

A Web Application for Predicting Diabetes Using Machine Learning Methods

Riyaz Ahmed^{1*}, Piyush Kumar Gupta²

Abstract

Diabetes is a long-term disease caused by high glucose quantity in the blood. It has the potential to result in serious health complications like heart disease, hypertension, and ocular damage. It is good to identify any health issues as early as possible to get the right medical treatment and make necessary lifestyle adjustments. One makes use of machine learning techniques to predict diabetes and develop treatment options using actual cases. The various methodologies to be applied in diverse models entail K-nearest neighbor (KNN), logistic regression (LR), decision tree (DT), support vector machine (SVM), (RF Computers can predict if someone has diabetes by recognizing signs. We will integrate the trained models into a web application that will connect the models to provide real-time predictions based on factors responsible for diabetes such as body mass index (BMI), age, and insulin levels. We utilized the Kaggle dataset to construct a machine learning-based model for predicting diabetes. We developed a Flask application for diabetes prediction to provide insights into health status and risk factors. We organized the app into modules, such as routes, templates, forms, and static assets, and used Flask's modular structure. Our site uses HTML, CSS, and JavaScript to create dynamic content and enable interactivity. Machine learning generates a diabetes risk score based on personal details. We have deployed this algorithm to the web with Flask (a Python framework). Predictive models in healthcare could lead to better patient outcomes. Health providers might employ them through predictive analytics within electronic medical records, mobile apps for monitoring individuals' wellness statuses, or population health management systems that aim at identifying those most in need so that interventions can be prioritized accordingly while keeping track over time. It would be essential for future research efforts towards the integration of various data sources if we are to obtain more accurate results.

Keywords: Machine learning, diabetes prediction, application technology, health, SVM model

INTRODUCTION

Diabetes mellitus is a severe and intricate condition characterized by the inability of the body to store or utilize sugar from food. When the sugar levels in the blood rise too high, it is known as blood glucose or sugar. Diabetes is a chronic progressive disease that continues to grow worldwide. This can lead to blindness, kidney failure, heart attack, stroke, and lower limb amputation. Depending on the cause, diabetes can be temporary or permanent. It is composed of β -cells within the pancreas (an organ just below the stomach and behind the liver). In individuals diagnosed with type 1 diabetes, the immune system targets and destroys these cells, resulting in the absence of insulin [1]. In type 2 diabetes, glucose build-up in the blood results from insulin resistance or an insufficient amount of insulin being produced. To effectively treat

*Author for Correspondence

Riyaz Ahmed

E-mail: riyazahmed8380@gmail.com

¹Student, Department of Computer Science and Engineering, School of Engineering Science and Technology, Jamia Hamdard University, New Delhi, India

²Assistant Professor, Department of Computer Science and Engineering, School of Engineering Science and Technology, Jamia Hamdard University, New Delhi, India

Received Date: May 18, 2024

Accepted Date: September 26, 2024

Published Date: October 07, 2024

Citation: Riyaz Ahmed, Piyush Kumar Gupta. A Web Application for Predicting Diabetes Using Machine Learning Methods. Journal of Artificial Intelligence Research & Advances. 2024; 11(3): 92–102p.

diabetes, restore glucose levels to normal, and prevent complications, proper lifestyle changes and treatment must accompany the disease.

Medical practitioners should simplify the concept of the complicated nature of diabetes to better explain it. Diabetes refers to a condition in which the human body cannot manage its sugar levels, and hence the accumulation of sugar in the blood. This condition affects a large number of people, and if not managed, it can lead to organ failure, such as cardiovascular conditions, nerve damage, and damage to the kidneys. Thus, early detection of diabetes is essential to prevent further development of health issues prominent in later years [2].

Doctors have now started using machine learning algorithms much more often to help them better diagnose diabetes. These algorithms, which are at the forefront of technology in this area, involve computers sifting through large amounts of medical data to identify trends and determine whether someone has the disease. Machine learning can detect subtle changes in healthcare numbers that might mean someone will soon develop diabetes; therefore, doctors must know about it as soon as possible [3].

One of the greatest benefits of using these new technologies is that they have created tools for diagnosing prediabetes in people who do not yet have any symptoms related to high blood sugar levels. Such tools use many different factors such as age, weight, body mass index (BMI), levels of insulin within a body, and their current value of pressure when providing an estimate of how likely someone is going to develop diabetes mellitus.

In addition, machine learning algorithms are used by healthcare providers for personalized care of patients with diabetes. These algorithms can be used to examine extensive medical records. It is easy to identify unique patterns associated with each patient, thereby enabling doctors to develop specific treatment plans based on these unique characteristics [4].

In conclusion, machine learning algorithms in healthcare have changed everything about diagnosing or managing diabetes. According only to this system will mean better accuracy in treatment provision but also personalized attention towards patients, thus improving outcomes from their health conditions at a large-scale level, if not for everybody worldwide, which might be quite impossible.

Millions of individuals worldwide are affected by serious health conditions known as diabetes. If not controlled, there are many problems, such as heart disease, kidney damage, or nerve damage, which are chronic diseases. To solve this problem, a new diagnostic tool has been created to provide people with an accurate and reliable diagnosis of the possibility of diabetes.

A cutting-edge device employs sophisticated algorithms that consider various health indicators, including blood sugar level, body mass index (BMI), and family history, which can determine a person's chances of developing this illness. With a simple user interface design, individuals can easily input their information and receive comprehensive reports on the screen.

These generated reports offer detailed analyses about one's health status, including risks for getting diabetic, while suggesting this path through which somebody can manage their situation, such as lifestyle adjustments, diet changes, or exercise routines. By doing so, it encourages people to be responsible for their well-being based on educated choices made after interpreting what the document contains about themselves.

Overall, this revolutionary tool is a huge step in the fight against diabetes as it equips mankind with what they need to take charge of their condition and lead healthy, productive lives. The text outlines a sophisticated, very effective tool created to help people deal with their diabetes symptoms and blood

sugar levels. With many different integrated features and functions, this state-of-the-art device is extremely user-friendly and intuitive for use by any person [5].

Users can obtain considerable information about their condition in addition to different resources for decision-making concerning health and lifestyle. This is because the program is meant to provide detailed and correct data that will always keep them updated to build informed choices regarding their well-being.

By helping individuals recognize potential diabetes-related health problems at an early stage, the importance of this tool cannot be overemphasized. In cases that are identified early enough, serious health complications are prevented, and when necessary, immediate medical attention is sought. Thus, through proactive utilization, people can justify good health, thereby avoiding the onset of serious obstacles that may arise as a result.

The product is a highly complex and easily usable tool that can help people and medical workers. This tool is able to enhance interactions among patients and medical caregivers, thereby fostering good patient outcomes. This is made possible by promoting timely diagnosis and control of preventive services, in addition to supporting a reactive fitness care system [6].

Moreover, this app also avails health status updates, inclusive of potential risks. With the use of big data analytics techniques alongside predictive modeling, one is able to comprehend their different risks very well and see what types of ailments they are likely to get before they even attack them. Armed with this information, an individual may decide to implement specific preventive measures that suit them best based on those findings together with professional medical advice from their healthcare provider.

Its importance cannot be overlooked when it comes to continuous monitoring, which provides real-time feedback and alerts on one's health condition to maintain wellness in situations where such action is needed, especially if performed early enough. Since it can provide personalized interventions in addition to supporting the administration of proactive healthcare, this innovation will likely change how we receive our services from the hospital to the doorstep delivery point and back [7].

Finally, the above application is indeed a tool that is highly advanced and appreciated to help both the healthcare provider and the individual work more efficiently to achieve better health results. If individuals can regularly monitor their blood pressure levels, or if healthcare providers can track changes in their patients' conditions over time, and then adjust treatments accordingly, they will have moved a step closer to improving health outcomes through the application. This is possible due to its ability to communicate between different parties involved in healthcare delivery, and because promoting preventive medicine also provides personalized care tools for proactive management, among others; hence, it is possible to change how we deliver or manage healthcare services.

WEB APPLICATION FOR HEALTHCARE

HTML and CSS web application technologies are more popular than standard desktop applications. Combining these two creates website content, design, and style in the front-end components of what are called web applications. Layouts, alignments of visual effects, or background colors can be changed using them.

When putting a prediction model into production, most developers use Flask, which is a Python-based web application structure developed specifically for this purpose, among others. To deploy the prediction model in production, a Python web application framework called Flask was used to build the back end of the web application. The application is pulled from GitHub using Git and deployed on the Google Cloud platform, which is a cloud platform to deploy, manage, and scale apps that can be accessed via web URL by doctors.

Doctors wishing to utilize predictive models should deploy within websites to make decisions based on the data produced thereof rather than relying solely on intuition or previous experience.

LITERATURE REVIEW

Kumari et al. [2] proposed a model for classifying diseases and diagnoses using various dimensionality reduction and classification algorithms. It has a wide range of applications in healthcare, such as medical imaging, medical knowledge extraction, medical decision support, overall patient management, and protein-protein interactions.

Hasan et al. [8] proposed a deep learning convolutional neural network that is mostly used for the classification of image datasets and achieved 99.67% accuracy. Their research demonstrated the efficacy of utilizing deep learning neural network algorithms for analyzing vital human data, enabling pre-diagnosis without the need for specialized medical knowledge [8].

Dutta et al. [4] implemented an ensemble method and trained multiple algorithms on the same learning. They used Voting Classification and Stacking and indicated significantly better results than other classifier models. However, the stacked generalization method showed superior results compared to the voting classifier model. The combined prediction provided by the stacked generalization method reaches an accuracy of 79.87%, F1-Score of 0.8, Cohen-Kappa Score of 0.51, and mean square error (MSE) of 0.2, which is quite good in terms of diabetic tendency prediction.

Chauhan et al. [9] used three classification models to predict diabetes using a web form. This model consists of a decision tree, neural network, and I-Bayes. To check the robustness of each model, the accuracy and receiver operating characteristic (ROC) curves were calculated and compared with others.

Ashiquzzaman et al. [10] employed a dropout method to address overfitting in the prediction of diabetes. They utilized deep learning neural networks, incorporating both fully connected layers and dropout layers and achieved an accuracy of 88.41%.

Wei et al. [7] forecasted diabetes mellitus by assessing the effectiveness of different classifiers, including decision trees, support vector machine (SVM), and deep neural networks (DNN), across various data preprocessing methods. They adjusted the parameters to enhance the accuracy, achieving a 77.86% accuracy rate through 10-fold cross-validation.

Alić et al. [5] conducted a comparative analysis utilizing machine learning techniques to classify diabetes and cardiovascular diseases (CVD), employing Artificial Neural Networks (ANNs) and Bayesian Networks (BNs). ANN showed better accuracy of 99.51% and 97.92% than BNs for the classification of diabetes and CVD, respectively.

Zheng et al. [6] identified type 2 diabetes from the Electronic Health Record (I) to automatically extract patterns of type 2 diabetes using machine learning classification models and used a cross-validation strategy to evaluate the performance of the classification models. The classification models implemented were KNN, Naïve Bayes, Decision Tree, RF, SVM, and logistic regression. The data contained 221 samples and 160 type 2 diabetes patterns were extracted [6].

Kumar et al. [11] implemented big data analytics to analyze and compare different machine learning algorithms, such as Random Forest (RF), SVM, KNN, CART, and Linear discriminant analysis (LDA), to identify the best prediction algorithm based on various metrics such as accuracy, kappa, precision, recall, sensitivity, and specificity.

Sisodia et al. [12] predicted diabetes at an early stage using three machine learning classification algorithms: decision tree, SVM, and I-Bayes. Experiments were performed using the Pima Indian Diabetes Database (PIDD). The performances of all algorithms were measured based on various measures, such as precision, accuracy, measure, and recall. I-Bayes outperformed the other algorithms with the highest accuracy of 76.30%.

Mhaskar et al. [13] used deep learning to predict blood glucose levels in diabetes using continuous glucose monitoring devices. It suggests a deep learning paradigm in which a solid mathematical theory and domain knowledge are used to solve this problem accurately, especially when the target function has a compositional structure.

METHODOLOGY

SVM Model

Among the classification methods used in medical diagnosis, we used the support vector machine. This type of supervised learning can also handle regression [14]. The main goal of this approach is to minimize the empirical errors of classification and increase the margins of geometry or distances; thus, it may be called the maximum margin classifier. The SVM training algorithm follows a general principle known as structural risk minimization, which guarantees risk bounds from statistical learning theory.

The kernel trick allows effective nonlinear classifications through SVMs. It operates by implicitly mapping inputs into high-dimensional feature spaces. Consequently, one can build a classifier without explicit learning of the feature space, making it useful in statistical analysis and machine learning.

$$M = \{(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)\}$$

where $y_n = 1/-1$, a constant that denotes the class to which that point is x_n = number of data points belonging to the sample. Each x_n is a p-dimensional vector. The SVM classifier initially transforms the input vectors into decision values and then applies a suitable threshold value. This process helps to delineate the hyperplane to visualize the training data.

$$\text{Mapping } w \cdot x + b = 0$$

In this context, “w” represents a weight vector with p-p dimensions, and “b” denotes a scalar value. Vector w indicates the direction perpendicular to the separating hyperplane, as shown in Figure 1. The bias term b enables margin expansion. In cases where the training data are linearly separable, we select these hyperplanes to ensure that no points lie between them while also aiming to maximize the distance between them. We set the distance between the hyperplanes to $2/|w|$. To minimize $|w|$, we must ensure that, for all,

$$w \cdot x_1 - b \geq 1 \text{ or } w \cdot x_1 - b \leq -1$$

In linear classification, a weight vector ‘w’ of p-dimensions and a scalar ‘b’ are used to obtain a separating hyperplane. The scalar ‘b’ enables us to increase the margin while the weight vector ‘w’ is perpendicular to the hyperplane.

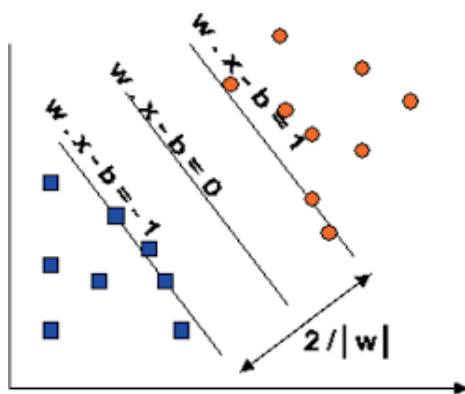


Figure 1. Maximum margin hyperplanes for SVM trained with samples from two classes.

When the training data are linearly separable, we select hyperplanes in a manner that ensures that there are no points between them and then strive to maximize the distance between these lines. The distance between the hyperplanes was defined as $2/|w|$.

Radial Basis Kernel Function

The SVM uses the Radial Basis Function (RBF) kernel as its classifier because the RBF kernel function can examine higher-dimensional data [15]. The output of this kernel depends on the distance x_j is from x_i (one should be a support vector, and the other will become a testing data point). In this case, γ determines how far into the data space this support vector reaches with its influence [15].

The kernel RBF function may be given by

$$k(x_j, x_i) = \exp(-\gamma \|x_i - x_j\|^2)$$

The training vector is x_i , and γ is a kernel parameter. A smoother decision face is produced by increasing γ , which makes the decision boundary more regular. This is because when RBF has a large value of γ , support vectors have a significant effect on large regions explosively. The training dataset was applied using stylish parameters to create a classifier. The testing dataset was classified by the created classifier to determine its vulnerability in the art of conception, as shown in Figure 2.

EXPERIMENTAL SETUP

Receiver operating characteristic curves were used to evaluate the models. It is a graph of true positive and false positive rates versus various cutoff points; each point on the graph represents a pair of sensitivity and specificity for a certain decision threshold.

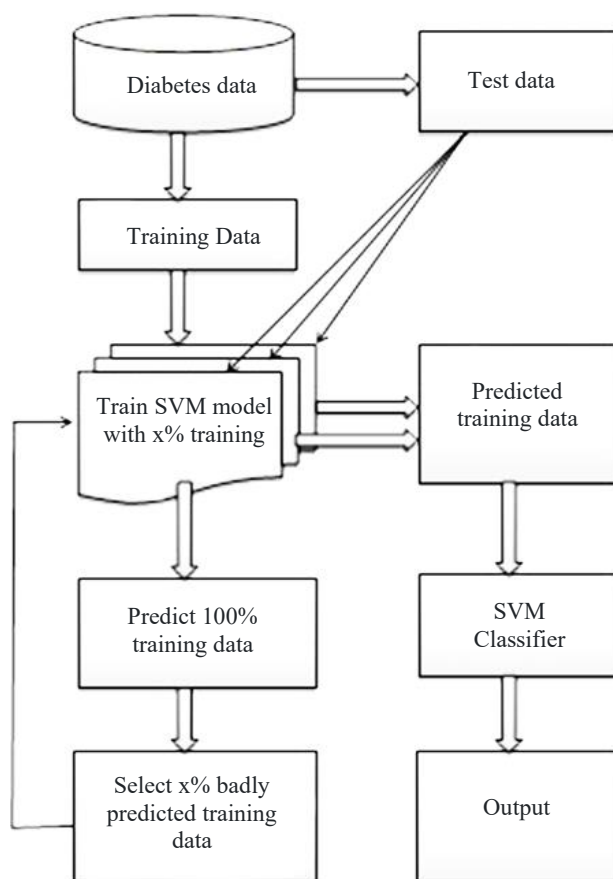


Figure 2. Using SVM disease classification.

Diabetes Disease Datasets

The Pima Indian diabetes dataset has 768 cases with medical reports on female patients aged at least 21 years who were of Pima Indian origin and lived in Phoenix, Arizona, USA [16]. The target variable in the dataset can have only two values '0' or '1.' A '1' means that one's test for diabetes is positive and '0' denotes negative tests. Of the 768 cases, there are 268 cases in class '1' whereas class '0' contains 500 instances. Variables were chosen based on significance criteria, whereby automatic selection was initially performed, followed by manual evaluation through fine-tuning parameters to further improve discriminative performance. Only variables with superior discriminative ability were retained.

There were eight numerical variables: (1) number of pregnancies; (2) glucose level; (3) diastolic blood pressure (mm Hg); (4) triceps skinfold thickness (mm); (5) 2-hour serum insulin (μ U/ml); (6) BMI; (7) Diabetes Pedigree Function; and (8) age (years). Although the dataset was labeled as having no missing values, some values were erroneously recorded as zero. Specifically, five cases had a glucose level of 0, 28 had a diastolic blood pressure of 0, 11 had a BMI of 0, 192 had skinfold thickness readings of 0, and 140 had serum insulin levels of 0. After removing these erroneous entries, 460 patients had complete data.

Training and Test Dataset Evaluation

Ten-fold cross-validation was conducted using the training dataset to gauge the resilience of the SVM classification models, a 10-fold cross-validation was conducted using the training dataset. Initially, the training dataset was divided into ten equally sized subsets. Each subset was subsequently utilized as a test dataset for a model trained in all instances, whereas an equal number of non-instances were randomly chosen from the remaining eight datasets. This cross-validation procedure was iterated ten times, with each subset serving once as the test dataset. The performances of the models were evaluated using these test datasets.

Now we are creating a web page application.

We selected a dataset of diabetes records of 768 Pima Indian Women. This dataset is valuable for analysis owing to the elevated incidence rate in this demographic, ensuring that it contains more representative diabetes features and symptoms compared to other datasets. The National Institute of Diabetes and Digestive and Kidney Diseases has extensively investigated eight features for examination since 1965. These eight features are outlined below.

1. Number of times pregnant
2. Glucose level
3. Diastolic blood pressure (mm. Hg)
4. Triceps skin fold thickness (mm)
5. 2-hour serum insulin (μ U/ml)
6. Body mass index (weight in kg/(height in m)²)
7. Diabetes pedigree function
8. Age (years)

RESULT ANALYSIS

Evaluating the ability of bracketing performance is done using measures of delicacy and AUC.

The classifier has four possible outcomes.

1. TP = (True Positive) is the number of instances correctly classified from the class.
2. TN = (True Negative) is the number of instances correctly rejected from the class.
3. FP = (False Positive) is the number of instances wrongly rejected from that class.
4. FN = (False Negative) is the number of instances incorrectly classified into that class.

We used the diabetes dataset for the bracket trials. An SVM classifier with an RBF kernel was used for bracketing. The dataset had 460 data points, of which 200 were used for training and the others were used for testing. Delirium was then calculated based on the diabetes dataset and the results are presented

in Table 1. Table 1 details the number of training and testing data points obtained from the diabetes database, the number of parameters used, and their sensitivity in terms of detection capability [16, 17].

To measure accuracy, sensitivity, and specificity, among other things, some equations are employed as a function to solve this problem through calculation depending on the outcome obtained.

- Accuracy = $(TP + TN)/(TP + TN + FP + FN)$
= $(TP + TN)/(TP + TN + FP + FN)$
- Recall (Sensitivity or true positive rate) = $TP/(TP + FN)$
- Precision = $TP/(TP + FP)$
= $TP/(TP + FN)$
= $TP/(TP + FP)$

The anticipated results are utilized to construct ROC curves, which provide a condensed depiction of the classifier performance. These curves quantify the balance between the true positive (TP) and false positive (FP) rates over various thresholds. Sensitivity is synonymous with the true positive rate, whereas specificity is synonymous with a false positive rate of 1. Sensitivity is the expectation of correctly identifying people with a disease (rate of disease detection), and specificity is the probability of correctly classifying people without the disease (false alarm rate). Sensitivity is the primary statistic to be maximized for accurate diagnosis, and specificity is the statistic that should be minimized, as shown in Table 2.

Table 1. Classification accuracy using SVM.

Data set	Sample	Training data	Testing data	Attributes	No. of classes	Using SVM
Diabetes	460	200	260	9	2	0.7555

Table 2. The performance of SVM classification.

Dataset	Accuracy	Sensitivity	Specificity
Diabetes	78%	80%	76.5%

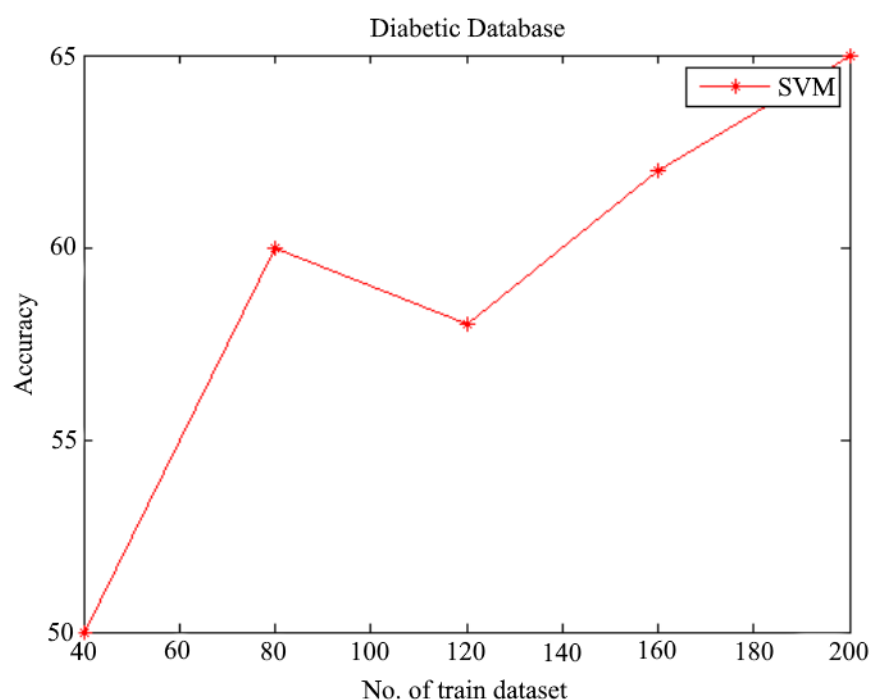


Figure 3. The training dataset accuracy of diabetes set.

The accuracy of the training dataset was 65.8%. The accuracy of the test dataset was 78.2% for it. The accuracy of the diabetes training set increased with an increase in the number of training data samples. The improvement in the accuracy was obvious. The ROC curve of the diabetes dataset is shown in Figures 3–5 and Table 3.

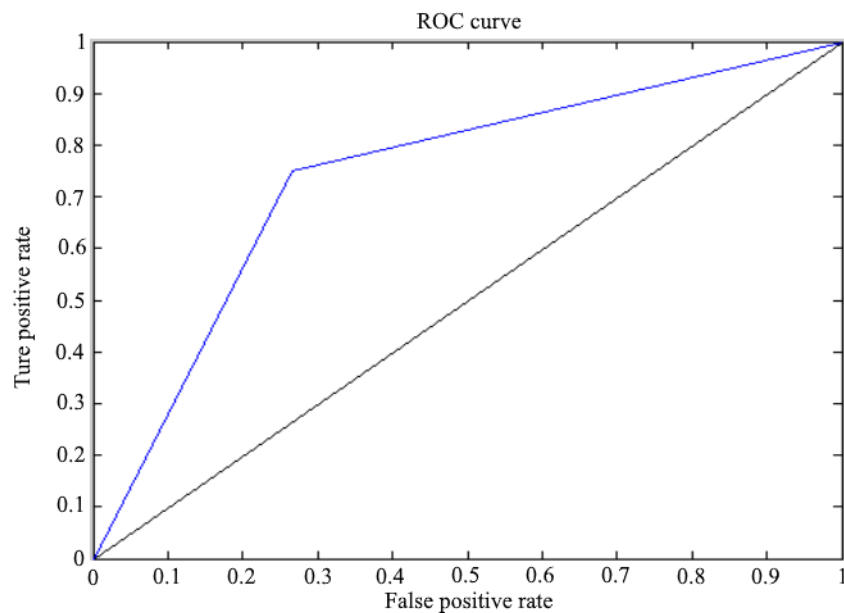


Figure 4. ROC of the SVM model using the diabetes data set.

The image shows a web-based form titled "Diabetes Prediction Jamia Hamdard By RIYAZ AHMED". The form is overlaid on a background image of a large white dome building, likely a mosque, with green trees in the foreground. The form includes the following input fields: "Number of Pregnancies:", "Glucose Level:", "Blood Pressure:", "Skin Thickness:", "Insulin Level:", "BMI:", "Diabetes Pedigree Function:", and "Age:". A green "Predict" button is located at the bottom of the form. A circular logo of Jamia Hamdard is visible in the top right corner of the form area.

Figure 5. Form of diabetes prediction.

Table3. The overview of the data set.

Pregnancies	Glucose	Blood pressure	Skin thickness	Insulin	BMI	Diabetes pedigree function	Age	Outcome
3	158	72	38	0	35.6	0.637	54	1
1	87	69	29	0	276	0.451	33	0
4	189	64	0	0	22.3	0.467	42	1
1	90	62	25	84	28.8	0.543	62	0
0	125	54	31	90	22.8	0.602	56	1

CONCLUSION

The accuracy of the SVM classifier for the diabetes prediction dataset was 78%, with a sensitivity of 80% and a specificity of 76.5%. Accuracy improved with an increase in the number of training data samples. The output shows that 1 patient had diabetes, but the result indicates that 0 patients had diabetes.

REFERENCES

1. Dey SK, Hossain A, Rahman MM. Implementation of a web application to predict diabetes disease: an approach using machine learning algorithm. 2018 21st International Conference of Computer and Information Technology (ICCIT), Dhaka, Bangladesh, 2018, pp. 1–5. DOI: 10.1109/ICCIT.2018.8631968.
2. Kumari VA, Chitra R. Classification of diabetes disease using support vector machine. Int J Eng Res Appl. 2013;3:1797–801.
3. Iswanto I, Laxmi Lydia E, Shankar K, Nguyen PT, Hashim W, Maselena A. Identifying diseases and diagnosis using machine learning. Int J Eng Adv Technol. 2019;8(6 Suppl 2):978–81. DOI: 10.35940/ijeat.F1297.0886S219.
4. Dutta S, Samir BK. Diabetes prediction using ensemble classifier. Int J Med Health Sci. 2020;9: 48–52.
5. Alić B, Gurbeta L, Badnjević A. Machine learning techniques for classification of diabetes and cardiovascular diseases. 2017 6th Mediterranean Conference on Embedded Computing (MECO), Bar, Montenegro. 2017. p. 1–4. DOI: 10.1109/MECO.2017.7977152.
6. Zheng T, Xie W, Xu L, He X, Zhang Y, You M, et al. A machine learning-based framework to identify type 2 diabetes through electronic health records. Int J Med Inform. 2017;97:120–7. DOI: 10.1016/j.ijmedinf.2016.09.014. PubMed: 27919371.
7. Wei S, Zhao X, Miao C. A comprehensive exploration to the machine learning techniques for diabetes identification. 2018 IEEE 4th World Forum on Internet of Things (WF-IoT), Singapore, 2018, pp. 291–5. DOI: 10.1109/WF-IoT.2018.8355130.
8. Hasan MK, Alam MA, Das D, Hossain E, Hasan M. Diabetes prediction using ensembling of different machine learning classifiers. IEEE Access. 2020;8:76516–31. DOI: 10.1109/ACCESS.2020.2989857.
9. Chauhan SC, Karvande V. Improve classification performance in diabetes prediction. Open Access Int J Sci Eng. 2019;4:29–33.
10. Ashiquzzaman A, Tushar AK, Islam MR, Shon D, Im K, Park JH, et al. Reduction of overfitting in diabetes prediction using deep learning neural network. In: IT Convergence and Security, Vol. 1. Springer: Singapore; 2018. p. 35–43. DOI: 10.1007/978-981-10-6451-7_5.
11. Kumar PS, Pranavi S. Performance analysis of machine learning algorithms on diabetes dataset using big data analytics. 2017 International Conference on Infocom Technologies and Unmanned Systems (Trends and Future Directions) (ICTUS), Dubai, United Arab Emirates. 2017, pp. 508–13. DOI: 10.1109/ICTUS.2017.8286062.
12. Sisodia D, Sisodia DS. Prediction of diabetes using classification algorithms. Procedia Comput Sci. 2018;132:1578–85. DOI: 10.1016/j.procs.2018.05.122.
13. Mhaskar HN, Pereverzyev SV, Van der Walt MD. A deep learning approach to diabetic blood glucose prediction. Front Appl Math Stat. 2017;3:14. DOI: 10.3389/fams.2017.00014.

14. Vapnik V. The nature of statistical learning theory. Springer Science+Business Media; 2013.
15. Park J, Sandberg IW. Universal approximation using radial-basis-function networks. *Neural Comput.* 1991;3(2):246–57. DOI: 10.1162/neco.1991.3.2.246. PubMed: 31167308.
16. Burges CJC. A tutorial on support vector machines for pattern recognition. *Data Min Knowl Discov.* 1998;2:121–67. DOI: 10.1023/A:1009715923555.
17. John (2018). Diabetes. [online] Kaggle.com. Available from: <https://www.kaggle.com/datasets/johndasilva/diabetes>.