

Implementation of the Naive Bayes Method for Stunting Classification in Children Under Five Years Old (*Balita*).

Sulthan Ahmad Ferdiansyah¹
Informatika

Universitas Pembangunan Nasional
"Veteran" Jawa Timur
Surabaya, Indonesia
20081010156@student.upnjatim.ac.id

Muhammad Farhan Maulana²
Informatika

Universitas Pembangunan Nasional
"Veteran" Jawa Timur
Surabaya, Indonesia
20081010159@student.upnjatim.ac.id

Akmal Aliffandhi Anwar³
Informatika

Universitas Pembangunan Nasional
"Veteran" Jawa Timur
Surabaya, Indonesia
20081010148@student.upnjatim.ac.id

Muhammad Rifki Bahrul Ulum⁴
Informatika

Universitas Pembangunan Nasional
"Veteran" Jawa Timur
Surabaya, Indonesia
20081010149@student.upnjatim.ac.id

I Gede Susrama Diyasa⁵

Magister Teknologi Informasi
Universitas Pembangunan Nasional
"Veteran" Jawa Timur
Surabaya, Indonesia
igsusrama.if@upnjatim.ac.id

*Vinza Hedi Satria⁶
Informatika

Universitas Pembangunan Nasional
"Veteran" Jawa Timur
Surabaya, Indonesia
vinzasatria.if@upnjatim.ac.id

**Correspondence Author*

Abstract— The issue of stunting in Indonesia has become a serious concern, drawing significant attention from the government. To address this problem, the government has set a target to reduce the stunting prevalence rate to 14% by 2024. As an initial step in supporting this goal, the present study aims to classify the nutritional status of children under five years old using the Naive Bayes method. The objective of this research is to evaluate the performance of the Naive Bayes algorithm in classifying the nutritional status of children under five, with a focus on body weight, height, and exclusive breastfeeding intake as predictors of stunting. The research process includes several stages, namely problem formulation, data collection, data preprocessing, data splitting, model construction, model training, model evaluation, and result analysis. The findings of this study indicate that the Naive Bayes method achieved an accuracy of 70% in classifying stunting among children under the age of five, with an F1-score evaluation of 70%.

Keywords— Classification, Naive Bayes, Stunting



© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution ShareAlike (CC BY SA) license (<https://creativecommons.org/licenses/by-sa/4.0/>).

I. INTRODUCTION

Stunting is a condition affecting children under five years of age, characterised by a height significantly lower than that of their peers, primarily caused by nutritional deficiencies. Stunting can have serious long-term consequences for a child's health and quality of life. Children who experience stunting are at a higher risk of developing health problems, including chronic diseases. Moreover, stunting can impair cognitive development, negatively affecting learning abilities and overall development. It is recognised as a critical public health challenge with harmful functional consequences, including poor cognitive and educational performance, reduced productivity, and increased risk of nutrition-related chronic diseases in adulthood [1].

The occurrence of stunting is influenced by multiple factors, such as the lack of exclusive breastfeeding, insufficient energy and protein intake, suboptimal complementary feeding practices (MP-ASI), low household income, incomplete immunisation, and a history of infectious diseases [2]. Inadequate parental care—often related to young maternal age or closely spaced pregnancies—is also a contributing factor. To address this issue, the Indonesian Ministry of Health (Kemenkes RI) has made efforts to improve nutritional status as part of its national health development priority programmes [3].

Naïve Bayes Classifier, as simple but efficient method in solving various real world problem [4]. As one of the oldest means to classify data, Naïve Bayes has been upgraded few times to increase it's accuracy [5] and has been used for many problem solving matter, such as research conducted by [6] that utilizes Naïve Bayes as decision-making system to decide PIP Awardee.

Beside general real world problem. Previous studies have implemented the Naive Bayes algorithm in the classification of stunting status among toddlers. For instance, the study conducted by [7] developed a web-based application using the Naive Bayes method to classify the nutritional status of toddlers into stunting categories. The research aimed to enhance knowledge of data management to support decision-making and outreach related to stunting cases. The study achieved an accuracy rate of 88% using 300 data entries with parameters including gender, age, weight, height, poverty status, and nutritional status category. However, the limitation of this study lies in its focus solely on accuracy, without discussing other potentially relevant aspects.

Similarly, the research by [8] applied the Naive Bayes algorithm with K-Fold Cross Validation for the classification of stunting nutritional status in toddlers. The results showed that, through 10-Fold Cross Validation, the highest accuracy was obtained in the 8th iteration at 95.14%, while the lowest accuracy was observed in the 3rd iteration at 81.73%. Overall, the average accuracy across all iterations was 88.53%, using

parameters such as age, gender, height, weight, and upper arm circumference.

Furthermore, the study by [9] implemented the Naive Bayes algorithm combined with Forward Selection for classifying stunting nutritional status at Puskesmas Pandanaran Semarang. The study compared the classification accuracy of Naive Bayes with and without Forward Selection. The findings indicated that integrating Forward Selection improved the classification accuracy from 83.33% to 86%. The accuracy improvement was achieved by reducing the number of features from six to four. However, the study did not specify which features were used in the process.

Based on these previous studies, the primary focus of the present research is to utilise the Naive Bayes algorithm for the classification of stunting among children under five years old.

II. LITERATUR REVIEW

A. Stunting

Stunting is a condition in which a child's growth and development are hindered or lag behind that of their peers, particularly during the critical early growth period from infancy to five years of age. This condition is caused by inadequate nutritional intake, recurrent infections, and various social factors that influence child development. Children who experience stunting are at a higher risk of developing health problems, including chronic diseases. In the long term, stunting can result in decreased intellectual capacity or cognitive ability in adulthood, leading to lower productivity. Neurological and brain cell impairments may occur, resulting in slower learning absorption and the emergence of diseases such as diabetes, heart disease, stroke, and hypertension [10], as well as an increased risk of obesity [11].

According to the Indonesian Ministry of Health Regulation (Permenkes) No. 2 of 2020 concerning Child Anthropometric Standards, anthropometric data can be used as a reference for determining a child's nutritional status, which is then compared with threshold values (z-scores). Several indices commonly used to assess a child's nutritional status include Weight-for-Age (W/A), Height-for-Age (H/A), and Weight-for-Height (W/H). The index used to determine stunting is Height-for-Age (H/A).

B. Klasifikasi

Classification is the process of assigning data into specific groups or categories based on a set of defined features. The purpose of classification is to understand and organise data, as well as to facilitate decision-making based on the grouped information. The classification process involves the use of algorithms or methods to assign data to appropriate categories. Typically, classification is carried out by training a model on labelled data (training data), followed by testing the model on unlabelled data (testing data). The results of classification can be utilised for various purposes, such as decision-making, pattern recognition, or gaining deeper insights into the data.

C. Naive Bayes

Naive Bayes is a classification algorithm that determines the category of an instance based on Bayes' theorem. This method assumes conditional independence among features, meaning that the features used for classification are considered independent of one another, even though this may

not fully hold in real-world scenarios. The algorithm estimates the probability that an object belongs to a particular category based on its features and assigns it to the class with the highest probability. The two most commonly used Naive Bayes models are Multinomial and Bernoulli. Although simple and easy to implement, Naive Bayes may not be suitable for datasets in which the features are highly dependent on one another.

In the field of mathematical statistical theory, the Naive Bayes approach enables the construction of uncertainty models for future events by combining general knowledge with observed data. The formula for Bayes' theorem is as follows:

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

Dimana,

X : Unknown Data Class

H : Data Hypothesis

P(H|X) : Probability for Hypothesis H based on X Condition

P(H) : Probability for H Hypothesis

P(X|H) : Probability Hypothesis X based on Condition on Hypothesis H

P(X) : Probability X

To illustrate the Naive Bayes Theorem, it is important to understand that in the classification process, a guideline is required to determine the most appropriate category for the sample being analysed. Therefore, Bayes' theorem is adapted as follows:

$$P(C|F_i) = \frac{P(C)P(F_i|C)}{P(X)} \quad (2)$$

In this context, the variable C represents the class, while the variables F_1 to F_n represent the feature attributes required for classification. The formula indicates that the probability of a sample with certain characteristics belonging to class C (posterior probability) is obtained by multiplying the probability of class C occurring prior to observing the sample (also referred to as the prior) with the likelihood of observing the given features in class C , and then dividing it by the overall probability of observing the given sample features (also known as the evidence). Therefore, the formula above can be expressed in a more simplified form as follows:

$$Posterior = \frac{Prior \times Likelihood}{Evidence} \quad (3)$$

The evidence value remains constant for all classes within a single sample. The posterior probabilities are subsequently compared across different classes to determine the class to which a sample should be assigned [12]. In the context of data classification, Naive Bayes often employs the Gaussian distribution (Gaussian density) to calculate probabilities. The Gaussian density is used to evaluate the likelihood that a particular data point originates from a specific Gaussian distribution.

$$P(x_i = x_i | Y = y_j) = \frac{1}{\sqrt{2\pi\sigma_{ij}}} e^{-\frac{(x_i - \mu_{ij})^2}{2\sigma_{ij}^2}} \quad (4)$$

III. RESEARCH METHOD

The conducted research consist of multiple steps. First, the problem formulation and literature review has been presented in the previous section. Next phase will be data collection, data pre-processing, data splitting, model formulation and lastly, evaluation of the previously formulated model.

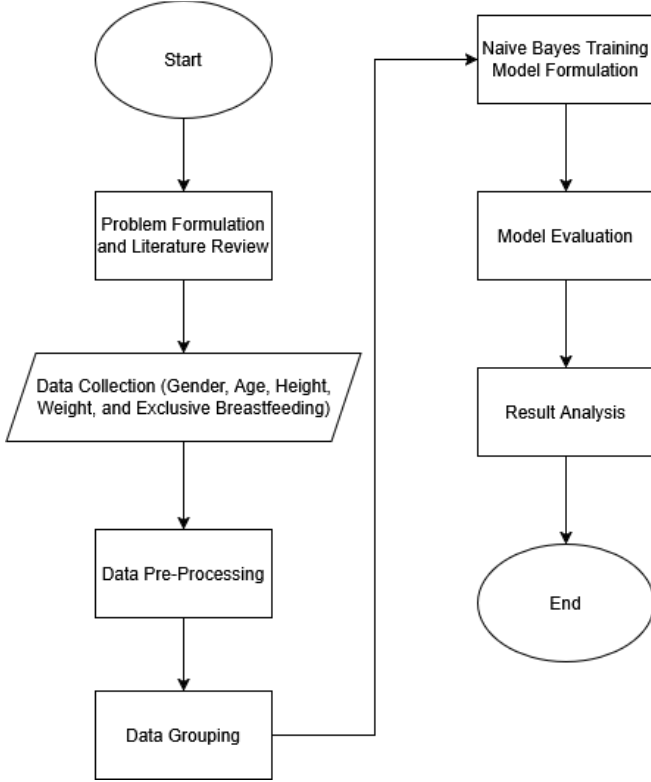


Fig. 1. Research Flowchart

A. Problem Formulation and Literature Review

At this stage, the authors identified the issues related to stunting in children under five years old (toddlers) and collected relevant journal references. The problem identification included analysing the factors that influence stunting among toddlers. The referenced journals provided a strong theoretical foundation. The Naive Bayes method was then employed to develop a classification model for stunting in children under five.

B. Data Collection

In this study, the data used were obtained from the Kaggle platform, consisting of a total of 6,500 records. Within the dataset, 3,312 entries were classified as stunting, while 3,188 entries were classified as normal. The dataset contains eight attributes, namely gender, age, birth weight and birth height, measurement weight and measurement height, exclusive breastfeeding status, and stunting status.

C. Data Pre-Processing

At this stage, four processes were carried out, namely data cleaning, transformation, and feature scaling. Data cleaning was performed to remove incomplete, incorrect, or missing

entries. Data transformation was conducted to modify the format or data types to meet the requirements of the study. Feature scaling was applied to ensure consistency across the features used.

D. Data Grouping

The dataset was divided into two parts, namely training data and testing data. The split was carried out proportionally, with 70% of the total data allocated for training the Naive Bayes model, while the remaining 30% was used to evaluate the model's performance. The division was conducted by taking into account the distribution of the stunting classes to maintain a balanced proportion between stunted and non-stunted children in each subset. This approach ensures that the model is trained and tested fairly without bias toward either class.

E. Naive Bayes Training Model Formulation

At this stage, the Naive Bayes algorithm model was constructed based on Bayes' Probability Theorem, assuming that each variable in the data is independent of one another with respect to the assigned class or label. The Naive Bayes process involves calculating the probability distributions of each feature (independent variable) and the target class (dependent variable). This includes estimating the probability of each feature value within each class, which is then used to compute the most likely class for a new instance based on its feature values.

The calculation employs the Bayesian formula as shown in Equation (1). The independent variables and/or required features include gender, age, birth weight, birth height, body weight, body height, and information regarding exclusive breastfeeding. The probability of the stunting class Y can be written as $P(X | Y_1)$, while the probability of the non-stunting class can be written $P(X | Y_2)$. Gender is categorized into two classes, namely Male and Female, and exclusive breastfeeding is also divided into two classes, namely Yes and No.

F. Model Evaluation

In this stage, the model was evaluated using a confusion matrix. Evaluation is a crucial step for gaining detailed insights into the performance of the classification model. The confusion matrix provides information on how accurately the model is able to predict the correct class and how frequently it makes errors.

Confusion Matrix		
	Prediksi	
	TRUE	FALSE
Aktual	TRUE	FALSE
TRUE	TP	FP
FALSE	FN	TN

Fig. 2. Confusion Matrix Table

Based on the figure, additional evaluation metrics such as precision, recall, and F1-score can be derived.

- Precision: The ability of the model to correctly identify the positive class out of all instances predicted as positive.

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

- Recall: The ability of the model to correctly identify the positive class out of all instances that are actually positive.

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

- F1-score: The combination of precision and recall, providing a holistic view of the model's performance.

$$F1-Score = \frac{2*Precision*Recall}{Precision+Recall} \quad (7)$$

IV. RESULT AND DISCUSSION

A. Research Variable

In studies related to stunting, the variables influencing this condition are generally divided into two main categories: the predictive variables (X) and the outcome variable (Y). The predictive variables refer to factors that can be used to estimate or forecast the likelihood of a child experiencing stunting, based on findings from previous research. Meanwhile, the outcome variable represents the actual stunting condition, typically defined using a binary classification of "Yes" or "No". This variable serves as the primary indicator for determining whether a child is classified as stunted or not.

TABLE I. RESEARCH VARIABLE

Variable	Name of Variable	Description
X1	Sex	1. Male 2. Female
X2	Age	Years Old
X3	Body Weight	Weight During Birth
X4	Body Len	Height During Birth
X5	Berat Badan	Weight Present
X6	Tinggi Badan	Height Present
X7	ASI Eksklusif	1. Yes 2. No
Y	Stunting	1. Yes 2. No

B. Naive Bayes Computation

TABLE II. TRAINING DATA

No	JK	U	BBL	TBL	BB	TB	ASI	Kat
1	M	16	2,4	48	6,6	76	Yes	Yes
2	F	47	2,5	49	10	91,5	No	Yes
3	F	31	3,3	49	9,1	90	Yes	No
4	F	25	2,6	49	9	90	No	No
5	M	4	2,7	47	9,8	69,6	Yes	No
6	F	12	3,6	50	8,2	70,5	Yes	No
7	F	3	2,8	48	6	54	No	Yes

8	M	55	3	49	8,5	81	No	No
9	F	12	2,9	49	5,8	69,5	Yes	No
10	M	19	3	52	7,7	57	Yes	Yes

Note : Gender (JK), Mace (M), Female (F), Age (U), Weight During Birth (BBL), Height During Birth (TBL), Weight Present (BB), Height Present (TB), Category (Kat)

The training data are then fed into the Naive Bayes algorithm to generate a predictive model for stunting status. To evaluate the Naive Bayes algorithm, the testing data are input into the prediction model. In this calculation, 10 sample records from the training data are used to determine the classification of the test data. The following are the steps for applying Naive Bayes using 5 test data samples.

TABLE III. TRAINING DATA

No	JK	U	BBL	TBL	BB	TB	ASI
1	F	14	2,9	49	6,5	66	No
2	M	11	2,9	49	7,7	66	No
3	F	20	1,8	48	7,3	73	Yes
4	M	9	2,9	50	7,3	62	No
5	M	53	2,9	49	15	96	Yes

Note : Gender (JK), Mace (M), Female (F), Age (U), Weight During Birth (BBL), Height During Birth (TBL), Weight Present (BB), Height Present (TB), Category (Kat)

1) Calculate the prior probability of stunting occurrence

- $P(Y1 = \text{Stunting}) = 4 / 10 = 0,4$
- $P(Y2 = \text{Not Stunting}) = 6 / 10 = 0,6$

2) Conditional probability according to class value in testing

Calculate $P(X | Y1)$ which is each attribute in the X data, then divided by the number of Y1 data.

- $P(\text{Female} | Y1) = 2 / 4 = 0.5$
- $P(\text{Male} | Y1) = 2 / 4 = 0.5$
- $P(\text{Breastfeed} = \text{Yes} | Y1) = 2 / 4 = 0.5$
- $P(\text{Breastfeed} = \text{No} | Y1) = 2 / 4 = 0.5$

Calculate $P(X | Y2)$ which is each attribute in the X data, then divided by the number of Y2 data.

- $P(\text{Female} | Y2) = 4 / 6 = 0.667$
- $P(\text{Male} | Y2) = 2 / 6 = 0.333$
- $P(\text{ASI} = \text{Yes} | Y2) = 4 / 6 = 0.667$
- $P(\text{ASI} = \text{No} | Y2) = 2 / 6 = 0.333$

3) If the data in the test is male and receiving breast milk

- $P(X|Y1) = P(\text{Breastfeed} = \text{Yes} | Y1) * P(\text{Male} | Y1)$
- $P(X|Y2) = P(\text{Breastfeed} = \text{Yes} | Y2) * P(\text{Male} | Y2)$

Search result of $P(X|Y1) * P(Y1)$ and $P(X|Y2) * Y2$

- $P(X|Y1) * P(Y1) = n$
- $P(X|Y2) * P(Y2) = n$

4) If the data in the test is male and does not receive breast milk

- $P(X|Y1) = P(\text{Breastfeed} = \text{No} | Y1) * P(\text{Male} | Y1)$
- $P(X|Y2) = P(\text{Breastfeed} = \text{No} | Y2) * P(\text{Male} | Y2)$

Search result of $P(X|Y1) * P(Y1)$ and $P(X|Y2) * Y2$

- $P(X|Y1) * P(Y1) = n$
- $P(X|Y2) * P(Y2) = n$

5) If the data in the test is female and receiving breast milk

- $P(X|Y1) = P(\text{Breastfeed} = \text{Yes} | Y1) * P(\text{Female} | Y1)$
- $P(X|Y2) = P(\text{Breastfeed} = \text{Yes} | Y2) * P(\text{Female} | Y2)$

Search for result $P(X|Y1) * P(Y1)$ and $P(X|Y2) * Y2$

- $P(X|Y1) * P(Y1) = n$
- $P(X|Y2) * P(Y2) = n$

6) If the data in the test is female and does not receive breast milk

- $P(X|Y1) = P(\text{Breastfeed} = \text{No} | Y1) * P(\text{Female} | Y1)$
- $P(X|Y2) = P(\text{Breastfeed} = \text{No} | Y2) * P(\text{Female} | Y2)$

Search for result $P(X|Y1) * P(Y1)$ and $P(X|Y2) * Y2$

- $P(X|Y1) * P(Y1) = n$
- $P(X|Y2) * P(Y2) = n$

Based on the calculations above, the five data are included in the stunting condition.

TABLE IV. PREDICTION RESULT

No	JK	U	BBL	TBL	BB	TB	ASI	Kat
1	M	16	2,4	48	6,6	76	Yes	Yes
2	F	47	2,5	49	10	91,5	No	Yes
3	F	31	3,3	49	9,1	90	Yes	Yes
4	F	25	2,6	49	9	90	No	Yes
5	M	4	2,7	47	9,8	69,6	Yes	Yes
Note : Gender (JK), Mace (M), Female (F), Age (U), Weight During Birth (BBL), Height During Birth (TBL), Weight Present (BB), Height Present (TB), Category (Kat)								

C. Evaluation

The validation stage employs the model's performance from the classification process using a Confusion Matrix to provide an overview of how well the model can distinguish between children who experience stunting and those who do not.

TABLE V. CONFUSION MATRIX

Actual	Prediction
--------	------------

	True	False
True	708 (TP)	319 (FP)
False	261 (FN)	662 (TN)
Note : True Positive (TP), False Positive(FP), False Negative (FN), False Positive (FP)		

In Table V, the actual class and the predicted class are presented. The actual class refers to labels that have been previously determined, while the predicted class represents the classification results produced by the Naive Bayes method. Based on the calculated confusion matrix, from the positive (stunting) perspective, the model successfully identified 708 samples as positive, whereas 261 samples that should have been classified as stunting were incorrectly predicted as negative. From the negative (non-stunting) perspective, the model accurately classified 662 samples as negative, but it also incorrectly predicted 319 samples as positive.

After calculating the Confusion Matrix, further computations were conducted to determine the precision, recall, and F1-score, the results of which are summarized in Table VI.

TABLE VI. RESULT OF PREDICTION

Performance	Not Stunting	Stunting	Accuracy
Precision	0.72	0.69	
Recall	0.67	0.73	
F1-Score	0.70	0.71	0.70

Based on Table VI, it can be observed that for class 0, the precision value of 0.72 indicates that 72% of the predicted positive instances for this class are truly relevant, while the recall value of 0.67 shows that the model successfully detected approximately 67% of the total cases that actually belong to class 0. The F1-score, which combines precision and recall, is approximately 0.70.

For class 1, a precision score of 0.69 indicates that 69% of the positive predictions are relevant, whereas a recall score of 0.73 suggests that the model was able to identify around 73% of the actual cases belonging to class 1. The F1-score for class 1 is approximately 0.71.

Overall, the classification model achieved an accuracy of 0.70, representing the proportion of correctly predicted instances across all classes.

V. ANALYSIS AND CONCLUSION

In the analysis of the classification model's performance, the information derived from the confusion matrix provides a deeper understanding of the model's predictive probability outcomes. From the confusion matrix, it can be observed that the model successfully identified 708 samples classified as positive within a high-probability interval (True Positive), indicating a strong level of confidence in correct classifications. Conversely, within lower-probability intervals, the model may exhibit uncertainty, as reflected in the 319 false positive cases. Nevertheless, the model was still able to detect the majority of actual positive cases.

The overall performance evaluation underscores the importance of maintaining model balance, as reflected by the high F1-score across both classes. A high F1-score demonstrates a strong alignment between the model's ability to correctly classify positive cases and its ability to minimize misclassification. With an accuracy of 0.70, it can be concluded that the model performs well in predicting classes overall. These results indicate that the model tends to produce correct and relevant outcomes at high confidence levels while also being capable of capturing most positive cases, even at lower confidence thresholds.

In a contextual sense, the Naive Bayes model appears to be suitable for classifying stunting cases in children under five years old. However, further validation and evaluation are necessary to ensure that the model is reliable and applicable to a wider range of datasets and variables.

References

- [1] "Stunting in a nutshell," World Health Organization, 19 November 2015. [Online].
- [2] E. Noorhasanah, N. I. Tauhidah and M. C. Putri, "FAKTOR-FAKTOR YANG BERHUBUNGAN DENGAN KEJADIAN STUNTING PADA BALITA DI WILAYAH KERJA PUSKESMAS TATAH MAKMUR KABUPATEN BANJAR," *Journal of Midwifery and Reproduction*, vol. 4 no. 1, pp. 13-20, 2020.
- [3] T. Prasetya, I. Ali, C. L. Rohmat and O. Nurdian, "Klasifikasi Status Stunting Balita di Desa Slangit Menggunakan Metode K-Nearest Neighbor," *Informatics for Educators and Professionals*, vol. 5 no. 1, pp. 93-104, 2020.
- [4] I. Wickramasinghe and H. Kalutarage, "Naive Bayes: applications, variations and vulnerabilities: a review of literature with code snippets for implementation," *Soft Computing*, vol. 25, no. 1, pp. 2277-2293, 2021.
- [5] O. Peretz, M. Koren and O. Koren, "Naive Bayes classifier – An ensemble procedure for recall and precision enrichment," *Engineering Applications of Artificial Intelligence*, vol. 136, no. 2, 2024.
- [6] A. Pebdika, R. Herdiana and S. Dodi, "KLASIFIKASI MENGGUNAKAN METODE NAIVE BAYES UNTUK MENENTUKAN CALON PENERIMA PIP," *JATI : Jurnal Mahasiswa Teknik Informatika*, vol. 7, no. 1, 2023.
- [7] M. Y. Titimeidara and W. Hadikurniawati, "Implementasi Metode Naive Bayes Classifier Untuk Klasifikasi Status Gizi Stunting Pada Balita," *Jurnal Ilmiah Informatika*, vol. 11 no. 1, pp. 54-59, 2021.
- [8] R. R. R. Arisandi, B. Warsito and A. R. Hakim, "APLIKASI NAIVE BAYES CLASSIFIER (NBC) PADA KLASIFIKASI STATUS GIZI BALITA STUNTING DENGAN PENGUJIAN K-FOLD CROSS VALIDATION," *Jurnal Gaussian*, vol. 11, pp. 130-139, 2022.
- [9] J. Zeniarja, K. Widia and R. Sani, "Penerapan Algoritma Naive Bayes dan Forward Selection dalam Pengklasifikasian Status Gizi Stunting pada Puskesmas Pandanaran Semarang," *JOINS (Journal of Information System)*, vol. 5 no. 1, pp. 2-9, 2020.
- [10] K. P. Bappenas, "PEDOMAN PELAKSANAAN INTERVENSI PENURUNAN STUNTING TERINTEGRASI DI KABUPATEN/ KOTA," 2018, pp. 1-51.
- [11] S. Hasanah, S. Handayani and I. R. Wilti, "Hubungan Sanitasi Lingkungan Dengan Kejadian Stunting Pada Balita di Indonesia (Studi Literatur)," *Jurnal Keselamatan Kesehatan Kerja Dan Lingkungan*, vol. 2 no. 2, pp. 83-94, 2021.
- [12] Bustami, "PENERAPAN ALGORITMA NAIVE BAYES UNTUK MENGLASIFIKASI DATA NASABAH ASURANSI.," *Jurnal Penelitian Teknik Informatika*, pp. 127-146, 2014.