

Prediction of heart attack risk using linear discriminant analysis methods

 Esra Sivari,  Selim Sürücü

Department of Computer Engineering, Çankırı Karatekin University, Çankırı, Turkey

Cite this article: Sivari E, Sürücü S. Prediction of heart attack risk using linear discriminant analysis methods. *J Comp Electr Electron Eng Sci.* 2023; 1(1): 5-9.

Corresponding Author: Esra Sivari, esrasivari@karatekin.edu.tr

Received: 24/02/2023

Accepted: 23/03/2023

Published: 28/04/2023

ABSTRACT

Aims: A heart attack is when blood flow to the heart muscle is interrupted. During a heart attack, the risk of permanent damage to the heart increases with every second that the heart tissue is not supplied with enough blood. If early and appropriate intervention is not carried out, the heart starts to fail and dies. The main risk factors for heart attack are smoking, cholesterol, diabetes, high blood pressure, age, obesity, genetics and high levels of certain substances produced in the liver. This study aims to use machine learning methods to predict the risk of heart attack using a dataset created by taking into account risk factors.

Methods: The performance of three types of linear discriminant analysis classifiers - Normal, Ledoit-Wolf and Oracle Shrinkage Approximating - was compared on the Cleveland data set.

Results: Normal linear discriminant analysis gave the best classification with an accuracy of 83.60% and performed better than regularised versions.

Conclusion: Linear discriminant analysis is a promising classifier for predicting myocardial infarction. It can be used in hospitals as an objective and automated system that reduces the workload of specialists and helps to reduce diagnostic costs.

Keywords: Machine learning, heart attack risk, linear discriminant analysis

INTRODUCTION

The heart is a vital organ located on the left side of the muscular chest and pumps almost 8000 litres of blood into the circulation, contracting on average 100,000 times a day. Cardiovascular diseases such as valvular disease, myocardial disease and myocardial infarction, which are related to the coronary arteries that supply blood to the heart tissue, can result from any disorder in the heart's structure. The WHO estimates that cardiovascular disease is the leading cause of death worldwide, accounting for 17.9 million deaths annually, and WHO research suggests that this number will rise to around 30 million by 2040. One third of deaths occur prematurely in middle-aged people, while four out of five deaths from cardiovascular disease are due to heart attacks and strokes.(Virani et al., 2021).

A heart attack is when the blood supply to the heart muscle is interrupted by blockage or narrowing of the coronary arteries, which are responsible for supplying the heart with oxygen and nutrients. During a heart attack, the risk of permanent damage increases with every second that the heart tissue is deprived of blood. Any sudden blockage in the arteries that supply blood to the heart can cause the heart muscle to be starved of oxygen, resulting in damage to the heart tissue. Loss of heart tissue occurs if the vessel is not opened early and with the correct procedure. The loss reduces the heart's ability to pump and heart failure occurs. Causes such as smoking, cholesterol, diabetes, high blood pressure, age, obesity, genetics

and elevated levels of certain substances produced in the liver are the main risk factors for heart attack. By taking these factors into account, the risk of heart attack can be predicted and people can be aware and take early precautions.

Today, prediction and classification applications are being developed for the early diagnosis of many diseases, thanks to machine learning techniques applied to medical data. These applications offer the benefits of objectivity in diagnosis and high and fast performance, reducing the workload and cost of the specialist. Machine learning studies conducted to predict whether an individual is at risk of a heart attack have focused on a 13-character Cleveland dataset (Detrano, 1989). In previous studies, many have proposed support vector machines (SVM) (Xing et al., 2008), neural networks (Ajam, 2015), decision trees (Thenmozhi& Deepika, 2014), K-Nearest Neighbour (KNN) (Wettschereck&Dietterich, 1995) and naive Bayes (Srinivas et al., 2010) methods to classify the Cleveland dataset and estimate the risk of heart attack for individuals, and experimental results have shown satisfactory accuracy. Another frequently used technique is the combination of Principal Component Analysis (PCA), one of the feature selection and size reduction methods, with machine learning algorithms (Aec et al., 2013). Fuzzy logic (Guidi et al., 2014), ensemble learning (Das et al., 2009) and deep learning (Manimurugan et al., 20-22) are recommended methods.

Looking at the studies conducted in recent years to classify the Cleveland dataset, machine learning algorithms based on statistical methods are recommended to achieve high performance. Ensemble techniques, such as Latha and Jeeva (2019) bagging and amplification in heart disease risk determination, provided a maximum accuracy increase of 7% in improving the predictive accuracy of weak classifiers. Bharti et al. (2021) achieved 94.2% accuracy with a hybrid method that included machine learning and deep learning algorithms. Santhanam and Ephzibah (2013) proposed the PCA technique for feature selection and achieved classification accuracies of 92 and 95.2% with combinations of PCA + regression and PCA + ANN. Dissanayake and Johar (2021) compared different feature selection algorithms (ANOVA, chi-square, mutual information, ReliefF, forward feature selection, backward feature selection, exhaustive feature selection, recursive feature elimination, lasso regression and ridge regression) on the dataset and the backward feature selection technique achieved the highest classification accuracy of 88.52%. Gárate-Escamila et al. (2020) achieved 98.7% accuracy on the Cleveland dataset using Chi-square and Principal Component Analysis (CHI-PCA) with Random Forests (RF). Katarya and Meena (2021) compared the performance of nine different machine learning classifiers and reported 95.60% accuracy with Random Forest.

For the detection of coronary artery disease, Singh et al. (2018) proposed the Fisher method and generalised discriminant analysis with binary classifiers for feature selection. Negi et al. (2016) used PCA independent linear discriminant analysis to study the electrocardiogram. Asl et al. (2008) proposed the combination of generalised discriminant analysis and SVM for arrhythmia classification.

However, to our knowledge, discriminant analysis has not been used in previous studies of the Cleveland data set. In this study, linear discriminant analysis techniques are proposed to predict the risk of myocardial infarction. The performances of three types of linear discriminant analysis classifiers, Normal, Ledoit-Wolf and Oracle Shrinkage Approximating (OAS), were compared on the Cleveland dataset. The main contribution of this study is that it is the first study to propose LDA algorithms for classification of the Cleveland dataset. Another contribution is the estimation of the risk of myocardial infarction with high performance without using feature selection algorithms with LDA algorithms.

METHODS

A Cleveland dataset consisting of 303 samples with 13 features was used for the application. The number of patients at risk of heart attack in the dataset is 165. The features and explanations of the dataset are given in Table 1. As the dataset did not contain any null values, the data was directly applied to the machine learning algorithms without the feature selection process. The method overview scheme is shown in Figure 1. The data was randomly divided into training and test sets, of which 61 were test data. With the 10-fold cross-validation technique, the training and validation phases were performed simultaneously. The training phase was carried out using Normal, Ledoit-Wolf and OAS LDA algorithms with Least Squares Optimisation (lsqr). After the training phase, the selected performance metrics and the classes of the test data were estimated and their performances were compared.

Table 1. The features and descriptions of the Cleveland dataset

Feature Names	Features Descriptions
Age	Age of the patient
Sex	Sex of the patient (1=male; 0=female)
exang	Exercise induced angina (1=yes; 0=no)
ca	Number of major vessels (0-3)
cp	Chest Pain type chest pain type (1=typical angina; 2=atypical angina; 3=non-anginal pain; 4=asymptomatic)
trtbps	Resting blood pressure (in mm Hg)
chol	Cholesterol in mg/dl fetched via BMI sensor
fbs	Fasting blood sugar > 120 mg/dl (1=true; 0=false)
rest_ecg	Resting electrocardiographic results (0=normal; 1=having ST-T wave abnormality (T wave inversions and/or ST elevation or depression of > 0.05 mV); 2=showing probable or definite left ventricular hypertrophy by Estes' criteria)
thalach	Maximum heart rate achieved
target	0= Less chance of heart attack 1= More chance of heart attack

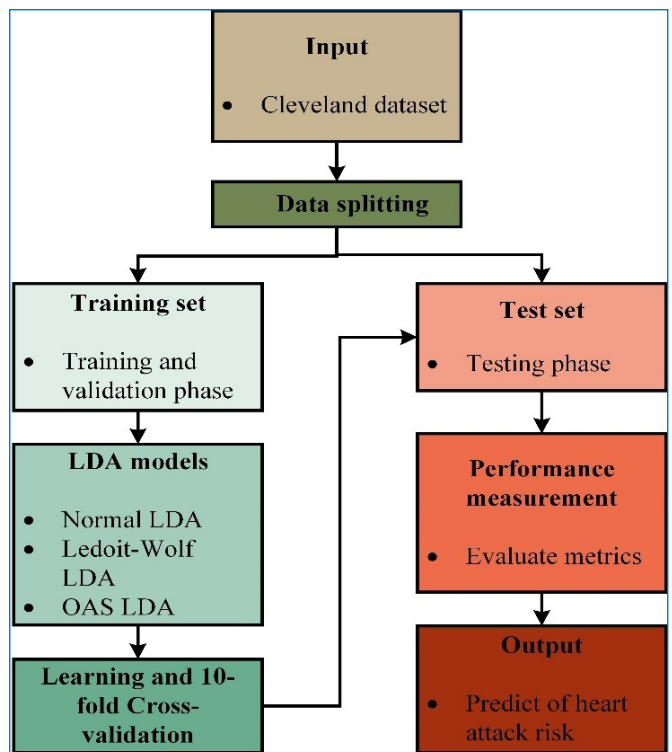


Figure 1. Overview of the method.

Linear discriminant analysis (LDA) is a machine learning technique used for classification, dimension reduction and feature extraction in pattern classification problems. Ronald A. Fisher developed the original linear discriminator in 1936 as a two-class technique (Izenman, 2013). As a supervised classification method, LDA makes a class prediction with a linear decision surface created by fitting conditional class densities to the data and using Bayes' rule (Equation (1)). In a case where there are two classes, i and j , P_{ji} in equation (1) is the base probability (prior probability) of each class observed in the training data, and $f_i(x)$ is the estimated probability of x for class i . Linear discriminant analysis (LDA) is a machine learning technique used for classification, dimension reduction, and feature extraction in pattern classification problems. Ronald A. Fisher developed the original linear discriminant in 1936 as a two-class technique (Izenman, 2013). As a supervised classification method, LDA makes a class prediction with a linear decision surface created by fitting conditional class densities to the data and using Bayes' rule (Equation (1)). If there are two classes, i and j , P_{ji} in Equation (1) is the base probability (prior probability) of each class observed in the training data, and $f_i(x)$ is the estimated

probability of x for class i . If there are two classes, i and j , P_{ji} is the base probability (prior probability) of each class observed in the training data.

$$P(y = i|x) = \frac{P(x|y = i)P(y=i)}{\sum_j \sum P_{ji} * f_i(x)} \quad (1)$$

The discriminant rule uses the probability of each class and the probability of the data belonging to each class. When the Gaussian distribution function is applied to the data, the distribution of x is characterised by its mean (μ) and covariance (Σ), and the discriminant function given in Equation (2) is obtained.

$$D_i(x) = \log f_i(x) + \log P_{ji} \quad (2)$$

The discriminant function calculates the probabilities of how much data x belongs to each class. The decision surface separating classes i and j is the set x , where the two discriminant functions for the two classes have the same value, and any data falling on the decision surface belongs equally to both classes (undefined). LDA occurs when we assume equal covariance between classes i and j ; all classes have the same covariance matrix. The equal covariance discriminant function of LDA is given in Equation (3). The mean, covariance and prior probability given in Equation (3) are estimated from the training data, and the probability of the class with the largest value determines the output estimate.

$$D_i(x) = x^T \mu_i \Sigma^{-1} - \frac{1}{2} \mu_i^T \mu_i \Sigma^{-1} + \log P_{ji} \quad (3)$$

Shrinkage is a penalty term, similar to applying L1 or L2 regularisation to the optimiser, to prevent overfitting in the LDA algorithm. When shrinkage is used as an estimator, a better variance estimator can be obtained than the original raw estimate. The gamma function is added to the estimator to obtain shrinkage. Gamma functions are any function that gives a shrinkage effect. The Ledoit-Wolf estimator (Ledoit & Wolf, 2004) is the most widely used shrinkage algorithm for regularising covariance matrices and improving their stability. The Ledoit-Wolf estimator asymptotically calculates the optimal shrinkage parameter by minimising a Mean Squared Error (MSE) criterion function. The Oracle Shrinkage Approximating (OAS) estimator (Chen et al., 2010) is another shrinkage algorithm that estimates the covariance matrix by iteratively approximating the shrinkage parameter.

RESULT

Learning curves are used to evaluate the results of the training and validation phases. Overfitting or underfitting situations can be observed in a learning curve graph by looking at the generalisation gap between the training and validation curves. A small generalisation gap is normal for an efficient training phase, because the training data is more than the validation data. Figure 2 shows the learning curves of the Normal, Ledoit-Wolf and OAS LDA classifiers. Ten points in the learning curves show cross-validation accuracy values according to the number of training data. The curves show that there is no overfitting or underfitting during the

training phase. However, as the number of training samples is increased, it is observed that the generalisation ability increases and the training and validation accuracy scores increase. According to the shading around the curves showing the standard deviation values, the training standard deviation is low and the validation standard deviation is high.

During the test phase, 61 examples were used, including 31 patients at low risk of myocardial infarction and 30 at risk of myocardial infarction. The confusion matrix gives the true positive (TP), false negative (FN), true negative (TN) and false positive (FP) values used to calculate the performance metrics. Figure 3 shows the confusion matrices of the LDA classifications. The Normal LDA made the highest prediction with 51 correct and 10 incorrect predictions. The OAS LDA made the lowest prediction with 44 correct and 17 incorrect predictions. Normal LDA and Ledoit-Wolf LDA made the same number of incorrect predictions from both classes, while OAS LDA made the same number of correct predictions from both classes. These equations result in performance metrics that are equal or very close to each other.

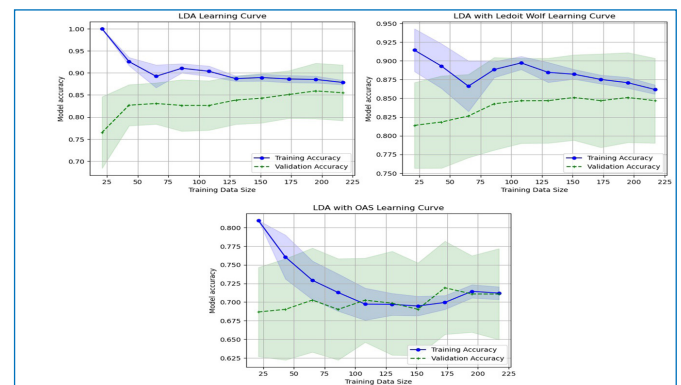


Figure 2. Learning curves of the LDA classifiers.

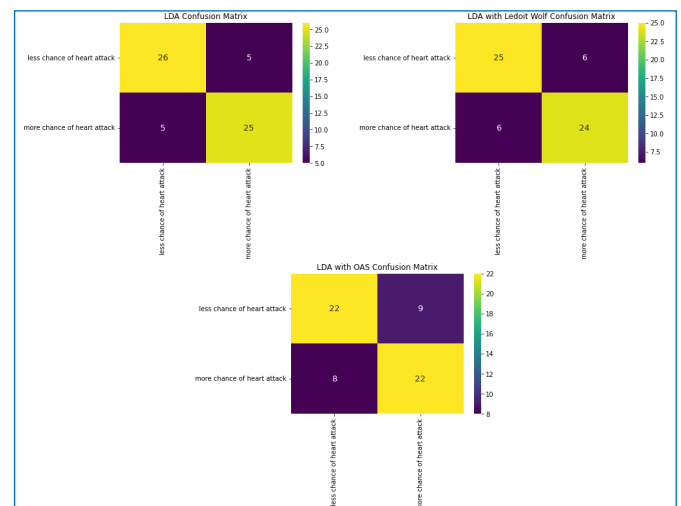


Figure 3. Confusion matrices of the LDA classifiers.

The performance measures and mathematical formulae used to calculate the test results are given in Table 2. Accuracy (ACC) is the ratio of correct predictions to all predictions. Precision is the ratio of true positives to all positives. Recall measures how accurately the true positives were predicted. The F_1 score is the harmonic mean of precision and recall. The area under the ROC curve (AUC) value measures how well two classes can be discriminated, and a value of 1 is desired for perfect discrimination. The F_1 and ACC metrics are positive class dependent, but the Matthews Correlation Coefficient (MCC) gives positive and negative class dependent results.

Table 2. Mathematical formulas of the performance metrics

Performance Metrics	Mathematical Formulas
Accuracy (ACC)	$\frac{TP + TN}{TP + FP + FN + TN} \quad (4)$
Precision	$\frac{TP}{TP + FP} \quad (5)$
Recall	$\frac{TP}{TP + FN} \quad (6)$
F ₁ Score	$2 \times \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (7)$
Matthews Correlation Coefficient (MCC)	$\frac{TP \times TN - FP \times FN}{\sqrt{(TP + FP)(TP + FN)(TN + FP)(TN + FN)}} \quad (8)$

Normal LDA gave the best classification with 83.60% ACC and 67.20% MCC. Ledoit-Wolf LDA was the second best with 80.32% ACC and 60.64% MCC. The OAS LDA performed poorly with 72.13% ACC and 44.30% MCC. The balance of the test set and the number of correct and incorrect predictions ensured that the ACC and AUC metrics had the same values, and the Precision, Recall and F₁ metrics had the same values. The F₁ values for Normal, Ledoit-Wolf and OAS LDA are 83.33%, 80% and 72.12% respectively. In medical applications, the precision metric must be higher than the recall. While the recall (73.33%) of OAS LDA is higher than the precision (70.96%), it is the same for the other LDA classifiers (Normal LDA: 83.33%, Ledoit-Wolf LDA: 80%).

Table 3. Test results of the LDA classifiers

Algorithm	ACC (%)	Precision (%)	Recall (%)	F ₁ Score (%)	AUC (%)	MCC (%)
Normal LDA	83.60	83.33	83.33	83.33	83.60	67.20
Ledoit-Wolf LDA	80.32	80.00	80.00	80.00	80.32	60.64
OAS LDA	72.13	70.96	73.33	72.12	72.15	44.30

DISCUSSION

LDA and its variants have been frequently proposed in the classification of various cancers using DNA microarray data. Dudoit et al (2011) classified two-class leukaemia DNA microarray data using Diagonal Linear Discriminant Analysis (DLDA) and Fisher Linear Discriminant Analysis (FLDA) methods. The authors reported the number of misclassified tumour samples as 0 for DLDA and 3 for FLDA. Li et al (2005) proposed Generalised Linear Discriminant Analysis (GLDA) for the detection of seven different cancer types using DNA microarray data. The average classification error rate for lymphoma was 0.05. Sharma and Paliwal (2008) misclassified only 2 samples in the three-class lung adenocarcinoma dataset using the gradient LDA method. A study by Huang et al (2009) compared LDA methods for cancer classification

based on gene expression data and reported that the mean test error of LDA modification methods was lower than that of the LDA method. These studies suggest that promising results can be obtained by applying LDA and its variants to various health data.

LDA is a supervised classification technique that is considered part of building competitive machine learning models, and its simplicity of implementation is an advantage. However, LDA can fail in some cases where the mean of the data distributions is shared and cannot create a new axis that makes both classes linearly separable. Linear decision boundaries may not sufficiently separate classes, and a regularised version of LDA or non-linear discriminant analysis can be used for more general boundaries.

For this application, Normal LDA outperformed the regularised versions. This result shows that linear separation may be more appropriate for the Cleveland dataset, but the overall classification performance is lower than studies in the literature. The reason for the poor performance is that no preprocessing or feature selection methods are applied to the data beyond the capabilities of the LDA classifier. The studies in the literature were pre-processed, and in this study the raw data was used in the classification of selected features. Considering the performance under these conditions, the LDA algorithm is very promising for the problem of measuring heart attack risk.

CONCLUSION

In this study, Normal LDA and the regularised versions Ledoit-Wolf and OAS LDA algorithms were used to automatically and objectively estimate the risk of myocardial infarction. According to the experimental studies on the Cleveland dataset, Normal LDA classified the raw data with 83.60% ACC. Raw data classification without data preprocessing and feature selection methods greatly affected the inability of the performance metrics to achieve high values. Suppose that experimental studies are developed by combining the proposed LDA algorithms with data preprocessing and feature selection methods. In this case, the applicability of this application in hospitals as an objective and automatic system to help experts reduce workload and costs can be suggested. In the future, such applications could be used to create mobile applications or devices that predict the risk of heart attack, independent of the specialist.

ETHICAL DECLARATIONS

Referee Evaluation Process: Externally peer-reviewed.

Conflict of Interest Statement: The authors have no conflicts of interest to declare.

Financial Disclosure: The authors declared that this study has received no financial support.

Author Contributions: All of the authors declare that they have all participated in the design, execution, and analysis of the paper, and that they have approved the final version.

REFERENCES

1. Aec, I., Akhil Jabbar, M., Deekshatulu, B. L., Priti, Akhil Jabbar, C. M., & Chandra, P. (2013). Classification of Heart Disease using Artificial Neural Network and Feature Subset Selection. Global Journal of Computer Science and Technology, 13(D3), 5–14.

2. Ajam, N. (2015). Heart Diseases Diagnoses using Artificial Neural Network. *Network and Complex Systems*, 5(4).
3. Asl, B. M., Setarehdan, S. K., & Mohebbi, M. (2008). Support vector machine-based arrhythmia classification using reduced features of heart rate variability signal. *Artificial Intelligence in Medicine*, 44(1), 51–64. <https://doi.org/10.1016/J.ARTMED.2008.04.007>
4. Bharti, R., Khamparia, A., Shabaz, M., Dhiman, G., Pande, S., & Singh, P. (2021). Prediction of Heart Disease Using a Combination of Machine Learning and Deep Learning. *Computational Intelligence and Neuroscience*, 2021. <https://doi.org/10.1155/2021/8387680>
5. Chen, Y., Wiesel, A., Eldar, Y. C., & Hero, A. O. (2010). Shrinkage algorithms for MMSE covariance estimation. *IEEE Transactions on Signal Processing*, 58(10), 5016–5029. <https://doi.org/10.1109/TSP.2010.2053029>
6. Das, R., Turkoglu, I., & Sengur, A. (2009). Effective diagnosis of heart disease through neural networks ensembles. *Expert Systems with Applications*, 36(4), 7675–7680. <https://doi.org/10.1016/J.ESWA.2008.09.013>
7. Detrano, R. (1989). Cleveland heart disease database. VA Medical Center, Long Beach and Cleveland Clinic Foundation.
8. Dissanayake, K., & Johar, M. G. M. (2021). Comparative study on heart disease prediction using feature selection techniques on classification algorithms. *Applied Computational Intelligence and Soft Computing*, 2021. <https://doi.org/10.1155/2021/5581806>
9. Dudoit, S., Fridlyand, J., & Speed, T. P. (2011). Comparison of Discrimination Methods for the Classification of Tumors Using Gene Expression Data. *Journal of the American Statistical Association*, 2011.97(457), 77–86. <https://doi.org/10.1198/016214502753479248>
10. Gárate-Escamila, A. K., Hajjam El Hassani, A., & Andrès, E. (2020). Classification models for heart disease prediction using feature selection and PCA. *Informatics in Medicine Unlocked*, 19, 100330. <https://doi.org/10.1016/J.IMU.2020.100330>
11. Guidi, G., Pettenati, M. C., Melillo, P., & Iadanza, E. (2014). A machine learning system to improve heart failure patient assistance. *IEEE Journal of Biomedical and Health Informatics*, 18(6), 1750–1756. <https://doi.org/10.1109/JBHI.2014.2337752>
12. Huang, D., Quan, Y., He, M., & Zhou, B. (2009). Comparison of linear discriminant analysis methods for the classification of cancer based on gene expression data. *Journal of Experimental and Clinical Cancer Research*, 28(1), 1–8. <https://doi.org/10.1186/1756-9966-28-149/FIGURES/2>
13. Izenman, A. J. (2013). Linear Discriminant Analysis. In *Modern Multivariate Statistical Techniques* 237–280. Springer, New York, NY. https://doi.org/10.1007/978-0-387-78189-1_8
14. Katarya, R., & Meena, S. K. (2021). Machine Learning Techniques for Heart Disease Prediction: A Comparative Study and Analysis. *Health and Technology*, 11(1), 87–97. <https://doi.org/10.1007/S12553-020-00505-7/TABLES/3>
15. Latha, C. B. C., & Jeeva, S. C. (2019). Improving the accuracy of prediction of heart disease risk based on ensemble classification techniques. *Informatics in Medicine Unlocked*, 16, 100203. <https://doi.org/10.1016/J.IMU.2019.100203>
16. Ledoit, O., & Wolf, M. (2004). A well-conditioned estimator for large-dimensional covariance matrices. *Journal of Multivariate Analysis*, 88(2), 365–411. [https://doi.org/10.1016/S0047-259X\(03\)00096-4](https://doi.org/10.1016/S0047-259X(03)00096-4)
17. Li, H., Zhang, K., & Jiang, T. (2005). Robust and accurate cancer classification with gene expression profiling. *Proceedings. IEEE Computational Systems Bioinformatics Conference*, 2005, 310–321. <https://doi.org/10.1109/CSB.2005.49>
18. Manimurugan, S., Almutairi, S., Aborokbah, M. M., Narmatha, C., Ganesan, S., Chilamkurti, N., Alzaheb, R. A., & Almoamari, H. (2022). Two-Stage Classification Model for the Prediction of Heart Disease Using IoMT and Artificial Intelligence. *Sensors* 2022,22(2). <https://doi.org/10.3390/S22020476>
19. Negi, S., Kumar, Y., & Mishra, V. M. (2016). Feature extraction and classification for EMG signals using linear discriminant analysis. *Proceedings - 2016 International Conference on Advances in Computing, Communication and Automation (Fall), ICACCA 2016*. <https://doi.org/10.1109/ICACCAF.2016.7748960>
20. Santhanam, T., & Ephzibah, E. P. (2013). Heart disease classification using PCA and feed forward neural networks. *Lecture Notes in Computer Science (Including Subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*, 8284 LNAI, 90–99. https://doi.org/10.1007/978-3-319-03844-5_10/COVER
21. Sharma, A., & Paliwal, K. K. (2008). Cancer classification by gradient LDA technique using microarray gene expression data. *Data & Knowledge Engineering*, 66(2), 338–347. <https://doi.org/10.1016/J.DATAK.2008.04.004>
22. Singh, R. S., Saini, B. S., & Sunkaria, R. K. (2018). Detection of coronary artery disease by reduced features and extreme learning machine. *Clujul Medical*, 91(2), 166. <https://doi.org/10.15386/CJMED-882>
23. Srinivas, K., Raghavendra Rao, G., & Govardhan, A. (2010). Analysis of coronary heart disease and prediction of heart attack in coal mining regions using data mining techniques. *ICCSE 2010 - 5th International Conference on Computer Science and Education, Final Program and Book of Abstracts*, 1344–1349. <https://doi.org/10.1109/ICCSE.2010.5593711>
24. Thenmozhi, K., & Deepika, P. (2014). Heart Disease Prediction Using Classification with Different Decision Tree Techniques. *International Journal of Engineering Research and General Science*, 2(6), 6–11.
25. Virani, S. S., Alonso, A., Aparicio, H. J., Benjamin, E. J., Bittencourt, M. S., Callaway, C. W., Carson, A. P., Chamberlain, A. M., Cheng, S., Delling, F. N., Elkind, M. S. V., Evenson, K. R., Ferguson, J. F., Gupta, D. K., Khan, S. S., Kissela, B. M., Knutson, K. L., Lee, C. D., Lewis, T. T., Tsao, C. W. (2021). Heart Disease and Stroke Statistics - 2021 Update: A Report From the American Heart Association. *Circulation*, 143(8), E254–E743. <https://doi.org/10.1161/CIR.0000000000000950/FORMAT/EPUB>
26. Wettschereck, D., & Dietterich, T. G. (1995). An experimental comparison of the nearest-neighbor and nearest-hyperrectangle algorithms. *Machine Learning*, 19(1), 5–27. <https://doi.org/10.1007/BF00994658>
27. Xing, Y., Wang, J., Zhao, Z., & Gao, andYonghong. (2008). Combination Data Mining Methods with New Medical Data to Predicting Outcome of Coronary Heart Disease. 868–872. <https://doi.org/10.1109/ICCIT.2007.204>

Esra Sivari

Esra Sivari received the B.S. degree in computer engineering from Selcuk University, Turkey, in 2017, and the M.S. degree in electrical and electronics engineering from Cankiri Karatekin University, Turkey, in 2021. She is currently a Ph.D. student at Ankara University Computer Engineering and a research assistant at Cankiri Karatekin University Computer Engineering. Her work focuses on the field of artificial intelligence in medicine.

