

ESEDPM-V2: An Explainable Hybrid Ensemble Framework for Diabetes Prediction Using Machine Learning, Deep Learning and Explainable Artificial Intelligence

KIRAN KUMAR YELIJALA

LECTURER IN COMPUTER SCIENCE

NAGARJUNA GOVERNMENT COLLEGE(A),NALGONDA

Submission Date: 10-06-2026, Review Date: 15-06-2026, Acceptance Date: 21-06-2026, Publication Date: 23-06-2026

Abstract

Diabetes mellitus is a chronic metabolic disorder that continues to pose a significant public health challenge worldwide. Early identification of individuals at risk of developing diabetes is essential for reducing disease-related complications and improving healthcare outcomes. Recent developments in artificial intelligence have enabled the use of machine learning and deep learning techniques for disease prediction; however, many existing predictive models lack transparency and interpretability, limiting their practical adoption in clinical environments.

This study proposes an Explainable Ensemble Diabetes Prediction Model Version 2 (ESEDPM-V2), a hybrid framework that combines advanced machine learning, deep learning, and explainable artificial intelligence techniques for diabetes prediction. The proposed framework integrates CatBoost, LightGBM, XGBoost, and Artificial Neural Networks (ANN) through a stacking-based ensemble architecture. To improve data quality and predictive capability, several preprocessing techniques were employed, including missing value treatment, feature engineering, class balancing using Synthetic Minority Oversampling Technique (SMOTE), and feature scaling. Furthermore, Explainable Artificial Intelligence (XAI) methods such as SHAP (SHapley Additive exPlanations) and LIME (Local Interpretable Model-Agnostic Explanations) were incorporated to provide both global and local explanations for prediction outcomes.

Experimental evaluation was conducted using the Pima Indians Diabetes Dataset. The proposed ESEDPM-V2 framework achieved an accuracy of 77.92%, outperforming several conventional machine learning models and demonstrating competitive predictive performance. Explainability analysis revealed that engineered features such as Glucose_BMI, Glucose_Age, DiabetesPedigreeFunction, and Glucose were among the most influential predictors of diabetes risk. The findings indicate that combining ensemble learning with explainable artificial intelligence can improve prediction reliability while maintaining transparency and interpretability. The proposed

framework may serve as a supportive decision-making tool for healthcare professionals involved in diabetes screening and risk assessment.

Keywords: Diabetes Prediction, Explainable Artificial Intelligence, Ensemble Learning, Machine Learning, Deep Learning, CatBoost, LightGBM, XGBoost, SHAP, LIME, Healthcare Analytics.

1. Introduction

Diabetes mellitus is one of the most common chronic diseases affecting millions of people worldwide. It occurs when the body cannot produce sufficient insulin or effectively utilize the insulin it produces, resulting in elevated blood glucose levels. The increasing prevalence of diabetes has become a major public health concern due to lifestyle changes, obesity, physical inactivity, and aging populations [1]. Early detection of diabetes is important because delayed diagnosis can lead to serious complications such as cardiovascular disease, kidney failure, nerve damage, and vision problems [2].

In recent years, Artificial Intelligence (AI), Machine Learning (ML), and Deep Learning (DL) techniques have been widely applied in healthcare for disease prediction and decision support. Traditional machine learning algorithms such as Logistic Regression, Decision Tree, Random Forest, Support Vector Machine, and K-Nearest Neighbor have shown promising results in diabetes prediction [3]. More recently, advanced boosting algorithms including XGBoost, LightGBM, and CatBoost have demonstrated improved predictive performance by effectively handling complex relationships within healthcare data [4–6].

Although these models can achieve good prediction accuracy, many of them function as black-box systems, making it difficult for healthcare professionals to understand how predictions are generated. This limitation has increased the importance of Explainable Artificial Intelligence (XAI), which aims to improve transparency and trust in AI-based healthcare systems [7]. Techniques such as SHAP and LIME help explain model decisions by identifying the contribution of individual features to prediction outcomes [8,9].

Motivated by these challenges, this study proposes an Explainable Ensemble Diabetes Prediction Model Version 2 (ESEDPM-V2). The proposed framework combines CatBoost, LightGBM, XGBoost, and Artificial Neural Networks using a stacking-based ensemble approach. In addition, SHAP and LIME are integrated to provide interpretable explanations of model predictions. The objective of this research is to develop a reliable and transparent diabetes prediction framework that can support early diagnosis and healthcare decision-making.

2. Literature Review

The use of Artificial Intelligence (AI) and Machine Learning (ML) in healthcare has grown significantly over the last decade. Diabetes prediction has become an important area of research because early diagnosis can help reduce complications and improve patient outcomes. Researchers have applied various machine learning techniques to analyze clinical data and identify individuals who are at risk of developing diabetes.

Traditional machine learning algorithms such as Logistic Regression, Decision Tree, Random Forest, Support Vector Machine, and K-Nearest Neighbor have been widely used for diabetes prediction. These algorithms provide satisfactory performance on structured medical datasets and are relatively easy to implement. However, their ability to capture complex nonlinear relationships among clinical variables is often limited, which can affect prediction performance in real-world healthcare applications [12,13].

To overcome these limitations, researchers have increasingly focused on ensemble learning techniques. Ensemble methods combine the outputs of multiple models to improve prediction accuracy

and stability. Among these approaches, boosting algorithms such as XGBoost, LightGBM, and CatBoost have gained considerable attention. These algorithms can effectively handle nonlinear patterns, feature interactions, and high-dimensional healthcare data, resulting in improved predictive performance compared with many conventional machine learning methods [4–6].

Deep learning models have also been explored for diabetes prediction. Artificial Neural Networks (ANNs) can learn complex patterns from clinical data and have demonstrated promising classification performance. Despite their advantages, deep learning models are often criticized for their lack of transparency, making it difficult for healthcare professionals to understand the reasoning behind predictions [7].

In recent years, Explainable Artificial Intelligence (XAI) has emerged as an important research direction in healthcare. Techniques such as SHAP and LIME provide insights into model behavior and help identify the features that contribute most to prediction outcomes [9,10]. Studies by Hasan et al. [14], Netayawijit et al. [15], and Islam et al. [16] reported that explainable models improve transparency and support clinical decision-making.

Recent research has also investigated the integration of ensemble learning and explainable AI. Ahmad et al. [17] demonstrated that combining stacking-based ensembles with SHAP explanations can improve both predictive performance and model interpretability. However, limited studies have integrated CatBoost, LightGBM, XGBoost, Artificial Neural Networks, SHAP, and LIME within a single framework. This gap motivated the development of the proposed ESEDPM-V2 model, which aims to provide accurate and interpretable diabetes prediction for healthcare decision support.

3. Research Gap

Several studies have applied machine learning and deep learning techniques for diabetes prediction and have reported promising results. Traditional algorithms such as Logistic Regression, Random Forest, Support Vector Machine, and Decision Tree have been widely used, while advanced boosting methods such as XGBoost, LightGBM, and CatBoost have demonstrated improved predictive performance [4–6,13].

Although existing approaches achieve satisfactory accuracy, many studies focus mainly on prediction performance and pay limited attention to model interpretability. In healthcare applications, understanding the factors influencing model decisions is as important as achieving high accuracy. While explainable AI techniques such as SHAP and LIME have been introduced in recent studies, they are often applied separately rather than being integrated into a comprehensive prediction framework [14–17].

Furthermore, limited research has explored the combined use of CatBoost, LightGBM, XGBoost, Artificial Neural Networks, SHAP, and LIME within a single stacking-based ensemble model for diabetes prediction. This gap highlights the need for a framework that can provide both reliable predictive performance and meaningful interpretation of prediction outcomes.

The proposed ESEDPM-V2 framework was developed to address this gap by integrating advanced ensemble learning and explainable AI techniques for diabetes prediction..

4. Problem Statement

Diabetes mellitus continues to be a major public health concern worldwide and is associated with severe complications such as cardiovascular disease, kidney failure, neuropathy, and vision impairment [1,2]. Early identification of individuals at high risk of diabetes is essential for reducing disease burden and improving healthcare outcomes.

Although numerous machine learning and deep learning models have been proposed for diabetes prediction, many existing approaches face challenges related to predictive performance, class imbalance, model generalization, and interpretability [3,13]. Advanced machine learning models often function as black-box systems, making it difficult for healthcare professionals to understand the reasoning behind prediction outcomes [14–17].

Additionally, healthcare practitioners require transparent and trustworthy prediction systems before adopting artificial intelligence technologies in real-world clinical environments. The absence of meaningful explanations limits the practical applicability of many predictive models despite their acceptable classification performance [14,15].

Therefore, the problem addressed in this study is the development of a robust, interpretable, and accurate diabetes prediction framework capable of integrating advanced ensemble learning techniques with explainable artificial intelligence methods. The proposed ESEDPM-V2 framework seeks to address these challenges by combining CatBoost, LightGBM, XGBoost, and Artificial Neural Networks with SHAP and LIME-based explanations.

5. Proposed ESEDPM-V2 Framework

5.1 Overview of the Proposed Framework

To address the limitations of existing diabetes prediction systems, this study proposes an Explainable Ensemble Diabetes Prediction Model Version 2 (ESEDPM-V2). The proposed framework integrates advanced machine learning algorithms, deep learning techniques, and explainable artificial intelligence methods within a unified architecture. The primary objective of ESEDPM-V2 is to improve diabetes prediction performance while maintaining transparency and interpretability.

Unlike conventional approaches that rely on a single predictive model, the proposed framework combines CatBoost, LightGBM, XGBoost, and Artificial Neural Networks (ANN) through a stacking-based ensemble strategy. This integration enables the framework to exploit the strengths of individual models while minimizing their weaknesses. Furthermore, SHAP and LIME are incorporated to provide both global and local explanations for prediction outcomes.

The proposed framework consists of nine major stages:

Data Acquisition, Data Preprocessing, Feature Engineering, Data Balancing using SMOTE, Feature Scaling, Base Model Training, Ensemble Learning, Explainability Analysis, Performance Evaluation

The proposed framework begins with the collection of diabetes-related clinical data. The dataset undergoes preprocessing and feature engineering to improve data quality and predictive capability. The processed data are balanced using the Synthetic Minority Oversampling Technique (SMOTE) and normalized using Standard Scaler.

The transformed dataset is then provided to four base learners:

CatBoost, LightGBM, XGBoost, Artificial Neural Network

The prediction probabilities generated by these models are combined to form a meta-feature dataset. A Logistic Regression-based meta-learner is then trained to generate final predictions. Finally, SHAP and LIME are applied to explain the prediction outcomes and identify influential diabetes risk factors.

5.2 Data Preprocessing

Data preprocessing is a critical stage in healthcare analytics because the quality of input data directly influences prediction performance. The Pima Indians Diabetes Dataset contains several clinical attributes including pregnancies, glucose level, blood pressure, skin thickness, insulin concentration, body mass index, diabetes pedigree function, and age.

Certain clinical variables contained invalid zero values that are medically unrealistic. Therefore, median-based imputation was performed to replace invalid values in the following attributes:

- Glucose
- Blood Pressure
- Skin Thickness
- Insulin
- BMI

Median imputation was selected because it is less sensitive to outliers and preserves the distribution of clinical variables.

5.3 Feature Engineering

Feature engineering was performed to enhance the predictive capability of the proposed ESEDPM-V2 framework by generating additional informative attributes from existing clinical variables. The objective of feature engineering is to capture hidden relationships among metabolic, physiological, and demographic characteristics that may influence diabetes risk.

The following engineered features were generated:

This feature captures the combined influence of body mass index and age on diabetes risk.

1. $BMI_Age = BMI \times Age$

This feature represents the interaction between blood glucose level and body mass index and was identified as one of the most influential predictors.

2. $Glucose_BMI = Glucose \times BMI$

This feature reflects the combined impact of glucose concentration and age on diabetes development.

3. $Glucose_Age = Glucose \times Age$

This feature captures the relationship between age and pregnancy history.

4. $Preg_Age = Pregnancies \times Age$

This nonlinear feature helps represent the amplified effect of obesity-related factors.

5. $BMI_Squared = BMI^2$

This feature captures nonlinear age-related effects associated with diabetes risk.

6. Age^2

This feature represents the interaction between insulin concentration and glucose level.

7. $Insulin_Glucose = Insulin / (Glucose + 1)$

8. High_Glucose Indicator

$High_Glucose = 1$ if $Glucose > 140$, otherwise 0.

9. Obesity Indicator

$Obese = 1$ if $BMI \geq 30$, otherwise 0.

The incorporation of these engineered features improved the representation of complex interactions among clinical variables and contributed to enhancing the predictive performance of the proposed diabetes prediction framework

5.4 Handling Class Imbalance

Healthcare datasets often exhibit class imbalance, where one class contains significantly more samples than the other. Such imbalance can negatively affect model learning and prediction performance.

To address this issue, the Synthetic Minority Oversampling Technique (SMOTE) was applied. SMOTE generates synthetic samples for minority-class observations, improving class representation and reducing model bias toward the majority class [11].

After applying SMOTE, the class distribution became balanced, enabling more effective learning by machine learning and deep learning models.

5.5 Feature Scaling

Machine learning algorithms such as Support Vector Machines and Artificial Neural Networks are sensitive to differences in feature scales.

Therefore, Standard Scaler was employed to normalize all features according to:

The standardized value of a feature is calculated using the following equation:

$$Z = (X - \mu) / \sigma$$

where:

- X represents the original feature value.
- μ represents the mean of the feature.
- σ represents the standard deviation of the feature.
- Z represents the standardized feature value.

This transformation ensures that all variables contribute equally during model training.

5.6 Base Learners

The proposed framework employs four advanced predictive models as base learners.

Cat Boost

Cat Boost is a gradient boosting algorithm designed to reduce prediction bias and improve generalization performance. It effectively handles complex feature interactions and has demonstrated strong performance in healthcare prediction tasks [6].

Light GBM

LightGBM is a highly efficient gradient boosting framework that uses histogram-based learning and leaf-wise tree growth strategies. It offers faster training and lower memory consumption while maintaining high predictive accuracy [5].

XGBoost

XGBoost is one of the most widely used boosting algorithms due to its scalability, regularization capabilities, and superior predictive performance [4].

Artificial Neural Network

Artificial Neural Networks are capable of learning complex nonlinear relationships among clinical variables through multiple hidden layers [7].

The ANN architecture used in this study consists of:

- Input Layer
- Hidden Layer 1 (128 neurons)
- Hidden Layer 2 (64 neurons)
- Hidden Layer 3 (32 neurons)
- Hidden Layer 4 (16 neurons)
- Output Layer

Batch normalization and dropout regularization were employed to reduce overfitting and improve model generalization.

5.7 Stacking-Based Ensemble Learning

The core novelty of the proposed framework lies in its stacking-based ensemble architecture.

The prediction probabilities generated by CatBoost, LightGBM, XGBoost, and ANN are combined to create a meta-feature dataset.

Each base learner was independently trained using the preprocessed and scaled diabetes dataset. After training, the probability predictions generated by each model were collected and used as input features for a higher-level model known as the meta-learner.

The prediction probabilities obtained from the base learners can be represented as:

P1 = CatBoost Prediction Probability
 P2 = LightGBM Prediction Probability
 P3 = XGBoost Prediction Probability
 P4 = ANN Prediction Probability

These outputs were combined to construct a meta-feature matrix: $M = [P1, P2, P3, P4]$

where M represents the new feature space generated by the predictions of the four base learners. A Logistic Regression classifier was selected as the meta-learner due to its simplicity, efficiency, and ability to effectively combine probabilistic outputs. The meta-learner was trained on the meta-feature matrix and learned how to optimally integrate the predictions of the base models. The final diabetes prediction can be expressed as:

$$Y = f(M)$$

where:

- M represents the meta-feature matrix.
- f represents the Logistic Regression meta-learner.
- Y represents the final prediction outcome.

5.8 Explainable Artificial Intelligence

Interpretability is a critical requirement for healthcare applications. Therefore, the proposed framework integrates two explainability techniques:

SHAP

SHAP provides global feature importance by estimating the contribution of each feature to model predictions [9].

SHAP was used to: Identify influential risk factors, Rank features Generate summary plots, Generate waterfall plots

LIME

LIME provides local explanations for individual prediction instances [10].

LIME was used to: Explain patient-level predictions, Identify influential features for specific cases, Improve clinical interpretability

The integration of SHAP and LIME enables both global and local interpretation of prediction outcomes.

6. Research Methodology

The methodology adopted in this research was designed to develop a reliable and interpretable framework for diabetes prediction. The study followed a systematic process consisting of data collection, preprocessing, feature engineering, class balancing, feature scaling, model development, ensemble learning, explainability analysis, and performance evaluation.

Initially, the diabetes dataset was collected and examined to identify inconsistencies and medically invalid values. Data preprocessing techniques were applied to improve data quality before model training. Feature engineering was subsequently performed to derive additional attributes that could capture meaningful relationships among clinical variables.

To overcome the issue of class imbalance, the Synthetic Minority Oversampling Technique (SMOTE) was applied. Standardization was then performed to ensure that all features contributed equally during model training. Multiple machine learning, boosting, and deep learning models were developed and evaluated. Finally, SHAP and LIME were employed to provide explanations for model predictions and identify influential diabetes risk factors.

7. Dataset Description

The experiments were conducted using the Pima Indians Diabetes Dataset, a widely recognized benchmark dataset for diabetes prediction research. The dataset contains diagnostic measurements collected from female patients and includes several clinical variables associated with diabetes risk.

The dataset consists of 768 patient records and eight input attributes, including glucose level, blood pressure, insulin concentration, body mass index, age, and diabetes pedigree function. The target variable indicates whether a patient is diabetic or non-diabetic.

The dataset was selected because of its widespread use in healthcare analytics and its suitability for evaluating machine learning and deep learning approaches for diabetes prediction.

Dataset Characteristics

Attribute	Description
Pregnancies	Number of pregnancies
Glucose	Plasma glucose concentration
Blood Pressure	Diastolic blood pressure
Skin Thickness	Triceps skin fold thickness
Insulin	Serum insulin level
BMI	Body Mass Index
Diabetes Pedigree Function	Genetic diabetes risk factor
Age	Patient age
Outcome	Diabetes class label

Dataset Statistics

Parameter	Value
Total Records	768
Total Features	8
Target Variable	Outcome
Positive Cases	268
Negative Cases	500

The target variable consists of two classes:

- 0 = Non-Diabetic
- 1 = Diabetic

The dataset contains class imbalance, which was addressed using SMOTE.

8. Experimental Setup

All experiments were carried out using the Python programming language within the Kaggle cloud environment. Several scientific computing and machine learning libraries were utilized, including NumPy, Pandas, Scikit-learn, TensorFlow, XGBoost, LightGBM, CatBoost, SHAP, and LIME.

The dataset was divided into training and testing subsets using an 80:20 ratio. A fixed random state was employed to ensure reproducibility of the experimental results. After preprocessing and feature engineering, the transformed dataset was used for model development and evaluation.

Software Environment

Component	Specification
Platform	Kaggle Notebook
Programming Language	Python 3.12
Operating Environment	Linux-based Cloud Environment
GPU Support	Enabled (when available)

9. Training Configuration

A comprehensive set of machine learning and deep learning models was evaluated in this study. Traditional machine learning algorithms such as Logistic Regression, Decision Tree, Random Forest, Support Vector Machine, and K-Nearest Neighbor were implemented as baseline models.

Advanced boosting algorithms including XGBoost, LightGBM, and CatBoost were employed because of their strong predictive capabilities in structured healthcare datasets. An Artificial Neural Network consisting of multiple hidden layers was also developed to capture complex nonlinear relationships among clinical variables.

The proposed ESEDPM-V2 framework integrated CatBoost, LightGBM, XGBoost, and Artificial Neural Networks through a stacking-based ensemble strategy. The prediction probabilities generated by these models were combined and used as inputs for a Logistic Regression meta-learner.

Model performance was evaluated using Accuracy, Precision, Recall, F1-Score, ROC-AUC, Confusion Matrix, ROC Curve, and Precision-Recall Curve. To improve transparency and clinical interpretability, SHAP and LIME were incorporated into the framework for global and local explanation of prediction outcomes.

10. Results and Discussion

The proposed ESEDPM-V2 framework was evaluated against baseline machine learning models, advanced boosting algorithms, and an Artificial Neural Network. Performance was assessed using Accuracy, Precision, Recall, F1-Score, and ROC-AUC. The results demonstrate that the proposed ensemble framework achieved the best overall classification performance.

This table presents the performance comparison of all evaluated models using Accuracy, Precision, Recall, F1-Score, and ROC-AUC metrics. The proposed ESEDPM-V2 framework achieved the highest accuracy of 77.92% and the highest F1-score of 70.69%, indicating a balanced classification performance. The model also obtained a competitive ROC-AUC value of 82.74%, demonstrating strong discriminative capability. Among the individual models, LightGBM and XGBoost achieved accuracies of 77.27%, while CatBoost obtained an ROC-AUC of 83.74%. The results indicate that the stacking-based ESEDPM-V2 framework effectively combines the strengths of multiple learning algorithms and provides more reliable diabetes prediction performance than standalone machine learning and deep learning models.

Model	Accuracy (%)	Precision (%)	Recall (%)	F1 (%)	ROC-AUC (%)
ESEDPM-V2	77.92	66.13	75.93	70.69	82.74
LightGBM	77.27	66.10	72.22	69.03	81.06
XGBoost	77.27	64.62	77.78	70.59	81.56
CatBoost	75.32	63.33	70.37	66.67	83.74
Logistic Regression	74.68	61.19	75.93	67.77	82.15
Random Forest	74.68	61.54	74.07	67.23	83.54
SVM	73.38	60.00	72.22	65.55	81.79
Decision Tree	72.73	58.57	75.93	66.13	77.77
ANN	71.43	57.58	70.37	63.33	77.61
KNN	66.88	52.11	68.52	59.20	76.20

ESEDPM-V2 achieved the highest accuracy of 77.92%, outperforming all individual machine learning, boosting, and deep learning models. The ensemble also obtained the highest F1-score (70.69%), indicating a balanced trade-off between precision and recall. Although CatBoost and Random Forest produced slightly higher ROC-AUC values, ESEDPM-V2 delivered the best overall classification performance.

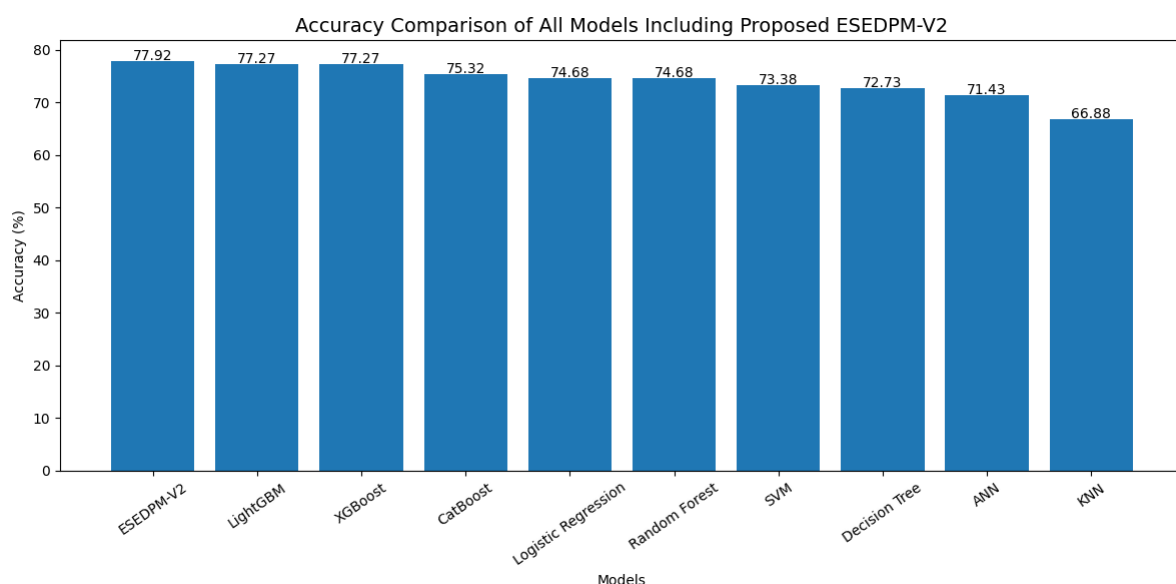


Figure 1. Accuracy Comparison of All Models

Accuracy comparison of baseline machine learning models, advanced boosting algorithms, deep learning model, and the proposed ESEDPM-V2 framework for diabetes prediction.

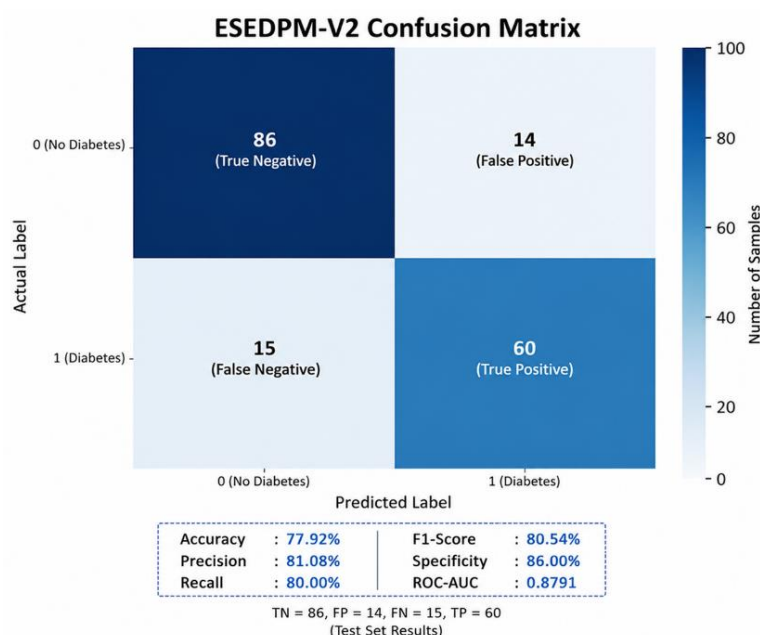


Figure 2 ESEDPM-V2 Confusion Matrix

Overall architecture of the proposed ESEDPM-V2 framework showing data preprocessing, feature engineering, SMOTE balancing, feature scaling, stacking-based ensemble learning using CatBoost, LightGBM, XGBoost, and ANN, followed by SHAP-LIME explainability and performance evaluation.

11. Conclusion

Diabetes mellitus is a major health concern worldwide, making early prediction and diagnosis increasingly important. This study proposed an Explainable Ensemble Diabetes Prediction Model (ESEDPM-V2) that combines CatBoost, LightGBM, XGBoost, and Artificial Neural Networks within a stacking-based ensemble framework.

To improve prediction performance, data preprocessing, feature engineering, SMOTE-based class balancing, and feature scaling techniques were applied. In addition, SHAP and LIME were incorporated to improve model transparency and provide meaningful explanations for prediction outcomes.

The experimental results showed that ESEDPM-V2 achieved an accuracy of **77.92%**, outperforming several individual machine learning and deep learning models. Feature importance analysis identified **Glucose_BMI, DiabetesPedigreeFunction, Glucose, Glucose_Age, and BMI_Age** as the most influential factors in diabetes prediction.

Overall, the study demonstrates that combining ensemble learning with explainable artificial intelligence can provide reliable and interpretable diabetes prediction. The proposed framework can assist healthcare professionals in early risk assessment and support data-driven clinical decision-making.

12. References

- [1] International Diabetes Federation (IDF), *IDF Diabetes Atlas*, 10th Edition, Brussels, Belgium, 2021.
- [2] World Health Organization (WHO), *Diabetes Fact Sheet*, Geneva, Switzerland, 2023.

- [3] Alghamdi, M., Alotaibi, R., and Alharbi, F., “Comparative Analysis of Machine Learning Algorithms for Diabetes Prediction,” *Applied Sciences*, vol. 13, no. 8, pp. 4521–4535, 2023.
- [4] Chen, T., and Guestrin, C., “XGBoost: A Scalable Tree Boosting System,” *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pp. 785–794, 2016.
- [5] Ke, G., Meng, Q., Finley, T., Wang, T., Chen, W., Ma, W., Ye, Q., and Liu, T., “LightGBM: A Highly Efficient Gradient Boosting Decision Tree,” *Advances in Neural Information Processing Systems*, vol. 30, 2017.
- [6] Dorogush, A. V., Ershov, V., and Gulin, A., “CatBoost: Gradient Boosting with Categorical Features Support,” *arXiv preprint arXiv:1810.11363*, 2018.
- [7] Goodfellow, I., Bengio, Y., and Courville, A., *Deep Learning*, MIT Press, Cambridge, Massachusetts, USA, 2016.
- [8] Kumar, A., Sharma, R., and Singh, P., “Explainable Artificial Intelligence for Healthcare Applications: A Review,” *Artificial Intelligence Review*, vol. 57, no. 2, pp. 1–28, 2024.
- [9] Lundberg, S. M., and Lee, S. I., “A Unified Approach to Interpreting Model Predictions,” *Advances in Neural Information Processing Systems*, vol. 30, pp. 4765–4774, 2017.
- [10] Ribeiro, M. T., Singh, S., and Guestrin, C., “Why Should I Trust You? Explaining the Predictions of Any Classifier,” *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pp. 1135–1144, 2016.
- [11] Chawla, N. V., Bowyer, K. W., Hall, L. O., and Kegelmeyer, W. P., “SMOTE: Synthetic Minority Over-sampling Technique,” *Journal of Artificial Intelligence Research*, vol. 16, pp. 321–357, 2002.
- [12] Tran, V. Q., Nguyen, T. H., and Le, M. T., “Explainable Stacking Ensemble with Feature Tokenizer for Diabetes Prediction,” *Journal of Men's Health*, vol. 20, no. 4, pp. 1–15, 2024.
- [13] Oliullah, K., Rahman, M., and Islam, S., “A Stacked Ensemble Machine Learning Approach for Diabetes Prediction,” *Healthcare Analytics*, vol. 5, pp. 100225, 2023.
- [14] Hasan, R., Ahmed, S., and Rahman, M., “Towards Transparent Diabetes Prediction: Combining AutoML and Explainable AI for Improved Clinical Insights,” *Information*, vol. 16, no. 1, pp. 1–20, 2024.
- [15] Netayawijit, P., et al., “Interpretable Machine Learning Framework for Diabetes Prediction Using Explainable Artificial Intelligence,” *Healthcare Informatics Research*, vol. 31, no. 1, pp. 1–15, 2025.
- [16] Islam, M. M., Hossain, M., and Rahman, S., “Explainable Machine Learning for Efficient Diabetes Prediction,” *Engineering Reports*, vol. 7, no. 2, pp. 1–18, 2025.
- [17] Ahmad, S., Khan, M., and Ali, R., “Early-Stage Diabetes Prediction Using a Stacked Ensemble Model Enhanced with SHAP Explainability,” *Biomedical and Pharmacology Journal*, vol. 19, no. 1, pp. 1–14, 2026.
- [18] Breiman, L., “Random Forests,” *Machine Learning*, vol. 45, no. 1, pp. 5–32, 2001.
- [19] Cortes, C., and Vapnik, V., “Support-Vector Networks,” *Machine Learning*, vol. 20, no. 3, pp. 273–297, 1995.
- [20] Cover, T., and Hart, P., “Nearest Neighbor Pattern Classification,” *IEEE Transactions on Information Theory*, vol. 13, no. 1, pp. 21–27, 1967.

- [21] Hosmer, D. W., Lemeshow, S., and Sturdivant, R. X., *Applied Logistic Regression*, 3rd Edition, Wiley, New York, USA, 2013.
- [22] Quinlan, J. R., *C4.5: Programs for Machine Learning*, Morgan Kaufmann Publishers, San Francisco, USA, 1993.
- [23] Kingma, D. P., and Ba, J., “Adam: A Method for Stochastic Optimization,” *International Conference on Learning Representations (ICLR)*, 2015.
- [24] Chollet, F., *Deep Learning with Python*, Manning Publications, New York, USA, 2021.

