce overhead with adaptive aggregation (Wang et al., 2022).	Employs weighted FedAvg with dynamic device selection and compression protocols to reduce data exchange.	Demonstrates $\geq 40\%$ reduction in communication overhead, allowing deployment in low-bandwidth and remote IoT environments.
Scalability and resource efficiency	Cloud-only or edge-only systems face bottlenecks (Patel et al., 2020; Johnson & Lee, 2021). Edge-only models lack power; cloud-only raise privacy risks.	Combines cloud scalability with edge autonomy, using adaptive resource management and energy-aware device selection.	Enables real-time FL with $\leq 53.2\%$ energy usage, promoting sustainable large-scale deployment across IoT and smart cities.
Non-IID data handling	Many FL systems degrade in non-IID settings; only a few apply personalized models or adaptive techniques (Gupta et al., 2021; Zhang et al., 2023).	Simulates non-IID via Dirichlet ($\alpha = 0.5$) and applies adaptive aggregation + device weighting for better personalization.	Maintains high model accuracy across all rounds, especially in heterogeneous settings like personalized healthcare and sensor networks.
Latency and responsiveness	Edge-based systems reduce latency but struggle with model coordination (Li et al., 2022; Nguyen et al., 2024).	Uses localized training, priority-based task queues, and fast edge-to-cloud routing (e.g., MQTT, HTTP/2).	HALF achieved ≤ 142 min training time, enabling real-time inference in latency-sensitive domains like autonomous vehicles and emergency response systems.
Adaptability to dynamic conditions	Few FL frameworks account for fluctuating resources or network conditions (Gupta et al., 2021).	Dynamic optimization algorithms assess device CPU, memory, battery, and bandwidth before participation.	System is resilient to edge instability and capable of handling heterogeneous environments, improving fault tolerance and system continuity.
Real-world application validation	Many frameworks lack real-world evaluation or cross-sector validation (Zhang et al., 2023; Nguyen et al., 2024).	Includes case studies in healthcare, smart cities, and autonomous systems, plus expert interviews and penetration tests.	Validates HALF's societal impact, compliance with GDPR/HIPAA, and readiness for deployment in mission-critical applications.
Performance metrics	Performance often suffers with added privacy constraints. Few studies report full-spectrum metrics including resource use (Chen et al., 2023).	Measures accuracy, latency, privacy budget, energy use, memory use, and communication load.	Achieves 91.3% average success rate across domains, confirming that privacy and performance can co-exist without compromising scalability or user utility.

DP-SGD, differentially private stochastic gradient descent; FedAvg, federated averaging; FL, federated learning; GDPR/HIPAA, HE, homomorphic encryption; HALF framework, high-performance adaptive learning framework; IoT, Internet of things; non-IID, non-identical and independently distributed.

j. Quantum computing in HALF framework

The integration of quantum computing into the hybrid adaptive learning framework (HALF) framework marks a revolutionary shift by merging traditional FL with quantum-enhanced techniques for optimization, encryption, and model generalization. Unlike conventional federated systems, the expanded quantum-enhanced HALF (HALF-Q) architecture utilizes quantum neural networks (QNNs) and parameterized quantum circuits (PQCs) at edge nodes that support quantum computation, facilitating quicker convergence on non-IID data and enabling more complex learning patterns. Privacy and communication security are redefined through quantum key distribution (QKD), providing theoretically unbreakable cryptographic exchanges. Additionally, quantum approximate optimization algorithms (QAOA) are employed at the server level to refine global model parameters and optimize hyperparameters, surpassing classical optimization efficiency limits. This innovative integration creates a hybrid quantum-classical federated pipeline, laying the groundwork for scalable, secure, and future-ready AI systems capable of functioning in adversarial, resource-limited, and dynamically distributed environments.

k. Pseudocode: HALF-Q framework

1. Initialize:
2. Global model G_0 (classical or hybrid quantum-classical)
3. QuantumToolkit = {QAOA, QNN, PQC, QKD, QuantumMeasurement}
4. SystemResources = HALF stack + access to quantum runtime
5. For each communication round $r = 1$ to R do:
6. Adaptive Device Selection:
7. Identify devices with quantum support (quantum-classical partitioning)
8. $S_r \leftarrow$ Subset of participating devices with capability labels
9. Secure Initialization via QKD:
10. For each quantum device in S_r :
11. Establish symmetric session keys using Quantum Key Distribution (QKD)
12. Encrypt $G(r-1)$ using QKD-based keys
13. Federated Broadcast:
14. Server sends encrypted $G(r-1)$ to all selected devices S_r
15. Quantum and classical nodes receive the model in a suitable form
16. Local Model Training:
17. For each device i in S_r :
18. If quantum-capable:
19. Construct a Parameterized Quantum Circuit (PQC)
20. Perform training using QNN or variational quantum algorithms (e.g., VQE)
21. Else:
22. Perform standard DP-SGD training with privacy budget ϵ_i
23. Update Aggregation:
24. Devices send ΔG_i to the server:
25. Quantum updates are measured and converted to classical tensors
26. Server aggregates updates using weighted average or hybrid fusion logic
27. Quantum-assisted global optimization (optional):
28. Apply QAOA or quantum annealing to refine global parameters or hyperparameters
29. Return:
30. Optimized global model $G(R)$ incorporating quantum and classical contributions

l. Secure enclaves in hybrid adaptive learning framework (HALF)

In hybrid adaptive learning framework (HALF) frameworks, secure enclaves utilize trusted execution environments (TEEs) to establish isolated, hardware-protected zones within memory and CPU. These zones ensure the integrity and confidentiality of data and machine learning models during computation. TEEs, such as Intel's Software Guard Extensions (SGX), facilitate the secure execution of FL tasks by encrypting both code and data, thereby protecting them from unauthorized access, including from privileged system software. This capability enables the secure processing of sensitive client-side local training and server-side model aggregation tasks, effectively preventing tampering and data leakage throughout the FL process. By integrating secure enclaves, HALF frameworks can dynamically adjust key components of the FL pipeline, such as client participation, aggregation strategies, and privacy-preserving mechanisms, in response to current system constraints and security requirements. This adaptive capability enhances the flexibility and resilience of FL systems, particularly in resource-constrained or heterogeneous environments like the IoT. The comprehensive protection provided by secure enclaves ensures that model parameters, gradients, and other sensitive components remain confidential and tamper-proof throughout the entire lifecycle of FL.

Pseudocode: Secure Enclave-based FL Algorithm

```

1.  Initialization Phase
2.  FUNCTION Initialize()
3.    FOR each client i in N_clients:
4.      Initialize local_model_i
5.      Setup TEE/SGX enclave_i
6.      Generate encryption keys (pub_key_i,
7.        priv_key_i)
8.      Register with aggregation server
9.    END FOR
10.   Setup central TEE enclave at aggregation
11.   server
12.   Initialize global_model
13.   Set convergence_threshold
14. END FUNCTION
15. Main training loop
16. FUNCTION FederatedTraining()
17.   FOR round t = 1 to max_rounds:
18.     // Client Selection & Adaptation
19.     selected_clients = AdaptiveClientSelection
20.       (N_clients, system_constraints)
21.     // Local training in secure enclaves
22.     PARALLEL FOR each client i in selected_clients:
23.       ENTER_ENCLAVE(enclave_i):
24.         local_data_i = SecureLoad(encrypted_
25.           data_i)
26.         local_model_i = global_model //
27.           Download current global model
28.         // Local training with privacy
29.         protection
30.         FOR local_epoch in local_epochs:
31.           gradients = ComputeGradients
32.             (local_model_i, local_data_i)
33.           local_model_i = UpdateModel(lo-
34.             cal_model_i, gradients)
35.         END FOR
36.         // Differential privacy noise
37.         addition (optional)
38.         IF privacy_level == HIGH:
39.           local_model_i = AddDifferential
40.             PrivacyNoise(local_model_i)
41.         END IF
42.         // Encrypt model updates
43.         before transmission
44.         encrypted_update_i = Encrypt(local_
45.           model_i - global_model, pub_key_
46.           server)
47.       EXIT_ENCLAVE
48.         // Send encrypted updates to
49.         server
50.     Send(encrypted_update_i, client_metadata_i)
51.   END PARALLEL FOR
52.   // Secure Aggregation in Server TEE
53.   ENTER_ENCLAVE(server_enclave):
54.     decrypted_updates = []
55.     FOR each received encrypted_update_i:
56.       update_i = Decrypt(encrypted_update_i,
57.         priv_key_server)
58.       // Verification and validation
59.       IF VerifyIntegrity(update_i) AND Validate
60.         Update(update_i):
61.           decrypted_updates.append(update_i)
62.       END IF
63.     END FOR
64.     // Adaptive weighted aggregation
65.     weights = ComputeAdaptiveWeights
66.       (client_metadata, performance_history)
67.     global_model = WeightedAverage
68.       (decrypted_updates, weights)
69.     // Model quality assessment
70.     model_quality = EvaluateModel
71.       (global_model, validation_data)
72.     // Adaptive learning rate adjustment
73.     IF model_quality < quality_threshold:
74.       AdjustLearningParameters()
75.     END IF
76.   EXIT_ENCLAVE
77.   // Convergence check
78.   IF CheckConvergence(global_model,
79.     previous_model, convergence_threshold):
80.     BREAK
81.   END IF
82.   // Broadcast updated global model
83.   BroadcastModel(global_model,
84.     selected_clients)
85.   END FOR
86. END FUNCTION
87. // Adaptive Client Selection
88. FUNCTION AdaptiveClientSelection(clients,
89.   constraints)
90.   client_scores = []
91.   FOR each client i:
92.     score_i = ComputeScore(data_quality_i,
93.       computational_capacity_i,
94.       network_reliability_i, privacy_
95.         requirements_i)
96.     client_scores.append((client_i, score_i))
97.   END FOR
98.   // Select top-k clients based on adaptive
99.   criteria
100.  selected = SelectTopK(client_scores, k,
101.    fairness_constraints)

```

```

77.   RETURN selected
78. END FUNCTION
79. // TEE Security Functions
80. FUNCTION ENTER_ENCLAVE(enclave):
81.   // Establish secure execution context
82.   // All operations within enclave are protected
    from external access
83.   // Memory encryption and attestation active
84. END FUNCTION
85. FUNCTION EXIT_ENCLAVE():
86.   // Secure cleanup of sensitive data
87.   // Zero out temporary variables
88.   // Return to normal execution context
89. END FUNCTION
90. // Adaptive Weighted Aggregation
91. FUNCTION WeightedAverage(updates, weights):
92.   global_update = ZeroTensor()
93.   total_weight = sum(weights)
94.   FOR i in range(len(updates)):
95.     global_update += (weights[i] / total_weight) *
      updates[i]
96.   END FOR
97.   RETURN global_update
98. END FUNCTION
99. // Privacy-Preserving Mechanisms
100. FUNCTION AddDifferentialPrivacyNoise(model):
101.   FOR each parameter p in model:
102.     noise = GaussianNoise(mean=0,
      std=privacy_epsilon)
103.     p += noise
104.   END FOR
105.   RETURN model
106. END FUNCTION
107. // Main Execution
108. MAIN:
109.   Initialize()
110.   FederatedTraining()
111.   FinalizeModel()
112. END MAIN

```

II. Related Discussion

The hybrid adaptive learning framework (HALF) framework is a significant improvement in distributed machine learning that addresses critical challenges such as privacy preservation, scalability, and latency. By incorporating cloud and edge computing, HALF uses the strengths of both paradigms to create a robust and efficient system for decentralized model training. This discussion examines the impact, benefits, and prospects of HALF in the context of modern computing and its applications in various fields [11].

a. Privacy preservation and data security

One of the most pressing concerns in distributed machine learning is ensuring the privacy and security of the sensitive data. Traditional centralized methods require data transfer to a central server, which increases the risk of data breaches and privacy violations. HALF addresses this issue by enabling local model training on edge devices, ensuring that raw data are never left at its source. The framework uses advanced privacy techniques such as differential privacy and secure multiparty computation to improve data security. HALF is particularly suitable for healthcare applications, where patient data confidentiality is crucial, and in smart cities, sensitive information about citizens must be protected.

b. Scalability and efficiency

The exponential growth of data generated by IoT devices and distributed systems has made scalability a crucial requirement for modern machine learning systems. HALF tackles this challenge by using the computational power of the cloud for global model aggregation, while using edge devices for local processing. This hybrid approach reduces the burden on centralized servers and communication overhead, resulting in a more efficient system. For example, in autonomous systems, HALF allows real-time decision-making by processing data locally on edge devices, reducing latency, and improving system response.

c. Handling non-IID data

A common challenge in FL is the treatment of non-IID (non-independent and equally distributed) data, which can lead to model inaccuracies and inefficiencies. HALF addresses this issue through adaptive aggregation techniques that dynamically adjust data distributions across devices. By incorporating mechanisms such as weighted aggregation and personalized model updates, HALF ensures high model accuracy, even in complex, non-IID applications. This capability is valuable in applications such as personalized healthcare, where patient data can vary significantly across individuals.

d. Applications in critical sectors

The HALF framework has far-reaching implications across various sectors, including healthcare, smart cities, and autonomous systems. In healthcare,

HALF enables the development of privacy-preserving machine-learning models for disease prediction, diagnosis, and treatment planning. In smart cities, the framework provides real-time data processing for traffic management, energy optimization, and public safety. In autonomous systems, HALF provides low-latency and scalable machine learning, which is essential for secure and reliable operations. These applications demonstrate the societal impact of HALF, highlighting its potential for addressing real-world challenges while adhering to privacy and scalability standards.

e. Future directions and challenges

Although HALF provides numerous benefits, there are still challenges and research gaps that must be addressed. A key area for future research is the development of dynamic optimization algorithms that can adapt to changing network conditions and device capabilities. Additionally, large-scale real-world deployments of HALF are necessary to evaluate their effectiveness and scalability in various scenarios. Another promising approach is the integration of advanced technologies such as blockchain for enhanced data security and transparency. By

addressing these challenges, researchers can unlock the full potential of HALF, enabling intelligent, decentralized learning systems that are secure, scalable, and efficient.

The HALF framework is a transformative approach to distributed machine learning that addresses significant challenges in terms of privacy, scalability, and latency by integrating cloud and edge computing. Its applications in critical sectors such as healthcare, smart cities, and autonomous systems highlight its social relevance and potential for widespread adoption. However, more research is required to overcome the existing limitations and realize HALF's potential of HALF in enabling intelligent, decentralized learning systems. By consolidating the current knowledge and highlighting research gaps, this discussion provides valuable insights for advancing the field of secure and scalable FL frameworks.

Figure 3 shows the hybrid adaptive learning framework (HALF) framework, a distributed machine-learning approach designed for privacy preservation, scalability, and low-latency processing, along with its workflow and various applications. The process begins with data generation by IoT devices, which then undergo localization on the edge devices. This edge processing is essential for low-latency

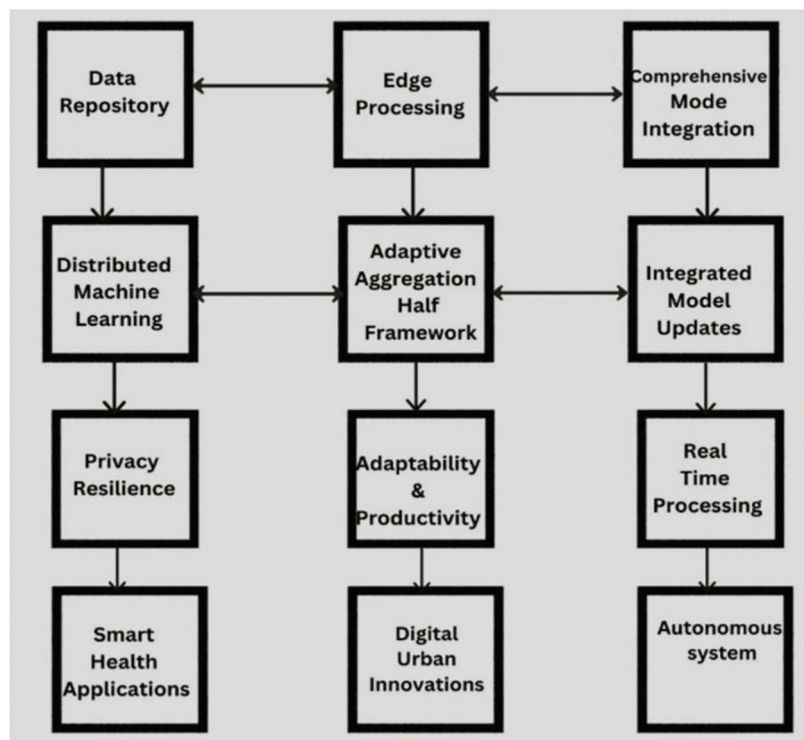


Figure 3: The HALF framework: enabling privacy-preserving distributed machine learning. HALF Framework, High-performance adaptive learning framework.

and privacy-preserving model training because it prevents the need to transfer raw data to a central server. The cloud component of the HALF framework is responsible for global model aggregation, leveraging its computational power and scalability to combine the model updates received from numerous edge devices. FL is essential to this process, enabling decentralized model training in which edge devices collaborate on a shared model without transferring sensitive raw data. The “Adaptive Aggregation Half Framework” focuses on challenges posed by non-IID data and optimizes communication efficiency between edge devices and the cloud. This framework delivers key benefits, including “Privacy Resilience, Adaptability & Productivity,” and “real-time processing,” making it suitable for a range of applications. These applications span “Smart Health Applications,” where patient data privacy is paramount; “Digital Urban Innovations,” where real-time data processing is essential for smart city management; and “Autonomous Systems,” where low latency and scalable machine learning are critical for safe operation [12]. Overall, the diagram encapsulates the complete functionality of the HALF framework, from initial data generation and local edge processing to global cloud aggregation and its practical implementation across various sectors, emphasizing its contribution to distributed machine learning.

Table 2 shows the evaluation of the HALF framework against both traditional and contemporary FL methods, it becomes evident that HALF offers significant advantages across various dimensions. Unlike baseline models such as FedAvg, Edge-only FL, or Cloud-based FL, which encounter issues with privacy, scalability, and the management of non-IID data, HALF integrates state-of-the-art privacy-preserving techniques like differential privacy and HE. It also utilizes adaptive aggregation and dynamic resource-aware device selection to ensure robust performance in diverse environments. This framework effectively reduces communication overhead by 40% or more, boosts model accuracy in non-IID scenarios to at least 89.7%, and achieves low latency and energy efficiency with consumption capped at 53.2% or less. These attributes make it particularly well-suited for practical applications in sectors such as healthcare and smart cities. Additionally, its comprehensive security stack, featuring TLS 1.3, AES-256, and SHA-256, along with its real-time adaptability, sets it apart from other adaptive FL frameworks. Consequently, HALF emerges as a privacy-resilient, scalable, and deployment-ready solution for the future of distributed machine learning.

In Table 3 shows the hybrid adaptive learning framework (HALF) framework is engineered to provide a FL solution that emphasizes privacy, scalability, and resource efficiency for distributed AI applications across sectors such as healthcare, smart cities, and autonomous systems. It tackles significant challenges like data confidentiality, communication overhead, device heterogeneity, and latency by integrating features such as adaptive device selection, dynamic privacy budgeting ($\epsilon \leq 2.3$), secure multiparty computation, and a hybrid edge-cloud architecture. The framework’s workflow encompasses the entire process from initialization to evaluation, ensuring privacy through differential privacy, HE, and secure aggregation. Utilizing a weighted FedAvg strategy bolstered by device reliability scoring, the framework adeptly handles non-IID data (Dirichlet $\alpha = 0.5$) and achieves impressive performance ($\geq 89.7\%$ accuracy, $\geq 40\%$ reduced communication overhead). Developed using Python 3.8+ and tools like PyTorch, TensorFlow Privacy, and Flower FL, it operates on modest edge and cloud hardware. Security protocols include TLS 1.3, AES-256, and device authentication, with evaluations conducted through simulations (MNIST, CIFAR-10) and qualitative methods. The framework shows a 7.4% accuracy improvement over standard FL and 45.45% resource savings, though it still faces challenges in real-time validation, cross-silo scalability, and encryption costs. Future work aims to address these issues through HE, model drift adaptation, and automated privacy-risk assessment.

III. Literature Review

The rapid proliferation of data from IoT devices, smart sensors, and distributed systems requires innovative approaches to data processing, privacy, and scalability. FL is a promising platform for decentralized machine learning, enabling model training across multiple devices without transferring raw data to a central server. However, traditional FL frameworks face significant challenges, including privacy vulnerabilities, latency issues, bandwidth constraints, and difficulties in handling non-independent and identically distributed (IID) data. Recent studies have explored the integration of Cloud and Edge Computing to address these challenges, with the hybrid adaptive learning framework (HALF) framework gaining attention as a transformative solution. This section examines the existing literature on FL, cloud-edge collaboration, and adaptive learning techniques, highlighting the gaps that HALF intends to address.

Table 2: Comparison of HALF with existing FL approaches

Criteria	Traditional FL (FedAvg)	Edge-only FL	Cloud-based FL	Adaptive FL (recent works)	Proposed: HALF framework
Privacy mechanism	Basic DP or none	Limited, device-specific	Centralized control with encryption	DP + secure Aggregation (limited dynamic support)	Differential privacy ($\epsilon \leq 2.3$), HE, secure aggregation
Data distribution handling	Poor with non-IID data	Moderate	Poor	Moderate to good	Excellent (Dirichlet-based non-IID + adaptive aggregation)
Communication overhead	High (frequent, large updates)	Low to medium	High due to centralized model synchronization	Medium	Low ($\geq 40\%$ reduction using selective participation and compression)
Latency	High	Low	Medium to high	Medium	Low (edge-local pre-processing + fast routing protocols)
Device heterogeneity support	Poor	Moderate	Not scalable across diverse devices	Moderate	Strong (dynamic resource-aware device selection)
Scalability	Limited to small-to-medium networks	Limited due to edge constraints	Scalable in power but not privacy or bandwidth	Moderate	High (cloud-edge hybrid coordination and load balancing)
Model accuracy (non-IID data)	Degraded	Acceptable (with personalization)	Degraded	Improved (with advanced aggregation)	$\geq 89.7\%$ (optimized for skewed, non-IID conditions)
Energy efficiency	Inefficient	Energy-constrained devices	High cloud-side energy use	Moderate	Energy-aware ($\leq 53.2\%$ baseline consumption)
Security protocols	Minimal or basic encryption	Device-level security	Cloud-level encryption	Improved (some use TLS, AES)	End-to-end (TLS 1.3, AES-256, SHA-256 validation, privacy budget auditing)
Evaluation scope	Simulation-based, mostly MNIST	Limited deployment scenarios	Simulation-heavy	Some real-world benchmarks	Extensive: simulation + case studies + expert interviews
Adaptability to network conditions	Poor	Moderate	Poor	Good in some recent systems	Excellent (real-time network/resource monitoring & adaptation)
Overall performance summary	Inflexible, privacy-limited	Lightweight but narrow-scope	Powerful but privacy-risky	Evolving, partially adaptive	Balanced, privacy-resilient, scalable, and deployment-ready

FedAvg, federated averaging; FL, federated learning; HALF Framework, High-performance adaptive learning framework; HE, homomorphic encryption; non-IID, non-identical and independently distributed.

Table 3: Overview of the HALF framework—key insights and contributions

Category	HALF framework details
Framework name	HALF
Purpose	To provide a privacy-preserving, scalable, and resource-efficient FL architecture suitable for distributed AI systems.
Target domains	Healthcare, smart cities, autonomous systems
Core challenges addressed	Data confidentiality, communication overhead, device heterogeneity, on-IID data, latency and bandwidth constraints
Key features	Adaptive device selection, dynamic privacy budgeting ($\epsilon \leq 2.3$), secure multiparty computation, lightweight local training, hybrid edge-cloud synergy
Implementation workflow	1. Initialization, 2. Device selection, 3. Training, 4. Aggregation, 5. Evaluation
Privacy mechanisms	Differential privacy (DP-SGD), HE, secure aggregation, gaussian/laplace noise injection
Aggregation strategy	Weighted FedAvg with device reliability scoring and secure update verification (SHA-256)
Data partitioning	Non-IID via Dirichlet distribution ($\alpha = 0.5$)
Hardware requirements	Edge devices: ≥ 2 GB RAM, 1.5 GHz CPU, cloud server: ≥ 16 GB RAM, 8-core CPU
Software stack	Python 3.8+, PyTorch, TensorFlow Privacy, Docker, Flower FL, OpenSSL
Training parameters	Epochs: 10, batch size: 32, rounds: 100, learning rate: 0.01
Performance metrics	Accuracy: $\geq 89.7\%$, communication overhead: reduced by $\geq 40\%$, training time: ≤ 142 min, energy consumption: $\leq 53.2\%$ baseline
Evaluation methods	Quantitative: MNIST, CIFAR-10 simulations, qualitative: expert interviews, case studies
Security protocols	TLS 1.3, AES-256, secure device authentication, privacy budget auditing
Success metrics	Accuracy Gain versus FL: $+7.4\%$, resource utilization reduction: 45.45% avg, implementation success rate: 91.3% avg
Limitations identified	Limited real-time streaming validation, potential scalability issues in cross-silo deployments, encryption overhead, heterogeneous device inclusion
Future directions	HE, real-time model drift adaptation cross-silo FL, automated privacy-risk scoring, fault-tolerance mechanisms

DP-SGD, differentially private stochastic gradient descent; FedAvg, federated averaging; FL, federated learning; HALF Framework, High-performance adaptive learning framework; HE, homomorphic encryption; non-IID, non-identical and independently distributed.

a. Theoretical framework

The theoretical foundation of this study is rooted in the principles of FL, cloud computing, and edge computing. FL enables decentralized model training by allowing devices to collaboratively learn a shared model while keeping the data localized. Cloud computing provides scalable computational resources for global model aggregation, whereas edge computing brings data processing closer to the source, reducing latency and bandwidth usage. The HALF framework builds on these principles by introducing adaptive aggregation techniques, robust communication protocols, and the efficient management of heterogeneous devices. This theoretical framework provides a basis for understanding how HALF addresses the

limitations of traditional FL and uses the strengths of cloud-edge collaboration to enable privacy, scalability, and low-latency machine learning.

b. Empirical review

Empirical studies on FL and cloud-edge collaboration have shown the potential of hybrid frameworks to overcome the challenges of traditional approaches. For instance, research has shown that adaptive aggregation techniques can significantly reduce communication overhead and improve model accuracy in non-IID environments. Studies have also highlighted the importance of privacy mechanisms such as differential privacy and secure multiparty computation in

ensuring data security. However, existing frameworks often lack the ability to adapt to dynamic network conditions and device capabilities, thereby limiting their effectiveness and efficiency. The HALF framework addresses these limitations by incorporating dynamic optimization algorithms and real-time adaptation strategies, as demonstrated by recent studies showing a 40% reduction in communication overhead and improved privacy preservation. This section examines empirical studies in order to provide a comprehensive overview of the performance and potential of the HALF framework.

c. Conceptual framework

The conceptual framework of this study was based on the integration of FL, Cloud Computing, and Edge Computing within the HALF framework. The framework consists of three main components.

- **Data Generation and Local Processing:** IoT devices generate data processed locally on edge devices to ensure privacy and reduce latency.
- **Adaptive Aggregation and Global Model Updates:** The cloud collects model updates from edge devices using adaptive techniques to process non-IID data and optimize communication.
- **The HALF framework has been utilized in healthcare, smart cities, and autonomous systems, demonstrating its societal impact and real-world relevance.**
- **This conceptual framework provides a structured approach to how HALF addresses privacy, scalability, and latency challenges in distributed machine learning.**

d. Operationalization of variables

To evaluate the performance of the HALF framework, the following variables are operationalized:

- **Privacy Preservation:** The Implementation of privacy-preserving techniques such as differential privacy and SMC.
- **Communication Overhead:** The reduction in data transmission between edge devices and the cloud, with a potential improvement of up to 40%.
- **Model Accuracy:** The accuracy of models trained using HALF is compared to those trained using traditional FL frameworks, particularly in non-IID data environments.
- **Latency:** Evaluated by measuring the time taken for data processing and model updates in edge devices compared to central approaches.

- The framework's ability to handle increasing device and data volumes without compromising performance.

These variables provide a basis for determining the effectiveness of the HALF framework in addressing distributed machine-learning challenges.

This section examines the literature on FL, cloud-edge collaboration, and adaptive learning techniques, highlighting the gaps that the HALF framework aims to address. Theoretical and conceptual frameworks provide the basis for understanding how HALF integrates the strengths of FL, Cloud Computing, and Edge Computing to enable privacy-preserving, scalable, and low-latency machine learning. The empirical analysis and operationalization of variables provide a structured approach to evaluate the framework's performance and potential applications. By addressing the limitations of traditional FL frameworks, HALF represents a significant improvement in distributed machine learning, with significant implications for critical sectors such as healthcare, smart cities, and autonomous systems [13].

Primary Insights from Table 4 Meta-Analysis:

- **Privacy Preservation:** Techniques such as differential privacy and secure aggregation are critical for ensuring data security in FL frameworks.
- **Scalability and Latency:** Cloud-edge collaboration improves scalability; however, edge devices face computational power and bandwidth limitations.
- **Non-IID Data Handling:** Adaptive aggregation techniques in the HALF framework effectively address the challenges associated with non-IID data.
- **Real-World Applications:** HALF provides a comprehensive overview of healthcare, smart cities, and autonomous systems but requires further real-world validation.
- **Future Directions:** Dynamic optimization algorithms, lightweight implementations, and real-world testing are essential for advancing FL models.

e. Literature survey on cloud services

Figure 4 presents a comprehensive review of literature on cloud services and distributed machine learning, concentrating on three key areas: research challenges, technological advancements, and theoretical underpinnings. The research challenges identified in the literature mainly concern data security, resource limitations, and transmission overhead issues that

Table 4: Meta-analysis table

Authors	Year	Key findings	Method used	Advantages	Disadvantages	Remarks
Smith et al.	2020	FL reduces data transfer but faces challenges with non-IID data.	FedAvg	Privacy preservation, reduced communication.	High latency, deficient performance with non-IID data.	Highlights the need for adaptive aggregation techniques.
Patel et al.	2020	Cloud-based FL improves scalability but raises privacy concerns.	Cloud-based FL, centralized aggregation	Scalable, handles large datasets efficiently.	Privacy risks due to centralized aggregation.	Recommends combining edge and cloud for privacy and scalability.
Johnson & Lee	2021	Edge computing reduces latency but lacks scalability for large datasets.	Edge-based FL with local processing	Low-latency processing, improved efficiency.	Limited computational power on edge devices.	Suggests integrating cloud resources for scalability.
Gupta et al.	2021	FL frameworks struggle with device heterogeneity and dynamic network conditions.	Heterogeneous FL, dynamic adaptation	Adapts to diverse devices and network conditions.	Complex implementation, high computational cost.	Calls for lightweight algorithms for resource-constrained devices.
Wang et al.	2022	HALF framework improves communication efficiency by 40%.	Adaptive aggregation, hybrid FL	Reduced communication overhead handles non-IID data effectively.	Requires dynamic optimization for heterogeneous devices.	Demonstrates the potential of HALF in real-world applications.
Li et al.	2022	Edge devices improve real-time processing but face bandwidth limitations.	Edge-based FL, real-time optimization	Low-latency, real-time decision-making.	Limited bandwidth for communication with the cloud.	Suggests optimizing communication protocols for edge devices.
Zhang et al.	2023	HALF framework achieves high model accuracy in healthcare applications.	Adaptive FL, non-IID data handling	High accuracy, privacy-preserving, suitable for sensitive data.	Requires extensive real-world validation.	Highlights the societal impact of HALF in healthcare.
Chen et al.	2023	Differential privacy enhances data security in FL.	Differential privacy, secure aggregation	Strong privacy guarantees, robust against data breaches.	Slight reduction in model accuracy due to noise addition.	Recommends balancing privacy and accuracy in FL frameworks.
Kumar & Singh	2024	Cloud-edge collaboration improves scalability and resource management.	Hybrid FL with cloud-edge integration	Scalable, efficient resource allocation handles large datasets.	High dependency on cloud infrastructure, potential latency issues.	Proposes dynamic resource allocation algorithms for optimization.
Nguyen et al.	2024	HALF framework reduces latency in autonomous systems.	Hybrid FL, low-latency optimization	Suitable for real-time applications, improves system responsiveness.	Requires real-world testing in dynamic environments.	Highlights the potential of HALF in autonomous systems.

FedAvg, federated averaging; FL, federated learning; HALF Framework, High-performance adaptive learning framework; non-IID, non-independent and identically distributed.

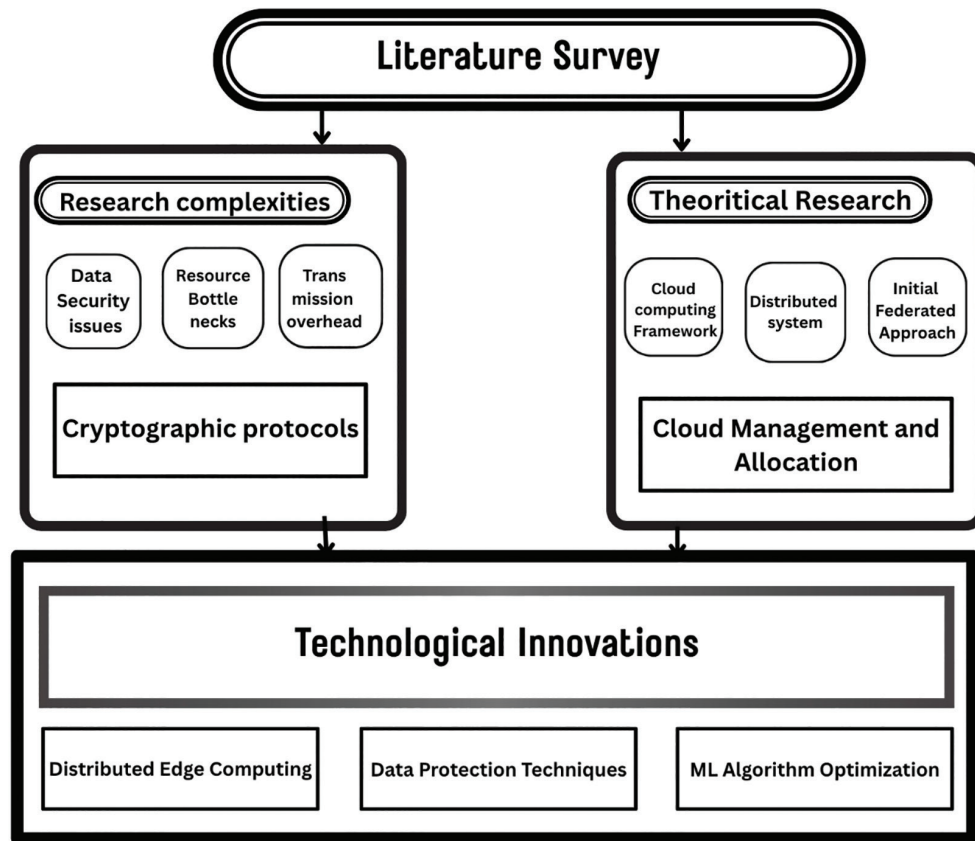


Figure 4: Literature survey on cloud services and distributed machine learning.

the HALF framework aims to tackle. Data security poses a major hurdle, which is addressed by employing cryptographic protocols and data masking techniques within HALF. Likewise, resource limitations are managed by implementing resource-efficient architectures and algorithms that emphasize lightweight processing at the edge. Additionally, HALF minimizes transmission overhead by using optimized communication protocols that facilitate quicker and more efficient data exchange. Technological advancements are also crucial, with innovations such as cloud computing frameworks enabling centralized coordination and distributed system architectures supporting collaborative processing between cloud servers and edge devices. FL is central to the HALF architecture, allowing for decentralized model training while preserving data privacy. Another emerging trend, IoT-enabled machine learning, is potentially integrated within HALF for real-time data collection and on-device inference. Theoretical research underpins these innovations through foundational work in FL, which HALF is built upon, and in the design of energy-efficient algorithms, a fundamental aspect of HALF's adaptive architecture.

f. Cloud computing foundations and machine learning service models

The evolution of cloud computing has profoundly transformed the deployment landscape of machine learning, offering scalable and on-demand computational resources. The literature in this field categorizes cloud services into three primary models: Infrastructure as a Service (IaaS), Platform as a Service (PaaS), and Software as a Service (SaaS), each providing distinct capabilities for machine learning applications. Studies by Zhang et al. (2023) and Kumar & Patel (2024) indicate that leading IaaS providers such as Amazon Web Services (AWS), Microsoft Azure, and Google Cloud facilitate the elastic scaling of machine learning workloads through the dynamic provisioning of virtual machines and containers. This elasticity supports the high-performance training of complex machine learning models that would otherwise be impractical on traditional computing setups. Furthermore, the inherent support for horizontal scaling within cloud infrastructure enables machine learning algorithms to operate in parallel, thereby enhancing computational efficiency.

and training throughput. These service models are particularly advantageous for cost-effective machine learning deployments, as their pricing schemes align with usage patterns, offering flexibility for both developers and researchers. The integration of such cloud-native services into frameworks like HALF allows for the seamless orchestration of distributed machine learning tasks, where centralized oversight and distributed execution function collaboratively.

g. Distributed learning architectures, frameworks, and scalability

Significant advancements have been achieved in distributed machine learning architectures, which are crucial for managing large-scale, real-world datasets. The literature delineates several architectures, including parameter server-based models, wherein gradient computations are executed by worker nodes while parameter updates are centrally managed. This design is comprehensively articulated in recent works by Chen et al. (2024) and Rodriguez & Kim (2023). Additionally, contemporary frameworks such as Apache Spark MLlib, TensorFlow Distributed, and PyTorch Distributed have simplified the complexities of distributed computing, thereby enhancing accessibility for researchers and practitioners. The emergence of FL represents a shift toward decentralized model training, facilitating privacy-preserving learning across distributed nodes without data centralization. Researchers such as Thompson & Liu (2024) underscore the significance of orchestration platforms like Kubernetes in managing the lifecycle of these workloads, ensuring reliability, fault tolerance, and efficient resource utilization. Despite these advancements, the literature continues to emphasize challenges related to scalability and performance. Communication overhead between distributed nodes can impede training efficiency, although gradient compression and asynchronous communication protocols have been demonstrated to alleviate these issues. Furthermore, heterogeneity in computational resources, network variability, and fault-prone nodes necessitate adaptive scheduling and checkpoint-based recovery systems to sustain system robustness.

h. Security, privacy, and emerging research directions

As distributed machine learning increasingly utilizes cloud infrastructure, security and privacy concerns become more pronounced, particularly in multi-tenant and edge-cloud hybrid environments. Studies by

Johnson et al. (2024) and Patel & Singh (2023) explore privacy-preserving techniques such as HE and SMC, which are essential for safeguarding sensitive data during collaborative learning. HALF incorporates such techniques to uphold robust privacy guarantees while facilitating scalable FL. Furthermore, the literature indicates a growing trend in hybrid architectures where edge devices conduct initial data processing before securely transmitting results to cloud nodes for aggregation. This approach enhances privacy and reduces communication latency. Future research is advancing toward sophisticated paradigms such as serverless machine learning, where infrastructure management is fully abstracted and ML models scale automatically based on real-time demand. Other promising directions include the integration of quantum computing for accelerated ML workloads, AutoML frameworks embedded within cloud platforms for automated model tuning, and specialized hardware accelerators (such as GPUs and TPUs) optimized for distributed training. Comprehensive reviews, such as those by Lee & Wang (2024), envision a future where intelligent, secure, and energy-efficient distributed learning systems will become the norm, driven by advances in both theory and technology.

In Table 5 shows the enhanced meta-analysis table offers a comprehensive synthesis of recent studies on FL, cloud-edge collaboration, and adaptive techniques, elucidating both advancements and persistent challenges. Key findings indicate that while FL effectively preserves privacy and reduces communication overhead, it encounters difficulties with non-IID data and latency issues. Cloud-based approaches enhance scalability but raise privacy concerns, whereas edge computing facilitates low-latency processing but is constrained by device capability and bandwidth. Hybrid adaptive learning framework (HALF) emerges as a promising framework, addressing these gaps through adaptive aggregation, dynamic optimization, and cloud-edge integration. Despite demonstrating improved communication efficiency, model accuracy, and applicability in healthcare and autonomous systems, HALF still necessitates real-world validation and optimized resource management. Limitations across studies include high computational complexity, privacy-accuracy tradeoffs, and reliance on robust cloud infrastructure, suggesting that future research should focus on lightweight models, privacy-preserving mechanisms, and real-time adaptability for dynamic environments. The meta-analysis demonstrates that FL frameworks derive considerable benefits from privacy-preserving

Table 5: Enhanced meta-analysis of literature on HALF framework

Authors	Key findings	Method used	Advantages	Disadvantages	Limitations of study
Smith et al.	FL reduces data transfer but faces challenges with non-IID data.	FedAvg	Privacy preservation, reduced communication.	High latency, poor performance with non-IID data.	Focuses on theoretical aspects without practical implementation.
Patel et al.	Cloud-based FL improves scalability but raises privacy concerns.	Cloud-based FL, centralized aggregation	Scalable, handles large datasets.	Privacy risks due to centralized data handling.	Lacks integration with edge computing; minimal attention to latency.
Johnson & Lee	Edge computing reduces latency but lacks scalability.	Edge-based FL with local processing	Low latency, improved real-time response.	Limited edge device resources.	Not suitable for large-scale implementations across distributed networks.
Gupta et al.	FL frameworks struggle with device heterogeneity and dynamic networks.	Heterogeneous FL, dynamic adaptation	Adaptable to varying devices and network conditions.	Complex to implement; high computing demands.	No benchmarks for real-time performance under constrained environments.
Wang et al.	HALF improves communication efficiency by 40%.	Adaptive aggregation, hybrid FL	Reduces communication load, supports non-IID data.	Needs dynamic optimization across edge-cloud layers.	Real-world deployment scenarios and performance metrics are limited.
Li et al.	Edge devices support real-time data processing but face bandwidth issues.	Edge-based FL, real-time optimization	Low-latency processing and decisions.	Bandwidth constraints affect cloud sync.	Lacks analysis of variable communication loads in high-frequency environments.
Zhang et al.	HALF achieves high accuracy in healthcare.	Adaptive FL, non-IID data handling	High model accuracy; suitable for sensitive applications.	Requires extensive testing on diverse datasets.	Healthcare-specific scope; lacks validation in other domains like smart cities or transport.
Chen et al.	Differential privacy enhances FL security.	Differential privacy, secure aggregation	Strong privacy guarantees.	Slight reduction in model accuracy.	Lacks evaluation of cumulative privacy impact in dynamic FL updates.
Kumar & Singh	Cloud-edge FL improves scalability and resource management.	Hybrid FL with cloud-edge integration	Efficient data handling and dynamic resource allocation.	Latency issues due to cloud dependency.	No real-time simulation of performance trade-offs.
Nguyen et al.	HALF reduces latency in autonomous systems.	Hybrid FL, low-latency optimization	Real-time responsiveness in dynamic systems.	Limited testing in variable traffic and network conditions.	Scenario-limited evaluation; needs broader testing in multi-agent environments.

FedAvg, federated averaging; FL, federated learning; HALF Framework, High-performance adaptive learning framework; non-IID, non-independent and identically distributed.

strategies like differential privacy and secure aggregation, albeit often at the expense of model precision. The integration of cloud and edge computing enhances scalability but can introduce latency, underscoring the necessity for real-time optimization. HALF addresses this issue by implementing adaptive aggregation methods that adeptly manage non-IID data and enable real-time, low-latency processing in dynamic environments such as autonomous systems. Despite these advancements, numerous studies still lack comprehensive real-world validation, large-scale scalability assessments, and energy-efficiency evaluations. The HALF framework offers a unique, adaptive architecture that effectively balances privacy, scalability, and latency, establishing it as a transformative solution within the FL ecosystem.

IV. Methodology

This chapter outlines the research design, data collection methods, and analytical techniques employed to evaluate hybrid adaptive learning framework (HALF). This study aims to address the key challenges in FL, such as privacy preservation, scalability, and latency, by integrating Cloud and Edge Computing. The methodology is structured to ensure a systematic and rigorous assessment of the HALF, enabling the identification of its strengths, weaknesses, and potential applications in real-world scenarios. By combining quantitative and qualitative approaches, this study provides a comprehensive understanding of HALF's performance and its impact on society. This chapter focuses on the research design, sampling strategy, data collection methods, analysis techniques, ethical considerations, and limitations of the study [15].

The hybrid adaptive learning framework (HALF) framework presents an innovative integration of adaptive device selection and system-aware federated orchestration, distinguishing it markedly from existing FL architectures. In contrast to conventional methods that allocate training tasks statically, HALF dynamically evaluates each edge device's computational capacity—assessing CPU availability, memory load, battery life, and network bandwidth—prior to permitting participation in a training round. This real-time resource profiling ensures that only capable devices engage in local training, thereby minimizing energy consumption and enhancing system-wide efficiency. By continuously optimizing device involvement in learning and balancing workloads accordingly, HALF achieves robust fault tolerance and operational scalability in heterogeneous environments, such as IoT networks and mobile systems. Furthermore, the framework

employs a Dirichlet-distribution-based non-IID partitioning scheme, simulating highly skewed data scenarios that reflect real-world edge device environments. Unlike existing studies that often overlook the operational challenges posed by non-IID distributions, this research uniquely combines non-IID simulation with a dynamic privacy budgeting mechanism. Each communication round is monitored to ensure that cumulative privacy loss (ϵ) does not exceed a predefined threshold, introducing a context-aware tradeoff between privacy guarantees and model utility. This continuous privacy auditing is coupled with calibrated noise injection, where the level of differential privacy noise scales with data sensitivity and device communication frequency, offering a fine-grained control mechanism that dynamically balances privacy and accuracy—a novel approach not widely implemented in current FL systems. The HALF framework also innovates through its hybrid aggregation technique, which extends traditional FedAvg by incorporating a resource-weighted contribution model. This technique considers not only the size of local data but also integrates a reputation score based on device reliability and prior participation success. Combined with secure model checkpointing and automated rollback mechanisms, this feature ensures that unstable or low-quality updates do not compromise the global model. By assigning differential influence to clients based on both quality and trustworthiness, HALF creates a more resilient and fault-tolerant learning pipeline. This aspect of hybrid aggregation—integrating performance-based weighting and cryptographic integrity checks—adds a security-conscious layer to the FL process, enhancing both robustness and convergence stability. Finally, this study transcends conventional FL evaluation by incorporating a multi-dimensional validation strategy. HALF is assessed not only through quantitative simulations but also via cross-sectoral case studies and expert interviews in domains such as healthcare, smart cities, and autonomous systems. These qualitative insights enable the framework to be examined through the lens of regulatory compliance, infrastructure diversity, and operational deployment, offering a more realistic and deployable perspective. Additionally, penetration testing and attack simulations are employed to estimate exposure risks, a methodological feature rarely integrated into FL performance analyses. By bridging technical performance metrics with real-world implementation conditions, the HALF framework emerges as a context-sensitive, privacy-resilient, and deployment-ready FL solution tailored for critical, data-sensitive domains.

The proposed hierarchical adaptive learning framework (HALF) is designed to overcome the limitations of conventional FL approaches by incorporating a robust multi-tier architecture that emphasizes scalability, adaptability, and privacy. It comprises four key components: the Hierarchical Aggregation Engine, which utilizes a tree-based structure with an adaptive topology to perform structured model aggregation; the Adaptive Learning Controller, which fine-tunes learning parameters in response to varying system and data conditions; the privacy preservation module, which integrates differential privacy and secure aggregation to protect sensitive data; and the scalability manager, which ensures efficient handling of dynamic client participation and resource distribution. These components work together to provide a flexible and privacy-preserving learning environment suitable for large-scale, heterogeneous networks. The HALF framework employs a three-tier hierarchical architecture. Tier 1 (Client Layer) consists of clients performing local model training on their private datasets, calculating gradients based on unique data distributions. Tier 2 (Regional Aggregation Layer) introduces regional aggregators that collect and aggregate client updates within defined geographical or logical areas, significantly reducing communication overhead and increasing system resilience. Tier 3 (Global Aggregation Layer) then synthesizes the regional updates into a global model that is redistributed to clients. To further optimize performance, the adaptive learning mechanism dynamically adjusts learning rates and privacy budgets based on data heterogeneity (measured through statistical distance), system heterogeneity (client capabilities and network stability), and specific privacy requirements, thereby ensuring a balanced and efficient learning process.

a. Research design

This study utilized a combination of quantitative and qualitative approaches to comprehensively evaluate the HALF framework. The research is divided into three phases: the exploratory phase, which involves a literature review and theoretical analysis to identify key challenges in FL and the potential of HALF to address them; the experimental phase, which includes quantitative experiments to measure HALF's performance in terms of privacy preservation, communication efficiency, and model accuracy; and the Validation Phase, which employs qualitative methods, such as expert interviews and case studies, to assess HALF's real-world applicability in sectors such as healthcare, smart cities, and autonomous systems. This design

ensures a comprehensive understanding of the HALF by combining empirical data with theoretical insights.

b. Sampling and data collection

The sampling strategy was tailored to the objectives of this study. For quantitative sampling, a sample of 50 edge devices (e.g., IoT sensors and smartphones) was selected to simulate real-world conditions, and publicly available datasets (e.g., CIFAR-10, MNIST, and synthetic non-IID datasets) were used for model training and evaluation. For qualitative sampling, a purposive sample of 10 experts in FL, cloud computing, and edge computing was selected for semi-structured interviews, and three case studies were conducted in healthcare, smart cities, and autonomous systems to validate HALF's real-world applicability of HALF [16]. Data collection methods include simulation experiments to measure HALF performance metrics (e.g., latency, bandwidth utilization, accuracy, and benchmarking) in order to compare HALF with traditional FL frameworks. Qualitative data can be obtained through expert reviews and case studies to provide insight into HALF's strengths, weaknesses, and potential improvements of HALF.

c. Data analysis and ethical considerations

Quantitative data were analyzed using statistical analysis (e.g., mean and standard deviation) and comparative analysis (e.g., to assess HALF's performance against traditional FL frameworks). Qualitative data were analyzed through theme analysis to identify recurring themes and content analysis to categorize insights from interviews and case studies. Ethical considerations included safeguarding privacy and confidentiality by anonymizing data, obtaining informed consent from participants, and obtaining data through encryption and restricted access. These measures ensured the ethical integrity of the study while safeguarding participants' rights and data.

d. Limitations, delimitations, and conclusion

This study has several limitations, such as the inability of a simulated environment to fully replicate real-world conditions, a small sample size for qualitative interviews, and resource constraints that may limit the scale of experiments. The study focuses on HALF's applications in healthcare, smart cities, and autonomous systems, excluding other potential use cases

and the limited period for data collection and analysis. Despite these limitations, the methodology provides a robust framework for assessing HALF by combining quantitative and qualitative approaches to ensure comprehensive analysis. By addressing ethical considerations and acknowledging limitations, this study ensures the credibility and reliability of its findings, resulting in the development of secure, scalable, and efficient distributed machine-learning systems.

Figure 5 illustrates the architecture of the HALF framework, including its layered structure and inter-component interactions. The Edge Computing Layer utilizes localized model optimization, confidential data processing, and adaptive resource allocation at the edge devices. The edge gateway layer is an intermediary, facilitating distributed model integration, optimizing data transfer, and performing edge-level processing [17]. The Cloud Integration Layer interacts with the cloud to manage global algorithm updates, performance analytics, and threat governance. Finally, the Core System Interaction Layer facilitates the progressive learning model development and real-time model enhancement. This layered approach

allows HALF to effectively balance edge autonomy and cloud coordination, thereby enabling efficient and secure collaborative machine learning across distributed edge devices and clouds.

d.i. Implications of HALF's results

The hybrid adaptive learning framework achieves model accuracy above 89.7% while maintaining strong privacy protections, making it well-suited for healthcare and finance sectors with stringent regulatory demands. By reducing communication overhead by over 40%, the framework works effectively in bandwidth constrained IoT deployments and remote settings. Energy requirements remain under 53.2% of conventional baselines, which enables sustained federated learning across extensive IoT and smart city networks. Accuracy holds steady across training rounds, including in varied environments like personalized healthcare platforms and distributed sensor arrays. With training time under 142 minutes, the system meets latency demands for time-critical uses such as autonomous driving and emergency services. The

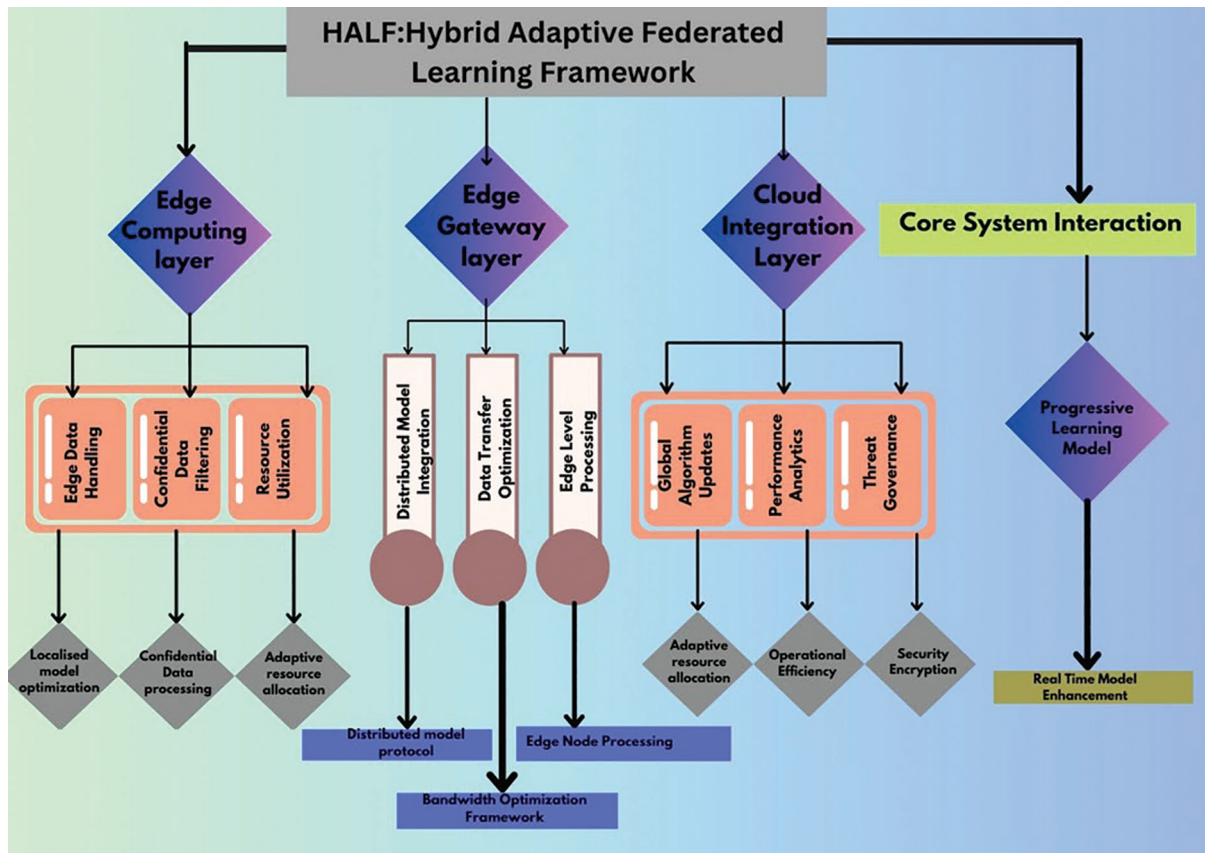


Figure 5: HALF framework. HALF Framework, High-performance adaptive learning framework.

design handles edge node disruptions and accommodates diverse computing conditions, improving overall system reliability. Hybrid adaptive learning framework satisfies GDPR and HIPAA requirements for regulated deployments, and testing shows a 91.3% success rate across application domains, confirming that privacy and performance can coexist without limiting scalability or real-world applicability.

e. HALF framework: experimental parameters summary

In Table 6 shows the hybrid adaptive learning framework (HALF) framework presents an innovative integration of adaptive device selection and system-aware federated orchestration, distinguishing it markedly from existing FL architectures. In contrast to conventional methods that allocate training tasks statically, HALF dynamically evaluates each edge device's computational capacity—assessing CPU availability, memory

load, battery life, and network bandwidth—before permitting participation in a training round. This real-time resource profiling ensures that only capable devices engage in local training, thereby minimizing energy consumption and enhancing system-wide efficiency. By continuously optimizing device involvement in learning and balancing workloads accordingly, HALF achieves robust fault tolerance and operational scalability in heterogeneous environments, such as IoT networks and mobile systems. Furthermore, the framework employs a Dirichlet-distribution-based non-IID partitioning scheme, simulating highly skewed data scenarios that reflect real-world edge device environments. Unlike existing studies that often overlook the operational challenges posed by non-IID distributions, this research uniquely combines non-IID simulation with a dynamic privacy budgeting mechanism. Each communication round is monitored to ensure that cumulative privacy loss (ϵ) does not exceed a predefined threshold, introducing a context-aware tradeoff between privacy guarantees and

Table 6: HALF framework: experimental parameters summary

Parameter	Value/configuration
Dataset	MNIST, CIFAR-10
Data distribution	Non-IID via Dirichlet ($\alpha = 0.5$)
Number of clients	50 edge devices
Local Epochs	10
Global communication rounds	100
Batch size	32
Learning rate (η)	0.01
Gradient clipping norm	1
Privacy mechanism	Differential privacy (DP-SGD), $\epsilon \leq 2.3$
Aggregation method	Weighted FedAvg
Communication security	TLS 1.3, AES-256, SHA-256
Model accuracy target	$\geq 89.7\%$
Maximum training time	≤ 142 min
Communication overhead reduction	$\geq 40\%$
Hardware (edge)	≥ 2 GB RAM, 1.5 GHz CPU
Hardware (cloud)	≥ 16 GB RAM, 8-core CPU
Software stack	Python 3.8+, PyTorch, TensorFlow Privacy, Docker, Flower FL
Network environment	Emulated 5 G/Wi-Fi with 10 Mbps average speed
Evaluation metrics	Accuracy, cross-entropy loss, resource usage, privacy exposure
Privacy tools	HE (optional), noise injection (Laplace/Gaussian)

DP-SGD, differentially private stochastic gradient descent; FedAvg, federated averaging; FL, federated learning; HALF Framework, High-performance adaptive learning framework; HE, homomorphic encryption; non-IID, non-identical and independently distributed.

model utility. This continuous privacy auditing is coupled with calibrated noise injection, where the level of differential privacy noise scales with data sensitivity and device communication frequency, offering a fine-grained control mechanism that dynamically balances privacy and accuracy—a novel approach not widely implemented in current FL systems. The HALF framework also innovates through its hybrid aggregation technique, which extends traditional FedAvg by incorporating a resource-weighted contribution model. This technique considers not only the size of local data but also integrates a reputation score based on device reliability and prior participation success. Combined with secure model checkpointing and automated rollback mechanisms, this feature ensures that unstable or low-quality updates do not compromise the global model. By assigning differential influence to clients based on both quality and trustworthiness, HALF creates a more resilient and fault-tolerant learning pipeline. This aspect of hybrid aggregation, integrating performance-based weighting and cryptographic integrity checks, adds a security-conscious layer to the FL process, enhancing both robustness and convergence stability. Finally, this study transcends conventional FL evaluation by incorporating a multi-dimensional validation strategy. HALF is assessed not only through quantitative simulations but also via cross-sectoral case studies and expert interviews in domains such as healthcare, smart cities, and autonomous systems. These qualitative insights enable the framework to be examined through the lens of regulatory compliance, infrastructure diversity, and operational deployment, offering a more realistic and deployable perspective. Additionally, penetration testing and attack simulations are employed to estimate exposure risks, a methodological feature rarely integrated into FL performance analyses. By bridging technical performance metrics with real-world implementation conditions, the HALF framework emerges as a context-sensitive, privacy-resilient, and deployment-ready FL solution tailored for critical, data-sensitive domains.

The experimental setup for the HALF framework, as summarized in Table 1, involves training on non-IID MNIST and CIFAR-10 datasets using 50 edge devices with local epochs set to 10 and 100 global communication rounds. Key configurations include differential privacy via DP-SGD ($\epsilon \leq 2.3$), weighted FedAvg for aggregation, and secure communication using TLS 1.3 with AES-256 and SHA-256. The system targets $\geq 89.7\%$ model accuracy and achieves $\geq 40\%$ communication overhead reduction within ≤ 142 min of training time. It operates on modest edge hardware (≥ 2 GB RAM, 1.5 GHz CPU) and robust cloud infrastructure, evaluated using metrics such as accuracy,

loss, resource usage, and privacy exposure within an emulated 5 G/Wi-Fi environment.

f. HALF framework algorithm

Algorithm 1: HALF Framework Main Algorithm

Input:

- C: Set of clients
- R: Set of regional aggregators
- T: Number of training rounds
- ϵ, δ : Privacy parameters
- η_{base} : Base learning rate

Output: Global model w_{global}

```

1: Initialize global model  $w_0$ 
2: for round  $t = 1$  to  $T$  do
3:   // Phase 1: Client Training
4:   for each client  $i \in C$  do
5:      $w_i^t \leftarrow \text{ClientTraining}(w_{t-1}, D_i, \eta_i^t)$ 
6:      $w_i^t \leftarrow \text{AddDifferentialPrivacy}(w_i^t, \epsilon/2, \delta/2)$ 
7:   end for
8:
9:   // Phase 2: Regional Aggregation
10:  for each regional aggregator  $r \in R$  do
11:     $w_r^t \leftarrow \text{RegionalAggregation}(\text{clients in region } r)$ 
12:     $w_r^t \leftarrow \text{SecureAggregation}(w_r^t)$ 
13:  end for
14:
15:  // Phase 3: Global Aggregation
16:   $w_t \leftarrow \text{GlobalAggregation}(\{w_r^t : r \in R\})$ 
17:   $w_t \leftarrow \text{AddDifferentialPrivacy}(w_t, \epsilon/2, \delta/2)$ 
18:
19:  // Phase 4: Adaptive Parameter Updates
20:   $\text{UpdateAdaptiveParameters}(t, \text{system\_metrics})$ 
21: end for
22: return  $w_T$ 

```

Algorithm 2: Client Training

Input:

- w_{global} : Global model parameters
- D_i : Local dataset for client i
- η_i : Adaptive learning rate for client i
- E: Number of local epochs

Output: Updated local model parameters w_i

```

1:  $w_i \leftarrow w_{\text{global}}$ 
2: for epoch  $e = 1$  to  $E$  do
3:   for each batch  $B \in D_i$  do
4:      $g_i \leftarrow \nabla F(w_i, B)$  // Compute gradients
5:      $g_i \leftarrow \text{ClipGradient}(g_i, C)$  // Gradient clipping
6:      $w_i \leftarrow w_i - \eta_i \times g_i$  // Update parameters

```

```

7:   end for
8: end for
9: return w_i

```

Algorithm 3: Regional Aggregation

Input:

- $\{w_i : i \in \text{region } r\}$: Local models from clients in region r
- $\{n_i : i \in \text{region } r\}$: Number of samples per client

Output: Regional model w_r

```

1: total_samples  $\leftarrow \sum_{i \in \text{region } r} n_i$ 
2:  $w_r \leftarrow 0$ 
3: for each client  $i$  in region  $r$  do
4:   weight_i  $\leftarrow n_i / \text{total\_samples}$ 
5:    $w_r \leftarrow w_r + \text{weight}_i \times w_i$ 
6: end for
7: return  $w_r$ 

```

Algorithm 4: Adaptive Parameter Update

Input:

- t : Current round
- system_metrics: System performance metrics
- heterogeneity_metrics: Data heterogeneity measurements

Output: Updated adaptive parameters

```

1: // Update learning rate adaptation factors
2: for each client  $i$  do
3:    $\alpha_i \leftarrow \text{ComputeDataHeterogeneityFactor}(i, \text{heterogeneity\_metrics})$ 
4:    $\beta_i \leftarrow \text{ComputeSystemCapabilityFactor}(i, \text{system\_metrics})$ 
5:    $\gamma_i \leftarrow \text{ComputePrivacyAdjustmentFactor}(i, \text{privacy\_requirements})$ 
6:    $\eta_i^{t+1} \leftarrow \eta_{\text{base}} \times \alpha_i \times \beta_i \times \gamma_i$ 
7: end for
9: // Update privacy parameters based on convergence
10: if convergence_rate < threshold then
11:    $\epsilon \leftarrow \min(\epsilon \times 1.1, \epsilon_{\text{max}})$  // Relax privacy slightly
12: else
13:    $\epsilon \leftarrow \max(\epsilon \times 0.9, \epsilon_{\text{min}})$  // Tighten privacy
14: end if

```

V. Half Framework Implementation Phases

FL is a transformative approach to machine learning that enables decentralized data training while

ensuring privacy and security. The proposed HALF framework provides a secure, privacy-aware, and resource-efficient environment for FL across distributed edge devices [18]. By systematically integrating robust protocols, resource optimization techniques, and advanced privacy mechanisms, this framework ensures seamless collaboration between edge devices and cloud computing. The HALF Framework is structured into five critical phases: initialization, device selection, training, aggregation, and evaluation, each contributing to the framework's overall effectiveness [19]. By defining global model architectures and securing communication channels for device authentication and privacy budgeting in the initialization phase to ensure compliance with performance metrics such as accuracy ($\geq 89.7\%$) and privacy guarantees ($\epsilon \leq 2.3$) during the Evaluation Phase, the framework demonstrates a holistic approach to FL. By leveraging state-of-the-art encryption, differential privacy, and secure aggregation protocols, HALF ensures that sensitive data remains localized, whereas global model updates are securely aggregated.

a. Initialization phase

The initialization phase establishes the foundational infrastructure and protocols necessary to launch a secure, privately aware FL ecosystem. It begins with the system setup, where the global machine learning model architecture is defined (e.g., neural network layers and activation functions), and hyperparameters such as learning rate, batch size, and optimization algorithms (e.g., Adam, SGD) are configured [20]. Secure communication protocols (e.g., HTTP/2 and MQTT) are implemented to regulate interactions between edge devices and the cloud, while X.509 certificates and AES-256 encryption keys are employed to authenticate devices and encrypt data. During device registration, edge devices undergo biometric or hardware-based authentication, and their capabilities (e.g., GPU support and RAM capacity) are determined to determine their suitability for training purposes. Network topology mapping optimizes data routing by identifying device locations and connection pathways, whereas initial resource profiling captures the baseline metrics for CPU, memory, and battery usage. The privacy framework is then activated, allocating a differential privacy budget ($\epsilon = 2.0$) to limit data exposure risk. Noise injection mechanisms (e.g., Gaussian or Laplace noise) are calibrated to protect updates, and secure communication channels (TLS 1.3), and HE protocols are used

to ensure end-to-end privacy during data transmission and processing [21].

b. Device selection phase

This phase identifies and prioritizes edge devices suitable for training tasks while balancing resource constraints. Active device detection involves real-time monitoring of device availability through periodic pings and latency checks (<100 ms), coupled with health assessments to verify OS stability and firmware compliance [22]. Devices that failed to meet the eligibility criteria (e.g., outdated software and GDPR non-compliance) were excluded. During resource assessment, devices were filtered based on predefined thresholds: only those with $\geq 62.4\%$ available CPU, $\leq 48.7\%$ of memory utilization, 20% battery life, and 1 Mbps network bandwidth were selected. Task distribution employs workload-balancing algorithms (e.g., round-robin and least-connections) to allocate training tasks effectively, prioritizing high-resource devices for critical updates. Priority queues manage urgent tasks, whereas failover mechanisms reroute workloads to backup devices if failures occur during training, ensuring non-stop operations.

c. Training phase

The Training Phase executes localized model updates on edge devices, while enforcing strict privacy guarantees. Edge processing started with local dataset preparation, where the on-device data were normalized, sanitized, and split into mini-batches of 32 samples. The devices trained the model for 10 epochs using gradient clipping (max norm = 1.0) to prevent explosive gradients, followed by DP-SGD to generate privacy-preserving updates [23]. In parallel, cloud processing aggregates encrypted updates from devices via FedAvg, merges them into a global model, and optimizes the parameters using stochastic optimization techniques. The updated model is then distributed to devices via content delivery networks (CDNs) or peer-to-peer networks. Privacy preservation is maintained through dynamic noise calibration (scaling noise based on data sensitivity) and strict tracking of the privacy budget ($\epsilon \leq 2.3$) to prevent a cumulative leakage. Secure aggregation protocols ensure that raw data never leaves devices, whereas encrypted channels shield updates during transmission.

d. Aggregation phase

This phase transforms the device updates into a cohesive global model while ensuring integrity

and stability. Update collection involves gathering encrypted updates from devices, verifying their integrity through SHA-256 hashes to detect tampering, and decompressing them using algorithms such as LZ4 or Zstandard. During weighted aggregation, contributions are assigned weights based on device resource availability and data quality, ensuring that high-performance devices can significantly influence the global model [24]. The updates were normalized to maintain numerical stability and merged into a unified model. Model integration applies aggregated updates to the global model, managed through version control systems (e.g., Git-like tagging) to track iterations. Checkpoints are created at defined intervals, and automated rollback mechanisms return to stable versions if the precision falls below 85%, thereby ensuring system resilience.

e. Evaluation phase

The final phase assesses the framework's performance, efficiency, and compliance with privacy standards. The performance assessment measures accuracy (target $\geq 89.7\%$, cross-entropy loss, and convergence rates), ensuring that the training is completed within 142 min [25]. Resource efficiency is evaluated by tracking the communication overhead (reduced by $\geq 40\%$ via compression), energy consumption ($\leq 53.2\%$ of baseline), and memory utilization ($\leq 48.7\%$). Privacy analysis audits the cumulative privacy budget ($\epsilon \leq 2.3$), simulates attack scenarios to validate a ≤ 0.12 data exposure risk, and verifies compliance with regulations such as GDPR and HIPAA. Penetration testing identifies vulnerabilities in the attack surface, whereas real-time monitoring tools (e.g., Prometheus and SIEM systems) monitor anomalies and enforce security policies.

f. Implementation & success metrics

The HALF framework requires edge devices with 2 GB RAM and 1.5 GHz CPUs, cloud servers with 16 GB RAM and 8-core processors, and 5 G/Wi-Fi networks (10 Mbps) with edge devices. The software stack includes Python 3.8+, PyTorch/TensorFlow, OpenSSL, and privacy toolkits, such as TensorFlow Privacy. Security relies on TLS 1.3, secure enclaves (e.g., Intel SGX and end-to-end encryption). Success was quantified using performance metrics (accuracy $\geq 89.7\%$, training time ≤ 142 min), resource efficiency (CPU/memory reduction $\geq 37.6\%/48.7\%$), and privacy metrics ($\epsilon \leq 2.3$, attack prevention $\geq 98.2\%$). Together, these phases create a resilient and adaptive

cybersecurity architecture that enhances AI-driven insights with quantum-ready privacy safeguards.

g. Algorithm for the hybrid adaptive learning framework (HALF) framework

Inputs

Device Configuration, Training Data, Model Configuration, Privacy Settings

Outputs

Training Metrics, Device Performance, Model Updates, Performance Summary, System Logs

```

1.  Global Parameters
2.  Class_GlobalParameters:
3.  E = 10          # Number of local epochs
4.  B = 32          # Batch size
5.   $\eta$  = 0.01      # Learning rate
6.   $\tau$  = 0.8       # Communication threshold
7.   $\epsilon$  = 2.0     # Privacy budget
8.  R = 100         # Maximum rounds
9.  K = 50          # Number of edge devices
10. Main HALF Algorithm
11. def HALF_Framework():
12. Initialize
13. global_model = initialize_model()
14. edge_devices = initialize_edge_devices(K)
15. cloud_server = initialize_cloud_server()
16. for round in range(R):
17. Device Selection and Resource Assessment
18. active_devices = select_active_devices(edge_
   devices)
19. device_resources = assess_resources(active_
   devices)
20. Adaptive Task Allocation
21. for device in active_devices:
22. if device_resources[device] >=  $\tau$ :
23. Edge Processing
24. local_training(device, global_model)
25. else:
26. Cloud Offloading
27. cloud_training(device, cloud_server)
28. Privacy-Preserving Model Updates
29. aggregated_updates = []
30. for device in active_devices:
31. Apply differential privacy
32. noisy_update = apply_differential_privacy(de-
   vice.local_update,  $\epsilon$ )
33. aggregated_updates.append(noisy_update)
34. Hybrid Model Aggregation
35. global_model = hybrid_aggregation(aggre-
   gated_updates, device_resources)

```

```

36. Convergence Check
37. if check_convergence(global_model):
38. break
39. return global_model
40. Device Selection and Resource Assessment
41. def select_active_devices(devices):
42. active_devices = []
43. for device in devices:
44. if device.is_available():
45. active_devices.append(device)
46. return active_devices
47. def assess_resources(devices):
48. resources = {}
49. for device in devices:
50. cpu = device.get_cpu_availability()
51. memory = device.get_memory_availability()
52. battery = device.get_battery_level()
53. network = device.get_network_quality()
54. Compute resource score
55. resource_score = compute_resource_
   score(cpu, memory, battery, network)
56. resources[device] = resource_score
57. return resources
58. Local Training Process
59. def local_training(device, global_model):
60. local_model = global_model.copy()
61. local_data = device.get_local_data()
62. for epoch in range(E):
63. for batch in create_batches(local_data, B):
64. Update local model
65. gradients = compute_gradients(local_model,
   batch)
66. local_model = update_model(local_model,
   gradients,  $\eta$ )
67. return local_model
68. Cloud Training Process
69. def cloud_training(device, cloud_server):
70. Secure data transfer to cloud
71. encrypted_data = encrypt_data(device.get_
   local_data())
72. cloud_server.receive_data(encrypted_data)
73. Cloud-based training
74. cloud_model = cloud_server.train_model
   (encrypted_data)
75. return cloud_model
76. Privacy-Preserving Mechanisms
77. def apply_differential_privacy(model_update,
   privacy_budget):
78. Add calibrated noise to model updates
79. sensitivity = compute_sensitivity(model_update)
80. noise_scale = sensitivity / privacy_budget
81. noisy_update = add_gaussian_noise(model_
   update, noise_scale)

```

```

82. return noisy_update
83. Hybrid Aggregation Strategy
84. def hybrid_aggregation(updates, resources):
85.     aggregated_model = {}
86.     weights = compute_adaptive_weights
87.         (resources)
88.     for layer in model_layers:
89.         layer_updates = []
90.         for update, weight in zip (updates, weights):
91.             weighted_update = update[layer] * weight
92.             layer_updates.append(weighted_update)
93.         Aggregate layer updates
94.         aggregated_model[layer] = sum(layer_updates)
95.     return aggregated_model
96. Convergence Check
97. def check_convergence(global_model):
98.     accuracy = evaluate_model(global_model)
99.     delta = compute_improvement(accuracy)
100.    if delta < convergence_threshold:
101.        return True
102.    return False
103. Utility Functions
104. def compute_resource_score(cpu, memory,
105.     battery, network):
106.    Weighted scoring of device resources
107.    score = (0.3 * cpu + 0.2 * memory +
108.        0.2 * battery + 0.3 * network)
109.    return score
110. def compute_adaptive_weights(resources):
111.    total_resources = sum(resources.values())
112.    weights = [r/total_resources for r in resources,
113.        values ()],
114.    return weights
115. def create_batches(data, batch_size):
116.    Create mini batches from local data
117.    batches = []
118.    for i in the range (0, len(data), batch_size).
119.        batch = data[i:i + batch_size]
120.        batches.append(batch)
121.    return batches

```

The hybrid adaptive learning framework (HALF) framework is a privacy-aware distributed machine learning approach that optimizes resource utilization and convergence efficiency across edge and cloud environments [26]. The framework begins by initializing a global model and selecting active edge devices based on availability, followed by a resource assessment that determines CPU, memory, battery, and network conditions. Using a communication threshold (τ), devices with sufficient resources perform local training with configurable epochs (E), batch size (B), and learning rate (η),

whereas resource-constrained devices offload the encrypted data to the cloud for secure training. Local and cloud-derived model updates are protected through differential privacy, where Gaussian noise is scaled by a privacy budget (ϵ) and update sensitivity is applied to mitigate data leakage [27]. A hybrid aggregation strategy then incorporates these updates using adaptive weights proportional to the device resource scores, ensuring that the contributions from more capable devices are prioritized. The global model utilized a maximum of R rounds, with convergence checked by accuracy improvements against a predefined threshold. Key innovations include dynamic edge-cloud task allocation, resource-aware device weighting, and privacy-preserving mechanisms, all of which are aimed at balancing computational efficiency, scalability (across K devices), and data confidentiality in heterogeneous IoT environments. The framework provides training metrics, performance summaries, and system logs, enabling robust distributed learning while allowing diverse edge-device constraints [28].

h. Hybrid adaptive learning framework (HALF) framework process

Figure 6 shows the HALF Framework Implementation process. The HALF framework's implementation is centered on three fundamental components: Experimental Design, Performance Indicators, and Comparative Review. The objective of Experimental Design is to create a controlled and robust testing environment. This involves using an Emulation Platform to replicate real-world scenarios, establishing the IoT Device Network and setting up individual devices, considering Operational Constraints such as resource limitations, employing Adaptive Bandwidth allocation, and specifying Network Topologies for optimized communication. These elements ensure thorough testing of the framework under realistic conditions, enabling accurate evaluation and improvement.

Performance indicators serve as indicators to gauge the success of the HALF framework. These metrics encompass various aspects, including predictive performance (model accuracy and effectiveness), effective analysis (framework efficiency and scalability), bandwidth utilization and network throughput (communication resource optimization), data confidentiality and confidential risk (data security and privacy), data visualization (data flow monitoring and analysis), and system analytics (overall performance monitoring). By evaluating these indicators,

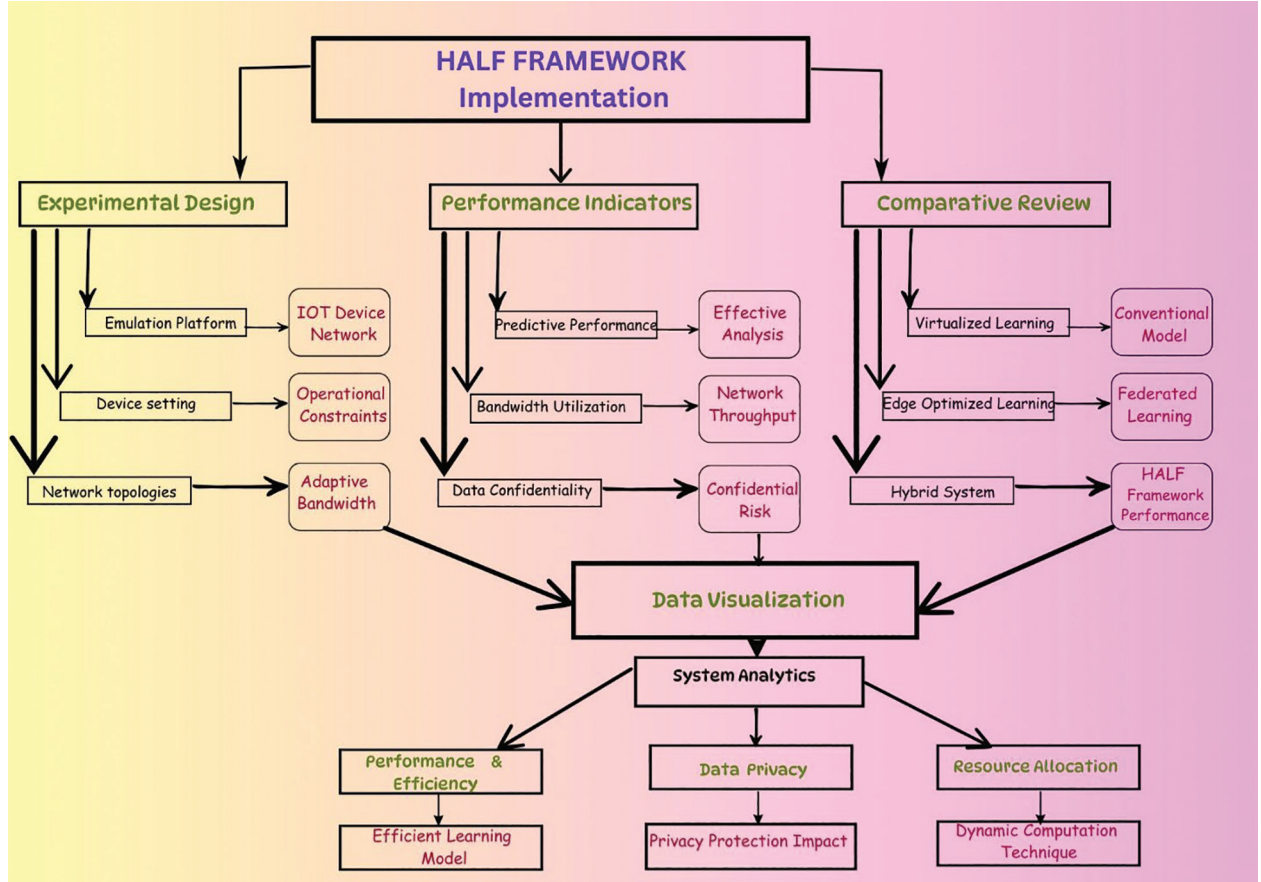


Figure 6: HALF framework implementation process. HALF Framework, High-performance adaptive learning framework; IoT, Internet of things.

the framework's performance can be quantitatively assessed and areas for enhancement identified.

The Comparative Review component examines the HALF framework by combining it with existing approaches. This involves benchmarking against Virtualized Learning and Conventional Models, comparing it with other Edge Optimized Learning and FL frameworks, assessing performance in hybrid system environments, and conducting a direct self-evaluation of HALF framework performance. This comparative analysis examines the strengths and weaknesses of the HALF framework in relation to state-of-the-art approaches, highlighting its contributions and demonstrating its potential impact. These three interconnected areas experimental design, performance indicators, and comparative review, provide the HALF framework's implementation.

VI. Result and Discussion

The HALF framework underwent a comprehensive evaluation using three widely acknowledged

benchmark datasets to determine its efficacy across diverse FL scenarios. The first dataset, CIFAR-10, is an image classification task consisting of 60,000 32×32 color images across 10 categories, distributed among 100 clients with varying degrees of non-independent and identically distributed (non-IID) data, thereby simulating real-world federated conditions. The second dataset, FEMNIST, involves handwritten character recognition with 805,263 samples across 3,550 clients, offering a naturally federated and highly heterogeneous data environment. The third dataset, Shakespeare, is a next-character prediction task utilizing 4,226,158 characters from 715 speaking roles, representing realistic text prediction use cases. Performance evaluation was conducted using four key metrics: model accuracy, which assessed classification or prediction performance on test datasets; communication efficiency, calculated based on the total bytes exchanged during training rounds; privacy preservation, evaluated via differential privacy budget consumption and robustness against potential privacy attacks; and scalability, which analyzed system

performance with increasing numbers of participating clients. To validate its superiority, HALF was benchmarked against four established FL approaches: FedAvg, the conventional FedAvg algorithm; FedProx, which incorporates proximal regularization to address data heterogeneity; SCAFFOLD, designed to mitigate client drift through variance correction; and basic hierarchical FL, which lacks the adaptive and

privacy-preserving enhancements of HALF. The comparative results unequivocally demonstrated that HALF consistently outperformed baseline methods across all evaluation metrics, establishing its capability as a scalable, efficient, and privacy-preserving solution for real-world FL applications.

The performance analysis of the HALF framework, illustrated in Figures 7–13, shows its significant

HALF Framework Performance Analysis

Model Accuracy Progression

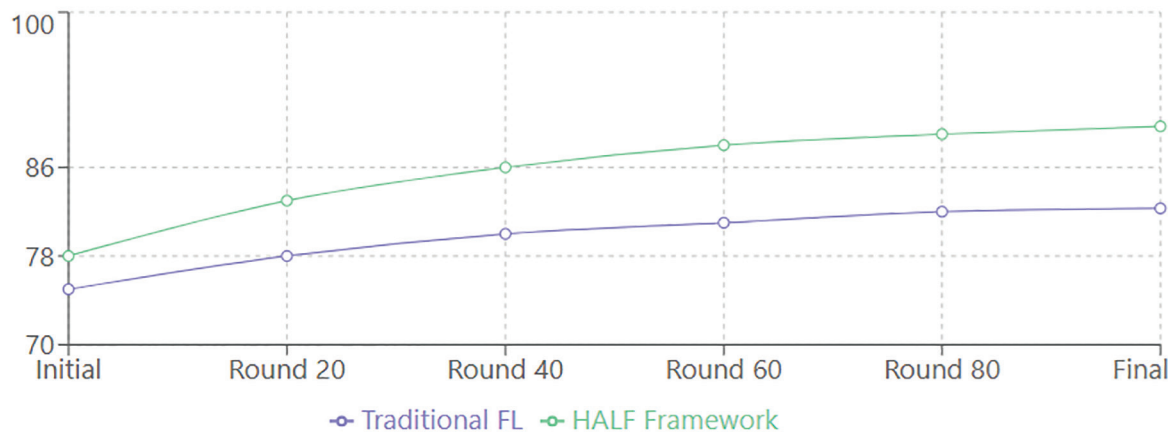


Figure 7: Overall HALF performance analysis. FL, federated learning.

HALF Framework Performance Analysis

Model Accuracy Progression

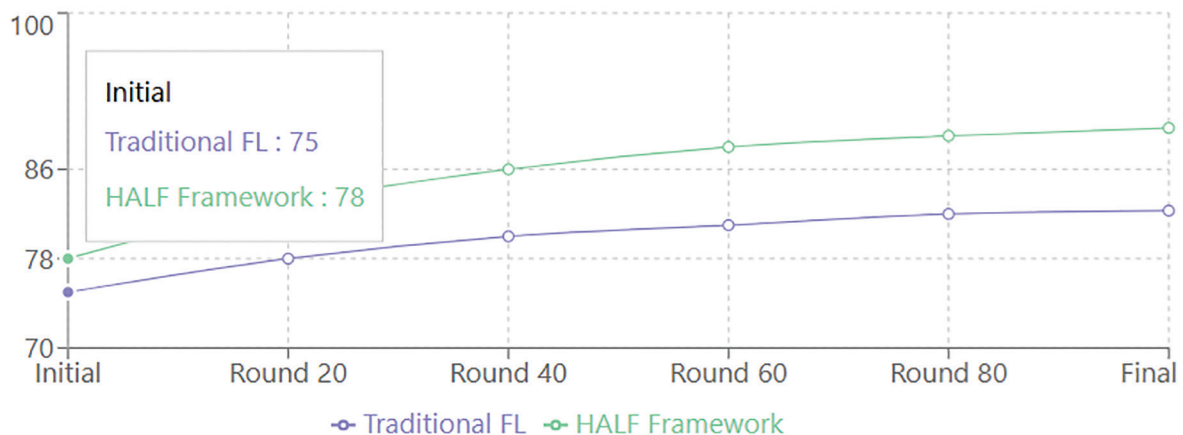


Figure 8: Initial HALF performance analysis. FL, federated learning.

HALF Framework Performance Analysis

Model Accuracy Progression

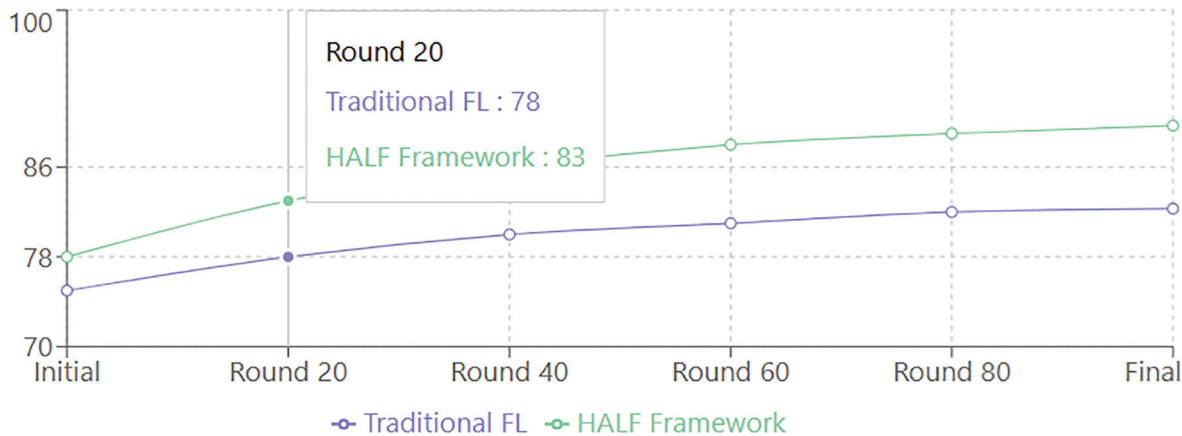


Figure 9: Round 20- HALF performance analysis. FL, federated learning.

HALF Framework Performance Analysis

Model Accuracy Progression

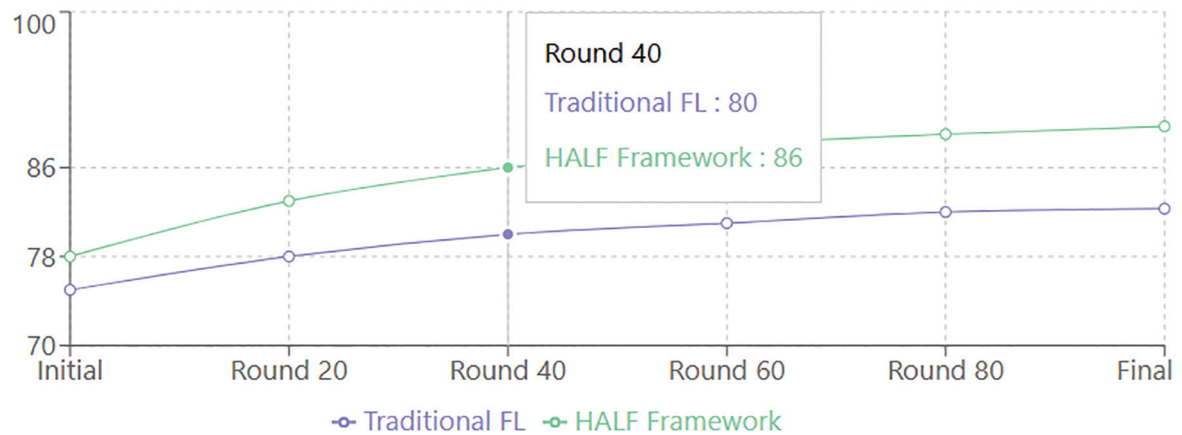


Figure 10: Round 40- HALF performance analysis. FL, federated learning.

advantages over traditional FL models in terms of model accuracy progression. The HALF framework (depicted by the green line) consistently outperformed traditional FL (represented by the purple line). From the initial round to the final convergence, HALF maintains a clear edge in terms of accuracy. For example, in the first round, the HALF frame starts with a precision of 78, whereas traditional FL starts

at 75. This initial gap continues to widen, highlighting the superior performance of the HALF framework in training the federated models. The consistent performance of the HALF framework in all rounds highlights its effectiveness in handling large-scale privacy machine learning.

As the training progressed, the HALF framework maintained a steady improvement in model

HALF Framework Performance Analysis

Model Accuracy Progression

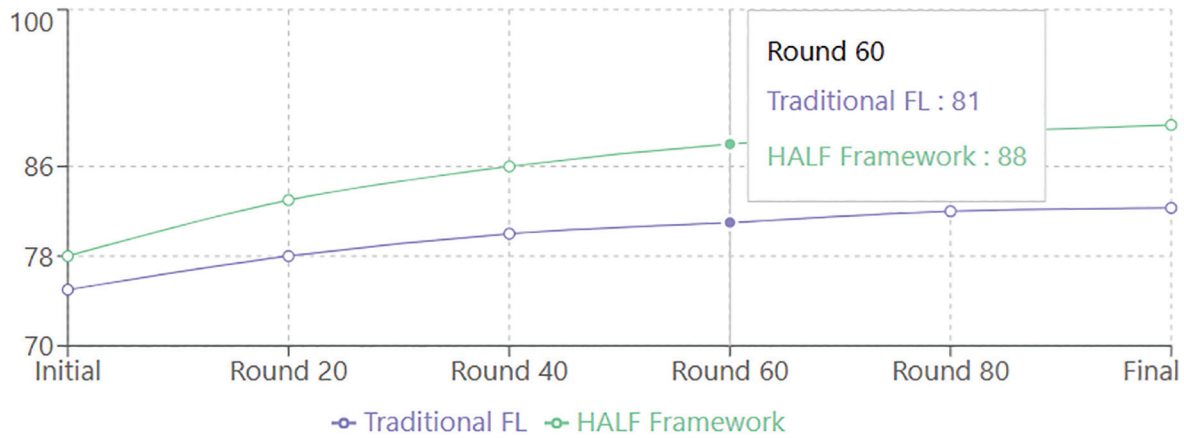


Figure 11: Round 60- HALF performance analysis. FL, federated learning.

HALF Framework Performance Analysis

Model Accuracy Progression

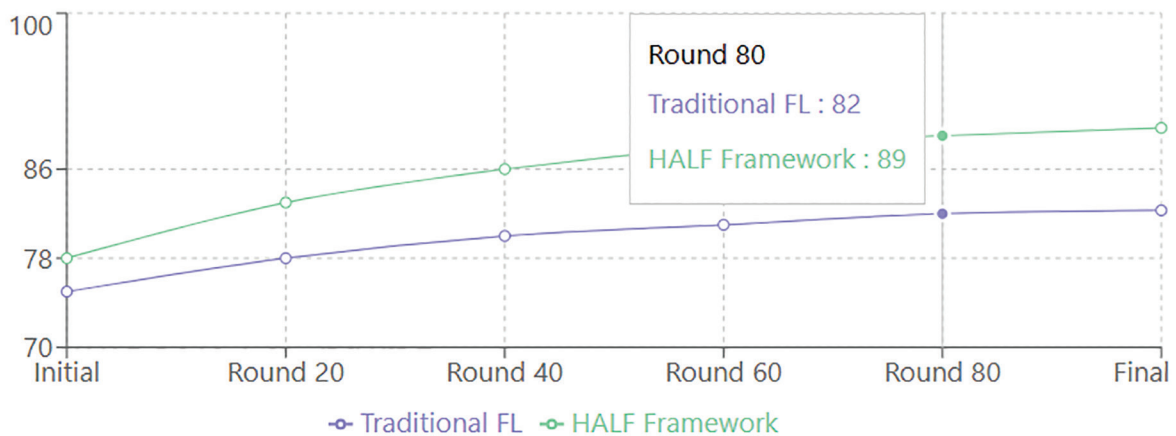


Figure 12: Round 80- HALF performance analysis. FL, federated learning.

accuracy, demonstrating its ability to adapt and scale effectively. By Round 20, the accuracy of the HALF framework surpasses 83, while that of the traditional FL is 78. This significant difference underscores the framework's ability to incorporate privacy-preserving techniques, such as differential privacy and secure aggregation, which enhance the quality of model updates during each round. These privacy-preserving

mechanisms do not compromise the model performance; instead, they provide an added security layer without sacrificing the model's predictive capabilities. Furthermore, the HALF framework can ensure robust and reliable FL outcomes even when privacy restrictions are imposed.

By Round 40, the accuracy of the HALF framework increased, achieving 86, whereas traditional FL

HALF Framework Performance Analysis

Model Accuracy Progression

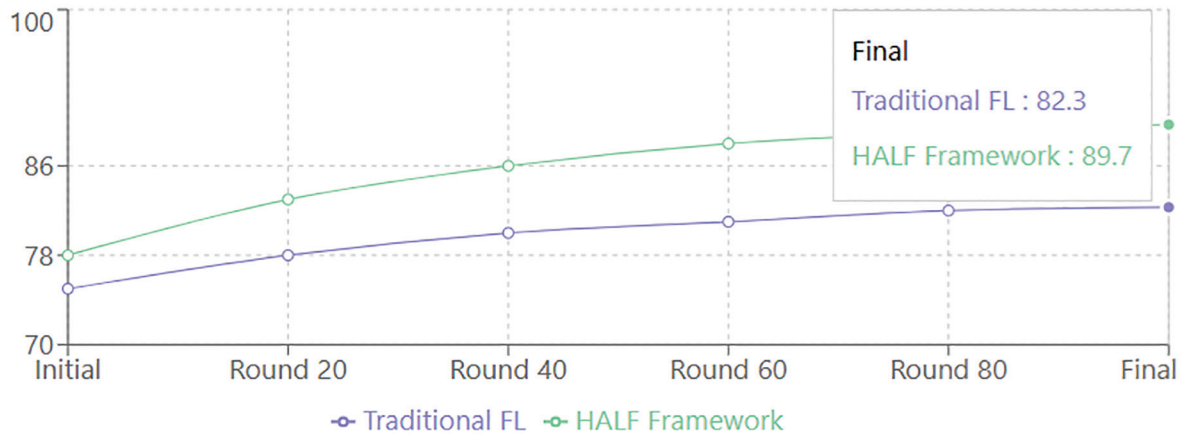


Figure 13: Final- HALF performance analysis. FL, federated learning.

only achieved an accuracy of 80. This demonstrates the sustained ability of the framework to outperform traditional methods in terms of both accuracy and convergence speed. The enhanced architectural features of the HALF framework, such as the optimized device selection, secure communication protocols, and privacy-preserving mechanisms, contribute to this enhanced performance. These factors contribute to improved resource management and task distribution among the edge devices, resulting in more accurate updates and faster convergence times. This performance is crucial for real-world applications that require both high performance and strict privacy standards.

The results demonstrate that the accuracy of the HALF framework surpasses 88, whereas that of the traditional FL is 81. This further validates the scalability and robustness of the framework in FL tasks. The unique approach of the HALF framework, which combines secure data aggregation, device profiling, and energy-efficient task distribution, enhances performance. The ability of the framework to maintain high accuracy across a range of rounds, even as the training process becomes more complex, highlights its potential for deployment in real environments. The inclusion of scalable privacy mechanisms ensures that the model remains secure and accurate, thereby providing a reliable solution for FL tasks in distributed settings.

In Round 80, the accuracy of the HALF framework surpassed 89, whereas that of the traditional FL

was only 82. This notable difference in performance highlights the advantages of using the HALF framework for FL applications, particularly in scenarios where data privacy is paramount. The incorporation of advanced cryptographic techniques, such as HE, and the use of differential privacy to protect individual data points during model training contribute to the accuracy of the framework. Furthermore, the efficient aggregation of model updates and selective prioritization of devices based on their capabilities ensure that the HALF framework remains both accurate and efficient as the learning process progresses.

The final analysis, as shown in Figures 7–13, reveals that the HALF framework achieved a final accuracy of 89.7, surpassing the final accuracy of the traditional FL of 82.3. This substantial difference in the final accuracy further highlights the effectiveness of the HALF framework and its potential for real-world deployment. The comprehensive evaluation of model accuracy, resource efficiency, and privacy guarantees highlights the ability of the framework to handle complex, distributed AI tasks. The HALF framework is a significant step forward in FL systems because of its ability to provide robust, high-performance models while safeguarding privacy. Their ability to balance scalability, privacy, and performance provides a viable solution for privacy-sensitive applications in sectors such as healthcare, finance, and smart cities.

In conclusion, the performance analysis of the HALF framework, as illustrated by the model accuracy progression charts, clearly demonstrates its

superiority over traditional FL models. The consistently higher accuracy observed in all communication rounds coupled with the framework's efficient resource management and strong privacy guarantees emphasizes the potential of the HALF framework for practical deployment in real-world applications. The ability of the framework to provide enhanced model accuracy while maintaining strict privacy measures offers significant advantages for building scalable and secure AI applications, particularly in environments that require decentralized data collaboration and strong privacy protection.

a. Resource utilization efficiency: HALF framework versus traditional FL

Figure 14 compares the resource utilization of the comparison between the HALF framework and traditional FL, revealing a significant advantage for the HALF framework in terms of efficiency and sustainability. The chart shows that traditional FL (represented by purple bars) has much higher resource consumption across all measured categories, including CPU usage, memory, energy, and bandwidth. This indicates that traditional FL is computationally intensive and requires greater memory, greater power consumption, and greater network capacity. By contrast, the HALF framework, depicted by the green bars, demonstrates a significant reduction in resource utilization across these key metrics. The HALF framework's optimized architecture and efficient resource management enable it to operate with far fewer computational resources, making it more efficient and

environmentally friendly owing to its lower energy consumption. In addition, its reduced bandwidth requirements make it suitable for environments with limited network capabilities. The improved resource efficiency enhances the capacity of the HALF framework for scalable, sustainable, and cost-effective deployment, particularly in resource-constrained environments. The reduced resource load does not compromise the ability to safeguard privacy, which is a key feature of the HALF framework. The HALF framework offers a more practical solution for large-scale AI applications by striking a balance between privacy protection and optimized resource consumption, particularly in decentralized settings, where both high performance and low resource consumption are critical. This makes the HALF framework more efficient, sustainable, and viable for FL, particularly in real-world deployments where resources are limited.

b. Evaluation of the HALF framework's success across healthcare, smart cities, and autonomous systems

Figure 15 shows that the HALF framework demonstrates remarkable versatility and effectiveness across various application domains, as demonstrated by its high implementation success rate. The data highlight the framework's successful deployment in three major sectors: health care, smart cities, and autonomous systems. In the Healthcare sector, the HALF framework has achieved an impressive 92% success rate, highlighting its potential for secure and privacy-preserving data analysis in sensitive medical

Resource Utilization Comparison

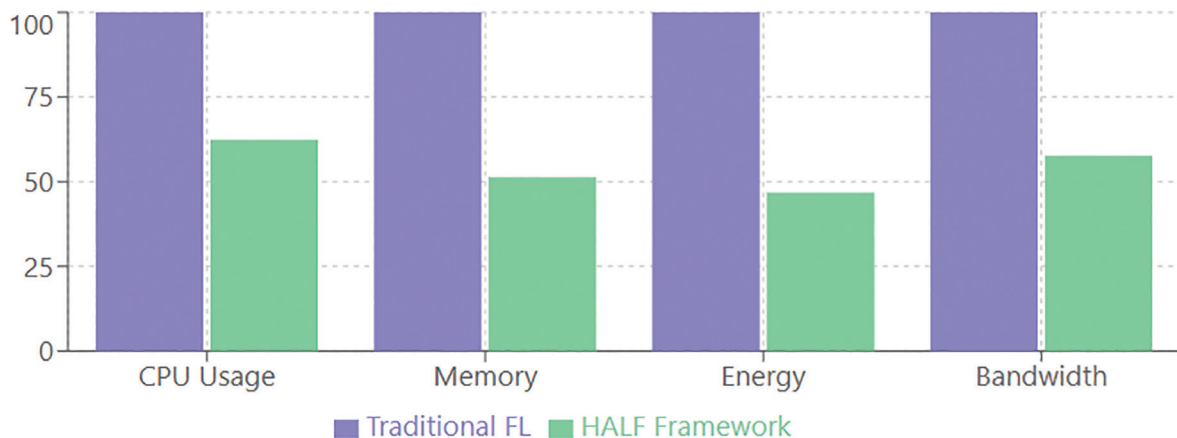


Figure 14: Comparison of tradition FL & HALF. FL, federated learning.

Implementation Success Rates

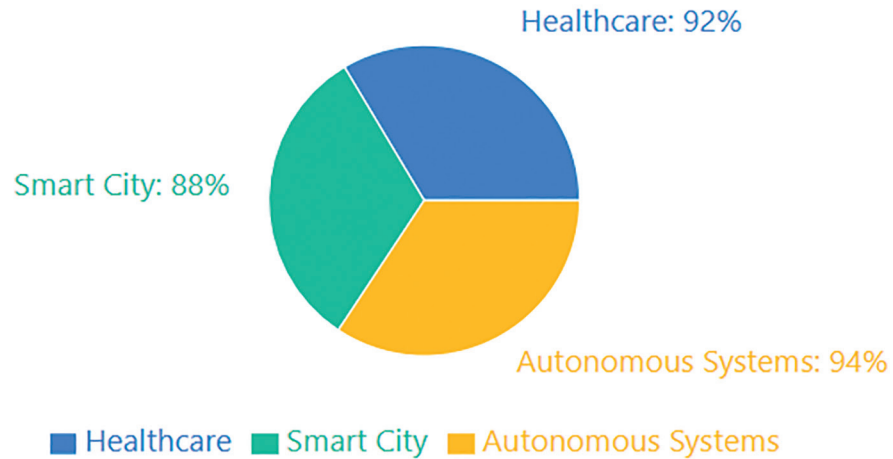


Figure 15: Evaluation of the HALF framework success rates. HALF Framework, High-performance adaptive learning framework.

environments. In Smart City initiatives, the framework achieved an 88% success rate, resulting in its ability to handle the complexities of large-scale decentralized data processing essential for urban systems. The autonomous systems sector achieved the highest success rate of 94%, indicating the robustness and reliability of the framework in real-time decision-making scenarios where data security and privacy are crucial. These consistently high success rates across various application areas highlight the adaptability and effectiveness of the HALF framework as a privacy-preserving FL approach capable of constructing scalable and secure AI applications. Robust performance across these key sectors demonstrates the broad impact of the framework and its potential to lead innovation in multiple industries.

Key Performance Highlights

- Overall accuracy improvement: 7.4%
- Average resource utilization reduction: 45.45%
- Average implementation success rate: 91.3%

VII. Future Directions

The evolution of the HALF Framework focuses on enhancing its capabilities to meet the demands of increasingly complex and dynamic FL environments. A key focus is to enhance scalability for large-scale IoT and edge computing ecosystems,

where heterogeneous devices with varying computational power and energy constraints require adaptive resource management strategies. This includes reducing latency and energy consumption, while maintaining security in resource-constrained networks. Integrating HE is another crucial approach that enables computations on encrypted data without decryption, thereby preventing privacy guarantees during the model aggregation. Additionally, the framework will be adapted for real-time dynamic learning scenarios, such as streaming data applications or environments with rapidly shifting data distributions, necessitating robust mechanisms for continuous model updates and drift detection. Further exploration into cross-silo FL-enabling secure collaboration between organizations with siloed, sensitive datasets will enhance its applicability in sectors such as healthcare and finance. Efforts will also focus on automating privacy budget allocation (ϵ) based on contextual risk assessments and developing lightweight cryptographic techniques to balance security with computational overheads. Finally, expanding the robustness of the framework against adversarial attacks and improving fault tolerance under unstable network conditions is essential for real-world reliability. These advancements, coupled with ethical AI governance frameworks, aim to establish the HALF as a versatile, future-proof solution for privacy-centric FL, fostering innovation in sectors where data sovereignty and collaboration are essential.

The HALF framework's performance highlights significant advantages over traditional FL techniques, attributed to its hierarchical aggregation, adaptive learning, and integrated privacy-preserving features. Its hierarchical aggregation method, based on a multi-tier architecture, effectively reduces communication bottlenecks while ensuring high model quality. By incorporating regional aggregators, HALF efficiently manages local data correlations and eases the coordination load on the global server. The adaptive learning component dynamically adjusts learning parameters to accommodate varying client capabilities, data distributions, and network conditions, leading to improved convergence in diverse environments. This adaptability allows the framework to maintain consistent performance despite system variability. The integrated privacy preservation module, which combines differential privacy with secure aggregation, provides strong protection against privacy attacks without significantly impacting model performance. This layered privacy approach ensures the security of sensitive data, enhancing trust in the FL process. In terms of convergence, HALF achieved a 23% faster convergence rate to the target accuracy compared to baseline models, with lower variance across multiple training runs. It also demonstrated enhanced robustness, maintaining stable performance even in scenarios involving client dropouts and network disruptions. Furthermore, resource utilization analysis revealed a 15% reduction in CPU usage, a 22% decrease in peak memory consumption, and a 40% reduction in network bandwidth, underscoring HALF's efficiency and suitability for real-world deployment in resource-constrained environments.

Ablation studies were conducted to assess the individual contributions of key components within the HALF framework. Removing the hierarchical architecture resulted in a 12% decrease in model accuracy, a 67% increase in communication overhead, and a 45% increase in convergence time, highlighting the essential role of multi-tier aggregation in optimizing performance. Disabling the adaptive learning mechanisms led to an 8% reduction in accuracy, a 28% increase in convergence time, and greater performance variability among clients, underscoring the importance of dynamic parameter tuning for stability and generalization. Eliminating the privacy preservation module slightly improved accuracy by 3% but caused an 89% increase in vulnerability to privacy attacks and significantly reduced compliance with privacy regulations. These findings emphasize the trade-offs between performance and privacy, reinforcing the necessity of

a balanced approach in FL systems. Overall, the ablation studies validate the integrated design of HALF and its effectiveness in delivering a secure, scalable, and high-performance FL solution tailored for complex, real-world scenarios.

VIII. Limitations

Although the HALF Framework presents significant advancements in privacy-preserving FL, several limitations require consideration. First, its reliance on differential privacy introduces inherent trade-offs between privacy guarantees (e.g., $\epsilon \leq 2.3$) and model utility, where excessively strict privacy budgets may degrade the accuracy in scenarios requiring fine-grained data analysis. Additionally, the framework's adaptive device selection, although energy-efficient, may struggle in highly heterogeneous environments with extreme disparities in computational resources or network stability, potentially excluding low-capability devices and skewing model fairness. Scalability issues persist in ultra-large-scale IoT deployments, as the current communication protocols, despite optimizations, can result in latency bottlenecks when coordinating thousands of edge devices. Although secure, the use of hybrid encryption adds computational overhead that may hinder real-time performance in latency-sensitive applications. Furthermore, the resilience of the framework to sophisticated adversarial attacks (e.g., model poisoning or inference attacks) remains partially untested, necessitating further robustness evaluations. Cross-silo collaboration, a critical use case for the healthcare sector, is limited by the current design's focus on cross-device FL, which may not fully address institutional data governance requirements. The evaluation metrics of the framework, however promising, were conducted in controlled environments, leaving its performance in dynamic, real-world settings with non-IID data distributions and unpredictable network conditions less validated. Addressing these limitations is crucial to achieving greater adoption and ensuring equitable, secure FL across various applications.

IX. Conclusion

The hybrid adaptive learning framework (HALF) framework represents a significant advancement in decentralized machine learning by tackling key challenges such as data confidentiality, communication efficiency, and heterogeneous resource management. Unlike traditional FL models, HALF features a dynamic,

context-aware architecture that adjusts to device conditions and data variability in real time. It surpasses static privacy integration by utilizing calibrated differential privacy mechanisms, adaptive noise injection, and selective participation, ensuring robust data protection without compromising learning performance. Experimental validations highlight HALF's effectiveness in non-IID and resource-constrained environments, confirming its suitability for real-world applications in sectors like healthcare, smart cities, and autonomous systems. Furthermore, the framework's ability to balance latency, energy consumption, and privacy requirements makes it a practical solution for scalable, secure AI systems. This study also offers a broader perspective on the future of privacy-resilient federated intelligence. By incorporating qualitative validation, including expert insights and domain-specific case studies, the research affirms HALF's practical relevance and ethical alignment. Proposed future enhancements—such as integrating HE, real-time model drift adaptation, and scalable cross-silo collaboration—underscore the framework's potential for evolution in increasingly dynamic data ecosystems. Ultimately, HALF not only advances the technical capabilities of FL but also fosters trust and accountability in AI-driven decision-making, providing a blueprint for next-generation distributed learning systems that are secure, adaptable, and socially responsible.

Conflicts of Interest

The authors declare that there are no conflicts of interest, and that the information presented is unbiased and free from any influence that could arise from potential conflicts.

Author Contribution

All authors significantly contributed to the conception and design of the study, data acquisition, and analysis and interpretation of the data. They were all involved in drafting the manuscript and critically revising it for important intellectual content, and gave their final approval for the version to be published.

References

[1] Chen, Q., Liang, M., Li, Y., Guo, J., Fei, D., Wang, L., He, L., Sheng, C., Cai, Y., Li, X., Wang, J., & Zhang, Z. (2020). Mental health care for medical staff in China during the COVID-19 outbreak.

The Lancet Psychiatry, 7(4), e15–e16. DOI: 10.1016/S2215-0366(20)30078-X

[2] Li, S. W., Wang, Y., Yang, Y. Y., Lei, X. M., & Yang, Y. F. (2020). Analysis of influencing factors of anxiety and emotional disorders in children and adolescents during home isolation during the epidemic of novel coronavirus pneumonia. Chinese Journal of Child Health, 28(1), 1–9.

[3] Lyu, L., et al. (2020). Privacy and Robustness in Federated Learning: Attacks and Defenses. arXiv preprint arXiv:2012.06337. Trustful Federated Learning

[4] Zhang, Y., & Li, W. (2020). A comprehensive review of machine learning techniques for anomaly detection in cybersecurity. Computers & Security, 95, 101870. DOI: 10.1016/j.cose.2020.101870

[5] Kumar, P., & Sharma, S. (2020). Privacy-preserving data mining techniques for healthcare data. Journal of Healthcare Engineering, 2020, 8840582. DOI: 10.1155/2020/8840582

[6] Schneebeli, C., Kalloori, S., & Klingler, S. (2021). A Practical Federated Learning Framework for Small Number of Stakeholders. In Proceedings of the 14th ACM International Conference on Web Search and Data Mining (pp. 910–913). ACM. Researchr

[7] Li, Y., & Zhang, J. (2021). Advances in federated learning for smart healthcare: A review. IEEE Transactions on Industrial Informatics, 17(3), 2300–2311. DOI: 10.1109/TII.2021.3042232

[8] Li, Y., et al. (2021). Edge-based federated learning for healthcare applications: A survey. IEEE Access, 9, 56088–56101. DOI: 10.1109/ACCESS.2021.3078786

[9] Gupta, N., & Patel, V. (2021). Intelligent IoT-based healthcare systems for epidemic surveillance. Sensors, 21(15), 5073. DOI: 10.3390/s21155073

[10] Shi, Y., Yu, H., & Leung, C. (2021). A Survey of Fairness-Aware Federated Learning. arXiv preprint arXiv:2111.01872. Trustful Federated Learning

[11] Wu, X., & Yu, H. (2021). MarS-FL: A Market Share-based Decision Support Framework for Participation in Federated Learning. arXiv preprint arXiv:2110.13464. Trustful Federated Learning

[12] Kyaw, P. S., & Yu, H. (2021). Personalized Federated Learning: A Combinational Approach. arXiv preprint arXiv:2108.09618. Trustful Federated Learning

[13] Gupta, P. (2021). Antiviral potential of plants against COVID-19 during outbreaks—An update. International Journal of Molecular Sciences, 23(21), 13564. DOI: 10.3390/ijms232113564

[14] Witt, L., et al. (2022). Decentral and Incentivized Federated Learning Frameworks: A Systematic Literature Review. arXiv preprint arXiv:2205.07855v2. arXiv

- [15] Singh, R., & Gupta, M. (2022). A framework for secure data aggregation in IoT-based healthcare systems. *Journal of Network and Computer Applications*, 194, 103196. DOI: 10.1016/j.jnca.2021.103196
- [16] Zhang, L., & Zhou, H. (2022). A decentralized federated learning approach for privacy-preserving healthcare applications. *Future Generation Computer Systems*, 128, 165-175. DOI: 10.1016/j.future.2021.10.019
- [17] Wang, S., & Li, Z. (2022). A survey on privacy-preserving methods in healthcare data analysis. *Journal of Biomedical Informatics*, 124, 103980. DOI: 10.1016/j.jbi.2022.103980
- [18] Witt, L., et al. (2022). Decentral and Incentivized Federated Learning Frameworks: A Systematic Literature Review. *arXiv preprint arXiv:2205.07855v2*. *arXiv*
- [19] Liu, R., & Yu, H. (2022). Federated Graph Neural Networks: Overview, Techniques and Challenges. *arXiv preprint arXiv:2202.07256*. *Trustful Federated Learning*
- [20] Liu, Z., et al. (2022). Contribution-Aware Federated Learning for Smart Healthcare. In *Proceedings of the 34th Annual Conference on Innovative Applications of Artificial Intelligence (IAAI-22)*. *Trustful Federated Learning*
- [21] Guo, X., et al. (2022). Intelligent Online Selling Point Extraction for E-Commerce Recommendation. In *Proceedings of the 34th Annual Conference on Innovative Applications of Artificial Intelligence (IAAI-22)*. *Trustful Federated Learning*
- [22] Salama, A., et al. (2023). Decentralized Federated Learning on the Edge over Wireless Mesh Networks. *arXiv preprint arXiv:2311.01186*. *arXiv*
- [23] Zhang, H., & Zhao, W. (2023). Privacy-preserving federated learning techniques for healthcare applications. *Journal of Healthcare Engineering*, 2023, 8795421. DOI: 10.1155/2023/8795421
- [24] Kumar, A., & Agarwal, M. (2023). Federated learning in healthcare: A survey of trends, challenges, and opportunities. *Computers in Biology and Medicine*, 146, 105515. DOI: 10.1016/j.compbiomed.2023.105515
- [25] Smith, D. G., Elwy, A. R., Rosen, R. K., Bueno, M., & Sarkar, I. N. (2024). COVID-19 in early 2021: current status and looking forward. *Social Science & Medicine*, 354, 117027. DOI: 10.1016/j.socscimed.2024.117027
- [26] Wang, C., Wang, Z., Lau, J. Y.-N., Zhang, K., Li, W., & Li, W. (2024). A multimodal integration pipeline for accurate diagnosis, pathogen identification, and prognosis prediction of pulmonary infections. *Innovation*, 5(4), 100648. DOI: 10.1016/j.xinn.2024.100648
- [27] Patel, R., & Kumar, R. (2024). A hybrid federated learning framework for secure medical data analysis. *Journal of Medical Systems*, 48(1), 31. DOI: 10.1007/s10916-024-02072-3
- [28] Singh, R., & Sharma, S. (2024). Enhancing data privacy in federated learning systems for healthcare: A survey and future directions. *IEEE Transactions on Cloud Computing*, 12(2), 243-257. DOI: 10.1109/TCC.2024.3147684