

RANDOM FOREST-BASED CARDIOVASCULAR DISEASE PREDICTION WITH SHAP-DRIVEN INTERPRETABILITY

Farrel Rafa Akbar¹, **Dina Tri Utari**^{2*}

^{1,2}Department of Statistics, Faculty of Mathematics and Natural Sciences, Universitas Islam Indonesia
Jln. Kaliurang Km 14.5, Sleman, Yogyakarta, 55584, Indonesia

Corresponding author's e-mail: * dina.t.utari@uii.ac.id

Article Info

Article History:

Received: 27th November 2025

Revised: 10th May 2026

Accepted: 06th June 2026

Available online: 01st July 2026

Keywords:

Cardiovascular disease prediction;
Clinical decision support systems;
Imbalanced data handling;
Model interpretability;
Random forest algorithm;
SHAP;
SMOTE.

ABSTRACT

Cardiovascular disease continues to be a significant worldwide health issue, where early detection is essential due to the frequently asymptomatic beginning of heart attacks. This research presents a model for predicting the risk of heart disease utilizing the Random Forest (RF) algorithm, trained on clinical data obtained from Zheen Hospital in Erbil, Iran. The data can be found online through the Mendeley Data website. The Synthetic Minority Oversampling Technique (SMOTE) was used to fix the problem of uneven class sizes, and Shapley Additive Explanations (SHAP) were used to explain how the model made its predictions. The RF model, improved with the best settings and evaluated with the F1-score, achieved impressive results, which are more than 99% for training, validation, and testing data. These results underscore its capacity to discover minority class patterns, crucial for recognizing rare yet significant occurrences. SHAP analysis identified troponin, creatine kinase-MB, and age as the primary predictors. The new idea is not just about individual methods but how they work together for predicting cardiovascular disease, especially by making AI easier to understand and dealing with uneven data using real clinical information. The subsequent study will aim to enhance robustness and generalizability across diverse patient populations.



This article is an open access article distributed under the terms and conditions of the [Creative Commons Attribution-ShareAlike 4.0 International License](https://creativecommons.org/licenses/by-sa/4.0/) (<https://creativecommons.org/licenses/by-sa/4.0/>).

How to cite this article:

F. R. Akbar and D. T. Utari., "RANDOM FOREST-BASED CARDIOVASCULAR DISEASE PREDICTION WITH SHAP-DRIVEN INTERPRETABILITY", *BAREKENG: J. Math. & App.*, vol. 20, no. 4, pp. 3545-3556, Dec, 2026.

Copyright © 2026 Author(s)

Journal homepage: <https://ojs3.unpatti.ac.id/index.php/barekeng/>

Journal e-mail: barekeng.math@yahoo.com; barekeng.journal@mail.unpatti.ac.id

Research Article · **Open Access**

1. INTRODUCTION

Cardiovascular disease remains a predominant global health challenge, accounting for approximately 17.9 million fatalities annually, as reported by the World Health Organization. In nations such as Indonesia, where early screening and healthcare infrastructure may be scarce, the imperative for timely and precise risk assessment becomes exceedingly critical. Recognizing high-risk individuals before clinical symptoms appear could markedly diminish mortality rates and enhance clinical outcomes. In this situation, machine learning has become a powerful tool for predicting health risks, as it can analyze complex clinical data and provide accurate risk predictions. These forecasts can guide healthcare providers in making informed decisions, prioritizing interventions, and allocating resources effectively. As a result, integrating machine learning into healthcare systems not only holds the promise of improved patient outcomes but also fosters a more proactive approach to disease management and prevention.

The Random Forest model is gaining a lot of attention in healthcare because it can handle complex relationships between different factors while still being accurate and not easily overfitting the data [1], [2]. Nonetheless, a persistent drawback of RF and analogous ensemble methodologies is their inherent lack of interpretability, commonly referred to as the black box phenomenon. In clinical settings, where decisions need to be explained and clear, not understanding why a model makes certain predictions can reduce trust and slow down its use [3].

Across the investigations, a common conclusion emerges machine learning models, particularly ensemble approaches like Random Forest, are outstanding in predicting cardiac disease. Employing varied datasets, these models regularly surpass individual classifiers for prediction accuracy [4], [5], [6], [7]. A study revealed that a hyperparameter-optimized Random Forest model attained elevated accuracy and AUC scores, underscoring its robust predictive capability [5]. A separate investigation revealed that the suggested PCA-BHHO method achieved superior accuracy on the Statlog dataset, exceeding that of the Boruta-RF model [4]. A comprehensive evaluation highlighted that, although numerous classical machine learning algorithms are practical, ensemble and deep learning techniques provide distinct benefits in enhancing prediction performance beyond conventional models [6]. Feature importance analyses indicated several factors as the most significant predictors in optimum Random Forest models [5]. These findings collectively highlight the efficacy of Random Forest and other ensemble methods as effective instruments for precise heart disease prediction.

To address this issue, researchers have started using SHAP, a clear method based on cooperative game theory that measures how much each feature contributes to a model's prediction. SHAP helps break down complex model results into simpler parts, making it easier to see which factors mainly affected a specific patient's predicted risk [8], [9]. Such transparency is paramount in healthcare environments where clinicians must not only depend on predictive analytics but also grasp the underlying rationale to guide patient dialogues, treatment decisions, and risk communication. While SHAP is increasingly prevalent in machine learning studies, its application to Random Forest models for cardiovascular disease prediction remains incomplete. Most extant research efforts have concentrated on enhancing model accuracy without fully harnessing interpretability tools to yield clinical insights. Also, the connection between the importance of features shown by SHAP and important demographic or clinical factors, such as age, gender, troponin levels, and CK-MB has not been fully studied. This results in a deficiency in understanding how predictive features fluctuate across diverse patient profiles and how such knowledge could be harnessed to bolster personalized treatment methodologies. Moreover, healthcare datasets frequently exhibit imbalance, with a significantly lower prevalence of positive cases (e.g., confirmed heart disease) compared to negative cases. This disparity can skew models toward the majority class, thereby diminishing sensitivity in the identification of high-risk patients. To solve this problem, methods like the SMOTE are used to create extra examples for the smaller group, which helps the model learn better from less common data. Even though SMOTE is commonly used, there is not enough information on how it works together with SHAP and RF in healthcare, especially regarding its effects on performance and understanding.

This study aims to fill the gaps by creating a combined predictive system that uses Random Forest, SMOTE, and SHAP methods to predict the risk of heart disease based on clinical data. This study is a novel contribution to the field of medical decision support, as no other research has integrated the three methods, Random Forest, SMOTE, and SHAP, on this clinical dataset or within this clinical setting. Employing a real-world dataset procured from Zheen Hospital in Erbil, Iraq, the research endeavors to construct a highly precise Random Forest model for heart disease prediction, implement SMOTE to rectify class imbalance, and

leverage SHAP to produce interpretable feature attributions that can be effectively communicated to healthcare professionals. Emphasis is placed on the influence of biochemical markers (e.g., troponin, CK-MB) and demographic variables (e.g., age, gender) in informing the model's decision-making processes. By focusing on predictive performance and interpretability, this study presents a clear and effective machine learning system for assessing cardiovascular risk. The findings are expected to enhance clinical decision-support tools in healthcare and to establish a foundation for future research examining system effectiveness, reliability, and human-centered design in medical applications.

2. RESEARCH METHODS

2.1 Random Forest Algorithm

Random Forest is a robust ensemble learning technique extensively employed for classification and regression tasks. Introduced by Breiman in 2001, it operates by constructing multiple decision trees and aggregating their outputs to improve accuracy and control overfitting [10]. The foundational premise of Random Forest is to leverage the concept of bagging (bootstrapped aggregation), where each tree is trained on a random subset of the training data, leading to diverse models that collectively enhance performance [11]. This aggregation reduces variance, thereby yielding more robust predictions compared to individual trees, which tend to be more susceptible to data noise [12].

The efficacy of Random Forests can be attributed to their capability to handle high-dimensional feature spaces, efficiently managing situations where the number of features exceeds the number of observations [11]. Moreover, the model is versatile and easily applicable to various domains, including biology, finance, and ecology [13], [14].

The Random Forest algorithm creates an ensemble of decision trees through a stochastic process, later amalgamating the predictions from these trees to produce the outcome [15]. Depending on whether the task is classification or regression, the Random Forest algorithm typically determines the outcome by a majority vote or by averaging the predictions. This approach enhances the model's accuracy and robustness, reducing the risk of overfitting that can occur with individual decision trees. Fig. 1 illustrates how Random Forest works.

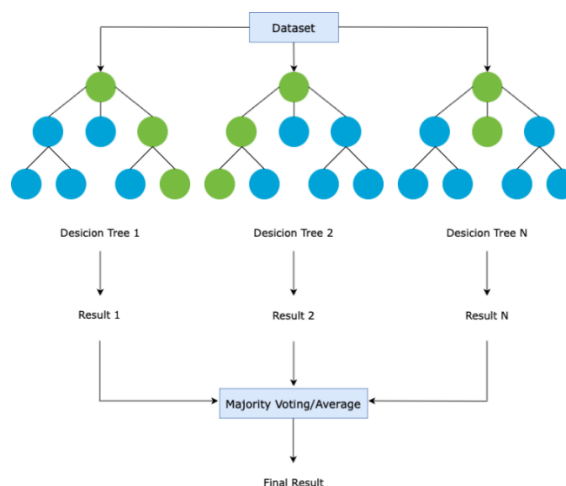


Figure 1. Random Forest Illustration

Let $D = \{(x_1, y_1), \dots, (x_N, y_N)\}$ represent the training dataset, where $x_i = (x_1, x_2, \dots, x_p)^T$. Then for j ranging from 1 to J :

1. Obtain a bootstrap sample, D_j , of size N from D .
2. Use the bootstrap sample D_j as the training dataset to build a decision tree by repeatedly splitting the data into two groups based on certain rules until you reach a stopping point. Start with all the data in one group.
 - a) Begin with all observations consolidated in a single node.

- b) Iteratively execute the subsequent steps for each unsplit node until the termination requirement is satisfied:
 - i. Select m predictors at random from the p available predictors.
 - ii. Identify the optimal binary partition among all binary partitions on the m predictors from step i.
 - iii. Divide the node into two child nodes utilizing the division from step ii [16].

For new data x , the prediction is determined based on the average or aggregation of the results from all trees.

$$\hat{f}(x) = \arg \max_y \sum_{j=1}^J I(\hat{h}_j(x) = y). \quad (1)$$

The result is the class most predicted by all trees (I is an indicator function that equals 1 if the j -th tree predicts class y and 0 otherwise).

The loss function $L(Y, f(X))$ quantifies the proximity of the projected outcome to the actual value of Y . This function imposes a penalty when the projected outcome significantly deviates from the actual value. A frequently employed loss function is zero-one loss:

$$L(Y, f(X)) = \begin{cases} 0, & \text{if } Y = f(X) \\ 1, & \text{if } Y \neq f(X) \end{cases}. \quad (2)$$

The optimal prediction function is the class with the highest probability ($P(Y = y|X = x)$) called Bayes' rule.

$$f(x) = \arg \max_{y \in Y} P(Y = y | X = x). \quad (3)$$

The Gini index for split voting uses the Eq. (4):

$$Q = \sum_{k \neq k'} \hat{p}_k \hat{p}_{k'}. \quad (4)$$

The proportion of data from k class in a node ($\hat{p}_k$) is calculated as:

$$\hat{p}_k = \frac{1}{n} \sum_{i=1}^n I(y_i = k). \quad (5)$$

Besides the Gini index, decision tree split selection may also employ entropy as a metric for node uncertainty. Entropy quantifies the irregularity in class distribution, where lower values signify more homogeneous nodes [17]. The entropy value is determined using the Eq. (6).

$$Entropy(N) = \sum_{i=1}^p -P_i * \log_2 P_i \quad (6)$$

Following the computation of class proportions within a node ($\hat{p}_k$), the recursive binary partitioning method for constructing a decision tree can be succinctly outlined as follows:

Let the training dataset be,

$$D = \{(x_1, y_1), \dots, (x_N, y_N)\} \text{ with } x_i = (x_{i1}, \dots, x_{ip})^T.$$

Algorithm:

1. Initialization: Place all observations $(x_1, y_1), \dots, (x_N, y_N)$ into a single root node.
2. Recursive splitting: Repeat for each node that is not yet terminal until a halting condition is satisfied:
 - a) Identify the binary split that results in the most significant reduction in impurity (or the highest improvement according to the selected criterion) among all possible binary splits across the p predictors.
 - b) Utilize this optimal split to partition the node into two offspring nodes.
3. To predict the response at a new point x :

- a) Cross the tree from the root, following the binary splits, until x falls into a terminal node k .
- b) Let the training responses in node k be $\{y_{k1}, \dots, y_{kn}\}$.
- c) The predicted response in classification task is given by the majority class (or class with maximum $\hat{p}k$).

2.2 SHapley Additive exPlanations (SHAP)

SHapley Additive Explanations (SHAP) is a method that is particularly influential for the interpretation of the predictions of machine learning models, particularly those that are classified as black boxes, such as random forests and gradient boosting. SHAP, which was created by Lundberg and Lee in 2017 [18], integrates machine learning principles with game theory to establish a unified framework for feature attribution across a variety of algorithms [19]. The fundamental concept of SHAP is to assess the relative contribution of each feature to the prediction of an instance by computing marginal contributions across all feasible feature combinations. This guarantees that the attributions are both globally consistent and locally accurate [20], [21].

SHAP's adaptability in simplifying intricate models into comprehensible insights is notably evident in a variety of applications, including environmental modeling and medical predictions. For example, SHAP has been implemented effectively to interpret machine learning models that predict ICU mortality in patients with sepsis-associated encephalopathy, thereby facilitating clinicians' comprehension of the features that have the most significant impact on model outputs [22]. In the same vein, it has been employed to analyze patient data to identify socio-reproductive factors that influence the risk of cancer in Bangladeshi women, demonstrating its adaptability to a variety of healthcare contexts [23].

Lloyd Shapley developed a mathematical formula for the values ϕ_j that satisfy these properties and demonstrated that these values are the sole ones capable of satisfying all of them [24]. These values are referred to as Shapley values.

Let N denote the complete set of features and let $j \in N$ be a particular feature. The mathematical formula for ϕ_j is as follows:

$$\phi_j = \sum_{S \subseteq N \setminus \{j\}} \frac{|S|! (|N| - |S| - 1)!}{|N|!} [f(S \cup \{j\}) - f(S)], \quad (7)$$

where:

$S \subseteq N \setminus \{j\}$: subset of features excluding j .
 $f(S)$: model prediction given only the features in S .
 $|S|$: cardinality of subset S .
 $|N|$: total number of features.

Shapley values for machine learning interpretability were first suggested by Štrumbelj and Kononenko in 2010 [25]. The authors demonstrated the application of Shapley values to fairly assess the contribution of features in a machine learning model. Lundberg et al. in 2017 subsequently popularized and expanded the method into what is known as SHAP analysis [18].

A fundamental aspect of SHAP values that is essential for their interpretation may be inferred. Consider the prediction for a sample i , $f(i)$. The SHAP values for all features j of sample i , denoted as $\phi_{i,j}$, using the formula:

$$f(i) = \phi_0 + \sum_{\text{features}} \phi_{i,j}, \quad (8)$$

where $f(i)$ represents model prediction for sample i , ϕ_0 is average prediction, and $\sum_{\text{features}} \phi_{i,j}$ declares sum of all SHAP values for sample i .

This implies that, for a given sample i , the SHAP values associated with each feature quantify the contribution of that feature to the deviation of the sample's prediction from the model's average (baseline) prediction [26].

SHAP calculation results can be illustrated using graphs such as force plots and summary plots, which elucidate the contribution of features to model predictions. Force plots specifically demonstrate the extent to

which each variable influences the predicted value relative to the model's baseline, with positive features augmenting the forecast and negative features diminishing it. This visualization facilitates user comprehension of feature influence on prediction outcomes, enhances model analysis and interpretation, and elucidates the factors affecting model decisions, particularly in significant contexts such as medical diagnoses or other critical determinations [27].

3. RESULTS AND DISCUSSION

This study employs the Heart Attack Dataset gathered at Zheen Hospital in Erbil, Iraq, from January to May 2019 [28]. The dataset comprises the following attributes: age, gender, heart rate, systolic blood pressure, diastolic blood pressure, blood sugar, CK-MB, and troponin levels, with the outcome variable denoting the occurrence or absence of a heart attack. Data preprocessing entailed the normalization of specific features: gender was encoded as one for male and zero for female; blood sugar was encoded as one if glucose levels surpassed 120 mg/dL and zero otherwise; and the outcome variable was denoted as 1 for positive (heart attack) and 0 for negative (no heart attack).

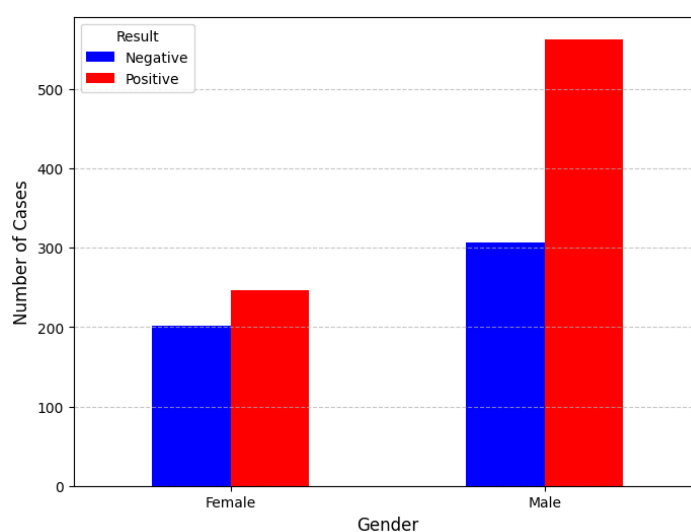


Figure 2. Distribution of Heart Attack Cases by Gender

The analysis of heart attack outcomes by gender in Fig. 2 uncovers numerous significant findings. Males exhibit a much greater prevalence of positive cases (563) than females (247), suggesting that men in this sample are over twice as likely to suffer a heart attack. Although males exceed females in negative situations (307 compared to 202), the gender disparity is significantly more prominent in the positive class. In terms of proportions, roughly 55% of female cases are positive, while this figure increases to around 65% for males, indicating a greater relative risk for men. These findings emphasize gender as a potentially significant variable in heart disease prediction and stress the necessity for predictive models to address dataset imbalance, as the predominance of males and favorable outcomes could skew model efficacy.

Table 1. Statistical Summary

Variables	Statistical Summary			
	Mean	Std	Min	Max
Age	56.192	13.647	14.000	103.000
Heart Rate	78.337	51.603	20.000	1111.000
Systolic	127.171	26.123	42.000	223.000
Diastolic	72.269	14.034	38.000	154.000
Blood Sugar	146.634	74.923	35.000	541.000
CK-MB	15.274	46.327	0.321	300.000
Troponin Levels	0.361	1.155	0.001	10.300

The statistical overview in Table 1 underscores numerous significant observations. The mean age of patients is 56 years ($SD \approx 13.6$), with a broad age range from 14 to 103, indicating the dataset encompasses both younger and older patients, while predominantly featuring middle-aged individuals. The heart rate has an atypically elevated standard deviation ($SD = 51.6$) and a maximum value of 1111, suggesting the potential

existence of outliers or data collection inaccuracies, as these figures are physiologically improbable. Blood pressure readings are within anticipated parameters: the average systolic pressure is approximately 127 mmHg, and the diastolic pressure is around 72 mmHg, but peak values (223/154 mmHg) indicate instances of severe hypertension. The mean blood sugar level is 146 mg/dL, surpassing the typical fasting threshold of approximately 126 mg/dL, indicating that a significant number of patients may have diabetes or impaired glucose regulation; the extensive range of 35–541 mg/dL further illustrates considerable variability. Cardiac biomarkers indicate an average CK-MB level of 15.3, exhibiting significant variability ($SD \approx 46.3$), with values reaching as high as 300, consistent with increased levels often associated with myocardial damage. Troponin levels, a vital marker of cardiac injury, average 0.36 ng/mL but can reach 10.3 ng/mL, exhibiting a high standard deviation concerning the mean, indicative of patients experiencing substantial cardiac events.

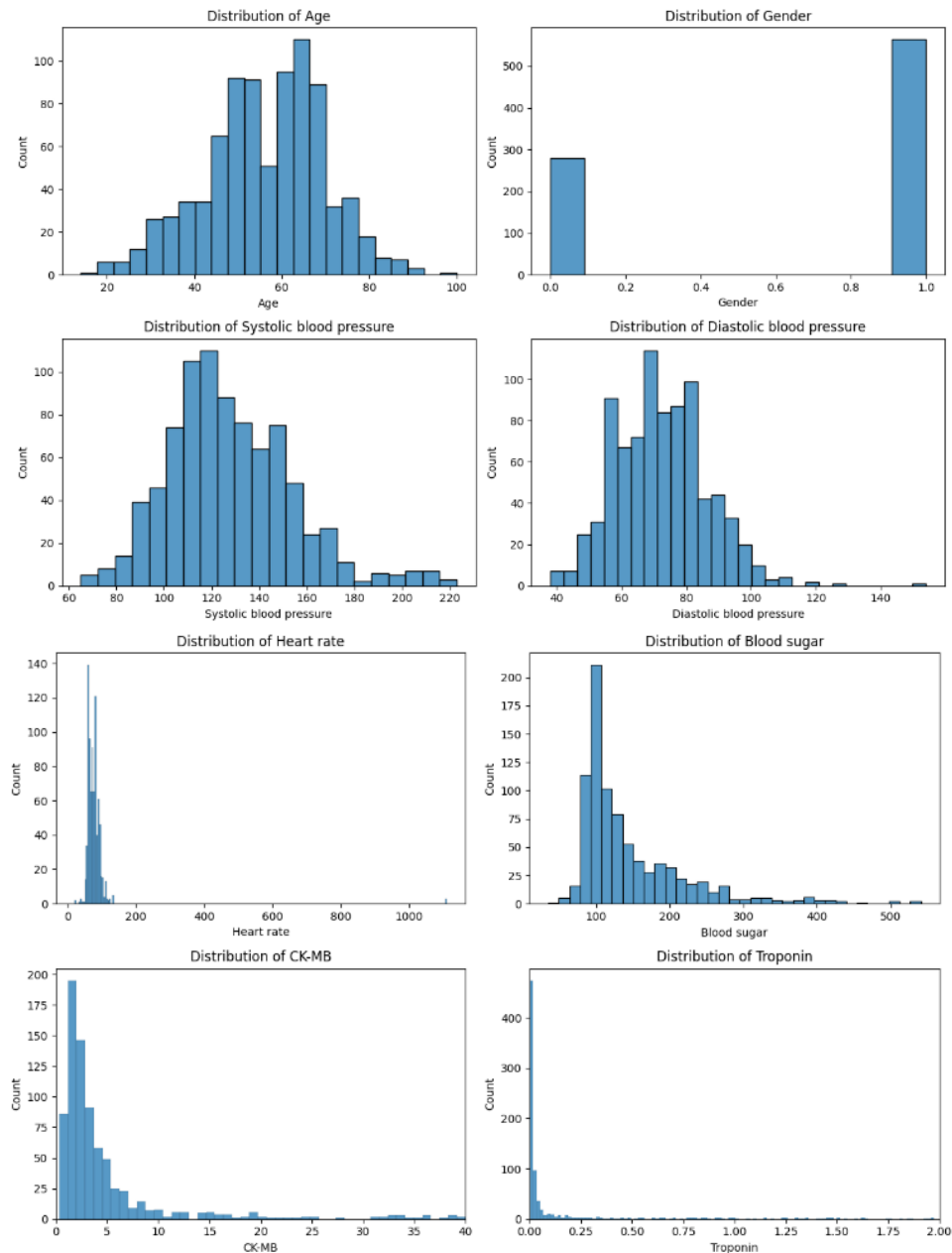


Figure 3. Distribution of All Variables

Fig. 3 explores significant patterns among demographic, clinical, and biochemical factors. The age distribution is approximately normal, with most patients aged between 50 and 65 years, indicating that middle-aged and older people constitute the principal at-risk demographic for heart disease, while the dataset also encompasses younger and elderly individuals. The gender distribution reveals a pronounced imbalance, with male patients far exceeding females, a consideration that could affect the efficacy of predictive models if inadequately managed. Blood pressure measurements (systolic and diastolic) predominantly fall within normal to mildly raised ranges (110–140 mmHg systolic, 60–90 mmHg diastolic), although have extended

upper tails reaching 220 and 150 mmHg, respectively, encompassing instances of severe hypertension. The heart rate primarily falls within the typical range of 60–120 bpm. However, severe outliers beyond 1000 bpm indicate potential data recording issues. Blood glucose levels exhibit a right-skewed distribution, with numerous patients presenting elevated values over 200 mg/dL, indicative of diabetes or impaired glucose metabolism, and infrequent instances surpassing 500 mg/dL. Cardiac biomarkers exhibit significantly skewed distributions: CK-MB levels predominantly fall below 10, occasionally exceeding 30, whereas troponin concentrations are typically near zero for many patients but exhibit a significant increase in a few, offering compelling diagnostic evidence of myocardial damage. These findings collectively emphasize the dataset's capacity to encompass both normal and pathological instances, while also highlighting the necessity of addressing gender imbalance, skewed distributions, and severe outliers in future predictive models.

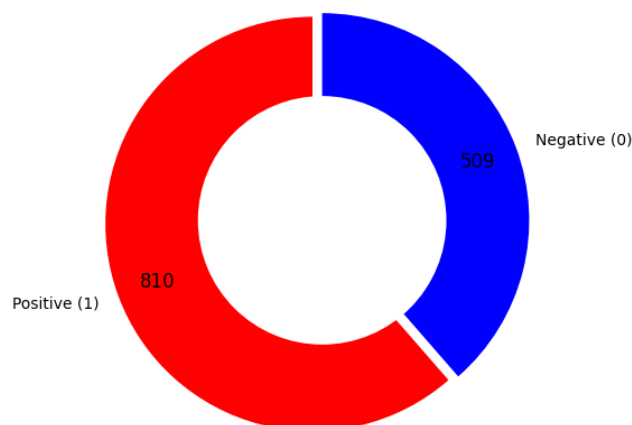


Figure 4. Distribution of Heart Attack Cases

According to Fig. 4, the imbalance ratio is computed at 1.591, signifying a disparity in the data, where the sample size of the majority class (positive) is about 1.591 times bigger than that of the minority class (negative). Consequently, data balancing was executed via SMOTE to enhance the representation of the minority class.

After the use of the SMOTE procedure, data distribution achieved a more equitable balance between the majority and minority classes, with each class having 810 cases. The outcome of this procedure indicates that the quantity of data points from both categories is now equilibrated. Following SMOTE, the dataset ($n = 1319$) was divided into training (80%, $n=1,055$) and testing (20%, $n = 264$) subsets. The dataset was partitioned into training (80%, $n = 844$) and validation (20%, $n = 211$) subsets to facilitate hyperparameter optimization and reduce overfitting. Using Google Colab programming [29], a Random Forest classifier was configured using $n_estimators = 100$, $criterion = 'entropy'$, $max_depth = 10$, $min_samples_split = 4$, $min_samples_leaf = 2$, $max_features = 'sqrt'$, and $random_state = 42$. We implement the Random Forest algorithm for its resilience to noisy information, capacity to model non-linear relationships, and efficacy in addressing class imbalance relative to singular classifiers. We set these parameterizations to equilibrate model complexity and generalization, guaranteeing reliable and resilient classification performance.

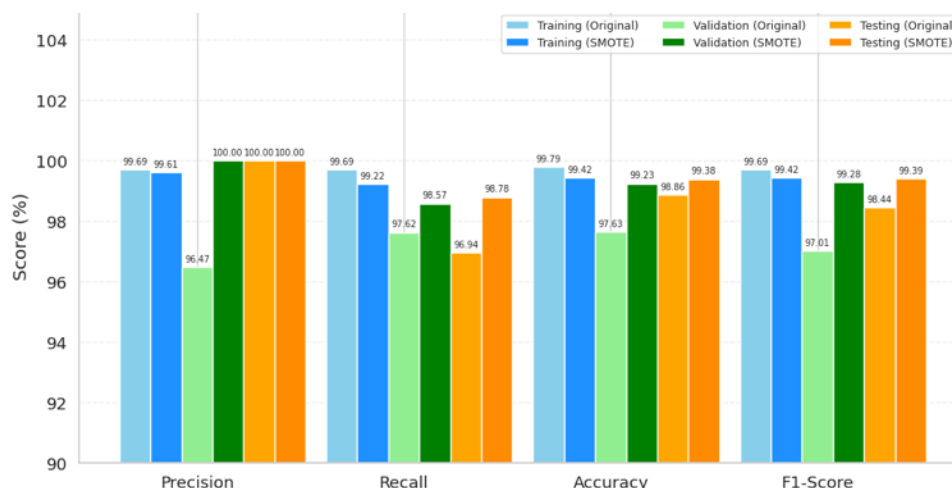


Figure 5. Model Performance Original Data and After SMOTE

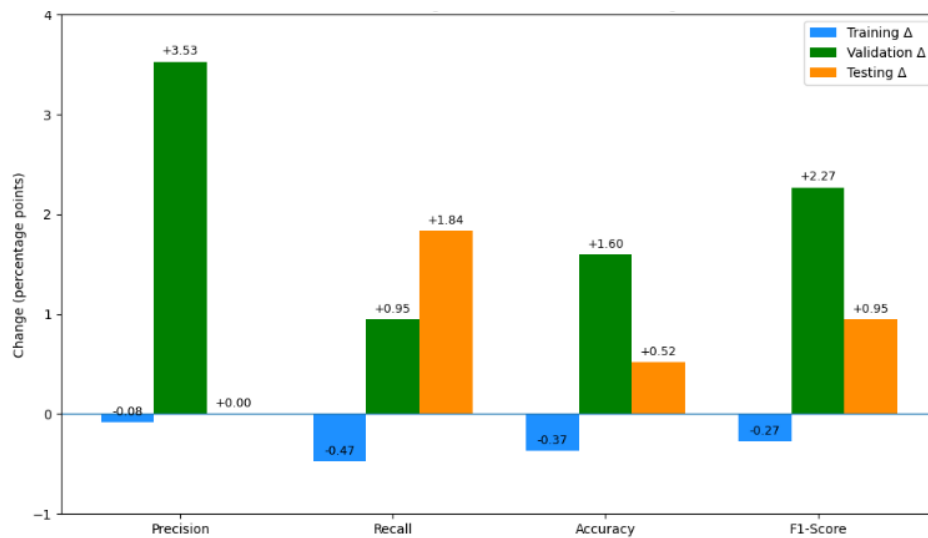


Figure 6. Performance Change from Original Data and After SMOTE

Figs. 5 and 6 show the comparative assessment of model performance pre- and post-SMOTE, indicating that oversampling significantly improved the generalization capacity of the Random Forest classifier, especially during the validation and testing stages. Although training measurements demonstrated only slight declines ($< 0.5\%$) in recall, accuracy, and F1-score, these diminutions indicate a decrease in overfitting and enhanced class balance. Significant enhancements were noted in the validation and testing sets, where SMOTE augmented precision, recall, and F1-score, with validation precision exhibiting the most considerable rise (+3.53 points), achieving 100%. Recall demonstrated constant improvement (+0.95 in validation and +1.84 in testing), underscoring the model's augmented capacity to identify positive cases accurately. The F1-score, which reconciles precision and recall, rose by 2.27 points in validation and 0.95 points in testing, while accuracy was enhanced by 1.60 and 0.52 points, respectively. These enhancements demonstrate that SMOTE not only alleviates class imbalance but also augments the model's efficacy in differentiating heart disease instances in novel data.

Subsequently, we perform SHAP analysis on the optimal model trained with the SMOTE-augmented dataset. The results were depicted through force charts, which demonstrate the contribution and relative significance of each variable in influencing the model's predictive outputs. SHAP force plots display the contributions of features to specific predictions, thus providing local interpretability of the model. Each attribute is categorized as either risk-enhancing (favoring the positive class, indicated in red) or risk-mitigating (favoring the negative class, indicated in blue). The base value, usually aligned with the mean model forecast, functions as the initial reference point. The aggregate contributions of individual characteristics adjust the prediction towards the ultimate output probability. This facilitates the precise identification of the most significant variables influencing a specific forecast, together with traits that mitigate the outcome. These visualizations are especially beneficial in clinical decision support, as they allow practitioners to comprehend both the final categorization and the model's underlying rationale at the individual patient level.

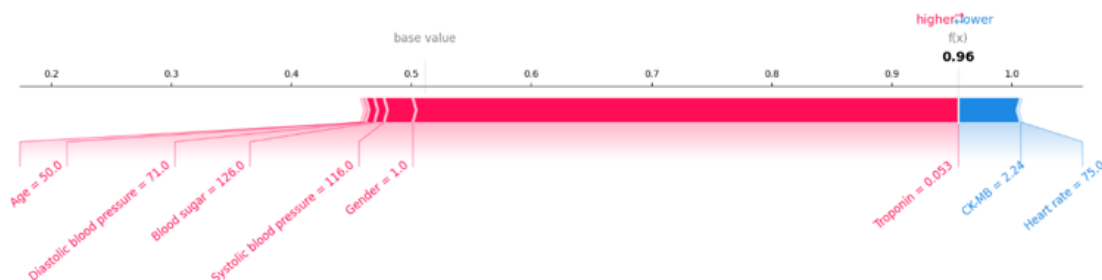


Figure 7. SHAP Force Plot of Individual Feature Contribution to Prediction for Patient 1

Fig. 7 SHAP force plot delineates the feature contributions influencing the model's prediction for a particular occurrence, resulting in a final output probability of 0.96, significantly preferring the positive class (heart attack). Positive contributions (in red) denote factors that enhanced the probability of a favorable

prediction, whereas negative contributions (in blue) diminished it. In this instance, age (59 years), diastolic blood pressure (71 mmHg), blood glucose (125 mg/dL), systolic blood pressure (116 mmHg), and gender (male) were the predominant risk-enhancing factors, thereby elevating the prediction over the baseline value of 0.5. In contrast, diminished troponin levels (0.063), CK-MB (2.24), and heart rate (75 bpm) served as mitigating variables, decreasing the likelihood of a favorable outcome. Nevertheless, the significant cumulative impact of the risk-increasing attributes surpassed the protective elements, resulting in the model forecasting a high likelihood of a heart attack.

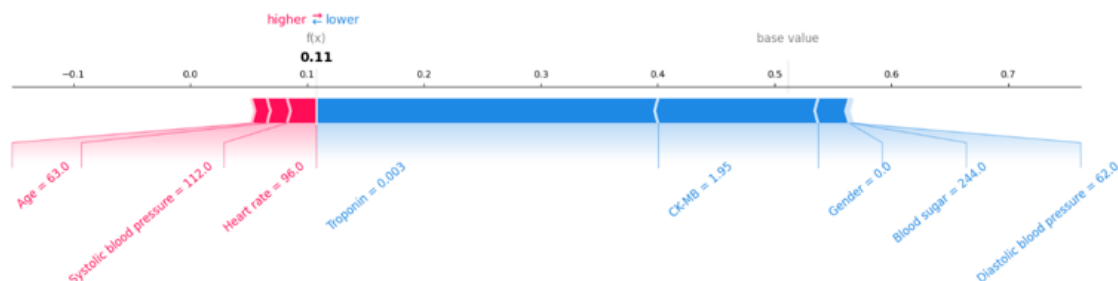


Figure 8. SHAP Force Plot of Individual Feature Contribution to Prediction for Patient 2

In Fig. 8, the baseline model prediction was almost 0.5. However, the final predicted probability was significantly reduced to 0.11, signifying a low risk of heart attack. Factors including age (63 years), systolic blood pressure (112 mmHg), and heart rate (96 bpm) contributed to an elevated prediction, resulting in a little increase in the estimated risk. Nevertheless, these effects were eclipsed by the more significant negative influences of other factors, particularly low concentrations of CK-MB (1.95) and troponin (0.013), as well as diastolic blood pressure (62 mmHg), blood glucose (214 mg/dL), and gender. These variables collectively diminished the estimated chance, resulting in a final classification aligned with a low-risk profile.

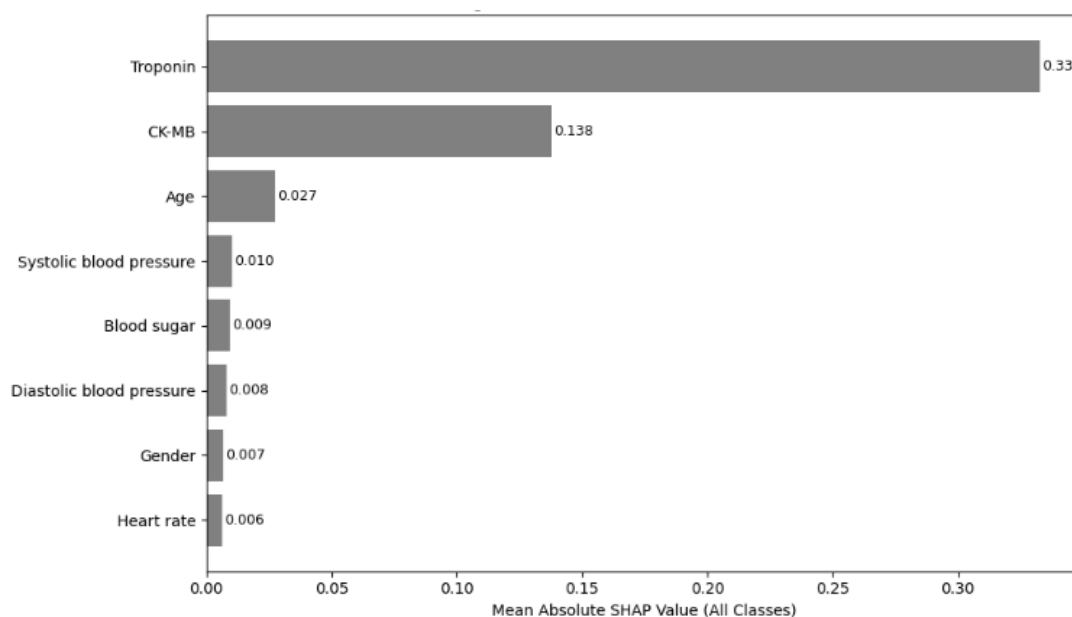


Figure 9. Average Features Contribution

As shown in Fig. 9, the SHAP feature importance analysis indicates that troponin levels are the primary predictor of heart attack risk, exhibiting the most significant mean absolute SHAP value (0.333), which signifies a significant influence on the model's classification outcomes. Subsequently, CK-MB (0.138) holds considerable significance, further underscoring the clinical relevance of cardiac biomarkers in forecasting unfavorable cardiac events. Age has a considerable contribution (0.027), underscoring its significance as a demographic risk factor. However, physiological variables, including systolic blood pressure (0.010), blood sugar (0.009), diastolic blood pressure (0.008), gender (0.007), and heart rate (0.006), exhibit comparatively minor effects. These findings underscore that although general clinical and demographic characteristics have additional predictive value, the model's decision-making process is predominantly influenced by biochemical indicators directly linked to heart attack. The study's conclusions are consistent with established medical data

and provide a more detailed, model-based perspective. They bolster the prevailing diagnostic significance of biochemical indicators and endorse the growing consensus in machine learning research that combining laboratory data significantly improves predictive accuracy. This study offers distinct insights by quantifying relationships within a SHAP-interpretable Random Forest framework, demonstrating that, in real-world clinical data from Zheen Hospital, biochemical markers retain superior predictive power even when integrated with a range of demographic and physiological characteristics.

4. CONCLUSION

This research assessed the application of data balancing and interpretable machine learning for predicting heart attacks. The dataset, comprising 810 positive and 509 negative instances, exhibited class imbalance, which was rectified through the application of SMOTE. A Random Forest classifier trained on balanced data demonstrated enhanced performance, especially in validation and testing, exhibiting elevated precision (up to 100%), recall, F1-scores, and accuracy, while concurrently mitigating overfitting. SHAP analysis was employed to guarantee interpretability, offering both local (patient-specific) and global elucidations of feature contributions. The results indicated that troponin and CK-MB were the primary predictors, followed by age, whilst other clinical and demographic factors had a negligible impact.

The findings possess significant clinical ramifications: the model's enhanced precision and recall diminish the likelihood of overlooked diagnoses, while SHAP-based elucidations augment transparency, connecting predictions with recognized medical knowledge. These methodologies can facilitate risk classification, prompt intervention, and hospital resource allocation. The study is constrained by its restricted dataset size, singular algorithmic methodology, and dependence on static clinical characteristics. Subsequent investigations should corroborate findings with bigger, more heterogeneous cohorts, examine alternative algorithms and multimodal data sources, and evaluate SHAP-based decision support in practical applications.

Author Contributions

Farrel Rafa Akbar: Data Curation, Formal Analysis, Project Administration, Visualization. Dina Tri Utari: Conceptualization, Formal Analysis, Investigation, Methodology, Supervision, Validation, Writing - Original Draft, Writing - Review and Editing. All authors discussed the results and contributed to the final manuscript.

Funding Statement

The authors received no financial support for the research, authorship or publication of this article.

Acknowledgment

The authors would like to express their sincere gratitude to the Department of Statistics, Universitas Islam Indonesia, for the continuous support, guidance, and resources provided during this work.

Declarations

The authors declare that there is no conflict of interest related to this work.

Declaration of Generative AI and AI-assisted technologies

The authors declare that no generative AI or AI-assisted technologies were used in the preparation of this manuscript, including for writing, editing, data analysis, or the creation of tables and figures.

REFERENCES

- [1] J. Huang and M. Li, "CLASSIFICATION AND PREDICTION OF CARDIOVASCULAR PATIENTS BASED ON OPTIMAL RANDOM FOREST ALGORITHM," *Cloud and Service-Oriented Computing*, vol. 2, no. 1, pp. 28–35, 2022, doi: <https://doi/10.23977/csoc.2022.020104>.
- [2] L. Yang, *et al.*, "STUDY OF CARDIOVASCULAR DISEASE PREDICTION MODEL BASED ON RANDOM FOREST IN EASTERN CHINA," *Sci Rep*, vol. 10, no. 1, pp. 5245, 2020, doi: <https://doi/10.1038/s41598-020-62133-5>.

- [3] C. Hu, et al., "INTERPRETABLE MACHINE LEARNING FOR EARLY PREDICTION OF PROGNOSIS IN SEPSIS: A DISCOVERY AND VALIDATION STUDY," *Infect Dis Ther*, vol. 11, no. 3, pp. 1117–1132, 2022, doi: <https://doi.org/10.1007/s40121-022-00628-6>.
- [4] S. J. Saba, E. A. Abd Al-Kareem, and M. R. Hameed, "OPTIMIZATION BASED APPROACH FOR HEART DISEASE CLASSIFICATION," *International Journal of Computational Methods and Experimental Measurements*, vol. 13, no. 1, pp. 149–156, 2025, doi: <https://doi.org/10.18280/ijcmem.130116>.
- [5] A. G. Abiodun, et al., "DETECTION OF HEART DISEASE USING BINARY CLASSIFICATION MACHINE LEARNING MODEL," *Ingenierie des Systemes d'Information*, vol. 30, no. 5, pp. 1111–1122, 2025, doi: <https://doi.org/10.18280/isi.300501>.
- [6] C. S. Chaithra, S. Siddesha, V. N. Manjunath Aradhya, and S. K. Niranjana, "A REVIEW OF MACHINE LEARNING TECHNIQUES USED IN THE PREDICTION OF HEART DISEASE," *Revue d'Intelligence Artificielle*, vol. 38, no. 1, pp. 201–212, 2024, doi: <https://doi.org/10.18280/ria.380120>.
- [7] S. Jaya, A. Yulianto, and E. Purwanto, "COMPUTATIONAL MODELING USING LINEAR REGRESSION AND RANDOM FOREST TO ANALYZE THE IMPACT OF WORKLOAD ON EMPLOYEE PERFORMANCE EVALUATIONS," *International Journal of Computational Methods and Experimental Measurements*, vol. 13, no. 2, pp. 227–237, 2025, doi: <https://doi.org/10.18280/ijcmem.130202>.
- [8] J. Ma, et al., "INTERPRETABLE MACHINE LEARNING ALGORITHMS REVEAL GUT MICROBIOME FEATURES ASSOCIATED WITH ATOPIC DERMATITIS," *Front Immunol*, vol. 16, pp. 1528046, 2025, doi: <https://doi.org/10.3389/fimmu.2025.1528046>.
- [9] I. Ahn, et al., "MACHINE LEARNING-BASED HOSPITAL DISCHARGE PREDICTION FOR PATIENTS WITH CARDIOVASCULAR DISEASES: DEVELOPMENT AND USABILITY STUDY," *JMIR Med Inform*, vol. 9, no. 11, pp. e32662, 2021, doi: <https://doi.org/10.2196/32662>.
- [10] L. Breiman, "RANDOM FORESTS," *Mach Learn*, vol. 45, no. 1, pp. 5–32, 2001, doi: <https://doi.org/10.1023/A:1010933404324>.
- [11] G. Biau, and E. Scornet, "A RANDOM FOREST GUIDED TOUR," *TEST*, vol. 25, no. 2, pp. 197–227, 2016, doi: <https://doi.org/10.1007/s11749-016-0481-7>.
- [12] R. Genuer, "VARIANCE REDUCTION IN PURELY RANDOM FORESTS," *J Nonparametr Stat*, vol. 24, no. 3, pp. 543–562, 2012, doi: <https://doi.org/10.1080/10485252.2012.677843>.
- [13] Y. Mishina, R. Murata, Y. Yamauchi, T. Yamashita, and H. Fujiyoshi, "BOOSTED RANDOM FOREST," *IEICE Trans Inf Syst*, vol. E98.D, no. 9, pp. 1630–1636, 2015, doi: <https://doi.org/10.1587/transinf.2014OPP0004>.
- [14] M. H. Hafezi, L. Liu, and H. Millward, "LEARNING DAILY ACTIVITY SEQUENCES OF POPULATION GROUPS USING RANDOM FOREST THEORY," *Transportation Research Record: Journal of the Transportation Research Board*, vol. 2672, no. 47, pp. 194–207, 2018, doi: <https://doi.org/10.1177/0361198118773197>.
- [15] F. David Krüger, and M. Nabeel, "HYPERPARAMETER TUNING USING GENETIC ALGORITHMS A STUDY OF GENETIC ALGORITHMS IMPACT AND PERFORMANCE FOR OPTIMIZATION OF ML ALGORITHMS," 2021.
- [16] A. Cutler, D. R. Cutler, and J. R. Stevens, "RANDOM FORESTS. IN: ENSEMBLE MACHINE LEARNING," New York, NY: Springer New York, pp. 157–175, 2012, doi: https://doi.org/10.1007/978-1-4419-9326-7_5.
- [17] P. Gulati, A. Sharma, and M. Gupta, "THEORETICAL STUDY OF DECISION TREE ALGORITHMS TO IDENTIFY PIVOTAL FACTORS FOR PERFORMANCE IMPROVEMENT: A REVIEW," *Int J Comput Appl*, vol. 141, no. 14, pp. 19–25, 2016, doi: <https://doi.org/10.5120/ijca2016909926>.
- [18] S. M. Lundberg and S.-I. Lee, "A UNIFIED APPROACH TO INTERPRETING MODEL PREDICTIONS," In: *Advances in Neural Information Processing Systems*, I. Guyon, U. Von Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan, and R. Garnett, Eds., Curran Associates, 2017, Inc. [Online].
- [19] V. Baklanova, A. Kurkin, and T. Teplova, "INVESTOR SENTIMENT AND THE NFT HYPE INDEX: TO BUY OR NOT TO BUY?," *China Finance Review International*, vol. 14, no. 3, pp. 522–548, 2024, doi: <https://doi.org/10.1108/CFRI-06-2023-0175>.
- [20] A. Stojić, M. Matek Sarić, and S. Herceg Romanić, "SHAPLEY ADDITIVE EXPLANATIONS OF INDICATOR PCB-138 DISTRIBUTION IN BREAST MILK," In: *Proceedings of the International Scientific Conference-Sinteza*, Beograd, Serbia: Singidunum University, 35–40, 2020, doi: <https://doi.org/10.15308/Sinteza-2020-35-40>.
- [21] L. V. Utkin and A. V. Konstantinov, "ENSEMBLES OF RANDOM SHAPS," 2021, [Online].
- [22] J. Guo, H. Cheng, Z. Wang, M. Qiao, J. Li, and J. Lyu, "FACTOR ANALYSIS BASED ON SHAPLEY ADDITIVE EXPLANATIONS FOR SEPSIS-ASSOCIATED ENCEPHALOPATHY IN ICU MORTALITY PREDICTION USING XGBOOST — A RETROSPECTIVE STUDY BASED ON TWO LARGE DATABASE," *Front Neurol*, vol. 14, 2023, doi: <https://doi.org/10.3389/fneur.2023.1290117>.
- [23] M. R. Islam, et al., "UNDERSTANDING CANCER RISK AMONG BANGLADESHI WOMEN: AN EXPLAINABLE MACHINE LEARNING APPROACH TO SOCIO-REPRODUCTIVE FACTORS USING TERTIARY HOSPITAL DATA," *Healthcare*, vol. 13, no. 12, pp. 1432, 2025, doi: <https://doi.org/10.3390/healthcare13121432>.
- [24] A. V. Ponce-Bobadilla, V. Schmitt, C. S. Maier, S. Mensing, and S. Stodtman, "PRACTICAL GUIDE TO SHAP ANALYSIS: EXPLAINING SUPERVISED MACHINE LEARNING MODEL PREDICTIONS IN DRUG DEVELOPMENT," *Clin Transl Sci*, vol. 17, no. 11, 2024, doi: <https://doi.org/10.1111/cts.70056>.
- [25] E. E. Erikštrumbelj and I. Kononenko, "AN EFFICIENT EXPLANATION OF INDIVIDUAL CLASSIFICATIONS USING GAME THEORY," 2010, [Online].
- [26] C. Molnar, "INTERPRETING MACHINE LEARNING MODELS WITH SAP: A GUIDE WITH PYTHON EXAMPLES AND THEORY ON SHAPLEY VALUES," 2023, Christoph Molnar c/o MUCBOOK, Heidi Seibold.
- [27] C. Molnar, "INTERPRETABLE MACHINE LEARNING A GUIDE FOR MAKING BLACK BOX MODELS EXPLAINABLE," 2019, Germany: Lean Publishing.
- [28] T. A. Rashid and B. Hassan, "HEART ATTACK DATASET," 2022, *Mendeley Data*, V1.
- [29] Google, GOOGLE COLABORATORY [Computer software], 2023