

Sentiment Analysis of YouTube Comments on the KDM Policy in Handling Juvenile Delinquency Using Naïve Bayes

Muhammad Afan Adi Saputra¹, Wildan Arosyid², Huda Mutamam³

¹Informatics Study Program, Universitas Nurul Huda

¹afanadisaputra5@email.com, ²wildanarosyid05@email.com, ³hudamutamam724@email.com

Abstract

The policy of West Java Governor Dedi Mulyadi (KDM) sending troubled youths to military barracks for character development amid juvenile delinquency triggered diverse public opinions in YouTube comments, which were difficult to analyze manually due to their volume and linguistic variation. This study aims to analyze the sentiment of YouTube comments toward the policy using the Naïve Bayes Classifier algorithm with Term Frequency-Inverse Document Frequency (TF-IDF) weighting, and to evaluate model performance based on accuracy, precision, recall, and F1-score metrics. The research data were obtained from a public Kaggle dataset comprising 7,875 comments, which after preprocessing resulted in 7,407 valid data with a proportion of 56.6% negative and 43.4% positive comments. The data were split using an 80:20 ratio with stratified sampling. Test results show that the Naïve Bayes model achieved an accuracy of 73.95%, precision of 84.64%, recall of 48.83%, and F1-score of 61.93%, with the negative class performing better (F1-score 80%) than the positive class (F1-score 62%) due to class imbalance in the dataset. Word cloud analysis revealed that negative sentiment was largely directed not at the KDM policy itself but at criticism toward the Indonesian Child Protection Commission (KPAI), offering an objective basis for evaluating juvenile delinquency policy.

Keyword : sentiment analysis, YouTube, Naïve Bayes, TF-IDF, public policy

1. INTRODUCTION

The development of information technology has driven the growth of digital platforms, including YouTube, which is the most widely accessed social media platform in Indonesia at 88% [1]. In addition to serving as a means of entertainment and information, YouTube also functions as a digital public space where people express opinions through the comment feature on various issues, including public policy [1].

One social issue that has drawn widespread attention in Indonesia is juvenile delinquency (juvenile delinquency), namely deviant behaviors such as brawls, truancy, drug abuse, and promiscuity triggered by internal and external factors such as identity crises, weak self-control, and family conditions [2]. To address this issue, West Java Governor Dedi Mulyadi (KDM) issued a policy of sending troubled youths to military barracks as a form of discipline-based character development. This policy drew both support and opposition from the public, particularly on the YouTube platform, generating thousands of comments that reflect the dynamics of public opinion. The large volume and linguistic variation of these comments make manual analysis inefficient, so a computational approach is needed [2].

Sentiment analysis is a Natural Language Processing (NLP) approach used to automatically classify opinions into positive or negative categories [3], and has been applied across various contexts, from political figures to government policies [4]. The Naïve Bayes algorithm is one of the most widely used methods due to its implementation simplicity, computational speed, and ability to process high-dimensional short text data such as social media comments [1]. The effectiveness of this algorithm has been demonstrated in various studies, such as sentiment classification of the Kampus Merdeka policy achieving 70% accuracy, a performance comparison between SVM and Gaussian Naïve Bayes on

YouTube Masterchef Indonesia comments that showed a performance improvement through oversampling techniques [4], as well as the application of Random Over Sampling (ROS) in sentiment analysis of a plastic-waste-reduction campaign, which was shown to improve the recall of the minority class [5].

A previous study that specifically examined sentiment toward the KDM policy used an Ensemble Majority Voting approach combining Naïve Bayes, KNN, dan Logistic Regression [2]. However, the pure performance characteristics, stability limitations, and computational efficiency of Naïve Bayes as a single classifier in this case study have not been comprehensively explored. Most other sentiment-analysis studies on YouTube comments also tend to focus on education, entertainment, and environmental issues [6], and generally combine Naïve Bayes with ensemble or oversampling techniques to improve performance [7],[8].

Therefore, the novelty of this study lies in testing the Naïve Bayes Classifier algorithm independently, without ensemble or hybrid methods, to determine how well a basic probabilistic model can classify sentiment in YouTube comments that are unstructured and contain non-standard language, slang, and regional language variation. This study aims to analyze the sentiment of YouTube comments on the KDM policy in handling juvenile delinquency using the Naïve Bayes Classifier, as well as evaluating its performance based on accuracy, precision, recall, and F1-score metrics, so as to provide the government with an objective picture for formulating data-based youth development programs.

2. METHOD

This study applies a sentiment analysis method to determine the tendency of YouTube users' opinions (positive/negative) toward Dedi Mulyadi's (KDM) policy in handling juvenile delinquency, using the Naïve Bayes Classifier algorithm. The research stages include data collection, data labeling, preprocessing, data splitting, feature weighting, classification, and model evaluation, as shown in Figure 1.

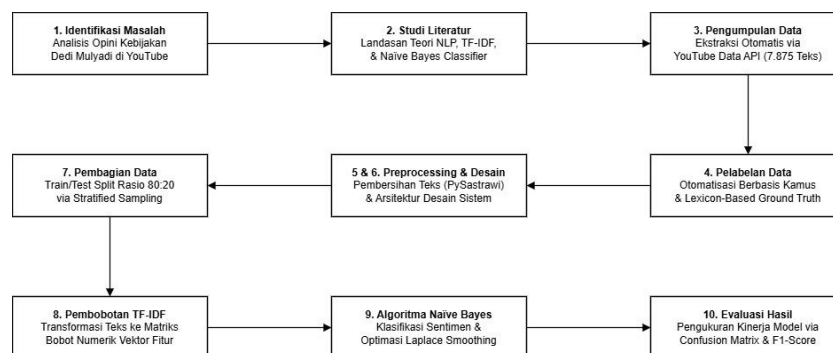


Figure 1. Research Methodology Flow

Problem Identification and Literature Study

The large number of YouTube comments discussing the KDM policy results in a diversity of opinions that is difficult to analyze manually, so automated sentiment analysis is needed to map public response tendencies in support of policy evaluation [9]. YouTube was chosen as the data source because it provides a comment section that the public actively uses to express opinions on public issues [8]. A literature review was conducted on journals, proceedings, and prior research related to sentiment analysis, NLP, TF-IDF, and Naïve Bayes, in order to understand the characteristics of processing informal Indonesian-language text containing a lot of slang [10]. Various studies show that the combination of TF-IDF and Naïve Bayes provides good performance in text-based sentiment classification, including on public policy issues [11].

Data Collection

Data were obtained from a public dataset on the Kaggle <https://www.kaggle.com/datasets/rizb11/indonesian-YouTube-public-policy-sentiment-dataset> platform, containing YouTube comments in Indonesian related to various public policy issues, collected using the YouTube Data API. From that dataset, specific filtering was carried out to extract comments related to the KDM policy on handling juvenile delinquency during the May–November 2025 period, yielding 7,875 comments in raw form. The use of a public dataset supports efficiency, transparency, and ease of verification and replication of the research [12],[13].

Data Labeling

The dataset has built-in labeling (ground truth) based on a dictionary-based (lexicon-based) approach, in which each comment is evaluated based on its sentiment score: a score < 0 is categorized as negative, and a score ≥ 0 is categorized as positive [14]. This approach ensures the objectivity of the labeling without any subjective input from the researcher [15].

Data Preprocessing

Preprocessing was carried out to clean and normalize the text through five stages [12]:

- a. **Case Folding**
Converting all letters to lowercase.
- b. **Cleaning**
Removing URLs, mentions, hashtags, numbers, and punctuation.
- c. **Tokenization**
Splitting text into word tokens.
- d. **Stopword Removal**
Removing common low-contribution words (e.g., “yang”, “dan”, “di”).
- e. **Stemming**
Converting affixed words to their root form using the library PySastrawi, library, which has proven accurate for Indonesian morphology and improves the quality of feature representation [11].

Data Splitting (Train/Test Split)

The dataset was split with an 80:20 ratio using the stratified sampling technique to maintain class proportions and avoid class imbalance [2]. This ratio is a common practice that has proven to provide stable evaluation results in similar studies [16].

TF-IDF Weighting

Text features were converted into numerical representations using Term Frequency-Inverse Document Frequency (TF-IDF), with the following equations:

$$TF(d, t) = f(d, t) \quad (1)$$

$$IDF(t) = \ln\left(\frac{D + 1}{df + 1}\right) + 1 \quad (2)$$

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

TF-IDF was chosen because it can highlight informative words and reduce the influence of words that appear too frequently, thereby producing better features than a simple bag-of-words approach [11]. The TF-IDF weighting process applies a fit-transform scheme, in which the vectorizer is fitted (learning

the vocabulary and IDF values) only on the training data, then applied (transformed) to the test data to prevent information leakage (data leakage) from the test data into the model training stage.

Naïve Bayes Algorithm

Classification was carried out using the Naïve Bayes Classifier based on Bayes' Theorem with the assumption of independence between features:

$$V_{MAP} = \arg \max_{v_j \in V} P(v_j) \prod_{i=1}^n P(x_i | v_j) \quad (4)$$

To address zero probability for words that do not appear in the training data, Laplace Smoothing was used with a smoothing parameter (alpha) of 1.0, the default value without additional tuning:

$$P(x_i | v_j) = \frac{\text{count}(x_i, v_j) + 1}{\text{count}(v_j) + |V|} \quad (5)$$

Naïve Bayes was chosen for its low complexity, fast computation, and competitive performance in Indonesian-language text classification, despite being simpler than deep learning approaches [16].

Evaluation Results

Evaluation was carried out using a confusion matrix and four metrics: accuracy, precision, recall, and F1-score [8]:

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} \quad (6)$$

$$\text{Precision} = \frac{TP}{TP + FP} \quad (7)$$

$$\text{Recall} = \frac{TP}{TP + FN} \quad (8)$$

$$F1 - \text{Score} = \frac{2 \times \text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (9)$$

3. RESULTS AND DISCUSSION

Sentiment Data Distribution

This study uses public data from the Kaggle platform, filtered to obtain comments related to Dedi Mulyadi's (KDM) policy on handling juvenile delinquency, yielding 7,875 raw comments. After data cleaning to remove duplicates and empty rows, 7,407 valid comments were obtained for use in the model training and testing stages. The data distribution is shown in Table 1.

Table 1. Sentiment Label Class Distribution of YouTube Comments

Sentiment Category (Label)	Initial Data Volume (Rows)	Final Clean Data Volume (Rows)	Final Percentage Proportion (%)
Negative (0)	4,282	4,194	56.62%
Positive (1)	3,593	3,213	43.38%
Total Dataset	7,875	7,407	100.00%

The larger number of negative-labeled comments indicates a tendency for the public to voice more criticism of the KDM policy through the YouTube comment section.

Data Preprocessing Results

Preprocessing was carried out through five stages (case folding, cleaning, tokenization, stopword removal, and stemming using PySastrawi), with the stemming process time for 7,407 comments being about 16 minutes 27 seconds. An example of the transformation results is shown in Table 2.

Table 2. Sample of Text Preprocessing Transformation Results on the Dataset

No	Text Clean (Initial Cleaned Data)	Tokenization & Stopword Removal	Stemming (Final Root Word)	Label
0	keren pidato kdm gubenur jabar	['keren', 'pidato', 'kdm', 'gubenur', 'jabar']	keren pidato kdm gubenur jabar	1
1	pejabat ningir mampu membangun rakyat	['pejabat', 'ningir', 'mampu', 'membangun', 'rakyat']	jabat ningir mampu bangun rakyat	0
2	palembang iri banget jabar semoga gubernur mengikutinya amin	['palembang', 'iri', 'banget', 'jabar', 'semoga', 'gubernur', 'mengikutinya', 'amin']	palembang iri banget jabar moga gubernur ikut amin	0
3	kdm gubernur keren	['kdm', 'gubernur', 'keren']	kdm gubernur keren	1
4	semoga ucapan pidatoya bpk menjadilah pelajaran dngpegawai indonesia semoga negara bersatulah mengerti mendidik negara amin	['semoga', 'ucapan', 'pidatoya', 'bpk', 'menjadilah', 'pelajaran', 'dngpegawai', 'indonesia', 'semoga', 'negara', 'bersatulah', 'mengerti', 'mendidik', 'negara', 'amin']	moga ucap pidatoya bpk jadi ajar dngpegawai indonesia moga negara satu erti didik negara amin	1

The stemming stage proved effective in simplifying affixed words; for example, “pejabat” and “membangun” became “jabat” and “bangun” (index 1), allowing words with the same meaning to be recognized by the model as a single feature.

Data Splitting

The dataset was split with an 80:20 ratio using stratified sampling, resulting in 5,925 training data and 1,482 testing data (Table 3), in line with common practice shown to produce stable evaluation results in similar studies [17].

Table 3. Training and Testing Data Split

Data Subset	Proportion	Data Count
Training Data	80%	5,925
Testing Data	20%	1,482
Total		7,407

TF-IDF Weighting Results

Feature extraction was carried out using the Term Frequency-Inverse Document Frequency (TF-IDF) method to convert text data into numerical vector representations that can be processed by the classification algorithm. TF-IDF gives higher weight to words that frequently appear in a document but rarely appear across all other documents, allowing words that truly distinguish a document to stand out more. The TF-IDF weighting results are shown in Table 4.

Table 4. TF-IDF Weighting Results

Description	Value
Number of Documents (Training Data)	5,925
Number of Features (Unique Vocabulary)	9,915

Model Evaluation Results and Confusion Matrix

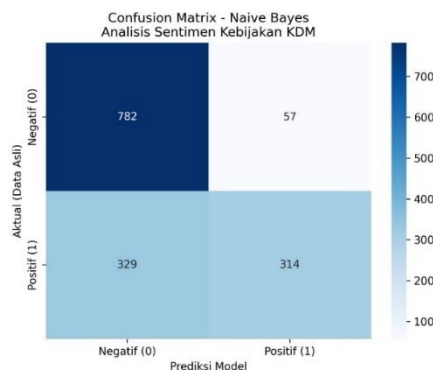


Figure 2. Naive Bayes Classification Confusion Matrix Chart

Table 5. Confusion Matrix Details

	Predicted Negative	Predicted Positive	
Actual Negative	782 (TN)	57 (FP)	Actual Negative
Actual Positive	329 (FN)	314 (TP)	Actual Positive

Of the 839 actual negative comments, 782 were correctly predicted (TN) and 57 were incorrectly predicted as positive (FP). Of the 643 actual positive comments, only 314 were correctly predicted (TP), while 329 were incorrectly predicted as negative (FN).

Table 6. Naive Bayes Model Evaluation Results

Performance Evaluation Metric	Decimal Score Value	Performance Value Percentage (%)
Accuracy (Accuracy)	0.7395	73.95%
Precision (Precision)	0.8464	84.64%
Recall	0.4883	48.83%
F1-Score	0.6193	61.93%

Table 7. Classification Report per Class

Class	Precision	Recall	F1-Score	Support
Negative (0)	70%	93%	80%	839
Positive (1)	85%	49%	62%	643
Macro Avg	78%	71%	71%	1,482

Model Naïve Bayes achieved an accuracy of 73.95%, precision of 84.64%, recall of 48.83%, and F1-score of 61.93%, meaning the model correctly classifies about 74 out of every 100 comments. There is a performance gap between the classes: the recall of the negative class reaches 93%, while the recall of the positive class is only 49%, so the F1-score of the positive class (62%) is much lower than that of the negative class (80%).

This gap is caused by two factors: (1) class imbalance in the dataset (negative 4,194 vs. positive 3,213), which causes the model to learn more from patterns in negative comments; and (2) the characteristics of positive comments on social media, which tend to be ambiguous and implicit, combined with the feature-independence assumption in Naïve Bayes, which limits the model's ability to fully capture sentence context. To address this limitation, future research could consider oversampling techniques such as SMOTE, or combining Naïve Bayes with other methods, as done in [2] and [4], which have been shown to improve classification performance.

Analysis of Public Opinion Findings Using Word Cloud



Figure 3. Word Cloud Visualization of the Positive Sentiment Document Corpus

Word cloud from 3,213 positive comments shows the dominance of the words kdm, pimpin, gubernur, bapa, pidato, and jabar, indicating public support and appreciation for KDM as the leader of West Java. The appearance of the words mantap, maju, bangga, and terima kasih confirms public approval of the policy.



Figure 4. Word Cloud Visualization of the Negative Sentiment Document Corpus

Word cloud from 4,194 negative comments is dominated by the words *kpai* and *kdm*, followed by *bagus*, *sedih*, *nyinyir*, *pidato*, and *tawa-tawa*. The high frequency of the words *kpai* and *nyinyir* indicates that the negative sentiment is largely not a rejection of the KDM policy itself, but rather criticism directed at the Indonesian Child Protection Commission (KPAI), which was seen as opposing the development program. The appearance of positively connoted words such as *bagus* and *pimpin* in negative comments occurs because of comparative sentence patterns, for example in the comment “Kebijakan KDM sudah bagus, KPAI tidak usah nyinyir” (“The KDM policy is already good, KPAI doesn’t need to nag”), the high negative weight of the word *nyinyir* causes the model to classify this comment as negative.

Discussion

Model Naïve Bayes Classifier as a single classifier without data-balancing techniques achieved an accuracy of 73.95% and a precision of 84.64%, indicating that when the model predicts a comment as positive, that prediction is mostly correct. However, the recall of the positive class is only 48.83% (F1-score 61.93%), due to class imbalance in the dataset (negative 56.62% vs. positive 43.38%), which makes the model more sensitive to negative patterns, as reflected in the negative class recall of 93%. The application of Laplace Smoothing proved effective in maintaining the stability of the probability calculations, so the model did not produce zero values when encountering new vocabulary in the test data. Quantitatively, this model’s accuracy (73.95%) is slightly higher than that of the Kampus Merdeka policy sentiment analysis study, which reported an accuracy of 70% [4], indicating that the Naïve Bayes algorithm as a single classifier, without ensemble or hybrid techniques, remains competitive for public-policy sentiment classification based on YouTube comments.

From the perspective of public opinion, the word cloud results reveal that public support for KDM (Figure 3) is fairly strong, while the negative sentiment (Figure 4) is largely not directed at the KDM policy itself, but rather at KPAI, which was seen as being overly vocal in criticizing the development program. This finding indicates that public perception of the KDM policy is substantively positive, even though it was recorded as negative sentiment due to the context of criticism directed at a third party.

4. CONCLUSION

This study aims to analyze the sentiment of YouTube comments on Dedi Mulyadi’s (KDM) policy in handling juvenile delinquency using the Naïve Bayes Classifier algorithm with TF-IDF weighting, as well as evaluating its performance based on accuracy, precision, recall, and F1-score metrics. Based on the research results, the following conclusions can be drawn.

First, of the 7,875 comments obtained from the public Kaggle dataset, 7,407 valid comments were obtained after preprocessing, with a distribution of 4,194 negative-sentiment comments (56.6%) and 3,213 positive-sentiment comments (43.4%). The word cloud visualization results show that the dominant negative sentiment is largely not a direct rejection of the KDM policy, but rather criticism and sarcasm directed at the Indonesian Child Protection Commission (KPAI), which was seen as opposing the youth development program.

Second, the Naïve Bayes Classifier model, implemented as a single classifier without ensemble methods or oversampling techniques, achieved an accuracy of 73.95%, precision of 84.64%, recall of 48.83%, and F1-score of 61.93%. The application of Laplace Smoothing proved effective in maintaining the stability of the probabilistic computation, so the model did not produce zero probability values when encountering new vocabulary in the test data. This result shows that Naïve Bayes is able to perform competitively on an unstructured Indonesian-language social media comment dataset, even when used without any additional algorithm combination.

Third, the model shows unbalanced performance between the two classes, with performance on the negative class (F1-score 80%) being much better than the positive class (F1-score 62%). This

condition is caused by class imbalance in the dataset as well as the limitations of the feature-independence assumption in Naïve Bayes, which cannot fully capture sentence context, especially in comments containing ambiguous words.

Based on these limitations, future research is recommended to apply class-imbalance handling techniques such as oversampling (SMOTE), as well as to consider using more complex methods such as ensemble learning or deep learning-based approaches (LSTM or BERT) that can better understand sentence context. The results of this study are also expected to serve as a consideration for the government in evaluating juvenile delinquency handling policies to make them more effective and data-driven.

5. BIBLIOGRAPHY

- [1] S. M. Harahap and R. Kurniawan, "Analisis Sentimen Komentar Youtube Terhadap Food Vlogger Dengan Menggunakan Metode Naïve Bayes," *MEANS (Media Inf. Anal. dan Sist.*, vol. 9, no. 1, pp. 87-96, 2024, doi: 10.54367/means.v9i1.3912.
- [2] R. B. Aprianto, A. Eviyanti, H. Setiawan, and I.R.I. Astutik, "Sentiment Analysis of YouTube Comments on Dedi Mulyadi's Policies Using Machine Learning Methods (Case Study: Juvenile Delinquency) [Analisis Sentimen Komentar Youtube Terhadap Kebijakan Dedi Mulyadi Dengan Metode Machine Learning (Studi Kasus: Kenakalan Remaja)]," pp. 1-10. doi: <https://doi.org/10.21070/ups.10439>.
- [3] M. R. Herdiansyah and A. Yuliana, "Analisis Sentimen Kebijakan Kampus Merdeka Menggunakan Naïve Bayes Berdasarkan Komentar Pada Youtube," vol. 8, no. 6, pp. 12454-12459, 2024. doi: <https://doi.org/10.36040/jati.v8i6.11963>.
- [4] D. Marganingsih, H. Oktavianto, and G. Abdurrahman, "Analisis Sentimen Komentar Youtube Masterchef Indonesia Menggunakan Algoritma Support Vector Machine dan Gaussian Naïve Bayes," *J. Inform. dan Teknol. Pendidik.*, vol. 5, no. 1, pp. 16-26, 2025, doi: 10.59395/jitp.v5i1.117.
- [5] N. D. H. Sadikin and S. Susanti, "Analisis Sentimen Publik Terhadap Kampanye Pengurangan Sampah Plastik Menggunakan Algoritma Naïve Bayes," *J. FASILKOM*, vol. 15, no. 2, pp. 202-212, 2025. doi: 10.37859/jf.v15i2.9574.
- [6] Z. Fatah and M. T. Ismardani, "Analisis Sentimen Komentar Pengguna YouTube Terhadap Kebijakan Penytiaan Tanah oleh Pemerintah Menggunakan Metode Naïve Bayes," *J. Inform. dan Sist. Inf.*, vol. 12, no. 1, pp. 59-66, 2026, doi: 10.37715/juisi.v12i1.6161.
- [7] A. Angdressey, L. Sitanayah, and I. L. H. Tangka, "Sentiment Analysis for Political Debates on YouTube Comments using BERT Labeling, Random Oversampling, and Multinomial Naïve Bayes," *J. Comput. Theor. Appl.*, vol. 2, no. 3, pp. 342-354, 2025, doi: 10.62411/jcta.11668.
- [8] M. Monica and A. Purwanto, "Evaluasi Komparatif Kinerja Algoritma Naïve Bayes dan Logistic Regression dalam Klasifikasi Sentimen Komentar YouTube Berbahasa Indonesia: Comparative Performance Evaluation of Naïve Bayes and Logistic Regression for Indonesian YouTube Comment Sentiment C," *MALCOM Indones. J. Mach. Learn. Comput. Sci.*, vol. 6, no. 2, pp. 934-943, 2026, [Online]. doi: <https://doi.org/10.57152/malcom.v6i2.2639>.
- [9] N. A. Eka Budiarti, "Sentiment Analysis of Indonesia's Free School Lunch Policy Using LSTM and Word2Vec on YouTube Comments," *J-KOMA J. Ilmu Komput. dan Apl.*, vol. 7, no. 2, pp. 28-36, 2025, doi: 10.21009/j-koma.v7i2.04.
- [10] W. B. Zulfikar, A. R. Atmadja, and S. F. Pratama, "Sentiment Analysis on Social Media Against Public Policy Using Multinomial Naive Bayes," *Sci. J. Informatics*, vol. 10, no. 1, pp. 25-34, 2023, doi: 10.15294/sji.v10i1.39952.
- [11] R. Wati, S. Ernawati, and H. Rachmi, "Pembobotan TF-IDF Menggunakan Naïve Bayes pada Sentimen Masyarakat Mengenai Isu Kenaikan BIPIH," *J. Manaj. Inform.*, vol. 13, no. 1, pp.

84-93, 2023, doi: 10.34010/jamika.v13i1.9424.

- [12] F. Yulistira and A. Rahman Isnain, "Analisis Sentimen Terhadap Seleksi CPNS Tahun 2024 Berbasis Media Sosial X Menggunakan Algoritma Naïve Bayes SENTIMENT ANALYSIS OF THE 2024 CPNS SELECTION BASED ON SOCIAL MEDIA X USING THE NAÏVE BAYES ALGORITHM," *J. Pendidik. dan Teknol. Indones.*, vol. 5, no. 3, pp. 887-897, 2025, [Online]. Available: <https://doi.org/10.52436/1.jpti.731>
- [13] I. Najiyah and M. Rizal, "Analisis Sentimen Media Sosial X Terhadap Kebijakan Presiden Republik Indonesia Prabowo Subianto," *J. Fasilkom*, vol. 15, no. 3, pp. 443-455, 2025, doi: 10.37859/jf.v15i3.10385.
- [14] F. Koto and G. Y. Rahmaningtyas, "Inset lexicon: Evaluation of a word list for Indonesian sentiment analysis in microblogs," *Proc. 2017 Int. Conf. Asian Lang. Process. IALP 2017*, vol. 2018-Janua, pp. 391-394, 2017, doi: 10.1109/IALP.2017.8300625.
- [15] C. H. Lin and U. Nuha, "Sentiment analysis of Indonesian datasets based on a hybrid deep-learning strategy," *J. Big Data*, vol. 10, no. 1, 2023, doi: 10.1186/s40537-023-00782-9.
- [16] R. Firdaus, R. Al Hariri, and H. F. Amran, "Sentimen Analisis Masyarakat Tentang Penetapan Hari Raya Idul Adha Tahun 2023 Pada Video Youtube Menggunakan Algoritma Random Forest dan Support Vector Machine," *J. Fasilkom*, vol. 14, no. 1, pp. 278-285, 2024, doi: 10.37859/jf.v14i1.7012.
- [17] N. Mahfudza and M. Ihksan, "Sentiment Analysis of Youtube Comments on Indonesian Presidential Candidates in 2024 using Naïve Bayes Classifier," *JURIKOM (Jurnal Ris. Komputer)*, vol. 12, no. 2, pp. 140-148, 2025, doi: 10.30865/jurikom.v12i2.8538.