

Analysis of the Application of Machine Learning Algorithms for Classification of Toddler Nutritional Status Based on Antropometric Data

Esti Yogyianti¹, Yuni Yamasari², and Ervin Yohannes³

Departement of Informatics Engineering, Universitas Negeri Surabaya, Surabaya, Indonesia

Abstract

The rapid advancement of technology has required appropriate strategies to achieve accurate and optimal results. Among these, machine learning has become one of the most widely applied technologies across various domains, including healthcare, due to its ability to process large volumes of data and produce reliable predictions. One critical health problem that can benefit from these approaches is malnutrition among toddlers, which continues to pose challenges to growth, development, and long-term well-being. This analysis aims to identify the most effective and efficient algorithms for classifying the nutritional status of toddlers based on anthropometric data. The review is grounded in relevant journal articles aligned with the research topic, which serve as the primary sources for synthesis. The selected studies underwent four stages of identification, selection, evaluation, and analysis to ensure both credibility and reliability. The analysis focuses on three main aspects: dataset characteristics, algorithms applied, and outcomes reported. Based on algorithm usage, three implementation strategies were identified: single model, multi-model, and model combination. The overall findings reveal that studies utilizing datasets with fewer than 500 records can effectively apply algorithms such as Random Forest, Decision Tree, and Naïve Bayes Classifier, which consistently achieve accuracy rates above 90%. For datasets exceeding 10,000 records, the XGBoost algorithm is recommended due to its scalability and efficiency in handling large-scale data. For datasets ranging between 500 and 10,000 records, hybrid approaches such as the C4.5 algorithm combined with Particle Swarm Optimization are preferable, with previous studies demonstrating an accuracy of 94.49%. This review highlights that algorithm selection should be adjusted according to dataset size and clinical needs. The findings provide valuable insights to support researchers, practitioners, and policymakers in developing accurate and effective solutions for toddler nutrition assessment.

Paper History

Received June 10, 2025

Revised August 10, 2025

Accepted Sept. 20, 2025

Published October 10, 2025

Keywords

Status Classification;

Machine Learning;

Prediction;

Toddler Nutrition;

Anthropometric Data.

Author Email

estiyogyianti187@gmail.com

yuniyamasari@unesa.ac.id

ervinyohannes@unesa.ac.id

1. Introduction

The rapid growth of digital technology has greatly changed many aspects of life, particularly in the healthcare field [1]. New technologies, such as large-scale data analytics and cloud-based infrastructures, have reshaped the landscape of information management by enhancing decision-making processes and improving the precision of medical diagnoses and health predictions [2]. Machine learning technology has shown great potential in revolutionizing modern healthcare services through large-scale data analysis. This technology can provide deeper insights, enhance the system's ability to predict outcomes, and support healthcare professionals in delivering more personalized and patient-centered care. Therefore, machine learning has now become a crucial tool in addressing various issues in the medical field, including the classification and prediction of nutritional status in early childhood [1].

The Indonesian government has emphasized the importance of nutritional development through national health policies [1]. According to Article 141 of Law No. 36

of 2009 on Health, nutritional development aims to improve the nutritional status of individuals and communities by promoting better dietary patterns based on the 13 Guidelines for Balanced Nutrition (PUGS) and by strengthening family awareness regarding nutrition [3].

One of the major health concerns in Indonesia is malnutrition among toddlers [4], [5], [6]. As stated by the WHO in 2010, a stunting prevalence between 30–39% is considered high, while a rate of 40% or above falls into the very high category. Indonesia is listed as one of the five countries with the highest stunting prevalence in the world, at 30.8%, following India, China, Nigeria, and Pakistan [7]. Additionally, the WHO reported in 2020 that over 47 million children globally suffer from underweight, while 38 million are classified as overweight [8], [9], [10]. As of 2022, an estimated 148.1 million children under five were classified as stunted, 45 million were experiencing wasting, and 37 million were categorized as overweight, based on height-for-age indicators [8].

Nutritional status in toddlers is a critical factor that influences their physical and cognitive development, with

Corresponding author: Esti Yogyianti, estiyogyianti187@gmail.com, Departement of Informatics Engineering, Surabaya State University, Surabaya Indonesia

Digital Object Identifier (DOI): <https://doi.org/10.35882/ijeeemi.v7i4.110>

Copyright © 2025 by the authors. Published by Jurusan Teknik Elektromedik, Politeknik Kesehatan Kemenkes Surabaya Indonesia. This work is an open-access article and licensed under a Creative Commons Attribution-ShareAlike 4.0 International License (CC BY-SA 4.0).

long-term effects on health, productivity, and well-being [9], [10], [11], [12]. It significantly influences the development of an individual's future potential. Consequently, toddler nutrition must be prioritized and addressed promptly and accurately as a public health concern [13], [14], [15].

Nevertheless, conventional approaches to evaluating nutritional status still depend largely on the manual gathering of anthropometric measurements, including weight, height, and body mass index (BMI) values. These methods are often inefficient, time-consuming, and prone to human error, and results can vary depending on the skills of healthcare personnel [16]. Therefore, implementing machine learning-based technology is necessary to enable real-time data monitoring, enhance the accuracy of analyses, and support more efficient and responsive healthcare services.

Several previous studies have explored the application of intelligent computing techniques, particularly those utilizing machine learning models, to address the issue of nutritional status classification in toddlers. Various algorithms, such as the Naïve Bayes Classifier, Support Vector Machine (SVM), Random Forest, and Logistic Regression, have been employed to analyze anthropometric data and other relevant indicators for classification purposes. The majority of these studies have reported that machine learning algorithms can yield high classification accuracy and are reliable for predicting nutritional status. Although some studies have focused on specific datasets, the findings consistently demonstrate the potential of machine learning in nutrition-related data analysis.

Nevertheless, there are still research gaps that need to be addressed, such as model optimization, parameter selection, and the integration of more complex data types. These gaps provide opportunities for further studies to improve the reliability and scalability of data-driven classification systems. Based on the above background, this study aims to conduct a systematic literature review on the application of machine learning algorithms for the classification of toddler nutritional status. This review will identify the most effective algorithms, types of datasets used, and performance outcomes. The findings are expected to contribute to the development of more accurate and efficient nutritional classification systems, ultimately supporting sustainable efforts in malnutrition prevention.

II. Materials And Method

A. Machine Learning

Machine learning is a technology that is widely applied in various research domains, as it enables systems to derive knowledge from data and enhance their capabilities without relying on manually written instructions. By integrating human creativity, machine learning is applied to solve complex problems, particularly those related to data mining and statistical modeling. The main advantage

of this technology lies in its ability to enhance system quality, efficiency, and performance while producing highly accurate solutions. Moreover, the outcomes generated through machine learning processes can be presented and interpreted in a straightforward manner, making them easy for users to understand [17], [18], [19], [20]. The Naïve Bayes algorithm can be formulated using Eq. (1) as follows [21], [22]:

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

where $P(C_i|X)$ represents the posterior probability of class C_i given data X . K-Nearest Neighbor uses the Euclidean distance to measure the similarity between data points, as shown in Eq. (2) as follows [23], [24]:

$$dis(x_1, x_2) = \sqrt{\sum_{i=0}^n (x_{1i} - x_{2i})^2} \quad (2)$$

where x_{1i} and x_{2i} are the feature values of the two data points being compared. The Support Vector Machine (SVM) aims to find the optimal separating hyperplane, which can be expressed using Eq. (3) as follows [25]:

$$f(x) = w \cdot x + b \quad (3)$$

where w is the weight vector and b is the bias term. The C4.5 algorithm determines the best attribute for splitting using the Gain Ratio, as presented in Eq. (4) as follows [26]:

$$GainRatio(s, j) = \frac{Gain(s, j)}{SplitInfo(s, j)} \quad (4)$$

Logistic regression estimates probabilities using the sigmoid function, which is defined in Eq. (5) as follows [25]:

$$f(x) = \frac{1}{1 + e^{-x}} \quad (5)$$

Adaptive boosting (AdaBoost) minimizes classification errors by adjusting weights, as expressed in Eq. (6) as follows [25]:

$$error(M_i) = \sum_{j=1}^d w_j \times err(X_j) \quad (6)$$

where w_j is the weight of instance j and $err(X_j)$ represents its error value.

B. Dataset

The dataset used in this study was obtained from published scientific journal articles relevant to the topic of toddler nutritional status classification using machine learning algorithms. The data analyzed included information available within the selected articles, such as the type of data (primary or secondary), sample size, data sources, and the parameters applied in each study. Some datasets are publicly available online, while others are institution-specific and can only be accessed through official health organizations. Examples of dataset sources referenced in the reviewed studies include:

- Kaggle: <https://www.kaggle.com/datasets>
- WHO Child Growth Standards: <https://www.who.int/tools/child-growth-standards>
- Indonesian Basic Health Research: <https://www.litbang.kemkes.go.id/riskesdas>

C. Data Collection

Data were gathered through a literature review conducted using various well-established academic sources,

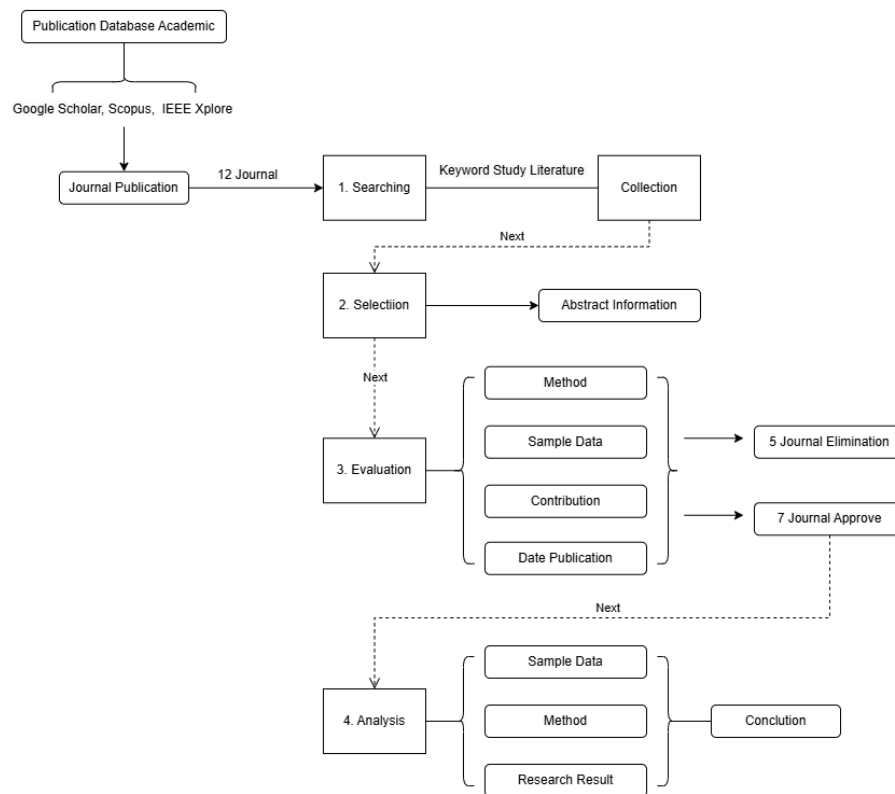


Fig. 1. Framework of the literature review process for data collection and analysis.

including Google Scholar, Scopus, and IEEE Xplore. The literature retrieval process utilized Boolean logic and specific keywords to identify relevant studies. From the initial search results, 12 journal articles were identified as potentially relevant based on their titles and research focus. A preliminary screening was then conducted by reviewing the abstracts of each article to assess their alignment with the objectives of this study. Articles that were not relevant, such as those that did not discuss toddler nutritional status or did not apply machine learning algorithms, were excluded. After this screening process, a total of seven (7) journal articles were selected for further analysis.

D. Data Processing

Data processing in this research was carried out through four main stages, as shown in Fig. 1. The four stages are searching, selection, evaluation, and analysis. The literature search stage used reputable academic databases such as Google Scholar, Scopus, and IEEE Xplore. The search focused on the topic of applying machine learning algorithms for the classification of nutritional status of toddlers with a publication year range between 2019 and 2025. Keywords used included Boolean combinations such as “machine learning” AND (“classification of nutritional status of toddlers” OR “nutrition prediction” OR “child nutrition”). Additional filters applied included full-text access as well as English-language articles. From this process, 12 articles were

collected, which were considered relevant to the research topic.

During the selection process, studies were screened based on predefined inclusion and exclusion criteria. Articles published from 2019 to 2025 and written in English were considered eligible for inclusion, specifically those that implemented machine learning algorithms for classifying the nutritional status of toddlers. Meanwhile, exclusion criteria included articles that only described nutritional measurements without classification, articles that did not present model evaluation metrics, and studies that did not have a clear methodology.

The articles that passed the selection were screened again using the PRISMA principles, with a focus on clarity of methods, adequate sample size, and completeness of reporting evaluation metrics. Of the 12 initial articles, five (5) were excluded because they did not meet evaluation standards, while seven (7) articles that met these criteria were analyzed further.

The analysis stage conducted on the seven (7) selected articles aimed to explore in more detail the machine learning algorithms applied, the sample data used, and the research results presented in each study. Each article was analyzed based on its methodological approach, classification accuracy, and contribution to the field of toddler nutritional status classification. The results of the analysis were then synthesized to draw conclusions that not only answer the research objectives but also

identify challenges, opportunities, and research gaps that can become the basis for future studies.

The data extraction process was carried out manually by the author independently, using a structured spreadsheet. The data collected includes the algorithm used, data source, number of samples, as well as evaluation metrics such as accuracy, precision, recall, and F1-score. To maintain consistency, any differences in interpretation between authors were resolved through discussion until consensus was reached. This approach was employed to ensure the validity and reliability of the data before it was synthesized.

III. Results

Here are seven articles that will be analyzed in more depth:

- 1) Machine Learning Method to Predict the Toddler's Nutritional Status [27].
- 2) Classification of Health and Nutritional Status of Toddlers Using the Naïve Bayes Classifier [28].
- 3) A Machine Learning Approach for Obesity Risk Prediction [25].
- 4) Prediction of Early Childhood Obesity with Machine Learning and Electronic Health Record Data [29].
- 5) Classification System of Toddler Nutrition Status using Naïve Bayes Classifier Based on Z-Score Value and Anthropometry Index [21].
- 6) Multiclass Classification of Toddler Nutritional Status using Support Vector Machine: A Case Study of Community Health Centers in Bangkalan, Indonesia [30].
- 7) Toddler Nutritional Status Classification Using C4.5 and Particle Swarm Optimization [26].

A. Data Used

Table 1. Dataset characteristics and parameters detail used in toddler nutritional status classification research.

No	Type of Research	Number of Samples	Type of Data	Data Source	Parameters Used
1	Quantitative	200 toddler data	Secondary	Kaggle [27]	Age, weight, height, BMI
2	Quantitative	21 toddler data	Primary	Anthropometric data from Posyandu [28]	Gender, age, weight, height
3	Quantitative	1100 toddler data	Primary	Data collected from various places and classes [25]	Parameters vary by the algorithm
4	Quantitative	850.520 patient data	Primary	Electronic Health Records (EHR) from Children's Hospital of Philadelphia [29]	Height, weight, BMI, age
5	Quantitative	225 toddler data	Primary	Anthropometric data [21]	Gender, age, weight, height
6	Quantitative	473 toddler data	Primary	Bangkalan District Health Office [30]	Weight, height, z-score
7	Quantitative	3961 toddler data	Primary	Nutritional status monitoring data from Riau Province [26]	Gender, age, height, and height measurement method (standing or lying down)

The analysis results in Table 1 show that all studies use a quantitative approach, because they rely on statistical processing of numerical data to produce measurable, objective, and accurate outcomes. There are three important characteristics of the dataset, namely the amount of data, data source, and type of parameters, which greatly influence the results of the machine learning algorithm in classifying the nutritional status of toddlers.

Research with small datasets (for example, only 21 toddler records) is suitable for using simple algorithms such as Naïve Bayes because it is easy to train and quite accurate, even though the data are limited, but the results can be less stable if tested on other data. Meanwhile, large datasets (such as 850,520 patient records from

hospitals) require algorithms such as XGBoost or Neural Network, which are strong and fast in processing large amounts of data. Apart from that, the data source also influences the quality and accuracy of the results. The use of primary data (collected directly from Posyandu or Puskesmas) tends to be more accurate and consistent because data collection is controlled by the researchers themselves. In contrast, research with secondary data from open sources such as Kaggle is more accessible but may contain errors or be unsuitable for local conditions. The type of parameters or features used determines the accuracy of the resulting classification. General parameters such as age, weight, and height are almost always used in every study. Some studies have added

Corresponding author: Esti Yogyanti, estiyogyanti187@gmail.com, Departement of Informatics Engineering, Surabaya State University, Surabaya Indonesia

Digital Object Identifier (DOI): <https://doi.org/10.35882/ijeemi.v7i4.110>

Copyright © 2025 by the authors. Published by Jurusan Teknik Elektromedik, Politeknik Kesehatan Kemenkes Surabaya Indonesia. This work is an open-access article and licensed under a Creative Commons Attribution-ShareAlike 4.0 International License (CC BY-SA 4.0).

features such as BMI and Z-score, which provide more specific information for nutritional status assessment, thereby increasing model accuracy. There are even studies that noted the method of height measurement (e.g., standing or lying), indicating attention to detail that could potentially influence classification results. Overall, the selection of an algorithm must consider the characteristics of the dataset.

B. Algorithm Used

Machine learning offers a variety of algorithms that can be applied to research related to the classification and prediction of toddler nutritional status. Some of the most commonly used algorithms include the Naïve Bayes Classifier, Random Forest, Support Vector Machine (SVM), Decision Tree, and K-Nearest Neighbor (KNN). Among these, the Naïve Bayes Classifier is frequently chosen due to its simplicity, computational efficiency, and ability to deliver good results even with relatively small datasets.

The analysis in Table 2 reveals three main approaches to the implementation of machine learning algorithms: single-model, multi-model, and hybrid-model approaches. The single-model approach is typically used in simple studies, where only one algorithm is implemented to classify or predict nutritional status. The multi-model approach involves comparing multiple algorithms within a single study to determine which one performs best based on evaluation metrics such as accuracy. Meanwhile, the hybrid-model approach combines classification algorithms with optimization techniques, such as Particle Swarm Optimization (PSO), to enhance model performance in terms of both accuracy and efficiency. These three approaches reflect the evolving and flexible nature of machine learning

applications, particularly in the context of toddler health data analysis, and demonstrate researchers' efforts to develop the most optimal and reliable classification models.

Table 2. Summary of machine learning algorithms applied in toddler nutritional status classification research.

No	Number of Algorithms	Detail Algoritma
1	6	Naïve Bayes Classifier, Linear Discriminant Analysis, Decision Tree, K-Nearest Neighbor, Random Forest, and Support Vector Machine
2	1	Naïve Bayes Classifier
3	9	k-Nearest Neighbor, Random Forest, Logistic Regression, Multilayer Perceptron, Support Vector Machine, Naïve Bayes Classifier, Adaptive Boosting, Decision Tree, and Gradient Boosting Classifier
4	7	Decision Tree, Gaussian Naive Bayes, Bernoulli Naïve Bayes, Logistic Regression, Neural Network, Support Vector Machine, dan XGBoost
5	1	Naïve Bayes Classifier
6	1	Support Vector Machine
7	1	Algoritma C4.5 dan Particle Swarm Optimization

C. Research Result

Table 3. Comparative evaluation of machine learning algorithms for toddler nutritional status classification across multiple performance metrics.

No	Evaluation Metrics	Results	Conclusion
1	Accuracy, sensitivity, specificity, Area Under the Curve (AUC), Cohen's Kappa Coefficient (CKC)	Random Forest : 97.37%, 95%, 98.81%, 99.9%, 96.09% Decision Tree : 97.37%, 95%, 98.81%, 96.67%, 96.09% Support Vector Machine : 96.09%, 85%, 96.7%, 95%, 88.05% Linear Discriminant Analysis : 92.11%, 88.53%, 96.96%, 99.31%, 92.11% Naïve Bayes Classifier : 89.47%, 86.26%, 96.2%, 98.01%, 84.46% k-Nearest Neighbor : 73.68%, 65.24%, 89.85%, 89.95%, 59.92%	The recommended algorithms based on high evaluation metric values are Random Forest and Decision Tree.
2	Accuracy	Naïve Bayes Classifier : 95.77%	The Naïve Bayes Classifier algorithm demonstrates good performance with an accuracy score close to 100%.

Corresponding author: Esti Yogyianti, estiyogyianti187@gmail.com, Departement of Informatics Engineering, Surabaya State University, Surabaya Indonesia

Digital Object Identifier (DOI): <https://doi.org/10.35882/ijeemi.v7i4.110>

Copyright © 2025 by the authors. Published by Jurusan Teknik Elektromedik, Politeknik Kesehatan Kemenkes Surabaya Indonesia. This work is an open-access article and licensed under a Creative Commons Attribution-ShareAlike 4.0 International License (CC BY-SA 4.0).

3	Accuracy, sensitivity, specificity, precision, recall, F_1 - score.	Logistic Regression : 97.09%, 100%, 100%, 97%, 97%, 97% Naïve Bayes Classifier : 86.04%, 100%, 100%, 86%, 86%, 86% k-Nearest Neighbor : 77.5%, 100%, 100%, 79%, 77%, 77% Random Forest : 72.3%, 94.11%, 100%, 57%, 72%, 63% Adaptive Boosting : 70.3%, 90.69%, 100%, 57%, 70%, 61% Decision Tree : 70.3%, 90.19%, 100%, 57%, 70%, 61% Multilayer Perceptron : 66.02%, 100%, 65.38%, 49%, 66%, 56% Support Vector Machine : 66.02%, 100%, nan, 53%, 66%, 56%. Gradient Boosting Classifier : 64.08%, 78.43%, 100%, 55%, 65%, 57%	Logistic Regression yielded the highest evaluation scores, making it the best-performing model in the study.
4	Accuracy, sensitivity, precision, F_1 - score.	XGBoost : 66.14%, 63.27%, 30.90%, 44.60% Decision Tree : 66.11%, 60.83%, 29.70%, 43.28% Support Vector Machine : 64.79%, 61.65%, 29.99%, 43.63% Logistic Regression : 64.81%, 61.68%, 30%, 43.65% Neural Network : 63.67%, 60.33%, 29.61%, 43.05% Gaussian Naïve Bayes : 63.23%, 59.76%, 29%, 45.59% Bernoulli Naïve Bayes : 61.76%, 58.11%, 28.06%, 41.47%	The XGBoost algorithm showed the highest values among all metrics evaluated in the study.
5	Accuracy	Naïve Bayes Classifier from 175 toddler data (Accuracy: 100%) – 44.58% classified as undernourished – 36.58% as normal – 18.86% as overweight	The Naïve Bayes Classifier algorithm is highly recommended, with the Z-score method achieving 100% accuracy.
6	Accuracy	Support Vector Machine : 76%	The SVM algorithm was able to classify the data with 76% accuracy.
7	Accuracy	Algoritma C4.5 dan Particle Swarm Optimization (PSO) : 94.49%	The combination of C4.5 and PSO optimization demonstrated strong performance with an accuracy score of 94.49%.

In Table 3, the most frequently used algorithm is the Naïve Bayes Classifier, which appears in four of the seven studies. Several studies report that this algorithm is able to achieve up to 100% accuracy on certain datasets, showing high performance, especially with data that have a limited number of samples. In addition, the Random Forest, Decision Tree, Logistic Regression, and C4.5 + PSO algorithms also show superior performance with accuracy above 94%, as well as relatively high sensitivity and precision. Fig. 2 presents a heatmap depicting the relative strength of each metric using a color gradient, with dark blue representing the highest values. Based on the heatmap, it appears that algorithms such as Random Forest and Decision Tree demonstrate superior

performance in terms of accuracy and sensitivity metrics. Meanwhile, Logistic Regression demonstrates superiority across almost all evaluation metrics. Conversely, algorithms such as Gradient Boosting, XGBoost, and Neural Network exhibit inconsistent performance, particularly in terms of the F1-score metric, indicating an imbalance between precision and recall. It should be noted that not all algorithms were reported with all five metrics simultaneously in every study. Some values in the heatmap are empty (zero values), not because of low performance, but because the metrics were not reported in the source article. This is one of the limitations of the synthesis process.

Corresponding author: Esti Yogyanti, estiyogyanti187@gmail.com, Departement of Informatics Engineering, Surabaya State University, Surabaya Indonesia

Digital Object Identifier (DOI): <https://doi.org/10.35882/ijeemi.v7i4.110>

Copyright © 2025 by the authors. Published by Jurusan Teknik Elektromedik, Politeknik Kesehatan Kemenkes Surabaya Indonesia. This work is an open-access article and licensed under a Creative Commons Attribution-ShareAlike 4.0 International License (CC BY-SA 4.0).

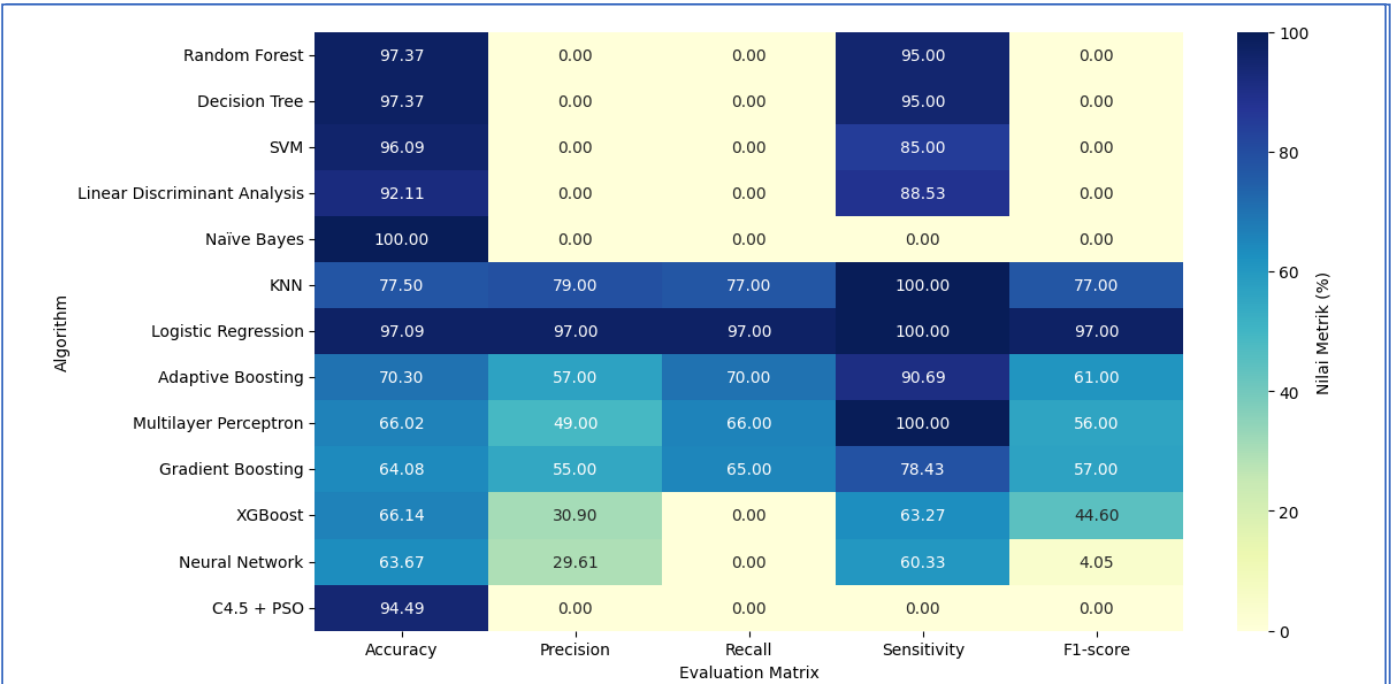


Fig. 2. Comparative performance heatmap of machine learning algorithms using accuracy, precision, recall, sensitivity, and f1-score.

A bar chart presented in Fig. 3 visualizes the primary performance metrics, including accuracy, precision, recall, sensitivity, and F1-score, for each algorithm. This visualization shows that Logistic Regression has a very stable performance across all metrics (around 97–100% each). In contrast, algorithms such as Neural Network and XGBoost show value inequality, with low precision and F1-score despite fairly high sensitivity. This indicates that although the model can detect minority classes, the prediction results are less precise (low precision).

By considering all the visual and numerical results presented, it can be concluded that the Random Forest, Decision Tree, and Logistic Regression algorithms are the most recommended models for classifying the nutritional status of toddlers, especially in the context of practical implementation in the health sector. Meanwhile, although Naïve Bayes shows good results in some studies, its performance is highly dependent on the characteristics of the dataset. Furthermore, complex algorithms such as XGBoost and Neural Network, despite their potential, require more attention to interpretability and metric stability, especially for clinical applications.

IV. Discussion

This review synthesizes findings from seven studies that applied machine learning (ML) algorithms to classify the nutritional status of toddlers based on anthropometric data. All evaluation metrics (accuracy, precision, recall/sensitivity, specificity, F1-score, AUC, and Cohen’s Kappa) were extracted verbatim from the original articles to maintain authenticity and ensure cross-study consistency. However, confidence intervals, standard

errors, and p-values were generally absent, therefore analyses are descriptive rather than meta-analytic.

The interpretation of results shows that model performance is strongly influenced by dataset size and metric completeness. In small- to medium-sized datasets, Random Forest and Decision Tree achieved 97.37% accuracy with 95% sensitivity and 98.81% specificity (n=200) [27]. In contrast, Logistic Regression demonstrated the highest stability in larger primary datasets (n≈1,100), achieving 97.09% accuracy, 100% sensitivity, and 100% specificity, thus outperforming Random Forest (72.3% accuracy) and Adaptive Boosting (70.3%) [25]. Naïve Bayes appeared in four out of seven studies, showing accuracy between 86.04% and 100%, including a case with 175 toddlers where it reached 100% accuracy using the Z-score method [21], [28]. Conversely, in very-large-scale electronic health record (EHR) data (n=850,520), XGBoost achieved only 66.14% accuracy with 30.90% precision and 44.60% F1-score, indicating challenges with class imbalance and noise. Hybrid approaches, such as C4.5 combined with Particle Swarm Optimization (PSO), reached 94.49% accuracy on 3,961 samples [26], highlighting the value of optimization methods for medium- to large-sized curated datasets.

When compared with similar studies, these findings are consistent. Logistic Regression achieved 96% accuracy in the classification of child nutrition in Malaysia as reported by Gustriansyah et al. [27], in line with the high stability found in these observations. Random Forest achieved 98% accuracy in monitoring child health with a dataset of more than 10,000 samples according to Pang et al. [29], confirming its scalability. This comparison supports that dataset size greatly influences model

effectiveness: simple models such as Naïve Bayes or Decision Tree excel on small datasets, whereas complex models such as XGBoost are better suited for large-scale data, as long as data imbalance and interpretability are taken into account.

Several limitations must be acknowledged. Some studies relied on very small datasets (as few as 21 samples) [28], which increase the risk of overfitting and reduce generalizability. Many studies also failed to report complete evaluation metrics, such as AUC or F1-score, thereby limiting cross-study comparability. In addition, only seven articles were included in this review, and most did not disclose preprocessing methods or statistical uncertainty measures. Potential publication bias and the lack of access to raw data further limit the validity of this synthesis.

The implications of this review are threefold. First, interpretable models such as Decision Tree and Naïve Bayes appear to be most suitable for community healthcare centers such as Posyandu, where transparency and accountability are crucial. Previous studies have already demonstrated accuracies between 90% and 100% in these contexts [21], [28], [30]. Second, larger healthcare institutions managing extensive EHRs may employ complex models such as XGBoost or Neural Network, but these should be accompanied by class imbalance handling and interpretability techniques such as SHAP or LIME to maintain trust. Third, future research should encourage open data policies, the complete and standardized reporting of evaluation metrics, and the exploration of hybrid approaches like C4.5 + PSO (which reached 94.49% accuracy) to improve both predictive performance and clinical applicability.

V. Conclusion

The findings of this study indicate a significant variation in how machine learning techniques are utilized to classify toddler nutritional status, particularly in terms of dataset sizes, selected parameters, and algorithm types. Sample sizes in the studies reviewed ranged from 10 to more than 850,000 records, with common parameters including gender, age, weight, height, body mass index (BMI), and Z-score values. The algorithms used included Naïve Bayes, Random Forest, Decision Tree, Logistic Regression, and SVM, as well as combinations of models such as C4.5 with PSO. Algorithm implementation was divided into three main approaches: single-model, multi-model comparison, and hybrid-model approaches. In general, algorithms such as Random Forest, Decision Tree, and Logistic Regression demonstrated high accuracy ($\geq 97\%$), while Naïve Bayes is able to achieve 100% accuracy on small and structured datasets. Combination models such as C4.5 + PSO have also proven effective with an accuracy of 94.49%. In contrast, algorithms such as K-Nearest Neighbor (KNN) and Gradient Boosting demonstrated lower performance in some studies. These findings suggest that algorithm

selection needs to be adapted to the characteristics of the dataset. For small datasets (<500 samples), algorithms such as Random Forest, Decision Tree, and Naïve Bayes are recommended. For medium datasets (500–10,000 samples), combination models such as C4.5 + PSO are more suitable. Meanwhile, for large datasets (>10,000 samples), XGBoost is recommended because of its ability to handle large-scale computational processes. To increase accuracy and generalization, it is recommended that future studies use primary data to ensure data quality and contextual relevance. With proper implementation, machine learning algorithms have the potential to become effective tools in decision-making for preventing malnutrition and improving the nutritional status of children.

References

- [1] H. A. Ganatra, "Machine Learning in Pediatric Healthcare: Current Trends, Challenges, and Future Directions," *J Clin Med*, 2025, doi: 10.3390/jcm14030807.
- [2] P. Apell and H. Eriksson, "Artificial intelligence (AI) healthcare technology innovations: the current state and challenges from a life science industry perspective," *Technol Anal Strateg Manag*, vol. 35, no. 2, pp. 179–193, 2023, doi: 10.1080/09537325.2021.1971188.
- [3] A. Asyraffi, O. Okfalisa, F. Insani, S. Agustian, and R. M. Candra, "Toddler nutritional status identification: Support vector machine (SVM) algorithm adoption," *Science, Technology and Communication Journal*, vol. 5, no. 2, p. 36, Feb. 2025, doi: 10.59190/stc.v5i2.282.
- [4] M. S. Khamdani, U. Habibah, N. Chamidah, and L. Muniroh, "Modeling the risk of Nutritional Status Wasting with Nonparametric Ordinal Logistic Approach Based on Spline," *International Journal of Multidisciplinary Research and Growth Evaluation*, vol. 6, 2025, [Online]. Available: www.allmultidisciplinaryjournal.com
- [5] M. Ula, A. Faridhatul Ulva, Mauliza, I. Sahputra, and Ridwan, "Implementation of Machine Learning in Determining Nutritional Status using the Complete Linkage Agglomerative Hierarchical Clustering Method," *Jurnal Mantik*, vol. 5, no. 3, 2021.
- [6] M. Ula, A. F. Ulva, M. Mauliza, M. A. Ali, and Y. R. Said, "Application Of Machine Learning in Determining the Classification of Children's Nutrition with Decision Tree," *Jurnal Teknik Informatika (Jutif)*, vol. 3, no. 5, pp. 1457–1465, Sep. 2022, doi: 10.20884/1.jutif.2022.3.5.599.
- [7] N. Siregar, R. Ratnawati, and atul Amini, "Nutritional status and toddler development: a relationship study," *Junal Kesehatan Ibu dan Anak*, vol. 14, no. 1, pp. 57–62, Jul. 2020, doi: 10.29238/kia.v14i1.611.
- [8] H. A. Santoso, N. S. Dewi, S. C. Haw, A. D. Pambudi, and S. A. Wulandari, "Enhancing

Corresponding author: Esti Yogyanti, estiyogyanti187@gmail.com, Departement of Informatics Engineering, Surabaya State University, Surabaya Indonesia

Digital Object Identifier (DOI): <https://doi.org/10.35882/ijeemi.v7i4.110>

Copyright © 2025 by the authors. Published by Jurusan Teknik Elektromedik, Politeknik Kesehatan Kemenkes Surabaya Indonesia. This work is an open-access article and licensed under a Creative Commons Attribution-ShareAlike 4.0 International License (CC BY-SA 4.0).

- nutritional status prediction through attention-based deep learning and explainable AI," *Intell Based Med*, vol. 11, 2025, doi: 10.1016/j.ibmed.2025.100255.
- [9] L. Cheikh Ismail *et al.*, "Nutritional status and adequacy of feeding Practices in Infants and Toddlers 0-23.9 months living in the United Arab Emirates (UAE): findings from the feeding Infants and Toddlers Study (FITS) 2020," *BMC Public Health*, no. 1, 2022, doi: 10.1186/s12889-022-12616-z.
- [10] B. Anita Ratna Etnis, W. Maria Prasetyo Hutomo, H. Maikel Su, I. Rahman, and E. Kolong, "Socioeconomic Factors and Its Correlation with Nutritional Status in Toddlers: A Study in Papua," *Journal of Health Sciences and Epidemiology*, vol. 2, no. 2, Aug. 2024, doi: <https://doi.org/10.62404/jhse.v2i2.49>.
- [11] R. Gustriansyah, N. Suhandi, S. Puspasari, A. Sanmorino, and D. Sartika, "Toddlers' Nutritional Status Prediction Using the Multinomial Logistics Regression Method," *Journal of Computer Networks, Architecture and High Performance Computing*, vol. 6, no. 1, Jan. 2024, doi: 10.47709/cnahpc.v6i1.3372.
- [12] S. Syafriani, H. Haryati, R. P. Fitri, and A. Afiah, "The Influence of Socio-Economic and Environmental Factors on the Nutritional Status of Toddlers in Urban and Rural Areas in Bangkinang Regency, Riau, Indonesia," *Journal of Health and Nutrition Research*, vol. 4, no. 1, Apr. 2025, doi: 10.56303/jhnrsearch.v4i1.348.
- [13] A. Ridwan and T. N. Sari, "The comparison of accuracy between naïve bayes classifier and c4.5 algorithm in classifying toddler nutrition status based on anthropometry index," *J Phys Conf Ser*, 2021, doi: 10.1088/1742-6596/1764/1/012047.
- [14] M. Fadil, A. Islamiyati, and S. A. Thamrin, "Classification of Nutritional Status in Support Vector Machine Method," *Communications in Mathematical Biology and Neuroscience*, 2025, doi: 10.28919/cmbn/9126.
- [15] M. Alif Alfiansyah, D. Novalendry, and Y. Hendriyani, "Clustering of Toddler Data at Puskesmas Nanggalo using the K-Means Algorithm," *urnal Ilmiah Sains dan Teknologi Scientica*, 2025.
- [16] M. Senbekov *et al.*, "The recent progress and applications of digital technologies in healthcare: A review," 2020, *Hindawi Limited*. doi: 10.1155/2020/8830200.
- [17] M. Noroozi *et al.*, "Machine and deep learning algorithms for classifying different types of dementia: A literature review," *Applied Neuropsychology: Adult*, 2024, doi: 10.1080/23279095.2024.2382823.
- [18] E. Miranda, M. Aryuni, A. Y. Zakiyyah, Y. E. Kurniawati, A. V. D. Sano, and M. Kumbangsila, "An early prediction model for toddler nutrition based on machine learning from imbalanced data," in *Procedia Computer Science*, Elsevier B.V., 2024, pp. 263–271. doi: 10.1016/j.procs.2024.10.251.
- [19] M. A. Talib, S. Majzoub, Q. Nasir, and D. Jamal, "A systematic literature review on hardware implementation of artificial intelligence algorithms," *Journal of Supercomputing*, 2021, doi: 10.1007/s11227-020-03325-8.
- [20] P. Hosseinzadeh Kasani, J. E. Lee, C. Park, C. H. Yun, J. W. Jang, and S. A. Lee, "Evaluation of nutritional status and clinical depression classification using an explainable machine learning method," *Front Nutr*, vol. 10, May 2023, doi: 10.3389/fnut.2023.1165854.
- [21] T. E. Putri, R. T. Subagio, Kusnadi, and P. Sobiki, "Classification System of Toddler Nutrition Status using Naïve Bayes Classifier Based on Z- Score Value and Anthropometry Index," *J Phys Conf Ser*, 2020, doi: 10.1088/1742-6596/1641/1/012005.
- [22] B. Phatcharathada and P. Srisuradetchai, "Randomized Feature and Bootstrapped Naive Bayes Classification," *Applied System Innovation*, vol. 8, no. 4, p. 94, Jul. 2025, doi: 10.3390/asi8040094.
- [23] J. B. Chandra and D. Nasien, "Application Of Machine Learning K-Nearest Neighbour Algorithm To Predict Diabetes," *International Journal of Electrical, Energy and Power System Engineering*, vol. 6, no. 2, Jun. 2023, [Online]. Available: <http://www.ijeepse.ejournal.unri.ac.id>
- [24] C. Wang, F. You, and Y. Wang, "Text intelligent classification model based on improved KNN algorithm in library service transformation," *Systems and Soft Computing*, vol. 7, 2025, doi: 10.1016/j.sasc.2025.200313.
- [25] F. Ferdowsy, K. S. A. Rahi, M. I. Jabiullah, and M. T. Habib, "A machine learning approach for obesity risk prediction," *Current Research in Behavioral Sciences*, vol. 2, Nov. 2021, doi: 10.1016/j.crbeha.2021.100053.
- [26] A. Nazir, A. Akhyar, Y. Yusra, and E. Budianita, "Toddler Nutritional Status Classification Using C4.5 and Particle Swarm Optimization," *Scientific Journal of Informatics*, vol. 9, no. 1, pp. 32–41, May 2022, doi: 10.15294/sji.v9i1.33158.
- [27] R. Gustriansyah, N. Suhandi, S. Puspasari, and A. Sanmorino, "Machine Learning Method to Predict the Toddlers' Nutritional Status," *JURNAL INFOTEL*, vol. 16, no. 1, Feb. 2024, doi: 10.20895/infotel.v15i4.988.
- [28] Y. Saputra, N. Khoirunnisa, and S. A. Haqq, "Classification of Health and Nutritional Status of Toddlers Using the Naïve Bayes Classifier," *CoreID Jurnal*, vol. 1, no. 2, pp. 49–57, Jul. 2023, doi: 10.60005/coreid.v1i1.8.
- [29] X. Pang, C. B. Forrest, F. Lê-Scherban, and A. J. Masino, "Prediction of early childhood obesity with

Corresponding author: Esti Yogyanti, estiyogyanti187@gmail.com, Departement of Informatics Engineering, Surabaya State University, Surabaya Indonesia

Digital Object Identifier (DOI): <https://doi.org/10.35882/ijeemi.v7i4.110>

Copyright © 2025 by the authors. Published by Jurusan Teknik Elektromedik, Politeknik Kesehatan Kemenkes Surabaya Indonesia. This work is an open-access article and licensed under a Creative Commons Attribution-ShareAlike 4.0 International License (CC BY-SA 4.0).

machine learning and electronic health record data," *Int J Med Inform*, vol. 150, Jun. 2021, doi: 10.1016/j.ijmedinf.2021.104454.

- [30] M. A. Syakur *et al.*, "Multiclass classification of toddler nutritional status using support vector machine: A case study of community health centers in Bangkalan, Indonesia," *BIO Web Conf*, 2024, doi: 10.1051/bioconf/202414601082.

Author Biography



Esti Yogyanti is a graduate of the Informatics Engineering Department from Universitas Negeri Surabaya (UNESA), where she received her bachelor's degree in 2024. Currently, she is pursuing a master's degree in Informatics at the same institution. Her research interests focus on the application of artificial intelligence and machine learning in the health sector, especially in nutritional analysis, early disease detection, and the development of digital health solutions. She is passionate about exploring innovative ways in which technology can be used to improve healthcare quality and accessibility. She aims to contribute to academic research as well as the practical implementation of machine learning in healthcare, particularly in addressing challenges related to child nutrition in Indonesia.



Yuni Yamasari is a lecturer at Universitas Negeri Surabaya (UNESA). She earned her bachelor's degree in Informatics Engineering from Institut Teknologi Sepuluh Nopember (ITS) in 2001,

continued her studies with a master's degree in the same field at ITS in 2008, and completed her doctorate in Electrical Engineering in 2020 at ITS. Her academic expertise lies in artificial intelligence, computational systems, and applied informatics. As a lecturer, she actively teaches, supervises student research, and engages in community service. She is also involved in academic publications and scientific collaborations. Yuni is dedicated to advancing informatics education and promoting the use of technology as a driver of innovation and social benefit, particularly in Indonesia's educational and health technology sectors.



Ervin Yohannes is a lecturer at the State University of Surabaya (UNESA) specializing in Computer Science. He obtained his Ph.D. in Computer Science from the National Central University, Taiwan, in 2023, after previously completing his bachelor's and master's degrees in Computer Science from leading universities in Indonesia. His research areas include artificial intelligence, machine learning, data science, and distributed systems. He has published academic works in both national and international journals and has participated in various collaborative projects. As an academic, he is dedicated to mentoring students, conducting innovative research, and applying computer science knowledge to develop practical solutions. His long-term vision is to contribute to the advancement of artificial intelligence and its applications in industry and society.