



A decision tree classifier algorithm for the automated analysis of electrocardiogram signals to identify heart arrhythmias.

Zavadjil M

Submitted: October 29, 2023, Revised: version 1, January 6, 2024
 Accepted: January 23, 2024

Abstract

Cardiovascular disease is one of the principal causes of death worldwide. Arrhythmias are one of the most common symptoms of most cardiovascular conditions, meaning that their quick and accurate diagnosis is paramount to the timely treatment of cardiovascular disease. In clinical practice, an electrocardiogram (ECG) is used to depict the activity of the heart, and a physician then identifies the arrhythmia based on this signal. This paper developed a computational approach for the automated classification of arrhythmias by training a decision tree classifier model in Python using features from the PTB-XL+ database, consisting of 21799 clinical 12-lead ECGs from 18869 patients of 10 second length each. The developed approach involved preprocessing the data and assessing which of 739 features from the database had the greatest significance using a wrapper method, training the model using a decision tree classifier optimised using a hyperparameter grid and a grid search, and evaluating the model using accuracy, precision, recall, f1-score, and a classification report, as well as a confusion matrix and a plot illustrating which features on the ECG played the greatest role in the determination of the diagnosis. The decision tree classifier model was compared to other supervised machine learning methods to ensure it resulted in the best performance, with a percent accuracy of 84%. The method was made as concise as possible to increase the efficiency of the diagnosis process, thereby increasing patient wellbeing and decreasing costs, both social and monetary, that society faces as a result of the considerable and increasing levels of heart disease globally.

Keywords

Electrocardiogram, Cardiac arrhythmia, Machine learning, Decision tree classifier, Classification, Computational biology, Bioinformatics

Mila Zavadjil, American International School Vienna, Kärntner Ring 15, 1010 Vienna, Austria. Milazavadjil@icloud.com

Introduction

According to the World Health Organization (WHO), cardiovascular diseases (CVDs) are the leading cause of death worldwide, generating nearly a third of all fatalities each year (1). The most common CVDs include heart attack and stroke, and the highest prevalence is in low- and middle-income countries. It is paramount that CVDs are detected as early as possible to enable appropriate treatment and minimize negative effects including disability and death, particularly as most are preventable through lifestyle changes (1).

A widespread tool for the diagnosis for CVDs in clinical practice is the electrocardiogram (ECG), as it is fast, inexpensive, and non-invasive (2-3). ECG signals relay the electrical activity of the heart and are comprised of the waveforms P, QRS, and T. Heart diseases can then be diagnosed by the shape and duration of these waveforms, as well as the distances between different peaks (4).

ECG signals consist of a periodic wave pattern and any deviation from normal may illustrate arrhythmias, which are abnormalities in the frequency, rhythm, origin, conduction velocity, or sequence of excitation of the electrical impulses of the heart that can be caused by conditions such as coronary artery disease, hypertension, cardiomyopathy, and valve disorders (5-6). Due to the periodic nature of ECG waves, the method that is currently most used for interpreting ECG signals and diagnosing arrhythmias is visual waveform analysis by physicians based on the morphological shape

observed. Usually, if the wave does not adhere to the healthy morphology, a CVD can be identified as the cause (5). Nevertheless, this method of analysis can be considerably time-consuming as well as subject to errors due to physician experience level and subjective biases (7-8). It can also be further complicated by the fact that ECGs are susceptible to noise from the instrument itself and from the environment, leading to unclear signals and the misdiagnosis of arrhythmias (9-10). This approach is also not applicable for analyzing sizeable quantities of ECG data (7).

Computer-aided analysis of ECG data could mitigate this predicament, since it provides a fast and automatic alternative to the manual classification of arrhythmias, improving efficiency and reliability, as well as being of considerable practical significance to economic and social development and patients' health (7,11). Its significance to economic and social development and to patients' health is especially vital considering the high rates of CVDs and the fact that they mostly occur in low- and middle-income countries, which do not have the up-to-date technological capacities.

In this work, a concise computational methodology employing machine learning for diagnosing arrhythmias has been constructed using Python to save time, money, and to benefit patients' health. The methodology is open-source, and the approach used in the developed methodology was data acquisition, data pre-processing, model training and classification, and finally performance evaluation (12).

ECG data, preprocessing, and feature selection

The PTB-XL ECG dataset is of considerable size, consisting of 21799 clinical 12-lead ECGs from 18869 patients of 10 second length each, collected as part of a long-term project at the Physikalisch-Technische Bundesanstalt (PTB), available on Physionet (13-14). Newly, the PTB-XL+ dataset was published as a supplement to the PTB-XL dataset containing ECG features extracted from the original signals using various algorithms, as well as providing automatic diagnosis statements. It turns the original dataset into a reference dataset for machine learning methods related to ECG data. The use of multiple commercially available algorithms for ECG feature extraction enables an immense number of features, 739, which would not be available to the general public otherwise. These features include amplitudes and intervals, onsets and offsets, heart rates and axes, areas, and

vectorcardiographic measurements among others (15). A full list of features and their descriptions can be found at https://physionet.org/content/ptb-xl-plus/1.0.1/features/feature_description.csv. This dataset was selected for this work for precisely these reasons, as such rigorous feature extraction would be not be possible alternatively, as well as since it was recently published and thus has not been as extensively employed (16). The implications of this for the developed methodology are that there will not be a need for manual feature extraction, thereby shortening and simplifying the code and thereby its implementation, in addition to providing an extensive range of features for the model to use.

The most significant features on an ECG signal consist of the P, QRS, and T peaks, in addition to the intervals between them, which are shown in Figure 1.

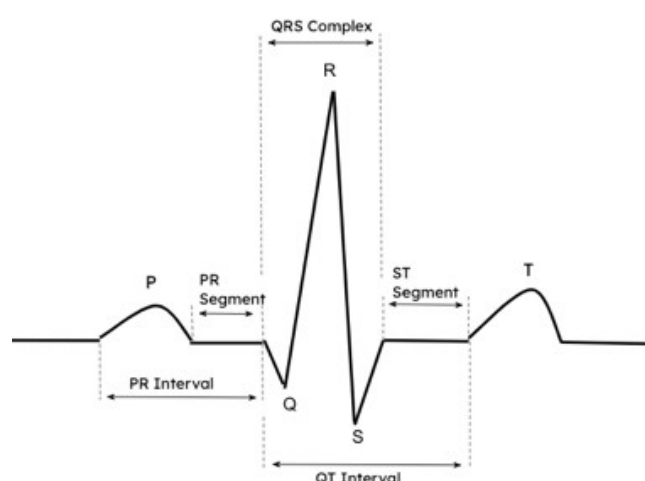


Figure 1. Annotated diagram of an ECG signal, including P, QRS, and T peaks, as well as the intervals between them.

The P-Peak represents the depolarization of the atria, triggered by the sinoatrial node. The PR Segment represents the delay of the atrioventricular node to permit filling of the

ventricles. The QRS Complex illustrates depolarization of the ventricles, which instigates the main pumping contractions. The ST Segment indicates the beginning of the repolarization of the ventricles. The T-Peak illustrates ventricular repolarization (17). The PTB-XL+ ECG dataset contains a wide array of different features based on

these peaks collected from the ECG signals, including times and areas of the P, QRS, and T peaks (18).

ECG classification

Figure 2 provides an overview of the specific machine learning model developed.

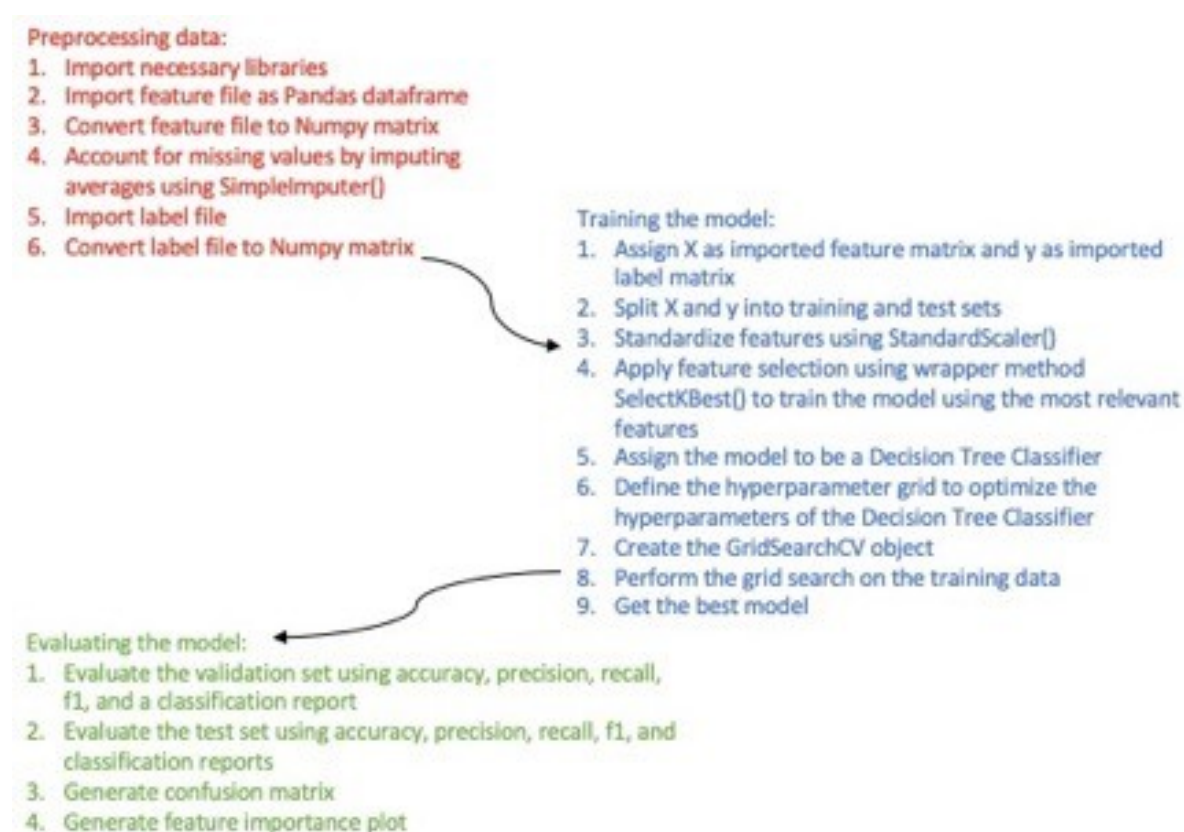


Figure 2. Pseudocode showing the methodology of the developed approach.

In this work, the ECG data was classified using machine learning, and in particular, supervised learning, where the model has both a known input and output when it is trained. In machine learning, classification is a predictive modeling problem where input data is automatically assigned a class label. It consists of predicting the value of a categorical attribute based on the value of

other attributes in the dataset it was trained with (19). This indicates that this technique is suitable for the evaluation of new data, as well as the processing of large amounts of data (20).

To find the model with the highest accuracy, several different widely used models were tested for their ability to use existing data to

classify unseen data, including a decision tree classifier, gradient boosting classifier, ridge regression, linear regression, logistic regression, linear discriminate analysis, naive Bayes, stochastic gradient descent, and support vector machine. The one found to be most effective was the decision tree classifier.

Decision trees are popular supervised learning algorithms for the classification of data. They function by building a tree structure where each internal node denotes a test for a specific attribute, each branch emerging from the node represents the outcomes of the test, and each terminal node, or leaf node, contains a class label. The tree is constructed by repeatedly splitting the training data based on its attributes until a specific class is reached (21).

The specific files from the dataset used in this work are 12sl_features.csv for the features, which contains a single row per ECG record, a column for the PTB-XL ECG identifier and a column for each ECG feature, and 12sl_statements.csv for the diagnostic statements, from which the original statements were used in this work.

These files were both imported into Python as Pandas data tables, and the feature data table was then preprocessed before being used to train the model. Namely, any missing values were removed used the SimpleImputer, which replaced missing or nonnumerical values with means of the other values in that column.

After the feature data was preprocessed, it was stored in a feature matrix in which each row represents a single patient, while each

column provides the values of a different feature, which would then be equated to the X variable in the machine learning model. 80% of the data was used for training, and 20% was used for testing, since this combination produced the most accurate results.

The StandardScaler from SciKit Learn was used to standardize the data to make individual features look like standard normally distributed data. Afterwards, features were selected by relevancy using the SelectKBest method from Scikit Learn. In feature selection, SelectKBest is a wrapper method, which function by selecting a subset of features through the evaluation of a model's performance using different subsets of features. Specifically, SelectKBest selects features according to the highest k scores of the scoring function, thereby reducing overfitting through less redundant data, improving accuracy through less misleading data, and reducing training time, particularly as the sizeable number of features would slow the model. It only left the columns of features that contribute the most to the final diagnosis.

During the development of the coding pipeline created in this work, K-fold validation was performed to demonstrate the absence of overfitting; however, it did not improve the performance measures of the model, instead resulting in slightly lower scores and only complicating the model and countering its purpose of being as simple and thus as quick to implement as possible. As a result, K-fold validation was not included in the final methodology.

To reach the optimal combination of hyperparameters in the decision tree model, a grid search was conducted consisting of a range of values to find the combination that yielded the most effective model. The parameters optimized included the strategy used to choose the function to measure the quality of a split (criterion), the split at each node (splitter), the maximum depth of the tree (maximum depth), the minimum number of samples required to split an internal node (minimum samples split), the minimum number of samples required to be at a leaf node (minimum samples leaf), the number of features to consider when looking for the best split (maximum features), the value at which a node would be split if the split induced a decrease of the impurity greater than or equal to this value (minimum impurity decrease), and the weights associated with classes (class weight).

The grid search found the optimal hyperparameters to be a class weight of none, meaning that all classes had been given equal weights, a criterion of entropy, meaning not all nodes belonged to the same class, a maximum depth of ten, meaning the tree was split ten times, none for maximum features, meaning all features were considered because they had already been selected using SelectKBest, a minimum impurity decrease of zero, meaning that all nodes were split, minimum samples for each leaf of one, meaning that at least one sample was required for each leaf node, minimum samples for each split of ten, meaning all ten features that SelectKBest selected as most crucial were considered, and a splitter of best split, meaning the algorithm found the best split to

separate the samples of a node into two groups by drawing random splits for each of the features and then selecting the best split among them.

Finally, the model was evaluated using accuracy, precision, recall, and f1-score, the formulas for which are given below:

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN}$$

$$Precision = \frac{TP}{TP + FP}$$

$$Recall = \frac{TP}{TP + FN}$$

$$F1 - score = 2 \times \frac{Precision \times Recall}{Precision + Recall}$$

In the formulas above, TP is the number of true positives, TN is the number of true negatives, FP is the number of false positives, and FN is the number of false negatives.

Results and discussion

The purpose of this work was to develop a concise machine learning model to classify arrhythmias from ECG signals which would decrease the time and money needed to classify an ECG signal, as well as decrease the likelihood of human error during the process. To evaluate how effectively it was achieved, a range of metric performance measures were applied.

Since the developed model is a classification model, the most common classification metrics were measured. The first metric performance measure employed to evaluate

the classification results was accuracy, which is the ratio of exact classification to the total classified data. For the developed approach, it was calculated to be 84.0%. The second performance measure used was precision, which is the ratio of true positive classifications to the total number of positive classifications. For the developed approach, it was calculated to be 81.4%. A further metric performance measure calculated was recall, also termed sensitivity, and is the ratio of true positively predicted samples to the total number of truly positive samples. It was calculated as 84.0%. The final measure was F1-score, which takes the harmonic mean of precision and recall to combine them. Here, it was found to be 82.1%.

Table 1. Classification report of the validation data set revealing precision, recall, f1-score, number of instances of each diagnosis, accuracy, macro average, and weighted average.

Validation Set				
Diagnosis	Precision	Recall	F1-Score	Number of signals assessed
ACCEL	0.00	0.00	0.00	1
AFIB	0.86	0.86	0.86	21
ARM	0.00	0.00	0.00	1
EABRAD	0.00	0.00	0.00	1
EAR	0.00	0.00	0.00	0
FLUT	0.50	1.00	0.67	1
MSBRAD	1.00	1.00	1.00	1
NSR	0.76	0.91	0.83	36
QCERR	0.11	0.03	0.04	40
SBRAD	0.83	0.83	0.83	36
SRTH	0.11	0.03	0.04	40
STACH	0.79	0.79	0.79	14
SVT	0.00	0.00	0.00	1
UR	0.00	0.00	0.00	3
VPR	0.00	0.00	0.00	1
Accuracy			0.75	300
Macro Avg	0.32	0.36	0.33	300
Weighted Avg	0.67	0.75	0.70	300

Table 2. Classification report of the test data set revealing precision, recall, f1-score, number of instances of each diagnosis, accuracy, macro average, and weighted average.

Test Set				
	Precision	Recall	F1-Score	Number of signals assessed
\$SIVR	0.00	0.00	0.00	1
ACCEL	1.00	0.50	0.67	2
AFIB	0.76	0.93	0.84	14
EAR	0.00	0.00	0.00	0
FLUT	0.67	1.00	0.80	4
MSBRAD	1.00	1.00	1.00	1
NSR	0.88	0.94	0.91	196
QCERR	0.00	0.00	0.00	1
SBRAD	0.86	0.90	0.88	41
SRTH	0.46	0.22	0.30	27
STACH	1.00	0.80	0.89	10
VPR	0.00	0.00	0.00	3
Accuracy			0.84	300
Macro Avg	0.51	0.48	0.48	300
Weighted Avg	0.82	0.84	0.82	300

The specific results for the training and test set were additionally illustrated using a classification report, which are shown in Tables 1 and 2 for the validation and test set, respectively. In these tables, the diagnoses are written using abbreviations - “ACCEL” indicates that the patient made a transition from sitting to standing, “AFIB” indicates atrial fibrillation, “ARM” indicates that inversion has occurred between the leads used to measure the ECG and that they have switched placed, “EABRAD” indicates ectopic atrial rhythm, “EAR” indicates ectopic atrial rhythm, “FLUT” indicates atrial flutter, “MSBRAD” indicates marked sinus bradycardia, “NSR” indicates a normal sinus rhythm, meaning there is no arrhythmia, “QCERR” indicates poor quality data that cannot be accurately assessed, “SBRAD” indicates sinus bradycardia, “SRTH” indicates a normal sinus rhythm, “STACH”

indicates sinus tachycardia, “SVT” indicates supraventricular tachycardia, “\$SIVR” indicates an idioventricular rhythm, “UR” indicates a rhythm whose mechanism cannot be determined from the ECG, and “VPR” indicates a ventricular paced rhythm.

An observation that can be made from the tables is that the evaluation metrics were significantly influenced by the number of signals in each class; i.e., the more examples were present and thus assessed in each class (a given diagnosis), the greater the precision, recall, and f1-score. This was also depicted by the fact that the weighted average, which considers how many examples having each label/diagnosis have been examined and therefore accounts for imbalanced classes, was considerably greater than the macro average, which does not consider the proportion of cases examined in each class,

only the absolute number of they were identified correctly. Even though this resulted in certain evaluation measures such as precision, recall, and F1-scores appearing low in the tables above, by reducing the number of cases included, the computation time was significantly reduced, and the algorithm could be implemented quickly and easily.

A further observation that can be made from the tables above is that the developed algorithm was better able to differentiate between diagnoses that differed from each other to a greater extent. For example, when two diagnoses had clear differences such as a

normal heart rhythm and atrial fibrillation, particularly when relatively many cases of each were examined, but when the two diagnoses did not have marked differences, such as normal sinus rhythm and sinus rhythm, it could result in confusion, even when the number of cases examined of each is not comparatively low. When there were very few instances of a particular diagnosis included in the validation or test set and the model was not able to identify them correctly, the result it would yield when it was presented with another case of the same class, was more unpredictable.

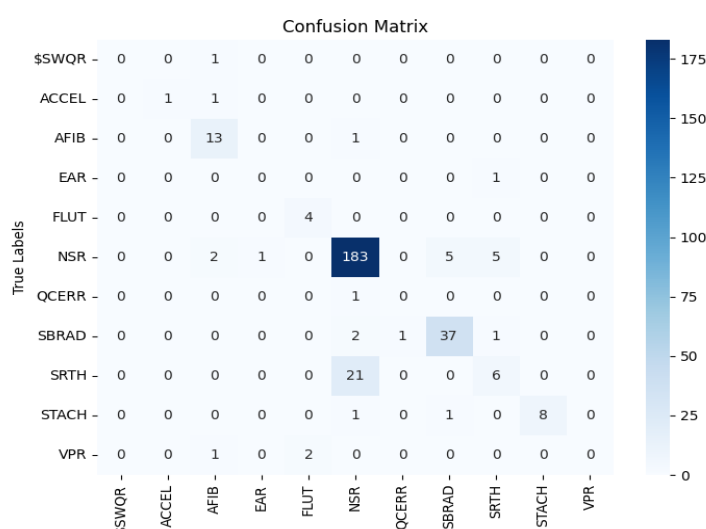


Figure 3. Multiclass confusion matrix of the test set of data demonstrating the true positive, true negatives, false positives, and false negatives.

To further evaluate the performance of the developed approach between classes, a multiclass confusion matrix was constructed, as can be seen in Figure 3. It provides a summary of the classes predicted by the model compared to the true classes, enabling

a visualization of the model's performance across different classes.

The confusion matrix displays the same conclusion as the previous evaluation metrics – whether a signal was classified correctly depended greatly on how many other

examples there were in the same class. The classes that were not correctly assigned were the ones with fewest instances such as “\$SWQR”, “EAR”, and “VPR”, as the model did not have enough data to learn from in these cases.

To visualize the relevance of different parameters on the class determination, a feature importance plot was created, as can be seen in Figure 4.

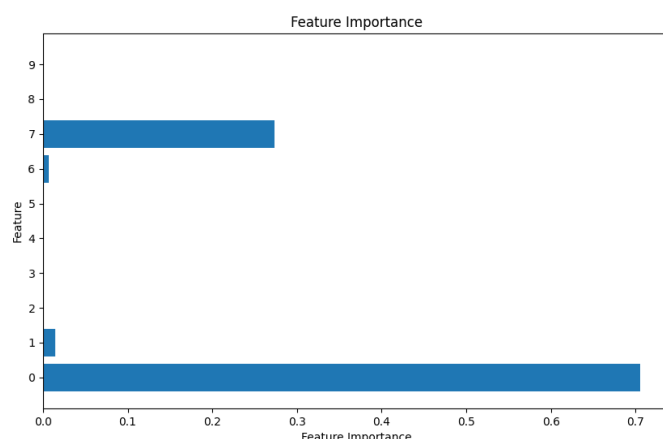


Figure 4. Plot illustrating how important each feature is for the outcome of the classification. Features were selected during data processing using the wrapper method SelectKBest by their significance to the class, and the top ten, labeled as 0-9, were then selected to train the model. The names of the features could not be accessed from SelectKBest.

Figure 4 shows that the developed methodology primarily based its classification on relatively few features. This indicates that although prioritizing a small number of features can still lead to a correct result, it may be more effective for the model to consider a larger number of features because one feature may not always be an accurate picture of the whole ECG signal.

Overall, an advantage of the developed approach is the large number of features from which the most relevant features were selected to train and test the mode. A further advantage is that the code is relatively short and can thus be executed relatively quickly.

Moreover, using decision tree models to classify data has various advantages. Namely, it can simplify complex relationships between input variables and target variables by dividing original input variables into significant subgroups, it is uncomplicated to understand and interpret, it has a non-parametric approach without distributional assumptions, it can easily handle missing values without needing to impute data, it can easily handle significantly skewed data without needing to transform it, and it is robust to outliers (21). Nonetheless, the primary advantage which was aimed for during the creation of the prediction model was for it encompass a small amount of code

and be able to run quickly, even on devices without a large capacity for data.

However, the developed approach also has limitations. The machine learning model in the developed approach was developed using only one data set with a limited number of different classes. To add on, these classes were imbalanced both due to the prevalence of different conditions, as well as the fact that a greater amount of data led to impractically large computation time that contradicted the intent of the work. Moreover, the model based its final decision making on only a small number of features despite the initial variety. Decision tree classifier models have disadvantages in addition to advantages. The main disadvantage is that the model can be subject to overfitting and underfitting, particularly when the data set is relatively small, which can limit its generalizability and robustness. Another potential problem is that strong correlation between different potential input variables may result in the selection of variables that improve the model statistics but are not causally related to the outcome of interest. Hence, the results of decision tree models must be interpreted with caution (21).

Conclusion

A concise automated computational approach to classifying arrhythmias from ECG signals is possible in Python. This open-source approach was able to successfully identify different types of arrhythmias using the features contained in the PTB-XL+ database taken from the PTB-XL database with 21799 recordings via a Decision Tree Classifier model. Even though models in literature have obtained higher F1-scores, they are far more complex and require more computational space, while this approach still enables an increase in efficiency, saving both time and money, in addition to still decreasing the risk of human error during the process of diagnosis (22). The next steps would be to improve the model by training it using ECG signals from multiple different databases, in addition to signals obtained from a greater number of ECG leads for enhanced accuracy. Another way to amplify the accuracy of the developed methodology more significantly would be to use deep learning as opposed to machine learning, so that features would not need to be extracted manually.

References

1. World Health Organization: WHO. "Cardiovascular Diseases (CVDs)." [www.who.int/en/news-room/fact-sheets/detail/cardiovascular-diseases-\(cvds\)](http://www.who.int/en/news-room/fact-sheets/detail/cardiovascular-diseases-(cvds))
2. Bogun, F., Anh, D., Kalahasty, G., Wissner, E., Serhal, CB., Bazzi, R., et al. (2004), "Misdiagnosis of Atrial Fibrillation and Its Clinical Consequences." *The American Journal of Medicine*, 117 (9), 636–642. <https://doi.org/10.1016/j.amjmed.2004.06.024>
3. Ullah, A., Rehman, S., Tu, S., Mehmood, RM., Ehatisham-ul-haq, M., (2021), "A Hybrid Deep CNN Model for Abnormal Arrhythmia Detection Based on Cardiac ECG Signal.", *Sensors*, 3, 951. <https://doi.org/10.3390%2Fs21030951>

4. Aziz, S., Ahmed, S., Alouni, MS. (2021), “ECG-based Machine-learning Algorithms for Heartbeat Classification.” *Scientific Reports*, 11, 18738, <https://doi.org/10.1038/s41598-021-97118-5>
5. Singh, AK., Krishnan, S. (2023), “ECG Signal Feature Extraction Trends in Methods and Applications.” *Biomedical Engineering*, 22, <https://doi.org/10.1186/s12938-023-01075-1>
6. Ahmad, Z., Tabassum, A., Guan, L., Khan, N. (2021), “ECG Heartbeat Classification Using Multimodal Fusion.” *IEEE Access*, 9, 100615–100626. <https://doi.org/10.1109/access.2021.3097614>
7. Mattioli, AV., Puviani, MB., Malagoli, A. (2020), “Quarantine and Isolation During COVID-19 Outbreak: A Case of Online Diagnosis of Supraventricular Arrhythmia Through Telemedicine.” *Journal of Arrhythmia*, 36 (6), 1114–1116. <https://doi.org/10.1002/joa3.12431>
8. Acharya, UR., Oh, SL., Hagiwara, Y., Tan, JH., Adam, M., Gertych, A., Tan, RS. (2017), “A Deep Convolutional Neural Network Model to Classify Heartbeats.” *Computers in Biology and Medicine*, 89, 389–396. <https://doi.org/10.1016/j.combiomed.2017.08.022>
9. Ahmed, AA., Ali, W., Abdullah, TAA., Malebary, SJ. (2023) “Classifying Cardiac Arrhythmia From ECG Signal Using 1D CNN Deep Learning Model.” *Mathematics*, 11(3), 562. <https://doi.org/10.3390/math11030562>
10. Huang, Y., Liu, H., Wu, C., Fang, L., Fang, Q., Wang, Q., et al. (2021), “Ventricular Arrhythmia Predicts Poor Outcome in Polymyositis/Dermatomyositis With Myocardial Involvement.” *Rheumatology*, 60 (8), 3809-3816. <https://doi.org/10.1093/rheumatology/keaa872>
11. Singh, M., Rao, R., Gupta, S. (2020), “KardiaMobile for ECG Monitoring and Arrhythmia Diagnosis.” *American Family Physician*, *American Academy of Family Physicians*, 102 (9), 562–64. <https://www.aafp.org/pubs/afp/issues/2020/1101/p562.html>
12. Zavadjil, Mila. ECG Signal Classifying Code. (2023) <https://github.com/MilaZavadjil/ECG-Classifier>
13. Wagner, P., Strodthoff, N., Bousseljot, R.-D., Kreiseler, D., Lunze, F.I., Samek, W., Schaeffter, T. (2020), PTB-XL: A Large Publicly Available ECG Dataset. *Scientific Data*. <https://doi.org/10.1038/s41597-020-0495-6>

14. Goldberger, A., Amaral, L., Glass, L., Hausdorff, J., Ivanov, P. C., Mark, R. et al. (2000). PhysioBank, PhysioToolkit, and PhysioNet: Components of a new research resource for complex physiologic signals. *Circulation*, 101 (23), e215–e220. <https://doi.org/10.1161/01.cir.101.23.e215>
15. Strodthoff, N., Mehari, T., Nagel, C., Aston, P.J., Sundar, A., Graff, C, et al. (2023). PTB-XL+, a comprehensive electrocardiographic feature dataset. *Sci Data* 10, 279. <https://doi.org/10.1038/s41597-023-02153-8>
16. Wagner, P., Strodthoff, N., Bousseljot, R., Samek, W., Schaeffter, T. (2022). PTB-XL, a large publicly available electrocardiography dataset (version 1.0.3). *PhysioNet*. <https://doi.org/10.13026/kfzx-aw45>
17. Naaz, A., Singh, S. (2014), “Feature Extraction and Analysis of ECG signal for Cardiac Abnormalities- A Review”, *International Journal of Eng. Res. Technol. (IJERT)*, 3 (11). <http://doi.org/10.17577/IJERTV3IS110909>
18. Strodthoff, N., Mehari, T., Nagel, C., Aston, P., Sundar, A., Graff, C., et al. (2023). PTB-XL+, a comprehensive electrocardiographic feature dataset (version 1.0.1). *PhysioNet*. <https://doi.org/10.13026/g6h6-7g88>
19. Priyam, A., Gupta, R., Rathee, A., Srivastava, S. (2013). “Comparative Analysis of Decision Tree Classification Algorithms.”, *International Journal of Current Engineering and Technology*, 3 (2), 334-337. <https://inpressco.com/wp-content/uploads/2013/03/Paper17334-3371.pdf>
20. Nikam, SS. (2017) “A Comparative Study of Classification Techniques in Data Mining Algorithms.” *International Journal of Modern Trends in Engineering and Research*, 4 (7), 58–63. <https://doi.org/10.21884/ijmter.2017.4211.vxayk>.
21. Song. YY., Lu, Y. (2015), “Decision Tree Methods: Applications for Classification and Prediction.” *Shanghai Arch Psychiatry*, 27 (2), 130-135. <https://doi.org/10.11919/j.issn.1002-0829.215044>.
22. Ansari, Y., Mourad, O., Qaraqe, K., Serpedin, E. (2023), “Deep Learning for ECG Arrhythmia Detection and Classification: An Overview of Progress for Period 2017–2023.” *Frontiers in Physiology*, 14, 1246746. <https://doi.org/10.3389/fphys.2023.1246746>.