

A VISUAL INTERFACE FOR EFFICIENT EVALUATION, VISUALIZATION, AND COMPARATIVE ANALYSIS OF MACHINE LEARNING MODELS

Simhadri Yeswanth¹

¹GMR Institute of Technology

ABSTRACT

A study to report a graphical tool for the reduction of hassle in the estimation of machine learning model performance in classifying problems. The approach is designed to help researchers decide which are optimal machine learning models to use with intuitive insights gained in handling datasets and identification of an effective algorithm. Decisions with respect to which of the data preprocessing methods should be used are supported; these are data standardization, normalization and, through PCA, dimensionality reduction. There are six supervised ML classifiers, namely, Logistic Regression, Support Vector Machines, Random Forest, K-Nearest Neighbour, Gaussian Naive Bayes, and Multilayer Perceptron, to be used for making the final predictions. This can be evidenced based on its capability to cope with diversified datasets meeting certain systems' needs, along with enabling the visualization of model performances and behavior using complete sets of metrics, including but not limited to, precision, accuracy, recall, F1 score, confusion matrices. Exploratory data analysis and model comparison within the GUI further give strength to compare classification models appropriately.

Keywords: Graphical tool, machine learning performance, data preprocessing, supervised classifiers, model comparison, visualization metrics, exploratory data analysis

1. INTRODUCTION

For several years, machine learning has become the tool for making data-driven decisions in almost all businesses. But what remains tricky for most of these practitioners is actually selecting the right model suited for the job and really understanding how the performance of these models would fare, particularly with more complicated datasets at hand. And with continued evolution in the field, it has only become essential to have an increasing array of tools evaluating and comparing model performance with greater ease, simplicity, and efficiency.

This paper proposes a graphical tool for better evaluation of classification tasks by machine learning models. The graphical tool provides a holistic, interactive interface through which one can carry out core functions like data preparation, model selection, and result display. It ensures the dataset is in the best form to be analyzed through various techniques such as standardization, normalization, and dimensionality reduction through PCA.

The GUI also supports six supervised machine learning algorithms, namely Logistic Regression, SVM, Random Forest, KNN, Gaussian Naive Bayes, and Multilayer Perceptron/Neural Network. So, the tool is based on an intuitive visualization of key performance metrics like accuracy, precision, recall, F1 score, and confusion matrices to give one a streamlined approach in terms of understanding model behavior and effectiveness. The ability to visualize model performance and compare different algorithms will help users make informed decisions concerning the best classifier for the specific datasets.

It actually brings together the gap that existed between theoretical model evaluation and practical application, creating a robust platform for use both by experienced data scientists and beginners in the field. It simplifies comparative analysis of machine learning models because it provides real-time insights into model performance, and it enhances the entire model selection process.

2. PREVIOUS WORKS

Yacoubi and Axman (2020) [1] propose probabilistic extensions for precision, recall, and F1 score for a higher depth of NLP classification model evaluation since standard measures are not sufficient, especially for imbalanced datasets; a contribution for the authors allows ignoring the uncertainty associated with the predictions and helps to better understand model performance. The authors demonstrate improved model assessment by applying these extended metrics to several NLP tasks, emphasizing the need for strong metrics that reflect the complexities of real-world classification challenges.

Wang et al. (2020) [2] present a survey on ML4VIS, addressing the question of how advanced machine learning can be applied to data visualization. The contribution classifies existing approaches within visual analytics, visualization design, and interactive visualizations and shows how the applications of machine learning optimize areas like design choices that can be

automated, even the data itself. This article discusses open challenges such as the model interpretability and the need for an effective user interface. This work is of most value to summarize the current techniques and future research direction on the enhancement of data visualization using machine learning techniques.

Ray (2019) [3] briefly discusses machine learning algorithms and groups them into three categories: supervised, unsupervised, and reinforcement learning. Key methods are outlined in this paper: linear regression, decision trees, and neural networks. They outline strengths, weaknesses, and appropriateness for each of these applications in the big data and cloud computing environments. The review will prove very useful for getting an understanding of the landscape of the machine learning algorithms and the practical applications of such an algorithm.

Moran et al. (2018) [4] present the use of machine learning in prototyping mobile app graphical user interfaces. They develop a framework that automates the design process by predicting user preferences and optimizing design elements based on interactions. The paper, therefore, explains how machine learning can bring efficiency, reduce development time, and improve user experience in mobile app design in the evolution of intelligent design tools in software engineering.

Maleki et al. (2020) [5] cover the basics of machine learning and classical methods with a focus on major algorithms, especially decision trees, support vector machines, and neural networks. The emphasis is on the importance of such foundational concepts to successful applications across disciplines, especially neuroimaging. This article thus provides a good source through which high-level concepts of machine learning and their implications in practice can be understood.

Oulasvirta et al. (2020) [6] examines combinatorial optimization for graphical user interface designs to present techniques in optimizing the layout with optimization methods. The study is set on evaluating design options that are dependent on user preference and usability metrics in making user-centered interfaces. The work indicates the importance of optimizing GUI components towards efficient and friendly interaction for the user, with some insights on applying optimization in the design process.

A study by Wardhani et al. (2021)[7] is on cross-validation metrics for evaluating classification performance on imbalanced datasets. The authors pointed out the weakness in traditional metrics, where they can misleadingly result due to class imbalance. The authors proposed some tailored cross-validation techniques and metrics improving evaluation of classification models, which are very helpful in gaining more reliable performance evaluation in machine learning applications.

[8] this work develops best practices on quality assurance and hopes to add to more reliable and certifiable systems of machine learning for sound AI applications. This work develops best practices on quality assurance and hopes to add to more reliable and certifiable systems of machine learning for sound AI applications.

Aria et al. (2021) [9] compare a few interpretive approaches of Random Forests with a focus on techniques that may make the model more interpretable. The paper evaluates how permutation importance, partial dependence plots, and SHAP can be used in assessing those techniques to figure out which one is more powerful to provide insight into users' understanding of predictions as well as feature importance. This paper forms a rich contribution towards the making of applications of Random Forest more transparent and trustworthy in machine learning.

Chen et al. (2020) [10] discussed critical feature selection in data classification through different methods of machine learning. According to them, feature selection is an essential aspect as it improves model performance as well as interpretation through identification of relevant variables. Here, the techniques such as filter, wrapper, and embedded methods are studied to see how they effectively improve classification accuracy and also reduce complexity. This research will serve as a guiding source for researchers and practitioners in the field of big data analytics.

RÁCZ et al. (2021) [11] perform a multilevel comparison of different machine learning classifiers along with the respective performance metrics. A few classifiers are compared one to another to see whether their performance in different sets and metrics of accuracy, precision, recall, or F1 score is efficient or not. In that context, authors emphasize the choice of proper metric which should be picked as per the type of the classification task and character of data. This work discusses some strengths and weaknesses of a few classifiers, thus giving this work as useful input for tool use in the model selection and evaluation especially in chemistry and molecular studies.

This work by Hasan and Abdulazeez (2021) [12] in reviewing the PCA algorithm gives its theoretical foundations, its mathematical formulation, and then its application in data analysis. Further, they exhibit that PCA has a power in simplifying a set while retaining essential information; therefore, it can be effective for visualization and noise reduction. The authors also elucidate some of the limits and areas for improvement in the algorithm, which shall present the basis for the application of PCA to the researcher and practitioner.

Çetin and Yıldız (2022) [13] reviewed the preprocessing techniques critical for improving the quality of the data and results of the analyses. Key methods, including cleaning, transformation, normalization, and feature selection, have been emphasized as playing crucial roles in preparing datasets for machine learning and statistical models. Challenges and best practices for data preprocessing are discussed with insight to effectively conduct preprocessing in research and practical applications.

Rao et al. (2020) [14] compared the difference in preprocessing techniques like normalization, standardization, encoding, and missing data handling to compare their impact on machine learning models. It has been identified how a choice of preprocessing methods significantly affects the performance of the model, and based on certain tasks and datasets, the author has provided insights on what technique to choose in advance.

[15] Yang integrates a GUI to explore the prediction of heart disease using decision trees and neural networks. The research highlights how the introduction of a GUI improves access and allows for real-time entry of data into a model for prediction: the possible marriage of machine learning with interactive interfaces to augment diagnoses in health

3. METHODOLOGY

3.1. Requirements and preprocessing of the datasets

The input datasets for the proposed system should be in the CSV-format file. Each row will be sample data, and columns denote the features (or attributes) of samples. In that case, a column can hold the labels or target classes for supervised models of learning. The header name for every feature has to be provided in the first row. Only numeric features are accepted that is, categorical data have to be explicitly transformed into numeric data before loading a dataset. Feature encoding of categorical data is not implemented in this system. It is advisable to clean data in advance by removing those columns that appear to be completely irrelevant or noisy for this task.

Now the loaded dataset can be split in order to prepare a training and test set, using a default 80/20 split between these, but the user has to type in the percent option; decimal percentages can also be typed in, too. The data points then randomly select for the distribution both on the training and on the test set.

3.2. Data Standardization

The system includes Z-score normalization as a preprocessing that will make sure all the features have zero mean and unit variance. It means to shift values of every feature so that it is around 0, and its range will be scaled based on the standard deviation of each feature. Thus, features can be scaled on a uniform scale irrespective of their actual ranges or units (such as centimeters, kilograms, etc.).

Most machine learning algorithms, including Logistic Regression, SVM and Neural Networks, will work much better if the features are scaled equally. The presence of some features with ranges many times larger than others may bias the learning in the model resulting in biased results.

This step further eliminates the curse of dimensionality, which comes from the sparsity of data expressed in a high-dimensional space and thus cannot be handled easily. Standardization has led to much more stable and efficient convergence when training, especially for sensitive models of feature scaling, like SVMs and KNN.

3.3. Dimensionality Reduction

The system makes use of PCA to simplify the complexity of computation for visualization in 2D space. PCA reduces the dimensionality of the data by selecting the components that hold the maximum variance in the data. The user may choose any of the following types of kernels for PCA:

Linear, Polynomial, Radial Basis Function (RBF), Sigmoid, Cosine.

By feature reduction to two dimensions, the system plots data into 2D graphs that describe how the model can easily separate classes.

3.4. Two-Dimensional Representation of the Model's Decision Boundary

One of the very striking features of the system is 2D visualization of the decision boundaries. The system plots reduced feature space with data points, colored by their respective class, after applying PCA. This plot shows a decision boundary in regions of different colors, pointing out the way in which classes can be separated by this classifier. This realization allows users to gain insight into the model's behavior and, hence, understand how it would classify their data. It plots only unique data points, that is, distinct feature vectors, to optimize memory usage. This increases the speed of visualization without losing the general form of the data in hand.

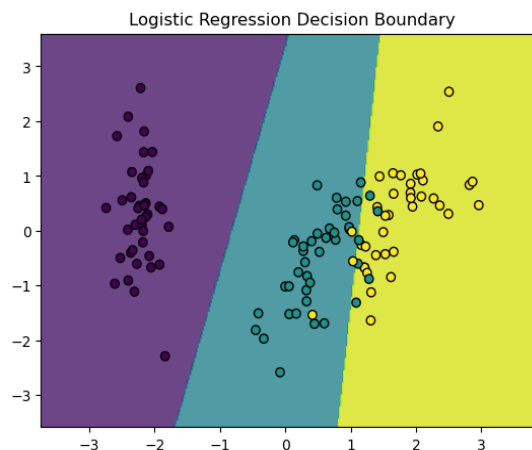


Fig .1. Model Decision Boundary for “iris” dataset

3.5. Classifier Choice

There are six supervised classifiers that are used to test the performance of the models. They are:

- **Logistic Regression:** This is a linear model that is used in solving classification problems in two classes.
- **Support Vector Machines- SVM** is a model that separates classes through the use of a hyperplane, but it also provides for both linear and nonlinear kernels.
- **Random Forest:** It is an ensemble method in which multiple decision trees are built and then aggregated to aid in prediction.
- **KNN:** It just indicates the class of the most majority nearest neighbor
- **Gaussian Naive Bayes:** A probability classifier that uses an assumption of independence between features that are distributed in a Gaussian manner
- **MLP:** A simple model of a neural network that can even handle non-linear relationships between features.
- These classifiers span a wide spectrum of methods, from simple linear models to complex non-linear and ensemble methods.

3.6. Model Performance Metrics

The application will evaluate each classifier for the following four key performance metrics:

Accuracy: Ratio of correctly classified samples of all samples.

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

Precision: the percentage of correct positive predictions out of the actual number of predicted positives, describes how the model can suppress false positives.

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

Recall: the percentage of positive predictions of the total to be actual positive, illustrating how good a model is at locating the positive samples.

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

F1-Score: it describes the harmonic mean of both precision and recall. Even if class imbalance occurs, the F1 score is balanced.

$$F1 \text{ Score} = 2 * \frac{Precision * Recall}{Precision + Recall}$$

All of these metrics can be available when the model outputs the results on the testing dataset. The user then compares all these metrics in these multi-classifiers in order to find the one which best fits his/her data set.

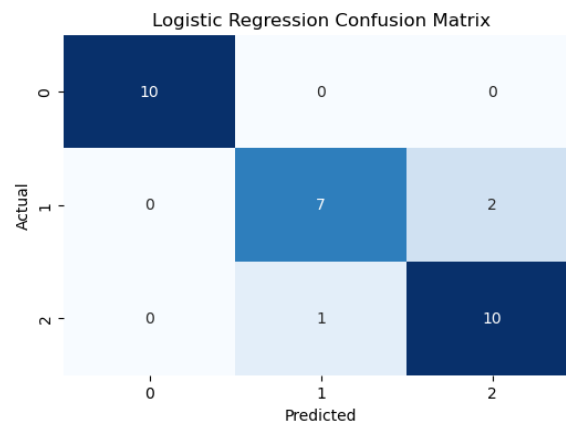


Fig. 2. Model performance matrix for “iris” dataset

3.7. System Flexibility and User Control

The system allows flexibility in training size, feature columns, and the classifier type to be used. The activation or deactivation of standardization, batch size settings, and feature selection are accessible by inputting the range of columns intended for use during training. It is user-friendly so that non-expert users can try several classifiers with different preprocessing steps without actually writing code.

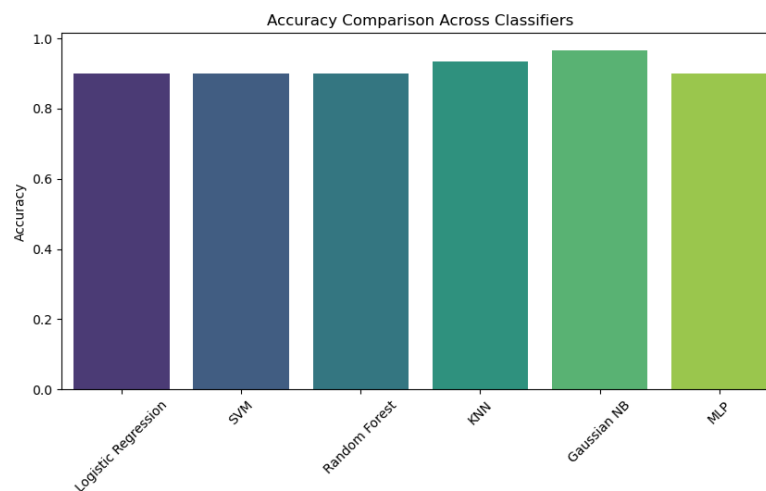


Fig. 3. Model's Accuracy

Flowchart for Machine Learning System

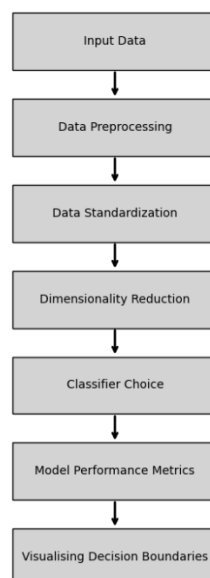


Fig. 4. Block diagram of the proposed system

4. EXPERIMENTAL RESULTS

Ref.no	Objectives	Limitations	Advantages	Gaps
[10]	Evaluate Random Forest for feature selection, compare it with other classifiers (SVM, KNN, LDA), and assess feature selection methods (varImp(), Boruta, RFE) across three datasets.	Focuses on a limited set of datasets and feature selection methods; may not generalize to all classification tasks or data	Demonstrates the outstanding performance of Random Forest in feature selection and classification, making provisions for a comparative analysis of classifiers and feature selection techniques.	Limited to specific datasets and classifiers, may not address all feature selection techniques or complex data environments
[9]	Survey interpretative methods for Random Forest, address its "Black Box" nature, and compare rule extraction methodologies (iTrees and NodeHarvest) across diverse datasets.	Focuses only on two rule extraction methods, may not address all interpretability techniques, limited to six specific datasets.	Provides a comparative analysis of rule extraction methods, enhances understanding of Random Forest interpretability, applies to real datasets with varied characteristics.	Limited exploration of other interpretability methods, may not cover all use cases or dataset types, focuses on Random Fores
[14]	Examine the impact of under- and over-preprocessing on classification performance, compare popular preprocessing techniques, and assess their effects on machine learning algorithms using the Wisconsin Diagnosis Breast Cancer dataset.	Concentrates on single dataset and cannot generalize to all types of data or classification problems.	Highlights the balance needed in preprocessing, provides insights into how preprocessing affects accuracy and precision	Limited to a specific dataset and techniques, may not explore preprocessing effects across diverse datasets or
[15]	Compare four machine learning models (DT, RF, KNN, SVM) for heart disease prediction, evaluate based on accuracy, recall rates, and F1 score, and integrate the best model with a GUI to enhance diagnostic interactivity.	Focuses only on four models, may not consider other advanced algorithms or real-world complexities in heart disease diagnosis.	SVM model achieves high recall (0.97) while maintaining accuracy, GUI integration improves user interaction and patient screening.	Limited to specific models and metrics, may not address other critical aspects of heart disease prediction or integration with clinical systems

Ref .no	Model	Accuracy	Precision	Recall	F1 Score
[10]	RF+SVM	0.89	0.9137	0.9791	-
	RF+LDA	0.9	0.9226	0.9679	-
	RF+KNN	0.886	0.9080	0.9691	-
	RF+RF	0.9099	0.9122	0.9810	-
[9]	Random Forest	0.9854	0.9911	-	0.9867
	nodeHarvest	0.9223	0.9409	-	0.9283
	inTrees	0.9684	0.9651	-	0.9714
[14]	Decision Tree	96.5	96	96	96
	Support Vector Machines	98.2	98	98	98
	MLP	99.1	99.1	99.1	99.1
	Logistic Regression	97.3	97	97	97
	K-Nearest Neighbors	97.4	97	97	97
	Random Forest	98.2	98	98	98
	Support Vector Machines	0.77	0.77	0.77	0.79
	K-Nearest Neighbors	0.83	0.83	0.83	0.84
[15]	Random Forest	0.83	0.83	0.83	0.84
	Decision Tree	0.78	0.78	0.78	0.78

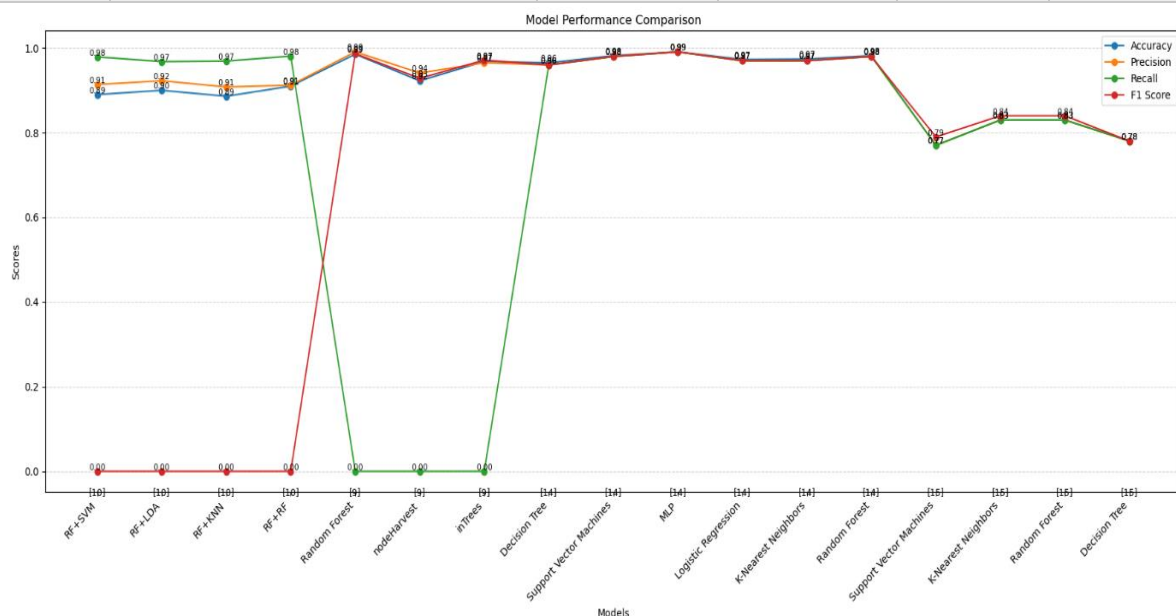


Fig. 5. Performance Comparison of ML Models by Metri

5. CONCLUSION

This review underscores a number of preprocessing techniques-like feature selection, standardizing, and dimensionality reduction-so that the performance of this machine learning model would likely be enhanced. There exist some algorithms like Random Forest, SVM, and even Neural Networks, that immensely benefit from these steps specially for linear and non-linear applications. Accuracy, precision, recall, and F1-score are some of the other performance metrics that would detail the evaluation, particularly so for imbalanced datasets.

A way of reducing dimensionality of the data is achieved through PCA, which also adds up to its interpretability while evading the curse of dimensionality. Good quality data and data preprocessing play a critical role in obtaining well-performing models. Going forward, the hybrid techniques have to be developed by bringing together strengths from single approaches to develop ML that is more reliable, accurate, and adaptable.

6. REFERENCE

- [1] Yacoubby, R., & Axman, D. (2020, November). Probabilistic extension of precision, recall, and f1 score for more thorough evaluation of classification models. In Proceedings of the first workshop on evaluation and comparison of NLP systems (pp. 79-91).
- [2] Wang, Q., Chen, Z., Wang, Y., & Qu, H. (2020). Applying machine learning advances to data visualization: A survey on ML4VIS. arXiv preprint arXiv:2012.00467, 1-18.
- [3] Ray, S. (2019). A quick review of machine learning algorithms. In 2019 International conference on machine learning, big data, cloud and parallel computing (COMITCon) (pp. 35-39). IEEE.
- [4] Moran, K., Bernal-Cárdenas, C., Curcio, M., Bonett, R., & Poshyvanyk, D. (2018). Machine learning-based prototyping of graphical user interfaces for mobile apps. IEEE Transactions on Software Engineering, 46(2), 196-221.
- [5] Maleki, F., Ovens, K., Najafian, K., Forghani, B., Reinhold, C., & Forghani, R. (2020). Overview of machine learning part 1: Fundamentals and classic approaches. Neuroimaging Clinics, 30(4), e17-e32.
- [6] Oulasvirta, A., Dayama, N. R., Shiripour, M., John, M., & Karrenbauer, A. (2020). Combinatorial optimization of graphical user interface designs. Proceedings of the IEEE, 108(3), 434-464.
- [7] Wardhani, N. W. S., Rochayani, M. Y., Iriany, A., Sulistyono, A. D., & Lestantyo, P. (2021, October). Cross-validation metrics for evaluating classification performance on imbalanced data. In 2021 International conference on computer, control, informatics and its applications (IC3INA) (pp. 14-18). IEEE.
- [8] Picard, S., Chapdelaine, C., Cappi, C., Gardes, L., Jenn, E., Lefèvre, B., & Soumarmon, T. (2020, October). Ensuring dataset quality for machine learning certification. In 2020 IEEE international symposium on software reliability engineering workshops (ISSREW) (pp. 275-282). IEEE.
- [9] Aria, M., Cuccurullo, C., & Gnasso, A. (2021). A comparison among interpretative proposals for Random Forests. Machine Learning with Applications, 6, 100094.
- [10] Chen, R. C., Dewi, C., Huang, S. W., & Caraka, R. E. (2020). Selecting critical features for data classification based on machine learning methods. Journal of Big Data, 7(1), 52.
- [11] Rácz, A., Bajusz, D., & Héberger, K. (2021). Multi-level comparison of machine learning classifiers and their performance metrics. Molecules, 24(15), 2811.
- [12] Hasan, B. M. S., & Abdulazeez, A. M. (2021). A review of principal component analysis algorithm for dimensionality reduction. Journal of Soft Computing and Data Mining, 2(1), 20-30.
- [13] Çetin, V., & Yıldız, O. (2022). A comprehensive review on data preprocessing techniques in data analysis. Pamukkale Üniversitesi Mühendislik Bilimleri Dergisi, 28(2), 299-312.
- [14] Rao, S., Poojary, P., Somaiya, J., & Mahajan, P. (2020). A comparative study between various preprocessing techniques for machine learning. International Journal of Engineering Applied Sciences and Technology, 5(3), 2455-2143.
- [15] Yang, Y. (2024). Heart disease prediction and GUI interaction based on machine learning. Transactions on Computer Science and Intelligent Systems Research, 5, 209-218.