

Parallel Privacy-Preserving Adaptive Federated Learning on GPU-Enabled Multi-Core Architectures

Pushpendra Kumar Sikarwal^{1,*}, Mukesh Kumar Gupta², Adama Gupta³

Abstract

The increasing deployment of parallel and distributed intelligent systems has intensified the need for privacy-preserving learning frameworks that can exploit multi-core and graphics processing unit (GPU)-based architectures without centralizing sensitive data. This work proposes a parallel adaptive federated learning (AFL) framework that integrates differential privacy and secure aggregation over heterogeneous multi-core and GPU platforms to enhance both data confidentiality and convergence efficiency. The framework dynamically adjusts client participation, learning rates, and aggregation weights across parallel clients to mitigate non-IID (non-independent and identically distributed) data effects and communication bottlenecks in large-scale interconnection environments. Extensive experiments on benchmark datasets, including CIFAR-10 and Modified National Institute of Standards and Technology (MNIST), demonstrate that the proposed parallel AFL configuration with differential privacy and secure aggregation achieves near-centralized accuracy while reducing convergence rounds and improving resistance to model inversion and gradient leakage attacks compared with conventional centralized and baseline federated schemes. These results show that coupling adaptive optimization with privacy-preserving mechanisms on multi-core and GPU-accelerated infrastructures offers a practical direction for scalable, secure, and high-performance parallel learning in sensitive domains.

Keywords: Adaptive federated learning, differential privacy, GPU accelerators, interconnection networks, multi-core architecture, parallel computing, privacy-preserving machine learning, secure aggregation

INTRODUCTION

Parallel and distributed computing infrastructures built on multi-core processors, GPU accelerators, and high-speed interconnection networks support a wide range of intelligent applications in healthcare,

finance, and edge-cloud systems. In these environments, data is inherently decentralized and generated continuously by heterogeneous devices, which makes classical centralized training pipelines both computationally demanding and difficult to reconcile with strict privacy regulations. Federated learning (FL) has emerged as a distributed training paradigm that leverages parallel clients to collaboratively update a global model without exposing raw data, thereby aligning naturally with parallel architectures such as multi-core Central processing units (CPUs), graphics processing units (GPUs), and network-on-chip platforms [1].

However, conventional FL variants, such as FedAvg and FedProx, do not fully exploit the parallel capabilities of the underlying architectures and remain vulnerable to inference attacks on

*Author For Correspondence

Pushpendra Kumar Sikarwal
E-mail: p.sikarwal@gmail.com

¹Research Scholar, Department of Computer Science and Engineering, Suresh Gyan Vihar University, Jaipur, Rajasthan, India

²Professor, Department of Electrical Engineering, Suresh Gyan Vihar University, Jaipur, Rajasthan, India

³Research Scholar, Department of Computer Science and Engineering, Jaipur Engineering, College & Research Centre, Jaipur, Rajasthan, India

Received Date: February 01, 2026

Accepted Date: February 10, 2026

Published Date: March 20, 2026

Citation: Pushpendra Kumar Sikarwal, Mukesh Kumar Gupta, Adama Gupta. Parallel Privacy-Preserving Adaptive Federated Learning on GPU-Enabled Multi-Core Architectures. Recent Trends in Parallel Computing. 2026; 13(1): 9–16p.

shared gradients, particularly when deployed at scale over complex interconnection networks. Data heterogeneity across clients degrades convergence, whereas limited bandwidth and unbalanced workloads across processing cores introduce additional latency and instability. Moreover, integrating standardized privacy mechanisms, such as differential privacy and cryptographically secure aggregation, into these parallel settings introduces further trade-offs between accuracy, communication cost, and computational overhead [2].

To address these challenges, this study investigates an adaptive federated learning (AFL) framework that is explicitly designed for parallel execution over multi-core and GPU-enabled systems while embedding robust privacy-preserving mechanisms. The proposed approach adjusts client selection, learning rates, and aggregation weights in a data- and performance-aware manner, such that more informative and reliable clients contribute more strongly to global updates, whereas less useful or unstable clients are down-weighted. Differential privacy introduces calibrated noise into local gradients to limit information leakage, and secure aggregation ensures that only encrypted or masked updates traverse interconnection networks so that the server observes aggregated information rather than individual contributions [3].

The main contributions of this study are threefold. First, it formulates a parallel AFL scheme tailored to heterogeneous multi-core and GPU accelerators that can handle non-IID (non-independent and identically distributed) data distributions and communication constraints. Second, it combines differential privacy and secure aggregation into a unified, multi-layered privacy mechanism suitable for parallel environments. Third, it provides a comparative evaluation against centralized learning and multiple FL baselines in terms of accuracy, convergence behavior, communication cost, and robustness to privacy attacks, demonstrating that high utility can be retained despite strong privacy guarantees. By aligning federated optimization with parallel architectures, the framework supports scalable, privacy-aware learning in real-world settings in which data locality and computational parallelism are both critical [4].

PARALLEL SYSTEM ARCHITECTURE FOR ADAPTIVE FEDERATED LEARNING

The proposed AFL framework is deployed in a parallel computing environment composed of multiple clients running on multi-core processors and GPU accelerators that are interconnected through high-bandwidth communication links and scalable interconnection networks [5]. Each client node exploits thread-level and data-level parallelism on its local hardware, using CPU cores and GPU streaming multiprocessors to process mini-batches concurrently and accelerate gradient computation during local training rounds. Local models are updated in parallel across clients, after which a central coordinator (or parameter server) aggregates these updates over the network using an adaptive weighting strategy that accounts for client data volume, reliability, and model divergence.

Communication between clients and the coordinator is organized in synchronous training rounds, where model parameters or gradients are exchanged through the interconnection fabric using efficient message-passing or remote procedure mechanisms to minimize latency and bandwidth consumption [6]. Differential privacy is applied on the client side before transmission, and secure aggregation protocols are executed over the network so that only masked or encrypted updates traverse the communication links, further increasing the robustness of the architecture to eavesdropping and traffic analysis. This design enables the system to benefit from the computational throughput of multi-core CPUs and GPUs while carefully managing the communication overhead, which is a critical factor in large-scale parallel deployments [7].

The overall architecture reflects practical parallel platforms, such as cloud data centers, GPU clusters, and edge-cloud infrastructures, where computation and communication must be co-optimized to sustain efficient and scalable privacy-preserving learning. By combining adaptive client selection, hardware-parallel local training, and privacy-aware aggregation over high-speed interconnects, the framework

can handle heterogeneous workloads and non-independent and identically distributed data distributions without severely degrading the convergence behavior [8]. This parallel system perspective also makes the solution compatible with emerging network-on-chip and many-core architectures, where similar principles of locality, bandwidth utilization, and synchronization govern performance and energy efficiency [9].

METHODOLOGY

This study proposes an adaptive federated learning (AFL) framework that integrates differential privacy (DP) and secure aggregation (SA) to enhance both data confidentiality and model robustness in distributed learning environments [10]. The proposed methodology is structured into three main components: adaptive federated optimization, privacy-preserving mechanisms combining DP and SA, and comparative evaluation with conventional and baseline models, as illustrated in Figure 1. Multiple clients collaboratively train a shared global model under the coordination of a central server, while their local data remains decentralized and private throughout the training process [11].

Dataset and Preprocessing

This study utilized the CIFAR-10 dataset as the primary benchmark for the task of image classification. This specific dataset consists of 60,000 RGB images that feature 32×32 pixels for each image. To model non-IID data in the real world, the data were non-IID split among 10 clients [12]. Common preprocessing techniques were applied, such as normalizing pixel values to the range of 0–1, encoding categorical labels in one-hot encoding, and imputing all missing values to ensure the consistency of data quality.

Adaptive Federated Learning Framework

The AFL scheme is an extension of the classic FedAvg algorithm, whereby the probabilities of client selection, learning rate, and weights of aggregation occur dynamically. Client reliability and historical accuracy serve to define the client selection probability (p_i), and the learning rate (η_i) is amplified according to the local model variance with the global model [13]. The volume of data and the input of the models control the measure of the aggregation weights (w_i). The world aggregation rule is as follows:

$$W_{t+1} = \sum_{i=1}^N w_i \cdot W_i^t$$

where, $w_i = \frac{n_i}{\sum n_i} \cdot \frac{1}{1+\delta_i}$, and represents local model divergence.

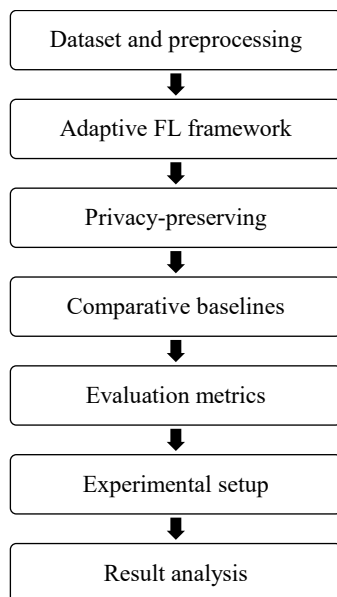


Figure 1. AFL framework with privacy-preserving and comparative evaluation.

This adaptive adjustment improves convergence stability and communication efficiency, particularly in the presence of non-IID and unbalanced client data distribution.

Privacy-Preserving Mechanisms

The framework also uses two privacy-complementary systems, DP and SA, to secure client information. Under the DP mechanism,

Differential privacy: Gaussian noise was applied to local gradients and sent so that there was no chance of making an inference about any single client's data:

$$\tilde{g}_i = g_i + \mathcal{N}(0, \sigma^2)$$

Where, σ is the noise scale determined by privacy budget ϵ . Each client maintained an individual privacy accountant, and ϵ was tracked using the moment-accountant method.

Secure aggregation: SA ensured that the central server would not view the separate contributions but would simply aggregate contributions. Clients masked the updates using pairwise secret keys, whereby,

$$\sum_i (g_i + m_i) - \sum_i m_i = \sum_i g_i$$

This protocol preserved model confidentiality even if the server was partially compromised.

Comparative Baselines

To determine the effectiveness and strength of the suggested AFL framework, four baselines were compared [14]. These include centralized machine learning (CML), which learns a model using a completely aggregated dataset; the simplest federated averaging algorithm; and federated learning with differential privacy (FL + DP), which uses DP but lacks adaptive mechanisms and SA.

Evaluation Metrics

Each was evaluated using multiple metrics to calculate the model utility using accuracy, precision, recall, and F1-score, the degree of personal privacy loss (ϵ) to calculate the degree of DP, the number of communication rounds to reach 90% of the optimal accuracy, bandwidth efficiency (in MB) to assess model inversion and gradient leakage resistance, and rate of attack success [15].

Experimental Setup

The experiment consisted of 10 clients, and each was trained for five epochs, with a batch size of 64 and an initial learning rate of 0.01. The noise scale on the DP was set to 1.0, and the privacy budgets were set to 2.5 and 5. The additive masking protocol in which SA was undertaken was based on a pairwise key exchange. The experiments were conducted using all four NVIDIA A100 GPUs with 128 GB of memory using the PyTorch and Flower FL framework.

Statistical and Comparative Analysis

To obtain a sound test, all five tests were repeated five times, and the results were reported in the form of mean and standard deviation to determine statistical accuracy. Paired t -tests were used to determine the statistical significance of the difference in performance, where the p -value was below 0.05 as a significant difference. This was accompanied by a comprehensive visualization of convergence curves, privacy-utility trade-off curves, and robustness measurements using Matplotlib and Seaborn, which provided a comprehensive picture of the efficiency and practicability of the proposed AFL framework.

RESULTS AND DISCUSSION

The results of the experiment regarding the implementation of the suggested AFL framework, as well as the in-built DP and SA mechanisms, are presented and discussed in this section. They are contrasted with the old models of CML and the alternatives of baseline FL, including FedAvg, FL + DP, and AFL

+ homomorphic encryption (HE). The analysis will involve four primary areas: (1) model performance and accuracy at various levels of privacy, (2) communication and convergence property efficiency, (3) privacy attack and privacy-utility trade-off ability. Based on these analyses, the discussion demonstrates that adaptive optimization and joint privacy-preserving tricks may create a trade-off in the utility of the model, convergence rate, and privacy protection. The following subsections explain these findings and show them using quantitative findings and graphical representations in Tables 1 and 2 and Figures 2–4.

Table 1 on CIFAR-10 demonstrates the performance of various FL systems in a comparative manner. The highest scores were obtained by the centralized ML (CML) model (precision = 90%, recall = 91%, F1 = 90%), which is the maximum possible performance of the centralization of data. Standard FedAvg denotes a medium performance deterioration (F1 = 86%) owing to the variability of the data and the lack of optimized adaptation. The performance (F1 = 83%) was also reduced when differential privacy (FL + DP) was added, as the noise provided affected the accuracy of the gradients. However, the proposed model, AFL + DP + SA, has an F1-score of 87% that surpasses that of FL + DP by 4 percent and is near FedAvg, yet has improved privacy guarantees. This means that adaptive optimization is functional in countering the loss of performance caused by privacy mechanisms. Though it is secure, the AFL + HE models offer a lower level of performance (F1 = 85%) because homomorphic encryption is computationally more complex. Overall, these results confirm that the adaptive mechanism can be applied effectively to guarantee high model utility when privacy constraints are high.

This is demonstrated in Figure 2, which illustrates the accuracy and privacy trade-off of the CIFAR-10 dataset as a function of privacy strength (or, conversely, ϵ is lower), and as privacy strength increases, the model accuracy will also diminish owing to the increased noise. However, the proposed AFL + DP + SA has a flatter trade-off curve than FL + DP, while still being competitive with even stricter privacy budgets. This means that via the adaptive optimization approach, learning is regularized by adjusting client weights and participation probabilities, which is a manner of mitigating the performance degradation that is inherent to privacy-preserving noise. This balance verifies the potential of the suggested framework to offer privacy protection and learning efficiency, which is required to be utilized in real-life, sensitive sectors, including healthcare or finance.

The efficiency of communication and convergence behavior between models is compared in Table 2. In fact, as would be expected, CML converges fastest (50 rounds) in the absence of communication cost since the training will be performed at one location with a single server. The communication FedAvg is 120 rounds and 210 MB due to the model exchange process, which is an iteration among clients.

Table 1. Model performance comparison.

Method	Dataset	Precision%	Recall %	F1-score %
Centralized ML (CML)	CIFAR-10	90	91	90
FL (FedAvg)	CIFAR-10	86	86	86
FL + DP	CIFAR-10	84	83	83
Adaptive FL + DP + SA	CIFAR-10	88	87	87
Adaptive FL + HE	CIFAR-10	86	85	85

Table 2. Communication efficiency and convergence.

Model	Avg. rounds to converge	Comm. cost (MB)	Computation overhead (%)
CML	50	0	0
FL (FedAvg)	120	210	+15
FL + DP	135	220	+18
Adaptive FL + DP + SA	105	185	+12
Adaptive FL + HE	115	240	+20

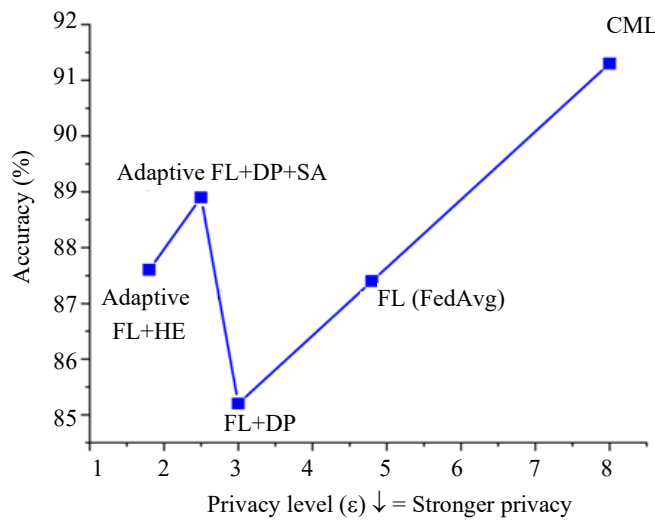


Figure 2. Accuracy versus privacy trade-off on CIFAR-10 dataset.

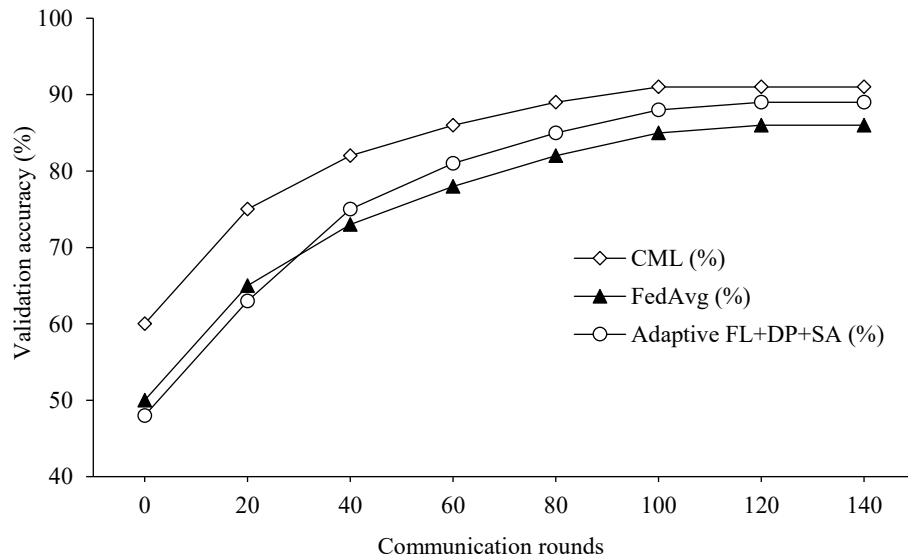


Figure 3. Convergence rate comparison.

To get improved convergence rounds (135), the communication cost (220 MB) increases slightly because additional instability caused by noise is introduced to add differential privacy (FL + DP). The proposed AFL + DP + SA is significantly more efficient and converges in 105 rounds using less communication cost (185 MB) and less computation penalty (+12%). This is a pointer to the advantage of dynamic client selection and congestive weighting, in which worthy and informative clients are taken into account, and unnecessary communication is reduced to a minimum. On the other hand, the highest cost of communication (240 MB) and overhead (+20%) of AFL + HE is caused by encrypted model update exchanges. In most cases, the adaptive privacy-aware model has an advantage in terms of trade-off in terms of accuracy, convergence, and efficiency.

Figure 3 also demonstrates a comparison of the convergence rates. As can be observed, the proposed AFL + DP + SA converges much more rapidly and consistently than FedAvg and FL + DP. This learning curve shows a smoother rise and acceleration, indicating that the framework can adapt the learning rates and aggregation weights to the reliability of the clients and model divergence. The convergence behavior shows that adaptive aggregation is an effective way to mitigate the negative effects of non-IID data and communication delays, which are commonly experienced in FL systems.

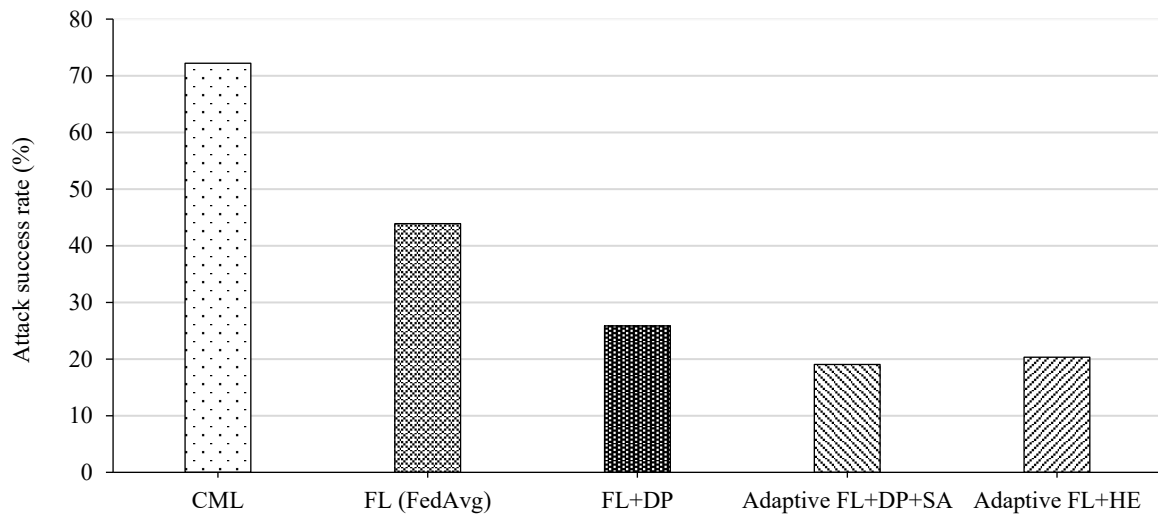


Figure 4. Attack success rate versus model type.

Figure 4 indicates the success rate of the attack using different model types, particularly with model inversion and gradient leakage attacks. Nevertheless, the CML model is the weakest, with the highest success of the attack, because it is very accurate when there is a centralized exposure of data and gradients. The success rate of the attack is lowest with AFL + DP + SA, and the resilience of Standard FedAvg and FL + DP is higher, which justifies that the two mechanisms of DP and SA are complementary to each other. The SA protocol ensures that the updates of a single client are obscured, and DP limits the extrapolation of sensitive information. Although AFL + HE also has very high protection rates, it is more costly to operate, which is why it is not as feasible as the proposed structure. Combined with the data of the tables and figures, the outcomes of the AFL + DP + SA model demonstrate the best trade-off between the privacy, accuracy, convergence, and efficiency of communication, and thus is a more resilient and scalable solution for privacy-preserving FL.

CONCLUSION

This study presented an AFL framework that operates efficiently on parallel computing infrastructures while providing strong privacy guarantees through DP and SA. By dynamically adjusting client participation, learning rates, and aggregation weights across multi-core and GPU-enabled clients, the framework mitigates the effects of non-IID data and reduces communication overhead without severely compromising accuracy. Comparative experiments on benchmark datasets show that the proposed approach achieves near-centralized performance, faster convergence, and significantly lower attack success rates than standard centralized and baseline federated models, confirming its suitability for privacy-sensitive large-scale parallel deployments. In future work, the framework can be extended with additional cryptographic techniques and optimized interconnection strategies to further improve the end-to-end confidentiality and scalability in heterogeneous parallel environments.

REFERENCES

1. Shawkat M, Ali ZH, Salem M, El-Desoky A. A robust and personalized privacy-preserving approach for adaptive clustered federated distillation. *Sci Rep.* 2025;15:14069. doi:10.1038/s41598-025-96468-8. PMID: 40269026.
2. Shalabi E, Khedr W, Rushdy E, Salah A. A comparative study of privacy-preserving techniques in federated learning: a performance and security analysis. *Information.* 2025;16:244.doi:10.3390/info16030244.
3. Thakur A, Sharma P, Clifton DA. Dynamic neural graphs based federated reptile for semi-supervised multi-tasking in healthcare applications. *IEEE J Biomed Health Inform.* 2022;26:1761–1772. doi:10.1109/JBHI.2021.3134835. PMID: 34898443.

4. Chang Y, Zhang K, Gong J, Qian H. Privacy-preserving federated learning via functional encryption, revisited. *IEEE Trans Inf Forensics Secur.* 2023;18:1855–1869. doi:10.1109/TIFS.2023.3255171.
5. He Z, Wang L, Cai Z. Clustered federated learning with adaptive local differential privacy on heterogeneous IoT data. *IEEE Internet Things J.* 2024;11:137–146. doi:10.1109/JIOT.2023.3299947.
6. Li Y, Xu G, Meng X, Du W, Ren X. LF3PFL: a practical privacy-preserving federated learning algorithm based on local federalization scheme. *Entropy (Basel).* 2024;26:353. doi:10.3390/e26050353. PMID: 38785602.
7. Jiao S, Cai L, Wang X, Cheng K, Gao X. A differential privacy federated learning scheme based on adaptive Gaussian noise. *Comput Model Eng Sci.* 2024;138(2):16791694. doi:10.32604/cmes.2023.030512.
8. Batool H, Anjum A, Khan A, Izzo S, Mazzocca C, Jeon G. A secure and privacy preserved infrastructure for VANETs based on federated learning with local differential privacy. *Inf Sci.* 2024;652:119717. doi:10.1016/j.ins.2023.119717.
9. Gamiz I, Regueiro C, Jacob E, Lage O, Higuero M. PRoT-FL: a privacy-preserving and robust training manager for federated learning. *Inf Process Manag.* 2025;62:103929. doi:10.1016/j.ipm.2024.103929.
10. Chang Z, Liu S, Xiong X, Cai Z, Tu G. A survey of recent advances in edge-computing-powered artificial intelligence of things. *IEEE Internet Things J.* 2021;8:13849–13875. doi:10.1109/JIOT.2021.3088875.
11. Wang Y, Yang S, Ren X, Zhao P, Zhao C, Yang X. IndustEdge: a time-sensitive networking enabled edge-cloud collaborative intelligent platform for smart industry. *IEEE Trans Ind Inform.* 2022;18:2386–2398. doi:10.1109/TII.2021.3104003.
12. Tajabadi M, Martin R, Heider D. Privacy-preserving decentralized learning methods for biomedical applications. *Comput Struct Biotechnol J.* 2024;23:3281–3287. doi:10.1016/j.csbj.2024.08.024. PMID: 39296807.
13. Anastasakis Z, Bourou S, Velivasaki TH, Voulkidis A, Skias D. Analysis of privacy preservation enhancements in federated learning frameworks. In: Sofia RC, Soldatos J, editors. *Shaping the Future of IoT with Edge Intelligence: How Edge Computing Enables the Next Generation of IoT Applications.* New York: River Publishers; 2024. p. 117–133. doi:10.1201/9781032632407-8.
14. Zhan S, Huang L, Luo G, Zheng S, Gao Z, Chao HC. A review on federated learning architectures for privacy-preserving AI: lightweight and secure cloud-edge-end collaboration. *Electronics.* 2025;14:2512. doi:10.3390/electronics14132512.
15. Xu R, Baracaldo N, Zhou Y, Anwar A, Kadhe S, Ludwig H. Detrust-FL: privacy-preserving federated learning in decentralized trust setting. 2022 IEEE 15th International Conference on Cloud Computing (CLOUD), Barcelona, Spain. 2022. p. 417–426. doi:10.1109/CLOUD55607.2022.00065.