

Under-Five Children's Nutritional Status Prediction Using Naïve Bayes and Decision Tree Based on Anthropometric Data and Mother–Child Class Participation

Tsirwatun Nisail Khasanah ^{1*}, Fajar Nugraha ², Yudie Irawan ³

^{1*,2,3} Information Systems Study Program, Faculty of Engineering, Universitas Muria Kudus, Kudus Regency, Central Java Province, Indonesia

Email: 202253136@std.umk.ac.id ^{1*}, fajar.nugraha@umk.ac.id ², yudie.irawan@umk.ac.id ³

Article history:

Received July 22, 2026

Revised August 3, 2026

Accepted August 8, 2026

Abstract

Nutritional status is an important indicator of the health and development of children under five, making early identification essential for supporting appropriate nutritional interventions. This study aimed to develop and compare the performance of the Naïve Bayes and Decision Tree algorithms in classifying the nutritional status of children under five based on anthropometric measurements and participation in the mother and Children Under Five Class program in Mlonggo District, Indonesia. A quantitative approach was applied using the Cross-Industry Standard Process for Data Mining (CRISP-DM) framework. The initial dataset contained 4,588 records, of which 4,563 valid records remained after the preprocessing stage. Model performance was evaluated using accuracy, precision, recall, F1-score, confusion matrix, ROC curve, and McNemar's test. The results showed that the Decision Tree algorithm achieved a higher cross-validation accuracy of 90.90% compared with 87.59% for Naïve Bayes. The testing results also demonstrated that Decision Tree consistently outperformed Naïve Bayes across the evaluation metrics. Therefore, Decision Tree was selected as the most suitable model for nutritional status classification. The model was subsequently implemented in a web-based application supporting individual and batch prediction, along with the presentation of prediction results and recommended health interventions. The system can support healthcare workers in conducting nutritional status assessments more efficiently and objectively.

Keywords:

Nutritional Status of Children Under Five; Data Mining; Naïve Bayes; Decision Tree; CRISP-DM.

1. INTRODUCTION

Based on these research gaps, this study compares the performance of the Naïve Bayes and Decision Tree algorithms for predicting the nutritional status of under-five children using anthropometric measurements and participation in the mother and Child Class program within the CRISP-DM framework. The best-performing model is subsequently implemented in a web-based application to support both individual and batch nutritional status prediction. The proposed system is expected to assist healthcare professionals in conducting nutritional assessments more efficiently, objectively, and accurately. The nutritional status of children under five is an important reflection of their health, growth, and development. Nutritional problems, including severe undernutrition, undernutrition, risk of overweight, overweight, and obesity, remain a concern because they may have long-term effects on children's physical and cognitive development. The 2023 Indonesian Nutritional Status Survey (SSGI) emphasizes the importance of identifying nutritional problems at an early stage so that interventions can be provided according to the needs of each child (Ministry of Health of the Republic of Indonesia, 2023). In practice, nutritional status is commonly determined through anthropometric measurements using the Child Growth Standards developed by the World Health Organization (WHO) (World Health Organization, 2022).

The growing amount of health data has opened opportunities for the use of data mining in health analysis and decision-making. Classification techniques can help identify patterns in existing data and use these patterns to predict nutritional conditions based on relevant characteristics. One approach commonly used to develop data mining models is the Cross-Industry Standard Process for Data Mining (CRISP-DM). This framework provides a structured process consisting of business understanding, data understanding, data preparation, modeling, evaluation, and deployment (Martínez-Plumed et al., 2021). The use of this approach can support a more organized analysis process and encourage data-driven decision-making (Kotu & Deshpande, 2022).

Several studies have examined machine learning methods for assessing the nutritional status of children under five. Naïve Bayes is a probability-based algorithm that is relatively efficient and has been applied to nutritional status classification (Nurainun et al., 2023). Its use has also been reported in research on the classification and monitoring of nutritional conditions in young children (Harliana et al., 2022). Another method, Decision Tree, is often considered useful because its tree-based structure allows the resulting classification rules to be interpreted more easily (Aryaningsih et al., 2025). Previous studies have investigated the use of Decision Tree for nutritional status classification and compared its performance with other machine learning algorithms (Adheilla, 2025). Research has also compared Random Forest and Decision Tree for classifying the nutritional status of children under five (Hidayat et al., 2026). These findings indicate that the performance of an algorithm may vary depending on the characteristics of the dataset. Therefore, comparing different algorithms is important when selecting a model for a specific research context.

In addition to anthropometric information, participation in the mother and Children Under Five Class program may provide additional information relevant to nutritional status. The program provides mothers with education related to nutrition, child growth and development, growth monitoring, and childcare. Participation in these activities may help improve mothers' understanding of children's nutritional needs and growth conditions (Hidayati, 2022). Previous research has also indicated that family assistance through mother-focused educational activities may contribute to improvements in children's weight and nutritional status (Elizar et al., 2025). Furthermore, maternal knowledge of nutrition has been associated with the nutritional status of children under five (Ertiana & Zain, 2023). These findings suggest that participation in the mother and Children Under Five Class program can be considered as an additional variable in the analysis.

Although machine learning has been widely applied to nutritional status assessment, studies that combine anthropometric measurements with participation in the mother and Children Under Five Class program remain limited. In addition, research that compares different classification algorithms and implements the selected model into a web-based prediction application is still relatively limited.

Therefore, this study compares the performance of the Naïve Bayes and Decision Tree algorithms in predicting the nutritional status of children under five using anthropometric measurements and participation in the mother and Children Under Five Class program. The analysis follows the CRISP-DM framework, and the model with the best performance is then implemented in a web-based application that supports individual and batch prediction. This application is expected to help healthcare workers conduct nutritional status assessments in a more efficient, objective, and accurate manner.

2. RESEARCH METHOD

This study applied a quantitative approach supported by data mining techniques to develop a model for predicting the nutritional status of children under five. The research process followed the Cross-Industry Standard Process for Data Mining (CRISP-DM) framework. CRISP-DM was selected because it provides a structured sequence for developing data mining models, from defining the research problem to implementing the resulting model. The framework consists of six stages: Business Understanding, Data Understanding, Data Preparation, Modeling, Evaluation, and Deployment (Martínez-Plumed et al., 2021).

2.1. Data Source and Collection

Secondary data for this study were obtained from Mlonggo Community Health Center, Jepara Regency, Indonesia. Initially, the dataset contained 4,588 records of children under five. Following data inspection, preparation, and quality checking, 4,563 records were retained as valid data for model development and evaluation.

Five predictor variables were used in the classification process: gender, age in months, body weight in kilograms, height in centimeters, and participation in the mother and Children Under Five Class program. Nutritional status served as the target variable and was divided into six categories: Severe Undernutrition, Undernutrition, Good Nutritional Status, Risk of Overnutrition, Overnutrition, and Obesity.

Table 1. Research Variables

Variable	Type	Description
Gender	Input	Male/Female
Age (Monts)	Input	Child's age (0-59 monts)
Body Weight	Input	Kilograms (kg)
Height	Input	Centimeters (cm)
Morher and Child Class Partisipation	Input	Yes/No
Nutritional Status	Output	Severely Underweight Underweight Normal Nutritional Status Risk of Overnutrition Overnutrition Obesity

Table 1 presents the variables used in the classification process. Gender represents the demographic characteristics of the children, while age, weight, and height provide anthropometric information for the prediction process. Participation in the mother and Children Under Five Class program was included as an additional predictor representing the mother's involvement in an educational and child assistance program. All five variables were used as input attributes, whereas nutritional status was defined as the target variable with six classification categories: Severe Undernutrition, Undernutrition, Good Nutritional Status, Risk of Overnutrition, Overnutrition, and Obesity.

2.2. Research and Procedure

This study followed the six stages of the CRISP-DM framework: Business Understanding, Data Understanding, Data Preparation, Modeling, Evaluation, and Deployment. The framework was used to organize the research process from problem identification and data analysis to model development, evaluation, and implementation. An overview of the research workflow is presented in Figure 1.

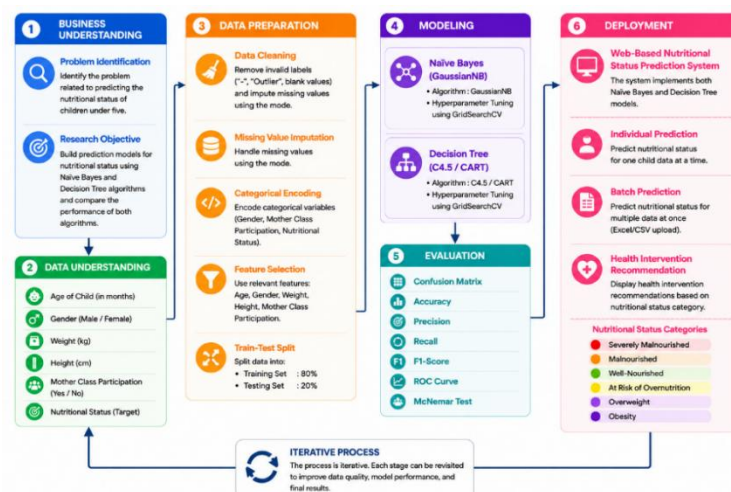


Figure 1. Research Workflow Based on the CRISP-DM Framework

The research procedure is described as follows.

- Business Understanding.** The first stage focused on defining the nutritional status problem among children under five in Mlonggo District. Nutritional status assessment involves several relevant factors, including anthropometric measurements and participation in the mother and Children Under Five Class program. Based on this problem, a prediction model was developed to classify the nutritional status of children under five. This study aimed to compare the performance of the Naïve Bayes and Decision Tree algorithms in classifying nutritional status. Following the comparison, the better-performing model was selected for implementation in a web-based application to support nutritional status assessment and decision-making.
- Data Understanding.** An initial examination was conducted to understand the structure and characteristics of the dataset before preprocessing. The analysis included the number of records, available attributes, data types, descriptive statistics, missing values, and the distribution of each variable. A total of 4,588 records of children under five were obtained from Mlonggo Community Health Center. Data inspection was performed to assess data quality and identify records that were unsuitable for use in the modeling process.
- Data Preparation.** Data preparation was performed to transform the raw dataset into a suitable format for analysis and model development. Records containing invalid target values, such as "Outlier" and "-", were

removed, along with incomplete records. Missing values in the participation variable for the mother and Children Under Five Class program were handled using mode imputation. Categorical variables were then converted into numerical values through encoding to ensure compatibility with the selected machine learning algorithms. Relevant variables were retained as features for the classification process. An 80:20 Stratified Split was subsequently applied to divide the dataset into training and testing subsets. This approach helped maintain a similar distribution of nutritional status classes in both subsets.

- d. **Modeling.** Two classification algorithms, Gaussian Naïve Bayes and Decision Tree, were applied during the modeling stage to compare their ability to classify nutritional status based on the selected predictor variables. Hyperparameter optimization was performed using GridSearchCV with Stratified 5-Fold Cross-Validation. This process was used to identify suitable parameter configurations for each algorithm and obtain a more reliable assessment of model performance based on the training data.
- e. **Evaluation.** Model performance was assessed using Accuracy, Precision, Recall, F1-score, Confusion Matrix, ROC Curve, and McNemar's Test. These measures provided a comprehensive basis for comparing the classification performance of the Naïve Bayes and Decision Tree algorithms. The evaluation focused on identifying the algorithm that provided the most suitable overall performance for the dataset, while also examining each model's ability to distinguish among the six nutritional status categories.
- f. **Deployment.** Following model evaluation, the trained Gaussian Naïve Bayes and Decision Tree models were integrated into a web-based application that supports both individual and batch nutritional status prediction. Prediction results are presented as nutritional status classifications and are accompanied by health intervention recommendations corresponding to the predicted category. Through this implementation, the model can function as part of a decision-support system to assist healthcare workers in assessing the nutritional status of children under five more efficiently and consistently.

2.3. Naïve Bayes Algorithm

Naïve Bayes is a probabilistic classification algorithm that applies Bayes' Theorem to estimate the probability of an observation belonging to a particular class. The method assumes conditional independence among the predictor variables. Despite this assumption, Naïve Bayes is commonly used because it offers efficient computation, straightforward implementation, and good performance across various classification tasks (Nurainun et al., 2023).

Gaussian Naïve Bayes was used in this study because the predictor variables include numerical attributes, such as age, weight, and height, which can be represented using a Gaussian distribution. Bayes' Theorem was applied to calculate the posterior probability for each nutritional status class, as presented in Equation (1).

$$P(C | X) = \frac{P(X | C) P(C)}{P(X)} \quad (1)$$

Where $P(C | X)$ represents the posterior probability of class C given the feature vector X , $P(X | C)$ denotes the likelihood of observing the feature vector X given class C , $P(C)$ is the prior probability of class C , and $P(X)$ is the probability of observing the feature vector X . The Gaussian Naïve Bayes model was trained using the training dataset, while the var_smoothing hyperparameter was optimized through GridSearchCV with Stratified 5-Fold Cross Validation to obtain the optimal model configuration.

2.4. Decision Tree Algorithm

Decision Tree is a supervised learning algorithm that uses a tree-based structure to represent classification decisions. Data are divided recursively according to selected attributes, with internal nodes representing decision criteria, branches showing possible outcomes, and leaf nodes representing the resulting classes. This structure makes the classification process relatively easy to interpret and is useful when transparent decision rules are required (Ramadhani & Ramadhanu, 2024).

In this study, the Entropy criterion was used to evaluate the quality of potential splits in the Decision Tree model. Entropy was calculated using Equation (2), while Information Gain was calculated using Equation (3) to identify the attribute that provided the most suitable division of the dataset.

$$Entropy(S) = - \sum_{i=1}^n p_i \log_2(p_i) \quad (2)$$

$$Gain(S, A) = Entropy(S) - \sum_{v \in Values(A)} \frac{|S_v|}{|S|} Entropy(S_v) \quad (3)$$

Where $Entropy(S)$ represents the entropy of dataset S , p_i denotes the proportion of instances belonging to class i , $Gain(S, A)$ represents the information gain of attribute A , S_v is the subset of data corresponding to

value v of attribute A , and $\frac{|Sv|}{|S|}$ represents the proportion of subset Sv relative to the entire dataset. Entropy measures the degree of impurity in the dataset, while Information Gain is used to determine the best attribute for splitting the data. The attribute with the highest Information Gain is selected to construct the decision tree. Furthermore, GridSearchCV with Stratified 5-Fold Cross Validation was applied to optimize the Decision Tree hyperparameters, including criterion, max_depth, and min_samples_split, in order to obtain the best classification performance.

2.5. Model Evaluation

Model evaluation was conducted to assess the ability of the classification algorithms to predict the nutritional status of children under five. The assessment used several evaluation measures, namely Accuracy, Precision, Recall, F1-score, Confusion Matrix, Receiver Operating Characteristic (ROC) Curve, and McNemar's Test. Applying multiple evaluation measures enabled a more comprehensive comparison of the two classification models.

Accuracy, Precision, and Recall were calculated using Equations (4), (5), and (6), respectively. Accuracy measures the proportion of correct predictions, Precision evaluates the correctness of positive predictions, and Recall measures the ability of a model to identify actual positive observations.

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (4)$$

$$Precision = \frac{TP}{TP + FP} \quad (5)$$

$$Recall = \frac{TP}{TP + FN} \quad (6)$$

Where TP represents the number of True Positive predictions, TN represents the number of True Negative predictions, FP represents the number of False Positive predictions, and FN represents the number of False Negative predictions.

The F1-score combines Precision and Recall into a single evaluation metric, making it particularly suitable for assessing classification performance on imbalanced datasets. The F1-score is calculated using Equation (7).

$$F1 = \frac{2 \times Precision \times Recall}{Precision + Recall} \quad (7)$$

In addition to these quantitative evaluation metrics, the Confusion Matrix was employed to examine the classification results for each nutritional status category, whereas the ROC Curve was used to assess the discriminative capability of the classification models. Furthermore, the McNemar Test was conducted to determine whether the performance difference between the Naïve Bayes and Decision Tree models was statistically significant.

3. RESULTS AND DISCUSSION

3.1. Dataset Preparation

Before applying the classification algorithms, the dataset underwent several preparation steps to ensure that the data were suitable for analysis. These steps included reviewing the data structure, cleaning the records, addressing missing values, removing invalid target labels, and encoding categorical variables into numerical values.

The dataset initially contained 4,588 records of children under five. During the cleaning process, 25 records were excluded due to invalid or incomplete information. After this process, 4,563 records remained and were used for subsequent analysis. In addition, 22 missing values were found in the dataset and addressed using mode imputation based on the relevant attributes.

Once the dataset had been cleaned and completed, it was separated into training and testing subsets using an 80:20 Stratified Split. This process produced 3,650 training records and 913 testing records. Stratification was applied to maintain a comparable distribution of nutritional status categories across both subsets.

Table 2. Dataset Preparation Result

Description	value
Initial Dataset	4588
Clean Dataset	4563
Removed Records	25
Missing Values	22
Imputation Method	Mode
Training Data	3650
Testing Data	913
Data Split Ratio	80:20

Table 2 summarizes the changes in the dataset during preparation. Of the 4,588 initial records, 25 were removed because of invalid or incomplete data, leaving 4,563 records for analysis. The 22 missing values identified during inspection were addressed using mode imputation. Afterward, the resulting dataset was divided into 3,650 training records and 913 testing records using an 80:20 split.

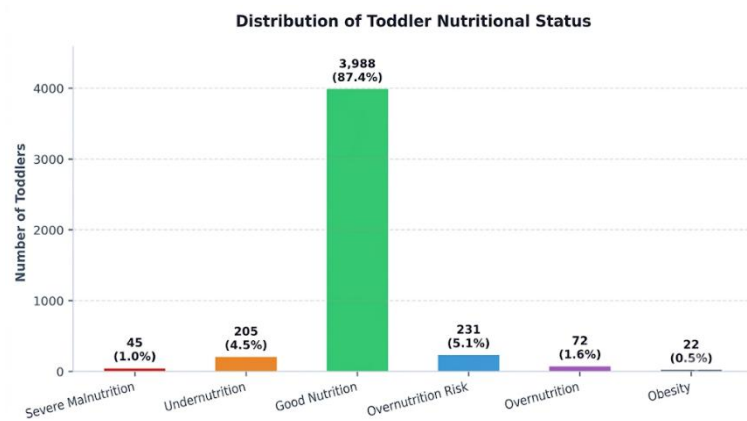


Figure 2. Data Preparation Result and Target Class Distribution

Figure 2 presents the results of the preparation process along with the distribution of nutritional status categories in the final dataset. The figure shows the number of records remaining after cleaning, the division between training and testing data, and the distribution of the six target classes used in the classification process. Most records belonged to the Good Nutritional Status category, which accounted for 3,988 records (87.4%). Risk of Overnutrition followed with 231 records (5.1%), while Undernutrition accounted for 205 records (4.5%). The remaining categories consisted of Overnutrition with 72 records (1.6%), Severe Undernutrition with 45 records (1.0%), and Obesity with 22 records (0.5%).

The distribution shows a considerable difference in the number of observations among the nutritional status categories. Good Nutritional Status dominated the dataset, whereas several other categories contained substantially fewer records. Such an imbalance may affect the ability of a classification model to identify minority classes correctly. Therefore, model assessment considered multiple evaluation measures, including Precision, Recall, F1-score, Confusion Matrix, and ROC Curve, rather than relying solely on Accuracy.

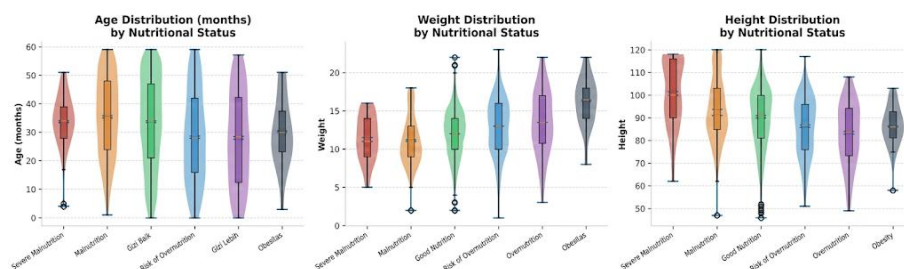


Figure 3. Dataset Feature Distribution

Figure 3 shows the distribution of continuous predictor variables used in the nutritional status classification process. Age, weight, and height are presented to describe the numerical and anthropometric characteristics of children included in the dataset. The distribution of these variables provides an overview of the data before the classification models are developed. Gender and participation in the mother and Children Under Five Class program were also included as predictor variables. However, both variables are categorical attributes that were converted into numerical values during the encoding process, so they are not included in the continuous variable distribution shown in Figure 3. Analysis of feature distributions was conducted to

provide an initial overview of the dataset before modeling. After data cleaning, missing-value handling, categorical encoding, and data splitting were completed, the processed dataset was ready for model development. Training data were used to develop and optimize the Gaussian Naïve Bayes and Decision Tree models, while testing data were kept separate to assess model performance on records that were not involved in the training process.

3.2. Modeling and Algorithm Optimization Results

After the dataset was prepared, the next step was to develop classification models for predicting the nutritional status of children under five. Gaussian Naïve Bayes and Decision Tree were selected as the algorithms for comparison. Both models used the same five predictor variables, consisting of gender, age, weight, height, and participation in the mother and Children Under Five Class program.

Model training was carried out using 3,650 training records. GridSearchCV and Stratified 5-Fold Cross-Validation were applied to search for the most appropriate hyperparameter configuration for each algorithm. The training data were divided into five subsets during cross-validation, allowing each subset to be used as validation data in turn.

Table 3. Hyperparameter Optimization and Cross-Validation Results

Model	Optimal Hyperparameter	Value	Cross-Validation Accuracy
Gaussian Naïve Bayes	var_smoothing	4.3288e-07	87.59%
Decision Tree	criterion	entropy	90.90%
Decision Tree	max_depth	15	90.90%
Decision Tree	min_samples_split	2	90.90%

The results of the optimization process are shown in Table 3. Gaussian Naïve Bayes obtained a cross-validation accuracy of 87.59% with a var_smoothing value of 4.3288e-07. For Decision Tree, the best configuration used the entropy criterion, max_depth of 15, and min_samples_split of 2, resulting in a cross-validation accuracy of 90.90%. The difference between the two models was 3.31 percentage points in favor of Decision Tree. This result shows that, during cross-validation, Decision Tree provided a better fit to the patterns contained in the training data. The optimized configurations were subsequently used for testing on data that had not been involved in model training.

3.3. Model Performance Evaluation

Model evaluation was performed using 913 testing records that were not involved in the training process. Both algorithms were assessed based on their ability to classify nutritional status into six categories: Severe Undernutrition, Undernutrition, Good Nutritional Status, Risk of Overnutrition, Overnutrition, and Obesity. Accuracy, Precision, Recall, F1-Score, Confusion Matrix, ROC Curve, and McNemar's Test were used to obtain a comprehensive comparison of model performance. Multiple evaluation measures were considered because the number of records varied considerably among the six nutritional status categories.

3.3.1. Comparison of Main Evaluation Metrics

Table 4. Comparison of Main Evaluation Metrics

Evaluation Metric	Naïve Bayes	Decision Tree	Difference (DT – NB)
Accuracy (%)	87.51%	91.13%	+3.61%
Macro Precision	0.2293	0.6981	+0.4687
Macro Recall	0.2083	0.6455	+0.4372
Macro F1-Score	0.2112	0.6654	+0.4542

Testing results presented in Table 4 show that Decision Tree achieved an Accuracy of 91.13%, whereas Naïve Bayes obtained 87.51%. Decision Tree also produced higher values for all macro-averaged evaluation measures, with Macro Precision of 0.6981, Macro Recall of 0.6455, and Macro F1-Score of 0.6654. Naïve Bayes recorded lower values, namely 0.2293 for Macro Precision, 0.2083 for Macro Recall, and 0.2112 for Macro F1-Score. These results indicate that Decision Tree provided more consistent classification performance across the nutritional status categories.

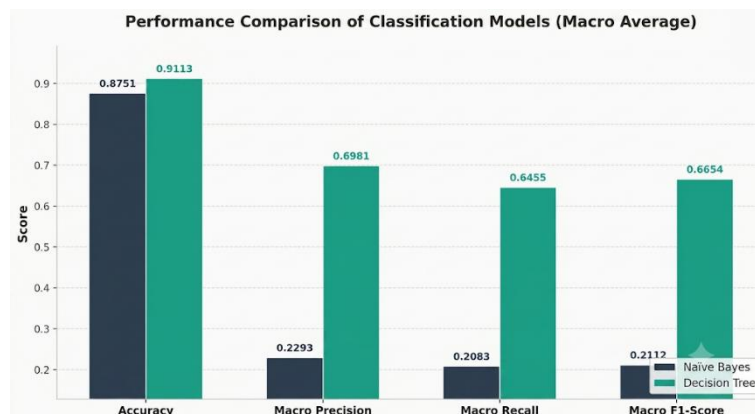


Figure 5. Comparison of Evaluation Metrics

Figure 5 compares the Accuracy, Macro Precision, Macro Recall, and Macro F1-Score obtained from both algorithms. Decision Tree achieved higher values for all four evaluation measures. This result is consistent with the numerical comparison presented in Table 4 and indicates stronger overall performance on the testing dataset.

3.3.2. Performance Comparison Based on Nutritional Status Classes

Table 5. Comparison of Classification Metrics by Nutritional Status Class

Nutritional Status Class	NB Precision	NB Recall	NB F1-Score	Support	DT Precision	DT Recall	DT F1-Score	Support
Severe Undernutrition	0.0000	0.0000	0.0000	9	0.8571	0.6667	0.7500	9
Undernutrition	0.0000	0.0000	0.0000	41	0.5000	0.4878	0.4938	41
Good Nutritional Status	0.8760	1.0000	0.9339	798	0.9541	0.9637	0.9589	798
Risk of Overnutrition	0.0000	0.0000	0.0000	46	0.6316	0.5217	0.5714	46
Overnutrition	0.0000	0.0000	0.0000	15	0.5789	0.7333	0.6471	15
Obesity	0.5000	0.2500	0.3333	4	0.6667	0.5000	0.5714	4
Macro Average	0.2293	0.2083	0.2112	913	0.6981	0.6455	0.6654	913

Table 5 presents the classification performance of both models for each nutritional status category. Naïve Bayes achieved its strongest performance for Good Nutritional Status, with a Recall of 1.0000 and an F1-Score of 0.9339. However, the model did not correctly identify the Severe Undernutrition, Undernutrition, Risk of Overnutrition, and Overnutrition categories. These four categories recorded Precision, Recall, and F1-Score values of 0.0000. Decision Tree was able to produce classification results for all six nutritional status categories. Its F1-Score reached 0.7500 for Severe Undernutrition, 0.4938 for Undernutrition, 0.9589 for Good Nutritional Status, 0.5714 for Risk of Overnutrition, 0.6471 for Overnutrition, and 0.5714 for Obesity.

Overall, Decision Tree showed a greater ability to recognize the different nutritional status categories. However, performance differences between the majority and minority classes were still observed. Good Nutritional Status achieved the highest performance, which is consistent with its substantially larger number of records in the dataset.

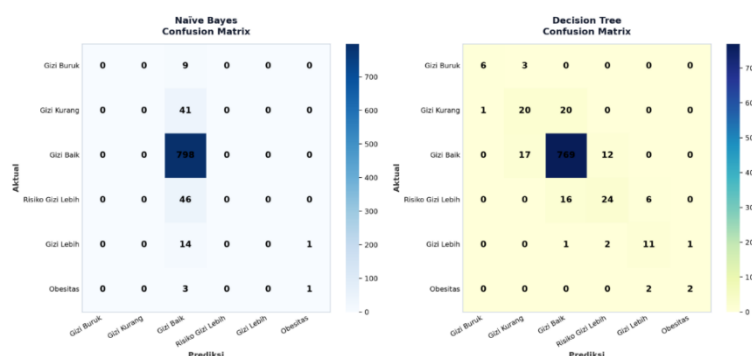


Figure 6. Confusion Matrix of the Naïve Bayes and Decision Tree Models

Figure 6 illustrates the relationship between the actual and predicted classes for both classification models. The Naïve Bayes confusion matrix shows a tendency to assign many records to Good Nutritional Status, resulting in limited recognition of several minority categories. In comparison, Decision Tree produced a more varied distribution of predictions and showed better separation among the nutritional status categories. These findings support the class-level results presented in Table 5, where Decision Tree demonstrated stronger performance in the multiclass classification task.

3.3.3. Evaluation of Model Discriminatory Ability

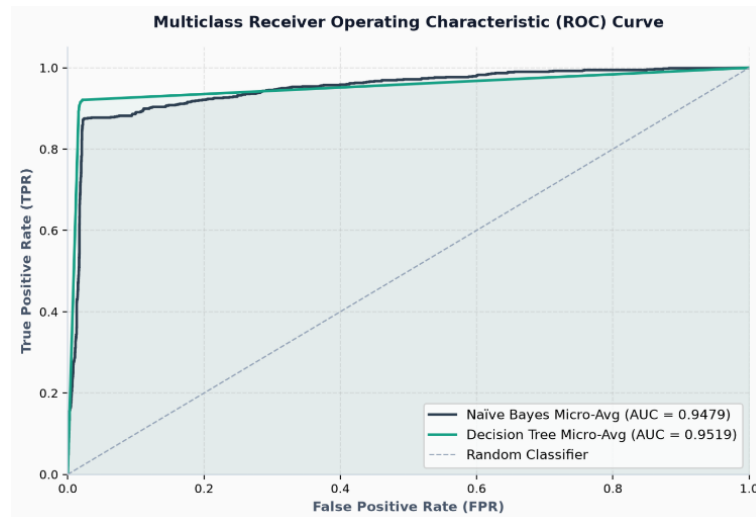


Figure 7. Multiclass ROC and AUC

Figure 7 presents the multiclass Receiver Operating Characteristic (ROC) curves using a One-vs-Rest approach. This method evaluates the ability of each model to distinguish one nutritional status category from the remaining categories. ROC and AUC analysis provides additional information regarding model discrimination beyond Accuracy, Precision, Recall, and F1-Score.

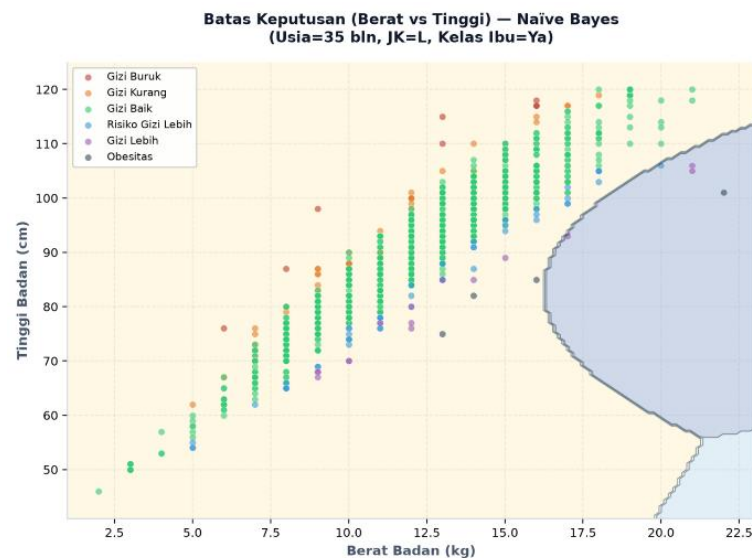


Figure 8. Model Decision Boundaries

Figure 8 illustrates the decision regions generated by the classification model using weight and height as the displayed variables. The visualization provides an overview of how different nutritional status categories are separated within the feature space. Only two variables are displayed for visualization purposes, whereas the actual model was trained using five predictor variables: gender, age, weight, height, and participation in the mother and Children Under Five Class program.

3.4. Web-Based System Implementation

The final prediction model was deployed in a web-based application designed to support nutritional status classification among children under five. The application provides several functions, including data

preprocessing, individual prediction, batch prediction, model selection, and visualization of prediction results. The web-based application incorporates both the optimized Naïve Bayes and Decision Tree models, enabling users to choose either algorithm for prediction while facilitating comparison of their classification outcomes. Users can enter predictor information, such as gender, age, weight, height, and participation in the mother and Children Under Five Class program. After the input data are submitted, the system processes the information and assigns the child to one of six nutritional status categories: Severe Undernutrition, Undernutrition, Good Nutritional Status, Risk of Overnutrition, Overnutrition, or Obesity. Batch prediction is also available, allowing multiple records to be processed simultaneously. Prediction results are presented together with supporting information, including the predicted nutritional status category, confidence level, and health intervention recommendations corresponding to the classification result. These outputs can provide an initial reference for healthcare workers in reviewing the nutritional status of children under five.



Figure 9. Nutritional Status Prediction Output in the System

Figure 9 presents examples of prediction results generated by the web-based application. The displayed outputs include the predicted nutritional status category, confidence level, and input information used during the classification process. The first output shows a Good Nutritional Status prediction generated by the Naïve Bayes model, while the second output indicates Risk of Overnutrition based on the Decision Tree model. These results illustrate how the system processes the entered data and generates nutritional status classifications through the applied prediction models.

In addition to displaying the predicted nutritional status, the application provides health intervention recommendations related to the classification result. This feature adds supporting information to the prediction output and may assist healthcare workers in reviewing the results as part of an initial nutritional status assessment. However, the prediction results are intended to function as supporting information and should be considered alongside professional assessment and other relevant health information. Overall, the implementation of the prediction models in a web-based application demonstrates the practical application of the research findings. By integrating data processing, classification, prediction results, and recommendations within a single system, the application provides a more structured approach to nutritional status assessment. Support for both individual and batch prediction also allows the system to accommodate different data processing needs.

4. CONCLUSION

This study developed a web-based prediction system for identifying the nutritional status of under-five children by applying the Naïve Bayes and Decision Tree algorithms within the Cross-Industry Standard Process for Data Mining (CRISP-DM) framework. The developed system demonstrates the applicability of data mining techniques in supporting the early identification of nutritional status based on anthropometric measurements and participation in the mother and Child Class program.

The findings indicate that both classification algorithms were successfully implemented and evaluated. During the model development stage, the Decision Tree algorithm achieved a higher average cross-validation accuracy (90.90%) than the Naïve Bayes algorithm (87.59%). The evaluation on the testing dataset also confirmed that the Decision Tree model achieved superior predictive performance in terms of accuracy, precision, recall, and F1-score. Both classification models were integrated into the web-based application, allowing users to select either algorithm for prediction. The system is capable of processing both individual and batch predictions while providing nutritional status classifications along with appropriate health intervention recommendations.

Despite these promising results, this study has several limitations. The prediction models were developed using data collected from a single study area and relied only on anthropometric variables and participation in the mother and Child Class program. Future studies are encouraged to employ more diverse datasets, include additional variables that may influence children's nutritional status, and evaluate other machine learning algorithms to further enhance prediction performance and model generalizability.

ACKNOWLEDGEMENTS

The authors would like to thank Mlonggo Community Health Center, Jepara Regency, for providing access to the nutritional status dataset and for its cooperation throughout the data collection and system development process. During the preparation of this manuscript, AI-assisted writing tools, including OpenAI ChatGPT, were used to assist with language refinement, manuscript organization, paraphrasing, and preparation of the research design figure based on the dataset, formulas, and results produced by the authors. All AI-assisted outputs were carefully reviewed, verified, and approved by the authors, who take full responsibility for the accuracy, integrity, and originality of the manuscript.

REFERENCES

- Adheilla, R. (2025). Comparison of Naïve Bayes, K-Nearest Neighbor, and Decision Tree C4.5 models for Children Under Five nutritional status classification (Bachelor's thesis, Universitas Malikussaleh). RAMA Repository Universitas Malikussaleh. <https://rama.unimal.ac.id/id/eprint/15940/>
- Aryaningsih, P., Triapanca, M. D., & Putra, T. D. (2025). Classification of Children Under Five nutritional status using the Decision Tree algorithm in RapidMiner in Gombang Village. In Proceedings of the 4th National Seminar on Research and Community Service in Computer Science (SENDIKO 2025). Universitas PGRI Madiun. <https://prosiding.unipma.ac.id/index.php/sendiko/article/view/6903>
- Elizar, Putri, H. W. K., Prihatin, N. S., Kartinazahri, Nurmila, Rosyita, & Jasmianti. (2025). The effect of family assistance through mother classes on improving body weight and nutritional status among Children Under Fives. *Dinamika Kesehatan: Jurnal Kebidanan dan Keperawatan*, 16(1), 11–19. <https://doi.org/10.33859/dksm.v16i1.1018>
- Ertiana, D., & Zain, S. B. (2023). Education and mothers' knowledge of nutrition are associated with the nutritional status of Children Under Fives. *Jurnal ILKES (Jurnal Ilmu Kesehatan)*, 14(1), 96–108. <https://doi.org/10.35966/ilkes.v14i1.279>
- Harliana, H., Yusron, R. D. R., & Machfud, I. (2022). Classification and monitoring of Children Under Five nutritional status through the application of a GIS-based Naïve Bayes classification method. *Jurnal Ilmiah Intech: Information Technology Journal of UMUS*, 4(2), 161–168. <https://doi.org/10.46772/intech.v4i02.869>
- Hidayat, T., Febriyanti, Z. R., Sukisno, Nugroho, A. H., Rizky, R., & Hakim, Z. (2026). Algoritma machine learning Random Forest dan Decision Tree dalam klasifikasi status gizi balita. *JUPI (Jurnal Ilmiah Penelitian dan Pembelajaran Informatika)*, 11(1), 47–58. <https://doi.org/10.29100/jupi.v11i1.7545>
- Hidayati, N. (2022). Hubungan keikutsertaan ibu dalam kelas ibu balita dengan status gizi balita di Desa Kedungsumber Kecamatan Temayang Kabupaten Bojonegoro (Bachelor's thesis, Institut Teknologi Sains dan Kesehatan Insan Cendekia Medika Jombang).
- Kotu, V., & Deshpande, B. (2022). *Data science: Concepts and practice* (3rd ed.). Morgan Kaufmann.
- Martínez-Plumed, F., Contreras-Ochando, L., Ferri, C., Hernández-Orallo, J., Kull, M., Lachiche, N., Ramírez-Quintana, M. J., & Flach, P. A. (2021). CRISP-DM twenty years later: From data mining processes to data science trajectories. *IEEE Transactions on Knowledge and Data Engineering*, 33(8), 3048–3061. <https://doi.org/10.1109/TKDE.2019.2962680>
- Ministry of Health of the Republic of Indonesia. (2023). *Indonesia Nutritional Status Survey (SSGI) 2023*. Ministry of Health of the Republic of Indonesia. <https://www.badankebijakan.kemkes.go.id/hasil-ski-2023/>
- Nurainun, Haerani, E., Syafria, F., & Oktavia, L. (2023). Application of the Naïve Bayes classifier algorithm for Children Under Five nutritional status classification using K-fold cross-validation. *Journal of Computer System and Informatics (JoSYC)*, 4(3), 578–586. <https://doi.org/10.47065/josyc.v4i3.3414>
- Ramadhani, & Ramadhanu. (2024). Machine learning methods for Children Under Five nutritional data classification using Naïve Bayes, K-nearest neighbors, and Decision Tree algorithms. *Jurnal SIMETRIS*, 15(1), 57–68. <https://jurnal.umk.ac.id/index.php/simet/article/view/10679>

World Health Organization. (2022). WHO child growth standards. <https://www.who.int/tools/child-growth-standards>