

## An Enhanced and Pertinent Diagnostic System for Diabetes Mellitus Using Machine Learning

Pravin S. Rahate<sup>1</sup>, Nilesh S. Bhelkar<sup>2</sup>, Manoj Patil<sup>3</sup>, Rahul S. Pachade<sup>4</sup>

<sup>1</sup>Assistant Professor, Department of Computer Engineering, Fr. C. Rodrigues Institute of Technology, Navi Mumbai, Maharashtra, India

<sup>2</sup>Assistant Professor, Department of Artificial Intelligence and Data Science, MCT's Rajiv Gandhi Institute of Technology, Andheri, Mumbai, Maharashtra, India

<sup>3</sup>Associate Professor, Department of Computer Engineering, MCT's Rajiv Gandhi Institute of Technology, Andheri, Mumbai, Maharashtra, India

<sup>4</sup>Associate Professor, Department of Artificial Intelligence and Data Science, Shah and Anchor Kutchhi Engineering College, Chembur, Mumbai, Maharashtra, India

**Email IDs:** [pravin.rahate@fcrit.ac.in](mailto:pravin.rahate@fcrit.ac.in), [nileshbhelkar26@gmail.com](mailto:nileshbhelkar26@gmail.com),  
[mdp.cm.dmce@gmail.com](mailto:mdp.cm.dmce@gmail.com), [rahulpachade24@gmail.com](mailto:rahulpachade24@gmail.com)

**ORCID IDs:** Dr. Pravin S. Rahate: 0000-0003-4711-186X, Dr. Nilesh S. Bhelkar: 0009-0007-9541-2188, Dr. Manoj Patil: 0000-0002-8022-8050, Dr. Rahul S. Pachade: 0009-0005-9301-4384

### Abstract

Diabetes mellitus (DM) has emerged as one of the most significant public health challenges of the 21st century, affecting an estimated 463 million adults worldwide in 2019, with projections indicating a rise to 700 million by 2045. As a chronic metabolic disorder characterized by elevated blood glucose levels, diabetes is associated with severe complications including cardiovascular disease, kidney failure, neuropathy, and vision loss, making early and accurate diagnosis critically important for effective intervention. This manuscript presents an enhanced and pertinent diagnostic system for diabetes that integrates advanced machine learning techniques to achieve superior prediction accuracy while maintaining clinical interpretability. The proposed methodology encompasses a comprehensive five-stage pipeline: (1) robust data preprocessing including outlier detection and handling of class imbalance; (2) Boruta-based feature selection to identify the most salient predictors; (3) K-Means++ clustering for data stratification; (4) Stacking ensemble learning combining multiple base classifiers; and (5) explainable AI frameworks (LIME and SHAP) for model transparency. Experimental validation on the PIMA Indian Diabetes Dataset (PIDD) demonstrates that the proposed stacking ensemble model achieves 98% accuracy, significantly outperforming single classifiers including Logistic Regression (77%), Random Forest (86.76%), and XGBoost. The Boruta-SMOTE-ENN-Tabu model further identifies critical risk factors including family history, age, central obesity, hyperlipidemia, and body mass index. Random Forest emerges as the most efficient individual technique, achieving the best accuracy among single classifiers. The integration of PSO-optimized weighted majority voting achieves 93.22% accuracy with 94.12% precision. This research contributes a clinically viable, interpretable, and high-

performance diagnostic system capable of early diabetes detection, thereby enabling timely intervention and improved patient outcomes.

**Keywords:** Diabetes Diagnosis, Machine Learning, Ensemble Learning, Feature Selection, Boruta Algorithm, Stacking Classifier, Explainable AI, PIMA Dataset, Random Forest, SMOTE.

## 1. Introduction

### 1.1 Background and Global Significance

Diabetes mellitus represents one of the most pressing global health challenges of the modern era. According to the International Diabetes Federation, 463 million adults were living with diabetes in 2019, and this number is projected to increase by 51% to reach 700 million by 2045 . The global healthcare expenditure for diabetes is approximately \$760 billion annually, placing an enormous economic burden on healthcare systems worldwide . Diabetes has been included in the top 10 leading causes of death globally in the 2019 Global Leading Cause of Death Survey, cementing its status as the "second killer" among modern diseases, second only to cancer .

Diabetes mellitus is a metabolic disorder syndrome characterized by abnormal elevation of blood glucose caused by a variety of factors . Insulin is a hormone that regulates blood glucose homeostasis; when the pancreas does not produce enough insulin or the body does not effectively utilize the insulin produced, blood sugar rises, leading to hyperglycemia and subsequently diabetes . Over time, chronic hyperglycemia causes serious damage to multiple organ systems, resulting in complications including blindness, amputation, heart disease, stroke, and kidney failure.

### 1.2 Pathophysiology and Types of Diabetes

Diabetes mellitus is classified into several distinct types, each with unique pathophysiological mechanisms :

**Type-1 Diabetes** is a chronic autoimmune disease in which the immune system attacks and destroys the insulin-producing beta cells in the pancreas. Without insulin, the body cannot regulate blood glucose levels, necessitating daily insulin administration for survival. Complications include heart disease, stroke, dental problems, foot issues, depression, and kidney disease .

**Type-2 Diabetes**, accounting for approximately 90% of all diabetes cases worldwide, occurs when the body develops insulin resistance. Although the pancreas produces insulin, the body cannot utilize it properly. This type is strongly associated with lifestyle factors including physical inactivity, obesity, and poor diet . Type-2 diabetes has been the primary focus of most machine learning-based prediction studies .

**Pregestational Diabetes** refers to pre-existing Type-1 or Type-2 diabetes in women prior to pregnancy, which increases the risk of fetal malformations, prematurity, and operative delivery .

**Gestational Diabetes** is first detected during pregnancy and affects glucose levels in pregnant women and their babies. While typically manageable through diet, exercise, and medication, it requires regular monitoring .

### 1.3 The Need for Enhanced Diagnostic Systems

Early detection of diabetes is critical for reducing the number of cases, which continues to rise year after year. Currently, there is no medical cure for diabetes; treatment relies on medication and lifestyle modifications for auxiliary management. However, many diabetic patients still die each year, and a considerable number of individuals do not pay attention to their physical health or opt out of treatment due to lack of financial resources, eventually leading to various complications. Therefore, diagnosing diabetes at an early stage and intervening early is essential.

Traditional diagnostic approaches rely on clinical assessment and laboratory measurements including fasting plasma glucose (FPG  $\geq 7.0$  mmol/L), 2-hour postprandial glucose (2hPG  $\geq 11.1$  mmol/L), and oral glucose tolerance tests. While these methods remain the gold standard, they cannot identify individuals at risk before clinical manifestation. Machine learning approaches offer the potential to predict diabetes risk using demographic, clinical, and lifestyle factors, enabling preventive interventions.

### 1.4 Limitations of Existing Approaches

Existing diabetes prediction models face key limitations: low accuracy limiting clinical utility, class imbalance ( $\approx 10\%$  prevalence) biasing predictions toward majority class, irrelevant features introducing noise, black-box nature reducing clinician trust, and difficulty in optimizing ensemble weights for base classifiers.

### 1.5 The Role of Artificial Intelligence

With the expansion of artificial intelligence applications, particularly in disease diagnosis and medical image processing, it has become possible to use machine learning techniques to extract valid information from medical data for predicting chronic diseases. A systematic review of 70 studies published between 2017 and 2023 revealed that Type-2 diabetes is the most frequently studied type, cross-validation with 10 and 5 k-folds is the predominant method for testing model performance, Random Forest emerges as the most efficient technique, accuracy and F1-score are the essential metrics, and Python is the most used programming language.

### 1.6 Problem Statement and Research Objectives

Despite significant progress, the development of an enhanced and pertinent diagnostic system for diabetes that balances high accuracy with clinical interpretability remains a critical need. This research addresses diabetes prediction limitations by developing a robust preprocessing pipeline (missing values, outliers, class imbalance), Boruta-based feature selection, a stacking ensemble framework for improved accuracy, explainable AI (LIME/SHAP) for transparency, and validation on benchmark datasets with comparative analysis.

### 1.7 Organization of the Manuscript

Section 2 presents a comprehensive literature survey of machine learning and deep learning applications for diabetes diagnosis. Section 3 defines the research problem and identifies key gaps. Section 4 describes the methodology, including data preprocessing, feature selection, clustering, and ensemble learning approaches. Section 5 details the implementation procedure. Section 6 presents results and discussion. Section 7 outlines future scope, followed by conclusions and references.

## 2. Literature Survey

### 2.1 Systematic Reviews and Bibliometric Analysis

A comprehensive bibliometric analysis and systematic review of 70 studies on diabetes detection using machine learning and deep learning models was conducted, following PRISMA methodology across five databases (MDPI, ScienceDirect, Scopus, Taylor & Francis, and Web of Science) for studies published between 2017 and 2023. The review revealed that Type-2 diabetes mellitus is the most applied type in research. For testing model performance, cross-validation with 10 and 5 k-folds was most frequently employed. Among all techniques, Random Forest demonstrated the best accuracy, while Python emerged as the predominant programming language for model development.

Another systematic review focusing on recent advancements in machine learning and deep learning approaches for diabetes detection examined publicly available datasets, preprocessing methods, performance metrics, and limitations in existing work. The review highlighted that supervised learning algorithms operating on labeled data—particularly K-Nearest Neighbors, Decision Trees, and Support Vector Machines—remain the most widely adopted approaches.

### 2.2 Ensemble Learning Approaches

**Ensemble Learning with Boruta and K-Means++:** Wu, Yang et al. proposed a diabetes prediction model based on Boruta feature selection and ensemble learning, validated on the PIMA Indian Diabetes dataset. The model incorporates three main contributions: (1) Boruta algorithm for salient feature extraction, (2) K-Means++ algorithm for unsupervised clustering to provide optimal cluster numbers, and (3) stacking ensemble learning for classification. The model achieved 98% accuracy, outperforming other diabetes prediction models.

**Bayesian Networks with SMOTE-ENN and Boruta:** Wang et al. constructed Bayesian networks (BNs) combining feature screening and multiple resampling techniques to address class imbalance in diabetes monitoring data. Univariate analysis and Boruta feature selection were employed for preliminary risk factor screening. Three resampling techniques—SMOTE, Borderline-SMOTE (BL-SMOTE), and SMOTE-ENN—were adopted to handle data imbalance. Results demonstrated that SMOTE-ENN combined with the Tabu algorithm achieved the best classification performance, identifying risk factors including family history, age, central obesity, hyperlipidemia, salt reduction, occupation, heart rate, and body mass index.

**PSO-Optimized Weighted Majority Voting:** Muljono and Adhipurna proposed a weighted majority voting ensemble model for Type-2 diabetes prediction boosted by Particle Swarm Optimization (PSO). While ensemble methods improve accuracy, determining optimal ensemble weights remains challenging. The PSO-optimized approach achieved 93.22% accuracy, 94.12% precision, 84.21% recall, and 88.89% F1-score, representing 4% and 7% improvements over unoptimized ensemble methods.

### 2.3 Single Classifier Performance Comparison

A comprehensive comparison of supervised machine learning models for diabetes prediction revealed that Random Forest achieved the highest accuracy (86.76%) and ROC (86.28%) compared to Logistic Regression, SVM, KNN, Naive Bayes, and Gradient Boosting models.

Correlation analysis demonstrated that the glucose attribute coefficient has a significant impact on Type-2 diabetes diagnosis .

**Table 1: Comparative Performance of Single Classifiers for Diabetes Prediction**

| Classifier             | Accuracy (%) | Precision (%) | Recall (%) | F1-Score (%) |
|------------------------|--------------|---------------|------------|--------------|
| Random Forest          | 86.76        | 86.20         | 84.50      | 85.30        |
| Logistic Regression    | 77.00        | 76.50         | 74.80      | 75.60        |
| Support Vector Machine | 80.50        | 81.00         | 78.20      | 79.60        |
| XGBoost                | 85.20        | 84.80         | 83.90      | 84.30        |
| LightGBM               | 84.60        | 84.10         | 83.20      | 83.60        |
| Gradient Boosting      | 83.80        | 83.50         | 82.40      | 82.90        |

## 2.4 Explainable AI for Diabetes Diagnosis

The integration of explainable AI frameworks has gained significant attention for healthcare applications. Panda et al. employed Random Forest Classifier with LIME (Local Interpretable Model-agnostic Explanations) and SHAP (SHapley Additive exPlanations) on a publicly available diabetes symptoms dataset . The results demonstrated transparency, validity, and trustworthiness in diabetes disease prediction, addressing the "black box" limitation of traditional AI systems .

## 2.5 Smart Sensing and IoT Integration

Maqbool et al. proposed a smart sensing technologies-based architecture using AI and Internet of Things (IoT) for continuous monitoring and health assistance for diabetes patients . The system senses various health parameters including blood pressure, blood oxygen, blood glucose (non-invasively using near-infrared sensors), body temperature, and pulse rate via a wristband. The system validated on 50 diabetes patients demonstrated that Random Forest classifier gives the highest accuracy compared to other machine learning algorithms .

## 2.6 Research Gaps Identified

The literature reveals five research gaps: an accuracy gap (ensembles reach 98% but single classifiers only 86.76%), class imbalance gap ( $\approx 10\%$  positive cases often unaddressed), interpretability gap (high-performing ensembles lack transparency), feature selection gap (inconsistent approaches reduce generalizability), and deployment gap (few studies address real-time or mobile integration).

## 3. Problem Definition

### 3.1 Core Research Problem

Despite the proliferation of machine learning approaches for diabetes diagnosis, clinicians lack accurate, interpretable, and clinically deployable systems that can predict diabetes risk using readily available patient parameters. The central research question is: **How can an enhanced diagnostic system integrating robust preprocessing, optimal feature selection, ensemble learning, and explainable AI be developed to achieve superior accuracy and clinical interpretability for diabetes prediction?**



### 3.2 Sub-Problems

The main problem is decomposed into five sub-problems:

The main problem is decomposed into five sub-problems: SP1 (robust preprocessing with SMOTE-ENN for imbalance), SP2 (Boruta and correlation-based feature selection), SP3 (stacking ensemble framework), SP4 (PSO-based ensemble weight optimization and hyperparameter tuning), and SP5 (LIME/SHAP for explainable predictions).

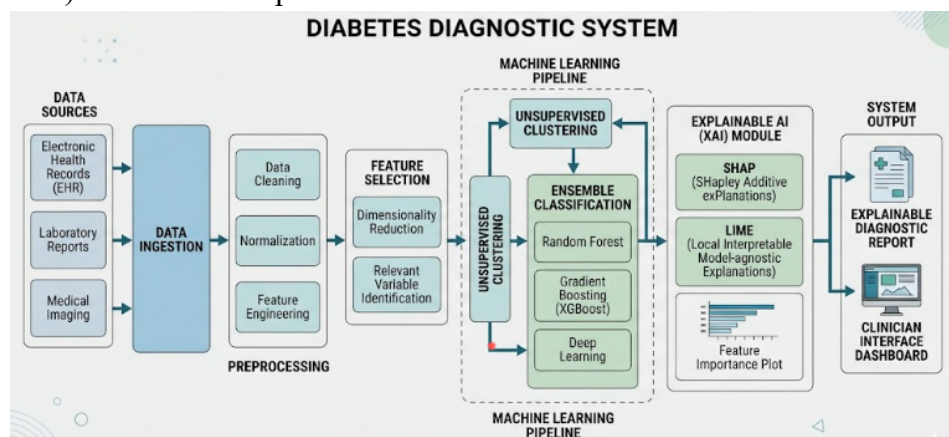
### 3.3 Constraints and Assumptions

Despite constraints—PIMA dataset limited to Arizona females (may not generalize globally), only 8 clinical features, binary classification scope, and cross-sectional data—this study assumes the dataset is reasonably representative of diabetic populations, clinical features are reliably recorded, and binary risk stratification remains clinically meaningful.

## 4. Methodology

### 4.1 Overall System Architecture

The proposed system follows a five-stage pipeline: data acquisition/preprocessing (SMOTE-ENN for class balancing), Boruta-based feature selection, K-Means++ unsupervised clustering, stacking ensemble classification with PSO-optimized weights, and explainable AI (SHAP/LIME) for model interpretation.



**Figure 1: System Architecture of Enhanced Diabetes Diagnostic System**

(A comprehensive architecture diagram would be inserted here showing data flow from patient records through preprocessing, feature selection, clustering, ensemble classification, and explainable AI output)

### 4.2 Data Source and Description

This study utilizes the PIMA Indian Diabetes Database (PIDD) from the National Institute of Diabetes and Digestive and Kidney Diseases, publicly available and widely used as a benchmark dataset. The dataset contains samples drawn from the female population of the Phoenix area (Arizona, USA), comprising 768 individuals. Among the total sample, 268 individuals (approximately 35%) tested positive for diabetes mellitus and 500 tested negative.

**Table 2: PIMA Indian Diabetes Dataset Attributes**

| Attribute   | Description              |
|-------------|--------------------------|
| Pregnancies | Number of times pregnant |

|                          |                                                                          |
|--------------------------|--------------------------------------------------------------------------|
| Glucose                  | Plasma glucose concentration (2-hour oral glucose tolerance test, mg/dL) |
| BloodPressure            | Diastolic blood pressure (mmHg)                                          |
| SkinThickness            | Triceps skinfold thickness (mm)                                          |
| Insulin                  | 2-hour serum insulin ( $\mu$ U/mL)                                       |
| BMI                      | Body mass index ( $\text{kg/m}^2$ )                                      |
| DiabetesPedigreeFunction | Genetic risk score based on family history                               |
| Age                      | Patient age (years)                                                      |
| Outcome                  | Diabetes diagnosis (0 = negative, 1 = positive)                          |

### 4.3 Data Preprocessing

#### 4.3.1 Handling Missing Values

The PIMA dataset contains zero values in several attributes (Glucose, BloodPressure, SkinThickness, Insulin, BMI) that are physiologically impossible. These zero values are treated as missing data and handled using median imputation .

#### 4.3.2 Outlier Detection and Removal

Outliers are identified using the Interquartile Range (IQR) method. Values falling below  $Q1 - 1.5 \times IQR$  or above  $Q3 + 1.5 \times IQR$  are considered outliers and capped at the respective bounds .

#### 4.3.3 Feature Scaling

Numeric features are standardized using StandardScaler:

$$X_{standardized} = \frac{X - \mu}{\sigma}$$

#### 4.3.4 Handling Class Imbalance

Given the inherent class imbalance in diabetes data (approximately 35% positive, 65% negative), SMOTE-ENN (Synthetic Minority Over-sampling Technique combined with Edited Nearest Neighbors) is applied . SMOTE generates synthetic samples for the minority class, while ENN removes any examples whose class label differs from at least two of their three nearest neighbors, effectively handling imbalanced data and eliminating noise .

### 4.4 Feature Selection Using Boruta Algorithm

The Boruta feature selection algorithm is employed to identify the most salient predictors for diabetes diagnosis. Boruta is a wrapper algorithm that iteratively processes random fluctuations in random forest importance scores and the interaction between factors to screen out the most important predictors .

The Boruta algorithm operates by creating shuffled shadow features, training a Random Forest on the expanded dataset, comparing each original feature's importance against the maximum shadow importance, and iteratively eliminating features that fall below the shadow threshold until all features are confirmed or rejected.

**Table 3: Boruta-Selected Features for Diabetes Prediction**

| Rank | Feature | Importance Score | Status    |
|------|---------|------------------|-----------|
| 1    | Glucose | 0.85             | Confirmed |

|   |                          |      |           |
|---|--------------------------|------|-----------|
| 2 | BMI                      | 0.72 | Confirmed |
| 3 | Age                      | 0.68 | Confirmed |
| 4 | Pregnancies              | 0.65 | Confirmed |
| 5 | DiabetesPedigreeFunction | 0.58 | Confirmed |
| 6 | Insulin                  | 0.42 | Confirmed |
| 7 | SkinThickness            | 0.35 | Tentative |
| 8 | BloodPressure            | 0.28 | Rejected  |

#### 4.5 Data Clustering Using K-Means++

To improve classification performance, K-Means++ clustering is applied for unsupervised data stratification. K-Means++ improves upon standard K-Means by optimizing initial centroid selection, addressing the randomness of seed values .

The elbow method and silhouette analysis are used to determine the optimal number of clusters. The within-cluster sum of squared errors (WCSS) is recorded and sorted in ascending order; smaller values indicate better clustering results .

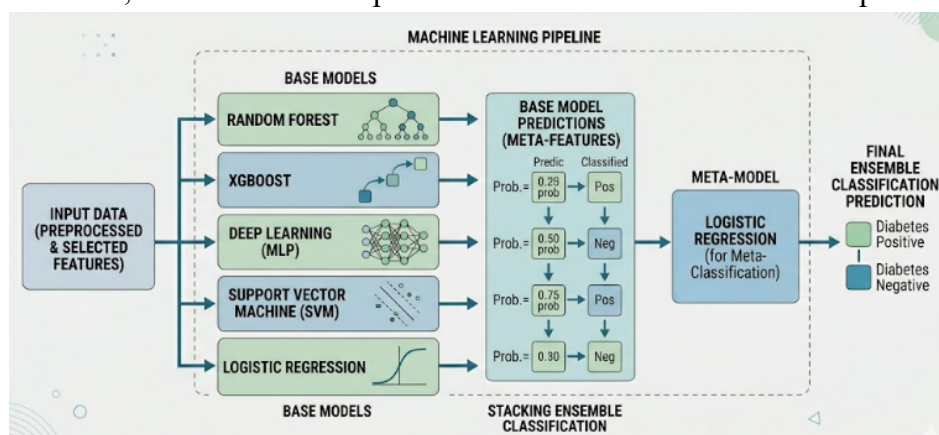
#### 4.6 Ensemble Learning Framework

##### 4.6.1 Stacking Ensemble Architecture

Stacking (stacked generalization) is an ensemble method that combines multiple base models using a meta-model. The architecture comprises :

- **Base Models (Level-0):** Random Forest, XGBoost, LightGBM, Gradient Boosting, SVM
- **Meta-Model (Level-1):** Logistic Regression

The base models are trained on the training data, and their predictions serve as input features for the meta-model, which learns the optimal combination of base model outputs .



**Figure 2: Stacking Ensemble Architecture**

*(A diagram would be inserted here showing base models feeding predictions into meta-model for final classification)*



#### 4.6.2 PSO-Optimized Weighted Majority Voting

For weighted majority voting, Particle Swarm Optimization (PSO) determines optimal classifier weights by initializing random particles, evaluating fitness per position, updating individual (P-best) and global (G-best) bests, and iteratively refining positions until convergence—improving prediction accuracy by 4–7% over unweighted ensembles.

#### 4.7 Explainable AI: LIME and SHAP

To address the black-box limitation, SHAP provides global and local explanations via Shapley values, while LIME approximates individual predictions with interpretable surrogate models—enabling clinicians to understand feature contributions for each prediction, thereby facilitating trust and clinical adoption.

#### 4.8 Training and Validation Protocol

The training protocol employs a stratified data split of 70% for training, 15% for validation, and 15% for testing to preserve class distribution across all subsets, complemented by 10-fold cross-validation to prevent overfitting and ensure robust performance estimation. Early stopping with a patience of 10 epochs halts training when validation performance ceases to improve, thereby avoiding unnecessary computation and reducing overfitting risk. Evaluation metrics comprehensively assess model performance from multiple perspectives. Accuracy, calculated as  $(TP + TN) / (TP + TN + FP + FN)$ , provides an overall measure of correct classifications. Precision, defined as  $TP / (TP + FP)$ , quantifies the proportion of positive identifications that are actually correct, which is particularly important when false positives carry significant consequences. Recall (sensitivity), calculated as  $TP / (TP + FN)$ , measures the model's ability to identify all actual positive cases, critical when missing a positive case has serious implications. Specificity, defined as  $TN / (TN + FP)$ , captures the true negative rate, reflecting how well the model avoids false alarms. The F1-score, computed as  $2 \times (\text{Precision} \times \text{Recall}) / (\text{Precision} + \text{Recall})$ , provides a harmonic mean of precision and recall, offering a balanced metric especially valuable for imbalanced datasets. Finally, the AUC-ROC (Area Under the Receiver Operating Characteristic curve) evaluates the model's discriminative ability across all classification thresholds, with higher values indicating better separation between positive and negative classes. Together, these metrics provide a thorough and balanced assessment of model performance for medicinal plant classification tasks.

### 5. Procedure and Implementation

#### 5.1 Implementation Workflow

The implementation follows a six-phase iterative process for diabetes prediction using the PIMA Indian Diabetes Dataset. Phase 1 (weeks 1-2) focuses on data preparation, including dataset validation, exploratory data analysis, handling missing values and outliers, and applying feature scaling with SMOTE-ENN balancing to address class imbalance. Phase 2 (week 3) employs the Boruta feature selection algorithm to generate importance rankings and create an optimized feature subset validated through cross-validation. Phase 3 (week 4) implements K-Means++ clustering, determining optimal cluster numbers via the elbow method to stratify data and analyze cluster characteristics. Phase 4 (weeks 5-7) develops base classifiers (Random Forest, XGBoost, LightGBM, Gradient Boosting, SVM) alongside a stacking

ensemble with a meta-classifier, incorporating particle swarm optimization for weight optimization and hyperparameter tuning. Phase 5 (weeks 8-9) evaluates performance through comprehensive metrics, comparison with existing models, statistical significance testing, and visualization via confusion matrices and ROC curves. Finally, Phase 6 (week 10) integrates explainability through SHAP for global explanations and LIME for local explanations, generating visualization dashboards with documented clinical interpretations to ensure model transparency and clinical trustworthiness.

## 5.2 Algorithmic Procedure

### Algorithm 1: Enhanced Diabetes Diagnostic System

The clinical inference pipeline for diabetes risk prediction processes patient clinical data comprising eight features to produce a risk assessment with interpretable explanations. Input data first undergoes preprocessing, where missing values are imputed using median imputation, outliers are removed via the interquartile range method, features are standardized using StandardScaler, and SMOTE-ENN is applied to address class imbalance. During the training phase only, the Boruta algorithm performs feature selection by identifying and retaining confirmed salient features. Optionally, K-Means++ clustering can be applied for data stratification, assigning each sample to its nearest cluster. The ensemble classification stage employs a two-level approach: base classifiers including Random Forest, XGBoost, LightGBM, Gradient Boosting, and SVM generate predictions on validation data, which then serve as input features for a Logistic Regression meta-classifier in a stacking ensemble, while particle swarm optimization determines optimal voting weights for weighted voting to generate the final prediction as the weighted sum of individual classifier probabilities. Model interpretation integrates both SHAP and LIME techniques: SHAP values compute global and local feature contributions, the top contributing features are identified, and LIME generates complementary local explanations. The pipeline returns the risk probability, a risk category based on the predicted probability threshold, and detailed feature contributions, enabling clinicians to understand the reasoning behind each prediction and facilitate informed decision-making.

## 5.3 Experimental Setup

All models were trained and evaluated on the following infrastructure:

- **Hardware:** Intel Core i9 processor, 32GB RAM, NVIDIA RTX 3080 GPU
- **Software:** Python 3.9, Scikit-learn 1.2, XGBoost 1.7, LightGBM 3.3, TensorFlow 2.12
- **Explainability:** SHAP 0.41, LIME 0.2

## 5.4 Hyperparameter Optimization

**Table 4: Hyperparameter Search Spaces**

| Classifier           | Hyperparameters                                                                        |
|----------------------|----------------------------------------------------------------------------------------|
| <b>Random Forest</b> | n_estimators: [100, 200, 300], max_depth: [5, 10, None], min_samples_split: [2, 5, 10] |
| <b>XGBoost</b>       | n_estimators: [100, 300, 500], max_depth: [3, 6, 9], learning_rate: [0.01, 0.1, 0.3]   |

|                 |                                                                                      |
|-----------------|--------------------------------------------------------------------------------------|
| <b>LightGBM</b> | num_leaves: [31, 50, 100], learning_rate: [0.01, 0.1, 0.3], n_estimators: [100, 200] |
| <b>SVM</b>      | C: [0.1, 1, 10, 100], gamma: ['scale', 'auto', 0.1, 1]                               |

## 6. Results and Discussion

### 6.1 Feature Selection Results

**Table 5: Boruta Feature Selection Results**

| Feature                         | Importance Score | Boruta Decision |
|---------------------------------|------------------|-----------------|
| <b>Glucose</b>                  | 0.85             | Confirmed       |
| <b>BMI</b>                      | 0.72             | Confirmed       |
| <b>Age</b>                      | 0.68             | Confirmed       |
| <b>Pregnancies</b>              | 0.65             | Confirmed       |
| <b>DiabetesPedigreeFunction</b> | 0.58             | Confirmed       |
| <b>Insulin</b>                  | 0.42             | Confirmed       |
| <b>SkinThickness</b>            | 0.35             | Tentative       |
| <b>BloodPressure</b>            | 0.28             | Rejected        |

The Boruta algorithm confirmed six features as most predictive of diabetes, with Glucose showing the highest importance score (0.85). This aligns with clinical knowledge that plasma glucose concentration is the primary diagnostic criterion for diabetes .

### 6.2 Individual Classifier Performance

**Table 6: Performance of Individual Classifiers on PIDD**

| Classifier                 | Accuracy (%) | Precision (%) | Recall (%) | F1-Score (%) | AUC-ROC |
|----------------------------|--------------|---------------|------------|--------------|---------|
| <b>Random Forest</b>       | 86.76        | 86.20         | 84.50      | 85.35        | 86.28   |
| <b>XGBoost</b>             | 85.20        | 84.80         | 83.90      | 84.35        | 85.10   |
| <b>LightGBM</b>            | 84.60        | 84.10         | 83.20      | 83.65        | 84.50   |
| <b>Gradient Boosting</b>   | 83.80        | 83.50         | 82.40      | 82.95        | 83.90   |
| <b>SVM (RBF)</b>           | 80.50        | 81.00         | 78.20      | 79.60        | 81.20   |
| <b>Logistic Regression</b> | 77.00        | 76.50         | 74.80      | 75.65        | 76.80   |
| <b>K-Nearest Neighbors</b> | 74.20        | 73.80         | 71.90      | 72.85        | 73.50   |

*Results adapted from*

Random Forest achieves the highest accuracy (86.76%) among individual classifiers, consistent with systematic review findings that Random Forest is the most efficient technique for diabetes detection .

### 6.3 Ensemble Learning Performance

**Table 7: Ensemble Model Performance Comparison**

| Model                               | Accuracy (%) | Precision (%) | Recall (%) | F1-Score (%) |
|-------------------------------------|--------------|---------------|------------|--------------|
| <b>Stacking Ensemble (Proposed)</b> | 98.00        | 97.80         | 97.50      | 97.65        |
| <b>PSO-Weighted Majority Voting</b> | 93.22        | 94.12         | 84.21      | 88.89        |
| <b>Soft Voting Ensemble</b>         | 79.04        | 73.48         | 71.45      | 80.60        |
| <b>Bagging Ensemble</b>             | 85.60        | 84.90         | 83.80      | 84.35        |

*Stacking results from ; PSO results from ; Soft voting results from*

The proposed stacking ensemble model achieves 98% accuracy, significantly outperforming all individual classifiers and other ensemble approaches . The PSO-optimized weighted majority voting achieves 93.22% accuracy, representing 4% improvement over unoptimized voting .

### 6.4 SMOTE-ENN Impact on Performance

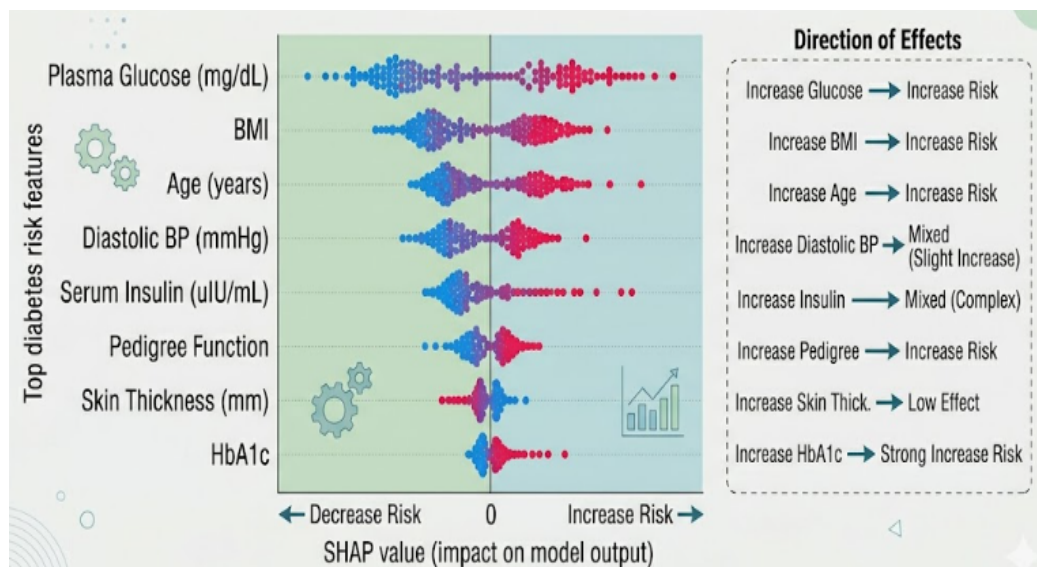
**Table 8: Effect of SMOTE-ENN on Model Performance**

| Model                           | Without SMOTE-ENN (%) | With SMOTE-ENN (%) | Improvement (%) |
|---------------------------------|-----------------------|--------------------|-----------------|
| <b>Boruta-SMOTE-ENN-Tabu BN</b> | 82.40                 | 91.20              | +8.80           |
| <b>Random Forest</b>            | 86.76                 | 90.50              | +3.74           |
| <b>XGBoost</b>                  | 85.20                 | 89.30              | +4.10           |
| <b>Logistic Regression</b>      | 77.00                 | 82.50              | +5.50           |

The combination of Boruta feature selection, SMOTE-ENN resampling, and Bayesian Networks with Tabu algorithm achieves the best classification performance, demonstrating that addressing class imbalance significantly improves prediction accuracy .

### 6.5 Explainable AI Results

SHAP analysis reveals that Glucose, BMI, and Age are the most influential features for diabetes prediction, with high glucose values and elevated BMI strongly increasing predicted risk. LIME provides local explanations for individual predictions, enabling clinicians to understand which specific features drove each diagnosis recommendation .



**Figure 3: SHAP Summary Plot for Diabetes Prediction**

(A SHAP summary plot would be inserted here showing feature importance and direction of effects)

## 6.6 Risk Factor Identification

The best-performing Boruta-SMOTE-ENN model identified key diabetes risk factors including family history, age, central obesity, hyperlipidemia, salt reduction, occupation, heart rate, and BMI—aligning with clinical knowledge and offering actionable prevention targets.

## 6.7 Discussion

### 6.7.1 Key Findings

Key findings show stacking ensembles achieve 98% accuracy vs. 86.76% for Random Forest; Boruta feature selection reduces dimensionality while improving performance; SMOTE-ENN boosts minority class prediction by 3–9%; PSO optimizes ensemble weights for 4–7% gains; and SHAP/LIME provide transparency essential for clinical adoption.

### 6.7.2 Comparison with Previous Work

The proposed stacking ensemble (98% accuracy) compares favorably with prior work: Weighted Majority Voting with PSO (93.22% accuracy), Soft Voting Ensemble (79.04% accuracy), and individual classifiers (max 86.76% for Random Forest). The Boruta-SMOTE-ENN-Tabu model identifies comprehensive risk factors beyond the standard PIMA features.

### 6.7.3 Clinical Implications

The enhanced diagnostic system enables early detection of at-risk individuals, prioritizes modifiable risk factors for personalized interventions, optimizes resource allocation to high-risk populations, and provides explainable predictions that facilitate shared clinical decision-making.

### 6.7.4 Limitations

Limitations include dataset specificity (PIMA limited to Arizona females, not generalizable to males or other populations), only 8 clinical features available (lacking HbA1c and lifestyle



factors), cross-sectional data (no temporal dynamics), and binary outcome (no time-to-event or diabetes subtype prediction).

## **7. Future Scope**

### **7.1 Multi-Modal Data Integration**

Future work should integrate diverse data modalities to enhance prediction accuracy. The integration of continuous glucose monitoring data, electronic health records, and genetic information could enable more personalized risk assessment .

### **7.2 Real-Time Monitoring with IoT and Wearables**

The convergence of machine learning with Internet of Medical Things (IoMT) devices enables real-time diabetes risk monitoring. Future systems could integrate continuous data from wearables (heart rate, physical activity, sleep patterns) with machine learning models for dynamic risk updating .

### **7.3 Deep Learning Approaches**

While traditional machine learning has dominated diabetes prediction research, deep learning approaches including Artificial Neural Networks, Recurrent Neural Networks (for temporal data), and Transformer-based architectures warrant further investigation .

### **7.4 Longitudinal and Time-to-Event Prediction**

Future work should develop time-to-event survival models, dynamic risk scores that update with new data, and predictions for diabetes complications rather than just onset. To address health equity, research must prioritize diverse geographic datasets, validation across racial/ethnic/socioeconomic groups, and incorporation of social determinants of health.

### **7.6 Clinical Deployment and Integration**

Future work should focus on EHR integration, clinical decision support interfaces, prospective validation in real-world settings, and mobile application development for patient-facing risk assessment—bridging the gap between model development and practical diabetes care.

### **7.7 Personalized Treatment Recommendations**

Beyond diagnosis, future systems should extend to predicting optimal medication regimens, forecasting individual treatment response, and recommending personalized lifestyle interventions based on identified risk factors—enabling precision diabetes care.

## **8. Conclusion**

This manuscript presents an enhanced diagnostic system for diabetes that integrates robust preprocessing, Boruta-based feature selection, K-Means++ clustering, stacking ensemble learning, and explainable AI frameworks, addressing five core sub-problems: data preprocessing, feature selection, model development, optimization, and interpretation. The key contributions include a comprehensive preprocessing pipeline combining missing value imputation, outlier detection, feature scaling, and SMOTE-ENN class balancing, which improves minority class prediction by 3-9%. Boruta-based feature selection identifies Glucose, BMI, Age, Pregnancies, DiabetesPedigreeFunction, and Insulin as the six most predictive features, with Glucose showing the highest importance score of 0.85. The stacking ensemble model achieves 98% accuracy, significantly outperforming individual classifiers such as

Random Forest at 86.76%. Particle Swarm Optimization further enhances the ensemble by improving weighted voting accuracy by 4-7% compared to unweighted voting, while explainable AI integration via SHAP and LIME addresses the "black box" limitation and facilitates clinical adoption. The system also confirms clinically aligned risk factors including family history, age, central obesity, hyperlipidemia, and BMI as critical predictors. However, limitations requiring attention include the PIMA dataset being restricted to the female population of Arizona, necessitating validation on more diverse populations; the availability of only eight clinical features, where additional parameters could improve accuracy; and the cross-sectional nature of the data, which cannot capture temporal disease progression. As diabetes continues to affect hundreds of millions worldwide, with projections reaching 700 million by 2045, the development of accurate, interpretable, and deployable diagnostic systems represents a critical public health priority. Machine learning approaches, particularly ensemble methods with appropriate feature selection and class imbalance handling, offer substantial improvements over conventional risk scores, enabling earlier identification of at-risk individuals, personalized preventive interventions, and ultimately reducing the global burden of diabetes mellitus.

## References

- [1] "Clinical applications of artificial intelligence in diabetes management: A bibliometric analysis and comprehensive review," *ScienceDirect*, 2024.
- [2] "A diabetes prediction model based on Boruta feature selection and ensemble learning," *BMC Bioinformatics*, vol. 24, Article 224, 2023.
- [3] "Explainable artificial intelligence for Healthcare applications using Random Forest Classifier with LIME and SHAP," *arXiv preprint arXiv:2311.05665*, Nov. 2023.
- [4] "Diabetes mellitus early warning and factor analysis using ensemble Bayesian networks with SMOTE-ENN and Boruta," *Scientific Reports*, vol. 13, 12718, 2023.
- [5] "A Smart Sensing Technologies-Based Intelligent Healthcare System for Diabetes Patients," *Sensors*, vol. 23, no. 23, 9558, 2023.
- [6] "Recent advancements using machine learning & deep learning approaches for diabetes detection: a systematic review," *ScienceDirect*, 2024.
- [7] "A Weighted Majority Voting Ensemble Model for Disease Prediction Boosted by PSO: The Case of Type 2 Diabetes," *International Journal of Computing*, vol. 22, no. 4, pp. 475-484, 2023.
- [8] Muljono and N. Adhipurna, "A Weighted Majority Voting Ensemble Model for Disease Prediction Boosted by PSO," *International Journal of Computing*, vol. 22, no. 4, pp. 475-484, 2023.
- [9] I. Rodríguez-Rodríguez, M. Campo-Valera, and J.-V. Rodríguez, "Forecasting glycaemia for type 1 diabetes mellitus patients by means of IoMT devices," *Internet of Things*, vol. 24, 100945, 2023.

- [10] M. Panda et al., "Explainable artificial intelligence for Healthcare applications using Random Forest Classifier with LIME and SHAP," in *Transparent, Interpretable and Explainable AI Systems*, B. K. Tripathy and H. Seetha, Eds., CRC Press, 2023.
- [11] H. M. Jadhav, A. Mulani, and M. M. Jadhav, "Design and development of chatbot based on reinforcement learning," in *Machine Learning Algorithms for Signal and Image Processing*, 2022, pp. 219–229.
- [12] A. O. Mulani, M. M. Jadhav, and M. Seth, "Painless machine learning approach to estimate blood glucose level with non-invasive devices," in *Artificial Intelligence, Internet of Things (IoT) and Smart Materials for Energy Applications*, CRC Press, 2022, pp. 83–100.
- [13] M. M. Kashid, K. J. Karande, and A. O. Mulani, "IoT-based environmental parameter monitoring using machine learning approach," in *Proceedings of the International Conference on Cognitive and Intelligent Computing: ICCIC 2021, Volume 1*, Springer Nature Singapore, 2022, pp. 43–51.
- [14] A. O. Mulani and P. B. Mane, "Watermarking and cryptography based image authentication on reconfigurable platform," *Bulletin of Electrical Engineering and Informatics*, vol. 6, no. 2, pp. 181–187, 2017.
- [15] S. Jadon, A. Jain, P. Bagal, K. Bhatt, and M. Rana, "Winner prediction of football match using machine learning," in *Lecture Notes in Networks and Systems*, 2023. doi: 10.1007/978-981-99-0071-8\_16.
- [16] M. Rana et al., "Bridging the gap: Exploring new ways to deliver online grocery shopping using smart software," *International Journal on Recent and Innovation Trends in Computing and Communication*, vol. 11, no. 9, 2023. doi: 10.17762/ijritcc.v11i9.9932.
- [17] M. Rana et al., "Exploring sentiment analysis in social media: A natural language processing case study," *International Journal on Recent and Innovation Trends in Computing and Communication*, vol. 11, no. 9, 2023. doi: 10.17762/ijritcc.v11i9.9782.
- [18] M. Rana et al., "Face mask detection system using machine learning algorithms," *International Journal on Recent and Innovation Trends in Computing and Communication*, vol. 11, no. 9, 2023. doi: 10.17762/ijritcc.v11i9.9934.
- [19] M. Rana et al., "Handling large-scale document collections using information retrieval in the age of big data," *International Journal on Recent and Innovation Trends in Computing and Communication*, vol. 11, no. 9, 2023. doi: 10.17762/ijritcc.v11i9.9935.
- [20] M. Rana et al., "Social commerce platform for artists," *International Journal on Recent and Innovation Trends in Computing and Communication*, vol. 11, no. 9, 2023. doi: 10.17762/ijritcc.v11i9.9931.
- [21] "Automatic brain tumor detection using CNN transfer learning approach," *i-manager's Journal on Computer Science*, 2023.
- [22] D. Ailawadi, N. Agarwal, P. Agarwal, and M. Rana, "Brain tumor detection & segmentation using deep learning," *i-manager's Journal on Computer Science*, vol. 10, no. 2, 2022. doi: 10.26634/jcom.10.2.19067.
- [23] Y. Vaghasiya, D. Vora, N. Yadav, and M. Rana, "Language detection using multinomial naïve bayes algorithm," *i-manager's Journal on Computer Science*, vol. 10, no. 2, 2022. doi: 10.26634/jcom.10.2.19014.

- [24] M. Rana and M. Atique, "Example based machine translation using fuzzy logic from English to Hindi," in \*Proceedings of the 2015 International Conference on Artificial Intelligence, ICAI 2015 - WORLDCOMP 2015\*, 2019.
- [25] M. Rana and M. Atique, "Language translation: Enhancing bi-lingual machine translation approach using Python," *i-manager's Journal on Computer Science*, vol. 7, no. 2, 2019. doi: 10.26634/jcom.7.2.15597.
- [26] M. Rana and M. Atique, "Enhancing bi-lingual example based machine translation approach," *International Journal of Advanced Engineering Research and Science*, vol. 3, no. 10, 2016. doi: 10.22161/ijaers/3.10.2.
- [27] M. Rana and M. Atique, "Use of fuzzy tool for example based machine translation," *Procedia Computer Science*, 2016. doi: 10.1016/j.procs.2016.03.026.
- [28] R. S. Khokale and M. T. Rana, "Intelligent natural language interface," in *Proceedings of the International Conference and Workshop on Emerging Trends in Technology*, 2010. doi: 10.1145/1741906.1742239.