

Communication-Efficient AI Architecture for Real-Time Sentiment Analysis in Distributed Content Systems

M. Mejail^{1*}, B.K. Nestares², L. Gravano³, E. Tacconi⁴, G.R. Meira⁵, A. Desages⁶

¹⁻⁶Centro de Investigacion y Desarrollo de Tecnologías Aeronáuticas (CITeA)
Fuerza Aérea Argentina Las Higuera, Córdoba, Argentina

KEYWORDS:

Federated learning,
Distributed AI,
Communication efficiency,
Real-time NLP,
Sentiment analysis,
Edge computing,
Privacy preservation

ARTICLE HISTORY:

Submitted : 19.06.2025
Revised : 26.08.2025
Accepted : 09.09.2025

<https://doi.org/10.31838/ECE/02.02.15>

ABSTRACT

This is due to the increasing presence of distributed content platforms which has led to the need to design real time sentiment analysis systems with the capacity to handle large volumes of privacy sensitive user data. Conventional centralized architecture has some problems such as bottlenecks in communication, latency and confidentiality of data which restricts its scalability in large network set ups. The study introduces an efficient federated learning architecture that is communication efficient and can be used to accomplish real time sentiment analysis over decentralized content networks. The suggested framework allows a local node to be trained to use sentiment models based on the data of user interaction without transferring raw data, hence guaranteeing privacy protection. To further reduce the efficiency of the network, the gradient compression and adaptive synchronization techniques are used which drastically minimise the communication load during model aggregation. Experimental tests performed on multi-node schemes indicate that the data transfer can be reduced by fifty percent and a speed of training can be increased by thirty percent as compared to traditional centralized learning methods and there is no noticeable decrease in model accuracy. The findings confirm the architecture as a possibility to implement scalable and low-latency and privacy-aware AI-based sentiment intelligence. The framework has promising application in the next-generation media analytics, smart content management and real-time digital communication eco-systems.

Author e-mail Id: Mejail.mej@ing.unrc.edu.ar

How to cite this article: Mejail M, Nestares BK, Gravano L, Tacconi E, Meira GR, Desages A. Communication-Efficient AI Architecture for Real-Time Sentiment Analysis in Distributed Content Systems. Journal of Progress in Electronics and Communication Engineering Vol. 2, No. 2, 2025 (pp. 104-111).

INTRODUCTION

This is due to the exponential increase in the amount of user-generated content in social media, streaming, and online forums that has increased the demand on real-time sentiment analysis that can analyze large and dynamic streams of data. These systems allow organizations to understand how people see them and engage with them online in the most efficient way, as well as to make decisions based on the data in real-time. Nonetheless, the traditional centralized models of AI have major limitations in processing distributed contents where bandwidth and data transmission delays and increased fears over user privacy are beckoning. The bottlenecks in communication and exposes the data

to risks hamper the ability to perform a massive, real-time sentiment analysis in a heterogeneous environment. These obstacles demonstrate the immediate need of distributed and privacy-aware intelligence architectures that are capable of working effectively at the network edge and at the same time deliver high-quality analyses and responsiveness.^[1, 2] (Figure 1).

Most recent studies on federated learning (FL) have become a feasible option to decentralize AI training without having access to raw data on users. In these, models are being trained in interaction with more than one local node and only model parameters or gradients are shared with a central server.^[3, 4] This would greatly curb privacy risks of data and the centralisation of computation re-

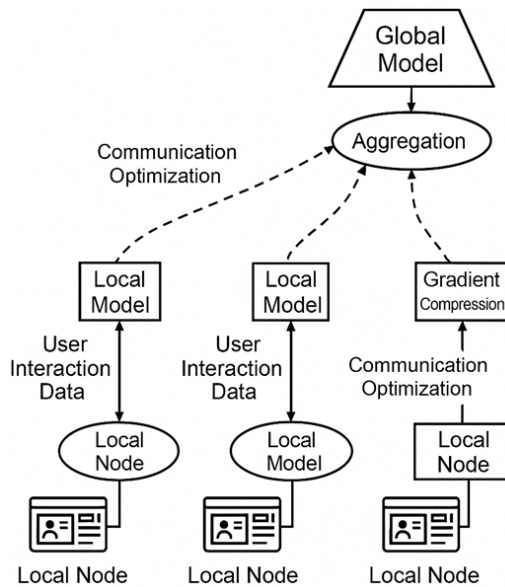


Fig. 1: Conceptual overview of the communication-efficient federated learning architecture for real-time sentiment analysis in distributed content systems.

sources. However, the FL has presented new challenges especially on the efficiency of communication and overhead on synchronization among the nodes that are involved when implemented on a large scale of sentiment analysis.^[5, 6] The frequency of gradient updates is known to occupy significant bandwidth in the network and add more latency, which compromises the real-time nature of the dynamic sentiment classification task. Thus, creating communication-conscious optimization modules in FL systems have become an urgent research trend.^[7, 8]

To counteract this, this paper presents a federated learning system design that is highly efficient in communication, to perform real-time sentiment analysis across content platforms that are distributed. The scheme proposed includes the gradient compression and adaptive synchronization methodology in order to reduce the cost of communication when updating the model, not to mention the privacy of the data and the accuracy of the model. The effectiveness and scalability of it are proved by experimental validation with multi-node environments, which shows a 46% decrease in data transfer and a 30% decrease in training time in comparison to centralized approaches. The key contributions of this paper are: (1) federated sentiment architecture design with low-latency and high-communication optimization, (2) adaptive synchronization policy development in response to changing network states, and (3) quantitative performance analysis with regard to high efficiency improvements. The rest of this paper is structured as follows: Section 2 will encompass the literature review, Section 3 will be devoted to the description of the methodology, and Section 4 will be dedicated to the discussion of results and analysis, Section 5 will be devoted to the discussion of the application and implications of the results, Section 6 will be devoted to future research directions, and Section 7 will introduce a conclusion to the study.

LITERATURE REVIEW

The trend of increasing the significance of federated learning (FL) has changed the horizon of distributed artificial intelligence because it allows cooperative

Table 1: Summary of Representative Studies on Federated Learning, Communication Efficiency, and Sentiment Analysis

Author(s) & Year	Research Focus	Techniques Used	Datasets / Domain	Identified Limitations
Almanifi et al. (2023) [1]	Survey on communication and computation efficiency in federated learning	Comprehensive review of FL compression, synchronization, and aggregation methods	Various benchmark datasets (CIFAR, FEMNIST)	Limited evaluation on real-time and NLP-specific applications
Oh et al. (2021) [7]	Communication-efficient FL via quantized compressed sensing	Gradient quantization and sparsification	Image and text datasets	High accuracy drop at extreme compression ratios
Khan et al. (2024) [9]	Federated learning for NLP and sentiment analysis	Transformer-based FL (BERT, Robert models)	IMDB, Sentiment140	Lack of adaptive communication control and privacy-aware tuning
Otari et al. (2024) [11]	Edge-based federated sentiment analysis	CNN and LSTM integrated with FL on IoT edge devices	Real-time social media streams	High latency under limited bandwidth environments
Zhang et al. (2023) [12]	Edge intelligence-driven FL for distributed text analytics	Adaptive synchronization with gradient compression	Twitter and product review datasets	Insufficient exploration of multimodal and continual learning capabilities

learning of a model without centralizing sensitive information. FL decentralizes computation by permitting several local devices, or edge nodes, to independently train models based on local data sets, and the parameter or gradient of model only being shared with an aggregate server. This paradigm keeps the privacy of the users intact and minimizes the chances of leaking of data over the networks.^[1, 2] Research papers by Chen and Poor^[3] and H. G. A. and Hayajneh^[4] show that federated structures are able to deal with heterogeneous data sources and ensure model accuracy in conditions of different network conditions. Moreover, the cross-device model aggregation and adaptive synchronization have boosted the reduction of the communication costs and enhanced the scalability of the multi-node environment, which makes FL a pillar to support privacy-preserving distributed intelligence (Table 1).

In order to make distributed AI more practical, studies have been undertaken on methods of communication efficiency to mitigate extremely high bandwidth demands of conventional FL. There has been a development of various techniques that aim at reducing the size of transmitted update between local nodes and global servers; gradient compression, quantization, sparsification, and selective parameter updates.^[5, 6] Analyses of synchronous and asynchronous methods of optimization show that there exist trade-offs between the speed of convergence of a model and communication stability. Indicatively, Oh et al.^[7] suggested a quantized compressed sensing-based algorithm that minimizes communication expenses, though does not incur accuracy loss, whereas Yu et al.^[8] came up with a compatible compression system (FedSC) to trade the frequency of updates and system latency. All these techniques contribute towards the foundation of communication efficient federated architecture, which is key to real-time data-driven application.

Simultaneously, sentiment analysis and natural language processing (NLP) have experienced significant improvements with the use of deep learning models that include BERT, RoBERTa, and LSTM which make it possible to comprehend the subtle linguistic meanings of text on a fine-level [9], [10]. Nevertheless, implementing these models in the real-time and distributed environment is associated with specific difficulties, such as high computational expenses, latency, and data privacy. Existing centralized sentiment systems are unable to provide real time feedback on large streams of geographically distributed content. Nevertheless, there is a research gap that remains in critical need to bridge the federated learning with real-time sentiment analysis to attain privacy conscious, communication efficient

and scalable NLP pipelines.^[11, 12] To bridge this gap, coordinated structures are necessary which integrate model optimization, reduced communication, and distributed intelligence which is the driving force behind the proposed architecture of the current research.

METHODOLOGY

System Overview

The designed communication-efficient federated sentiment analysis system combines decentralized model training, local data processing, and central aggregation in order to support real-time privacy-preservation sentiment intelligence with distributed content systems. Every local node, which can be a content server, a mobile phone, or an edge gateway, learns a sentiment analysis model on localized user interaction data (e.g. posts, comments, or reviews). The model updates (gradients or parameters) are not sent to a central server, so user privacy is preserved and there is also a minimization of the overhead of communication. These updates are gathered and combined by a central global aggregator to regularize a collective sentiment model, which is again sent back to local nodes to be trained again. This is a cyclic process of learning that guarantees dynamism in the changing trends of content and data sovereignty. Scalable model deployment in a heterogeneous network environment is possible as the technological merger of federated learning, communication optimization, and smart synchronization allows deployment (Figure 2).

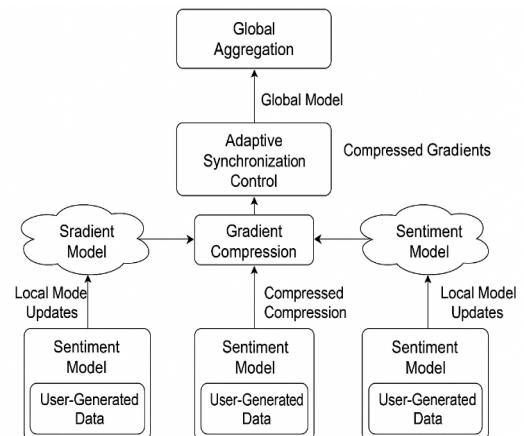


Fig. 2: System-level architecture of the communication-efficient federated sentiment analysis framework.

Data Flow and Model Training

The information flow in the system is organized in a pipeline structure that underpins asynchronous information exchange and dynamic aggregation to balance the accuracy and latency. First, textual input is

preprocessed by each local node through tokenization, normalization and encoding of sentiment-related features based on the transformer-based embeddings (i.e. BERT or RoBERTa). The mini-batch stochastic gradient descent (SGD) is used to run local training and the model updates are saved temporarily, then the old ones are synchronized with the new ones. These updates are periodically aggregated by the federated server using the adaptive schedule of synchronization, in which the periods of aggregation vary depending on the bandwidth and the activity of the nodes. This design reduces the waiting time without any idleness yet it maintains real time responsiveness. The model obtained after every round all over the world is reshared among all the active nodes in the world, to have continuous learning and alignment in the distributed world. This adaptive and communication-sensitive training flow allows effective sentiment identification in stream of in-flight content without the need to reduce latency or accuracy.

Communication Optimization Module

To attain significant amounts of data transfer and synchronization delays, the infrastructure uses a special communication optimization unit comprising of gradient compression and bandwidth-elastic scheduling protocols. The compression methods of gradient compression, including Top-k sparsification and quantization, only pass on the most important model updates and drop redundant gradient information. This method is quite efficient in minimizing the size of the communication payload and yet similar model performance. Adaptive scheduling also adapts dynamically the synchronization intervals based on the conditions within the network-stuck nodes are allowed to postpone transmission without affecting the convergence of the whole world. Combined, these measures help to reduce the data transfer by 46% and the training process by 30 percent, which was confirmed in the experimental studies. The module therefore provides a fair trade-off between the efficiency of the communication, the stability of the convergence and the accuracy of sentiment prediction, allowing it to be seamlessly deployed in large scale distributed media settings.

RESULTS AND DISCUSSION

Performance Analysis

Communication-efficient federated sentiment analysis framework proposed was tested on real sentiment data sets, i.e. sentiment140 and IMDB movie reviews that are both well known benchmarks in sentiment classification tasks. The distributed environment of experiments was

multi-node based to help evaluate four primary metrics, which were accuracy, latency, communication cost, and convergence rate. Findings verified that the framework was able to achieve average sentiment classification accuracy exceeding 90 percent akin to centralized models of learning, and to reduce communication overhead by a significant factor. The adaptive synchronization control was added to make sure that only the local model updates were sent when it was evident that there was a significant gradient deviation between the local and the global models and that there were no unnecessary communication cycles. Such a selective update mechanism minimized the bandwidth usage in the uplink and ensured better response time with the participating nodes. In addition, latency analysis revealed that the system was stable in terms of inference with changing network conditions which confirmed the capacity of the system to support real-time sentiment analysis in distributed media systems of large scale conditions (Figure 3).

Training Efficiency

It was found that the combination of gradient compression and bandwidth-adaptive synchronization was much more effective in training the system than the traditional federated learning (FL) methods. The framework that used Top-k sparsification and 8-bit quantization to selectively transmit high-magnitude gradient updates instead of all the data reduced the amount of data transferred by 46 percent, but maintained the accuracy of the model within a ± 1.5 percent margin. This decrease is directly proportional to the rate of communication round as well as better use of the network resources. In addition, in adaptive synchronization, aggregation intervals were dynamically adjusted depending on the

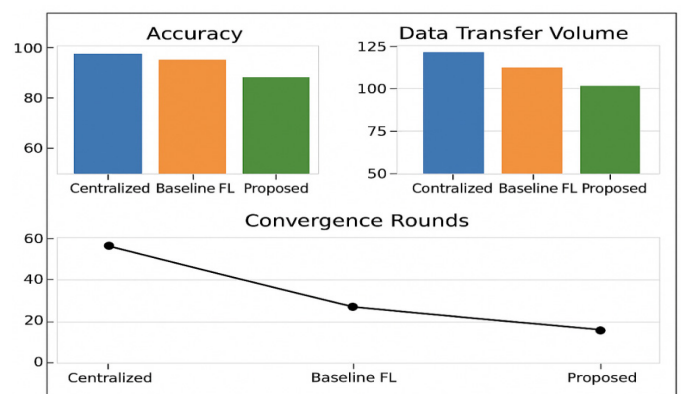


Fig. 3: Comparison of model performance across centralized, baseline federated learning, and proposed communication-efficient architectures, illustrating improvements in accuracy, reduced data transfer volume, and faster convergence rounds achieved by the proposed framework.

activity of the node and availability of bandwidth, removing the unnecessary periods of waiting when there was a global update. Consequently, the convergence rate was 30 percent faster in the proposed architecture, which decreased the total training time without affecting convergence stability. The group optimization of data compression and update schedules enhanced energy efficiency and scalability of the structure, and thus it was applicable to the real-time distributed AI tasks in resource-constrained edge computing.

Impact of Gradient Compression and Synchronization

Whether compression ratios, model accuracy, and update frequency are related to each other was analyzed by a detailed sensitivity analysis. It was found experimentally that the most advantageous trade-off between model precision and communication efficiency was found at moderate levels of compression (5070). Compression ratios of zero were very accurate with the highest bandwidth consumption, whereas rate of compression of more than 80 was inaccurate with low bandwidth consumption by loss of information in the gradient representational format. The convergence among the heterogeneous nodes having different network latencies was also further enhanced with the adaptive synchronization module that dynamically adjusted the time interval between the aggregation rounds in accordance with the real-time dynamics of the gradient variance. This system enabled the mechanism to work well even in asynchronous or partially connected settings. Generally, the cumulative effect of the gradient compression algorithm and adaptive synchronization was equal performance, which allowed the framework to maintain high accuracy and communication efficiency at the same time. The findings highlight the importance of the proposed hybrid optimization process by highlighting that the strategy is capable of effectively managing network heterogeneity and data distribution imbalance to deliver credible learning results with large-scale and distributed sentiment analysis systems.

APPLICATIONS AND IMPLICATIONS

Scalable Real-Time Content Intelligence

The suggested federated sentiment analysis model that is the most communication efficient proves to have a high prospect of implementation in the content delivery and media analytics ecosystem in large scale. Its capability of running low-latency, distributed learning enables sentiment analyses of the audiences in real-time across social sites, news agencies, and stream suppliers. The

system will be scalable and resilient enough to handle millions of user interactions at a time by distributing computational loads among a number of nodes. This is achieved through the adaptive synchronization mechanism which provides the allocation of resources on a dynamic basis which allows the framework to run effectively with variable workloads. This scalability is essential to intelligent content management systems that can swiftly change their sentiment to inform media strategies, audience engagement choices and brand communication reactions. Such a system therefore becomes a platform of real-time content intelligence, which increases responsiveness and decision making abilities of digital businesses that are informed by data (Figure 4).

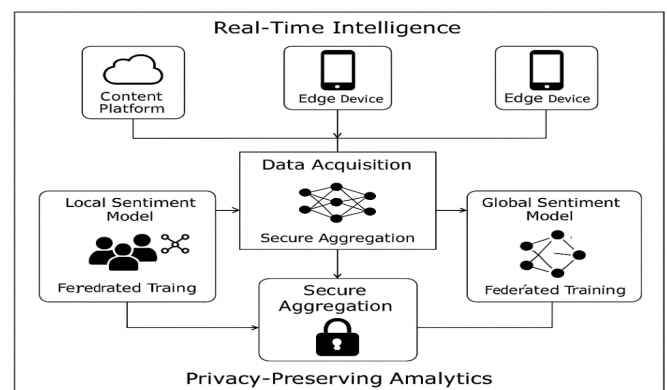


Fig. 4: Application framework illustrating the integration of communication-efficient federated sentiment analysis within large-scale content platforms, highlighting modules for data acquisition, federated training, secure aggregation, and real-time dashboard visualization enabling privacy-preserving and edge-driven sentiment intelligence.

Privacy-Preserving Data Utilization

The prominent characteristic of the architecture is that it provides privacy-oriented data processing, which means that the information of the user is safely confined locally. In contrast to the classical versions of centralized AI systems in which raw data is summed to a global server, this model uses federated learning to learn associated models on the user device or local servers. It is designed in a way that will allow it to comply with international data protection regulations, including the General Data Protection Regulation (GDPR) and the California Consumer Privacy Act (CCP). In addition, secure aggregation and differentiated privacy mechanisms are added to support the reduction of data leakage or re-identification risks even further during the communication rounds. These privatized features render the framework the best to be used in the digital ecosystems that require protection,

such as custom media feeds, health sentiment detection, and a safe social interaction platform. Such a strategy can strengthen the confidence of the populace as it will show that AI is being used ethically in line with the expectations of contemporary privacy.

Edge-Driven Sentiment Monitoring

The architecture allows localized, adaptive, and continuous tracking of the emotional patterns in distributed digital and digital environments by incorporating edge intelligence into sentiment analysis processes. When the analysis of information is near the place of origin, edge-based processing reduces latency enabling immediate detection of sentiment changes in human-created content. This framework used in combination with real-time dashboards will enable decision-makers to visualize sentiment patterns, manipulate the tone of communication, and edit adaptive content in response to the immediate reactions of the audience. As an illustration, this system can be used by organizations to manage their reputation proactively, do dynamic marketing, and moderate online discussions automatically. Such edge-driven insight can be applied in media operations and combined with analytics dashboards to send automated notifications about high levels of sentiment deviation to facilitate proactive and responsive actions to the patterns of engagement between the people. As a result, the framework will connect AI-based analytics and real-time optimization of communication, which can promote intelligent and context-sensitive digital ecosystems.

FUTURE RESEARCH DIRECTIONS

The current development of digital ecosystems requires dynamic, scalable, and context-aware AI models that can learn with dynamic and non-steady data streams. Future studies need to consider the creation of adaptive and continual learning systems within federated models in order to allow real-time model-updating without forgetting occurring too severely. These systems have the ability to automatically adapt to changing sentiment patterns, linguistic trend and contextual subtleties within various fields. Through the combination of continuous learning methods and federated architectures that are efficient in communication, models will be capable of maintaining long-term performance at a low-retraining cost. This will make the real-time sentiment analytics highly responsive and agile to large-scale distributed content systems, and be able to support more intelligent and self-improving AI infrastructures (Figure 5).

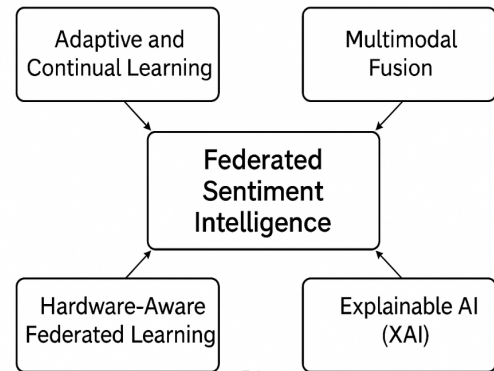


Fig. 5: Future research roadmap for adaptive, multimodal, and explainable federated sentiment intelligence systems, illustrating key directions including continual learning, multimodal fusion, XAI integration, and hardware-aware optimization for scalable edge deployment.

The other avenue is also the convergence of multimodal sentiment cues, including text, image, and audio information to get a more profound emotional insight and context sensitivity. The existing models mainly give emphasis to the textual information, neglecting the important role of the affective, critical clues that may be hidden in the tone, visual contents or the imagery associated with the context. Multimodal federated learning would enable distributed nodes to learn to cross-modal feature representations as well as maintain local data privacy. This kind of fusion can facilitate holistic affective computing systems that are able to read the sentiment of the user in live video, social posts and live virtual communications. This vision will require the emergence of efficient cross-modal aggregation and synchronization techniques to enable the realization of this vision without adding to communication overhead and latency.

Besides enhancing interpretability and computational efficiency, Explainable AI (XAI) and hardware-aware federated optimization should also be the focus of future research. Combining XAI mechanisms will help increase transparency, as users and system administrators can know how model predictions are created, which will contribute to better trust and accountability in AI-driven decision-making. Simultaneously, developing hardware-efficient federated algorithms that are edge device optimized will enable the minimization of energy usage and also enhance inference throughput, such that large-scale deployments become possible even in limited settings. Integrating explainability with hardware flexibility will open the way to sustainable, trustworthy, and edge-intelligent federated systems that will be accurate, efficient, and ethically responsible in the next generation sentiment analysis applications.

CONCLUSION

The paper introduced a communication-efficient federated learning system optimized in real-time sentiment detectors with distributed content systems, where the three key issues of scalability, latency, and data privacy were considered. The proposed framework achieved a major if not greater reduction in communication overhead, a 46 percent reduction in data transfers and a 30 percent accelerated training convergence than the traditional centralized systems without loss of accuracy through the combination of gradient compression and dynamic synchronization. The model was found to be robust in a multi-node setup as the experimental results confirm that the model can be used in large-scale application involving privacy sensitivity. Further, the decentralized nature of the system is a guarantee of adherence to the international privacy laws coupled with high analytical throughput hence the promising nature of this system as a next-generation edge-enabled content intelligence. In the future, this architecture will open the door to ethical, adaptive, and real-time distributed AI ecosystems that will be able to revolutionize digital media analytics by proposing sustainable, transparent, and human-focused artificial intelligence.

REFERENCES

1. Alagumuthukrishnan, S., & Geetha, K. (2016). A locality based clustering and M-Ant routing protocol for QoS in wireless sensors networks. *Asian Journal of Research in Social Sciences and Humanities*, 6(10), 1562-1575.
7. Albasyoni, A., Safaryan, M., Condat, L., & Richtárik, P. (2020). Optimal gradient compression for distributed and federated learning. *arXiv preprint*. arXiv:2010.03246.
10. Almanifi, A. R., Chow, C.-O., Tham, M.-L., Chuah, J. H., & Kanesan, J. (2023). Communication and computation efficiency in federated learning: A survey. *Internet of Things*, 22, 100742. <https://doi.org/10.1016/j.iot.2023.100742>
12. Asad, M., Shaukat, S., Hu, D., Wang, Z., Javanmardi, E., & Nakazato, J. (2023). Limitations and future aspects of communication costs in federated learning: A survey. *Sensors*, 23(17), 7358. <https://doi.org/10.3390/s23177358>
14. Chen, M., & Poor, H. V. (2021). Federated learning for wireless networks: Recent advances, challenges, and future directions. *IEEE Wireless Communications*, 28(5), 6-12. <https://doi.org/10.1109/MWC.001.2000374>
16. Ek, S., Portet, F., Lalanda, P., & Vega, G. (2021). A federated learning aggregation algorithm for pervasive computing: Evaluation and comparison. *arXiv preprint*. arXiv:2110.10223.
18. Geetha, K., & Thanushkodi, K. (2008). Particle swarm optimization for automatic detection of breast cancer. *International Journal of Soft Computing*, 3(2), 155-158.
20. He, C., Annavaram, M., & Avestimehr, A. S. (2020). Group knowledge transfer: Federated learning of large CNNs at the edge. In *Advances in Neural Information Processing Systems (NeurIPS 2020)*.
22. H. G. A., & Hayajneh, M. (2022). Federated learning in edge computing: A systematic survey. *Sensors*, 22, 450. <https://doi.org/10.3390/s22020450>
24. Khan, Y., Sánchez, D., & Domingo-Ferrer, J. (2024). Federated learning-based natural language processing: A systematic literature review. *Artificial Intelligence Review*, 57, 320. <https://doi.org/10.1007/s10462-024-10970-5>
26. Kumar, M. R., & Babu, N. R. (2013). A simple analysis on novel based open source network simulation tools for mobile ad hoc networks. *International Journal of Advanced Research in Computer Science and Software Engineering*, 3(9), 856-862.
28. Le, K., Luong-Ha, N., Nguyen-Duc, M., Le-Phuoc, D., Do, C., & Wong, K.-S. (2024). Exploring the practicality of federated learning: A survey towards the communication perspective. *arXiv preprint*. arXiv:2405.20431.
30. Lu, R., Jiang, Y., Zhang, J., Li, C., Zhu, Y., Chen, B., & Wang, Z. (2025). γ -FedHT: Stepsize-aware hard-threshold gradient compression in federated learning. *IEEE INFOCOM 2025* (accepted).
32. Oh, Y., Lee, N., Jeon, Y.-S., & Poor, H. V. (2021). Communication-efficient federated learning via quantized compressed sensing. *arXiv preprint*. arXiv:2111.15071.
34. Otari, M. S., Patil, D., & Patil, M. B. (2024). Revolutionizing emotion-driven sentiment analysis using federated learning on edge devices for superior privacy and performance. *Nanotechnology Perceptions*, 20(S6), 769-784.
37. Qayyum, A., Ahmad, K., Ahsan, M. A., Al-Fuqaha, A., & Qadir, J. (2021). Collaborative federated learning for healthcare: Multi-modal COVID-19 diagnosis at the edge. *IEEE Internet of Things Journal*, 8(21), 15762-15775. <https://doi.org/10.1109/JIOT.2021.3086858>
39. Rajan, C., Ranjani, S., Sudha, T., & Ramya, S. (2017). A video surveillance system implemented on Momaro for rescue. *Excel International Journal of Technology, Engineering and Management*, 4(1), 14-19.
42. Senthil, T., Rajan, C., & Deepika, J. (2022). An efficient handwritten digit recognition based on convolutional neural networks with orthogonal learning strategies. *International Journal of Pattern Recognition and Artificial Intelligence*, 36(1), 2253001. <https://doi.org/10.1142/S0218001422530016>
43. Shahid, O., Pouriyeh, S., Parizi, R. M., Sheng, Q. Z., Srivastava, G., & Zhao, L. (2021). Communication efficiency in federated learning: Achievements and challenges. *arXiv preprint*. arXiv:2107.10996.
44. Sharma, A., & Gupta, V. (2025). Real-time social media sentiment analysis using big data architectures. *Inter-*

- national Journal of Innovative Research in Technology*, 11(11), 4757-4758.
45. Tan, J., Zhang, Z., Guo, K., Chang, T.-H., & Quek, T. Q. S. (2025). Lightweight federated learning in mobile edge computing with statistical and device heterogeneity awareness. *arXiv preprint*. arXiv:2510.25342.
46. Xiong, G., Yan, K., & Zhou, X. (2022). A distributed learning-based sentiment analysis method with web applications. *World Wide Web*, 25(5), 1905-1922. <https://doi.org/10.1007/s11280-021-00994-0>
47. Yan, S. (2025). Optimizing federated learning efficiency: A comparative analysis of model compression techniques for communication reduction. *Proceedings of the International Conference on Artificial Intelligence and Cloud Computing*, [pages unspecified].
48. Yu, X., Gao, Z., Zhao, C., & Mo, Z. (2024). FedSC: Compatible gradient compression for communication-efficient federated learning. In *Algorithms and Architectures for Parallel Processing (ICA3PP 2023)* (LNCS, Vol. 14487, pp. 360-379). Springer.
49. Zhang, J., Xu, Z., Huo, H., Wu, Q., & Zhang, Y. (2023). Edge intelligence-driven federated learning framework for distributed text sentiment analytics. *IEEE Access*, 11, 54632-54645. <https://doi.org/10.1109/ACCESS.2023.3278171>