

Evaluating an Ensemble Machine Learning Based Diabetic Medical Illness Analytics Platform for E-Healthcare

Dodla S. Rao¹, Mudra.Narasimharao², Bashir N. Jamadar³, Anil A. Powar⁴, B. B. Vhankhande⁵

^{1, 3, 5} Department of Humanities and Sciences, Walchand Collage of Engineering Sangli, Vishrambag, Sangli, Maharashtra, India

² ITER, S'O'A University, Department of ECE, ITER, Bhubaneswar, Odisha, India

⁴ Department of Humanities and Sciences, Walchand Collage of Engineering Sangli, Vishrambag, Sangli, Maharashtra, India

Ph. D, Research Scholar, JJTU, Jhunjhunu, Rajasthan, India

Email Id: ¹ sivanageswararao.dodla@walchandsangli.ac.in, ² narasimharaomudra@gmail.com

Abstract

Diabetes Mellitus (DM) is a glucose metabolism illness that is globally prevalent, long-lasting, slow poison, and a danger to public health. Since there isn't a long-term solution for diabetics, accurate early identification is essential. To determine which machine learning (ML) method predicts slightly earlier mellitus the most accurately, this study compares the Bagging and Boosting Machine Learning models. The algorithms used were RF, Ada, Gradient Boost, XGBoost, some Normal algorithms linear regression, and support vector machines (SVM), all of which exhibited average accuracy. Hence, additional studies are needed to determine diabetic conditions owing to an absence of endeavor and low accuracy. So, to increase the effectiveness of the combined adaptable classification method and reduce the possibility of misdiagnosing a specific example, research work devised an ensemble method known as the "Bagging and Boosting Classifier." Random Forest and XGBoost are employed as conceptual. The proposed Bagging and Boosting classifiers outperform existing models for diabetes diagnosis, such as LR, SVM, Ada, GB, and CatBoost. Accuracy levels for Random Forest (RF) and XGBoost were 99.00% and 99.00%, respectively, which is extremely acceptable. As a result, the accuracy, sensitivity, precision, F1-score, specificity, and ROC AUC metrics are used to evaluate the effectiveness of the Machine learning techniques, and a visual comparison of all the Machine learning techniques is presented as a result. Particularly in impoverished nations, effective diabetes diagnosis using ML systems can considerably lower the yearly death rate. To effectively treat Chronic Diseases, practitioners may benefit significantly from this work. As a result, an effective e-healthcare system can be built in the coming years.

Keywords: Machine Learning Algorithms, Analysis, Random Forest (RF), XGBoost, Data Mining.

1. Introduction

Diabetes mellitus is a group of metabolism illnesses that affects a large portion of the population today and causes several issues inside the human system. Adolescents often develop Type 1 diabetes, also known as insulin-dependent diabetes mellitus (IDDM), due to inherited abnormalities that prevent the production of a hormone enough blood sugar. Meanwhile, Type 2 diabetes, also known as Non-Insulin-Dependent Diabetes Mellitus NIDDM, is typically seen in the age group of patients. In these situations, the human system is unable to utilize the insulin generated internally, cannot generate enough insulin, or both NIDDM [1]. Additionally, diabetes mellitus can appear

as a subsequent disorder brought on by a primary illness like pancreas illness, a hereditary trait like muscular dystrophy, or medications like corticosteroids. Pregnancy-related diabetes mellitus is a transient disorder. The blood sugar levels in this case rise through gestation but often revert to normal following delivery [2]. According to data from the 2014 World Health Organization Diabetes, it is projected that 42.2 million individuals worldwide, roughly a large percentage of the population have diabetes. A significant amount of data is produced throughout this era of big data, while algorithms have emerged as a crucial tool for analyzing the complexities of the information produced. Many methods, such as

classification algorithms and classification ensembles, have been used in big data analysis for clinical conditions [3]. Diabetes mellitus is a disease brought on by insufficient insulin levels in the pancreas or inefficient insulin utilization by the body. Blood glucose level is controlled by the hormone insulin. Metabolic derangement, or insulin levels, is a common side effect of high blood glucose levels that over time seriously harms a variety of biological methods, such as the brain and vital tissues. Adults 18 and older who might have diabetes made up 8.5% of the population in 2014. And to the total of 150 lakh deaths were directly related to diabetes in 2019, and 48 percent of these patients died under the age of 70. Early fatality rates (deaths before the age of 70) from diabetes increased by 5% between 2000 and 2016. From 2000 to 2010, the premature death rate in greater nations fell, but from 2010 to 2016, it rose. Both times, the early rate of death from Mellitus grew in lower-middle-income countries. In contrast, between 2000 and 2016, there was an 18% global decline in the rate of death between the ages of 30 and 70 from any of the four major diabetes complications (cardiac diseases, or metabolic disorders). Insulin deficiency production is a hallmark of T1dm, sometimes referred to as insulin-dependent, juvenile, or early life, which necessitates daily insulin therapy. T1dm affected 9 million people in 2017, the majority of whom reside in strong nations. Its etiology and methods of prevention are unknown. Polyuria (abnormal outflow of pee), dryness (excessive sweating), excessive hunger, loss of weight, visual abnormalities, and exhaustion are some of the signs. These signs could appear out of nowhere. It may be feasible to lower diagnosis mode, enhance the accuracy of health information, and decrease health expenses with ML algorithms. In contrast to other traditional techniques, such features are thus commonly utilized to examine diagnostics analyses. Just one way to lower the rates of death imposed on illnesses would be by rapid diagnosis and efficient therapies. The innovations in analysis methods in illness prediction are interesting to the majority of clinical professionals. A crucial strategy for using the vast amounts of diabetes data that are currently accessible for information retrieval is to utilize information mining and machine learning techniques in the study. Due to the serious socioeconomic consequences of the particular disease, DM is one of the top priorities in clinical scientific study, which invariably produces enormous quantities of information. The goal of this research is to evaluate the

efficacy of the best machine-learning methods for predicting the development of diabetic illnesses. The proposed method was evaluated on the dataset utilized in the system retrieved from the Kaggle Diabetes types dataset [4]. The diabetic dataset obtained from the hospital in Frankfurt Hospital, Germany will be used in this study. The data collection details contain 2000 instances, with 9 features. In this dataset, 1316 parameters belong to non-diabetic patients remaining 684 Parameters were diabetic patients. This study used all the available Attributes in all Research experiments. The Frankfurt Hospital, Germany Data was divided into trainsets (65.80%) and Parts tested 34.20% of the time. Furthermore, these were random selections, and to ensure that the method was not biased, the trials were run times. The Random Forest (RF), support vector machine (SVM), Gradient Boost, Catboost, AdaBoost, XGBoost, and Logistic Regression algorithms are utilized for databases to predict diabetes. However, based on the findings, it can be shown that when compared to the rest classification methods utilized in the suggested method, the suggested framework with RF and XGBoost produced the highest accuracy.

The diabetes prediction system using the Ensemble methods workflow system is shown in Figure 1.

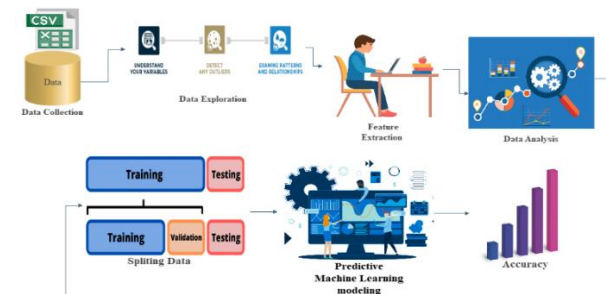


Figure 1: Proposed workflow for Diabetic Prediction Dataset system.

2. Materials work & Related Work

To improve the accuracy of the diagnosis process, diabetes training data depending on cutting-edge machine learning techniques is suggested in this study. As a result, this part will describe relevant test samples for diabetes, describe all of the segmentation techniques utilized in this study, and then provide the assessment score.

The literature survey regarding many machine learning methods, evaluated the accuracy of random forests, naive Bayes, and neural networks. A Matthews correlation analysis was utilized by the authors to assess various ML techniques [5].

The support vector machines classifier outperformed the Naive Bayes method after the authors used multiple classifiers using the k-fold cross-validation method for train & test [6].

This Pima Database was the subject of research that used Naïve Bayes, Random Forest, and Logistic Regression. They discussed those 3 techniques including the evolutionary algorithm and found that the algorithm performed much better accuracy [7].

This paper included a vital to enhance the feature extraction approach. For training and testing, a few Machine Learning algorithms, Radial Basis Function, K Nearest Neighbor (KNN), and Ada Boost, have also been used in the data [8].

Methods for supporting diabetic nephropathy were proposed that operate instantaneously to use classifiers, capturing variants from the problems identified and reflecting the skills of the medical professionals who believe that a high relationship between the health consequences of some diabetes complications as well as the blood sugar percentage. These survey implications could do something about simply classifying insulin treatments. Thus, the following are the primary duties: It uses a few unrestricted parameters [9].

Using deep neural networks on the Mellitus data set from PIMA. The Radial Basis Function (RBF), Multi-Layer Perceptron (MLP), and linear regression, classification techniques together up the deep computational model. Since the deep neural networks performed the filters, the database was purposefully left unprocessed. the writers' accuracy rate [10].

Developed a technique in which the reviewers determined the level of relationship between the attributes and substituted the missing data with the missing value. Another of the characteristics can be eliminated while the second is included for classifiers if there is a significant level of relationship between the 2 features. The models are used to estimate the relationship between characteristics and then build a correlogram matrix. The approach has an extremely long time to process [11].

Diabetes Dataset

Within that section, a summary of the information utilized will be provided. The diabetes information was downloaded from the Kaggle machine-learning library to be classified. There are examples of each numbered feature throughout each dataset. Together at the German Frankfurt Hospital, this information was acquired [12]. 2000 objects and 9 characteristics make up the dataset.

Table 1: Represents the Attributes, range, and statistical support for the dataset

Sl. No.	Attributes	Range	Description
11	Pregnancies	0-17	Number of women who are pregnancy
22	Glucose	0-199	Glucose is measured every 2 Hours in an oral glucose tolerance test
33	Blood Pressure	0-122	Diastolic (mm-Hg)
44	Skin Thickness	0-99	Triceps skin (mm)
55	Insulin	0-846	Every 2 hours insulin serum (uU/ml)
66	BMI	0-67.1	Weight and Height of the Body
77	Diabetes Pedigree Function	0.078-2.42	Family Heredity
88	Age	21-81	Year
99	Outcome	0 or 1	1-Diabetic 0-non-diabetic

Each sample is given a number in the column of the table, and the classifier that indicates the person's diagnosis of diabetes is defined in the final entry. Class factor 1 denotes a person with a disease, and class factor 0 denotes people without diabetes. Table 1

presents the cases, features, and several statistical support for the dataset.

Data statistics

The dataset is to explain a patient's Diabetes Mellitus depending on a few diagnosis attributes

included within the data samples. The data together with several health techniques. Information on data

statistics like mean, standard deviation, min, and max this detailed and shown in Table 2

Table 2: Detailed information of Dataset.

Attributes	Count	Mean	SD	Min	Max
Pregnancies	2000.0	3.703	3.306	0.000	17.00
Glucose	2000.0	121.1	32.06	0.000	99.00
Blood Pressure	2000.0	69.14	19.18	0.000	122.0
Skin Thickness	2000.0	20.93	16.10	0.000	110.0
Insulin	2000.0	80.25	111.1	0.000	744.0
BMI	2000.0	32.19	8.149	0.000	80.60
Diabetes Pedigree Function	2000.0	0.470	0.323	0.078	2.42
Age	2000.0	33.09	11.78	21.000	81.00
Outcome	2000.0	0.342	0.474	0.000	1.00

DPF in Mathematically

$$Pedigree = \frac{\sum_i K_i (88 - ADM_i) + 20}{\sum_j K_j (ALC_j - 14) + 50} \quad (1)$$

Measurement of performance

In this determination of performance in the data of this method, several common evaluation criteria like prediction, and correlation matrix have been taken into the data. The correlation coefficient is shown in Figure 2.

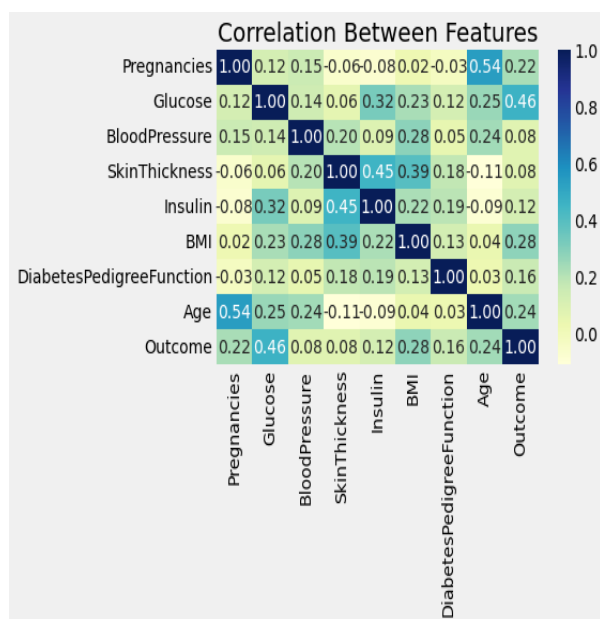


Figure 2: The Correlation map between features.

Presentation of Diabetes diagnosis

The presentation proposed a framework for the analysis of diabetes however, the dataset somewhat altered the placement, meaning there are approximately 1316 classes written as 0 indicating no diabetes, and 684 as 1 indicating diabetics. The number of count values is shown in Figure 3.

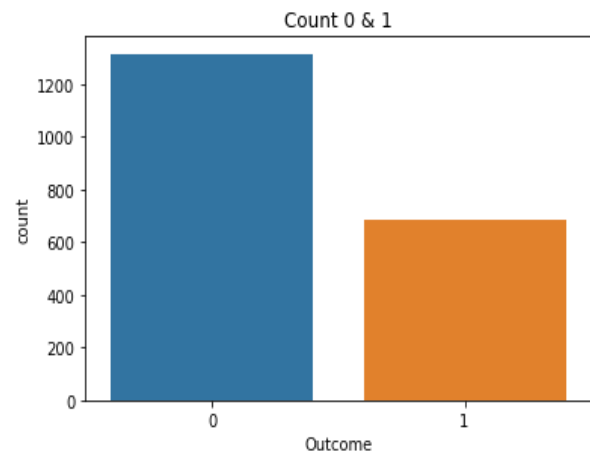


Figure 3: The number of diabetes and non-diabetic count values in the dataset.

Data missing values or null Values

The dataset includes 2000 patient records in total, 9 of which have some null data. The remaining 2000 patient records were utilized for pre-treatment after those 9 records were eliminated from the data. Null values that are missing in the given dataset are shown in Table 3.

Table 3: Missing Number of Units in Dataset.

Sl. No.	Attributes	Missing Values
1	Pregnancies	0
2	Glucose	13
3	Blood Pressure	90
4	Skin Thickness	573
5	Insulin	956
6	BMI	28
7	Diabetes Pedigree Function	0
8	Age	0
9	Outcome	0

3. Proposed Methodology

To create a clustered transformation matrix, the combining method that multiple models acquire is integrated through ensemble learning.

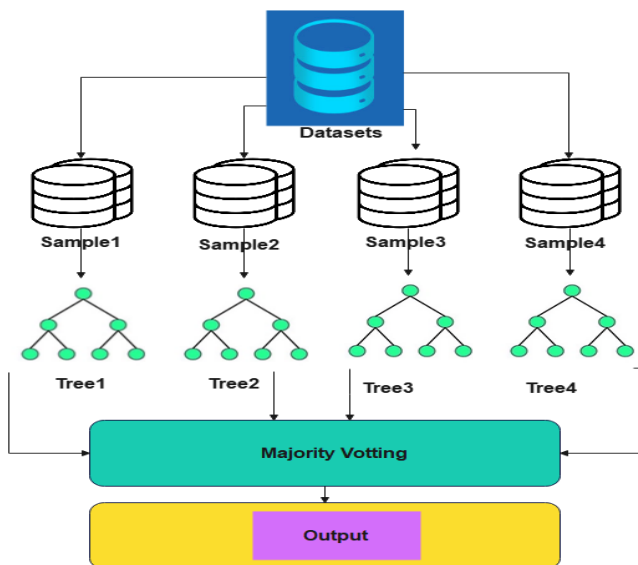


Figure 4: Workflow of the Bootstrapping Ensemble method.

To determine this combination, numerous methods have been developed over the years utilizing different strategies. The most well-known methods that are frequently used in research are discussed.

[13]. Every dataset should be taken to samples and the Bootstrapping methods to determine the dataset output as shown in Figure 4.

Among the earlier ensemble techniques to be presented is known as the Bagging ensemble methodology, which stands for Bootstrapping aggregating. The term Bootstrap testing refers to the

process of creating subsets out of a database for this approach. To put it in simple terms, replacement creates randomized picks from a collection of information, which could lead to a large number of instances that include identical information. Such instances shall be outfitted with different machine-learning approaches since they were treated as distinct datasets. While evaluating, every comparable algorithm's prediction on different subsets of comparable information shall be taken into consideration [14]. Utilizing a collection, the ultimate prediction is computed through processes like averaging, weighted averaging, etc.

The Glucose feature was where the source node of this graph was divided.

Samples: This represents the number of samples with a glucose level below 154.5.

Value: Provides the combined sample count for outcomes 0 and 1. For instance, the tree is categorized as belonging to class 0 since its value of 400 (which denotes category '0') is higher than its value of 214 (which denotes category '1').

Class: Diabetic - 1 or non-diabetic - 0.

The boosted ensemble technique and the method of the bagging process operate quite separately. Instead of processing the information in addition, it is treated consecutively in this instance. The initial algorithm receives the entire dataset, and the outcomes are evaluated. The samples close to the selection border of the training dataset, where Model 1 failed to predict accurately are sent to the 2nd Model.

The tree structure of each attribute value is determined as shown in Figure 5.

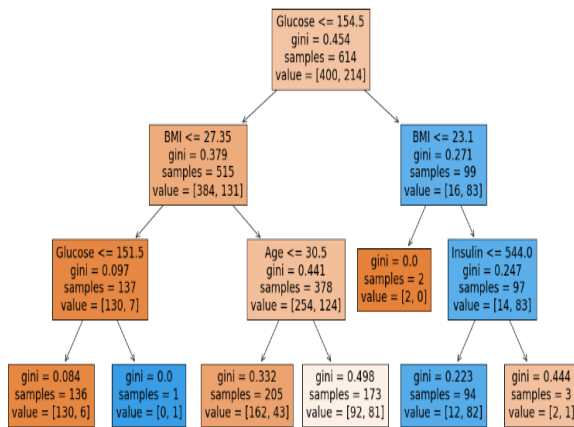


Figure 5: The tree structure of the dataset attributes value information.

This is done so that Model 2 can concentrate its attention on the characteristic design's difficult regions and learn the proper binary Model. Similar to this, more stages of the same concept are used, and the ultimate forecast on the testing data is derived using the ensembles of each of these prior models. Every dataset should be taken to samples and the Bagging methods to determine the dataset output as shown in Figure 6.

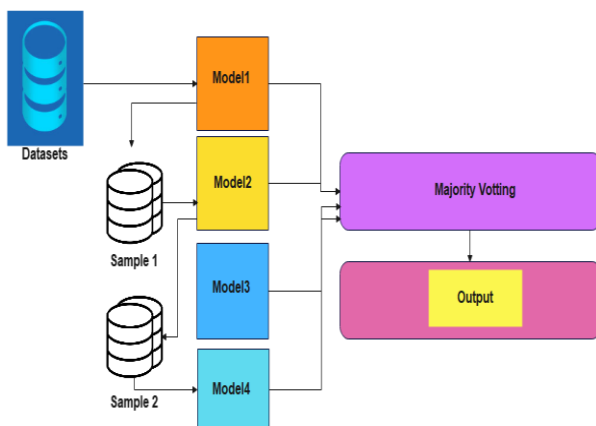


Figure 6: Workflow of the Bagging Ensemble method.

Data analysis and getting insight into the data are being carried out in the organization that uses data mining, among the most important and remarkable techniques. A range of data mining techniques, such as machine learning, statistics, and artificial intelligence, are used in prediction analytics. In this instance of work, disease estimation is performed using an Ensemble machine-learning technique. Ensemble ML provides a range of tools and methods that let computers turn unprocessed data into insightful, relevant information. The image illustrates the various types of machine learning (ML) techniques that are currently in use. These categories of machine learning algorithms are depicted. Classification and regression issues arise

during Ensemble learning. It is mostly used for forecasting since it builds a structure from information. The data also includes what was said or the outcomes. the model is trained on data with labels. The method of learning is used when a classification system is built but the results or responses are unknown. Utilizing data for training This type of learning is typically used for recognizing patterns and description analysis. Utilizing findings collected through interactions with the outside environment, the ultimate goal of aggregate machine-learning instruction is to make decisions that optimize reward or minimize danger.

These Ensemble machine-learning algorithms are Random Forest and XGBoost.

4. Implementation and Simulation

Random Forest Model

Some of the vector mathematical formulations use random forests, which are predictors made up of several well-chosen trees. This is one of those terms that is frequently used to express things like classification and prediction. It may be used to evaluate the importance of the various aspects.

The Random Forest is the best split using the Gini cost function by:

$$Gini = \sum_{z=1}^n (m_z \times (1 - m_z)) \quad (2)$$

Where z = each class and m = proportion of training instances.

XGBoost Model

A networked, modular gradient-boosted decision tree (GBDT) ensemble learning system is called Xtreme Gradient Boosting (XGBoost). It enables simultaneous tree boosting and is the best computational framework for prediction, classification, and scoring problems. Before comprehending XGBoost, it is crucial to comprehend pre-trained methods, decision networks, classification models, and gradient boosting the fundamental computing ideas and methods. While only a few attributes need to be chosen for prediction purposes, every one of the models has been applied because so far, the diagnosis of chronic disease is concentrated mostly on attribute selection strategic plan as then choose machine learning techniques LR, SVM, Ada boost, Cat Boost, Gradient Boost, and ensemble learning Random Forest, XGBoost [15].

Each tree of choices in the ensemble it creates fixes the mistakes of the ones before it. Allow me to explain XGBoost's mathematical foundation, which includes the model representation and objective function. In XGBoost, the model is an ensemble of trees, where each tree is added iteratively to reduce the residual errors from previous trees.

The model representation output:

$$\hat{z} = \sum_{k=1}^k f_k(x) \quad (3)$$

Where $\hat{z}$ = predicted value $f_k(x)$ = prediction value from the k-th Tree, K = Total number of trees.

The goal of XGBoost is to minimize the objective function, which combines a loss function (to measure the error between the predicted and actual values) and a regularization term (to prevent overfitting by penalizing overly complex models).

The objective function is given by:

$$L(\theta) = \sum_{i=1}^N l(z_i, \hat{z}_i) + \sum_{k=1}^K \Omega(f_k) \quad (4)$$

Where $l(z_i, \hat{z}_i)$ = loss function for the i-th function, $\Omega(f_k)$ = regularization term for the k-th tree, N = number of training examples, K = Total number of trees. All the proposed Models in the preceding study are developed utilizing Py 3.10 using a combination of Keras and Py APIs.

5. Results and Discussion

The efficacy of the approaches' classification in this section is given to assess every technique's performance, as shown in Figure 7.

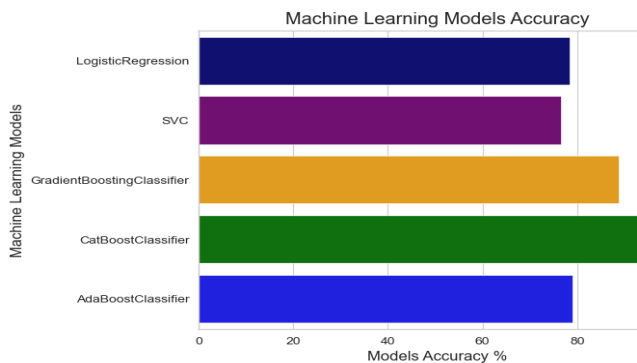


Figure 7: Performance Model Accuracy score for various ML algorithms.

The Ensemble ML classifier accuracy is shown in Figure 8.

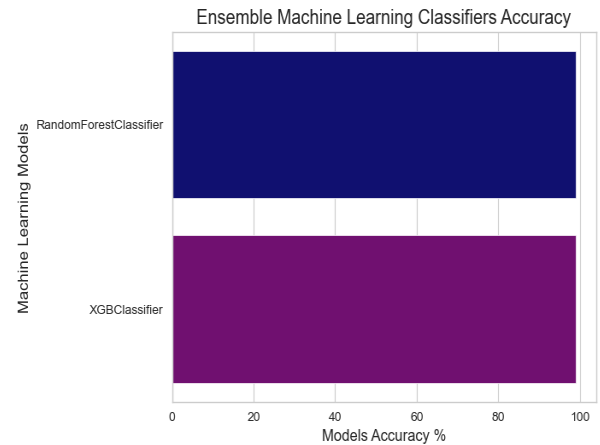


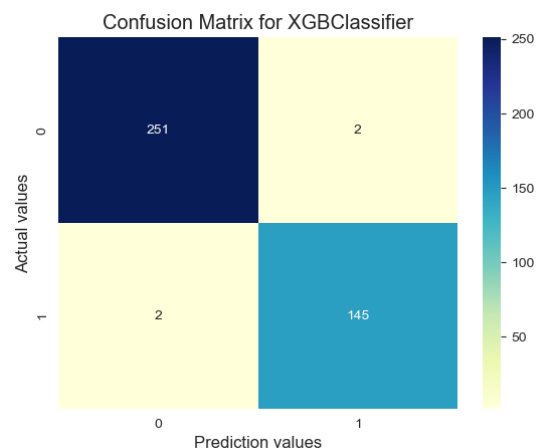
Figure 8: Performance Model Accuracy score for Ensemble algorithms.

Compared to other research as shown in Table 4.

Table 4: Comparison to other research methods.

Model name	Evaluation parameters	Accuracy	Ref
SVM, RF, NB, DT, KNN	Accuracy.	80%	[7]
XGB, ML-PR	ROC.	80%	[8]
SVM, RF	Accuracy.	89.86%	[10]
Neural Networks	Precision, Recall, F1-measure, ROC.	75.03 %	[11]
RF, XGB	ACC, Precision, Recall, F1-measure	99%	Proposed

These figures explain the confusion matrix present in the dataset. If the value of 'actual value' and 'prediction value' is the diabetic dataset is shown in Figure 9.



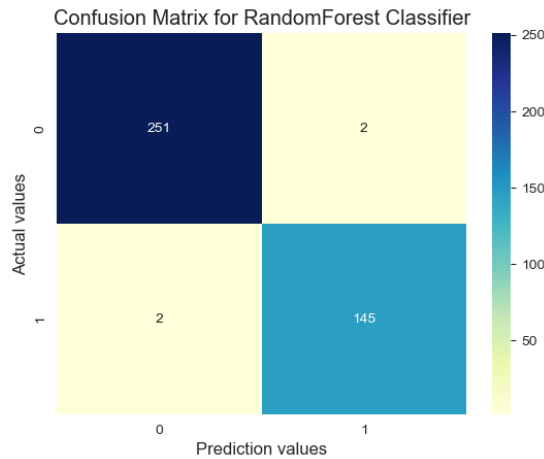


Figure 9: The Confusion Matrix for Ensemble Classifiers.

This figure explains the ROC Curve present in the dataset as shown in Figure 10.

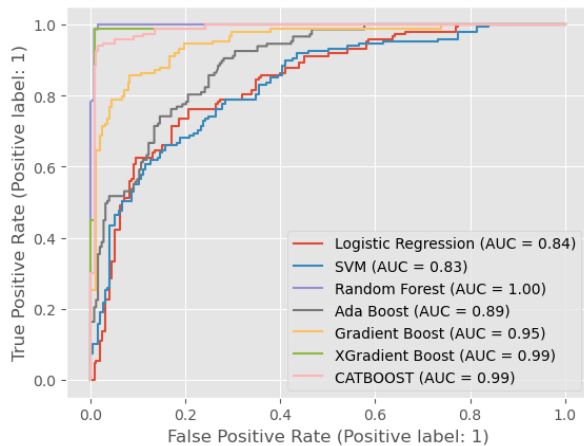


Figure 10: Simulation Results for ML Models AUC curve.

The ROC Curve of the Ensemble ML Models is present in the dataset. Random forest is the Bagging classifier and XGBoost is the Boosting Classifier as shown in Figure 11.

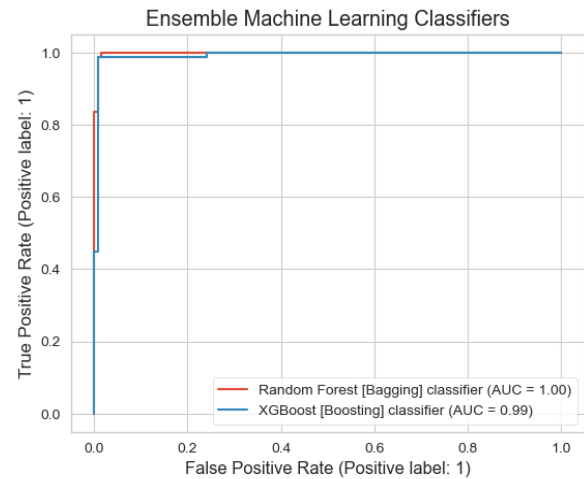


Figure 11: Simulation Results for Ensemble Classifiers AUC curve.

The Recall of the Ensemble ML Models present in the dataset. Random forest is the Bagging classifier and XGBoost is the Boosting Classifier as shown in Figure 12.

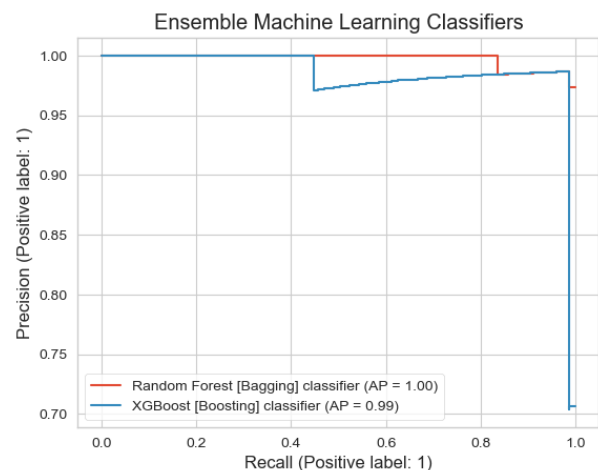


Figure 12: Simulation Results for Ensemble Classifiers Recall Curve.

Evaluation metrics:

This outcome is a True Positive (TP) if the method finds that the unit's actual value in the framework is True and thus it is True. Nevertheless, the forecast is a False Negative (FN) if the model of prediction predicts that it is going to be incorrect. Accordingly, the prediction may be True Negative (TN) if the approach forecasts that the division's real value in the whole thing is False but that value is True. It is False Positive (FP) if the learning process finds that the forecast is true. The figure clearly illustrates the confusion matrix, which makes it easy to assess the performance of the built prediction system. The evaluation techniques and the suggested evaluation were completed.

$$\text{Precision} = \frac{N_{TP}}{N_{TP} + N_{FP}} \quad (5)$$

$$\text{Accuracy} = \frac{N_{TP} + N_{TN}}{N_{TP} + N_{TN} + N_{FP} + N_{FN}} \times 100 \quad (6)$$

$$\text{Recall} = \frac{N_{TP}}{N_{TP} + N_{FN}} \quad (7)$$

$$F1\text{-Score} = \frac{2 \times \text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (8)$$

$$\text{Missed Detection Rate} = \frac{N_{FN}}{N_{TP} + N_{TN}} \quad (9)$$

$$\text{Diagnostic Odds Ratio} = \frac{N_{TP} \times N_{TN}}{N_{FP} \times N_{FN}} \quad (10)$$

6. Conclusion

Predictions and scholarly experts in the field may change how they assess health records if information processing is used more frequently in the medical sector. Well-known machine methods were used for massive data evaluation. These ensemble methods include XGBoost and RF. Additionally, they consist of LR, SVM, GB, CatBoost, and ADA. A diagnosis of diabetes was generated using the 2000 record dataset. For this type of training and test data of the predictions, 9 attributes were used. The results of the study show that the most trustworthy method for predicting diabetes is RF and XGBoost. Each of these strategies has a 99% accuracy rate, making it the most accurate method utilized in this research. As a result, it is possible to assert that RF and XGBoost are suitable for forecasting the onset of diabetes. The large size of the database and the attribute values that lack details are weaknesses in this research. There must be millions of data without any missing value to build Diabetes classification techniques with a 99.99% accuracy. Our future work will be heavily focused on the addition of new techniques to the present model to improve the precision of its characteristics. Finally, forecasts going to once evaluated, be greater in accuracy and provide a wide range of extra data to considerable data with few or no incomplete original datasets.

References

[1] Tejaswini, V., Nishat, Mirza Muntasir, et al. "Performance assessment of different machine learning algorithms in predicting diabetes

mellitus." *Biosc. Biotech. Res. Comm* 14.1 (2021): 74-82.

- [2] Abdulkader, Rasha Mahdi, and Assist Prof Dr Abdulbasit K. Alazzawi. "A Comparison of Five Machine Learning Algorithms in the Classification of Diabetes Dataset." *Turkish Journal of Computer and Mathematics Education (TURCOMAT)* 12.14 (2021): 1244-1259.
- [3] WHO, *Global Health Risks Report 2016*, <https://www.who.int/diabetes/global-report>, World Health Organisation, Geneva, Switzerland, 2020, <https://www.who.int/diabetes/global-report>
- [4] <https://www.kaggle.com/datasets/johndasilva/diabetes/diabetes.csv>
- [5] Dohare, S., Pamulaparthi, L., Abdufattokhov, S., Naga Ramesh, J. V., El-Ebiary, Y. A. B., & Thenmozhi, E. (2024). Enhancing Diabetes Management: A Hybrid Adaptive Machine Learning Approach for Intelligent Patient Monitoring in e-Health Systems. *International Journal of Advanced Computer Science & Applications*, 15(1).
- [6] Bhat, S. S., Ansari, G. A., & Ansari, M. D. (2024). Performance analysis of machine learning based on optimized feature selection for type II diabetes mellitus. *Multimedia Tools and Applications*, 1-20.
- [7] Uddin, M. A., Islam, M. M., Talukder, M. A., Hossain, M. A. A., Akhter, A., Aryal, S., & Muntaha, M. (2024). Machine learning based diabetes detection model for false negative reduction. *Biomedical Materials & Devices*, 2(1), 427-443.
- [8] Zou, X., Luo, Y., Huang, Q., Zhu, Z., Li, Y., Zhang, X., & Ji, L. (2024). Differential effect of interventions in patients with prediabetes stratified by a machine learning-based diabetes progression prediction model. *Diabetes, Obesity and Metabolism*, 26(1), 97-107.
- [9] Narasimharao, M., Swain, B., Nayak, P. P., & Bhuyan, S. (2023, November). Enhanced Diabetic Retinopathy Detection through Convolutional Neural Networks for Retinal Image Classification. In *2023 2nd International Conference on Ambient Intelligence in Health Care (ICAHC)* (pp. 1-6). IEEE.
- [10] Salih, M. S., Ibrahim, R. K., Zeebaree, S. R., Asaad, D., Zebari, L. M., & Abdulkareem, N. M. (2024). Diabetic prediction based on machine learning using PIMA indian dataset. *Communications on Applied Nonlinear Analysis*, 31(5s), 138-156.

- [11] Reza, M. S., Amin, R., Yasmin, R., Kulsum, W., & Ruhi, S. (2024). Improving diabetes disease patients classification using stacking ensemble method with PIMA and local healthcare data. *Heliyon*, 10(2).
- [12] Malik, Sumbal, Saad Harous, and Hesham El-Sayed. "Comparative analysis of machine learning algorithms for early prediction of diabetes mellitus in women." *Modelling and Implementation of Complex Systems: Proceedings of the 6th International Symposium, MISC 2020, Batna, Algeria, October 24-26, 2020* 6. Springer International Publishing, 2021.
- [13] Khan, Asfandiyar, et al. "Cardiovascular and Diabetes Diseases Classification Using Ensemble Stacking Classifiers with SVM as a Meta Classifier." *Diagnostics* 12.11 (2022): 2595.
- [14] Gollapalli, Mohammed, et al. "A novel stacking ensemble for detecting three types of diabetes mellitus using a Saudi Arabian dataset: pre-diabetes, T1DM, and T2DM." *Computers in Biology and Medicine* 147 (2022): 105757.
- [15] M. Narasimharao, P. P. Nayak, B. Swain and S. Bhuyan, "Performance Evaluation of Machine Learning Algorithms for Arrhythmia Classification in ECG Signals," 2023 International Conference in Advances in Power, Signal, and Information Technology (APSIT), Bhubaneswar, India, 2023, pp. 54-59, doi: 10.1109/APSIT58554.2023.10201737.