

Evaluating Data-Protection Mechanisms in Federated Learning: A Performance-Centric Analysis

Mitende Nicholas Nyapete, Richard Omolo, Newton Masinde

Department of Computer Science and Software Engineering, Jaramogi Oginga Odinga University of Science and Technology, Bondo, Kenya

Email: mnyapete@gmail.com, comolo@gmail.com, amphinewt@gmail.com

How to cite this paper: Nyapete, M.N., Omolo, R. and Masinde, N. (2026) Evaluating Data-Protection Mechanisms in Federated Learning: A Performance-Centric Analysis. *Journal of Information Security*, 17, 393-405.

<https://doi.org/10.4236/jis.2026.174018>

Received: July 11, 2026

Accepted: August 11, 2026

Published: August 14, 2026

Copyright © 2026 by author(s) and Scientific Research Publishing Inc. This work is licensed under the Creative Commons Attribution International License (CC BY 4.0).

<http://creativecommons.org/licenses/by/4.0/>



Open Access

Abstract

Background: Federated learning enables multiple participants to train ML models collaboratively without exposing or sharing raw data, making FL attractive approach for privacy-sensible applications. To reduce residual leakage risks in shared model updates, studies have proposed several data-protection mechanisms, including differential privacy, homomorphic encryption, secure multi-party computation, secure aggregation protocols, federated averaging with secure aggregation, model distillation, regularization, and client-side anonymization. While these mechanisms are regularly evaluated for the strength of guarantees they provide, their performance implications: computational overhead, communication cost, model accuracy, and scalability are comparatively under-examined as a harmonized basis for practical selection. **Method:** The study conducted a systematic comparative analysis of eight data-protection mechanisms in federated learning, evaluating them using performance metrics. Drawing on a structured analysis of the most recent innovations, the study developed a comparison framework, applied it consistently across mechanisms, and presented the resulting trade-offs in tabular and narrative form. **Results:** The study found that lightweight mechanisms: secure aggregation, model distillation, and regularization techniques offer the most favourable performance profiles for resource-constrained, large-scale deployments, while the cryptographic approach: homomorphic encryption and secure multi-party computation impose substantial computational and communication burdens that limit technique scalability. **Conclusion:** The study concludes that, with practical guidance for mechanism selection based on deployment context, it outlines directions for future performance-oriented research in federated learning.

Keywords

Differential Privacy, Homomorphic Encryption, Secure Multi-Party Computation, Secure Aggregation, Federated Averaging, Model Distillation, Data Anonymization

1. Introduction

Federated Learning has recently been adopted in decentralised machine learning to enable multiple participants to collaboratively train a shared model without revealing or transferring their raw data to a central location. By keeping data local and exchanging only model updates or gradients, federated learning reduces many of the data-exposure risks associated with centralised training. However, federated learning is not inherently immune to leakage: adversaries can exploit shared updates or gradients to infer sensitive information about the underlying data. To close this gap, researchers have proposed a variety of data-protection mechanisms, including differential privacy, homomorphic encryption, secure multi-party computation, secure aggregation protocols, federated averaging with secure aggregation, model distillation, regularisation, and client-side data anonymisation.

These mechanisms address data exposure and entail distinct operational costs. Homomorphic encryption and secure multi-party computation have offered strong guarantees at the expense of computation and communication, thereby impacting the performance of resource-constrained devices such as smartphones and IoT devices. Statistical approaches such as differential privacy are lightweight but introduce a trade-off between the strength of protection and model utility. Secure aggregation, model distillation, and regularisation remain lightweight with their own limitations regarding dropout handling, teacher-model dependency, and the risk of underfitting.

A large body of literature evaluates these mechanisms based on the strength and formality of the protection they provide. However, less emphasis on systematic comparison on performance metrics: computational overhead, communication cost, impact on model accuracy, and scalability they impose on federated learning system that is directly useful for practitioners selecting mechanism for a given deployment setting, the paper proposed to address this gap by conducting a comparative analysis of the performance of existing data protection techniques and propose mechanism that can be adopted in order to improve the data protection performance in federated learning systems. The analysis was guided by the following questions: 1) How commonly are data protection mechanisms used in federated learning systems compared in terms of computational overhead? 2) How do existing mechanisms compare in terms of communication overhead? 3) What is the impact of each mechanism on model accuracy and utility? Moreover, 4) How well does each mechanism scale to large, heterogeneous, and resource-con-

strained federated learning deployments?

2. Methods

The study adopted a systematic comparative analysis of the literature; the approach was chosen because the mechanisms under review span heterogeneous implementations, hardware assumptions, and federated learning frameworks, making a single controlled experimental comparison across all the data protection mechanisms within the study's scope impractical. The comparative approach synthesizes performance-relevant findings reported across studies and carefully evaluates each mechanism, considering the common performance criteria.

2.1. Threat Model and Protection Objective

The protection objective is to prevent adversaries from interfering with the client raw data training during the federated learning, model updates, gradients, and aggregated parameters. Three threat actors were considered in this review: 1) an honest-but-curious central server, the actor follows protocol correctly and attempts to infer client data from received updates; 2) an external eavesdropper which intercepts client-server communication; and 3) a malicious client attempting to corrupt and infer information about participants. The attacks relevant to "protection" were update and gradient leakage, model inversion, and membership inference, which determines whether a sample belongs to the client's training set. Poisoning and backdoor attacks were referenced whenever they intersect with a mechanism's confidentiality guarantees.

2.2. Mechanism Selection

A total of eight data-protection mechanisms were selected for analysis based on their frequent use in federated learning and their conceptual distinctness. Mechanisms analysed were differential privacy, homomorphic encryption, secure multi-party computation, secure aggregation protocols, FedAvg with secure aggregation, model distillation, regularization techniques, and client-side data anonymization. Mechanisms were selected because they represent major categories of data-protection approaches used in federated learning, such as statistical approaches (differential privacy), cryptographic approaches (homomorphic encryption, secure multi-party computation, secure aggregation), architectural/algorithmic approaches (FedAvg, model distillation, regularization), and data-level approaches (client-side anonymization). FedAvg, model distillation, and regularization are optimization methods: developed to improve convergence, communication efficiency, and generalization, while the methods can incidentally change the content of the client information being transmitted, none provide a formal confidentiality guarantee comparable to differential privacy's (ϵ , σ) bound or cryptographic guarantees underlying homomorphic encryption, secure multi-party computation, and secure aggregation. **Table 1** summarise classification of the mechanisms.

Table 1. Mechanism classification by protection type.

Mechanism	Category	Protection Guarantee	Design Purpose
Differential Privacy	Data protection (statistical)	Yes— (ϵ, σ) bound	Bound information leakage via calibrated noise
Homomorphic Encryption	Data protection (cryptographic)	Yes—computational hardness	Computation on encrypted data (Not exposing raw data)
Secure Multi-Party Computation	Data protection (cryptographic)	Yes—computational hardness/information-theoretic	Joint computation without revealing individual inputs
Secure Aggregation	Data protection (cryptographic/architectural)	Yes, conditional on protocol and non-collusion assumptions	Prevent server from observing individual plaintext updates
Client-Side Data Anonymization	Data protection (data-level)	No formal guarantee; heuristic	Obscure identifying features in data prior to training
FedAvg	Training/optimization method (baseline)	None	Efficient aggregation of local model updates
Model Distillation	Training/optimization method	None	Transfer knowledge from teacher to student model; reduce communication
Regularization	Training/optimization method	None	Improve generalization, reduce overfitting

2.3. Literature Sources

Sources that address data protection were identified through a literature search across Google Scholar, Semantic Scholar, PubMed, Springer Nature, Research Gate, ScienceDirect, IEEE, Scilit, ACM digital library, Wiley online library, SciSpace, National foundation (.gov), HAL open science database, open Ukrainian citation index, open review, Iniria, and nature.com from inception to March 2026. These sources were selected based on their relevance to the operational behaviour of each mechanism, with greater emphasis on the reported effects on computation time, communication overhead, model accuracy, and behaviour under scale and/or heterogeneity. Search terms combined the federated learning with mechanism-specific term and a performance-related term described in the pattern: (“federated learning”) AND (“overhead” OR “performance” OR “scalability” OR “communication cost” OR “computational cost”) The eight mechanism-specific term sets used were: “differential privacy”; “homomorphic encryption”; “secure multi-party computation” OR “SMPC”; “secure aggregation”; “FedAvg” OR “federated averaging”; “model distillation” OR “knowledge distillation”; “regularization”; and “data anonymization” OR “client-side anonymization.”

Inclusion criteria: 1) Source that address federated learning explicitly; 2) source that reports at least one of the four performance dimensions (computational overhead, communication overhead, accuracy/utility impact, or scalability) for one of the eight mechanisms, either empirically or through direct technical/theoretical analysis; 3) the source was a peer-reviewed journal or conference article, or an arXiv preprint reporting a complete technical contribution rather than a work-in-progress abstract.

Exclusion criteria: Articles addressing only federated learning without a reported performance evaluation were excluded.

Titles and abstracts for articles and reports were screened, followed by full-text eligibility assessment. A total of 100 records were identified across the four mechanism-and-performance search patterns; after de-duplication and eligibility screening, 38 unique sources met the inclusion criteria and were cited (**Table 2**). This is a targeted, illustrative synthesis of the literature and should be read as representative of reported performance patterns, not as a complete census of the literature.

2.4. Comparison Framework

Each mechanism was evaluated against four performance metrics identified to reflect the practical issues that are most relevant to federated learning deployment. The performance metrics considered were:

- 1) Computational overhead, which is the additional processing burden the mechanism places on client devices and/or the aggregation server.
- 2) Communication overhead, which is the additional bandwidth/number of communications rounds the mechanism requires relative to unprotected federated learning.
- 3) Impact on model accuracy or utility, which is the degree to which the mechanism degrades the predictive performance of the trained model.
- 4) Scalability, which describes how the mechanism's overhead behaves as the number of participating clients grows and as client capability and connectivity become more heterogeneous.

Reported findings across sources were synthesized into a qualitative rating (low, moderate, high, or very high overhead; poor, moderate, or good Scalability) and avoiding the use of a single quantitative benchmark figure since the underlying studies use different models, datasets, and hardware, making it not directly numerically comparable. The qualitative synthesis was presented in (**Table 2**).

2.5. Rating and Coding Rule

Each mechanism was described with a qualitative overhead rating for **Table 2** using a fixed coding rule applied consistently across the four performance dimensions.

- 1) Computational and communication overhead ratings
 - i) Low: No additional cryptographic computation or communication is required beyond FedAvg aggregation. Synthesized sources reported negligible overhead under approximately 10% relative to baseline.
 - ii) Moderate: No superlinear increase (bounded) in computation and communication (noise calibration, masking, and key exchange) with client count or model size. Synthesized sources reported overhead in the range of 10% - 50% relative to baseline or qualitatively described as “moderate” or “manageable”.
 - iii) High: Linear increase of intensive operations with the number of partici-

pants. Synthesized sources reported an overhead of 50% relative to baseline, and others described it as a primary practical deployment barrier.

iv) Very High: Prohibitive overhead for real-time or resource-constrained deployment in the majority of the sources analyzed for that mechanism (due to large ciphertext expansion factors in homomorphic encryption).

2) Scalability rating

i) Poor: overhead increases with client count, and the mechanism also requires a trusted central coordinator, which can become a bottleneck as the number of participants increases.

ii) Moderate: overhead grows linearly and is manageable at the cross-setting scale, but is not extensively validated at the cross-device scale in the cited sources.

iii) Good: overhead reported as sub-linear or constant in client count, or explicitly validated at a scale exceeding 100 simulated clients in the cited sources.

Resolving Conflicting Findings

Where the source analyzed reported different overhead levels for the same mechanism under different conditions (for instance in secure aggregation's communication overhead varying with client dropout rate), **Table 2** rating reflects the overhead reported under typical, stable operating conditions, low-to-moderate dropout, cross-setting, or moderately sized cross-device settings as described by the majority of cited sources for that mechanism. Conditions under which overhead diverges substantially from this typical case (for instance, high client churn, very large ciphertexts) were described in Sections 3 and 4 so that rating can reflect a central tendency rather than a best or worst-case figure.

3. Results

This section presents a summarized comparative performance profile of the eight mechanisms across the four-evaluation metrics: computation overhead, communication overhead, accuracy and utility, and scalability of the mechanisms.

3.1. Computational Overhead

The literature analysis shows a clear separation between cryptographic and non-cryptographic mechanisms in terms of computational cost. Homomorphic encryption was found to impose the highest computational burden among the mechanisms reviewed in the study due to the complexity of performing arithmetic directly on ciphertext, which is limiting on resource-constrained client devices, including smartphones and IoT devices. Secure multi-party computation incurs high computational costs due to secret sharing and garbled-circuit operations. Differential privacy, secure aggregation protocols, FedAvg, model distillation, regularization, and client-side anonymization have low-to-moderate computational overhead given that their operation of noise addition, masking, local averaging, knowledge transfer, penalty terms, and data generalization, respectively, are lightweight, which does not require many rounds of computations (**Table 2**).

3.2. Communication Overhead

Communication costs across the data protection mechanisms show a largely common pattern: homomorphic encryption, for instance, ranks highest when encrypted payloads are larger than plaintext model updates. Secure multi-party computation incurs high communication costs due to multiple rounds of interaction among participants. In contrast, secure aggregation protocols and FedAvg have a small communication overhead, which is attributed to key exchange and masking rounds, and/or to the size and frequency of model update transmissions. On the other hand, client-side anonymization, regularization, and model distillation have a lower communication overhead among the mechanisms reviewed, since they either transmit smaller artifacts (distilled knowledge, sparse updates) or incur no additional transmission cost (Table 2).

3.3. Impact on Model Accuracy and Utility

Regarding model accuracy and utility, data protection mechanisms show greater divergence, where differential privacy has a direct, well-documented utility trade-off with smaller privacy budgets (ϵ , σ), providing stronger data protection but degrading model performance, particularly on small and/or non-independent and identically distributed (non-IID) datasets. On client-side anonymization, adding noise and generalization obscure data patterns that are important for learning. Homomorphic encryption and secure multi-party computation have minimal direct effect on model accuracy, given that they preserve exact computation over protected data without approximation. Model distillation's effect on accuracy depends on the quality of the teacher model, and regularization can improve generalization and, in some instances, cause underfitting if applied aggressively (Table 2).

3.4. Scalability

Table 2. Comparative performance of data-protection mechanisms in federated learning.

Mechanism	Computational Overhead	Communication Overhead	Impact on Model Accuracy/Utility	Scalability
Differential Privacy (DP) [1]-[7]	Low-moderate (noise injection, gradient clipping)	Moderate (noisy updates, clipping)	Degrades with smaller privacy budget (ϵ); worse on small/non-IID datasets	Moderate; harder to tune uniformly across heterogeneous clients
Homomorphic Encryption (HE) [8]-[12]	Very high (complex cryptographic operations)	Very high (ciphertext far larger than plaintext)	Minimal direct accuracy loss, but limited to supported operations	Poor in decentralized settings; key management does not scale well
Secure Multi-Party Computation (SMPC) [5] [13]-[15]	High (secret sharing, garbled circuits)	High (frequent multi-round interactions)	Minimal direct accuracy loss	Poor; overhead grows exponentially with participants
Secure Aggregation Protocols [10] [16]-[18]	Moderate (masking/key exchange)	Moderate (multi-round key sharing/unmasking)	Minimal accuracy loss under stable participation	Good for large-scale FL, but degrades with dropout and heterogeneity
Federated Averaging with Secure Aggregation (FedAvg) [19]-[24]	Low on clients (local computation)	Moderate-high (frequent large model transmissions)	Robust to non-IID data; No confidentiality guarantee	Good; widely used baseline

Continued

Model Distillation [25]-[31]	Low on client (smaller student model)	Low (smaller updates transmitted)	Dependent on teacher model quality; can lose nuance	Good; well suited to resource-constrained devices
Regularization Techniques [32]-[34]	Low-moderate (extra tuning cost)	Low (can shrink update size, for instance L1 sparsity)	Improves generalization but risks underfitting if over-applied	Good; lightweight to deploy
Client-Side Data Anonymization [35]	Low-moderate (masking/generalization)	Low	Can obscure useful data patterns, reducing accuracy	Moderate; local computation strains weaker devices

Analysis shows that scalability varies widely across all data protection mechanisms. Secure multi-party computation and homomorphic encryption were found to scale poorly. Overhead grows substantially, and in secure multi-party computation grows exponentially as the number of participants increases, making decentralized key management increasingly difficult without a central trusted authority. Secure aggregation protocols scale favourably; however, their sensitivity remains to client dropout and device heterogeneity. FedAvg, model distillation, and regularization are reported as the most scalable of the mechanisms analysed, which is consistent with their design that emphasizes lightweight, resource-conscious operation, which is suited to large and heterogeneous client populations (Table 2).

4. Discussion

The results points to a consistent pattern: mechanisms that provide the strong formal or cryptographic protection: homomorphic encryption, and secure multi-party computation, which tend to impose great performance cost, while lightweight architectural and statistical mechanisms: secure aggregation, FedAvg, model distillation, and regularization offer more favourable performance profiles but achieve protection through different, generally less exhaustive, means [36]. This is not a new observation in isolation, but framing it explicitly along four performance dimensions clarifies where the trade-offs are sharpest and where they are more manageable.

4.1. Cryptographic Mechanisms: Strong Guarantees, High Cost

Homomorphic encryption and secure multi-party computation both emerge from this analysis as the least performance-friendly mechanisms, principally because their protection guarantees depend on computationally expensive operations encrypted arithmetic in the case of homomorphic encryption, and secret-sharing protocols in the case of secure multi-party computation. Both also face acute scalability problems in decentralized FL settings, where there is no single trusted authority to coordinate key management. These findings suggest that, from a purely performance-oriented standpoint, cryptographic mechanisms are best suited to smaller-scale, higher-value deployments for example, cross-silo FL among a small number of well-resourced institutional participants rather than large-scale, cross-device deployments involving thousands of resource-constrained clients.

4.2. Statistical and Architectural Mechanisms: Lower Cost, Different Trade-Offs

Differential privacy has comparatively low computational and communication overhead, making the mechanism attractive. It introduces a direct, tunable trade-off against model utility. This makes practitioners scruple over the difficulty of tuning the privacy budget correctly under non-IID data distributions, where the noise disproportionately affects clients with smaller datasets or, in some instances, an imbalanced dataset. Secure aggregation protocols, FedAvg, model distillation, and regularization represent the most performance-favourable cluster of mechanisms analysed. Secure aggregation achieves meaningful protection at a fraction of the computational and communication cost of full cryptographic methods: homomorphic encryption and secure multi-party computation, but remains sensitive to client dropout, which is a challenge in cross-device federated learning. Model distillation and regularization introduce small overhead across the mechanisms analysed in this study and directly reduce communication costs by producing smaller distilled models or sparser regularized updates. However, they depend on auxiliary quality factors: teacher-model quality and careful hyperparameter tuning captured by overhead metrics.

4.3. Practical Implications for Mechanism Selection

The study's findings suggested that mechanism selection in FL should be dictated by deployment settings rather than a single protection-strength metric, especially in large-scale, cross-device deployments on heterogeneous, resource-constrained clients such as mobile and IoT devices. Findings also indicate that providing better data protection, lightweight mechanisms, secure aggregation, FedAvg, model distillation, or regularization implemented individually or combined, cryptographic on the other hand, is sufficient where implementation is in smaller-scale, cross-silo deployments with well-resourced participants, where the highest level of confidentiality is paramount, communication and compute budgets are less constrained. Therefore, cryptographic mechanisms: homomorphic encryption and secure multi-party computation are viable despite their overhead. Differential privacy offers a middle path, at a cost of requiring careful, context-specific tuning of the privacy budget to avoid unacceptable utility loss.

4.4. Limitations

The analysis was based on a synthesis of the reported findings across heterogeneous studies and the qualitative ratings in [Table 1](#) should be interpreted as directional comparisons rather than precise, universally applicable figures. Reported overheads depend heavily on implementation details, model architecture, dataset characteristics, and hardware, all of which vary across the underlying literature. In addition, several mechanisms are frequently deployed in combination for instance differential privacy being layered secure aggregation. Future work applying a controlled, common testbed across mechanisms would help validate and refine

the directional findings presented in this study.

5. Conclusion

The study carried out systematic comparative analysis of eight data-protection mechanisms in Federated learning, evaluated based on their performance metrics, including computational overhead, communication overhead, model accuracy, and scalability. Study analysis found that there is a consistent trade-off pattern where cryptographic mechanisms: homomorphic encryption, secure multi-party computation, provide a strong protection at high computational, communication, and scalability cost, while secure aggregation, FedAvg, model distillation, and regularization offer more favourable performance, which is suited to large-scale and resource-constrained deployments. Differential privacy has low overhead but comes with a direct, tunable trade-off against model utility. The findings offer practical guidance for practitioners selecting data-protection mechanisms for deployment and also highlight the need for standardized, controlled benchmarking across mechanisms to complement literature-based approaches. Future work can conduct empirical evaluations of mechanisms for the hybrid combination across varied federated learning systems, including cross-device deployments, to validate and extend the directional findings of this study.

Acknowledgements

The author would like to thank the School of Informatics and Innovative Systems at Jaramogi Oginga Odinga University of Science and Technology for providing a conducive environment for conducting this research. Richard Omolo Newton Masinde for their constructive suggestions and comments.

Author Contributions

Mitende Nicholus Nyapete: Conceptualization, Data curation, Formal Analysis, Funding acquisition, Investigation, Methodology, Resources, Software, Validation, Visualization, Writing—original draft, Writing review & editing.

Richard Omolo: Supervision, Writing—review & editing.

Newton Masinde: Supervision, Writing—review & editing.

Conflicts of Interest

The authors declare no conflicts of interest regarding the publication of this paper.

References

- [1] Banse, A., Kreischer, J. and Jürgens, X.O.I. (2024) Federated Learning with Differential Privacy. arXiv: 2402.02230. <http://arxiv.org/abs/2402.02230>
- [2] Fu, J., Hong, Y., Ling, X., Wang, L., Ran, X., Sun, Z., Wang, W.H., Chen, Z. and Cao, Y. (2024) Differentially Private Federated Learning: A Systematic Review. arXiv: 2405.08299.
- [3] Ibrahim Khalaf, O., Ashokkumar, S.R., Algburi, S., Anupallavi, S., Selvaraj, D., Sharif,

- M.S., *et al.* (2024) Federated Learning with Hybrid Differential Privacy for Secure and Reliable Cross-IoT Platform Knowledge Sharing. *Security and Privacy*, **7**, e374. <https://doi.org/10.1002/spy2.374>
- [4] Ren, X., Yang, S., Zhao, C., McCann, J. and Xu, Z. (2024) Belt and Braces: When Federated Learning Meets Differential Privacy. *Communications of the ACM*, **67**, 66-77. <https://doi.org/10.1145/3650028>
 - [5] Saha, S., Hota, A., Chattopadhyay, A.K., Nag, A. and Nandi, S. (2024) A Multifaceted Survey on Privacy Preservation of Federated Learning: Progress, Challenges, and Opportunities. *Artificial Intelligence Review*, **57**, Article No. 184. <https://doi.org/10.1007/s10462-024-10766-7>
 - [6] Wang, B., Li, H., Ren, X. and Guo, Y. (2023) An Efficient Differential Privacy-Based Method for Location Privacy Protection in Location-Based Services. *Sensors*, **23**, Article 5219. <https://doi.org/10.3390/s23115219>
 - [7] Wei, K., Li, J., Ding, M., Ma, C., Yang, H.H., Farhad, F., Jin, S., Quek, T.Q.S. and Poor, H.V. (2019) Federated Learning with Differential Privacy: Algorithms and Performance Analysis. arXiv: 1911.00222. <http://arxiv.org/abs/1911.00222>
 - [8] Gong, Y., Chang, X., Mišić, J., Mišić, V.B., Wang, J. and Zhu, H. (2024) Practical Solutions in Fully Homomorphic Encryption: A Survey Analyzing Existing Acceleration Methods. *Cybersecurity*, **7**, Article No. 5. <https://doi.org/10.1186/s42400-023-00187-4>
 - [9] Hussien, N., Hussien, N.M., Salman, S.A. and Aljanabi, M. (2023) Secure Federated Learning with a Homomorphic Encryption Model. *International Journal Paper Advance and Scientific Review*, **4**, 1-7. <https://doi.org/10.47667/ijpasr.v4i3.235>
 - [10] Moshawrab, M., Adda, M., Bouzouane, A., Ibrahim, H. and Raad, A. (2024) Securing Federated Learning: Approaches, Mechanisms and Opportunities. *Electronics*, **13**, Article 3675. <https://doi.org/10.3390/electronics13183675>
 - [11] Munjal, K. and Bhatia, R. (2022) A Systematic Review of Homomorphic Encryption and Its Contributions in Healthcare Industry. *Complex & Intelligent Systems*, **9**, 3759-3786. <https://doi.org/10.1007/s40747-022-00756-z>
 - [12] Nikfam, F., Casaburi, R., Marchisio, A., Martina, M. and Shafique, M. (2023) A Homomorphic Encryption Framework for Privacy-Preserving Spiking Neural Networks. *Information*, **14**, Article 537. <https://doi.org/10.3390/info14100537>
 - [13] Byrd, D. and Polychroniadou, A. (2020) Differentially Private Secure Multi-Party Computation for Federated Learning in Financial Applications. *Proceedings of the First ACM International Conference on AI in Finance*, New York, 15-16 October 2020, 1-9. <https://doi.org/10.1145/3383455.3422562>
 - [14] Hosseini, S.M., Sikaroudi, M., Babaei, M. and Tizhoosh, H.R. (2022) Cluster Based Secure Multi-Party Computation in Federated Learning for Histopathology Images. In: Albarqouni, S., *et al.*, Eds., *Distributed, Collaborative, and Federated Learning, and Affordable AI and Healthcare for Resource Diverse Global Health*, Springer, 110-118. https://doi.org/10.1007/978-3-031-18523-6_11
 - [15] Liu, F., Zheng, Z., Shi, Y., Tong, Y. and Zhang, Y. (2023) A Survey on Federated Learning: A Perspective from Multi-Party Computation. *Frontiers of Computer Science*, **18**, Article No. 181336. <https://doi.org/10.1007/s11704-023-3282-7>
 - [16] Biswas, S., Kermarrec, A.M., Pires, R., Sharma, R. and Vujasinovic, M. (2024) Secure Aggregation Meets Sparsification in Decentralized Learning. arXiv: 2405.07708.
 - [17] Pasquini, D., Francati, D. and Ateniese, G. (2022) Eluding Secure Aggregation in Federated Learning via Model Inconsistency. *Proceedings of the 2022 ACM SIGSAC*

- Conference on Computer and Communications Security*, Los Angeles, 7-11 November 2022, 2429-2443. <https://doi.org/10.1145/3548606.3560557>
- [18] Zhang, Y., Behnia, R., Yavuz, A.A., Ebrahimi, R. and Bertino, E. (2024) Uncovering Attacks and Defenses in Secure Aggregation for Federated Deep Learning. 2024 *IEEE International Conference on Data Mining Workshops (ICDMW)*, Abu Dhabi, 9 December 2024, 650-656. <https://doi.org/10.1109/icdmw65004.2024.00090>
 - [19] Casella, B. and Fonio, S. (2023). Architecture-Based FedAvg for Vertical Federated Learning. *Proceedings of the IEEE/ACM 16th International Conference on Utility and Cloud Computing*, Taormina, 4-7 December 2023, 1-6. <https://doi.org/10.1145/3603166.3632559>
 - [20] Wan, W., Hu, S., Lu, J., Zhang, L.Y., Jin, H. and He, Y. (2022) Shielding Federated Learning: Robust Aggregation with Adaptive Client Selection. *Proceedings of the Thirty-First International Joint Conference on Artificial Intelligence*, Vienna, 23-29 July 2022, 753-760. <https://doi.org/10.24963/ijcai.2022/106>
 - [21] Sun, T., Li, D. and Wang, B. (2023) Decentralized Federated Averaging. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, **45**, 4289-4301. <https://doi.org/10.1109/tpami.2022.3196503>
 - [22] Zhang, H., Wu, T., Cheng, S. and Liu, J. (2024) CC-FedAvg: Computationally Customized Federated Averaging. *IEEE Internet of Things Journal*, **11**, 4826-4841. <https://doi.org/10.1109/jiot.2023.3300080>
 - [23] Zhou, Y., Qing, Y. and Lv, J. (2021) Communication-Efficient Federated Learning with Compensated Overlap-FedAvg. arXiv: 2012.06706.
 - [24] Moshawrab, M., Adda, M., Bouzouane, A., Ibrahim, H. and Raad, A. (2023) Reviewing Federated Learning Aggregation Algorithms; Strategies, Contributions, Limitations and Future Perspectives. *Electronics*, **12**, Article 2287. <https://doi.org/10.3390/electronics12102287>
 - [25] Jiang, D., Shan, C. and Zhang, Z. (2020) Federated Learning Algorithm Based on Knowledge Distillation. 2020 *International Conference on Artificial Intelligence and Computer Engineering (ICAICE)*, Beijing, 23-25 October 2020, 163-167. <https://doi.org/10.1109/icaice51518.2020.00038>
 - [26] Li, D. and Wang, J. (2019) FedMD: Heterogenous Federated Learning via Model Distillation. arXiv: 1910.03581.
 - [27] Mora, A., Tenison, I., Bellavista, P. and Rish, I. (2024) Knowledge Distillation in Federated Learning: A Practical Guide. *Proceedings of the Thirty-Third International Joint Conference on Artificial Intelligence*, Jeju, 3-9 August 2024, 8188-8196. <https://doi.org/10.24963/ijcai.2024/905>
 - [28] Song, C., Saxena, D., Cao, J. and Zhao, Y. (2024) FedDistill: Global Model Distillation for Local Model De-Biasing in Non-IID Federated Learning. arXiv: 2404.09210.
 - [29] Wu, C., Wu, F., Lyu, L., Huang, Y. and Xie, X. (2022) Communication-Efficient Federated Learning via Knowledge Distillation. *Nature Communications*, **13**, Article No. 2032. <https://doi.org/10.1038/s41467-022-29763-x>
 - [30] Yang, G. and Tae, H. (2023) Federated Distillation Methodology for Label-Based Group Structures. *Applied Sciences*, **14**, Article 277. <https://doi.org/10.3390/app14010277>
 - [31] Boix-Adsera, E. (2024) Towards a Theory of Model Distillation. arXiv: 2403.09053.
 - [32] Kotsilieris, T., Anagnostopoulos, I. and Livieris, I.E. (2022) Special Issue: Regularization Techniques for Machine Learning and Their Applications. *Electronics*, **11**, Article 521. <https://doi.org/10.3390/electronics11040521>

- [33] Liu, L. and Zhou, D. (2025) Analysis of Regularized Federated Learning. *Neurocomputing*, **611**, Article ID: 128579. <https://doi.org/10.1016/j.neucom.2024.128579>
- [34] Lv, Y., Ding, H., Wu, H., Zhao, Y. and Zhang, L. (2023) FedRDS: Federated Learning on Non-IID Data via Regularization and Data Sharing. *Applied Sciences*, **13**, Article 12962. <https://doi.org/10.3390/app132312962>
- [35] Xu, Q., Cohn, T. and Ohrimenko, O. (2023) Fingerprint Attack: Client Deanonimization in Federated Learning. *ECAI 2023: 26th European Conference on Artificial Intelligence*, Kraków, 30 September-4 October 2023, 2792-2801.
- [36] Ma, C., Li, J., Ding, M., Yang, H.H., Shu, F., Quek, T.Q.S. and Poor, H.V. (2020) On Safeguarding Privacy and Security in the Framework of Federated Learning. arXiv: 1909.06512. <http://arxiv.org/abs/1909.06512>