

Sentiment Analysis of Twitter Towards the Free Lunch Program Using the C4.5 Algorithm

Allamanda Salsabila Hanin¹, Maryam¹

¹Department of Informatics Engineering, Faculty of Communication and Informatics, Muhammadiyah University of Surakarta

Article Info

Article history:

Received Jan 04, 2025

Revised Feb 18, 2025

Accepted March 13, 2025

Keywords:

Sentiment analysis

Twitter

Free lunch program

C.45 algorithm

Classification

ABSTRACT

This study analyzes public sentiment towards the Free Lunch Program on social media X using the C4.5 algorithm. This program, which was initiated as a campaign promise in the 2024 Election, aims to provide free nutritious food for school students in Indonesia. Given the high public interaction on social media, this study was conducted to determine the public response to the program, which can be positive, neutral, or negative sentiment. The methods used include data collection from social media X, text pre-processing, sentiment labeling, application of Term Frequency-Inverse Document Frequency (TF-IDF), and model evaluation with accuracy metrics. The dataset consists of 3,344 tweets which are then classified using the C4.5 algorithm. Based on the evaluation results, it produces an average precision value of 79%, recall of 76%, F1-score of 77%, and is able to provide an accuracy of 78%. Thus, this model shows effective performance in classifying public sentiment. This study can contribute to the use of social media sentiment analysis as a tool for public policy evaluation.

This is an open access article under the [CC BY-SA](#) license.



Corresponding Author:

Allamanda Salsabila Hanin,
Department of Informatics Engineering,
Muhammadiyah University of Surakarta,
A. Yani Street 57162 Makamhaji, Central Java
Email: allamandasalsabilahanin@gmail.com

1. INTRODUCTION

The social media platform X has become a primary medium for the public to express their opinions, both positive and negative, on various issues, including government programs. One program garnering attention is the "Free Lunch Program" initiated by the presidential and vice-presidential candidates Prabowo Subianto and Gibran Rakabuming Raka. This program aims to provide free lunches for school children across Indonesia, yet it has sparked public debate regarding its benefits, risks, and impact on the national budget [1].

Concerns have also arisen among teachers about the potential allocation of part of the BOS (School Operational Assistance) funds to support this program [3]. With the growing influence of social media X, public opinions on this program are often conveyed in text form (tweets). To understand public sentiment with high accuracy, this study utilizes the Decision Tree C4.5 algorithm, renowned for its ability to classify data into an easily interpretable tree structure [9].

Previous studies have demonstrated that the Decision Tree C4.5 algorithm outperforms the Naïve Bayes method in sentiment analysis. Research by Hidayat et al. [11] recorded an accuracy of 73.91% for the Decision Tree compared to 63.77% for Naïve Bayes, while Rahmayanti et al. [2] reported an accuracy of 90% for the Decision Tree, higher than the 85% achieved by Naïve Bayes.

Based on these findings, the Decision Tree C4.5 algorithm has proven superior to Naïve Bayes in the context of sentiment analysis. Similarly, Rahmayanti et al. found that C4.5 is superior in interpretability compared to Support Vector Machines (SVM), but less efficient on larger datasets [2]. These findings further justify the choice of C4.5 for sentiment classification in this study.

The C4.5 algorithm has several advantages, including its ability to handle both numerical and categorical data, as well as its easy-to-interpret decision tree model. It can also process missing values. However, it tends to overfit, requires pruning to reduce complexity, and is less efficient with large datasets. Additionally, C4.5 is sensitive to data imbalance, which can negatively affect its performance.

Therefore, this study employs the Decision Tree C4.5 algorithm to classify tweets on social media X into three categories: positive, neutral, and negative, regarding the free lunch program. This research aims to draw conclusions about public opinion trends and evaluate the accuracy of the Decision Tree C4.5 algorithm in sentiment analysis.

2. RESEARCH METHOD

This study employs a scientific method in data mining and Knowledge Discovery in Database (KDD) to analyze public sentiment toward the free lunch program using the C4.5 algorithm. The stages include data collection, sentiment labeling, preprocessing, TF-IDF application, evaluation, and result visualization.

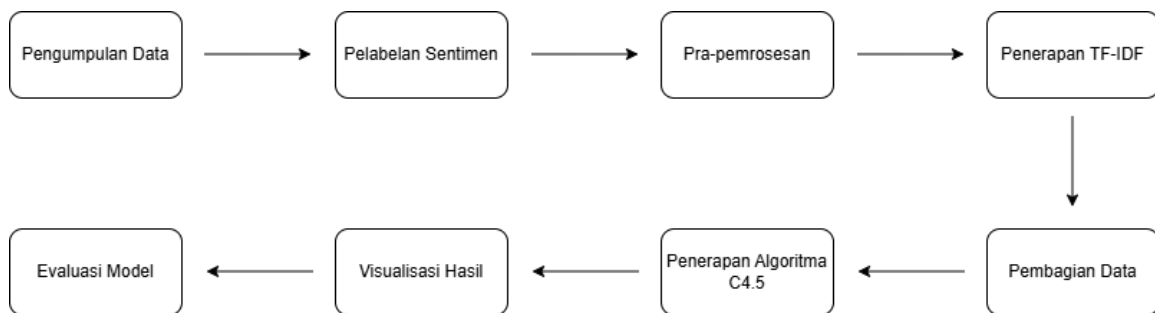


Figure 1. Research Stages

Data were collected from social media platform X using keywords such as #MakanSiangGratis and #ProgramPresiden. The information gathered includes tweet text, timestamp, location, number of retweets, and likes. A total of 3,344 tweets were stored in CSV format for further analysis [10].

Table 1. Data Sample

<i>conversation id str</i>	<i>created at</i>	<i>favorite count</i>	<i>full text</i>
1847292883310874827	Fri Oct 18 15:16:10 +0000 2024	0	“@Puspen_TNI min ksh tau pak prabowo kl program makan siang gratis mending micro managementnya dihandle sm ibuk2 bhayangkari yg pangkatny belum perwira. supaya supaya rapih & teratur karna ini amanah. sehingga anggaran & manfaat yg diterima peserta didik lebih sesuai harapan. saran aja”

Sentiment labeling involves assigning labels to data based on the text content. Tweets were labeled as positive, negative, or neutral sentiments using a lexicon-based dictionary containing a list of words with specific polarity scores. This process matches the words in the tweets with the dictionary entries, and the total score determines the sentiment polarity.

Kamus Leksikon Sentimen Negatif			Kamus Leksikon Sentimen Positif		
	A	B		A	B
1	word	weight	1	word	weight
2	putus tali g	-2	2	hai	3
3	gelebah	-2	3	merekam	2
4	gobar hati	-2	4	ekstensif	3
5	tersentuh (	-1	5	paripurna	1
6	isak	-5	6	detail	2
7	larat hati	-3	7	pernik	3
8	nelangsa	-3	8	belas	2
9	remuk reda	-5	9	welas	4
10	tidak segar	-2	10	kabung	1
11	gemar	-1	11	rahayu	4
12	tak segan	-1	12	maaf	2
13	sesal	-4	13	hello	2
14	pengen	-2	14	promo	3
15	penghayati	-2	15	terimakasih	5

Figure 2. Negative & Positive Sentiment Lexicon Dictionary

To enhance accuracy, the results of automated labeling were manually verified, particularly in cases where contextual inconsistencies were found. Updates to the dictionary were made to ensure consistency in the labeling process.

Table 2. Sentiment Labeling

Full_text	Sentimen
“@Puspen_TNI min ksh tau pak prabowo kl program makan siang gratis mending micro managementnya dihandle sm ibuk2 bhayangkari yg pangkatny belum perwira. supaya supaya rapih & teratur karna ini amanah. sehingga anggaran & manfaat yg diterima peserta didik lebih sesuai harapan. saran aja”	positif

Data preprocessing aims to clean raw data to make it more suitable for analysis. This stage is crucial for reducing noise in the data and ensuring more accurate analysis results. The steps in this stage include:

1) Converting Text to Lowercase

Text is converted to lowercase to avoid variations caused by capitalization, such as "Makan" and "makan".

Table 3. Lowercased Text Results

Full_text	Cleaned_text
“@Puspen_TNI min ksh tau pak prabowo kl program makan siang gratis mending micro managementnya dihandle sm ibuk2 bhayangkari yg pangkatny belum perwira. supaya supaya rapih & teratur karna ini amanah. sehingga anggaran & manfaat yg diterima peserta didik lebih sesuai harapan. saran aja”	“@puspen_tni min ksh tau pak prabowo kl program makan siang gratis mending micro managementnya dihandle sm ibuk2 bhayangkari yg pangkatny belum perwira. supaya supaya rapih & teratur karna ini amanah. sehingga anggaran & manfaat yg diterima peserta didik lebih sesuai harapan. saran aja”

2) Removing URLs and Mentions

This step removes URLs and mentions (@username) that frequently appear in data from X. URLs and mentions typically do not carry relevant sentiment meaning, so they are removed to focus on the main text [14].

Table 4. Removal of URLs and Mentions from Text

<i>Full_text</i>	<i>Cleaned_text</i>
“@Puspen_TNI min ksh tau pak prabowo kl program makan siang gratis mending micro managementnya dihandle sm ibuk2 bhayangkari yg pangkatny belum perwira. supaya supaya rapih & teratur karna ini amanah. sehingga anggaran & manfaat yg diterima peserta didik lebih sesuai harapan. saran aja”	“min ksh tau pak prabowo kl program makan siang gratis mending micro managementnya dihandle sm ibuk2 bhayangkari yg pangkatny belum perwira. supaya supaya rapih & teratur karna ini amanah. sehingga anggaran & manfaat yg diterima peserta didik lebih sesuai harapan. saran aja”

3) Noise Removal

Noise removal eliminates irrelevant elements such as punctuation, numbers, and special symbols to minimize distractions in the data.

Table 5. Noise Removal Results

<i>Full_text</i>	<i>Cleaned_text</i>
“@Puspen_TNI min ksh tau pak prabowo kl program makan siang gratis mending micro managementnya dihandle sm ibuk2 bhayangkari yg pangkatny belum perwira. supaya supaya rapih & teratur karna ini amanah. sehingga anggaran & manfaat yg diterima peserta didik lebih sesuai harapan. saran aja”	“min ksh tau pak prabowo kl program makan siang gratis mending micro managementnya dihandle sm ibuk bhayangkari yg pangkatny belum perwira supaya supaya rapih amp teratur karna ini amanah sehingga anggaran amp manfaat yg diterima peserta didik lebih sesuai harapan saran aja”

4) Tokenization

The cleaned data is broken down into word units to facilitate processing. Each word is sorted alphabetically, allowing for the identification of the weight of each word in a tweet [7].

Table 6. Tokenization Results

<i>Cleaned_text</i>	<i>Tokens</i>
“min ksh tau pak prabowo kl program makan siang gratis mending micro managementnya dihandle sm ibuk bhayangkari yg pangkatny belum perwira supaya supaya rapih amp teratur karna ini amanah sehingga anggaran amp manfaat yg diterima peserta didik lebih sesuai harapan saran aja”	“[min, ksh, tau, pak, prabowo, kl, program, makan, siang, gratis, mending, micro, managementnya, dihandle, sm, ibuk, bhayangkari, yg, pangkatny, belum, perwira, supaya, supaya, rapih, amp, teratur, karna, ini, amanah, sehingga, anggaran, amp, manfaat, yg, diterima, peserta, didik, lebih, sesuai, harapan, saran, aja]”

5) Stop-word Removal

Stop-word removal is the process of deleting common words that do not provide significant meaning in sentiment analysis, such as "dan," "di," "yang," or similar words. By removing stop-words, the text becomes more focused on words that impact sentiment [4].

Table 7. Stop-word Removal Results

<i>Tokens</i>	<i>stop</i>
“[min, ksh, tau, pak, prabowo, kl, program, makan, siang, gratis, mending, micro, managementnya, dihandle, sm, ibuk, bhayangkari, yg, pangkatny, belum, perwira, supaya, supaya, rapih, amp, teratur, karna, ini, amanah, sehingga, anggaran, amp, manfaat, yg, diterima, peserta, didik, lebih, sesuai, harapan, saran, aja]”	“[min, ksh, tau, pak, prabowo, kl, program, makan, siang, gratis, mending, micro, managementnya, dihandle, sm, ibuk, bhayangkari, yg, pangkatny, perwira, rapih, amp, teratur, karna, amanah, anggaran, amp, manfaat, yg, diterima, peserta, didik, lebih, sesuai, harapan, saran, aja]”

6) Stemming and Lemmatization

Stemming and Lemmatization transform words into their root forms. For instance, the words “berlari,” “lari,” and “berlarian” are processed into the root word “lari.” This step simplifies variations of words that essentially carry the same meaning.

Table 8. Stemming and Lemmatization Results

Stop	stemmed
“[min, ksh, tau, pak, prabowo, kl, program, makan, siang, gratis, mending, micro, managementnya, dihandle, sm, ibuk, bhayangkari, yg, pangkatny, perwira, rapih, amp, teratur, karna, amanah, anggaran, amp, manfaat, yg, diterima, peserta, didik, lebih, sesuai, harapan, saran, aja]”	“[min, ksh, tau, pak, prabowo, kl, program, makan, siang, gratis, mending, micro, managementnya, dihandle, sm, ibuk, bhayangkari, yg, pangkatny, perwira, rapih, amp, atur, karna, amanah, anggar, amp, manfaat, yg, terima, serta, didik, lebih, sesuai, harap, saran, aja]”

7) TF-IDF Application

After preprocessing, the next stage is the application of TF-IDF ("Term Frequency-Inverse Document Frequency"). TF-IDF assigns weights to words based on their frequency in a document and inversely to their frequency across the entire document collection. This algorithm helps identify the most influential words in each analyzed tweet [6].

a. TF-IDF Formula

$$TF - IDF(t, d) = TF(t, d) \times IDF(t)$$

b. TF Formula

$$TF(t, d) = \frac{ft, d}{Nd}$$

Where:

ft, d : frequency of term t in document d

Nd : total number of terms in document d

c. Rumus IDF

$$IDF(t) = \log \frac{N}{ft}$$

Where:

N : total number of documents

ft : number of documents containing term t

8) Data Splitting

In the data splitting stage, the dataset is divided into training data (80%) and testing data (20%). The training data is used to build the classification model, while the testing data evaluates the model's performance. Previous studies have shown that an 80:20 ratio yields the highest accuracy when tested using a confusion matrix for three sentiments (negative, positive, and neutral) [5].

9) Implementation of the C4.5 Algorithm

The C4.5 algorithm is a machine learning technique for building decision trees by selecting the best attribute using gain ratio calculations [8]. This algorithm can handle both numerical and categorical data, involving steps such as setting an attribute as the root, creating branches for each value, and splitting cases into branches [13]. The development of C4.5 includes capabilities to handle missing values, improve pruning, and process recurring data [12].

$$\text{Entropy}(S) = \sum_{i=1}^n -p_i * \log_2 p_i$$

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

10) Model Evaluation

Model evaluation tests the model's performance using testing data with metrics such as accuracy, precision, recall, and F1-score. These metrics are used because each provides different insights into the model's ability to classify correctly [15].

a. Accuracy Formula

Accuracy measures how well the model correctly classifies data, with higher values indicating better model performance [15]. However, accuracy is less effective on imbalanced datasets as it does not reflect the model's ability to handle minority classes. The formula for accuracy is:

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN}$$

b. Precision Formula

Precision measures the proportion of correctly predicted positive cases out of all predicted positive cases. This metric evaluates how accurately the model classifies positive data. In sentiment classification, precision measures how many tweets classified as positive truly have positive sentiment.

$$\text{Precision} = \frac{TP}{TP + FP}$$

c. Recall Formula

Recall evaluates the model's ability to identify actual positive data. This metric is crucial to determine how much positive data is correctly identified, even if the model produces incorrect negative predictions [15].

$$\text{Recall} = \frac{TP}{TP + FN}$$

Where:

TP: True Positive (the number of correctly classified positive data points).

TN: True Negative (the number of correctly classified negative data points).

FP: False Positive (the number of negative data points incorrectly classified as positive).

FN: False Negative (the number of positive data points incorrectly classified as negative).

d. F1-Score Formula

The F1-score is a metric that describes the balance between precision and recall. In classification tasks, the F1-score demonstrates how effectively the model combines precision and recall, providing an overview of its ability to classify opinions on Twitter. The formula is:

$$F1 = 2 \times \frac{\text{Presisi} \times \text{Recall}}{\text{Presisi} + \text{Recall}}$$

e. Results Visualization

Results visualization is the final stage of the research, where model evaluation outcomes are presented using graphs, tables, or diagrams. This visualization aims to simplify the understanding and interpretation of the analysis results.

3. RESULTS AND DISCUSSION

3.1. Data Collection

A total of 3,344 tweets were collected using the previously mentioned keywords, including tweet text, timestamps, retweets, locations, and like counts. This dataset will be analyzed to uncover public sentiment regarding the free lunch program, with the results discussed in this section to provide a deeper understanding.

3.2. Sentiment Labeling

The collected 3,344 data entries from the social media platform X were successfully labeled automatically. The data was categorized into three main sentiment groups:

- 1) Positive: Data reflecting supportive responses or praise for the free lunch program.
- 2) Negative: Data indicating criticism or dissatisfaction with the program.
- 3) Neutral: Data that is informational or does not convey any specific emotion.

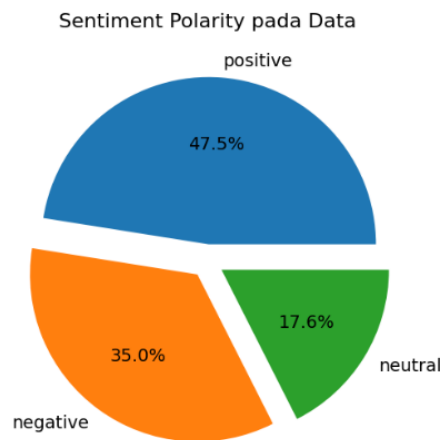


Figure 3. Sentiment Labeling Data Pie Chart

3.3. Results of TF-IDF Application

After sentiment labeling and preprocessing, the TF-IDF method was applied to calculate the word weights in each tweet. The combination of Term Frequency (TF) and Inverse Document Frequency (IDF) resulted in TF-IDF values, indicating the significance of words in the tweets. The TF, IDF, and TF-IDF calculations for a sample of the data are shown in Table 8-10.

Table 8. Term Frequency (TF)

	Kata	Nilai TF
2294	makan	2824
2750	program	2743
705	gratis	2628
679	gizi	2211
2307	makansiinggratis	1682
...	...	...
1438	httpstcopusafoii	1
1434	httpstcoptmoxpqk	1
1433	httpstcopslizvsjk	1
1432	httpstcoprvqhkz	1

3327	zulkifli	1
------	----------	---

[3328 rows x 2 columns]

The Term Frequency (TF) table shows the frequency of word occurrence in each tweet.

Table 9. Inverse Document Frequency (IDF)

	IDF
Aas	8.422075
httpstcozwpoqejsp	8.422075
Httpstcozqntomff	8.422075
Httpstcozunaiatg	8.422075
Httpstcozvcawb	8.422075
...	...
Presidenprabowo	1.703062
Makansiaggratis	1.686889
Program	1.258515
Gratis	1.253880
Makan	1.232907

[3328 rows x 1 columns]

The Inverse Document Frequency (IDF) indicates the importance of each word across the entire collection of tweets.

Table 10. TF-IDF

	aas	abah	abahwhahaha	abal	abidzar	abjodohcomeback	abomination	\
0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	
1	0.0	0.0	0.0	0.0	0.0	0.0	0.0	
2	0.0	0.0	0.0	0.0	0.0	0.0	0.0	
...	...	...	...	...	...	...	...	
3341	0.0	0.0	0.0	0.0	0.0	0.0	0.0	
3342	0.0	0.0	0.0	0.0	0.0	0.0	0.0	
3343	0.0	0.0	0.0	0.0	0.0	0.0	0.0	

	abuqecana	ac	acara	...	yg	young	yovst	yudsky	yuk	\
0	0.0	0.0	0.0	...	9.944174	0.0	0.0	0.0	0.0	
1	0.0	0.0	0.0	...	9.944174	0.0	0.0	0.0	0.0	
2	0.0	0.0	0.0	...	0.000000	0.0	0.0	0.0	0.0	

...	...	...	...	...	...	...	...	...	...
3341	0.0	0.0	0.0	...	0.000000	0.0	0.0	0.0	0.0
3342	0.0	0.0	0.0	...	0.000000	0.0	0.0	0.0	0.0
3343	0.0	0.0	0.0	...	0.000000	0.0	0.0	0.0	0.0

	yungalah	yusuf	zaini	Zhndrew	zulkifli
0	0.0	0.0	0.0	0.0	0.0
1	0.0	0.0	0.0	0.0	0.0
2	0.0	0.0	0.0	0.0	0.0
...	...	...	...	...	...
3341	0.0	0.0	0.0	0.0	0.0
3342	0.0	0.0	0.0	0.0	0.0
3343	0.0	0.0	0.0	0.0	0.0

[3344 rows x 3328 columns]

The TF-IDF table is the result of combining both values, highlighting the significance of a word within a specific document.

3.4. Data Splitting Results

After preprocessing and sentiment labeling, the dataset is split with an 80:20 ratio, where 80% is used for training and 20% for testing. The processed data is used to train the C4.5 model, while the remaining data is used to test the model's performance in classifying tweet sentiments as positive, negative, or neutral. This split allows for the evaluation of the model's performance on unseen data.

Table 11. Dataset Splitting

Jenis Data	Positif	Negatif	Netral	Jumlah Data	Presentase
Data Latih	1324	962	389	2675	80%
Data Uji	351	211	107	669	20%

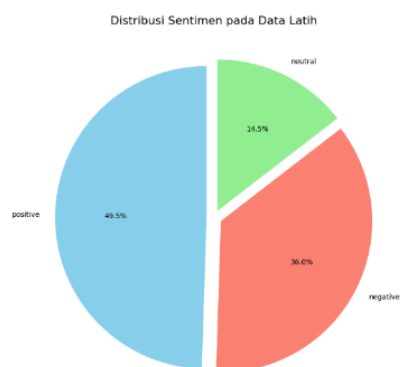


Figure 1. Training Data Pie Chart

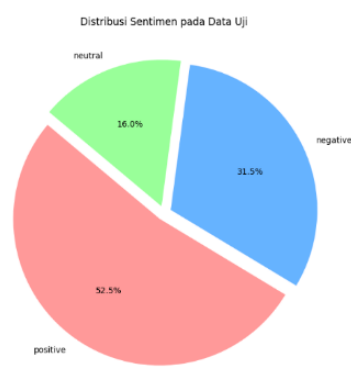


Figure 2. Testing Data Pie Chart

3.5. Model Evaluation Results

This study employs the C4.5 algorithm to analyze public sentiment regarding the Free Lunch Program proposed by a pair of presidential and vice-presidential candidates. Based on the research findings, out of 669 tweets analyzed, 351 contained positive sentiment, 211 contained negative sentiment, and the remaining 107 tweets were neutral. Although the number of positive sentiments is higher, it does not necessarily mean that the program receives full support from the public, as the number of negative sentiments is also significant.

Puad et al. (2024) explained in their study that negative sentiments on social media often stem from the public's overly high expectations. When a program fails to meet public expectations or faces issues during its implementation, public criticism becomes more visible on social media [16]. Conversely, positive sentiments reflect how the government successfully convinces the public of its planned initiatives. In other words, the government has successfully campaigned for the program, as noted in the study by Azhar et al. (2022), which highlighted how public perceptions are heavily influenced by the government's approach to promoting or implementing its policies.

In this study, the C4.5 algorithm achieved an accuracy of 78%, indicating that the model performs reasonably well in classifying sentiments. However, the relatively low recall value of 73% suggests that the model has limitations in identifying all positive sentiments. This finding aligns with the observations of Hidayat et al. (2024), who noted that data imbalance can affect model performance. This imbalance might hinder the algorithm's ability to fully recognize the characteristics of positive sentiments, resulting in the low recall value.

Table 13. Classification

	Precision	Recall	F1-Score	Support
Negative	0.67	0.74	0.70	211
Neutral	0.89	0.73	0.80	107
Positive	0.82	0.81	0.82	351
Accuracy			0.78	669
Macro avg	0.79	0.76	0.77	669
Weighted avg	0.78	0.78	0.78	669

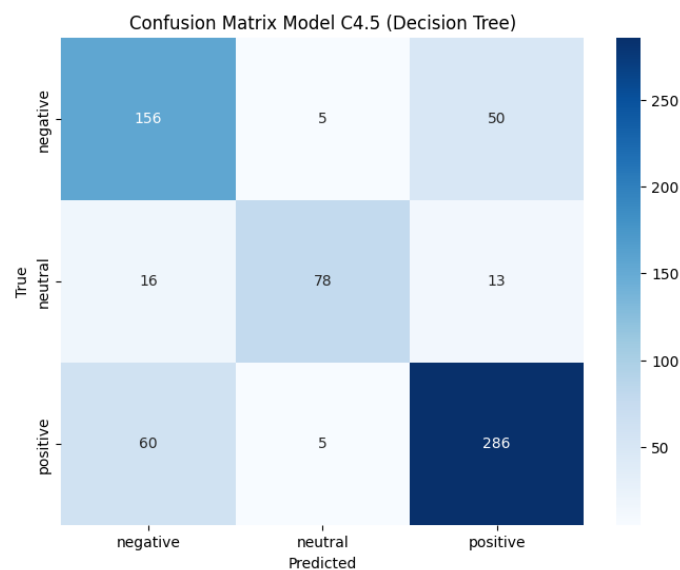


Figure 6. Confusion Matrix

This study also indicates that better communication from the government can help reduce negative sentiments. As noted by Muchram (2024), public perception of government programs is significantly influenced by how these programs are promoted and introduced to the public. If the government improves transparency and enhances public communication regarding the goals and benefits of the Free Lunch Program, it is highly likely that public acceptance of the program will increase [17].

This finding is supported by the research of Wijaya (2021), which states that programs backed by effective communication tend to receive more positive responses from the public. The government should evaluate the implementation of the Free Lunch Program and consider the feedback received from the public [18]. As recommended by Priyanto et al. (2024), program evaluation based on public sentiment can serve as an effective tool to adjust policies to better align with the needs and expectations of the public [19]. By doing so, the program can achieve its objectives more effectively and improve the welfare of vulnerable groups, who are its primary targets.

3.6. Result Visualization

The final stage of the research involved the visualization of results. The image illustrates the comparison of precision, recall, and F1-score values for each sentiment, with an accuracy of 78% (negative, neutral, and positive). Figure 7 contains a graph of sample data by vocabulary, ranging from the highest to the lowest frequency, aligned with the word cloud in Figure 8. In the word cloud, the most frequent vocabulary is displayed with the largest font, and the size decreases according to the number of samples.

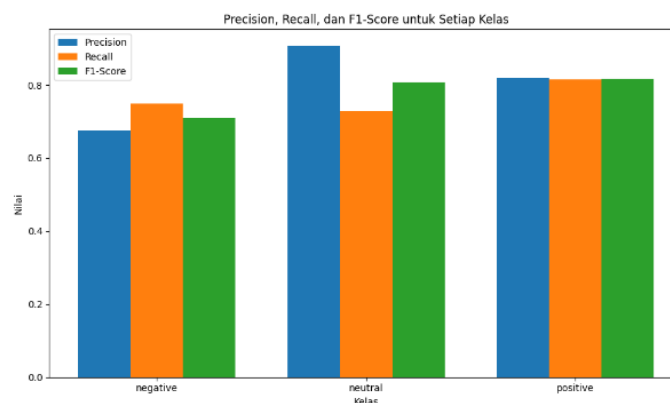


Figure 7. Comparison of Each Class



Figure 8. Word Cloud

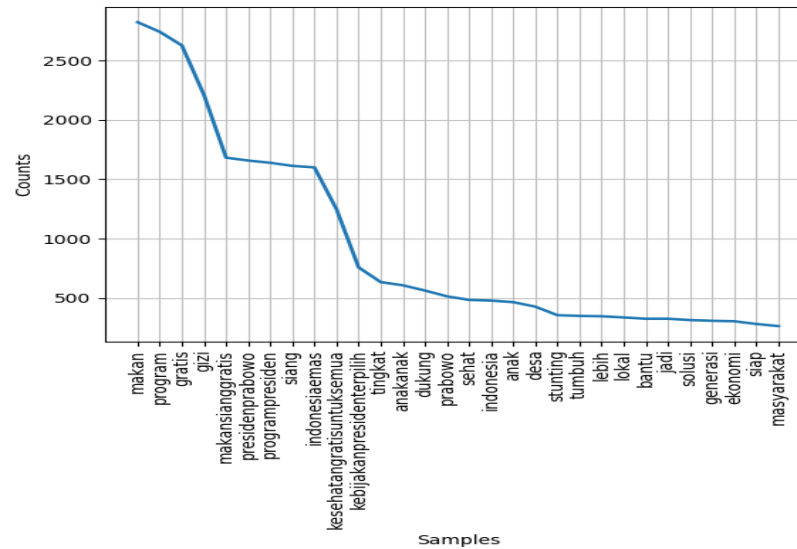


Figure 9. Frequently Occurring Words Count

3.7. Sentiment Analysis Interface Implementation

This research also produced an implementation in the form of an interface to facilitate user interaction with the system. The interface was developed to help the government understand public sentiment towards the free lunch program, enabling them to evaluate the program based on public sentiment on social media X or other platforms.

The sentiment analysis interface was developed using Streamlit in Visual Studio Code (VSCode) and includes two main features. The first feature, Sentiment Analysis, allows users to analyze the sentiment category (positive, neutral, or negative) of a specific tweet based on a dataset embedded in the code. The second feature, Dataset Upload, enables users to analyze a pre-labeled dataset, displaying the number of tweets in each sentiment category through a distribution graph. Streamlit was chosen as the framework due to its ability to build interactive web interfaces simply, without requiring expertise in complex frontend technologies.

This interface leverages Streamlit as the primary framework for interactive interfaces. Data processing is performed using Python with the pandas library for data manipulation and scikit-learn for classification using the C4.5 algorithm. Data visualization is achieved using Matplotlib and Streamlit Charts to present sentiment distribution graphs interactively. The application operates through two main workflows. In the Sentiment Analysis feature, users input a specific text or tweet, which is analyzed using the C4.5 model to determine the sentiment category (positive, neutral, or negative).

The classification results are displayed immediately on the screen. In the Dataset Upload feature, users upload a pre-labeled dataset file, and the application calculates the number of tweets in each sentiment category and presents the results as a bar chart. The process for analyzing a single tweet involves copying the tweet, pasting it into the "Masukkan tweet untuk dianalisis" field, and then clicking the "Submit" button (Figure 10). The results of the analysis are displayed in Figure 11, including the sentiment category (negative, positive, neutral) and the sentiment score in the form of a graph.

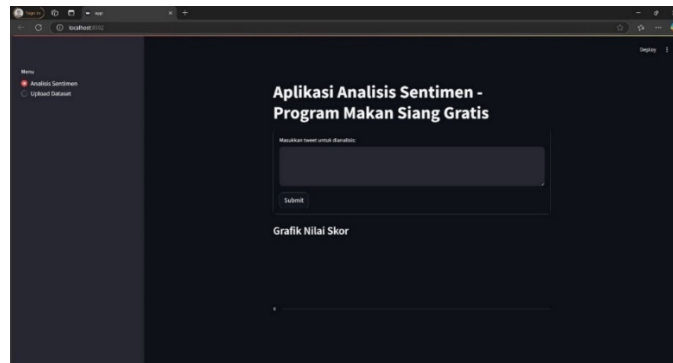


Figure 10. Sentiment Analysis

The main features of the interface include Sentiment Analysis, which allows users to analyze the sentiment of individual tweets and displays the classification results directly on the screen, categorized as positive, neutral, or negative.

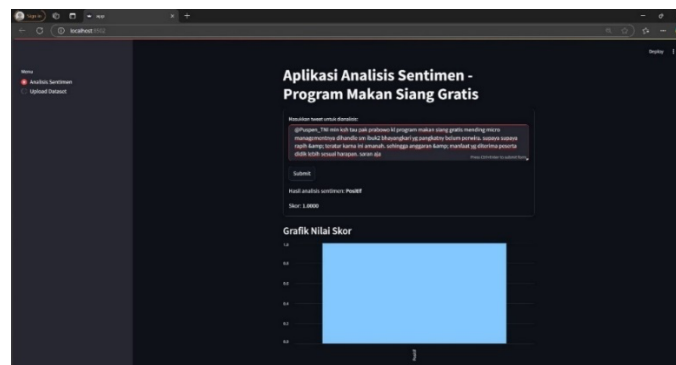


Figure 11. Sentiment Analysis Results

The "Upload Dataset" feature is used to determine the number of each sentiment category from the labeled tweet dataset. The process starts with uploading the dataset in the "Upload file hasil labeling" field (Figure 12). The results are displayed in Figure 13 as a graph showing the number of data for each sentiment category.

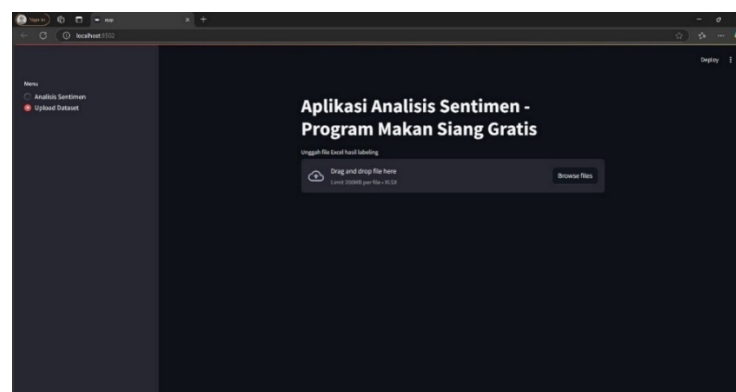


Figure 12. Upload Dataset

The Dataset Upload feature allows users to upload a pre-labeled dataset to visualize the distribution of sentiment categories. It displays the number of tweets in each category as a bar chart.

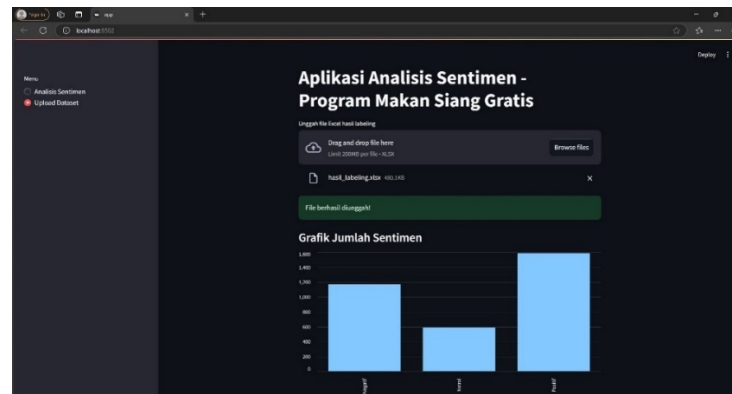


Figure 13. Upload Dataset Result

The Streamlit-based sentiment analysis interface is user-friendly, intuitive, and provides real-time results with features like sentiment graphs and data export. Limitations include support only for pre-labeled datasets and reduced performance on large datasets. Future improvements could add automatic labeling, better models like BERT, and real-time data integration.

4. CONCLUSION

The sentiment analysis of public opinion on social media X regarding the free lunch program using the C4.5 algorithm involved 3,344 data entries themed "Program Makan Siang Gratis." The data was processed through pre-processing steps, including lowercasing, removal of URLs and mentions, noise removal, tokenization, stop-word removal, stemming, and lemmatization. With an 80:20 ratio for training and test data, the prediction results showed 351 positive data, 211 negative, and 107 neutral, with an accuracy of 78%. It can be concluded that the free lunch program received a positive response from the public on social media X. The majority of the public supports this program, as evidenced by the dominance of positive sentiment compared to neutral and negative sentiments.

REFERENCES

- [1] A. Prasetyantoko, "Makan Siang Gratis dan Beban Fiskal," 2024, Mar. 12.
- [2] A. Rahmayanti, L. Rusdiana, and S. Suratno, "Perbandingan Metode Algoritma C4.5 Dan Naive Bayes Untuk Memprediksi Kelulusan Mahasiswa," *Walisongo Journal of Information Technology*, vol. 4, no. 1, pp. 11–22, 2022. <https://doi.org/10.21580/wjit.2022.4.1.9654>
- [3] A. T. Setyawan, "Program Makan Siang Gratis: Memperbaiki Gizi Atau Mengancam Keuangan Negara?" 2024, Jul. 15.
- [4] F. N. Zaman, M. A. Fadhillah, M. A. Ulinuha, and K. Umam, "MENGANALISIS RESPONS NETIZEN TWITTER TERHADAP PROGRAM MAKAN SIANG GRATIS MENERAPKAN NLP METODE NAÏVE BAYES," *Jurnal UMJ*, vol. 14, no. 3, 2024. <https://jurnal.umj.ac.id/index.php/just-it/index>
- [5] I. Vusuvangat, "Penerapan Algoritma C4.5 Untuk Klasifikasi Sentimen Masyarakat Terhadap #RUUKUHP pada Twitter," 2024.
- [6] J. Supriyanto, D. Alita, and A. R. Isnain, "Penerapan Algoritma K-Nearest Neighbor (K-NN) Untuk Analisis Sentimen Publik Terhadap Pembelajaran Daring," *Jurnal Informatika Dan Rekayasa Perangkat Lunak*, vol. 4, no. 1, pp. 74–80, 2023. <https://doi.org/10.33365/jatika.v4i1.2468>
- [7] L. Nursingah, R. Ruuhwan, and T. Mufizar, "ANALISIS SENTIMEN PENGGUNA APLIKASI X TERHADAP PROGRAM MAKAN SIANG GRATIS DENGAN METODE NAÏVE BAYES CLASSIFIER," *Jurnal Informatika Dan Teknik Elektro Terapan*, vol. 12, no. 3, 2024. <https://doi.org/10.23960/jitet.v12i3.4336>
- [8] M. Adriansa, L. Yulianti, L. Elfianty, U. Dehasen Bengkulu, and J. Meranti Raya, "Analisis Kepuasan Pelanggan Menggunakan Algoritma C4.5," 2022.
- [9] N. P. Panjaitan, S. Z. Harahap, and R. M. Ah, "Analisis Minat Masyarakat Menggunakan Media Sosial Menggunakan Algoritma C4.5 dan Metode Naïve Bayes," *INFORMATIKA Universitas Labuhanbatu*, vol. 12, no. 3, 2024.
- [10] R. Azhar, A. Surahman, and C. Juliane, "Analisis Sentimen Terhadap Cryptocurrency Berbasis Python TextBlob Menggunakan Algoritma Naïve Bayes," *Jurnal Sains Komputer & Informatika (J-SAKTI)*, vol. 6, no. 1, 2022.

- [11] R. Hidayat, A. B. Hakim, and R. Nugraha, "Perbandingan Metode Naïve Bayes Dan Decision Tree C4.5 untuk Analisis Sentimen Produk Es Teh Indonesia di Media Sosial Twitter," 2024.
- [12] R. Puspita and A. Widodo, "Perbandingan Metode KKN Decision Tree dan Naive Naves Terhadap Analisi Sentimen Pengguna Layanan BPJS," *Jurnal Informatika Universitas Pamulang*, 2021.
- [13] S. W. Siahaan, K. D. R. Sianipar, and P. P. A. R.H Zer, "Penerapan Algoritma C4.5 Dalam Meningkatkan Kemampuan Bahasa Inggris Pada Mahasiswa," *Petir*, 2020.
- [14] T. A. Dewi and E. Mailoa, "Perbandingan Implementasi Metode Smote pada Agoritma Support Vector Machine (SVM) dalam Analisis Sentimen Opini Masyarakat Tentang Mixue," *Jurnal Indonesia: Manajemen Informatika Dan Komunikasi*, vol. 4, no. 3, pp. 849–855, 2023. <https://doi.org/10.35870/jimik.v4i3.289>
- [15] T. Program, M. Siang, G. Tundo, and D. N. Rachmawati, "Implementasi Algoritma Naive Bayes untuk Analisis Sentimen," *Jurnal Indonesia: Manajemen Informatika dan Komunikasi (JIMIK)*, vol. 5, no. 3, 2024. <https://journal.stmiki.ac.id>
- [16] S. Puad, G. Garno, and A. S. Y. Irawan, "Analisis sentimen masyarakat pada Twitter terhadap Pemilihan Umum 2024 menggunakan algoritma Naïve Bayes," *JATI (Jurnal Mahasiswa Teknik Informatika)*, vol. 7, no. 3, pp. 1560–1566, 2023. DOI: <https://doi.org/10.36040/jati.v7i3.6920>.
- [17] M. Muchram, "Analisis hubungan latar belakang pendidikan, visi misi dan berita media online capres 2024 terhadap pasar modal," *Jurnal Media Akademik (JMA)*, vol. 2, no. 2, 2024. DOI: <https://doi.org/10.62281/v2i2.154>.
- [18] T. N. Wijaya, R. Indriati, and M. N. Muzaki, "Analisis sentimen opini publik tentang Undang-Undang Cipta Kerja pada Twitter," *Jambura Journal of Electrical and Electronics Engineering*, vol. 3, no. 2, pp. 78–83, 2021. DOI: <https://doi.org/10.37905/jjee.v3i2.10885>.
- [19] I. Priyanto, E. M. Dewanti, T. Tundo, M. Nurdin, and R. Kasiono, "Penerapan algoritma metode Naïve Bayes untuk penentuan penerimaan bantuan Program Indonesia Pintar (PIP)," *Jurnal Manajemen Informatika Jayakarta*, vol. 4, no. 2, pp. 162–172, 2024. DOI: <https://doi.org/10.52362/jmijayakarta.v4i2.1355>.