

Privacy-Preserving Federated Learning for Science: Challenges and Research Directions

Kibaek Kim^{*1}, Krishnan Raghavan^{*1}, Olivera Kotevska^{†1}, Matthieu Dorier^{*2}, Ravi Madduri^{*2}, Minseok Ryu^{§2}, Todd Munson^{*3}, Rob Ross^{*3}, Thomas Flynn^{†3}, Ai Kagawa^{†3}, Byung-Jun Yoon^{†3}

Christian Engelmann^{‡3}, Farzad Yousefian^{¶3}

^{*}Argonne National Laboratory, Lemont, IL

[†]Brookhaven National Laboratory, Upton, NY

[‡]Oak Ridge National Laboratory, Oak Ridge, TN

[§]Arizona State University, Tempe, AZ

[¶]Rutgers University, Piscataway, NJ

Abstract—This paper discusses the key challenges and future research directions for privacy-preserving federated learning (PPFL), with a focus on its application to large-scale scientific artificial intelligence models, in particular, foundation models (FMs). PPFL enables collaborative model training across distributed datasets while preserving privacy—an important collaborative approach for science. We discuss the need for efficient and scalable algorithms to address the increasing complexity of FMs, particularly when dealing with heterogeneous clients. In addition, we underscore the need for developing advance privacy-preserving techniques, such as differential privacy, to balance privacy and utility in large FMs emphasizing fairness and incentive mechanisms to ensure equitable participation among heterogeneous clients. Finally, we emphasize the need for a robust software stack supporting scalable and secure PPFL deployments across multiple high-performance computing facilities. We envision that PPFL would play a crucial role to advance scientific discovery and enable large-scale, privacy-aware collaborations across science domains.

Index Terms—federated learning, privacy preservation, foundation models, distributed computing

I. INTRODUCTION

As artificial intelligence (AI) models grow larger and more sophisticated, the need for innovative methods to train these models efficiently and securely has become increasingly important. Traditional, centralized approaches to model training, which require collecting data at a single location, are often impractical or infeasible in sectors such as healthcare, material science, and energy systems due to data sharing restrictions. Federated learning (FL) offers a distributed approach where model training occurs across multiple decentralized devices or servers, allowing for collaborative learning without the need for centralized storage. This is particularly valuable in fields requiring multi-institutional collaboration. Privacy-preserving FL (PPFL), built on FL; incorporating privacy mechanisms such as differential privacy (DP) will play an important role in such collaboration by enabling large-scale model training while protecting sensitive information.

¹These authors contributed equally to this work as major contributions.

²These authors contributed equally to this work as significant contributions.

³These authors contributed equally to this work.

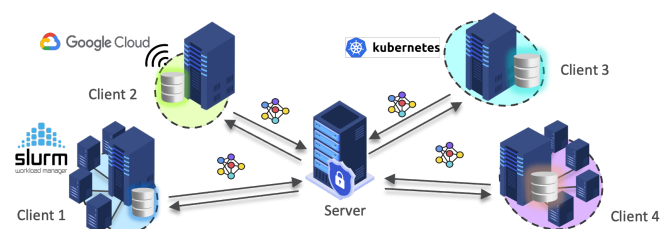


Fig. 1. Illustration of federated learning across four heterogeneous computing facilities and a central server, each of which facilitates a large amount of data without sharing. Only model updates are exchanged between clients and the server potentially with privacy-preserving mechanism.

The U.S. Department of Energy (DOE) has recognized the significance of AI in advancing scientific research and recently awarded \$68 million in funding for AI for science initiatives, aimed at leveraging AI and ML to drive scientific discoveries. One of the key research areas is the development of PPFL techniques to learn foundation models (FMs)—large-scale, pre-trained models that can be adapted and fine-tuned for specific downstream tasks—across distributed datasets and computing environments. Training FMs in scientific domains requires computational efficiency and often privacy protection, as sensitive and proprietary data can be spread across geographically distributed institutions, each with their own privacy requirements. Therefore, the ability to train FMs efficiently and reliably across distributed computing environments is crucial for accelerating scientific advancements and optimized operations in science and engineering domains.

PPFL has the potential to advance scientific collaboration by enabling secure, large-scale learning without moving sensitive data outside institutions. For instance, in domains such as biomedicine and energy systems, where data is distributed across hospitals, laboratories, or electric utility companies, PPFL enables to train models from a large amount of decentralized data while maintaining privacy guarantees. However, scaling PPFL to train FMs in science applications presents unique challenges: efficient and scalable algorithms are necessary to handle the large size and complexity of FMs, particularly with heterogeneous client environments See Figure 1.

Another crucial challenge is privacy preservation. Existing techniques such as DP, that protecting privacy-sensitive data; introduce the well-known trade-off between model performance and privacy guarantee increasing computation and communication overheads when applied to large models with billions of parameters. More crucially, ensuring privacy with optimal-but resource efficient performance when collaborating institutions have different privacy requirements and computing resources.

Deploying PPFL in high-performance computing (HPC) environments is a challenging task, particularly in the context of scientific research where datasets and computing resources are distributed across multiple secure locations with strict security protocols and firewall policies. The development of a deployable, robust software stack capable of scaling PPFL across heterogeneous large computing facilities, including DOE supercomputers and the National AI Research Resource (NAIRR) Pilot resources, is crucial to support the massive computational demands of scientific FMs to ensure data privacy and compliance with institutional security policies.

This vision paper explores the key challenges in leveraging PPFL for scientific applications, focusing on the need for efficient and scalable algorithms, advanced privacy-preserving techniques, fairness and incentive structures, and the development of a robust software stack to support distributed computing across heterogeneous computational environments.

II. STATE OF THE ART AND CHALLENGES

PPFL is essential in various fields such as healthcare, finance, and energy systems, where privacy-sensitive data cannot be shared directly. Current PPFL methods, while effective in controlled settings, face several limitations when scaling to larger AI models or diverse datasets. In this section, we discuss state-of-the-art approaches and various challenges in PPFL.

A. Efficient and Scalable Algorithms

FL faces challenges in scaling to FMs. Current methods such as quantization [1], sparsification [2], and model compression [3] have been introduced to reduce the communication overhead involved in sharing model updates between clients and the server. However, these techniques struggle to maintain efficiency as model sizes grow exponentially. Such limitations become challenging with heterogeneous clients with different network, computational capabilities and datasets.

Moreover, while memory-efficient techniques such as low-rank adaptation (LoRA) [4] and model pruning [5] have shown potential in reducing memory usage, they are not yet fully enabled for FL. For instance, LoRA's strategy of adapting in the low dimensional space does not account for clients with heterogeneous hardware capabilities as the low dimensional space is not same for all clients, requiring algorithms to dynamically adjust their memory and computation settings to remain effective. The advancement of FL algorithms must reduce communication and memory requirements and dynamically adapts to the computing power of each client to improve efficiency across a broad range of scientific applications.

B. Privacy Preservation for Large Models

Applying DP in the training of large-scale models intensifies challenges related to execution time, memory consumption, and accuracy [6]. The large number of model parameters significantly increases training costs, with DP often extending training times substantially [7]. Memory requirements are also elevated, as per-example gradients are necessary, and model accuracy tends to degrade as more noise is introduced for privacy [8]. Striking a balance between privacy and model performance remains a major challenge.

DP fine-tuning methods are introduced in [6] and applied to transformer-based LLMs for task-specific fine-tuning using smaller, private datasets. To minimize runtime and memory overheads, DP has also been integrated into various parameter-efficient fine-tuning methods [9]. Utilizing the DP stochastic gradient descent (DP-SGD) algorithm, FL faces the challenge of balancing privacy with utility due to the vast model size. Additionally, different FL clients may have varying privacy needs due to diverse policies or user preferences, resulting in heterogeneous or personalized privacy budgets across institutions. Currently this challenge has been studied in centralized settings [10], but only a few works exist as simple heuristics in FL settings [11]. However, some efforts have been done to enhance privacy protection with high utility but they are not designed for clients with varying privacy needs [12].

C. Incentives and Fairness

Ensuring client participation in FL is a critical challenge, particularly when clients have different levels of data quality, computational power, network connectivity, and privacy requirement. Traditional incentive structures, often based on contribution metrics like data volume or computation cycles, can lead to unfair imbalances, where clients with more power and resources dominate the training process. This can result in biased global models that overly learn data from these dominant clients, leading to poor generalization across the entire client base. In scientific collaborations where data heterogeneity is common, the lack of proper incentives can discourage participation from smaller clients, ultimately reducing the diversity of data and degrading model performance.

In addition to incentives, fairness in model performance is a key concern in FL, especially in cross-silo settings where institutions contribute multimodal data with heterogeneous distributions. State-of-the-art techniques (e.g., [13], [14]) aim to adjust the influence of clients based on the size or quality of their data. However, these methods often struggle in the presence of highly imbalanced datasets or heterogeneous client participation and lack the flexibility to handle real-time fairness adjustments as client contributions and data distributions evolve. The key challenge is developing fairness-aware optimization techniques that embed fairness constraints into the training process while maintaining computational efficiency. Balancing fairness with scalability and model accuracy, particularly in large-scale scientific collaborations, remains an open problem; addressed to ensure that FL systems deliver FMs that generalize equitably across all clients.

D. Continually-but Federated Learning Foundation Models

It is critical to ensure that a FM trained using PPFL maintains optimal performance throughout its life cycle. However, when these models are fine tuned on many downstream tasks, they exhibit a phenomenon known as catastrophic forgetting [15]. This phenomenon is an important challenge to overcome. A trivial but infeasible solution to this problem is to ensure that the FM does not forget prior information by simply retraining using all data. This solution is infeasible due to memory and computational requirements. Moreover, the FM needs to ensure capability to learn new downstream tasks while selectively remembering relevant prior tasks. This stability-plasticity trade-off [16] is an important challenge to consider. This particular challenge is further complicated by the requirements of FL and the need for privacy guarantees. Upholding such requirements across distributed-but heterogeneous clients throughout the lifelong of the model with low memory and compute footprints is necessary.

E. Software Stack

Scaling FL across high-performance computing (HPC) environments like DOE supercomputers presents several unique challenges, primarily related to security, data management, and communication efficiency. Existing open-source FL frameworks lack the necessary flexibility to operate in highly secure and distributed environments, where strict firewall policies often prevent direct communication across multiple HPC sites. These security constraints make it difficult to deploy FL experiments across distributed systems.

Another major challenge is the integration of efficient data management and communication tools capable of handling the scale and complexity of FL across HPC systems. While tools like Globus Compute (previously known as FuncX [17]) offer secure data transfer capabilities for FL experiments (e.g., APPFL [18], [19] and Flight [20]), managing the continuous flow of model updates across multiple compute nodes remains a challenging task. Message broker services like Kafka have been used for FL [21] but lack the capabilities to exploit the specialized hardware of supercomputers.

The complexity increases with the need for low-latency, real-time messaging systems that can ensure efficient synchronization of model updates across geographically distributed clients. Moreover, the traditional data management systems used for handling tokenized models and their updates need to be enhanced for large-scale distributed environments, ensuring that data transfers occur smoothly without overwhelming the network. This requires the development of advanced, scalable communication infrastructures capable of maintaining performance while adhering to the strict security and bandwidth constraints typical of HPC settings. Addressing these challenges is critical for enabling scalable and efficient FL across HPC systems, and will require the integration of more dynamic, flexible software solutions tailored to such environments.

III. FUTURE DIRECTIONS

This section outlines key future directions, focusing on enhancing efficiency, privacy, fairness, and the development of a robust software stack.

A. Developing Efficient and Scalable Algorithms

Communication efficiency is critical for scaling FL, particularly applied to large FMs. While techniques such as model quantization and sparsification help reduce communication overhead, they struggle to keep pace with the increasing complexity and size of FMs. A key future direction involves adaptive algorithms that adjust compression schemes based on network conditions and client resources. For instance, layer-wise compression could selectively compress less important model layers more aggressively, enabling efficient communication without sacrificing model performance.

Another critical challenge is heterogeneous client environments. FL systems must deal with clients that have different computational capabilities, ranging from large computing facilities and data centers to low-power edge devices. This heterogeneity can lead to performance bottlenecks, particularly when slower devices, or stragglers, delay the global model update process. Future work must explore *heterogeneity-aware algorithms*, such as Federated Dropout [22] and asynchronous FL [23], [24], which allows clients to contribute only partial model updates based on their computational capacity. Efficiently balancing workload distribution, enabling faster global updates minimally affected by stragglers is a must.

Memory efficiency is also critical when dealing with FMs. Existing techniques such as LoRA and model pruning need to be further optimized for FL settings [25]. For example, state-of-the-art federated pruning (e.g., [2], [26]) could further be improved by allowing clients to individually prune their local models based on the available resources, while the server assembles a globally pruned model that maintains performance. Layer-wise adjustment techniques (e.g., by either freezing and partially aggregating) could be explored; providing efficiency in computation, communication, and memory usage in FL.

Energy efficiency will become increasingly important as FL scales up, particularly when deploying models in environments that span the computing continuum, from edge devices to HPC clusters. We envision that energy-aware scheduling algorithms will play a crucial role in optimizing the trade-off between computational load and energy consumption. These algorithms could orchestrate client participation based on their energy profiles and computing power, allowing more energy-efficient distribution of workload among clients.

B. Enhancing Privacy Preservation for Large Models

Managing heterogeneous privacy requirements in FL is a critical and under-explored research area. Current methods often apply simplistic strategies, such as enforcing the strictest privacy settings to all clients or introducing biases through weighted averaging. Both approaches can significantly degrade performance and fail to accommodate diverse privacy needs. A promising future direction lies in building on techniques

like projected DP-SGD [27], which balance between privacy and utility using projection-based strategies.

One potential approach is the development of multilevel projection methods that dynamically adapt to the changing needs of the training process. These projections could be applied in stages to capture important parameter variations while filtering out noise more effectively. Such methods could address multiple groups of clients with different privacy budgets, ensuring that their diverse requirements are met without compromising the overall training process. Research in this area could focus on optimizing projection strategies for more robust, scalable, and computationally efficient FL with diverse privacy requirements.

Another promising direction is the use of shuffling models [28], which enhance privacy by incorporating a randomizer and a shuffler. Tailored shuffling models for FL could focus on the communication messages exchanged between the server and clients. Client-side protocol could be designed such that data is randomized and permuted before transmission during communication rounds. These mechanisms could add an extra layer of privacy while addressing critical challenges, such as improving scalability and resilience to failures in distributed systems [29], reduce communication costs [30], and mitigate risks of unintended information leakage [31].

C. Fostering Incentives and Fairness

Balancing different heterogeneous clients with varying data quality, computational resources, and privacy requirements warrants novel research directions. Modelling incentive and fairness structures across clients as a game with evolving strategies to capture interactions across and among clients is important. Designing and developing precise fairness metrics to ensure equitable contributions and performance across all clients while understanding their impact on the global model is a way to describe fairness metric. Another important direction is to investigate and develop fairness mechanisms, ideally in real time by adjusting client contributions dynamically during the training process. This requires development of adaptive optimization and dynamic programming algorithms to adjust aggregation weights and other hyperparameters based on the evolving contributions of clients. Game-theoretic approaches could also be applied incentivize clients for maintaining fairness in model performance across all participants.

D. Continually Adapted Federated Learning

Adapting the FM throughout the lifetime of the models' relevance is crucial due to reasons discussed in section II-D. we must address many of the challenges with fairness and incentives discussed in section II-C and section II-A in an environment with evolving distributions. This evolving heterogeneity necessitates asynchronous updates. This problem will take the shape of a dynamically evolving Nash game which will model a non-stationary trade-off between different clients ensuring fairness while dealing with the heterogeneity between the data distributions. Within the context of this evolving game, different client might behave as players describing the

varying needs of FM with an extremely flexible framework would be extremely flexible that identifies the different trade-offs, enables privacy [16].

E. Building a Robust Software Stack

To address the challenges of deploying FL across HPC environments, future software frameworks must focus on building scalable, secure communication infrastructures. As outlined in Section II-E, existing frameworks struggle with managing real-time communication and data synchronization due to stringent firewall policies and security protocols in DOE supercomputers. Future developments must focus on creating low-latency communication systems capable of efficiently handling large FMs across geographically distributed clients, while ensuring compliance with security requirements. Moreover, these systems must be designed to remain resilient in the face of network disruptions and client dropouts, ensuring reliable updates to global model in dynamic, distributed settings.

In addition to communication, the data management layer of the software stack must be significantly improved to handle the tokenization, transfer, and updating of large models. Current systems, as discussed in Section II-E, face significant bottlenecks when dealing with model updates across multiple HPC nodes. Future frameworks will need to implement real-time data processing and streamlined model versioning systems to reduce the computational and bandwidth overhead associated with updating models in federated environments. Furthermore, these systems must incorporate intelligent caching and data locality strategies to minimize redundant data transfers, when dealing with massive datasets across various HPC sites.

Moreover, while some existing FL frameworks already embed DP and secure aggregation, future development must prioritize scalability and efficiency, particularly for large-scale models in HPC environments. Future software stacks should focus on improving the scalability of privacy-preserving tools, ensuring they can handle the computational and communication demands of HPC systems without compromising performance. This requires not only enhancing the flexibility of privacy mechanisms to accommodate varying client needs but also tightly integrating them with data management and communication layers to streamline the entire FL process. The next generation of FL software must ensure that privacy-preserving features, such as DP, are implemented in a way that minimizes overhead and maximizes compatibility with HPC-specific security and performance constraints.

In summary, future FL frameworks for HPC environments will require a highly integrated software stack that combines scalable communication, efficient data management, and robust privacy mechanisms. By addressing the challenges discussed in Section II-E, these future software ecosystems enable FL with scaling across the computing continuum enabling large-scale, privacy-conscious scientific applications.

IV. CONCLUSIONS

Recent advancements in FL have demonstrated its potential to train models across distributed data sources while preserv-

ing data privacy. However, as we move toward larger, more complex models such as FMs, several challenges arise. While promising, state-of-the-art approaches show limitations when scaling to the diverse and heterogeneous settings of scientific applications. Key issues such as communication bottlenecks, resource imbalances among clients, and integrating privacy-preserving mechanisms into large FMs will require careful design of algorithms and robust software stack.

Addressing these challenges requires interdisciplinary collaboration across multiple fields. Scalable optimization algorithms, fairness-aware mechanisms, and privacy-preserving techniques developed through collaborative efforts between experts in optimization, computer science, and domain-specific fields. In particular, collaborations with scientists and engineers in energy systems, healthcare, climate science, and experimental facilities are crucial to advance and adapt FL for these scientific domains.

A key research direction lines in the development of adaptive FL systems that can better handle large-scale model training across heterogeneous clients in terms of datasets, computational resources, and privacy requirements. Advances in game-theoretic approaches and fairness-aware learning algorithms could be promising to address the challenges with client incentives and fairness. Moreover, novel privacy-preserving techniques are important for model training at scale, without sacrificing the training efficiency and model performance.

ACKNOWLEDGMENT

This work was supported by the U.S. Department of Energy, Office of Science, Advanced Scientific Computing Research, under Contract DE-AC02-06CH11357.

REFERENCES

- [1] K. Gupta, M. Fournarakis, M. Reisser, C. Louizos, and M. Nagel, "Quantization robust federated learning for efficient inference on heterogeneous devices," *Transactions on Machine Learning Research*, 2023.
- [2] G. Bai, Y. Li, Z. Li, L. Zhao, and K. Kim, "FedSpaLLM: Federated pruning of large language models," *arXiv preprint arXiv:2410.14852*, 2024.
- [3] G. Wilkins, S. Di, J. C. Calhoun, Z. Li, K. Kim, R. Underwood, R. Mortier, and F. Cappello, "FedSZ: Leveraging error-bounded lossy compression for federated learning communications," in *2024 IEEE 44th International Conference on Distributed Computing Systems (ICDCS)*. IEEE, 2024, pp. 577–588.
- [4] E. J. Hu, Yelong shen, P. Wallis, Z. Allen-Zhu, Y. Li, S. Wang, L. Wang, and W. Chen, "LoRA: Low-rank adaptation of large language models," in *International Conference on Learning Representations*, 2022.
- [5] G. Bai, Y. Li, C. Ling, K. Kim, and L. Zhao, "SparseLLM: Towards global pruning of pre-trained language models," in *The Thirty-eighth Annual Conference on Neural Information Processing Systems*, 2024.
- [6] D. Yu, S. Naik, A. Backurs, S. Gopi, H. A. Inan, G. Kamath, J. Kulkarni, Y. T. Lee, A. Manoel, L. Wutschitz, S. Yekhanin, and H. Zhang, "Differentially private fine-tuning of language models," in *The Tenth International Conference on Learning Representations, ICLR*, 2022.
- [7] P. Subramani, N. Vadivelu, and G. Kamath, "Enabling fast differentially private SGD via just-in-time compilation and vectorization," *Advances in Neural Information Processing Systems*, vol. 34, 2021.
- [8] R. Bassily, A. Smith, and A. Thakurta, "Private empirical risk minimization: Efficient algorithms and tight error bounds," in *IEEE 55th annual symposium on foundations of computer science*, 2014, pp. 464–473.
- [9] R. Karimi Mahabadi, J. Henderson, and S. Ruder, "Compacter: Efficient low-rank hypercomplex adapter layers," *Advances in Neural Information Processing Systems*, vol. 34, pp. 1022–1035, 2021.
- [10] B. Niu, Y. Chen, B. Wang, J. Cao, and F. Li, "Utility-aware exponential mechanism for personalized differential privacy," in *2020 IEEE Wireless Communications and Networking Conference*. IEEE, 2020, pp. 1–6.
- [11] A. K. Chathoth, A. Jagannatha, and S. Lee, "Federated intrusion detection for IoT with heterogeneous cohort privacy," *arXiv preprint arXiv:2101.09878*, 2021.
- [12] A. Girgis, D. Data, S. Diggavi, P. Kairouz, and A. T. Suresh, "Shuffled model of differential privacy in federated learning," in *International Conference on Artificial Intelligence and Statistics*, 2021.
- [13] T. Li, M. Sanjabi, A. Beirami, and V. Smith, "Fair resource allocation in federated learning," in *International Conference on Learning Representations*, 2020.
- [14] W. Chu, C. Xie, B. Wang, L. Li, L. Yin, A. Nourian, H. Zhao, and B. Li, "Focus: Fairness via agent-awareness for federated learning on heterogeneous data," in *International Workshop on Federated Learning in the Age of Foundation Models in Conjunction with NeurIPS*, 2023.
- [15] V. V. Ramasesh, A. Lewkowycz, and E. Dyer, "Effect of scale on catastrophic forgetting in neural networks," in *International Conference on Learning Representations*, 2021.
- [16] K. Raghavan and P. Balaprakash, "Formalizing the generalization-forgetting trade-off in continual learning," *Advances in Neural Information Processing Systems*, vol. 34, pp. 17 284–17 297, 2021.
- [17] R. Chard, Y. Babuji, Z. Li, T. Skluzacek, A. Woodard, B. Blaiszik, I. Foster, and K. Chard, "funcX: A federated function serving fabric for science," in *Proceedings of the 29th International Symposium on High-Performance Parallel and Distributed Computing*, 2020, p. 65–76.
- [18] M. Ryu, Y. Kim, K. Kim, and R. K. Madduri, "APPFL: open-source software framework for privacy-preserving federated learning," in *2022 IEEE International Parallel and Distributed Processing Symposium Workshops (IPDPSW)*. IEEE, 2022, pp. 1074–1083.
- [19] Z. Li, S. He, Z. Yang, M. Ryu, K. Kim, and R. Madduri, "Advances in APPFL: A comprehensive and extensible federated learning framework," *arXiv preprint arXiv:2409.11585*, 2024.
- [20] N. Hudson, V. Hayot-Sasson, Y. Babuji, M. Baughman, J. G. Pauloski, R. Chard, I. Foster, and K. Chard, "Flight: A FaaS-based framework for complex and hierarchical federated learning," *arXiv preprint arXiv:2409.16495*, 2024.
- [21] A. J. Chaves, C. Martín, and M. Díaz, "Towards flexible data stream collaboration: Federated learning in Kafka-ML," *Internet of Things*, vol. 25, p. 101036, 2024.
- [22] S. Caldas, J. Konečný, H. B. McMahan, and A. Talwalkar, "Expanding the reach of federated learning by reducing client resource requirements," *arXiv preprint arXiv:1812.07210*, 2018.
- [23] Z. Li, P. Chaturvedi, S. He, H. Chen, G. Singh, V. Kindratenko, E. A. Huerta, K. Kim, and R. Madduri, "FedCompass: Efficient cross-silo federated learning on heterogeneous client devices using a computing power-aware scheduler," in *The Twelfth International Conference on Learning Representations*, 2024.
- [24] C. Iakovidou and K. Kim, "Asynchronous federated stochastic optimization with exact averaging for heterogeneous local objectives," *arXiv preprint arXiv:2405.10123*, 2024.
- [25] Z. Li, S. He, P. Chaturvedi, V. Kindratenko, E. A. Huerta, K. Kim, and R. Madduri, "Secure federated learning across heterogeneous cloud and high-performance computing resources—a case study on federated fine-tuning of llama 2," *Computing in Science & Engineering*, 2024.
- [26] K. Yi, N. Gazagnadou, P. Richtárik, and L. Lyu, "FedP3: Federated personalized and privacy-friendly network pruning under model heterogeneity," in *The Twelfth International Conference on Learning Representations*, 2024.
- [27] J. Liu, J. Lou, L. Xiong, J. Liu, and X. Meng, "Projected federated averaging with heterogeneous differential privacy," *Proceedings of the VLDB Endowment*, vol. 15, no. 4, pp. 828–840, 2021.
- [28] A. Cheu, "Differential privacy in the shuffle model: A survey of separations," *arXiv preprint arXiv:2107.11839*, 2021.
- [29] Y. Wei, J. Jia, Y. Wu, C. Hu, C. Dong, Z. Liu, X. Chen, Y. Peng, and S. Wang, "Distributed differential privacy via shuffling vs aggregation: A curious study," *IEEE Transactions on Information Forensics and Security*, 2024.
- [30] A. M. Girgis and S. Diggavi, "Multi-message shuffled privacy in federated learning," *IEEE Journal on Selected Areas in Information Theory*, vol. 5, pp. 12–27, 2024.
- [31] T. Qi, H. Wang, and Y. Huang, "Towards the robustness of differentially private federated learning," in *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 38, no. 18, 2024, pp. 19911–19919.