

Privacy-Preserving Machine Learning at the Edge: A Comparative Study of Federated Learning, Differential Privacy, and Secure Aggregation

Dr. Pankaj Kumar

Associate Professor in Computer Science, Maharani Padmavati Government College for Women, Shahzadpur (Ambala), Haryana

Abstract - The increasing use of artificial intelligence on smartphones, Internet of Things devices, edge servers, and other distributed platforms has created new opportunities for intelligent applications but has also intensified concerns regarding the privacy of machine-learning data. Conventional centralized machine learning requires sensitive data to be collected and transferred to a central server, creating potential risks of unauthorized access, data leakage, and misuse. Privacy-Preserving Machine Learning (PPML) addresses these concerns by introducing techniques that allow useful models to be trained while reducing exposure of sensitive information. Among the most important approaches for edge environments are Federated Learning (FL), Differential Privacy (DP), and Secure Aggregation (SA). Federated Learning keeps raw training data on participating devices and transfers model updates rather than the original data [1]. Differential Privacy provides a mathematical framework for limiting the privacy loss associated with individual data records or users [2,3]. Secure Aggregation uses cryptographic techniques so that a coordinating server can obtain an aggregate of client updates without directly observing individual contributions [4]. This paper presents a comparative analysis of these three approaches with respect to privacy protection, accuracy, communication overhead, computational requirements, scalability, and suitability for edge computing. Published results indicate that federated learning can reduce communication rounds substantially compared with conventional synchronized training, while practical secure aggregation has demonstrated measurable but manageable communication expansion under representative settings [1,4]. The study further develops a multi-criteria evaluation framework and discusses the advantages of combining FL, DP, and SA rather than relying on a single privacy mechanism. The paper concludes that layered privacy protection offers a stronger foundation for trustworthy edge intelligence, although challenges involving non-IID data, client heterogeneity, privacy-utility trade-offs, malicious participants, communication costs, and deployment complexity remain significant research opportunities.

Keywords: Privacy-Preserving Machine Learning, Federated Learning, Differential Privacy, Secure Aggregation, Edge Computing, Internet of Things, Mobile Computing, Data Privacy, Secure AI, Federated Analytics

1. INTRODUCTION

Machine learning has traditionally relied on centralized data collection. In this architecture, information from multiple users or devices is transferred to a central server where a model is trained. Although centralized training can simplify data management and computation, it can also create a large concentration of sensitive information. Smartphones, wearable devices, healthcare sensors, smart-home systems, and industrial IoT devices continuously generate data that may reveal personal behavior, preferences, locations, health information, or organizational information.

Edge computing provides an alternative architecture by moving computation closer to the location where data is generated. This creates an opportunity to perform machine-learning operations locally while reducing the need to transfer raw data to centralized infrastructure.

Federated Learning is one of the most influential approaches in this area. McMahan et al. proposed Federated Averaging as a practical method in which clients retain local training data and communicate model updates to a coordinating server [1]. Their experiments reported a reduction of approximately 10–100 times in the number of communication rounds compared with synchronized stochastic gradient descent under their evaluated settings [1]. This result is particularly important for mobile and edge environments where network communication can be expensive.

However, Federated Learning by itself does not guarantee complete privacy. Model updates can potentially reveal information about local training data. NIST notes that trained models and federated-learning systems can be exposed to privacy attacks, making additional protections important [5].

Differential Privacy and Secure Aggregation address different parts of this problem. Differential Privacy provides a formal way to quantify and control privacy loss, while Secure Aggregation prevents the server from inspecting individual client updates during aggregation [3,4].

This paper therefore focuses on the following research question:

How do Federated Learning, Differential Privacy, and Secure Aggregation compare as privacy-preserving mechanisms for machine learning at the edge, and what advantages can be obtained by combining them?

2. PRIVACY CHALLENGES IN EDGE MACHINE LEARNING

Edge environments differ considerably from centralized cloud environments. Devices may have limited memory, processing power, battery capacity, and network reliability. At the same time, edge devices may hold highly personal or sensitive information.

A privacy-preserving architecture must therefore address several risks simultaneously. These include unauthorized access to raw data, reconstruction of sensitive information from model updates, malicious clients, inference attacks, model poisoning, communication interception, and accidental disclosure by the aggregation server.

Table 1. Major Privacy and Security Risks in Edge Machine Learning

Risk	Description	Potential Impact
Raw-data exposure	Sensitive information is transferred to centralized infrastructure	High
Gradient leakage	Model updates may reveal information about training examples	High
Membership inference	Attacker attempts to determine whether a person's data participated in training	Medium–High
Model inversion	Attacker attempts to reconstruct information about training data	High
Malicious client	Participant intentionally sends harmful updates	High
Eavesdropping	Communication between device and server is intercepted	Medium–High
Server curiosity	Server attempts to inspect individual client contributions	High
Device compromise	Attacker gains access to a participating edge device	High
Model poisoning	Malicious updates manipulate the global model	High

Table 1 presents that the privacy risks in edge machine learning are not restricted to the storage of raw data. Even when the original dataset never leaves a device, information may potentially be inferred from model updates or the final trained model. NIST specifically highlights that trained models can leak information from their training data [5]. Therefore, a complete privacy architecture requires protection during data generation, local training, communication, aggregation, and model release.

3. FEDERATED LEARNING

Federated Learning changes the traditional machine-learning workflow by keeping training data distributed across clients. Instead of uploading raw datasets, participating devices download a model, perform local training, and return model updates. The server aggregates these updates to produce an improved global model [1].

A simplified federated-learning cycle consists of four stages:

1. The server distributes the current global model.
2. Selected edge devices train the model using local data.
3. Devices send model updates to the server.
4. The server aggregates the updates and produces a new global model.

The process is repeated over several rounds.

Table 2. Conventional Centralized ML versus Federated Learning

Feature	Centralized ML	Federated Learning
Raw data location	Central server	Local devices
Local computation	Limited	High
Server receives raw data	Yes	Normally no

Communication requirement	Data transfer	Model/update transfer
Privacy exposure	Higher	Reduced
Device heterogeneity	Less important	Important
Non-IID data	Usually controlled	Major challenge
Offline/local processing	Limited	Strong potential
Scalability	High with centralized infrastructure	Depends on client participation
Main privacy limitation	Centralized data collection	Update/model leakage

The Table 2 shows the central advantage of Federated Learning is data localization. The raw training data remains on participating devices. However, the exchanged model updates can themselves contain information about local data. Consequently, FL should be regarded as a privacy-enhancing architecture rather than a complete privacy guarantee.

McMahan et al. demonstrated that FL can work with unbalanced and non-IID data distributions and reported 10–100× reductions in communication rounds compared with synchronized SGD in their experiments [1].

4. DIFFERENTIAL PRIVACY

Differential Privacy provides a formal mathematical framework for limiting the amount of information that an algorithm reveals about an individual's participation or data. In machine learning, a common approach is to restrict the influence of individual records or users and then add carefully calibrated random noise.

Abadi et al. introduced a practical method for training deep neural networks with Differential Privacy based on noisy gradient processing and privacy accounting [2]. Differential Privacy can therefore be applied to federated learning to reduce the possibility that an individual user's information can be inferred from the trained model.

NIST's 2025 SP 800-226 provides guidelines for evaluating Differential Privacy guarantees and emphasizes that privacy should be evaluated through formal parameters and careful consideration of implementation risks [3].

Table 3. Main Characteristics of Differential Privacy

Characteristic	Description
Privacy mechanism	Mathematical privacy guarantee
Main technique	Controlled randomization/noise
Privacy parameter	Commonly represented using ϵ and δ
Primary objective	Limit information leakage
Main advantage	Formal privacy guarantee
Main limitation	Noise may reduce utility
Computational effect	Can increase training complexity
Edge suitability	Good when carefully optimized
Compatibility with FL	High
Compatibility with Secure Aggregation	High

The Table 3 presents that the differential Privacy provides something that ordinary data minimization does not: a formal way to quantify privacy loss. However, stronger privacy generally requires greater randomization, which can reduce model accuracy or utility. TensorFlow Federated supports differentially private aggregation using clipping and Gaussian noise, illustrating the practical integration of DP into federated learning [6].

5. SECURE AGGREGATION

Secure Aggregation is a cryptographic mechanism designed to prevent the server from viewing individual client updates. Instead, the server obtains only an aggregate, such as the sum of the updates.

In a federated-learning system, this is important because simply not transmitting raw data does not prevent the server from inspecting model updates. Secure Aggregation addresses this exposure by cryptographically masking individual contributions.

Bonawitz et al. developed a practical secure-aggregation protocol for mobile devices that was designed to tolerate client dropouts and protect client contributions from the server [4].

Table 4. Secure Aggregation Characteristics

Feature	Secure Aggregation
Raw data protection	Yes, when raw data remains local
Individual update visibility to server	Designed to prevent it
Cryptographic protection	Yes
Aggregated result available	Yes
Protection against curious server	Strong
Communication overhead	Additional protocol messages
Computation overhead	Additional cryptographic processing
Dropout handling	Supported by practical protocols
Compatibility with FL	Very High
Compatibility with DP	Very High

The Table 4 presents that the secure Aggregation is complementary to Federated Learning rather than an alternative to it. Federated Learning determines where training data remains, while Secure Aggregation determines what the server can observe about client updates. TensorFlow Federated describes secure aggregation as a mechanism in which the server can obtain a sum while individual updates remain hidden [7].

The original practical protocol reported concrete communication expansion. For 16-bit input values, the authors reported approximately $1.73\times$ communication expansion for 2^{10} users and 2^{20} -dimensional vectors, and approximately $1.98\times$ for 2^{14} users and 2^{24} -dimensional vectors [4]. These figures demonstrate that privacy protection has an infrastructure cost, but also show that the overhead can remain manageable in large-scale settings.

6. COMPARATIVE STUDY OF FL, DP AND SECURE AGGREGATION

Federated Learning, Differential Privacy, and Secure Aggregation solve related but different privacy problems. FL decentralizes the training data, DP limits information leakage, and SA protects individual updates during aggregation.

Table 5. Comparative Evaluation of the Three Techniques

Evaluation Factor	Federated Learning	Differential Privacy	Secure Aggregation
Keeps raw data local	Excellent	Not necessarily	Not necessarily
Formal privacy guarantee	No by itself	Yes	Cryptographic security guarantee
Protects individual updates	Limited	Partially	Strong
Accuracy impact	Usually application-dependent	Possible reduction	Usually low direct effect
Communication overhead	Medium	Low–Medium	Medium–High
Computation overhead	Distributed local training	Noise/clipping/privacy accounting	Cryptographic operations
Handles curious server	Partially	Depends on deployment	Strong
Handles model leakage	Limited	Stronger protection	Does not by itself protect final model
Scalability	High potential	High potential	Requires protocol engineering
Best role	Decentralized training	Privacy guarantee	Hidden aggregation

The table 5 demonstrates that no single method completely solves every privacy problem. FL provides the architectural foundation, SA protects the aggregation process, and DP provides a formal mechanism for limiting information leakage. Their combination can therefore create a layered privacy architecture.

7. EXPERIMENTAL FRAMEWORK

To convert the comparative framework into an empirical research study, the three techniques can be evaluated using a common dataset and model architecture.

A controlled experiment should establish a non-private federated baseline and then progressively add privacy mechanisms. For example, four configurations can be evaluated:

- **Configuration A:** Standard Federated Learning
- **Configuration B:** FL + Differential Privacy
- **Configuration C:** FL + Secure Aggregation
- **Configuration D:** FL + Differential Privacy + Secure Aggregation

The same hardware, dataset partition, optimizer, number of clients, client participation rate, and number of communication rounds should be used wherever possible.

Table 6. Proposed Experimental Design

Parameter	Proposed Setting
Dataset	FEMNIST, CIFAR-10, or another federated benchmark
Number of clients	100–1,000 simulated clients
Client participation	10–100 clients/round depending on experiment
Baseline	Federated Averaging
Privacy methods	DP, SA, and DP+SA
Model	Lightweight CNN or compact neural network
DP measurement	ϵ , δ and utility
Accuracy metric	Accuracy/F1-score
Communication	MB per round
Computation	Training time per round
Energy	Joules per client round where hardware measurement is available
Security evaluation	Privacy leakage and attack success rate
Repetitions	At least 3–5 independent runs

The Table 6 presents the proposed methodology provides a controlled basis for comparison. A researcher can start with simulation and subsequently validate the best configurations on actual edge devices. For stronger publication quality, the experiment should report confidence intervals, standard deviations, client heterogeneity, and non-IID data distributions.

8. DISCUSSION

The comparative study demonstrates that Federated Learning, Differential Privacy, and Secure Aggregation should not be viewed as competing technologies. Their functions are complementary.

Federated Learning primarily addresses data centralization. It enables a shared model to be trained while raw data remains distributed [1]. Differential Privacy addresses information leakage by providing a formal mathematical mechanism for limiting privacy loss [2,3]. Secure Aggregation addresses individual update exposure, preventing the server from directly examining client contributions [4].

The strongest architecture therefore combines all three mechanisms.

However, stronger privacy can introduce additional overhead. Differential Privacy may reduce model utility because of noise, while Secure Aggregation introduces cryptographic computation and communication. The original secure-aggregation study provides evidence that this overhead can be practical, but the exact cost depends on the number of clients and vector dimensions [4].

The appropriate solution should consequently depend on the application. A healthcare system may prioritize privacy more heavily than a low-risk recommendation application. Similarly, a battery-powered IoT sensor may require aggressive communication and computation optimization.

9. CONCLUSION

Privacy-Preserving Machine Learning at the Edge is becoming increasingly important as intelligent applications move from centralized cloud environments toward smartphones, IoT devices, and distributed edge platforms. The central challenge is to obtain the benefits of collaborative machine learning without creating unnecessary exposure of sensitive user data.

This paper compared three major privacy-preserving approaches: Federated Learning, Differential Privacy, and Secure Aggregation. Federated Learning keeps raw training data distributed, Differential Privacy provides a formal framework for limiting information leakage, and Secure Aggregation protects individual model updates during aggregation.

Published research provides quantitative evidence supporting the practical potential of these methods. Federated Averaging reduced communication rounds by approximately 10–100× compared with synchronized SGD in the original experiments [1]. Practical Secure Aggregation reported communication expansion of 1.73× and 1.98× in two representative large-scale configurations [4]. Research on differentially private language models has also demonstrated that useful predictive performance can be achieved under user-level privacy protection in suitable settings [9].

The analysis suggests that no individual privacy mechanism is sufficient for every edge-learning threat. A layered architecture combining FL, DP, and SA can provide stronger protection than relying on Federated Learning alone. Nevertheless, this protection must be balanced against accuracy, communication, computational requirements, energy consumption, and scalability.

Future research should therefore focus on adaptive, energy-aware, secure, and hardware-conscious privacy-preserving federated learning. In particular, the integration of lightweight models, Differential Privacy, Secure Aggregation, robust aggregation, and real-device benchmarking represents a promising direction for building trustworthy edge intelligence.

REFERENCES

- [1] B. McMahan, E. Moore, D. Ramage, S. Hampson, and B. Agüera y Arcas, “Communication-Efficient Learning of Deep Networks from Decentralized Data,” *Proceedings of the 20th International Conference on Artificial Intelligence and Statistics*, PMLR, vol. 54, pp. 1273–1282, 2017.
- [2] M. Abadi, A. Chu, I. Goodfellow, H. B. McMahan, I. Mironov, K. Talwar, and L. Zhang, “Deep Learning with Differential Privacy,” *Proceedings of the ACM Conference on Computer and Communications Security*, pp. 308–318, 2016.
- [3] J. Near, D. Darais, N. Lefkowitz, and G. Howarth, *Guidelines for Evaluating Differential Privacy Guarantees*, NIST Special Publication 800-226, National Institute of Standards and Technology, 2025. NIST describes Differential Privacy as a mathematical framework for quantifying privacy loss and identifies common implementation hazards.
- [4] K. A. Bonawitz et al., “Practical Secure Aggregation for Privacy-Preserving Machine Learning,” *Proceedings of the 2017 ACM SIGSAC Conference on Computer and Communications Security*, 2017. The study presents a practical secure-aggregation protocol and reports concrete communication expansion for large client/vector configurations.
- [5] National Institute of Standards and Technology, “Protecting Trained Models in Privacy-Preserving Federated Learning,” NIST Cybersecurity Insights, 2024. The article discusses privacy attacks against trained models and the role of Differential Privacy in protecting training information.
- [6] TensorFlow Federated, “Differential Privacy in TFF,” TensorFlow Documentation. The documentation demonstrates user-level Differential Privacy in federated learning and explains the privacy-utility trade-off.
- [7] TensorFlow Federated, “Secure Aggregator,” TensorFlow Documentation. The secure aggregator is designed so that the server receives aggregated information while individual client updates remain hidden.
- [8] TensorFlow Federated, “Tuning Recommended Aggregations for Learning,” TensorFlow Documentation. The documentation discusses Differential Privacy, secure aggregation, clipping, compression, and the associated utility and communication considerations.
- [9] B. McMahan, D. Ramage, K. Talwar, and L. Zhang, “Learning Differentially Private Recurrent Language Models,” *International Conference on Learning Representations (ICLR)*, 2018. The study reports user-level Differential Privacy for recurrent language models and discusses the computational and utility implications of privacy protection.
- [10] J. Bell, K. A. Bonawitz, A. Gascon, T. Lepoint, and M. Raykova, “Secure Single-Server Vector Aggregation with (Poly)Logarithmic Overhead,” *Research/Conference Publication*. The work investigates secure aggregation constructions with lower asymptotic communication and computation overhead than earlier approaches.