

AI-Driven Privacy Audit Automation and Data Provenance Tracking in Large-Scale Systems

Chen-Wei Chang*

Independent Researcher, WA 98011, USA

**Corresponding author*

Keywords: Machine Learning, Data Security, Privacy Protection, Distributed Training, Homomorphic Encryption

Abstract: With the rapid development of mobile Internet, big data, supercomputing, sensor networks, brain science and other technologies, machine learning has entered a period of accelerated development to promote economic and social progress. However, its entire life cycle (data preprocessing, model training, model reasoning) is facing increasingly complex data security and privacy challenges, which is difficult for traditional protection technologies to cope with. This study proposes the following innovative solutions to address this issue: in the data preprocessing stage, a secure feature extraction scheme based on a single cloud server called SeiFS is developed, which integrates cryptographic primitives such as obfuscation circuits, unintentional transmission, and secret sharing to achieve end-to-end privacy protection; In the model training phase, two privacy preserving distributed training schemes (PEFL and PEFLimd) are designed, which combine additive homomorphic encryption and robust aggregation strategy to automatically filter poisoning gradients and protect gradient privacy[1-3]; In the model inference stage, an efficient convolution evaluation scheme based on homomorphic encryption (supporting large kernel and large step convolution) and a CryptoGT scheme for Transformer graph neural networks are proposed to reduce computational overhead and solve the "neighbor explosion" problem through fast convolution algorithms and hierarchical evaluation protocols. All schemes have passed formal security proofs and experimental verification. SeiFS achieves efficient security feature extraction in cloud computing scenarios, PEFL series schemes improve training efficiency while filtering poisoning attacks[4-5], CNN inference schemes significantly reduce homomorphic computational complexity, and CryptoGT effectively supports security evaluation of complex nonlinear functions, solving the scalability challenges of modern neural architectures and ensuring privacy as a whole. In the future, it is necessary to break through key issues such as lightweight and verifiable security feature extraction, universal security attack and defense system, non leakage model inference technology that preserves accuracy, and full lifecycle privacy protection for new generation models (such as big language models and video generation systems)[6-7].

1. Introduction

The rapid development of mobile Internet, big data, supercomputing, sensor networks and brain

science has pushed machine learning (ML) into an unprecedented accelerated development stage, promoting technological breakthroughs in medical, financial, autonomous driving and other fields. However, its entire life cycle (data preprocessing, model training, model reasoning) is facing increasingly complex security and privacy challenges - traditional technologies (such as differential privacy, basic security multi-party computing) are difficult to balance security, efficiency and scalability[8-9], especially in resource constrained scenarios such as cloud computing and edge devices. Although existing research proposes solutions to local problems, there are still key shortcomings: security feature extraction schemes in the data preprocessing stage may lack end-to-end privacy protection due to high computational costs or inability to seamlessly integrate with federated learning frameworks; Distributed solutions during the model training phase (such as federated learning) or sacrificing efficiency and robustness (such as vulnerability to poisoning attacks), or excessively sacrificing performance through cryptographic protocols; In the model inference stage, there are common issues with the security protocols of Convolutional Neural Networks (CNN) and Transformer based Graph Neural Networks (GNN), such as high communication costs and insufficient support for complex operations (such as large-sized convolution kernels and nonlinear activation functions)[10-13].

This article aims to construct a privacy protection framework that covers the entire lifecycle of machine learning, with a focus on addressing three core issues: designing a secure and efficient feature extraction scheme suitable for cloud assisted environments (compatible with horizontal federated learning), proposing an anti attack gradient aggregation strategy to enhance the robustness and efficiency of distributed training, and developing a low-cost secure inference protocol for modern neural architectures such as CNN and Transformer GNN (supporting linear and nonlinear operations). Specifically, a single cloud assisted feature extraction framework called SeiFS is proposed for horizontal federated learning scenarios, which integrates obfuscation circuits, unintentional transmission, and secret sharing techniques[14-15]. By balancing binary tree encoding and secure feature selection mechanisms, it achieves seamless integration with existing federated learning protocols while ensuring user data privacy; For the distributed training phase, a PEFL protocol combining additive homomorphic encryption and robust gradient aggregation strategy was designed. By automatically filtering malicious user submitted poisoning gradients and optimizing aggregation rules based on gradient sign and numerical correlation, the reliability of model training was improved. At the same time, a lightweight version of PEFLimd was introduced, which significantly reduced training latency by reducing redundant calculations and merging communication rounds. Its performance advantages were verified through convergence analysis; In the model inference stage, a dedicated secure inference framework CryptoGT for Transformer GNN was developed, which includes a homomorphic encryption convolution evaluation scheme (using a fast convolution algorithm to convert convolution operations into element wise multiplication and addition, avoiding ciphertext rotation, supporting large-sized convolution kernels and large stride lengths) and a high-order polynomial security evaluation protocol based on vector unintentional linear evaluation (VOLE) (for nonlinear functions such as Softmax, LayerNorm, GeLU, etc., complex calculations are completed through a single protocol call, significantly reducing communication complexity). All schemes have been verified through formal security proofs (semi honest models) and real dataset experiments: SeiFS achieves 1.13 times higher feature extraction efficiency than baseline federated learning in cloud scenarios; PEFL/PEFLimd reduces training time by up to 10.49 times while filtering 95% of the poisoning gradient; CryptoGT reduces communication overhead by 44.61 times compared to existing solutions in GNN inference tasks.

2. Correlation Theory

2.1 Machine Learning Full Lifecycle Data Security and Privacy Protection Technology

The current research on machine learning data security and privacy protection focuses on the key stages of the entire lifecycle (data preprocessing, model training, model inference), proposing diverse technical solutions based on the characteristics of each stage, but all face the challenge of balancing efficiency, security, and functionality. The data preprocessing stage focuses on secure feature extraction, and existing solutions mainly rely on homomorphic encryption or secure multi-party computation (MPC). Methods based on homomorphic encryption, such as Rao et al.'s research, are only suitable for binary features or small-scale datasets due to their high computational complexity; The secure multi-party computation scheme (such as Li et al.'s vertical federated learning feature selection) requires modifications to the training process, which limits compatibility. Some studies, such as SeiFS, have achieved distributed feature extraction through obfuscation circuits and secret sharing, eliminating dependence on non collusive servers. However, they lack lightweight verifiable mechanisms, making it difficult to resist malicious adversary attacks and have security weaknesses. The model training phase focuses on improving the efficiency of poisoning attack defense and privacy protection. The technical path can be divided into three categories: first, secure aggregation protocols (such as Krum, multi Krum), which screen abnormal gradients through distance metrics, but need to obtain real gradient information, which conflicts with privacy protection targets; The second is differential privacy schemes (such as Phan et al.'s adaptive noise injection), which protect privacy by adding noise, but require a compromise between model accuracy and security strength; The third is the combination of homomorphic encryption and MPC schemes (such as Aono et al.'s encryption gradient aggregation). Although it avoids directly exposing gradients, encryption operations (such as modular exponentiation) lead to a significant increase in computational overhead and rely on non realistic assumptions (such as non collusive servers). In addition, most schemes are only applicable to shallow models and lack support for deep networks (such as models with convolutional layers). The technical paths in the model inference stage include homomorphic encryption (HE), secure multi-party computation (MPC), and trusted execution environment (TEE). HE schemes (such as Cryptonets and CryptoDL) protect privacy through ciphertext computation, but nonlinear functions (such as ReLU and Softmax) require polynomial approximation, resulting in accuracy loss; MPC schemes (such as GAZELLE and DELPHI) balance efficiency and accuracy through hybrid protocols, but rely on the assumption of non collusion among multiple parties and have high communication costs; TEE solutions (such as SGX) provide hardware level isolation, but have poor scalability and are vulnerable to side channel attacks. For complex models such as Transformers, existing solutions simplify nonlinear operations (such as replacing GeLU with ReLU) or approximate exponential functions (such as Chen et al.'s polynomial fitting), but at the cost of sacrificing model expression ability, some studies even lead to a decrease in accuracy of more than 2%. Limitations: Lack of end-to-end solutions that balance security and efficiency in the preprocessing stage; It is difficult to maintain model accuracy while protecting privacy during the training phase; The support for new complex models (such as big language models) in the inference stage is insufficient, and existing technologies rely heavily on strong assumptions (such as non conspirators and trusted hardware). In the future, breakthroughs are needed in lightweight verifiable feature extraction, universal security defense frameworks, and privacy protection technologies that support dynamic model structures to balance security, efficiency, and functionality requirements.

2.2 Privacy Protection Technology in the Inference Stage of Machine Learning

The entire lifecycle of machine learning covers three core stages: data preprocessing, model training, and model inference. Each stage faces security and privacy challenges and has given rise

to diverse technical solutions.

The data preprocessing stage focuses on extracting security features, with the goal of screening key feature subsets from high-dimensional data. The existing methods are mainly divided into three categories: filtering based (such as feature selection based on Gini impurity, using formulas)

$$G_I(I_j^I) = 1 - \sum_{z=1}^{\phi} (\text{Pr}_z(I_j^I))^2$$

and

$$Gs(f_i) = \frac{1}{m} \sum_{j=1}^h |I_j^I| \times G_I(I_j^I)$$

quantifying feature importance, independent of classifiers and efficient, but lacking defense against malicious attacks), encapsulated (combining machine learning algorithms to evaluate feature subset performance, such as sequence forward/backward selection, with high accuracy but significant computational cost), embedded (embedding feature selection into model training, such as directly optimizing the objective function, balancing efficiency and accuracy). However, existing schemes rely on homomorphic encryption or secure multi-party computation (MPC), which have high computational complexity (such as Rao et al.'s homomorphic encryption scheme only applicable to small-scale data), limited compatibility (such as Li et al.'s vertical federated learning requiring modification of the training process), or security shortcomings (such as SeiFS eliminating non collusive server dependencies but lacking lightweight verification mechanisms).

The model training phase focuses on distributed training (such as horizontal federated learning), which includes: the central server selects users and issues global models $w^{(t-1)}$, and users calculate gradients based on local datasets D_j

$$G_j^{(t-1)} = \frac{1}{|D_j|} \sum_{(x_i, y_i) \in D_j} \frac{\partial L_f(w^{(t-1)}; x_i, y_i)}{\partial w^{(t-1)}}$$

after uploading the gradient, the server aggregates and updates the model

$$w^{(t)} = w^{(t-1)} - \eta \sum_{j \in U^{(t)}} \frac{1}{k} G_j^{(t-1)}$$

Security threats include poisoning attacks (requiring filtering of abnormal gradients), privacy breaches (requiring protection of gradient information), and free riding attacks (reducing user computational contributions). There are three types of existing technologies: secure aggregation protocols (such as Krum filtering gradients through distance metrics, but requiring real gradient information that conflicts with privacy goals), differential privacy (protecting privacy by adding noise, but compromising between model accuracy and security strength), and homomorphic encryption+MPC (avoiding direct exposure of gradients, but encryption operations (such as modular exponentiation) lead to a significant increase in computational overhead and rely on non realistic assumptions (such as non collusive servers)). Most solutions are only applicable to shallow models and lack support for deep networks (such as models with convolutional layers).

The technical paths in the model inference stage include homomorphic encryption (HE), secure multi-party computation (MPC), and trusted execution environment (TEE). HE schemes (such as Cryptonets) protect privacy through ciphertext computation, but nonlinear functions (such as ReLU, Softmax) require polynomial approximation, resulting in accuracy loss; MPC schemes (such as GAZELLE) balance efficiency and accuracy through a hybrid protocol (combining obfuscation

circuits and homomorphic encryption), but rely on the assumption of non collusion among multiple parties and have high communication costs; TEE solutions (such as SGX) provide hardware level isolation, but have poor scalability and are vulnerable to side channel attacks. For complex models such as Transformers, existing solutions simplify nonlinear operations (such as replacing GeLU with ReLU) or approximate exponential functions (such as Chen et al.'s polynomial fitting), but at the cost of sacrificing model expression ability, some studies even lead to a decrease in accuracy of more than 2%. Limitations: Lack of end-to-end solutions that balance security and efficiency in the preprocessing stage; It is difficult to maintain model accuracy while protecting privacy during the training phase; The support for new complex models (such as big language models) in the inference stage is insufficient, and existing technologies rely heavily on strong assumptions (such as non conspirators and trusted hardware). In the future, breakthroughs are needed in lightweight verifiable feature extraction, universal security defense frameworks, and privacy protection technologies that support dynamic model structures to balance security, efficiency, and functionality.

3. Research Method

3.1 Research on Efficient and Secure Feature Extraction Scheme (SeiFS) in Distributed Scenarios

In distributed scenarios, this study proposes an efficient and secure feature extraction scheme (SeiFS) to address the contradiction between the efficiency of continuous value feature extraction and privacy protection. The traditional feature extraction algorithm based on Gini impurity requires sorting the feature values to determine the segmentation point, which has a high computational complexity. However, the Mean Split Gini method divides continuous features into two groups by using the feature mean as a threshold, reducing complexity. However, the comparison protocol under the secure computing framework still has high overhead issues. To this end, the SeiFS scheme achieves efficient feature extraction under privacy protection through three core sub processes: local feature encoding, security feature measurement, and security feature selection. In the local feature encoding stage, a balanced binary tree (BBT) is used to encode feature values into an array, and deep first search is used to store node information. Combined with virtual node completion and redirection operations, the complexity of the comparison protocol is optimized to, and the feature access mode is hidden to prevent statistical information leakage; The security feature measurement stage integrates security aggregation, security comparison, security read-write, and security division/multiplication protocols. The security comparison protocol reduces encryption table entry overhead by optimizing the circuit structure, while the security aggregation protocol achieves efficient aggregation of multi-party data based on the double mask mechanism and secret sharing; The security feature selection stage completes feature screening and data extraction through security protocols and secure reading protocols. In terms of security, SeiFS assumes that the adversary can access all intermediate data under an honest and curious threat model, but proves the security of each sub protocol (such as FAgg, FCmp, etc.) through a hybrid model security lemma (F-hybrid model), ensuring the privacy of the original data combined with the selected feature set. Performance evaluation in LAN and WAN environments (as shown in Table 1)

It indicates that SeiFS has significantly lower running time than traditional obfuscation circuit implementations when dealing with different numbers of clients, data volumes, feature dimensions, and classification labels (by 2.3-169.3 times), and the acceleration ratio under multi-threaded optimization can reach up to 64.5 times (LAN) and 46 times (WAN), while maintaining high accuracy (such as accuracy improvement of 0.6% -7.1% on COG, LSVT, and SPEED datasets). This study balances binary tree encoding, protocol optimization, and approximation algorithms to achieve efficient feature extraction while ensuring privacy, providing theoretical support and

practical reference for privacy preserving machine learning in distributed scenarios.

Table 1 Multi-threaded Performance Comparison: SeiFS vs. Baseline

Thread Count	LAN Runtime (s)		Speedup Ratio	WAN Runtime (s)		Speedup Ratio
	Baseline	SeiFS		Baseline	SeiFS	
1	142.48	2.21	64.5	3013.91	65.51	46.0
2	131.58	2.09	62.9	2910.96	62.74	46.4
4	118.01	1.94	60.8	2761.84	60.02	46.0
8	107.36	1.80	59.6	2685.33	57.85	46.4
16	93.32	1.63	57.3	2577.81	53.48	48.2
32	82.19	1.61	51.1	2433.78	51.82	47.1
64	77.05	1.51	51.0	2329.26	50.77	45.9
80	72.01	1.49	48.3	2072.60	47.48	43.7

3.2 Theoretical Analysis and Implementation of Secure Federated Learning Training Program

This study proposes a federated learning training scheme (PEFL) based on additive homomorphic encryption and secure computing protocol, aimed at resisting poisoning attacks initiated by malicious users while ensuring user data privacy. The system model consists of four core entities: the Key Generation Center (KGC) is responsible for key management, the data owner (user) trains the model locally and uploads encryption gradients, the Service Provider (SP) coordinates the aggregation process, and the Cloud Platform (CP) assists in completing ciphertext calculations. The threat model assumes that the proportion of malicious users does not exceed half of the total number of users, and that SP and CP are honest and curious adversaries who follow the protocol but attempt to obtain user privacy information, with no collusion between entities. The plan achieves its goals through two stages: the local training stage, where users use Momentum Stochastic Gradient Descent (MomentumSGD) to update the gradient, and upload it encrypted with the public key of CP, combined with ciphertext packaging and SIMD technology to reduce communication overhead; In the robust aggregation stage, SP and CP calculate the median gradient as the benchmark through the SecMed protocol, and then evaluate the Pearson correlation between the user gradient and the benchmark through the SecPear protocol. Low correlation gradients are given zero weight to resist poisoning attacks. The security analysis is based on the mixed argumentation method, which proves that under the semi honest adversary model, the scheme meets the IND-CPA security standards, and malicious users cannot obtain the privacy of other users; The convergence analysis shows that the scheme inherits the convergence properties of the SGD algorithm, and the convergence rate can reach communication and computational overhead at a constant learning rate. The basic PEFL is improved through protocol optimization (such as merging data transmission), significantly reducing communication epochs and complexity. Compared with existing solutions such as Krum and Bulyan, PEFL demonstrates advantages in computational efficiency (linear complexity), privacy protection (homomorphic encryption and dual server design), and no need to predict the number of drug users. The experiment verified that the scheme effectively balances computational efficiency and privacy protection requirements while resisting poisoning attacks.

3.3 Performance Evaluation and Experimental Verification of PEFL Scheme

This study comprehensively evaluated the performance of the PEFL scheme through real dataset experiments, with a focus on verifying its accuracy, robustness, and defense capabilities. The experiment was conducted on the MNIST (handwritten digit recognition) and CIFAR-10 (image recognition) datasets, with model architectures consisting of two fully connected networks and two convolutional layers plus three fully connected layers. Attack methods include label flipping attack (source class label 1 changed to target class 9) and backdoor attack (adding a 5×5 pixel trigger in the bottom right corner of the image and modifying the label). The experimental parameters are set to 51 users, batch size 128, momentum 0.9, initial learning rate 0.1, and the average of five experiments is taken as the result. Accuracy evaluation shows that as the proportion of poisonings increases, the performance of all defense strategies decreases, but PEFL performs the best at different poisoning ratios (0-50%). For example, in the MNIST tag flipping attack, the accuracy of PEFL still exceeded 96% even with a 50% proportion of poisonings, significantly better than Krum (89%), Bulyan (94%), and Trimmed Mean (92%). The analysis of iteration times shows that PEFL converges faster, with an accuracy rate of 96.5% for 300 iterations under MNIST backdoor attack and 62% for 1000 iterations under CIFAR-10 backdoor attack, which is superior to other schemes. The robustness evaluation was verified through the success rate of attacks, and the highest success rate of PEFL in label flipping attacks was only 0.04. Under backdoor attacks, the proportion of source classes with triggers being misclassified as target classes was as low as 0.04, significantly lower than the baseline model (100%) and other schemes (such as Krum's highest 95%). Compared with median based schemes (GeoMed, MarMed, MeaMed), PEFL does not require prior knowledge of the number of malicious users and maintains high accuracy even at a 50% poisoning rate, while GeoMed suffers from convergence issues due to its reliance on a single user gradient. PEFL effectively balances computational efficiency and privacy protection requirements while resisting poisoning attacks, and experimental results validate its superior performance in practical scenarios.

4. Results and Discussion

4.1 Privacy Protection Model Inference Scheme Based on Fast Convolution Algorithm

This article proposes an efficient privacy protection model inference scheme based on fast convolution algorithm. By optimizing the traditional convolution evaluation process and replacing the `Img2col` method of dependency matrix multiplication with the "input/convolution kernel transformation Hadamard product output transformation" mode, the rotation operation overhead in ciphertext computation is significantly reduced. In response to the limitation of the original fast convolution algorithm (such as Winograd algorithm) that is only applicable to 3×3 convolution kernels with a stride of 1, the scheme introduces a decomposable fast convolution algorithm, which uses a divide and conquer strategy to split large-sized convolution kernels into multiple small-sized sub kernels, and converts convolutions with stride greater than 1 into operations with equivalent stride of 1, thereby expanding the applicability of the algorithm to any convolution kernel size and stride scenario. Under the framework of secure computing, this solution achieves secure two-dimensional convolution evaluation through a six step process of "data reassembly segmentation transformation Hadamard product inverse transformation result merging": the cloud platform locally completes the reassembly, segmentation, and transformation of convolution kernels; Calculate the recombination, segmentation, and transformation of input data by both parties based on the secret sharing feature; Implementing secure Hadamard product operations using Package Addition Homomorphic Encryption (PAHE) based on BFV scheme; Finally, the original output is restored through inverse transformation and result merging, which improves computational efficiency while ensuring privacy.

4.2 Model Experiment

Under the framework of secure computing, this solution achieves secure two-dimensional convolution evaluation through a six step process of "data reassembly segmentation transformation Hadamard product inverse transformation result merging": the cloud platform locally completes the reassembly, segmentation, and transformation of convolution kernels; Calculate the recombination, segmentation, and transformation of input data by both parties based on the secret sharing feature; Implementing secure Hadamard product operations using Package Addition Homomorphic Encryption (PAHE) based on BFV scheme; Finally, the original output is restored through inverse transformation and result merging, which improves computational efficiency while ensuring privacy. The security of this scheme is rigorously proven through a series of formal theorems, among which Theorem 5.9 establishes the security of the OPE protocol in a semi honest adversary model. By decomposing an n -order polynomial into n linear polynomials for parallel evaluation and utilizing the established security of the VOLE protocol, it ensures that adversaries cannot distinguish between real polynomial coefficients and random values; Theorems 5.10 and 5.11 further extend security to the OMFE and ONFE protocols. The former achieves security evaluation by publicly generating segmented polynomial approximations and calling OPE_SS, while the latter combines dynamic order adjustment and finite field optimization to reduce the number of segments; Theorems 5.12 and 5.13 respectively verify the secure implementation of the Softmax and Norm functions. By combining components such as ArgMax, OMFE, multiplication protocol, and reciprocal square root evaluation, it ensures that the input and output only exist in a secret shared form, and that the intermediate calculation steps do not leak any sensitive information. Experimental results have shown that the framework achieves a balance between privacy protection and computational efficiency by improving polynomial evaluation efficiency by at least 2.6 times and reducing the number of segments by 4.7 times through parallelization design and dynamic optimization strategies, while meeting security requirements.

4.3 Effect Analysis

This chapter proposes CryptoGT, an efficient and secure two party computation framework designed specifically for secure inference in Transformer based graph neural networks. In response to the "neighbor explosion" problem in graph neural networks, this framework proposes a layer by layer inference strategy, which significantly reduces duplicate computations and improves efficiency by merging the common computing tasks of all target nodes within the same layer. In terms of nonlinear function evaluation, CryptoGT adopts segmented polynomial approximation and dynamic adjustment strategy, supports flexible adjustment of approximation function complexity according to accuracy requirements, and introduces a highly parallel secure polynomial evaluation protocol (OPE). This protocol converts the evaluation of n -order polynomials into parallel computation of n linear polynomials through the Horner's rule, and combines the selected input vector Unintentional Linear Evaluation Protocol (VOLE) to achieve polynomial evaluation in a single call, which is at least 2.6 times faster than traditional iterative multiplication schemes. For the scenario of secret sharing, it is further extended to the OPE_SS protocol to maintain equivalent security and computational efficiency as the original OPE. For complex nonlinear functions such as Sigmoid, Tanh, GeLU, the framework optimizes segmented polynomial approximation through finite field replacement and dynamic order adjustment, reducing the number of segments by 4.7 times. It also combines custom protocols OMFE and ONFE to cover common univariate function evaluation requirements, while supporting efficient and secure implementation of multivariate functions such as Softmax and Normalization (as shown in Figure 1).

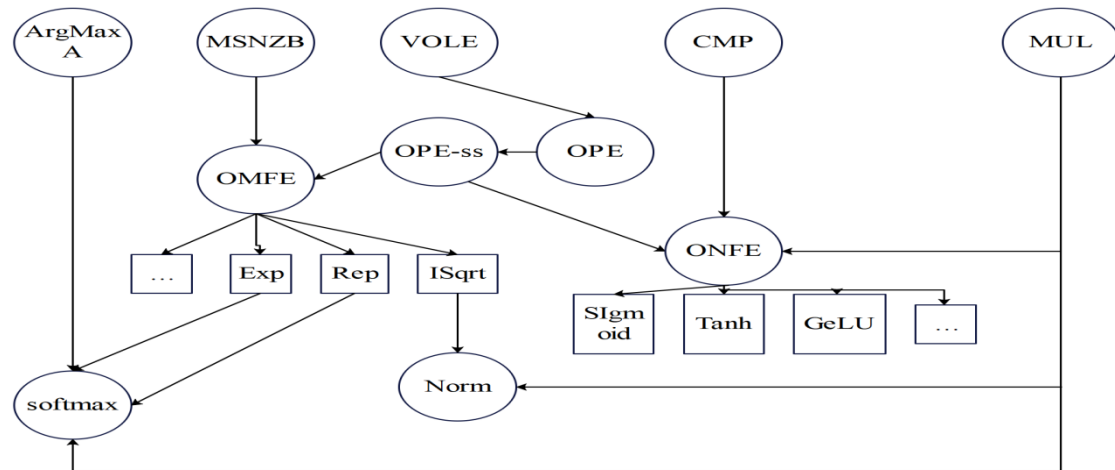


Figure 1 Comparison of OMFE and ONFE protocol computing architectures

In terms of security, CryptoGT rigorously proves its security through a series of formal theorems. Theorem 5.9 establishes the security of the OPE protocol in a semi honest adversary model, which focuses on the parallel evaluation of n -order polynomials decomposed into n linear polynomials, and utilizes the established security of the VOLE protocol to demonstrate through simulators that adversaries cannot distinguish between real polynomial coefficients and random values. Theorems 5.10 and 5.11 further extend security to the OMFE and ONFE protocols, where OMFE achieves security evaluation by publicly generating segmented polynomial approximations and calling OPE_SS, while ONFE combines dynamic order adjustment and finite field optimization to reduce the number of segments while maintaining approximation accuracy. Theorems 5.12 and 5.13 respectively verify the secure implementation of Softmax and Norm functions. The former ensures the secret sharing of exponential operation and normalization by combining ArgMax, OMFE, and multiplication protocol, while the latter uses OMFE to evaluate the reciprocal of the square root and cooperates with multiplication protocol to complete standardized calculation. All protocols are based on a semi honest security model, which constructs composite security by mixing basic security components such as VOLE, multiplication protocol, comparison protocol, etc., ensuring that input and output only exist in a secret shared form, and that intermediate computing steps do not leak any sensitive information. In the experimental section, CryptoGT was comprehensively tested on a cloud server equipped with Intel Xeon E5-2680 v4 processor and 48GB RAM. Local area network (LAN, bandwidth 3GBps, latency 0.3ms) and wide area network (WAN, bandwidth 400MBps, latency 40ms) environments were simulated using C++ language combined with SIRNN library and BFV encryption scheme. Compared with existing solutions such as SIRNN, NFGGen, and CipherGPT, CryptoGT exhibits significant advantages in nonlinear function evaluation: in LAN environments, OMFE evaluates Rec, Exp, and ISqrt with running times 75.60 times, 14.27 times, and 44.61 times faster than SIRNN, respectively, and reduces communication overhead by 4.20 times, 1.21 times, and 3.23 times (Table 2)

ONFE evaluated that the running time of Sigmoid, Tanh, and GeLU is 2.07 times, 2.21 times, and 1.46 times faster than NFGGen, and the communication overhead is reduced by 6.09 times, 4.71 times, and 2.27 times. For multivariate functions (Softmax, Norm), CryptoGT runs 2.08 times and 8.19 times faster than SIRNN on LAN and WAN, respectively, and reduces communication overhead by 1.72 times and 1.32 times. In end-to-end testing, CryptoGT showed significantly better running time than SIRNN on real models such as GraphGPS, Gophormer, and GTRS, especially when processing Gophormer (Cora dataset) in WAN environment with a speed increase of 5.67 times, and when processing GTRS (Human3.6M dataset) with a speed increase of 10.81 times (Table 3)

Table 2 Performance comparison of secure mathematical function evaluation protocols

Function	Environment	OMFE (Time /ms)	Poly_Eval (Time /ms)	SIRNN (Time /ms)	OMFE (Comm ./MB)	Poly_Eval (Comm ./MB)	SIRNN (Comm ./MB)	Speedup (LAN)	Speedup (WAN)
Rec	LAN	2.12	6.02	160.28	0.41	4.31	1.72	2.83 ×	12.04 ×
Rec	WAN	82.15	221.53	989.26	1.06	1.52	1.71	75.60 ×	10.51 ×
Exp	LAN	10.62	13.77	151.54	1.06	1.52	1.71	1.20 ×	2.43 ×
Exp	WAN	407.89	494.34	989.17	1.21	1.43	1.61	14.27 ×	1.43 ×
ISqrt	LAN	10.75	13.86	479.61	1.41	4.56	6.67	1.29 ×	7.68 ×
ISqrt	WAN	417.01	499.02	3204.22	1.20	3.23	4.73	44.61 ×	3.23 ×

Table3 CryptoGT vs. SIRNN: Real-World E2E Performance Benchmark

Model	Dataset	CryptoGT LAN (min)	SIRNN LAN (min)	Speedup (LAN)	CryptoGT WAN (h)	SIRNN WAN (h)	Speedup (WAN)	CryptoGT Comm. (GB)	SIRNN Comm. (GB)	Comm. Reduction
GraphGPS	ZINC	4.49	5.07	1.13 ×	0.44	0.49	1.12 ×	0.20	0.21	1.07 ×
Gophermer	Cora	10.09	117.08	11.60 ×	1.76	9.98	5.67 ×	4.17	4.80	1.15 ×
Gophermer	DBLP	12.81	87.94	6.86 ×	2.33	13.31	5.73 ×	3.46	4.23	1.22 ×
GTRS	Human3.6M	33.71	353.73	10.49 ×	5.45	58.96	10.81 ×	12.95	14.51	1.12 ×

The experiment also showed that as the data dimension increased (210 to 220), CryptoGT's advantage in running time further expanded. Through parallel design and dynamic optimization strategies, CryptoGT achieves a balance between privacy protection and computational efficiency while meeting security requirements.

5. Conclusion

With the rapid development of mobile Internet, big data and other technologies, machine learning is profoundly affecting economic and social development. However, its data security and privacy protection throughout its life cycle (data preprocessing, model training, reasoning) are facing complex challenges, which are difficult for traditional technologies to cope with. This article focuses on this issue and proposes a series of innovative solutions: in the preprocessing stage, a SeiFS scheme is designed for horizontal federated learning, which integrates cryptographic techniques such as obfuscation circuits to achieve secure and efficient feature extraction; Combining additive homomorphic encryption and robust aggregation techniques during the training

phase, PEFL and PEFL imd schemes are proposed to enhance training privacy and reliability by filtering malicious gradients, optimizing aggregation strategies, and improving communication efficiency; In the inference stage, homomorphic encryption convolution evaluation and CryptoGT scheme are developed. The former converts ciphertext operations to reduce computational burden, while the latter designs secure polynomial protocols for Transformer graph neural networks to alleviate "neighbor explosion" and reduce communication complexity. All schemes have been proven to be safe and experimentally validated, with significantly better performance than existing methods. Future research will focus on lightweight security algorithms, universal attack and defense systems, inference techniques that balance accuracy and efficiency, and full lifecycle protection solutions for complex models (such as graph neural networks and large language models) to meet the needs of trustworthy applications of artificial intelligence technology.

References

- [1] Zhu, Z. (2025). *Application of Database Performance Optimization Technology in Large-Scale AI Infrastructure*. *European Journal of Engineering and Technologies*, 1(1), 60-67.
- [2] Huang, W., Zhang, Z., Zhao, W., Peng, J., Xu, W., Liao, Y., ... & Wang, Z. (2025). *Auditing privacy budget of differentially private neural network models*. *Neurocomputing*, 614, 128756.
- [3] An, C. (2025). *Exploration of Data-Driven Capital Market Investment Decision Support Model*. *European Journal of Business, Economics & Management*, 1(3), 31-37.
- [4] Pan Y. *Research on the Design of a Real-Time E-Commerce Recommendation System Based on Spark in the Context of Big Data*[C]//2025 IEEE International Conference on Electronics, Energy Systems and Power Engineering (EESPE). IEEE, 2025: 1028-1033.
- [5] Lai L. *Data-Driven Credit Risk Assessment and Optimization Strategy Exploration*[J]. *European Journal of Business, Economics & Management*, 2025, 1(3): 24-30.
- [6] Ullah, F., Pun, C. M., Mohmand, M. I., Mahendran, R. K., Khan, A. A., Alhammad, S. M., ... & Farouk, A. (2025). *Privacy-aware secure data auditing for cloud-based intelligence of things environment*. *IEEE Internet of Things Journal*.
- [7] Li, W. (2025). *Discussion on Using Blockchain Technology to Improve Audit Efficiency and Financial Transparency*. *Economics and Management Innovation*, 2(4), 72-79.
- [8] Tang X, Wu X, Bao W. *Intelligent Prediction-Inventory-Scheduling Closed-Loop Nearshore Supply Chain Decision System*[J]. *Advances in Management and Intelligent Technologies*, 2025, 1(4)
- [9] Sheng C. *Research on AI-Driven Financial Audit Efficiency Improvement and Financial Report Accuracy*[J]. *European Journal of Business, Economics & Management*, 2025, 1(2): 55-61.
- [10] Zhang, Q., Qian, S., Cui, J., Zhong, H., Wang, F., & He, D. (2025). *Blockchain-Based Privacy-Preserving Deduplication and Integrity Auditing in Cloud Storage*. *IEEE Transactions on Computers*.
- [11] Wu X, Bao W. *Research on the Design of a Blockchain Logistics Information Platform Based on Reputation Proof Consensus Algorithm*[J]. *Procedia Computer Science*, 2025, 262: 973-981.
- [12] Yang D, Liu X. *Collaborative Algorithm for User Trust and Data Security Based on Blockchain and Machine Learning*[J]. *Procedia Computer Science*, 2025, 262: 757-765.
- [13] Wei, X. (2025). *Practical Application of Data Analysis Technology in Startup Company Investment Evaluation*. *Economics and Management Innovation*, 2(4), 33-38.
- [14] Fallatah, E. (2025). *Ensuring Compliance: Data Privacy Audits Under Global Privacy Regulations*. *International Journal of Applied Economics, Finance and Accounting*, 22(2), 133-144.

- [15] Wang, C. (2025). *Exploration of Optimization Paths Based on Data Modeling in Financial Investment Decision-Making*. *European Journal of Business, Economics & Management*, 1(3), 17-23.