



A gradient boosting–based platform with fuzzy linguistic representation for cardiovascular disease risk prediction

Amir Saleh¹, Fadhillah Azmi^{*2}

Department of Computer Engineering, Politeknik Negeri Medan, Medan, Indonesia¹

Department of Electrical Engineering, Universitas Medan Area, Medan, Indonesia²

Article Info

Keywords:

Cardiovascular Disease, Gradient Boosting, Fuzzy Linguistic Representation, Risk Prediction, Interpretability

Article history:

Received: January 21, 2026

Accepted: June 01, 2026

Published: August 01, 2026

Cite:

Amir Saleh and F. Azmi, "A Gradient Boosting–Based Platform with Fuzzy Linguistic Representation for Cardiovascular Disease Risk Prediction", *KINETIK*, vol. 11, no. 3, Aug. 2026.
<https://doi.org/10.22219/kinetik.v11i3.2699>

*Corresponding author.

Fadhillah Azmi

E-mail address:

azmi.fadhillah007@gmail.com

Abstract

Cardiovascular disease (CVD) remains one of the leading causes of mortality globally, thus early risk prediction critical for preventative healthcare. Although machine learning methods have shown promising results in CVD prediction, many existing models are difficult to interpret in clinical practice. This study proposes a CVD risk prediction platform that combines gradient boosting (GB) with fuzzy linguistic representation to improve both predictive performance and interpretability. approach has numerous preprocessing phases, including data cleaning, normalization, outlier handling, and recursive feature elimination (RFE) for feature selection. Numerical clinical attributes are transformed into fuzzy linguistic variables to provide more intuitive risk interpretation for healthcare professionals. The gradient boosting model is trained using both original and fuzzy-transformed feature representations to improve model generalization. The proposed approach is evaluated against several baseline machine learning models using accuracy, precision, recall, and F1-score. Experimental results show that the proposed framework achieves better performance than conventional models, with a maximum accuracy of 94.30%. In addition, the developed platform provides visual risk interpretation and decision-support insights that may assist healthcare practitioners in evaluating cardiovascular risk more effectively.

1. Introduction

Cardiovascular disease (CVD) ranks among the foremost causes of mortality globally [1]. Digital healthcare systems have become useful tools for tracking and diagnosing health issues because of the advancement of digital technology and easier access to health data [2]–[4]. The adoption of machine learning (ML) methods in digital health systems provides significant potential for improving disease prediction accuracy and strengthening clinical decision-making processes [5]. Compared to conventional approaches, ML algorithms can analyze intricate and detailed medical data and provide more accurate predictive models [6]. Through the utilization of big data and contemporary computational capabilities, ML can reveal complex patterns and interdependencies that are not readily detected by conventional analysis methods.

However, despite the enormous potential of ML in health prediction, its application in clinical practice still faces several challenges. One of these is the requirement for a platform that can combine several machine learning algorithms and provide an easy-to-use user interface for medical professionals and patients [7]. Such a platform must be able to process medical data from multiple sources, perform predictive analysis, and present results in a way that is simple to understand and actionable. Human oversight of AI-based decisions is essential to ensure reliability, security, and confidentiality in clinical diagnosis. Furthermore, efficient signal processing and AI-based automated classification can improve cost efficiency by reducing reliance on manual analysis, which can significantly reduce the time and resources needed for tasks such as hospital triage processes and remote patient health monitoring [8]. These technologies support various healthcare activities, including hospital triage, remote patient monitoring, physician supervision, and self-monitoring through mobile devices [9].

Numerous research have demonstrated the predictive capacity of ML algorithms in CVD risk assessment. The predictive performance for CVD was examined using three classification models, including k-nearest neighbors (KNN), naïve Bayes (NB), and support vector machines (SVM). Based on the experimental results, the SVM method has the highest accuracy value, whereas the NB approach has the lowest [10]. Another study applied the decision tree (DT) algorithm to a large CVD dataset, achieving an accuracy of 73% [11]. Despite these results, these approaches still have limitations in terms of interpretability and practical implementation in real-world healthcare systems. In addition, the adoption of advanced ML techniques, including AdaBoost and KNN, resulted in measurable gains in sensitivity (3–4%), specificity (4–5%), and accuracy (3–4%) over previous methods [12].

Despite these advancements, previous studies primarily focus on improving predictive accuracy and often overlook the interpretability of the models, which is essential in clinical decision-making. In addition, most existing approaches evaluate ML models independently without integrating them into a practical platform that can be directly used by healthcare professionals or patients. Furthermore, limited research has explored the use of fuzzy-based data transformation to enhance the clinical meaning of numerical medical attributes, which could improve the interpretability of machine learning models in clinical settings. Therefore, there remains a research gap in developing an integrated, interpretable, and application-oriented ML framework for CVD risk prediction. In particular, there is still a lack of studies that combine predictive accuracy, interpretability, and system implementation within a single unified framework.

To address these gaps, this study aims to develop an integrated CVD risk prediction platform by combining multiple ML methods and evaluating their performance. The platform utilizes a publicly available Kaggle dataset containing clinical attributes related to cardiovascular risk assessment and is designed to support practical clinical decision-making.

This study applies several preprocessing techniques, including feature normalization and feature selection, and utilizes a gradient boosting (GB) model to improve prediction performance. In addition, fuzzy logic transformation is applied to selected attributes, such as blood pressure, to convert numerical values into meaningful clinical categories. This transformation enhances the interpretability and clinical relevance of the prediction outcomes. Similar fuzzy-based transformation techniques have been adopted in previous studies and have demonstrated competitive predictive performance and improved decision-support capabilities [13], [14].

Unlike previous studies, this study contributes by integrating multiple ML algorithms within a unified platform, incorporating fuzzy logic transformation to enhance interpretability, and systematically evaluating model performance using original, fuzzy-transformed, and hybrid feature representations. Furthermore, the proposed framework emphasizes practical implementation, making it suitable for real-world clinical decision-support systems.

2. Research Method

This research develops an ML-based platform for CVD risk prediction through several methodological steps, which can be seen in Figure 1.

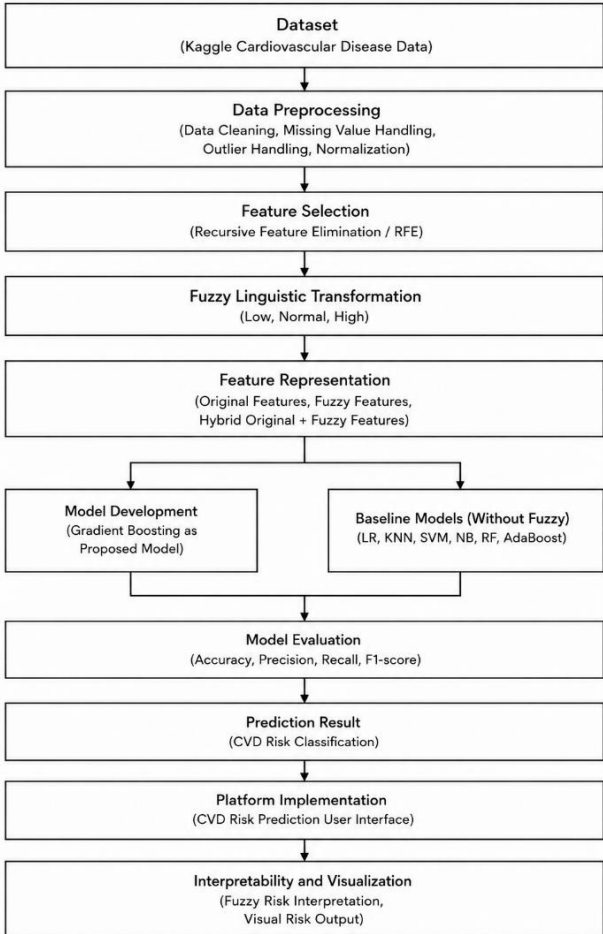


Figure 1. Proposed Methodology for CVD Risk Prediction

Based on Figure 1, the stages carried out in this research can be described as follows:

2.1 Data Collection

The dataset used in this research was obtained from a Kaggle dataset entitled “[Risk Factors for CVD](#)”. The data consists of 70,000 rows with 12 features. Information about the dataset used can be seen in Table 1 below.

Table 1. Dataset Description

Features	Description	Data type
age	Patient age expressed in days	Numerical
gender	Biological sex of the patient (male or female)	Categorical
height	Patient body height measured in centimeters	Numerical
weight	Patient body weight measured in kilograms	Numerical
ap_hi	Systolic blood pressure measurement	Numerical
ap_lo	Diastolic blood pressure reading of the patient	Numerical
cholesterol	Patient cholesterol category	Categorical
gluc	Patient glucose level category	Categorical
smoke	Smoking behavior status	Categorical
alco	Alcohol consumption behavior	Categorical
active	Level of physical activity	Categorical
cardio	Indicator of cardiovascular disease presence or absence	Binary

To assess the models, the data in Table 1 are used for testing and training ML models. A training set of 70% and a testing set of 30% are randomly selected from the original dataset.

2.2 Data Preprocessing

To ensure the readiness of the data for analysis and ML model development, the preparation stage involves cleaning the data, handling missing values, normalizing it, and transforming its features [15]. Filling in missing values in the dataset with values deemed appropriate can be done using the mean method in Equation 1 [16].

$$\bar{x} = \frac{1}{N} \sum_{i=1}^N x_i \quad (1)$$

where $\bar{x}$ denotes the mean of the dataset, x_i represents the value of the i -th observation, and N indicates the total number of data points.

Selecting data and eliminating outliers will be the next phase. The purpose of this is to gather relevant data that aligns with the research's goals. This process is increasingly important in the health sector, where the need for effective decision-making is very high [17]. The technique used in this research is k-means clustering, as described in Equation 2 and Equation 3 [18].

$$d(x_i, \mu_j) = \sqrt{\sum_{m=1}^M (x_{im} - \mu_{jm})^2} \quad (2)$$

$$\mu_j = \frac{1}{|C_j|} \sum_{x_i \in C_j} x_i \quad (3)$$

where $d(x_i, \mu_j)$ denotes the distance between data point x_i and centroid μ_j ; M represents the dimensionality of the feature space; μ_j corresponds to the feature vector of the j -th centroid; and x_i denotes the feature vector of the i -th observation. Furthermore, C_j refers to the set of data points assigned to cluster j , and $|C_j|$ indicates the number of data points within that cluster.

The procedure is repeated until a convergence criterion is satisfied, namely when the centroids exhibit negligible changes or a predefined maximum number of iterations is reached. Upon completion of the k-means clustering process, data points that do not conform to any cluster (outliers) can be identified and removed, ensuring that only data that are relevant and aligned with the research objectives are retained for subsequent analysis. The final stage of this process is normalization, where this step changes the feature scale to a certain range. This is done to ensure that all features have an equal contribution to the model. This process is carried out using Z-score normalization in Equation 4 [19].

$$x' = \frac{x - \mu}{\sigma} \quad (4)$$

where μ corresponds to the mean value and σ indicates the standard deviation of the data.

2.3 Feature Selection

This procedure is applied to identify and extract relevant features from the dataset through feature selection techniques. Removing redundant and irrelevant data with feature selection is an efficient method of reducing dimensionality [20]. This research applies recursive feature elimination (RFE) and evaluates its performance against alternative feature selection methods. Through this approach, the most informative features for heart disease prediction are identified and selected to improve the effectiveness and accuracy of the learning algorithm. In the feature selection process, RFE is used to assess each feature's contribution to the target variable in Equation 5 [21].

$$RFE(f) = \arg \min_{\text{subset}(S) \text{ of } f} \left(\frac{1}{K} \sum_{i=1}^K CVscore(S) \right) \quad (5)$$

where:

$RFE(f)$ is the RFE function that returns the best set of features.

f is the set of features being evaluated.

$CVscore(S)$ is the cross-validation score.

K is the number of folds in cross-validation.

$f[i]$ is the feature removed at iteration i .

This formula illustrates how RFE improves model performance by gradually removing less important data while retaining just the most valuable information related to the target variable.

2.4 Machine Learning Models

This research employs gradient boosting (GB), an ensemble learning technique widely used for classification and regression tasks. This method works by combining a large number of simple predictive models called weak learners or base learners, such as decision trees (DTs), into a more powerful predictive model [22]. In GB, the following is used to update the prediction $F(x)$ with a weak learner $h(x)$ in Equation 6 [23].

$$F_m(x) = F_{m-1}(x) + \alpha_m h_m(x) \quad (6)$$

where $F_m(x)$ represents the ensemble model's prediction after the m -th iteration; $F_{m-1}(x)$ represents the previous ensemble prediction; α_m is the learning rate to control the contribution of each weak learner, and $h_m(x)$ is a weak learner (for example, a DT) selected at the m -th iteration, which tries to improve the residuals from the previous prediction.

This equation reflects the GB iterative procedure, which adds each weak learner to the ensemble model's total predictions by taking into account residuals from prior predictions. Furthermore, multiple ML models are utilized in this research, namely logistic regression (LR), k-nearest neighbors (KNN), support vector machines (SVM), AdaBoost, naïve Bayes (NB), and random forest (RF). In this study, baseline machine learning models are trained using the original feature representation, while the proposed gradient boosting model is evaluated using original, fuzzy-transformed, and hybrid feature representations. In order to achieve robust and accurate CVD prediction, this research seeks to identify the optimal model based on comparative evaluation using accuracy, precision, recall, and F1-score metrics.

2.5 Model Evaluation Metrics

The performance of ML models is assessed using a range of evaluation metrics, namely accuracy, precision, recall, and F1-score [24]. This evaluation method aids in assessing the relative performance of each model, allowing the selection of the most appropriate model for applied data analysis purposes. The formulas for calculating accuracy, precision, recall, and F1-score are presented in Equation 7, Equation 8, Equation 9, and Equation 10 [25].

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (7)$$

$$Precision = \frac{TP}{TP + FP} \quad (8)$$

$$Recall = \frac{TP}{TP + FN} \quad (9)$$

$$F1 - score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (10)$$

where FN , TP , TN , and FP correspond to false negatives, true positives, true negatives, and false positives, respectively.

2.6 Fuzzy Logic Transformation

To make blood pressure measurements obtained from clinical measurement devices easier for healthcare professionals to interpret, this research employs the fuzzy logic method. This research aims to provide a more interpretable representation of measurement findings using fuzzy logic, allowing healthcare professionals to make better decisions based on the available information [26]. In this research, fuzzy sets are defined with a membership function that describes the extent to which an element is included in the set. For each variable that represents blood measurements, a fuzzy set can be defined with a membership function, as shown in Table 2 below [27].

Table 2. Membership Functions for Blood Pressure Measurements

Membership function	Systolic BP (mmHg)	Diastolic BP (mmHg)
Low	[90, 100, 110]	[60, 70, 80]
Normal	[100, 120, 130]	[70, 80, 90]
High	[120, 140, 160]	[80, 90, 110]

In Table 2, systolic and diastolic blood pressure values were transformed into fuzzy linguistic variables using triangular membership functions. Each fuzzy set is defined by three parameters representing the lower bound, peak, and upper bound.

2.7 Platform Development

This CVD risk prediction platform was developed using several key features to support effective medical data analysis and management. Integration of data from multiple medical sources ensures the platform has access to comprehensive and up-to-date information. The intuitive user interface is specifically designed to meet the needs of medical professionals and make interaction with patients easier. This platform also implements selected ML models to perform CVD risk predictions based on available data and provides prediction results through clear reports and informative data visualizations. Thus, this platform not only supports better diagnosis but also facilitates informed decision-making in CVD risk management. The stages used to build the proposed platform can be seen in Figure 2 below:

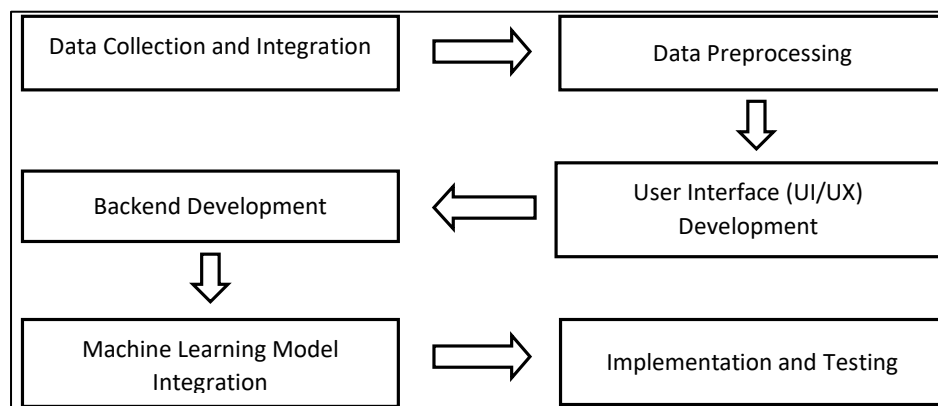


Figure 2. Development of the Proposed CVD Platform

In Figure 2, the process of developing a CVD risk prediction platform involves several key stages. First, the medical data obtained is collected and integrated to ensure comprehensive and up-to-date information is available. Next, the data are processed, including cleaning, handling missing values, and data transformation using fuzzy logic for more informative blood pressure categories. The user interface was then developed to be intuitive and user-friendly, optimized for interactions between medical professionals and patients, and features were designed to clearly display patient information and predict results. Meanwhile, the application backend is built to process data, integrate ML models, and store user data. ML models that have been trained are evaluated using relevant performance metrics. The

application is then implemented and tested using selected testing data to ensure that the proposed system operation is responsive and reliable.

3. Results and Discussion

This section presents and analyzes the experimental results of the proposed CVD risk prediction platform using advanced ML methods. The process includes feature selection, model evaluation, and system validation to ensure reliable performance.

3.1 Feature Selection Results

Feature selection constitutes a critical stage in model development to ensure that ML models utilize the most relevant and informative data. In this study, multiple feature selection techniques were employed to identify the most influential variables for CVD risk prediction. The evaluated methods include principal component analysis (PCA), analysis of variance (ANOVA), recursive feature elimination (RFE), and mutual information (MI). The feature selection results are presented in Table 3.

Table 3. Feature Selection Results Using Feature Selection Techniques

Methods	Features
RFE	'f_age', 'f_height', 'f_weight', 'f_ap_hi', 'f_ap_lo', 'f_cholesterol', 'f_gluc', 'f_active'
PCA	'f_age', 'f_cholesterol', 'f_gluc', 'f_ap_hi', 'f_active', 'f_ap_lo', 'f_weight', 'f_alco'
ANOVA	'f_age', 'f_height', 'f_weight', 'f_ap_hi', 'f_ap_lo', 'f_cholesterol', 'f_gluc', 'f_smoke'
MI	'f_age', 'f_gender', 'f_weight', 'f_ap_hi', 'f_ap_lo', 'f_cholesterol', 'f_gluc', 'f_active'

According to Table 3, each feature selection method identified a subset of eight features as the most significant variables for CVD prediction. The selected features were then used in the training and testing process to optimize model performance. This indicates that key clinical features, such as blood pressure, cholesterol, and age, consistently contribute to CVD risk prediction.

3.2 Machine Learning Performance

The predictive capability of the proposed ML model was evaluated through performance analysis. Several algorithms were implemented and compared to evaluate the proposed approach. The models were trained and validated using the preprocessed dataset, which was partitioned into training and testing subsets. The outcomes of feature selection and ML model performance are presented in Table 4.

Table 4. Performance of ML Algorithms in Predicting CVD Risk

ML methods	Feature selection methods	Performance Results			
		Accuracy	Precision	Recall	F1-score
LR	None	0.8790	0.8790	0.8790	0.8789
	RFE	0.8794	0.8794	0.8794	0.8793
	PCA	0.8789	0.8789	0.8789	0.8788
	ANOVA	0.8798	0.8798	0.8798	0.8797
	MI	0.8793	0.8793	0.8793	0.8792
KNN	None	0.8699	0.8699	0.8699	0.8697
	RFE	0.8748	0.8749	0.8748	0.8747
	PCA	0.8799	0.8799	0.8799	0.8798
	ANOVA	0.8755	0.8756	0.8755	0.8754
	MI	0.8794	0.8794	0.8794	0.8793
SVM	None	0.8826	0.8827	0.8826	0.8824
	RFE	0.8848	0.8849	0.8848	0.8846
	PCA	0.8849	0.8851	0.8849	0.8848
	ANOVA	0.8850	0.8852	0.8850	0.8848
	MI	0.8850	0.8852	0.8850	0.8848
NB	None	0.8301	0.8307	0.8301	0.8302
	RFE	0.8327	0.8333	0.8327	0.8328
	PCA	0.8327	0.8333	0.8327	0.8329
	ANOVA	0.8322	0.8328	0.8322	0.8324
	MI	0.8338	0.8344	0.8338	0.8339
AdaBoost	None	0.8939	0.8941	0.8939	0.8937
	RFE	0.8937	0.8940	0.8937	0.8935

ML methods	Feature selection methods	Performance Results			
		Accuracy	Precision	Recall	F1-score
RF	PCA	0.8930	0.8932	0.8930	0.8928
	ANOVA	0.8937	0.8939	0.8937	0.8935
	MI	0.8932	0.8935	0.8932	0.8930
	None	0.8897	0.8903	0.8897	0.8895
	RFE	0.8862	0.8867	0.8862	0.8860
	PCA	0.8770	0.8771	0.8770	0.8769
	ANOVA	0.8872	0.8876	0.8872	0.8869
	MI	0.8801	0.8802	0.8801	0.8799
	None	0.8947	0.8968	0.8947	0.8942
	RFE	0.8965	0.8985	0.8965	0.8960
GB	PCA	0.8951	0.8971	0.8951	0.8946
	ANOVA	0.8959	0.8979	0.8959	0.8954
	MI	0.8963	0.8983	0.8963	0.8958

The results in Table 4 show that various ML algorithms achieve competitive accuracy in predicting CVD risk. Furthermore, the results demonstrate that the performance of these algorithms varies depending on the complexity and effectiveness of the applied feature selection. Algorithms such as GB, AdaBoost, and RF are more prominent, with accuracies reaching 89.47%, 89.39%, and 88.97%, respectively, when using all features or without feature selection. On the other hand, performance can improve for certain classifiers when feature selection techniques such as MI, ANOVA, PCA, and RFE are applied, although the effect varies across models.

The superior performance of ensemble-based models, such as GB and AdaBoost, can be attributed to their ability to capture complex non-linear relationships and reduce prediction errors through iterative learning. In contrast, simpler models such as NB show lower performance due to their assumptions of feature independence, which may not hold in medical datasets.

Based on the results, gradient boosting (GB) consistently demonstrates the highest performance across evaluation metrics, particularly when combined with RFE feature selection. Similarly, AdaBoost also shows strong performance without requiring additional feature selection. These findings suggest that ensemble-based methods are highly effective for CVD prediction tasks, particularly when handling complex and high-dimensional medical data.

The selection of the best method depends on factors such as interpretability, computational efficiency, and prediction reliability. GB and AdaBoost demonstrate superior performance in terms of prediction accuracy. Therefore, the selection of the most suitable method should consider the specific requirements of the developed platform.

Among the evaluated models, GB combined with RFE feature selection provides the best performance, indicating that selecting relevant features significantly improves model accuracy and generalization. This combination is therefore considered the most suitable for implementation in the proposed CVD prediction platform.

3.3 Platform Testing Results

On the developed platform, accuracy, interpretability, and the ability of the model to handle data complexity are critical. The combination of GB for ML models and RFE for feature selection is an option that can provide optimal and interpretable prediction results. This enables the system to generate accurate and interpretable predictions and aligns well with the objectives of the developed CVD prediction platform.

These results indicate that the developed platform is reliable and capable of supporting clinical decision-making through accurate and interpretable predictions. This process comprises unit testing, integration testing, and system testing, and is conducted to ensure that all developed components function correctly and produce the expected outcomes.

3.3.1 Platform Functionality Testing Results

Unit testing is performed for each component of the platform, including data input modules, data processing modules, ML modules, and result display modules. The components and test results of this process are presented in Table 5 below.

Table 5. Unit Test Results of the Developed Platform

Components	Features	Test results
Data Input Module	Data input format validation	Data is received and validated correctly
Data Processing Module	Data normalization and standardization	Data is processed according to standards
Machine Learning Module	Risk prediction using a pre-trained model	Predictions are accurate according to test data
Result Display Module	Displays prediction results to the UI	Results are displayed correctly

Integration testing is performed to ensure that each component works harmoniously when combined. The results of integration testing on the developed platform can be seen in Table 6 below.

Table 6. Integration Testing Results of the Developed Platform

Components	Description	Test results
Data Input Module - Data Processing	The entered data is sent for processing	The data is processed and forwarded correctly
Data Processing - Machine Learning	The processed data is sent to the machine learning module	The model receives the data and generates predictions
Machine Learning - Display of Prediction	Results from the model are sent to the UI	Prediction results are displayed correctly

System testing is carried out to ensure that the entire platform functions as expected in real usage scenarios. Table 7 below displays the results of system testing on the developed platform.

Table 7. System Test Results of the Developed Platform

Test scenario	Description	Test results
Input Flow to Results Display	User enters data, data is processed, predictions are made, results are displayed	All steps run smoothly and results are displayed correctly
Invalid Data Handling	Enter invalid data to test input validation	Errors are handled correctly and an error message is displayed
UI Responsiveness	Ensures the UI is responsive and can be used with various devices	The UI is responsive and functions well on various devices

The platform functionality testing results show that all components, their integration, and the entire system function well. All tests produced satisfactory results without any significant problems. This indicates that the platform demonstrates reliable functionality for experimental and clinical support scenarios. These testing results indicate that the platform can provide accurate and reliable predictions for CVD risk assessment. This confirms that the system architecture is robust and capable of handling end-to-end data processing and prediction tasks without functional errors.

3.3.2 Testing Results with Machine Learning

The use of the CVD prediction platform begins with the selection of a user-oriented menu tailored to the needs of end users, such as patients. This menu provides options to initiate CVD risk prediction based on the input health data, view the history of prior predictions, and configure user preferences or platform settings. The menu interface is illustrated in Figure 3.

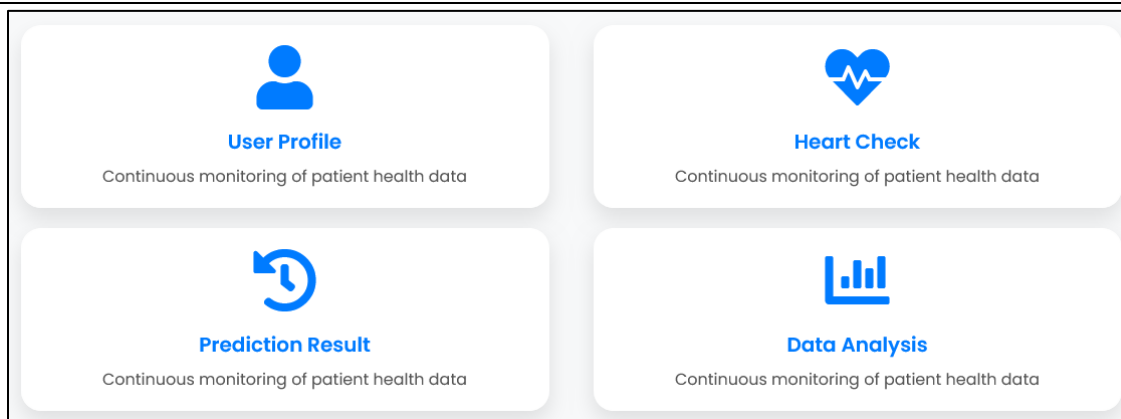


Figure 3. Menu Display for CVD Prediction

After a menu option is selected, users input relevant health information, including demographic characteristics, medical history, and diagnostic test results. The entered data are then submitted to the system for subsequent processing. At this stage, preprocessing techniques are applied to prepare the data for input into the machine learning model. The corresponding interface is shown in Figure 4 below.

Data Detail			Back
No	Check Result	Result	
1	Patient Name	ID_Patient_00209	
2	Date	2024-06-26	
3	Age	42	
4	Height	165.00 cm	
5	Weight	60.00 Kg	
6	Cholesterol Category	1	
7	Glucose Category	1	
8	Physical Activity Level	Not Active	
9	Systolic Blood Pressure	120	
10	Diastolic Blood Pressure	80	
Process to Prediction			

Figure 4. Display of Patient Data for CVD Prediction

Following data processing, an evaluation stage is conducted to assess the model's performance, including predictive accuracy and interpretability. The generated prediction outcomes are examined to ensure their reliability and clinical relevance. Finally, the prediction results are presented to end users in a clear and comprehensible format, thereby supporting informed medical decision-making. The prediction results are illustrated in Figure 5 below.

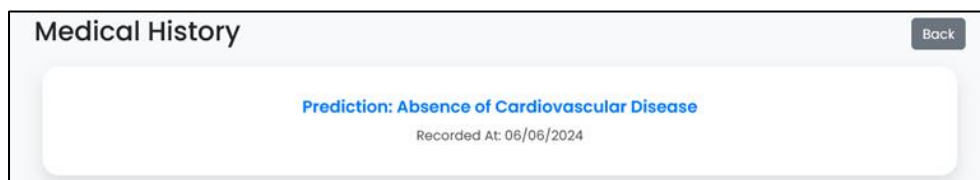


Figure 5. Display of Prediction Results for CVD

This platform also provides fuzzy transformations to simplify understanding prediction results. Fuzzy logic is applied to transform numerical values into intuitive categories (e.g., low, normal, high), improving user understanding and clinical interpretability. For example, blood pressure results can be classified into categories such as "low," "normal,"

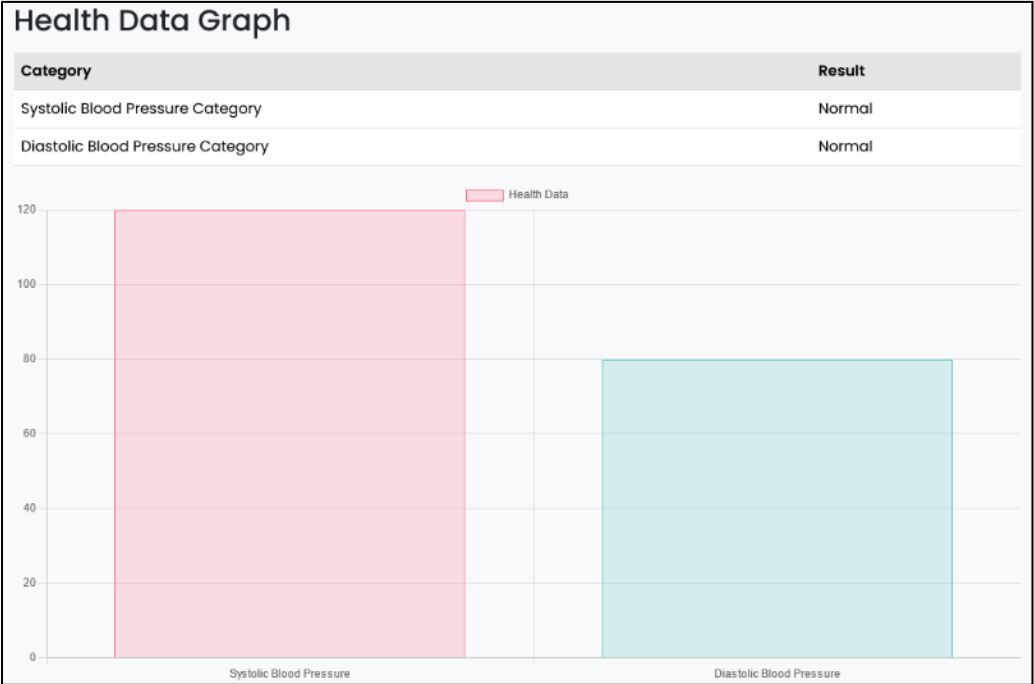


Figure 6. Fuzzy Transformation Display on Blood Pressure Measurements

Once the data entry and storage processes were confirmed to operate correctly, the next step involved validation using 1,000 randomly selected samples from the testing dataset. The validation results indicate that the proposed approach produces reliable predictions. The performance results presented in Table 4 correspond to baseline experiments conducted using original feature representations and conventional feature selection techniques. In contrast, Table 8 presents the final validation results of the proposed framework after incorporating fuzzy transformation into the gradient boosting model.

Table 8. Validation Results for CVD Prediction	
Methods	Accuracy
GB without fuzzy transformation	0.9280
GB with fuzzy transformation	0.9430

Table 8 shows the results of two prediction methods using GB. The first method, which uses GB without fuzzy transformation, achieves an accuracy of 92.80%. In this method, the original feature representation is directly used to train the prediction model. Although this accuracy is quite high, there is room for improvement. The second method, which uses GB with fuzzy transformation, achieves an accuracy of 94.30%. Fuzzy transformation helps deal with uncertainty and variability in medical data, which contributes to improving the performance of prediction models. These findings indicate that fuzzy transformation can improve CVD prediction accuracy by reducing errors and increasing model generalization to new data.

In general, the incorporation of fuzzy transformations into CVD prediction systems has demonstrated effectiveness in improving both predictive accuracy and model robustness. This indicates that the proposed approach is a dependable choice for this application, offering the highest accuracy and potentially providing the most reliable CVD risk predictions. Compared to previous studies that reported accuracy in the range of 73%–89%, the proposed approach achieves higher accuracy, demonstrating the effectiveness of combining GB with fuzzy transformation in improving prediction performance. This improvement highlights the effectiveness of fuzzy transformation in enhancing model performance.

Despite the promising results, this study has several limitations. The model performance is highly dependent on the quality and characteristics of the dataset used. In addition, the incorporation of fuzzy transformation increases computational complexity, which may affect system efficiency in large-scale implementations. Furthermore, the study is limited to a single dataset, and additional validation using diverse datasets is required to ensure generalizability.

4. Conclusion

The designed CVD prediction platform demonstrates a systematic and structured workflow, starting from user input to prediction result evaluation. The platform enables end users to input relevant health data and obtain accurate CVD risk estimates. The integration of feature selection techniques and fuzzy logic transformation has proven effective. This approach simplifies the interpretation of results and supports clinical decision-making.

The experimental results indicate that the proposed approach achieves improved performance. The GB model with fuzzy transformation attains an accuracy of 94.30%, compared to 92.80% without fuzzy transformation. This demonstrates that fuzzy transformation effectively addresses data uncertainty and variability, enhances model generalization, and improves prediction accuracy.

Overall, this study successfully achieves its objective by developing an integrated, interpretable, and practical CVD prediction system that combines multiple machine learning models with fuzzy logic transformation. Compared with previous studies reporting accuracy values ranging from 73% to 89%, the proposed framework demonstrated promising performance improvements together with enhanced interpretability.

However, this study has several limitations, including dependency on a single dataset and increased computational complexity due to fuzzy transformation. Future research should focus on evaluating the proposed approach using diverse datasets and improving computational efficiency for large-scale implementation.

References

- [1] F. F. Firdaus, H. A. Nugroho, and I. Soesanti, "A Review of Feature Selection and Classification Approaches for Heart Disease Prediction," *IJITEE (International J. Inf. Technol. Electr. Eng.)*, vol. 4, no. 3, p. 75, 2021. <https://doi.org/10.22146/ijitee.59193>
- [2] Z. Ahmed, K. Mohamed, S. Zeeshan, and X. Q. Dong, "Artificial intelligence with multi-functional machine learning platform development for better healthcare and precision medicine," *Database*, vol. 2020, pp. 1–35, 2020. <https://doi.org/10.1093/database/baaa010>
- [3] F. García-Peñalvo et al., "KoopML: A Graphical Platform for Building Machine Learning Pipelines Adapted to Health Professionals," *Int. J. Interact. Multimed. Artif. Intell.*, vol. In Press, no. In Press, p. 1, 2023. <https://doi.org/10.9781/ijimai.2023.01.006>
- [4] Y. Liu, Z. Ling, B. Huo, B. Wang, T. Chen, and E. Mouine, "Building A Platform for Machine Learning Operations from Open Source Frameworks," *IFAC-PapersOnLine*, vol. 53, no. 5, pp. 704–709, 2020. <https://doi.org/10.1016/j.ifacol.2021.04.161>
- [5] G. Quer, R. Arnaout, M. Henne, and R. Arnaout, "Machine Learning and the Future of Cardiovascular Care: JACC State-of-the-Art Review," *J. Am. Coll. Cardiol.*, vol. 77, no. 3, pp. 300–313, 2021. <https://doi.org/10.1016/j.jacc.2020.11.030>
- [6] S. Zeadally, F. Siddiqui, Z. Baig, and A. Ibrahim, "Smart healthcare: Challenges and potential solutions using internet of things (IoT) and big data analytics," *PSU Res. Rev.*, vol. 4, no. 2, pp. 149–168, 2020. <https://doi.org/10.1108/PRR-08-2019-0027>
- [7] S. Nashif, M. R. Raihan, M. R. Islam, and M. H. Imam, "Heart Disease Detection by Using Machine Learning Algorithms and a Real-Time Cardiovascular Health Monitoring System," *World J. Eng. Technol.*, vol. 06, no. 04, pp. 854–873, 2018. <https://doi.org/10.4236/wjet.2018.64057>
- [8] O. Faust, N. Lei, E. Chew, E. J. Ciaccio, and U. R. Acharya, "A smart service platform for cost efficient cardiac health monitoring," *Int. J. Environ. Res. Public Health*, vol. 17, no. 17, pp. 1–18, 2020. <https://doi.org/10.3390/ijerph17176313>
- [9] C. A. Gómez-García, M. Askar-Rodríguez, and J. Velasco-Medina, "Platform for Healthcare Promotion and Cardiovascular Disease Prevention," *IEEE J. Biomed. Heal. Informatics*, vol. 25, no. 7, pp. 2758–2767, 2021. <https://doi.org/10.1109/JBHI.2021.3051967>
- [10] A. Damayunita, R. S. Fuadi, and C. Juliane, "Comparative Analysis of Naive Bayes, K-Nearest Neighbors (KNN), and Support Vector Machine (SVM) Algorithms for Classification of Heart Disease Patients," *J. Online Inform.*, vol. 7, no. 2, pp. 219–225, 2022. <https://doi.org/10.15575/join.v7i2.919>
- [11] R. Waigi, S. Choudhary, P. Fulzele, and G. Mishra, "Predicting The Risk Of Heart Disease Using Advanced Machine Learning Approach," *Eur. J. Mol. Clin. Med.*, vol. 7, no. May, p. 2020, 2020.
- [12] P. Kumar and A. Kumar, "Heart Disease Classification and Recommendation by Optimized Features and Adaptive Boost Learning," *Int. J. Adv. Comput. Sci. Appl.*, vol. 14, no. 3, pp. 909–914, 2023. <https://doi.org/10.14569/IJACSA.2023.01403103>
- [13] J. B. Jane and E. N. Ganesh, "A review on big data with machine learning and fuzzy logic for better decision making," *Int. J. Sci. Technol. Res.*, vol. 8, no. 10, pp. 1221–1225, 2019.
- [14] L. Maretti, M. Faccio, and D. Battini, "A Multi-Criteria Decision-Making Model Based on Fuzzy Logic and AHP for the Selection of Digital Technologies," *IFAC-PapersOnLine*, vol. 55, no. 2, pp. 319–324, 2022. <https://doi.org/10.1016/j.ifacol.2022.04.213>
- [15] C. Fan, M. Chen, X. Wang, J. Wang, and B. Huang, "A Review on Data Preprocessing Techniques Toward Efficient and Reliable Knowledge Discovery From Building Operational Data," *Front. Energy Res.*, vol. 9, no. March, pp. 1–17, 2021. <https://doi.org/10.3389/fenrg.2021.652801>
- [16] K. Maharana, S. Mondal, and B. Nemade, "A review: Data pre-processing and data augmentation techniques," *Global Transitions Proceedings*, vol. 3, no. 1, pp. 91–99, 2022. <https://doi.org/10.1016/j.gltp.2022.04.020>
- [17] H. Lattar, A. Ben Salem, and H. H. Ben Ghezala, "Does data cleaning improve heart disease prediction?," *Procedia Comput. Sci.*, vol. 176, pp. 1131–1140, 2020. <https://doi.org/10.1016/j.procs.2020.09.109>
- [18] L. Fanani and N. Priandani, "Data Cleaning and Prototyping Using K-Means to Enhance Classification Accuracy," *Int. J. Appl. Eng. Res.*, vol. 12, pp. 5242–5247, Jan. 2017.
- [19] H. A. Prihanditya, "The Implementation of Z-Score Normalization and Boosting Techniques to Increase Accuracy of C4.5 Algorithm in Diagnosing Chronic Kidney Disease," *J. Soft Comput. Explor.*, vol. 1, no. 1, pp. 63–69, 2020. <https://doi.org/10.52465/josce.v1i1.8>
- [20] C. Kuzudisli, B. Bakir-Gungor, N. Bulut, B. Qaqish, and M. Yousef, "Review of feature selection approaches based on grouping of features," *PeerJ*, vol. 11, 2023. <https://doi.org/10.7717/peerj.15666>
- [21] H. Jeon and S. Oh, "Hybrid-recursive feature elimination for efficient feature selection," *Appl. Sci.*, vol. 10, no. 9, pp. 1–9, 2020. <https://doi.org/10.3390/app10093211>
- [22] R. Suhendra et al., "Cardiovascular Disease Prediction Using Gradient Boosting Classifier," *Infolitika J. Data Sci.*, vol. 1, no. 2, pp. 56–62, 2023. <https://doi.org/10.60084/ijds.v1i2.131>
- [23] S. P. Nainggolan and A. Sinaga, "Comparative Analysis of Accuracy of Random Forest and Gradient Boosting Classifier Algorithm for Diabetes Classification," *Sebatik*, vol. 27, no. 1, pp. 97–102, 2023. <https://doi.org/10.46984/sebatik.v27i1.2157>
- [24] S. A. Hicks et al., "On evaluation metrics for medical applications of artificial intelligence," *Sci. Rep.*, vol. 12, no. 1, pp. 1–9, 2022. <https://doi.org/10.1038/s41598-022-09954-8>

- [25] V. Sharma and S. Singh Samant, "Health Recommendation System by Using Deep Learning and Fuzzy Technique," *SSRN Electron. J.*, no. Aece, pp. 72–78, 2022. <https://doi.org/10.2139/ssrn.4157328>
- [26] D. F. Hasan and A. S. M. Khidhir, "Toward enhancement of deep learning techniques using fuzzy logic: a survey," *Int. J. Electr. Comput. Eng.*, vol. 13, no. 3, pp. 3041–3055, 2023. <https://doi.org/10.11591/ijece.v13i3.pp3041-3055>
- [27] A. A. Nancy, D. Ravindran, P. M. D. Raj Vincent, K. Srinivasan, and D. Gutierrez Reina, "IoT-Cloud-Based Smart Healthcare Monitoring System for Heart Disease Prediction via Deep Learning," *Electron.*, vol. 11, no. 15, 2022. <https://doi.org/10.3390/electronics11152292>