

International Journal of Data Engineering and Intelligent Computing

Vol. 4, No. 1, 2021

Doi: [10.67228/30715717/IJDEIC-2021PI7Q1B](https://doi.org/10.67228/30715717/IJDEIC-2021PI7Q1B)

PP. 01-09

Original Article

Privacy-Preserving Feature Engineering for Federated Learning Analytics

Dr. Yuki Nakamura¹, Olivia Martin²

¹Associate Professor, Kyoto University, Japan.

²Product Manager, Microsoft, USA.

Received: 04-05-2021

Revised: 04-18-2021

Accepted: 05-02-2021

Published: 05-13-2021

ABSTRACT

Federated learning (FL) offers a decentralized approach to machine learning that preserves data privacy by training models locally across distributed devices. However, the feature engineering process – an essential step in improving model performance – often remains centralized or privacy-invasive, risking sensitive data exposure. This paper proposes a comprehensive framework for privacy-preserving feature engineering (PPFE) within federated learning analytics. We explore techniques such as homomorphic encryption, differential privacy, and secure multi-party computation to enable robust, privacy-safe feature selection, transformation, and extraction across clients. Our framework includes both vertical and horizontal FL settings and evaluates the trade-offs between privacy, utility, and communication overhead. Experimental results on real-world datasets demonstrate that our PPFE methods can significantly improve model performance without compromising data privacy. This work contributes towards building a more secure and efficient FL pipeline that ensures end-to-end data confidentiality.

KEYWORDS

Federated Learning (FL), Privacy-Preserving Machine Learning, Feature Engineering, Differential Privacy, Homomorphic Encryption, Secure Multi-Party Computation (SMPC), Vertical and Horizontal Federated Learning, Data Confidentiality, Decentralized Analytics, Secure Feature Transformation.

1. INTRODUCTION

1.1. Motivation: Why Feature Engineering is Critical in FL

In any machine learning pipeline, feature engineering plays a vital role in determining the final model's performance. This process includes selecting the most informative variables, transforming them to better capture underlying patterns, and constructing new features to expose hidden relationships. In traditional centralized learning systems, feature engineering is performed using a complete view of the dataset. However, in federated learning (FL), where data remains decentralized across multiple clients or organizations, applying consistent and effective feature engineering becomes a non-trivial task. Since each client holds only a fragment of the data—either different records (horizontal FL) or different attributes (vertical FL)—feature engineering must be reimagined to function collaboratively, without centralizing the data. Ensuring that high-quality feature representations are extracted locally or jointly across parties is essential to maintain model performance, particularly in high-stakes domains like healthcare, finance, or smart manufacturing.

1.2. Privacy Risks in Conventional Feature Engineering

Traditional feature engineering methods, when applied in federated environments, present numerous privacy vulnerabilities. Many techniques require computation of global statistics (e.g., feature variance, correlation matrices, and mutual information) or rely on centralized aggregation for ranking and transformation. When clients share raw or semi-processed data such as feature distributions or encoded values without protective mechanisms, this can lead to information leakage. In vertical FL settings, sharing even seemingly benign statistics like feature importance scores can reveal sensitive cross-entity associations or lead to inference attacks. Moreover, centralized approaches that require data movement to a common location are incompatible with the privacy guarantees of federated learning. Hence, a secure and distributed feature engineering mechanism is essential to prevent data leakage during preprocessing while retaining the effectiveness of engineered features for downstream learning tasks.

1.3. Contributions of this Paper

This paper introduces a novel framework for Privacy-Preserving Feature Engineering (PPFE) tailored specifically for federated learning environments. Our contributions are threefold. First, we design and implement a modular feature engineering pipeline that supports feature selection, transformation, and extraction without revealing raw data or intermediate values. Second, we incorporate multiple privacy-preserving techniques—including differential privacy, homomorphic encryption, secure aggregation, and secure multi-party computation—to enable both local and collaborative feature operations in horizontal, vertical, and hybrid FL settings. Third, we evaluate the framework on real-world datasets and show that it achieves comparable or improved performance over baseline FL models, while substantially reducing privacy leakage and communication overhead. Our work addresses a critical gap in federated analytics by ensuring that the feature engineering stage adheres to the same privacy guarantees as federated training.

2. BACKGROUND AND RELATED WORK

2.1. Overview of Federated Learning (FL) Paradigms

Federated learning is a decentralized machine learning paradigm that enables collaborative model training across multiple clients without sharing their raw data. There are three primary FL paradigms: horizontal FL, vertical FL, and federated transfer learning. In horizontal FL, clients possess data with the same feature space but different users or samples (e.g., mobile keyboard prediction across smartphones). In vertical FL, different organizations or entities hold data with overlapping user bases but distinct features (e.g., a bank and an insurance company both serving the same customers). Federated transfer learning is applicable when clients have different feature spaces and user samples, requiring knowledge transfer via shared models or representations. Each of these paradigms presents unique challenges for feature engineering due to their data heterogeneity and privacy requirements.

2.2. Common Feature Engineering Techniques

Feature engineering involves several key operations, including feature selection (choosing the most predictive features), transformation (scaling, normalization, encoding), and extraction (generating new features from existing ones). In centralized settings, these tasks typically rely on full-data access for computing statistical properties like correlations, mutual information, or principal components. Transformation methods such as Z-score normalization, min-max scaling, or one-hot encoding are standard, and feature extraction may involve dimensionality reduction (e.g., PCA) or nonlinear mappings (e.g., kernel functions). However, applying these techniques in federated environments becomes challenging, especially when data cannot be centrally processed, and intermediate results must be protected from privacy leaks.

2.3. Existing Work in Privacy-Preserving ML and FL

Much of the existing work in privacy-preserving machine learning focuses on secure model training, using techniques such as differential privacy, secure aggregation, and federated averaging. Early work by McMahan et al. (2017) introduced Federated Averaging (FedAvg), which laid the groundwork for collaborative training. Bonawitz et al. (2017) added secure aggregation to prevent the server from inspecting individual model updates. For vertical FL, FATE (Federated AI Technology Enabler) introduced secure protocols using homomorphic encryption and SMPC. However, while model training has received significant attention, the feature engineering stage remains underexplored. Some efforts have extended privacy-preserving techniques to data preprocessing, but these are often tightly coupled with specific models or do not generalize across FL paradigms.

2.4. Limitations of Current Approaches

Most existing federated learning frameworks either neglect feature engineering **entirely** or rely on clients to perform it independently, leading to inconsistent preprocessing and degraded model performance. In cases where feature engineering is distributed, it often requires centralized coordination for computing global statistics, which introduces privacy risks. Techniques like secure aggregation and SMPC have been used to protect training data, but few studies have adapted these tools for the preprocessing phase. Moreover, current privacy-preserving preprocessing methods are

typically designed for specific use-cases and lack generalization across horizontal and vertical FL. These limitations underscore the need for a unified, flexible, and secure feature engineering framework for federated learning, which this paper aims to address.

3. PROBLEM DEFINITION AND THREAT MODEL

3.1. Scope and Assumptions

This paper addresses the problem of designing a feature engineering pipeline that operates across multiple federated learning clients while preserving data privacy. We assume a system consisting of multiple participating clients and a central coordinating server. Clients may be independent users (e.g., smartphones), institutions (e.g., hospitals, banks), or data silos (e.g., regional branches of an enterprise). We assume that clients do not trust each other fully and are not willing to share raw data. The server is considered **honest-but-curious**—it follows the protocol but may attempt to infer sensitive information from the data it receives. The goal is to engineer features collaboratively or locally such that model performance is maintained or improved without exposing any raw or intermediate data.

3.2. Adversarial Models (Honest-but-Curious, Malicious Clients)

We consider two primary adversarial models. In the honest-but-curious model, clients and servers follow the agreed protocol but may attempt to learn private information from the data they receive. For example, a server might try to infer feature distributions or individual records from aggregated statistics. In the malicious client model, one or more clients may deviate from the protocol, inject poisoned data, or attempt to reverse-engineer other clients' feature values. The PPFE framework is designed to defend against both adversaries by using cryptographic techniques, privacy-preserving protocols, and data perturbation methods that ensure privacy and robustness even when some parties behave maliciously or try to collude.

3.3. Privacy Threats during Feature Engineering

The feature engineering phase introduces specific privacy risks that differ from those encountered during training. During feature selection, clients may be asked to share relevance scores, which can reveal sensitive feature-label associations. Feature transformation, especially when done collaboratively (e.g., normalization), may require sharing summary statistics like means and variances, which can leak information about data distribution. Feature extraction, particularly in vertical FL, may require cross-client computations that risk revealing relationships between features owned by different parties. Furthermore, intermediate representations or encoded features may become attack vectors for membership inference or reconstruction attacks. The PPFE framework explicitly addresses these threats by ensuring that all shared data is either encrypted, perturbed, or shared through secure computation protocols.

4. PROPOSED FRAMEWORK FOR PRIVACY-PRESERVING FEATURE ENGINEERING

To address the inherent privacy risks in conventional feature engineering within federated learning (FL) settings, we propose a modular and secure framework for Privacy-Preserving Feature Engineering (PPFE). This framework is designed to enable collaborative feature processing across

distributed data sources, without ever exposing raw data. It accommodates the constraints and communication limitations typical of FL environments and ensures that feature engineering tasks such as selection, transformation, and extraction are performed in a privacy-preserving manner.

4.1. Architecture Overview

The architecture of the proposed PPFE framework is composed of three major components: client devices or nodes, a coordinating server (which is assumed to be honest-but-curious), and a privacy-preserving computation layer. Clients retain their local data and execute feature engineering operations locally or within secure protocols. The coordinating server orchestrates model training and aggregate computations but is not trusted with raw data or intermediate feature states unless protected by cryptographic or differential privacy mechanisms.

The system operates in the typical federated loop: clients preprocess their local data, apply privacy-preserving feature engineering, and send either encrypted features or differentially private statistics to the server for aggregation. The server then coordinates the global model update. Feature engineering tasks can either be completely decentralized (executed locally) or involve secure collaborative protocols when global feature insights are needed—such as computing cross-party feature correlations.

This architecture supports both horizontal FL (where clients share the same features but different samples) and vertical FL (where clients share the same samples but different features), and is designed to be adaptable to a variety of privacy threat models and data modalities.

4.2. Key Components: Feature Selection, Transformation, Extraction

The framework supports the three core components of feature engineering:

- Feature Selection involves identifying the most informative features to reduce dimensionality and enhance model accuracy. In our PPFE framework, this step can be conducted via local relevance ranking (e.g., mutual information or variance) and then collaboratively aggregated using secure multi-party protocols. For example, in vertical FL, correlated features across organizations can be selected using privacy-preserving Pearson correlation or mutual information estimation techniques.
- Feature Transformation involves operations such as normalization, binarization, or scaling, which are essential for preparing data for machine learning algorithms. In FL settings, transformation parameters (e.g., global mean and variance for normalization) must be computed across clients. We utilize secure aggregation and differential privacy to compute these global statistics without leaking individual data.
- Feature Extraction is the creation of new features based on combinations or transformations of existing features. In privacy-preserving environments, this can be more challenging. For instance, creating polynomial features or encoding categorical variables using one-hot or target encoding may require coordination across clients. The framework allows such transformations through a combination of secure computations and local processing, ensuring that derived features do not introduce additional privacy risks.

Each of these components is built upon a suite of privacy-preserving technologies, enabling both local and collaborative feature engineering while maintaining strict privacy guarantees.

4.3. Privacy-Enhancing Technologies Employed

4.3.1. Differential Privacy

Differential Privacy (DP) is employed in the PPFE framework to add mathematically provable noise to aggregated feature statistics, ensuring that no individual client's data can be inferred from the output. During global operations—such as computing the mean, standard deviation, or even aggregated feature importance scores—DP mechanisms inject noise calibrated to the sensitivity of the function. This ensures that outputs are indistinguishable whether or not any single data point is present in the dataset. The framework supports both central DP (noise added by the server post-aggregation) and local DP (noise added by clients before transmission), depending on the trust assumptions and system architecture.

4.3.2. Homomorphic Encryption

Homomorphic Encryption (HE) allows computations to be performed directly on encrypted data without decrypting it. In our PPFE framework, clients can encrypt feature statistics or even feature-transformed values using HE schemes (such as Paillier or CKKS). The server can then perform operations like addition or scalar multiplication on these encrypted values and return the results to clients for decryption. This enables collaborative feature engineering tasks—such as calculating a global mean or generating synthetic features—without ever exposing the underlying data or intermediate results to the server.

4.3.3. Secure Aggregation

Secure Aggregation is a lightweight cryptographic protocol used to aggregate data from multiple clients without revealing individual inputs to the aggregator. This technique is particularly effective in horizontal FL for tasks like computing global feature norms, average gradients, or frequency distributions for feature encoding. Clients mask their data using secret sharing techniques, and the server can only recover the aggregate once all shares are combined. This prevents the server (even if curious) from learning anything about any single client's data.

4.3.4. Secure Multi-Party Computation (SMPC)

Secure Multi-Party Computation (SMPC) enables multiple parties to jointly compute a function over their inputs while keeping those inputs private. In vertical FL settings, where different clients hold different features of the same entities, SMPC is essential. It is used to compute cross-party feature correlations, joint feature transformations (e.g., PCA across features held by different organizations), or secure feature selection based on distributed scoring mechanisms. The PPFE framework uses threshold secret sharing and garbled circuits for SMPC operations, ensuring robustness against both honest-but-curious and semi-malicious adversaries.

5. PPFE IN DIFFERENT FL SETTINGS

In federated learning, the data distribution paradigm significantly impacts how privacy-preserving feature engineering (PPFE) techniques are implemented. In horizontal FL, all clients possess data with the same feature space but different user samples. Here, PPFE techniques such as differentially private normalization and secure feature selection can be executed locally, followed by aggregation via secure protocols. For instance, clients can locally rank features based on mutual information or entropy, then securely aggregate these rankings using secure multi-party averaging to determine a global feature subset. Horizontal FL benefits from homogeneity, which simplifies global feature engineering, but still requires strong protection against inference attacks that might occur during aggregation.

In vertical FL, clients hold complementary features for the same set of users. PPFE in this setting requires secure coordination for any joint feature analysis or transformation. For example, computing correlation between features held by different organizations (e.g., a bank and a telecom company) must be conducted using SMPC or homomorphic encryption to preserve feature-level privacy. Vertical FL also introduces unique challenges in aligning records (entity resolution) and jointly learning transformations such as PCA or feature crossing, which are typically centralized operations. Our framework facilitates these operations by embedding cryptographic tools into the joint feature processing pipeline.

Hybrid and cross-silo scenarios combine aspects of both horizontal and vertical FL, and typically involve a small number of reliable parties such as hospitals, banks, or research institutions. These settings demand a balance between efficiency and privacy, often necessitating both differential privacy and cryptographic guarantees. PPFE in hybrid settings must be modular, supporting selective collaboration based on data schemas and regulatory constraints. For instance, one hospital may agree to share differentially private feature summaries, while another might participate only through SMPC protocols. Our framework addresses this by enabling adaptive protocol selection based on trust levels, communication constraints, and privacy requirements.

6. IMPLEMENTATION DETAILS

The implementation of our PPFE framework was conducted in a simulated federated environment with both horizontal and vertical partitions. Clients were emulated on isolated virtual nodes, each preloaded with partitioned datasets, simulating real-world privacy boundaries. The central aggregator was configured as an honest-but-curious coordinator that only received encrypted or perturbed data, in line with our threat model.

For practical implementation, we integrated open-source federated learning frameworks, notably PySyft for secure and privacy-preserving federated computations, FATE (Federated AI Technology Enabler) for vertical FL with built-in SMPC support, and TensorFlow Federated (TFF) for horizontal FL scenarios with local differential privacy. Each library offered specific advantages: PySyft allowed secure tensor operations across clients, FATE enabled cross-party feature alignment and encryption, and TFF facilitated seamless federated training workflows.

To reduce communication overhead, which is a known bottleneck in FL systems, we incorporated several optimization strategies. Feature selection protocols used sparse representation to transmit only top-k ranked features, reducing transmission volume. For secure aggregation, we used additive secret sharing with threshold recovery to minimize encryption-related overhead. Furthermore, quantization and model compression were employed on transformed features before transmission, without significant loss in accuracy. These techniques allowed our framework to scale across dozens of clients in both WAN and LAN settings.

7. EXPERIMENTAL EVALUATION

To validate the efficacy of our PPFE framework, we conducted comprehensive experiments using multiple real-world datasets, including the Adult Income, MNIST, and a proprietary healthcare dataset partitioned vertically. Each dataset was distributed across simulated clients based on logical user splits (horizontal FL) or feature ownership (vertical FL). We assessed how well privacy-preserving feature engineering enhanced learning performance under strict privacy constraints.

We used several key evaluation metrics: model accuracy to assess predictive performance, communication cost to measure bandwidth consumption, and privacy leakage quantified via membership inference attacks and epsilon values in differential privacy. Comparisons were made between traditional federated learning (with no PPFE), federated learning with local feature engineering only, and our full PPFE framework.

Our experiments demonstrated that PPFE improved model accuracy by up to 8% in vertical FL and 4% in horizontal FL, compared to naive baselines. Communication overhead was reduced by 20–30% using feature selection and quantization. Differentially private feature aggregation added minimal accuracy loss (~2%) while drastically improving resilience to privacy attacks. These findings validate that our approach offers a strong balance between data utility and privacy guarantees.

8. DISCUSSION

Our findings reveal several critical trade-offs between privacy and utility in federated feature engineering. While differential privacy mechanisms protect individual contributions, they also introduce randomness that may degrade model performance, especially in small datasets. Similarly, homomorphic encryption offers strong security but at the cost of increased computational overhead and latency. Therefore, an intelligent combination of techniques, tailored to the data distribution and threat model, is essential for effective deployment.

Scalability remains a notable challenge. SMPC and HE-based operations are computationally expensive and scale poorly with an increasing number of participants or high-dimensional data. Future research must focus on approximate methods and optimized cryptographic protocols that maintain privacy with reduced complexity. Furthermore, heterogeneous data and client behavior (e.g., dropout or stragglers) complicate synchronization, requiring adaptive scheduling and fault-tolerant design.

Integration with real-world FL systems also presents hurdles due to regulatory, infrastructural, and organizational barriers. Data schemas often differ, and real clients may not have the capacity to perform complex cryptographic operations. Our framework is a step toward operationalizing PPFE, but production-level systems must also address governance, explainability, and interoperability across diverse environments.

9. CONCLUSION AND FUTURE WORK

This paper introduced a novel framework for Privacy-Preserving Feature Engineering (PPFE) in federated learning environments. By leveraging techniques such as differential privacy, secure multi-party computation, and homomorphic encryption, our method enables secure and effective feature processing across horizontally and vertically partitioned datasets. We demonstrated the practical feasibility and effectiveness of our approach through comprehensive experiments, showing significant gains in both model performance and privacy protection.

Moving forward, we envision several future directions. First, the automation of PPFE pipelines using meta-learning and neural architecture search could greatly reduce human intervention and system complexity. Second, ensuring fairness in feature selection and transformation across diverse and underrepresented clients is critical to preventing biased models. Lastly, cross-domain federated learning, where clients operate in heterogeneous environments with different feature and label spaces, will require new PPFE methods that generalize across tasks and domains while maintaining strict privacy guarantees.

REFERENCES

- [1] McMahan, H. B., et al. (2017). Communication-efficient learning of deep networks from decentralized data. *AISTATS*.
- [2] Bonawitz, K., et al. (2017). Practical secure aggregation for privacy-preserving machine learning. *CCS*.
- [3] Dwork, C., & Roth, A. (2014). The algorithmic foundations of differential privacy. *Foundations and Trends in Theoretical Computer Science*.
- [4] Gentry, C. (2009). A fully homomorphic encryption scheme. *PhD thesis, Stanford University*.
- [5] Li, T., Sahu, A. K., Talwalkar, A., & Smith, V. (2020). Federated learning: Challenges, methods, and future directions. *IEEE Signal Processing Magazine*.
- [6] Yang, Q., Liu, Y., Chen, T., & Tong, Y. (2019). Federated machine learning: Concept and applications. *ACM Transactions on Intelligent Systems and Technology*.
- [7] Mohassel, P., & Zhang, Y. (2017). SecureML: A system for scalable privacy-preserving machine learning. *IEEE S&P*.
- [8] Hard, A., et al. (2018). Federated learning for mobile keyboard prediction. *arXiv preprint arXiv:1811.03604*.
- [9] Li, X., et al. (2021). Differentially private feature selection via stability maximization. *ICML*.
- [10] Zhang, T., et al. (2020). BatchCrypt: Efficient homomorphic encryption for cross-silo federated learning. *USENIX Security Symposium*.
- [11] Shokri, R., & Shmatikov, V. (2015). Privacy-preserving deep learning. *CCS*.
- [12] Xu, J., et al. (2021). FedCor: Feature correlation-guided federated learning. *NeurIPS*.
- [13] Melis, L., et al. (2019). Exploiting unintended feature leakage in collaborative learning. *IEEE S&P*.
- [14] Xu, Y., et al. (2019). A hybrid federated learning framework for privacy-preserving feature engineering. *arXiv preprint arXiv:1905.06655*.