

Sentiment Analysis of Youtube Comments on Indonesian Presidential Candidates in 2024 using Naïve Bayes Classifier Method

Nurbaiti Mahfudza*, Muhammad Ikhsan

Science and Technology, Computer Science, State Islamic University of North Sumatra, Medan, Indonesia

Email: ¹mahfudzanurbaiti@gmail.com ²mhd.ikhsan@uinsu.ac.id

Email Penulis Korespondensi: mahfudzanurbaiti@gmail.com

Submitted 14-04-2025; Accepted 30-04-2025; Published 30-04-2025

Abstract

The 2024 Indonesian presidential election is one of the most talked about topics on various social media platforms, including YouTube. The comments that appear on political-themed videos can reflect public opinion towards presidential candidates. This research aims to conduct sentiment analysis of YouTube comments related to Indonesian presidential candidates in 2024 using the Naïve Bayes Classifier method. This method was chosen due to its ability to classify text data effectively and efficiently. Data was collected from a number of relevant Kompas tv videos on YouTube, then text preprocessing stages such as data cleaning, tokenization, and stemming were performed. Next, the data was classified into three sentiment categories, namely positive, negative, and neutral. The research shows that the Naïve Bayes model is able to classify sentiment with sufficient accuracy. This finding can provide an overview of public perceptions of each presidential candidate as well as input for interested parties in the fields of politics and public communication. The results of this study show that the naïve bayes classifier algorithm can analyze with an accuracy of 61 % in the evaluation process using confusion matrix. The results of this study indicate that the naïve bayes classifier algorithm can be an effective alternative for analyzing the sentiment of YouTube comments on presidential candidates.

Keywords: Sentiment Analysis; Youtube; Presidential Election; Naïve Bayes Classifier; Public Opinion

1. INTRODUCTION

Indonesia is a democratic country where the role of the people is very decisive in choosing the leader of the country which will be carried out through general elections or what is often referred to as elections. A politician who will run for president will certainly check or consider his popularity based on public opinion.[1] With the current technological advances, it has made it easier for people to get information very quickly, the internet is a proof of the development of technology which currently has a lot of influence on society. Since the announcement of the name of the Indonesian presidential candidate for 2024, the names have begun to be widely discussed by the public, especially through social media. One of the social media that can be used by the public to express their opinions is through the social media Youtube[2].

Youtube is one of the platforms or applications available on social media as a means of communication and information that is audio-visual in nature by presenting diverse content. Through Youtube users can watch, save, download, upload, and even share videos with others[3]. The facilities provided by Youtube for users to be able to respond to a video are by giving opinions in the comments column. It is through comments on Youtube that people can express their support, criticism, suggestions and also their hopes for presidential candidates.

Sentiment analysis is a natural language processing technique used to extract and analyze sentiment or opinion from text documents such as news, social media posts, or product evaluations with the aim of detecting the emotions implied in a text. In sentiment analysis, there are three main forms of sentiment: positive, negative and neutral[4]. Sentiment analysis can be used to measure public sentiment or opinion on a particular topic[5]. The results of sentiment analysis can be used to find out the opinions given by the public towards Indonesian presidential candidates in 2024. Sentiment analysis is usually done using machine learning methods such as K-Nearest Neighbors, Support Vector Machine, Naïve Bayes Classifier[6].

The K-Nearest Neighbors method is a method for classifying objects based on the learning data that is closest to the object. The accuracy of K-Nearest Neighbors is greatly affected by choosing the number of K nearest neighbors. K-Nearest Neighbors is easy to interpret and the mathematical calculations are easy to understand, but the prediction phase is slow for larger data sets because accurate distance calculation plays a major role in determining the accuracy of the algorithm[7].

The Support Vector Machine method is a machine learning algorithm that works on the principle of Structural Risk Minimization (SRM) with the aim of finding the best hyperplane that separates two classes and input spaces. Support Vector Machine is a very efficient and relatively easy method to use, but it is difficult to use in large-scale problems[8].

Metode Naïve Bayes Classifier merupakan metode klasifikasi yang berakar pada teorema bayes dengan menggunakan metode probabilitas dan statistik yang dikemukakan oleh ilmuwan Inggris Thomas Bayes yaitu memprediksi peluang di masa depan berdasarkan pengalaman di masa sebelumnya[9]. Kelebihan dari metode Naïve Bayes Classifier adalah tidak memerlukan jumlah data yang banyak dan perhitungannya cepat dan efisien serta memiliki nilai akurasi yang tinggi[10].

In this research, the Naïve Bayes Classifier method is the most suitable method used to determine and classify public opinion into negative sentiment and positive sentiment. The Naïve Bayes Classifier method will produce a technique to find the largest conditional probability value of each class[11]. Various kinds of research have been conducted to compare the capabilities of these algorithms[12]. In research conducted by Septika Sari (2022), it shows

that Naïve Bayes Classifier has a higher accuracy of 74% than Lexiconn Based of 71%. In other research conducted by Imran et al (2024), it shows that Naïve Bayes Classifier has a higher accuracy of 97% than Support Vector Machine of 95%. Previous research only used 2 variables, namely positive and negative variables, therefore in this research based on these problems, the authors decided to take the title “Sentiment Analysis of Youtube Comments on Indonesian Presidential Candidates in 2024 Using the Naïve Bayes Classifier Method” by using 3 supporting variables, namely positive, negative and neutral variables. so that this research is expected to be able to produce a classification of community sentiment towards Indonesian presidential candidates in 2024.[13].

2. RESEARCH METHODOLOGY

Naïve Bayes Classifier is a classification method rooted in classification using probability and statistical methods proposed by British scientist Thomas Bayes, which predicts future opportunities based on previous experience, known as Bayes' theorem. The Naïve Bayes Classifier method works very well compared to other classification models. The main feature of the Naïve Bayes Classifier is the very strong assumption of independence of each condition or event[14]. The data used in this study uses 100 data sets that have been collected through the Kompas TV Youtube channel. Besides that, the literature technique is also carried out by looking for references to printed books, research journals, and websites related to this research, both objects and methods as material for discussing the results of the research.

2.1 Research Framework

In the research framework, a research stage is needed which contains a research model. This is done so that this research is more structured[15]. The stages of this research can be seen in Figure 1 Research Framework

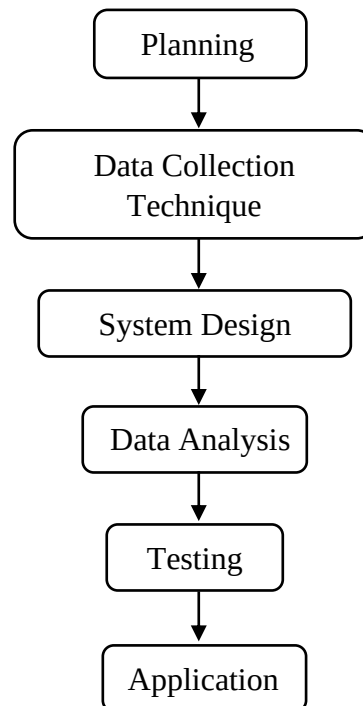


Figure 1. Research Framework

- a. Planing
He initial stage of this research process will begin with planning in determining the topic to be discussed. The topic of this research is about sentiment analysis of Youtube comments on Indonesian presidential candidates in 2024 using the Naïve Bayes Classifier method.
- b. Data Collecction Technique
Data collection is the most important stage in conducting a study to collect the necessary data related to the research to be carried out
- c. System Design
After the Youtube comment data was collected, the researcher designed a system known as a process flowchart
- d. Data Analysis
By presenting the data, it will be easier to examine and organize the data. After the comment data is collected, class labeling will be carried out according to the provisions
- e. Testing

The testing process on the system built has the aim of ensuring that the system runs properly. The testing process on the system built to analyze community sentiment is carried out by collecting all data and proceeding to the preprocessing stage and then stored.

f. Application

The application of this research begins with data collection from Youtube comments regarding public opinion about the event to be classified.

2.2 System Design

After the Youtube comment data was collected, the researcher designed a system known as a process flow diagram. This process flow diagram includes three components, namely input, process and output. Figure 2.2 Flowchart of Research System Design

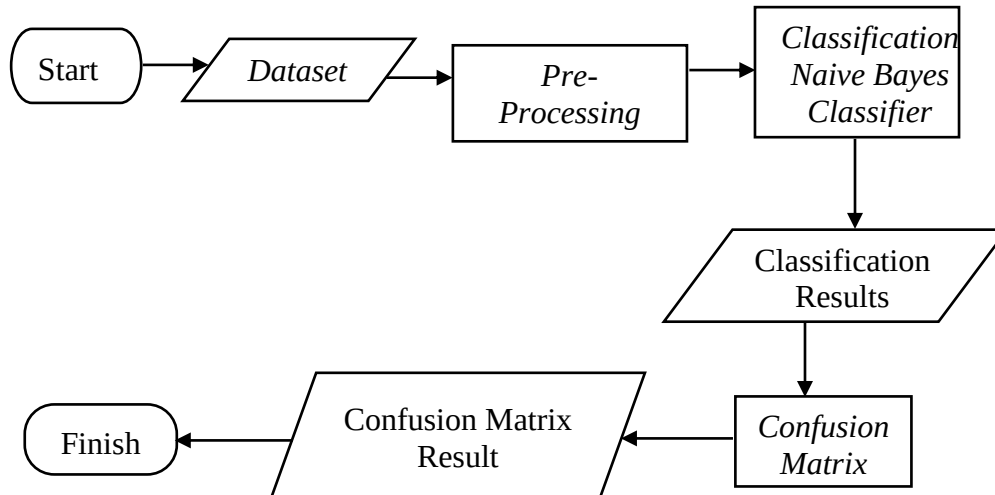


Figure 2. Flowchart of Research System Design

a. Start

This is the first step to start processing

b. Dataset

At the data collection stage in this study, the authors collected data manually, namely collecting each review into a dataset. The variables used in this study are reviews of Kompas TV comments from the YouTube application.

c. Preprocessing

Data preprocessing is a technique used to enhance the algorithm's performance by transforming the unprocessed data into an appropriate form for efficient handling of data by the machine learning algorithm. In this stage the reviews are selected so that they are more structured and the information obtained is clearer. There are several stages in the processing process, namely[16]:

1. Cleaning : In this step, the data will be cleaned or unnecessary characters will be removed, such as punctuation, numbers, URLs, usernames.
2. Case Folding : this stage will change the text to lowercase, remove punctuation, and remove numbers in the review.
3. Normalization: this stage the abbreviated words are changed into words that comply with the KBBI.
4. Tokenizing: this stage the text will be broken down into words.
5. Stopword: this stage filters words that do not have significant meaning in a document.
6. Stemming: at this stage, affixing words in the text will be remove

d. Classification Naive Bayes Classifier

At this stage, the training data that has been trained by the model will be filled with test data where the data has gone through a previous process so that it can be grouped into certain categories based on the words contained in it. grouped into certain categories based on the words contained in it[17].

e. Classification Result

at this stage is obtained from the classification results to find out which class or criteria the tested data is entered into.

f. Confusion Matrix

Testing the accuracy, precision, and recall of the previous results will be carried out using the formula at this stage. The accuracy of the results is measured by the recall, precision, and accuracy values. Where recall (True Positive Rate) is the ratio of true positive identification compared to all true positive data, precision (Positive Predictive Value) is the ratio of true positive identification to all positive identification results, and accuracy is the ratio of true positive identification to all data[18].

Table 1. Confusion Matrix

Actual data	prediction data		
	TRUE	FALSE	TOTAL
True	TP	FN	P
False	FP	TN	N
Totally	P'	N'	P+N

Information :

TP : Correctly classified positive data.

TN : Correctly classified negative data.

FP : Negative data classified as positive.

FN : Positive data classified as negative

g. Confusion Matrix Result

at this stage is done to find out the negative and positive predictions of the data used

h. Finish

End processing and get results in real numbers

3. RESULT AND DISCUSSION

In conducting this research, the researcher faced various challenges that required special attention. Obstacles such as difficulties in collecting sufficient data from various sources, the need to clean and format data before analysis, and challenges in optimal model design were crucial factors that could affect the results. Researchers must also carefully interpret the results to avoid bias and provide useful insights for stakeholders[19]. By managing every aspect of the research carefully and effectively, it is hoped that the results obtained can make a significant contribution to society, as well as add to the understanding of political dynamics through sentiment analysis on digital platforms.[20]

3.1 Data Representation

Data for sentiment analysis of youtube comments on Indonesian presidential candidates in 2024 using the naïve bayes classifier method was collected through crawling using a python library called youtube-comment-downloader/egbertbouman. The data collected is not used immediately, but is pre-processed first, so that it becomes 1158 comment data on presidential and vice presidential candidates. The feedback variable will be used as output, which is categorized into positive and negative comments from presidential candidates and random comments.

Table 2. Dataset

Author	Komentar
@subhan6132	Kelihatan sekali calon yg berkwalitas, punya kapasitas dan memiliki kompetensi, integritas dan layak sbg panutan
@KoparPa r-x8u	Iniharitulisawrangkanbukanakanmukasimykamalmuoksipakahejenbuatsiapakahtipusiapakahringkismoh heheawrangapajalantaaraptodakampanikurupsimuhehesiapakuejentamantamubybysakidsaciamu
@KoparPa r-x8u	WowsiapakurinduGanjarjalanapasalkiprankurinduiapasalrindubuatssaiyasamuatawlawngbuleluuharikabu atkanapasalsakarangmohlakimykamalmukasiasikancintaawrangmukasitokasiapakahtipubuatssudakuciariku BatulawrangmohmuGanjarbuatawragkanmeledudutmuphotomohlacapakbuatangmuhahalumamserta mu
@Falentin usTinu	Kasian pak ganjar
@semkara sander4517	Buktinya ganjar anis kmu kalah pilpres
@jon_8889	Anis cuma pinter ngomong waktu jadi gubernur DKI banyak ngibul tidak sesuai fakta. Tapi sekarang Kabar nya Anis banyak hutang nyungsep ke got jadi gembel politik 🤔😏
@ArlanTimur	1❤️1
@hotibiho tibi9339	Siapa yg menonton vidio ini setelah bapak prabowo terpilih❤️
@JauharA hnafFaliHzzuddin	aku cinta prabowo kiw kiw
@Siti-c8f	Assalamualaikum ys betul harus di selesaikan 😊 abah luar biasa 😊
@Siti-c8f	sabar sabar 😊
@Siti-c8f	Assalamualaikum 😊😊😊

@muham madjalalu ddin2101 Apapun yg terbaik utk memimpin RI sebenarnya hanya Bpk Prof Dr Anes Rasyid Baswedan dan yg terburuk adalah Prabowo.
@Endang 123-h8n Ganjar menghilang.nyentap d telam bumi

3.2 Transform Dataset

In the dataset transformation stage, data that has gone through the pre-processing process is further processed through the tokenization process. Tokenization is done to break down each comment into smaller word units or tokens, so that they can be analyzed more systematically in the modeling process. In addition, each comment is labeled with a sentiment according to the predefined categories, namely positive sentiment, negative sentiment, neutral sentiment and random comments. This labeling is done manually based on the context and meaning of each comment that has been collected. This step aims to ensure that the data used in training the Naïve Bayes Classifier model has a structure suitable for sentiment analysis. With this transformation process, the dataset becomes more structured and ready to be used in the classification process, so that it can improve the accuracy of the model in identifying sentiment patterns in YouTube comments related to Indonesia's presidential and vice presidential candidates in 2024. Contains the results of application implementation or program results (which are important only), or results from method testing.

Figure 3. Transform dataset

Words	Label
['kelihatan', 'sekali', 'calon', 'yg', 'berkwalitas', 'punya', 'kapasitas', 'dan', 'memiliki', 'kompetensi', 'integritas', 'dan', 'layak', 'sbg', 'panutan']	random_Netral
['iniharitulisawrangkanbukanakanmukasimykamalmuoksiapakahejenbuatsiapakahtipusiapaka hringkismohheawrangapajalantaaraptodakampanikurupsimuhehesiapakuejentamantamubyb ysakidsaciamu']	random_Netral
['wowsiapakurinduganjarjalanapasalkiprankurinduiapasalrindubuatsaiyasamuatawlawngbulelu uharikabuatkanapasalsakarangmohlakimykalmukasiasikancintaawrangmukasitokasiapakahti pubuatsudakuciarikubatulawrangmohmuganjarbuatawlangkanmeledudutmuphotomohlacekap kabuatangmuhahalumamsertamu']	ganjar_Netral
['kasian', 'pak', 'ganjar']	ganjar_Netral
['buktinya', 'ganjar', 'anis', 'kmu', 'kalah', 'pilpres']	anis_Netral
['anis', 'cuma', 'pinter', 'ngomong', 'waktu', 'jadi', 'gubernur', 'dki', 'banyak', 'ngibul', 'tidak', 'sesuai', 'fakta', 'tapi', 'sekarang', 'kabar', 'nya', 'anis', 'banyak', 'hutang', 'nyungsep', 'ke', 'got', 'jadi', 'gembel', 'politik', '🤔🤔']	anis_Netral
['1♥1']	random_Netral
['siapa', 'yg', 'menonton', 'vidio', 'ini', 'setelah', 'bapak', 'prabowo', 'terpilih♥']	prabowo_Netral
['aku', 'cinta', 'prabowo', 'kiw', 'kiw']	prabowo_Netral
['assalamualaikum', 'ys', 'betul', 'harus', 'di', 'selesaikan', '😊', 'abah', 'luar', 'biasa', '😊']	random_Netral

3.3 Probability of testing data

After the probability calculation on the training data is completed, the next step is to calculate the probability on the testing data, using the probability results from the training data as a reference to the testing data.

Figure 4. Probability of testing data

No	Words	Feedback
1	['pak', 'anis', 'presiden', '2024', 'amin', 'ya', 'rab']	?
2	['anis', 'cuma', 'omong', 'besar', 'aja', 'bukti', 'zonk']	?
3	['andai', 'aja', 'gibran', 'ama', 'anies', 'udah', 'kompli', 'sama', 'sama', 'jenius']	?
4	['betul', 'itu', 'pak', 'prabowo', 'penentu', 'suara', 'adalah', 'rakyat']	?
5	['setuju', 'sama', 'pak', 'prabowo', 'saya', 'lebih', 'suka', 'nada', 'berbicara', 'pak', 'prabowo']	?

3.4 Model Implementation in Python

This code imports various libraries used in data analysis and visualization as well as building text classification models. The pandas library is used to read and process the data, while sklearn model_selection train_test_split divides the data into two parts, the training set and the testing set. For text representation in numerical form, sklearn feature_extraction text TfidfVectorizer is used, which converts text into Term Frequency-Inverse Document Frequency (TF-IDF) based vectors. The classification model applied is Naive Bayes, implemented through sklearn naive_bayes MultinomialNB, which is specifically used for text classification. Model evaluation is performed using the sklearn metrics library, which provides metrics such as accuracy and confusion matrix. For the purpose of visualizing the analysis results, the matplotlib pyplot and seaborn libraries were used. In addition, for the analysis of words in the dataset, the collections Counter library

is used to count the frequency of occurrence of words in each label, while wordcloud.WordCloud is used to create word cloud visualizations based on the words in the dataset.

```
1 import pandas as pd
2 from sklearn.model_selection import train_test_split
3 from sklearn.feature_extraction.text import TfidfVectorizer
4 from sklearn.naive_bayes import MultinomialNB
5 from sklearn.metrics import classification_report, confusion_matrix, accuracy_score
6 import matplotlib.pyplot as plt
7 import seaborn as sns
8 from collections import Counter
9 from wordcloud import WordCloud
```

Figure 3. Model Implementation in python

After the MultinomialNB() model predicts the test data, the classification results are saved into an Excel file for further analysis. This process is done using the pandas library, where the data containing the original text, the actual label (y_test), and the prediction results are organized into a DataFrame. Then, this DataFrame is exported into Excel format using the .to_excel() function. Storing the classification results in Excel allows users to perform manual inspection, visualization, or statistical analysis of the model's performance.

```
1 conf_matrix = confusion_matrix(y_test, y_pred)
2 accuracy = accuracy_score(y_test, y_pred)
3 error_score = 1 - accuracy
4
5 plt.figure(figsize=(10, 8))
6 sns.heatmap(conf_matrix, annot=True, fmt='d', cmap='Blues', xticklabels=mnb.classes_, yticklabels=mnb.classes_)
7 plt.title('Confusion Matrix Accuracy: {accuracy:.1f}, Error: {error_score:.1f}')
8 plt.xlabel('Predicted')
9 plt.ylabel('Actual')
10 plt.show()
```

Figure 4. Model evaluation with confusion matrix

3.5 Confusion Matrix

To evaluate the performance of the MultinomialNB() model, a confusion matrix is used, which shows the number of correct and incorrect predictions for each class. The confusion matrix is calculated using the confusion_matrix() function of the scikit-learn library, which compares the predicted results of the model with the true label (y_test). To facilitate interpretation, seaborn was used to create a heatmap of the confusion matrix, so that the prediction error patterns could be more easily analyzed visually.

In addition, two main metrics were used to measure model performance, namely accuracy and error rate. Accuracy is calculated as the ratio between the number of correct predictions and the total number of samples, while error rate shows the percentage of prediction errors. A high accuracy indicates that the model is able to perform well in classification, while the error rate gives an idea of the level of error that still occurs.

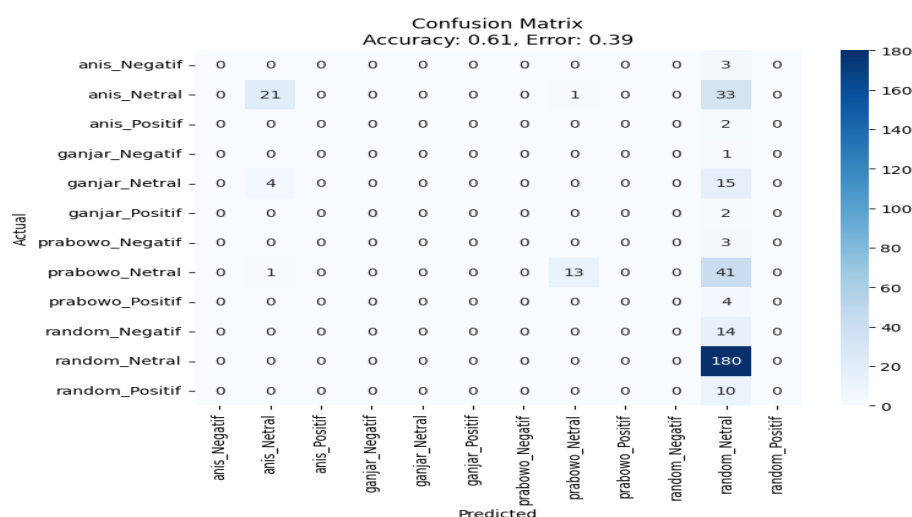


Figure 5. Confusion Matrix

3.6 Evaluation with precision, recall and F1-Score

Untuk mengevaluasi performa model lebih lanjut, dihitung tiga metrik utama: *precision*, *recall*, dan *F1-score* untuk setiap kelas menggunakan fungsi *classification_report()* dari pustaka *scikit-learn*.

- Precision mengukur seberapa banyak prediksi positif yang benar dibandingkan dengan total prediksi positif yang dibuat model.
- Recall menunjukkan seberapa banyak sampel positif yang benar-benar terdeteksi oleh model dibandingkan dengan total sampel positif yang ada.
- F1-score* adalah rata-rata harmonik dari precision dan recall, yang memberikan keseimbangan antara keduanya.

Untuk visualisasi hasil evaluasi, dibuat bar chart menggunakan pustaka *matplotlib* dan *seaborn*. Setiap metrik (*precision*, *recall*, dan *F1-score*) diplot dalam bentuk batang untuk setiap kelas, sehingga memudahkan analisis terhadap performa model pada masing-masing kategori

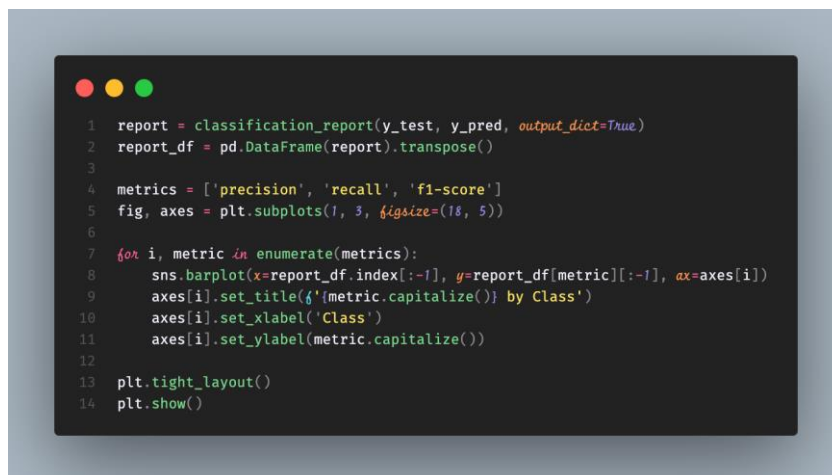


Figure 6. Evaluation with precision, recall and F1-Score

To understand the distribution of data before training the model, the seaborn library was used to display the amount of data by label in the form of a bar chart. This visualization helps in analyzing the class balance in the dataset. If the amount of data in each label is imbalanced (imbalanced dataset), then the model may tend to be more accurate in predicting classes with more data, while the performance may decrease for classes with less data. A bar graph is created using the *sns.countplot()* function, where the horizontal axis represents the label categories, while the vertical axis shows the number of samples in each category. With this visualization, adjustments such as resampling techniques or class weighting can be made if needed to improve the model's balance.

To analyze the most frequently occurring words in each label, the *collections.Counter* library was used to count the frequency of words in each category. The steps are as follows:

- Concatenate the text within each label: All texts in a category are concatenated into a single string.
- Tokenization and text cleaning: The text is cleaned of punctuation, converted to lowercase, and broken down into words (tokens).
- Counting word frequency: Using Counter to get a list of the most frequently occurring words in each label.
- Selecting the top 10 words: The 10 words with the highest frequency from each label are taken.

After getting the list of most frequently occurring words, visualization is done with bar charts using seaborn. Each bar chart represents the frequency of occurrence of words in a label, making it easy to analyze the typical words that frequently appear in a particular category.

To understand the distribution of words in each label, a word cloud visualization with the *wordcloud* library is used. Word clouds display the most frequently occurring words in each category, where the font size represents the frequency of occurrence of the word. The creation process begins by merging text based on labels, then cleaning the text from punctuation, numbers, and common words (stopwords) so that only meaningful words remain. After that, *WordCloud()* from the *wordcloud* library is used to generate a visual representation of the most frequently occurring words. The results are then displayed using *matplotlib* for each label in the dataset. With this visualization, the dominant word patterns in each category can be analyzed more intuitively, thus helping to understand the structure of text data and improving the interpretation of the classification model.



Figure 7. Output word cloud untuk label ganjar_positif

The generation process begins by merging the text based on the labels, and then cleaning the text from punctuation, numbers, and stopwords so that only meaningful words remain. After that, WordCloud() from the wordcloud library is used to generate a visual representation of the most frequently occurring words. The results were then displayed using matplotlib for each label in the dataset. With this visualization, the dominant word patterns in each category can be analyzed more intuitively, thus helping in understanding the structure of the text data and improving the interpretation of the classification model.

3.7 Discussion

Testing in this study was conducted using the blackbox testing method, which is a common approach used in software evaluation to assess the functionality of a system without considering its internal structure. This method focuses on analyzing the extent to which program units can meet the functional requirements that have been set, based on clear and detailed system specifications. The specifications cover various important aspects, including the expected functions, the menu options provided, and the level of compatibility of the model used in this study. Thus, blackbox testing allows researchers to evaluate software from the end-user perspective, ensuring that all features function according to user expectations and needs.

In carrying out this testing, the program that has been developed is executed in various usage scenarios to identify potential problems or deficiencies. This process involves careful observation to ensure that the results obtained are in accordance with predetermined requirements and criteria. Through this approach, it is expected to produce a comprehensive evaluation of system performance and reliability. As part of this report, below is a table showing the system performance metrics measured during the testing process. This table provides a clear picture of the test results and helps in further analysis of the effectiveness and efficiency of the system being tested.

4. CONCLUSION

This research has successfully implemented the naïve bayes classifier method to perform sentiment analysis on YouTube comments about Indonesia's 2024 presidential candidates. Through the stages of data collection, pre-processing, and classification, this research shows that the naïve bayes classifier is a reliable and efficient method for handling large amounts of text data. This method is able to provide an accurate picture of public perception of presidential candidates through sentiment comments on social media, especially YouTube. During the research, it was found that positive and negative comments can be clearly distinguished based on the key words that appear, such as love, win, good, etc. for positive sentiment and worst, corruption, breaking the law, etc. for negative sentiment. The data is then divided into training data and testing data with a ratio of 70:30 to validate the accuracy of the built model. The use of preprocessing techniques such as removal of irrelevant words and converting category variables to numerical form plays an important role in improving the performance of the model. This research successfully proved that sentiment analysis through naïve bayes classifier can be used as an effective tool to monitor public opinion, especially in the political context. However, several challenges were encountered, such as the difficulty in collecting diverse data and the need to thoroughly clean the data to ensure valid and reliable results.

REFERENCES

- [1] A. Y. Simanjuntak, I. S. S. Simatupang, and Anita, "Implementasi Data Mining Menggunakan Metode Naïve Bayes Classifier Untuk Data Kenaikan Pangkat Dinas," *J. Sci. Soc. Res.*, vol. 4307, no. 1, pp. 85–91, 2022.
- [2] Rayuwati, Husna Gemasih, and Irma Nizar, "IMPLEMENTASI ALGORITMA NAIVE BAYES UNTUK MEMPREDIKSI TINGKAT PENYEBARAN COVID," *Jurnal Ris. Rumpun Ilmu Tek.*, vol. 1, no. 1, pp. 38–46, 2022, doi: 10.55606/juritek.v1i1.127.
- [3] J. Sihombing, "Klasifikasi Data Antropometri Individu Menggunakan Algoritma Naïve Bayes Classifier," *BIOS J. Teknol. Inf. dan Rekayasa Komput.*, vol. 2, no. 1, pp. 1–10, 2021, doi: 10.37148/bios.v2i1.15.

- [4] S. R. Cholil, T. Handayani, R. Prathivi, and T. Ardianita, "Implementasi Algoritma Klasifikasi K-Nearest Neighbor (KNN) Untuk Klasifikasi Seleksi Penerima Beasiswa," *IJCIT (Indonesian J. Comput. Inf. Technol.*, vol. 6, no. 2, pp. 118–127, 2021, doi: 10.31294/ijcit.v6i2.10438.
- [5] W. S. Dharmawan, "INFORMATIKA Dalam Prediksi Penyakit Jantung," *J. Inform. Manaj. dan Komput.*, vol. 13, no. 2, pp. 31–41, 2021.
- [6] B. Delvika, N. Abror, and U. R. Gurning, "Perbandingan Algoritma NBC dan C4. 5 Dalam Analisa Sentimen Pemilihan Presiden 2024 Pada Twitter: Comparison of the NBC and C4. 5 Algorithms in Sentiment ...," *SENTIMAS Semin. Nas. ...*, pp. 41–48, 2023, [Online]. Available: <https://journal.irpi.or.id/index.php/sentimas/article/view/548%0Ahttps://journal.irpi.or.id/index.php/sentimas/article/download/548/336>
- [7] S. Adelia, E. Milanda, J. Santari, D. T. Kesuma, E. Silvia, and F. Kurniawan, "Analisis Sentimen Belajar Programming Pada Media Sosial Youtube Menggunakan Algoritma Klasifikasi Naive Bayes," *J. Inf. Technol. Ampera*, vol. 4, no. 3, pp. 254–264, 2023, [Online]. Available: <https://journal-computing.org/index.php/journal-ita/article/view/430>
- [8] G. Rininda, I. Hartami Santi, and S. Kirom, "Penerapan Svm Dalam Analisis Sentimen Pada Edlink Menggunakan Pengujian Confusion Matrix," *JATI (Jurnal Mhs. Tek. Inform.*, vol. 7, no. 5, pp. 3335–3342, 2024, doi: 10.36040/jati.v7i5.7420.
- [9] S. F. Pane and M. S. Amrullah, "Systematic Literature Review: Analisa Sentimen Masyarakat terhadap Penerapan Peraturan ETLE," *J. Appl. Comput. Sci. Technol.*, vol. 4, no. 1, pp. 65–74, 2023, doi: 10.52158/jacost.v4i1.493.
- [10] Endrik, A. Nugroho, and A. T. Zy, "Penerapan Algoritma Naive Bayes dan PSO pada Analisis Sentimen Kandidat Calon Presiden 2024," *Remik Ris. dan E-Jurnal Manaj. Inform. Komput.*, vol. 7, no. 3, pp. 1367–1380, 2023.
- [11] A. R. Abdillah and F. N. Hasan, "Analisis Sentimen Terhadap Kandidat Calon Presiden Berdasarkan Tweets Di Sosial Media Menggunakan Naive Bayes Classifier," *Smatika J.*, vol. 13, no. 01, pp. 117–130, 2023, doi: 10.32664/smatika.v13i01.750.
- [12] M. K. Insan, U. Hayati, and O. Nurdiawan, "Analisis Sentimen Aplikasi Brimo Pada Ulasan Pengguna Di," *J. Mhs. Tek. Inform.*, vol. 7, no. 1, pp. 478–483, 2023.
- [13] M. F. Londjo, "Implementasi White Box Testing Dengan Teknik Basis Path Pada Pengujian Form Login," *J. Siliwaangi*, vol. 7, no. 2, pp. 35–40, 2021.
- [14] K. S. Putri, I. R. Setiawan, and A. Pambudi, "Analisis Sentimen Terhadap Brand Skincare Lokal Menggunakan Naive Bayes Classifier," *Technol. J. Ilm.*, vol. 14, no. 3, p. 227, 2023, doi: 10.31602/tji.v14i3.11259.
- [15] I. Nursyamsi and A. Momon, "Analisa Pengendalian Kualitas Menggunakan Metode Seven Tools untuk Meminimalkan Return Konsumen di PT. XYZ," *J. Serambi Eng.*, vol. 7, no. 1, pp. 2701–2708, 2022, doi: 10.32672/jse.v7i1.3878.
- [16] D. ARIFIN, "Universitas negeri medan," *Temat. Univ. Negeri Medan*, vol. 11, no. 1, pp. 26–36, 2012, [Online]. Available: <https://jurnal.unimed.ac.id/2012/>
- [17] Mustika et al., *Data Mining dan Aplikasinya*. 2021. [Online]. Available: <https://repository.penerbitwidina.com/uk/publications/351768/data-mining-dan-aplikasinya>
- [18] A. F. Setyaningsih, D. Septiyani, and S. R. Widiyari, "Implementasi Algoritma Naive Bayes untuk Analisis Sentimen Masyarakat pada Twitter mengenai Kepopuleran Produk Skincare di Indonesia," *J. Teknol. Inform. dan Komput.*, vol. 9, no. 1, pp. 224–235, 2023, doi: 10.37012/jtik.v9i1.1409.
- [19] H. H. Zain, R. M. Awannga, and W. I. Rahayu, "Perbandingan Model Svm, Knn Dan Naive Bayes Untuk Analisis Sentiment Pada Data Twitter: Studi Kasus Calon Presiden 2024," *JIMPS J. Ilm. Mhs. Pendidik. Sej.*, vol. 8, no. 3, pp. 2083–2093, 2023, [Online]. Available: <https://jim.usk.ac.id/sejarah>
- [20] A. Setiawan, A. T. Prastowo, and D. Darwis, "Sistem Monitoring Keberadaan Posisi Mobil Berbasis Gps Dan Penyadap Suara Menggunakan Smartphone," *J. Tek. dan Sist. Komput.*, vol. 3, no. 1, pp. 35–44, 2022, doi: 10.33365/jtikom.v3i1.1644.