

# Analyzing the major contribution factor for heart attack using ML algorithms and LIME

Bishal Ghimire<sup>1</sup>Manoj Rai<sup>2</sup>Nishant Uprety<sup>3</sup>Rubim Shrestha<sup>4</sup>

<sup>a</sup> b4shal@gmail.com, <sup>b</sup>manoj.kiranti@gmail.com, <sup>c</sup>me.nishant2002@gmail.com, <sup>d</sup>rubimshrestha@kcc.edu.np  
Bishal Ghimire, Kantipur City College, Kathmandu, 44600, Nepal  
Manoj Rai, Kantipur City College, Kathmandu, 44600, Nepal  
Nishant Uprety, Thapathali Engineering Campus, Kathmandu, 44600, Nepal  
Rubim Shrestha, Kantipur City College, Kathmandu, 44600, Nepal

---

## Abstract

Heart disease, a leading cause of global mortality, can significantly impact overall health. Timely prediction and identification of key risk factors are essential yet challenging. This study employs Machine Learning and Explainable Artificial Intelligence (XAI) to predict and identify major contributors to heart disease. Using a publicly available cardiovascular disease (CVD) dataset from Kaggle, consisting of 1319 samples with eight variables—age, gender, heart rate, systolic and diastolic blood pressure, blood sugar, CK-MB, and Test-Troponin—we trained and evaluated models for heart attack prediction. Supervised learning techniques, including SVM, Decision Tree, and Random Forest, were compared. The LIME technique was used to elucidate the influence of factors in different test cases. The study aimed to enhance classification and detection of key features leading to heart attacks. Model performance was assessed using metrics such as confusion matrix, accuracy, precision, recall, and F1 score. In the medical context, minimizing false negatives (missed actual positive cases) is critical. Random Forest emerged as the most effective model for reducing false negatives, making it a promising tool for heart disease prediction.

*Keywords: Accuracy; Cardiovascular Disease (CVD); Decision Tree; False Negatives; Random Forest; SVM*

---

## 1. Introduction

The heart, a vital organ responsible for pumping blood throughout the body, is susceptible to diseases that pose significant threats to human life. Cardiovascular diseases (CVDs), marked by symptoms such as breathlessness, weakness, and fatigue, are primarily caused by factors like high blood pressure, smoking, elevated cholesterol levels, and physical inactivity. Annually, approximately 18 million deaths are attributed to CVDs worldwide.

---

1  
2  
3  
4

To address this, a Computer-Aided Decision Support System utilizing data mining techniques has been developed to enhance the prediction accuracy of heart disease. This study employs advanced machine learning algorithms, including Support Vector Machines (SVM), Logistic Regression, Naive Bayes, k-Nearest Neighbors (KNN), and Decision Trees, to extract meaningful patterns from cardiovascular datasets. SVM and Logistic Regression are particularly noted for their high accuracy in predicting heart failure.

Recent advancements in machine learning highlight the effectiveness of discriminative classifiers such as SVM, Random Forest (RF), Logistic Regression (LR), Back Propagation Neural Network (BPNN), and Multilayer Perceptron (MLP) in detecting cardiac diseases. Additionally, hybrid models combining techniques like Random Forest with linear models, Multilayer Perceptron, Bayes Net, and Naive Bayes have shown promising results in improving prediction accuracy.

This study focuses on predicting heart attacks using RF and SVM, employing the LIME technique for factor identification. Explainable Artificial Intelligence (XAI) is used to explore key features through the LIME algorithm. The analysis is conducted on a publicly available Kaggle dataset, utilizing Python on Google Colab for computation, preprocessing, and visualization.

### 1.1. *Statement of Problem*

Cardiovascular diseases (CVDs), or heart conditions, often manifest through symptoms such as chest pain, difficulty breathing, fatigue, and palpitations. Despite advances in medical research, heart attacks remain a leading cause of global mortality. Although age, gender, and lifestyle are known risk factors, the complex interplay and relative importance of these factors in contributing to heart attacks require further investigation.

Deep learning (DL) and machine learning (ML) methodologies have been explored to aid medical practitioners in identifying patterns indicative of heart attacks. These approaches aim to enhance diagnostic capabilities in cardiology. However, challenges persist in determining the major contributing factors and ensuring the explainability of these models. Enhancing the interpretability of ML models is crucial for a comprehensive understanding of the underlying factors leading to heart attacks, thereby facilitating more informed clinical decision-making.

### 1.2. *Research Objectives*

- To evaluate and compare the predictive performance of machine learning algorithms, specifically Support Vector Machine (SVM), Decision Tree, and Random Forest, in determining the likelihood of a heart attack based on various input parameters.
- To analyze the major contributing factors for heart attacks using the LIME technique.

### 1.3. *Research Questions*

- How do machine learning algorithms, specifically Support Vector Machine (SVM), Decision Tree, and Random Forest, compare in predicting the likelihood of heart attacks, and what are the major contributing factors as identified by the LIME technique?

## 2. Literature Review

Krishnaiah et al. (2016) highlights the transformative impact of Computer-Aided Decision Support Systems and data mining techniques, particularly Fuzzy Intelligent Techniques, in enhancing heart disease prediction accuracy. Amin et al. (2018) This research identifies crucial features and optimal data mining techniques for heart disease prediction, achieving 87.4% accuracy with a hybrid Vote technique. Haq et al. (2018) Introducing machine learning for noninvasive heart disease diagnosis, this study uses seven algorithms and three feature selection methods, achieving high classification accuracy and reliability. Jan et al. (2018) This study proposes an Ensemble model combining five classifiers, demonstrating superior predictive accuracy and reliable diagnostic performance for cardiovascular disease recurrence. Sharmila et al. (2018) A Big Data-based model using Hadoop and MapReduce with Support Vector Machine (SVM) is proposed to enhance heart disease prediction accuracy amid varying data challenges. Alotaibi et al. (2019) Employing multiple machine learning approaches, this study enhances heart failure prediction accuracy using the UCI heart disease dataset, surpassing previous accuracy scores. Uddin et al. (2019) Reviewing 48 articles, this study finds that Random Forest and SVM algorithms lead in accuracy for disease prediction, highlighting trends in supervised machine learning applications. Motarwar et al. (2020) A machine learning framework integrating five algorithms, including Random Forest and SVM, achieves the highest accuracy in predicting heart disease using the Cleveland dataset. Pires et al. (2020) This analysis demonstrates that Decision Tree and Support Vector Machine achieve the best results, with 87.69% accuracy in both 20-fold and 10-fold cross-validation. Shah et al. (2020) This study employs data mining and supervised learning algorithms, with K-nearest neighbor achieving the highest accuracy score for predicting heart disease in the UCI dataset. Alqahtani et al. (2022) This study employs an ensemble approach combining machine learning and deep learning models, achieving 88.70% accuracy in predicting cardiovascular disease using a public dataset. Drożdż et al. (2022) Exploring the link between metabolic-associated fatty liver disease and cardiovascular risk, this study uses logistic regression classifiers to achieve high performance in identifying high-risk patients based on key clinical variables. Hasan et al. (2020) Comparing feature selection algorithms, this study finds Random Forest and Support Vector Classifier to be the most accurate for cardiovascular disease prediction. Mayank Agrawal et al. (2023) This study compares Random Forest, Decision Tree, and AdaBoost for chronic kidney disease prediction, with Random Forest achieving the highest accuracy of 99.98%. Rasheed Omobolaji Alabi et al. (2023) Using LIME and SHAP techniques, this study develops a prognostic system for Nasopharyngeal Cancer, achieving 85.9% accuracy and identifying key factors influencing patient survival.

## 3. Research Methodology

### 1.4. Data Collection

The dataset used for this research comprises of 1319 samples, each with nine fields. Eight fields are input features: age, gender, heart rate, systolic blood pressure, diastolic blood pressure, blood sugar, CK-MB, and Test-Troponin. The ninth field is the output, indicating the presence or absence of a heart attack, categorized as "negative" (no heart attack) or "positive" (heart attack). This dataset facilitates the classification task of predicting heart attack likelihood based on the input features.

### 1.5. Data Preprocessing

The preprocessing steps included dropping the target column from the categorical data and assigning binary labels to each value. Correlation analysis was performed to identify similarities and dissimilarities among the categorical data. Missing values were filled by calculating and assigning the median of the respective fields. Once missing values were addressed, the dataset was split into training and testing sets with a ratio of 4:5 for

training (80%) and 1:5 for testing (20%). The training set was used to train the model, while the testing set evaluated the model's generalization to unseen data, ensuring accurate assessment of predictive capabilities and overall performance.

#### 1.6. *Model Development*

In this phase, several machine learning algorithms, including Support Vector Machine (SVM), Random Forest, and Decision Tree (DT), were utilized. Each algorithm was paired with the Explainable AI approach, Local Interpretable Model-agnostic Explanations (LIME). LIME approximates any black box machine learning model with a local, interpretable model to explain individual predictions. This integration enhances the transparency of complex models, providing insights into the reasoning behind each prediction, and contributes to a more understandable and interpretable machine learning system.

#### 1.7. *Local Interpretable Model-agnostic Explanations (LIME)*

In our experimentation, we generated 150 perturbations using random ones and zeros, forming a matrix with perturbations as rows and superpixels as columns. Each superpixel was either ON (one) or OFF (zero), corresponding to the number of superpixels in the image. The test image underwent perturbations based on this vector and predefined superpixels.

Ensuring the understandability of AI models is crucial for building user trust. AI interpretability reveals the inner workings of these systems, aiding in identifying issues such as causality, information leakage, model bias, and robustness. LIME (Local Interpretable Model-agnostic Explanations) offers a framework to interpret black-box models, elucidating the reasoning behind AI predictions or suggestions.

Mathematically, local surrogate models with interpretability constraints are expressed as:

$$\text{Explanation}(x)_n = \arg \min L(F, g, \pi_x + \Omega(g))$$

The training process for local surrogate models involves:

1. Selecting the instance of interest for which an explanation is required.
2. Perturbing the dataset and obtaining black-box predictions for these new points.
3. Weighting the new samples based on their proximity to the instance of interest.
4. Training a weighted, interpretable model on the perturbed dataset.
5. Explaining the prediction by interpreting the local model.

This approach creates interpretable models, shedding light on the decision-making process of complex AI systems, thus enhancing transparency and user trust.

## 4. **Results and Analysis**

For Class 0, the SVM algorithm achieved an accuracy of 92.05%, with precision, recall, and F1-score values of 87.74%, 92.08%, and 89.86%, respectively. The Decision Tree and Random Forest algorithms demonstrated higher performance with accuracies of 97.73% and 98.11%, respectively. In Class 1, all algorithms exhibited strong performance, with the Random Forest algorithm achieving the highest accuracy of 98.11% and precision, recall, and F1-score values of 98.17%, 98.77%, and 98.47%, respectively.

The models' accuracy was tested using cross-validation, revealing variations in performance based on the data subsets from each fold. SVM and Decision Tree showed the most variation in their performance metrics across different folds.

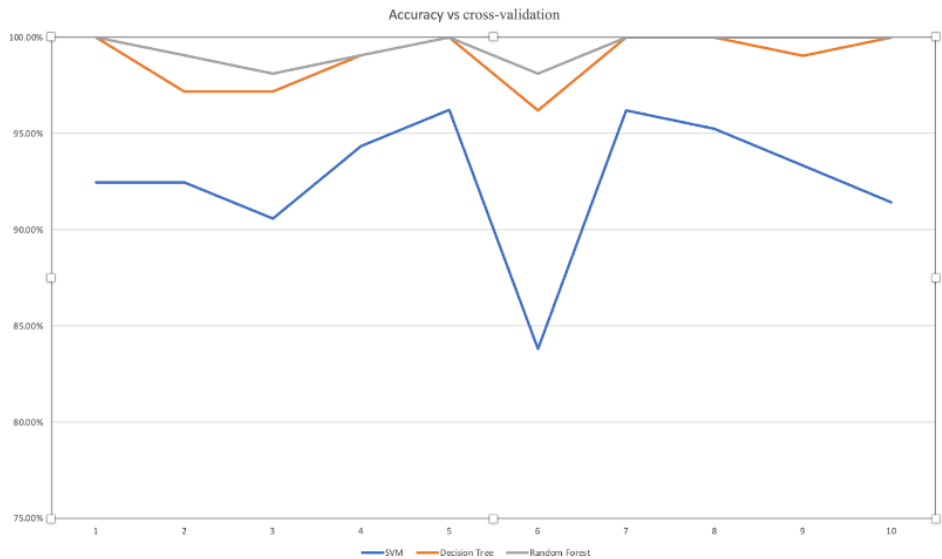


Figure 1: Cross Validation for different Machine Learning Models

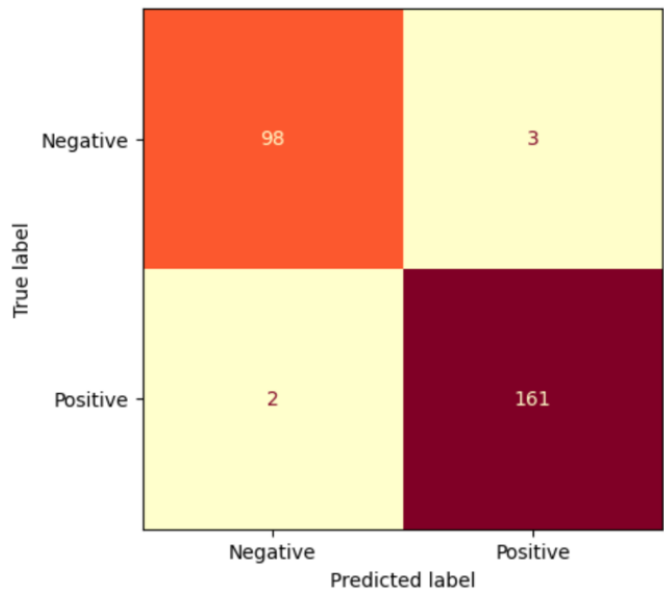


Figure 2: Confusion Matrix for Random Forest Model

The area under the precision-recall curve (PR curve) indicates how well a model balances precision and recall. Among the Random Forest, Decision Tree, and SVM models, Random Forest exhibited the highest area under the curve, indicating superior performance. The Random Forest model achieved a better trade-off between precision and recall, making it the preferred choice. The higher area under the PR curve for Random Forest

signifies its enhanced ability to maintain precision while effectively identifying positive cases, demonstrating robustness and balance compared to Decision Tree and SVM.

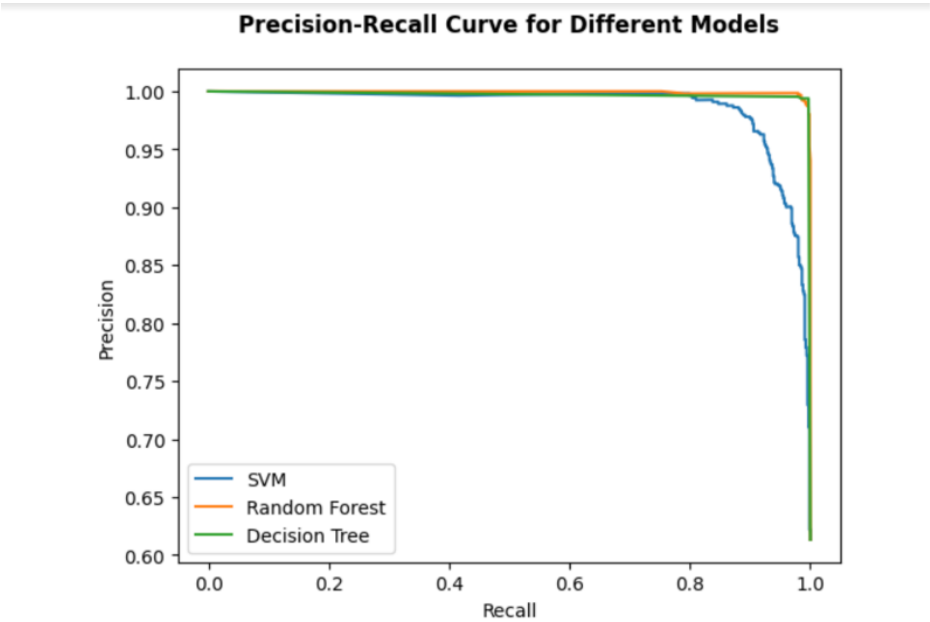


Figure 3: Precision Recall Curve for Different Models

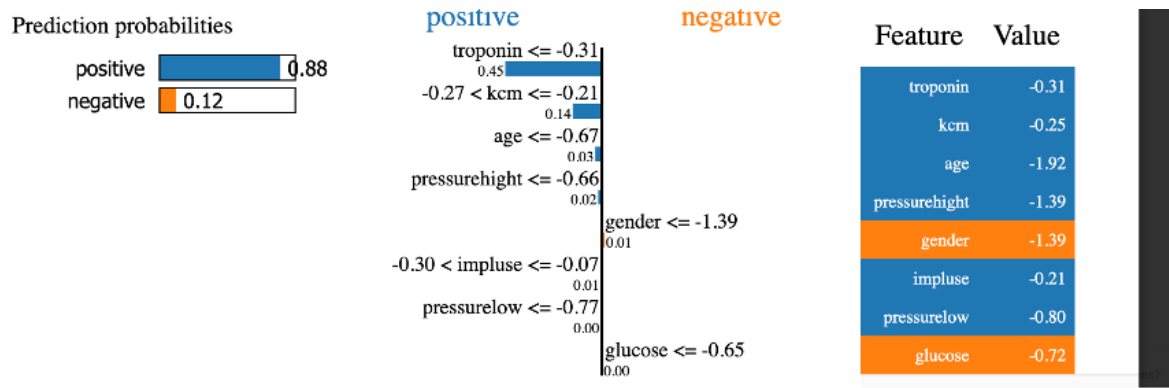


Figure 4: LIME Explanation for Random Forest Model

### Features pushing the prediction towards positive (blue color):

- **troponin:** This feature with value -0.31, with a weight of 0.45, has the strongest positive influence on the positive prediction.
- **kcm:** The kcm of -0.25 also contributes to the positive prediction with a weight of 0.14.
- **age:** The age of -1.92 also contributes to the positive prediction with a small weight of 0.03.
- **pressurehigh:** A high pressure reading of -1.39 adds slightly to the positive prediction weight of 0.02.
- **impulse:** A impulse reading of -0.21 adds slightly to the positive prediction weight negligibly.
- **pressurelow:** A low pressure reading of -0.80 adds slightly to the positive prediction weight of 0.02.

### Features pushing the prediction towards negative (red color):

- **gender:** A gender reading of -1.39 adds slightly to the negative prediction weight of 0.01
- **glucose:** A glucose reading of -0.72 adds slightly to the negative prediction weight negligibly.

## 5. Conclusion

In the medical domain, particularly where misclassifying an unwell patient as healthy has significant consequences, selecting a model with high recall is crucial. Recall, or sensitivity, measures the model's ability to correctly identify all positive instances. Precision and recall are pivotal in evaluating a model's effectiveness, with a preference for high recall to reduce false negatives—instances where the model fails to recognize actual positive cases (ill patients).

Upon assessing the decision tree and random forest models, the random forest consistently demonstrated a higher average recall compared to the decision tree. Additionally, the random forest model outperformed the decision tree across various performance metrics, including precision and accuracy. Most notably, the random forest model exhibited a lower number of false negatives, which is critical in medical applications.

Overall, the analysis suggests that the random forest is a more robust and effective model, particularly in minimizing false negatives and enhancing predictive performance. This underscores the importance of evaluating models using multiple metrics to gain a comprehensive understanding of their effectiveness in medical contexts.

## Acknowledgements

I would like to extend my heartfelt appreciation to my exceptional supervisor, Rubim Shrestha, whose unwavering support and encouragement inspired me to pursue this paper. His invaluable guidance and advice were instrumental in making this achievement possible. I am also deeply grateful to Dr. Gajendra Sharma for his continuous feedback and support. Furthermore, I wish to thank my parents and sister for their moral support, and my friend Nishant Uprety for his continuous assistance and expertise throughout the research timeframe.

## References

- [1] V. Krishnaiah, G. Narsimha, and N. Subhash, "Heart Disease Prediction System using Data Mining Techniques and Intelligent Fuzzy Approach: A Review," *International Journal of Computer Applications*, vol. 136, pp. 43–51, Feb. 2016, doi: 10.5120/ijca2016908409.
- [2] M. Amin, Y. K. Chiam, and K. Varathan, "Identification of significant features and data mining techniques in predicting heart disease," *Telematics and Informatics*, vol. 36, Nov. 2018, doi: 10.1016/j.tele.2018.11.007.
- [3] A. U. Haq, J. P. Li, M. H. Memon, S. Nazir, and R. Sun, "A Hybrid Intelligent System Framework for the Prediction of Heart Disease Using Machine Learning Algorithms," *Mobile Information Systems*, vol. 2018, p. e3860146, Dec. 2018, doi: 10.1155/2018/3860146.
- [4] M. Jan, A. A. Awan, M. S. Khalid, and S. Nisar, "Ensemble approach for developing a smart heart disease prediction system using classification algorithms," *RRCC*, vol. 9, pp. 33–45, Dec. 2018, doi: 10.2147/RRCC.S172035.
- [5] R. Sharmila and S. Chellammal, "A conceptual method to enhance the prediction of heart diseases using big data Techniques," *International Journal of Computer Sciences and Engineering*, 2018.
- [6] F. S. Alotaibi, "Implementation of Machine Learning Model to Predict Heart Failure Disease," *International Journal of Advanced Computer Science and Applications (IJACSA)*, vol. 10, no. 6, 2019, doi: 10.14569/IJACSA.2019.0100637.
- [7] S. Uddin, A. Khan, M. E. Hossain, and M. A. Moni, "Comparing different supervised machine learning algorithms for disease prediction," *BMC Med Inform Decis Mak*, vol. 19, no. 1, p. 281, Dec. 2019, doi: 10.1186/s12911-019-1004-8.
- [8] P. Motarwar, A. Duraphe, G. Suganya, and M. Premalatha, "Cognitive Approach for Heart Disease Prediction using Machine Learning," in *2020 International Conference on Emerging Trends in Information Technology and Engineering (ic-ETITE)*, Feb. 2020, pp. 1–5. doi: 10.1109/ic-ETITE47903.2020.242.
- [9] I. M. Pires, G. Marques, N. M. Garcia, and V. Ponciano, "Machine learning for the evaluation of the presence of heart disease," *Procedia Computer Science*, vol. 177, pp. 432–437, Jan. 2020, doi: 10.1016/j.procs.2020.10.058.
- [10] D. Shah, S. Patel, and D. Bharti, "Heart Disease Prediction using Machine Learning Techniques," *SN Computer Science*, vol. 1, Nov. 2020, doi: 10.1007/s42979-020-00365-y.
- [11] A. Alqahtani, S. Alsubai, M. Sha, L. Vilcekova, and T. Javed, "Cardiovascular Disease Detection using Ensemble Learning," *Computational Intelligence and Neuroscience*, vol. 2022, p. e5267498, Aug. 2022, doi: 10.1155/2022/5267498.
- [12] K. Drożdż et al., "Risk factors for cardiovascular disease in patients with metabolic-associated fatty liver disease: a machine learning approach," *Cardiovasc Diabetol*, vol. 21, no. 1, Nov. 2022, doi: 10.1186/s12933-022-01672-9.
- [13] N. Hasan and Y. Bao, "Comparing different feature selection algorithms for cardiovascular disease prediction," *Health and Technology*, vol. 11, Oct. 2020, doi: 10.1007/s12553-020-00499-2.
- [14] M. Agrawal and V. Jain, "Chronic Kidney Disease Prediction Using Random Forest, Decision Tree and Ada Boost Classifier," in *2023 4th International Conference on Smart Electronics and Communication (ICOSEC)*, Sep. 2023, doi: 10.1109/icosec58147.2023.10276324.
- [15] R. O. Alabi, M. Elmusrati, I. Leivo, A. Almangush, and A. Mäkitie, "Machine learning explainability in nasopharyngeal cancer survival using LIME and SHAP," *Scientific Reports*, vol. 13, no. 1, Jun. 2023, doi: 10.1038/s41598-023-35795-0.