

International Journal of
Engineering Research and Science & Technology



ISSN : 2319-5991

www.ijerst.com

Email: editor@ijerst.com or editor.ijerst@gmail.com

SECURE FEDERATED LEARNING WITH BLOCKCHAIN AND SMPC: PROTECTING HEALTHCARE SYSTEMS FROM POISONING ATTACKS

¹K. Prashanth, MCA Student, Department of MCA

²Dr.K.Pavan Kumar, Ph.D, Associate Professor, Department of MCA

¹²Dr KV Subba Reddy Institute of Technology, Dupadu, Kurnool

ABSTRACT

Federated learning (FL) has been used to several domains, such as smart cities, intelligent healthcare systems, and sectors reliant on the internet of things (IoT). This is in response to the growing concern about privacy and security in machine learning applications. Clients may work together to train a global model using FL, even when they don't have access to the same data as before. Nevertheless, adversarial attacks may exploit the weaknesses of existing FL methods. Because of its design, protecting against harmful model changes and identifying them are both made difficult. Not enough has been investigated in the most current research on protecting the model's privacy while detecting FL from harmful updates. This study presented a solution to the problem of poisoning assaults on healthcare systems: federated learning based on blockchain technology with SMPC model verification. We start by removing the compromised model after verifying it via an encrypted inference procedure using the FL participants' machine learning models. After each participant's local model is checked, it is safely forwarded to the blockchain node for aggregation. In order to test our proposed framework, we used a variety of medical datasets in a number of studies.

I. INTRODUCTION

The healthcare industry is one of the many services that have made use of the Internet of Things (IoT). The Internet of Medical Things (IOMT) is another name for the integration of IOT in the healthcare sector. Many healthcare

equipment are now linked because to the development of IOMT, which enables information sharing between medical professionals and AI-based applications. Hospitals and other healthcare facilities benefit from this interconnectedness by increasing the effectiveness and calibre of their offerings. In the realm of medical diagnostics, medical imaging technology help medical personnel diagnose and treat patients earlier.

This interconnectedness makes it simple to get medical images, producing vast amounts of data with a broad range of variances. As a result, medical picture analysis has become difficult for professionals in the field and is prone to human mistake. A notable advancement in medical picture classification tasks has been made in recent years because to Deep Learning's (DL) success in computer vision tasks. Numerous DL investigations in the realm of medical imaging have shown encouraging outcomes by offering precise and effective diagnosis [1].

One model that arose to address the availability of computational and storage resources is cloud computing, as seen in Figure 1. As a result, the DL model is often deployed for training and data inference on the cloud. However, it will be highly costly to transport the IOMT cluster's raw data to the cloud. In order to process the data before sending it to the cloud, edge computing—such as edge servers—will be useful in this situation.

<https://doi.org/10.62643/ijerst.2025.v21.i2.pp1032-1042>

Vol. 21, Issue 2, 2025

It is well recognised that a big and varied dataset is necessary for training a Deep Learning (DL) model with excellent performance. In the healthcare sector, this extensive dataset is often gathered via voluntary data exchange and multi-institutional or multi-national data aggregation. While sharing patient data raises privacy issues and related rules like the General Data Protection Regulation (GDPR) and Health Insurance Portability and Accountability Act (HIPAA), enormous data gathering is necessary for the deep learning process. Healthcare organisations may not be able to share their medical datasets because of the growing worries. In some situations when sharing is feasible, limitations are imposed, which leads to insufficient data exchange.

A federated learning strategy that enables cooperative model training by exchanging local model updates with a parameter server was developed in recent works by [2]. Since machine learning models learn from healthcare IOMT data rather than depending on a third-party cloud to store their data, this approach seems safer than centralised training [3]. Federated learning does, however, have several drawbacks that could restrict its use in practical situations. Federated learning, for instance, is still susceptible to a number of attacks that might lead to the poisoning of the learning model [5] or the disclosure of private information [4]. Furthermore, the participants in the present FL setting are unable to confirm the machine learning model's legitimacy. The current defence strategy primarily focusses on maintaining the secrecy of the machine learning gradients in order to safeguard the privacy of FL participants. One of the often used techniques to protect the privacy of the learning model is Differential Privacy (DP) [6], [7]. The privacy of the participant models may be enhanced by integrating DP into a federated learning

scenario. Nevertheless, the accuracy of the learning model will be decreased if noise is introduced into machine learning gradients [7]. A flawed global model results from DP's inability to prevent poisoning assaults while preserving model performance. Existing research on anomaly detection [8], [9] has been investigated to address the poisoning assault. Nevertheless, the current techniques are unable to eradicate every contaminated model, which lowers the global model's accuracy. Additionally, they use a plaintext model to apply the anomaly detection approach. This will result in an additional problem where the attacker may execute a membership inference attack [11] and a parameter theft attack [10]. For federated learning applications, a verified and safe anomaly detection technique is also required.

In order to remove poisoned local models in a federated learning situation, this study suggests a privacy-preserving verification technique. The suggested approach uses an SMPC-based encrypted inference procedure to remove the compromised local model while ensuring the privacy of the local model's parameters. The confirmed portion of the local model is sent to the blockchain for the aggregation procedure once the local model has been validated. The safe aggregation between the hospital and the blockchain is carried out via SMPC-based aggregation. Following the aggregation procedure, the global model is kept in a secure location. Afterwards, the blockchain sends the global model to each hospital, which then confirms its legitimacy. The following is a summary of our work's contributions:

In order to guarantee the security of the global model used for illness classification, suggest a novel block chain-based federated learning architecture for healthcare systems.

Using SMPC as the underlying technology, develop a privacy-preserving technique for local

<https://doi.org/10.62643/ijerst.2025.v21.i2.pp1032-1042>

model anomaly detection in a federated learning environment. Our encrypted model verification technique removes the tainted model while shielding the local model privacy against parameter theft and membership inference attacks.

_ We provide a verifiable machine learning model for federated learning participants utilising block chain in the IOMT scenario. _ We suggest an SMPC-based secure aggregation in the block chain as a platform to decentralise the aggregation process.

This is how the remainder of the paper is structured. The issue and design objectives are defined in Section II. The related work is discussed in Section III. In Section IV, we then introduce the suggested frameworks and explain the system architecture. In Section V, we next go into the experimental setup and assessment findings for the suggested study. Lastly, Section VII draws a conclusion.

II. PROBLEM SCENARIO AND DESIGN GOALS

A. Issue Situation

We utilise a hospital scenario with IOMT enabled to explain and illustrate the existing problems with federated learning (see Fig. 2). Suppose that a number of smart hospitals with diverse patient demographics and illnesses are situated in various locations. Every smart hospital has a collection of IOMT devices. To find a serious illness, the patient will be scanned using the IOMT equipment. Given their limited resources and inability to execute any machine learning algorithms, IOMT devices will serve as data sources in the present IOMT scenario. As a result, every hospital has a computing-capable edge server that can do machine learning tasks using local datasets. However, the machine learning model accuracy produced from the local datasets is very poor because of dataset restrictions. Consequently, each hospital's edge

server takes part in the federated learning platform. Without transferring private datasets to the cloud provider, the federated learning platform gathers and aggregates locally trained models from the hospital's edge server to create a highly accurate machine learning model. The global or aggregated model is then returned to the edge server for further federated learning procedures. The global model will be used to more precisely identify the illness once it achieves the required level of accuracy.

The federated learning scenario outlined above increases machine learning accuracy overall, however it has the following security flaws:

_ Local model security risks: Each participant that submits a local model to the cloud for aggregation without verifying the model's validity is the present configuration of federated learning. The danger of poisoning a local model is introduced by this conventional FL technique. A flawed local model might result, for instance, from a poisoning attack in which the attacker uses tainted data to train the model. Sending unencrypted local models to the cloud may provide privacy problems since healthcare data is so important. Therefore, in order to protect the local model from different security concerns, it must be validated and secured.

_ Dangers of creating a biased aggregated model: The local model's aggregation process is carried out on cloud services, which are susceptible to manipulation and may result in a biased global model. For instance, during the aggregation phase, an attacker may introduce a poisoned local model, which might result in a misleading classification for the global model. Therefore, to address the existing security issue, a safe aggregation technique is needed.

_ Risk of obtaining a flawed global model: Under the current federated learning approach, each hospital's edge server will get the global

<https://doi.org/10.62643/ijerst.2025.v21.i2.pp1032-1042>

model that was created in the cloud. The hospital is unable to validate the global model they were given, however. The global model may be intercepted and changed by the attacker. Consequently, the hospital was given a flawed worldwide model. To guarantee the integrity of the global model, a global model verification technique is needed in light of this issue.

B. Design Objectives Given the aforementioned dangers and threats, our objectives for protecting privacy in Federated Learning may be broken down into the following three parts:

The suggested approach should be robust enough to stop the attacker from tainting federated learning. This enables the federated learning participant to increase the accuracy of their model by learning from a benign global model. To protect the aggregation process against an attacker, a strong aggregation technique must also be created.

_ Privacy: Previous research [12] has shown that an attacker may use a poisoning attack to misclassify the machine learning model, hence reducing the global model accuracy. Verifying the participant's local learning model while preserving the privacy of the local model is crucial for safeguarding the federated learning participants.

_ Verifiability: The machine learning model, more especially the global model, should be verifiable by the defined technique. because the global model might be tainted or changed by the enemy. The participant in the present federated learning scenario obtained the global model from the cloud without being aware of the model's legitimacy.

II. LITERATURE SURVEY

"Analysis, technology, and applications of edge-cloud computing and artificial intelligence in the Internet of Medical Things,"

Y. Guo, H. Ren, X. Jiang, and L. Sun

The rapid growth of medical data and its heterogeneity have created significant challenges for data access, security, and privacy, as well as information processing in the Internet of Medical Things (IoMT). This is due to the booming development of medical informatisation and the pervasive connections in the fifth generation mobile communication technology (5G) era. Under the limitations of the current medical environment and medical-related technology, this paper offers a thorough evaluation of how to achieve the fast processing and analysis of medical big data as well as the sinking of high-quality medical resources. Our primary emphasis is on the benefits that cloud computing, edge computing, and artificial intelligence technologies provide to the Internet of Medical Things. In order to provide patients access to high-quality medical care, we also investigate ways to rationalise the use of medical resources as well as the protection and privacy of medical data. Lastly, we go over the existing issues and potential future research paths in the IoMT field, which is connected to edge-cloud computing and artificial intelligence.

"On fedavg convergence on non-iid data,"

Z. Zhang, S. Wang, W. Yang, K. Huang, and X. Li,

Federated learning allows several edge computing devices to train a model together without exchanging data. Federated Averaging (FedAvg), a prominent technique in this context, averages the sequences only occasionally while performing Stochastic Gradient Descent (SGD) in parallel on a small fraction of all devices. Although it is straightforward, it does not provide theoretical assurances in practical situations. In this work, we examine FedAvg's convergence on non-iid data and determine that, for smooth and highly convex problems, the convergence rate is $O(1/T)$, where T is the number of SGD's. Crucially, our bound illustrates a trade-off between convergence rate

<https://doi.org/10.62643/ijerst.2025.v21.i2.pp1032-1042>

and communication efficiency. We explore various averaging algorithms and loosen the assumption of complete device involvement to partial device participation since user devices may be disconnected from the server. This allows for a low device participation rate without significantly slowing down learning. Our findings support empirical observations by showing that data heterogeneity slows down convergence. Additionally, we provide a prerequisite for FedAvg on non-iid data: even when full-gradient is used, the learning rate η must decrease; otherwise, the solution will be $\Omega(\eta)$ from the optimum.

"Disease prediction research using enhanced deepfm and iomt,"

Z. Yu, M. Alhussein, S. U. Amin, and Z. Lv,

Deep Learning has started to gain popularity in recent years as computer processing power has increased. In almost every industry, its ability to learn non-linear feature combinations has been crucial in a way that classical machine learning cannot. Because Deep Learning and Factorisation Machine (FM) can learn high-order and low-order feature combinations, respectively, and because FM's hidden vector system allows it to learn information from sparse data, the use of Deep Learning has also fuelled the development of FM in the field of recommendation systems. Many academics have taken an interest in their integration. For the Click-Through-Rate (CTR) issue, they have studied a number of traditional models, including Factorization-supported Neural Network (FNN), Product-based Neural Networks (PNN), Inner PNN (IPNN), Wide&Deep, Deep&Cross, DeepFM, etc., and their performance keeps improving. Because of the aforementioned benefits, this kind of model is also appropriate for use in meteorology, agriculture, disease prediction, and other domains. We apply small enhancements and parameter changes to DeepFM and use it to

forecast the occurrence of hepatitis in each sample in the structured illness prediction data of the 2020 Artificial Intelligence Challenge Preliminary Competition. The enhanced DeepFM performs very well in terms of AUC when compared to other models. By using this study to electronic medical data, physicians may concentrate on samples with higher anticipated occurrence rates and lessen their burden. We can provide the Internet of Medical Things (IoMT) for certain shifting data, including blood pressure, height, weight, cholesterol, etc. To make sure that the illness can be anticipated in advance, just in case, IoMT's sensors may be employed to conduct transmission. A healthcare system that is better at predicting and time performance is created after joining IoMT.

"A framework for assessing federated learning's client privacy leaks,"

Wei W., Liu L., Loper M., Chow K.-H., Gursoy M. E., Truex S., and Wu Y.

An innovative distributed machine learning framework for cooperative model training with a network of clients (edge devices) is called federated learning (FL). By enabling clients to store sensitive information on local devices and to only exchange local training parameter changes with the federated server, FL provides default client privacy. Even transmitting local parameter updates from a client to the federated server, however, may be vulnerable to gradient leaking attacks and violate the client's privacy with respect to its training data, according to current research. We provide a guiding framework in this study for assessing and contrasting various types of client privacy leakage attacks. Using formal and practical analysis, we first demonstrate how adversaries may simply analyse the shared parameter update from local training (e.g., local gradient or weight update vector) to reconstruct the private local training data. Next, we examine the potential effects of various attack algorithm settings and

<https://doi.org/10.62643/ijerst.2025.v21.i2.pp1032-1042>

federated learning hyperparameter setups on attack cost and attack efficacy. Additionally, our system uses communication efficient FL protocols to test, assess, and analyse the efficacy of client privacy leakage attacks under various gradient compression ratios. In order to emphasise the significance of offering a methodical attack evaluation framework towards a thorough comprehension of the different types of client privacy leakage threats in federated learning and creating theoretical underpinnings for attack mitigation, our experiments also include some initial mitigation techniques.

"A survey on federated learning security and privacy,"

G. Srivastava, A. Dehghantanha, Y. Huang, S. Pouriyeh, R. M. Parizi, and V. Mothukuri,

A new kind of artificial intelligence (AI) called federated learning (FL) expands on decentralised data and training to deliver learning to the edge or right onto the device. Because of its (unknown) security and privacy consequences, FL, a new study topic sometimes described to as a new dawn in AI, is still in its infancy and has not yet acquired significant community confidence. The FL approach's security and privacy issues must first be recognised, assessed, and recorded in order to forward the field's research and achieve widespread use and mass acceptance. An implementer or adopter of FL may effectively create a safe environment and provide researchers with a clear picture of potential study fields by having a clear perspective and awareness of risk factors. FL is favoured in use-cases where security and privacy are the primary considerations. In order to close the gap between the current level of federated AI and a future when broad adoption is feasible, this article attempts to provide a thorough analysis of FL's security and privacy features. We examine the current issues in FL and provide an illustrated explanation of techniques and different implementation styles.

We also offer a comprehensive evaluation of security and privacy considerations that must be taken into consideration in a clear and comprehensive context. Our study's findings imply that, on the whole, FL has fewer privacy-specific risks than security risks. Backdoor attacks, poisoning, and communication bottlenecks are presently the most specific security issues, while inference-based assaults pose the greatest danger to FL's privacy. We provide much-needed research suggestions at the end of the study to help FL adapt to real-world situations.

III. SYSTEM ANALYSIS AND DESIGN EXISTING SYSTEM

By delivering the model to the customer and doing local training, FL ensures data protection. The central server will later gather the locally trained model and combine it into a global model. Participants in this strategy did not submit any datasets; they merely shared the local model. FL by itself, however, is insufficient to provide privacy. To protect the FL architecture, some research has been done. By include noise in the local datasets, the authors in [6] and [7] improve the data privacy in FL using differential privacy (DP). Additionally, add a proxy server in [7] to make the end-user anonymous. Nonetheless, the experiment's findings indicate a significant decline in accuracy. Since accuracy is necessary for the inference process, this privacy-preserving approach is inappropriate for FL in healthcare systems.

Zhang et al. [13] use a batch encryption technique called fully homomorphic encryption (FHE) to carry out training and aggregation procedures. However, since training takes a long period, none of the homomorphic encryption techniques can be used in healthcare settings. An adversarial attack on the FL architecture has been successfully carried out by the authors in [14], [15], and [16]. A poisoning attack on the

<https://doi.org/10.62643/ijerst.2025.v21.i2.pp1032-1042>

local client's datasets has been shown by the authors. The global model will be impacted by the created poisoned model. The DP and FHE methods are not enough to prevent the poisoning attack based on the current assault. A privacy-enhanced FL against poisoning adversaries was presented by the author in [17]. They use linear homomorphic encryption to encrypt the machine learning model for security. The training procedure will take longer than with standard machine learning since they encrypt the model from the first FL round. Following the completion of the encrypted training, The server will get the local model for encrypted aggregation. Their aggregation approach lowers the machine learning model's accuracy, according to the findings of their trials.

Blockchain technology is utilised for tamper-proof storage and is renowned for its immutability. Blockchain technology may be used to monitor the local or global model for audibility. FL and blockchain together can guarantee the integrity of the machine learning model. Verifiable aggregation for FL was suggested by the author in [18]. Their approach is based on the blockchain idea, where they calculate the digest for verification using the hash. However, only one server is used for the hashing and aggregation procedure. The issue may be resolved with the right use of blockchain technology. [19] used blockchain-enabled FL to suggest decentralised privacy as a solution to the problem. Although the participant is linked to the same blockchain, they store the model on it and utilise cross-validation to validate it. Participants in their framework have the ability to use other people's local models, which raises privacy concerns.

A smart contract is used in the work on [20] to validate the global model. Smart contracts may be used to verify that the global model is

legitimate. They did not, however, verify either the local or global model. Furthermore, no audit procedure can be carried out since the local model is not transferred to the blockchain. They are unable to manage any poisoning attack based on the planned work.

Disadvantages

- The system didn't implement a verifiable Federated Learning (FL) scenario that leverages SMPC to perform an encrypted local model verification process and secure aggregation on the blockchain node.
- The federated learning scenario not allows each participant to collaboratively train the machine learning model locally with their local datasets. Later the machine learning model will send to the cloud for the model aggregation process.

PROPOSED SYSTEM

In a federated learning situation, the system suggests a privacy-preserving verification technique to get rid of poisoned local models. The suggested approach uses an SMPC-based encrypted inference procedure to remove the compromised local model while ensuring the privacy of the local model's parameters. The confirmed portion of the local model is sent to the blockchain for the aggregation procedure once the local model has been validated. The safe aggregation between the hospital and the blockchain is carried out via SMPC-based aggregation. Following the aggregation procedure, the global model is kept in a secure location. After receiving the global model via the blockchain, each hospital confirms the global model's legitimacy.

Advantages

- Propose a new blockchain-based federated learning architecture for healthcare systems to ensure the security

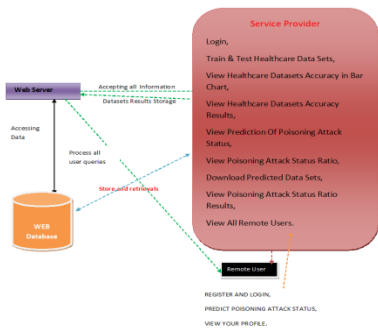
<https://doi.org/10.62643/ijerst.2025.v21.i2.pp1032-1042>

of the global model used for classifying disease.

Design a privacy-preserving method for local model anomaly detection in a Federated learning scenario with SMPC as the underlying technology. Our encrypted model verification method eliminates the poisoned model while protecting the local model privacy from membership inference attacks and parameter stealing.

- Propose an SMPC-based secure aggregation in the blockchain as a platform to decentralize the aggregation process.
- We present a verifiable machine learning model for federated learning participants using blockchain in the IoMT scenario.

SYSTEM ARCHITECTUTRE



IMPLEMENTATION

Modules

Service Provider

The Service Provider must use a working user name and password to log in to this module. He may do many tasks after successfully logging in, including training and testing healthcare data sets, View the accuracy of healthcare datasets in a bar chart, view the accuracy results, view the prediction of the status of poisoning attacks, view the ratio of poisoning attacks, and download the predicted data sets. View All

Remote Users and the Poisoning Attack Status Ratio Results.

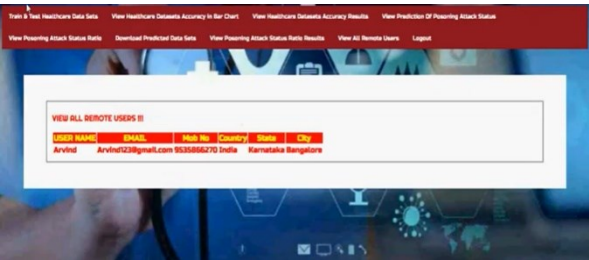
View and Authorize Users

The administrator may see a list of all registered users in this module. Here, the administrator may see the user's information, like name, email, and address, and they can also grant the user permissions.

Remote User

A total of n users are present in this module. Before beginning any actions, the user needs register. Following registration, the user's information will be entered into the database. Following a successful registration, he must use his password and authorised user name to log in. Following a successful login, the user may do tasks including registering and logging in, predicting the status of poisoning attacks, and seeing their profile.

IV. SCREEN SHOTS



<https://doi.org/10.62643/ijerst.2025.v21.i2.pp1032-1042>

	Age	anemia	hypertension	diabetes	asthma	high_cholesterol	stroke	heart_failure	kidney_failure	liver_failure	blood_glucose_level	blockchain_code_label
1	25	0	0	0	0	0	0	0	0	0	100	1000000000
2	35	0	0	0	0	0	0	0	0	0	100	1000000000
3	45	0	0	0	0	0	0	0	0	0	100	1000000000
4	55	0	0	0	0	0	0	0	0	0	100	1000000000
5	65	0	0	0	0	0	0	0	0	0	100	1000000000
6	75	0	0	0	0	0	0	0	0	0	100	1000000000
7	85	0	0	0	0	0	0	0	0	0	100	1000000000
8	95	0	0	0	0	0	0	0	0	0	100	1000000000
9	105	0	0	0	0	0	0	0	0	0	100	1000000000
10	115	0	0	0	0	0	0	0	0	0	100	1000000000
11	125	0	0	0	0	0	0	0	0	0	100	1000000000
12	135	0	0	0	0	0	0	0	0	0	100	1000000000
13	145	0	0	0	0	0	0	0	0	0	100	1000000000
14	155	0	0	0	0	0	0	0	0	0	100	1000000000
15	165	0	0	0	0	0	0	0	0	0	100	1000000000
16	175	0	0	0	0	0	0	0	0	0	100	1000000000
17	185	0	0	0	0	0	0	0	0	0	100	1000000000
18	195	0	0	0	0	0	0	0	0	0	100	1000000000
19	205	0	0	0	0	0	0	0	0	0	100	1000000000
20	215	0	0	0	0	0	0	0	0	0	100	1000000000
21	225	0	0	0	0	0	0	0	0	0	100	1000000000
22	235	0	0	0	0	0	0	0	0	0	100	1000000000
23	245	0	0	0	0	0	0	0	0	0	100	1000000000
24	255	0	0	0	0	0	0	0	0	0	100	1000000000
25	265	0	0	0	0	0	0	0	0	0	100	1000000000
26	275	0	0	0	0	0	0	0	0	0	100	1000000000
27	285	0	0	0	0	0	0	0	0	0	100	1000000000
28	295	0	0	0	0	0	0	0	0	0	100	1000000000
29	305	0	0	0	0	0	0	0	0	0	100	1000000000
30	315	0	0	0	0	0	0	0	0	0	100	1000000000

REGISTER

First Name:

Last Name:

Mobile Number:

Country:

City:

State:

Postcode:

Login

Username:

Password:

Forgot Password

User Login

View Healthcare Datasets Trained and Tested Results

Model Type	Accuracy
Convolutional Neural Network (CNN)	93.33333333333333
SVM	96.66666666666667
Logistic Regression	96.66666666666667
Decision Tree Classifier	96.66666666666667

ENTER HERE ALL HEALTH DETAILS IN

Enter PId:

Enter Age:

Enter anemias:

Enter hypertension:

Enter diabetes:

Enter asthma:

Enter high_cholesterol:

Enter stroke:

Enter heart_failure:

Enter kidney_failure:

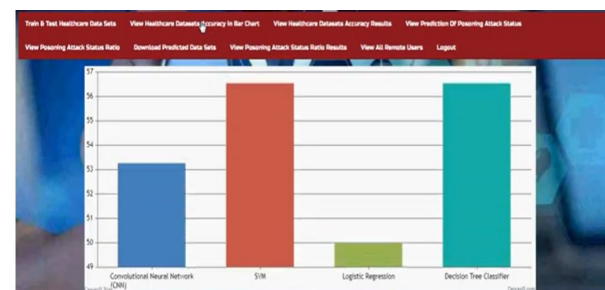
Enter liver_failure:

Enter blood_glucose_level:

Enter blockchain_code_label:

Predict

PREDICTED POISONING ATTACK STATUS



View Predicted Poisoning Attack Status Details II

anemia	hypertension	diabetes	asthma	high_cholesterol	stroke	heart_failure	kidney_failure	liver_failure	blood_glucose_level	blockchain_code_label	Prediction
0	0	0	0	0	0	0	0	0	100	1000000000	Poisoning Attack Found
1	0	0	0	0	0	0	0	0	100	1000000000	No Poisoning Attack Found
1	0	0	0	0	0	0	0	0	100	1000000000	No Poisoning Attack Found
1	0	0	0	0	0	0	0	0	100	1000000000	No Poisoning Attack Found



V. CONCLUSION

In order to safeguard healthcare systems, this article suggests federated learning based on blockchain technology with secure model verification. The primary goal is to guarantee that the local model is free of poisoning while protecting participant privacy and federated learning participants' verifiability.

<https://doi.org/10.62643/ijerst.2025.v21.i2.pp1032-1042>

In this approach, we verify the local model in a privacy-preserving manner prior to the aggregation process. The verification is carried by using an encrypted inference enabled by the SMPC protocol in order to protect privacy on the local model. Using encrypted models and photos, the verifier may verify the model using this technique. The block chain node receives the confirmed share of the local model when it has been validated. SMPC-based secure aggregation will be carried out by the hospital and the block chain. The global model is saved in the block chain after the majority of nodes have the same outcome. Every hospital that participates in the federated learning cycle will thereafter get the updated global model from the tamper-proof storage.

In the experiment, we create local models and aggregate them under FL circumstances using Convolutional Neural Network (CNN) based algorithms using several medical datasets. The findings of our experiment demonstrate that the model encrypted verification procedure may remove all of the participants' tainted models while preserving the local model's privacy. Furthermore, the accuracy of the global model may be recovered by up to 25%. Notably, our secure inference processing time is almost identical to that of the original inference procedure.

We want to provide an effective consensus method for block chain-based aggregation in the future. In this study, we assume that all hospitals construct their local models using the same setup and the homogeneous model. In the future, nevertheless, we want to expand our work to accommodate a heterogeneous model in block chain-based federated learning.

REFERENCES

- [1] L. Sun, X. Jiang, H. Ren, and Y. Guo, "Edge-cloud computing and artificial intelligence in internet of medical things: Architecture, technology and application," *IEEE Access*, vol. 8, pp. 101 079–101 092, 2020.
- [2] X. Li, K. Huang, W. Yang, S. Wang, and Z. Zhang, "On the convergence of fedavg on non-iid data," *arXiv preprint arXiv:1907.02189*, 2019.
- [3] Z. Yu, S. U. Amin, M. Alhussein, and Z. Lv, "Research on disease prediction based on improved deepfm and iomt," *IEEE Access*, vol. 9, pp. 39 043–39 054, 2021.
- [4] W. Wei, L. Liu, M. Loper, K.-H. Chow, M. E. Gursoy, S. Truex, and Y. Wu, "A framework for evaluating client privacy leakages in federated learning," in *European Symposium on Research in Computer Security*. Springer, 2020, pp. 545–566.
- [5] V. Mothukuri, R. M. Parizi, S. Pouriyeh, Y. Huang, A. Dehghantanha, and G. Srivastava, "A survey on security and privacy of federated learning," *Future Generation Computer Systems*, vol. 115, pp. 619–640, 2021.
- [6] Y. Zhao, J. Zhao, M. Yang, T. Wang, N. Wang, L. Lyu, D. Niyato, and K.-Y. Lam, "Local differential privacy-based federated learning for internet of things," *IEEE Internet of Things Journal*, vol. 8, no. 11, pp. 8836–8853, 2021.
- [7] B. Zhao, K. Fan, K. Yang, Z. Wang, H. Li, and Y. Yang, "Anonymous and privacy-preserving federated learning with industrial big data," *IEEE Transactions on Industrial Informatics*, vol. 17, no. 9, pp. 6314–6323, 2021.
- [8] X. Wang, S. Garg, H. Lin, J. Hu, G. Kaddoum, M. Jalil Piran, and M. S. Hossain, "Toward accurate anomaly detection in industrial internet of things using hierarchical federated learning," *IEEE Internet of Things Journal*, vol. 9, no. 10, pp. 7110–7119, 2022.

<https://doi.org/10.62643/ijerst.2025.v21.i2.pp1032-1042>

- [9] V. Mothukuri, P. Khare, R. M. Parizi, S. Pouriyeh, A. Dehghantanha, and G. Srivastava, “Federated-learning-based anomaly detection for iot security attacks,” *IEEE Internet of Things Journal*, vol. 9, no. 4, pp.2545–2554, 2022.
- [10] B. Wang and N. Z. Gong, “Stealing Hyperparameters in Machine Learning,” in 2018 IEEE Symposium on Security and Privacy (SP), 2018, pp. 36–52.
- [11] M. Nasr, R. Shokri, and A. Houmansadr, “Comprehensive privacy analysis of deep learning: Passive and active white-box inference attacks against centralized and federated learning,” in 2019 IEEE Symposium on Security and Privacy (SP), 2019, pp. 739–753.
- [12] V. Tolpegin, S. Truex, M. E. Gursoy, and L. Liu, “Data poisoning attacks against federated learning systems,” in *European Symposium on Research in Computer Security*. Springer, 2020, pp. 480–501.
- [13] C. Zhang, S. Li, J. Xia, W. Wang, F. Yan, and Y. Liu, “Batchcrypt: Efficient homomorphic encryption for cross-silo federated learning,” in 2020 fUSENIXg Annual Technical Conference (fUSENIXgfATCg 20), 2020, pp. 493–506.
- [14] J. Zhang, B. Chen, X. Cheng, H. T. T. Binh, and S. Yu, “PoisonGAN: Generative poisoning attacks against federated learning in edge computing systems,” *IEEE Internet of Things Journal*, vol. 8, no. 5, pp. 3310–3322, 2021.