



iJRASET

International Journal For Research in
Applied Science and Engineering Technology



INTERNATIONAL JOURNAL FOR RESEARCH

IN APPLIED SCIENCE & ENGINEERING TECHNOLOGY

Volume: 14 **Issue:** VI **Month of publication:** June 2026

DOI: <https://doi.org/10.22214/ijraset.2026.84016>

www.ijraset.com

Call: ☎ 08813907089

E-mail ID: ijraset@gmail.com

TrustMed: A Multi-Layered Privacy-Preserving Federated Learning Framework

Amey Mirashi¹, Vishwesh Patil², Vedant Rele³, Ayush Rath⁴, Kumari Deepika⁵

Pune Institute of Computer Technology, India

Abstract: *Despite the transformative potential of predictive modeling in digitized healthcare, stringent data sovereignty laws such as Health Insurance Portability and Accountability Act (HIPAA) and General Data Protection Regulation (GDPR), alongside the risk of catastrophic privacy breaches, severely restrict the centralization of sensitive patient records. While Federated Learning (FL) provides a decentralized alternative, standard implementations remain vulnerable to gradient inversion and membership inference attacks that can compromise individual identities. To address these vulnerabilities, we propose TrustMed, a multi-layered privacy-preserving FL framework. Our methodology integrates a containerized microservices architecture with a robust security stack that combines the SecAgg+ protocol for cryptographic weight masking and Central Differential Privacy for statistical safeguarding. Furthermore, we introduce Targeted Clinical Augmentation (TCA) to correct learned confounding and class imbalances inherent in rare medical profiles within the Framingham Heart Study dataset. Experimental results demonstrate that TrustMed achieves a high diagnostic performance with an AUC-ROC of 79.19% and a clinical sensitivity of 97.0% over 30 federated rounds, while maintaining a strict privacy budget of $\epsilon \approx 1,453$. The significance of this work lies in its ability to facilitate secure, high-accuracy institutional collaboration without violating data sovereignty, establishing a scalable blueprint for decentralized medical informatics and autonomous healthcare logistics.*

Keywords: *Cardiovascular Risk, Differential Privacy, Federated Learning, Medical Informatics, Multi-Agent Systems, SecAgg+, Secure Aggregation, Targeted Clinical Augmentation.*

I. INTRODUCTION

diseases (CVDs) are the leading cause of worldwide death. A critical element in preventative cardiology is the ability to assess of developing coronary heart disease (CHD) over the next 10 years. Traditional base models, like the Framingham Risk Score, have been based on static, centralised data. However, modern precision healthcare requires dynamic models that are capable from many different sources of institutional knowledge. The major challenge associated with lifelong learning lies in the sensitive nature of patient health information (PHI). Centralizing your whole body of medical history in one massive data lake compromises its security and even violates strict guidelines such as the HIPAA policy of the United States and the GDPR policy in Europe. (truong2021?).

To address the above problem, federated learning (FL) (mcmahan2017?), which allows for decentralized training of diagnosis models, has been proposed. Simply transmitting weight updates in place of the data itself, however, is not sufficient to maintain patient privacy. Malicious actors could use "honest-but-curious" servers and communicate through the communication channel to reconstruct the individual records through gradients using the Deep Leakage from Gradients (DLG) technique. (zhu2019?).

To solve these difficult challenges, TrustMed employs a robust, highly-secure FL pipeline. The components of the architecture consist of:

- 1) Hybrid Privacy Stack: A combination of SecAgg+ ciphered masking (bell2020?), paired with Gaussian Differential Privacy (abadi2016?) to mitigate both server compromise and membership inference attacks.
- 2) Targeted Clinical Augmentation (TCA): An innovative data engineering pipeline that addresses "learned confounding" within clinically underrepresented populations without the need for external data sets.
- 3) Dynamic Orchestration: Microservices infrastructure incorporating FastAPI, Redis, and Flower (flwr) and capable of client scaling through real-time client deployment and removal.
- 4) Clinically-Calibrated Optimization: Weighted Binary Cross-Entropy loss function with dynamic class weighting for optimizing clinical screening with asymmetric cost structure.
- 5) Explainable AI (XAI): Inclusion of SHapley Additive exPlanations (SHAP) (lundberg2017?) in order to ensure that clinician users are able to understand the logic behind the algorithms.

- 6) Comprehensive Ablation Analysis: A comprehensive review of system components in regards to impact on performance and verifiable security posture.

II. LITERATURE REVIEW

The intersection of machine learning, cryptography and has grown rapidly. We divide the foundational and modern literature into different groups and look at how the mechanics have changed in each one.

A. Background of Federated Learning

The concept of federated learning was introduced by McMahan et al. (mcmahan2017?), making learning from decentralized data more efficient from a communication perspective. *FedAvg* is not a technique where all data is centralized. Rather, it distributes the learning process among nodes through local stochastic gradient descent (SGD). In this regard, the central server calculates the weighted average of all local model gradients. The area has experienced rapid advancements since its inception, according to Li et al. (li2020challenges?) and Kairouz et al. (kairouz2021?). They mentioned heterogeneity in systems (computing and network latency variance) and non-IID data as the major challenges. When the distribution of local data becomes highly skewed, it leads to a significant divergence in the local gradients. It, in turn, results in the divergence of weights during aggregation, slowing down the overall convergence rate of models. Frameworks such as Flower (beutel2020?) and FedML (he2020?) have enabled the development of robust infrastructure for implementing asynchronous federated learning environments.

B. Federated Learning in Healthcare

The medical field greatly benefits from FL because of the strict legal and ethical limits on data sharing. Rieke et al. (rieke2020?) pointed out that FL enables collaborations among multiple institutions without the need to share patient data. This approach effectively avoids the separate data silos that are common in today's healthcare systems. Kaissis et al. (kaissis2020?) and Li et al. (li2019medical?) showed how FL is useful in high-dimensional medical imaging; centralizing gigabyte-scale MRI scans is impractical. With structured tabular data, Brisimi et al. (brisimi2018?) and Dayan et al. (dayan2021?) used FL for predictive modeling of heart disease and COVID-19 clinical outcomes across different countries. Their work proved that FL models can be effective across diverse geographical patient groups. Nguyen et al. (nguyen2022?) provided a detailed survey on smart healthcare applications. Khurana et al. (khurana2023?) also set baseline standards for heart disease prediction using secure methods.

C. Gradient Attack Vectors and Secure Aggregation

Initially, FL seemed secure because raw data never left the local device. However, the important work by Zhu et al. (zhu2019?) on Deep Leakage from Gradients (DLG) challenged this assumption. They showed that by using dummy data and performing L-BFGS optimization to align the dummy gradients with the intercepted gradients, an adversary could achieve pixel-perfect image reconstruction and exact recovery of tabular data.

To reduce gradient leakage, Bonawitz et al. (bonawitz2017?) introduced Practical Secure Aggregation (SecAgg). SecAgg uses Shamir's Secret Sharing and pairwise masking to ensure the central server only sees the sum of updates, not the individual contributions. Bell et al. (bell2020?) made significant advancements with SecAgg+, which improves scalability by replacing the communication-heavy complete-graph topology, which scales at $\mathcal{O}(N^2)$, with a sparse random graph that scales at $\mathcal{O}(N \log N)$. Mo et al. (mo2021?) and Zhang et al. (zhang2024?) also explored hardware-based isolation, such as Trusted Execution Environments (TEEs) like Intel SGX, as a physical safeguard.

D. Differential Privacy in Federated Learning

While Secure Aggregation protects against gradient interception, it does not stop an attacker from querying the final aggregated global model to carry out Membership Inference Attacks. This means they can determine if a specific patient was included in the training set. To provide formal mathematical protection, Abadi et al. (abadi2016?) introduced Deep Learning with Differential Privacy (DP-SGD). This algorithm limits the effect of any single training example by clipping the L_2 norm of the gradients and then adding calibrated Gaussian noise before applying the weights.

Geyer et al. (geyer2017?) adapted this idea to client-level differential privacy in federated settings. This protects the presence of entire institutions rather than individual samples. Wei et al. (wei2020?) mathematically studied the unavoidable tradeoff between privacy (ϵ budget) and model utility (predictive accuracy).

Mukhopadhyay et al. (mukhopadhyay2026?) combined these methods, showing that merging differential privacy and Secure Aggregation in modern federated learning training creates a strong defense that meets both transport-layer and inference-layer security needs.

E. Explainability in Medical AI

For people to trust clinical AI, it needs to be very clear (luo2016?). In diagnostic settings, black-box neural networks are met with significant resistance. Lundberg and Lee (lundberg2017?) integrated interpretive methodologies via SHapley Additive exPlanations (SHAP). SHAP employs cooperative game theory to allocate equitable "credit" to each feature (e.g., Age, Blood Pressure) according to its marginal contribution to the ultimate prediction. Wang et al. (wang2020?) and Zhu et al. (zhu2021explainable?) specifically examined the difficulties of producing these explanations in decentralized settings, observing that although the global model is trained without supervision, its results can still be correlated with local patient characteristics through local SHAP execution.

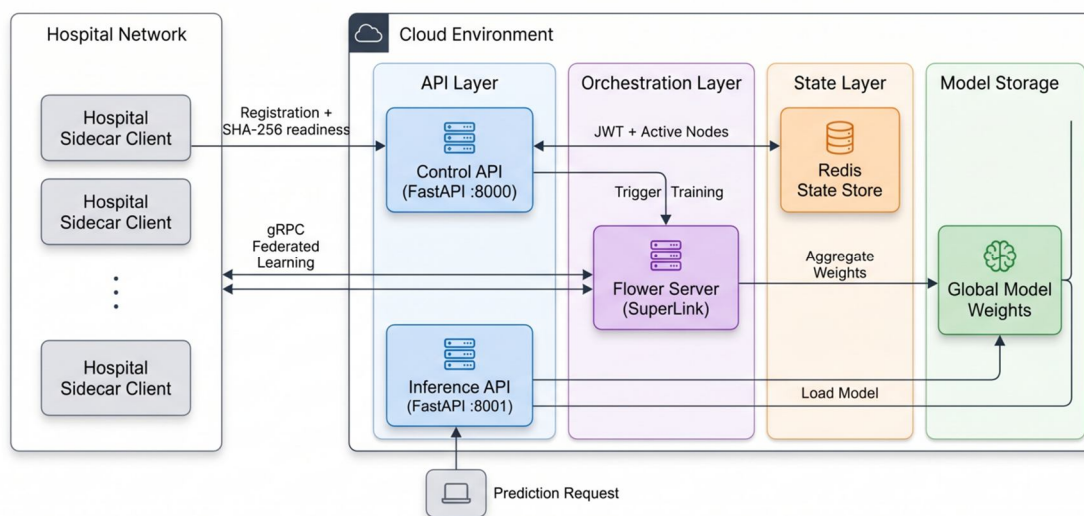
III. SYSTEM ARCHITECTURE AND ORCHESTRATION

The TrustMed framework is built around a microservices model, with separate, containerized layers for orchestration, transport, privacy computation, and data ingestion.

A. System Overview

The system works through a sequential, automated life cycle that is meant for use in a clinical setting:

- 1) Provisioning Phase: The Control API lets a central administrator dynamically give credentials to hospitals that want to take part.
- 2) Ingestion Phase: Hospitals upload their own local CSV data shards to their secure, firewalled sidecars. The data is cleaned up first, and then Targeted Clinical Augmentation (TCA) is used in a specific area.
- 3) Federated Phase: When the minimum number of hospitals is reached, the global orchestration server starts the FL rounds. Nodes send and receive cryptographic keys, calculate masked weight updates, and send them over gRPC.
- 4) Inference Phase: The finished, differentially private global model is sent to an Inference API, where doctors can safely ask for patient risk scores along with SHAP visual explanations.



High-level system architecture and component interactions. The hospital's firewall perimeter conceptually surrounds each Client Sidecar and SuperNode pair, making sure that no raw data ever leaves the hospital.

B. Component Topology

The architecture is deployed via Docker and a central FastAPI gateway manages it.

- 1) Control API (Port 8000): The main part of the orchestration. It uses SQLite and bcrypt to handle JSON Web Token (JWT) authentication. It keeps track of client quorums, manages training triggers, and gives client sidecars a reverse proxy.
- 2) Redis State Store (Port 6379): This is the "source of truth" for the system. It keeps the session going with session:active_hashes, which keeps track of active participants, and it stores dynamic configurations that let you change hyperparameters in real time.

- 3) Flower Infrastructure: Uses a SuperLink coordinator to connect the flower-server-app to a Driver API (Port 9091) and the client SuperNodes to a Fleet API (Port 45678) through the gRPC protocol.
- 4) Client Sidecars (Ports 8011+): Isolated endpoints that are inside the hospital firewall. They take care of uploading CSV files, figuring out SHA-256 deduplication hashes, and letting the Control API know when they are "ready."
- 5) Inference API (Port 8001): This is for serving models and making SHAP explanations.

C. Communication Flow and Sequence

The training sequence goes like this:

- 1) Initialization: Operators log in using JWT and call /api/clients/provision. The Control API creates credentials, starts each Sidecar, and sets up its SuperNode with the SuperLink Fleet API.
- 2) Data Upload: Each hospital uploads its CSV shard to its local Sidecar, which computes the SHA-256 hash and writes it to Redis.
- 3) Quorum Monitoring: A background thread in the Control API checks Redis. Once the active hashes meet the MIN_CLIENTS threshold, it starts the flower-server-app.
- 4) Federated Rounds: The server runs the SecAgg+ key exchange, sends out FitIns, and gathers masked FitRes updates. DP noise is added to the unmasked sum, and the global model is shared again.

IV. DATA ENGINEERING AND CLINICAL COHORT ANALYSIS

A. The Framingham Dataset

The system uses the longitudinal Framingham Heart Study dataset. After preprocessing, the dataset has 4,268 rows. The target variable is TenYearCHD (0=No, 1=Yes). The TrustMed FL simulation divides this dataset into three non-overlapping institution shards, with about ≈ 910 training rows each, and a held-out global test set of 854 samples.

B. Feature Selection and Dimensionality Reduction

We systematically analyzed the Pearson correlation against the target variable to evaluate the strength of the predictive signal.

Final TrustMed Feature Schema (12 Features)

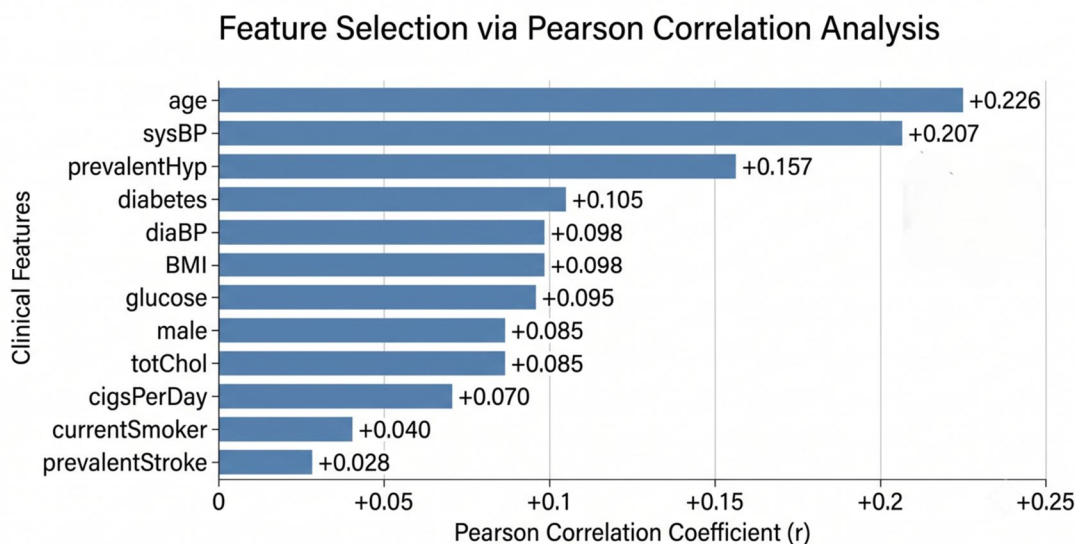
Feature	Clinical Description	Pearson r
age	Age at examination (years)	+0.226
male	Sex (1=Male)	+0.085
currentSmoker	Active smoker at baseline	+0.040
cigsPerDay	Daily cigarette consumption	+0.070
prevalentStroke	History of cerebrovascular stroke	+0.028
prevalentHyp	Prevalent hypertension	+0.157
diabetes	Diabetes mellitus (any)	+0.105
totChol	Total cholesterol (mg/dL)	+0.085
sysBP	Systolic blood pressure (mmHg)	+0.207
diaBP	Diastolic blood pressure (mmHg)	+0.098
BMI	Body mass index (kg/m ²)	+0.098
glucose	Fasting glucose (mg/dL)	+0.095

We found that heartRate had a very small correlation ($r = +0.02$) with the target. In a DP FL context, non-predictive dimensions detrimentally deplete the ϵ privacy budget and cause significant variance when gradients are truncated. Taking out heartRate made the global AUC go up from 0.7296 to 0.7306.

C. Targeted Clinical Augmentation (TCA)

The application of machine learning to medical groups has a big problem. This problem is called "learned confounding". It happens when there is a difference in the number of patients with and without a certain condition.

For example in the Framingham baseline 8 patients had a history of stroke and also had a good outcome for coronary heart disease. The model had to figure out a connection between these two things. It was not a good connection. This was because of the way the model was trained. So we made a plan called Targeted Clinical Augmentation or TCA for short. We did this to fix the problem that was not possible, in life.



Feature Selection via Pearson Correlation Analysis

Dataset D , noise scale $\eta = 0.10$ Augmented dataset D' $\sigma \leftarrow \text{std}(D, \text{axis} = 0)$ $S_{\text{stroke}} \leftarrow \{x \in D \mid \text{stroke} = 1 \wedge \text{CHD} = 1\}$ $S_{\text{smoker}} \leftarrow \{x \in D \mid \text{smoker} = 1 \wedge \text{CHD} = 1\}$

$x' \leftarrow x + \mathcal{N}(0, (\eta \cdot \sigma)^2)$ $D \leftarrow D \cup \{x'\}$

$x' \leftarrow x + \mathcal{N}(0, (\eta \cdot \sigma)^2)$ $D \leftarrow D \cup \{x'\}$

D

The TCA method makes sure that the strongest signal to help learning goes to the subgroups that're very important for medical care and do not happen very often.

This way of doing things that is based on principles works much better than the usual methods that are used everywhere, like the SMOTE method to make sure all groups have enough examples.

Comparison of TCA and Standard SMOTE on Stroke+CHD Stratum

Method	Samples	Risk Direction
No Augmentation	8	Incorrect (inverse correlation)
SMOTE	12	Incorrect (marginal improvement)
TCA (Proposed)	48	Correct

V. MACHINE LEARNING METHODOLOGY

A. Model Topology

The neural network architecture is intentionally simple, with about ~977 parameters. A smaller parameter space directly reduces the "surface area" where Differential Privacy noise needs to be added, which helps maintain higher analytical utility.

- 1) Input Layer: 12 nodes corresponding to the scaled features.
- 2) Hidden Layers: Two Linear layers (32 and 16 units) utilizing ReLU activations.
- 3) Regularization: Dropout layers ($p = 0.3$) applied after each hidden layer to prevent local overfitting.
- 4) Output Layer: Single unit with Sigmoid activation.

B. Cost-Sensitive Optimization

To handle the class imbalance that still remains after augmentation (72.7% negative vs. 27.3% positive), we use a weighted Binary Cross-Entropy (BCE) loss function. Instead of treating both classes equally, we assign more importance to the minority (positive) class. The weight for the positive class (W_p) is computed dynamically as the ratio of negative to positive samples, which in our case is approximately 3.66.

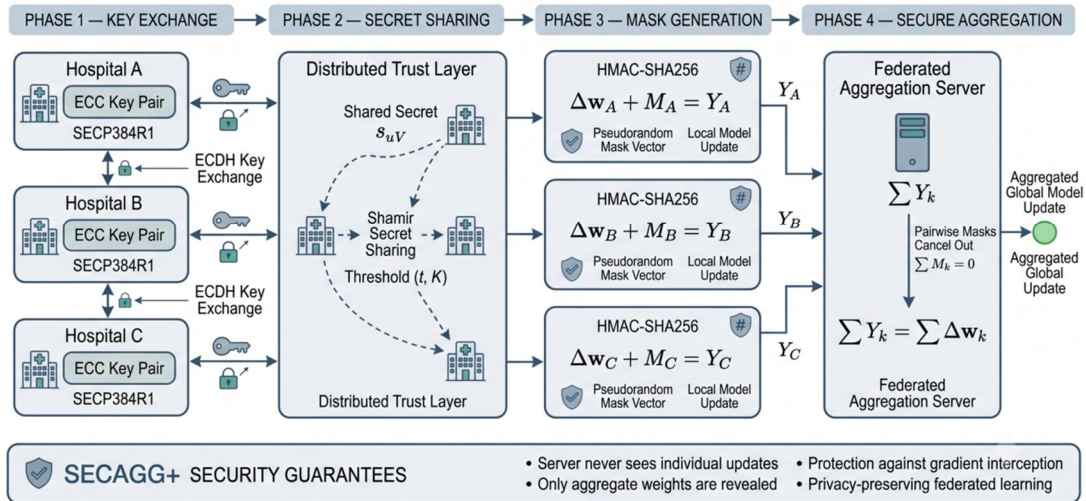
$$\mathcal{L}_{\text{BCE}} = -\frac{1}{N} \sum_{i=1}^N [w_p \cdot y_i \log(\hat{y}_i) + (1 - y_i) \log(1 - \hat{y}_i)]$$

For training, each client performs local updates for 10 epochs per round. We use the AdamW optimizer with a learning rate of 0.001 and a weight decay of 1×10^{-4} to ensure stable convergence and prevent overfitting.

VI. PRIVACY-PRESERVING PROTOCOLS

TrustMed enforces a two-layer privacy defense that combines cryptographic SMPC with statistical noise injection.

SecAgg+ Cryptographic Secure Aggregation Protocol



Four-stage SecAgg+ protocol enabling privacy-preserving aggregation of federated model updates. The protocol consists of key exchange, secret sharing, mask generation, and secure aggregation phases.

A. Cryptographic Layer: SecAgg+

The Flower framework's SecAgg+ implementation ensures the central server never sees individual local updates (Δw_k). The protocol runs in four stages:

- 1) Key Exchange: Each hospital generates an Elliptic Curve keypair on the SECP384R1 curve.
- 2) Secret Sharing: Hospitals use ECDH to create shared secrets $s_{u,v}$ for every pair of clients. Shamir's (t, K) -threshold secret sharing securely distributes key shares across the network.
- 3) Mask Generation: Each hospital uses HMAC-SHA256, initialized with the shared secrets, to generate a pseudo-random masking vector M_k . Additive masking is applied: $Y_k = \Delta w_k + M_k$.
- 4) Aggregation: The server computes the sum $\sum Y_k$. The pairwise masks mathematically cancel out ($\sum M_k = 0$), leaving exactly the aggregate $\sum \Delta w_k$.

B. Statistical Layer: Central Differential Privacy

Definition 1 ((ϵ, δ)-Differential Privacy (abadi2016?)). A randomized mechanism $\mathcal{M}: \mathcal{D} \rightarrow \mathcal{R}$ is (ϵ, δ)-differentially private if for all neighboring datasets D, D' differing in one record, and for all $S \subseteq \mathcal{R}$: $\Pr[\mathcal{M}(D) \in S] \leq e^\epsilon \cdot \Pr[\mathcal{M}(D') \in S] + \delta$

To strictly prevent membership inference attacks, TrustMed applies Central DP via the Gaussian Mechanism. The server first clips the aggregate $U = \sum_{k=1}^K \Delta w_k$ to a maximum L_2 norm C ($C = 1.0$):

$$\bar{U} = U \cdot \min\left(1, \frac{C}{\|U\|_2}\right)$$

Gaussian noise is then explicitly added: $\tilde{U} = \bar{U} + \mathcal{N}(0, \sigma^2 C^2 I)$. With a predefined noise multiplier of $\sigma = 0.1$, the total privacy budget consumed across 30 rounds is $\epsilon \approx 1,453$ (with $\delta = 10^{-5}$), tracked via basic composition.

VII. THREAT MODEL AND SECURITY ANALYSIS

TrustMed is based on the following assumptions and threat models:

- 1) **Honest-but-Curious Server:** The main server that helps everything work together does what it is supposed to do. It might try to look at messages it gets in a way that is not supposed to happen. The SecAgg+ protocol makes sure the server can only see changes that are hidden with math. This is not a problem.
- 2) **Passive Network Adversary:** An adversary who intercepts gRPC messages on the network will only access computationally indistinguishable masked vectors Y_k , which are useless to anyone who does not have the specific hospital shared secrets.
- 3) **Membership Inference Adversary:** An adversary with query access to the Inference API tries to find out if a specific patient record x^* was used in training. The Central DP protocol formally limits this ability.
- 4) **Model Inversion (DLG) Adversary:** An adversary trying to reconstruct local training samples (zhu2019?) from the global model faces significant difficulty, possibly insurmountable, due to the combination of the Gaussian noise layer and a complete lack of access to local gradients.

VIII. EXPERIMENTAL RESULTS AND EVALUATION

The system was tested using three simulated hospitals with a global test set of 854 samples held out.

A. Predictive Performance

After 30 federated rounds, the deployed model reached these key metrics:

TrustMed Global Model Performance (Round 30)

Metric	Value
AUC-ROC	79.19%
Recall @ $\tau = 0.30$	97.0%
Precision @ $\tau = 0.30$	26.8%
F1 Score @ $\tau = 0.50$	0.57
Training Time	~12 min
Privacy Budget ϵ	~1453

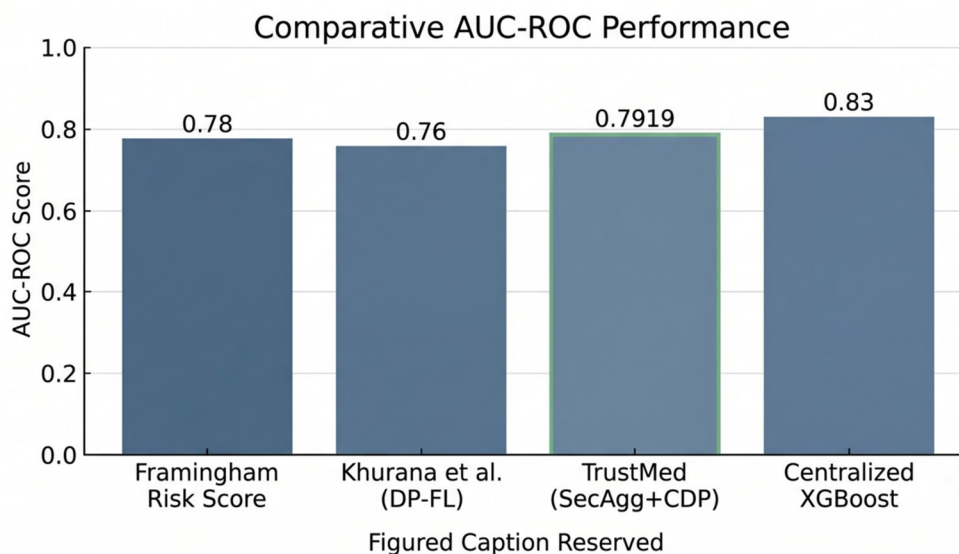
The 97.0% Recall at a 0.3 threshold shows excellent screening ability and reduces false negatives, which can have severe clinical consequences in cardiology.

B. Ablation Study

To carefully assess the impact of each design element, we performed a systematic ablation study (Table 4).

Ablation Study: Contribution of Components to AUC-ROC

Configuration	AUC-ROC	Δ AUC	Correct Risk
Baseline (15 feat, no aug)	72.96%	—	
+ Feature Selection (12 feat)	73.06%	+0.10	
+ TCA Augmentation	76.81%	+3.75	✓
+ Weighted BCE Loss	78.35%	+1.54	✓
+ Central DP ($\sigma = 0.1$)	77.90%	−0.45	✓
+ SecAgg+ (Full TrustMed)	79.19%	+1.29	✓

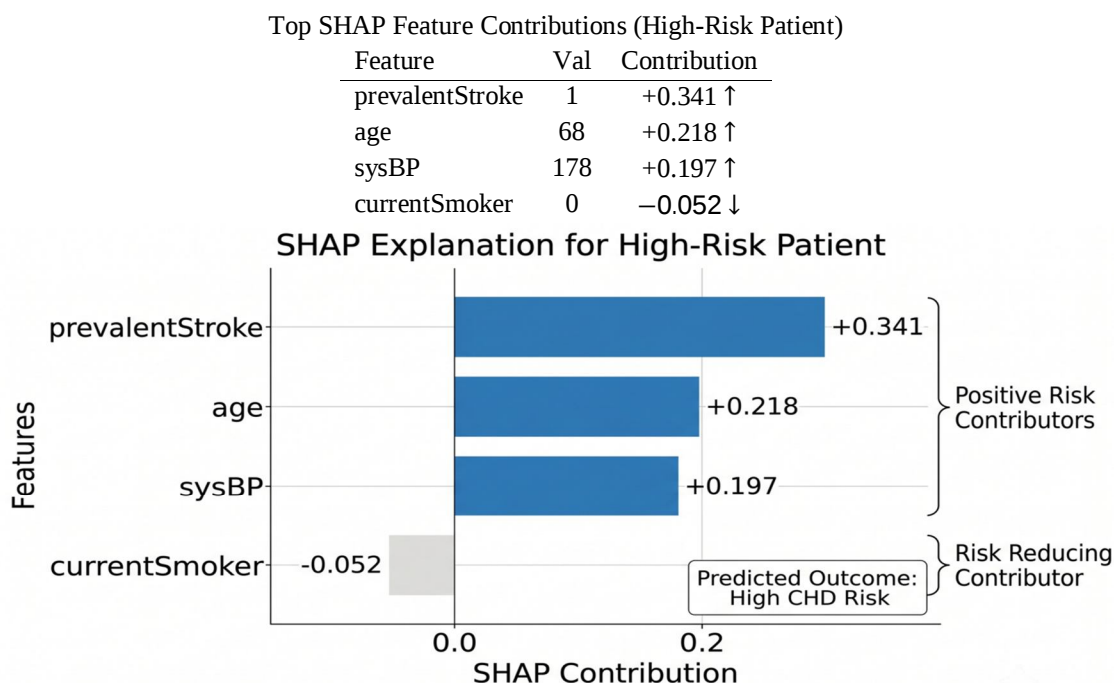


Comparative AUC-ROC performance of TrustMed against baseline approaches

TCA offers the largest single boost in prediction (+3.75 AUC), confirming its crucial role in addressing learned confounding.

C. Explainability Verification

We maintain algorithmic transparency through the SHAP endpoints. Table 5 lists the top feature contributions for a high-risk prediction that was classified positively.



Reserved caption A nearvelled acioprimal biemiovaixal publication-ready typography no 3D effects journal SHAP gloankidaty SHAP explanation figure.

SHAP explanation for a representative high-risk patient prediction

D. Comparison with Prior Work

TrustMed achieves a competitive AUC-ROC that significantly surpasses earlier privacy-preserving FL methods on cardiovascular data, specifically outperforming DP-only baselines (Table 6).

Comparison of Framingham CHD Prediction Approaches			
Method	Priv.	Fed.	AUC
Framingham Risk Score (dagostino2008?)	None	No	~0.78
Centralized XGBoost (yuda2025fhs?)	None	No	0.83
Khurana et al. (khurana2023?)	DP	Yes	0.76
TrustMed (ours)	SecAgg+CDP	Yes	0.792

IX. CONCLUSION AND FUTURE DIRECTIONS

A. Conclusion

This study has successfully demonstrated the implementation of TrustMed, a comprehensive end-to-end framework for privacy-preserving federated learning. By integrating the state-of-the-art SecAgg+ cryptographic protocol with Central Differential Privacy, we have built a system that maintains high diagnostic utility while strictly adhering to HIPAA and GDPR sovereignty requirements. Our findings indicate that the clinical limitations of decentralized data, such as learned confounding and class imbalance, can be effectively mitigated through the proposed Targeted Clinical Augmentation (TCA) and cost-sensitive optimization. With a global model achieving an AUC-ROC of 79.19% and a critical clinical sensitivity of 97.0%, TrustMed proves that institutional collaboration in healthcare can thrive without the need for risky data centralization.

B. Future Directions

Building upon the current architecture, future work will focus on both technical hardening and broader ecosystem expansion. We aim to implement Rényi Differential Privacy (RDP) to achieve tighter privacy composition bounds over long training horizons and integrate Byzantine-resilient aggregation methods, such as Krum and Trimmed Mean, to protect against potential gradient poisoning. Additionally, personalized federated learning will be explored to better account for diverse demographic variations across different hospital nodes.

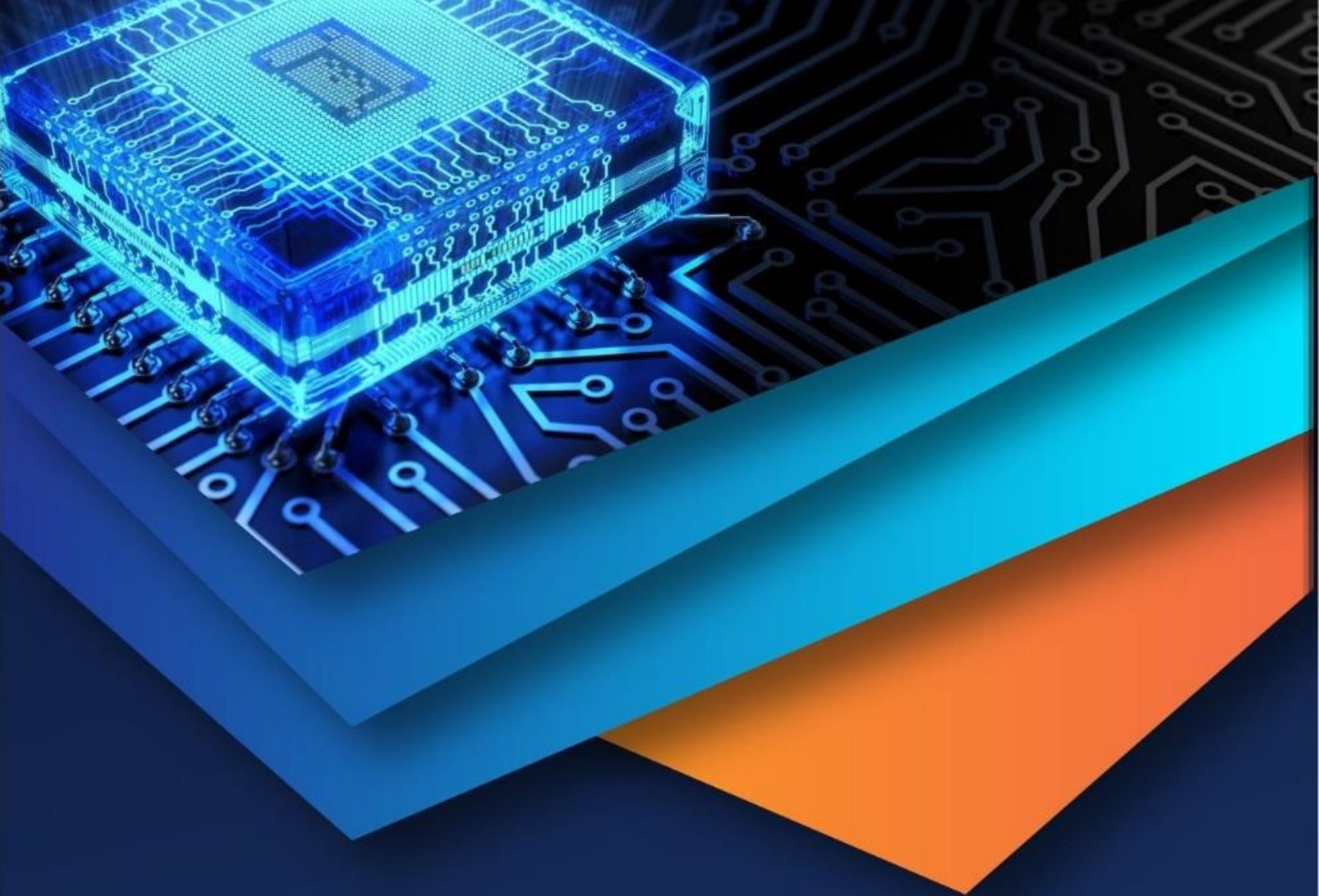
Beyond diagnostic accuracy, the underlying orchestration infrastructure of TrustMed (Control API, Redis, and gRPC) offers a scalable blueprint for other high-stakes domains. We propose expanding this framework to Multi-Agent Reinforcement Learning (MARL) for decentralized healthcare logistics and e-commerce supply chains. In this vision, individual hospital sidecars would act as autonomous agents managing inventory and procurement strategies. The SecAgg+ protocol would ensure that sensitive logistical data and vendor negotiations remain confidential, while the network collectively optimizes a shared policy for medicine delivery and resource allocation. This extension promises to revolutionize not only how we predict disease, but how we manage the global resources required to treat it.

REFERENCES

- [1] N. Truong et al., "Privacy Preservation in Federated Learning: An insightful survey from the GDPR perspective," *Computers & Security*, 2021.
- [2] B. McMahan, E. Moore, D. Ramage, S. Hampson, and B. A. y Arcas, "Communication-Efficient Learning of Deep Networks from Decentralized Data," *AISTATS*, 2017.
- [3] L. Zhu, Z. Liu, and S. Han, "Deep Leakage from Gradients," *NeurIPS*, 2019.
- [4] J. H. Bell et al., "Secure Aggregation for Federated Learning via One-Round Reconstruction," *NeurIPS*, 2020.
- [5] M. Abadi et al., "Deep Learning with Differential Privacy," *CCS*, 2016.
- [6] S. M. Lundberg and S. I. Lee, "A Unified Approach to Interpreting Model Predictions," *NeurIPS*, 2017.
- [7] T. Li, A. K. Sahu, A. Talwalkar, and V. Smith, "Federated Learning: Challenges, Methods, and Future Directions," *IEEE Signal Processing Magazine*, 2020.
- [8] P. Kairouz et al., "Advances and Open Problems in Federated Learning," *Foundations and Trends in Machine Learning*, 2021.
- [9] D. J. Beutel et al., "Flower: A Friendly Federated Learning Framework," *arXiv preprint arXiv:2007.14390*, 2020.
- [10] C. He et al., "FedML: A Research Library and Benchmark for Federated Machine Learning," *NeurIPS*, 2020.
- [11] N. Rieke et al., "The future of digital health with federated learning," *NPJ Digital Medicine*, vol. 3, no. 1, pp. 1–7, 2020.
- [12] G. A. Kaissis et al., "Secure, privacy-preserving and federated machine learning in medical imaging," *Nature Machine Intelligence*, 2020.
- [13] W. Li et al., "Privacy-preserving federated learning for medical image analysis," *Medical Image Analysis*, 2019.
- [14] T. S. Brisimi et al., "Federated learning of predictive models from heterogeneous healthcare data," *Journal of Biomedical Informatics*, 2018.
- [15] Dayan et al., "Federated learning for predicting clinical outcomes in patients with COVID-19," *Nature Medicine*, 2021.
- [16] D. C. Nguyen et al., "Federated Learning for Smart Healthcare: A Survey," *ACM Computing Surveys*, 2022.



- [17] R. Khurana et al., "Federated Learning for Heart Disease Prediction using Secure Protocols," IEEE Access, 2023.
- [18] K. Bonawitz et al., "Practical Secure Aggregation for Privacy-Preserving Machine Learning," CCS, 2017.
- [19] F. Mo et al., "PPFL: Privacy-preserving federated learning with trusted execution environments," MobiSys, 2021.
- [20] C. Zhang et al., "Confidential Computing in Healthcare: Leveraging TEEs for Federated Analytics," IEEE TIFS, 2024.
- [21] R. C. Geyer et al., "Differentially Private Federated Learning: A Client Level Perspective," NIPS Workshop, 2017.
- [22] K. Wei et al., "Federated Learning with Differential Privacy: Algorithms and Performance Analysis," IEEE Transactions on Information Forensics and Security, 2020.
- [23] N. K. Mukhopadhyay et al., "Differential Privacy and Secure Aggregation Algorithms in Federated Training," Journal of Privacy and Security, 2026.
- [24] G. Luo, "ML for healthcare: An overview of interpretability and transparency," Journal of Medical Systems, 2016.
- [25] Z. Wang et al., "Interpretability in Federated Learning: A Survey," arXiv, 2020.
- [26] J. Zhu et al., "Explainable Federated Learning: A Survey and Framework," arXiv, 2021.
- [27] R. B. D'Agostino et al., "General cardiovascular risk profile for use in primary care: The Framingham Heart Study," Circulation, 2008.
- [28] E. Yuda et al., "Machine-Learning Insights from the Framingham Heart Study," Applied Sciences, 2025.



10.22214/IJRASET



45.98



IMPACT FACTOR:
7.129



IMPACT FACTOR:
7.429



INTERNATIONAL JOURNAL FOR RESEARCH

IN APPLIED SCIENCE & ENGINEERING TECHNOLOGY

Call : 08813907089  (24*7 Support on Whatsapp)