



Public Sentiment Analysis on Bullying Cases of Binus Serpong Students Using the Naive Bayes Method

Anggelia Dwi Rahayu ^{*1}, Erik Iman Heri Ujianto ¹

¹ Dept. of Informatics, University of Technology Yogyakarta. Yogyakarta, Indonesia.

KEYWORDS

Sentiment Analysis
Bullying
Public Perception
Naive Bayes Classifier
YouTube

ARTICLE HISTORY

Received 1 November 2024
Received in revised form
20 November 2024
Accepted 21 November 2024
Available online 27 November
2024

ABSTRACT

Sentiment analysis has evolved into a vital tool for comprehending public perceptions of social issues, including bullying cases in educational settings. This research focuses on sentiment analysis of public perceptions regarding bullying cases involving students at High School Binus Serpong. The dataset used consists of 2,259 comments collected from YouTube, categorized into three sentiment types : positive, negative, and neutral. The method applied in this study is the Complement Naive Bayes algorithm and multinomial Naive Bayes integrated with TF-IDF within a classification pipeline. The analysis results indicate that the classification model achieved an accuracy of 80.47%. Specifically, the evaluation results show that negative sentiment contributes 72.2%, positive sentiment 20.7%, and neutral sentiment 0.7%. Furthermore, the evaluation results indicate that negative sentiment has a precision of 0.81, recall of 0.93, and f1-score of 0.87, positive sentiment has a precision of 0.67, recall of 0.40, and f1-score of 0.50; while neutral sentiment has a precision of 0.70, recall of 0.43, and f1-score of 0.53.

© 2024 The Authors. Published by Penteract Technology.

This is an open access article under the CC BY-NC 4.0 license (<https://creativecommons.org/licenses/by-nc/4.0/>).

1. INTRODUCTION

Bullying is a recurring aggressive behaviour directed towards individuals, particularly prevalent among teenagers. Although it is common, bullying should not be taken lightly as it can lead to long-term negative consequences for the victims [1]. This phenomenon can occur in various settings, including schools, workplaces, and social media platforms[2]. Adolescents are especially vulnerable to bullying, which is often influenced by factors such as egocentrism and aggression [3].

The main types of bullying identified include physical, verbal, relational, and cyberbullying, with verbal bullying being the most common[4]. This behaviour is regarded as part of adolescent delinquency and is exacerbated by the egocentric nature typical of teenage years. Bullying significantly disrupts adolescents' social interactions, resulting in isolation, decreased self-esteem, and hindering their ability to form healthy social relationships [5]. Contributing factors include parental influence, individual characteristics, peer environments, and a lack of supervision in schools. The impacts of bullying are often reflected in declining academic performance and social ostracism [6].

In light of these issues, this research aims to explore and assess public emotions related to bullying, with a particular focus on interactions on social media [7]. Social media serves as an effective and efficient communication tool; however, the freedom of expression on these platforms faces significant challenges, particularly regarding the rapid spread of negative news[8]. Therefore, sentiment analysis will be employed to differentiate between negative, positive, and neutral public responses, providing a clearer understanding of their perspectives and attitudes towards bullying through comments on platforms such as YouTube [9].

YouTube allows users to express their opinions through comments, which can be analysed using text mining techniques[10]. Sentiment analysis will be used to categorise these opinions and assess sentiment towards the topics or sources discussed [11]. A deeper understanding of these reactions is expected to contribute significantly to the development of more effective and comprehensive bullying prevention strategies, both in social media and in real life [12].

To achieve this, the research will utilise the Naive Bayes algorithm, a classification method commonly used in data mining and text mining[13]. This method is based on Bayes'

*Corresponding author:

E-mail address: Anggelia Dwi Rahayu <anggelidwurahayu@gmail.com>.

<https://doi.org/10.56532/mjsat.v4i4.412>

2785-8901/ © 2024 The Authors. Published by Penteract Technology.

This is an open access article under the CC BY-NC 4.0 license (<https://creativecommons.org/licenses/by-nc/4.0/>).

theorem, where all features are assumed to contribute equally and independently to determining a specific class [14]. The Naive Bayes algorithm is a straightforward classification approach that calculates probabilities from the dataset and will be used in this study to classify public opinions [15]. This method was chosen for its effectiveness in producing highly accurate results, making it a popular choice among researchers [16].

Previous research employing the Naive Bayes method for sentiment analysis has been conducted, such as the study by Erina Undamayanti, Teguh Iman Hermanto, and Ismi Kaniawulan (2022), which showed that the Merdeka Belajar Kampus Merdeka (MBKM) programme was well received by Twitter users, particularly students, with positive sentiment reaching 61.92% and negative sentiment at 38.08% [17]. Additionally, research by Syahjuddin Azra and Rizky Tahara Shita (2023) analysed customer reviews on Tokopedia using the Naive Bayes method, achieving an accuracy of 86.96%, with precision at 90.48% and recall at 95% [18]. Furthermore, Noor Aliyah Susanti, Miftahul Walid, and Hoiriyah (2022) analysed hate speech on Twitter using the Multinomial Naive Bayes method, yielding a classification accuracy of 69.23% [19]. Research conducted by Fajar Sodik Pamungkas and Iqbal Kharisudin (2021) demonstrated that in the sentiment analysis of public reactions in Indonesia to the Covid-19 pandemic on Twitter, the Support Vector Machine (SVM) algorithm achieved the highest accuracy of 90.01%. Following this, Naive Bayes achieved an accuracy of 79.20%, while K-Nearest Neighbour (KNN) recorded an accuracy of 62.10% [20]. Research by Tobby Wiratama Putra, Agung Triayudi, and Andrianingsih (2021) analysed sentiments regarding online learning using Twitter data and compared three algorithms: Naive Bayes, K-Nearest Neighbour (KNN), and Decision Tree. The Decision Tree approach demonstrated the highest performance, achieving an accuracy of 61.92%, precision of 73.63%, and recall of 11.42% [21].

By employing sentiment analysis methods, this research aims to uncover patterns of attitudes and opinions within society, categorising responses to comments as positive, negative, or neutral. Through the collection and analysis of data from various sources, including YouTube videos that have garnered public attention, this study seeks to provide a more comprehensive perspective on the bullying situation among students at Binus Serpong High School.

2. METHODOLOGY

This study performs sentiment analysis, which involves gathering and assessing opinions obtained from various social media platforms, particularly YouTube. This method allows for the gathering of public opinions through social networks, where discussions about public services and current issues occur. In this context, sentiment analysis involves extracting public opinions from comments on the YouTube platform regarding bullying cases that took place at High School Binus Serpong, utilizing unstructured text data.

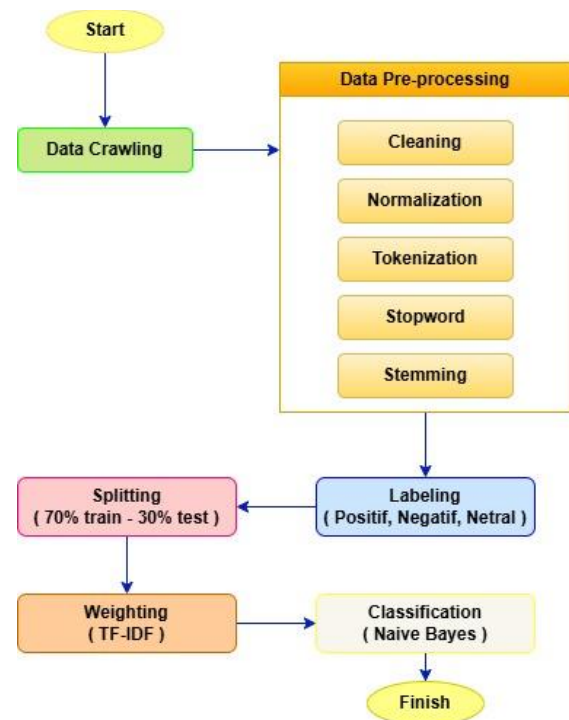


Fig. 1. Research Method

The data used in this research consists of comments from YouTube. Collected from March 3, 2024, to September 30, 2024, using a Python-based data crawling method. The collected data then underwent a preprocessing stage, which included cleaning, normalization, tokenization, stopwords, and stemming, all performed using Google Colab and Python.

For the labeling process, the comments were categorized based on sentiment: positive, negative, and neutral. Afterward, the data was split, with 70% allocated for training and 30% for testing. Next, a data weighting procedure was applied using the TF-IDF method. The final step involved modeling with the Multinomial and Complement Naïve Bayes algorithms, resulting in the evaluation and accuracy metrics for the models used in this study.

2.1 Data Collection Process

The data collection process for this research utilized the Google Cloud Console, a platform designed to gather relevant data from social media, specifically YouTube (X). To access the YouTube API, it is necessary to create an account to obtain an API Key. This API Key serves as an authentication key, allowing the user to interact with the YouTube API and granting permission to access the required data from the platform. For this study, a total of 2,259 comment data entries were collected through a data crawling process. The template is used to format your paper and style the text. All margins, column widths, line spaces, and text fonts are prescribed; please do not alter them. You may note peculiarities. For example, the head margin in this template measures proportionately more than is customary. This measurement and others are deliberate, using specifications that anticipate your paper as one part of the entire proceedings, and not as an independent document. Please do not revise any of the current designations.

Table 1. Crawled Data Results

	published At	authorDisplay Name	textDisplay	likeCount
0	2024-04-20T15:01:29Z	@wahdangun	Kasusnya mandek neh anak artisnya kayaknya bebas	0
1	2024-04-14T14:55:06Z	@cheeseD-wr8ug	nama nya juga geng tai kelakuan nya juga kayak...	0
2	2024-04-14T02:15:03Z	@hendraserdadudisc3440	0.09 Terkadang org tua tersangka malah berupaya men...	0
3	2024-03-27T23:49:13Z	@bambangsupriyadi7988	Ortu elit ndidik anak sulit	0
4	2024-03-13T22:04:23Z	@ayujohan5655	Dlm kasus perundungan biasanya ada satu otakny...	0

temannya sy amati dg sigap dan langsung sy tarik jegernya ke tengah lapangan yg lain semua langsung pda berhenti	otaknya preman dalam satu kelompok itu pengalaman saya jadi guru pernah menjumpai anak didik yang keroyok temannya saya amati dengan sigap dan langsung saya tarik preman ke tengah lapangan yang lain semua langsung pada berhenti
--	---

The first column, Index, shows the sequential number for each entry. The cleaned_text column displays the original text that has been cleaned of typographical errors and non-standard language usage, while the norm_text column presents the text that has been normalised into a more precise and clearer form.

For example, the phrase "kasusnya mandek neh anak artisnya kayaknya bebas" has been changed to "kasusnya berhenti lagi anak artisnya kayaknya bebas," making it easier to understand. Additionally, complex long sentences such as "terkadang org tua tersangka malah berupaya mengajak damaibrbagi ku itu malah tdk mendidik" have been simplified and restructured to enhance clarity.

This normalisation process is crucial for rectifying ambiguities in the text, resulting in more consistent and accurate data for further sentiment analysis. Thus, normalisation provides a solid foundation for better understanding public sentiment regarding the issue of bullying.

- c) Tokenization : Dividing text into smaller components (tokens) like words or phrases.

Table 3. Tokenization Result

index	norm text	token
0	kasusnya berhenti lagi anak artisnya kayaknya bebas	kasusnya,berhenti,lagi,anak,artisnya,kayaknya,bebas
1	nama juga kelompok kotoran kelakuan juga seperti kotoran	nama,juga,kelompok,kotoran,kelakuan,juga,seperti,kotoran
2	Terkadang orang tua tersangka malah berupaya mengajak damai bagi saya itu malah tidak mendidik seakan akan orang tua tersangka uang adalah segalanya dan menganggap hal itu masalah sepele dengan dalih bentuk kasih sayang thdp anakbrkalo ingin mendidik anak biarkan hukum berjalan dan semua tersangka bertanggung jawab dari apa yang mereka perbuat	terkadang,orang,tua,tersangka,malah,berupaya,mengajak,damai,bagi,saya,itu,malah,tidak,mendidik,seakan,akan,orang,tua,tersangka,uang,adalah,segalanya,dan,menganggap,hal,itu,masalah,sepele,dengan,dalih,bentuk,kasih,sayang,terhadap,anak,jika,ingin,mendidik,anak,biarkan,hukum,berjalan,dan,semua,tersangka,bertanggung,jawab,dari,apa,yang,mereka,perbuat
3	orang tua berada mendidik anak sulit	orang,tua,berada,mendidik,anak,sulit
4	dalam kasus perundungan biasanya ada satu otaknya preman dalam satu kelompok itu pengalaman saya jadi guru pernah menjumpai anak didik yang keroyok temannya saya	dalam,kasus,perundungan,biasanya,ada,satu,otaknya,preman,dalam,satu,kelompok,itu,pengalaman,saya,jadi,guru,pernah,menjumpai,anak,anak,keroyok,temannya,saya

2.2 Data Preprocessing

- Data Cleaning : Removing specific characters or symbols, numbers, and converting all letters to lowercase.
- Normalization : Ensuring data consistency.

Table 2. Normalization Result

index	cleaned_text	norm_text
0	kasusnya mandek neh anak artisnya kayaknya bebas	kasusnya berhenti lagi anak artisnya kayaknya bebas
1	nama nya juga geng tai kelakuan nya juga kayak tai	nama juga kelompok kotoran kelakuan juga seperti kotoran
2	terkadang org tua tersangka malah berupaya mengajak damaibrbagi ku itu malah tdk mendidik seakanakan org tua tersangka quotuang adalah segalanya dan menganggap hal itu masalah sepele dg dalih bentuk kasih sayang thdp anakbrkalo ingin mendidik anak ya biarkan hukum berjalan dan semua tersangka bertanggung jawab dr apa yg mereka perbuat	Terkadang orang tua tersangka malah berupaya mengajak damai bagi saya itu malah tidak mendidik seakan akan orang tua tersangka uang adalah segalanya dan menganggap hal itu masalah sepele dengan dalih bentuk kasih sayang terhadap anak jika ingin mendidik anak biarkan hukum berjalan dan semua tersangka bertanggung jawab dari apa yang mereka perbuat
3	ortu elit ndidik anak sulit	orang tua berada mendidik anak sulit
4	pernah menjumpai anak didik yg memgeroyok	dalam kasus perundungan biasanya ada satu

amati dengan sigap dan langsung saya tarik preman ke tengah lapangan yang lain semua langsung pada berhenti	yang,keroyok,temannya, saya,amati,dengan,sigap, dan,langsung,saya,tarik, preman,ke,tengah,lapangan,yang,lain,semua,langsung,pada,berhenti
---	---

The table consists of an index indicating the sequential number for each entry, a column for normalised text that has been cleaned for clarity, and a column of individual tokens derived from that text. For example, the phrase "kasusnya berhenti lagi anak artisnya kayaknya bebas" is tokenised into "kasusnya, berhenti, lagi, anak, artisnya, kayaknya, bebas," while "nama juga kelompok kotoran kelakuan juga seperti kotoran" becomes "nama, juga, kelompok, kotoran, kelakuan, juga, seperti, kotoran," illustrating the redundancy present. More complex sentences are also broken down into manageable components, allowing for a detailed analysis. Ultimately, tokenisation enables a more nuanced examination of the text, facilitating the identification of sentiment-bearing words and enhancing the understanding of public sentiment towards bullying.

- a) *Stopword : Removing frequently used words to concentrate on more significant terms.*

Table 4. Stopword Result

index	token	stopword
0	kasusnya,berhenti,lagi,anak,artisnya,kayaknya,bebas	kasusnya berhenti anak artisnya kayaknya bebas
1	nama,juga,kelompok,kotoran,kelakuan,juga,seperti,kotoran	nama kelompok kotoran kelakuan kotoran
2	terkadang,orang,tua,tersangka,malah,berupaya,mengajak,damai,bagi,saya,itu,malah,tidak,mendidik,seakan,akan,orang,tua,tersangka,uang,adalah,segalanya,dan,menganggap,hal,itu,masalah,sepele,dengan,dalih,bentuk,kasih,sayang,terhadap,anak,jika,ingin,mendidik,anak,biarkan,hukum,berjalan,dan,semua,tersangka,bertanggungjawab,dari,apa,yang,merdeka,perbuat	terkadang orang tua tersangka berupaya mengajak damai mendidik seakan orang tua tersangka uang menganggap sepele dalih bentuk kasih sayang anak mendidik anak biarkan hukum berjalan tersangka bertanggung perbuat
3	orang,tua,berada,mendidik,anak,sulit	orang tua mendidik anak sulit
4	dalam,kasus,perundungan,biasanya,ada,satu,otaknya,preman,dalam,satu,kelompok,itu,pengalaman,man,saya,jadi,guru,pernah,menjumpai,anak,didik,yang,keroyok,temannya,saya,amati,dengan,sigap,dan,langsung,saya,tarik,preman,ke,tengah,lapangan,yang,lain,semua,langsung,pada,berhenti	perundungan otaknya preman kelompok pengalaman guru menjumpai anak didik keroyok temannya amati sigap langsung tarik preman lapangan langsung berhenti

The original tokens include terms such as "kasusnya," "berhenti," "lagi," "anak," "artisnya," and "kayaknya," among others. These tokens, while structurally important, contain

many stopwords like "juga," "dan," "itu," "yang," and "adalah," which can clutter the analysis by introducing noise without adding significant meaning.

After removing stopwords, the remaining tokens will focus on the more substantive content of the comments, allowing for a clearer understanding of public sentiment regarding the issue of bullying. This process aids in enhancing the accuracy of sentiment analysis by concentrating on words that contribute more directly to the conveyed emotions and opinions.

- b) *Stemming : Transforming words into their root form (stem) to categorize words with similar meanings.*

Table 5. Stemming Result

index	data komentar
0	kasus henti anak artis seperti bebas
1	nama kelompok kotor laku kotor
2	terkadang orang tua sangka upaya ajak damai didik akan orang tua sangka uang anggap sepele dalih bentuk kasih sayang anak didik anak biar hukum jalan sangka tanggung buat
3	orang tua didik anak sulit
4	rundung otak preman kelompok alam guru jumpa anak didik keroyok teman amat sigap langsung tarik preman lapang langsung henti

The stemming process simplifies words to their base or root form, consolidating variations into a single representation, which is beneficial for text analysis. For example, the comment "kasusnya berhenti lagi anak artisnya kayaknya bebas" is stemmed to "kasus henti anak artis seperti bebas," retaining its core meaning while eliminating extraneous forms. Similarly, "nama juga kelompok kotoran kelakuan juga seperti kotoran" is reduced to "nama kelompok kotor laku kotor," focusing on essential content.

A complex sentence is transformed into "terkadang orang tua sangka upaya ajak damai didik akan orang tua sangka uang anggap sepele dalih bentuk kasih sayang anak didik anak biar hukum jalan sangka tanggung buat," capturing the main ideas without variations. The phrase "orang tua berada mendidik anak sulit" becomes "orang tua didik anak sulit," emphasizing the primary action and subjects. Lastly, a lengthy comment is condensed to "rundung otak preman kelompok alam guru jumpa anak didik keroyok teman amat sigap langsung tarik preman lapang langsung henti," simplifying the phrasing while preserving the narrative.

2.3 Data Labeling

Labeling In sentiment analysis, labeling refers to the process of categorizing textual data into classifications such as positive, negative, or neutral to identify opinions or emotions expressed in the text. This procedure allows algorithms to detect sentiment patterns and forecast sentiments in new data, making it crucial for training machine learning models. Labeled data is utilized to train algorithms for the effective classification of text based on sentiment. A negative label signifies criticism or dissatisfaction, such as condemnation of bullying behavior, while a positive label indicates support or appreciation for the victim or the handling of the case. In contrast, a neutral label is applied to comments that convey information without strong emotions. By employing labeling, researchers can conduct a deeper analysis of public perception regarding bullying cases.

2.4 Data Splitting

Once the labeling is complete, the dataset is split into two parts : 70% allocated for train and 30% for test. The training data is used to teach the machine learning model how to recognize patterns in the labeled data. In contrast, the testing data evaluates the model's performance on previously unseen data, which is essential for confirming its accuracy and reliability in real-world scenarios.

2.5 Data Weighting Process Using TF-IDF

After splitting the dataset into training and testing sets, the subsequent step involved applying the Term Frequency-Inverse Document Frequency (TF-IDF) method. TF-IDF assesses the importance of terms within a document relative to a set of documents, consisting of two components : Term Frequency (TF) and Inverse Document Frequency (IDF)[22].

$$TF(t, d) = \frac{\text{Number of times term } t \text{ appears in document } d}{\text{Total number of terms in document } d} \quad (1)$$

This quantifies how often a term t appears in document d .

$$IDF(t) = \log \left(\frac{N}{DF(t)} \right) \quad (2)$$

Where N is the total number of documents, and $DF(t)$ is the count of documents containing term t . IDF decreases the weight of common terms that appear across many documents, focusing on more distinctive and relevant words[23].

$$TF - IDF(t, d) = TF(t, d) \times IDF(t) \quad (3)$$

This equation combines TF and IDF, producing a score that indicates the significance of each word in relation to the entire dataset. In this study, utilizing TF-IDF on the labeled YouTube comments transforms the text into numerical vectors, highlighting important terms[24]. This enables models such as Multinomial Naive Bayes and Complement Naive Bayes to classify sentiments (positive, negative, neutral) more efficiently based on the relevance of the terms.

2.6 Model Selection

This study employs two classification algorithms : Multinomial Naive Bayes and Complement Naive Bayes, both based on Bayes' Theorem and assuming feature independence within the dataset. Multinomial Naive Bayes is particularly well-suited for sentiment analysis of bullying-related comments, as it calculates the probability of each word based on its frequency, making it effective for multi-category text data. In contrast, Complement Naive Bayes addresses the limitations of Multinomial Naive Bayes, especially in handling class imbalance by comparing word frequencies in bullying-related comments with those in unrelated comments, thus improving performance in cases of uneven distribution between positive and negative sentiments. The choice of these algorithms is informed by their efficiency in managing large datasets and their complementary approaches in analyzing public sentiment about bullying, with the expectation that their combination will yield accurate sentiment classifications for comments collected from social media platforms like YouTube.

3. RESULTS AND DISCUSSION

In this study, the researchers obtained sentiment labels from a total of 2,259 comments, classified as follows, Negative : 1.626, Positive : 617, and Neutral : 16.

Comment Sentiment Distribution

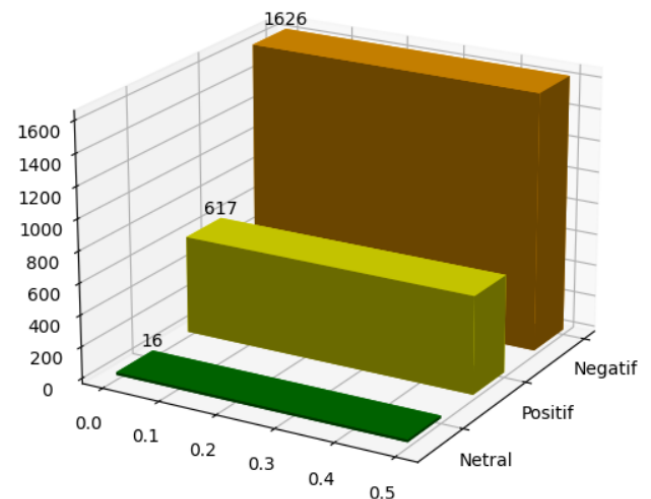


Fig. 2. Sentiment Distribution

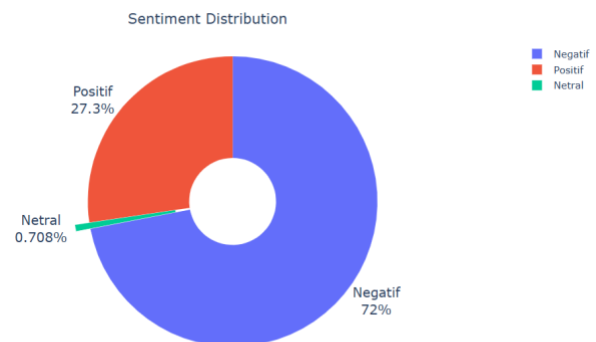


Fig. 3. Percentage of Sentiment Distribution

Based on the analysis of 2,259 YouTube comments related to the issue of bullying, the results of the data labeling distribution indicate that 1,626 comments (71.98%) were classified as negative, 617 comments (27.36%) as positive, and only 16 comments (0.71%) as neutral.

The high percentage of negative comments reflects public concern and dissatisfaction regarding the issue of bullying, which is regarded as a serious social phenomenon detrimental to victims. These negative comments can be interpreted as an emotional response from the community, expressing concern about the psychological and social impacts of bullying. This indicates a collective awareness of the need for improved handling of bullying issues across various environments, including schools and communities.

On the other hand, the number of positive comments, which totalled 617 (27.36%), indicates support for the efforts made to address the issue of bullying. These positive

comments may reflect the public's hope and desire to see significant changes in the prevention and management of bullying cases. Although the proportion of positive comments is smaller compared to negative ones, their existence is still important as an indication that some members of society are endeavouring to provide support and solutions to the problems at hand.

The minimal number of neutral comments, with only 16 comments (0.71%), suggests that the majority of respondents have a clear perspective on the issue of bullying, whether in the form of criticism or support. This indicates that bullying is a topic that cannot be overlooked and requires serious attention from all parties involved.

3.1 Distribution Data Splitting

Once the sentiment results for positive, negative, and neutral comments are classified, the next step involves dividing the labeled data into two subsets: 70% is designated as the training data, while 30% serves as the testing data.

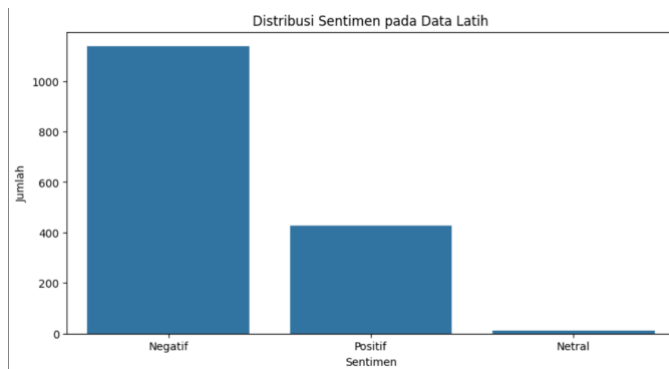


Fig. 4. Distribution of Training Data Sentiment

A total of 70% from the 2,259 labeled comments comprising 1,626 negative comments, 617 positive comments, and 16 neutral comments is utilized to train the model. This training data will be employed to develop and optimize the classification algorithms used in this research.

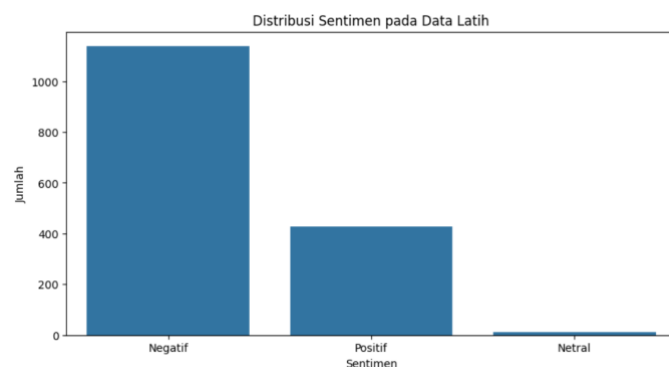


Fig. 5. Distribution of Testing Data Sentiment

The remaining 30% of the 2,259 labeled comments is set aside as testing data. This testing data will be utilized to assess the performance of the trained model, focusing on measuring the accuracy and effectiveness of the algorithms in sentiment classification of the comments.

3.2 Confusion Matrix

Additionally, data weighting is conducted using the TF-IDF method to evaluate the significance of words in a document in relation to a broader collection.. This ensures that relevant words receive higher weights in sentiment analysis.

Classification modeling is conducted with two algorithms: Multinomial Naïve Bayes (NB) and Complement Naïve Bayes (NB). Multinomial NB achieves 75,70% accuracy, indicating effective word identification but with room for improvement. In contrast, Complement NB, which addresses class imbalance, reaches 80,47% accuracy, demonstrating greater effectiveness. To evaluate performance, a confusion matrix illustrates how the model classifies data based on actual categories and analyzes classification errors, providing insights into correct and incorrect predictions for each sentiment category positive, negative, and neutral thereby highlighting the strengths and weaknesses of each algorithm.

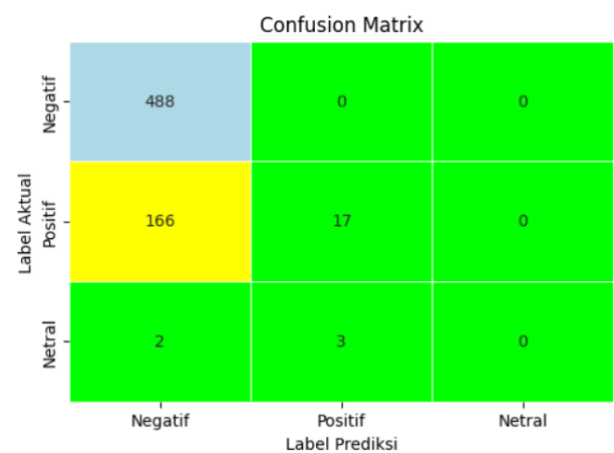


Fig. 6. Confusion Matrix of Multinomial Naive Bayes

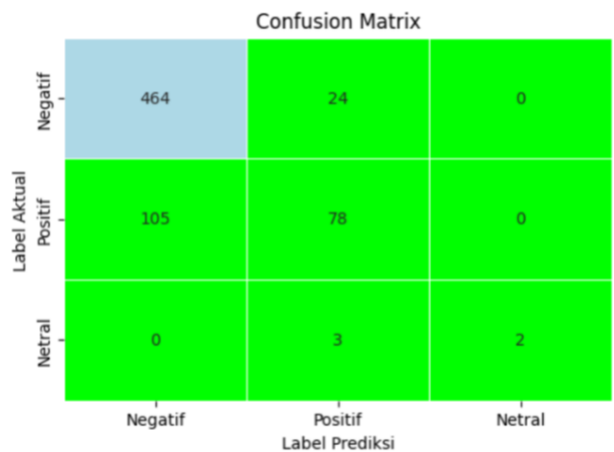


Fig. 7. Confusion Matrix of Complement Naive Bayes

The confusion matrix results for the Multinomial Naïve Bayes and Complement Naïve Bayes algorithms highlight significant differences in performance.

For the Multinomial Naïve Bayes, the F1-score for the negative category is 0.85, while it struggles with the neutral category, achieving an F1-score of 0.00. The positive category has an F1-score of 0.17, resulting in an overall accuracy of 0.75. While it effectively identifies negative comments, it fails to classify neutral and positive comments accurately.

In contrast, the Complement Naïve Bayes shows improved performance with an F1-score of 0.88 for the negative category. The overall accuracy of this model is 0.80, effectively identifying negative comments, but it still encounters challenges in classifying positive comments. This comparison illustrates that the Complement Naïve Bayes algorithm performs better overall, particularly in handling negative sentiments.

3.3 Classification Method

Table 6. Classification Report for Multinomial Naïve Bayes Results

	Precision	Recall	F1-Score	Support
Negatif	0.74	0.95	0.85	488
Positif	0.00	0.45	0.00	5
Netral	0.85	0.43	0.17	183
accuracy			0.75	676
macro avg	0.53	0.36	0.34	676
weighted avg	0.77	0.75	0.66	676

The Multinomial Naïve Bayes model achieves an overall accuracy of 75.70%, with strong performance in identifying negative comments (F1-score of 0.85), but poor results for neutral (F1-score of 0.00) and positive sentiments (F1-score of 0.17), indicating a significant challenge in accurately classifying these categories.

Table 7. Classification Report for Complement Naïve Bayes Results

	Precision	Recall	F1-Score	Support
Negatif	0.82	1.00	0.88	488
Positif	1.00	0.00	0.57	5
Netral	0.74	0.09	0.54	183
accuracy			0.80	676
macro avg	0.85	0.59	0.66	676
weighted avg	0.80	0.80	0.78	676

The Complement Naïve Bayes model achieves an overall accuracy of 80.47%, with strong performance in identifying negative comments (F1-score of 0.88) and positive sentiments (F1-score of 0.54), while it shows limitations in the neutral category (F1-score of 0.57), indicating better overall classification performance compared to the Multinomial Naïve Bayes model.

Complement Naive Bayes (CNB) was identified as the best-performing algorithm in this study, offering several advantages relevant to sentiment analysis in bullying cases. One of the primary reasons for its effectiveness is its ability to handle imbalanced class distributions. The dataset used in this research showed a significant dominance of negative sentiment at 72.2%, with positive sentiment at 20.7% and neutral sentiment at only 0.7%. CNB is specifically designed to address this imbalance by comparing word frequencies across classes, making it more sensitive to minority classes compared to other algorithms, such as Multinomial Naive Bayes (MNB).

Moreover, CNB focuses on reducing misclassification errors for minority classes, which in this study are positive and neutral sentiments. This capability significantly improves evaluation metrics such as precision, recall, and F1-score for minority classes, overcoming a common limitation of other algorithms. In the context of sentiment analysis for bullying

cases, this is particularly important to ensure a balanced representation of public perceptions.

CNB also demonstrates high relevance in analyzing bullying cases. Sentiment analysis in such cases is often dominated by negative sentiment, as public comments tend to express empathy, anger, or criticism regarding the incidents. CNB's mechanism, which accounts for imbalanced class distributions, allows it to capture these sentiment patterns more effectively. The findings of this study revealed that CNB achieved an accuracy of 80.47%, with strong evaluation metrics for the negative sentiment class, including a precision of 0.81, recall of 0.93, and F1-score of 0.87. These results indicate that CNB can effectively identify the majority negative sentiment without compromising performance for the minority classes.

Another advantage of CNB is its optimal performance when combined with the Term Frequency-Inverse Document Frequency (TF-IDF) method. TF-IDF assigns appropriate weights to important words in the text, enabling CNB to accurately distinguish sentiment patterns among different classes. This combination further enhances the algorithm's effectiveness in sentiment analysis for the dataset used in this study.

With these various advantages, CNB is the most suitable algorithm for this study. Its ability to handle imbalanced class distributions, its focus on improving the performance of minority classes, and its capability to capture dominant sentiment patterns in social issues such as bullying make it a reliable method for understanding public perceptions of bullying cases.

3.4 Visualization WordCloud

Then, the visualization results for the most frequently appearing words are presented in the following Wordcloud

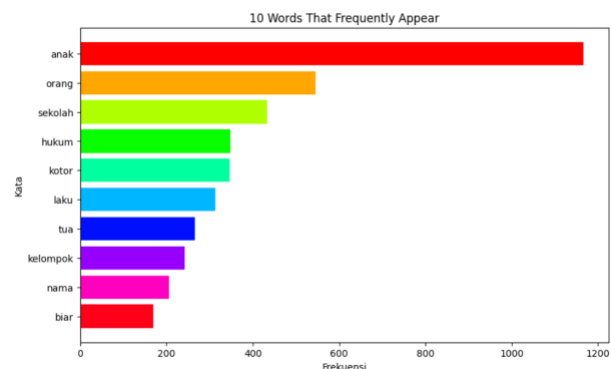


Fig. 8. Word Frequency Distribution



Fig. 9. WordCloud

4. CONCLUSION

This conclusion summarizes the key findings, strengths, limitations, and future research directions in the sentiment analysis of bullying cases at High School Binus Serpong, offering insights into the effectiveness of various classification algorithms.

- The study effectively analyzed 2,259 comments, utilizing Multinomial; Naïve Bayes and Complement Naïve Bayes algorithms to achieve significant results.
- The Multinomial Naïve Bayes model reached an accuracy of 74.70%, with a top F1-score of 0.85 for the negative category but struggled with neutral (0.00) and positive (0.17) categories.
- In contrast, the Complement Naïve Bayes model demonstrated enhanced performance, attaining an accuracy of 80.47% and an F1-score of 0.88 for the negative category, though it encountered challenges in the positive category (0.54).
- Limitations include difficulties in classifying neutral and positive sentiments, indicating the need for improvements.
- Future research should explore advanced NLP techniques, utilize larger and more diverse datasets, and tackle sentiment classification challenges related to bullying.
- The findings enhance understanding of sentiment analysis methodologies and lay the groundwork for future improvements in sentiment identification for sensitive issues like bullying.

REFERENCES

- [1] A. Mardhiah, A. Muana Husniati, C. Andyna, and C. Puspasari, "Penguatan Karakter Diri Sebagai Kunci Mengatasi Perundungan di Lingkungan Sman 7 Lhokseumawe," Penguatan Karakter Diri (Ainol Mardhiah, Dkk.) [353 Jurnal Malikussaleh Mengabdikan, vol. 2, no. 2, pp. 2829–6141, 2023, doi: 10.29103/jmm.v2n2.13349. <https://doi.org/10.29103/jmm.v2n2.13349>
- [2] M. Alfarizi, M. Rizqy, R. I. Ghufroni, D. Fathurahman, R. D. Saputra, and F. Kurniawan, "Analisis Sentimen Persepsi Publik Terhadap Kasus Bullying Siswa Cilacap Menggunakan Pendekatan Machine Learning," Journal of Information Technology Ampera, vol. 4, No. 3, pp. 2774–2121, 2023. [Online]. <https://journal-computing.org/index.php/journal-ita/index>
- [3] E. Agisyaputri, N. A. Nadhirah, and I. Saripah, "Identifikasi Fenomena Perilaku Bullying Pada Remaja," Jurnal Bimbingan Konseling dan Psikologi, vol. 3, no. 1, 2023. <https://jurnal.stkipmb.ac.id/index.php/jubikops/article/view/201>
- [4] J. Keperawatan Profesional, M. Tri Bagas Romadhoni, M. Junnatul Azzizah Heru, A. Rofiqi, Z. Warquatul Hasanah, and V. Anda Yani, "Pengaruh Perilaku Bullying Terhadap Interaksi Sosial Pada Remaja," Jurnal Keperawatan Profesional (JKP), vol. 11, pp. 2685–1830, 2023. <https://ejournal.unuja.ac.id/index.php/jkp/article/view/5545>
- [5] H. S. Rifai, S. Febrianti, and I. Santoso, "Analisis Sentimen Tanggapan Masyarakat Terhadap Cyberbullying di Media Sosial Menggunakan Algoritma Naïve Bayes (NB)." Jurnal IKRAITH-INFORMATIKA, vol. 7 no. 2, pp. 2654–8054, 2023 [Online]. Available: <https://journals.upi-yai.ac.id/index.php/ikraith-informatika/issue/archive>
- [6] D. Kurniawan and D. M. Yasir, "Optimization Sentiment Analysis Using Crisp-Dm And Naïve Bayes Methods Implemented On Social Media", Jurnal Pendidikan Teknologi Informasi, vol. 6, no. 2, pp. 2597–9671, 2022. <https://jurnal.ar-raniry.ac.id/index.php/cyberspace/article/view/12793>
- [7] R. Hilma, M. Ula, and S. Fachrurrazi, "Analisis Sentimen Cyberbullying pada Media Sosial Twitter Menggunakan Metode Support Vector Machine dan Naïve Bayes Classifier", Jurnal Teknik Informatika (TECHSI), vol. 14, no. 2, 2023. <https://ojs.unimal.ac.id/techsi/article/view/12103/0>
- [8] A. Safira, A. S. Masyarakat, and F. N. Hasan, "Analisis Sentimen Masyarakat Terhadap Paylater Menggunakan Metode Naive Bayes Classifier," Jurnal Sistem Informasi, vol. 5, no. 1, 2023. <https://journal.unilak.ac.id/index.php/zn/article/view/12856>
- [9] R. A. Fauzan, "Analisis Sentimen Komentar Youtube Program Kampus Merdeka Berbasis Web Menggunakan Algoritma Multinomial Naïve Bayes", "Seminar Nasional Mahasiswa Fakultas Teknologi Informasi (SENAFTI), 30 Agustus 2023-Jakarta", vol. 2, no. 2, pp. 2962–8628 2023. <https://senafiti.budiluhur.ac.id/index.php/senafiti/article/view/929>
- [10] D. N. Larasakti, A. Aziz, and D. Aditya, "Analisis Sentimen Komentar Video Youtube Dengan Metode K-Nearest Neighbor," Jurnal Ilmiah Wahana Pendidikan, vol. 2023, no. 5, pp. 132–142, doi: 10.5281/zenodo.7728573. <https://doi.org/10.5281/zenodo.7728573>
- [11] J. Khatib Sulaiman, F. Baehaqi, N. Cahyono, and U. Amikom Yogyakarta, "Analisis Sentimen Terhadap Cyberbullying Pada Komentar di Instagram Menggunakan Algoritma Naïve Bayes," Indonesian Journal of Computer Science, vol. 13, no. 1, pp. 2549–7286, 2024. [Online]. doi: 10.33022/ijcs.v13i1.3301 <https://doi.org/10.33022/ijcs.v13i1.3301>
- [12] A. Putri and A. Muzakir, "How to cite: Analisis Sentimen Cyberbullying Kpop di Media Sosial Twitter Menggunakan Metode Naive Bayes," Jurnal Ilmiah Indonesia, vol. 7, no. 9, 2022. jurnal.syntaxliterat.co.id
- [13] D. Darwis, N. Siskawati, and Z. Abidin, "Penerapan Algoritma Naïve Bayes untuk Analisis Sentimen Review Data Twitter BMKG Nasional," Jurnal Tekno Kompak vol. 15, no. 1, pp. 2656–3525, 2021. <https://ejournal.teknokrat.ac.id/index.php/teknokompak/article/view/744>
- [14] B. Z. Ramadhan, I. Riza, and I. Maulana, "Analisis Sentimen Ulasan Pada Aplikasi E-Commerce Dengan Menggunakan Algoritma Naïve Bayes," Journal of Applied Informatics and Computing (JAIC), vol. 6, no. 2, 2022. [Online]. <http://jurnal.polibatam.ac.id/index.php/JAIC>
- [15] K. Verena, S. Toy, Y. A. Sari, and I. Cholissodin, "Analisis Sentimen Twitter menggunakan Metode Naive Bayes dengan Relevance Frequency Feature Selection (Studi Kasus: Opini Masyarakat mengenai Kebijakan New Normal)," Jurnal Pengembangan Teknologi Informasi dan Ilmu Komputer vol. 5, no. 11 2021. [Online] <https://j-ptiik.ub.ac.id/index.php/j-ptiik/article/view/10172>
- [16] S. Kusuma Wardani and Y. Arum Sari, "Analisis Sentimen menggunakan Metode Naïve Bayes Classifier terhadap Review Produk Perawatan Kulit Wajah menggunakan Seleksi Fitur N-gram dan Document Frequency Thresholding," Jurnal Pengembangan Teknologi Informasi dan Ilmu Komputer vol. 5, no. 12 2021. [Online]. <https://ejournal.teknokrat.ac.id/index.php/teknokompak/article/view/744>
- [17] E. Undamayanti et al., "Analisis Sentimen Menggunakan Metode Naive Bayes Berbasis Particle Swarm Optimization Terhadap Pelaksanaan Program Merdeka Belajar Kampus Merdeka," Jurnal Sains Komputer & Informatika (J-SAKTI), vol. 6, no. 2, 2022. <http://ejournal.tunasbangsa.ac.id/index.php/jsakti/article/view/502>
- [18] S. Azra and R. Tahara Shita, "Analisis Sentimen Menggunakan Metode Naïve Bayes Terhadap Produk Pt.Imin Technology Berdasarkan Ulasan dari Tokopedia", "Seminar Nasional Mahasiswa Fakultas Teknologi Informasi (SENAFTI), 30 Agustus 2023-Jakarta," 2023. <https://senafiti.budiluhur.ac.id/index.php/senafiti/article/view/753>
- [19] N. A. Susanti And M. Walid, "Klasifikasi Data Tweet Ujaran Kebencian di Media Sosial Menggunakan Naive Bayes Classifier," Jurnal Mahasiswa Teknik Informatika (JATI), vol. 6 no. 2 2022. [Online]. doi: <https://doi.org/10.36040/jati.v6i2.5174>
- [20] F. S. Pamungkas and I. Kharisudin, "Analisis Sentimen dengan SVM, Naive Bayes dan KNN untuk Studi Tanggapan Masyarakat Indonesia Terhadap Pandemi Covid-19 pada Media Sosial Twitter" Prosiding Seminar Nasional Matematika (PRISMA), vol. 4, pp. 628–634, 2021, [Online]. Available: <https://journal.unnes.ac.id/sju/index.php/prisma/>
- [21] T. Wiratama Putra and A. Triayudi, "Analisis Sentimen Pembelajaran Daring Menggunakan Metode Naïve Bayes, KNN, dan Decision Tree," Jurnal Teknologi Informasi dan Komunikasi, vol. 6, no. 1, p. 2022, 2022, doi: 10.35870/jti. <http://journal.lembagakita.org/index.php/jtik>
- [22] M. Diqi, D. Rhesa, R. Marselina, E. Hiswati, W. Ordiyasa, and I. Hafizah, "Digital Democracy: Analyzing Political Sentiments through Multinomial Naive Bayes in Election Campaign Ads," Jurnal Sistem

Cerdas, vol. 07, no. 02, 2024. <https://apic.id/jurnal/index.php/jsc/article/view/379>

- [23] D. Sri, R. R. Novita, T. Khairil, and A. Zarnelly, "Sentiment Analysis Chatgpt Using The Multinomial Naïve Bayes Classifier (Nbc) Algorithm," *Jurnal Sistem Cerdas* (2024) vol. 07, no. 01, 2024. <https://apic.id/jurnal/index.php/jsc/article/view/388>
- [24] N. A. Rice, N. Mustakim, and M. Afdal, "Sentiment Analysis on the Impact of Artificial Intelligence (AI) Development to Determine Technology Needs," *Jurnal Sistem Cerdas*, vol. 07, no. 02, 2024. doi : <https://doi.org/10.37396/jsc.v7i2.404>