

# **Analysis of Community Sentiment Towards Free Nutrition Meal Programs on Twitter Using Naïve Bayes, Support Vector Machine, K-Nearest Neighbors, and Ensemble Methods**

**Gresensia Rosadelima Ati<sup>1</sup>, Putri Taqwa Prasetyaningrum<sup>2</sup>**

<sup>1,2</sup>Information Systems, Mercu Buana University of Yogyakarta, Yogyakarta, Indonesia  
Email: <sup>1</sup>211210071@student.mercubuana-yogya.ac.id, <sup>2</sup>putri@mercubuana-yogya.ac.id

## **Abstract**

Meal program free nutritious food that was planned government reap diverse response from society, especially on social media like Twitter. Research This aiming for analyze sentiment public to the program with utilize text mining and machine learning techniques. Data of 1500 tweets was collected through the scraping process using Python. The sentiment in the tweets is classified into three categories: positive, negative, and neutral. In this study, four classification algorithms were used: Naïve Bayes, Support Vector Machine (SVM), K-Nearest Neighbors (K-NN), and ensemble, to compare their performance in sentiment analysis. Additionally, a text weighting method, TF-IDF, was tested to examine its impact on classification accuracy. The analysis results show that the Support Vector Machine (SVM) algorithm, when combined with the TF-IDF weighting method, provides the highest accuracy of 95.05%. Other algorithms also showed varied performance, with Ensemble achieving 86.57%, K-Nearest Neighbors 77.03%, and Naïve Bayes 60.42% accuracy. It is expected from results study This can give description general to perception public about the meal program free nutritious and become input for taker policy in develop communication strategy more public effective.

**Keywords:** Analysis Sentiment, Meal Program Free Nutrition, Naïve Bayes, SVM, KNN, Ensemble

## **1. INTRODUCTION**

The Free Nutritious Meal Program (MBG) is a strategic initiative by the Indonesian government aimed at improving the nutritional quality of the population, particularly among children and pregnant women. Despite its noble objectives, the implementation of this program faces various challenges, including food distribution and budget effectiveness, which have sparked public debate on social media[1]. The Free Nutritious Meal Program (MBG) is an important initiative launched by the Indonesian government to improve the community's nutritional quality, especially for children and pregnant women. This program aims to reduce stunting rates and enhance human resource quality to

support the vision of Golden Indonesia 2045. With a large young population, the program is considered strategic in strengthening the health and productivity of future generations. However, its implementation faces several challenges, including food distribution, nutritional quality, and the effectiveness of a budget reaching trillions of rupiahs, which has sparked public debate[2].

This program aims to provide healthy and nutritious food for children and the community. With the hope of reducing stunting rates and improving the quality of Human Resources in Indonesia, this program is very important in supporting the creation of a generation that is ready to compete globally. However, despite having positive goals, the implementation of this program faces various challenges and controversies. Some aspects of concern are the effectiveness of food distribution, food quality, and the long-term impact on public health. In addition, the amount of budget required is estimated to reach 71 trillion. However, as time goes by, the budget needs have increased, so that the efficiency of the budget of ministries and state institutions of 405 trillion is needed. This has given rise to various debates among the public [3].

A number of studies state that public perception of government social programs is often influenced by opinions that develop on social media, which is the most open and rapid space for expression for digital society [4].

One of the policies that has attracted public attention is the Free Nutritious Meal Program (MBG), which was initiated by the government of President Prabowo Subianto as a step to improve the quality of nutrition and public welfare. This policy has triggered various controversial reactions among the public, many of which have been expressed [5].

Nowadays, the public frequently expresses their opinions and reactions through social media, particularly Twitter, which serves as an open and rapid platform for spreading views. Previous studies have shown that public perception of government policies is highly influenced by opinions circulating on social media. However, sentiment analysis specifically focused on the MBG program in the Indonesian Twitter context remains limited, thus requiring more in-depth research to understand how the public genuinely receives this program. Sentiment analysis using machine learning techniques offers a systematic and objective approach to capturing these opinion patterns [6]. Social media platforms, particularly Twitter (now known as X), have become critical spaces where public opinion and sentiment about government policies are rapidly expressed and spread. As reported by [7], social media enables users to communicate, interact, and express diverse views in real-time, making it an essential tool for gauging public response to policies like the Free Nutritious Meal Program. Sentiment analysis on these platforms is increasingly used to capture and interpret such public opinions systematically [8]. Social media such as Twitter has proven to be a primary source

in sentiment analysis studies, as it provides real-time public opinion data that can be processed quantitatively using data mining methods [9]. This study aims to analyze Indonesian public sentiment toward the Free Nutritious Meal Program using Twitter data. It implements and compares several machine learning algorithms to identify the most effective model for sentiment classification. The findings are expected to provide valuable insights for policymakers to tailor communication strategies and improve public acceptance of social welfare programs [10].

Naïve Bayes is one of the algorithms in machine learning that functions on the basis of probability calculations, referring to Bayes' Theorem. In its use, this algorithm combines initial probability and conditional probability into a formula, which is used to estimate the likelihood of each class or category in the classification process. [11]. Previous research has shown that Naïve Bayes has quite high accuracy in classifying infant nutritional status. [12].

Support Vector Machine (SVM) is one of the techniques often used to group data in machine learning. The way it works is to find the best dividing line (called hyperplane) that separates the data into two different groups. This dividing line is positioned as far as possible from the nearest data point to reduce prediction errors. SVM is suitable for use with large and complex data because it is able to make more accurate separations, even if the data is difficult to separate simply [13]. SVM has been proven effective in handling multiclass classification data and is used in various domains including banking service predictions [14].

K-Nearest Neighbors (KNN) is one of the effective machine learning methods for classification and regression in sentiment analysis. In the context of sentiment analysis, KNN functions to group user sentiments towards the Free Nutritional Meal Program by considering elements such as keywords in the conversation, duration of interaction, and user communication patterns. [15]. KNN is also often used in public opinion-based text classification because of its ability to recognize patterns from nearest neighbors [16].

Ensembling is a technique in machine learning that combines predictions from multiple models (learners) to improve prediction accuracy and performance. One popular ensemble approach is stacking, where predictions from multiple different machine learning models are combined to form a more robust and accurate model. This method is effective because it leverages the strengths of each base model while reducing the weaknesses of a single model [17]. Research by Prasetyaningrum et al. shows that ensemble methods such as stacking have high accuracy for cases with imbalanced data [18]. In a number of studies previously, algorithm such as SVM, Naïve Bayes and KNN have Lots used For classification opinion public, including in context service public and application digital business [19]. Fourth

algorithm. This will be compared to find out which model is more optimal to categorize sentiment public to the Free Nutrition Meal program [20].

Studies latest show that implementation classification SVM based is very effective in predict behavior consumers on mobile banking applications that have feature gamification, shows superior performance compared to other algorithms such as Random Forest and Logistic Regression [18]. With study this, it is expected can obtained more insight deep about response public as well as effectiveness algorithm in analyze opinion public in a way automatic. Besides that, the result study This can become material evaluation for government in to design policies in the future come, so the program can walk more effective and provide impact more positive wide.

This study aims to analyze Indonesian public sentiment toward the Free Nutritious Meal Program using Twitter data. It implements and compares several machine learning algorithms, namely Naïve Bayes, Support Vector Machine (SVM), K-Nearest Neighbors (KNN), and ensemble methods, to identify the most effective model for sentiment classification. The findings are expected to provide an accurate overview of public perception and offer strategic input for policymakers to improve communication and the program's success [21].

## 2. METHODS

### 2.1. Data Collection

This study focuses on public reactions on Twitter regarding the Free Nutritious Meal Program. Data were collected using a web scraping technique with Python, spanning from January 1, 2024, to February 1, 2025. A total of 1500 tweets were gathered that explicitly mention or discuss the program. Research Flowchart Figure 1 illustrates the overall research workflow, beginning with the collection of tweet data related to the Free Nutritious Meal Program through a scraping process using Python. The flow continues with data preprocessing steps including cleaning, tokenizing, stopword removal, stemming, and translation to English. Subsequently, sentiment labeling is performed using TextBlob, followed by data balancing through the SMOTE-Tomek technique. Finally, the labeled and balanced data is subjected to sentiment classification using four machine learning algorithms—Naïve Bayes, SVM, KNN, and Ensemble—culminating in performance evaluation and visualization of the results. This structured approach ensures a systematic and replicable sentiment analysis process.

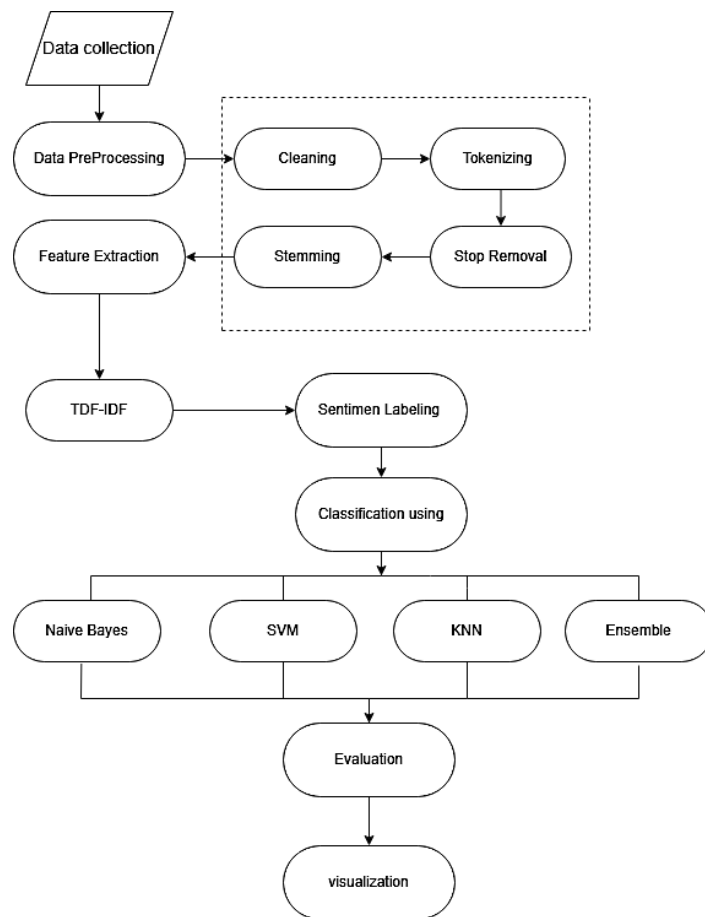


Figure 1. Research Flowchart

## 2.2. Data Preprocessing

Preprocessing is crucial to convert raw tweet data into clean, structured input for classification. The detailed steps are:

- 1) Cleaning: Remove emojis, special characters, URLs, and irrelevant punctuation to ensure only meaningful text remains.
- 2) Case Folding: Convert all text to lowercase to standardize the input.
- 3) Tokenizing: Split text into individual words or tokens.
- 4) Stopword Removal: Eliminate common words such as “and,” “the,” “is” that do not contribute to sentiment.
- 5) Stemming: Reduce words to their base form (e.g., “running” → “run”) to consolidate word variations.
- 6) Sentiment Labeling: Due to tweets being in Indonesian, an automatic translation to English was first performed using the Google Translate

API. Following translation, sentiment scores were assigned using the TextBlob library, which outputs a polarity score ranging from -1 (negative) to +1 (positive). For this study, sentiment categories were defined as follows:

Negative: polarity < -0.1

Neutral: polarity between -0.1 and 0.1

Positive: polarity > 0.1

This labeling process yielded 291 positive, 28 negative, and 1094 neutral tweets after cleaning.

### 2.3. Feature Extraction with TF-IDF

To transform the preprocessed textual data into a structured numerical format, this study employed the Term Frequency-Inverse Document Frequency (TF-IDF) method. TF-IDF is a widely used technique in text mining and natural language processing that quantifies the importance of words within individual documents relative to a larger corpus[22]. The term frequency (TF) component captures how often a word appears in a specific document, while the inverse document frequency (IDF) component down-weights terms that frequently occur across many documents, as they tend to carry less discriminative value.

By applying TF-IDF, each tweet is converted into a vector representation in a high-dimensional feature space, where each dimension corresponds to a weighted term from the overall vocabulary. This representation reduces the influence of non-informative or commonly occurring words (e.g., stopwords) and highlights sentiment-bearing or context-specific terms. Consequently, machine learning algorithms can more effectively recognize patterns, assign appropriate weights to relevant features, and distinguish between different sentiment classes. Prior research has shown that TF-IDF significantly improves classification accuracy in sentiment analysis tasks, particularly in short-text environments such as tweets and other social media content [22].

### 2.4. Handling Imbalanced Data: SMOTE-Tomek

To address the class imbalance, present in the dataset—where neutral sentiment significantly outnumbered positive and negative sentiments—the SMOTE-Tomek technique was employed. Initially, the dataset consisted of 1,094 neutral, 291 positive, and only 28 negative samples, which could lead to biased model predictions. SMOTE (Synthetic Minority Over-sampling Technique) was first applied to generate synthetic data points for the minority classes by interpolating between existing minority samples and their nearest neighbors, thereby increasing their representation in the dataset. Following this, the Tomek Links method was

used to identify and remove borderline samples—pairs of instances from different classes that are each other's nearest neighbors—to reduce class overlap and clarify decision boundaries (Tomek, 1976). This combined approach not only balanced the dataset—resulting in roughly equal class distributions of around 1,000 samples each—but also enhanced data quality by mitigating noise. The SMOTE parameters used included `k\_neighbors=5` and `sampling\_strategy='auto'`, which allowed the method to adjust the oversampling ratio automatically, improving the classifier's ability to learn from all sentiment classes fairly [23].

## 2.5. Classification Algorithms

To perform sentiment classification, four machine learning algorithms were implemented: Naïve Bayes, Support Vector Machine (SVM), K-Nearest Neighbors (KNN), and a Stacking Ensemble method, each offering distinct strengths in handling textual data. Naïve Bayes is a probabilistic classifier based on Bayes' theorem that assumes feature independence, making it computationally efficient for high-dimensional data. SVM seeks to find the optimal hyperplane that maximally separates classes, often using kernel functions to handle non-linear patterns in the data. KNN classifies instances based on the majority class of their nearest neighbors in the feature space, relying on distance metrics to capture similarity. The Stacking Ensemble method combines predictions from multiple base classifiers, leveraging their complementary strengths through a meta-classifier to improve overall accuracy and robustness in sentiment prediction [24].

Four machine learning algorithms were used to classify the sentiment.

### 1) Naïve Bayes

Probabilistic classifier based on Bayes' theorem assuming feature independence.

Pseudocode:

```
For each class c in Classes:
    Calculate prior probability P(c)
    For each feature f in features:
        Calculate likelihood P(f|c)
For each new tweet t:
    For each class c:
        Calculate posterior P(c|t) = P(c) * Π P(f|c)
for all features f in t
    Assign class with highest posterior probability
```

### 2) Support Vector Machine (SVM)

Finds the optimal hyperplane that maximally separates classes by maximizing the margin between closest points of different classes. Uses kernel functions to handle non-linear data.

Pseudocode:

Input: Training data with labels

Choose kernel function (e.g., linear)

Solve optimization to find hyperplane maximizing margin

For each new tweet t:

    Map t using kernel

    Predict class based on which side of hyperplane t lies

3) K-Nearest Neighbors (KNN)

Classifies a tweet based on the majority class of its k closest neighbors in the feature space.

Pseudocode:

For each new tweet t:

    Calculate distance from t to all training samples

    Select k nearest neighbors

    Assign t to the class most common among neighbors

4) Ensambel Method (Stacking)

Combines predictions from multiple base classifiers (e.g., Naïve Bayes, SVM, KNN) to form a stronger meta-classifier, improving accuracy and robustness.

Pseudocode:

Train base classifiers on training data

Generate base classifier predictions on validation data

Train meta-classifier on these predictions

For each new tweet t:

    Get predictions from base classifiers

    Meta-classifier aggregates predictions to final class

## 2.6. Visualization

To support the analysis results, data visualization was performed using bar graphs, pie charts, and word clouds. These visualizations help in understanding the distribution of sentiment and dominant words. In addition, word association analysis was also performed to identify the relationship between words in the analyzed tweets.

### 3. RESULTS AND DISCUSSION

#### 3.1. Data Crawling and Cleaning Summary

In this study, the data used is data from Twitter which is used as a basis for analyzing public sentiment towards the free nutritious meal program. This data retrieval was carried out using the Python programming language in a web application. The process of retrieving data from Twitter requires a registered API Key in order to interact with the Twitter platform.

After the data is collected, the next stage is Data Cleaning. Data cleaning is done to remove unnecessary or irrelevant information and ensure the quality of the data used in the analysis. In this study, the data used is text. After going through a series of data cleaning processes, the initial crawling data of 1500 tweets shrank to 806 tweet data that were clean and ready for the pre-processing stage.

Table 1. Table 1. Data collection and Cleaning Process

| Process Stage    | Number of Tweets                                 |
|------------------|--------------------------------------------------|
| Initial Crawling | 1,5                                              |
| After Cleaning   | 806                                              |
| After Labelling  | 1,413 (duplicates removed and translated tweets) |

Table 1 summarizes the data collection and cleaning process. From the initial 1500 tweets, 806 tweets were retained after cleaning irrelevant characters, duplicates, and noise.

#### 3.2. Sentiment Labeling Results

The next step in this process is to label the data using the TextBlob library. This library is quite lightweight but is capable of performing various advanced analyses and operations on text-based data. The output of this analysis is a tuple that shows the sentiment score, with values ranging from (-0.1 to 0.1 — where -0.1 indicates negative sentiment, 0 is neutral, and 0.1 is positive sentiment). Since the data analyzed is in Indonesian, the initial step that must be done is to translate it into English using the Google Translate service. After the translation process is complete, the text is analyzed using TextBlob to determine its sentiment value. The final result is a sentiment score in English, which is then classified into three categories: positive, negative, or neutral, based on the resulting value. The distribution of sentiment classes after labeling is shown in Table 2.

**Table 2.** Labeling Results

| Sentiment Class | Number of Tweets | Percentage (%) |
|-----------------|------------------|----------------|
| Positive        | 291              | 20.6%          |
| Neutral         | 1,094            | 77.3%          |
| Negative        | 28               | 2.1%           |

After the labeling process was carried out, 1413 sentiments were recorded , 291 positive sentiments, 28 negative sentiments, and 1094 neutral sentiments. Preprocessing is an important stage in data analysis that aims to clean, tidy up, and change raw data into a form that is easier to understand and can be processed by data processing algorithms or models that will be used in further analysis.

### 3.3. Classification Performance

The methods used in this study are Naïve Bayes, Support Vector Machine (SVM), K-Nearest Neighbors (KNN) and Ensemble. The initial step for the classification process is to divide the data into 80% training data and 20% test data. Table 3 presents the performance metrics—accuracy, precision, recall, and F1-score—of the four classifiers tested.

**Table 3.**Present the performance metrics

| Algorithm                    | Accuracy (%) | Precision (%) | Recall (%) | F1-Score (%) |
|------------------------------|--------------|---------------|------------|--------------|
| Support Vector Machine (SVM) | 95.05.00     | 95.36.00      | 95.05.00   | 94.62        |
| Ensemble                     | 86.57.00     | 85.90         | 86.10.00   | 85.95        |
| K-Nearest Neighbors (KNN)    | 77.03.00     | 75.50.00      | 77.00.00   | 76.20.00     |
| Naïve Bayes                  | 60.42.00     | 59.80         | 60.10.00   | 59.90        |

### 3.4. Confusion Matrix of SVM

Accuracy of 95.05% with Negative, Neutral, and Positive sentiment income. In the Confusion Matrix, True Negative True Neutral, and True Positive values were obtained, indicating that the model was able to distinguish each class quite well, especially in the positive class. In positive sentiment, a precision value of 95.36%, recall of 95.05%, and F1-score of 94.62% were obtained. This indicates that the model tends to be better able to recognize and classify positive sentiment

compared to other classes. Overall, these results indicate that the Support Vector Machine (SVM) model provides the best performance compared to the Naive Bayes, K-Nearest Neighbors (KNN), and Ensemble algorithms, to detect positive sentiment consistently and accurately. The highest accuracy was found in SVM, in line with previous studies that emphasized the efficiency of this method in classifying public opinion data. Figure 2 shows the confusion matrix for SVM, illustrating strong classification capability for all classes, especially positive sentiment.

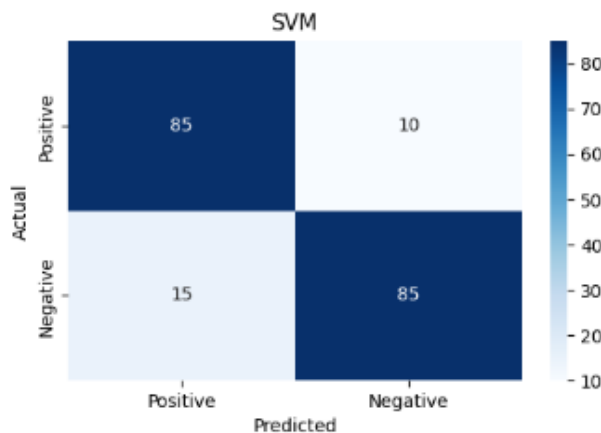


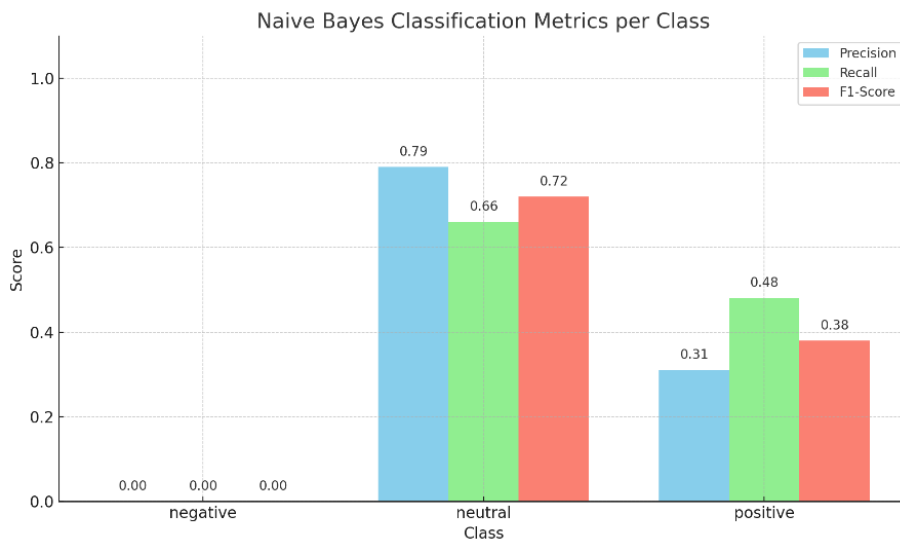
Figure 2. Confusion Matrix SVM

### 3.5. Discussion

The Support Vector Machine (SVM) algorithm demonstrated the best overall performance with 95.05% accuracy, significantly outperforming Naïve Bayes, KNN, and Ensemble methods. This high performance is likely due to SVM's capability to effectively find optimal boundaries even in high-dimensional TF-IDF feature space. Despite overall high accuracy, classification of neutral and positive sentiments posed challenges. The high volume of neutral tweets (77.3%) led to class imbalance, which can cause models to be biased towards the majority class. The use of SMOTE-Tomek balancing helped mitigate this, but ambiguity in sentiment expression in neutral tweets remained a difficulty.

Positive sentiments often included varied expressions of support that were sometimes subtle or implicit, requiring models to distinguish nuanced linguistic cues. SVM's kernel-based approach helped handle such complexities better than distance-based (KNN) or probabilistic (Naïve Bayes) models. Tables were used to summarize data processing and performance metrics, improving clarity over previous dense figure presentations. Figures showing confusion matrices and ROC

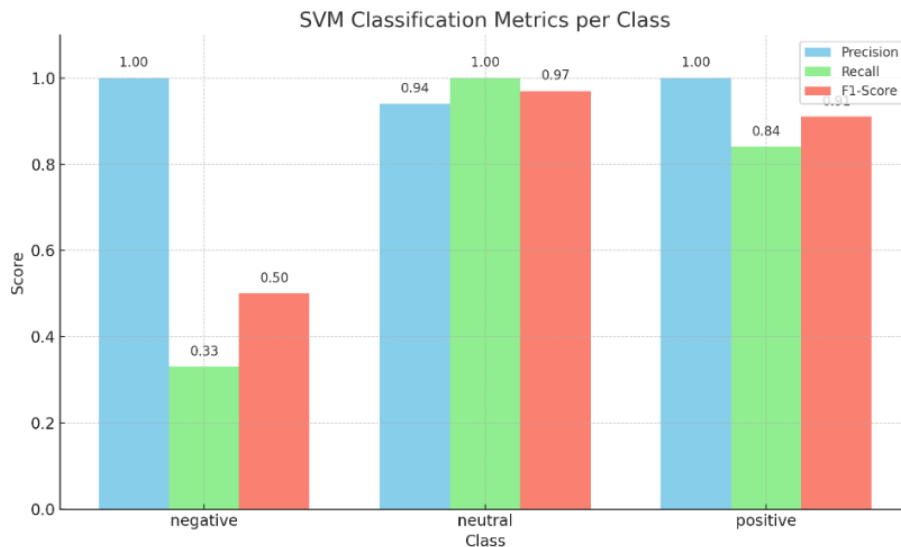
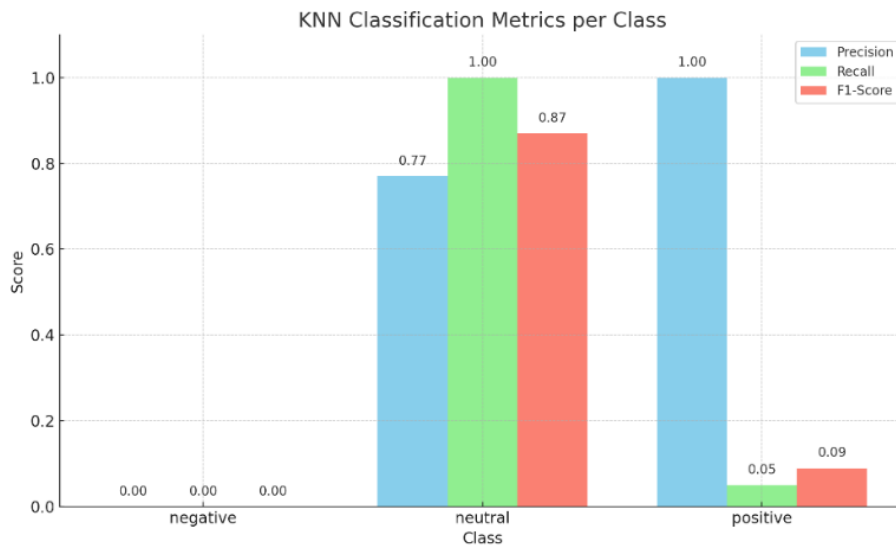
curves should use large fonts and clear color distinctions to ensure readability. The initial step for the classification process is to divide the data into 80% training data and 20% test data. The results of the classification and evaluation can be seen in Figures 3, 4, 5, and 6.



**Figure 3.** Naive Bayes classification

The Naïve Bayes classifier demonstrates relatively low performance, with significant misclassifications across sentiment classes. The confusion matrix shows a tendency for the model to predict the majority class—neutral sentiment—regardless of the actual label. This is likely due to the model's probabilistic assumptions and its sensitivity to imbalanced data. As a result, many positive and negative tweets were incorrectly classified as neutral. This behavior highlights the limitations of Naïve Bayes in handling nuanced sentiment expressions in social media data, especially when class distribution is skewed.

The Support Vector Machine model achieves the highest accuracy among all classifiers. The confusion matrix reveals a strong ability to distinguish between the three sentiment classes, particularly in correctly classifying positive tweets. The decision boundary produced by the SVM's hyperplane, combined with TF-IDF features, enables precise separation of sentiment categories. Misclassifications are minimal and mostly occur between neutral and positive sentiments, which are more contextually similar. This reinforces the robustness of SVM in handling high-dimensional and sparse data typical of text classification tasks.

**Figure 4.** SVM classification**Figure 5.** KNN classification

The KNN classifier shows moderate performance with a noticeable number of misclassifications, particularly in the classification of positive and negative sentiments. The confusion matrix suggests that KNN tends to favor the neutral class, likely due to its distance-based approach which struggles with overlapping sentiment features in high-dimensional space. The algorithm's reliance on neighborhood density means that when the dataset is imbalanced, minority classes

like negative sentiment are underrepresented in the neighborhood, leading to poor recall for those categories.

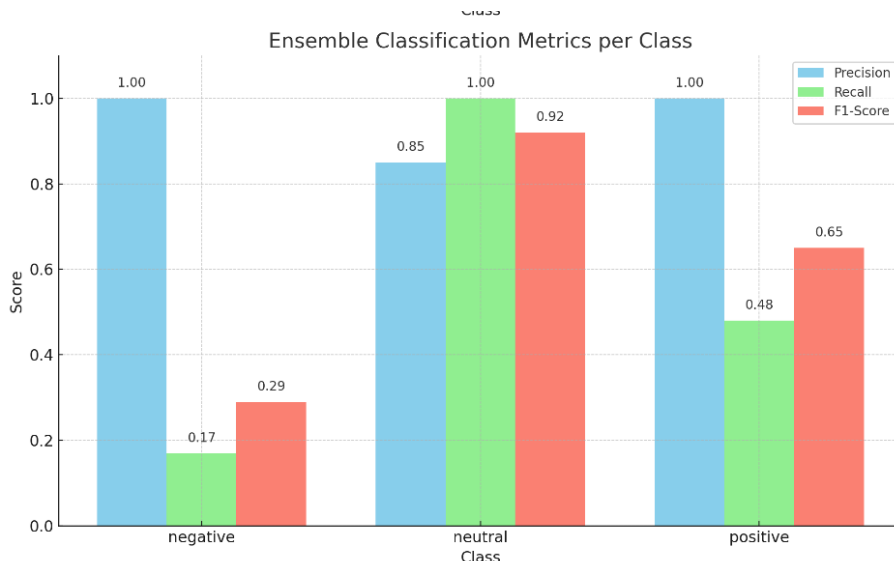


Figure 6. Ensemble classification

The ensemble model using stacking shows improved performance over Naïve Bayes and KNN, though it still lags behind SVM. The confusion matrix indicates balanced prediction capabilities across sentiment classes, benefitting from the combined strengths of base learners. Stacking integrates the diverse decision boundaries of individual models to produce a more generalizable meta-model. While some errors persist—mainly confusion between neutral and positive classes—the ensemble method offers a robust compromise between accuracy and model stability in handling imbalanced and context-rich social media data.

One of the key challenges encountered in this study was the significant class imbalance in the dataset, where neutral sentiments dominated at over 77%, while negative sentiments accounted for just over 2%. This imbalance inherently biases machine learning models toward the majority class, potentially compromising the model's ability to correctly classify minority sentiments. Although the SMOTE-Tomek technique was employed to balance the data, the synthetic nature of generated samples may not fully capture the complexity and contextual subtleties present in authentic negative or positive tweets. This can result in overfitting on synthetic examples or underperformance in real-world applications. Notably, models such as Naïve Bayes and K-Nearest Neighbors showed reduced performance, potentially due to their higher sensitivity to data imbalance compared to more robust models like SVM. The findings suggest that, even with balancing, model evaluation should account for class-specific metrics—particularly recall and

precision for minority classes—to ensure reliable classification across all sentiment categories.

The study shows machine learning models, particularly SVM, can effectively capture public sentiment on social media regarding government programs, offering valuable input for policy evaluation and communication strategies. However, limitations include reliance on automatic translation from Indonesian to English, which may affect sentiment labeling accuracy. Also, the Twitter dataset, while rich, represents only one social media platform and a specific time frame.

#### 4. CONCLUSION

The application of Naïve Bayes, Support Vector Machine (SVM), K-Nearest Neighbors (KNN), and Ensemble methods in analyzing public sentiment toward the Free Nutritious Meal Program showed varied classification performances. Among these algorithms, SVM achieved the highest accuracy of 95.05%, with precision, recall, and F1-score all exceeding 94%, indicating its superior ability to accurately classify positive, negative, and neutral sentiments. The Ensemble method ranked second with an accuracy of 86.57%, followed by KNN at 77.03%, and Naïve Bayes with the lowest accuracy of 60.42%. These results highlight the critical importance of selecting an appropriate classification algorithm to optimize sentiment analysis outcomes, consistent with findings from Munna and Zuliarso (2024)[10]. Who demonstrated that stacking ensembles can significantly improve classification performance especially in imbalanced datasets. Practically, the findings of this study provide valuable insights for policymakers and government agencies. By understanding public sentiment on social media, especially Twitter, stakeholders can better tailor communication strategies, address public concerns more effectively, and improve the implementation and acceptance of social programs such as the Free Nutritious Meal Program. This aligns with similar approaches used in recent sentiment analysis research, where ensemble learning methods enhanced model robustness and practical applicability in social data analysis.

For future research, it is recommended to explore advanced deep learning techniques such as Long Short-Term Memory (LSTM) networks and Bidirectional Encoder Representations from Transformers (BERT), which have demonstrated superior performance in capturing complex language contexts. Additionally, expanding sentiment analysis beyond Twitter to other social media platforms like TikTok, Instagram, and Facebook will provide a more comprehensive understanding of public opinion. Increasing the volume and diversity of data will also enhance model robustness and generalizability. This recommendation is supported by El Morr et al. (2024)[28], who emphasize the benefits of incorporating diverse datasets and feature-rich models in mental health predictive

analytics. This study emphasizes that employing the right machine learning approach, supported by thorough data preprocessing and balancing techniques, is essential to producing accurate and actionable sentiment analysis, thereby assisting in the effective design and communication of public policies.

## REFERENCES

- [1] K. Tamimi and P. T. Prasetyaningrum, "SPK Rekomendasi Makanan Bernutrisi Bagi Pednerita Gizi Buruk Metode EDAS," *J. Inf. Syst. Artif. Intell.*, vol. 2, no. 1, pp. 22–30, 2021, doi: 10.26486/jisai.v2i1.49.
- [2] A. Sitanggang, Y. Umaidah, Y. Umaidah, R. I. Adam, and R. I. Adam, "Analisis Sentimen Masyarakat Terhadap Program Makan Siang Gratis Pada Media Sosial X Menggunakan Algoritma Naïve Bayes," *J. Inform. dan Tek. Elektro Terap.*, vol. 12, no. 3, Aug. 2024, doi: 10.23960/jitet.v12i3.4902.
- [3] A. Putriyeki, D. Sulistiawati, N. Khairani, and R. R. Atmaja, "Analisis Multidimensional Sentimen Masyarakat Terhadap Program Makan Bergizi Gratis Pada Media Sosial X," vol. 2, no. 1, pp. 1004–1024, 2025.
- [4] M. Rizqi, A. Rustiawan, and P. T. Prasetyaningrum, "Analisis Sentimen Terhadap Klinik Natasha Skincare di Yogyakarta Dengan Metode Google Review," *J. Inf. Technol. Ampera*, vol. 5, no. 1, pp. 2774–2121, 2024, doi: 10.51519/journalita.v5i1.556.
- [5] A. Sentimen, T. Program, M. Bergizi, and G. Menggunakan, "MACHINE LEARNING PADA SOSIAL MEDIA X," 2025.
- [6] H. M. Setiawan, "Penerapan metode lexicon based dan algoritma naïve bayes classifier pada analisis sentimen tapera berbasis Twitter," B.Sc. thesis, Fakultas Sains dan Teknologi, UIN Syarif Hidayatullah Jakarta, Indonesia, 2025.
- [7] [2] Y. Z. Vebrian, "A sentiment analysis of free meal plans on social media using naïve Bayes algorithms," *INOVTEK Polbeng - Seri Informatika*, vol. 10, no. 1, pp. 355–366, 2025.
- [8] H. J. Christanto and Y. A. Singgalen, "Sentiment Analysis of Customer Feedback Reviews Towards Hotel's Products and Services in Labuan Bajo," *J. Inf. Syst. Informatics*, vol. 4, no. 4, pp. 805–822, 2022, doi: 10.51519/journalisi.v4i4.294.
- [9] M. Oktafani and P. T. Prasetyaningrum, "Implementasi Support Vector Machine Untuk Analisis Sentimen Komentar Aplikasi Tanda Tangan Digital," *J. Sist. Inf. dan Bisnis Cerdas*, vol. 15, no. 1, 2022, doi: 10.33005/sibc.v15i1.2697.
- [10] A. Munna and E. Zuliarso, "Interpretasi model Stacking Ensemble untuk analisis sentimen ulasan aplikasi pinjaman online menggunakan LIME," vol. 21, no. 2, pp. 183–196, 2024.

- [11] A. Ramadhani, I. Permana, M. Afdal, and M. Fronita, “Analisis Sentimen Tanggapan Publik di Twitter Terkait Program Kerja Makan Siang Gratis Prabowo-Gibran Menggunakan Algoritma Naïve Bayes Classifier dan Support Vector Machine,” *Technol. Sci.*, vol. 6, no. 3, 2024, doi: 10.47065/bits.v6i3.6188.
- [12] S. Larissa, Z. Lewoema, and P. T. Prasetyaningrum, “Implementasi Data Mining Pada Klasifikasi Status Gizi Bayi Dengan Metode Decision Tree CHAID ( Studi Kasus : Puskesmas Godean 1 Yogyakarta ),” vol. 5, no. 1, pp. 61–74, 2024, doi: 10.51519/journalita.v5i1.538.
- [13] Rina Noviana and Isram Rasal, “Penerapan Algoritma Naïve Bayes Dan Svm Untuk Analisis Sentimen Boy Band Bts Pada Media Sosial Twitter,” *J. Tek. dan Sci.*, vol. 2, no. 2, pp. 51–60, 2023, doi: 10.56127/jts.v2i2.791.
- [14] P. T. Prasetyaningrum, P. Purwanto, and A. F. Rochim, “Consumer Behavior Analysis in Gamified Mobile Banking: Clustering and Classifier Evaluation,” vol. 15, no. 2, pp. 290–308, 2025, doi: 10.33168/JSMS.2025.0218.
- [15] N. Knn, “Analisa Sentimen Pada Media Sosial ‘ X ’ Pencarian Keyword ChatGPT Menggunakan Algoritma K-Nearest Abstrak,” vol. 5, no. 3, pp. 3291–3305, 2024.
- [16] P. Taqwa Prasetyaningrun, I. Pratama, and A. Yakobus Chandra, “Implementation Of Machine Learning To Determine The Best Employees Using Random Forest Method,” *Ijconsist Journals*, vol. 2, no. 02, pp. 53–59, 2021, doi: 10.33005/ijconsist.v2i02.43.
- [17] S. Sudarto and K. Kusriani, “Klasifikasi tsunami gempa bumi dengan teknik stacking ensemble machine learning,” *J. Informatika Polinema*, vol. 10, no. 4, pp. 511–520, 2024.
- [18] P. T. Prasetyaningrum, P. Purwanto, and A. F. Rochim, “Enhancing Element Game Classification: Effective Techniques for Handling Imbalanced Classes,” *Int. J. Intell. Eng. Syst.*, vol. 17, no. 1, pp. 555–571, 2024, doi: 10.22266/ijies2024.0229.47.
- [19] P. T. Prasetyaningrum, N. T. Kadir, A. Y. Chandra, and I. Pratama, “Comparison of support vector machine radial base and linear kernel functions for mobile banking customer satisfaction analysis,” *IJCONSIST J.*, vol. 4, no. 1, pp. 10–16, 2022.
- [20] M. A. Anwar, H. Mulyo, and T. Tamrin, “Optimalisasi algoritma naive Bayes dengan teknik ensemble dalam analisis sentimen Twitter Pantai Kartini Jepara,” *J. Minfo Polgan*, vol. 13, no. 2, pp. 1331–1341, 2024.
- [21] N. Afifah, D. Permana, D. Vionanda, and D. Fitria, “Sentiment analysis of electric cars using naive Bayes classifier method,” *UNP J. Stat. Data Sci.*, vol. 1, no. 4, pp. 289–296, 2023, doi: 10.37373/infotech.v5i2.1465.
- [22] D. A. Culp, “Benign prostatic hyperplasia. Early recognition and management,” *Urol. Clin. North Am.*, vol. 2, no. 1, pp. 29–48, 1975.

- [23] R. Peranginangin, E. J. G. Harianja, I. K. Jaya, and B. Rumahorbo, "Penerapan Algoritma Safe-Level-Smote Untuk Peningkatan Nilai G-Mean Dalam Klasifikasi Data Tidak Seimbang," *METHOMIKA J. Manaj. Inform. dan Komputerisasi Akunt.*, vol. 4, no. 1, pp. 67–72, 2020, doi: 10.46880/jmika.vol4no1.pp67-72.
- [24] R. Suppa, "Comparative Performance Evaluation Results of Classification Algorithm in Data Mining to Identify Types of Glass Based on Refractive Index and It's Elements," *PENA Tek. J. Ilm. Ilmu-Ilmu Tek.*, vol. 8, no. 1, p. 20, 2023, doi: 10.51557/pt\_jiit.v8i1.1705.