

A Review of the Framework for Privacy-Preserving Sentiment Analysis in Quantum-Enhanced Federated Learning

Hongzhen Yu ^{1,2*}, Hazirah Bee Binti Yusof Ali ¹ and M. Kazem Chamran ¹

¹Faculty of Information Technology, City University Malaysia, Darul Ehsan, Malaysia
²Business College, Jiangxi Institute of Fashion Technology, Nanchang, Jiangxi Province, China
* Correspondence: Hongzhen Yu

Abstract: The rapid growth of real-time consumer sentiment analysis on e-commerce, retail, and social media platforms has intensified the conflict between computational efficiency and data privacy. Federated learning enables decentralized privacy-preserving training, but it still faces communication overhead, unstable convergence, and non-IID data challenges. Quantum-enhanced Transformers show strong potential for high-dimensional NLP processing, yet their deployment is limited by quantum noise, scalability issues, and security risks. This paper systematically reviews the integration of federated learning, quantum machine learning, and sentiment analysis. It discusses key privacy-preserving techniques, including differential privacy, secure multi-party computation, and quantum homomorphic encryption, while analyzing quantum-driven NLP innovations and hybrid federated quantum architectures. The review identifies major gaps in real-time applicability, quantum circuit security, and the trade-off between encryption strength and system latency. Based on these findings, the paper proposes future directions such as optimized quantum homomorphic encryption, domain-specific modular architectures, and adaptive federated protocols. It provides a roadmap for developing efficient, scalable, and secure real-time hybrid quantum federated sentiment analysis systems.

Keywords: first term; second term; third term; fourth term; fifth term; sixth term

1. Introduction

Real-time emotion analysis constitutes a pivotal application across diverse domains, including social media monitoring, customer feedback analytics, and financial market sentiment assessment. It aims to dynamically detect the affective valence—particularly emotional polarity and intensity—embedded in textual data, thereby enabling organizations to track, interpret, and respond to evolving public sentiment with temporal precision. In contrast to conventional sentiment analysis approaches—which typically operate on static, batch-processed corpora—real-time emotion analysis processes high-velocity, continuous data streams, performing inference

incrementally as textual inputs arrive. This capability proves especially critical in time-sensitive contexts characterized by rapid sentiment shifts, such as breaking news events, product releases, or live political discourse. A key technical challenge remains the scalable and low-latency processing of voluminous, heterogeneous, and often noisy real-time data—without compromising analytical accuracy or interpretability. With the acceleration of digital transformation, the applicability and strategic importance of real-time emotion analysis continue to expand across various sectors. Table 1 summarizes its application scope.

Table 1. Scope of application for real-time emotional analysis

Field of Application	Specific Application Scenarios and Examples
Business & Customer Service	-Customer service systems: real-time dialogue sentiment recognition to mitigate customer complaints. -E-commerce live streaming: analyzing live chat messages to detect sentiment and optimize operations and scripts. -Offline retail: real-time customer satisfaction tracking via interactive feedback devices. -Social media platforms: NLP-based sentiment classification of user comments for brand reputation and public opinion management.
Social Governance & Public Safety	-Public opinion monitoring: integrates multi-source online data for rapid sentiment analysis and early warning of public opinion risks. -Mental health and behavioral risk monitoring: uses audio and video recognition to detect psychological risks, applicable in campuses and judicial systems. -Smart city governance: analyzes public feedback from governmental and social media data to evaluate public sentiment and support decision-making.
Financial Market Regulation	-Quantitative trading and investment: extracts sentiment signals from news, financial reports, and social media to support sentiment-driven trading strategies. -Credit risk control: analyzes sentiment in online behavior and reviews of SMEs to complement traditional credit assessment and identify potential risks.
Education	-Student mental health screening: uses emotion recognition tools in campus environments for rapid psychological assessment. -Teaching and remote evaluation: analyzes real-time emotional states from video platforms to assess attention and psychological status. -Personalized learning path recommendation: adapts learning difficulty and resources based on students' emotional feedback.
Healthcare	-Remote diagnostic assistance: uses speech emotion analysis to support physicians in evaluating patients' psychological states. -Early cognitive impairment screening: analyzes speech, facial expressions, and language patterns for early detection of Alzheimer's disease and related disorders.

Sentiment analysis, as a core branch of natural language processing, has shifted from the traditional batch

processing mode to a highly dynamic real-time analysis paradigm. Traditional methods rely on rule-based approaches and classical machine learning models, such as support vector machines (SVM) and naive Bayes classifiers. Although these models are effective in basic scenarios, they have difficulties in capturing complex language nuances and context dependencies. Modern real-time sentiment analysis largely relies on natural language processing techniques, especially Transformer-based models such as BERT and GPT.

Recent advancements in Transformer architectures have significantly improved sentiment analysis performance. These models effectively overcome several limitations of previous methods. BERT can process text in both directions, capturing contextual relationships from preceding and succeeding words, enabling it to more effectively understand subtle emotional cues and ambiguities within the text. GPT is mainly designed for text generation but also performs well in sentiment analysis due to its ability to model long-term dependencies in text data. After fine-tuning for specific domain datasets, these models achieve remarkable accuracy in sentiment classification tasks. For instance, BERT achieved an 94% accuracy rate in the sentiment classification task involving synthetic text generated by GPT-2, demonstrating its robustness in extracting emotional-related features [1]. Additionally, studies have shown that the performance of BERT is highly sensitive to fine-tuning parameters, while GPT-2 exhibits greater stability during training and fine-tuning [2]. These insights highlight the adaptability of Transformer-based models in handling sentiment analysis across different datasets and contexts, making them a key component of modern sentiment analysis frameworks.

With growing concerns over data privacy and the explosive increase in user data volumes, coupled with stringent privacy regulations such as GDPR, traditional machine learning approaches that centrally process data are increasingly exposed to drawbacks in privacy, efficiency, and scalability. Federated learning has emerged as a solution, enabling local data processing and distributed training without requiring access to raw data, thereby effectively mitigating privacy and security risks.

However, federated learning still faces several challenges. The rapid growth of internet data far exceeds current processing capabilities, resulting in an increasingly severe problem of information overload [3]. In addition, the frequent exchange of model updates leads to high communication costs, particularly in large-scale federated learning systems. This communication overhead also slows model convergence. Another major challenge is the heterogeneity of client devices and local datasets. Differences in computing resources and the non-independent and identically distributed (non-IID) nature of local data can significantly reduce model accuracy and training efficiency. These issues limit the scalability and practical deployment of federated learning in real-world applications.

Quantum-enhanced Transformers leverage fundamental principles of quantum computation, including superposition, entanglement, and parallel state evolution,

to accelerate high-dimensional tensor operations, reduce training latency, and improve hardware efficiency. Their compatibility with federated learning architectures has been empirically demonstrated [4]. In this hybrid paradigm, the quantum-enhanced Transformer functions as a high-capacity local learner, enabling effective representation learning over heterogeneous, high-dimensional, and non-IID data through parameterized quantum circuits that facilitate parallel feature encoding and efficient gradient optimization. Meanwhile, federated learning preserves data locality by enabling decentralized model training and collaborative optimization without exposing raw data, thereby providing strong privacy guarantees while alleviating key challenges of conventional distributed learning, including communication overhead, straggler effects, and degraded convergence under statistical heterogeneity. Recent advances in quantum circuit compilation, adaptive gate scheduling, and error-mitigation techniques have further improved the robustness, scalability, and practical feasibility of quantum inference, paving the way for deployment on large-scale, resource-constrained federated systems. By tightly integrating quantum acceleration with privacy-preserving distributed learning, this hybrid framework establishes a scalable and trustworthy architecture for computationally intensive and latency-sensitive applications, such as real-time multimodal sentiment analysis, intelligent recommendation systems, and streaming high-density sensor analytics, representing a promising direction for next-generation distributed artificial intelligence.

In conclusion, federated learning provides a strong foundation for privacy-preserving decentralized machine learning. Integrating it with quantum-enhanced Transformers forms a hybrid framework that has the potential to address major challenges, including communication overhead, slow model convergence, and computational inefficiency. This framework combines the privacy-preserving and decentralized characteristics of federated learning with the computational advantages of quantum-enhanced models. As a result, it enables more efficient, scalable, and secure real-time sentiment analysis. Beyond sentiment analysis, this hybrid approach also promotes the convergence of quantum computing, machine learning, and privacy-preserving artificial intelligence. It offers promising applications in domains such as e-commerce, social media, and digital marketing. Moreover, it helps bridge the gap between privacy protection and computational efficiency in modern AI systems.

This paper presents a comprehensive review of quantum federated learning for privacy-preserving sentiment analysis. Section 2 introduces the fundamental concepts of federated learning, quantum machine learning, and their integration, together with their respective strengths and limitations. Section 3 reviews recent advances and discusses the key challenges of quantum federated learning for privacy-preserving sentiment analysis. Finally, Section 4 summarizes the main findings and outlines future research directions and potential application prospects.

2. Federated Learning and Quantum-enhanced Transformer

2.1. Federated Learning

Federated learning is a distributed machine learning paradigm designed to address data privacy concerns while enabling collaborative model training. Unlike traditional centralized approaches, it does not require raw data to be collected on a central server. Instead, the training process is distributed across multiple edge devices, such as smartphones, IoT devices, and local servers, while the data always remain on the local devices. During training, each client updates the model using its local dataset. Only the model parameters, such as weights or gradients, are transmitted to a central server for aggregation. The server combines these local updates using the Federated Averaging (FedAvg) algorithm to produce an improved global model [5]. Because raw data are never shared, this approach effectively reduces the risks of privacy leakage and data misuse. The global model is then redistributed to the participating clients for the next round of local training. This iterative process continues until the model reaches the desired performance [6]. The overall architecture of federated learning is illustrated in Figure 1.

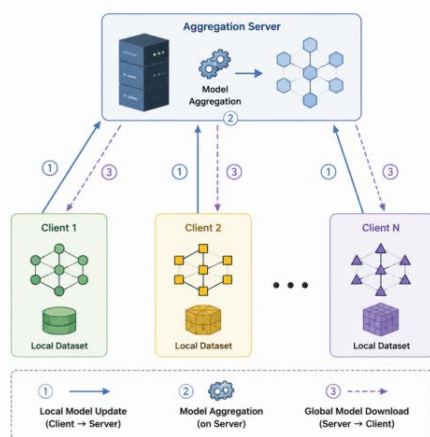


Figure 1. Federated learning

Table 2. Characteristics and applications of quantum machine learning

Feature / Application	Description	Example / Application Scenario
Data Processing Capability	Utilizes quantum superposition and quantum parallelism to efficiently process and analyze large-scale and high-dimensional datasets that are difficult for classical computers to handle.	High-dimensional image recognition, genomic sequence analysis, large-scale network traffic pattern mining.
Optimization Acceleration	Quantum algorithms (such as quantum annealing and variational quantum eigensolvers) can accelerate the solution of complex optimization problems in machine learning model training.	Quantum support vector machines (QSVM), parameter optimization in quantum neural networks (QNN).
High-Dimensional Feature Mapping	Represents complex data in Hilbert space using quantum states, enabling implicit mapping into exponentially high-dimensional feature spaces that are difficult for classical methods to capture, thereby enhancing feature representation capability.	Quantum kernel methods: using quantum inner products to implicitly map data into high-dimensional feature spaces for improved classification performance.
Model Acceleration and Compression	Quantum circuits reduce the number of model parameters or computational steps, thereby decreasing training time and computational complexity.	Quantum-enhanced deep learning models, such as quantum convolutional neural networks and quantum Transformers.

2.3. Transformer and Quantum-enhanced Transformer

The Transformer model, proposed by Vaswani et al. [9], is a sequence modeling architecture entirely based on the

attention mechanism. Compared with traditional centralized machine learning, federated learning offers several significant advantages. First, data remain on local devices throughout the training process, which greatly reduces the risk of data leakage and improves compliance with privacy regulations. Moreover, its decentralized architecture minimizes the security risks associated with centralized data repositories, which are common targets of cyberattacks [7]. Federated learning also provides excellent scalability. Because model training is performed in parallel across multiple edge devices, the framework can efficiently support large-scale deployments. To further strengthen privacy protection, federated learning incorporates advanced techniques such as differential privacy and secure aggregation.

Differential privacy protects sensitive information by injecting carefully calibrated noise into model updates. This mechanism prevents attackers from inferring individual user data from shared gradients or model parameters [8]. Secure aggregation provides an additional layer of protection. It enables the central server to access only the aggregated model updates rather than the contributions of individual clients, thereby preserving user privacy and anonymity. Another key advantage of federated learning is its ability to train models under non-independent and identically distributed (non-IID) data settings. In real-world applications, the data collected by different devices are often heterogeneous and follow different distributions. Federated learning is specifically designed to accommodate this heterogeneity, making it well suited for practical deployment across diverse user environments.

2.2. Quantum Machine Learning

Quantum machine learning is an interdisciplinary field that combines quantum computing and machine learning. It aims to optimize the efficiency and complexity of traditional machine learning algorithms by leveraging quantum superposition and entanglement properties. It is mainly applied to large-scale data processing and high-dimensional optimization problems. The main characteristics and applications of quantum machine learning are summarized in Table 2.

attention mechanism. Unlike traditional recurrent neural networks, which rely on step-by-step recursive computation over time, the Transformer directly models dependencies between any positions in a sequence through

the self-attention mechanism, thereby enabling parallel computation of global contextual information. Its core operation is scaled dot-product attention, which computes similarity scores among queries (Q), keys (K), and values (V), and then normalizes these scores to obtain weighted representations. This mechanism allows the model to dynamically focus on important features, significantly improving training efficiency and enhancing its ability to model long-range dependencies [9].

Although the Transformer has achieved remarkable success in various domains, the spatiotemporal complexity of its self-attention mechanism grows quadratically with the sequence length, resulting in extremely high computational and memory costs. This limitation has motivated researchers to explore new computational paradigms, such as integrating quantum computing mechanisms into Transformer architectures. Related studies have shown that quantum representations can achieve more compact feature encoding in high-dimensional Hilbert spaces. In addition, parameterized quantum circuits can enhance model expressiveness. These findings provide a theoretical foundation for the development of quantum-enhanced Transformers [10].

The Transformer excels in sequence modeling due to its self-attention mechanism, but this mechanism's computational complexity increases quadratically with the length of the input sequence, severely limiting its ability to handle long texts. This bottleneck is particularly significant in real-time sentiment analysis scenarios. Mainstream sentiment analysis tasks such as long social media comments, live chat bulletins, and financial report analysis all require capturing long-range sentiment dependencies, and thus have extremely high requirements for the model's ability to process long sequences.

Quantum computing has a deep mathematical alignment with the Transformer's self-attention mechanism, providing a new direction for solving this problem: On one hand, the large-scale matrix multiplication and inner product operations at the core of self-attention are types of operations that quantum computing theoretically has a natural advantage in accelerating; on the other hand, variational quantum circuits can efficiently approximate and represent high-dimensional functions, providing a feasible technical path for the implementation of quantum attention mechanisms.

In summary, the quadratic computational complexity of the classical Transformer makes it difficult to meet the scalability requirements of long text sentiment analysis tasks. However, the unique advantages of quantum computing in accelerating matrix operations and mapping high-dimensional features provide a practical and feasible technical solution to break through this bottleneck. The following text will focus on introducing the specific architecture of quantum-enhanced Transformer and its application in sentiment analysis.

2.4. Quantum-enhanced Federated Learning

Despite having great potential, quantum-enhanced Transformers still face challenges related to hardware. Current noisy medium-scale quantum devices are plagued

by quantum noise, limited number of qubits, and gate errors, which affect the reliability and scalability of the models [11]. Moreover, encoding large datasets into quantum states is still computationally intensive and resource-constrained for practical deployment. In the context of federated learning, quantum-enhanced Transformers offer promising solutions to address communication overhead, non-independent and non-identically distributed data distributions, and model convergence challenges. These Transformers can perform computationally intensive tasks locally on quantum devices, minimizing communication rounds and optimizing aggregation processes. Additionally, they demonstrate improved model convergence on heterogeneous datasets, suitable for distributed sentiment analysis and recommendation systems [12].

However, combining the two also introduces new complexities. The synchronization across distributed quantum nodes is more challenging due to the stability and speed limitations imposed by quantum communication protocols. Studies have shown that quantum-enhanced Transformers are particularly effective in handling nonlinear data distributions, which is a common issue in federated learning environments. The quantum circuits embedded in the Transformer architecture allow for parallel processing and enhanced feature extraction, thereby improving model accuracy and reducing latency [13]. However, to realize its full potential, further progress is still needed in quantum communication protocols and hardware stability.

3. Progress and Limitations

3.1. Research Progress

3.1.1. Background and motivation

Today, industries such as retail, entertainment, and e-commerce heavily rely on real-time data analysis to provide personalized recommendations to customers. This requires the collection of a large amount of consumer information, including purchasing habits, preferences, and personal opinions. Ensuring that this information is properly protected and complies with privacy policies such as GDPR has become crucial. The massive user data generated by online social networks, e-commerce platforms, and entertainment requires that artificial intelligence systems be able to analyze emotions in real time, which is essential for marketing strategies, customer service, and recommendation systems. However, the traditional centralized data collection and analysis model has raised deep concerns about data security vulnerabilities, intrusions, and abuse [14]. This has made users more concerned about privacy and even may refuse services that require their information. Therefore, federated learning provides an active solution that allows model training to be conducted without transmitting data to a central server, with the data remaining on local devices. This decentralized method only shares model updates, fundamentally reducing the possibility of exposing the original data [15]. However, at the same time, federated learning also faces challenges such as high communication

costs between nodes and inconsistent model performance on different devices [16].

Given these limitations, the necessity of introducing quantum-enhanced Transformers to improve system efficiency and scalability has become increasingly evident. Quantum computing utilizes the principles of superposition and entanglement, providing new opportunities for accelerating machine learning and enabling complex computations to be executed more efficiently than classical systems. For instance, quantum-based Transformers, as a form of quantum machine learning models, have shown the potential to reduce computation time and enhance performance in natural language processing tasks such as sentiment analysis.

Quantum-enhanced models can handle massive amounts of data in a single pass, which is highly valuable in real-time applications that require rapid processing and response. Integrating quantum-enhanced Transformers into the federated learning framework has the potential to overcome the core challenges faced by real-time sentiment analysis and personalized recommendations. Specifically, quantum computing can address the computational complexity brought by consumer sentiment data, and the quantum-enhanced Transformer will reduce latency in real-time processing and improve data utilization efficiency, thereby providing more timely and accurate recommendations. This integration aims to design a unified framework that balances privacy, accuracy, and computational efficiency to solve the current problem where real-time sentiment analysis technologies often need to make trade-offs among these three aspects.

3.1.2. Key technological advances

In recent years, significant progress has been made in three main areas of related research:

(1) Privacy Protection and Performance Verification of Federated Emotional Learning

Federated learning has been proven to be an effective technique for protecting privacy in machine learning, and its application in emotional analysis has been preliminarily verified. By processing data using local devices, federated learning reduces the possibility of exposing sensitive consumer information and significantly enhances privacy. Studies have shown that in the distributed learning framework, the accuracy of emotional analysis and the privacy protection ability for different emotional types have both been strengthened [17] further explored the feasibility of extending large-scale language models such as BERT through federated learning, confirming that even for such complex models, federated learning can complete training without significant performance loss, laying the foundation for developing real-time emotional analysis systems that maintain accurate results while ensuring data confidentiality.

(2) Model Innovation of Quantum-Enhanced and Quantum-Improved Transformer

Quantum computing offers new opportunities for enhancing the efficiency and speed of machine learning. The Quantum Enhanced Transformer integrates quantum features such as quantum attention mechanisms, etc., while

maintaining or even improving accuracy, significantly reducing the computing time for tasks like sentiment analysis. On this basis, the concept of Quantum Improved Transformer was proposed, which focuses on optimizing the design itself to achieve practicality and scalability. This involves improving quantum circuits, optimizing quantum gates, and applying quantum algorithms to enhance the scalability and robustness of the model, making it more reliable in different data environments. For example, the Transformer quantum state model demonstrates how this improvement enhances the model's adaptability to diverse tasks and datasets. The difference between the two lies in that the Quantum Enhanced Transformer focuses on introducing basic quantum features to enhance performance, while the Quantum Improved Transformer pays more attention to ensuring the practicality and scalability of these features in real-world applications.

(3) Preliminary exploration of quantum federated hybrid architecture

Combining quantum-enhanced Transformer with federated learning methods to solve real-time sentiment analysis problems is an emerging hybrid research direction. Chen and Yoo [16] explored the concept of federated quantum machine learning, where quantum neural networks adopt distributed training to protect data privacy and improve training efficiency. This hybrid approach aims to combine the speed advantage of quantum computing with the data security characteristics of federated learning, which is particularly beneficial in areas such as call centers, e-commerce recommendation systems, and social media platform monitoring that require rapid processing of massive data streams. Preliminary studies have shown that embedding quantum-enhanced models in federated learning systems can handle data more effectively locally, reduce the number of communication rounds for model updates, and overall improve system performance. Moreover, these models perform well in enhancing the performance of tasks such as sentiment analysis, which is crucial for providing personalized recommendations that accurately reflect user interests [18].

3.2. Limitations

Despite the aforementioned advancements, when integrating these cutting-edge technologies into real-time sentiment analysis, a series of fundamental limitations arise due to the combined effects of the inherent characteristics of the technologies and the specific requirements of the sentiment analysis task. This section will closely follow the three major progress directions outlined in the previous section, analyze their inherent limitations one by one, and specifically explain how these limitations constrain the performance and practicality of real-time privacy-preserving sentiment analysis systems.

3.2.1. Privacy-efficiency and statistical heterogeneity challenges of federated learning in sentiment analysis

Federated sentiment learning has achieved comparable performance to centralized training while ensuring privacy. However, this advancement encounters severe "privacy-efficiency" trade-offs and statistical heterogeneity

challenges when dealing with the complexity of real-world sentiment data.

The inherent heterogeneity of emotional expression exacerbates the data statistical challenges in federated learning. The text data processed by sentiment analysis is highly personalized and context-dependent. The same word may express completely opposite emotions in different regions, age groups, or sub-cultures. This cross-user and cross-device distribution of emotional data is naturally extremely non-independent and non-identically distributed. Traditional federated learning relies on the federated averaging algorithm. When aggregating the updates of such heterogeneous models, the global model is easily skewed by the dominant local data distribution, resulting in a significant decline in the recognition accuracy of "minority" or "abnormal" emotions [19]. This means in real-time public opinion monitoring or financial sentiment warning, the system may fail to notice the non-mainstream but crucial group emotional signals, missing the early warning window.

The conflict between the privacy protection mechanism and the effectiveness of the sentiment analysis model becomes more pronounced in a federated environment. To protect users' emotional privacy, federated learning systems typically introduce differential privacy. The core lies in adding noise to the model gradients. However, the sentiment analysis task requires the model to precisely capture subtle emotional cues and semantic differences. A strong privacy budget will directly blur these crucial gradient update signals, resulting in a significant reduction in the global model's ability to identify the strength of emotions. Research progress has not demonstrated a mature solution that can perfectly retain the fine-grained sentiment discrimination ability under strong privacy guarantees, which is precisely the technical bottleneck behind the user feedback "The recommended content doesn't understand me" in practical deployments.

3.2.2. Computational advantages and realistic gaps of quantum-enhanced models in real-time emotional understanding

The theoretical potential of quantum-enhanced Transformers in accelerating NLP tasks. However, when these models need to be deployed to support real-time emotional analysis, there is a significant gap between their computational advantages and the current hardware reality.

Quantum model updates may become a new attack vector for reconstructing sensitive sentiment data. Although federated learning only shares model updates, the privacy concerns discussed in Section 3.1.1 may be further amplified in this setting. Existing studies have shown that classical gradients can be exploited to reconstruct original inputs. In hybrid architectures, however, clients may upload quantum model parameters or hybrid gradients generated at the classical-quantum interface. Owing to their high-dimensional and entangled characteristics, such updates may encode richer local sentiment-related features than classical gradients. Therefore, inverse attacks against these new types of gradients are theoretically possible, yet they remain largely

underexplored. If such vulnerabilities are exploited, not only could individual user comments be leaked, but long-term emotional states implicitly captured during local training, such as depressive tendencies, may also be inferred.

The hardware noise in the NISQ era directly undermines the precise representation of model-based emotional semantics. Emotional analysis is highly dependent on the accuracy of semantics. Current medium-scale quantum devices with noise suffer from qubit decoherence and gate errors [20]. This means that high-dimensional entangled states representing complex emotional semantics in quantum circuits are highly vulnerable to noise damage, leading to distortion in emotional representation. For example, a previously positively polarized "mixed positive-negative" sarcastic comment may, due to gate errors, cause its quantum state in the Hilbert space to collapse towards the negative pole, resulting in a completely incorrect emotional judgment. The performance breakthroughs mentioned in Section 3.1 were mostly achieved in simulation environments or very small-scale ideal datasets, far from meeting the requirements for stable processing of massive heterogeneous emotional data on real noisy hardware.

3.2.3. The high complexity of emotional data privacy risks in hybrid architectures

The quantum federated hybrid architecture holds great promise, but this integration process has pushed the privacy risks faced by emotional data to an unprecedentedly complex level. This is not merely the simple addition of the risks of the two systems, but has generated new and coupled vulnerabilities.

First, the quantum model update itself could potentially serve as a new vector for reverse attacks on emotional data. Federated learning only shares model updates, but the privacy issues discussed in Section 3.1.1 have now been elevated to a higher dimension. Research has shown that classical gradients can be used to reconstruct the original input. In a hybrid architecture, if the client uploads either the traditional quantum model parameters or the mixed gradients generated at the classical-quantum interface, their high-dimensional and entangled characteristics may encode far more local emotional data features than classical gradients. Reverse attacks on this new type of gradient are theoretically possible, but there is currently a lack of research in this area. Once the vulnerability is exploited, not only could individual users' single comments be leaked, but their long-term emotional states hidden during local training could also be stolen.

Secondly, the existing privacy technologies are unable to meet the extremely high real-time requirements of quantum hybrid systems. To address these threats, quantum homomorphic encryption or post-quantum cryptography methods can be adopted. However, such operations have huge computational costs [21], and on top of the communication delays caused by federated learning, they will further prolong each training cycle. This is fatal for public safety scenarios that require a minute-level response to sudden public opinions. At the same time, the

current post-quantum cryptography is mainly designed to defend against attacks from future quantum computers, but it has no protection against real-time information leakage caused by side channels or classical interfaces during the operation of current quantum devices. This forms a deadlock between security and performance, and is a key obstacle that must be overcome before the practical application of the hybrid architecture.

Therefore, these limitations are not isolated. Instead, they form an interconnected network of constraints that hinders the transition of privacy-preserving sentiment analysis from theory to robust real-world practice. Federated learning addresses the privacy issues of centralized data storage, but introduces challenges in model utility degradation. Quantum models promise computational acceleration, but remain constrained by hardware limitations and data encoding challenges. Although hybrid architectures appear to combine the advantages of both, they in fact introduce new and more complex trade-offs between privacy and efficiency. This clearly points to a key research gap. Developing a system that can process highly non-IID heterogeneous sentiment data streams with millisecond-level latency, while ensuring provable security guarantees and resistance to quantum attacks, remains an open and significant challenge.

4. Summary and Outlook

This paper systematically reviews the relevant literature on federated learning, quantum-enhanced Transformer, and their integrated applications in privacy-preserving sentiment analysis. The comprehensive analysis indicates that although the combination of federated learning and quantum-enhanced Transformer has made significant progress, there is still a significant gap among these three technologies in the field of real-time consumer sentiment analysis that requires strong anonymity, efficient real-time processing, and high accuracy. The current platforms or frameworks fail to fully integrate these technologies to simultaneously address the dual challenges of privacy and computational efficiency.

A key gap lies in achieving provable and quantum-resistant anonymity in quantum-enabled federated learning systems. On one hand, federated learning protects data privacy by decentralizing information processing; but on the other hand, the introduction of quantum-enhanced Transformers brings entirely new vulnerabilities. Emerging solutions such as quantum homomorphic encryption can help mitigate risks, but they come with huge resource costs and are difficult to meet real-time requirements. There are also obstacles in terms of model accuracy and real-time processing. Quantum-enhanced Transformers perform well in specific tasks, but they still struggle with model convergence when dealing with typical non-independent and identically distributed data in a federated environment. The frequent interactions required in the federated learning process also result in inflated communication costs and delays. Moreover, the noise and error issues of existing quantum hardware limit

its practicality in real-time emotional assessment scenarios that require high stability.

Looking towards the future, the key directions for breaking the current deadlock may include:

Developing lightweight quantum-secure privacy protection protocols: Research resource-efficient quantum homomorphic encryption schemes or differential privacy mechanisms suitable for the quantum environment, providing end-to-end quantum-secure protection for model updates and inference processes within acceptable performance overhead.

Future research may develop an adaptive and modular hybrid federated architecture that responds to heterogeneous data distributions, network conditions, and privacy requirements. Such an architecture could integrate classical and quantum components according to domain-specific sentiment analysis needs, including social media posts, e-commerce reviews, financial news, and online medical consultations. A promising direction is to introduce a lightweight emotion domain adaptation layer on local clients, enabling rapid adaptation to new sentiment domains through few-shot meta-learning without retraining the entire global model. Given the substantial differences across platforms in vocabulary distribution, implicit semantic expression, and polarity representation, this design may improve generalization under non-IID federated settings. While domain-adaptive modules have been explored in classical federated learning, their application in quantum federated learning remains largely unexplored.

Constructing noise-robust efficient quantum neural networks: Conduct in-depth research on error mitigation techniques for variational quantum circuits at the algorithmic level, and overcome the limitations of current NISQ devices at the software level to enhance the training stability and convergence speed of the model in a distributed heterogeneous environment.

Establishing standardized evaluation benchmarks and theoretical frameworks: Build a standardized dataset and evaluation metrics for quantum federated learning sentiment analysis, and develop a theoretical framework for unified analysis of the convergence, privacy loss, and communication complexity of mixed systems, providing rigorous guidance for future system design.

The construction of the benchmark should fully consider the specificity of the sentiment analysis task, including but not limited to: specialized datasets covering complex sub-tasks such as irony detection, fine-grained emotion classification, and sentiment analysis of multilingual mixed text; and a systematic evaluation protocol for the sentiment classification performance of the model under non-IID degree gradient control.

Through continuous exploration in these directions, the hybrid paradigm combining quantum-enhanced Transformer and federated learning is expected to ultimately solve the impossible triangle of privacy, efficiency, and accuracy in real-time sentiment analysis, paving the way for the construction of the next generation of trustworthy artificial intelligence systems.

References

- [1] Gupta, P. Evaluating the BERTScore of synthetic text and its sentiment analysis. *Research Square* 2023, preprint. <https://doi.org/10.21203/rs.3.rs-3248507/v1>.
- [2] Qian, T.; Xie, A.; Bruckmann, C. Sensitivity analysis on transferred neural architectures of BERT and GPT-2 for financial sentiment analysis. *arXiv* 2022. <https://doi.org/10.48550/arXiv.2207.03037>.
- [3] Chen, C.; Wang, Z.; Yang, Y. An enhanced social matrix factorization model for recommendation. *Journal of Information Science* 2020, 46, 312–328. <https://doi.org/10.1177/0165551519876724>.
- [4] Sipio, R.; Huang, J.-H.; Chen, S.Y.-C.; Mangini, S.; Worring, M. The dawn of quantum natural language processing. In *Proceedings of the ICASSP 2022 – 2022 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*; IEEE: Piscataway, NJ, USA, 2022; pp. 8612–8616. <https://doi.org/10.1109/ICASSP43922.2022.9747675>.
- [5] Lee, R.; Kim, M.; Rezk, F.; Li, R.; Venieris, S.I.; Hospedales, T. FedP EFT: Federated learning to personalize PEFT for multilingual LLMs. In *Proceedings of the AAAI Conference on Artificial Intelligence*; 2026; Volume 40, pp. 22805–22813. <https://doi.org/10.1609/aaai.v40i27.39443>.
- [6] Kairouz, P.; McMahan, H.B.; Avent, B.; Bellet, A.; Bennis, M.; Bhagoji, A.N.; et al. Advances and open problems in federated learning. *Foundations and Trends in Machine Learning* 2021, 14, 1–210. <https://doi.org/10.1561/22000000083>.
- [7] Jiang, Y.; Yuan, X.; Zhang, W.; Ke, W.; Lam, C.-T.; Im, S.-K. PFL-ALP: Personalized federated learning against backdoor attacks via attention-based local purification. *IEEE Transactions on Information Forensics and Security* 2025, 20, 12995–13010. <https://doi.org/10.1109/TIFS.2025.3639936>.
- [8] Dwork, C.; Roth, A. The algorithmic foundations of differential privacy. *Foundations and Trends in Theoretical Computer Science* 2014, 9, 211–407. <https://doi.org/10.1561/04000000042>.
- [9] Vaswani, A.; Shazeer, N.; Parmar, N.; Uszkoreit, J.; Jones, L.; Gomez, A.N.; Kaiser, Ł.; Polosukhin, I. Attention is all you need. *Advances in Neural Information Processing Systems* 2017, 30, 5998–6008.
- [10] Li, Z.; Li, C.; Li, Z.; Weng, J.; Huang, F.; Zhou, Z. PPMGS: An efficient and effective solution for distributed privacy-preserving semi-supervised learning. *Information Sciences* 2024, 120934. <https://doi.org/10.1016/j.ins.2024.120934>.
- [11] Broadbent, A.; Jeffery, S. Quantum homomorphic encryption for circuits of low T-gate complexity. In *Advances in Cryptology – CRYPTO 2015*; Gennaro, R.; Robshaw, M., Eds.; Springer: Berlin/Heidelberg, Germany, 2015; Volume 9216, pp. 609–629. https://doi.org/10.1007/978-3-662-48000-7_30.
- [12] Xie, D.; Li, D. Privacy-preserving and verifiable personalized federated learning. *Symmetry* 2025, 17, 361. <https://doi.org/10.3390/sym17030361>.
- [13] Zeguendry, A.; Jarir, Z.; Quafafou, M. Quantum machine learning: A review and case studies. *Entropy* 2023, 25, 287.
- [14] Alzahrani, M.E.; Aldhyani, T.H.; Alsubari, S.N.; Althobaiti, M.M.; Fahad, A. Developing an intelligent system with deep learning algorithms for sentiment analysis of e-commerce product reviews. *Computational Intelligence and Neuroscience* 2022, 2022, 3840071. <https://doi.org/10.1155/2022/3840071>.
- [15] Bansal, S.; Singh, M.; Bhadauria, M.; Adalakha, R. Federated learning approach towards sentiment analysis. In *Proceedings of the 2022 2nd International Conference on Technological Advancements in Computational Sciences (ICTACS)*; IEEE: Piscataway, NJ, USA, 2022; pp. 717–724. <https://doi.org/10.1109/ICTACS56270.2022.9987996>.
- [16] Chen, S.Y.-C.; Yoo, S. Federated quantum machine learning. *Entropy* 2021, 23, 460. <https://doi.org/10.3390/e23040460>.
- [17] Hilmkil, A.; Callh, S.; Barbieri, M.; Süßfeld, L.R.; Listo Zec, E.; Mogren, O. Scaling federated learning for fine-tuning of large language models. In *Lecture Notes in Computer Science*; Springer: Cham, Switzerland, 2021; pp. 15–23. https://doi.org/10.1007/978-3-030-80599-9_2.
- [18] Huang, R.; Tan, X.; Xu, Q. Quantum federated learning with decentralized data. *IEEE Journal of Selected Topics in Quantum Electronics* 2022, 28, 1–10. <https://doi.org/10.1109/JSTQE.2022.3170150>.
- [19] Gholamiangonabadi, D.; Grolinger, K. Federated learning for sentiment analysis in presence of non-IID data: Sensitivity of deep learning models. *IEEE Access* 2024, 12, 128049–128060. <https://doi.org/10.1109/ACCESS.2024.3453068>.
- [20] Preskill, J. Quantum computing in the NISQ era and beyond. *Quantum* 2018, 2, 79. <https://doi.org/10.22331/q-2018-08-06-79>.
- [21] Alebouyeh, Z.; Bidgoly, A.J. Benchmarking robustness and privacy-preserving methods in federated learning. *Future Generation Computer Systems* 2024, 155, 18–38. <https://doi.org/10.1016/j.future.2024.01.009>.