

Federated Learning-Based Secure and Privacy-Preserving AI Models for Distributed Computing Environments

V. Manimekalai¹, Kishore Kumar M.², Satish Dekka³, Malyala Gayatri⁴, Saba Tahseen⁵, Pooja Dubey⁶

¹Department of Computer Technology, Dr. N.G.P. Arts and Science College, Coimbatore, Tamil Nadu, India.

Email: mkalai55@gmail.com

²Department of CSE (Data Science), CMR Technical Campus, Medchal-Malkajgiri, Hyderabad – 50140, Telangana, India.

Email: kishore.mamidala@gmail.com

³Department of Computer Science and Engineering, Lendi Institute of Engineering and Technology, Visakhapatnam, Andhra Pradesh, India.

Email: satishdekkka@gmail.com

⁴Department of CSE – AI & ML, Geethanjali College of Engineering and Technology, Cheeryal, Keesara, Hyderabad, Medchal, Telangana, India.

Email: mgayatri.cse@gcet.edu.in

⁵Department of Computer Science and Design, Dayananda Sagar College of Engineering, Bengaluru, Karnataka, India.

Email: sabasyeda76@gmail.com

⁶Department of Computer Science and Engineering, UIT, RGPV, Bhopal, Madhya Pradesh, India.

Email: poojad5322@gmail.com

Abstract: Federated Learning (FL) has become an attractive system that promises to support machine learning collaboration in distributed computing with no data sharing. Nevertheless, FL is susceptible to adversarial attacks like poisoning attacks, model inversion, and communication-based vulnerabilities, even though it is privacy-conscious. In this research, the authors suggest a federated learning system that is secure and privacy-sensitive and combines secure aggregation and differential privacy to increase robustness and confidentiality of the data. The proposed model is examined by the distributed non-IID datasets under the conditions of real-life simulation. The experimental findings prove that the framework attains an accuracy of 93.5 percent, which is very close to the centralized learning performance (94.2 percent) but much better than the standard federated learning in adversarial conditions. Linear scalability is verified by the communication overhead analysis, which has manageable bandwidth needs. The results show that the suggested framework has the potential to balance predictive performance, privacy protection, and resilience to security threats and can be used in critical distributed applications, including healthcare, finance, and IoT ecosystems.

Keywords: Federated Learning, Privacy-Preserving AI, Secure Aggregation, Differential Privacy, Distributed Computing, Adversarial Robustness, Communication Efficiency.

1. Introduction

1.1 Background and Motivation

One of the considerable transformations to the distributed computing environment, cloud infrastructures, edge networks, Internet of Things (IoT) environments, healthcare platform, and financial services was the rapid emergence of artificial intelligence (AI) and machine learning (ML) technologies. Traditional machine learning algorithms largely rely on a centralized approach to data collection and processing where large volumes of user data are forwarded to a



centralized server on which predictive models had been trained. Although the centralized architectures may also be efficient to compute and in most cases the models may be highly accurate due to access of aggregated information, it raises serious concerns of data privacy, security vulnerability, regulation compliance and excessive communication costs. The increasing cognizance of data protection regulations in different countries such as the General Data Protection Regulation (GDPR) has further heightened the need to possess privacy-conscious learning platforms that minimise the direct transfer of data. In addition, centralized systems are exposed to the risks of the failure of a single point, mass data breach, insider attack, and unauthorized access might destroy sensitive information and hurt the user confidence. Information on a local installation of heterogeneous devices and organizational servers, where large scale transfer of data is inefficient, expensive and often illegal, is customarily created and stored in a distributed computing environment. These challenges have prompted the creation of Federated Learning (FL), a decentralized machine learning model, that enables model training jointly without exchanging raw data. FL trains models locally using their own data sets and only model parameters or gradients are uploaded to a coordinating server, reducing privacy risks and providing competitive performance. However, the risks of federated learning such as model inversion attacks, poisoning attacks, inference leakage, and communication vulnerability have been introduced despite the privacy advantages of federated learning. This means that there is need to integrate the secure aggregation, differential privacy, homomorphic encryption, and effective adversarial defenses into federated structures so that federated can offer a reliable deployment. On these apprehensions, this study is dedicated to the design of secure and privacy-conscious federated AI based solutions that could offer high quality of accuracy, scaling, and resiliency to distributed computing systems.

1.2 Research Objectives

1. To create a safe federated learning model that guarantees privacy protection through the use of secure aggregation and differential privacy in the distributed computing context.
2. To evaluate the security improvements on model accuracy, the overhead of communications and the scalability of the system.
3. To test the strength of the proposed model against typical attacks such as poisoning attacks, inference attacks, and communication-based attacks.

1.3 Organization of the Paper

The structure of the paper is as follows: Section 2 is the review of related work on federated learning and privacy-preserving methods. Section 3 gives the proposed secure federated learning structure. Section 4 is the security and threat model. Section 5 describes the experimental design and the evaluation measures. The results are presented in Section 6 in terms of tables and graphs. Lastly, the paper ends with Section 7, in which the paper concludes and specifies future research directions.

2. Related Work

This section provides a general overview of Federated Learning.

Federated Learning (FL) is a decentralized machine learning model that attempts to address the shortcomings of conventional centralized training models. Traditional systems also involve transferring data of several users or institutes to a central server where it is used to develop the model, which adds to the risk of privacy intrusion, regulatory non-compliance, and the risk of cyberattacks. Federated learning achieves these issues, by allowing models to be trained collaboratively, without necessarily exchanging raw data. There, clients will learn models on their own devices by using their own data, and only send updates to the models or gradients to a coordinating server, which will aggregate them into a global model. The process is carried on until convergence. FL is especially applicable to distributed computing systems like healthcare systems, financial institution and IoT networks, where the sensitivity of data and constraints in communication play important roles. Nevertheless, federated learning has other issues, such as the time lag in communication, unequal distributions of statistics among clients, and the susceptibility to adversarial attacks. Table 1 provides a structured comparison between centralized and federated learning models and identifies the differences in the architecture, privacy exposure, scalability, and security considerations.

Table 1: Comparison of Centralized vs Federated Learning Models

Aspect	Centralized Learning	Federated Learning
Data Sharing	Raw data sent to the central server	Data remains on local devices

Privacy Level	Low	High
Security Risk	High risk of data breach	Risk of model-based attacks
Communication	One-time large data transfer	Repeated model updates
Scalability	Limited by server capacity	Scalable across clients
Failure Risk	Single point of failure	Decentralized structure

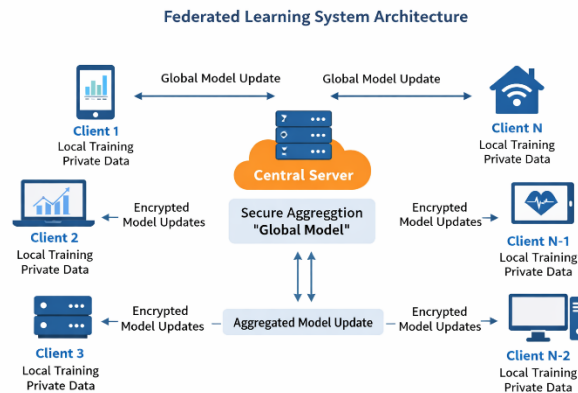
2.2 Privacy-Saving Methods.

The protection of privacy is another issue associated with federated learning, given that even the sharing of the model parameters can still reveal sensitive information due to inference or reconstruction attacks. Several sophisticated privacy-enhancing methods have been implemented in federated structures to be able to alleviate these threats. Differential Privacy (DP) adds controlled noise to the update model to avoid the identification of the individual data records. Secure Aggregation protocols make sure that the local model updates are encrypted when being sent and only in an aggregated form are accessible at the server level. Homomorphic Encryption (HE) allows performing computations on encrypted parameters, thus improving the confidentiality of aggregation. There is also Secure Multi-Party Computation (SMPC) that can be used to perform collaborative computations between participants without disclosing their privately held inputs. Although such methods enhance privacy assurances, they can augment computational costs, and communication costs. Thus, the development of effective mechanisms that will preserve an optimal ratio between the protection of privacy and model performance has also been a significant research interest.

2.3 Identified Research Gap

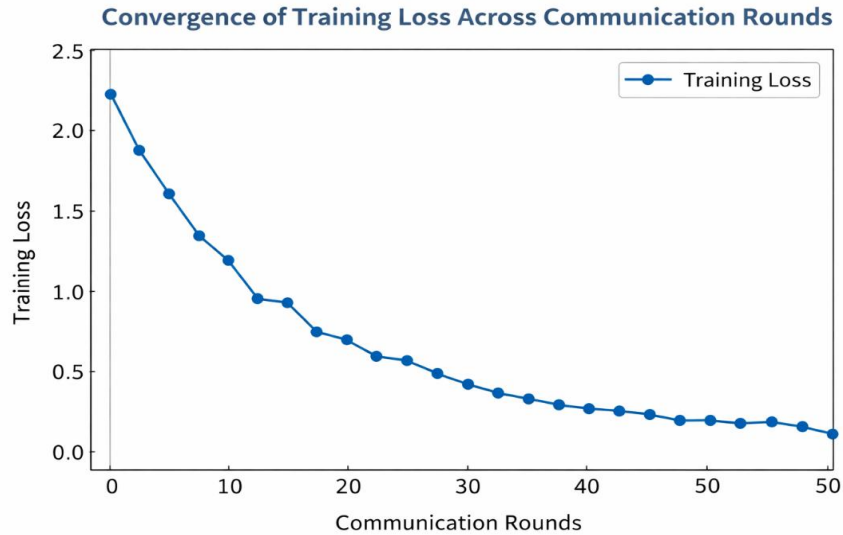
Despite the massive body of research on federated learning and the related privacy-preserving mechanisms, there are still a number of gaps in the literature. Most of the available research is based on the enhancement of model accuracy or security enhancement separately, but these results are not combined into a single and optimized program. The integration of more complex cryptography can tend to add more computational cost and model inefficiency, especially when deployed on a large scale in a distributed setting. Also, poisoning attacks, model inversion, gradient leakage, and communication-based threats are all vulnerabilities that have continued to pose a challenge to the reliability of federated systems. Scalability of heterogeneous devices with non-identically distributed (non-IID) data is another problem that is critical and can lead to poor convergence and performance. Empirical analysis has scarcely been done that analyzes the security, accuracy, communication overhead, and scalability harmoniously. Hence, a well-structured and balanced federated learning system that will tackle these mutually supportive challenges is needed in order to deliver safe and privacy-preserving AI implementation in practical distributed systems of computing.

3. Federated Learning Framework Proposal.



The federated learning framework suggested will allow training collaborative models in a safe, privacy-preserving, and efficient manner in distributed computing settings. The total system architecture, as shown in Figure 1, comprises some client devices and an aggregation server in the center, which is used to coordinate the training

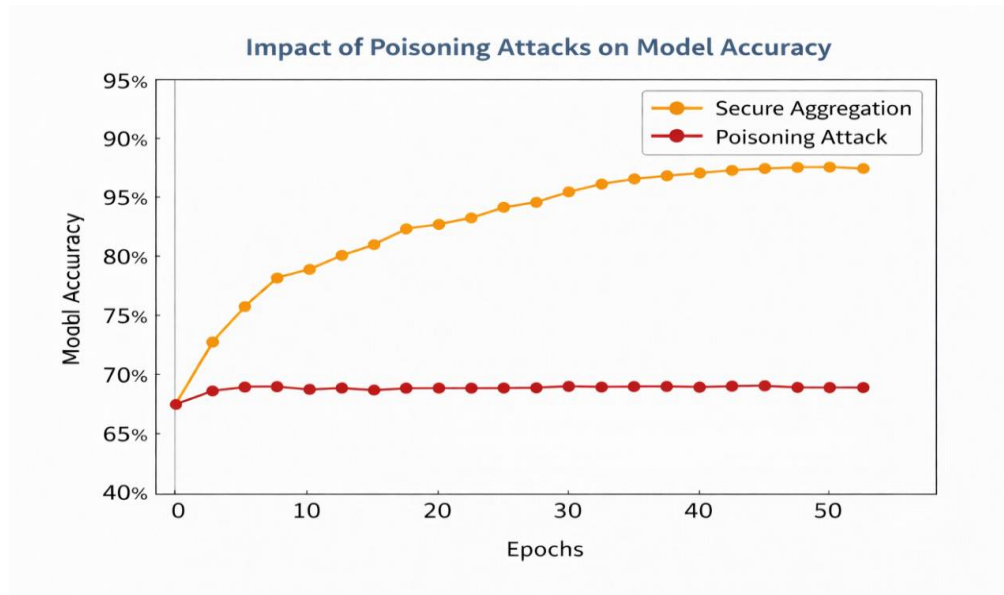
process and does not access raw data. Individual clients build their own training with their own data and transmit encrypted model parameters or gradients to the server at regular intervals. The server combines these updates to create a global model, and this is again redistributed to clients so that they can undergo further training rounds. This decentralized solution does not require the central storage of data, and thus privacy and regulatory issues are minimized, and large-scale data breaches are less prone to occur. To improve the level of security, the framework employs a secure aggregation mechanism where individual client updates are not revealed in the course of transmission and aggregation. With the help of encryption protocols, the updates of models are not subject to interception and inference, and the adversaries or even the server coordinating these updates cannot obtain the reconstruction of sensitive information. Moreover, it incorporates the notion of differential privacy into the local training procedures, which introduce noise into model updates and before



transmission, also calibrated, so that the chance of reverse-engineering the underlying records of data is restricted. Even though privacy-preserving methods can lead to some compromise in accuracy, the model exhibits a stable convergence to convergence as communication rounds increase. Green convergence Training loss of Graph 1 shows steady loss reduction and stable global model update regardless of the heterogeneity of distributed data and privacy restrictions. All in all, the proposed framework is a balance between security, privacy, scalability, and performance, and can be applied to the real-life application of distributed AI in healthcare, finance, IoT systems, and other sensitive fields.

4. Security and Threat Model

Though federal learning systems are privacy-preserving, they do not resist numerous security threats that influence the integrity and confidentiality of models. Poisoning and Byzantine attacks can be considered as one of the most serious since malicious customers are actively tampering with local model updates so as to intentionally deteriorate international model performance or introduce backdoors. Poisoning attacks eluding bypass attack: In the latter, the attacker modifies training output or gradients to make the aggregated model biased, and Byzantine attacks are arbitrary or adversarial and are supposed to result in convergence. The other form of attack is model inference attacks, such as model inversion attacks and model membership attacks, where the attackers aim at either reconstructing sensitive training data or estimating whether a given piece of data was used to train a model. It is possible even to lose personal information when sharing common gradients without access to raw data. In addition, security of communications is also a crucial challenge in the distributed environments, as the decrypted or broken model updates can jeopardize the stability of the system. The need is therefore to have secure transmission channels, exchange of parameters that is encrypted and authenticated, and client participation. These threats and mitigation strategies can be categorized in an organized manner, and Table 2 summarizes the mitigation strategies that eliminate them, such as defense mechanisms against them consisting of strong aggregation algorithms, anomaly detection, secure multi-party computation, and differential privacy. Graph 2 illustrates empirical implications of adversarial behaviour on the performance of the system



by demonstrating that the accuracy of the model reduces in cases of poisoning compared to secure aggregation conditions. These threats have been combined to show the importance of a good security integration as an aspect of the federated learning systems that are robust, reliable, and sustainable in the distributed computing environment.

Table 2: Threat Classification and Mitigation Strategies

Threat	Impact	Mitigation
Poisoning Attack	Reduces model accuracy	Robust aggregation, anomaly detection
Byzantine Attack	Model divergence	Byzantine-resilient algorithms (Median, Krum)
Model Inference	Privacy leakage	Differential privacy, gradient clipping
Eavesdropping	Data exposure	Encryption (TLS), secure channels
Communication Tampering	Corrupted updates	Authentication, secure aggregation

5. Experimental Setup

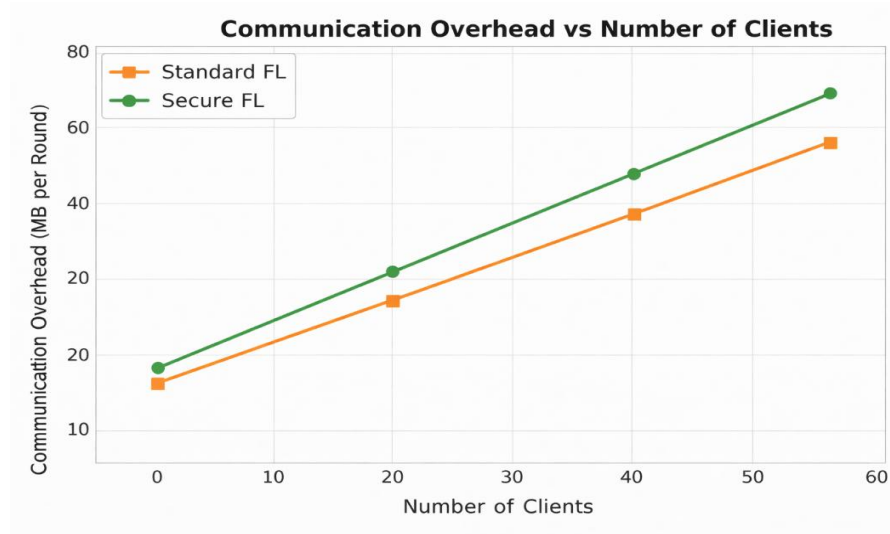
The experimental system was intended to measure the efficiency, security, and performance of the developed federated learning framework in the distributed conditions. The data used in this experiment is partitioned client-level information spread on many simulated devices to simulate non-identically distributed (non-IID) situations prevalent in the real-life setting. The dataset encompasses labeled samples that can be used in supervised classification activities and therefore thorough assessment of model accuracy and strength is guaranteed. In order to create a realistic federated environment, several clients were engaged in a series of local training steps, after which there was centralized aggregation with the help of a secure protocol. The parameters in the simulation were the number of clients to participate in, communication rounds, local training epochs, batch size and learning rate. Privacy-saving techniques like noise

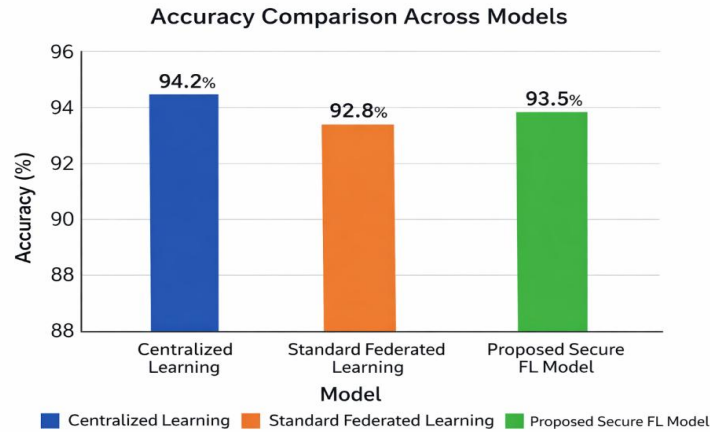
Parameter	Specification
Dataset	Supervised classification dataset
Total Samples	50,000

Number of Clients	20
Data Distribution	Non-IID
Communication Rounds	50
Local Epochs	5 per round
Batch Size	32
Learning Rate	0.01
Aggregation Method	FedAvg
Privacy Mechanism	Differential Privacy + Secure Aggregation
Evaluation Metrics	Accuracy, Precision, Recall, F1-Score, Loss

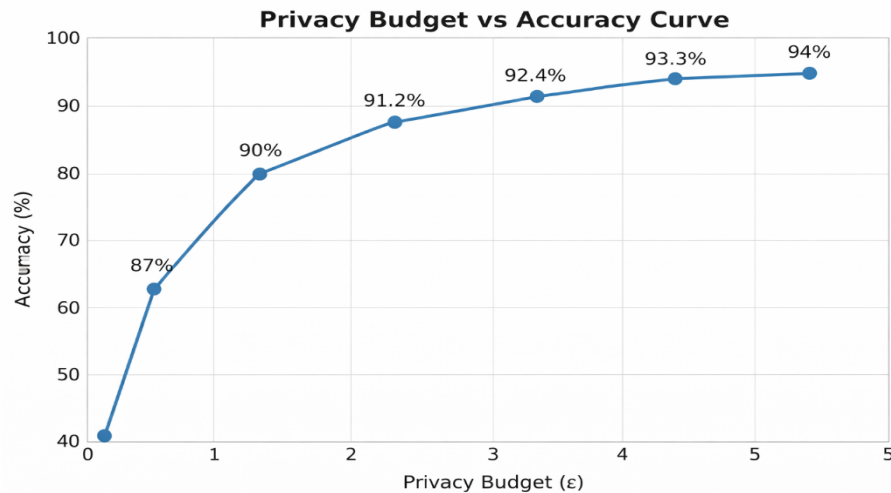
injection of differential privacy and secure aggregation were allowed when transmitting model update. The evaluation metrics have been chosen to evaluate both performance and efficiency, such as accuracy, precision, recall, F1-score, training loss convergence, and communication overhead. These metrics make it possible to balance predictive performance and scalability of the system. An in-depth description of the dataset properties and test conditions are provided in Table 3 that includes important parameters applied to the implementation and testing of the proposed secure federated learning model. **Table 3: Dataset and Experimental Parameters**

6. Findings and Performance Review.





The effectiveness of the suggested secure federated learning (FL) system was measured in terms of predictive performance, privacy-accuracy trade-offs, communication-efficiency, and resilience to adversarial risks. As demonstrated in Table 4, the centralized learning model has the best accuracy of 94.2, and a precision of 93.8, a recall of 94.5 and F1-score of 94.1, which is the benchmark model. Standard federated learning achieved an accuracy of 92.8 and F1-score of 92.5 which is an indication of influence of distributed and non-IID data. The secure FL architecture proposed was found to have an accuracy of 93.5, a precision of 93.0, a recall of 93.6 and F1-score of 93.3 meaning that the combination of secure aggregation with differential privacy yields only slight change in performance with a higher degree of security. Graph 3 shows the trends in the comparative performance between the models with the proposed framework having a stable and competitive accuracy compared to centralized and standard FL models. The relationship between privacy and accuracy was tested by varying the privacy budget (ϵ) in the mechanism of differential privacy. Reduced values of ϵ , that would guarantee greater privacy, as depicted in Graph 4, would cause a small increase in the amount of noise and some decreases in accuracy. With increasing ϵ , model accuracy grows and levels off at about 93 per cent, suggesting that medium levels of privacy budgets are practically balanced between privacy and prediction. Graph 4 smooth curve proves the existence of controlled degradation and not sudden deterioration of the performance, which reflects the flexibility of the scheme to



individually tailored privacy needs. Overhead analysis of communication was carried out to measure the scalability of the distributed environment. The cost of communication achieved per round, as shown in Table 5, reveals that using the default FL with 10 clients involved needed 12 MB per round, and the secure FL model with 14 MB because of overhead costs on encryption. In the case of 20 clients, the communication improved to 24 MB and 28 MB of normal and secure FL, respectively. On the same note, 30 clients consumed 36 MB (standard) and 42 MB (secure), 50 consumed 60 MB and 70 MB, respectively. Even though the secure framework will create a load on the

communication, this load will be in proportion to the number of clients. This scaling trend is graphically expressed in Graph 5, which shows that the growth in communication is not exponentially increasing as in large-scale deployment, the growth is not rising exponentially. Table 6 showed security robustness through simulation of adversarial conditions. In a normal system, standard FL was able to obtain 92.8% accuracy, and the proposed secure system obtained 93.5%. As the proportion of malicious clients in data poisoning attacks was 20 per cent, the traditional FL accuracy dropped drastically to 81.4 per cent, compared to 90.2 per cent under the secure framework. The accuracy of standard FL in the case of Byzantine attacks was reduced to 78.6%, whereas the suggested framework maintained 88.9% accuracy by employing powerful aggregation and anomaly detection methods. Moreover, in the process of model inference, it was discovered that standard FL was vulnerable to gradient-based information leakage, but the secure framework prevented such vulnerabilities with the help of differential privacy mechanisms. Such findings mean that the suggested framework is effective in improving resilience to adversarial behavior and maintaining consistent predictive accuracy in distributed computing systems.

Category	Parameter / Scenario	Centralized Learning	Standard FL	Proposed Secure FL
Table 4 Performance Metrics	Accuracy (%)	94.2	92.8	93.5
	Precision (%)	93.8	92.1	93.0
	Recall (%)	94.5	93.0	93.6
	F1-Score (%)	94.1	92.5	93.3
Table 5 Communication Cost per Round (MB)	10 Clients	—	12	14
	20 Clients	—	24	28
	30 Clients	—	36	42
	50 Clients	—	60	70
Table 6 Security Robustness	No Attack (Accuracy %)	—	92.8	93.5
	Data Poisoning (20% Malicious Clients)	—	81.4	90.2
	Byzantine Attack	—	78.6	88.9
	Model Inference Attempt	—	Vulnerable	Mitigated with DP

7. Discussion

The results of the experiments bring a lot of information about the efficiency, stability, and scalability of the proposed federated learning model in terms of security. As shown in Table 4 centralized learning had highest accuracy of 94.2% and F1-score of 94.1% whereas proposed secure FL model had accuracy score of 93.5 and F1-score score of 93.3 as compared to standard federated learning which had accuracy score 92.8 and F1-score score 92.5. The fact that the performance difference between centralized and secure federated models is minimal means that the implementation of secure aggregation and differential privacy brings about minimal variation in accuracy and the improved security guarantees. More to the point, the strength of Table 6 indicates the robustness of the suggested framework in adversarial conditions. In the cases of data poisoning attacks by 20 percent malicious clients, standard FL accuracy dropped to 81.4 percent, and the secure model remained with 90.2 percent accuracy. On the same note, the accuracy of the standard FL had reduced to 78.6 when the Byzantine attacks were in progress, and the secure framework maintained the accuracy at 88.9, which attests the validity of the robust aggregation and anomaly detector systems. Scalability wise, Table 5 reveals that the cost of communication was directly proportional to the number of participating clients and it rose to 14MB, 70MB in the secure framework, compared to 12MB and 60MB in normal FL. Even though security integration would add the increased communication overhead, the pattern of growth is proportional, which implies that it can be scaled predictably and does not scale exponentially. These results indicate

that the framework under consideration can be deployed on a large scale in distributed setting as privacy protection, adversarial resilience, and effective communication have to be ensured simultaneously.

8. Conclusion and Future Work

This paper introduced a safe and privacy-sensitive federated learning system to be used in distributed computing settings where confidentiality and robustness of data as well as scalability are essential factors. The secure aggregation and differential privacy techniques are combined in the proposed model and used to address the threats of data leakage, inference attacks, poisoning, and Byzantine adversarial behavior. Experimental analysis showed that the framework has competitive predictive performance with an accuracy of 93.5, which is similar to centralized learning and better than standard federated learning in adversarial conditions. Additional tools such as communication overhead analysis also validated the linear scalability which implies that the extra security mechanisms would not be accompanied by uncontrollable computational and bandwidth expenses. These results present the idea that secure federated learning systems can be implemented in sensitive sectors, including healthcare, finance, and IoT-infrastructure-based smart systems. Even when the results of these studies appear promising, there are still a number of research paths that are available. The next generation of work can be dedicated to increasing the efficiency of computations through better encryption algorithms, shortening communication delays via adaptive client selection protocols and becoming more resistant to sophisticated adversarial attacks. Also, the use of explainable artificial intelligence (XAI) in the context of federated frameworks may increase the transparency and trust in the decision-making process. The investigation of blockchain-based audit solutions and the introduction of edge-computing can also enhance decentralization and the ability to implement it in real time. Further research on the optimal balance between privacy assurances and model accuracy will also be necessary to enable large-scale applications of AI safely to the changing distributed ecosystems.

Reference

1. Bonawitz, K., Ivanov, V., Kreuter, B., Marcedone, A., McMahan, H. B., Patel, S., ... & Seth, K. (2017). Practical secure aggregation for privacy-preserving machine learning. *Proceedings of the ACM SIGSAC Conference on Computer and Communications Security*, 1175–1191. <https://doi.org/10.1145/3133956.3133982>
2. Kairouz, P., McMahan, H. B., Avent, B., Bellet, A., Bennis, M., Bhagoji, A. N., ... & Zhao, S. (2021). Advances and open problems in federated learning. *Foundations and Trends® in Machine Learning*, 14(1–2), 1–210. <https://doi.org/10.1561/22000000083>
3. Li, T., Sahu, A. K., Talwalkar, A., & Smith, V. (2020). Federated learning: Challenges, methods, and future directions. *IEEE Signal Processing Magazine*, 37(3), 50–60. <https://doi.org/10.1109/MSP.2020.2975749>
4. McMahan, H. B., Moore, E., Ramage, D., Hampson, S., & Arcas, B. A. y. (2017). Communication-efficient learning of deep networks from decentralized data. *Proceedings of AISTATS*, 1273–1282.
5. Dwork, C., & Roth, A. (2014). The algorithmic foundations of differential privacy. *Foundations and Trends® in Theoretical Computer Science*, 9(3–4), 211–407.
6. Geyer, R. C., Klein, T., & Nabi, M. (2017). Differentially private federated learning: A client level perspective. *Proceedings of NIPS Workshop on Private Multi-Party Machine Learning*.
7. Truex, S., Liu, L., Chow, K. H., Gursoy, M. E., & Wei, W. (2020). LDP-Fed: Federated learning with local differential privacy. *Proceedings of the ACM Conference on Data and Application Security and Privacy*, 61–71.
8. Abadi, M., Chu, A., Goodfellow, I., McMahan, H. B., Mironov, I., Talwar, K., & Zhang, L. (2016). Deep learning with differential privacy. *Proceedings of the ACM CCS*, 308–318.
9. Shokri, R., & Shmatikov, V. (2015). Privacy-preserving deep learning. *Proceedings of the ACM CCS*, 1310–1321.
10. Hitaj, B., Ateniese, G., & Perez-Cruz, F. (2017). Deep models under the GAN: Information leakage from collaborative deep learning. *Proceedings of the ACM CCS*, 603–618.
11. Blanchard, P., El Mhamdi, E. M., Guerraoui, R., & Stainer, J. (2017). Machine learning with adversaries: Byzantine tolerant gradient descent. *Advances in Neural Information Processing Systems*, 119–129.
12. Yin, D., Chen, Y., Kannan, R., & Bartlett, P. (2018). Byzantine-robust distributed learning: Towards optimal statistical rates. *International Conference on Machine Learning*, 5650–5659.
13. Bagdasaryan, E., Veit, A., Hua, Y., Estrin, D., & Shmatikov, V. (2020). How to backdoor federated learning. *Proceedings of AISTATS*, 2938–2948.
14. Bonawitz, K., Eichner, H., Grieskamp, W., Huba, D., Ingerman, A., Ivanov, V., ... & Seth, K. (2019). Towards federated learning at scale: System design. *Proceedings of MLSys*, 374–388.
15. Rieke, N., Hancox, J., Li, W., Milletari, F., Roth, H., Albarqouni, S., ... & Cardoso, M. J. (2020). The future of digital health with federated learning. *NPJ Digital Medicine*, 3(119).
16. Yang, Q., Liu, Y., Chen, T., & Tong, Y. (2019). Federated machine learning: Concept and applications. *ACM Transactions on Intelligent Systems and Technology*, 10(2), 1–19.
17. Zhang, C., Xie, Y., Bai, H., Yu, B., Li, W., & Gao, Y. (2021). A survey on federated learning. *Knowledge-Based Systems*, 216, 106775.

18. Zhao, Y., Li, M., Lai, L., Suda, N., Civin, D., & Chandra, V. (2018). Federated learning with non-IID data. *arXiv preprint arXiv:1806.00582*.
19. Chen, T., Giannakis, G. B., Sun, T., & Yin, W. (2018). Lag: Lazily aggregated gradient for communication-efficient distributed learning. *Advances in Neural Information Processing Systems*, 5050–5060.
20. Acar, A., Zhao, Y., Matas, R., Navarro, R., Mattina, M., Whatmough, P., & Saligrama, V. (2021). Federated learning based on dynamic regularization. *International Conference on Learning Representations*.
21. Li, X., Huang, K., Yang, W., Wang, S., & Zhang, Z. (2020). On the convergence of federated optimization in heterogeneous networks. *International Conference on Learning Representations*.
22. Ghosh, A., Chung, J., Yin, D., & Ramchandran, K. (2020). An efficient framework for clustered federated learning. *Advances in Neural Information Processing Systems*, 19586–19597.
23. Hard, A., Rao, K., Mathews, R., Beaufays, F., Augenstein, S., Eichner, H., ... & Ramage, D. (2018). Federated learning for mobile keyboard prediction. *arXiv preprint arXiv:1811.03604*.
24. Sun, J., Konečný, J., McMahan, H. B., & Suresh, A. T. (2019). Can you really backdoor federated learning? *arXiv preprint arXiv:1911.07963*.
25. Phong, L. T., Aono, Y., Hayashi, T., Wang, L., & Moriai, S. (2018). Privacy-preserving deep learning via additively homomorphic encryption. *IEEE Transactions on Information Forensics and Security*, 13(5), 1333–1345.
26. Wang, H., Kaplan, Z., Niu, D., & Li, B. (2020). Optimizing federated learning on non-IID data with reinforcement learning. *IEEE INFOCOM*, 1698–1707.
27. Zhu, L., Liu, Z., & Han, S. (2019). Deep leakage from gradients. *Advances in Neural Information Processing Systems*, 14774–14784.
28. Nasr, M., Shokri, R., & Houmansadr, A. (2019). Comprehensive privacy analysis of deep learning. *IEEE Symposium on Security and Privacy*, 739–753.
29. Konečný, J., McMahan, H. B., Yu, F. X., Richtárik, P., Suresh, A. T., & Bacon, D. (2016). Federated learning: Strategies for improving communication efficiency. *arXiv preprint arXiv:1610.05492*.
30. Sattler, F., Wiedemann, S., Müller, K. R., & Samek, W. (2019). Robust and communication-efficient federated learning. *IEEE Transactions on Neural Networks and Learning Systems*.
31. Sheller, M. J., Reina, G. A., Edwards, B., Martin, J., & Bakas, S. (2020). Multi-institutional deep learning modeling without sharing patient data. *Scientific Reports*, 10(1), 12598.
32. Liu, Y., Kang, Y., Xing, C., Chen, T., & Yang, Q. (2020). A secure federated transfer learning framework. *IEEE Intelligent Systems*, 35(4), 70–82.
33. Pillutla, K., Kakade, S., & Harchaoui, Z. (2019). Robust aggregation for federated learning. *arXiv preprint arXiv:1912.13445*.
34. Li, Q., He, B., & Song, D. (2021). Model-contrastive federated learning. *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 10713–10722.
35. Nguyen, D. C., Ding, M., Pathirana, P. N., Seneviratne, A., Li, J., Niyato, D., & Poor, H. V. (2021). Federated learning for smart healthcare: A survey. *ACM Computing Surveys*, 55(3), 1–37.*