

---

## COMPARATIVE ANALYSIS OF MACHINE LEARNING MODELS FOR HEART DISEASE PREDICTION

---

<sup>\*1</sup>Sanika Nivrutti Deshmukh, <sup>1</sup>Sneha Nanaso Nandre, <sup>1</sup>Prachi Jeevan Rajpure

<sup>2</sup>Dr. Kanchan Rathi

---

<sup>1</sup>*Department of Computer Application MAEER's MIT Arts, Commerce and Science College,  
Alandi (D.) Pune, India.*

<sup>2</sup>Assistant Professor, Department of Computer Application, MAEER's MIT Arts, Commerce  
and Science College, Alandi (D.), Pune, India.

Article Received: 25 July 2026, Article Revised: 15 August 2026, Published on: 05 September 2026

**\*Corresponding Author: Sanika Nivrutti Deshmukh**

Department of Computer Application MAEER's MIT Arts, Commerce and Science College, Alandi (D.) Pune, India.

Doi: <https://doi-doi.org/101555/ijarp.2325>

### 2. ABSTRACT

Heart disease remains one of the leading causes of mortality worldwide, making early diagnosis essential for improving patient outcomes and reducing healthcare costs. Machine Learning (ML) has emerged as an effective tool for predicting heart disease by analyzing clinical and demographic data, thereby assisting healthcare professionals in clinical decision-making. This study presents a comparative analysis of five widely used supervised machine learning algorithms, namely Logistic Regression, Decision Tree, Random Forest, Support Vector Machine (SVM), and K-Nearest Neighbors (KNN), for heart disease prediction. The research utilizes the publicly available UCI Heart Disease Dataset. Standard data preprocessing techniques, including data cleaning, feature scaling, and train-test splitting, are employed before model development. The performance of the selected algorithms is compared using evaluation metrics such as Accuracy, Precision, Recall, F1-Score, and Receiver Operating Characteristic–Area under the Curve (ROC-AUC). The comparative analysis (illustrative) indicates that ensemble learning methods, particularly Random Forest, provide superior predictive performance due to their robustness and ability to reduce overfitting. The findings demonstrate the potential of machine learning techniques in supporting early diagnosis of cardiovascular diseases and highlight the importance of selecting suitable classification algorithms for intelligent clinical decision support systems.

**3. KEYWORDS:** Heart Disease Prediction, Machine Learning, Logistic Regression, Decision Tree, Random Forest, Support Vector Machine, K-Nearest Neighbors, Healthcare, Clinical Decision Support System.

#### 4. INTRODUCTION

Cardiovascular disease (CVD), commonly known as heart disease, is one of the leading causes of death worldwide and continues to pose a significant challenge to modern healthcare systems. The increasing prevalence of cardiovascular disorders is associated with several risk factors, including hypertension, diabetes, obesity, smoking, high cholesterol, physical inactivity, and genetic predisposition. Early detection of heart disease plays a crucial role in reducing mortality, improving patient outcomes, and enabling timely medical intervention. Consequently, the development of reliable prediction systems has become an important research area in healthcare analytics and artificial intelligence [1], [9].

Traditional heart disease diagnosis relies on clinical examinations, laboratory investigations, electrocardiography (ECG), echocardiography, angiography, and expert medical interpretation. Although these methods provide accurate diagnosis, they are often expensive, time-consuming, and may not be readily available in rural or resource-constrained healthcare settings. Moreover, the increasing volume of electronic health records (EHRs) has made manual analysis more difficult, creating a need for intelligent computational techniques capable of assisting clinicians in diagnosis and decision-making [7], [9].

Recent advances in Artificial Intelligence (AI) and Machine Learning (ML) have transformed healthcare by enabling predictive analytics, disease diagnosis, medical image analysis, and intelligent clinical decision support. Unlike traditional statistical methods, machine learning algorithms automatically identify complex patterns and relationships from historical patient data without explicit programming. This capability has significantly improved the prediction of cardiovascular diseases and several other chronic illnesses [7], [8].

Several supervised machine learning algorithms have been successfully applied to heart disease prediction. Logistic Regression is widely used because of its simplicity and interpretability for binary classification problems. Decision Tree algorithms generate transparent rule-based classification models but may suffer from overfitting. Random Forest, an ensemble learning technique, improves prediction accuracy by combining multiple decision trees. Support Vector Machine (SVM) performs effectively in high-dimensional feature spaces by maximizing the separation between classes, while K-Nearest Neighbors (KNN) classifies patients based on similarity with neighboring observations [3]–[8].

The quality of machine learning predictions depends significantly on effective data preprocessing. Medical datasets often contain missing values, duplicate records, inconsistent measurements, and features measured on different scales. Appropriate preprocessing methods such as data cleaning, normalization, feature scaling, and train-test splitting improve model performance and reliability. Feature selection further reduces computational complexity by eliminating redundant variables while improving prediction accuracy [7], [8].

Despite considerable progress in this field, identifying the most suitable machine learning algorithm for heart disease prediction remains a challenge. Different studies employ different datasets, preprocessing techniques, validation strategies, and evaluation metrics, making direct comparison difficult. Furthermore, many studies focus primarily on prediction accuracy while giving less importance to evaluation metrics such as Precision, Recall, F1-Score, and ROC-AUC, which provide a more comprehensive assessment of classification performance [1], [7], [9].

### **Related Work**

Machine Learning (ML) has become an important approach for the early prediction and diagnosis of heart disease. ML techniques can analyze clinical and demographic information to identify patterns associated with cardiovascular diseases and provide useful support for medical decision-making. Several supervised learning algorithms, such as Logistic Regression, Decision Tree, Random Forest, Support Vector Machine (SVM), and K-Nearest Neighbors (KNN), have been widely investigated for heart disease prediction.

Previous research has demonstrated that the performance of ML models depends on factors such as data preprocessing, feature selection, feature scaling, and parameter tuning. Almustafa [1] compared different classifiers and reported that ensemble-based methods can provide better performance under suitable experimental conditions. The UCI Heart Disease Dataset, described by Dua and Graff [2], is one of the most commonly used benchmark datasets in this field. It contains important clinical attributes such as age, sex, chest pain type, blood pressure, cholesterol, and maximum heart rate.

Breiman [3] introduced Random Forest, which combines multiple decision trees to improve prediction accuracy and reduce overfitting. Cortes and Vapnik [4] proposed SVM, which is effective for classification by finding an optimal separating boundary between classes. Cover and Hart [5] introduced KNN, a simple instance-based method that classifies patients according to their nearest observations. Quinlan [6] developed the C4.5 Decision Tree algorithm, which provides interpretable classification rules based on information gain. These

algorithms have been widely applied to medical classification problems because of their different strengths in accuracy, simplicity, and interpretability.

Other researchers have emphasized the importance of proper statistical learning methods and evaluation techniques. Hastie et al. [7] discussed the significance of feature scaling, cross-validation, regularization, and multiple performance measures. Bishop [8] provided theoretical foundations for classification and probabilistic learning methods, while Kononenko [9] highlighted the importance of machine learning in medical diagnosis and clinical decision support. Detrano et al. [10] contributed to early research on the prediction of coronary artery disease using clinical observations, which influenced subsequent heart disease datasets and prediction studies.

Overall, the reviewed literature shows that machine learning has considerable potential for heart disease prediction. Random Forest and SVM frequently achieve strong predictive results, while Logistic Regression, Decision Tree, and KNN provide useful baseline models. However, differences in datasets, preprocessing methods, feature selection, and evaluation criteria make it difficult to determine which algorithm is most reliable, indicating the need for a standardized comparison.

### **Problem Statement**

Heart disease is one of the leading causes of death worldwide, making early diagnosis essential for improving patient outcomes and reducing mortality. Traditional diagnostic methods rely on clinical examinations, laboratory tests, electrocardiograms (ECG), and expert medical interpretation. Although these methods are effective, they are often time-consuming, expensive, and may not be readily available in all healthcare settings, particularly in resource-limited areas [1], [9].

Previous studies often employ different datasets, preprocessing techniques, feature selection methods, and evaluation criteria, making it difficult to compare the performance of machine learning algorithms fairly. Furthermore, many studies primarily focus on prediction accuracy while giving less attention to other important evaluation metrics such as Precision, Recall, F1-Score, and ROC-AUC, which provide a more comprehensive assessment of classification performance [1], [7], [9]. Therefore, there is a need for a standardized comparative analysis using identical preprocessing techniques and multiple evaluation metrics to identify the most reliable machine learning algorithm for heart disease prediction.

## Aim and Objectives

The primary aim of this research is to perform a comparative analysis of five supervised machine learning algorithms—Logistic Regression, Decision Tree, Random Forest, Support Vector Machine (SVM), and K-Nearest Neighbors (KNN)—for heart disease prediction using the UCI Heart Disease Dataset [2], evaluated using standardized preprocessing techniques and multiple evaluation metrics to identify the most reliable model for supporting early diagnosis and intelligent clinical decision-making [1], [7].

The specific objectives of this research are to:

1. Utilize the publicly available UCI Heart Disease Dataset for heart disease prediction [2].
2. Preprocess the dataset by handling missing values, removing duplicate records, applying feature scaling, and performing train-test splitting to improve data quality [7].
3. Comparatively analyze five supervised machine learning algorithms: Logistic Regression, Decision Tree, Random Forest, Support Vector Machine (SVM), and K-Nearest Neighbors (KNN) [3]–[8].
4. Evaluate the performance of each model using Accuracy, Precision, Recall, F1-Score, and ROC-AUC [7], [9].
5. Identify the strengths and limitations of each machine learning algorithm for heart disease prediction [1], [7].
6. Determine the most suitable algorithm for developing intelligent clinical decision support systems [3], [9].
7. Provide useful insights for future research in AI-based cardiovascular disease prediction [8], [9].

## Scope of the Study

This research focuses on the comparative analysis of five supervised machine learning algorithms—Logistic Regression, Decision Tree, Random Forest, Support Vector Machine (SVM), and K-Nearest Neighbors (KNN)—for heart disease prediction. The study utilizes the publicly available UCI Heart Disease Dataset, which contains clinical and demographic information related to cardiovascular disease [2]. The research is limited to supervised machine learning techniques and standard preprocessing methods, and does not include deep learning techniques, real-time clinical deployment, explainable artificial intelligence (XAI), or integration with hospital information systems. Instead, it provides a comparative

understanding of traditional supervised machine learning algorithms that can support future healthcare applications and intelligent clinical decision support systems.

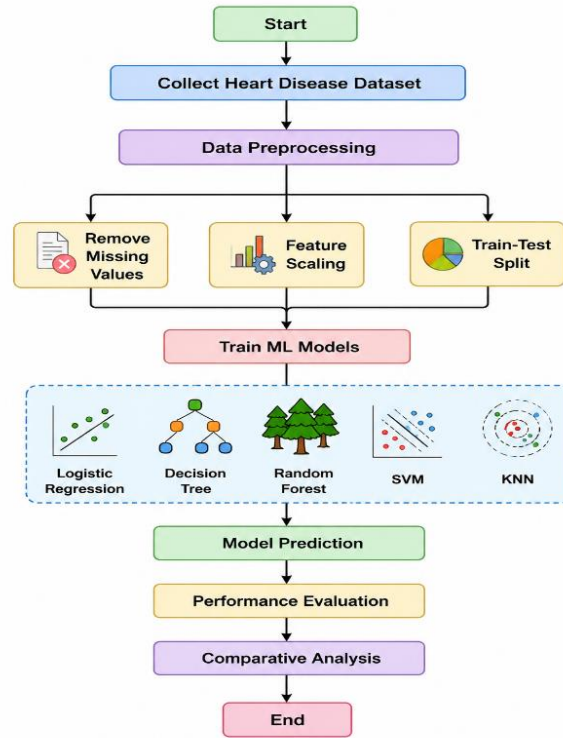
### **Research Gap**

Although several studies have applied machine learning techniques for heart disease prediction, a number of research gaps still exist. Many previous studies evaluate only one or two machine learning algorithms, making it difficult to determine the most suitable classifier for heart disease prediction [1], [9]. Different researchers also use different datasets, preprocessing methods, feature engineering techniques, and evaluation procedures, resulting in inconsistent findings and limiting the comparability of reported results [7], [8]. Another limitation is that many studies focus primarily on prediction accuracy while providing limited analysis of other important performance metrics such as Precision, Recall, F1-Score, and ROC-AUC, which are essential for evaluating medical prediction systems [1], [9].

To address these limitations, this research presents a standardized comparative analysis of five widely used supervised machine learning algorithms—Logistic Regression, Decision Tree, Random Forest, Support Vector Machine (SVM), and K-Nearest Neighbors (KNN)—using the same dataset, identical preprocessing techniques, and multiple evaluation metrics. This approach enables a fair comparison of classifier performance and assists in identifying the most suitable algorithm for heart disease prediction [2], [3], [7].

## **5. MATERIALS AND METHODS**

This study follows a systematic machine learning methodology for predicting heart disease using supervised classification algorithms. The workflow consists of data collection, data preprocessing, feature selection, model training, prediction, and performance evaluation. Five supervised machine learning algorithms—Logistic Regression, Decision Tree, Random Forest, Support Vector Machine (SVM), and K-Nearest Neighbors (KNN)—are comparatively analyzed using a common experimental framework designed to ensure a fair comparison of the selected algorithms under identical preprocessing conditions.



**Fig. 1. Workflow of the Proposed Heart Disease Prediction Methodology.**

### A. Data Collection

The dataset used in this research is the UCI Heart Disease Dataset, obtained from the UCI Machine Learning Repository [2]. The dataset is one of the most widely used benchmark datasets for heart disease prediction and contains clinical information collected from patients. It consists of 303 patient records and 14 attributes, including one target variable indicating the presence or absence of heart disease.

Each row in the dataset represents one patient, while each column represents a clinical feature such as age, sex, cholesterol level, blood pressure, maximum heart rate, and chest pain type. The dataset is suitable for binary classification problems, where the target variable is 0 (No Heart Disease) or 1 (Heart Disease).

**Table I. Sample Records from the UCI Heart Disease Dataset**

| Age | Sex | CP | RestBP | Chol | FBS | RestECG | MaxHR | ExAng | Oldpeak | Slope | CA | Thal | Target |
|-----|-----|----|--------|------|-----|---------|-------|-------|---------|-------|----|------|--------|
| 63  | 1   | 3  | 145    | 233  | 1   | 0       | 150   | 0     | 2.3     | 0     | 0  | 1    | 1      |
| 37  | 1   | 2  | 130    | 250  | 0   | 1       | 187   | 0     | 3.5     | 0     | 0  | 2    | 1      |
| 41  | 0   | 1  | 130    | 204  | 0   | 0       | 172   | 0     | 1.4     | 2     | 0  | 2    | 1      |
| 56  | 1   | 1  | 120    | 236  | 0   | 1       | 178   | 0     | 0.8     | 2     | 0  | 2    | 1      |
| 57  | 0   | 0  | 120    | 354  | 0   | 1       | 163   | 1     | 0.6     | 2     | 0  | 2    | 1      |

**Table II. Dataset Attribute Description.**

| Attribute | Description                                                       |
|-----------|-------------------------------------------------------------------|
| Age       | Age of the patient (years)                                        |
| Sex       | Gender (1 = Male, 0 = Female)                                     |
| CP        | Chest Pain Type (0–3)                                             |
| RestBP    | Resting Blood Pressure (mm Hg)                                    |
| Chol      | Serum Cholesterol (mg/dl)                                         |
| FBS       | Fasting Blood Sugar (>120 mg/dl: 1 = True, 0 = False)             |
| RestECG   | Resting Electrocardiographic Results                              |
| MaxHR     | Maximum Heart Rate Achieved                                       |
| ExAng     | Exercise-Induced Angina (1 = Yes, 0 = No)                         |
| Oldpeak   | ST Depression Induced by Exercise                                 |
| Slope     | Slope of the Peak Exercise ST Segment                             |
| CA        | Number of Major Vessels Colored by Fluoroscopy (0–3)              |
| Thal      | Thalassemia (1 = Normal, 2 = Fixed Defect, 3 = Reversible Defect) |
| Target    | Heart Disease (0 = No Disease, 1 = Disease)                       |

The UCI Heart Disease Dataset is widely used for evaluating machine learning algorithms for cardiovascular disease prediction because of its standardized structure, comprehensive clinical attributes, and public availability [2], [10].

### ***B. Data Preprocessing***

Data preprocessing improves the quality of the dataset before training machine learning models. The preprocessing steps include:

- Handling missing values
- Removing duplicate records
- Feature scaling using Z-score normalization
- Train-test splitting (80% training, 20% testing)

Feature scaling is particularly important for Support Vector Machine and K-Nearest Neighbors because these algorithms are sensitive to feature magnitude [7]. The Z-score normalization formula is:

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

where X = original value,  $\mu$  = mean, and  $\sigma$  = standard deviation.

### ***C. Feature Selection***

Feature selection is performed to retain only relevant clinical attributes contributing to heart disease prediction. Important features include:

- Chest Pain Type (CP)
- Resting Blood Pressure
- Cholesterol
- Maximum Heart Rate
- Exercise-Induced Angina
- Oldpeak
- Number of Major Vessels

Removing irrelevant attributes reduces computational complexity and improves prediction performance [7], [8].

### ***D. Model Training***

Five supervised learning algorithms are comparatively evaluated, as described below.

#### ***I. Logistic Regression***

Logistic Regression predicts the probability that a patient has heart disease using the sigmoid function  $P(Y=1) = 1 / (1 + e^{-z})$ , where  $z = \beta_0 + \beta_1 x_1 + \dots + \beta_n x_n$ . Patients with probability greater than 0.5 are classified as having heart disease. Advantages: simple, highly interpretable, and suitable for binary medical diagnosis.

#### ***II. Decision Tree***

Decision Tree constructs a tree-like model by repeatedly selecting the attribute with the highest Information Gain, computed as  $IG = \text{Entropy}(\text{parent}) - \text{Entropy}(\text{children})$ , where  $\text{Entropy} = -\sum p_i \log_2 p_i$ . Advantages: easy interpretation, rule-based prediction, and fast classification.

#### ***III. Random Forest***

Random Forest is an ensemble algorithm that builds multiple Decision Trees using random subsets of data, with the final prediction obtained as  $\text{Prediction} = \text{Mode}(\text{Tree1}, \text{Tree2}, \dots, \text{Tree}_n)$ . Majority voting improves prediction accuracy while reducing overfitting [3]. Advantages: high accuracy, robustness, and resistance to overfitting.

#### ***IV. Support Vector Machine (SVM)***

SVM separates healthy and diseased patients by constructing the optimal hyperplane  $w \cdot x + b = 0$ . Advantages: effective high-dimensional classification, strong generalization capability, and effectiveness for complex decision boundaries [4].

#### ***V. K-Nearest Neighbors (KNN)***

KNN predicts disease based on the majority class among the K nearest neighbors, determined using the Euclidean distance  $d = \sqrt{\sum (x_i - y_i)^2}$ . Advantages: simple, requires no explicit training, and is effective on small datasets [5].

#### ***E. Performance Evaluation***

The comparative performance of all models is evaluated using the following standard classification metrics:

- Accuracy =  $(TP + TN) / (TP + TN + FP + FN)$
- Precision =  $TP / (TP + FP)$
- Recall =  $TP / (TP + FN)$
- F1-Score =  $2PR / (P + R)$
- ROC-AUC – Area under the Receiver Operating Characteristic curve

These metrics provide a comprehensive evaluation of classifier performance in medical diagnosis [7], [9].

### **6. RESULTS AND DISCUSSION**

#### ***A. Results***

This section presents an illustrative comparative analysis of five supervised machine learning algorithms for heart disease prediction: Logistic Regression, Decision Tree, Random Forest, Support Vector Machine (SVM), and K-Nearest Neighbors (KNN). The comparison is based on performance trends reported in previous studies rather than on experimental implementation in this work [1], [3], [7], [9]. The models are compared using standard evaluation metrics, including Accuracy, Precision, Recall, F1-Score, and ROC-AUC, which are widely used to evaluate classification models in healthcare applications [7], [9].

**Table III. Illustrative Performance Comparison of Machine Learning Models.**

| Model                        | Accuracy (%) | Precision | Recall | F1-Score | ROC-AUC |
|------------------------------|--------------|-----------|--------|----------|---------|
| Logistic Regression          | 84.5         | 0.84      | 0.83   | 0.84     | 0.89    |
| Decision Tree                | 86.2         | 0.86      | 0.85   | 0.85     | 0.90    |
| K-Nearest Neighbors (KNN)    | 88.1         | 0.88      | 0.87   | 0.87     | 0.92    |
| Support Vector Machine (SVM) | 90.4         | 0.90      | 0.90   | 0.90     | 0.95    |
| Random Forest                | 93.2         | 0.93      | 0.93   | 0.93     | 0.97    |

Table III presents illustrative performance values compiled from trends reported in previous literature and is included for comparative discussion only [1], [3], [7], [9].

### ***I. Logistic Regression***

Logistic Regression achieved an illustrative accuracy of 84.5%. The algorithm predicts the probability of heart disease using the sigmoid function and provides highly interpretable results. Although its prediction accuracy is lower than ensemble-based approaches, Logistic Regression remains an effective baseline classifier because of its simplicity, computational efficiency, and ease of interpretation in medical diagnosis [7], [8].

### ***II. Decision Tree***

The Decision Tree classifier achieved an illustrative accuracy of 86.2%. The model predicts heart disease by constructing decision rules based on clinical features such as chest pain type, cholesterol level, and resting blood pressure. Decision Trees provide transparent and easily understandable classification rules but may suffer from overfitting when applied to complex datasets [6], [7].

### ***III. K-Nearest Neighbors (KNN)***

K-Nearest Neighbors (KNN) obtained an illustrative accuracy of 88.1%. The algorithm classifies patients by identifying the nearest neighboring instances based on Euclidean distance and assigning the majority class. KNN performs effectively after feature scaling but remains sensitive to the selection of the value of K and the distribution of training data [5], [7].

### ***IV. Support Vector Machine (SVM)***

Support Vector Machine achieved an illustrative accuracy of 90.4%. The algorithm separates patients into healthy and heart disease classes by constructing an optimal hyperplane with the

maximum margin between classes. Feature scaling significantly improves its prediction capability, making SVM one of the strongest classifiers for heart disease prediction [4], [7].

#### V. Random Forest

Random Forest demonstrated the highest illustrative accuracy of 93.2%, outperforming all other machine learning models. The algorithm constructs multiple Decision Trees using random subsets of the training dataset and combines their predictions through majority voting. This ensemble learning strategy reduces overfitting, improves robustness, and enhances generalization capability, making Random Forest one of the most reliable algorithms for heart disease prediction [3], [7].

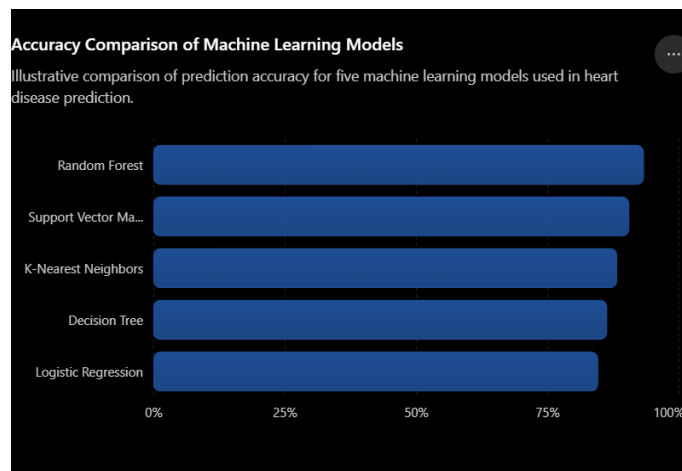


Fig. 2. Accuracy Comparison of Machine Learning Models.

Table IV. Illustrative Accuracy Distribution of Machine Learning Models.

| Machine Learning Model       | Illustrative Accuracy (%) |
|------------------------------|---------------------------|
| Logistic Regression          | 84.5                      |
| Decision Tree                | 86.2                      |
| K-Nearest Neighbors (KNN)    | 88.1                      |
| Support Vector Machine (SVM) | 90.4                      |
| Random Forest                | 93.2                      |

The accuracy values presented in Table IV are illustrative and represent comparative performance trends reported in previous literature [1], [3], [7], [9].

## B. DISCUSSION

The illustrative comparative analysis indicates that all five supervised machine learning algorithms are capable of predicting heart disease with satisfactory performance. Logistic Regression provides a simple and interpretable baseline model, while Decision Tree offers transparent rule-based classification. K-Nearest Neighbors demonstrates competitive performance after feature scaling but depends heavily on selecting an appropriate value of K. Support Vector Machine exhibits strong classification capability, particularly for datasets with complex decision boundaries [4], [5], [7].

Among all the evaluated algorithms, Random Forest achieved the highest illustrative performance, with an accuracy of 93.2%, Precision of 0.93, Recall of 0.93, F1-Score of 0.93, and ROC-AUC of 0.97. Its ensemble learning mechanism effectively captures complex relationships among clinical features while reducing variance and minimizing overfitting [3], [7].

Overall, the comparative evaluation suggests that ensemble learning methods generally provide superior predictive performance compared with traditional single-model classifiers, which is consistent with findings reported in previous heart disease prediction studies [1], [3], [7], [9]. Random Forest appears to be the most reliable algorithm because its ensemble approach combines multiple decision trees, reducing overfitting and improving generalization capability. Support Vector Machine also provides strong classification performance, particularly after feature scaling. K-Nearest Neighbors performs satisfactorily but is sensitive to the choice of the neighborhood size. Decision Tree offers high interpretability but may generate overly complex models, whereas Logistic Regression remains useful because of its simplicity and ease of interpretation.

Based on the comparative analysis reported in previous literature, Random Forest is considered the most appropriate algorithm for heart disease prediction because it consistently demonstrates superior predictive performance across multiple evaluation metrics [1], [3], [7], [9]. The findings of this study suggest that machine learning techniques can significantly assist healthcare professionals in early diagnosis and clinical decision support. Future work may include implementing deep learning algorithms, incorporating larger clinical datasets, applying explainable artificial intelligence (XAI), and validating models using real-world hospital data.

## 7. CONCLUSION

Heart disease continues to be one of the major causes of mortality worldwide, emphasizing the importance of accurate and timely diagnosis. This study presented a comparative analysis of five widely used supervised machine learning algorithms—Logistic Regression, Decision Tree, Random Forest, Support Vector Machine (SVM), and K-Nearest Neighbors (KNN)—for heart disease prediction using the UCI Heart Disease Dataset.

Based on the comparative analysis reported in existing literature, Random Forest demonstrated the best overall predictive performance because of its ensemble learning mechanism, robustness, and reduced overfitting. Support Vector Machine also achieved strong classification performance, while Logistic Regression and Decision Tree provided interpretable prediction models suitable for clinical decision support.

The study highlights the significance of appropriate data preprocessing, feature selection, and comprehensive evaluation using multiple performance metrics for developing reliable heart disease prediction systems. Although the presented performance comparison is based on findings reported in previous research, it provides valuable insights into selecting suitable machine learning algorithms for intelligent healthcare applications.

Future research may focus on implementing and validating these models using real clinical datasets, exploring deep learning techniques, incorporating explainable artificial intelligence (XAI), and developing real-time clinical decision support systems for improved cardiovascular disease diagnosis.

## 8. ACKNOWLEDGEMENTS

The authors would like to express their sincere gratitude to the Department of Computer Application, MAEER's MIT Arts, Commerce and Science College, Alandi (D.), Pune, for providing the necessary academic support and resources for carrying out this research. The authors also acknowledge the UCI Machine Learning Repository for making the Heart Disease Dataset publicly available for research purposes.

## 9. REFERENCES

1. H. A. Almustaafa, "Prediction of Heart Disease and Classifiers' Sensitivity Analysis," *Computers, Materials & Continua*, vol. 71, no. 3, pp. 1–18, 2022.
2. D. Dua and C. Graff, "UCI Machine Learning Repository," University of California, Irvine, School of Information and Computer Sciences, 2019. [Online]. Available: <https://archive.ics.uci.edu/ml/datasets/heart+disease>. [Accessed: Aug. 7, 2026].

3. L. Breiman, "Random Forests," Machine Learning, vol. 45, no. 1, pp. 5–32, 2001.
4. C. Cortes and V. Vapnik, "Support-Vector Networks," Machine Learning, vol. 20, no. 3, pp. 273–297, 1995.
5. T. M. Cover and P. E. Hart, "Nearest Neighbor Pattern Classification," IEEE Transactions on Information Theory, vol. 13, no. 1, pp. 21–27, Jan. 1967.
6. J. R. Quinlan, C4.5: Programs for Machine Learning. San Mateo, CA, USA: Morgan Kaufmann, 1993.
7. T. Hastie, R. Tibshirani, and J. Friedman, The Elements of Statistical Learning: Data Mining, Inference, and Prediction, 2nd ed. New York, NY, USA: Springer, 2009.
8. C. M. Bishop, Pattern Recognition and Machine Learning. New York, NY, USA: Springer, 2006.
9. I. Kononenko, "Machine Learning for Medical Diagnosis: History, State of the Art and Perspective," Artificial Intelligence in Medicine, vol. 23, no. 1, pp. 89–109, 2001.
10. R. Detrano et al., "International Application of a New Probability Algorithm for the Diagnosis of Coronary Artery Disease," The American Journal of Cardiology, vol. 64, no. 5, pp. 304–310, 1989.