

Predicting Underweight Toddlers in Gorontalo Province Using Supervised Learning Algorithms

Muhajir Yunus^{*1}, St Suryah Indah Nurdin², Fitriah³

¹Digital Business, Universitas Muhammadiyah Gorontalo, Indonesia

²Midwifery, Universitas Muhammadiyah Gorontalo, Indonesia

³Informatics, Universitas Muhammadiyah Bengkulu, Indonesia

Email: ¹muhajiryunus@umgo.ac.id

Received: Sep 25, 2025; Revised: Oct 14, 2025; Accepted: Oct 16, 2025; Published: Oct 23, 2025

Abstract

Malnutrition in toddlers, notably underweight, remains a critical public health issue in Indonesia. According to the 2023 Indonesian Health Survey, the prevalence of underweight among toddlers has reached 15.9%. This condition has a significant impact on children's physical growth, cognitive development, and overall quality of life. This study aims to develop a predictive model for early detection of toddler nutritional status using three supervised machine learning algorithms: Decision Tree C4.5, K-Nearest Neighbor, and Naïve Bayes. The dataset consisted of 9,284 toddler records from Gorontalo Province, comprising eight attributes and one class label indicating nutritional status. Evaluation results showed that the Decision Tree C4.5 algorithm delivered the best performance with 98.56% accuracy. The K-Nearest Neighbor model achieved an accuracy of 97.99%, while the Naïve Bayes model obtained 96.96%. These findings demonstrate that machine learning can be an effective tool for identifying toddlers at risk of undernutrition early in their development. Beyond individual predictions, the proposed model represents a significant advancement in health informatics by providing a scalable decision-support system. This system can enhance the efficiency and precision of public health interventions, enabling faster, data-driven responses to combat malnutrition and improve child health outcomes across broader populations.

Keywords: Decision Tree C4.5, K-Nearest Neighbor, Naïve Bayes, Supervised Learning, Toddlers, Underweight.

This work is an open access article and licensed under a Creative Commons Attribution-Non Commercial 4.0 International License



1. INTRODUCTION

Malnutrition among toddlers in Indonesia remains a pressing public health concern, significantly manifested through acute and chronic forms, such as stunting, wasting, and underweight. As reported in the 2023 Indonesia Health Survey, the national prevalence of underweight in this demographic is 15.9% [1]. This situation is echoed in studies from various regions that highlight the high incidence of both acute and chronic malnutrition, emphasizing its detrimental impact on children's health outcomes. Underweight is a critical health condition that signifies a body weight below the expected range for an individual's height and age. This condition can arise from various factors, including inadequate nutritional intake, underlying medical issues, and diverse physiological or environmental elements. In the context of young children, underweight is a significant concern as it is often associated with acute malnutrition, particularly in the early stages of life when nutritional needs are paramount for growth and development [2]. Individuals with excessively low body weight face a multitude of health complications that encompass both physical and psychological dimensions.

The condition is typically categorized as underweight, defined by a body mass index (BMI) that falls below the accepted threshold for age and height [3]. Physical health implications include a significant increase in susceptibility to infectious diseases, as underweight individuals often demonstrate compromised immune responses. Chronic underweight conditions can lead to profound malnutrition,

which is associated with significant morbidity. For instance, in pediatric populations, underweight status is linked with a greater risk of stunting and wasting [4]. Children who are underweight not only experience impaired growth but are also more likely to suffer from gastrointestinal disturbances and respiratory issues due to insufficient caloric and nutrient intakes [5]. Furthermore, studies indicate that hospitalization of malnourished children can result in prolonged stays, increased hospital costs, and elevated morbidity and mortality risks exacerbated by both acute and chronic malnutrition [6].

The occurrence of underweight, particularly in children and young adults, can be influenced by a range of etiological factors, including genetic predispositions, hormonal imbalances, and insufficient nutritional intake. Each of these factors plays a significant role in the development and persistence of underweight status. Genetic predisposition is a notable factor contributing to underweight. Studies have shown that certain genetic variations can influence an individual's body mass index (BMI) and their susceptibility to weight gain or loss [7]. An example can be found in research linking genetic factors to conditions like anorexia nervosa, where individuals may have a familial tendency to be underweight, as seen in genome-wide association studies detecting negative correlations between body mass index and genetic markers associated with eating disorders [8]. Such evidence supports the notion that genetics can complicate nutritional management and contribute to sustained underweight.

Body weight is recognized as a key indicator of health status in Toddlers. It serves as a critical reflection of both nutritional adequacy and overall growth patterns. The assessment of body weight in this demographic not only provides insight into immediate health conditions but also lays the groundwork for potential future health outcomes. Nutritional status, as measured by body weight, is intrinsically linked to various health-related factors, including dietary intake, nutritional quality, and socioeconomic conditions. A study by Bekele and Fetene highlights the association between age and underweight, suggesting that children in older age groups within the under-five bracket are at a higher risk of experiencing underweight status due to compounded nutritional deficiencies over time [9].

The prevalence of undernutrition among toddlers old remains alarmingly high, particularly within rural communities and among populations with restricted access to adequate health services and nutrition. This issue is exacerbated by various socio-economic factors that contribute to malnutrition in these vulnerable populations. Recent studies indicate a significant prevalence of undernutrition indicators, such as stunting and wasting, particularly in South Asian countries like Bangladesh and India. For instance, a community health survey in Pakistan revealed that the prevalence of undernutrition among toddlers was as high as 52% when combining various forms of anthropometric failures [10]. This aligns with findings from another study across 26 countries in Sub-Saharan Africa, which demonstrated that the prevalence of undernutrition among toddlers was 34.59%, often compounded by coexisting forms [11]. Contributing factors to these high rates include parental education level and socio-economic status, which significantly affect nutritional outcomes [12]. Efforts to address this issue require more effective and proactive strategies, including the integration of modern technologies for the early detection of underweight risk in toddlers.

Machine learning (ML) has emerged as a transformative technology with the ability to revolutionize various domains, especially in healthcare settings where it is increasingly utilized for predictive analytics. The unique capacity of ML methodologies to process large and complex datasets enables a nuanced identification of patterns that surpass the limitations of traditional analytics. This ability is crucial in predicting malnutrition among toddlers, as it allows for timely interventions that can prevent more severe health complications in the future. Several studies elucidate the effectiveness of ML in the healthcare landscape, particularly focusing on malnutrition among children. For instance, Saleem et al. investigated the prediction of child malnutrition in Pakistan, highlighting the intricate interplay of sociodemographic factors that affect malnutrition rates. Their findings emphasize the potential of ML to identify critical risk factors, thus providing a foundation for targeted intervention

strategies aimed at reducing malnutrition prevalence among infants and toddlers [13]. Similarly, Rahman et al. examined the risk factors contributing to stunting and underweight among young children in Bangladesh, demonstrating how ML algorithms can successfully predict malnutrition outcomes based on these risk factors. Their research indicated significant improvements in predictive accuracy compared to earlier approaches [14].

The general applicability of ML in predicting health outcomes is reinforced by broader studies within the field. For example, Arueyingho et al. discuss various ML models and their applications in improving healthcare management, thereby emphasizing the critical role of these technologies in enhancing disease prediction and personalized medical care [15]. Specifically, the implementation of predictive models such as logistic regression and other sophisticated ML techniques illustrates the versatility of ML in capturing complex interactions within healthcare data, leading to enhanced predictive capabilities.

Based on the above explanation, nutritional problems among toddlers, particularly underweight, remain a serious concern with significant implications for physical growth, cognitive development, and long-term quality of life. Although various interventions have been implemented, early detection of underweight risk remains suboptimal, as the complex and non-linear risk patterns are difficult to identify using conventional methods applied by healthcare providers. This study focuses on developing an accurate prediction model employing supervised machine learning methods, namely Naïve Bayes, K-Nearest Neighbors, and the C4.5 Decision Tree algorithm.

This study draws upon several previous works relevant to both the topic and the methods employed. One such study investigated the prediction of undernutrition status using the Random Forest algorithm, achieving an accuracy of 76.2% [16]. Another study examined the determinants of malnutrition in women using machine learning approaches, in which the Random Forest algorithm was identified as the most effective model, yielding an accuracy of 81.4% [17]. Another study explored the prediction of determinants of child undernutrition in Ethiopia, where the eXtreme Gradient Boosting (XGBoost) algorithm demonstrated superior predictive performance [18]. A subsequent study on the prediction of low birth weight (LBW) in Ethiopia reported that the Random Forest algorithm was the most effective model, achieving an accuracy of 91.60% [19]. A previous study on the classification of stunting among children aged 0–60 months found that age was a significant determinant of stunting, with the model achieving an accuracy of 61.82% [20].

2. METHOD

The research methodology adopted in this study is presented in the schematic diagram below.

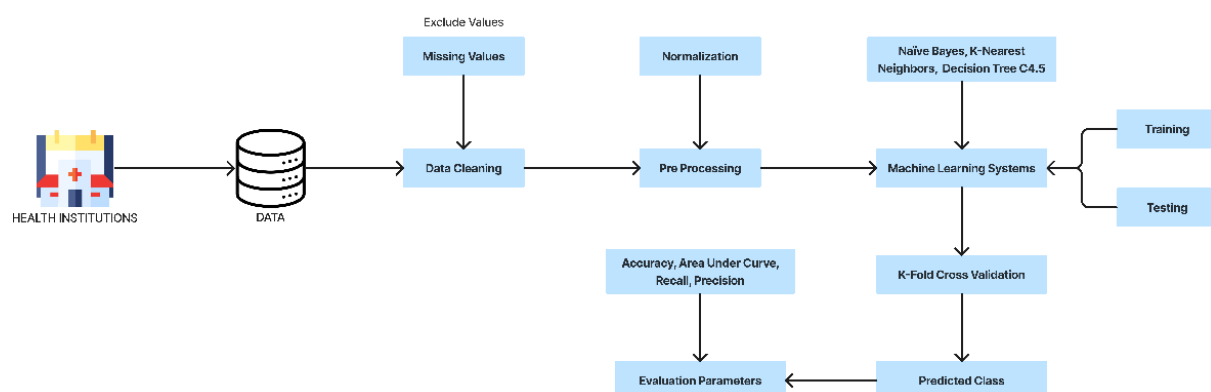


Figure 1. Research Method Flow

Figure 1 illustrates that the research process in this study is divided into several stages. First, the required data were collected, consisting of anthropometric measurements. Second, data preprocessing was conducted, which involved handling missing values, transforming attributes, and correcting data inconsistencies. Third, the dataset was divided into two subsets, training and testing data, with an 80:20 ratio. Fourth, predictive models were developed using the Naïve Bayes, K-Nearest Neighbors, and C4.5 Decision Tree algorithms implemented in Python. These models were trained on the training dataset to identify patterns and rules that would enable accurate prediction of underweight cases. Fifth, model validation was performed using the k-fold cross-validation method to ensure that the models trained on the training data also demonstrated strong performance on previously unseen data. Finally, model evaluation was carried out to assess the predictive performance of the algorithms using accuracy, precision, and F1-score metrics. These metrics provide a comprehensive assessment of the models' ability to predict underweight conditions among toddlers. The overall problem-solving approach is presented in the following figure.

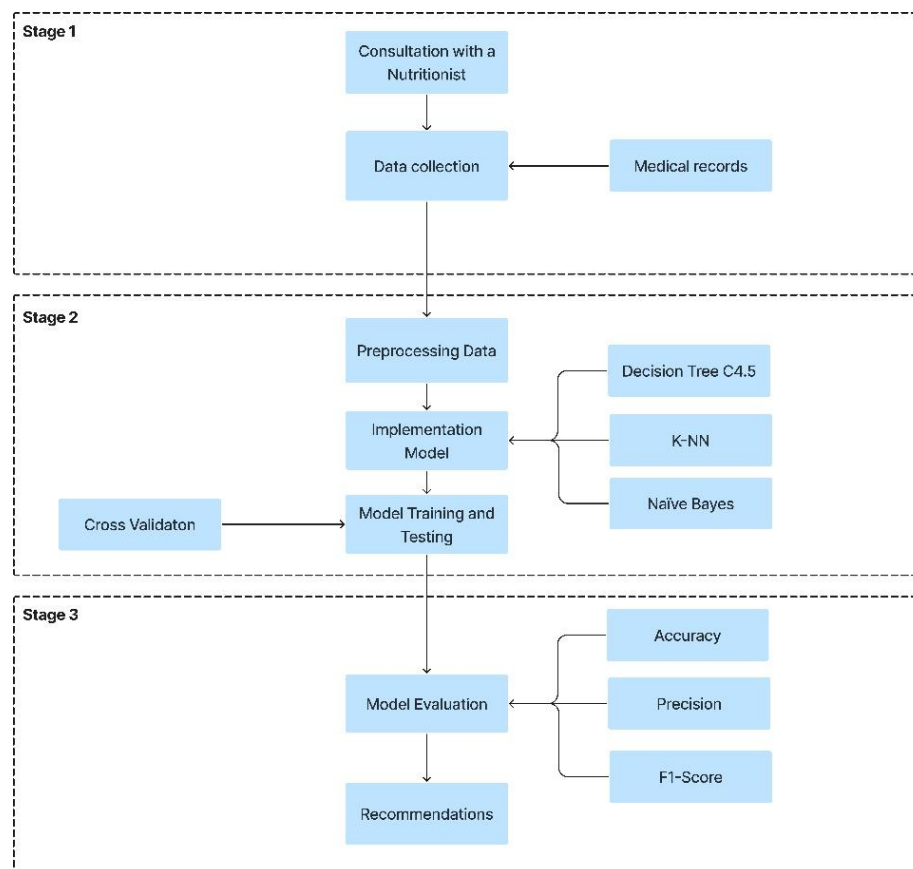


Figure 2. Troubleshooting Flow

Figure 2 illustrates the key stages involved in predicting underweight using machine learning methods, encompassing data collection, preprocessing, model implementation, evaluation, and policy recommendation. The first stage involves data collection from multiple sources, with collaboration with nutrition experts to obtain essential guidance and information before initiating the process. The second stage involves data preprocessing to ensure readiness for subsequent analysis, followed by model implementation, which includes training and testing with the application of cross-validation techniques. The third stage involves evaluating the model using predetermined performance metrics, such as accuracy, precision, and F1-score. Based on the evaluation results, policy recommendations can then be formulated to support decision-making and guide subsequent interventions.

2.1. Dataset

This study employed by-name by-address (BNBA) underweight data collected from Gorontalo Province in June 2025. The dataset comprised a total of 9,284 records, each containing eight attributes and one label. The attributes included: 1. Sex, 2. Birth Weight, 3. Birth Length, 4. Body Weight, 5. Body Height, 6. Measurement Method, 7. Weight-for-Age Z-Score (WAZ), 8. Weight-for-Age (W/A). The label variable was Predicted Body Weight, which categorized the records into two classes: Underweight (n = 4,642) and Normal (n = 4,642). A detailed overview of the data types for each attribute is provided in Table 1.

Tabel 1. Dataset Underweight

No	Sex	Birth Weight	Birth Length	Body Weight	Body Height	Measurement Method	Weight-for-Age Z-Score	Weight-for-Age (W/A)
1	1	3.00	48.0	13.3	104.5	1	-2.43	Kurang
2	1	3.00	47.0	13.7	100.2	1	-2.12	Kurang
3	1	2.50	45.0	12.4	98.0	1	-2.96	Kurang
....								
9282	1	2.80	48.0	16.5	99.0	1	-0.29	Normal
9283	0	3.10	48.0	16.5	99.0	1	-0.16	Normal
9284	0	3.05	46.0	13.0	96.0	1	-1.88	Normal

Information: Kurang = Underweight

2.2. Decision Tree C4.5

The C4.5 Decision Tree algorithm, developed by Ross Quinlan, represents a significant evolution from its predecessor, ID3, particularly in its ability to enhance the construction and efficacy of decision trees in various applications. One of the notable advancements of C4.5 is its capability to handle both numerical and categorical data. This flexibility allows the algorithm to effectively categorize attributes without requiring data preprocessing steps that might compromise the dataset's integrity [21]. Unlike ID3, which only used information gain as a splitting criterion, C4.5 employs the Gain Ratio criterion, which mitigates the bias toward attributes with many distinct values. This shift enhances the robustness of the decision tree by helping to prevent overfitting to noise present in the data [22]. Furthermore, C4.5's capability to manage missing values provides a critical practical advantage in real-world scenarios. The ability to handle such data ensures that valuable information is not discarded but rather used for informed decision-making, thus increasing the decision tree's overall accuracy. Calculating the value of C4.5 gain can use the following equation [23]:

$$Gain(S, A) = Entropy(S) - \sum_{i=1}^n \frac{S_i}{S} * Entropy(S_i) \quad (1)$$

Information

S : Case set (Underweight/Normal)

A : Attribution

S_i : the number of cases on partition to i

S : Number of cases in S

2.3. K-Nearest Neighbor

The K-Nearest Neighbor (KNN) algorithm is recognized for its straightforward implementation and effectiveness in both classification and regression tasks. Its nature as a lazy learning algorithm

makes it particularly attractive, as it does not require explicit training once the dataset is provided; instead, it classifies new data points based on the proximity of training data points in the feature space [24]. This characteristic allows KNN to adapt readily to a variety of data types, including non-linear and multi-class datasets, which is an essential consideration in machine learning applications [25][26]. Several studies confirm the efficacy of KNN in various domains. For instance, KNN has demonstrated significant performance improvements when combined with techniques such as Boosting, enhancing its accuracy by up to 75% compared to other algorithms like Least Squares Support Vector Machine (LS-SVM) [27]. Furthermore, research on blood donor data classification indicates that KNN can achieve accuracy rates exceeding 80% in several predictive scenarios, typically performing slightly lower than conventional decision tree algorithms [5]. In the context of medical applications, KNN's non-parametric nature empowers it to serve effectively in identifying health-related outcomes, such as predicting the recurrence of atrial fibrillation. In general, the Euclidean distance formula is as follows [28]:

$$d(x, y) = \sqrt{\sum_{i=1}^k (x_i - y_i)^2}$$

Information

x : New data

y : Training data

i : Data dimensions or attributes

k : Number of data dimensions

2.4. Naïve Bayes

The Naïve Bayes algorithm is widely recognized for its ability to produce probabilistic outputs, making it exceptionally valuable for decision-making processes that involve confidence levels. This algorithm operates on the fundamental assumption of feature independence, which facilitates rapid calculations and scalability. As a result, Naïve Bayes remains a relevant and frequently utilized machine learning algorithm, especially in contexts that require fast, computationally efficient predictions. The simplicity of Naïve Bayes allows for straightforward implementation, which is significant in both educational and professional settings, where speed and ease of interpretation are paramount [29]. Notably, the versatility of Naïve Bayes has led to its application across various domains, including text classification and health diagnostics. For instance, it has been effectively employed in identifying sentiment in customer reviews and classifying levels of consumer satisfaction based on various datasets [30]. Its probabilistic approach demonstrates efficacy in handling diverse datasets, showcasing its adaptability across different fields [31]. The equation of Bayes' theorem is based on the following formula [32]:

$$P(H|X) = \frac{P(X|H) * P(H)}{P(X)} \quad (3)$$

Information

X : Data with unknown classes

H : The hypothesis of data X is a specific class

$P(H|X)$: Probability of hypothesis H based on condition X (posteriori probability)

$P(H)$: Probability hypothesis H (prior probability)

$P(X|H)$: Probability X based on the condition on hypothesis H

$P(X)$: Probability X

2.5. Cross Validation

Cross-validation has emerged as a crucial technique in the realm of machine learning, serving as a robust method for model evaluation and ensuring more reliable performance metrics. This approach, which systematically partitions the dataset into multiple subsets for training and testing, enhances the accuracy, stability, and objectivity of model assessments compared to relying on a single train-test split. The application of cross-validation techniques, such as k-fold cross-validation, is widespread and supported by studies highlighting its benefits in various domains, including medical imaging and neuroscience [33]. One of the primary strengths of cross-validation is its potential to optimize the use of available data, particularly in instances where datasets are limited. The iterative process of training the model on different subsets enables a comprehensive evaluation, thereby reducing the risk of overfitting, where a model performs well on training data but poorly on unseen data [34]. Studies emphasize the significant contribution of cross-validation to model performance evaluation, which accounts for variance and ensures the generalizability of results across different data samples [35].

2.6. Confusion Matrix

The confusion matrix is a handy tool that helps us understand how well a classification model is really performing. Instead of just giving us an overall accuracy, it breaks down the predictions into categories like true positives (TP), false positives (FP), true negatives (TN), and false negatives (FN). This detailed view is especially useful when the data isn't evenly balanced say, if one class is much more common than another since it helps highlight where the model might be biased or making mistakes. By looking at these different types of errors, we can get a clearer picture of what's going wrong and how to improve the model. It's a straightforward yet powerful way to guide better decision-making in real-world situations [36][37]. Take healthcare, for example. Researchers often use confusion matrices to assess the accuracy of their disease detection systems. In one recent study on spotting tooth decay, the matrix was key to measuring how well the model performed, based on the counts of TP, FP, and FN. These numbers helped researchers assess effectiveness and identify areas for refinement, which ultimately influences patient care and treatment decisions [38].

2.7. Supervised Learning

Supervised learning algorithms present several notable advantages that underpin their widespread adoption across diverse domains. First, when sufficient labeled data are available, they achieve high levels of predictive accuracy, rendering them suitable for both classification and regression tasks [39]. Second, these algorithms demonstrate strong generalization capabilities to previously unseen data, particularly when combined with robust validation strategies such as cross-validation [40]. Third, the interpretability of many supervised models such as decision trees, support vector machines, and linear models facilitates the identification of feature importance and supports transparent decision-making processes [41]. Specific supervised approaches, including Naïve Bayes and logistic regression, enable probabilistic prediction and uncertainty estimation, a feature of particular importance in risk-sensitive fields such as healthcare [42]. Finally, the widespread availability of labeled datasets, together with well-established evaluation metrics such as accuracy, precision, recall, and F1-score, further reinforces the practical utility and relevance of supervised learning in both research and applied contexts [43].

3. RESULT

The evaluation results demonstrate the predictive performance of underweight classification using machine learning algorithms, as depicted in the performance graph below.

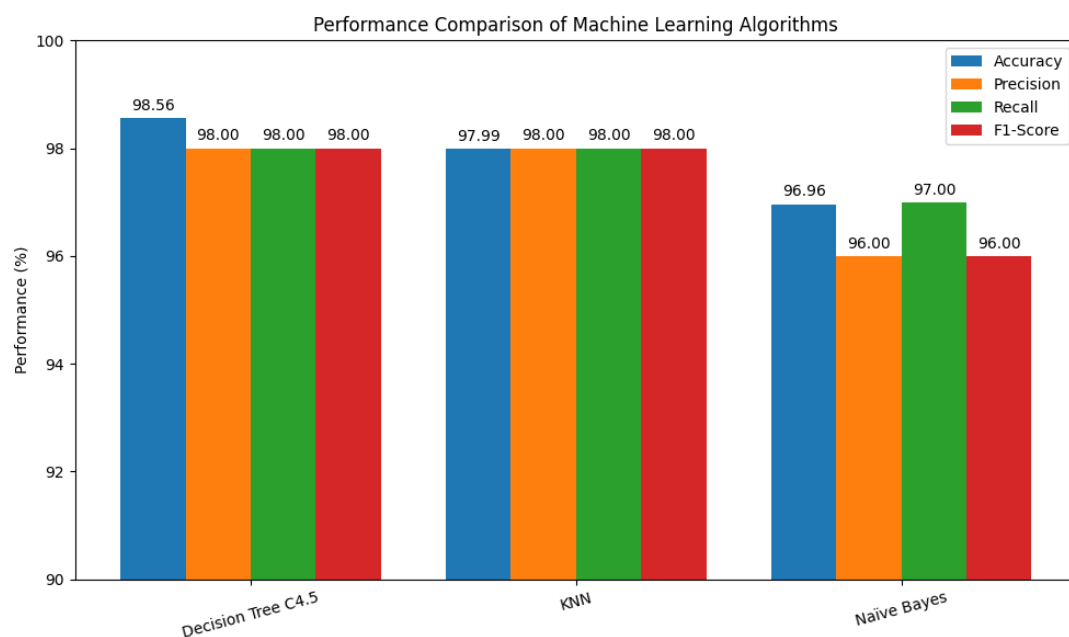


Figure 3. Machine Learning Model Performance

Figure 3 illustrates the comparative performance of the three machine learning algorithms in predicting underweight status. Among them, the C4.5 Decision Tree demonstrated the highest predictive capability, achieving an accuracy of 98.56%, followed by K-Nearest Neighbors at 97.99% and Naïve Bayes at 96.96%. Notably, all models yielded precision, recall, and F1-scores above 95%, indicating strong reliability and consistency in classifying the nutritional status of toddlers. The superior performance of the C4.5 Decision Tree may be attributed to its ability to effectively handle both categorical and numerical attributes while maintaining interpretability, making it particularly suitable for health-related datasets. These findings reinforce the potential of supervised learning algorithms as robust tools for early detection of underweight, thereby supporting timely interventions in child nutrition.

4. DISCUSSIONS

In the context of BNBA data containing mixed types, the superiority of the C4.5 algorithm lies in its gain ratio-based splitting mechanism, which reduces bias toward attributes with many categories, its ability to handle missing values, and post-training pruning that minimizes overfitting. These characteristics enable the decision tree to capture clinically meaningful numeric thresholds (e.g., weight-for-age or height-for-age cutoffs) while still utilizing categorical indicators without requiring excessive preprocessing. Compared to KNN, which is sensitive to scaling and the parameter, or Naïve Bayes, which assumes feature independence, C4.5 effectively models simple non-linear interactions through its hierarchical branches while maintaining rule-based interpretability. This explains the obtained accuracy of 98.56%, with precision and recall exceeding 95%, indicating strong generalization across both classes (underweight vs normal).

When compared with previous studies, the model's performance exceeds that of Random Forest, which achieves 76.2% accuracy in similar undernutrition classification tasks. This disparity can be attributed to (i) the better match between C4.5's architecture and the structure of mixed numeric-categorical attributes, (ii) balanced data and well-curated features (e.g., highly informative indicators like WAZ and W/A), and (iii) the robust validation strategy applied. In other words, "a more complex model does not necessarily mean a better one"; in this dataset, the alignment of C4.5's inductive bias with simple, clinically grounded decision patterns provides a tangible advantage over more opaque or hyperparameter-heavy approaches.

The implications for health informatics are substantial. Decision trees produce interpretable IF–THEN rules that can be easily audited, communicated, and translated into operational guidelines for local health centers (e.g., “If $WAZ < X$ and $BB < Y$ for age Z , prioritize home visits and nutritional supplementation”). Such transparency strengthens policy accountability by reducing the “black-box” nature of algorithmic decisions and supports scalability: the model can be embedded in regional nutritional surveillance systems, deployed on low-resource devices, and integrated with dashboards for resource prioritization (nutrition staff, supplementation, referrals). With its combination of interpretability and high accuracy, C4.5 stands out as a production-ready decision-support model for early detection of child malnutrition risks, accelerating data-driven responses, and contributing to improved public health outcomes.

4.1. Result Model Decision Tree C4.5

The detailed evaluation metrics of the predictive model developed using the C4.5 Decision Tree algorithm are presented in Figure 4, providing a comprehensive overview of its classification performance. The model achieved an accuracy of 98.56%, indicating a high level of correctness in predicting underweight status among toddlers. Furthermore, the precision, recall, and F1-score values all exceeded 95%, underscoring the model's robustness and reliability across different evaluation dimensions. High precision reflects the model's ability to minimize false positives, while high recall demonstrates its effectiveness in correctly identifying underweight cases. The F1-score, as a harmonic mean of precision and recall, further confirms the model's balanced performance. These results collectively affirm that the C4.5 Decision Tree algorithm is a highly effective approach for underweight prediction, offering both accuracy and interpretability, which are essential for practical application in public health decision-making.

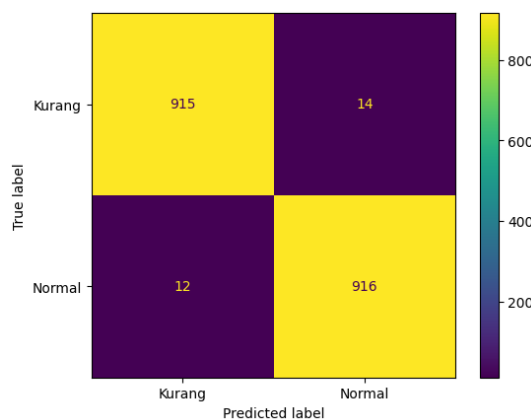


Figure 4. Evaluation Decision Tree C4.5 Model

The evaluation of the C4.5 Decision Tree model, as shown in Figure 4, highlights its strong performance in classifying toddlers as underweight. The model achieved an overall accuracy of 98.56%, supported by a confusion matrix that revealed a balanced distribution of correct and incorrect classifications. Specifically, 915 underweight cases were correctly classified as True Positives, while only 14 were misclassified as Normal (False Negatives). Similarly, 916 Normal cases were correctly identified as True Negatives, with only 12 instances misclassified as Underweight (False Positives). The precision, recall, and F1-scores for both classes were nearly identical, each exceeding 0.98, underscoring the robustness and reliability of the model. These results indicate that the Decision Tree not only minimizes the risk of misclassifying underweight children, which is critical in early intervention contexts, but also maintains strong accuracy in recognizing Normal cases. The balanced performance

across both classes demonstrates the model's suitability for real-world public health applications, where both sensitivity and specificity are essential for effective decision-making.

4.2. Result Model K-Nearest Neighbors

The evaluation metrics of the predictive model developed using the K-Nearest Neighbors algorithm are summarized in Figure 5. The KNN model achieved an accuracy of 97.99%, reflecting strong predictive capability in classifying underweight status among toddlers. Precision, recall, and F1-score values were consistently above 95%, indicating reliable performance across multiple evaluation criteria. The high precision highlights the model's ability to reduce false classifications of typical cases as underweight, while the substantial recall value demonstrates its effectiveness in capturing actual underweight cases. The F1-score further validates the balanced trade-off between precision and recall. These results confirm that KNN is a robust and effective algorithm for underweight prediction, although slightly outperformed by the C4.5 Decision Tree.

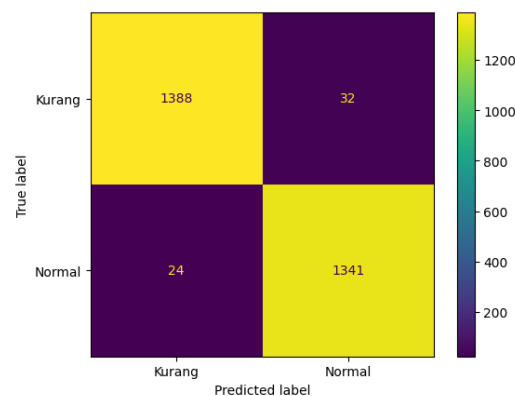


Figure 5. Evaluation K-NN Model

The evaluation of the K-Nearest Neighbors (KNN) model, as shown in Figure 5, demonstrates strong predictive performance with an overall accuracy of 97.99% on the test data. The confusion matrix reveals that the model correctly classified 1,388 underweight cases and 1,341 Normal cases. In comparison, 32 underweight cases were misclassified as Normal (False Negatives) and 24 Normal cases were misclassified as Underweight (False Positives). Although the number of misclassifications was slightly higher compared to the Decision Tree C4.5, the KNN model maintained reliable performance across both classes. The balance between correctly identified underweight and Normal cases indicates that KNN effectively captures similarity-based patterns within the dataset. These findings underscore the robustness of KNN as a predictive model for classifying underweight individuals. However, its sensitivity to parameter selection and data distribution may explain the marginally lower accuracy compared to the Decision Tree C4.5.

4.3. Result Model Naïve Bayes

The evaluation outcomes of the Naïve Bayes algorithm are presented in Figure 6. The model attained an accuracy of 96.96%, which, while slightly lower than the other two algorithms, still reflects high predictive reliability. Similar to KNN and C4.5, the precision, recall, and F1-score values for Naïve Bayes all exceeded 95%, demonstrating its competence in distinguishing between underweight and normal cases. The probabilistic nature of Naïve Bayes contributes to efficient computation and rapid prediction, making it particularly advantageous for scenarios requiring lightweight yet practical models. Although its overall performance was marginally lower compared to the Decision Tree and KNN

algorithms, Naïve Bayes remains a valuable approach, particularly as a baseline model for classification tasks involving underweight data.

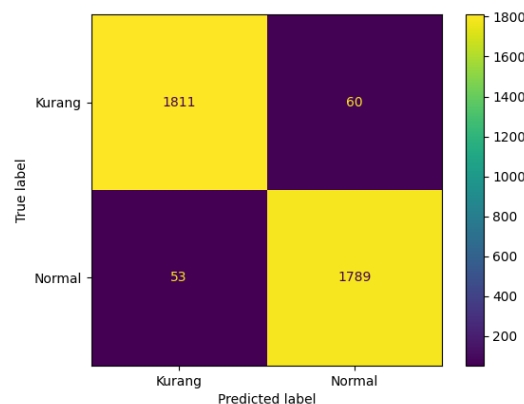


Figure 6. Evaluation Naïve Bayes Model

The evaluation of the Naïve Bayes model, as illustrated in Figure 6, demonstrates solid predictive performance, with an overall accuracy of 96.96%, as well as precision and F1-score. The confusion matrix shows that the model correctly classified 1,811 underweight cases (True Positives) and 1,789 Normal cases (True Negatives), while misclassifying only 60 Normal cases as underweight (False Positives) and 53 underweight cases as Normal (False Negatives). These results indicate that Naïve Bayes was able to maintain a balanced trade-off between precision and recall, minimizing both types of errors. Although its performance was slightly lower than that of the C4.5 Decision Tree and KNN models, the algorithm remains a reliable and computationally efficient approach for underweight prediction. Its probabilistic framework also provides interpretable outputs, which can be valuable for supporting decision-making in nutritional risk assessment.

5. CONCLUSION

The comparative analysis confirms that all three supervised learning algorithms Decision Tree 4.5, K-Nearest Neighbors, and Naïve Bayes are capable of accurately predicting underweight status in toddlers. Among them, the Decision Tree C4.5 proved to be the most effective, not only due to its superior predictive accuracy (98.56%) but also because of its ability to handle both categorical and numerical attributes while maintaining interpretability through rule-based outputs. These findings have significant implications for the field of health informatics. Integrating predictive models, such as C4.5, into public health information systems can enable the early detection of undernutrition risk, support faster interventions, and facilitate data-driven decision-making at the primary care level. Future research should explore ensemble methods, such as Random Forest to further enhance the robustness and accuracy of predictions. Expanding the dataset to include other regions of Indonesia is also recommended to test model generalizability and gain deeper insights into regional determinants of toddler nutrition.

CONFLICT OF INTEREST

The authors declare that there is no conflict of interest between the authors or with research object in this paper.

ACKNOWLEDGEMENT

The authors would like to thank the Kementerian Pendidikan Tinggi, Sains dan Teknologi Republik Indonesia (Kemendiktisaintek) for providing research funds through the Penelitian Dosen Pemula (PDP) scheme grant with contract number 013/PK-Pen/LPPM-UMGO/V/2025. The authors would also like to thank Lembaga Penelitian dan Pengabdian Kepada Masyarakat Universitas Muhammadiyah Gorontalo (LPPM UMGO) for providing support and administrative assistance during this research.

REFERENCES

- [1] M. R. Nanda, "Malnutrition and Pneumonia among Under-five Children in Sadewa Maternal and Child Hospital, Yogyakarta, Indonesia," *Biosci. Med. J. Biomed. Transl. Res.*, vol. 6, no. 18, hal. 2980–2984, Feb 2023, doi: 10.37275/bsm.v6i18.742.
- [2] S. Thurstans *dkk.*, "The relationship between wasting and stunting in young children: A systematic review," *Matern. Child Nutr.*, vol. 18, no. 1, Jan 2022, doi: 10.1111/mcn.13246.
- [3] G. A. Odei Obeng-Amoako *dkk.*, "Concurrently wasted and stunted children 6-59 months in Karamoja, Uganda: prevalence and case detection," *Matern. Child Nutr.*, vol. 16, no. 4, Okt 2020, doi: 10.1111/mcn.13000.
- [4] M. R. K. Chowdhury, M. S. Rahman, B. Billah, R. Kabir, N. K. P. Perera, dan M. Kader, "The prevalence and socio-demographic risk factors of coexistence of stunting, wasting, and underweight among children under five years in Bangladesh: a cross-sectional study," *BMC Nutr.*, vol. 8, no. 1, hal. 84, Agu 2022, doi: 10.1186/s40795-022-00584-x.
- [5] N. Kapci, M. Akcam, T. Koca, S. Dereci, dan M. Kapci, "The nutritional status of hospitalized children: Has this subject been overlooked?," *Turkish J. Gastroenterol.*, vol. 26, no. 4, Apr 2020, doi: 10.5152/tjg.2015.0013.
- [6] R. Amin, T. Akter, R. Ferdous, R. Akter, M. M. Shahriar, dan M. A. Islam, "Joint modelling of anthropometric child undernutrition indicators to identify their risk factors in Bangladesh: evidence from the Multiple Indicator Cluster Survey (MICS) 2019," *BMJ Open*, vol. 15, no. 7, hal. e092536, Jul 2025, doi: 10.1136/bmjopen-2024-092536.
- [7] W. Su *dkk.*, "Underweight Predicts Greater Risk of Cardiac Mortality Post Acute Myocardial Infarction," *Int. Heart J.*, vol. 61, no. 4, hal. 658–664, Jul 2020, doi: 10.1536/ihj.19-635.
- [8] T. Peters *dkk.*, "Reasons for admission and variance of body weight at referral in female inpatients with anorexia nervosa in Germany," *Child Adolesc. Psychiatry Ment. Health*, vol. 15, no. 1, hal. 78, Des 2021, doi: 10.1186/s13034-021-00427-w.
- [9] S. A. Bekele dan M. Z. Fetene, "Modeling non-Gaussian data analysis on determinants of underweight among under five children in rural Ethiopia: Ethiopian demographic and health survey 2016 evidences," *PLoS One*, vol. 16, no. 5, hal. e0251239, Mei 2021, doi: 10.1371/journal.pone.0251239.
- [10] O. S. Balogun, A. M. Asif, M. Akbar, C. Chesneau, dan F. Jamal, "Prevalence and Potential Determinants of Aggregate Anthropometric Failure among Pakistani Children: Findings from a Community Health Survey," *Children*, vol. 8, no. 11, hal. 1010, Nov 2021, doi: 10.3390/children8111010.
- [11] M. G. Worku, I. Mohanty, Z. Mengesha, dan T. Niyonsenga, "Risk Factors of Standalone and Coexisting Forms of Undernutrition Among Children in Sub-Saharan Africa: A Study Using Data from 26 Country-Based Demographic and Health Surveys," *Nutrients*, vol. 17, no. 2, hal. 252, Jan 2025, doi: 10.3390/nu17020252.
- [12] M. Z. Nahar dan M. S. Zahangir, "The role of parental education and occupation on undernutrition among children under five in Bangladesh: A rural-urban comparison," *PLoS One*, vol. 19, no. 8, hal. e0307257, Agu 2024, doi: 10.1371/journal.pone.0307257.
- [13] M. U. Saleem, M. U. Aslam, A. G. Khatir, dan Q. Jiang, "Exploring Risk Factors and Predictive Modeling of Child Malnutrition in Pakistan Using Machine Learning," *Am. J. Hum. Biol.*, vol. 37, no. 9, Sep 2025, doi: 10.1002/ajhb.70124.
- [14] S. M. J. Rahman *dkk.*, "Investigate the risk factors of stunting, wasting, and underweight among

- under-five Bangladeshi children and its prediction based on machine learning approach,” *PLoS One*, vol. 16, no. 6, hal. e0253172, Jun 2021, doi: 10.1371/journal.pone.0253172.
- [15] O. V Arueyingho, A. Al-Taie, dan C. McCallum, “Scoping review: Machine learning interventions in the management of healthcare systems,” *Digit. Heal.*, vol. 10, Jan 2024, doi: 10.1177/20552076221144095.
- [16] H. M. Fenta, T. Zewotir, dan E. K. Muluneh, “A machine learning classifier approach for identifying the determinants of under-five child undernutrition in Ethiopian administrative zones,” *BMC Med. Inform. Decis. Mak.*, vol. 21, hal. 2–12, Des 2021, doi: 10.1186/s12911-021-01652-1.
- [17] M. M. Islam dkk., “Application of machine learning based algorithm for prediction of malnutrition among women in Bangladesh,” *Int. J. Cogn. Comput. Eng.*, vol. 3, hal. 46–57, Jun 2022, doi: 10.1016/j.ijcce.2022.02.002.
- [18] F. H. Bitew, C. S. Sparks, dan S. H. Nyarko, “Machine learning algorithms for predicting undernutrition among under-five children in Ethiopia,” *Public Health Nutr.*, vol. 25, no. 2, hal. 269–280, Feb 2022, doi: 10.1017/S1368980021004262.
- [19] W. T. Bekele, “Machine learning algorithms for predicting low birth weight in Ethiopia,” *BMC Med. Inform. Decis. Mak.*, vol. 22, no. 1, hal. 2–16, Des 2022, doi: 10.1186/s12911-022-01981-9.
- [20] M. Yunus, M. K. Biddinika, dan A. Fadlil, “Classification of Stunting in Children Using the C4.5 Algorithm,” *J. Online Inform.*, vol. 8, no. 1, hal. 99–106, Jun 2023, doi: 10.15575/join.v8i1.1062.
- [21] R. Mondal dkk., “Decision tree learning in Neo4j on homogeneous and unconnected graph nodes from biological and clinical datasets,” *BMC Med. Inform. Decis. Mak.*, vol. 22, no. S6, hal. 347, Mar 2023, doi: 10.1186/s12911-023-02112-8.
- [22] B. Liu dan Z. Sun, “Global Economic Market Forecast and Decision System for IoT and Machine Learning,” *Mob. Inf. Syst.*, vol. 2022, hal. 1–12, Apr 2022, doi: 10.1155/2022/8344791.
- [23] S. S. Hajimirzaie dkk., “Predicting postpartum female sexual interest/arousal disorder via adiponectin and biopsychosocial factors: a cohort-based decision tree study,” *Sci. Rep.*, vol. 15, no. 1, hal. 27354, Jul 2025, doi: 10.1038/s41598-025-12025-3.
- [24] A. S. Syahab, A. Widiyanto, L. R. Anuggilarso, dan A. B. Wijaya, “Comparison of Machine Learning Algorithms for Classification of Ultraviolet Index,” *J. Teknol. Inf. dan Pendidik.*, vol. 15, no. 2, hal. 132–146, Mar 2023, doi: 10.24036/jtip.v15i2.692.
- [25] K. Wang dkk., “Improving Risk Identification of Adverse Outcomes in Chronic Heart Failure Using SMOTE+ENN and Machine Learning,” *Risk Manag. Healthc. Policy*, vol. Volume 14, hal. 2453–2463, Jun 2021, doi: 10.2147/RMHP.S310295.
- [26] M. David dan M. A. Firdaus, “Defect Rate Prediction in Manufacturing Process Using K-Nearest Neighbor Algorithm,” *Int. J. Informatics, Econ. Manag. Sci.*, vol. 3, no. 2, hal. 161, Agu 2024, doi: 10.52362/ijiem.v3i2.1599.
- [27] V. S. Kumar dan K. Vidhya, “Heart Plaque Detection with Improved Accuracy using K-Nearest Neighbors classifier Algorithm in comparison with Least Squares Support Vector Machine,” *CARDIOMETRY*, no. 25, hal. 1590–1594, Feb 2023, doi: 10.18137/cardiometry.2022.25.15901594.
- [28] H. A. A. Al-Khamees, N. S. Sani, A. S. Gifal, L. X. W. Liu, dan M. I. Esa, “A dynamic model using k-NN algorithm for predicting diabetes and breast cancer,” *Comput. Biol. Med.*, vol. 192, hal. 110276, Jun 2025, doi: 10.1016/j.combiomed.2025.110276.
- [29] P. S. Bakti dan E. Eliyani, “Application of J48 and Naïve Bayes Algorithms to Predict Ream Bookings at PT. Nippon Presisi Teknik,” *Eduvest - J. Univers. Stud.*, vol. 3, no. 6, hal. 1047–1060, Jun 2023, doi: 10.59188/eduvest.v3i6.834.
- [30] F. F. Hasibuan, M. H. Dar, dan G. J. Yanris, “Implementation of the Naïve Bayes Method to determine the Level of Consumer Satisfaction,” *Sinkron*, vol. 8, no. 2, hal. 1000–1011, Apr 2023, doi: 10.33395/sinkron.v8i2.12349.
- [31] C. B. I. Hulu dan H. T. Sihotang, “Plant Disease Diagnosis Expert System Cardamom (Ammomum Cardamomum L.) Using The Naive Bayes Method Web-Based,” *J. Mandiri IT*, vol. 10, no. 2, hal. 68–73, Jan 2022, doi: 10.35335/mandiri.v10i2.94.
- [32] Y. Zhao dkk., “Prediction of Drug-Induced Liver Injury: From Molecular Physicochemical

- Properties and Scaffold Architectures to Machine Learning Approaches,” *Chem. Biol. Drug Des.*, vol. 104, no. 2, Agu 2024, doi: 10.1111/cbdd.14607.
- [33] S. Gitto *dkk.*, “CT and MRI radiomics of bone and soft-tissue sarcomas: an updated systematic review of reproducibility and validation strategies,” *Insights Imaging*, vol. 15, no. 1, hal. 54, Feb 2024, doi: 10.1186/s13244-024-01614-x.
- [34] F. Schroeder, S. Fairclough, F. Dehais, dan M. Richins, “The impact of cross-validation choices on pBCI classification metrics: lessons for transparent reporting,” *Front. Neuroergonomics*, vol. 6, Jul 2025, doi: 10.3389/fnrgo.2025.1582724.
- [35] J. Cui *dkk.*, “Development and validation of an explainable machine learning model to predict Delphian lymph node metastasis in papillary thyroid cancer: a large cohort study,” *J. Cancer*, vol. 16, no. 6, hal. 2041–2061, Mar 2025, doi: 10.7150/jca.110141.
- [36] S. Sathyanarayanan, “Confusion Matrix-Based Performance Evaluation Metrics,” *Ajbr*, hal. 4023–4031, 2024, doi: 10.53555/ajbr.v27i4s.4345.
- [37] N. K. E. Sapitri, U. Sa’adah, dan N. Shofianah, “Knowledge Discovery From Confusion Matrix of Pruned CART in Imbalanced Microarray Data Ovarian Cancer Classification,” *Sci. J. Informatics*, vol. 11, no. 1, hal. 227–236, 2024, doi: 10.15294/sji.v11i1.50077.
- [38] E. Asci *dkk.*, “A Deep Learning Approach to Automatic Tooth Caries Segmentation in Panoramic Radiographs of Children in Primary Dentition, Mixed Dentition, and Permanent Dentition,” *Children*, vol. 11, no. 6, hal. 690, 2024, doi: 10.3390/children11060690.
- [39] J. E. van Engelen dan H. H. Hoos, “A survey on semi-supervised learning,” *Mach. Learn.*, vol. 109, no. 2, hal. 373–440, Feb 2020, doi: 10.1007/s10994-019-05855-6.
- [40] A. Tiwari, “Supervised learning: From theory to applications,” in *Artificial Intelligence and Machine Learning for EDGE Computing*, Elsevier, 2022, hal. 23–32. doi: 10.1016/B978-0-12-824054-0.00026-5.
- [41] J. Gui *dkk.*, “A Survey on Self-Supervised Learning: Algorithms, Applications, and Future Trends,” *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 46, no. 12, hal. 9052–9071, Des 2024, doi: 10.1109/TPAMI.2024.3415112.
- [42] J. M. Duarte dan L. Berton, “A review of semi-supervised learning for text classification,” *Artif. Intell. Rev.*, vol. 56, no. 9, hal. 9401–9469, Sep 2023, doi: 10.1007/s10462-023-10393-8.
- [43] X. Wang, X. Lin, dan X. Dang, “Supervised learning in spiking neural networks: A review of algorithms and evaluations,” *Neural Networks*, vol. 125, hal. 258–280, Mei 2020, doi: 10.1016/j.neunet.2020.02.011.