

UNDERSTANDING PUBLIC OPINION ON POLITICAL CANDIDATES THROUGH TWITTER SENTIMENT ANALYSIS: A COMPARATIVE STUDY OF FEATURE EXTRACTION

Amelia Devi Putri Ariyanto^{1*}, Fari Katul Fikriah²

^{1,2}Information Systems and Technology Study Program, Widya Husada University, Semarang, Indonesia

*Email: ¹ameliadev26@gmail.com, ²farichatulfikriyah45@gmail.com

(Received: 01 June 2025, Revised: 19 June 2025, Accepted: 22 July 2025)

Abstract

Presidential elections are crucial in a country's political dynamics and are increasingly discussed on social media platforms like Twitter. However, sentiment analysis of public opinion on these platforms faces significant challenges, such as large data volumes, diverse formats, and the complexity of informal language. The key challenge is choosing the most appropriate feature extraction technique and classification algorithm to address the unique characteristics of Indonesian-language tweets in the context of presidential elections. This study aims to compare the effectiveness of two feature extraction approaches—semantic based on BERT (Bidirectional Encoder Representations from Transformers) and statistical based on TF-IDF (Term Frequency-Inverse Document Frequency)—in sentiment analysis of Indonesian-language tweets related to the presidential election, using four classification algorithms: Support Vector Machine (SVM), Naive Bayes, K-Nearest Neighbors, and Decision Tree. The experimental results demonstrate that the combination of TF-IDF with SVM provides the best performance, with an accuracy of 85.1% and a macro f1-score of 0.81, outperforming the BERT approach used statically. These findings indicate that statistical approaches such as TF-IDF remain relevant and practical for short social media texts and emphasize the importance of choosing a method that suits the characteristics of the data and the context of the analysis.

Keywords: *BERT, Machine Learning, Presidential Election, Sentiment Analysis, Twitter*

This is an open access article under the [CC BY](#) license.



**Corresponding Author: Amelia Devi Putri Ariyanto*

1. INTRODUCTION

The presidential election is a crucial moment in the political life of a country. This process reflects the people's will and society's social, economic, and cultural conditions in a specific period [1]. With the development of digital technology, social media, such as Twitter, has become the leading platform for people to express their opinions, discuss, and voice their political aspirations [2]. Sentiment analysis on Twitter data related to the presidential election can provide in-depth insights into people's views on candidates, key issues, and developing political dynamics.

However, understanding patterns of public sentiment through social media is not a simple task because of several fundamental challenges that must be overcome. First, the enormous and growing volume of data on social media, such as Twitter, requires a capable computing infrastructure for real-time data collection, processing, and analysis. Second, the

diversity of data formats is an obstacle because data on social media consists of short texts, images, videos, or even memes, each of which has a different context and nuance in conveying messages. Third, there are limitations in understanding the context of language because the language used on Twitter is often informal, abbreviated, or uses slang and emoticons [3], [4]. These characteristics of Twitter text make sentiment analysis complex because Natural Language Processing (NLP) models must be able to capture these nuances and contexts [5], [6]. In addition, the use of sarcasm, irony, and humor often makes direct sentiment analysis inaccurate.

Several previous studies have conducted sentiment analysis on the Indonesian language. Research by Cahyanti et al. [7] who conducted sentiment analysis related to the election of presidential candidates in 2024 using TF-IDF feature extraction and several classifier algorithm models such as Naïve Bayes, Random Forest, and SVM, where the

highest F1 Score value was obtained when using the Random Forest algorithm (89.84%). Research by Firmansyah et al. [8] has also conducted sentiment analysis for the 2019 presidential election using TF-IDF with classifier algorithms such as KNN and SVM. Research by Firdaus et al. [9] conducted sentiment analysis related to the 2024 presidential election using SVM, obtaining the highest accuracy value of 79%.

This study is different from previous studies because it compares two feature extraction approaches for presidential election sentiment analysis: the transformer-based semantic approach using BERT and the statistical approach using TF-IDF. The BERT approach utilizes a bidirectional model that can understand the context of words by looking at the entire sentence from the left and right sides, thus capturing semantic meaning in more depth [10]. In contrast, TF-IDF only measures the frequency and importance of words in a document without considering the semantic context. This study also explores several machine learning methods based on distance, probability, function, and decision trees for sentiment classification to determine the most effective feature extraction and classification methods in the context of presidential election sentiment analysis.

The main contribution of this research is to provide a comprehensive understanding of the effectiveness of two different feature extraction approaches—TF-IDF and IndoBERT—in the context of sentiment analysis of Indonesian-language tweets, as well as to examine the performance of various classical classification algorithms in handling political opinion data from social media. This study also provides practical insights for researchers and developers of public opinion analytics systems in selecting the appropriate method based on data characteristics and analysis objectives.

2. RESEARCH METHOD

Figure 1 is a flowchart of the sentiment analysis process on Twitter text data that will be carried out in this study. The overall flow of the method for sentiment analysis on Indonesian Twitter text, including data preparation, data preprocessing, feature extraction, sentiment analysis, and evaluation results, will be explained in depth in the form of sub-sections.

2.1 Data Preparation

The research dataset uses a public dataset [11], which is a dataset for sentiment analysis related to the 2024 presidential candidates taken from the Twitter platform. This dataset contains several presidential candidates, such as Ganjar Pranowo, Prabowo Subianto, and Anies Baswedan. The range of tweet text retrieval was carried out from October 2022 to April 2023, with 29,731 tweets used in this study. A total of 21,654 are included in the positive sentiment

label, while 8,074 are classified as negative labels, with examples of sentiment labels as in Table 1.

2.2 Data Preprocessing

Preprocessing or text processing is converting unstructured data into structured data that can be adjusted to needs so that it will be easier for the data to be processed to the next stage [12]. The preprocessing consists of several stages, namely tokenization and removing stopwords. Tokenization is converting documents into words by word by removing spaces. Stopword removal, which is removing words that often appear in a document but do not have informative or significant value to the document, will be removed at this stage [13].

Table 1. Example of Sentiment Labels

Twitter text in Indonesian	Twitter text in English	Sentiment Labels
<i>lanjut pak anies kita kawal sampai jadi presiden</i>	continue pak anies we escort until become president	Positive
<i>anies mundur dari calon presiden menyerahkan sepenuhnya pada kpu siapa yg akan dijadikan presiden</i>	anies resigns from being a presidential candidate, leaving it entirely up to the kpu who will be made president	Negative

2.3 Feature Extraction

This study will compare two main approaches in sentiment analysis: the semantic feature-based approach using BERT and the statistical feature-based approach using TF-IDF (Term Frequency-Inverse Document Frequency), which will measure how important a word is in a document relative to the entire corpus. TF-IDF combines two main components: Term Frequency (TF), which reflects the frequency of occurrence of a word in a document, and Inverse Document Frequency (IDF), which calculates how rarely the word appears across documents [4]. Thus, TF-IDF gives higher weight to unique and important words to detect sentiment.

Meanwhile, BERT-based features are extracted using a feature-based approach, which utilizes the vector representation of a pre-trained BERT model without retraining (fine-tuning). This study uses Simple Transformers, a Python library built on Hugging Face's Transformers Library, to facilitate this BERT model. Simple Transformers simplifies the implementation of Transformer models, including IndoBERT, which has been specially trained on Indonesian text data. With Simple Transformers, extracting semantic features from IndoBERT becomes more practical and can be easily used as input for sentiment classification models. By comparing semantic features from BERT (with Simple Transformers) and statistical features from TF-IDF, this study is expected to determine the most effective approach for sentiment analysis in Indonesia [14].

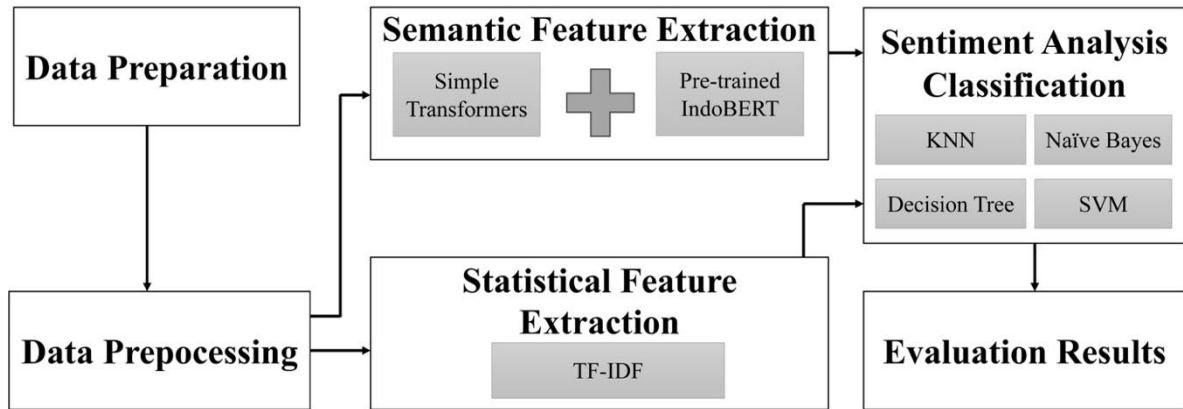


Figure 1 Flowchart of Presidential Election Sentiment Analysis

Table 2. Evaluation of the test set for semantic features with IndoBERT

Model	Precision	Recall	F1-Score	Accuracy
NB	0.730	0.700	0.710	0.700
SVM	0.810	0.810	0.810	0.810
DT	0.780	0.770	0.770	0.770
KNN	0.790	0.800	0.790	0.800

Table 3. Evaluation of the test set for statistical features with TF-IDF

Model	Precision	Recall	F1-Score	Accuracy
NB	0.774	0.784	0.755	0.784
SVM	0.848	0.851	0.849	0.851
DT	0.824	0.821	0.822	0.821
KNN	0.772	0.627	0.646	0.628

2.4 Sentiment Analysis

This study uses four machine learning algorithms to classify sentiment analysis in the context of the presidential election. The algorithms used include decision tree-based approaches (Decision Tree), probability-based (Naïve Bayes), and distance-based (K-Nearest Neighbors). The Decision Tree algorithm breaks the dataset into smaller subsets through decision branches. The algorithm selects the most relevant features at each node to separate the data into different classes. This process continues until each branch contains homogeneous data. To make a prediction, the model follows the path from the root to the tree's leaves according to the input data's characteristics until the predicted class is obtained [15].

Meanwhile, the K-Nearest Neighbors (KNN) algorithm stores the entire training dataset and utilizes the proximity of new data to the training data to predict its class. The distance between data is usually measured using Euclidean Distance. The predicted class is determined based on the majority of labels from the K nearest neighbors of the new data [16].

The Naïve Bayes algorithm calculates the probability of data belonging to a particular class based on the distribution of its features. This model assumes that each feature is independent of each other, meaning that the value of a feature does not affect the value of another feature. Thus, using Bayes' Theorem, the algorithm calculates the posterior probability of each class and assigns the class with the highest probability as the prediction [17].

Support Vector Machine (SVM) works by finding a hyperplane that optimally separates data from two classes in the feature space [18]. This algorithm tries to maximize the margin between the two classes to minimize the classification error. Thus, SVM is very effective for handling high-dimensional datasets and is often used in various text applications, including sentiment analysis. Using these four algorithms, this study compares the performance of each approach to determine the most appropriate classification model for analyzing sentiment toward the presidential election.

2.5 EVALUATION METRICS

The model's performance in detecting sentiment was assessed using several evaluation metrics, namely precision, recall, and F1 score. Precision measures how accurate the positive predictions produced by the model are, namely, how many positive predictions are positive classes. Recall evaluates how well the model finds all the positive data in the dataset. Meanwhile, the F1 score is calculated as the harmonic mean between precision and recall, thus balancing the two metrics [19]. Using these three metrics, this study can comprehensively assess the effectiveness of the model in classifying sentiment.

3. RESULT AND DISCUSSION

Implementing sentiment analysis for the 2024 Indonesian presidential election was carried out using a two-stage experimental approach: feature extraction and classification. In the first stage, two feature extraction approaches were carried out: a transformer-based semantic representation with IndoBERT and a TF-IDF-based statistical approach. IndoBERT was chosen because it is a BERT-based pre-trained model

optimized for Indonesian. It can capture nuances of meaning in short texts full of informal language variations, as commonly found on Twitter [20]. The implementation was carried out by loading the `indobenchmark/indobert-base-pl` model from the Huggingface Transformers library, then feature extraction was carried out by taking the embedding in the [CLS] token as a sentence representation. This process was carried out on training and test data using GPU devices for computational efficiency.

The dataset used is public [11], consisting of 29,731 labeled tweets, with a class distribution of 21,654 positive and 8,074 negative tweets. This data reflects public opinion on three presidential candidates, namely Prabowo Subianto, Anies Baswedan, and Ganjar Pranowo. Furthermore, the extracted features are used as input for four classic machine learning models, namely Naive Bayes (NB), Support Vector Machine (SVM), Decision Tree (DT), and K-Nearest Neighbor (KNN). These models are tested using testing data (10% of the data). For the TF-IDF approach, vectorization is performed with `TfidfVectorizer` using unigrams and bigrams, with a maximum of 5,000 features. The training and evaluation process follows the same procedure as the IndoBERT approach. Furthermore, the extracted features are used as input for four classical machine learning models: Naive Bayes (NB), Support Vector Machine (SVM), Decision Tree (DT), and K-Nearest Neighbor (KNN). All models were implemented with default parameters from the Scikit-learn library, except for SVM, which used a linear kernel and parameter $C=1.0$, and KNN, which used $k=5$ (default), as empirically, this value tends to provide stable performance on high-dimensional text data.

The test results demonstrate that the IndoBERT + SVM approach provides very competitive accuracy, with an accuracy value on the test set reaching 81% (Table 2), with a positive class f1-score of 0.88 and a negative class of 0.62. This indicates that the model can effectively identify positive tweets, which are the majority class, but also maintains decent performance in the negative class. The KNN and Decision Tree models with IndoBERT features also showed stable performance, with 80% and 77% accuracy, respectively. In contrast, Naive Bayes' performance was suboptimal on BERT features due to its limitations in handling negative feature values—a characteristic common in BERT's dense, high-dimensional embedding representations. Despite absolute conversion of feature values, Naive Bayes tends to assume independence between features and a Gaussian or multinomial data distribution, which is inconsistent with BERT's semantic vector output. This is one of the reasons for the low accuracy of NB compared to other models.

For the TF-IDF-based approach, the most prominent results were obtained from the combination of TF-IDF + SVM, with a test set accuracy of 85.1% (Table 3). The positive class F1-score reached 0.89,

and the negative class was 0.72, giving an average f1-score of 0.81 overall. This shows that although TF-IDF does not capture semantic context as deeply as IndoBERT, the optimal combination of n-grams and dominant opinion text characteristics can be modeled quite effectively by SVM. The Decision Tree model with TF-IDF also gave satisfactory results with an accuracy of 82% and a macro f1-score of 0.78. However, the performance of KNN in this approach decreased drastically, only achieving an accuracy of 62.7%, with a significant imbalance in the classification of the two classes (negative class recall is very high, but positive recall is low). This suggests that KNN's sensitivity to feature distribution and scale makes it less stable when the representation is not properly normalized.

Although theory and many literatures demonstrate that semantic representations such as IndoBERT have advantages in capturing the context of meaning, syntactic structure, and complex language nuances, experimental results show that in this case, statistical features based on TF-IDF combined with the SVM algorithm can achieve higher accuracy (85.1%) compared to the IndoBERT approach (81%). This phenomenon can be explained from several technical perspectives, as well as the characteristics of the data used.

First, the characteristics of language on Twitter, which tend to be short, direct, and explicit, make many opinions expressed with fairly repetitive and standardized words, such as "*dukung* (support)," "*pilih* (choose)," "*bagus* (good)," "*buruk* (bad)," "*korup* (corrupt)," or "*amanah* (trustworthy)." These patterns can be very well captured by TF-IDF, which measures the weight of the importance of a word in a document relative to the entire corpus. In this context, words frequently appearing in positive and negative tweets become strong sentiment markers, allowing models such as SVM to form very sharp classification boundaries.

Second, the large amount of data (almost 30 thousand tweets) with a strong dominance of the positive class also benefits the statistical approach. The TF-IDF feature can identify the most frequently occurring keywords in the majority class and then utilize this information to guess the class of new tweets efficiently. This contributes to improving the performance of metrics such as precision and recall, especially in the dominant positive class.

Third, although IndoBERT offers the power of semantic representation, this model produces high-dimensional and complex vectors, which do not always match optimally with classical models such as SVM or Decision Tree without fine-tuning. In this experiment, the IndoBERT embedding is used statically (without fine-tuning), so it cannot fully capture the context of specific domains such as political language or public opinion in elections. In contrast, TF-IDF does not experience a decrease in

performance because it does not require context adjustment.

Fourth, the accuracy of all models has not yet reached above 90%, which is generally the benchmark for advanced classification systems. This can be explained by several factors: (1) the presence of imbalanced data, where the positive class dominates, which can lead to model bias; (2) variations in informal language, abbreviations, and sarcasm in tweets that are difficult to capture by simple classification methods; and (3) the lack of advanced preprocessing, such as contextual fine-tuning for the political domain, which could increase the model's sensitivity to data nuances.

Fifth, although SVM and KNN tend to produce better results than Decision Tree and Naive Bayes, this is due to SVM's ability to form optimal classification margins on high-dimensional data such as TF-IDF or BERT, while KNN can capture local patterns if the data distribution is sufficiently balanced. On the other hand, Decision Tree is prone to overfitting on unstructured features, and Naive Bayes has a too strong assumption of independent distributions, which is unrealistic for text data.

Thus, although conceptually semantic features are superior in understanding the meaning of sentences, in the context of this experiment, the suitability between data characteristics, feature extraction methods, and classification algorithms is the dominant factor that causes the TF-IDF feature to produce higher performance. It is also important to note that selecting features and models is not always absolute but very contextual to the data type, domain, and analysis objectives.

4. CONCLUSION

This research conducted sentiment analysis on Twitter data related to the 2024 presidential election using two feature extraction approaches: semantic-based using IndoBERT and statistical-based using TF-IDF. Experimental results showed that the combination of TF-IDF with the Support Vector Machine (SVM) algorithm yielded superior classification performance compared to the IndoBERT feature in the classical model, with an accuracy of 85.1% and a macro f1-score of 0.81. This superior performance is influenced by the characteristics of Twitter language, which tends to be short, explicit, and contains repetitive and standardized words, such as "dukung (support)," "pilih (choose)," "bagus (good)," or "korup (corrupt)." These words are very effectively captured by the TF-IDF approach, which weights word frequency proportionally across the entire corpus, making it easier for the SVM to form sharp classification boundaries.

Furthermore, the large amount of data with a predominance of positive classes strengthens the effectiveness of the statistical approach, where keywords in the majority class serve as strong indicators for the model's predictions. In contrast,

IndoBERT features, while conceptually superior in capturing semantic nuances, produce high-dimensional and complex representations that are suboptimal when used statically without fine-tuning, especially when combined with classical models such as SVM or Decision Tree.

However, this study has several limitations. First, IndoBERT was used without fine-tuning, thus unable to fully capture the context of the political domain. Second, all models were developed with default parameters without exploratory hyperparameter tuning. Third, Naive Bayes did not perform optimally on IndoBERT embeddings due to limitations in handling negative values, while Decision Tree tended to overfit unstructured text features. Fourth, the accuracy of all models did not reach 90%, which could be due to data imbalance, the use of informal language, and the lack of advanced preprocessing such as lemmatization, slang normalization, or handling of sarcasm and emojis.

Furthermore, we acknowledge that this study did not apply imbalance handling techniques or any feature selection process. Both aspects are essential for improving model generalizability and performance, especially when working with real-world social media datasets. Their exclusion is a known limitation of this research, primarily due to scope constraints. Based on these findings, we recommend that future research explore the fine-tuning of Transformer-based models like IndoBERT for domain-specific adaptation. It is also worth considering ensemble or hybrid approaches that combine the strengths of semantic and statistical features. In addition, incorporating imbalance dataset handling techniques (e.g., resampling or cost-sensitive learning) and feature selection methods could significantly enhance model performance. Finally, future work should consider multimodal data (e.g., images, hashtags, videos) and develop robust NLP pipelines that better capture the informal, sarcastic, and dynamic nature of language on social media platforms.

Acknowledgment

This study was funded by a research grant from Widya Husada University for the academic year 2024/2025. The authors gratefully acknowledge the support of the Institute for Research and Community Service (LPPM) of Widya Husada University, whose assistance made this research possible.

5. REFERENCE

- [1] A. N. Ma'aly, D. Pramesti, A. D. Fathurahman, and H. Fakhurroja, "Exploring Sentiment Analysis for the Indonesian Presidential Election Through Online Reviews Using Multi-Label Classification with a Deep Learning Algorithm," *Information (Switzerland)*, vol. 15, no. 11, pp. 1–33, 2024, doi: 10.3390/info15110705.

- [2] A. D. P. Ariyanto, D. Purwitasari, B. Amaliah, C. Fatichah, M. G. Taqiuddin, and Haikal, "Annotated Data for Semantic Role Labeling of Crisis Events in Indonesian Tweets," *Data Brief*, p. 111688, 2025, doi: <https://doi.org/10.1016/j.dib.2025.111688>.
- [3] F. W. Edlim *et al.*, "Urgency Detection of Events Through Twitter Post: A Research Overview," *ICECOS 2024 - 4th International Conference on Electrical Engineering and Computer Science, Proceeding*, pp. 406–411, 2024, doi: [10.1109/ICECOS63900.2024.10791202](https://doi.org/10.1109/ICECOS63900.2024.10791202).
- [4] A. D. P. Ariyanto, F. K. Fikriah, and A. F. Setyawan, "Impact of Statistical and Semantic Features Extraction for Emotion Detection on Indonesian Short Text Sentences," *COMMIT (COMMUNICATION AND INFORMATION TECHNOLOGY) JOURNAL*, vol. 19, no. 1, pp. 1–13, 2025.
- [5] A. D. P. Ariyanto, D. Purwitasari, and C. Fatichah, "A Systematic Review on Semantic Role Labeling for Information Extraction in Low-Resource Data," *IEEE Access*, vol. 12, no. April, pp. 57917–57946, 2024, doi: [10.1109/ACCESS.2024.3392370](https://doi.org/10.1109/ACCESS.2024.3392370).
- [6] A. D. P. Ariyanto, C. Fatichah, and D. Purwitasari, "Semantic Role Labeling for Information Extraction on Indonesian Texts: A Literature Review," in *2023 International Seminar on Intelligent Technology and Its Applications (ISITIA)*, IEEE, Jul. 2023, pp. 119–124. doi: [10.1109/ISITIA59021.2023.10221008](https://doi.org/10.1109/ISITIA59021.2023.10221008).
- [7] R. Cahyanti, D. N. Maftuhah, A. B. Santoso, and I. Budi, "Twitter Sentiment Analysis Towards Candidates of the 2024 Indonesian Presidential Election," *JURNAL RESTI (Rekayasa Sistem dan Teknologi Informasi)*, vol. 10, no. 4, pp. 516–524, 2024, [Online]. Available: <http://jurnal.iaii.or.id>
- [8] F. Firmansyah *et al.*, "Comparing Sentiment Analysis of Indonesian Presidential Election 2019 with Support Vector Machine and K-Nearest Neighbor Algorithm," *6th International Conference on Computing, Engineering, and Design, ICCED 2020*, pp. 3–8, 2020, doi: [10.1109/ICCED51276.2020.9415767](https://doi.org/10.1109/ICCED51276.2020.9415767).
- [9] Asno Azzawagama Firdaus, Anton Yudhana, and Imam Riadi, "Analisis Sentimen Pada Proyeksi Pemilihan Presiden 2024 Menggunakan Metode Support Vector Machine," *Decode: Jurnal Pendidikan Teknologi Informasi*, vol. 3, no. 2, pp. 236–245, 2023, doi: [10.51454/decode.v3i2.172](https://doi.org/10.51454/decode.v3i2.172).
- [10] A. F. Setyawan, A. Devi, P. Ariyanto, F. K. Fikriah, and R. I. Nugraha, "Analisis Sentimen Ulasan iPhone di Amazon Menggunakan Model Deep Learning BERT Berbasis Transformer," *JURNAL ELEKTRONIKA DAN KOMPUTER (ELKOM)*, vol. 17, no. 2, pp. 447–452, 2024.
- [11] A. A. Firdaus, A. Yudhana, I. Riadi, and Mahsun, "Indonesian presidential election sentiment: Dataset of response public before 2024," *Data Brief*, vol. 52, p. 109993, 2024, doi: [10.1016/j.dib.2023.109993](https://doi.org/10.1016/j.dib.2023.109993).
- [12] A. T. Ni'mah and A. Z. Arifin, "Perbandingan Metode Term Weighting terhadap Hasil Klasifikasi Teks pada Dataset Terjemahan Kitab Hadis," *Rekayasa*, vol. 13, no. 2, pp. 172–180, 2020.
- [13] A. D. P. Ariyanto, L. A. A. Z. A. M. Maryamah, R. W. S, and R. I, "Metode Pembobotan Kata Berbasis Cluster Untuk Perangkingan Dokumen Berbahasa Arab," *Techno.Com*, vol. 20, no. 2, pp. 259–267, May 2021, doi: [10.33633/tc.v20i2.4357](https://doi.org/10.33633/tc.v20i2.4357).
- [14] A. D. P. Ariyanto, F. K. Fikriah, and A. F. Setyawan, "Emotion Detection Using Contextual Embeddings for Indonesian Product Review Texts on E-commerce Platform," *JURNAL ILMIAH KOMPUTER GRAFIS*, vol. 17, no. 1, pp. 179–185, 2024.
- [15] N. A. Maulidiyyah, T. Trimono, A. T. Damaliana, and D. A. Prasetya, "Comparison of Decision Tree and Random Forest Methods in the Classification of Diabetes Mellitus," *JIKO (Jurnal Informatika dan Komputer)*, vol. 7, no. 2, pp. 79–87, 2024, doi: [10.33387/jiko.v7i2.8316](https://doi.org/10.33387/jiko.v7i2.8316).
- [16] M. T. Nawawi and A. Suhendar, "IMPLEMENTATION OF MSME CREDIT LOAN DETERMINATION USING MACHINE LEARNING TECHNOLOGY WITH K-NN ALGORITHM," *JIKO (Jurnal Informatika dan Komputer)*, vol. 7, no. 3, pp. 217–221, 2024, doi: [10.33387/jiko.v7i3.9064](https://doi.org/10.33387/jiko.v7i3.9064).
- [17] K. Yulawan and S. Murib, "Comparison of Decision Tree and Naïve Bayes Algorithms in Predicting Student Graduation At Ypk Junior High School, Nabire Regency," *JIKO (Jurnal Informatika dan Komputer)*, vol. 7, no. 2, pp. 117–122, 2024, doi: [10.33387/jiko.v7i2.8506](https://doi.org/10.33387/jiko.v7i2.8506).
- [18] P. Wicaksono and S. Sriani, "Application of Support Vector Machine Algorithm for Students' Final Assignment Stress Classification," *JIKO (Jurnal Informatika dan Komputer)*, vol. 7, no. 2, pp. 138–144, 2024, doi: [10.33387/jiko.v7i2.8618](https://doi.org/10.33387/jiko.v7i2.8618).
- [19] F. A. Ramadhan and P. H. Gunawan, "Sentiment Analysis of 2024 Presidential Candidates in Indonesia: Statistical Descriptive and Logistic Regression Approach," *2023 International Conference on Data Science and Its Applications, ICoDSA 2023*, pp. 327–332, 2023, doi: [10.1109/ICoDSA58501.2023.10276417](https://doi.org/10.1109/ICoDSA58501.2023.10276417).
- [20] F. Koto, A. Rahimi, J. Lau, and T. Baldwin, "IndoLEM and IndoBERT: A Benchmark Dataset and Pre-trained Language Model for Indonesian NLP," in *Proceedings of the 28th International Conference on Computational Linguistics*, 2020, pp. 757–770. doi: [10.18653/v1/2020.coling-main.66](https://doi.org/10.18653/v1/2020.coling-main.66).