



Research article

Public Sentiment Analysis from Twitter (X) About Dissolving the House With *Naive Bayes* & *Svm*

M. Restu Habiburrahman¹, Mohd. Fitra Ramadhan², Al Fauzan Hakiki³

Information Systems, Faculty of Engineering and Computer Science, Indragiri Islamic University, Tembilahan City, Indragiri Hilir, Riau, Indonesia

E-mail: ¹restuhabiburrahman@gmail.com, ²mohdfitrarmdn13@gmail.com, ³alfauzanhakiki2004@gmail.com

* Correspondence

ARTICLE INFO

Article history:

Received November 26, 2025

Revised December 23, 2025

Accepted June 1, 2026

Available online June 2, 2026

Keywords:

Sentimen Analysis

Naive Bayes

Twitter (X)

Text Classification

ABSTRACT

Twitter (X) has become a major platform for the expression of public opinion, including political issues such as hashtags #BubarkanDPR reflecting public dissatisfaction with the legislature. This study aims to (1) analyze public sentiment on the issue of "Dissolve the DPR" and (2) compare the performance of the Naive Bayes (NB) and Support Vector Machine (SVM) classification methods in the context of Indonesian texts. This study used 1989 tweet data obtained through *web scraping*. Data through the *preprocessing* stage (cleaning, normalization, *stopword removal*, *stemming*) and automatic labeling (positive, negative, neutral) using the *RoBERTa pre-trained model*. Features were extracted using TF-IDF, then the data was divided into 80% of the training data and 20% of the test data to train the NB and SVM models. The results of the evaluation showed that the Support Vector Machine (SVM) provided superior performance with an accuracy of 79%, compared to Naive Bayes (NB) with an accuracy of 74%. SVM also showed a more balanced F1-Score for negative (82%) and neutral (80%) classes. Both models showed significant difficulty in classifying positive sentiment (F1-Score SVM 12%; NB 0%), which was identified as a result of severe data imbalance in the dataset. In conclusion, SVM is more effective than Naive Bayes for this classification task, and the analysis shows that public sentiment towards #BubarkanDPR issue is dominated by negative and neutral polarity.

1. Introduction

In the rapidly growing digital era, media transformation and communication dynamics have become the main spotlight in the communication discipline. The development of information and communication technology has changed the way we interact, communicate, and access information [1]. Platforms such as Twitter (X), Facebook, and Instagram allow people to express their opinions openly and in real-time on ongoing issues. Among the various platforms, Twitter (X) has become the most popular platform in the world. A social channel of choice that can be used to extract sentiment-rich data. First launched in 2006, Twitter (X) is currently the most popular and influential microblogging service connecting millions of people around the world. Twitter (X) is a social networking microblogging media service that is available for free, which allows registered users to broadcast short messages or posts called "tweets" to other registered users in real-time [2].

Social phenomena that are widely discussed on Twitter (X) are often related to political issues and public policy. One of the highlights in recent years is the emergence of the hashtag #BubarkanDPR, which describes the expression of public dissatisfaction with the legislature in Indonesia. This issue is rooted in the public perception that the performance of the House of Representatives is often not in line with the interests of the people, especially when there is a polemic in the process of passing laws or cases of ethical violations among council members. This condition is reinforced by the characteristics of Twitter (X) users who tend to be active, responsive, and fond of arguments. User X plays an important role in disseminating information, voicing opinions, and shaping public discourse, especially regarding trending socio-political issues [3].

To understand the tendency of public opinion, a scientific approach is needed that is able to transform unstructured text data into meaningful information. One of the most widely used approaches is sentiment analysis, which is an approach that uses Natural Language Processing (NLP) to extract, convert, and interpret opinions from a text and classify it [4]. Sentiment analysis allows the identification of people's emotional patterns (positive, negative, or neutral) towards an event or figure, so that it can be interpreted in a broader social and political framework. This approach has been widely used in various fields such as digital marketing, public services, and

government policies. Through sentiment analysis, government agencies and researchers can understand public perception objectively based on real data taken from social media.

In the context of academic research, various previous studies have proven the effectiveness of sentiment analysis in understanding public opinion. Go et al introduce the *distant supervision* to classify tweets based on emoticons, thus allowing the formation of large datasets without manual labeling. Pak and Paroubek then used Twitter (X) as a corpus for opinion analysis, proving that the informal language and abbreviations typical of Twitter users (X) are a challenge in the text classification process. Meanwhile, Pang and Lee developed a theoretical model in *opinion mining* which explains how language patterns can represent a person's emotions and attitudes towards an entity [5].

Commonly used methods in sentiment analysis are the Naive Bayes Classifier (NBC) and the Support Vector Machine (SVM). These two algorithms are included in the approach *supervised learning*, where the model is trained with labeled data to recognize text patterns that indicate specific sentiments. Naïve Bayes has proven to be accurate in determining the true polarity of a sentence, even in an unbalanced dataset. In addition, the model has a high bias and low variance classifier that works well even in small datasets [6]. Meanwhile, the implementation of the Support Vector Machine (SVM) method in the sentiment analysis study of the Digital Korlantas Polri application with 1,200 reviews showed good performance, with an accuracy rate of 82% in a data ratio of 90:10 [11]. Both studies show that Naïve Bayes and SVM have an effective ability to classify sentiment [7].

Various studies have compared the performance of the two methods. The research conducted by Moh. Aminullah Al Fachri and Umami Athiyah to compare the two algorithms in conducting sentiment analysis on the issue of high cooking oil prices using 9,194 datasets obtained an accuracy of 81.6% for SVM and 74.06% for Naïve Bayes. Previous research used customer opinions from tweet data based on positive or negative sentiment to compare Naive Bayes (NB) and Support Vector Machine (SVM) classification algorithms. The output of this research is to find out the best algorithm based on the highest accuracy value for sentiment analysis of the use of marketplace platforms [8]. However, Naive Bayes generally excels in speed and accuracy for text data [9]. In the context of the Indonesian language, the main challenge lies in the complexity of morphology and slang variation (*slang*), which often affects the results of the classification. Juwiantho proposes the use of word-based representations *Word2Vec* with model *Convolutional Neural Network (CNN)* to improve the results of the classification of Indonesian texts, but the method requires higher computation and large datasets [10].

This research will focus on the analysis of public sentiment on the issue of "Dissolve the House of Representatives" on Twitter (X) using the Naive Bayes and SVM methods. The goal is to find out how the public's perception of the DPR institution, whether dominated by positive, negative, or neutral sentiments, as well as compare the performance of the two methods in classifying tweets. The analysis is carried out with *stages of data crawling, preprocessing*, data labeling, classification, and model evaluation using metrics such as accuracy, precision, and recall.

This research is expected to contribute both from a theoretical and practical perspective. From the theoretical side, this study expands the application of *machine learning algorithms* in the socio-political context of Indonesian. From a practical perspective, the results of the research can be used by academics, governments, and research institutions to understand public opinion of state institutions in a more objective and data-based manner. Thus, this study not only provides an overview of public perception of the DPR, but also becomes an example of the application of data analysis technology in supporting transparency and evaluation of public policies in the digital era.

2. Research Methods

2.1. Sentimen Analysis

Sentimen analysis (*Sentiment Analysis*) is a way of gathering the opinion of the general public using social networks in which there are public services and current issues [11]. Sentiment analysis is included in one of the fields of Natural Language Processing (NLP) and is a process used to help identify the content of a dataset in the form of opinions or views (sentiments) in the form of text on an issue or event that is positive, negative or neutral [12].

2.2. Naïve Bayes

Naïve Bayes (NB) is one of the machine learning methods that uses probability calculations. The basic concept used by Naïve Bayes is Bayes' Theorem, which is classification by calculating probability values [13].

This method uses the concept of probability to classify data on a specific class. In general, each attribute is almost certain to be dependent on the others, but assuming this makes the Naive Bayes method easier in terms of calculations. Naive Bayes is a method commonly used to find the highest probability value in the process of classifying test data in each category of the most appropriate class, which can be expressed in the following formula equation [14].

$$P(X|Y) = \frac{P(Y|X) \cdot P(X)}{P(Y)}$$

Where:

- y = unknown class data
- x = hypothesis data y
- $P(x|y)$ = the probability of hypothesis x based on condition y
- $P(x)$ = probabilities hypothesis x
- $P(y|x)$ = probability y based on conditions on hypothesis x
- $P(y)$ = probability of y

Advantages of Naïve Bayes:

1. It can be used for quantitative and qualitative data.
2. It doesn't require a large amount of data.
3. There is no need to do *a lot* of training data.
4. If there is a missing value, it can be ignored in the calculation.
5. The calculations are quick and efficient to understand.

2.3. Support Vector Machine (SVM)

Support Vector Machine (SVM) is a model derived from statistical learning theory that will provide better results than other methods [15]. SVM works with a learning system that uses hypothetical space in the form of linear functions in a high-dimensional feature space, only in the classification method SVM can only classify data into two classes [16]. SVM seeks the best equivalent hyperplane by maximizing the margin or distance of two sets of different classes [17]. This can be formulated in the SVM optimization problem for linear classification, as below:

$$f(x_d) = \sum_{i=1}^{ns} a_i y_i \vec{x}_i \vec{x}_d + b$$

Where:

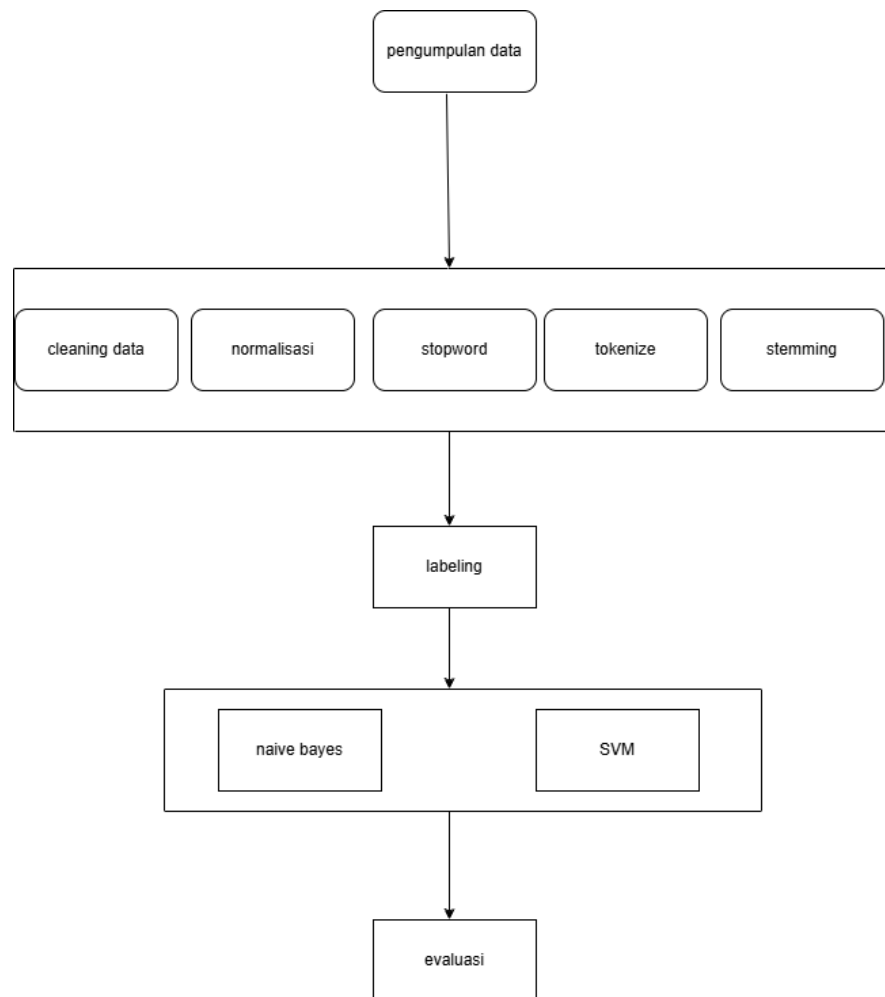
- ns = Jumlah *support vector*
- a_i = Value of the weight of each data point
- y_i = Data class
- $x_{>i}$ = Variabel *support vector*
- $x_{>d}$ = Data to be classified
- b = Error or *bias value*

SVMs have advantages in high data generalization and can build good classification models even when trained with relatively little data. However, this method is difficult to apply to data sets with very large samples and dimensions. This shows that SVM is able to produce good performance even with a relatively small amount of data.

2.4. Research Stage

This study uses a systematic *Text Mining* framework to analyze public sentiment on Twitter (X). This process is divided into four main phases: Data Acquisition, Pre-processing, Classification Modeling, and Evaluation

Figure 1 Research Flow



1. Data Acquisition and Collection

The first stage is *Data Acquisition*. The data used is primary data (*tweets*) obtained from the Twitter platform (X) with a focus on the issue of "Dissolve the DPR". Data is collected using *the Twitter (X) API* over a specific period. *Tweets* that are successfully collected include text, dates, and user information, and are then further processed.

2. Pre-processing of Data

Once the data is collected, a series of *preprocessing stages* are performed to clean and standardize the text, ensuring that the data is ready to be processed by machine learning algorithms.

a. Cleaning Data

The Data Cleaning stage aims to remove *noise* or irrelevant characters from *tweets*. This includes the *removal of* URLs, user *mentions* (@user), punctuation, symbols, numbers, and other special characters that do not contribute to the meaning of the sentiment.

b. Normalisasi (Normalization)

Normalization aims to change informal words or alay into standard forms. Although not explicitly mentioned in your step, *case folding* is an important part of normalization, called *Case Transformation* in this context. Normalization can also include *general typo correction*.

c. Tokenize

Tokenizing is the process of breaking down each sentence in *a tweet* into units of words (called *tokens*). This is the initial stage where the text is transformed into a structure that can be further processed by the computer.

d. Stopword Removal

The Stopword Removal stage (*Filter Stopwords*) aims to remove common words that have a high frequency but low informative value to sentiment polarity (e.g.: 'which', 'and', 'is', 'in'). Stopword removal reduces feature dimensions and speeds up the model training process.

e. *Vote*

Stemming is the process of converting suffix words (including prefixes, inserts, and suffixes) into root *words*. For example, the word 'dissolve' is changed to 'dissolve'. In Indonesian, *voting* is often done using *the Literary* library.

3. Pelabelan (*Labeling*)

The Labeling stage is done after the data has been cleaned (or may be done before *preprocessing*). *Cleaned tweet data* is manually labeled as sentiment polarity by researchers or annotators. The labels used are Positive, Negative, or Neutral. This labeling is crucial because it creates training data that will be learned by the *supervised learning model*.

4. Classification Modeling

The pre-processed *labeled data* is then converted into numerical vectors (e.g., using TF-IDF) before being entered into the classification algorithm.

a. Naive Bayes (NB)

The Naive Bayes algorithm is applied to classify sentiment. This algorithm works on the basis of Bayes' Theorem and assumes the independence of features (words) from each other. The main advantage of Naive Bayes is its speed of calculation.

b. Support Vector Machine (SVM)

Support Vector Machine (SVM) is also applied as a comparative classification algorithm. SVM works by looking for the optimal *hyperplane* that separates the sentiment class by the largest *margin*. SVMs are known to perform well in text classification, especially data with high dimensions and relatively small sample counts.

5. Evaluation

The last stage is Evaluation which aims to measure and compare the performance of the Naive Bayes and SVM models. The metrics used to measure the effectiveness of the model are Accuracy, Precision, and Recall. To ensure the validity of the results, the Cross Validation method is used (e.g., with k-folds 10). The results of this evaluation determine which model is the most optimal in analyzing the sentiment of the "Dissolve the DPR" issue.

3. Results and Discussion

This section presents the results of the sentiment classification process using two *machine learning* models, namely Naive Bayes and Support Vector Machine (SVM). The data used were 1989 tweets in Indonesian related to the keyword "Dissolve the DPR" which had gone through *the stages of automatic preprocessing, voting, and sentiment labeling*.

3.1 Distribution of Sentiment Data

Before going into modeling, the data is divided into 80% training data and 20% test data. It is important to note that the distribution of sentiment in the dataset is highly imbalanced. As seen in the data visualization of Figure 2, the dataset is dominated by negative and neutral sentiment. Positive sentiment has a very small sample count, which significantly affects the training and evaluation of the model.

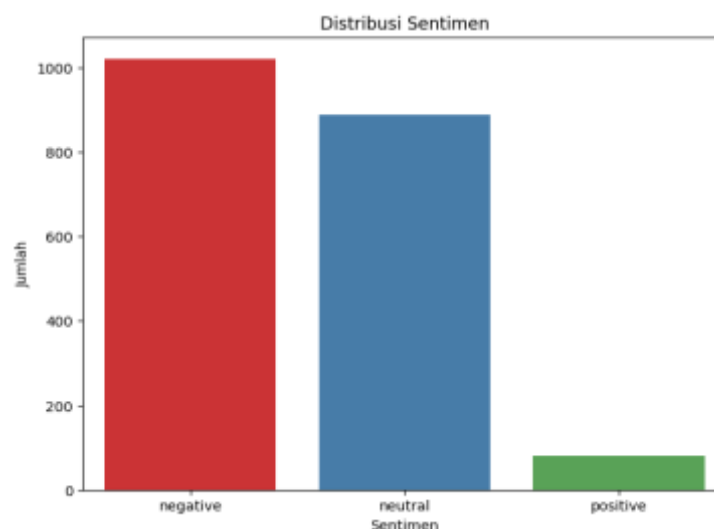


Figure 2. Visualization of Test Data Sentiment Distribution

3.2 Labeling

The labeling process is done manually to determine that the tweet contains positive sentiments that contain compliments such as happy, happy, or satisfied, while negative sentiments contain criticism, complaints, insults, upset and disappointed. The labeling can be seen in Table 1.

Table 1. Pelabelan

Label Sentimen	Definition/Criteria of Labeling (Tweet Content)
Positive	Tweets that contain expressions of support, praise, or feelings of pleasure, happiness, or satisfaction with the House of Representatives or related issues.
Negatives	Tweets that contain expressions of dissatisfaction, criticism, complaints, insults, annoyance, or disappointment with the House of Representatives or related issues.
Neutral	Tweets that don't explicitly show positive or negative sentiment are often factual information, non-partisan questions, or non-partisan attitudes.

3.3 Naive Bayes Modeling Results

The first model tested was the Multinomial Naive Bayes (MNB), which is often used as a robust *baseline* for text classification. The model was trained on trained data that had been digitized using TF-IDF. The results of the evaluation on the test data showed that the MNB model achieved an accuracy of 74%. Detailed performance for each class is presented in *the Classification Report Table 2* and *the Confusion Matrix Figure 3*.

Table 2. Classification Report Model Naive Baye

Label	Precision	Recall	F1-Score	Support
Negative	0.68	0.95	0.79	203
Neutral	0.87	0.57	0.69	179
Positive	0.00	0.00	0.00	16
Accuracy			0.74	398
Macro Avg	0.52	0.51	0.49	398
Weighted Avg	0.74	0.74	0.71	398

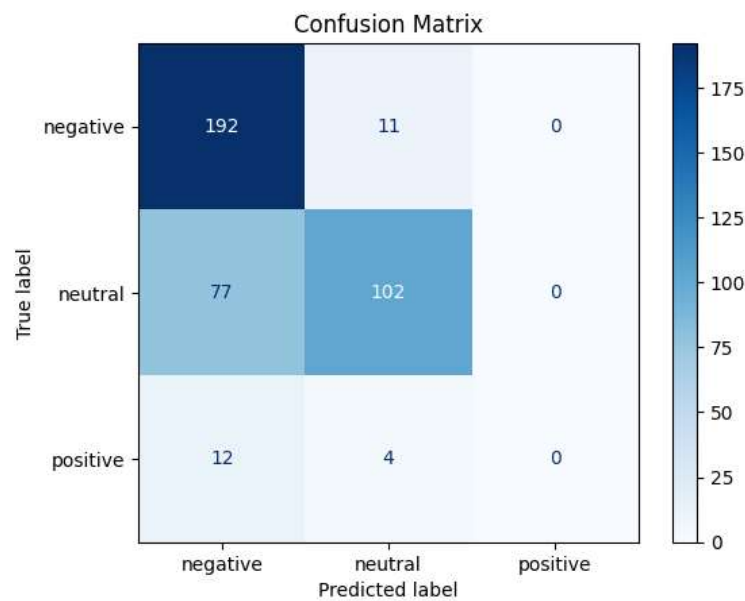


Figure 3. Confusion Matrix Model Naive Bayes

Naive Bayes analysis: Based on Table 2, the performance of MNB is highly unbalanced. This model is excellent at identifying negative sentiment (Recall 0.95), but at the expense of other classes. The precision for neutral sentiment is quite high (0.87), but *the recall* is low (0.57), which indicates that many neutral sentiments are misclassified.

The most significant weakness is the total failure of the model to identify positive sentiment (F1-Score 0.00). This is most likely due to the very small amount of positive data (only 16 samples) in the test data, so the model does not have enough data to study the pattern.

3.4 Support Vector Machine (SVM) Modeling Results

The second model tested was the Support Vector Machine with a linear kernel (LinearSVC), which is known to be reliable for high-dimensional text data. The SVM model showed superior performance by achieving an accuracy of 79%, or 5% higher than MNB. Detailed performance results are presented in Table 3 and Figure 4.

Table 3. *Classification Report Model Support Vector Machine (SVM)*

Label	Precision	Recall	F1-Score	Support
Negative	0.77	0.87	0.82	203
Neutral	0.83	0.77	0.80	179
Positive	1.00	0.06	0.12	16
Accuracy			0.79	398
Macro Avg	0.87	0.57	0.58	398
Weighted Avg	0.80	0.79	0.78	398

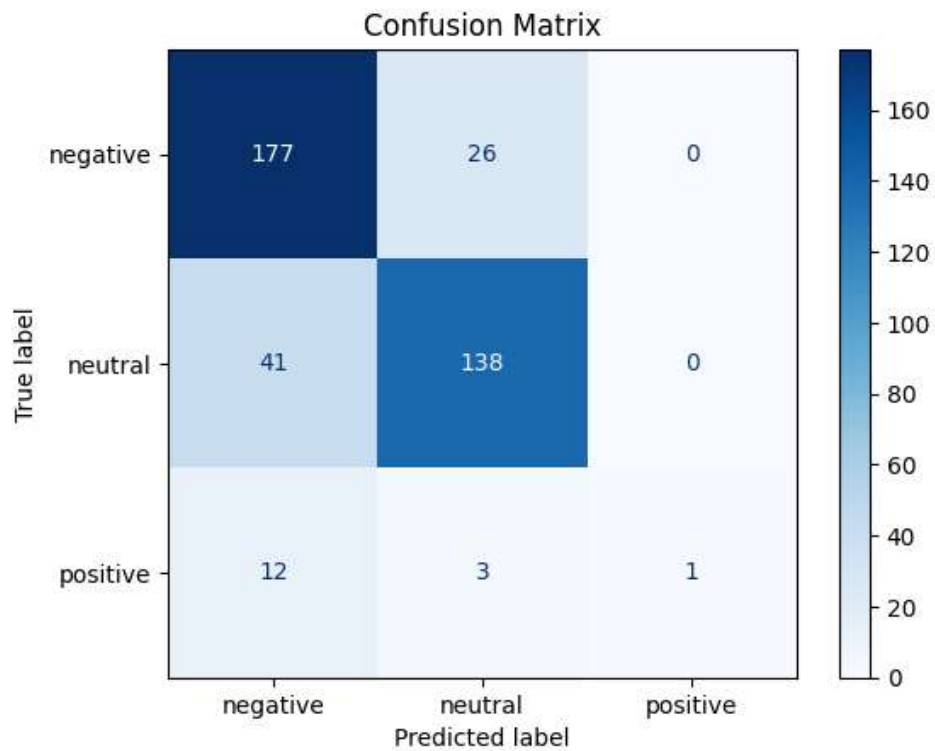


Figure 4. *Confusion Matrix Model Support Vector Machine (SVM)*

SVM Analysis: The SVM model shows a much more balanced and better performance on the two dominant classes. The F1-Score for negative (0.82) and neutral (0.80) sentiment is high and close to each other. This shows that SVM is more effective in distinguishing these two classes than MNB.

Although it was still very difficult with the positive class (F1-Score 0.12), the SVM managed to correctly identify one positive sample (Precision 1.00, Recall 0.06). Like MNB, the poor performance in this positive class is a direct reflection of the problem of extreme data imbalance.

3.5 Discussion

Methodologically, the Support Vector Machine (SVM) proved to be a classification model that was definitively superior to the Naive Bayes (NB) for this dataset. SVM's excellence is quantitatively measured through the achievement of 79% accuracy, which significantly surpasses Naive Bayes (74%).

This advantage doesn't just stop at accuracy. SVM shows a much more balanced and reliable performance in classifying the two dominant sentiment classes. This is evidenced by a high and consistent F1-Score for negative (82%) and neutral (80%) sentiments. SVM's ability to find the optimal *hyperplane* has proven to be effective in handling high-dimensional text data (TF-IDF). In contrast, Naive Bayes, although fast, exhibits an unbalanced

performance; the model is very sensitive in detecting negative sentiment (Recall 95%) but has difficulty recognizing neutral sentiment (Recall 57%), so its F1-Score for the neutral class is much lower (69%).

The most critical finding of this study is the impact of extreme *data imbalance* on positive classes. With only 16 positive samples (about 4%) in the test data (398 samples), these two *supervised learning models* were essentially "starved" of data and failed miserably in studying patterns for positive sentiment. This failure is evident in the F1-Score metric, where Naive Bayes recorded 0.00% and SVM was only able to record 0.12%.

Themathematically, the failure of this model is not a computational weakness, but a strong statistical confirmation. This shows that in the public discourse on Twitter (X) regarding "#BubarkanDPR", the sentiments expressed are almost exclusively negative and neutral. This sentiment analysis objectively maps that public opinion on the issue on the Twitter platform (X) is massively dominated by dissatisfaction (negative) and non-partisan (neutral) attitudes, with almost no expression of support (positive).

4. Conclusion

This study concluded that the *Support Vector Machine* (SVM) method showed superior performance to *Naive Bayes* in classifying public sentiment regarding the issue of "Dissolve the DPR", with an accuracy of 79% compared to 74%. SVM has been shown to be more stable and balanced in distinguishing the dominant sentiment classes (negative and neutral), while *Naive Bayes* tends to have a high bias towards negative classes but is weak in recognizing neutral classes. In addition to the technical aspects, this analysis reveals the empirical fact that public opinion on the Twitter platform (X) towards the legislature is strongly dominated by negative sentiment (dissatisfaction) and neutral attitudes, where the failure of the two algorithms in detecting positive sentiment is not due to a computational error, but real statistical evidence of a lack of public support or a *severe data imbalance* in the discourse. Based on the limitations found, further research is recommended to apply *imbalanced data handling* techniques such as *the oversampling* method (SMOTE) or data augmentation so that the classification in the minority class (positive) can be properly identified. In addition, future model development can be expanded by comparing *Deep Learning approaches* such as *Long Short-Term Memory* (LSTM) or BERT which are able to capture semantic context more deeply, as well as refining the pre-processing stage with a more comprehensive slang normalization dictionary to address the high variation of informal language on social media.

Acknowledgment

Praise be to Allah SWT for all His blessings so that this research can be completed. The author would like to thank **Indragiri Islamic University** for the support of academic facilities provided during the preparation of this study.

References

- [1] P. R. Sari, D. R. Indah, E. Rasywir, M. A. Firdaus, and G. Athalina, "Comparison of Naive Bayes and SVM Algorithms for Sentiment Analysis of PUBG Mobile on Google Play Store," *System*, vol. 13, no. 6, p. 2767, 2024, doi: 10.32520/stmsi.v13i6.4814.
- [2] A. Kumar, "Systematic literature review of sentiment analysis on Twitter using soft computing techniques," no. November 2018, 2019, doi: 10.1002/cpe.5107.
- [3] I. Syapriani, D. Putriana, R. P. Aula, G. S. Sembiring, V. Aulia, and S. Khoirira, "Journal of Nusantara Research on Social Media User Communication Strategy in Managing Language, Ethics, and Social Resistance in the Comment Room of Digital Platform X (Twitter)," vol. 1, pp. 360–368, 2025.
- [4] Z. Drus, H. Khalid, Z. Drus, and H. Khalid, "ScienceDirect ScienceDirect Sentiment Analysis in Social Media and Its Application: Systematic Sentiment Analysis in Social Media and Its Application: Systematic Literature Review Literature Review," *Computation. Sci.*, vol. 161, pp. 707–714, 2019, doi: 10.1016/j.procs.2019.11.174.
- [5] A. P. Widyassari, D. Salsabilla, and M. A. Amrozi, "Analysis of public sentiment on twitter towards president Prabowo's inauguration using the Naïve Bayes algorithm," vol. 10, no. 1, pp. 13–24, 2025.
- [6] C. Villavicencio, J. J. Macrohon, X. A. Inbaraj, J. H. Jeng, and J. G. Hsieh, "Twitter sentiment analysis towards covid-19 vaccines in the Philippines using naïve bayes," *Inf.*, vol. 12, no. 5, 2021, doi: 10.3390/info12050204.
- [7] P. R. Sari *et al.*, "Comparison of Naive Bayes and SVM Algorithms for Sentiment Analysis of PUBG Mobile on Google Play Store," vol. 13, pp. 2767–2779, 2024.
- [8] W. Ningsih, B. Alfiana, and D. Wulandari, "Comparison of Naive Bayes and SVM Algorithms in Twitter Sentiment Analysis on Electric Car Use in Indonesia," vol. 4, no. April, pp. 556–562, 2024.
- [9] M. Wahyudi, M. S. Novelan, U. Development, and P. Budi, "Comparison Of K-Nn And Naive Bayes Performance In Information Technology Governance To Analyze The Level Of Satisfaction Of Shopee Application Users," vol. 4307, no. August, pp. 3424–3431, 2025.
- [10] H. Juwiantho *et al.*, "Sentiment Analysis Twitter Bahasa Indonesia Berbasis Word2vec Menggunakan Deep Convolutional Neural Network Sentiment Analysis Twitter Indonesian Language Word2vec-Based Using Deep Convolutional Neural Network," vol. 7, no. 1, pp. 181–188, 2020, doi: 10.25126/jtiik.202071758.
- [11] A. D. Adhi Putra, "Sentiment Analysis on User Reviews of Bibit and Bareksa Applications with KNN

- Algorithm," *JATISI (Tek Journal. Inform. and Sist. Information)*, vol. 8, no. 2, pp. 636–646, 2021, doi: 10.35957/jatisi.v8i2.962.
- [12] M. I. Fikri, T. S. Sabrila, and Y. Azhar, "Comparison of Naïve Bayes and Support Vector Machine Methods in Twitter Sentiment Analysis," *Smatika J.*, vol. 10, no. 02, pp. 71–76, 2020.
- [13] H. Susana and N. Suarna, "application of the naive bayes method classification model to the use of internet access stmik ikmi cirebon jl perjuangan no 10b kesambi cirebon city 3) software engineering study program stmik ikmi Cirebon Jl Perjuangan No 10B Kesambi Cirebon City 4) Computerized Accounting Study Program STMIK IKMI Cirebon Jl Perjuangan No 10B Kesambi Cirebon City," *J. Sist. Inf. and Technology. Information*, vol. 4, no. 1, pp. 1–8, 2022.
- [14] gigih ,putra Kawani, "Journal of Informatics, Information System, Software Engineering and Applications," *J. Informatics, Inf. Syst. Softw. Eng. Appl.*, vol. 1, no. 2, pp. 73–081, 2019.
- [15] Wikipedia, "Corruption Eradication Commission of the Republic of Indonesia," *Wikipedia*, vol. 7, no. 1, p. 1, 2024.
- [16] D. Alita, Y. Fernando, and H. Sulistiani, "Implementation of the Svm Multiclass Algorithm on Indonesian Public Opinion on Twitter," *J. Compact Techno*, vol. 14, no. 2, p. 86, 2020, doi: 10.33365/jtk.v14i2.792.
- [17] A. S. Ritonga and E. S. Purwaningsih, "The Application of the Support Vector Machine (SVM) Method in the Quality Classification of Smaw (Shield Metal Arc Welding)," *Ilm. Edutic*, vol. 5, no. 1, pp. 17–25, 2018.