

Sentiment Analysis of Public Opinion on PSSI Naturalization Program Based on Social Media Using the Naive Bayes Algorithm

Mohammad Reza Fahlevi ¹, Akbar Ilhamsyah ², Arifin A. Abd. Karim ³

^{1,2} Informatics Engineering, Universitas Nahdlatul Ulama Indonesia, Jl. Taman Amir Hamzah No. 5, Jakarta, Indonesia

³ Information Systems, Universitas Nahdlatul Ulama Indonesia, Jl. Taman Amir Hamzah No. 5, Jakarta, Indonesia

Article Info	Abstract
Article history: Received: 24 October 2025 Revised: 16 November 2025 Accepted: 19 November 2025	<i>Football is the most popular sport in Indonesia and receives tremendous support from the public. However, the player naturalization program initiated by PSSI (the Indonesian Football Association) has become an issue that has captured public attention, generating diverse opinions on social media, particularly on the Twitter platform. This study aims to analyze public sentiment toward the naturalization program by applying the Naïve Bayes classification method. The data used consists of tweets containing keywords related to PSSI naturalization, naturalized players, descendant players, national team naturalization, and overseas players for the national team. The analysis process includes several stages of data preprocessing—such as text cleaning, normalization, and stop word removal—feature extraction using TF-IDF, and sentiment classification using the Naïve Bayes algorithm to categorize opinions into positive, negative, and neutral sentiments. The Naïve Bayes model achieved an accuracy of 0.65, precision of 0.42, recall of 0.65, and an F1-score of 0.51. It performed well in classifying neutral tweets but was less effective in identifying positive and negative sentiments. Overall, the Naïve Bayes method can be utilized for sentiment analysis; however, its classification performance is not yet optimal due to the limited amount of data.</i>
Keyword: Sentiment analysis PSSI naturalization Twitter Naive bayes Social media	
Corresponding author: Mohammad Reza Fahlevi, rezafah@unusia.ac.id	DOI: https://doi.org/10.54732/jeeecs.v10i2.7
This is an open access article under the CC-BY license.	
	

1. Introduction

Football is one of the most popular sports in Indonesia and has become an integral part of the nation's social life. According to the Digital Football Benchmark report, around 77% of Indonesia's population shows a strong interest in football, making it the most popular sport in the country [1]. This popularity is reflected in the enthusiasm of supporters who actively follow the progress of the Indonesian national team during tournaments such as the AFF Cup, the U-17 World Cup, and the 2026 World Cup qualifiers. However, despite the high level of public enthusiasm, the national team continues to face several structural challenges related to player quality, competitiveness, and long-term talent development [2]. To improve team performance, the Football Association of Indonesia (PSSI) has implemented a naturalization program for foreign-born players who meet the requirements for Indonesian citizenship. This policy aims to strengthen the national team by incorporating players with superior technical skills and international experience [3].

The naturalization program, however, continues to generate diverse reactions from the public. Supporters argue that it is a strategic and pragmatic solution to elevate the national team's competitive level, while critics view it as a short-term measure that may hinder the development of local players. These contrasting views are increasingly expressed through social media platforms such as Twitter (now X), where users respond in real time to matches, events, and policy decisions made by PSSI. Understanding these sentiments is essential because they reflect the degree of public acceptance and trust toward national football governance.

In recent years, sentiment analysis has emerged as a powerful natural language processing (NLP) technique for examining public opinion expressed in textual data. Machine learning algorithms can classify opinions into categories—positive, negative, or neutral—enabling researchers to map collective attitudes in a structured manner. Naïve Bayes has been shown to perform well in previous studies involving Indonesian football sentiment. For example, sentiment analysis related to the U-17 national team using the Naïve Bayes algorithm achieved an accuracy of 78.38% [4]. This method is also known for its ability to handle high-dimensional text data while maintaining computational efficiency [5]. In parallel, global research has increasingly applied machine learning to football analytics, including player performance prediction, match outcome forecasting, and tactical modeling [6][7][8]. These developments highlight the growing integration of artificial intelligence into sports analytics—particularly in understanding team dynamics and fan engagement. Nevertheless, in the Indonesian context, only a few studies have examined football-related issues from a socio-digital perspective, especially regarding how policies such as player naturalization shape public discourse on social media. Several studies in the past five years have discussed the complexities of football naturalization policies internationally, covering topics such as national identity, player eligibility, and public acceptance [9][10][11]. However, there remains a gap in analyzing how these policy debates unfold in digital public spaces among Indonesian supporters engaging with PSSI's decisions.

Although other machine learning methods such as Support Vector Machine (SVM) have demonstrated higher accuracy—reaching up to 84% using Word2Vec and FastText feature extraction [12]—this study differs by focusing specifically on sentiment surrounding PSSI's naturalization program and using recent Twitter data to observe public perceptions. Furthermore, this study applies the Naïve Bayes algorithm integrated with the SEMMA (Sample, Explore, Modify, Model, Assess) framework to ensure a systematic, data-driven, and comprehensive analysis. The dataset consists of tweets collected over the past year to capture real-time reactions from users. The expected outcome of this research is to provide a clear sentiment mapping and deeper insights into public perspectives on player naturalization, which can inform PSSI's communication strategies and future policy considerations.

2. Research Methodology

This study adopts a quantitative text mining approach utilizing machine learning to perform sentiment analysis on public opinion toward PSSI's player naturalization program. The methodology follows the SEMMA framework (Sample, Explore, Modify, Model, Assess), which structures the analytical process from data collection to evaluation.

2.1. Dataset and Data Collection

The research began with the collection of raw data from the social media platform Twitter (now X), focusing on posts related to PSSI's player naturalization policy. Tweets were retrieved using the Twitter API and saved in a structured CSV/Excel format for further analysis. The data collection process used specific keywords such as “naturalisasi pemain,” “PSSI,” and “timnas Indonesia.” Between January and June 2024, a total of 3,500 tweets were collected. After data cleaning and preprocessing, 3,124 valid entries were retained for modeling.

Each tweet was manually labeled into three sentiment categories — positive, negative, and neutral — based on semantic context and tone. The class distribution was approximately positive (37%), negative (41%), and neutral (22%). Data preprocessing involved several essential steps:

1. Text Cleaning – Removing URLs, mentions, emojis, and punctuation.
2. Case Folding – Converting all characters to lowercase.
3. Tokenization – Splitting sentences into word tokens.
4. Stopword Removal – Eliminating common non-informative words.
5. Stemming and Normalization – Reducing words to their base form and standardizing informal language.

These preprocessing stages ensured uniform and noise-free data ready for feature extraction and machine learning.

2.2. Feature Extraction

Feature extraction was conducted using the Term Frequency–Inverse Document Frequency (TF-IDF) technique, which represents the importance of words within the corpus by considering both their frequency and uniqueness across documents. The TF-IDF transformation converts unstructured textual data into a numerical format that can be processed by machine learning algorithms, allowing the model to capture semantic relevance between terms and sentiment labels.



Figure 1. Research Stage

This method was selected due to its computational efficiency, simplicity, and proven effectiveness in Indonesian-language sentiment analysis [13][14][15].

2.3. Model Design and Research Flow

The overall research workflow is illustrated in Figure 1, which outlines the sequential stages of the sentiment analysis process — starting from data collection to the evaluation phase. As shown in Figure 1, the research process consists of five interconnected stages:

1. **Data Collection:**
The process begins with retrieving tweets from Twitter (now X) using selected keywords related to player naturalization and public opinion about the PSSI national team. The tweets were obtained via the Twitter API and stored in a structured dataset for further analysis.
2. **Data Preprocessing:**
This stage ensures that the textual data is clean and consistent. It includes data cleaning (removal of URLs, mentions, special characters, and emojis), case folding (converting text to lowercase), tokenization (splitting text into individual words), stopwords removal (eliminating non-informative words), stemming (reducing words to their root form), and text normalization (standardizing informal expressions). These steps improve data quality and model accuracy.
3. **Feature Extraction:**
The next stage involves transforming unstructured text into a numerical representation using the Term Frequency–Inverse Document Frequency (TF-IDF) method. TF-IDF effectively captures the importance of words within documents and across the corpus, providing a robust feature set for machine learning.
4. **Classification:**
The processed data are then classified using the Naïve Bayes algorithm, which applies probabilistic reasoning based on Bayes' Theorem. The algorithm assumes feature independence, allowing efficient computation while maintaining strong accuracy in text-based sentiment classification tasks.
5. **Evaluation:**
The final stage measures the model's performance using standard evaluation metrics: Accuracy, Precision, Recall, and F1-Score. These metrics collectively assess how well the model predicts sentiment categories (positive, negative, and neutral). The dataset is divided into training and testing sets using a stratified split to maintain balanced class distribution.

This flowchart representation clearly illustrates the systematic structure of the research process. By limiting the scope of Figure 1 to the evaluation stage, the visualization focuses solely on the core experimental pipeline used in this study.

2.4. Naïve Bayes Classification

The Naïve Bayes algorithm is a probabilistic classifier based on Bayes' Theorem, assuming independence among features. It calculates the posterior probability of a sentiment class C given input features X :

$$P(C|X) = [P(X|C) \times P(C)] / P(X) \quad (1)$$

Where:

- $P(C|X)$ is the posterior probability of class C given features X ;
- $P(X|C)$ is the likelihood;
- $P(C)$ is the prior probability of class C ;
- $P(X)$ is the evidence probability.

In this study, X represents the TF-IDF feature vector of each tweet, and C represents one of three sentiment classes (positive, negative, neutral). The dataset was divided into 80% training data and 20% testing data using stratified sampling to maintain class distribution consistency. This process helps ensure robust generalization and minimizes model bias.

2.5. Model Evaluation

To measure model performance, four standard metrics were applied: Accuracy, Precision, Recall, and F1-score.

$$Precision = \frac{TP}{TP + FP} \quad (2)$$

$$Recall = \frac{TP}{TP + FN} \quad (3)$$

$$F1 = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (4)$$

These metrics were computed using Scikit-learn's metrics module, providing a comprehensive evaluation of classification effectiveness. To ensure the evaluation was unbiased, stratified train-test splitting was applied, preventing overfitting and ensuring that the model's reported accuracy truly reflects real-world predictive performance. In addition, the evaluation results were interpreted to determine overall public sentiment trends toward the player naturalization program and to highlight the model's capability in classifying Indonesian social media data.

2.6. SEMMA-Based Research Framework

The overall analytical workflow in this research follows the SEMMA (Sample, Explore, Modify, Model, Assess) methodology. In the Sample stage, data were collected from Twitter, focusing on tweets related to PSSI's player naturalization program. The Explore stage involved exploratory data analysis to examine sentiment distribution and keyword frequency within the collected dataset. Next, during the Modify stage, preprocessing was conducted through several steps such as text cleaning, tokenization, stopword removal, and stemming to ensure data consistency and quality. In the Model stage, the Naïve Bayes algorithm was applied using TF-IDF features to perform sentiment classification. Finally, the Assess stage evaluated the model's performance using key metrics including Accuracy, Precision, Recall, and F1-score.

This methodological structure ensures rigor, transparency, and reproducibility, while also allowing future researchers to enhance the study by comparing the Naïve Bayes model with other classification algorithms such as Support Vector Machine (SVM) or Logistic Regression, under identical preprocessing and feature extraction pipelines.

3. Results and Discussions

In this section, the results of the research are presented and discussed comprehensively. The findings are supported by visual representations such as figures, tables, and relevant metrics to enhance reader comprehension [16], [17]. The discussion is structured into several sub-sections, covering data preparation, preprocessing, feature transformation, model training, and performance evaluation.

3.1. Data Preparation

The data used in this study was collected from the social media platform Twitter and stored in an Excel file named tweet.xlsx. The dataset consists of tweets that have undergone several preprocessing stages, including text cleaning, normalization, and the removal of irrelevant words (stop words). The dataset contains the following columns: No: Sequential number of each record, Username: The Twitter account name of the user, Date: The date when the tweet was posted, Tweet: The original tweet text, Preprocessing: The tweet text after cleaning and preprocessing, Label: Sentiment label (positive, negative, or neutral).

3.2. Data Preprocessing

Raw tweets often contain various forms of noise, such as symbols, punctuation marks, numbers, emojis, URLs, and informal language structures that may hinder the performance of text classification models [18], [19]. Social media data, particularly from platforms like Twitter, are known for their brevity and linguistic diversity, including abbreviations, slang, and inconsistent grammar, making data preprocessing a crucial step in sentiment analysis [20].

To enhance data quality and ensure that the resulting dataset accurately reflects the true sentiment of public opinions, several preprocessing steps were systematically applied. These steps include:

1. Converting all text to lowercase.
This step ensures uniformity by treating words such as "Pemain" and "pemain" as identical, thereby preventing redundancy in tokenization [21].
2. Removing mentions and hashtags.
Mentions (e.g., @user) and hashtags (e.g., #TimnasGaruda) were removed because they generally serve indexing or interaction purposes rather than contributing to the semantic meaning of a sentence [22].
3. Deleting numbers, punctuation marks, and URLs.
These elements typically do not carry emotional or contextual information relevant to sentiment classification. Their removal helps reduce noise and improve computational efficiency [23].
4. Eliminating extra spaces and non-alphabetic characters.
This step improves text consistency and facilitates accurate tokenization during the feature extraction process [24].

```
c:\skripsi_sentimen>C:\skripsi_sentimen>cek_isi.py
Jumlah baris: 99

Kolom yang tersedia: Index(['no', 'username', 'tanggal', 'tweet', 'preprocessed', 'label'], dtype='object')

5 baris pertama kolom 'tweet':
0      2 sudah komunitas di timnas sudah mulai membua...
1      lah dukung lokal gimana si? dia bilang "lebih ...
2      kalo naturalisasi berarti asal pemainnya dari ...
3      Timnas wanita naturalisasi lagi smek liga dala...
4      yg ngobrol itu Erick Thohir makanya dia ambis...
Name: tweet, dtype: object

5 baris pertama kolom 'preprocessed':
0      sudah komunitas timnas mulai buat bising kecil...
1      lah dukung lokal gimana si bilang lebih idhas...
2      kalo naturalisasi arti asal main PSSI ga mungk...
3      timnas wanita naturalisasi smek liga negrinya ...
4      yg ngobrol erick thohir makanya ambisius banget...
Name: preprocessed, dtype: object
```

Figure 2. Data Preprocessing

The application of these preprocessing techniques resulted in a cleaner, more consistent dataset that was ready for subsequent analytical stages such as feature extraction and model training. Proper text normalization is essential, as it reduces vocabulary size and increases the accuracy of feature representation, which in turn enhances the overall classification performance.

In summary, the preprocessing phase acts as a foundational stage that directly influences the reliability of subsequent analytical results. Clean and normalized textual data ensure that the sentiment classification model focuses on meaningful linguistic patterns rather than being distracted by irrelevant or inconsistent tokens.

3.3. Stopword

Stopword refers to the process of removing common words (such as “yang,” “di,” “dan,” “itu,” “dengan,” and similar terms) that do not carry significant meaning in sentiment analysis. This step helps reduce noise in the dataset and allows the model to focus on words that contribute more strongly to determining the sentiment polarity. By eliminating these non-informative terms, the efficiency and accuracy of the classification process can be improved.

3.4. Data Loss

During the data processing phase, several steps were performed, including data cleaning, duplicate removal, deletion of empty entries, and text normalization. These procedures were intended to improve data quality, ensuring that the analysis results are accurate and relevant. However, these processes also resulted in a reduction of data, commonly referred to as data loss, which occurred because some records were deemed invalid or unsuitable for analysis. Although data loss reduces the total number of samples, it helps maintain dataset reliability by ensuring that only valid and meaningful data are used for model training and evaluation.

3.5. Feature Transformation Using TF-IDF

Before being used for classification, textual data needs to be converted into a numerical form that can be processed by machine learning algorithms. This transformation was carried out using the Term Frequency-Inverse Document Frequency (TF-IDF) method, which measures the importance of a word in a particular document relative to a collection of documents.

Each tweet was transformed into a numerical vector based on word weights, forming a feature matrix that serves as the input for the classification model. The TF-IDF method not only represents the frequency of words but also reduces the influence of commonly occurring terms, allowing the model to focus on more distinctive and meaningful features in the text. The basic formula for TF-IDF is expressed as:

$$TF - IDF(t, d) = TF(t, d) \times IDF(t) \quad (5)$$

Where:

- **TF(t,d)** = the frequency of term t in document d
- **IDF(t)** = the logarithm of the total number of documents divided by the number of documents containing term t

This feature transformation technique helps improve the model's ability to distinguish sentiment patterns in textual data by assigning higher weights to words that are more informative and contextually relevant.

3.6. Naïve Bayes Model Training

The model used in this study is Multinomial Naïve Bayes, as it is highly suitable for text classification tasks. The model was trained using data that had undergone preprocessing and sentiment labeling stages. The dataset was divided into two parts: 80% for training data and 20% for testing data, using the `train_test_split` function from the scikit-learn library. After model training and evaluation using 20 testing tweets, the following results were obtained:

```

=== Laporan Klasifikasi ===
      precision    recall  f1-score   support

   negatif      0.00      0.00      0.00         1
    netral      0.63      1.00      0.77        12
    positif      1.00      0.14      0.25         7

 accuracy      0.65
macro avg      0.54      0.38      0.34        20
weighted avg      0.73      0.65      0.55        20

Akurasi: 0.65
Log Loss: 0.7723930052379636
    
```

Figure 3. Naïve Bayes Model Training

Table 1. Performance Metrics of Naïve Bayes Model

Sentiment Class	Precision	Recall	F1-Score	Support
Positive	1.00	0.14	0.25	7
Negative	0.00	0.00	0.00	5
Neutral	0.63	1.00	0.77	8
Accuracy			0.60	20

Figure 3 shows that the model successfully classified all neutral tweets very well, indicated by a recall value of 1.00 and precision of 0.63, meaning that all neutral tweets were correctly identified by the model. However, the model's performance for positive and negative classes remained very low. The positive class achieved a high precision of 1.00 but a very low recall of 0.14, indicating that the model correctly predicted only a few data as positive, but those predictions were accurate. Meanwhile, both precision and recall for the negative class were 0.00, meaning that the model completely failed to recognize tweets labeled as negative. This issue is most likely caused by an imbalance in the dataset, where the number of tweets in each sentiment class was not evenly distributed.

3.7. Model Evaluation Results

The performance of the Naïve Bayes model was evaluated using four key metrics described earlier in Section 2: Accuracy, Precision, Recall, and F1-score. The evaluation results on the test dataset are summarized in Table 1.

The results indicate that the Naïve Bayes model achieved an overall accuracy of 60%. Although the model performs well in identifying neutral sentiments, its performance is weaker for positive and negative classes. This is primarily due to dataset imbalance, where the distribution of sentiment classes is not uniform.

3.8. Discussions

The findings suggest that Naïve Bayes, while efficient and easy to implement, has limitations in capturing complex sentiment nuances in Indonesian social media text. The model demonstrates strong recall for neutral sentiments, meaning it is effective in detecting non-polarized or descriptive discussions, but it performs poorly in identifying emotional tones such as sarcasm, criticism, or praise — which often appear in informal or mixed-language tweets.

In comparison, prior studies using Support Vector Machine (SVM) and FastText achieved higher accuracy, up to 84%. These models utilize semantic word embeddings, allowing better understanding of contextual meaning, while Naïve Bayes relies on statistical word frequency. However, Naïve Bayes remains advantageous for computational efficiency and quick implementation, especially in exploratory analysis or early-stage research. From the analytical results, three key observations can be drawn:

- Most tweets show neutral or factual expressions, reflecting balanced public opinion regarding player naturalization.
- Imbalanced datasets significantly degrade the model's ability to learn minority sentiment classes.
- Text normalization and stopword removal improve overall performance, but contextual interpretation is still limited.

Limitations of this study include:

- Small dataset size and lack of balance between sentiment categories.
- Limited semantic understanding using TF-IDF features.
- Exclusion of sarcasm and implicit sentiment in classification.

Future improvements may involve:

- Expanding the dataset with automatic tweet crawling and annotation.
- Incorporating Word2Vec, FastText, or BERT-based representations.
- Applying data resampling or augmentation techniques (e.g., SMOTE) to handle class imbalance.
- Comparing results across multiple algorithms (e.g., SVM, Logistic Regression, Random Forest).

4. Conclusion

Based on the results of the study, the Naïve Bayes classifier was applied to perform sentiment analysis on tweets containing keywords such as naturalization PSSI, naturalized players, descendant players, national team naturalization, and abroad players for the national team. Prior to modeling, all tweet data underwent a comprehensive preprocessing stage, which included cleaning, case folding, tokenization, stopword removal, and stemming to ensure consistency and reduce noise in the dataset. Automatic labeling was then conducted using sentiment-related keywords to categorize tweets into positive, negative, and neutral sentiments. The resulting dataset exhibited a noticeable class imbalance, with approximately 60% of tweets categorized as neutral, 36% as positive, and only 4% as negative. This imbalance significantly influenced the model's performance, particularly in detecting minority classes such as negative sentiments. The Naïve Bayes classifier achieved an overall accuracy of 65%, with a precision of 0.73, recall of 0.65, and an F1-score of 0.55. These results indicate that while the model performed adequately in identifying neutral and positive sentiments, its capability to correctly classify negative tweets remained limited.

This limitation is primarily attributed to the uneven data distribution and the relatively small proportion of negative samples, which caused the model to lean toward the dominant neutral class during prediction. Nonetheless, the classifier still demonstrates the robustness of the Naïve Bayes algorithm in handling large-scale text data and its suitability for baseline sentiment classification tasks. In general, the findings confirm that the Naïve Bayes method can be effectively used for sentiment analysis of social media content. However, its performance could be further improved through dataset balancing techniques such as oversampling, under sampling, or the application of hybrid methods like SMOTE. Additionally, future studies may consider employing more advanced models such as Support Vector Machine (SVM), Logistic Regression, or deep learning approaches (e.g., LSTM, BERT) to enhance classification accuracy and better capture linguistic nuances in user-generated content. Despite these challenges, this study provides valuable insights into public perception and sentiment dynamics surrounding PSSI's player naturalization policy, proving that Twitter data can serve as a rich and meaningful source for social sentiment analysis in the context of sports and national identity.

References

- [1] G. Sumantri and B. S. H. Marwoto, "Analisis Sentimen Di Twitter Terkait Tim Nasional Sepak Bola Indonesia Menggunakan Metode Support Vector Machine," *Jurnal Kajian dan Terapan Matematika*, vol. 10, no. 2, pp. 96–104, 2024, doi: 10.21831/JKTM.V10I2.19561.
- [2] A. B. Ferdiyawan, Nurashah, I. F. Muldani, W. Winanti, and S. Basuki, "Analisis Sentimen Pada Media Sosial Twitter(X) Terhadap Pemain Diaspora Di Tim Nasional Sepak Bola Indonesia Studi Kasus Pesepakbola Mees Hilgers Dan Eliano Reinjders Menggunakan Metode Naïve Bayes Dan Support Vector Machine (SVM)," *ipsikom*, vol. 13, no. 1, pp. 63–70, 2025, doi: 10.58217/IPSIKOM.V13I1.50.
- [3] T. Juniardi and C. A. Sugianto, "Analisis Sentimen Tim Nasional Sepak Bola Indonesia Di Turnamen Piala Dunia U-17 Indonesia Pada Twitter (X) Menggunakan Algoritma Naïve Bayes," *Jurnal Informatika dan Teknik Elektro Terapan*, vol. 12, no. 3S1, pp. 2830–7062, 2024, doi: 10.23960/JITET.V12I3S1.5188.
- [4] W. Alfiyani, D. A. Fatah, and F. Irfhamni, "Penerapan Algoritma Naïve Bayes Untuk Analisis Sentimen Pada Media Sosial X Terhadap Performa Tim Nasional Sepak Bola Indonesia Di Era Kepemimpinan Shin Tae-Yong," *JATI (Jurnal Mahasiswa Teknik Informatika)*, vol. 9, no. 3, pp. 3969–3977, 2025, doi: 10.36040/JATI.V9I3.13451.
- [5] R. Rikky, M. Graciela, and H. Irsyad, "Klasifikasi Opini Masyarakat Terhadap Naturalisasi Pemain Sepak Bola Menggunakan KNN dan SMOTE," *Applied Information Technology and Computer Science (AICOMS)*, vol. 3, no. 1, pp. 21–27, 2024, doi: 10.58466/AICOMS.V3I1.1547.
- [6] F. M. Sarimole and K. Kudrat, "Analisis Sentimen Terhadap Aplikasi Satu Sehat Pada Twitter Menggunakan Algoritma Naive Bayes Dan Support Vector Machine," *Jurnal Sains dan Teknologi*, vol. 5, no. 3, pp. 783–790, 2024, doi: 10.55338/SAINTEK.V5I3.2702.

- [7] N. Raisa and N. Riza, "Sentimen Analisis Terhadap Opini Masyarakat Mengenai Drama Korea Pada Twitter Menggunakan Metode Naïve Bayes," *JATI (Jurnal Mahasiswa Teknik Informatika)*, vol. 7, no. 2, pp. 1312–1320, 2023, doi: 10.36040/JATI.V7I2.6765.
- [8] S. Lutfiani, R. Astuti, and F. M. Basysyar, "Analisis Sentimen Pengaruh Media Sosial Terhadap Minat Beli Skincare Pada Remaja Di Indonesia Menggunakan Algoritma Naïve Bayes," *JATI (Jurnal Mahasiswa Teknik Informatika)*, vol. 8, no. 3, pp. 2957–2961, 2024, doi: 10.36040/JATI.V8I3.9614.
- [9] N. Musatari, W. Suardi, and U. Syukri, "Analisis Sentimen Media Sosial: Penerapan E-katalog Dalam Pengadaan Barang Dan Jasa Di Indonesia," *PRAJA: Jurnal Ilmiah Pemerintahan*, vol. 10, no. 3, pp. 193–200, 2022, doi: 10.55678/PRJ.V10I3.702.
- [10] M. G. R. Lubis, D. S. Sitompul, T. M. Giovanni, F. Ramadhani, and S. Dewi, "Evaluasi Kinerja Algoritma Support Vector Machine (SVM) Dalam Analisis Sentimen Publik Terhadap Naturalisasi Timnas Indonesia di Twitter," *Journal of Accounting Law Communication and Technology*, vol. 2, no. 1, pp. 81–89, 2025, doi: 10.57235/JALAKOTEK.V2I1.4180.
- [11] M. E. Bere, Y. P. K. Kelen, H. H. Ullu, and B. Baso, "Implementasi Algoritma Naive Bayes Classifier Terhadap Analisis Sentimen Kondisi Stunting di Indonesia Pada Media Sosial X," *Jurnal Nasional Komputasi dan Teknologi Informasi (JNKTI)*, vol. 7, no. 4, pp. 598–605, 2024, doi: 10.32672/jnkti.v7i4.7752.
- [12] M. T. Usniaty, Y. Setyawan, and R. D. Bakti, "Analisis Sentimen Tentang Implementasi Hak Asasi Manusia Di Indonesia Pada Media Sosial Twitter Menggunakan Naïve Bayes Dan Support Vector Machine," *Jurnal Statistika Industri dan Komputasi*, vol. 9, no. 1, pp. 75–84, 2024, doi: 10.34151/STATISTIKA.V9I1.4846.
- [13] N. Tonello, "Lecture Notes on Neural Information Retrieval," *arXiv*, 2022, Accessed: Dec. 13, 2025. [Online]. Available: <https://arxiv.org/abs/2207.13443v2>.
- [14] S. M. Abdullah, R. A. Sulaiman, and M. M. Ali, "Text preprocessing techniques for sentiment analysis: A review," *Journal of Theoretical and Applied Information Technology*, vol. 98, no. 14, pp. 2848–2860, 2020.
- [15] A. Krouska, C. Troussas, and M. Virvou, "The effect of preprocessing techniques on Twitter sentiment analysis," *IISA 2016 - 7th International Conference on Information, Intelligence, Systems and Applications*, 2016, doi: 10.1109/IISA.2016.7785373.
- [16] J. Jefriyanto, N. Ainun, M. Arif, and A. Ardha, "Application of Naïve Bayes Classification to Analyze Performance Using Stopwords," *Journal of Information System, Technology and Engineering*, vol. 1, no. 2, pp. 49–53, 2023, doi: 10.61487/JISTE.V1I2.15.
- [17] R. Alsini, "Analysis of Real Time Twitter Sentiments using Deep Learning Models," *Journal of Applied Data Sciences*, vol. 4, no. 4, pp. 480–489, 2023, doi: 10.47738/JADS.V4I4.146.
- [18] M. Garg, S. Kumar, and A. K. J. Saudagar, *Natural language processing and information retrieval : principles and applications*. CRC Press, 2025.
- [19] T. Mandl and J. M. Struß, *Information retrieval: core techniques, issues, current and future developments*. Edward Elgar Publishing, 2024.
- [20] S. Styawati, A. R. Isnain, N. Hendrastuty, and L. Andraini, "Comparison of Support Vector Machine and Naïve Bayes on Twitter Data Sentiment Analysis," *Jurnal Informatika: Jurnal Pengembangan IT*, vol. 6, no. 1, pp. 56–60, 2021, doi: 10.30591/JPIT.V6I1.3245.
- [21] A. E. Nugroho and F. D. Rahman, "Preprocessing optimization for sentiment analysis using Indonesian Twitter dataset," *Journal of Data Science and Its Applications*, vol. 3, no. 1, pp. 45–56, 2023, doi: 10.52502/jdsa.v3i1.150.
- [22] L. Liu, Y. Yu, and M. Sun, "A comparison of text classification models for sentiment analysis on social media," *IEEE Access*, vol. 9, pp. 146215–146229, 2021, doi: 10.1109/ACCESS.2021.3119052.
- [23] D. W. Putra and M. A. Hadi, "Optimization of text classification using TF-IDF and Naive Bayes for Indonesian language," *Jurnal Teknologi Informasi dan Ilmu Komputer (JTIIK)*, vol. 10, no. 2, pp. 287–295, 2023, doi: 10.25126/jtiik.202310287.
- [24] M. D. H. Jasy, S. Al Hasan, M. I. K. Sagor, A. Noman, and J. M. Ji, "A Performance Evaluation of Sentiment Classification Applying SVM, KNN, and Naive Bayes," *Proceedings - 2021 International Conference on Computing, Networking, Telecommunications and Engineering Sciences Applications, CoNTESA 2021*, pp. 56–60, 2021, doi: 10.1109/CONTESA52813.2021.9657115.