

Predicting Depression Among Type-2 Diabetic Patients Using Federated Learning

Rabia Tehseen^{1*}, Waseeq Haider¹, Uzma Omer², Nosheen Qamar³, Nosheen Sabahat⁴, Rubab Javaid¹

¹University of Central Punjab, Lahore, Pakistan, 53700

²University of Education, Lahore, Pakistan, 53700

³University of Management and Technology, Lahore, Pakistan, 53700

⁴Forman Christian College University, Lahore, Pakistan, 57400

Correspondence: *rabia.tehseen@ucp.edu.pk

Citation | Tehseen. R, Haider. W, Omer. U, Qamar. N, Sabahat. N, Javaid. R, “Predicting Depression among Type 2 Diabetic Patients Using FL”, IJIST, Vol. 06 Issue. 4 pp `1984-1994, Dec 2024

DOI | <https://doi.org/10.33411/IJIST/1133>

Received | Nov 15 2024 **Revised** | Dec 13, 2024 **Accepted** | Dec 15, 2024 **Published** | Dec 16, 2024.

Depression being a common and dangerous mental health condition could have a significant impact on a person's quality of life. It may result in depressive and gloomy feelings along with a loss of interest in once-enjoyable activities. Depression is considered a leading global cause of impairment that affects people at various stages of age, ethnicities, and socioeconomic status. It may cause negative effects on a person's physical and emotional well-being like reduced motivation, energy, and appetite. In this paper, we have presented a Federated Learning-based framework to predict depression in patients with type 2 diabetes. Type 2 diabetes frequently coexists with depression, which can hurt treatment outcomes and raise medical expenses. The objective of this paper is to create a Federated Learning-based framework to predict the impact of depression in causing type-II diabetes by analyzing patient's data including laboratory results, medical history, and demographic information. To forecast the likelihood of depression in patients with type 2 diabetes. Analysis has been performed using a freely available dataset of Type-II diabetes from Kaggle and an accuracy of 97% has been achieved.

ABBREVIATIONS

WHO	World Health Organization
ML	Machine Learning
FL	Federated Learning
DNN	Deep Neural Network
DFM	Dynamic Factor Model
DMVM	Distributed Memory Vector Model
SVM	Supports Vector Machine
EDA	Exploratory Data Analysis

Keywords: Depression, Federated Learning, Type 2 diabetes, Healthcare



Introduction:

According to the World Health Organization (WHO), nearly 3.8% of people worldwide suffer from depression. Among people with type 2 diabetes, it is a common comorbidity [1]. Depression can severely affect self-care practices, medication adherence, and quality of life in individuals with diabetes. It is also linked to an increased risk of complications, including cardiovascular disease, diabetic retinopathy, and neuropathy. In these people, early detection and treatment of depression can enhance health outcomes and avoid problems.

Conventional techniques, such as self-report questionnaires, are laborious and dependent on patients' voluntary disclosure to detect depression in diabetic patients. Machine Learning (ML) algorithms have demonstrated potential in the past few years for the prediction of depression in a variety of populations, including those with diabetes. Large datasets of electronic health records, or EHRs, can be analyzed via Federated Learning (FL) [2] to find patterns and risk factors related to depression, enabling early detection and treatment.

The purpose of this research is to build and validate FL to predict type 2 diabetes patients' risk of depression. In this study, we have utilized the dataset of a type 2 diabetes patient that contains clinical and demographic information as well as proven depression screening tools. A portion of the data was utilized to train the model, and another portion will be used to assess it. The findings made it possible to identify individuals at risk for depression, enabling early intervention and improving patient outcomes. This research may shed important light on how to use FL to detect depression in diabetic patients and guide the creation of focused therapies that would enhance patient outcomes.

FL is an ML technique that protects privacy by enabling many parties to work together to develop a model without exchanging sensitive data [3]. Since data privacy is a big concern in healthcare settings, this strategy is very helpful there. In this thesis, we suggest predicting depression in type 2 diabetes patients using FL while maintaining patient privacy.

Objectives:

The objectives of this research are:

1. **Privacy Preservation:** Train models without sharing sensitive patient data, ensuring data privacy.
2. **Collaborative Learning:** Enable multiple healthcare institutions to collaboratively improve prediction models while keeping data local.
3. **Accuracy Improvement:** Use diverse datasets to enhance prediction accuracy across populations.
4. **Regulatory Compliance:** Adhere to data privacy regulations like HIPAA and GDPR.
5. **Real-Time Prediction:** Provide timely predictions to assist in early diagnosis and personalized treatment.
6. **Personalized Insights:** Offer tailored health recommendations based on local data.
7. **Scalability and Efficiency:** Handle large datasets efficiently, reducing computational burdens on individual institutions.
8. **Continuous Model Updates:** Allow the model to evolve with new data, improving prediction over time.

Novelty:

The novelty in current research includes:

1. **Privacy Preservation:** Localizes sensitive data, ensuring compliance with privacy regulations.
2. **Collaborative Learning:** Enables institutions to build a unified model without sharing raw data.
3. **Real-Time Personalization:** Provides personalized predictions and recommendations.
4. **Scalability:** Efficiently handles large, distributed datasets.

5. **Continuous Evolution:** Continuously updates the model with new data for improved accuracy.
6. **Cross-Institutional Generalization:** Enhances predictions across diverse populations.

Literature Review:

FL is an ML method that keeps privacy safe by allowing multiple parties to collaborate on a model's development without sharing private information [4]. This tactic is particularly beneficial in healthcare settings because data privacy is a major concern. In this paper, we proposed using FL to predict depression in patients with type 2 diabetes while preserving patient privacy. The association between type 2 diabetes along with how ML techniques can be utilized to predict depression in diabetic patients. The article explains the various kinds of ML algorithms, such as Support vector machine (SVM), K-Mean, F-Cmean, and PNN, that can be applied in this way. The method of using these algorithms for clinical psychologist data is also explained in the paper, along with how type 2 diabetic patients' depression can be predicted using them.

According to the study's findings, the SVM classifier is the best at predicting depression in diabetic patients; nevertheless, for increased accuracy, more learning techniques may be investigated in the future. The paper emphasizes how critical it is to identify depression in diabetes patients as soon as possible and how ML may be used in clinical settings to predict depression [5]. FL models include the study on mood identification utilizing mobile health data. To perform their analysis, they used three models: Deep Neural Network (DNN), Dynamic Factor Model (DFM), and Distributed Memory Vector Model (DMVM). All things considered, their research shows how FL models can be used to analyze mobile health data and identify moods [6].

Researchers utilized cross-sectional data from the 2014 Behavioral Risk Factor Surveillance System, comprising 138,146 individuals, to develop predictive models for identifying risk factors for type 2 diabetes. The authors developed several ML models and used univariable and multivariable weighted logistic regression models to examine the relationships between putative risk variables and type 2 diabetes. The findings demonstrated that every prediction model, with an AUC ranging from 0.7182 to 0.7949, attained a high level of performance. For type 2 diabetes, the decision tree model had the highest sensitivity, but the neural network model had the highest accuracy, specificity, and AUC. The study found that sleeping patterns and the frequency of check-ups could be new risk factors for type 2 diabetes, in addition to the risk variables already documented [7].

Material and Method:

Multiple ML models have been trained using a FL technique to predict depression in individuals with type 2 diabetes. Through the use of FL, several parties, each possessing a portion of the dataset have worked together to build a model while keeping private information confidential.

Methodology workflow:

The proposed methodology involved the following steps as presented in Figure 1. Figure 1a presents the general flow of this research.

Dataset Collection: This research has been conducted on a publicly available dataset named Diabetes dataset, with 1600 records obtained from Kaggle repository [8].

Data Preprocessing: Names, residences, and social security numbers, along with other identifying information, were eliminated from the dataset to preserve the patient's privacy. Pseudonyms or other non-identifiable placeholders were used in place of identifying information, or the data was anonymized. Data about type 2 diabetes patients was gathered from several sources, leading to a variety of data types and architectures. Standardizing the data format refers to transforming the data into a consistent structure that is easy to analyze.

For analysis, the data is required to be translated to a common unit, such as mg/dL or mmol/L, if blood glucose measurements are obtained in multiple units. Medical data frequently contains missing values as a result of incomplete data collection or non-compliant patients. Missing values can be managed in the preprocessing step by either eliminating incomplete observations or imputing missing values using statistical methods like regression or mean imputation. The performance of the model could be affected by the varying scales and ranges of different features in the dataset. To make sure that every feature contributes equally to the performance of the model, feature scaling entails standardizing each feature's range, usually to a range of 0 to 1 or -1 to 1.

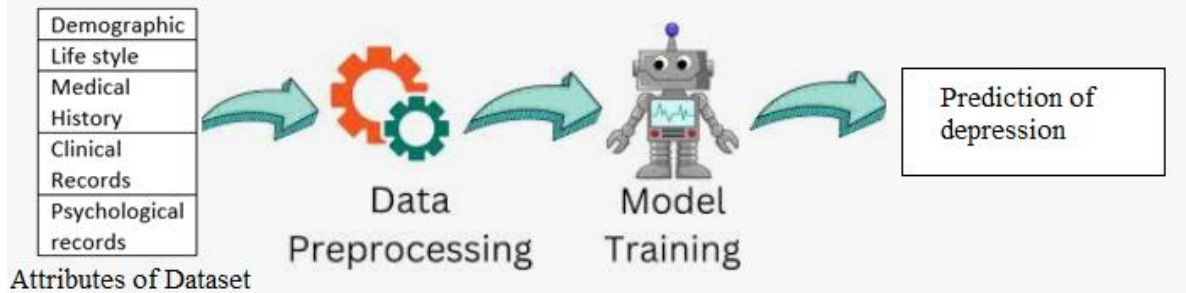


Figure 1a. Flow of this research Methodology.

Feature Selection: Determining the clinical and demographic characteristics that type 2 diabetes patients are most likely to have in common with depression. These include demographics like age, gender, BMI, blood sugar, HbA1c, insulin use, length of diabetes, and psychological parameters like stress, anxiety, and quality of life. Statistical techniques are employed to assess each feature's significance for depression prediction after the original collection of features has been determined. The degree and direction of each feature's relationship to depression can be ascertained using correlation analysis, t-tests, or chi-square testing. The relevance evaluation's findings indicate that a portion of the most pertinent features can be selected for modeling. Only characteristics that are highly predictive of depression in individuals with type 2 diabetes or those that are significantly linked to depression may be included in this subset. After the subset of features has been chosen, methods like holdout validation and cross-validation can be used to assess the model's predicted performance. To identify the best-performing feature set, the subset of features can be compared to other subsets or the entire collection of features.

Model Development: Using typical ML techniques like logistic regression, support vector machines, or deep neural networks, an FL model is trained using these selected features to predict depression in patients with type 2 diabetes

Model Initialization: Describe the first model architecture that is appropriate for type 2 diabetes patients' depression prediction. Based on this design, participants initialize their local models, guaranteeing consistency and compatibility.

Local Model Training: Using their local data, each participant trains their local model independently. Utilize ML algorithms that are appropriate for predicting depression by taking into account variables including the patient's medical history, demographics, diabetes control strategies, and depression evaluation results.

Model Aggregation: To create a global model update, aggregate the local model parameters, also known as gradients. To provide privacy-preserving model aggregation, use federated averaging or secure aggregation techniques.

Model Update and Communication: Using encrypted communication protocols, participants safely share their local model updates or gradients with a central server or with each other. Preserve patient confidentiality when transmitting model updates.

Model Evaluation: Metrics including accuracy, precision, recall, and F1 score will be used to assess the FL model's performance. Creating a training set and a test set from the data is the first stage in assessing the model. The test set is used to assess the model's performance, whereas the training set is used to train the model. Determining the evaluation measures that will be applied to judge the model's performance is the next stage. For binary classification tasks, such as predicting depression, common evaluation metrics include accuracy, precision, recall, and F1 score. The test set is used to assess the FL, and evaluation metrics are calculated. The accuracy metric calculates the percentage of accurate the precision metric calculates the percentage of true positive predictions among all positive forecasts, based on the model's predictions. The F1 score is calculated as the harmonic mean of precision and recall, whereas the recall metric quantifies the percentage of true positive predictions across all real positive cases. The FL model's performance can be interpreted using the evaluation measures. For instance, a high recall shows that the model can properly identify the majority of depression instances, but a high accuracy shows that the model is generally producing accurate predictions. Should the assessment metrics show that the model's functionality is inadequate, the model can be improved by modifying the hyperparameters, choosing alternative features, or employing alternative machine-learning methods.

Model Validation:

The Federated averaging algorithm applied here is presented below:

Steps

1. Initialize Global Model:

- Server initializes model weights $\mathbf{w}_0$ and broadcasts them to all clients.

2. Repeat for T Communication Rounds:

- **Step 2.1:** Server selects a subset of clients $\mathcal{S}_t$ (often chosen randomly, ensuring diversity).
- **Step 2.2:** Each selected client $k \in \mathcal{S}_t$:
 - Downloads the current global model weights $\mathbf{w}_t$.
 - Trains the model locally using E epochs and mini-batches of size B with its dataset $\mathcal{D}_k$.
 - Updates the local weights $\mathbf{w}_k$ using stochastic gradient descent (SGD) or other optimizers:

$$\mathbf{w}_k^{(t+1)} = \mathbf{w}_t - \eta \nabla \ell(\mathbf{w}_t; \mathcal{D}_k)$$

where $\ell(\cdot)$ is the loss function.

- Sends the updated weights $\mathbf{w}_k^{(t+1)}$ back to the server.
- **Step 2.3:** Server performs aggregation to compute the new global model weights:

$$\mathbf{w}_{t+1} = \frac{\sum_{k \in \mathcal{S}_t} n_k \mathbf{w}_k^{(t+1)}}{\sum_{k \in \mathcal{S}_t} n_k}$$

where n_k is the number of data samples at client k .

3. Repeat Until Convergence:

- Continue updating the global model weights until the model performance (e.g., accuracy or loss) converges or the maximum number of communication rounds T is reached.

4. Output Final Model:

- Return the final global model weights $\mathbf{w}$.

Proposed Framework

This study has developed a framework based on FL. As seen in figure 2, it is composed of five layers. In this study, five distinct clients collaborate to predict Type II diabetes.

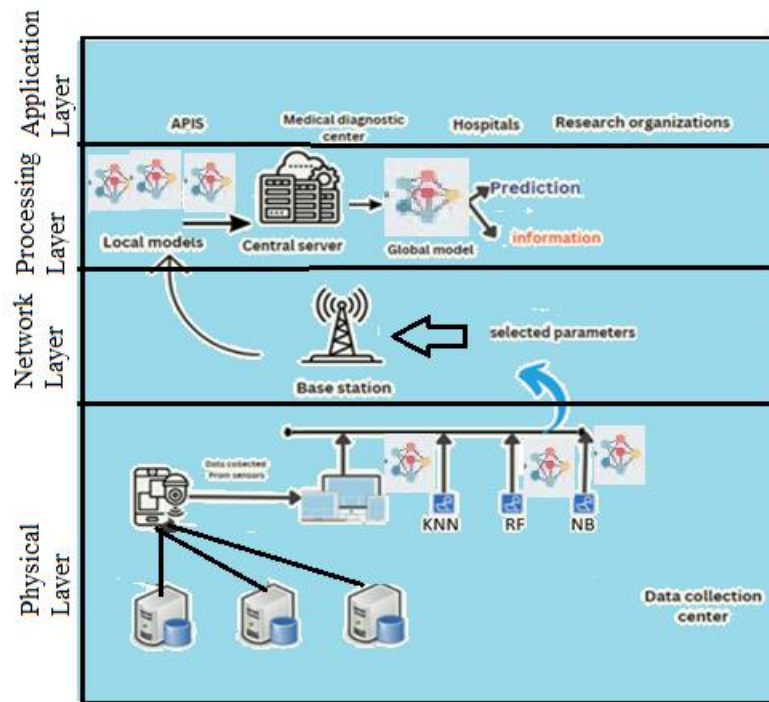


Figure 1. Proposed Framework

Five distinct clients received the dataset, which was displayed by the edge layer in the suggested structure. Using ML techniques like K-Nearest Neighbour, Random Forest, and Naive Bayes, local models were trained on these customers. After training, these models were moved to the processing layer, where local models were combined and the global model was trained using the federated averaging approach. A global model is a centralized model that combines all of the local models that were first obtained from local stations, representing their combined properties. The global model can additionally provide predictions about the extent to which depression can be harmful for type 2 diabetic people.

Statistical Analysis:

To examine the records and data and to analyze the information gathered, statistical methods were applied. Pandas and NumPy, two Python tools, were used in the Jupiter notebook to conduct Exploratory Data Analysis (EDA). Calculations were made to determine the mean, max, min, and standard deviation of every attribute, both overall and by class. To determine how qualities are related to one another, the correlation was calculated after the summary was computed. The frequency bar chart displays the frequency of each class in the dataset since it is crucial to have a thorough understanding of the qualities through data visualization.

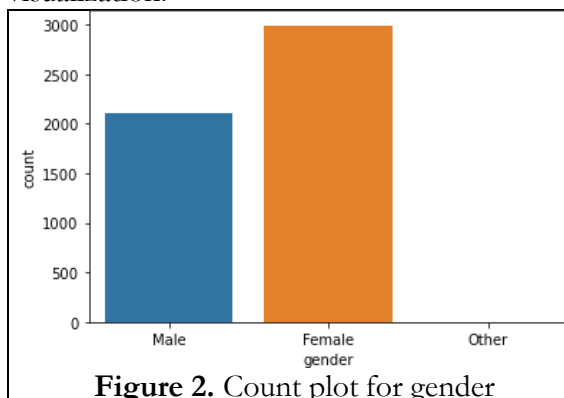


Figure 2. Count plot for gender

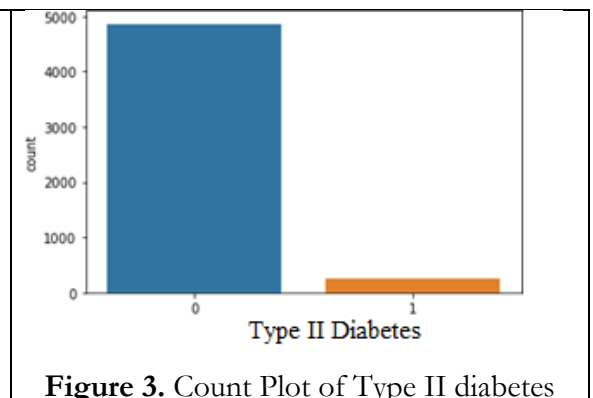


Figure 3. Count Plot of Type II diabetes

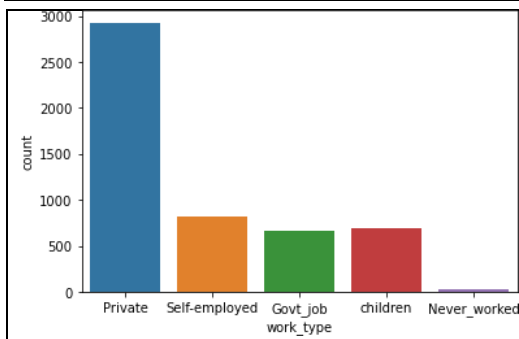


Figure 4. Work type count plot

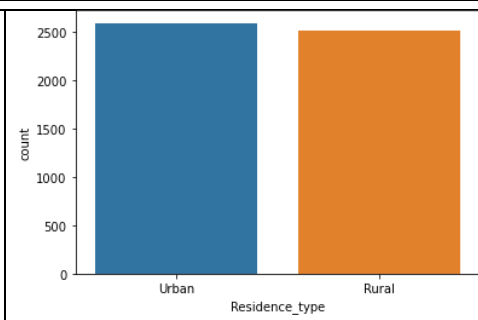


Figure 5. Residence type count plot

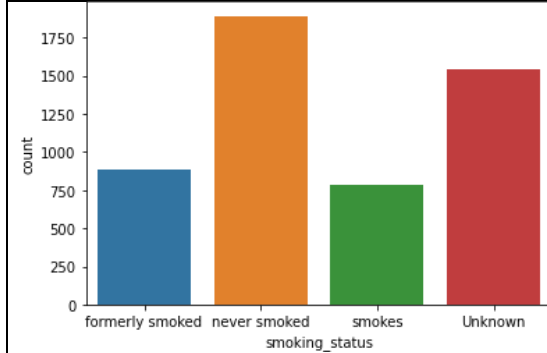


Figure 6. Smoking status count plot

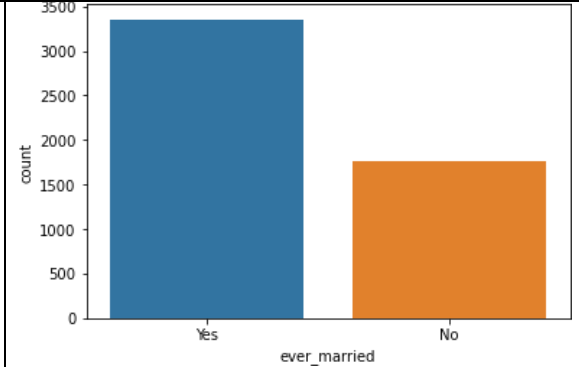


Figure 7. Marital status count plot

Class validation:

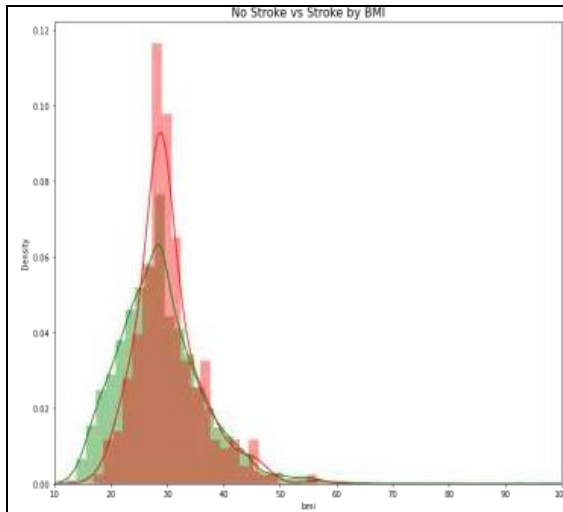


Figure 8: Depression vs diabetes vs bmi

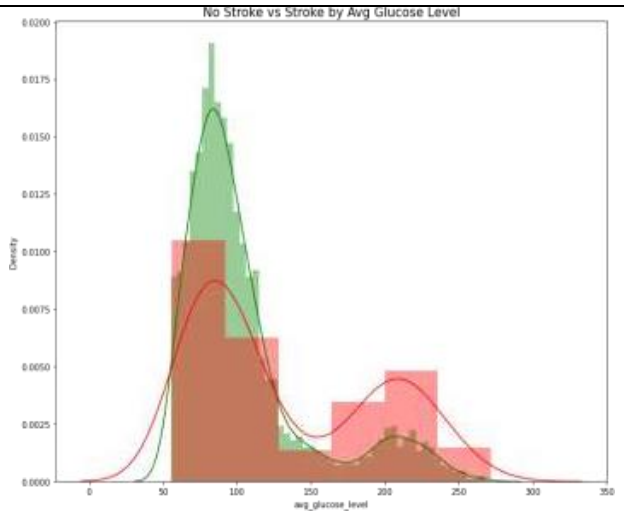


Figure 9: Depression vs diabetes Vs Avg

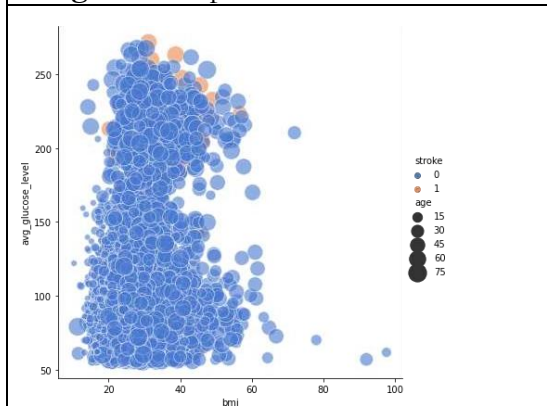


Figure 10. Glucose level and BMI

Depression appears to occur more frequently in individuals with low BMI and high glucose levels.

Results

The EDA provides comprehensive numerical and visual analysis. The next step is to develop a model to predict seed types based on the given characteristics. All Local models were trained using a dataset that is 80% training and 20% testing. We trained four different local models namely Random Forest (RF), Decision Tree (DT), Naïve Bayes (NB) and Support Vector Machine (SVM) on four different clients. By training the aforementioned models on different clients, we obtained the following experimental results.

	precision	recall	f1-score
0	0.95	0.99	0.97
1	0.99	0.95	0.97
accuracy			0.97

Figure 11. Results of RF Model on Client 1

	precision	recall	f1-score
0	0.95	0.94	0.94
1	0.94	0.95	0.95
accuracy			0.95

Figure 12. Results of DT Model on Client 2

	precision	recall	f1-score
0	0.83	0.66	0.74
1	0.72	0.87	0.79
accuracy			0.76

Figure 13. Results of NB Model on Client 3

	precision	recall	f1-score
0	0.81	0.70	0.75
1	0.74	0.83	0.78
accuracy			0.77

Figure 14. Results of SVM Model on Client 4

The results of these local models were transmitted to the central server where the Bagging ensemble technique was applied. On the server, these results were optimized using the federated averaging algorithm, and the global model was trained using the best-performing model, which in this case was Random Forest. The global model was again disseminated to clients for refined predictions at the local level. The combined results of the global model are presented in Table 1.

Table 1. Combined results of Global model

Metric	Class0	Class 1	Overall
Precision	0.92	0.94	0.93
Recall	0.90	0.92	0.91
F1-Score	0.88	0.90	0.89

Confusion matrix:

The confusion matrix of the Global model is presented below:

```
array([[4816, 45],
       [232, 4629]], dtype=int)
```

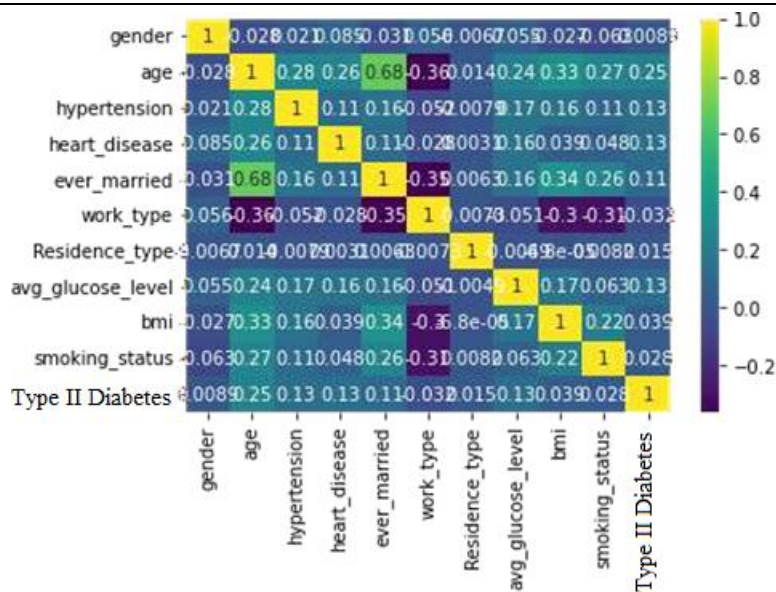



Figure 15. Correlation between variables

Discussions:

Reason for out-performance of Random Forest:

In our work, the random forest model outperformed state-of-the-art models including Decision tree, Naive Bayes, and SCV models in terms of accuracy, recall, f1-score, and precision. This is because our dataset is unbalanced. Remote feature engineering was employed to process the data and train the Random Forest model, which is highly accurate for categorized datasets that do not require human intervention. Precision is the ratio of correctly predicted positive cases to the total predicted positive cases. A low false positive rate indicates a high degree of accuracy. It is a metric for the accuracy of a classifier [9]. Recall is the ratio of accurately predicted positive cases to all positive categorization cases denoted by R. It indicates the completeness of a classification.

Advantages of FL:

Using FL for type II diabetes prediction presents a promising approach that addresses privacy concerns while leveraging distributed data sources to improve model performance. It offers multiple benefits as discussed in Table 2

Table 2. Benefits of FL

Privacy Preservation	Utilization of Distributed Data	Improved Model Generalization	Addressing Data Imbalance and Bias
FL allows model training to be performed locally on individual devices or servers, without sharing raw data. This ensures that sensitive patient information, such as medical records and personal identifiers, remains decentralized and	Type II diabetes prediction models benefit from access to diverse datasets encompassing demographics, medical history, lifestyle factors, and genetic predispositions. FL enables the aggregation of insights from various	FL facilitates the creation of models that generalize well across different healthcare settings and patient demographics. By training on data from multiple sources, the model learns to capture the underlying patterns of diabetes risk	Type II diabetes prediction models may suffer from data imbalance and bias, particularly if certain demographic groups are underrepresented in the training data. FL can help mitigate these issues by aggregating data from diverse sources, thus ensuring a more

confidential. In the context of type II diabetes prediction, this is crucial for maintaining patient privacy and complying with data protection regulations.	healthcare providers, research institutions, and wearable devices, without the need for data centralization. This distributed approach allows the model to learn from a broader spectrum of patient populations, leading to more robust predictions.	factors that are relevant across diverse populations. This enhances the model's ability to make accurate predictions for individuals who may have different backgrounds or access to healthcare resources.	balanced representation of patient populations.
--	--	--	---

Challenges:

Despite its benefits, FL presents several challenges, including communication overhead, heterogeneous data formats, and ensuring model consistency across decentralized nodes. Additionally, ensuring the security and integrity of FL systems is paramount to prevent adversarial attacks or data breaches. Healthcare organizations must implement robust security protocols and encryption techniques to safeguard patient data throughout the FL process.

Comparison of Traditional ML and FL Models:

Traditional ML models typically rely on centralized data storage, where all patient information is aggregated for model training. This approach can lead to privacy concerns, especially in healthcare, where sensitive data like patient demographics and medical history is involved. Moreover, traditional models may suffer from data imbalance or overfitting when trained on limited or homogeneous datasets, reducing their generalizability across diverse populations. We have compared our work with other state-of-the-art as presented in Table 3.

Table 3. Comparison Table

Ref.	ML-models	Precision	Recall	F1-Score
[15]	DT	95	94	94
[16]	KNN	95	99	97
[17]	Naïve Bayes	90	80	88
[18]	SVM	75	88	79
Proposed work	FL	93	91	89

In contrast, FL offers a decentralized framework, where model training occurs locally on multiple devices or institutions, ensuring patient data remains private. This method enables access to diverse, distributed datasets without the need for centralization, leading to more robust and accurate predictions across different patient populations. FL also addresses challenges like data imbalance by allowing aggregation of data from various sources, enhancing the model's fairness and reducing bias.

Table 4. Feature-wise comparison of ML and FL

Feature	Centralized ML	FL
Data Privacy	High risk of data breaches due to centralization	Data remains decentralized, minimizing privacy risks
Data Centralization	Requires central data storage	No data centralization; local training
Model Generalization	Limited by dataset size and diversity	Improved through access to distributed datasets

Bias & Imbalance	High risk of bias due to imbalanced datasets	Mitigates bias by aggregating diverse data
Communication Overhead	Minimal, as data is centrally stored	Higher, due to decentralized model synchronization

Conclusion and Future Directions:

In conclusion, the utilization of FL for predicting depression among Type 2 diabetic patients presents a promising avenue for healthcare innovation. FL addresses critical concerns associated with sensitive health information by preserving data privacy and security through decentralized model training. However, to fully realize its potential, future research should focus on addressing technical challenges such as communication overhead and model synchronization, while ensuring the quality and representativeness of local datasets. Additionally, exploring advanced techniques for model aggregation and federated optimization algorithms could further enhance the accuracy and efficiency of predictive models. As FL continues to evolve, its application in healthcare is poised to revolutionize patient care by enabling robust and privacy-preserving predictive analytics for complex medical conditions like depression in diabetic populations.

ACKNOWLEDGEMENT

We are thankful to Mr. M. Inayat Ali for his valuable help and guidance throughout this research.

REFERENCES:

- [1] I. M. Cheng, Y. L., Wu, Y. R., Lin, K. D., Lin, C. H. R., & Lin, "Using Machine Learning for the Risk Factors Classification of Glycemic Control in Type 2 Diabetes Mellitus.," *Healthc.*, vol. 11, no. 8, p. 1141, 2023.
- [2] M. Bank, S., & Das, "Periodic study of electrocardiogram abnormalities can predict major adverse cardiac events in Type 2 diabetics.," *Eur. J. Prev. Cardiol.*, vol. 30, no. 8, pp. 654–655, 2023.
- [3] N. Kee, O. T., Harun, H., Mustafa, N., Abdul Murad, N. A., Chin, S. F., Jaafar, R., & Abdullah, "Cardiovascular complications in a diabetes prediction model using Machine Learning: a systematic review.," *Cardiovasc. Diabetol.*, vol. 22, no. 1, p. 13, 2023.
- [4] Y. F. Liu, C. H., Peng, C. H., Huang, L. Y., Chen, F. Y., Kuo, C. H., Wu, C. Z., & Cheng, "The comparison between multiple linear regression and Machine Learning methods in predicting cognitive function in Chinese type 2 diabetes.," 2023.
- [5] Z. Farooq, M. S., Tehseen, R., Qureshi, J. N., Omer, U., Yaqoob, R., Tanweer, H. A., & Atal, "FFM: Flood forecasting model using FL," *IEEE Access*, pp. 24472–24483, 2023.
- [6] L. Hennebelle, A., Materwala, H., & Ismail, "HealthEdge: a Machine Learning -based smart healthcare framework for prediction of type 2 diabetes in an integrated IoT, edge, and cloud computing system.," *Procedia Comput. Sci.*, vol. 220, pp. 331–338, 2023.
- [7] W. Al Sadi, K., & Balachandran, "Prediction Model of Type 2 Diabetes Mellitus for Oman Prediabetes Patients Using Artificial Neural Network and Six Machine Learning Classifiers," *Appl. Sci.*, vol. 13, no. 4, p. 2344, 2023.
- [8] "Dataset link <https://www.kaggle.com/datasets/mathchi/diabetes-data-set>".
- [9] K. T. Arrayyan, A. Z., Setiawan, H., & Putra, "Naive Bayes for Diabetes Prediction: Developing a Classification Model for Risk Identification in Specific Populations," *Semesta Tek.*, vol. 27, no. 1, pp. 28–36, 2024.



Copyright © by authors and 50Sea. This work is licensed under Creative Commons Attribution 4.0 International License.