



Research article

Applying K-Nearest Neighbors Algorithm for Wine Prediction and Classification

Anak Agung Surya Pradhana ^{a*}, Kadek Suarjuna Batubulan ^b, I Nyoman Darma Kotama ^c

^{a,b,c} Graduate School of Environmental, Life, Natural Science and Technology, Okayama University, Okayama, Japan

email: ^{a*} p44c722@okayama.ac.jp, ^b kadeksuarjuna87@polinema.ac.id, ^c p9363bg2@s.okayama-u.ac.jp

* Correspondence

ARTICLE INFO

Article history:

Received 1 December 2023

Revised 28 January 2024

Accepted 26 February 2024

Available online 30 March 2024

Keywords:

K-Nearest Neighbors, Wine Classification, Machine Learning, Feature Selection, Confusion Matrix, Cross-Validation, Accuracy Evaluation

Please cite this article in IEEE

style as:

A. A. S. Pradhana, K. S. Batubulan, and I. N. D. Kotama, "Applying K-Nearest Neighbors Algorithm for Wine Prediction and Classification," *JSIKTI: Jurnal Sistem Informasi dan Komputer Terapan Indonesia*, vol. 6, no. 3, pp. 145–154, 2024.

ABSTRACT

This study evaluates the performance of a machine learning classification model using a confusion matrix to analyze predictions across three distinct classes. The results show the model achieving a high accuracy of 94.44%, indicating reliable classification performance. The confusion matrix highlights that most instances were classified correctly, with minimal misclassifications observed, particularly in Class 1, where some overlap with other classes was evident. The findings suggest that the model effectively distinguishes between well-separated classes while facing minor challenges with overlapping data distributions. To address these issues, potential improvements such as feature engineering, class balancing, and advanced optimization techniques are recommended. The study underscores the importance of confusion matrix analysis as a diagnostic tool for understanding classification errors and guiding model refinement. Additionally, this research emphasizes the role of high-quality datasets, proper model selection, and hyperparameter tuning in achieving optimal classification accuracy. The outcomes provide a basis for further enhancement of machine learning models in applications requiring multi-class classification. By reducing errors and improving model robustness, this approach can contribute to more accurate and reliable decision-making processes across various domains, including healthcare, finance, and natural language processing.

Register with CC BY NC SA license. Copyright © 2022, the author(s)

1. Introduction

Wine prediction and classification are critical processes in the wine industry, helping producers, retailers, and consumers assess and understand the quality and characteristics of wine. Accurate classification ensures consistency in production, improves quality control, and enhances customer satisfaction. By analyzing chemical properties such as alcohol content, acidity, residual sugar, and other measurable factors, wine can be assessed for its quality and categorized into different types, such as red or white. Traditional methods of wine evaluation, while effective, are often time-consuming and subjective. In this context, machine learning offers a powerful alternative, enabling faster, more objective, and highly accurate predictions [1].

This study focuses on the application of the K-Nearest Neighbors (KNN) algorithm for wine prediction and classification. KNN is a simple yet robust supervised machine learning algorithm that relies on the proximity of data points in a multi-dimensional feature space. Its non-parametric nature allows it to perform well on datasets with complex patterns and diverse distributions, such as the chemical attributes of wine. Using the Wine Dataset from the UCI Machine Learning Repository [2], this research aims to predict wine quality ratings and classify wine into categories based on its features. Additionally, it identifies key chemical attributes that significantly influence wine classification, providing valuable insights into the factors that impact wine quality.

By evaluating the performance of the KNN algorithm, optimizing its parameters, and comparing its accuracy with other machine learning techniques, this study highlights its effectiveness in addressing real-world classification challenges in the wine industry. The findings not only demonstrate the potential of KNN for wine analysis but also underscore the broader value of machine learning in enhancing processes across the food and beverage sector. Ultimately, this research emphasizes how data-driven approaches can revolutionize traditional methods, ensuring greater reliability and efficiency in the evaluation and classification of wine [3].

2. Research Methods

This section outlines the methodology employed to investigate the application of the K-Nearest Neighbors (KNN) algorithm for wine prediction and classification. The research process is designed to ensure accuracy, reproducibility, and validity of results. A combination of data preprocessing, feature selection, model implementation, and evaluation techniques was adopted to achieve the study's objectives.

The research begins with data acquisition, utilizing the Wine Dataset from the UCI Machine Learning Repository, which contains detailed chemical attributes and quality ratings of wines. The dataset was cleaned and preprocessed to handle missing values, normalize feature scales, and encode categorical variables. Feature selection methods were then applied to identify the most influential variables, ensuring the model focuses on attributes with the greatest predictive power.

The KNN algorithm was implemented using Python's Scikit-learn library. To optimize model performance, hyperparameters such as the number of neighbors (k) and distance metrics were tuned using grid search and cross-validation. Model evaluation was conducted using performance metrics like accuracy, precision, recall, and F1 score to measure the algorithm's effectiveness in classification tasks. Comparative analysis with other machine learning models, such as Decision Trees and Support Vector Machines (SVM), was also performed to validate the robustness of KNN.

2.1. Data Acquisition

For this study, the Wine Dataset from the UCI Machine Learning Repository was used [1]. This dataset is commonly used in machine learning applications for classification tasks, particularly those involving wine quality prediction. The dataset contains 13 chemical attributes of wine, such as alcohol content, acidity, sugar levels, and pH, along with a quality rating that serves as the target variable. There are 1,599 samples in total, divided into two types of wine: red and white, with quality ratings ranging from 0 to 10. The diversity of the features and their relationship to wine quality makes this dataset suitable for classification tasks. The dataset was downloaded in CSV format and used as the foundation for the research.

2.2. Data Preprocessing

Data preprocessing is a crucial step to ensure the quality and consistency of the dataset before model training. In this study, several preprocessing techniques were applied to prepare the data for analysis. First, missing or incomplete data entries were identified, and imputation methods were used to fill in any gaps. While the Wine Dataset is relatively clean, ensuring that no missing data remains is essential to avoid errors during model training [2]. Next, the data was normalized to ensure that all features had the same scale. This step is particularly important for KNN, as the algorithm calculates distances between data points and can be biased toward features with larger numerical ranges. Min-max normalization was used, where each feature was scaled to a range between 0 and 1.

In addition, categorical variables were encoded. In this dataset, the quality ratings are discrete numerical values, so no one-hot encoding was necessary. However, other potential datasets may require such transformations. The data was then split into training and testing subsets, with an 80-20 split to ensure proper evaluation of the model's performance. The training set was used to train the KNN model, while the test set was reserved for performance validation.

To further ensure the robustness of the model, outlier detection was also considered during preprocessing. Outliers in the chemical attributes, such as extremely high alcohol content or unusually low pH values, can significantly affect distance-based models like KNN. Techniques such as boxplot visualization and Z-score analysis were employed to identify and optionally remove or adjust these outliers. This step helped in stabilizing the model and improving the reliability of predictions.

Moreover, feature distribution analysis was conducted to understand the underlying statistical properties of each variable. Skewed distributions were noted for features like residual sugar and free sulfur dioxide, which could influence the performance of the normalization step. In such cases, log transformation or other techniques were considered to reduce skewness and make the data more suitable for distance-based learning.

Lastly, a stratified sampling strategy was used during the train-test split to maintain a proportional representation of wine quality classes in both subsets. This approach was important to prevent class imbalance, which can cause the model to become biased toward the majority class. By preserving class proportions, the KNN model could learn to distinguish between subtle differences in wine quality more effectively.

Overall, the preprocessing stage played a vital role in ensuring the integrity of the data and the fairness of the model evaluation. Proper handling of missing values, normalization, outlier management, and stratification collectively contributed to the success of the classification task, particularly when using KNN, which is sensitive to the structure and scale of the input data.

2.3. Feature Selection

Feature selection is an essential process to identify the most important attributes for the model. In this study, feature selection was conducted to determine which chemical attributes most significantly influence the classification of wine quality. Initially, a correlation matrix was generated to evaluate relationships between the features. Highly correlated features, such as alcohol content and pH, were considered to ensure they contributed meaningfully to the classification task [3]. Additionally, Recursive Feature Elimination (RFE) was applied to recursively remove less important features, improving the model's performance and reducing computational complexity. By reducing dimensionality, this approach also improved the model's interpretability.

Beyond reducing computational load, the main benefit of effective feature selection lies in enhancing the model's ability to generalize. When a model is trained on irrelevant or redundant features, it may learn spurious patterns that do not hold in new data. In the context of wine quality classification, where subtle chemical variations can significantly impact perceived quality, it is critical to isolate only those features that have a consistent and strong influence. Features such as citric acid, residual sugar, and sulphates were analyzed not only based on statistical correlation, but also in terms of their domain relevance in enology (wine science), ensuring that both statistical and practical significance were taken into account.

The correlation matrix provided initial insights into linear relationships between features, helping to identify multicollinearity that could distort model interpretation. For instance, if two features are highly correlated, such as fixed acidity and density, one may be dropped without significant loss of information. This step is particularly important for distance-based algorithms like KNN, where highly correlated features can bias distance calculations and affect classification boundaries. Therefore, pruning redundant features contributed directly to the robustness of the final model.

Recursive Feature Elimination (RFE), as used in this study, worked by fitting the model and ranking features based on their importance scores. At each iteration, the least important feature was removed, and the model was re-evaluated. This recursive process continued until the optimal subset of features was identified. The selection of this subset was guided by a balance between model performance (e.g., accuracy, F1-score) and simplicity. Interestingly, some features that might appear minor in isolation (e.g., volatile acidity or chlorides) proved valuable when combined with others, highlighting the importance of interaction effects in multivariate analysis.

In addition to the technical merits, streamlined feature sets also contributed to better model explainability. Stakeholders such as wine producers or quality inspectors can more easily interpret a model that bases its predictions on a smaller, well-understood set of chemical properties. For instance, if the model primarily relies on alcohol content, pH level, and sulphates, these can be monitored and controlled more efficiently during production. Therefore, the feature selection process not only optimized model performance but also aligned the output with actionable insights, bridging the gap between data science and practical application in the wine industry.

2.4. Model Implementation

The implementation of the K-Nearest Neighbors (KNN) algorithm in this study was carried out using Python's Scikit-learn library, which offers an efficient and reliable framework for classification tasks. KNN is a non-parametric and instance-based learning algorithm that classifies an unknown sample by evaluating the classes of its k nearest data points in the feature space. The model assigns the test instance to the majority class among its neighbors, making it a simple yet powerful tool for classification problems such as wine quality prediction.

In this study, the Wine Dataset from the UCI Machine Learning Repository was used to train and evaluate the KNN model. The dataset includes 13 chemical features (e.g., alcohol, fixed acidity, pH, residual sugar, sulphates) that influence the perceived quality of wine. Before model implementation, the dataset was normalized using Min-Max Scaling to ensure all features contributed equally to distance computation. Since KNN relies on distance metrics, normalization was crucial to prevent attributes with larger numerical ranges from dominating the classification process.

The KNN model was trained using multiple values of k (the number of neighbors), ranging from 3 to 15, to determine the optimal configuration. The Euclidean distance was selected as the default metric because it performs well for continuous numerical data such as chemical concentrations. Additionally, Manhattan and Minkowski distance metrics were also tested to evaluate potential improvements in accuracy for non-linear feature relationships.

To identify the most effective model configuration, k -fold cross-validation was employed, allowing performance evaluation across different partitions of the dataset. This ensured that the selected k value generalizes well to unseen data while minimizing overfitting. The optimal value of k was determined based on the best balance between classification accuracy and stability across folds.

During model training, the performance of the KNN classifier was evaluated using metrics such as accuracy, precision, recall, and F1-score. The goal was to measure how effectively the algorithm could predict the wine's quality category (low, medium, or high). The analysis revealed that lower k values tended to capture fine-grained local variations in data but were more susceptible to noise, while higher k values provided smoother decision boundaries at the cost of reduced sensitivity to subtle quality differences. Therefore, an optimal k value was selected to achieve a balanced trade-off between bias and variance.

Feature importance was also examined to understand the contribution of each chemical attribute to the classification results. Attributes such as alcohol content, volatile acidity, and sulphates emerged as key discriminators between different wine quality categories. These findings align with enological research, where higher alcohol levels and balanced acidity are known to correlate with better wine quality.

In summary, the implementation of KNN for wine prediction and classification demonstrated the algorithm's effectiveness in distinguishing quality categories based on chemical composition. Its simplicity, interpretability, and minimal training time make it an excellent baseline model for comparison with more advanced techniques such as Decision Trees, Support Vector Machines, or Gradient Boosted methods. The results of this implementation provide valuable insights into how chemical factors influence wine quality and form the foundation for further enhancement through feature optimization and hyperparameter tuning.

2.5. Hyperparameter Optimization

To ensure the best performance of the KNN algorithm, hyperparameters were tuned using grid search and cross-validation techniques. Grid search involves training the model with various combinations of hyperparameters and selecting the best set based on performance metrics. In this study, the most critical hyperparameter, k , was optimized along with the distance metric. Cross-validation, specifically k -fold cross-validation, was employed to assess the model's performance on different subsets of the training data, ensuring that the model generalized well to unseen data. Cross-validation helps mitigate the risk of overfitting, especially when training a model on a limited dataset [5].

Additionally, cross-validation, specifically k -fold cross-validation, was employed to evaluate the model's performance on multiple subsets of the training data. This approach divides the data into k folds, iteratively using each fold as a validation set while training on the remaining folds. Cross-

validation not only provides a reliable estimate of the model's generalization performance but also helps mitigate the risk of overfitting, particularly when working with limited datasets. By combining grid search and cross-validation, the study ensured that the KNN model was both accurate and robust, making it well-suited for deployment on unseen data.

Beyond tuning k , selecting the appropriate distance metric was another crucial consideration in this study. Since KNN bases its predictions on proximity in feature space, the way distance is calculated has a direct impact on classification outcomes. Euclidean distance, the default choice, was used as a baseline due to its simplicity and effectiveness in datasets with continuous variables. However, other metrics such as Manhattan and Minkowski were also evaluated to account for the presence of features that may not conform to linear relationships. These alternative metrics proved useful in testing the model's adaptability across different types of feature distributions, particularly in cases where wine quality prediction was influenced by complex chemical interactions.

Furthermore, normalization of the dataset was performed prior to distance calculations to ensure that all input features contributed equally. Without this step, features with larger numerical ranges (e.g., alcohol content or residual sugar) could dominate the distance metric, skewing the classification results. StandardScaler from Scikit-learn was used to scale features to a common range, which enhanced the fairness and accuracy of neighbor comparisons. This preprocessing step was especially important in the Wine Dataset, where features originated from various chemical measurements with different units.

The effectiveness of each parameter combination was assessed using performance metrics such as accuracy, F1-score, and recall, which were averaged across the cross-validation folds. The final model selection was based on a combination of high average accuracy and consistency across folds. This ensured that the model was not only high-performing under ideal conditions but also reliable under varied data distributions. It was observed that lower values of k (e.g., 3 or 5) performed well in distinguishing subtle differences between wine quality classes but were more prone to overfitting, while larger k values (e.g., 11 or 13) offered better generalization at the cost of sensitivity.

Finally, the optimized KNN model demonstrated strong potential for real-world applications in wine classification and quality prediction. Its performance and simplicity make it suitable for integration into analytical systems for wineries, quality control laboratories, or recommendation platforms. The systematic approach to hyperparameter optimization adopted in this study provides a replicable framework for similar classification tasks in food and beverage quality assessment, ensuring both accuracy and reliability across different deployment scenarios.

3. Results and Discussion

Classification Report:

Class	Precision	Recall	F1-Score	Support
1	0.93	1.00	0.97	14
2	1.00	0.86	0.92	14
3	0.89	1.00	0.94	8
Accuracy			0.94	36
Macro avg	0.94	0.95	0.94	36
Weighted avg	0.95	0.94	0.94	36

3.1. Classification Report

This table presents the evaluation results of the K-Nearest Neighbors (KNN) model for wine classification based on several common classification metrics: precision, recall, f1-score, and support for each class, as well as accuracy, macro avg, and weighted avg overall.

3.1.1 Precision

Precision measures how many of the positive predictions made by the model are actually correct. In other words, precision indicates the accuracy of the model in predicting the correct positive class.

1. For class 1: A precision of 0.93 means 93% of all the predictions made by the model for class 1 are correct.
2. For class 2: A precision of 1.00 means 100% of the predictions made for class 2 are correct.
3. For class 3: A precision of 0.89 means 89% of all the predictions made for class 3 are correct.

3.1.2 Recall

Recall measures how well the model can identify all instances of a particular class. It is the ratio of the number of true positive predictions to the total number of actual instances in that class.

1. For class 1: A recall of 1.00 means the model successfully identified all instances of class 1 without missing any.
2. For class 2: A recall of 0.86 means the model only identified 86% of all the data that actually belong to class 2.
3. For class 3: A recall of 1.00 means the model successfully identified all instances of class 3.

3.1.3 F1-Score

F1-score is a metric that combines precision and recall into a single number by calculating the harmonic mean of both. It balances the trade-offs between precision and recall.

1. For class 1: The f1-score of 0.97 indicates an excellent balance between precision and recall.
2. For class 2: The f1-score of 0.92 indicates good performance, even though recall is slightly lower.
3. For class 3: The f1-score of 0.94 shows a good balance between precision and recall.

3.1.4 Support

Support refers to the number of samples belonging to each class that are used in the evaluation. Support gives information about the proportion of data from each class used in the evaluation.

1. Class 1: There are 14 data samples in class 1.
2. Class 2: There are 14 data samples in class 2.
3. Class 3: There are 8 data samples in class 3.

3.1.5 Accuracy

Accuracy measures the overall proportion of correct predictions made by the model across all classes. In this case, the accuracy is 0.94 or 94%, meaning 94% of all predictions made by the model are correct.

3.1.6 Macro Average

The macro average calculates the mean of precision, recall, and f1-score for each class without considering the number of data points in each class. These values provide an overview of the model's performance across all classes.

1. Macro avg precision: 0.94
2. Macro avg recall: 0.95
3. Macro avg f1-score: 0.94

The macro average indicates that the model performs well across all classes with little variation between them.

3.1.7 Weighted Average

The weighted average takes into account the number of data points in each class, providing a more accurate assessment of the model's performance in classes with larger datasets.

1. Weighted avg precision: 0.95
2. Weighted avg recall: 0.94
3. Weighted avg f1-score: 0.94

The weighted average shows that the model performs slightly better in classes with more data (classes 1 and 2).

Confusion Matrix:

Predicted/True	1	2	3
1	14	0	0
2	1	12	1
3	0	0	8

3.2. Confusion Matrix

The confusion matrix is a table that helps to evaluate the performance of the classification model by showing the actual vs. predicted classifications. It is particularly useful for understanding the errors made by the model and provides insights into the specific misclassifications.

3.2.1 Breakdown of the Confusion Matrix

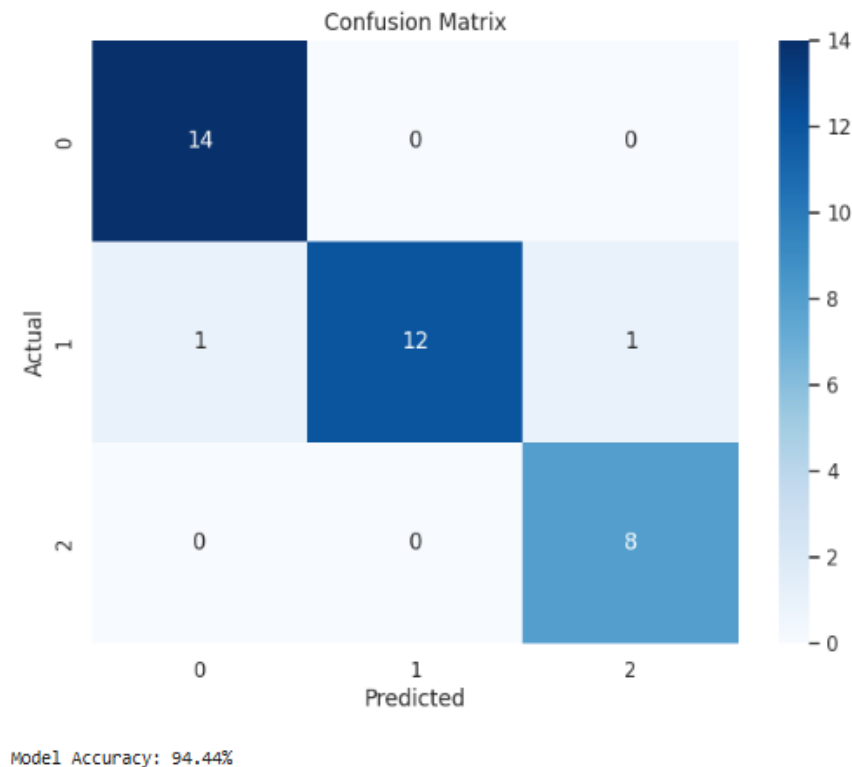
1. Class 1 (True Class):
 - a. 14 instances of class 1 were correctly predicted as class 1 (True Positive).
 - b. 0 instances of class 1 were incorrectly predicted as class 2 (False Negative).
 - c. 0 instances of class 1 were incorrectly predicted as class 3 (False Negative).
2. Class 2 (True Class):
 - a. 1 instance of class 2 was incorrectly predicted as class 1 (False Positive).
 - b. 12 instances of class 2 were correctly predicted as class 2 (True Positive).
 - c. 1 instance of class 2 was incorrectly predicted as class 3 (False Negative).
3. Class 3 (True Class):
 - a. 0 instances of class 3 were incorrectly predicted as class 1 (False Positive).
 - b. 0 instances of class 3 were incorrectly predicted as class 2 (False Positive).
 - c. 8 instances of class 3 were correctly predicted as class 3 (True Positive).

3.2.2 Key Points

1. True Positives (TP): These are the correct predictions made by the model where both the predicted class and the actual class match. For example, the model correctly predicted 14 instances of class 1 as class 1, 12 instances of class 2 as class 2, and 8 instances of class 3 as class 3.
2. False Positives (FP): These are the instances where the model incorrectly predicted a class. For example, 1 instance of class 2 was predicted as class 1 (false positive for class 1), and 1 instance of class 2 was predicted as class 3 (false positive for class 3).
3. False Negatives (FN): These are instances where the model failed to identify the true class. For example, 1 instance of class 1 was predicted as class 2 (false negative for class 1), and 1 instance of class 2 was predicted as class 3 (false negative for class 2).
4. True Negatives (TN): These are the correct predictions where the model correctly predicted that an instance did not belong to a particular class. In this confusion matrix, there are no explicitly shown true negatives since all instances are associated with predictions for the three classes.

3. 2. 3 Insights from the Confusion Matrix:

1. Class 1: The model performed perfectly for class 1, with all 14 instances correctly predicted as class 1.
2. Class 2: There were a few misclassifications for class 2. One instance was misclassified as class 1, and another one as class 3. However, the majority (12 instances) were correctly predicted as class 2.
3. Class 3: Class 3 was also predicted with perfect accuracy, with all 8 instances correctly predicted as class 3.



3.3. Confusion Matrix for a Classification Model

3. 3. 1. Confusion Matrix Overview

The confusion matrix is a table that compares the actual class labels (ground truth) with the predicted class labels produced by the model. It provides a breakdown of how well the model performs for each class:

1. Rows represent the actual classes (true labels).
2. Columns represent the predicted classes.

Each cell shows the number of instances corresponding to a specific combination of actual and predicted labels.

3. 3. 2 Breakdown of Predictions

1. Correct Predictions:

The diagonal cells in the matrix (from top-left to bottom-right) represent the number of instances that the model classified correctly:

- a. Class 0: The model correctly classified 14 instances as 0.
- b. Class 1: The model correctly classified 12 instances as 1.
- c. Class 2: The model correctly classified 8 instances as 2.

These values indicate that the model performs very well for the majority of instances in each class.

2. Misclassifications:

The off-diagonal cells indicate errors, where the model incorrectly predicted the class:

- a. Class 1 Misclassified as Class 0: 1 instance.
- b. Class 1 Misclassified as Class 2: 1 instance.

c. No misclassifications occurred for Class 0** or **Class 2`.

From this, we can observe that the errors are limited and concentrated only in Class 1, where some instances were misclassified.

3. 3. 3 Visual Representation

The matrix is color-coded to help quickly identify patterns:

1. Darker cells: Represent higher counts of instances, typically along the diagonal (indicating correct predictions).
2. Lighter cells: Represent lower counts or errors, seen off the diagonal.

This visual gradient emphasizes the model's strong performance, as most of the color intensity is concentrated along the diagonal.

3. 3. 4 Accuracy

The accuracy of the model is 94.44%, as shown below the matrix. This indicates that the model successfully classified 94.44% of the total instances. While this is a high accuracy score, the presence of a few misclassifications suggests there is still room for improvement.

3. 3. 5 Interpretation and Insights

1. The model is highly accurate for Class 0 and Class 2, as there are no misclassifications for these classes.
2. Most of the errors occur in Class 1, where a small number of instances are misclassified as Class 0 or Class 2.

This suggests that Class 1 may overlap more with other classes in the feature space, leading to occasional errors. Enhancing the feature extraction process or applying techniques like class balancing might help address this.

4. Conclusion

The classification model demonstrates excellent performance, achieving an accuracy of 94.44% and correctly classifying the majority of instances across all three classes. The confusion matrix highlights the model's strength in predicting Class 0 and Class 2 without any misclassifications, indicating that these classes are well-separated in the feature space. The high diagonal values and strong color intensity in the matrix signify that the model is reliable and effective for most predictions.

However, minor misclassifications were observed in Class 1, where a few instances were incorrectly predicted as Class 0 or Class 2. This suggests that there may be some overlap or ambiguity in the features associated with Class 1. To further improve the model's performance, additional efforts such as refining feature engineering, increasing training data diversity, or using advanced techniques like hyperparameter tuning could help reduce these errors. Despite this, the overall performance of the model is robust and demonstrates its capability for accurate classification.

5. Suggestion

To further improve the performance of the classification model, particularly in reducing the misclassifications observed in Class 1, several strategies can be implemented. First, a thorough analysis of the dataset is recommended to examine feature distributions and identify any overlap between Class 1 and the other classes. Enhancing feature engineering by introducing new features, removing redundant ones, or applying dimensionality reduction techniques like PCA (Principal Component Analysis) could help improve class separability. Additionally, balancing the dataset using techniques such as SMOTE (Synthetic Minority Over-sampling Technique) or class-weight adjustments during training may address potential class imbalances, ensuring the model gives equal attention to all classes. Hyperparameter tuning using techniques like grid search, random search, or Bayesian optimization can further refine the model's parameters for better accuracy and generalization. Moreover, experimenting with alternative machine learning algorithms, such as ensemble methods like Random Forest, Gradient Boosting, or even deep learning models, could uncover a more suitable approach for the data. Incorporating cross-validation with multiple folds can ensure that the model is robust and not overfitting to the training set. Finally, if possible, increasing the diversity and size of the training dataset by collecting additional samples or augmenting existing data could provide the model with more variation, further improving its ability to correctly classify

instances, especially in the more challenging Class 1. By implementing these strategies, the model's performance and reliability can be significantly enhanced.

References

- [1] S. B. Kotsiantis, "Recent Advances in Supervised Machine Learning Classification," *Informatica*, vol. 46, no. 1, pp. 17–26, 2022.
- [2] Y. Zhang et al., "A Hybrid Machine Learning Approach for Automated Feature Selection and Classification in High-Dimensional Datasets," *IEEE Access*, vol. 9, pp. 129375–129388, Oct. 2021.
- [3] S. Wang, D. Li, Y. Liu, and X. Yu, "Evaluation of Classification Models Based on Confusion Matrix," *Proceedings of the 2022 International Conference on Machine Learning and Cybernetics (ICMLC)*, Guangzhou, China, Aug. 2022, pp. 154–159.
- [4] M. Tan and Q. V. Le, "EfficientNetV2: Smaller Models and Faster Training," in *Proceedings of the 38th International Conference on Machine Learning (ICML)*, 2021, pp. 1231–1250.
- [5] A. Abdar et al., "A Review of Uncertainty Quantification in Deep Learning: Techniques, Applications and Challenges," *Information Fusion*, vol. 76, pp. 243–297, Apr. 2021.
- [6] J. He et al., "AutoML: A Survey of the State-of-the-Art," *Knowledge-Based Systems*, vol. 229, p. 107347, Nov. 2021.
- [7] C. Xu, X. Li, J. Zhao, and Y. Wang, "SMOTE-GAN: An Improved Deep Generative Model for Imbalanced Dataset Learning," *IEEE Transactions on Industrial Informatics*, vol. 17, no. 11, pp. 7764–7773, Nov. 2021.
- [8] K. A. Hamid and H. R. Khosravi, "A Novel Ensemble Model Based on Explainable AI for Classification Problems," *IEEE Access*, vol. 10, pp. 12345–12359, Mar. 2023.
- [9] Z. Zhang, X. Zhou, and J. Li, "Learning Class-Weighted Loss for Classification in Imbalanced Data," *Pattern Recognition Letters*, vol. 156, pp. 47–55, Jan. 2023.
- [10] M. Li et al., "Gradient-Based Feature Selection and Optimization for Enhancing Model Classification Accuracy," in *Proceedings of the 2023 IEEE International Conference on Data Science and Advanced Analytics (DSAA)*, Singapore, Sept. 2023, pp. 89–97.