



In-hospital mortality prediction model of heart failure patients using imbalanced registry data: A machine learning approach

Hadi Sabahi ^a, Mansour Vali ^{a,*}, Davood Shafie ^b

a. Department of Biomedical Engineering, K.N. Toosi University of Technology, Tehran, Iran.

b. Heart Failure Research Center, Cardiovascular Research Institute, Isfahan University of Medical Sciences, Isfahan, Iran.

* Corresponding author: mansour.vali@eetd.kntu.ac.ir (M. Vali)

Received 22 December 2022; Received in revised form 17 May 2023; Accepted 17 July 2023

Keywords

Heart Failure (HF);
In-hospital mortality;
Registry data;
Imbalanced data;
Machine learning.

Abstract

Heart Failure (HF) is a cardiac dysfunction disease with a high mortality rate that is mostly calculated via registry data. The objective of this work was to predict in-hospital mortality in patients hospitalized with HF utilizing their pre-hospitalization registry data. The data include 3968 HF records extracted from the Persian Registry of cardio Vascular disease (PROVE)/HF registry. We proposed a method that contains an imbalanced ensemble probabilistic model which using registry data to predict HF patients who die during hospitalization from those who survive. The suggested ensemble model uses machine learning models, namely Decision Tree, Random Forest, Linear Discriminant Analysis (LDA), Logistic Regression, Support Vector Machine (SVM), K-Nearest Neighbors (KNN), and Extreme Gradient Boosting (XGBoost), which were evaluated. We also used feature importance analysis to identify the important features and reduce the complexity. The results illustrated that the proposed method can predict in-hospital mortality of HF patients using XGBoost that outperformed all others. The feature importance ranking obtained by XGBoost demonstrated that the proposed method can achieve acceptable performance with the first 18 important features and XGBoost (accuracy: $76.4\% \pm 1.6\%$, sensitivity: $76.8\% \pm 6.9\%$, specificity: $76.4\% \pm 1.8\%$). Moreover, statistical analysis presented significant predictors of in-hospital mortality (P -value < 0.01). In conclusion, the proposed method can effectively predict in-hospital mortality of HF patients using the imbalanced data.

1. Introduction

Heart Failure (HF) is one of the prevalent causes of hospitalization and mortality worldwide. Despite advances in diagnosis and treatment, HF still has a high mortality rate, resulting in a growing burden on health providers [1]. Evidence indicates that 40% of hospitalized patients with HF die or are hospitalized again within one year [2]. The increasing mortality of HF has changed it into a life-threatening disease. Therefore, identifying the prevalent predictors and prediction of mortality has been the main focus of current studies [3].

Early recognition of risk factors of the disease can improve prognosis and will be used as predictors of mortality to help in decision-making. Several clinical predictors such as age and depression are associated with increased HF

mortality [2]. There is little research on risk prediction models for elderly patients; however, age is an independent predictor in HF patients [4].

Accurately predicting mortality allows for effective risk classification and provides more appropriate medical care. Calculating the mortality rate of HF around the world is usually based on registry systems [3]. Miró et al. predicted mortality of Acute Heart Failure (AHF) patients using the Epidemiology of AHF in emergency department registry data [5].

The use of machine learning techniques for predicting in-hospital mortality of hospitalized patients has been considered a helpful solution in recent years [6]. Fonarow et al. proposed a regression tree using Acute Decompensated

To cite this article:

H. Sabahi, M. Vali, D. Shafie "In-hospital mortality prediction model of heart failure patients using imbalanced registry data: A machine learning approach", *Scientia Iranica* (2025) 32(17): 7412. <https://doi.org/10.24200/sci.2023.61637.7412>

Heart Failure National Registry (ADHERE) registry data to predict in-hospital mortality probability in patients hospitalized with HF [7]. König et al. developed a reliable algorithm to calculate expected in-hospital mortality in HF cohorts based on routine administrative data by comparing regression analysis with four machine learning models [8]. Luo et al. constructed a risk stratification method using an Extreme Gradient Boosting (XGBoost) algorithm and available clinical data to predict the in-hospital mortality of hospitalized HF patients in Intensive Care Units (ICUs) [9].

Although registry systems provide informative data regarding diseases in society, they are usually divided into imbalanced classes that often result in low sensitivity when ordinary machine learning algorithms are applied. There are many approaches that address imbalanced classification problems. The most common methods consist of oversampling and undersampling, which are relatively able to improve classification performance. The undersampling and ensembling approach has been proven to be advantageous for imbalanced classification [10], in which some classifiers are trained on the minority class and an undersampled majority class. Then, they are combined into an ensemble model. Therefore, the undersampling and ensembling approach could overcome imbalanced classification problems, but their performance is still not suitable for all imbalanced datasets, and they can be improved on a registry dataset for mortality prediction.

Because of the existence of irrelevant and correlated features to the target in actual data, feature importance analysis is usually employed to address dimensionality challenges and to improve the system generalization [11]. Alizadehsani et al. ranked all features of coronary artery disease datasets based on their clinical importance to select the best ones with machine learning techniques [12].

In this research, we want to predict in-hospital mortality of HF patients using their pre-hospitalization imbalanced registry data. To address this issue, we proposed an imbalanced ensemble probabilistic model to predict in-hospital mortality of HF patients using imbalanced registry data. We showed that the proposed model with XGBoost can accurately classify both minority and majority classes with higher performance in comparison with conventional classifiers.

In this research, we investigated the importance of the features to find the influential ones and to reduce the complexity of the proposed model. The use of the identified important features will reduce the cost and time of registering HF patients to predict in-hospital mortality. Furthermore, we also found significant predictors resulting from the statistical analysis of the pre-hospitalization registry data that are helpful for health providers to forecast mortality and better manage resources. In addition, we have used a Decision Tree algorithm to extract special rules from a subset of data.

In the remainder of this paper, we will present the materials and methods, then the results and relevant discussion, and finally, the conclusion will be stated.

2. Material and methods

2.1. Data description

The data for this work included records of patients hospitalized with decompensated or acute HF from March 2015 until October 2018, using data extracted from the Persian Registry Of cardio Vascular disease (PROVE). This is the first registry program for cardiovascular diseases that was

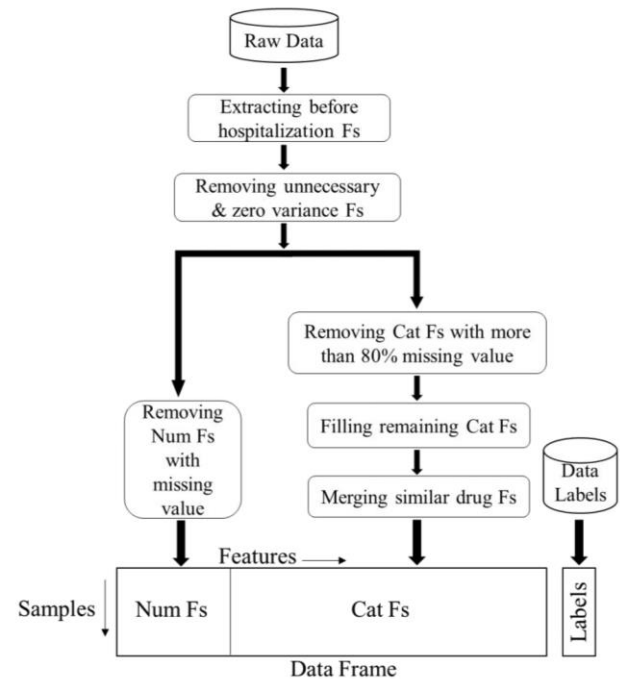


Figure 1. Preprocessing of HF registry data (features, categorical, numerical).

launched as a pilot study in Isfahan (Iran) in 2014. PROVE registry was for patients with stroke, acute coronary syndrome, atrial fibrillation, ST-elevation myocardial infarction (MI), HF, Percutaneous Coronary Intervention (PCI), congenital heart disease, familial hypercholesterolemia, and chronic ischemic cardiovascular disease [13-14]. In this study, informed consent forms were obtained from all patients [14].

PROVE/HF is part of the PROVE registry that registers hospitalized HF patients. The collected data consisted of demographic data, underlying diseases, comorbidities, signs and symptoms, physical examination results, diagnoses, paraclinical tests, treatments, and medications. All the gathered data were related to the period before and during hospitalization as well as the discharge time of the patients. The PROVE/HF registry was followed at 3, 6, and 12 months after the first admission, as needed.

The PROVE/HF registry data included 3968 records belonging to 2918 patients (male: 60.52%, age: 68.97 ± 13.26 years, female: 39.48%, age: 73.27 ± 11.66 years); some patients had more than one admission at different times. Totally, 606 features related to before and during hospitalization, discharge time, and three consecutive follow-ups were registered for each patient.

Given the aim of this study, before-hospitalization features were only used to predict in-hospital mortality of HF patients.

2.2. Preprocessing of data

After data acquisition, preprocessing plays a vital role in data mining, transforming raw data into appropriate forms for subsequent uses. Figure 1 depicts all preprocessing steps of raw PROVE/HF registry data. As mentioned before, only pre-hospitalization features were used in this study to predict in-hospital mortality of patients.

Therefore, these features should first be extracted from the registry. There are many unnecessary features that should

be removed, such as dates of procedures. In addition, we removed some features that were the same for all patients and had no variance. The data features are two types; the first type is categorical, which describes categories or groups such as the “cigar status” of the patient. The second type is numerical, which takes numerical values and represents a measurement such as the “weight” of the patient. Since some categorical features have a lot of missing values, we removed those features that had missing values above an arbitrary threshold depending on the importance of the features. Since we wanted to assess the effect of more categorical features on the prediction results, we removed only the features with more than 80% missing values [15]. We filled the remaining categorical features after consultation with cardiologists and specialists. There were some drugs in the data that belonged to the same type, and we merged them as a feature. We also removed some numerical features with many missing values that did not exist in patients’ medical records. Finally, each sample was labeled according to the mortality status of the patient. If a patient dies during hospitalization, his records are labeled as ‘1’, otherwise, as ‘0’.

After all preprocessing steps, the HF registry data comprise 3252 samples (class ‘0’=3070, class ‘1’=182) and 42 features (categorical=36, numerical=6). The features between patients who died in the hospital and those who survived were compared using the χ^2 test and t -test for categorical and numerical features, respectively. A P -value less than 0.01 was statistically considered significant. Table 1 shows all remaining features after preprocessing, including 8 different groups: Demographic, aetiology, medical history, vital signs, physical examination, procedures, medications, and biomarkers. Numerical features are presented as mean $\pm$ SD (Standard Deviation), and categorical features are shown as n (%) (number (percentage)). Most of the categorical features have two states. For instance, 453 patients in class ‘0’ had Chronic Obstructive Pulmonary Disease (COPD), out of 3070 patients (14.8%), and the others did not. Figures 2 and 3 show bar plots of the categorical features and error bars of the numerical features, respectively.

2.3. Method

In this section, we describe the proposed method to predict in-hospital mortality of HF patients using their pre-hospitalization features of the preprocessed data. The structure of the proposed method is shown in Figure 4.

As reflected in Figure 4, the obtained features should be normalized after preprocessing the raw data. Since there are two types of categorical and numerical features in our data, we used two different methods for each one. Categorical and numerical features are normalized using the min-max scaling and standard scaling methods, respectively [16]:

Min-max scaling:

$$X_n = \frac{X - X_{\min}}{X_{\max} - X_{\min}}. \quad (1)$$

Standard scaling:

$$X_n = \frac{X - X_{\text{mean}}}{SD}. \quad (2)$$

This work aimed to predict in-hospital mortality of HF patients with a low in-hospital mortality rate of 5.6%. Therefore, we encounter an imbalanced classification problem with two classes, in which surviving patients during hospitalization are the majority class. Class imbalance problems frequently occur in the field of medical data processing [17]. Class imbalance causes most algorithms to assign all samples of both classes to the majority one to achieve a high accuracy [18]. We proposed an imbalanced ensemble probabilistic classifier model to distinguish HF patients who die during hospitalization from those who survive using the imbalanced data. Figure 5 illustrates the structure of the imbalanced ensemble probabilistic model.

The proposed method uses the “undersampling and ensembling” strategy to undersample the majority class samples, together with the minority class samples, to train some classifiers [10]. In this method, the majority class samples are undersampled to the same size as the minority class ones. The undersampled data of the majority class are not put back again. Whenever the number of majority class samples is not enough to undersample, the needed number of samples is randomly selected from the majority class sample subsets. Hence, the structure of the imbalanced ensemble probabilistic model is created using each undersampled subset of the majority class in conjunction with all samples of the minority class. The total number of created training subsets is:

$$N = \left\lceil \frac{n_j}{n_m} \right\rceil + 1,$$

where n_j and n_m indicate the number of majority and minority class samples in the training set, respectively. Each training subset is used to train a classifier, and each classifier has its own hyperparameters. Most machine learning models have parameters known as hyperparameters [19] that need to be fixed before training [20]. As Figure 5 shows (Clf tuners), this work tunes the hyperparameters of each classifier based on the corresponding training subset and then trains them using the best-found hyperparameters and the training subset. To find the best hyperparameters, we used the basic grid search technique, in which all possible permutations of the hyperparameters of a model are applied to build the models. The best model is selected after the evaluation of the performance of each one. To evaluate the built model, the grid search technique uses the 5-fold cross-validation method, which is a resampling method used for evaluating machine learning models [21]. In addition, we designated ‘Accuracy’ as an evaluation strategy for the performance of the cross-validated model on the test set. Figure 6 summarizes the grid search to find the best hyperparameters of the models.

Following training each classifier with the best found hyperparameters and the training subset, the trained classifier predicts the probability of test set samples. Afterward, the mean probability of the output of all classifiers is computed for each test set sample. Then, a decision threshold of 0.5 is considered to calculate the predicted label for each test set sample, where probability values less than 0.5 are assigned to class ‘0’, otherwise to class ‘1’.

To assess the performance of the model, some metrics are calculated as follows:

$$\text{Accuracy} = \frac{TN + TP}{TN + TP + FN + FP}, \quad (3) \quad \text{Specificity} = \frac{TN}{TN + FP}, \quad (5)$$

$$\text{Sensitivity} = \frac{TP}{TP + FN}, \quad (4) \quad F1\text{-score} = \frac{2TP}{2TP + FP + FN}, \quad (6)$$

Table 1. Features of the preprocessed HF registry data, including categorical (*n* (%)) and numerical (mean $\pm$ SD) features in 8 different groups. *P*-value<0.01 is considered statistically significant. Index: MI of a patient, COPD, SBP, DBP, JVP, CPO, Coronary Artery Bypass Grafting (CABG), PCI, CRT-D, ICD, NIV, Hb, BUN.

#	Group features	Type	Features	Class '0', <i>n</i> =3070	Class '1', <i>n</i> =182	<i>P</i> -value
1	Demographic	Categorical	Index	650(21.1)	47(25.8)	0.516
2	Aetiology	Categorical	Primary hypertension	2072(67.5)	116(63.7)	0.294
3		Categorical	Primary heart ischemic	2623(85.4)	148(81.3)	0.128
4		Categorical	Primary valvular heart	1234(40.2)	88(48.4)	0.03
5	Medical history	Categorical	Hypertension	2063(67.2)	117(64.3)	0.417
6		Categorical	Arrhythmia	574(18.7)	39(21.4)	0.36
7		Categorical	Diabetes	1454(47.4)	85(46.7)	0.863
8		Categorical	COPD	453(14.8)	29(15.9)	0.664
9		Categorical	Thyroid	224(7.3)	22(12.1)	0.018
10		Categorical	Stroke	137(4.5)	14(7.7)	0.044
11		Categorical	Anemia	332(10.8)	51(28)	<0.001
12		Categorical	Kidney Disease	789(25.7)	73(40.1)	<0.001
13	Vital sign	Numerical	SBP	131.99 $\pm$ 26.95	115.81 $\pm$ 25.65	<0.001
14		Numerical	DBP	80.77 $\pm$ 15.86	74.26 $\pm$ 15.7	<0.001
15		Numerical	HR	88.43 $\pm$ 20.9	92.73 $\pm$ 25.92	0.008
16	Physical examination	Categorical	Edema	1464(47.7)	106(58.2)	0.006
17		Categorical	JVP	594(19.3)	51(28.0)	0.004
18		Categorical	Crackle	1992(64.9)	140(76.9)	0.001
19		Categorical	CPO	62(2.0)	12(6.6)	<0.001
20	Procedures	Categorical	CABG	404(13.2)	27(14.8)	0.517
21		Categorical	PCI	770(25.1)	29(15.9)	0.005
22		Categorical	CRT-D	47(1.5)	4(2.2)	0.482
23		Categorical	ICD	183(6.0)	11(6.0)	0.963
24		Categorical	Hemodialysis	85(2.8)	21(11.5)	<0.001
25		Categorical	NIV	73(2.4)	65(35.7)	<0.001
26	Medications	Categorical	Captopril	424(13.8)	23(12.6)	0.655
27		Categorical	Losartan	1348(43.9)	56(30.8)	0.001
28		Categorical	Metoral	1349(43.9)	69(37.9)	0.111
29		Categorical	Hydrochlorothiazide	150(4.9)	19(10.4)	0.001
30		Categorical	Furosemide	1399(45.6)	92(50.5)	0.190
31		Categorical	Spironolactone	779(25.4)	47(25.8)	0.892
32		Categorical	Digitalis	832(27.1)	51(28.0)	0.786
33		Categorical	Atorvastatin	1336(43.5)	64(35.2)	0.027
34		Categorical	Nitrocountine	1357(44.2)	72(39.6)	0.220
35		Categorical	Warfarin	590(19.2)	35(19.2)	0.997
36		Categorical	ASA	1759(57.3)	79(43.4)	<0.001
37		Categorical	Plavix	453(14.8)	26(14.3)	0.862
38	Biomarker	Numerical	Hb	11.70 $\pm$ 4.39	11.11 $\pm$ 4.61	0.081
39		Numerical	Creatinine	1.476 $\pm$ 1.26	1.963 $\pm$ 1.36	<0.001
40		Numerical	BUN	25.85 $\pm$ 19.08	40.1 $\pm$ 29.02	<0.001
41	Biomarker	Categorical	Troponin	0 = Not done	514(16.7)	41(22.5)
				1 = Positive	333(10.8)	54(29.7)
				2 = Negative	2223(72.4)	87(47.8)
42	Demographic	Categorical	Sex	Male	1914(62.3)	106(58.2)
				Female	1156(37.7)	76(41.8)

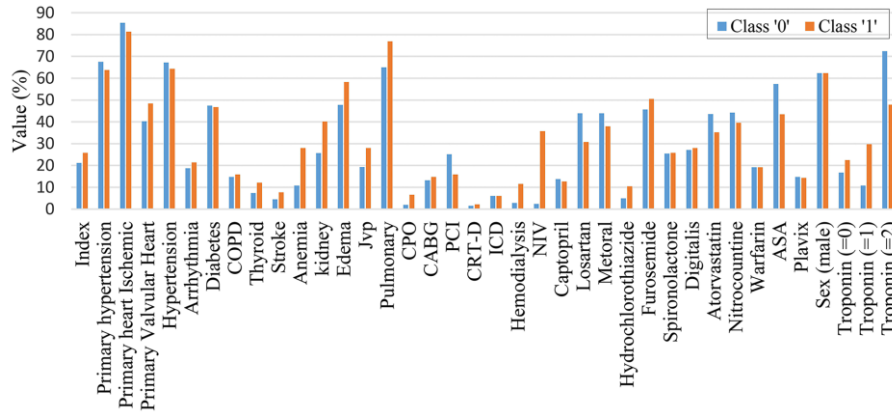


Figure 2. Categorical features.

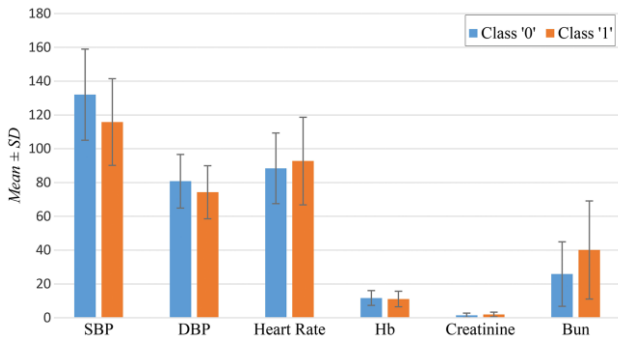


Figure 3. Numerical features (standard deviation).

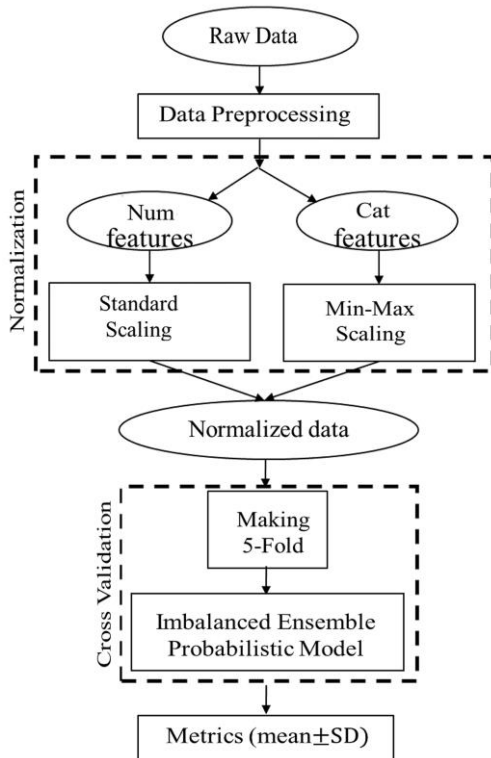


Figure 4. The overall structure of the proposed method (standard deviation).

$$\text{Matthews Correlation Coefficient (MCC)} = \frac{TP \times TN - FP \times FN}{\sqrt{(TP + FP)(TP + FN)(TN + FP)(TN + FN)}}, \quad (7)$$

where True Negative (TN) and True Positive (TP) respectively show the numbers of negative and positive samples that are correctly diagnosed. Moreover, False Negative (FN) and False Positive (FP) respectively represent the numbers of negative and positive samples that are not correctly diagnosed. In addition, we used two other valuable metrics, which include ROC_auc , which calculates the area under the Receiver Operating Characteristic (ROC) curve, and PR_auc , which calculates the area under the PR curve from predicted labels [22].

According to our proposed model in Figure 4, we used the 5-fold cross-validation method with 10 repetitions to evaluate our model. The final results are computed using the average of all different metrics for model evaluation. The most crucial advantage of the 5-fold cross-validation method is its lower variance than the single hold-out set method. Therefore, its sensitivity to any partitioning bias on the dataset is less than the single hold-out set method. Besides, 5-fold cross-validation is a more robust method than the single hold-out method that randomly splits data into training and testing sets [23].

In practice, machine learning models have different performances for various datasets because their characteristics are different. Therefore, in this study, we evaluated and compared several models to select the one that has the best performance for our purpose. In particular, the seven different evaluated classification models are Decision Tree, Random Forest, Linear Discriminant Analysis (LDA), Logistic Regression, Support Vector Machine (SVM), K-Nearest Neighbor (KNN), and XGBoost [6,8,15].

The model was fitted with the best-found hyperparameters using the described method above. The technical hyperparameters of the considered models to find the best ones using the basic grid search technique are listed in Table 2. Additionally, we performed hierarchical clustering analysis over models based on the FN and FP values.

2.4. Feature importance analysis

In this study, we used feature importance analysis to reduce the number of features and complexity of the proposed model. Feature importance includes techniques that assign a score to each input feature based on how they are effective at the classification performance of a target variable [24]. There are many types of importance scores such as statistical

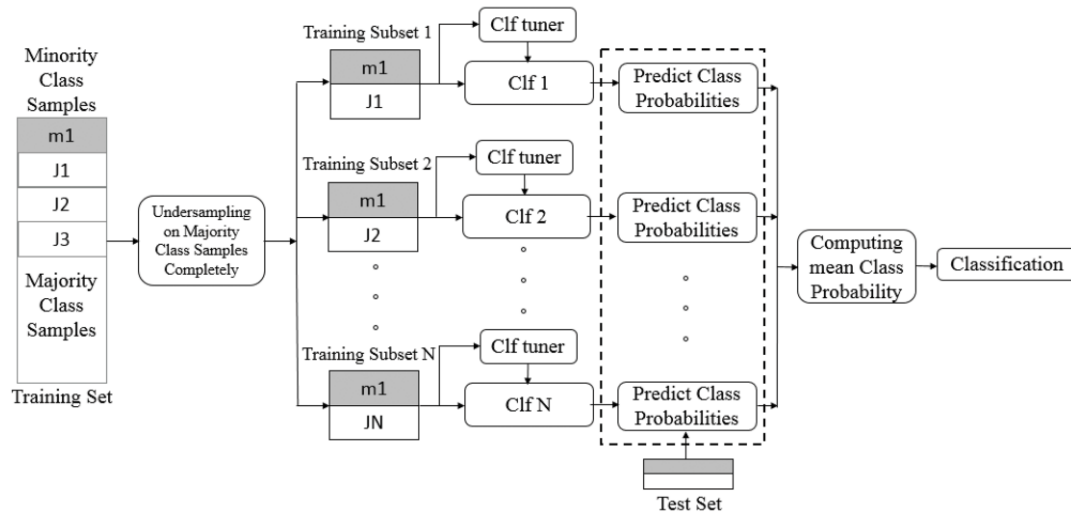


Figure 5. The structure of the imbalanced ensemble probabilistic model (classifier).

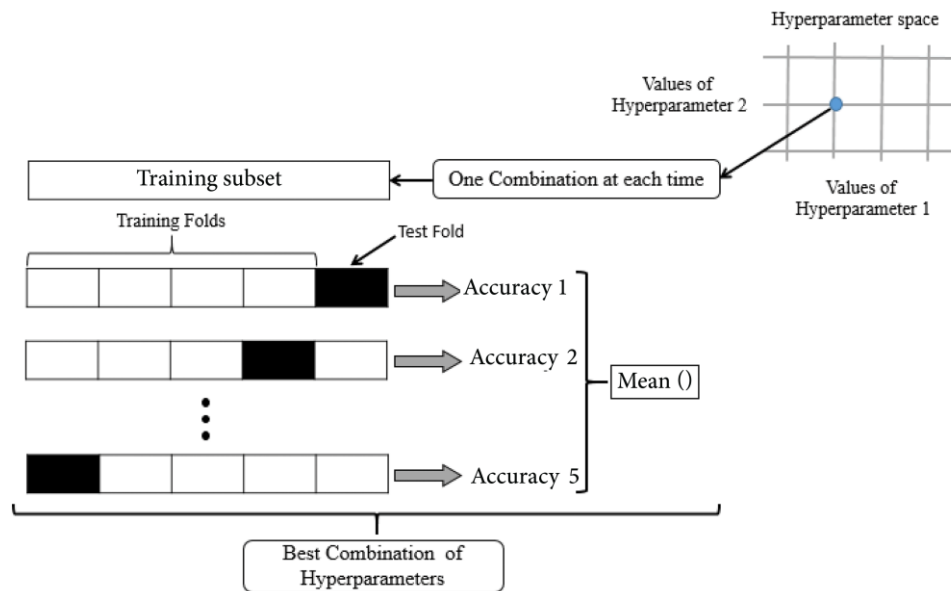


Figure 6. The grid search for hyperparameter tuning with 5-fold cross-validation.

Table 2. Technical hyperparameters of the classification models to find the best ones using the basic grid search technique.

Model	Hyperparameters
	Criterion (to measure the split quality)
Decision Tree	Min_samples_split (minimum number of samples to split an internal node)
	min_samples_split
Random Forest	Max_features (the number of features for the best split)
LDA	Solver (algorithm to use in optimization problem)
Logistic Regression	Solver C (regularization)
SVM	Kernel, C (regularization)
KNN	k (number of neighbors) Weights (weight function) Metric (distance metric)
XGBoost	n_estimators (number of boosting stages) learning_rate Subsample (the fraction of samples to fit the individual learners) max_depth (the maximum depth of individual learners)

Table 3. The performance of the proposed model on the test sets for various classifiers used to predict the in-hospital mortality of HF patients, mean $\pm$ SD (%) (standard deviation).

Model	Accuracy	Sensitivity	Specificity	F1-score	ROC_auc	PR_auc	MCC
Decision Tree	76 $\pm$ 2.2	75.3 $\pm$ 6.4	76 $\pm$ 2.5	26 $\pm$ 2.2	83.4 $\pm$ 2.6	31.3 $\pm$ 6	26.7 $\pm$ 3.2
Random Forest	75.2 $\pm$ 1.9	76.8 $\pm$ 6	75.1 $\pm$ 2	25.8 $\pm$ 2	84.2 $\pm$ 2.6	32.1 $\pm$ 6	26.7 $\pm$ 3
LDA	77.7 $\pm$ 1.4	66.9 $\pm$ 8.1	78.3 $\pm$ 1.5	25.1 $\pm$ 2.5	81.4 $\pm$ 3.3	29.5 $\pm$ 6	24.3 $\pm$ 3.9
Logistic Regression	77 $\pm$ 1.4	71.2 $\pm$ 8.2	77.3 $\pm$ 1.6	25.7 $\pm$ 2.4	82.8 $\pm$ 3.1	27.7 $\pm$ 6	25.6 $\pm$ 3.9
SVM	75.5 $\pm$ 1.7	72.8 $\pm$ 7.5	75.7 $\pm$ 2	25 $\pm$ 2.1	82.5 $\pm$ 2.9	27.7 $\pm$ 6	25.1 $\pm$ 3.4
KNN	83.7$\pm$1.6	52.3 $\pm$ 8.2	85.5$\pm$1.9	26.4 $\pm$ 3.4	77.9 $\pm$ 3.1	22.5 $\pm$ 5	23.4 $\pm$ 4.4
XGBoost	76.7 $\pm$ 1.9	77.3$\pm$6.9	76.6 $\pm$ 2.2	27.1$\pm$2.1	84.7$\pm$2.9	34.6$\pm$6.7	28.2$\pm$3.1

correlation coefficient scores, Decision Trees, and permutation scores. Decision Tree algorithms suggest importance scores based on the decrease in the criterion of the split points, like Entropy or Gini. This approach can be applied to ensembles of trees such as the Random Forest and XGBoost algorithms. The models Decision Tree, Random Forest, and XGBoost allowed the importance of features to be derived during the training of the models. Linear machine learning algorithms such as Logistic Regression calculate coefficient statistics of each feature and target variable in order to apply in a weighted sum to make a prediction. These coefficients can be utilized directly as a feature importance score. The models LDA, SVM, and KNN use permutation importance scores, where feature importance is difficult to extract, and they do not support native importance scores. Briefly, the model is trained on the data, then it is applied to classify the data while the values of a feature have been scrambled. This is repeated for every feature, and the whole process is repeated several times. The result would be a mean importance score for every feature. The accuracy of the model is considered as a basis for the importance score. Obviously, the effect of scrambling the feature values is small for unimportant features but is considerable for important ones, which will reduce the model's accuracy.

2.5. Rule extraction

Finally, in this work, we extracted some significant rules using the Decision Tree. Due to the imbalanced data, the majority class samples are undersampled to match the minority class ones. The Decision Tree model is fitted on the constructed data so that both classes have the same size, and then all possible rules are extracted. In order to extract the confident rules, we chose the rules with an accuracy of 100% that are supported by at least ten samples.

All of the experiments of the current study were conducted using the scikit-learn [25] machine learning library in Python (version 3.7.0) and SPSS Statistics for Windows, version 15 (SPSS Inc., Chicago, Ill., USA).

3. Results

At first, we illustrate the results of the statistical analysis using the X^2 test for categorical features and the t -test for numerical features in patients who died in hospital (class '1') and those who did not (class '0'). Table 1 lists all used features with their P -value which less than 0.01 considered statistically significant. Numerical features are presented as mean $\pm$ SD, and categorical features are shown as numbers and percentages. All features with P -value<0.01 in Table 1 are statistically significant and can be considered as predictors of in-hospital mortality for HF patients.

After the statistical analysis of the features, we present the obtained results of the proposed method to predict the in-hospital mortality of HF patients using preprocessed imbalanced registry data. As mentioned before, we used the 5-fold cross-validation method with 10 repetitions to evaluate our model, while each set includes the same percentage of each target class as the complete dataset. Then, each set is given to the imbalanced ensemble probabilistic model as the input for training and testing. According to the number of training samples, the total number of created training subsets equals 17, where each training subset contains 292 samples, equally split between both classes. Table 3 shows the performance of the proposed model on the test sets for all classifier models used to predict the in-hospital mortality of HF patients.

As Table 3 demonstrates, using KNN as a classifier in the structure of the imbalanced ensemble probabilistic model of the proposed method gives the best accuracy and specificity (83.7% and 85.5%, respectively). In this case, however, the corresponding sensitivity has the lowest value (50%). XGBoost achieves the best sensitivity, $F1$ -score, ROC_auc, PR_auc, and MCC of 77.3%, 27.1%, 84.7%, 34.6%, and 28.2%, respectively. According to Table 3, XGBoost has the highest number of top metrics and, therefore, it outperforms all other classifiers. The average ROC and Precision-Recall (PR) curves with the 5-fold data resampling and 10 repetitions are depicted in Figure 7. Most classifiers have AUC values above 80%, but the value of KNN is lower (78%). We also used the AUC value as the criterion for the PR curve. The lowest and the highest AUC values of AUC of PR curves are for KNN and XGBoost, respectively (22%, 35%).

In this study, the hierarchical clustering analysis was applied to cluster the seven classifiers using the FN and FP values from a random sampling. Figure 8 shows the hierarchical clustering analysis that represents similar classifiers culminating in similar results; for instance, tree-based classifiers Decision Tree, Random Forest, and XGBoost are clustered closely.

As previously mentioned, we used feature importance analysis to reduce the number of features and complexity of the proposed model in the current study. We applied three different types of importance scores, including tree-based scores for the Decision Tree, Random Forest, and XGBoost models (Figure 9), statistical correlation coefficient scores for the Logistic Regression model (Figure 10), and permutation scores for LDA, SVM, and KNN models (Figure 11).

As can be seen, each classifier model creates its feature importance scores based on the corresponding method. In order to reduce the dimension of the features and complexity of the proposed model, we should use a specific model to compute its feature importance scores. According to Table 3,

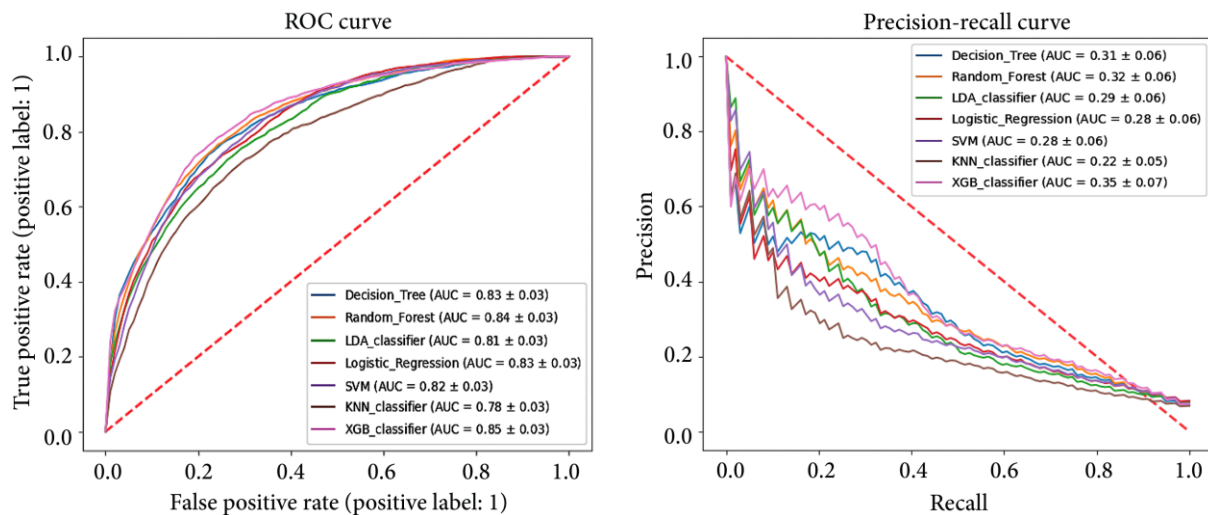


Figure 7. ROC curve (left) and PR curve (right) of the proposed model.

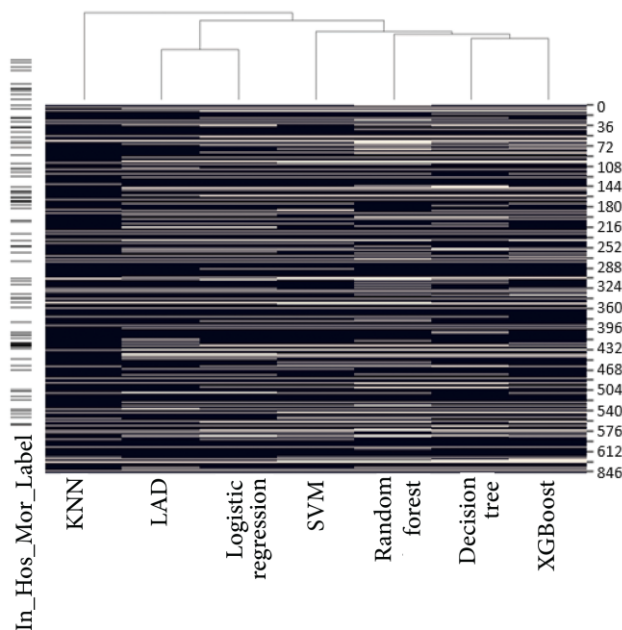


Figure 8. Hierarchical clustering analysis of the classifier models.

XGBoost outperforms all other models; therefore, it can be used to find the important features. Then, we apply different numbers of sorted important features to the proposed model with XGBoost and figure out the ROC_auc metric. The result is presented in Figure 12.

As Figure 12 shows, at first, the ROC_auc increases with the number of important features. But after 18 features, the ROC_auc does not have significant changes. Therefore, we can consider that at least the first 18 important features are required to achieve acceptable model performance. The list of the first 18 important features can be found from the feature importance diagram of XGBoost in Figure 9. Table 4 shows the results of using the first 18 important features as the input of the proposed model for all classifiers used to predict the in-hospital mortality of HF patients.

According to Table 4, XGBoost achieves the best ROC_auc, PR_auc, and MCC of 84.9%, 34.6%, and 27.7%, respectively. Therefore, XGBoost has the highest number of top metrics and outperforms all other classifier models with the first 18 important features. A comparison of the results of the proposed model with XGBoost in both cases (using all of

the features (Table 3) or the first 18 important features (Table 4)) indicates that the model can perform slightly better with all features, but the improvement is not significant. Therefore, in order to reduce the complexity of the model, we can only use the first 18 important features that are extracted by XGBoost.

Finally, we describe the significant extracted rules. Based on the Decision Tree, seven rules, which are illustrated as IF (Antecedent) and Then (Consequent) in Table 5, were generated with an accuracy of 100% and at least ten samples. The presence of high Blood Urea Nitrogen (BUN), low Heart Rate (HR), low Hemoglobin (Hb), and None-Invasive Ventilation (NIV) usage was associated with HF patients who died during hospitalization. On the other hand, the normal range of BUN, Hb, Systolic Blood Pressure (SBP), and not using NIV was associated with those who did not die during hospitalization. However, more investigation with larger data sets and more features is still required.

4. Discussion

The present research applied the data of the first Iranian national registry of cardiovascular diseases and, besides the statistical analysis of the significant features, proposed a model to predict the in-hospital mortality of HF patients. As mentioned before, some of the features were statistically significant ($P\text{-value} < 0.01$). The results of the statistical analysis are in line with previous related studies reporting the predictors of mortality rate and morbidity in HF patients. According to Table 1, anemia and kidney disease have a significant relationship to in-hospital mortality, which is consistent with the previous findings in other countries [26]. The effect of anemia on HF mortality has been clarified in several studies that have recommended treatment of anemia as a preventive factor to reduce HF mortality [27]. According to Table 1, the average Systolic and Diastolic Blood Pressures (SBP and DBP) of patients who died during hospitalization are lower than those who did not die. This issue induces hypotension (low blood pressure) in patients with severe HF. The higher mortality rate of HF patients with hypotension is consistent with our result. It is completely a logical finding because the low blood pressure in HF disease is related to cardiogenic shock, which shows pump failure of the heart and a lower Left Ventricular Ejection Fraction (LVEF) [26]. On the other hand, if the blood pressure is low,

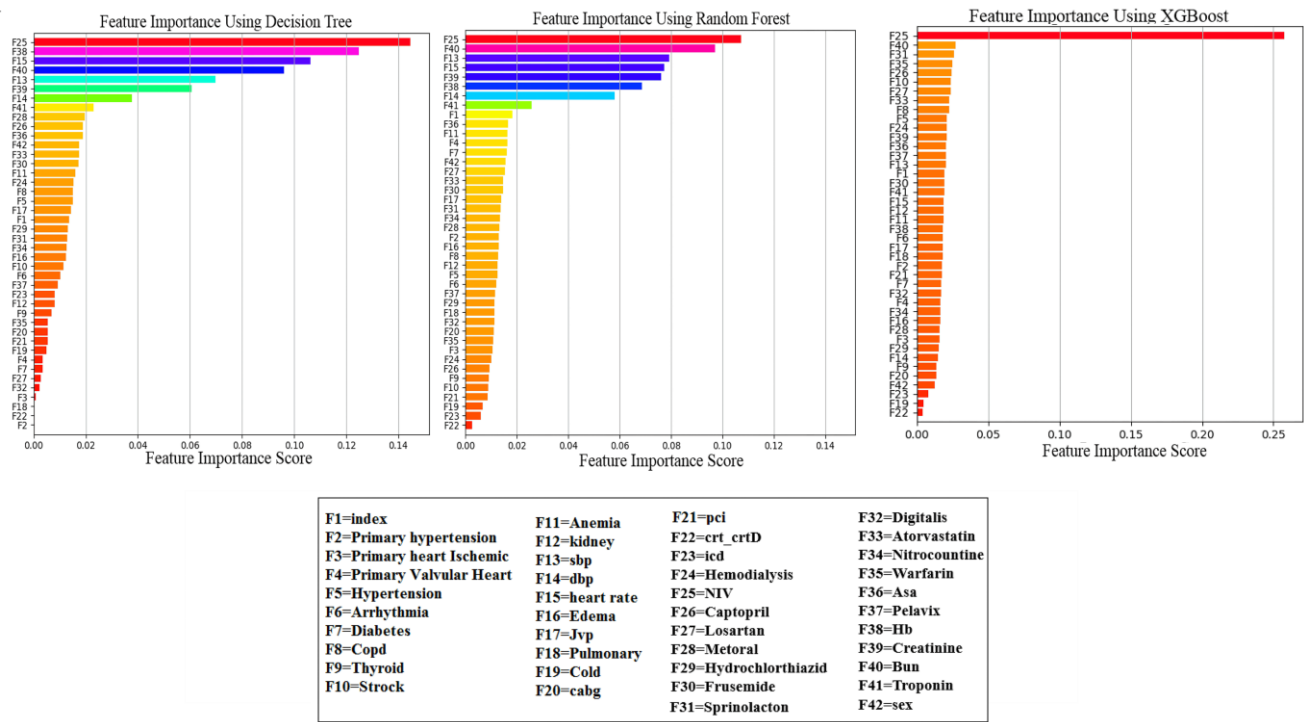


Figure 9. Feature importance bar chart of the Decision Tree, Random Forest, and XGBoost models.

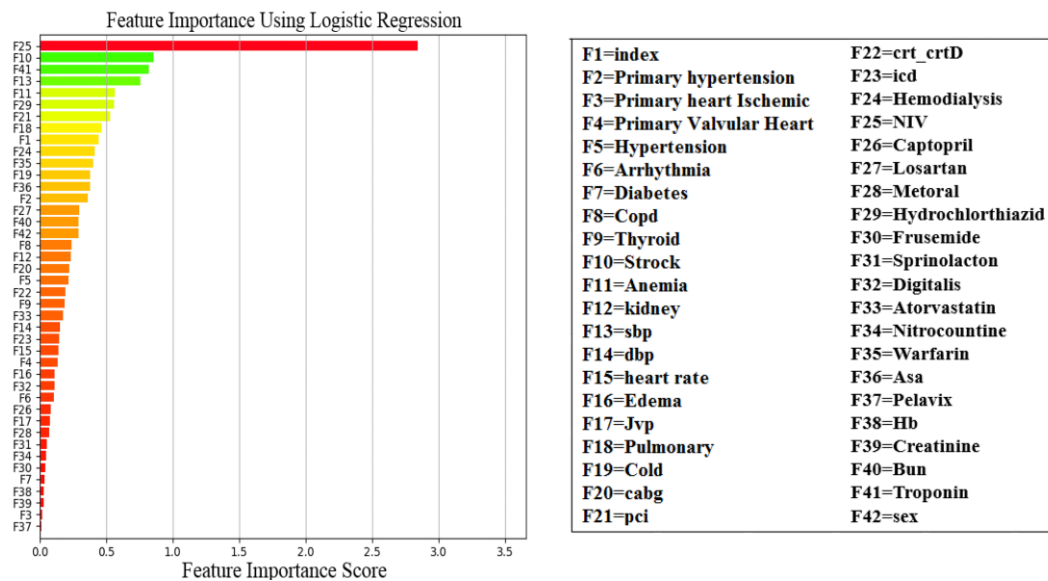


Figure 10. Feature importance bar chart of the Logistic Regression model.

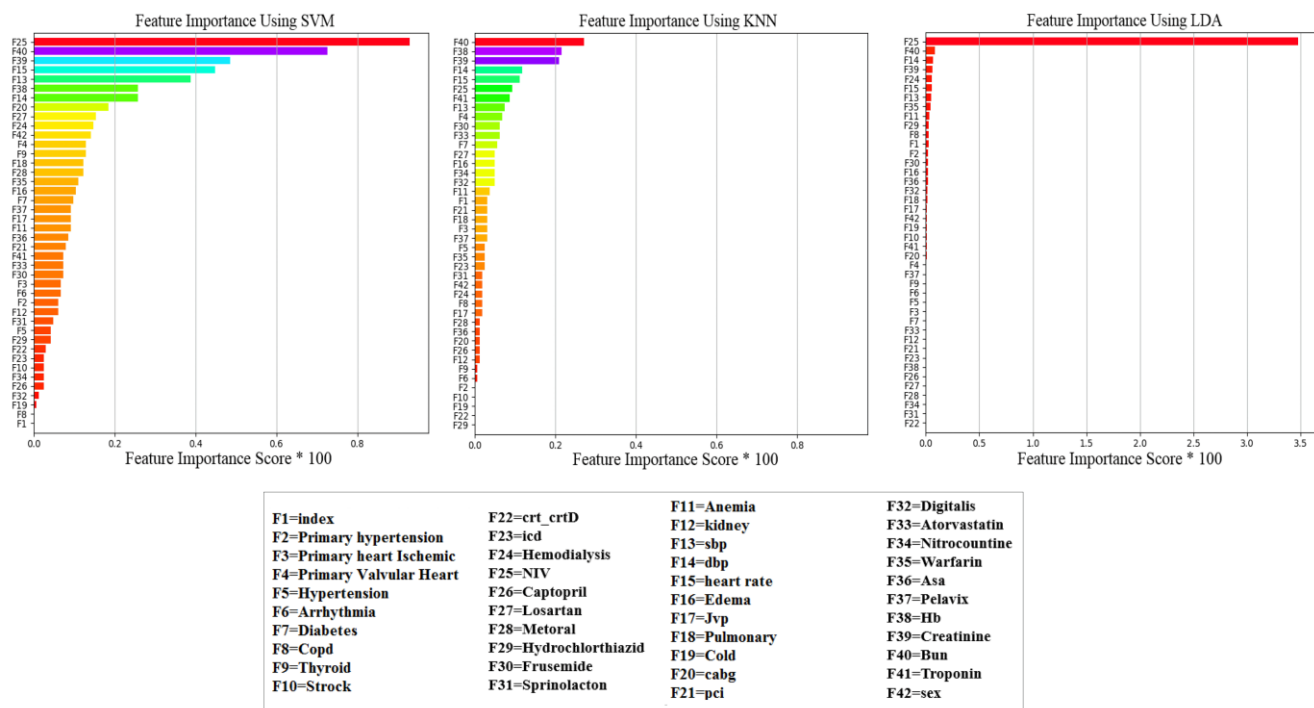
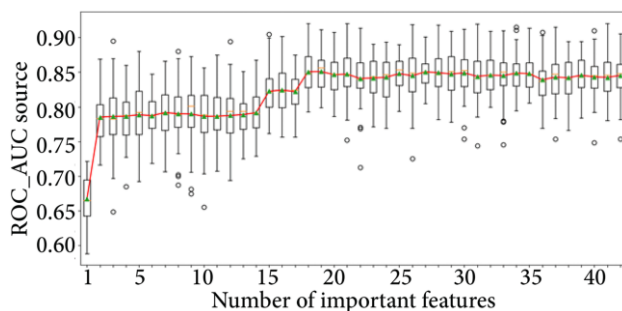
the heart will struggle to deliver enough oxygenated blood to the cells; therefore, the body will increase the HR to push more oxygen-rich blood to the cells [28]. Our finding of the HR of severe HF patients justifies this result again. In spite of advances in technology, physical examination remains essential in the management of HF patients. Edema, Jugular Venous Pulse (JVP), Crackles, and Cold Peripheral Organs (CPO) are all statistically significant physical examinations of the current study that can be considered as predictors of in-hospital mortality. Patients with any type of procedure in their medical history are mostly at a severe stage of the disease.

In these patients, the higher mortality rate is a sensible result, demonstrated by a statistically significant relationship to PCI, hemodialysis, and NIV therapy. The investigation of

the predictive role of medications shows that losartan as an Angiotensin Receptor Blocker (ARB) and ASA as an antiplatelet were more often used in the patients who did not die during hospitalization, while Hydrochlorothiazide as a diuretic was utterly vice versa. Therefore, losartan and Acetyl Salicylic Acid (ASA) can be considered preventive factors for mortality. There is also disagreement on the role of diuretic medications on the mortality of HF patients in the results of various studies. Some studies suggested higher mortality rates of HF patients with diuretic use, which is in line with the outcome of the current research[29], while others offered protective effects [30]. The effect of diuretics on the mortality of HF patients has to be further studied. The crucial role of biomarkers is increasingly recognized in HF management, diagnosis, and screening of severe patients [31].

Table 4. The performance of the proposed model on the test sets with the first 18 important features for various classifiers used to predict the in-hospital mortality of HF patients, mean $\pm$ SD (%) (standard deviation).

Model	Accuracy	Sensitivity	Specificity	F1-score	ROC_auc	PR_auc	MCC
Decision Tree	76.5 $\pm$ 1.8	74.9 $\pm$ 7.9	76.5 $\pm$ 2	26.3 $\pm$ 2.4	84 $\pm$ 3.5	33.1 $\pm$ 6.6	26.9 $\pm$ 3.8
Random Forest	75.9 $\pm$ 1.9	77.6$\pm$7.1	75.8 $\pm$ 2.1	26.5 $\pm$ 2.1	84.9 $\pm$ 3.1	34.3 $\pm$ 6.5	27.7 $\pm$ 3.3
LDA	80.2 $\pm$ 1.6	65.8 $\pm$ 8.1	81.1 $\pm$ 1.9	27.1 $\pm$ 2.4	83.3 $\pm$ 3.3	31.6 $\pm$ 6.5	26.2 $\pm$ 3.6
Logistic Regression	78.7 $\pm$ 1.5	71.6 $\pm$ 7.2	79.2 $\pm$ 1.8	27.4$\pm$2.1	83.8 $\pm$ 3.3	30.7 $\pm$ 6	27.5 $\pm$ 3.3
SVM	77.4 $\pm$ 1.7	74.2 $\pm$ 7.3	77.6 $\pm$ 2	26.9 $\pm$ 1.9	83.9 $\pm$ 3	28.8 $\pm$ 5.7	27.4 $\pm$ 3.1
KNN	82.9$\pm$1.4	57.3 $\pm$ 7.3	84.4$\pm$1.5	27.3 $\pm$ 3	80.7 $\pm$ 3.2	24.7 $\pm$ 5.4	25 $\pm$ 4
XGBoost	76.4 $\pm$ 1.6	76.8 $\pm$ 6.9	76.4 $\pm$ 1.8	26.7 $\pm$ 1.9	84.9$\pm$2.8	34.6$\pm$7.5	27.7$\pm$3.1

**Figure 11.** Feature importance bar chart of the SVM, KNN, and LDA models.**Figure 12.** ROC_auc metric of the proposed model with XGBoost for different numbers of important features.

An increased level of serum creatinine during HF hospitalization is associated with worse outcomes [32]. This issue is in line with our result about the higher creatinine levels in patients who died in hospital. A persistently high level of BUN is also associated with an increased risk of cardiovascular readmission and death [10], which is close to

our result according to Table 1. Another important primary biomarker is cardiac troponin, whose level can be elevated in HF patients [33]. This issue is demonstrated in our findings in Table 1 which patients who died in hospital have higher troponin than others. In the current data, the in-hospital mortality rate of HF patients is 5.6%, which is in the range of several published registries (4%-7%) [34,35]. To reduce the mortality of HF patients, on the one hand, healthcare staff and physicians should pay more attention to the predictors and treatment of underlying diseases (such as anemia and hypotension). On the other hand, patients should ADHERE to medications (especially ARB and ASA).

Besides the statistical analysis of the features, we proposed a new model to predict in-hospital mortality of the HF patients using the imbalanced registry data. Most algorithms frequently obtain poor performance with imbalanced datasets because they tend to get high accuracy and assign the most samples to the majority class, which causes low sensitivity.

Table 5. Significant extracted rules with an accuracy of 100% and at least ten samples using the Decision Tree.

Antecedent	Consequent
(NIV=No), (Bun<25), (Troponin>1.5), (Creatinine<=1.45), (SBP>110.32), (Hb<=15.39)	Class '0'
(NIV=Yes), (DBP>40.44), (10.49< Hb<=18.44), (Bun<84)	Class '1'
(NIV=No), (Bun>25), (SBP<=119.81), (3.69<Hb<17), (Troponin<=1.5)	Class '1'
(NIV=No), (Troponin<=1.5), (SBP>132.31), (Bun<=13.12), (Heart Rate<=165.61)	Class '0'
(NIV=Yes), (DBP>40.44), (Bun<84), (Hb<10.5), (Troponin<=1.5)	Class '1'
(NIV=No), (SBP<=99.83), (Heart Rate>52.6), (Bun<=20.52), (Hb<=15.74)	Class '1'
(NIV=No), (SBP<=119.8), (Hb<17), (Troponin>1.5), (heart rate<=122.6), (Bun>30.2), (DBP>64.4)	Class '1'

There are many approaches that address imbalanced classification problems. Oversampling and undersampling are the most commonly used approaches. These methods improve the overall performance of the classification. However, Oversampling may increase the likelihood of occurring overfitting, especially at higher rates of oversampling. Furthermore, it will increase the computational effort and decrease the classifier performance. In undersampling, a large number of data points are discarded. This can be very problematic, as the elimination of such data may make the decision boundary between majority and minority classes harder to learn, resulting in high variance and performance loss.

The undersampling and ensembling approach was shown to be more effective than others for imbalanced classification [10], which trains several classifiers using the minority and the undersampled majority class samples and then combines the output of classifiers into an ensemble structure. Ensemble methods will reduce the variance of the results by aggregating the prediction performance of the classifiers [18]. Inspired by this approach, we proposed a new ensemble model to predict the in-hospital mortality of HF patients using imbalanced registry data.

In the suggested model, the class probability of each test set sample is calculated after training the classifier. Then, the mean class probability of all classifiers is computed for labeling the sample. In the proposed model, we compute the class probability of each test set sample instead of predicting the class label for each classifier directly. This method will provide a more accurate class probability for each test set sample and, therefore, will reduce the *FP* and *FN* that will cause an increase in the performance of the classification.

In the proposed model, different classifiers are used, and each one has its own performance. According to Tables 3 and 4, although KNN has the highest accuracy and specificity among all used classifiers, its main drawback is the great difference between sensitivity and specificity, which shows it cannot equally classify both minority and majority classes. According to Table 3, XGBoost has the highest sensitivity with all features, but according to Table 4, Random Forest has the highest sensitivity with only the first 18 important features. These tables demonstrate that the *F1-scores* of all classifiers are close to each other, but XGBoost has the highest *F1-score* when all features are used, and Logistic Regression has the highest one when only the first 18 important features are used. According to Tables 3 and 4, Random Forest and XGBoost have the best ROC_auc, PR_auc, and MCC among all classifiers, and their values are so close to each other. However, XGBoost is preferred because its values are a bit better than Random Forest. In

XGBoost, the variance of ROC_auc and MCC is lower, and the mean of PR_auc is higher than the values of Random Forest. Also, the sensitivity and specificity of XGBoost are so close to each other that it can classify both classes similarly. According to Tables 3 and 4, XGBoost has the highest number of top metrics and, therefore, it outperforms all other classifiers.

The number of training samples is a fundamental determinant of classification accuracy, and they are correlated with classification accuracy [36]. In the current study, we assessed the effect of the training sample size on the classification performance for all applied classifier models.

To evaluate the effect of the training size on the model performance, we examined the proposed model for different numbers of training sets that are applied in the structure of the imbalanced ensemble probabilistic model (Figure 5). As Figures 13-15 demonstrate, the more training sets, the better performance the proposed model can achieve for all classifier models. A larger training sample size improves the learning process, and classifiers can classify the new samples with more accuracy. Therefore, if we have more data, we will obtain higher classification performance. In addition, these figures restate that KNN has the lowest sensitivity for different numbers of training sets despite having the best accuracy and specificity; therefore, KNN is not a proper classifier for our purpose. According to Figures 13-15, the SD of sensitivity is greater than the SD of accuracy and specificity, which is because our data are highly imbalanced and the number of patients who died in hospital (class '1') is much less than the number of survivors (class '0').

Feature importance provides insight into the model and data, and it is the basis for feature selection, which can improve the performance of a predictive model. In short, the feature Importance score is used to perform feature selection. Supervised feature selection methods include filter methods, wrapper methods, and embedded methods; each one includes different techniques. The filter methods choose the best subset of the feature space immediately before feeding it to a learning algorithm. The remaining two approaches, embedded and wrapper, create the optimal subset of features in conjunction with the learning algorithm. Contrary to the other methods, the embedded methods put together the advantages of both the wrapper and filter methods. Dissanayake and Md Johar have conducted an experimental evaluation of the performance of models created for heart disease prediction using various feature selection techniques such as ANOVA, Chi-square, mutual information, Relief algorithm, forward and backward feature selection, and so on [37]. Finally, the feature subset achieved by the backward technique that belongs to wrapper methods led to the highest classification accuracy. In this study, we

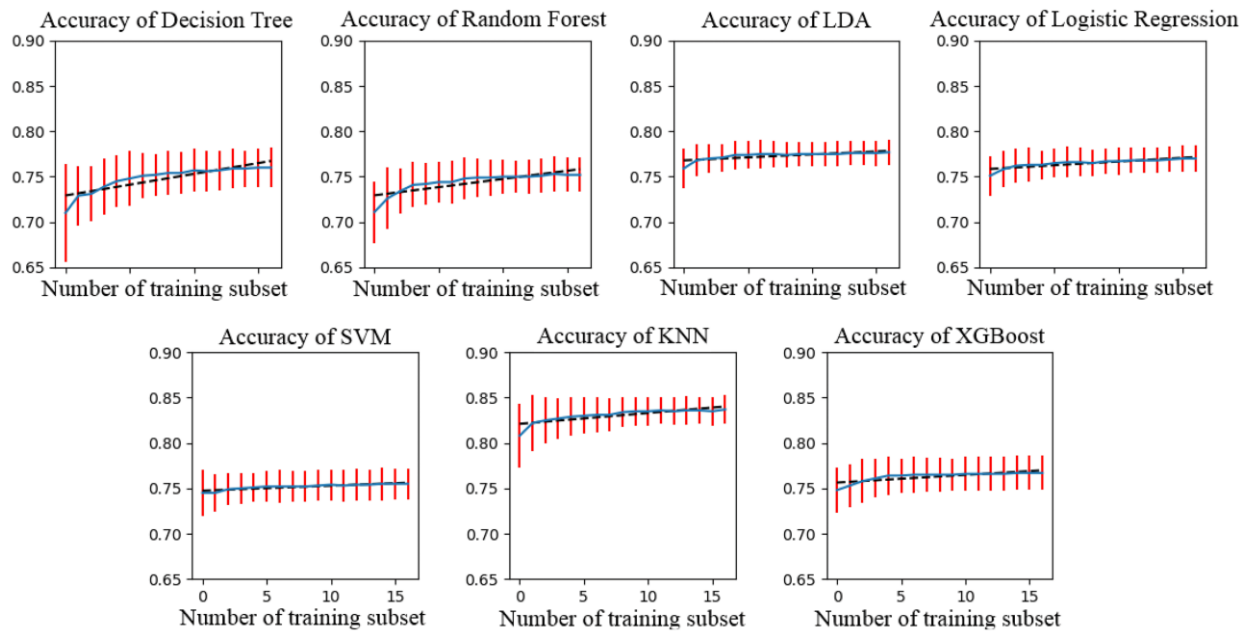


Figure 13. Accuracy of the proposed model and the number of the training subsets for different classifiers. Dashed lines: Linear regression, Vertical lines: Standard deviation.

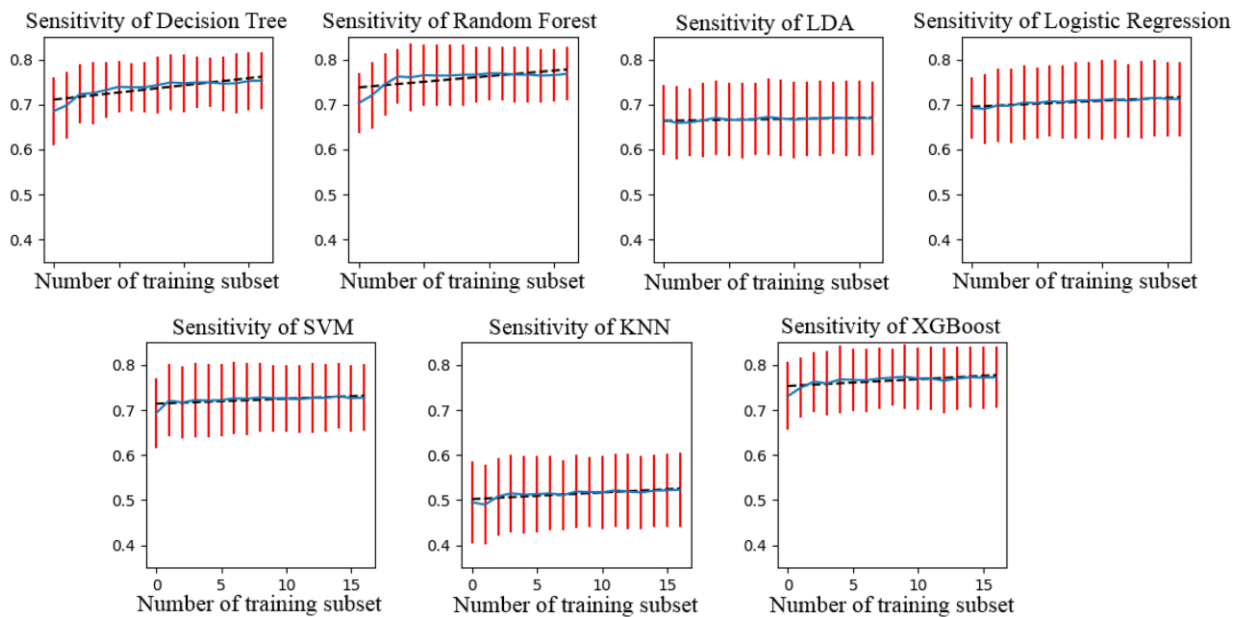


Figure 14. Sensitivity of the proposed model and the number of the training subsets for different classifiers. Dashed lines: Linear regression, Vertical lines: Standard deviation.

used the feature importance analysis to reduce the number of features and complexity of the proposed model. We applied the classifier models to extract the feature importance scores, and the results of each classifier completely differed in the order of the features. Because XGBoost had the best performance, we chose its feature importance scores, and only the first 18 important features had a significant effect on the XGBoost performance. According to the feature importance analysis, for the most models, some features, including NIV, Hb, HR, BUN, SBP, DBP, Creatinine, and Troponin, have high importance scores and can affect the model performance considerably. On the other hand, there are some features such as Cardiac Resynchronization Therapy-Defibrillator (CRT-D), Primary Heart Ischemic, CPO, Captopril, and Implantable Cardioverter Defibrillator

(ICD) which showed lower importance scores and would not affect the performance significantly. NIV was shown with the highest importance score in all seven models; therefore, it is a notable predictor of in-hospital mortality for HF patients. NIV can help to decrease respiratory effort and will improve gas exchange and cardiac output [38]. Hence, HF patients who are in a severe stage could use NIV to provide relief from HF symptoms, and for this reason, in our data, the percentage of dead patients who used NIV is higher than that of others. We aim to add more data to decrease the effect of imbalanced data in the future. In addition, using the registry data to predict mortality of HF patients during 3, 6, and 12 months after discharge can be investigated in future studies. It would be also interesting to develop the proposed model for imbalanced multiclass classification problems.

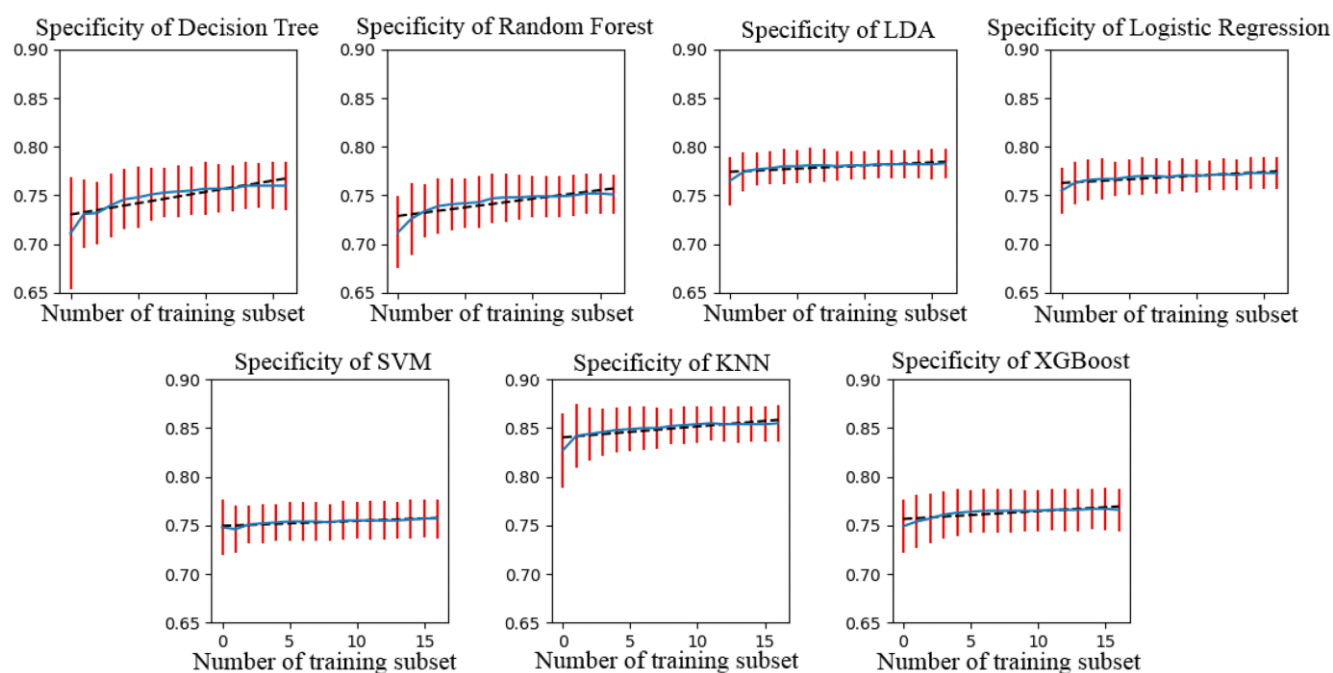


Figure 15. Specificity of the proposed model and the number of the training subsets for different classifiers. Dashed lines: Linear regression, Vertical lines: Standard deviation.

5. Conclusion

In this work, we proposed a method to predict in-hospital mortality of Heart Failure (HF) patients using PROVE/HF imbalanced registry data of hospitalized patients with HF. The method contains an imbalanced ensemble probabilistic model that, using an undersampling and ensembling approach, can distinguish HF patients who die during hospitalization from those who do not. The suggested model uses machine learning models, and among the various models evaluated, Extreme Gradient Boosting (XGBoost) could outperform all others with higher performance. Feature importance analysis using XGBoost showed the proposed method could achieve acceptable performance with fewer but important features (accuracy: $76.4\% \pm 1.6\%$), which can reduce the system complexity considerably. In addition, the statistical analysis of the features suggests predictors that can be used by health providers to determine the required medical resources to reduce the in-hospital mortality of HF patients.

Declaration of competing interests

The authors have no conflict of interest to disclose.

Funding

This research did not receive any specific grant from funding agencies in the public, commercial, or not-for-profit sectors.

Conflicts of interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Authors contribution statement

First Author: Research, Methodology, Data Processing, Validation, Writing-review and Editing;

Second Author: Research, Methodology, Data Processing,

Validation, Writing-review and Editing;

Third Author: Data Gathering, Medical Consultation.

References

1. Yancy, C.W., Jessup, M., Bozkurt, B., et al. "2013 ACCF/AHA guideline for the management of heart failure: A report of the American College of cardiology foundation/American heart association task force on practice guidelines", *Journal of the American College of Cardiology*, **62**(16), pp. 147-239 (2013).
<https://doi.org/10.1016/j.jacc.2013.05.019>
2. Pocock, S.J., Wang, D., Pfeffer, M.A., et al. "Predictors of mortality and morbidity in patients with chronic heart failure", *European Heart Journal*, **27**(1), pp. 65-75 (2006).
<https://doi.org/10.1093/eurheartj/ehi555>
3. Mahmood, S.S., Levy, D., Vasan, R.S., et al. "The Framingham heart study and the epidemiology of cardiovascular diseases: A historical perspective", *Lancet*, **383**(9921), pp. 999-1008 (2014).
[https://doi.org/10.1016/s0140-6736\(13\)61752-3](https://doi.org/10.1016/s0140-6736(13)61752-3)
4. Azad, N. and Lemay, G. "Management of chronic heart failure in the older population", *Journal of Geriatric Cardiology: JGC*, **11**(4), pp. 329-337 (2014).
<https://doi.org/10.11909/j.issn.1671-411.2014.04.008>
5. Miró, Ò., Rossello, X., Gil, V., et al. "Predicting 30-day mortality for patients with acute heart failure in the emergency department: A cohort study", *Annals of Internal Medicine*, **167**(10), pp. 698-705 (2017).
<https://doi.org/10.7326/m16-2726>

6. Seki, T., Kawazoe, Y., and Ohe, K. "Machine learning-based prediction of in-hospital mortality using admission laboratory data: A retrospective, single-site study using electronic health record data", *PLOS One*, **16**(2), 0246640 (2021).
<https://doi.org/10.1371/journal.pone.0246640>
7. Fonarow, G.C., Adams, K.F., Abraham, W.T., et al. "Risk stratification for in-hospital mortality in acutely decompensated heart failure: classification and regression tree analysis", *JAMA*, **293**(5), pp. 572-580 (2005).
<https://doi.org/10.1001/jama.293.5.572>
8. König, S., Pellissier, V., Hohenstein, S., et al. "Machine learning algorithms for claims data-based prediction of in-hospital mortality in patients with heart failure", *ESC Heart Failure*, **8**(4), pp. 3026-3036 (2021).
<https://doi.org/10.1002/ehf2.13398>
9. Luo, C., Zhu, Y., Zhu, Z., et al. "A machine learning-based risk stratification tool for in-hospital mortality of intensive care unit patients with heart failure", *Journal of Translational Medicine*, **20**, 138 (2022).
<https://doi.org/10.1186/s12967-022-03340-8>
10. Wallace, B.C., Small, K., Brodley, C.E., et al. "Class imbalance, redux", *Paper presented at: 2011 IEEE 11th International Conference on Data Mining*, pp. 754-763 (2011).
<https://doi.org/10.1109/ICDM.2011.33>
11. Wojtas, M. and Chen, K. "Feature importance ranking for deep learning", *arXiv* (2020).
<https://doi.org/10.48550/arXiv.2010.08973>
12. Alizadehsani, R., Khosravi, A., Roshanzamir, M., et al. "Coronary artery disease detection using artificial intelligence techniques: A survey of trends, geographical differences and diagnostic features 1991-2020", *Computers in Biology and Medicine*, **128**, 104095 (2021).
<https://doi.org/10.1016/j.combiomed.2020.104095>
13. Sakata, Y. and Shimokawa, H. "Epidemiology of heart failure in Asia", *Circulation Journal*, **77**(9), pp. 2209-2217 (2013).
<https://doi.org/10.1253/circj.cj-13-0971>
14. Givi, M., Sarrafzadegan, N., Garakyaraghi, M., et al. "Persian Registry of cardioVascular diseases (PROVE): Design and methodology", *ARYA Atherosclerosis*, **13**(5), pp. 236-244 (2017).
15. Hachesu, P.R., Ahmadi, M., Alizadeh, S., et al. "Use of data mining techniques to determine and predict length of stay of cardiac patients", *Healthcare Informatics Research*, **19**(2), pp. 121-129 (2013).
<https://doi.org/10.4258/hir.2013.19.2.121>
16. Freedman, D., Pisani, R., and Purves, R., *Statistics: Fourth International Student Edition*, W.W. Norton and Company (2007).
17. Guo, X., Yin, Y., Dong, C., et al. "On the class imbalance problem", *Paper presented at: 2008 Fourth International Conference on Natural Computation* (2008).
<https://doi.org/10.1109/ICNC.2008.871>
18. Wang, Y., Wang, D., Ye, X., et al. "A tree ensemble-based two-stage model for advanced-stage colorectal cancer survival prediction", *Information Sciences*, **474**, pp. 106-124 (2019).
<https://doi.org/10.1016/j.ins.2018.09.046>
19. Probst, P., Boulesteix, A.L., and Bischl, B. "Tunability: Importance of hyperparameters of machine learning algorithms", *The Journal of Machine Learning Research*, **20**(1), pp. 1934-1965 (2019).
<https://doi.org/10.48550/arXiv.1802.09596>
20. Weerts, H.J., Mueller, A.C., and Vanschoren, J. "Importance of tuning hyperparameters of machine learning algorithms" *arXiv*, 2007.07588 (2020).
<https://doi.org/10.48550/arXiv.2007.07588>
21. Daghistani, T.A., Elshaw, R., Sakr, S., et al. "Predictors of in-hospital length of stay among cardiac patients: A machine learning approach", *International Journal of Cardiology*, **288**, pp. 140-147 (2019).
<https://doi.org/10.1016/j.ijcard.2019.01.046>
22. Flach, P. and Kull, M. "Precision-recall-gain curves: PR analysis done right", *Advances in Neural Information Processing Systems*, **28**(1), pp. 838-846 (2015).
<https://doi.org/10.5555/2969239.2969333>
23. Refaailzadeh, P., Tang, L., and Liu, H. "Cross-validation", *Encyclopedia of Database Systems*, **5**, pp. 532-538 (2009).
https://doi.org/10.1007/978-0-387-39940-9_565
24. Malliaris, M.E. and Pappas, M. "Revenue generation in hospital foundations: Neural network versus regression model recommendations", *International Journal of Management and Information Systems (IJMIS)*, **15**(1), pp. 59-66 (2011).
<https://doi.org/10.19030/ijmis.v15i1.1596>
25. Pedregosa, F., Varoquaux, G., Gramfort, A., et al. "Scikit-learn: Machine learning in Python", *the Journal of Machine Learning Research*, **12**, pp. 2825-2830 (2011).
<https://doi.org/10.5555/1953048.2078195>
26. Mosterd, A. and Hoes, A.W. "Clinical epidemiology of heart failure", *Heart*, **93**(9), pp. 1137-1146 (2007).
<https://doi.org/10.1136/hrt.2003.025270>
27. Kaldara-Papatheodorou, E.E., Terrovitis, J.V., and Nanas, J.N. "Anemia in heart failure: Should we supplement iron in patients with chronic heart failure?", *Polskie Archiwum Medycyny Wewnętrznej*,

- 120(9), pp. 354-360 (2010).
<https://doi.org/10.20452/pamw.967>
28. Norcliffe-Kaufmann, L., Kaufmann, H., Palma, J.A., et al. "Orthostatic heart rate changes in patients with autonomic failure caused by neurodegenerative synucleinopathies", *Annals of Neurology*, **83**(3), pp. 522-531 (2018).
 29. Domanski, M., Tian, X., Haigney, M., et al. "Diuretic use, progressive heart failure, and death in patients in the DIG study", *Journal of Cardiac Failure*, **12**(5), pp. 327-332 (2006).
<https://doi.org/10.1016/j.cardfail.2006.03.006>
 30. Ahmed, A., Husain, A., Love, T. E., et al. "Heart failure, chronic diuretic use, and increase in mortality and hospitalization: an observational study using propensity score methods", *European Heart Journal*, **27**(12), pp. 1431-1439 (2006).
<https://doi.org/10.1093/eurheartj/ehi890>
 31. Spoletini, I., Coats, A.J., Senni, M., et al. "Monitoring of biomarkers in heart failure", *European Heart Journal Supplements*, **21**, pp. M5-M8 (2019).
 32. Shlipak, M.G., Chertow, G.C., and Massie, B.M. "Beware the rising creatinine level", *Journal of Cardiac Failure*, **9**(1), pp. 26-28 (2003).
<https://doi.org/10.1054/jcaf.2003.10>
 33. Wettersten, N. and Maisel, A. "Role of cardiac troponin levels in acute heart failure", *Cardiac Failure Review*, **1**(2), pp. 102-106 (2015).
<https://doi.org/10.15420/cfr.2015.1.2.102>
 34. Kociol, R.D., Hammill, B.G., Fonarow, G.C., et al. "Generalizability and longitudinal outcomes of a national heart failure clinical registry: Comparison of Acute Decompensated Heart Failure National Registry (ADHERE) and non-ADHERE Medicare beneficiaries", *American Heart Journal*, **160**(5), pp. 885-892 (2010).
<https://doi.org/10.1016/j.ahj.2010.07.020>
 35. Abraham, W.T., Fonarow, G.C., Albert, N.M., et al. "Predictors of in-hospital mortality in patients hospitalized for heart failure: Insights from the Organized Program to Initiate Lifesaving Treatment in Hospitalized Patients with Heart Failure (OPTIMIZE-HF)", *Journal of the American College of Cardiology*, **52**(5), pp. 347-356 (2008).
<https://doi.org/10.1016/j.jacc.2008.04.028>
 36. Ramezan, C.A., Warner, T.A., Maxwell, A.E., et al. "Effects of training set size on supervised machine-learning land-cover classification of large-area high-resolution remotely sensed data", *Remote Sensing*, **13**(3), 368 (2021).
<https://doi.org/10.3390/rs13030368>
 37. Dissanayake, K. and Md Johar, G.M. "Comparative study on heart disease prediction using feature selection techniques on classification algorithms", *Applied Computational Intelligence and Soft Computing*, **2021**, 5581806 (2021).
<https://doi.org/10.1155/2021/5581806>
 38. Baram, M. "BiPAP and the relief of CHF symptoms" *Critical Care*, **1**(3000) (2001).
<https://doi.org/10.1186/ccf-2001-3003>

Biographies

Hadi Sabahi is a PhD Candidate in the field of medical data analysis at the Biomedical Engineering Department, K.N. Toosi University of Technology, under the supervision of Professor Mansour Vali. His research interests lie in the fields of large data processing in medical applications using artificial intelligence, as well as medical image processing. He has focused on the analysis of medical registry data of cardiovascular diseases using machine learning approaches to predict mortality and morbidity.

Mansour Vali received his PhD degree in 2006 from Amirkabir University of Technology in Biomedical Engineering, Tehran, Iran. He has been working as an Assistant Professor and a faculty member at the Department of Biomedical Engineering, K.N. Toosi University of Technology since 2013. His main research interests are "sound and speech processing in medical and psychological assessments". Also, he is working on large data processing in medical applications.

Davood Shafie received his Diploma of Medicine in 2002 from Isfahan University of Medical Sciences, Isfahan, Iran. Due to his interest in studying cardiovascular disease, he studied cardiology at the Isfahan University of Medical Sciences. He completed his professional training in the field of HF and transplantation at the Iran University of Medical Sciences. He currently serves as an HF specialist at Shahid Chamran Cardiovascular, Medical & Research Center in Isfahan. He has also been the deputy director of the Heart Failure Research Center since 2014. Over the years, he has also coordinated the HF registry program in Isfahan. He is also a member of the Isfahan Cardiovascular Research Center (WHO Collaborating Center in EMR).