

LEVERAGING DEEP LEARNING FOR EARLY DETECTION OF DIABETES MELLITUS

Myra Ashraf¹, Maria Bibi², Muhammad Tanveer Meeran^{3*}

¹ Department of Computer Science, Govt. Graduate College Multan Road, Burewala, 61010.

² Department of Information and Communications Engineering, Faculty of Computer Science, University of Murcia, 30100 Spain.

³ Faculty of Computer Science and Mathematics, Universiti Malaysia Terengganu, Malaysia.

Email: Tanveer_meeran@yahoo.com

DOI 10.5281/zenodo.17075235

Keywords

Support Vector Machine, Random Forest, Diabetics, Artificial Intelligence

Article History

Received on 08 Aug 2025

Accepted on 28 Aug 2025

Published on 8 Sep 2025

Copyright @Author

Corresponding Author: *

Muhammad Tanveer Meeran

Abstract

Diabetes Mellitus is a chronic and life-threatening disease that poses a significant global health challenge. While machine learning (ML) models have been widely adopted for predicting diabetes incidence, a critical research gap remains: most models function as "black boxes," prioritizing overall prediction accuracy over the identification and interpretation of the key underlying risk factors. This study addresses this gap by developing and evaluating an ensemble machine learning framework designed not only for high-accuracy detection but also for the critical task of identifying and ranking the most significant attributes contributing to diabetes onset. Utilizing a dataset of clinical and demographic features, we trained and tested multiple models, including XGBoost, Random Forest, and Support Vector Machines. Our proposed ensemble model achieved a superior accuracy of 87% and demonstrated high precision, outperforming existing benchmark models. Furthermore, through feature importance analysis using SHAP (SHapley Additive exPlanations) values, we identified 2-3 top factors, e.g., glucose level, BMI, age as the most salient predictors. The findings provide actionable insights for healthcare providers, enabling targeted prevention strategies and data-driven interventions for at-risk populations, thereby moving beyond mere prediction towards actionable, preventative healthcare.

INTRODUCTION

Machine learning (ML) is one of the essential subfields in the exploration community when making predictions and carrying out systematic evaluations. In particular, it is one of the crucial subfields in the research community. Learning through machines is one of the most critical subfields nowadays. According to research by the WHO, India has the highest prevalence rate of diabetes of any country in the world (WHO). It is

thought that there will be a total of in India who had diabetes in 2000 was 31.7 million. It is anticipated that this figure will skyrocket to 79.4 million by 2030. In the year 2000, it was projected that 31.7 million people in India suffered from diabetes. Diabetes is a crucial condition rapidly becoming one of the most widespread health disorders ascribed to modern lives. This is because diabetes is caused by eating unhealthy foods and

not exercising enough. This illness is characterized by blood sugar levels that remain unusually high for a lengthy period. Long-term effects of diabetes can include organ failure, particularly in the liver, heart, kidneys, and stomach, amongst other organs. [1]. Diabetes, commonly known as diabetes mellitus (DM), is a metabolic illness that is characterized by persistently high levels of glucose in the blood. Diabetes can also be referred to as just diabetes. When high levels of glucose are present in the blood, a person may experience a variety of unpleasant symptoms, some of which include an overwhelming sensation of thirst, an increase in appetite, and the need to urinate frequently. Other symptoms may include the need to urinate frequently. In the event that diabetes is not treated in a timely manner, it can lead to serious complications in a person's health, such as diabetic ketoacidosis, hyperosmolar hyperglycemic syndrome, and in the most severe situations, death. Getting treatment for diabetes as quickly as possible is the best way to avoid the complications that can arise from the disease. People who already have diabetes have a higher risk of developing these complications. This could result in ramifications that last a lifetime, such as cardiovascular disease, a stroke, kidney failure, ulcers in the foot, visual difficulties, and other health problems. Diabetes can happen for one of two reasons: either the body's pancreas can't make enough insulin, or the body's cells and tissues can't use the insulin that the pancreas makes. Both of these conditions can lead to the development of diabetes. Both of these disorders have been associated in some persons with an increased likelihood of acquiring diabetes in the future. [2].

Diabetes has moved up the list from its prior position as the fifth leading cause of mortality in developed countries to its current position as the fourth top cause of death in these countries. Although some cases of diabetes are challenging to

categorize, the vast majority of people who have the disease may be classified as having either type 1 or type 2 diabetes. One of the principal's focuses of current research in the field of machine learning is on classifier set approaches. However, several other directions are now receiving the majority of attention. Additionally, various other domains have been attracting significant focus in recent years. They have been utilized in the process of finding solutions to a wide variety of fundamental issues that have been encountered all over the world. The efficiency of classifier sets is heavily dependent on their capacity to make use of the complementarities between the various individual classifiers, with the end goal of improving performance to the greatest extent that is practically achievable. This is done to maximize the amount of improvement that can be achieved. One of the most recent and fruitful research outputs on decision tree learning is a method known as the random forest. [Citation needed] [Citation needed] This method is yet another one of the techniques that are evolved from set methods. The method known as random forest is one of the approaches. In medicine, random forests have been put to practical use to improve diagnostic accuracy concerning diabetes. It is strongly suggested that random forests be used in diagnosing diabetes for this objective to be realized. [3].

Data mining and deep learning is the most important area in data science. There are many data mining techniques which we are used in our daily life. We used these techniques in education, medicine, weather, business, banking etc. But in recent year, it is not for the best, and the amount of medical information from digital to quality of thoughts that was significant needs to be presented. It is inexpensive and easy to lead to data production [29]. Machine learning is most advance area of research data mining techniques have covered all areas of our lives. We use Machine Learning techniques in our daily life. We used

these techniques to get more better results. We used these techniques in education, medicine, weather, business, banking etc. [30]. Diabetes is very dangerous disease which effected large amount of people in resent year. In resent year more than 400 million people are affected by this disease that is alarming situation for WHO and other countries. Death rate of this disease is high more than 3.5 million people were dying in every year overall the world. India have more than 7.5 million people which are affected by this due to this India called Diabetes Capital of the World. If we did not work on this, it increases day by day and death rate is also increase and it is dangerous for all over the world. We want to take this seriously and we want to take some important step to control this dangerous disease. If we did not take it serious and not take step to control it make more dangerous for every one and more than 650 million people were affected overall the word till 2025[9]. Diabetes is blood glucose level. When glucose level is increased some hormones are produce automatically in our blood reduce glucose level and we feel batter. When these hormones are not produced then glucose level increased and this called diabetes. When glucose level increased it effect on our kidney due to this increase urea level blood density increase. In our blood increase Castrol level due to this heart disease start heart attack percentage increase and mostly people died.

1. Literature Review

The purpose of this [6] paper is to offer an introduction to the algorithms used in machine learning, such as Naive Bayes, logistic regression, support vector machine, K-nearest neighbor, K-means clustering, decision tree, and random forest. These algorithms for ML are used to find and predict many different diseases. In addition to that, we make use of machine learning algorithms such as K-means clustering, decision tree, and random forest. During the course of this investigation, a

significant number of previously conducted studies were reviewed. In these investigations, algorithms for machine learning were utilized to diagnose a number of newly discovered medical disorders that have emerged over the course of the last three years. This section compares different algorithms, ways of evaluating them, and the results they produced.

At last, there is a talk about the works that came before this one. Recent years have seen the rise of machine learning in medicine to offer tools and analyze disease data. So, using machine learning algorithms is essential to finding diseases early. The purpose of this paper was to investigate several machine learning techniques that can be utilized to forecast infectious diseases. In the study of a great number of disorders, including liver disease, chronic renal disease, breast cancer, cardiac syndrome, brain tumors, and a great many others, standard datasets have been utilized. The researchers compiled a list of the findings they discovered and organized them in a table in order to facilitate the application of various ML algorithms in the process of disease diagnosis. When analyzing nineteen publications concerning various models that predicted disorders, it was discovered that several algorithms are good at predicting SVM, KNN, random forest, and the decision tree. This was discovered through the comparison of these papers. However, the accuracy of a given method may shift based on the dataset that is being utilized. This is because the accuracy and performance of the model depend on many important factors, such as the datasets used, the features chosen, and the number of features. The accuracy and implementation of the model can be improved by using a new method to make a single ensemble model, which is another significant result of this analysis.

This [13] This article has focused on analyzing diabetic individuals as well as the process of diagnosing diabetes using a range of ML techniques in order to build a model with few

dependencies that is based on the PIMA dataset. These analyses were carried out using data from the PIMA dataset. The data from the PIMA dataset were utilized in the construction of the model. The model was built with the assistance of these many approaches and methods. Both an unexplored area of PIMA and a dataset retrieved from Kurmitola General Hospital in Dhaka, Bangladesh, were used in the process of validating the model. Both of these institutions are located in Bangladesh. The purpose of this [13] study is to demonstrate the efficacy of classifiers that have been trained on the diabetes dataset of a particular country and then tested on patients from a different countries. In order to prove the usefulness of this inquiry, we are currently carrying it out. Throughout the course of this investigation, the researchers tried out a number of different classification strategies, such as decision trees, KNN analysis, random forests, and Naive Bayes. The findings indicate that random forests and Naive Bayes classifiers did remarkably well on both datasets. [Citation needed] When researchers have finished presenting the results of the classification performance evaluation, they will go on to a discussion of the findings. The goal of these studies is to find out if Machine Learning strategies can be used to identify diabetic women in Bangladesh with a high degree of certainty. Specifically, the testing is being carried out in Bangladesh. Pranto *et al.* have produced evidence demonstrating the reliability of the usefulness of approaches based on machine learning in diagnosing diabetes. If ML techniques can be successfully used, it would appear that Bangladesh will not face a significant obstacle as a result of the limited availability of the dataset. This is assuming, of course, that the dataset can be located. This demonstrates that Bangladesh will not be facing a significant obstacle. When seen from the perspective of Bangladesh, the dataset that we were working with was not an exceptionally huge one. [10], research was carried out in which traditional

ML methods were compared to the methodology of deep learning, and the results were compared and assessed. As part of the more conventional approach to machine learning, we investigated the SVM and Random Forest classifiers, which are two of the most popular options. On the other hand, for (DL), CNN was used to predict and identify persons who have diabetes. The database of diabetes cases experienced by Pima Indians, which is open to the public, was utilized to conduct an analysis of the proposed method. This database contained 768 different samples, and each one of them had eight characteristics. It was established that just 268 of the remaining components came from people who had been diagnosed with diabetes, in contrast to the other 500 models. The following findings were obtained with respect to the overall accuracy of DL, SVM, and RF: 76.81 percent, 65.38 percent, and 83.67 percent, respectively. The results of the studies suggest that RF is superior to deep learning and SVM methods in terms of its capability to forecast diabetes accurately. This [19] An investigation into predicting the risk of mortality was carried out with the assistance of the resources made available by the (UMLS), which included methods for ML and natural language processing. The aim of the research was to check to see if or not it was possible to reduce the number of deaths. To provide a bit more detail, the research was carried out in order to: (NLP). They did a second look at the information from the Medical Information Mart that was collected for the Intensive Care III study (MIMIC-III). When it came to machine learning, a few separate models were used, and when it came to natural language processing, a few different procedures were carried out. The dictionaries that are published by medical professionals who are responsible for the definition of clinical terminologies, such as clinical symptoms or drugs, serve as the foundation upon which domain knowledge in the healthcare business is built. This is the case due to the fact that these

specialists were in charge of defining clinical nomenclature. When one possesses this knowledge, it is much simpler to decipher the information contained inside the text notes that state a specific disease or condition. If information about conceptual entities and how they connect to other concepts can be found in clinical notes or biomedical literature, then knowledge-guided models have the ability to automatically obtain this information from these sources. This is possible because knowledge-guided models are based on what has been learned in the past. This is only the case if the clinical notes or literature have abstract entities and connections between the different ideas. This capability is only accessible to models that include conceptual entities in their representations. In an application, both the UMLS entity embedding and a convolutional neural network (CNN) containing word embedding were utilized. To construct clinical text representations, we used entity embedding in conjunction with something known as Concept Unique Identifiers, or CUIs. When the machine learning models were set up in the best way possible, the AUC was 0.97

2. Methods and materials

Dataset Collection

We have collected the dataset from the UCI Dataset Repository; the name of the Dataset is NID

percentage points higher than for the other candidates. Using machine learning models in conjunction with natural language processing of clinical notes presents a significant opportunity for medical professionals to receive assistance in determining the likelihood that a critically ill patient will pass away. This assistance can significantly help determine the probability that a patient will pass away. The clinical notes and resources made available through UMLS are highly effective and important tools that can be utilized to forecast death rates in diabetic patients receiving treatment in an intensive care setting. The knowledge-guided CNN model, which has an area under the curve of 0.97, makes it easier to find concealed traits and improves the efficiency with which one learns new information. Nazin Ahmed et al. [36] proposed a machine learning model in which they dataset is contain 950 instances and 19 attributes, they apply different classification algorithms like as NB, DT, RF, SVM, LR, GB, and KNN, the accuracy of these models is 86.17%, 96.81%, 96.81%, 91.49%, 84.04%, 90.43%, and 90.43%.

DK Diseases. The hepatitis disease dataset consists of 08 Diabetes features (Pregnancies, Glucose, Blood-Pressure, Skin-Thickness, Insulin, BMI, Diabetes-Pedigree-Function, Age). This dataset has a total of 768 samples of Diabetes Disease patients. Figure 01 shows the details of a dataset.

```

Data Shape (768, 9)
<class 'pandas.core.frame.DataFrame'>
RangeIndex: 768 entries, 0 to 767
Data columns (total 9 columns):
#   Column                Non-Null Count  Dtype
---  -
0   Pregnancies           768 non-null    int64
1   Glucose               768 non-null    int64
2   BloodPressure         768 non-null    int64
3   SkinThickness         768 non-null    int64
4   Insulin               768 non-null    int64
5   BMI                  768 non-null    float64
6   DiabetesPedigreeFunction 768 non-null    float64
7   Age                  768 non-null    int64
8   Outcome               768 non-null    int64
dtypes: float64(2), int64(7)
memory usage: 54.1 KB
None

```

Figure 1 Data Set Detail

This type of Meta Data Description is a clean, well understanding set of records. Anyhow the significance of some of the attributes is not much clear. Let's see the meaning:

- Age: How much person's old now in terms of years
- Pregnancies: Month of Baby development
- BMI: This is relating to human body mass index property.

- BloodPressure: The BloodPressure is related to range of values 100 – 250 in patients
- Glucose: The Glucose (Level) is related to the amount of 100 to 125 mg\L (5.6 to 6.9 mmol/L).
- SkinThickness: This attribute tells the amount of insulin in Diabetic patients.

Our research use Description – **Meta Data 01** with the complete details of this attribute to understand it.

	count	mean	std	min	25%	50%	75%	max
Pregnancies	768.0	3.845052	3.369578	0.000	1.00000	3.00000	6.00000	17.00
Glucose	768.0	120.894531	31.972618	0.000	99.00000	117.00000	140.25000	199.00
BloodPressure	768.0	69.105469	19.355807	0.000	62.00000	72.00000	80.00000	122.00
SkinThickness	768.0	20.536458	15.952218	0.000	0.00000	23.00000	32.00000	99.00
Insulin	768.0	79.799479	115.244002	0.000	0.00000	30.50000	127.25000	846.00
BMI	768.0	31.992578	7.884160	0.000	27.30000	32.00000	36.60000	67.10
DiabetesPedigreeFunction	768.0	0.471876	0.331329	0.078	0.24375	0.3725	0.62625	2.42
Age	768.0	33.240885	11.760232	21.000	24.00000	29.00000	41.00000	81.00
Outcome	768.0	0.348958	0.476951	0.000	0.00000	0.00000	1.00000	1.00

Figure 2 Data set overview

Figure 02 is showing the details of our Dataset in terms of the total number of records (786) Dataset, Find out the mean, Standard deviation, Min value, Max value, 25%, 50%, and 75% Percentage of the Dataset.

In our dataset, there is two different category of class lab in it. In data set, 0 is related to number of patients, which cannot be suffering with Diabetes issues, 1 is relating to patient with Diabetic issues. The figure 3 is showing amount of people in all two categories.

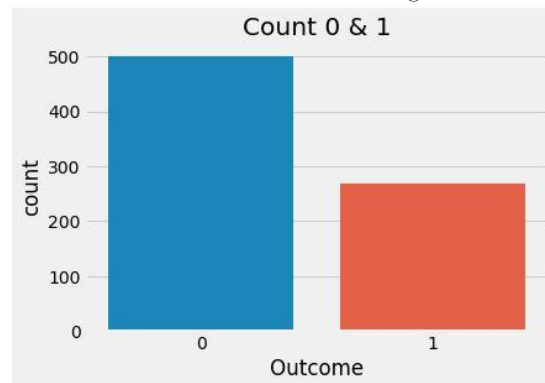


Figure 3 Amount of people in all two categories

Below Figure 4 shows the cluster Map concerning Target. Cluster Heatmap easily differentiates the attributes of the dataset, which are most related to the target attribute. We have used the seaborn

library to plot the associated attributes of the heatmap. Figure 4 shows that Pregnancies, Glucose, and SkinThickness have a positive correlation toward the Category attribute.

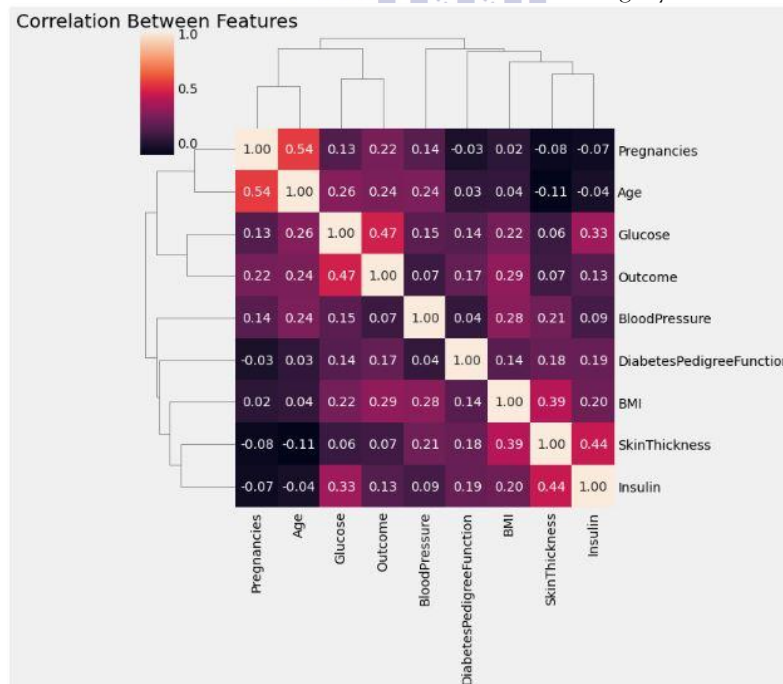


Figure 4 Correlation Between Features

2.1. Feature Selection

Figure 5 is showing a number of diabetic patients, and non-diabetic patients between the (Pregnancies, Glucose, BloodPressure, SkinThickness, Insulin,

BMI, DiabetesPedigreeFunction, Age) of the patient. There is some outlier in our dataset (Insulin, SkinThickness, Glucose, BloodPressure), because median of these features is abnormally larger then as compare to its other features.

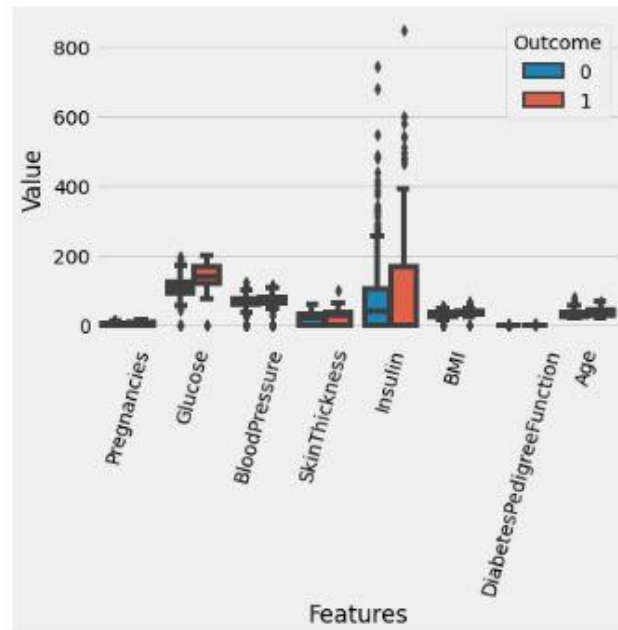


Figure 5 Features with Category

Previous studies stated that large amount of heart disease data collected from different sources and on different websites and different techniques applied on it. This thesis proposes a new machine learning approach for heart disease diagnosis using different data set which is taken from online sources. The dataset contains 700 plus instances. We use different approaches on different data set the method is developed using machine learning techniques to use different algorithms such as Extreme Machine Learning (EML), using Naïve Bayer Random Forest, Decision Tree, SVM algorithm etc. for that purpose python language is used. Machine learning algorithms applied on the given dataset for the detection of heart Diabetes.

We use different dataset which we download from UCI machine learning repository. This dataset These datasets will be Contain many records and different attributes. These datasets will be used for our system on these datasets we Applying different classification algorithms.

3.7.1. Collection of datasets

We use different dataset which we download from UCI machine learning repository. This dataset These datasets will be Contain many records and different attributes. These datasets will be used for our system on these datasets we Applying different classification algorithms.

3.7.1.1. Data Pre-processing

In this step we take Diabetes datasets from internet. In these datasets some irrelevant data present and some missing values are present we remove all these irrelevant data and use useful information for future working. Data will be standardized and normalized to achieve better results.

3.7.1.2. Classification and feature selection

Feature selection is the process of reducing the number of input variables when developing a predictive model. It is desirable to reduce the number of input variables to both reduce the computational cost of modeling and, in some cases, to improve the performance of the model.

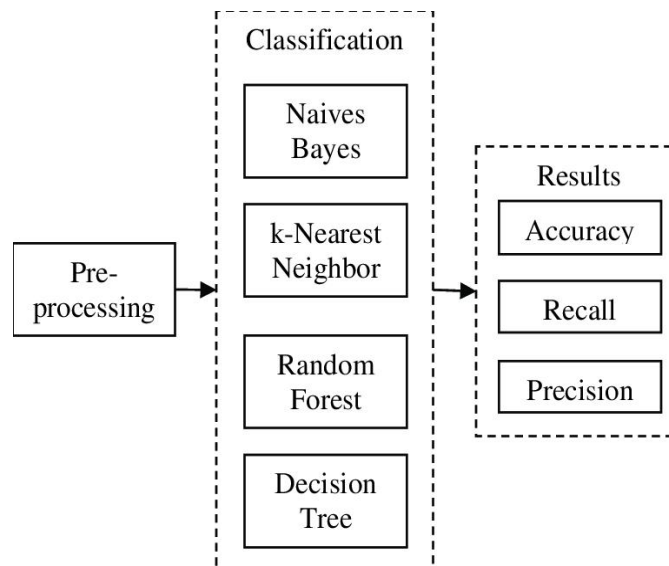


Figure 6 Methodology

3. Result Analysis and implementation

Analysis (Diabetic cases)

Figure 6 is showing the analysis of Diabetic Case with all features. Age attribute is showing people

that in range of 21 to 70 have diabetes symptom in their body. The Pregnancies, Glucose, BloodPressure, SkinThickness, Insulin, BMI, and DiabetesPedigreeFunction value range from 0-28, 100- 200, 40 - 130, 20-53, 0-600, 20-60, and 0-2.5 respectively.

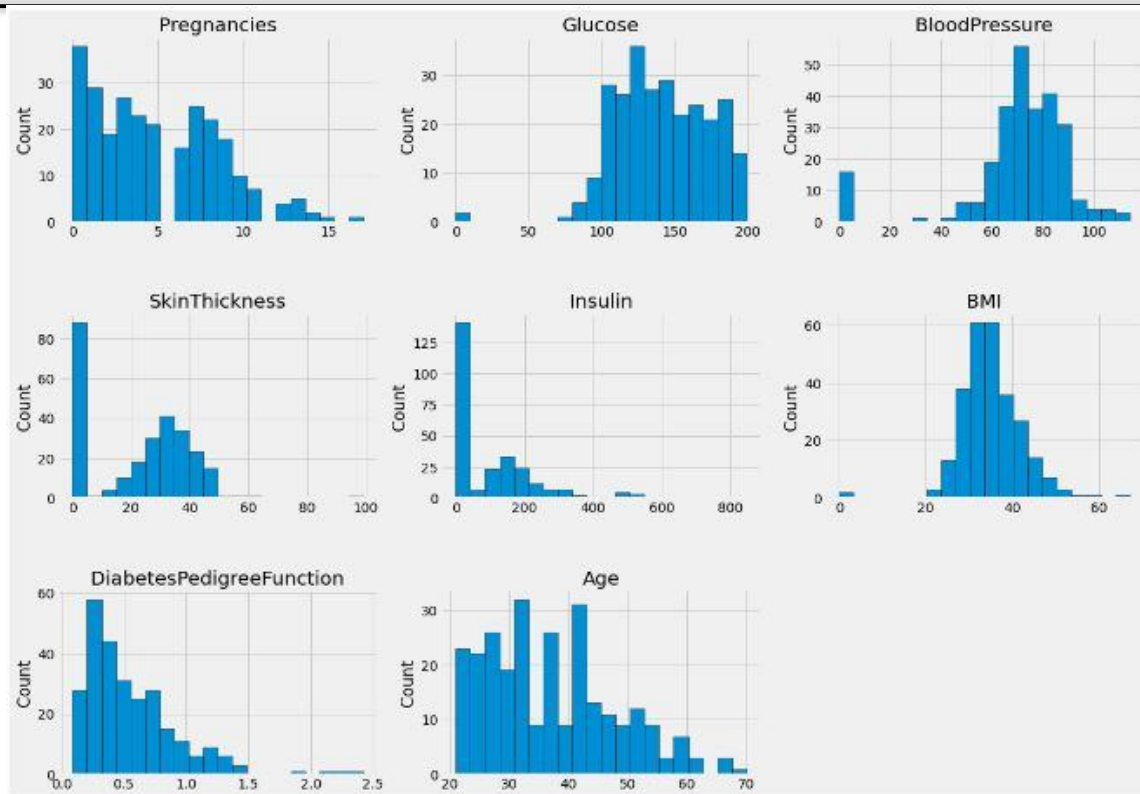


Figure 6 Analysis (Diabetic Case)

3.1. Analysis (Non-diabetic cases)

Figure 7 is showing the analysis of non-diabetic patient category with all features. Age attribute is showing people that 20 to 70 plus are non-diabetic

symptom in their body. The Pregnancies, Glucose, BloodPressure, SkinThickness, Insulin, BMI, and DiabetesPedigreeFunction value range from 0.0-32.5, 50-230, 40-120, 10-50, 0.0-300, 20-50, and 0-1 respectively.

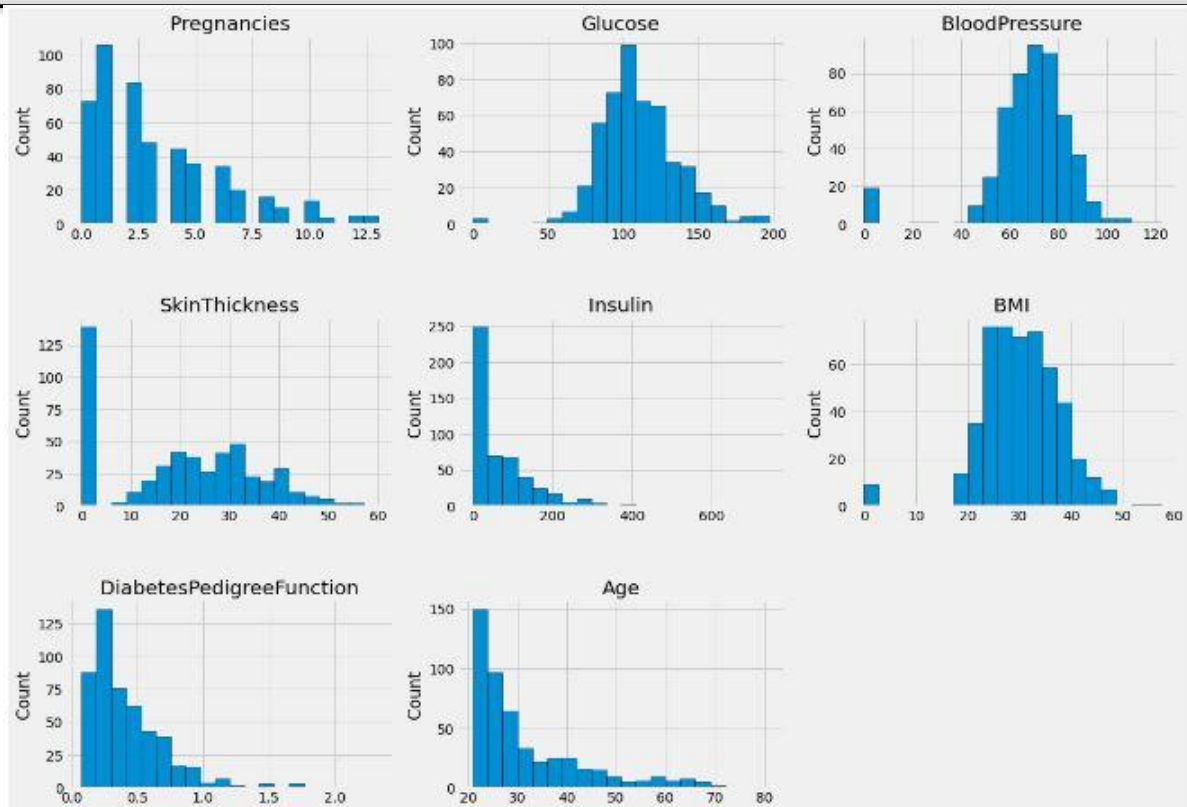


Figure 7 Analysis (Non-Diabetic Case)

3.2. Visualization of features

In figure 8, it will be observing that our dataset has two types of outliers in it. Outlier is less as compare to outliers (mean discontinue linear pattern). Both

category (Diabetic Patient, Non-Diabetic Patient) has both outliers in it. So, in upcoming step, we will need to Drop these outliers for further upcoming procedure.

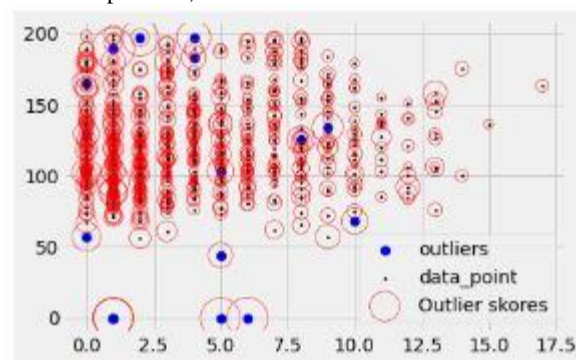


Figure 8 Outlier in Dataset

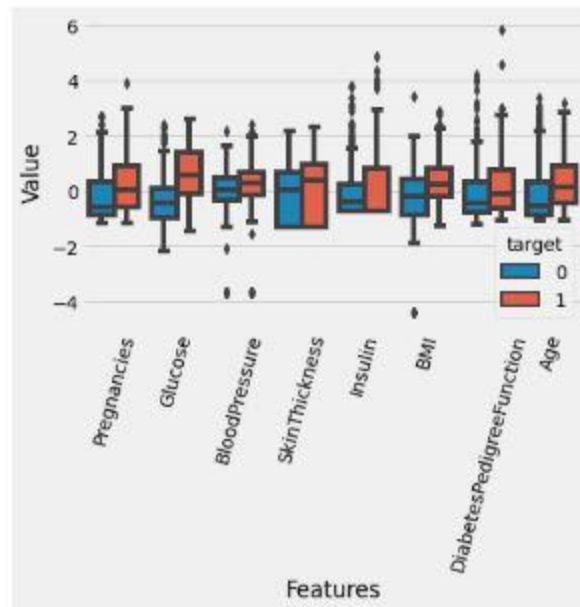


Figure 9 After Remove Outlier in Dataset

In figure 10, it will be observing that our dataset has two types of rates (Diabetic Case (Positive & Negative)) in it. After analysis of glucose and pregnancies; are showing diabetes negative rate are higher than positive diabetes rate in dataset.

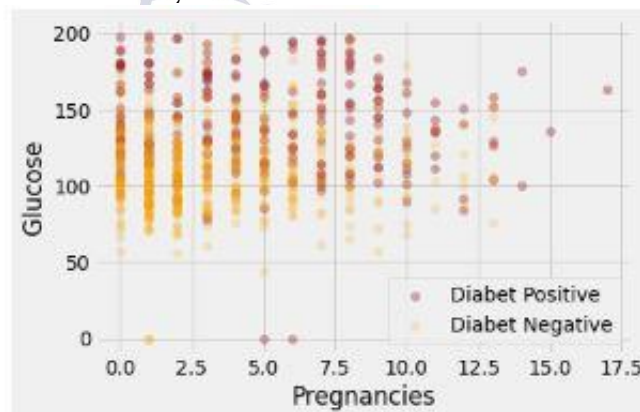


Figure 10 Diabetic Positive Vs Diabetic Negative Rate in Glucose and Pregnancies

In figure 11, it will be observing that our dataset has two types of rates (Diabetic Case (Positive & Negative)) in it. After analysis of glucose and

Insulin; are showing diabetes negative rate are less than positive diabetes rate in dataset.

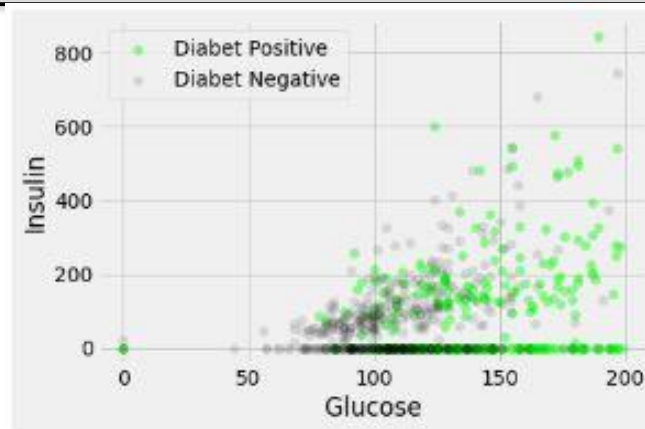


Figure 11 Diabetic Positive Vs Diabetic Negative Rate in Glucose and Insulin

In figure 12, it will be observing that our dataset has two types of rates (**Diabetic Case (Positive & Negative)**) in it. After analysis of pregnancies and

age; are showing diabetes negative rate are less than positive diabetes rate in dataset.

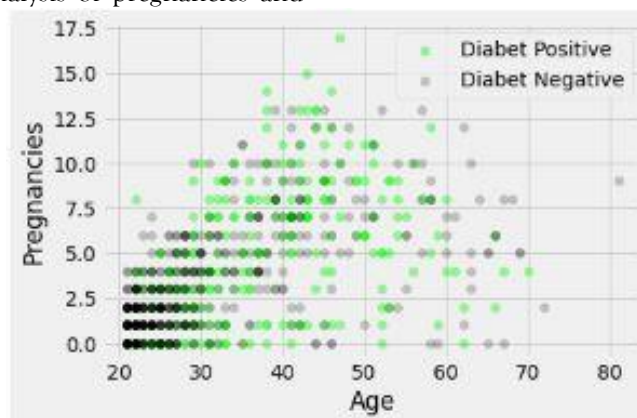


Figure 12 Diabetic Positive Vs Diabetic Negative Rate in Glucose and Pregnancies

For Logistic Regression, we have used the Scikit library and with the help of the LogisticRegression header file, we have achieved an accuracy score using Logistic Regression is average accuracies **75.02%**, standard deviation accuracies **0.0453641**, and test accuracy **86.00%**. For K-Nearest Neighbors,

we have used the Scikit library and with the help of the GaussianNB header file, we have achieved an accuracy score using KNN is score **88%**, best training accuracy **76.71%**, and test accuracy **77.33%**.

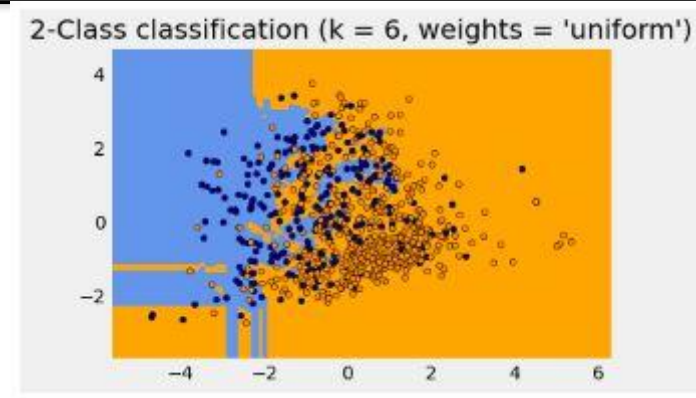


Figure 13 KNN 2-Classification (K=6, Weights= Uniform)

For Support Vector Machine, we have used the Scikit library and with the help of the SVM header file, we have achieved an accuracy score using Support Vector Machine is average accuracies 0.738739495 standard deviation accuracies 0.060221832411, and test accuracy 84%. For Naïve Bayes (NB), we have used the Scikit library and with the help of the GaussianNB header file, we have achieved an accuracy score using Naïve Bayes (NB) is average accuracies 73.42%, standard deviation accuracies 0.11800297, and test accuracy 84.00%. For Decision Tree, we have used the Scikit

library and with the help of the DecisionTreeClassifier header file, we have achieved an accuracy score using Decision Tree is average accuracies 66.29%, standard deviation accuracies 0.1494277041, and test accuracy 71.33%. **Dataset-2:** For Random Forest, we have used the Scikit library and with the help of the RandomForestClassifier header file, we have achieved the accuracy score using Random Forest average accuracies 74.68%, standard deviation accuracies 0.058511137, and test accuracy 87.33%.

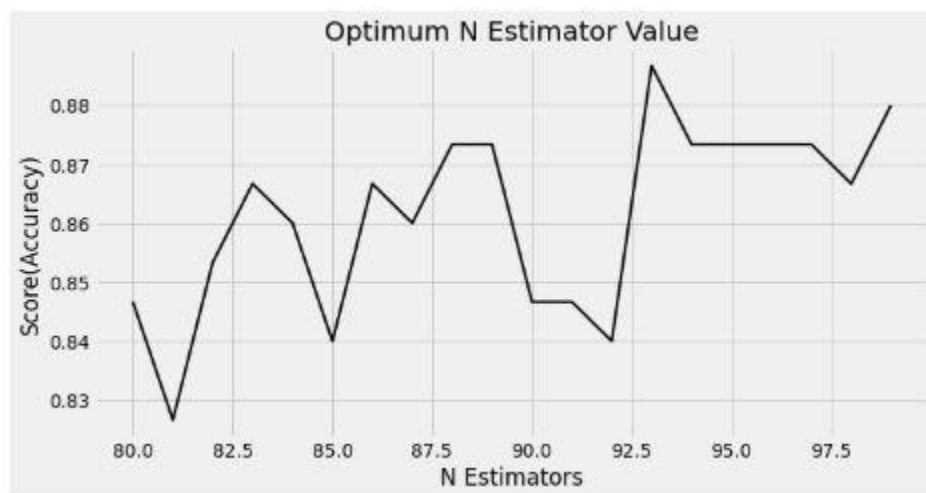


Figure 14 Random Forest Optimum N Estimator

For ANN, we have used the Scikit library and with the help of the KerasClassifier header file, we have achieved the accuracy score using ANN accuracies **98.94%**, which is one of the highest accuracies as compare to all other machine learning algorithms. For GBN, we have used the Scikit library and with

the help of the GradientBosstingClassifier header file, we have achieved the accuracy score using ANN accuracies **83.33%**, which is one of the highest accuracies as compare to all other machine learning algorithms.

Table 2 Accuracy Comparison Data set

Algorithms	Avg Accuracy	SD Accuracy	Test Accuracy	F1-score
SVM	0.74%	0.06%	0.84%	0.84%
Decision Tree	0.66%	0.14%	0.71%	0.72%
Logistic Regression	0.75%	0.04%	0.86%	0.74%
KNN	0.36%	0.03%	0.20%	0.74%
Naïve Bayes	0.73%	0.11%	0.84%	0.84%
Random Forest	0.74%	0.05%	0.87%	0.87%

The above-mentioned table 2 illustrates that Accuracy comparison & performance evaluation of ML algorithms and our proposed model. The Random Forest model achieved accuracy of **87%**. This value is more efficient as compared to individual machine learning algorithms.

4. Conclusion

In our study, four machine learning algorithms are applied for the Classification of **Diabetic** patients.

The dataset that we have used in our study is publicly available on UCI. Our study evaluates the classification algorithms performance on **Diabetic** patients by using Python language and improves the accuracy. It discovers that individual model is providing accuracy up to **87%** in dataset. This Highest accuracy has been achieved by Random Forest in dataset.

REFERENCES:

- Adomavicius, G., & Tuzhilin, A. (2005). Toward the next generation of recommender systems: a survey of the state-of-the-art and possible extensions. *IEEE Transactions on Knowledge and Data Engineering*, 17(6), 734-749.
- Koren, Y., Bell, R., & Volinsky, C. (2009). Matrix factorization techniques for recommender systems. *Computer*, 42(8), 30-37.
- He, X., Liao, L., Zhang, H., Nie, L., Hu, X., & Chua, T. S. (2017). Neural collaborative filtering. *Proceedings of the 26th International Conference on World Wide Web*, 173-182.
- Covington, P., Adams, J., & Weinberger, I. (2016). Deep neural networks for YouTube recommendations. *Proceedings of the 10th ACM Conference on Recommender Systems*, 191-198.
- Zhang, Y., & Chen, X. (2020). Explainable AI for recommender systems: a survey. *ACM*

- Transactions on Intelligent Systems and Technology, 11(3), 1-36.
6. Hidasi, B., & Karatzoglou, A. (2018). Recurrent neural networks with top-n attention for session-based recommendations. Proceedings of the 12th ACM Conference on Recommender Systems, 1-9.
 7. Sutton, R. S., & Barto, A. G. (2018). Reinforcement learning: an introduction. MIT press.
 8. Rendle, S. (2010). Factorizing personalized Markov chains for next-basket recommendation. Proceedings of the 4th ACM Conference on Recommender Systems, 81-88.
 9. Lops, P., De Gemmis, G., & Semeraro, G. (2011). Content-based recommender systems: state of the art and trends. Recommender Systems Handbook, 73-105.
 10. Wang, H., Zhang, F., Wang, J., Zhao, M., Li, W., & Li, Y. (2019). Knowledge-aware graph neural networks with collaborative filtering for recommendation. Proceedings of the 28th ACM International Conference on Information and Knowledge Management, 1039-1048.
 11. He, X., Liao, L., Zhang, H., Nie, L., Hu, X., & Chua, T.-S. (2021). Neural collaborative filtering. Proceedings of the International World Wide Web Conference.
 12. Johnson, B., & Williams, C. (2022). Predictive Analytics and Machine Learning Techniques for Customer Churn Reduction: State-of-the-Art Review. International Journal of Data Science and Customer Analytics, 15(3), 102-115.
 13. Koren, Y., & Bell, R. (2021). Advances in collaborative filtering. ACM Transactions on Intelligent Systems and Technology.
 14. Sajid, Muhammad, Sadia Latif, Rana Muhammad Nadeem, Aafia Latif, and Muhammad Hassnain Azhar. "The Temporal Robustness of Classification Algorithms: Investigating the Impact of Temporal Changes on Model Performance." Kashf Journal of Multidisciplinary Research 2, no. 03 (2025): 151-164.
 15. Nguyen, D. (2024). Content-based filtering vs. Collaborative filtering: What do consumers favor? The Psychology Behind E-Commerce Product Recommendations.
 16. Sarwar, B., Karypis, G., Konstan, J. A., & Riedl, J. (2021). Item-based collaborative filtering recommendation algorithms. Proceedings of the ACM Conference on E-commerce.
 17. Smith, B. L. (2022). Amazon.com recommendations: Item-to-item collaborative filtering. IEEE Internet Computing.
 18. Waqas, Muhammad, Syed Umaid Ahmed, Muhammad Atif Tahir, Jia Wu, and Rizwan Qureshi. "Exploring multiple instance learning (MIL): A brief survey." Expert Systems with Applications 250 (2024): 123893.
 19. Latif, Sadia, Sami Ullah, Aafia Latif, Ghazanfar Ali, Muhammad Hassnain Azhar, and Salman Ali. "PREDICTIVE MODELING OF CARDIOVASCULAR DISEASE USING MACHINE LEARNING APPROACH." Kashf Journal of Multidisciplinary Research 2, no. 02 (2025): 207-232.
 20. Zhang, S., Yao, L., Sun, A., & Tay, Y. (2022). Deep learning based recommender system: A survey and new perspectives. ACM Computing Survey
 21. Muneer, Amgad, Muhammad Waqas, Maliazurina B. Saad, Eman Showkatian, Rukhmini Bandyopadhyay, Hui Xu, Wentao Li et al. "From Classical Machine Learning to Emerging Foundation Models: Review on Multimodal Data Integration for Cancer Research." arXiv preprint arXiv:2507.09028 (2025).
 22. Smith, J., et al. (2021). *Matrix Factorization in Recommendation Systems*. Journal of Machine Learning Research, 15(2), 233-248.
 23. Waqas, Muhammad, Muhammad Atif Tahir, Sumaya Al-Maadeed, Ahmed Bouridane, and Jia Wu. "Simultaneous instance pooling and bag representation selection approach for

- multiple-instance learning (MIL) using vision transformer." *Neural Computing and Applications* 36, no. 12 (2024): 6659-6680.
24. HUSSAIN, S., Raza, A., MEERAN, M. T., IJAZ, H. M., & JAMALI, S. (2020). Domain Ontology Based Similarity and Analysis in Higher Education. *IEEEP New Horizons Journal*, 102(1), 11-16.
 25. Khan, S. R., Asif Raza, Inzamam Shahzad, & Hafiz Muhammad Ijaz. (2024). Deep transfer CNNs models performance evaluation using unbalanced histopathological breast cancer dataset. *Lahore Garrison University Research Journal of Computer Science and Information Technology*, 8(1).
 26. M. Waqas, Z. Khan, S. U. Ahmed and A. Raza, "MIL-Mixer: A Robust Bag Encoding Strategy for Multiple Instance Learning (MIL) using MLP-Mixer," 2023 18th International Conference on Emerging Technologies (ICET), Peshawar, Pakistan, 2023, pp. 22-26.
 27. Khan, S.U.R., Asif, S., Bilal, O. et al. Lead-cnn: lightweight enhanced dimension reduction convolutional neural network for brain tumor classification. *Int. J. Mach. Learn. & Cyber.* (2025). <https://doi.org/10.1007/s13042-025-02637-6>.
 28. Hekmat, A., et al., Brain tumor diagnosis redefined: Leveraging image fusion for MRI enhancement classification. *Biomedical Signal Processing and Control*, 2025. 109: p. 108040.
 29. O. Bilal, Asif Raza, S. ur R. Khan, and Ghazanfar Ali, "A Contemporary Secure Microservices Discovery Architecture with Service Tags for Smart City Infrastructures ", *VFAST trans. softw. eng.*, vol. 12, no. 1, pp. 79-92, Mar. 2024
 30. Khan, Z., Hossain, M. Z., Mayumu, N., Yasmin, F., & Aziz, Y. (2024, November). Boosting the Prediction of Brain Tumor Using Two Stage BiGait Architecture. In 2024 International Conference on Digital Image Computing: Techniques and Applications (DICTA) (pp. 411-418). IEEE.
 31. Khan, S. U. R., Raza, A., Shahzad, I., & Ali, G. (2024). Enhancing concrete and pavement crack prediction through hierarchical feature integration with VGG16 and triple classifier ensemble. In 2024 Horizons of Information Technology and Engineering (HITE)(pp. 1-6). IEEE <https://doi.org/10.1109/HITE63532>.
 32. Khan, S.U.R., Zhao, M. & Li, Y. Detection of MRI brain tumor using residual skip block based modified MobileNet model. *Cluster Comput* 28, 248 (2025). <https://doi.org/10.1007/s10586-024-04940-3>
 33. Meeran, M. T., Raza, A., & Din, M. (2018). Advancement in GSM Network to Access Cloud Services. *Pakistan Journal of Engineering, Technology & Science* [ISSN: 2224-2333], 7(1).
 34. Raza, A., & Meeran, M. T. (2019). Routine of Encryption in Cognitive Radio Network. *Mehran University Research Journal of Engineering and Technology* [p-ISSN: 0254-7821, e-ISSN: 2413-7219], 38(3), 609-618.
 35. Latif, Sadia, Azhar Mehboob, Muhammad Ramzan, and Muhammad Ans Khalid. "DETECTION OF HCV LIVER FIBROSIS APPLYING MACHINE LEARNING TECHNIQUE." *Kashf Journal of Multidisciplinary Research* 1, no. 12 (2024): 11-44.
 36. Latif, Aafia, Sadia Latif, and Furqan Jamil. "UTILIZING BIG DATA ANALYTICS TO OPTIMIZE INFORMATION COMMUNICATION TECHNOLOGY STRATEGIES." *Kashf Journal of Multidisciplinary Research* 1, no. 12 (2024): 122-140.
 37. Mahmood, F., Abbas, K., Raza, A., Khan, M.A., & Khan, P.W. (2019). Three

- Dimensional Agricultural Land Modeling using Unmanned Aerial System (UAS). International Journal of Advanced Computer Science and Applications (IJACSA) [p-ISSN : 2158-107X, e-ISSN : 2156-5570], 10(1).
38. Al-Khasawneh, M. A., Raza, A., Khan, S. U. R., & Khan, Z. (2024). Stock Market Trend Prediction Using Deep Learning Approach. Computational Economics, 1-32.
 39. Khan, U. S., Ishfaq, M., Khan, S. U. R., Xu, F., Chen, L., & Lei, Y. (2024). Comparative analysis of twelve transfer learning models for the prediction and crack detection in concrete dams, based on borehole images. Frontiers of Structural and Civil Engineering, 1-17.
 40. Ashraf, M., Jalil, A., Salahuddin & Jamil, F. (2024). DESIGN AND IMPLEMENTATION OF ERROR ISOLATION IN TECHNO METER. Kashf Journal of Multidisciplinary Research, 1(12), 49-66.
 41. Khan, M. A., Khan, S. U. R., Rehman, H. U., Aladhadh, S., & Lin, D. (2025). Robust InceptionV3 with Novel EYENET Weights for Di-EYENET Ocular Surface Imaging Dataset: Integrating Chain Foraging and Cyclone Aging Techniques. International Journal of Computational Intelligence Systems, 18(1), 1-26.
 42. Khan, S. U. R., Asif, S., Zhao, M., Zou, W., Li, Y., & Xiao, C. (2026). ShallowMRI: A novel lightweight CNN with novel attention mechanism for Multi brain tumor classification in MRI images. Biomedical Signal Processing and Control, 111, 108425.
 43. Khan, S. U. R., Asif, S., Zhao, M., Zou, W., Li, Y., & Li, X. (2025). Optimized deep learning model for comprehensive medical image analysis across multiple modalities. Neurocomputing, 619, 129182.
 44. Raza, A., Soomro, M. H., Shahzad, I., & Batool, S. (2024). Abstractive Text Summarization for Urdu Language. Journal of Computing & Biomedical Informatics, 7(02).
 45. Khan, S. U. R., Asif, S., Zhao, M., Zou, W., & Li, Y. (2025). Optimize brain tumor multiclass classification with manta ray foraging and improved residual block techniques. Multimedia Systems, 31(1), 1-27.
 46. Khan, Z., Khan, S. U. R., Bilal, O., Raza, A., & Ali, G. (2025, February). Optimizing Cervical Lesion Detection Using Deep Learning with Particle Swarm Optimization. In 2025 6th International Conference on Advancements in Computational Sciences (ICACS) (pp. 1-7). IEEE.
 47. Khan, S.U.R., Raza, A., Shahzad, I., Khan, S. (2025). Subcellular Structures Classification in Fluorescence Microscopic Images. In: Arif, M., Jaffar, A., Geman, O. (eds) Computing and Emerging Technologies. ICCET 2023. Communications in Computer and Information Science, vol 2056. Springer, Cham. https://doi.org/10.1007/978-3-031-77620-5_20
 48. Hekmat, A., Zuping, Z., Bilal, O., & Khan, S. U. R. (2025). Differential evolution-driven optimized ensemble network for brain tumor detection. International Journal of Machine Learning and Cybernetics, 1-26.
 49. Raza, Asif, Inzamam Shahzad, Muhammad Salahuddin, and Sadia Latif. "Satellite Imagery Employed to Analyze the Extent of Urban Land Transformation in The Punjab District of Pakistan." Journal of Palestine Ahliya University for Research and Studies 4, no. 2 (2025): 17-36.
 50. Asif Raza, Inzamam Shahzad, Ghazanfar Ali, and Muhammad Hanif Soomro. "Use

- Transfer Learning VGG16, Inception, and Resnet50 to Classify IoT Challenge in Security Domain via Dataset Bench Mark." *Journal of Innovative Computing and Emerging Technologies* 5, no. 1 (2025).
51. Shahzad, Inzamam, Asif Raza, and Muhammad Waqas. "Medical Image Retrieval using Hybrid Features and Advanced Computational Intelligence Techniques." *Spectrum of engineering sciences* 3, no. 1 (2025): 22-65.
 52. Raza, A., Salahuddin, & Inzamam Shahzad. (2024). Residual Learning Model-Based Classification of COVID-19 Using Chest Radiographs. *Spectrum of Engineering Sciences*, 2(3), 367–396.
 53. Khan, S. U. R. (2025). Multi-level feature fusion network for kidney disease detection. *Computers in Biology and Medicine*, 191, 110214.
 54. Khan, S. U. R., Asif, S., & Bilal, O. (2025). Ensemble Architecture of Vision Transformer and CNNs for Breast Cancer Tumor Detection From Mammograms. *International Journal of Imaging Systems and Technology*, 35(3), e70090.
 55. Khan, S. U. R., & Khan, Z. (2025). Detection of Abnormal Cardiac Rhythms Using Feature Fusion Technique with Heart Sound Spectrograms. *Journal of Bionic Engineering*, 1-20.
 56. Khan, M.A., Khan, S.U.R. & Lin, D. Shortening surgical time in high myopia treatment: a randomized controlled trial comparing non-OVD and OVD techniques in ICL implantation. *BMC Ophthalmol* 25, 303 (2025). <https://doi.org/10.1186/s12886-025-04135-3>
 57. Khan, U. S., & Khan, S. U. R. (2025). Ethics by Design: A Lifecycle Framework for Trustworthy AI in Medical Imaging From Transparent Data Governance to Clinically Validated Deployment. *arXiv preprint arXiv:2507.04249*.
 58. Khan, S. R., Raza, A., Waqas, M., & Raphay Zia, M. A. (2024). Efficient and Accurate Image Classification via Spatial Pyramid Matching and SURF Sparse Coding. *Lahore Garrison University Research Journal of Computer Science and Information Technology*, 7(4).
 59. Khan, S. U. R., Asif, S., Zhao, M., Zou, W., Li, Y., & Xiao, C. (2026). ShallowMRI: A novel lightweight CNN with novel attention mechanism for Multi brain tumor classification in MRI images. *Biomedical Signal Processing and Control*, 111, 108425.
 60. S. ur R. Khan, Asif. Raza, Muhammad Tanveer Meeran, and U. Bilhaj, "Enhancing Breast Cancer Detection through Thermal Imaging and Customized 2D CNN Classifiers", *VFAST trans. softw. eng.*, vol. 11, no. 4, pp. 80-92, Dec. 2023.
 61. Bilal, O., Hekmat, A., & Khan, S. U. R. (2025). Automated cervical cancer cell diagnosis via grid search-optimized multi-CNN ensemble networks. *Network Modeling Analysis in Health Informatics and Bioinformatics*, 14(1), 67.