

Analyzing Public Sentiment Toward Probolinggo UMKM Products on TikTok Through Naïve Bayes Classification and Data Visualization

Umi Diantika Susilowati
Informatics

Hafshawaty Zainul Hasan University
Probolinggo, Indonesia
Umidiantika05@gmail.com

Anisa Nurul Wilda
Informatics

Hafshawaty Zainul Hasan University
Probolinggo, Indonesia
Anisanurulwilda0@gmail.com

Muhammad Ichsan
Informatics

Hafshawaty Zainul Hasan University
Probolinggo, Indonesia
ichsan29061997@gmail.com

Syifa Ayu Mika Bahrul
Informatics

Hafshawaty Zainul Hasan University
Probolinggo, Indonesia
Syifaayuvia@gmail.com

Wahyu Nofiyani Hadi
Informatics

Hafshawaty Zainul Hasan University
Probolinggo, Indonesia
navoleo7@gmail.com

Moh. Fadel
Informatics

Hafshawaty Zainul Hasan University
Probolinggo, Indonesia
Cuexsade007@gmail.com

Rojil Ghufon
Informatics

Hafshawaty Zainul Hasan University
Probolinggo, Indonesia
rojilghufon77@gmail.com

Arya Dwi Nugraha
Informatics

Hafshawaty Zainul Hasan University
Probolinggo, Indonesia
aryadwinugraha555@gmail.com

Sischa Wahyuning Tyas
Data Science

Universitas Pembangunan Nasional “Veteran” Jawa Timur
Surabaya, Indonesia
sischa_wahyuning.sada@upnjatim.ac.id

The development of social media, particularly TikTok, has become one of the main platforms for the public to provide opinions and reviews on MSME products. This research aims to analyze public sentiment towards MSME products in Probolinggo based on user comments on TikTok using the Naïve Bayes Classification method. Research data was obtained through scraping TikTok comments, followed by preprocessing steps including cleaning text, case folding, tokenizing, stopword removal, and stemming. The data was then labeled as positive, neutral, and negative sentiment before feature extraction using TF-IDF. The classification process was carried out using the Naïve Bayes algorithm combined with SMOTE on training data to overcome class imbalance and evaluated using pure test data. The results showed that the Naïve Bayes method was able to classify sentiment with a final accuracy rate of 66.67%. The analysis showed that neutral sentiment dominated public opinion with 365 comments (47.04%), followed by positive sentiment with 341 comments (43.94%) and negative sentiment with 70 comments (9.02%). Data visualization indicated that most users responded very positively to MSME Probolinggo products on TikTok. This research

is expected to help MSME actors understand consumer perceptions and improve digital marketing strategies based on social media

Keywords: Sentiment Analysis, Naïve Bayes, TikTok, MSME Probolinggo, Text Mining, TF-IDF, SMOTE



© 2026 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution ShareAlike (CC BY SA) license (<https://creativecommons.org/licenses/by-sa/4.0/>).

I. INTRODUCTION

The development of digital technology and social media has transformed the way people interact when expressing opinions about a product or service. One of the social media platforms that has experienced rapid growth is TikTok. TikTok is not only used as an entertainment medium but is also utilized

as a means of promotion and digital marketing by Micro, Small, and Medium Enterprises (MSMEs). The interactive short video content on TikTok increases user engagement and influences consumers' purchasing decisions [1].

MSMEs play an important role in supporting local economic growth, including in Probolinggo Regency [2]. In today's digital era, many MSME actors are starting to use TikTok as a promotional medium to expand their product marketing reach. User comments on the TikTok platform can be a valuable source of information because they reflect consumer opinions, satisfaction, and criticism of the marketed products. However, the large number of comments makes manual analysis ineffective and time-consuming [3].

Sentiment analysis is one of the techniques in text mining used to identify a person's opinion or emotion toward an object based on textual data. Through sentiment analysis, user comments can be classified into positive, negative, or neutral categories, which can help business actors understand consumer perceptions of their products. A commonly used method in sentiment classification is Naïve Bayes, as it offers a simple, fast computational process and can provide good classification performance for text data [4].

Several previous studies have shown that the Naïve Bayes method is quite effective for sentiment analysis on social media platforms. Research on sentiment analysis of the TikTok application using the Naïve Bayes algorithm achieved high classification performance in categorizing user reviews [1]. Other digital review studies also showed good classification results by utilizing preprocessing and TF-IDF stages for text feature extraction [4]. The problem of class imbalance, which is often found in raw review data from e-commerce and social media, can also be addressed using the integration of the SMOTE algorithm to optimally enhance the sensitivity of minority class prediction [5].

Based on these issues, this study aims to analyze public sentiment toward MSME products in Probolinggo on the TikTok platform using the Naïve Bayes Classification method. Research data was obtained through scraping TikTok comments, which were then processed using preprocessing stages such as cleaning text, normalized repeat, case folding, tokenizing, stopword removal, and stemming. Feature extraction using TF-IDF and the application of the SMOTE technique were then used to overcome data imbalance. The classification results are visualized to identify the dominant public sentiment toward MSME Probolinggo products. This research is expected to help MSME actors understand consumer opinions and support more effective digital marketing strategy decision-making.

II. LITERATURE REVIEW

A. Sentiment Analysis

Sentiment analysis is a branch of text mining and Natural Language Processing (NLP) used to identify, extract, and classify a person's opinion or emotion regarding a topic based on textual data. Sentiment analysis is widely applied to social

media because users often provide responses, comments, or reviews about a product or service. The results of sentiment analysis are generally categorized into positive, negative, and neutral sentiment. This technique helps companies or businesses understand consumer perceptions of the products being marketed [6].

Social media such as TikTok is a potential data source for sentiment analysis due to its very high user interaction. User comments can be used to determine the level of satisfaction, criticism, or public response to a product. Therefore, sentiment analysis on the TikTok platform is important to help business actors evaluate their digital marketing strategies [7].

B. Naïve Bayes Classification

Naïve Bayes Classification is one of the machine learning algorithms that is widely used in text classification. This algorithm works based on Bayes' Theorem with the assumption that each feature is independent of the others. Naïve Bayes has several advantages, including fast computation, simplicity, and effectiveness in processing large amounts of textual data [8]. Therefore, this method is widely applied in sentiment analysis research on social media.

In this study, the Naïve Bayes method was selected because it is capable of providing good classification performance on textual data with a relatively limited dataset size. In addition, Naïve Bayes has lower computational complexity, allowing the model training process to be carried out more efficiently [9] [10]. In sentiment analysis research, Naïve Bayes is often combined with TF-IDF to improve classification accuracy. Previous studies on TikTok sentiment analysis using the Naïve Bayes method have shown that the algorithm can produce good classification performance in distinguishing positive, negative, and neutral sentiments. The formula can be seen in Equation 1.

$$P(C_j | W_i) = P(C_j) \prod_{1 \leq i \leq n} P(W_i | C_j) \quad (1)$$

Where:

$P(C_j | W_i)$: Posterior calculates the probability of the occurrence of class C_j , where $j = 1, 2, 3, \dots$

$P(W_i | C_j)$: Conditional probability calculates the probability of the occurrence of a term/word in class j .

$P(C_j)$: Prior calculates the probability of the occurrence of documents in class j in the training data.

C. Word Weighting (TF-IDF)

At the word weighting stage using TF-IDF (Term Frequency-Inverse Document Frequency), textual data is converted into numerical data based on the frequency of words in the document. TF measures how often a word appears in each document, while IDF calculates how many documents contain the word. IDF uses a logarithmic function to reduce the effect of words that frequently appear in many documents, according to the formula in Equation 2 [11].

$$TF - IDF_{t,d} = TF_{t,d} \times IDF_t \quad (2)$$

Where:

IDF_t : $Log \frac{N}{DF_t}$
 t : considered words;
 d : sentence weight (d);
 $TF-IDF_{t,d}$: weight (d) applied to word (t);
 $TF_{t,d}$: *TermFrequency*;
 IDF_t : *InverseDocumentFrequency*;
 N : total sentences;
 DF_t : number of occurrences of a word

D. SMOTE (Synthetic Minority Oversampling Technique)

Data imbalance is a common problem in sentiment analysis because the amount of data in each sentiment class is often unequal. This can cause the classification model to predominantly predict the majority class. To address this issue, the SMOTE (Synthetic Minority Oversampling Technique) technique is used. If this class imbalance is not addressed or is ignored, the model can become significantly biased toward the majority class. This can result in high overall accuracy due to the majority class dominance, but the predictive results for the minority class will be very low [5]. The SMOTE technique formula is in Equation 3.

$$X_{syn} = X_i + (X_{knn} - X_i) \times \delta \quad (3)$$

Where:

X_{syn} : Synthetic data
 X_i : Data to be replicated
 X_{knn} : Data that is the closest neighbor to X_i
 δ : Random value between 0 and 1

III. RESEARCH METHODOLOGY

Fig. 1 is the research methodology of this study.

A. Data Collection

The data used in this study consists of TikTok user comments related to Probolinggo MSME products. The data was obtained by scraping comments on TikTok videos discussing culinary products and local MSMEs from Probolinggo. Data collection was carried out using Google Colab with the help of the Python programming language and several supporting libraries such as pandas and numpy.

The collected data consists of comments in Indonesian. After a process of selection and data cleaning, a total of 776 comments were used in this study. The dataset consists of two main columns: the 'comment' column containing user comments and the 'label' column containing the sentiment category. Sentiment labels are divided into three categories: positive, neutral, and negative.

B. Data Preprocessing

Preprocessing is a crucial step in text mining aimed at cleaning and preparing text data before the classification process. The preprocessing stages in this study include cleaning, normalized repeat, case folding, tokenizing, stopwords removal, and stemming

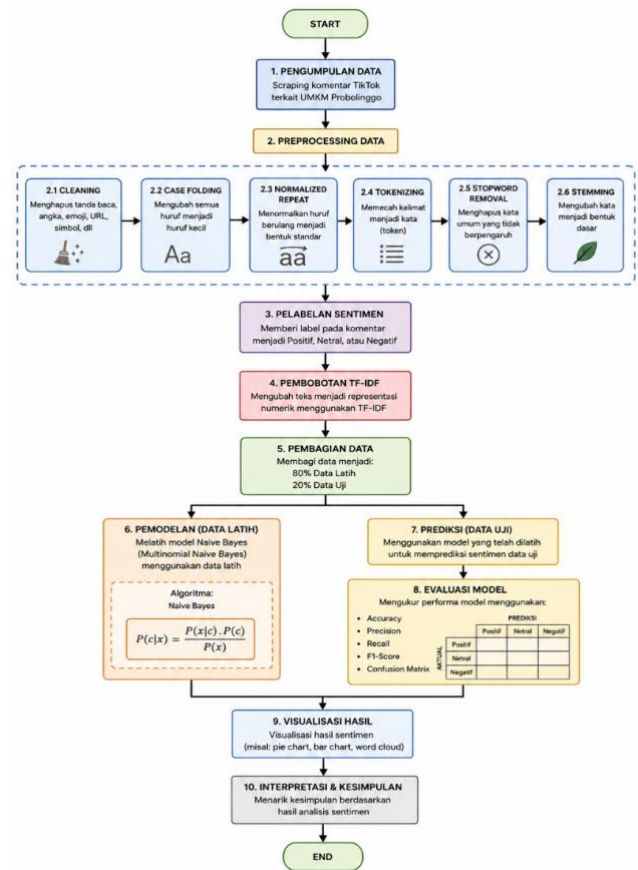


Fig. 1. Research Methodology

• Cleaning

Cleaning is performed to remove unnecessary characters in the text such as punctuation, numbers, emojis, URLs, repeated letters, and other irrelevant symbols. This process aims to reduce noise in the dataset, thereby improving the quality of analysis. TABLE 1 is the result of data cleaning process.

TABLE 1. EXAMPLE OF REVIEW DATA BEFORE AND AFTER CLEANING

Before Cleaning	After Cleaning
Tahu petisnya enak banget 😊😊	Tahu petis nya enak banget
Tiap lewat mesti ngantri 🙄	Tiap lewat mesti ngantri
Selera kak, monggo dicoba dulu kak 🙄 nanti balik sini lagi 😊	Selera kak monggo dicoba dulu kak nanti balik sini lagi

• Normalized Repeat

Normalized repeat is the process of normalizing repeated characters in words. Social media users often add repeated letters to express emotion, such as “enaaakkk”, “mantaapp”, or “baguuuss”. These repeated characters can cause words to be perceived as different even though they have the same meaning. Therefore, normalized repeat is applied to standardize such words. TABLE 2 is the normalization result.

TABLE 2. EXAMPLE OF REVIEW DATA BEFORE AND AFTER NORMALIZED REPEAT

Before Normalized Repeat	After Normalized Repeat
Tahu petisnya enak banget	Tahu petisnya enak banget
Tiap lewat mesti ngantri	Tiap lewat mesti ngantri
Selera kak monggo dicoba dulu kak nanti balik sini lagi	Selera kak monggo dicoba dulu kak nanti balik sini lagi

- Case Folding

Case folding is the process of converting all uppercase letters to lowercase so that the text is uniform. This helps avoid analysis discrepancies caused using uppercase and lowercase letters. TABLE 3 is the example of case folding result.

TABLE 3. EXAMPLE OF REVIEW DATA BEFORE AND AFTER CASE FOLDING

Before Case Folding	After Case Folding
Tahu petisnya enak banget	tahu petisnya enak banget
Tiap lewat mesti ngantri	tiap lewat mesti ngantri
Selera kak monggo dicoba dulu kak nanti balik sini lagi	selera kak monggo dicoba dulu kak nanti balik sini lagi

- Tokenizing

Tokenizing is the process of breaking sentences into individual words or tokens. This step helps the classification model understand each word separately so that sentiment can be better recognized. TABLE 4 is the result of tokenizing process.

TABLE 4. EXAMPLE OF REVIEW DATA BEFORE AND AFTER TOKENIZING

Before Tokenizing	After Tokenizing
tahu petisnya enak banget	['tahu', 'petisnya', 'enak', 'banget']
tiap lewat mesti ngantri	['tiap', 'lewat', 'mesti', 'ngantri']
selera kak monggo dicoba dulu kak nanti balik sini lagi	['selera', 'kak', 'monggo', 'dicoba', 'dulu', 'kak', 'nanti', 'balik', 'sini', 'lagi']

- Stopword Removal

Stopword removal is the process of removing common words that do not have significant impact on sentiment analysis, such as conjunctions or frequently appearing words with little meaningful value. TABLE 5 is the stopwords removal result.

TABLE 5. EXAMPLE OF REVIEW DATA BEFORE AND AFTER STOPWORD REMOVAL

Before Stopword Removal	After Stopword Removal
['tahu', 'petisnya', 'enak', 'banget']	['tahu', 'petisnya', 'enak']
['tiap', 'lewat', 'mesti', 'ngantri']	['lewat', 'ngantri']
['selera', 'kak', 'monggo', 'dicoba', 'dulu', 'kak', 'nanti', 'balik', 'sini', 'lagi']	['selera', 'monggo', 'dicoba']

- Stemming

Stemming is the process of reducing words to their root form by removing prefixes or suffixes. This step aims to simplify words with similar meanings into a single root form. The result of stemming process can be seen in TABLE 6.

TABLE 6. EXAMPLE OF REVIEW DATA BEFORE AND AFTER STEMMING

Before Stemming	After Stemming
['tahu', 'petisnya', 'enak']	['tahu', 'petis', 'enak']
['lewat', 'ngantri']	['lewat', 'antri']
['selera', 'monggo', 'dicoba']	['selera', 'mongo', 'coba']

- Sentiment Labeling

After the preprocessing stage, the next process is sentiment labeling which can be seen in TABLE 7. Sentiment labeling aims to classify each comment into one of three sentiment categories: positive, negative, or neutral. In this study, the labeling process was conducted manually using a lexicon-based approach by considering the presence of sentiment-bearing words and the contextual meaning of each comment.

TABLE 7. SENTIMENT LABELING RESULTS

Data	Label
tahu petis enak	Positive
lewat antri	Negative
selera mongo coba	Neutral

Positive sentiment labels were assigned to comments containing expressions of satisfaction, appreciation, support, or positive opinions toward MSME products in Probolinggo. Examples of positive expressions found in the dataset include words such as “*enak*” (delicious), “*mantap*” (great), “*murah*” (cheap), “*rekomen*” (recommended), and “*favorit*” (favorite). These comments generally indicate consumer satisfaction with product quality, taste, service, or price.

Negative sentiment labels were assigned to comments containing dissatisfaction, complaints, criticism, or negative opinions regarding the products or services. Examples of negative expressions identified in the dataset include words such as “*mahal*” (expensive), “*lama*” (slow), “*kurang enak*” (not tasty), “*antri*” (queue), and “*jauh*” (far). Negative comments usually reflect consumer disappointment related to product quality, service, waiting time, or accessibility.

Meanwhile, neutral sentiment labels were assigned to comments that did not contain clear emotional polarity, such as questions, informational statements, promotional responses, or ordinary interactions between users. Neutral comments generally do not explicitly express either satisfaction or dissatisfaction toward the products.

To maintain consistency in the labeling process, each comment was reviewed repeatedly by considering both lexical indicators and sentence context. This approach was applied to reduce subjectivity and improve the quality of sentiment classification before entering the TF-IDF weighting and modeling stages.

- TF-IDF Weighting

After sentiment labeling, the data is converted into numerical form using the TF-IDF (Term Frequency–Inverse Document Frequency) method. This method is used to calculate the importance of a word in a document based on its frequency of occurrence in the dataset.

- Data Balancing Using SMOTE

The dataset in this study has an imbalance in the amount of data in each sentiment class, especially in the negative class. Therefore, the SMOTE (Synthetic Minority Oversampling Technique) method is used to balance the amount of data in each class. This technique works by creating synthetic data in the minority class so that the classification model can work more optimally.

C. Modeling

After preprocessing, the data is split into training data and testing data. The split is performed using the `train_test_split` method with a composition of 80% training data and 20% testing data. The classification process in this study uses the Multinomial Naïve Bayes algorithm, as it is suitable for text classification and performs well in sentiment analysis. The classification process starts from TF-IDF weighting, model training using training data, then prediction on test data. Dataset is classified into three sentiment categories: Positive sentiment, Neutral sentiment, and Negative sentiment.

D. Evaluation

The evaluation stage is carried out to measure the performance of the classification model that has been built. In this study, evaluation is performed using several metrics: Accuracy, Precision, Recall, F1-Score, and Confusion Matrix.

IV. RESULT AND DISCUSSION

Based on the labeling process using a lexicon-based approach on the entire original dataset of 776 comments, an overview of public perception on the TikTok platform was obtained. The results of sentiment label distribution visualization can be seen in Fig. 2. The figure shows that out of the total 776 comments analyzed, public opinion is dominated by Neutral Sentiment with 365 comments (47.04%), closely followed by Positive Sentiment with 341 comments (43.94%), while Negative Sentiment has the least number with only 70 comments (9.02%).

The high number of neutral sentiment (365 data) represents that public interaction on social media is dominated by audiences seeking operational information, such as questions about product price details, available local culinary menu variations, store operating hours, and the route to the physical location of the MSMEs. On the other hand, the very high number of positive sentiment (341 data) proves that digital marketing content based on short videos on TikTok successfully triggers satisfaction and a positive emotional response from consumers towards MSME Probolinggo products.

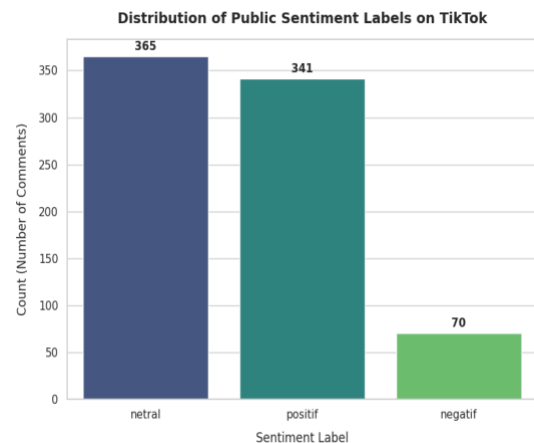


Fig. 2. The Sentiment Visualization

After the dataset was divided using a ratio of 80% for training data and 20% for pure test data (without synthetic data intervention), a total of 156 test data was obtained. To address class imbalance, the SMOTE method was specifically applied to the training data before modeling. Testing the Multinomial Naïve Bayes model resulted in a final accuracy rate of 66.67%. Details of the precision, recall, and f1-score values are shown in the classification report Fig. 3.

```
... Akurasi Akhir: 0.6667
```

Classification Report:				
	precision	recall	f1-score	support
negatif	0.21	0.43	0.28	14
netral	0.78	0.59	0.67	73
positif	0.76	0.80	0.78	69
accuracy			0.67	156
macro avg	0.58	0.60	0.58	156
weighted avg	0.72	0.67	0.68	156

Fig. 3. Classification Report Naïve Bayes

To deepen the analysis of model performance, evaluation is also complemented by a Confusion Matrix, representing the number of correct and incorrect predictions on pure test data. Fig. 4. Based on the matrix, the diagonal shows the number of data correctly classified by the model, namely 6 negative comments, 43 neutral comments, and 55 positive comments. A more detailed analysis of the matrix describes the classification phenomena as follows:

Positive Sentiment Classification Strength: The model proved very strong in recognizing positive sentiment, correctly predicting 55 out of 69 actual data (recall 0.80). The lexical pattern of positive reviews on TikTok tends to be explicit and repetitive (such as the use of the words "delicious", "great", "cheap"), so the TF-IDF weighting of these text features is very optimal for the model to learn.

Neutral Class Ambiguity: The largest classification error occurs in the actual neutral category, where there are 14 neutral

comments predicted as negative and 16 neutral comments predicted as positive. This ambiguity arises because neutral reviews on social media often include specific product names (for example: "tahu petis") which in the training data have positive content, thus confusing the independence assumption of the Naïve Bayes algorithm features.

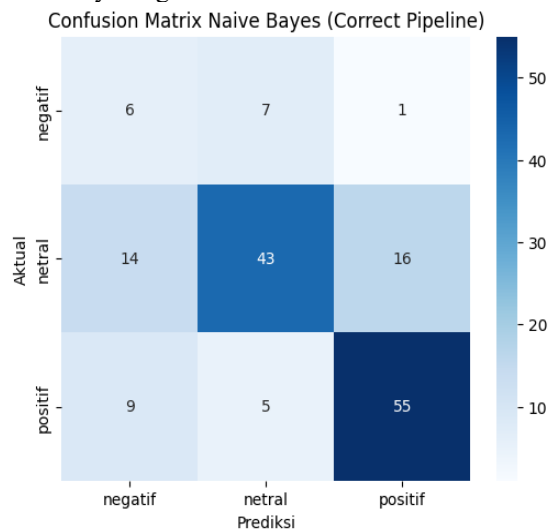


Fig. 4. Confusion Matrix

SMOTE Limitation on Real Negative Data: Out of 14 actual negative data, the model was only able to identify 6 precisely. This objectively proves that synthetic data balancing through SMOTE in the training data has not fully captured the characteristics of real negative reviews, which are often conveyed implicitly (sarcasm) or using informal vocabulary.

V. CONCLUSION

Based on the results of the experiments and discussions that have been conducted, it can be concluded that sentiment analysis of public opinion toward Probolinggo MSME products on the TikTok platform resulted in neutral sentiment being the majority opinion of the public with a percentage of 47.04% (365 comments), followed by positive sentiment at 43.94% (341 comments), and negative sentiment at 9.02% (70 comments). The implementation of the Multinomial Naïve Bayes algorithm combined with TF-IDF weighting and class imbalance handling using SMOTE resulted in a final accuracy of 66.67%.

This model proved to have the highest performance in classifying positive sentiment (f1-score 0.78) but still has limitations in recognizing negative sentiment (f1-score 0.28) due to the limited variety of real data in the field. This research demonstrates that Probolinggo MSME products have a very good digital impression on the TikTok platform, where most public interactions are in the form of positive support and further inquiries regarding the operations of local business products.

Future research can be expanded by using larger datasets collected from various social media platforms such as Instagram, Facebook, and X (Twitter). In addition, further

studies may implement other classification methods or deep learning-based models to improve sentiment analysis performance. Further development can also be conducted by applying aspect-based sentiment analysis to identify consumer opinions more specifically regarding certain aspects, such as product quality, pricing, service, and MSME promotional strategies. Therefore, the analysis results are expected to provide deeper insights for MSME actors in improving service quality and digital marketing strategies.

ACKNOWLEDGEMENT

The authors would like to express their highest gratitude to Universitas Hafshawaty Zainul Hasan for funding and fully supporting this research through the Internal Research Grant scheme (*Hibah Internal*). We also extend our appreciation to the Department of Informatics for providing the necessary facilities and environment to complete this study.

REFERENCES

- [1] M. T. F. Maulana, B. I. Nugroho, and E. U. S. Utami, "Analisis Sentimen Ulasan Aplikasi TikTok di Google Playstore Menggunakan Algoritma Naive Bayes," *RIGGS*, vol. 4, no. 3, pp. 6627–6635, Oct. 2025, doi: 10.31004/riggs.v4i3.2962.
- [2] U. D. Susilowati and N. A. Khatijah, "Pelatihan digital marketing dan pemanfaatan e-commerce untuk optimalisasi penjualan UMKM Keripik Sukun Holifah Sari Pajajaran," *Jurnal Dedikasi Masyarakat*, vol. 7(2), pp. 99–104, Mar. 2024.
- [3] K. Fadli and F. N. Hasan, "Analisis Sentimen Tanggapan Masyarakat Terhadap Penutupan TikTok Shop Menggunakan Metode Naive Bayes," *josh*, vol. 6, no. 1, pp. 396–405, Oct. 2024, doi: 10.47065/josh.v6i1.6060.
- [4] A. Komarudin and A. M. Hilda, "Analisis Sentimen Ulasan Aplikasi Identitas Kependudukan Digital Pada Play Store Menggunakan Metode Naive Bayes," *co-science*, vol. 4, no. 1, pp. 28–36, Jan. 2024, doi: 10.31294/coscience.v4i1.2955.
- [5] M. P. Pulungan, A. Purnomo, and A. Kurniasih, "Penerapan SMOTE untuk Mengatasi Imbalance Class dalam Klasifikasi Kepribadian MBTI Menggunakan Naive Bayes Classifier," *JTIK*, vol. 11, no. 5, pp. 1033–1042, Oct. 2024, doi: 10.25126/jtik.1077989.
- [6] Friska Aditia Indriyani, Ahmad Fauzi, and Sutan Faisal, "Analisis sentimen aplikasi tiktok menggunakan algoritma naïve bayes dan support vector machine," *tekno*, vol. 10, no. 2, pp. 176–184, Jul. 2023, doi: 10.37373/tekno.v10i2.419.
- [7] N. Cahyono and Anggista Oktavia Praneswara, "Analisis Sentimen Ulasan Aplikasi TikTok Shop Seller Center di Google Playstore Menggunakan Algoritma Naive Bayes," *ijcs*, vol. 12, no. 6, Dec. 2023, doi: 10.33022/ijcs.v12i6.3473.
- [8] S. K. Wardani and Y. A. Sari, "Analisis Sentimen menggunakan Metode Naive Bayes Classifier terhadap

- Review Produk Perawatan Kulit Wajah menggunakan Seleksi Fitur N-gram dan Document Frequency Thresholding”.
- [9] R. P. Dwitama and I. V. Paputungan, “Analisis Perbandingan Naive Bayes Dan Svm Dalam Klasifikasi Gender Di Platform Twitter”.
- [10] I. Ernawati, “Naive Bayes Classifier Dan Support Vector Machine Sebagai Alternatif Solusi Untuk Text Mining,” *JTIP*, vol. 12, no. 2, pp. 32–38, Dec. 2019, doi: 10.24036/tip.v12i2.219.
- [11] S. A. S. Mola, P. R. Lete, B. J. A. J. A. Pa, Triyanto, and T. Widiastuti, “Analisis Sentimen Menggunakan Metode Naive Bayes Dan Metode Support Vector Machine Pada Kasus Pelantikan Artis Sebagai Anggota Anggota Dpr Ri Tahun 2024,” *HOAQ*, vol. 15, no. 1, pp. 22–32, May 2024, doi: 10.52972/hoaq.vol15no1.p22-32.