

Large Language Model-Driven Federated Learning Framework for Adaptive Traffic Forecasting in Next-Generation Wireless Networks

Jerome C. Phillips

Department of Computer Science, University of Houston, Houston, TX, USA.
contactjerome@uh.edu

Arthur Ortega

Department of Computer Science and Engineering, University of Nevada, Reno, Reno, NV, USA.
ortega943@unr.edu

Matteo M. Reed

Department of Computer Science, University of Central Florida, Orlando, FL, USA.
matteowork@ucf.edu

Abstract

Next-generation wireless networks will operate in highly dynamic, heterogeneous, and ultra-dense environments that demand real-time traffic forecasting for resource allocation, energy optimization, and quality-of-service assurance. Conventional centralized machine learning approaches face fundamental tensions between predictive accuracy, communication overhead, and data privacy, while distributed solutions such as federated learning offer a partial resolution yet often struggle with non-IID data distributions and limited semantic understanding of network context. This paper proposes a large language model-driven federated learning framework that integrates the representational and reasoning capabilities of pre-trained language models into a privacy-preserving, collaborative learning architecture tailored for adaptive traffic forecasting. The framework employs language models as meta-coordinators capable of interpreting multimodal network metadata, guiding local model adaptation, and facilitating zero-shot and few-shot transfer across diverse base station environments without sharing raw traffic measurements. System-level design principles, architectural trade-offs, infrastructure implications, sustainability constraints, robustness guarantees, and fairness considerations are examined comprehensively. The analysis situates the framework within the broader landscape of edge intelligence, telecommunication governance, and sustainable AI, providing a forward-looking perspective on how large language models can reshape the forecasting pipeline while aligning with the operational realities of future wireless systems. Through deep conceptual exploration, this work articulates a pathway toward cognitively adaptive, ethically grounded, and environmentally conscious traffic intelligence for sixth-generation networks.

Keywords

federated learning, large language models, traffic forecasting, next-generation wireless networks, edge intelligence, sustainable AI, network governance

1 Introduction The proliferation of beyond-fifth-generation and sixth-generation wireless networks is characterized by extreme device densities, heterogeneous service requirements, and stringent

latency and reliability constraints. In this landscape, accurate and adaptive traffic forecasting becomes a cornerstone for proactive resource orchestration, load balancing, spectrum sharing, and energy-aware operation. Traditional forecasting models, ranging from statistical time-series methods to deep learning architectures, have demonstrated substantial progress, yet their deployment in next-generation infrastructures is challenged by the inherent distribution, privacy sensitivity, and semantic richness of network data. Centralizing raw traffic traces from thousands of base stations poses severe risks to user privacy, incurs prohibitive communication overhead, and conflicts with emerging data sovereignty regulations. Federated learning has emerged as a powerful paradigm to train machine learning models across decentralized data silos without direct data sharing, offering a compelling foundation for privacy-respecting traffic intelligence [1], [2], [3]. However, purely statistical federated models often lack the capacity to interpret high-level network context such as service-level agreements, application types, emergency events, or evolving user behavior patterns, which are critical for accurate forecasting under distributional shifts. Large language models (LLMs) have recently exhibited remarkable capabilities not only in natural language understanding but also in reasoning over structured data, time-series modeling, and cross-modal knowledge transfer [6], [13]. By encoding vast amounts of world knowledge and demonstrating emergent few-shot learning abilities, LLMs can serve as a bridge between raw numerical traffic measurements and the semantic fabric of network operations. Integrating such models into a federated forecasting framework introduces a new class of system architectures where the LLM functions as a meta-learner, contextual enhancer, or coordinator, thereby enabling adaptive, context-aware predictions without compromising the locality principle of federated learning. This paper presents a theoretical and system-level blueprint for an LLM-driven federated learning framework for adaptive traffic forecasting. Rather than focusing on algorithmic implementation details, the discussion emphasizes architectural design choices, structural trade-offs, system governance, infrastructure sustainability, robustness, fairness, and policy implications. The contribution lies in synthesizing insights across distributed machine learning, foundation models, wireless networking, and socio-technical system design to delineate a research trajectory that is both technically viable and societally responsible. The remainder of this paper is organized as follows. Section 2 reviews related work and background across federated learning, traffic forecasting, and the emerging intersection of large language models with time-series and network intelligence. Section 3 presents the system architecture and core design principles of the proposed framework. Section 4 examines the integration of large language models within federated learning pipelines, emphasizing semantic augmentation and adaptation strategies. Section 5 discusses infrastructure deployment models, sustainability trade-offs, and energy implications. Section 6 addresses robustness, security, and fairness concerns alongside governance and regulatory considerations. Section 7 offers a discussion of open challenges and future directions, and Section 8 concludes the paper.

2 Related Work and Background

Federated learning was introduced as a communication-efficient and privacy-preserving framework wherein multiple clients collaboratively train a shared model under the orchestration of a central server, while keeping training data localized [1]. Subsequent research thoroughly mapped the challenges of statistical heterogeneity, system constraints, and adversarial threats, framing federated learning as a multi-dimensional optimization problem spanning communication, computation, and accuracy [2], [3]. In the wireless domain, a growing body of literature has demonstrated the potential of federated learning for tasks such as channel estimation, user clustering, and demand prediction, highlighting both opportunities and wireless-specific bottlenecks including limited uplink bandwidth and energy budgets of edge devices [9], [14]. Parallel to

advances in federated learning, traffic forecasting in cellular and wireless networks has evolved from autoregressive integrated moving average models to deep spatio-temporal architectures, including convolutional and graph neural networks that exploit spatial correlations among base stations [4], [5]. These methods typically assume abundant centralized data and can be vulnerable to distributional shifts caused by sudden changes in user mobility, application mix, or network reconfiguration. Federated learning-based traffic prediction has been proposed to address privacy and data locality concerns, yet it largely inherits the limitations of pure numerical models, producing forecasts that are agnostic to semantic context such as local events, policy changes, or application semantics [12], [15]. In parallel, the emergence of LLMs has reshaped the landscape of time-series analysis. Researchers have demonstrated that pre-trained language models can serve as powerful zero-shot forecasters after minimal adaptation, leveraging their internal representations of sequential structures and world knowledge to model numerical dynamics [6]. Domain-specific explorations have extended LLM-based forecasting to network traffic, with initial efforts incorporating spatial-temporal graph structures and prompting strategies to capture spatio-temporal dependencies [13]. The convergence of LLM capabilities with networking tasks has been further surveyed, revealing potential for tasks extending beyond forecasting to configuration generation, anomaly explanation, and automated troubleshooting [7]. These developments suggest that embedding language model intelligence into a federated framework could overcome the context-blindness of traditional statistical learners, enabling a new breed of cognitively adaptive traffic forecasting systems. The framework proposed herein builds upon these advancements by articulating a system-level integration that respects privacy, scale, and sustainability constraints.

3 System Architecture and Design Principles

The architecture envisions a hierarchical federation of wireless edge nodes, each associated with one or more base stations collecting downlink and uplink traffic measurements, control-plane events, and metadata logs. At each edge node, a lightweight local forecasting model, potentially a compact recurrent or attention-based neural network, predicts future traffic loads over short- and medium-term horizons. A logically centralized federated aggregation server, which may be geographically distributed across cloud and edge tiers, periodically collects encrypted model updates and orchestrates global aggregation using secure aggregation protocols without accessing raw data [1], [10]. Crucially, the architecture introduces an LLM-enabled meta-coordination layer that operates alongside the federated aggregation server or in a decentralized council configuration. This layer ingests high-level metadata such as application identifiers, quality-of-service class labels, network slice descriptors, event logs, and textual descriptions of incident reports, which are far less privacy-sensitive than raw traffic payloads and can be shared under well-defined governance policies. The LLM processes these semantic signals to generate context vectors, adaptation hints, or even modulated global model initialization that is broadcast to clients. By separating the numerical forecasting model training from the semantic reasoning pathway, the architecture preserves the privacy advantages of federated learning while enriching the learning process with interpretable knowledge. A second design principle involves dual-path model architecture, in which the local forecasting model and the LLM-based context encoder operate in a loosely coupled manner. The local model produces a base forecast from historical numerical patterns, while the LLM-derived embeddings condition or recalibrate predictions by adjusting internal gating mechanisms or generating residual correction terms. This separation ensures that even if the LLM pathway faces latency or availability constraints, the basic forecasting capability remains intact, sustaining operational reliability. Moreover, the LLM does not require fine-tuning on raw traffic data from individual clients; instead, it can be deployed as a frozen, pre-

instructed model that generates guidance based on metadata, drastically reducing the computational burden on edge nodes and respecting energy budgets. This design aligns with observations that LLMs can effectively serve as reasoning engines over structured side-information without heavy retraining [6], [21]. At the system level, the communication protocol must accommodate two distinct data streams: the high-frequency model update stream following standard federated optimization rounds, and a lower-frequency semantic inference stream where edge nodes optionally submit anonymized metadata summaries. To maintain scalability, the LLM-coordination messages are aggregated and processed in batch, with adaptive scheduling that accounts for network congestion and the differential freshness requirements of context information. The aggregation server also implements cross-client fairness-aware weighting, ensuring that regions with richer semantic data do not dominate the global reasoning process, a topic that is further explored in Section 6. This layered communication architecture enables the framework to operate across heterogeneous cellular networks, including dense urban deployments and rural macro-cells, where metadata availability and latency characteristics vary markedly.

4 Integration of Large Language Models in Federated Traffic Prediction

Incorporating LLMs into federated traffic forecasting requires careful mediation between the nature of language models, which are pre-trained on vast textual corpora, and the structured, temporal nature of wireless traffic data. A promising integration strategy is to treat the LLM as a contextual reasoning engine that translates high-level network semantics into actionable numerical modulation signals. For instance, when a base station experiences an unusual traffic surge due to a localized event such as a sports match or emergency evacuation, the numerical forecaster may fail to extrapolate unless it has encountered similar patterns during training. An LLM, prompted with textual descriptions of the event category, time of day, and affected cell identities, can infer likely traffic shape characteristics based on its accumulated world knowledge, generating a contextual adjustment that steers the local model toward more plausible predictions, even in few-shot or zero-shot regimes [6], [22]. This process does not require the LLM to store or process raw subscriber-level data, preserving privacy and aligning with data minimization principles. Another design pattern involves transferable prompt engineering, whereby standardized prompt templates are constructed from network metadata and broadcast by the aggregation server. These prompts might include information such as “urban macro cell, weekday evening, video streaming application dominant, moderate user density” and are processed by the LLM to generate embedding vectors that are then sent to clients as conditioning signals. Clients map these embeddings onto their local model parameters via small adapter networks, achieving domain adaptation without raw data exchange. The prompts can be generated in a differentially private manner by the aggregation server using aggregated metadata statistics, adding a further layer of protection. This plug-and-play style of LLM integration reduces the need for per-client fine-tuning and allows the framework to scale to thousands of edge nodes. Federated fine-tuning of the LLM itself is conceivable but introduces significant system-level trade-offs. Fully updating or even parameter-efficient fine-tuning of an LLM at client sites is computationally prohibitive for most edge devices and raises concerns about model inconsistency across rounds. An alternative that balances capability and cost is to perform federated distillation, whereby local client models are treated as students that learn to mimic the behavior of a centrally fine-tuned LLM teacher on the traffic forecasting task, while the LLM remains stationary. Distillation can take the form of output logit matching or feature alignment between the lightweight local model and the LLM’s internal representations, and the federated aggregation process then merges the student updates [8]. This approach capitalizes on the LLM’s rich representations without requiring LLM inference at the edge

during deployment, an advantageous property for latency-sensitive forecasting operations where model inference must occur within strict time budgets. However, such distillation must carefully handle the non-IID nature of both local traffic data and the semantic contexts encountered across the network, ensuring that the global student model does not collapse to a mode that ignores rare but critical contextual signals.

5 Infrastructure, Deployment, and Sustainability Considerations

Deploying an LLM-driven federated forecasting framework in operational wireless networks entails navigating complex infrastructure hierarchies that span centralized cloud data centers, regional edge computing nodes, and far-edge base stations. The LLM component, due to its memory and compute footprint, is naturally hosted at aggregation points with access to accelerated hardware, such as GPUs or specialized AI accelerators. The local forecasting models, in contrast, can be executed on embedded processors or network interface cards within virtualized radio access network stacks. This tiered deployment model aligns with contemporary edge intelligence architectures where heavy inference tasks are strategically placed at higher tiers to balance latency and energy consumption [14], [20]. System designers must decide on the division of labor between edge and cloud: a thicker edge that performs local inference and light adaptation offers lower latency and greater resilience against backhaul disruptions, while a thicker cloud-centric approach with richer LLM reasoning provides more contextually nuanced adjustments at the expense of higher communication delay. Sustainability emerges as a critical evaluation dimension. Federated learning inherently reduces the carbon footprint associated with data transmission compared to centralized cloud training, as only model updates traverse the network [2]. However, the addition of LLM inference for context reasoning introduces a non-negligible energy cost. The life-cycle carbon impact of the framework must be assessed holistically, considering the amortized costs of pre-training large language models, which themselves carry significant embodied emissions [19]. Optimizations such as quantized inference, mixture-of-experts routing, and sparse activation can substantially reduce per-inference energy consumption. Moreover, the framework can implement carbon-aware scheduling, delaying or throttling LLM context refresh cycles during periods of high grid carbon intensity, while allowing the local numerical forecasters to maintain base operation. This type of green AI design not only aligns with emerging operator sustainability mandates but also creates a research nexus between federated learning efficiency and environmentally adaptive computing. Cross-domain analogies with federated smart grid load forecasting, where similar privacy-sustainability tensions exist, offer instructive design patterns for carbon-responsive orchestration [11]. The operational sustainability also depends on hardware lifecycle considerations. Rapid obsolescence of high-end accelerators needed for LLM inference could force frequent infrastructure refresh, with attendant environmental costs. Therefore, the framework should support modular replacement of the LLM backbone, allowing newer, more efficient foundation models or even distantly related small-language models to be swapped in without disrupting the federated training pipeline. Such modularity necessitates carefully designed interface contracts between the context module and the forecasting core, an architectural discipline that pays dividends in long-term maintainability and regulatory compliance across jurisdictions with varying carbon accounting standards.

6 Robustness, Security, Fairness, and Governance

Wireless traffic data exhibit pronounced non-IID characteristics due to differences in population density, mobility patterns, device types, and application usage across geographic regions and demographic groups. Federated learning in such settings is susceptible to convergence degradation and fairness erosion, where globally aggregated models may systematically underperform on minority clients [16], [17]. The LLM-driven framework must incorporate mechanisms to detect and mitigate such

disparities. On the robustness front, the LLM context pathway can inadvertently amplify biases present in its pre-training data, for example, by associating certain districts or event types with distorted traffic priors. Auditing the embeddings and prompts for fairness requires transparent governance practices, including client-wise performance monitoring and fairness metrics that are themselves aggregated via secure federated protocols. Techniques such as fair model aggregation, reweighting, and local debiasing procedures derived from fairness-aware federated learning literature can be integrated into the coordinator logic [17]. Security threats in the proposed framework are multifaceted. The federated aggregation pathway is vulnerable to model poisoning attacks and inference attacks on gradient updates [18]. The LLM context pathway introduces new attack surfaces, such as prompt injection via compromised metadata streams or adversarial perturbations of textual descriptions that mislead the contextual reasoning. Robust defense requires a combination of secure aggregation, differential privacy noise injection on metadata summaries, and anomaly detection on semantic inputs. The dual-path design provides a measure of resilience, because even if the LLM pathway is compromised, the numerical forecasting baseline continues to operate, limiting the blast radius of an attack. Additionally, the framework can adopt zero-trust principles at the semantic layer, validating the provenance and integrity of metadata before allowing it to influence the predictions. Governance of an LLM-driven federated forecasting ecosystem cannot be treated merely as an afterthought. The involvement of large pre-trained models, which are typically developed by a small set of technology companies, introduces concentration risks and dependencies that could conflict with the multi-stakeholder nature of network operations where operators, virtual network providers, and regulator bodies coexist. To foster trust, the framework should be underpinned by model auditing, interpretability requirements on the semantic adjustments, and clear contractual boundary conditions delineating responsibilities for model outputs. The recent introduction of the TIDES framework exemplifies the growing interest in leveraging advanced language models within spatial-temporal prediction tasks, highlighting that similar governance models will be needed as such tools move from research to operational deployment in critical infrastructure [13].

Policy frameworks should require disclosure of what semantic information is used, how it is safeguarded, and under what conditions the LLM component may be updated or decommissioned. Additionally, the federation governance structure must decide on aggregation policies for context embeddings: should a single global context model be mandated, or should regions be allowed to maintain culturally and operationally specific semantic modules? This question touches upon the broader principle of subsidiarity in distributed AI systems and demands careful regulatory alignment.

7 Discussion and Future Directions

The proposed framework opens numerous avenues for future research and industrial exploration. As sixth-generation networks begin to materialize, the integration of LLMs will likely extend beyond forecasting into closed-loop network automation, where predictions directly trigger spectrum reallocation, sleep mode orchestration, and virtualized network function scaling. To realize this vision, the LLM-driven federated learning pipeline must evolve to support ultra-low-latency decision threads, requiring advances in on-device LLM inference and novel communication-computation co-design strategies. Edge-optimized language models that operate within milliwatt power budgets are beginning to appear, and their alignment with federated distilling techniques could bring semantic reasoning directly to the far edge without compromising sustainability targets. A second forward-looking theme is the fusion of multiple foundation models, where a traffic forecasting LLM collaborates with a vision-language model analyzing camera feeds or a graph neural foundation model processing topology, all within a federated multi-modal intelligence fabric. The orchestration of such a

multi-model federation would require advanced meta-learning controllers capable of dynamically configuring model roles based on contextual demand and resource availability, posing a rich set of open problems at the intersection of systems, learning theory, and networking. Accountability and legal compliance in this ecosystem will demand robust provenance tracking for every forecasting decision. Explainability methods that trace how an LLM’s semantic inference influenced a particular load prediction are still in their infancy, yet will be indispensable for dispute resolution, regulatory audits, and public trust. Similarly, sustainability-driven carbon accounting across the federated system must be standardized, giving operators the ability to report not only the operational energy but also the amortized pre-training carbon cost of the LLM backbone, a practice that aligns with emerging AI environmental disclosure frameworks. Finally, the socioeconomic implications of predictive intelligence in wireless infrastructure must be acknowledged. Improved traffic forecasting can lead to more efficient spectrum use, lower operational expenditures, and potentially lower consumer prices, yet it may also exacerbate digital divides if advanced semantic models are only deployed in affluent urban areas while rural regions rely on degraded fallback modes. Mitigating such divides demands deliberate policy interventions, including open-source foundation model initiatives and public-private partnerships that subsidize semantic infrastructure in underserved areas. The framework proposed here is thus not merely a technical construct but a socio-technical system whose design must be continuously negotiated among engineers, policymakers, and civil society. 8

Conclusion This paper has presented a comprehensive system-level vision for a large language model-driven federated learning framework for adaptive traffic forecasting in next-generation wireless networks. By interweaving privacy-preserving collaborative learning with the contextual reasoning power of pre-trained language models, the framework addresses the critical limitations of current forecasting approaches that operate blindly within purely numerical spaces. The analysis spanned architectural decomposition, integration strategies, infrastructure deployment, sustainability trade-offs, robustness and fairness safeguards, and governance imperatives, revealing a rich landscape of interdependent design decisions that must be balanced. The dual-path separation of numerical and semantic pathways, the emphasis on modularity and sustainability, and the foregrounding of equity and governance distinguish the proposed framework from ad-hoc integrations of LLMs into wireless systems. As the wireless industry advances toward increasingly autonomous and cognition-driven architectures, the principles articulated in this work offer a guiding template for building forecasting systems that are not only accurate and efficient but also trustworthy, equitable, and environmentally responsible. The convergence of federated learning and large language models for network intelligence thus represents a fertile frontier where systems research, artificial intelligence, and policy innovation must progress in concert.

References

1. McMahan, B., Moore, E., Ramage, D., Hampson, S., & Arcas, B. A. y. (2017). Communication-efficient learning of deep networks from decentralized data. In *Proceedings of the 20th International Conference on Artificial Intelligence and Statistics (AISTATS)*, 1273–1282.
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.

3. Niknam, S., Dhillon, H. S., & Reed, J. H. (2020). Federated learning for wireless communications: Motivation, opportunities, and challenges. *IEEE Communications Magazine*, 58(6), 46–51.
4. Bega, D., Gramaglia, M., Fiore, M., Banchs, A., & Costa-Pérez, X. (2019). DeepCog: Cognitive network management in sliced 5G networks with deep learning. In *IEEE INFOCOM 2019*, 280–288.
5. Xu, Z., Tang, J., Yin, C., Wang, Y., & Xue, G. (2020). Traffic prediction in cellular networks: A survey. *IEEE Access*, 8, 148601–148618.
6. Gruver, N., Finzi, M., Qiu, S., & Wilson, A. G. (2023). Large language models are zero-shot time series forecasters. In *Advances in Neural Information Processing Systems*, 36.
7. Chen, Y., Zhang, R., Han, Z., Li, Y., & Zhang, Y. (2023). Large language models for wireless networks: An overview. *IEEE Vehicular Technology Magazine*, 18(4), 72–81.
8. 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.
9. 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.
10. Bonawitz, K., Ivanov, V., Kreuter, B., Marcedone, A., McMahan, H. B., Patel, S., ... & Seth, K. (2019). Towards federated learning at scale: System design. In *Proceedings of Machine Learning and Systems*, 1, 374–388.
11. Wu, Q., Chen, X., Zhou, Z., & Zhang, J. (2022). Federated learning in edge computing: A systematic survey. *IEEE Communications Surveys & Tutorials*, 24(2), 905–937.
12. Wang, S., Tuor, T., Salonidis, T., Leung, K. K., Makaya, C., He, T., & Chan, K. (2019). Adaptive federated learning in resource constrained edge computing systems. *IEEE Journal on Selected Areas in Communications*, 37(6), 1205–1221.
13. Zhang, C., Zhang, H., Qiao, J., Li, Z., & Alouini, M. S. (2025). TIDES: Traffic Intelligence with DeepSeek Enhanced Spatial Temporal Prediction. *IEEE Journal on Selected Areas in Communications*.
14. Shi, Y., Yang, K., Jiang, T., Zhang, J., & Letaief, K. B. (2020). Communication-efficient edge AI: Algorithms and systems. *IEEE Communications Surveys & Tutorials*, 22(4), 2167–2191.
15. Hard, A., Rao, K., Mathews, R., Beaufays, F., Augenstein, S., Eichner, H., ... & Ramage, D. (2019). Federated learning for mobile keyboard prediction. *arXiv preprint arXiv:1811.03604*.
16. Li, X., Huang, K., Yang, W., Wang, S., & Zhang, Z. (2020). On the convergence of FedAvg on non-IID data. In *International Conference on Learning Representations (ICLR)*.
17. Zhang, Y., Xia, Y., Yang, Z., & Wen, C. (2021). Fairness in federated learning. *ACM Computing Surveys*, 54(9), 1–35.
18. Bagdasaryan, E., Veit, A., Hua, Y., Estrin, D., & Shmatikov, V. (2020). How to backdoor federated learning. In *Proceedings of the 23rd International Conference on Artificial Intelligence and Statistics (AISTATS)*, 2938–2948.

19. Patterson, D., Gonzalez, J., Le, Q., Liang, C., Munguia, L., Rothchild, D., ... & Dean, J. (2021). Carbon emissions and large neural network training. arXiv preprint arXiv:2104.10350.
20. Zhou, Z., Chen, X., Li, E., Zeng, L., Luo, K., & Zhang, J. (2019). Edge intelligence: Paving the last mile of artificial intelligence with edge computing. *Proceedings of the IEEE*, 107(8), 1738–1762.
21. Laird, B. S., Kocielnik, R., Weld, D. S., & Krishnamurthy, A. (2023). What can large language models do for time series? In *Proceedings of the IEEE International Conference on Big Data (Big Data)*, 3793–3798.
22. Liu, Y., Han, T., Xu, M., & Zhu, X. (2024). Prompting large language models for time series forecasting: A bootstrapping approach. *Proceedings of the AAAI Conference on Artificial Intelligence*, 38(13), 14267–14274.
23. Zhang, J., Kailkhura, B., & Han, T. Y. J. (2023). Mix-calibrated inference: Generalization with confidence in federated learning. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 45(7), 8992–9007.