

Comparative Analysis of Naïve Bayes Variants for Predicting Stunting-Risk Families

Alwas Muis^{1*}, Zila Razilu¹, Umy Ramadhani Senga¹, Hutri Wulandari¹, Reva Andriyani¹

¹Information Technology Education Study Program, Universitas Muhammadiyah Kendari
Jl. KH. Ahmad Dahlan, No. 10, Kota Kendari, Sulawesi Tenggara, Indonesia

Correspondence e-mail: alwas.muis@umkendari.ac.id

Submission:	Revision:	Acceptance:	Available Online:
23-01-2026	15-04-2026	17-04-2026	22-04-2026

Abstract - Stunting is a chronic nutritional condition that adversely affects children's physical growth and cognitive development, highlighting the need for effective early detection, particularly at the household level. This study proposes a comparative analysis of three Naïve Bayes variants Gaussian, Multinomial, and Bernoulli to identify families at risk of stunting using machine learning techniques. The dataset used in this study consists of family-level records obtained from the National Population and Family Planning Agency (BKKBN) of Southeast Sulawesi Province, comprising demographic, socioeconomic, and health-related attributes. Data preprocessing involved handling missing values, removing irrelevant attributes, and transforming categorical variables. The dataset was divided into training and testing sets using an 80:20 ratio. The main contribution of this study lies in evaluating the effectiveness of different Naïve Bayes variants for family-based stunting risk prediction, which has been rarely explored in previous studies. Model performance was evaluated using accuracy, precision, recall, and F1-score. The results indicate that the Bernoulli Naïve Bayes model achieved the best performance, with an accuracy of 88% and balanced evaluation metrics across both classes. These findings suggest that the Bernoulli Naïve Bayes model is the most suitable approach for predicting family-level stunting risk and can support data-driven early intervention strategies.

Keywords: Machine Learning, Naïve Bayes Variants, Stunting Risk Prediction, Family-Level Classification, Public Health Analytics

1. Introduction

Stunting remains a significant public health problem at both global and national levels, particularly in developing countries (Maryati et al., 2023; Hussain et al., 2025). It is defined as a condition of impaired linear growth in children caused by chronic malnutrition, especially during the first thousand days of life, which is a critical period for physical and cognitive development (Raiten & Bremer, 2020). Children affected by stunting are at higher risk of experiencing long-term consequences, including reduced cognitive performance, lower educational attainment, and decreased productivity in adulthood. In Indonesia, the prevalence of stunting remains relatively high and continues to be a major concern in national development, as it reflects the overall quality of human resources (Beal et al., 2018; Ahmed et al., 2023).

Stunting is widely recognized as a multifactorial condition influenced by a complex interaction of health, socioeconomic, and environmental factors (Laksono et al., 2022). At the family level, several determinants have been identified, including parental education, household income, access to adequate nutrition, sanitation, and clean water, as well as parenting practices and maternal health conditions during pregnancy (Fauziah, 2023; Atamou et al., 2023; Mulyani et al., 2025; Siswati et al., 2022). These findings indicate that the family plays a central role in shaping children's growth and development outcomes. Therefore, identifying families at risk of stunting is a crucial step toward implementing early and targeted interventions.

Despite various government programs aimed at reducing stunting prevalence, current detection approaches are still largely reactive. In most cases, stunting is identified only after children exhibit growth failure based on anthropometric measurements collected through periodic monitoring (Huru et al., 2024). Such approaches limit the effectiveness of early prevention strategies, as interventions are implemented after the condition has already occurred. In addition, manual data processing and reporting systems often require substantial time and resources, which may delay decision-making and reduce responsiveness. These limitations highlight the need for more proactive, efficient, and data-driven approaches to support early identification of stunting risk, particularly at the family level.

With the rapid advancement of information technology, machine learning techniques have been increasingly applied in the health domain to support predictive analysis and decision-making (Reddy et al., 2024). Machine learning enables the processing of large-scale and multidimensional data, including demographic,

socioeconomic, and health-related variables, to uncover hidden patterns that are difficult to detect using conventional statistical methods (Resti et al., 2023; Vaccari et al., 2022; Shen et al., 2023). In the context of stunting, several studies have demonstrated that machine learning models can improve prediction accuracy and support early detection of at-risk populations (Khanapi et al., 2021). However, most existing studies focus primarily on predicting stunting at the individual child level rather than at the family level, which may limit the effectiveness of preventive interventions.

Furthermore, although the Naïve Bayes algorithm is known for its simplicity, efficiency, and strong performance in classification tasks, its application in stunting prediction studies is often limited to a single variant without comprehensive comparison. Previous research has shown that different variants of Naïve Bayes, such as Gaussian, Multinomial, and Bernoulli, have distinct characteristics depending on the type and distribution of the data (Rintyarna et al., 2022). However, there is still a lack of studies that systematically evaluate and compare these variants in the context of family-level stunting risk prediction.

Based on the aforementioned limitations, two main research gaps can be identified. First, most prior studies emphasize individual-level prediction, while the family as a fundamental unit of intervention has received limited attention. Second, there is a lack of comparative analysis of different Naïve Bayes variants to determine the most suitable model for handling family-level health and socioeconomic data. To the best of our knowledge, no previous study has specifically addressed these two aspects simultaneously. Therefore, this study offers a novel contribution by integrating family-level analysis with a comparative evaluation of multiple Naïve Bayes variants for stunting risk prediction.

Accordingly, this study aims to develop and evaluate a machine learning-based classification model for predicting families at risk of stunting using three variants of the Naïve Bayes algorithm: Gaussian, Multinomial, and Bernoulli. The main contribution of this study lies in providing a systematic comparison of these variants to identify the most effective model for family-level prediction. In addition, this study demonstrates how machine learning can be utilized to transform family health data into actionable insights for early detection. From a practical perspective, the findings are expected to support policymakers, health institutions, and related stakeholders in designing more targeted, efficient, and data-driven stunting prevention programs.

2. Research Methods

This study adopts a quantitative research design with an experimental approach, applying data mining techniques to develop and evaluate classification models for predicting families at risk of stunting (Sutarmi et al., 2023; Gan et al., 2021). The experimental design enables a systematic comparison of multiple machine learning models under the same conditions to ensure objective performance evaluation.

The overall research process is illustrated in Figure 1, which consists of five main stages: literature review, data collection, data preprocessing, model development, and performance evaluation. To improve clarity, each stage is described in detail in the following subsections.

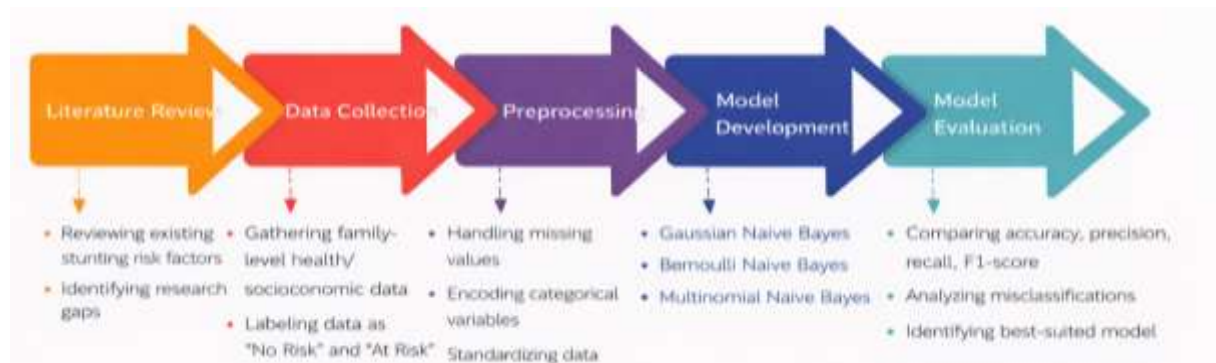


Figure 1. Research stages

2.1. Literature Review

The literature review was conducted to establish a theoretical foundation related to stunting, its risk factors, and the application of machine learning techniques in health data analysis. Relevant sources were collected from peer-reviewed journals, conference proceedings, and official reports. This stage also aimed to identify research gaps and determine appropriate variables and modeling approaches for the study.

2.2. Data Collection

The dataset used in this study was obtained from the National Population and Family Planning Agency (BKKBN) of Southeast Sulawesi Province. The data consist of family-level records containing demographic, socioeconomic, and health-related attributes associated with stunting risk.

After the initial cleaning process, the dataset comprised 1,479 family records with a combination of numerical and categorical variables. The target variable is the stunting risk status, categorized into two classes: *At Risk* and *No Risk*. The input variables include indicators such as parental education level, household economic status, maternal health conditions, and access to sanitation and clean water.

2.3. Preprocessing

Data preprocessing was performed to ensure data quality and improve model performance. Several steps were applied systematically:

1. Data Cleaning: Missing values were handled using appropriate imputation techniques (e.g., mode imputation for categorical variables and mean imputation for numerical variables). Duplicate records were removed to avoid bias.
2. Feature Selection: Irrelevant attributes, including personal identification data, were removed to maintain privacy and reduce noise in the dataset.
3. Data Transformation: Categorical variables were converted into numerical representations using encoding techniques such as label encoding and binary encoding, depending on the nature of the variables.
4. Outlier Handling: Extreme values in numerical variables were identified using statistical thresholds and treated to minimize their impact on model performance.
5. Data Splitting: The dataset was divided into training and testing sets using an 80:20 ratio to ensure unbiased model evaluation.

2.4. Model Development Process

In this study, the Naive Bayes algorithm is employed as the classification method to predict families at risk of stunting. Naive Bayes is a probabilistic algorithm based on Bayes' Theorem and operates under the assumption of conditional independence among input variables. Despite this simplifying assumption, Naive Bayes has demonstrated strong performance in various classification tasks, particularly in health and social data analysis. This study compares three variants of the Naive Bayes model: Gaussian Naive Bayes, which is suitable for normally distributed numerical data; Multinomial Naive Bayes, commonly applied to frequency-based data; and Bernoulli Naive Bayes, which is appropriate for binary-valued data. The comparison aims to identify the most optimal model for classifying family-level stunting risk based on the characteristics of the dataset.

The model development process involves training and testing each Naive Bayes variant. The preprocessed dataset is divided into training and testing sets using an 80:20 ratio, with 80% of the data used for training and 20% reserved for testing. The training data are used to build the Gaussian Naive Bayes, Multinomial Naive Bayes, and Bernoulli Naive Bayes models, where each model learns the relationship between input variables and family stunting risk status. Subsequently, the testing data are used to objectively evaluate the performance of each model under the same testing scheme. The prediction results from the three models are then compared to determine the model with the best performance in classifying families at risk of stunting.

2.5. Model Performance Evaluation

Model performance evaluation was conducted to assess the ability of each Naive Bayes algorithm variant to classify families at risk of stunting. The evaluation process was performed using the testing dataset, which comprised 20% of the total dataset and was not involved in the model training phase. Model performance was analyzed using a confusion matrix to represent the classification results between the at-risk and not-at-risk classes. To quantitatively measure classification performance, this study employed several evaluation metrics commonly used in binary classification problems, namely accuracy, precision, recall, and F1-score. These metrics were calculated based on the values obtained from the confusion matrix, which consists of True Positive (TP), True Negative (TN), False Positive (FP), and False Negative (FN).

The TP value represents the number of families that are truly at risk of stunting and are correctly classified as at risk by the model, while TN represents the number of families that are not at risk and are correctly classified as not at risk. Meanwhile, FP indicates the number of families that are actually not at risk but are incorrectly classified as at risk, and FN represents families that are actually at risk but are incorrectly classified as not at risk.

Accuracy is used to measure the overall correctness of the model in classifying the testing data. This metric represents the proportion of correct predictions out of all evaluated instances. Accuracy is calculated using Equation (1).

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (1)$$

Furthermore, precision measures the accuracy of the model in predicting the class of families at risk of stunting. Precision indicates the proportion of predicted at-risk cases that are truly at risk. This metric is important for minimizing incorrect positive predictions (false positives). Precision is calculated using Equation (2).

$$Precision = \frac{TP}{TP + FP} \quad (2)$$

Recall is used to measure the model's ability to detect all families that are truly at risk of stunting. This metric indicates the extent to which the model can minimize incorrect negative predictions (false negatives), which is particularly important in the context of public health. Recall is calculated using Equation (3).

$$Recall = \frac{TP}{TP + FN} \quad (3)$$

To provide a balance between precision and recall, the F1-score is used, which represents the harmonic mean of precision and recall. The F1-score is particularly useful when the class distribution is imbalanced, as it considers both metrics simultaneously. F1-score is calculated using Equation (4).

$$F1-Score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (4)$$

By using these four evaluation metrics, this study obtains a more comprehensive understanding of the performance of each variant of the Naive Bayes algorithm. This evaluation not only assesses the overall accuracy of the models but also considers their ability to accurately and consistently identify families at risk of stunting, thereby providing a strong basis for determining the most suitable model for family-level stunting risk prediction.

3. Result and Discussion

3.1. Result

All datasets were consistently processed through the same preprocessing stages to ensure data quality and uniformity prior to model development. The initial step involved removing family identification data to maintain data confidentiality and to avoid irrelevant bias in the classification process. Subsequently, variables that were considered unrelated to the research objectives or did not contribute significantly to the prediction process were removed from the dataset.

In the next stage, missing values were handled to prevent distortion of the analysis results and to improve the stability of the developed models. In addition, the identification and treatment of outliers were performed to minimize the influence of extreme values that could potentially reduce classification accuracy. The application of uniform preprocessing procedures aims to ensure that the performance comparison among Gaussian Naive Bayes, Multinomial Naive Bayes, and Bernoulli Naive Bayes is conducted fairly and objectively, so that the evaluation results validly represent the capabilities of each model.

1. Gaussian NB

Based on the testing results, the Gaussian Naïve Bayes model achieved an accuracy of 69.26%, indicating moderate overall classification performance. However, accuracy alone does not fully reflect the model's effectiveness, particularly in the context of imbalanced class importance in public health applications such as stunting risk prediction.

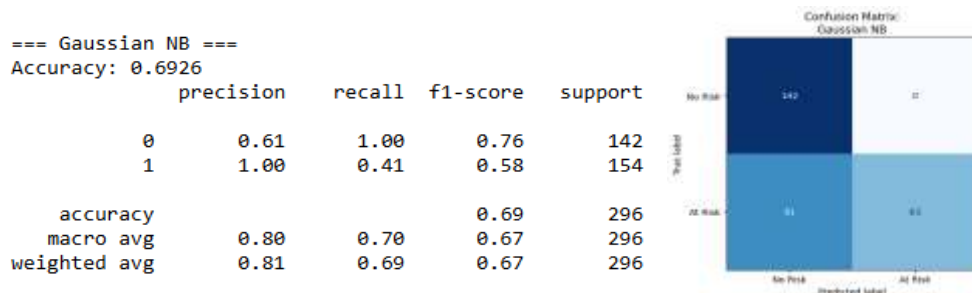


Figure 2. Gaussian model results

Figure 2 illustrates the confusion matrix of the Gaussian Naïve Bayes model. The model demonstrates a strong ability to classify the *No Risk* class, correctly identifying nearly all non-risk families. This suggests that the model effectively captures the general characteristics of families that are not at risk of stunting. However, this strong performance is not consistent across classes.

In contrast, the model shows limited capability in identifying families belonging to the *At Risk* class. A substantial number of at-risk families are misclassified as *No Risk*, indicating that the model fails to adequately learn the distinguishing patterns associated with stunting risk. This imbalance in performance between classes reflects a bias toward the majority or more easily separable class. From a probabilistic perspective, this behavior may be influenced by the assumption of normal distribution in Gaussian Naïve Bayes, which may not fully align with the distribution of the dataset, particularly if several features are categorical or non-Gaussian in nature. As a result, the model becomes less sensitive to patterns that characterize at-risk families.

Although the model achieves perfect precision for the *At Risk* class, this result indicates that predictions labeled as at risk are highly reliable. However, this comes at the expense of very low recall, meaning that many actual at-risk families are not detected. This trade-off suggests that the model is overly conservative in assigning the *At Risk* label.

This issue is particularly critical in a public health context. In stunting prevention, failing to identify at-risk families (*false negatives*) is more problematic than incorrectly labeling non-risk families as at risk (*false positives*). Low recall in the *At Risk* class implies that a significant proportion of vulnerable families may not receive early intervention, potentially reducing the effectiveness of prevention programs.

The F1-score further confirms this imbalance, showing that while the model performs relatively well for the *No Risk* class, its performance for the *At Risk* class remains suboptimal. This indicates that the model lacks robustness in handling class-specific prediction tasks.

Overall, the Gaussian Naïve Bayes model demonstrates limited suitability for this classification problem due to its inability to balance performance across classes. Compared to other Naïve Bayes variants, this result suggests that Gaussian assumptions may not be appropriate for the underlying data characteristics, highlighting the importance of selecting model variants that align with feature distributions and data representation.

2. Bernoulli NB

The Bernoulli Naïve Bayes model achieved an accuracy of 87.50%, indicating a substantial improvement over the Gaussian variant. More importantly, this model demonstrates a well-balanced performance across both classes, which is critical for reliable stunting risk prediction.

=== Bernoulli NB ===

Accuracy: 0.8750

	precision	recall	f1-score	support
0	0.86	0.89	0.87	142
1	0.89	0.86	0.88	154
accuracy			0.88	296
macro avg	0.87	0.88	0.87	296
weighted avg	0.88	0.88	0.88	296

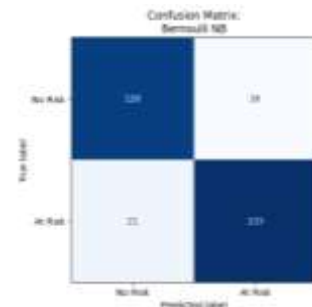


Figure 3. Bernoulli model results

Figure 3 presents the confusion matrix of the Bernoulli Naïve Bayes model. Unlike the Gaussian model, which shows a bias toward the No Risk class, the Bernoulli model is able to maintain a relatively even distribution of correct predictions for both No Risk and At Risk classes. This balance is reflected in the comparable precision and recall values across classes, indicating that the model is neither overly conservative nor biased toward a specific class.

From an analytical perspective, the superior performance of the Bernoulli Naïve Bayes model can be attributed to its suitability for handling binary or discretized feature representations. In this study, several variables were transformed into categorical or binary forms during preprocessing. As a result, the Bernoulli assumption where features are treated as binary indicators aligns more closely with the underlying data characteristics compared to the Gaussian assumption of continuous normal distribution. This alignment enables the model to better capture the presence or absence of risk-related attributes within family data.

Furthermore, the relatively high recall for the At Risk class indicates that the model is capable of identifying the majority of families that require early intervention. This is particularly important in a public health context, where minimizing false negatives is essential to ensure that vulnerable populations are not overlooked. At the same time, the model maintains strong precision, suggesting that the identified at-risk families are indeed relevant targets for intervention programs.

The consistency between precision, recall, and F1-score across both classes also suggests that the Bernoulli Naïve Bayes model is more robust and stable compared to the other variants. This balance indicates that the model performs reliably under the given data distribution without favoring one class over the other.

Overall, the results confirm that Bernoulli Naïve Bayes is the most suitable model for this classification task. Its ability to align with the binary nature of the dataset and to provide balanced predictive performance makes it particularly effective for family-level stunting risk prediction.

3. Multinomial NB

The Multinomial Naïve Bayes model achieved an accuracy of 70.95%, indicating moderate classification performance, although still lower than that of the Bernoulli variant. While the model demonstrates a relatively balanced ability to classify both *No Risk* and *At Risk* classes, a considerable number of misclassifications remain, particularly in identifying at-risk families.

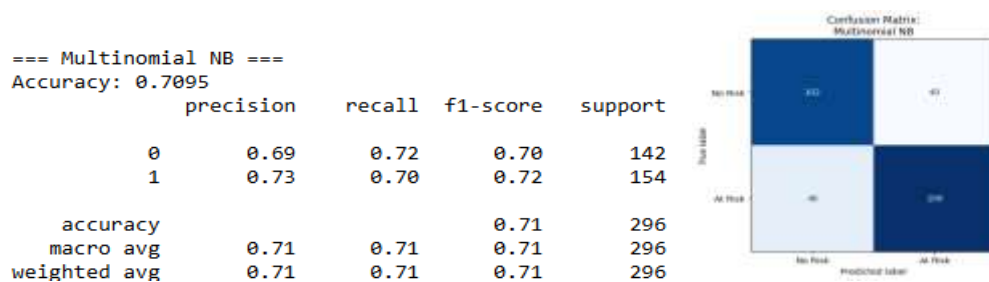


Figure 4. Multinomial model results

Figure 4 presents the confusion matrix of the Multinomial Naïve Bayes model. Overall, the model shows comparable precision and recall across both classes, suggesting no strong bias toward a specific class. However, this apparent balance is accompanied by a relatively high error rate, indicating that the model struggles to clearly distinguish between the two classes. From an analytical perspective, the limitations of the Multinomial Naïve Bayes model can be attributed to the mismatch between its underlying assumptions and the characteristics of the dataset. The Multinomial model is typically well-suited for frequency-based data, such as word counts in text classification. In contrast, the dataset used in this study consists primarily of categorical and binary features rather than frequency-based representations. As a result, the model may not effectively capture the presence or absence of critical risk indicators, leading to overlapping class boundaries and increased misclassification.

This limitation is particularly evident in the *At Risk* class, where a notable portion of at-risk families is incorrectly classified as *No Risk*. In the context of stunting prevention, this issue is critical, as false negatives may result in missed opportunities for early intervention. Compared to the Bernoulli Naïve Bayes model, which explicitly models binary feature presence, the Multinomial approach appears less sensitive to the underlying structure of the data, reducing its effectiveness in detecting subtle risk patterns.

Furthermore, although the F1-scores indicate a relatively consistent balance between precision and recall, this consistency does not necessarily reflect optimal performance. Instead, it suggests that the model performs moderately across both classes without achieving strong discriminative power. This highlights the importance of aligning model assumptions with data representation to improve classification performance.

Overall, the results indicate that the Multinomial Naïve Bayes model is less suitable for this classification task compared to the Bernoulli variant. Its inability to effectively represent binary and categorical features leads to higher misclassification rates, particularly for the *At Risk* class, limiting its applicability in early stunting risk prediction.

3.2. Discussion

Based on the comparative evaluation of the three Naïve Bayes variants Gaussian, Bernoulli, and Multinomial significant differences are observed in their ability to classify *No Risk* and *At Risk* families. Table 1 summarizes the performance of each model using accuracy, precision, recall, and F1-score. However, beyond these quantitative metrics, a deeper analysis reveals that the effectiveness of each model is strongly influenced by the compatibility between model assumptions and the underlying data characteristics.

The Gaussian Naïve Bayes model demonstrates a clear imbalance in classification performance across classes. While it performs well in identifying *No Risk* families, it fails to adequately detect *At Risk* cases, as reflected in its low recall for this class. This indicates that the model is biased toward the dominant or more easily separable class. From a methodological perspective, this limitation may arise from the Gaussian assumption of normally distributed continuous features, which does not fully align with the dataset that includes categorical and discretized variables. As a result, the model becomes less sensitive to patterns associated with stunting risk, leading

to a high number of false negatives

Table 1. Presents a comparison of the model testing results

Naïve Bayes	Class	Precision	Recall	F1-Score	Accuracy
Gaussian	No Risk	61	100	76	69%
	At Risk	100	41	58	
Bernoulli	No Risk	86	89	87	88%
	At Risk	89	86	88	
Multinomial	No Risk	69	72	70	71%
	At Risk	73	70	72	

In contrast, the Bernoulli Naïve Bayes model achieves the most balanced and consistent performance across both classes. This superior performance can be explained by its suitability for binary feature representation. During preprocessing, several variables were transformed into categorical or binary forms, making the Bernoulli assumption more appropriate for capturing the presence or absence of risk-related attributes. Consequently, the model is better able to distinguish between *At Risk* and *No Risk* families, resulting in both high recall and precision. This balance is particularly important in classification tasks where both types of errors false positives and false negatives must be carefully controlled.

The Multinomial Naïve Bayes model, although relatively stable, shows moderate performance with a higher misclassification rate compared to the Bernoulli variant. This can be attributed to the fact that the Multinomial model is designed for frequency-based data, whereas the dataset in this study primarily consists of categorical and binary variables rather than count-based features. This mismatch reduces the model's ability to effectively capture the distribution of risk-related attributes, leading to less accurate classification, particularly for the *At Risk* class.

From a public health perspective, the implications of misclassification are critical. In particular, false negatives in the *At Risk* class where families at risk are incorrectly classified as not at risk can result in missed opportunities for early intervention. Such errors may delay preventive actions and reduce the overall effectiveness of stunting reduction programs. Therefore, models with higher recall for the *At Risk* class, such as Bernoulli Naïve Bayes, are more desirable in this context, even if slight trade-offs in other metrics occur.

Overall, this study highlights that model performance is not solely determined by accuracy but also by how well the model aligns with data representation and the specific objectives of the application domain. The findings suggest that Bernoulli Naïve Bayes is the most suitable approach for family-level stunting risk prediction, as it provides a balance between predictive performance and practical relevance. From a computational perspective, its simplicity, efficiency, and compatibility with binary features make it a reliable choice for implementation in real-world decision support systems. At the same time, from a practical standpoint, the model has the potential to support policymakers and health practitioners in identifying high-risk families more effectively and enabling timely, targeted interventions.

4. Conclusion

This study evaluated the performance of three Naïve Bayes variants Gaussian, Multinomial, and Bernoulli in predicting families at risk of stunting using family-level health and socioeconomic data. The results indicate that each model exhibits different classification characteristics, largely influenced by the compatibility between model assumptions and data representation. The Gaussian Naïve Bayes model shows a clear imbalance in performance, with strong capability in identifying *No Risk* families but limited effectiveness in detecting *At Risk* cases. This limitation reduces its suitability for early risk identification. The Multinomial Naïve Bayes model provides more balanced performance; however, its effectiveness remains moderate due to its limited ability to capture the underlying structure of categorical and binary features in the dataset. Among the three models, the Bernoulli Naïve Bayes model demonstrates the best overall performance, achieving high and balanced precision, recall, and F1-score across both classes. This indicates that the model is not only accurate but also consistent in identifying at-risk families, which is essential for supporting early intervention strategies. The superior performance of this model is closely related to its compatibility with binary feature representation, which aligns well with the characteristics of the preprocessed dataset. From a practical perspective, the findings highlight the importance of selecting classification models that align with data characteristics, particularly in public health applications. The ability of the Bernoulli Naïve Bayes model to minimize false negatives makes it especially valuable for stunting prevention, as it supports more effective identification of families requiring early intervention.

For future work, further improvements can be achieved by applying advanced preprocessing and feature selection techniques, as well as evaluating additional classification methods such as Support Vector Machines, Random Forests, and ensemble approaches. The use of larger and more diverse datasets, along with cross-validation strategies, is also recommended to enhance model robustness and generalizability in real world applications.

Acknowledgement

The authors would like to express their sincere gratitude to the Directorate of Research, Technology, and Community Service (DRTPM) of Universitas Muhammadiyah Kendari for the support and funding provided, which made this research possible.

References

- Ahmed, K. Y., Ross, A. G., Hussien, S. M., Agho, K. E., Olusanya, B. O., & Ogbo, F. A. (2023). Mapping local variations and the determinants of childhood stunting in Nigeria. *International Journal of Environmental Research and Public Health*, 20(3250), 1–16. <https://doi.org/10.3390/ijerph20043250>
- Atamou, L., Rahmadiyah, D. C., & Hassan, H. (2023). Analysis of the determinants of stunting among children aged below five years in stunting locus villages in Indonesia. *Healthcare*, 11, 1–12. <https://doi.org/10.3390/healthcare11060810>
- Beal, T., Tumilowicz, A., Sutrisna, A., Izwardy, D., & Neufeld, L. M. (2018). A review of child stunting determinants in Indonesia. *Maternal and Child Nutrition*, 14(4), 1–10. <https://doi.org/10.1111/mcn.12617>
- Fauziah, S. R. (2023). Family household characteristics and stunting: An updated scoping review. *Nutrients*, 15, 1–17. <https://doi.org/10.3390/nu15010233>
- Gan, S., Shao, S., Chen, L., Yu, L., & Jiang, L. (2021). Adapting hidden Naïve Bayes for text classification. *Mathematics*, 9, 1–14. <https://doi.org/10.3390/math9192378>
- Huru, M. M., Mamoh, K., & Mangi, J. L. (2024). Pemberdayaan kader kesehatan dan orang tua asuh anak stunting dalam pencegahan dan penatalaksanaan stunting. *JMM (Jurnal Masyarakat Mandiri)*, 7(6), 5365–5373. <https://doi.org/10.31764/jmm.v7i6.17638>
- Hussain, I., Chauhadry, I. A., Umer, M., Nisa, N., Nadeem, S., Hassan, M., Sattar, A. A., Habib, M. A., Ariff, S., Jahangir, A., Hudspeth, C., Soofi, S. B., & Bhutta, Z. A. (2025). The impact of the Central Asia stunting initiative on stunting among children under five years old. *Nutrients*, 17(3255), 1–12. <https://doi.org/10.3390/nu17203255>
- Khanapi, M., Ghani, A., Noma, N. G., Mohammed, M. A., Abdulkareem, K. H., Garcia-Zapirain, B., Maashi, M. S., & Mostafa, S. A. (2021). Innovative artificial intelligence approach for hearing-loss symptoms identification using machine learning techniques. *Sustainability*, 13(5406), 1–30. <https://doi.org/10.3390/su13105406>
- Laksono, A. D., Edi, N., Sukoco, W., Rachmawati, T., & Wulandari, R. D. (2022). Factors related to stunting incidence in toddlers with working mothers in Indonesia. *International Journal of Environmental Research and Public Health*, 19, 1–9. <https://doi.org/10.3390/ijerph191710654>
- Maryati, I., Annisa, N., & Amira, I. (2023). Faktor dominan terhadap kejadian stunting balita. *Jurnal Obsesi: Jurnal Pendidikan Anak Usia Dini*, 7(3), 2695–2707. <https://doi.org/10.31004/obsesi.v7i3.4419>
- Mulyani, A. T., Khairinisa, M. A., Khatib, A., & Chaerunisaa, A. Y. (2025). Understanding stunting: Impact, causes, and strategies to accelerate reduction. *Nutrients*, 17(9), 1–29. <https://doi.org/10.3390/nu17091493>
- Pan, Y. (2018). Identification of bacteriophage virion proteins using Multinomial Naïve Bayes with g-gap feature tree. *International Journal of Molecular Sciences*, 19(1779), 1–20. <https://doi.org/10.3390/ijms19061779>
- Raiten, D. J., & Bremer, A. A. (2020). Exploring the nutritional ecology of stunting: New approaches to an old problem. *Nutrients*, 12(371), 1–13. <https://doi.org/10.3390/nu12020371>
- Reddy, G. P., Rohan, D., Venkata, Y., Kumar, P., Prakash, K. P., & Srikanth, M. (2024). Artificial intelligence-based detection of Parkinson's disease using voice measurements. *Engineering Proceedings*, 82, 1–8. <https://doi.org/10.3390/ecsai-11-20481>
- Resti, Y., Irsan, C., Neardiaty, A., Annabila, C., & Yani, I. (2023). Fuzzy discretization on the Multinomial Naïve Bayes method for multiclass classification. *Mathematics*, 11, 1–21. <https://doi.org/10.3390/math11081761>
- Rintyarna, B. S., Kuswanto, H., Sarno, R., & Rachmaningsih, E. K. (2022). Modelling service quality of internet service providers during COVID-19. *Informatics*, 9(11), 1–12. <https://doi.org/10.3390/informatics9010011>
- Shen, H., Zhao, H., & Jiang, Y. (2023). Machine learning algorithms for predicting stunting among under-five children. *Children*, 10(1638), 1–18. <https://doi.org/10.3390/children10101638>
- Sin, M. P., Forsberg, B. C., Peterson, S. S., & Alfv, T. (2024). Assessment of childhood stunting prevalence and risk factors in Bangladesh. *Children*, 11(650), 1–15. <https://doi.org/10.3390/children11060650>
- Siswati, T., Iskandar, S., Pramestuti, N., Raharjo, J., & Rubaya, A. K. (2022). Drivers of stunting reduction in Yogyakarta, Indonesia: A case study. *International Journal of Environmental Research and Public Health*, 19(16497), 1–14. <https://doi.org/10.3390/ijerph192416497>
- Sutarmi, S., Warijan, W., Indrayana, T., B, D. P. P., & Gunawan, I. (2023). Machine learning model for stunting prediction. *Jurnal Health Sains*, 4(9), 10–23. <https://doi.org/10.46799/jhs.v4i9.1073>
- Vaccari, I., Carlevaro, A., Narteni, S., & Mongelli, M. (2022). Explainable and reliable machine learning in data analytics. *IEEE Access*, 10, 83949–83970. <https://doi.org/10.1109/ACCESS.2022.3197299>