



## Public Sentiment Classification of Danantara in Social Media X Using Support Vector Machine and Random Forest

Primandani Arsi<sup>1)</sup>✉, Pungkas Subarkah<sup>1)</sup>, and Ranggi Praharaningtyas Aji<sup>2)</sup>

<sup>1</sup> Informatics, Faculty of Computer Science, Universitas Amikom Purwokerto, Banyumas, 53127, Indonesia

<sup>2</sup> Information System, Faculty of Computer Science, Universitas Amikom Purwokerto, Banyumas, 53127, Indonesia

### Article Info

#### Article History:

Received: 8 December 2026

Revised: 5 March 2026

Accepted: 23 April 2026

#### Keywords:

Danantara, Random Forest, Sentiment Analysis, Social Media X, Support Vector Machine.

### Abstract

The increasing use of social media as a platform for public discourse provides valuable data for understanding societal responses to national strategic policies. One prominent example is the establishment of Danantara (Daya Anagata Nusantara), a sovereign wealth fund launched by the Indonesian government in February 2025. This study aims to analyze public sentiment toward Danantara using Indonesian-language posts collected from social media platform X and to comparatively evaluate the performance of Support Vector Machine (SVM) and Random Forest (RF) algorithms. A dataset of 1,434 public tweets was collected through web scraping and processed using text preprocessing techniques, including cleaning, tokenization, stopword removal, stemming, and TF-IDF feature extraction. Sentiment labels were generated using an Indonesian RoBERTa model and validated by a linguistic expert. Class imbalance was addressed using the Synthetic Minority Oversampling Technique (SMOTE). Model performance was evaluated using 5-fold stratified cross-validation with accuracy, precision, recall, and F1-score metrics. Experimental results show that Random Forest achieved slightly superior performance, reaching an average accuracy of 91.47%, compared to 91.06% obtained by SVM. Confusion matrix analysis indicates that RF better distinguishes neutral sentiment, while SVM performs competitively in identifying strong sentiment polarity. This study contributes by providing the first empirical comparison of classical machine learning approaches for analyzing public sentiment toward Indonesia's sovereign wealth fund discourse, offering methodological insights and practical implications for data-driven policy evaluation using social media analytics.

## INTRODUCTION

The rapid development of information technology has significantly transformed the way people express opinions and interact, particularly through social media platforms such as X (formerly Twitter) (Arsi et al., 2021). With its open and real-time nature, X has become a dynamic medium for capturing public sentiment toward current issues, ranging from entertainment and politics to government policies (Firdaus et al., 2024). One of the recent national strategic policies that has drawn substantial public attention is the establishment of Danantara (Daya Anagata Nusantara), a sovereign wealth fund launched in February 2025 to manage state-owned assets and strengthen sustainable economic growth (Dongoran et al., 2025). While the initiative has been welcomed by some as an innovation in national asset management, others have raised concerns regarding its transparency and effectiveness, making Danantara a subject of extensive public debate on social media (Ayu et al., 2025).

In recent years, public sentiment expressed through social media has increasingly influenced policy perception, trust in government institutions, and public communication strategies. Understanding these perceptions is particularly important for newly introduced national economic institutions, where public acceptance may affect policy legitimacy and long-term sustainability. Despite the strategic importance of sovereign wealth funds, empirical studies examining public reactions toward such institutions using computational approaches remain limited, especially in the Indonesian context.

Sentiment analysis provides a systematic method for capturing and classifying public opinions into categories such as positive, negative, and neutral, enabling objective and data-driven insights into societal attitudes (Dayani et al., 2024). Previous studies have demonstrated its relevance in various policy-related contexts. For instance, an analysis of public opinion on the Personal Data Protection Law (UU PDP) using Support Vector Machine (SVM) achieved 75.89% accuracy across three sentiment categories (Abdul et al., 2024). Similarly, research on the relocation of Indonesia's capital city (IKN) compared SVM and Naïve Bayes, in which SVM achieved higher accuracy and an F1-score of 84% (Setiawan et al., 2024). Random Forest has also been applied to policy discourse, achieving 76% accuracy in classifying tweets about IKN and 77% accuracy when analyzing TikTok comments related to IKN development (Herdiyani et al., 2022). These

studies suggest that SVM often performs well in text classification tasks involving policy debates, whereas Random Forest remains effective for larger datasets with more diverse class distributions.

Beyond policy-related issues, SVM and Random Forest have also demonstrated strong performance in non-policy domains, such as application review analysis. Studies on ChatGPT reviews from the Google Play Store reported an SVM accuracy of 92.72% (Ernawati et al., 2024), while Random Forest achieved an average accuracy of 90% on reviews of the MyPertamina application (Indrayanto et al., 2023). These findings underscore the consistency of both algorithms in handling sentiment classification tasks across different domains. The evidence indicates that SVM is particularly suitable for binary classification tasks involving high-dimensional data, while Random Forest offers stability and robustness when working with imbalanced datasets. Therefore, both algorithms remain relevant for evaluation in new contexts such as public opinion on Danantara (Rahayu et al., 2024).

Although sentiment analysis has been widely applied to policy debates such as the Personal Data Protection Law and Indonesia's capital relocation, existing studies primarily focus on established political issues rather than newly formed economic institutions. Moreover, prior research rarely evaluates how different classical machine learning algorithms perform in analyzing public discourse surrounding sovereign wealth funds. Therefore, a research gap exists in understanding both public sentiment dynamics and algorithmic performance in the context of emerging national financial institutions such as Danantara. This study addresses this gap by conducting a comparative evaluation of Support Vector Machine and Random Forest models for sentiment classification of Danantara-related discussions on social media X, providing methodological and practical insights for policy communication analysis.

## RESEARCH METHODS

This study was conducted over a six-month period, from February to July 2025. The research employs a text mining and machine learning approach, with the workflow comprising problem identification, data collection, sentiment labeling, preprocessing, feature extraction, data balancing, model development, and performance evaluation.

The primary objective of this research is to systematically analyze public perceptions of Danantara using sentiment analysis and to

evaluate the effectiveness of Support Vector Machine (SVM) and Random Forest (RF) in classifying opinions into positive, negative, and neutral categories.

#### A. Data Collection

Data were collected through web scraping on the social media platform X using Python scripts, targeting public tweets containing the keyword *Danantara*. The scraping process was conducted from 24 February to 30 March 2025, yielding a total of 1,434 Indonesian-language tweets. The collected data were stored in CSV format for subsequent analysis. Additionally, a literature review was carried out to strengthen the theoretical framework of the study.

This study utilized publicly available data collected from social media platform X. No private or personally identifiable information was accessed or disclosed. Usernames and personal identifiers were anonymized during preprocessing to ensure data privacy and ethical compliance. The research follows ethical guidelines for social media research involving publicly accessible online content.

#### B. Labeling

Sentiment labeling was performed automatically using the Indonesian RoBERTa pre-trained model available on HuggingFace, which classified the tweets into positive, negative, or neutral categories (Hidayat et al., 2025). Each tweet was assigned a sentiment label along with a corresponding confidence score (Permana et al., 2025). To ensure the reliability of the automated labeling, the results were validated by a language expert with a background in Indonesian linguistics (Meriana et al., 2023).

#### C. Preprocessing

The preprocessing stage prepared the raw text for analysis by applying the following steps: Cleaning and Case Folding: Removing links, emoticons, hashtags, punctuation, numbers, and special characters, and converting all text to lowercase (Astuti et al., 2023). Tokenization: Splitting sentences into individual words or tokens (Setiawan et al., 2024). Stopword Removal: Eliminating common Indonesian words (e.g., *dan*, *yang*, *di*) that do not contribute significantly to sentiment classification (Halim et al., 2023). Stemming: Reducing words to their root form (e.g., *membeli*, *dibeli*, *pembelian* → *beli*) (Purnama et al., 2022).

#### D. Feature Extraction

The preprocessed text was converted into numerical vectors using the Term Frequency–

Inverse Document Frequency (TF-IDF) method (Rihastuti et al., 2025). This weighting technique measures the importance of a word within a document relative to its frequency across the entire corpus, enabling more effective representation of textual features for machine learning models (Putra et al., 2025).

#### E. Data Balancing

The initial dataset exhibited class imbalance, consisting of 645 neutral, 385 negative, and 138 positive tweets. To address this issue, the Synthetic Minority Over-sampling Technique (SMOTE) was applied to generate synthetic samples for the minority classes, resulting in a more balanced dataset for model training (Umam et al., 2024).

#### F. Modelling

Two classification algorithms were implemented in this study: Support Vector Machine (SVM), A supervised learning algorithm effective for high-dimensional feature spaces and capable of identifying optimal separating hyperplanes for sentiment classes (Hadi et al., 2025). Random Forest (RF): An ensemble-based algorithm that mitigates overfitting and demonstrates robustness when dealing with noisy or imbalanced datasets (Fachreza et al., 2024).

SVM was selected due to its effectiveness in high-dimensional sparse text representations, which are common in TF-IDF-based sentiment classification. Random Forest was included as a comparative ensemble-based approach known for robustness against noise and class imbalance. Comparing these two algorithms enables evaluation between margin-based learning and ensemble decision strategies within Indonesian social media discourse.

#### G. Evaluation

Model performance was evaluated using 5-fold stratified cross-validation, which ensured that each fold preserved the original distribution of sentiment labels (Firmansyah et al., 2023). The evaluation metrics included accuracy, precision, recall, and F1-score. The final performance values were obtained by averaging the results across all folds, while the best-performing fold was also reported to highlight the maximum potential of each algorithm (Wijayanto et al., 2024).

### RESULT AND DISCUSSION

#### A. Dataset Overview

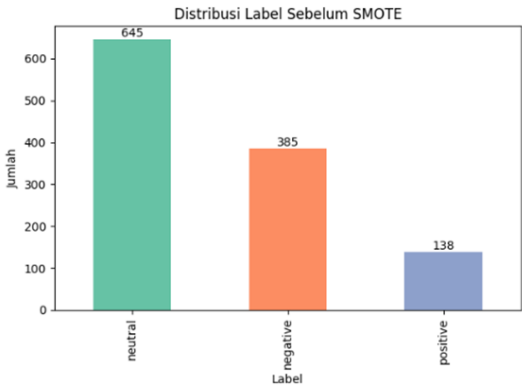
This study utilized 1,434 Indonesian-language tweets related to *Danantara*, collected from the social media platform X between 24

February and 30 March 2025. Following preprocessing and automatic sentiment labeling using the RoBERTa model, the dataset exhibited an imbalanced distribution, with neutral sentiment dominating (645 tweets), followed by negative (385 tweets) and positive (138 tweets).

Tables 1. Dataset Distribution Before SMOTE

| Sentiment | Number of Tweets |
|-----------|------------------|
| Positive  | 138              |
| Negative  | 385              |
| Neutral   | 645              |
| Total     | 1.168            |

To address this class imbalance, the Synthetic Minority Over-sampling Technique (SMOTE) was applied, resulting in a more balanced distribution across the sentiment categories.



The figures 1. Label distribution before SMOTE

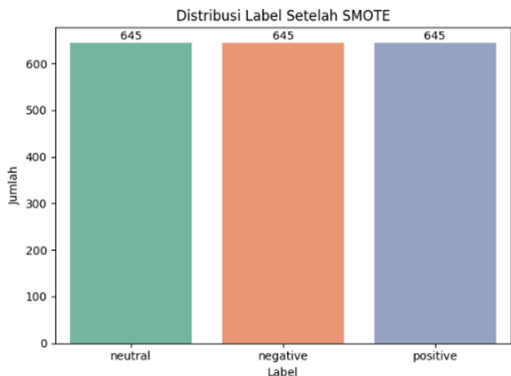


Figure 2. Label distribution after SMOTE

B. Model Evaluation

The balanced dataset was classified using two machine learning algorithms: Support Vector Machine (SVM) and Random Forest (RF). Model evaluation was performed using 5-fold stratified cross-validation, with performance measured through accuracy, precision, recall,

and F1-score. The average results across all folds are summarized in Table 2, indicating that Random Forest consistently outperformed SVM across all evaluation metrics.

Table 2. Average Performance Metrics for SVM and Random Forest Models

| Metric    | SVM (%) | Random Forest (%) |
|-----------|---------|-------------------|
| Accuracy  | 91.06   | 91.47             |
| Precision | 91.56   | 91.67             |
| Recall    | 91.06   | 91.47             |
| F1-score  | 90.90   | 91.44             |

In addition to the average results, the best-performing fold for each model is presented in Table 3. This provides insight into the maximum performance achieved by the algorithms during cross-validation.

Table 2. Best Fold Results of SVM and Random Forest

| Model         | Metric    | Best Fold | Score  |
|---------------|-----------|-----------|--------|
| SVM           | Accuracy  | 3         | 92.76% |
| SVM           | Precision | 3         | 93.16% |
| SVM           | Recall    | 3         | 92.76% |
| SVM           | F1-score  | 3         | 92.70% |
| Random Forest | Accuracy  | 3         | 93.80% |
| Random Forest | Precision | 3         | 93.95% |
| Random Forest | Recall    | 3         | 93.80% |
| Random Forest | F1-score  | 3         | 93.80% |

As shown in Table 3, Random Forest again achieved better results, with accuracy, precision, recall, and F1-score all reaching around 93.8%, compared to SVM's 92.7% accuracy.

C. Confusion Matrix

To provide a clearer picture of classification performance, the best confusion matrix (Fold 3 – highest accuracy) from each algorithm is presented in Figures 3 and 4. The SVM confusion matrix (Figure 2) shows generally strong performance, but it misclassified several neutral tweets as negative and some neutral tweets as positive. This indicates that SVM has greater difficulty separating the neutral class from the others.

In contrast, the Random Forest confusion matrix (Figure 3) demonstrates better classification capability, with fewer misclassifications across sentiment categories. Notably, Random Forest performed particularly well in distinguishing neutral tweets, which contributed to its superior overall performance.

Overall, the findings confirm that Random Forest (RF) outperformed Support Vector Machine (SVM) in classifying public sentiment

toward Danantara. Based on the 5-Fold Stratified Cross-Validation, RF consistently achieved slightly higher accuracy, precision, recall, and F1-score compared to SVM, as shown in the average evaluation results.

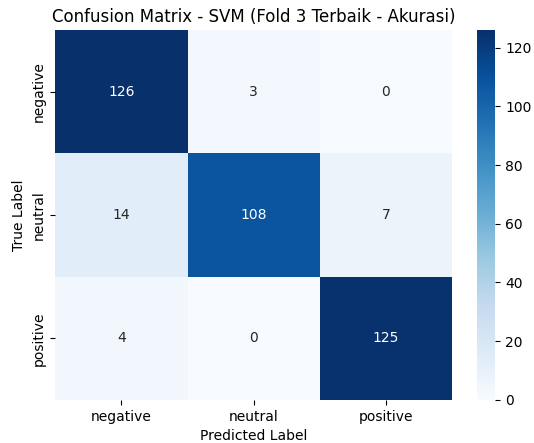


Figure 3. Confusion matrix of SVM (Fold 3 – best accuracy)

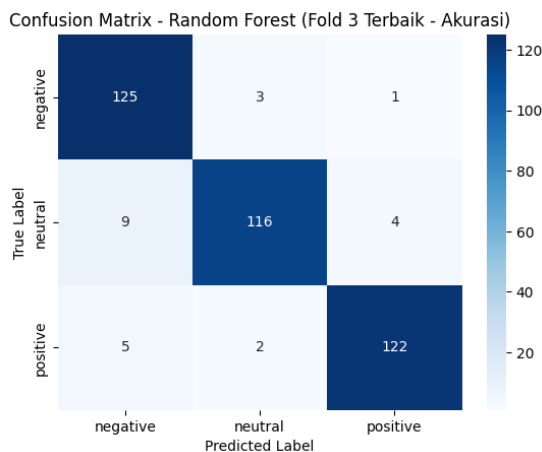


Figure 4. Confusion matrix of random forest (Fold 3 – best accuracy)

In addition, the best-performing fold (Fold 3) further highlighted RF's superiority, with accuracy, precision, recall, and F1-score all reaching around 93.8%, surpassing SVM's best accuracy of 92.7%. The confusion matrix analysis reinforced these results: SVM tended to misclassify neutral tweets as negative or positive, while RF demonstrated stronger capability in distinguishing between sentiment classes, particularly in the neutral category.

These results demonstrate that RF is a more effective and reliable algorithm for sentiment analysis of Indonesian social media data. Its consistent performance across folds and robustness in handling imbalanced textual data

make RF well-suited for analyzing public responses to national strategic policies such as Danantara.

#### D. Discussion

The findings of this study indicate that both Support Vector Machine (SVM) and Random Forest (RF) achieved strong performance in classifying public sentiment toward Danantara, with accuracy levels exceeding 90% following preprocessing, TF-IDF feature extraction, and data balancing using SMOTE. Although the difference in average performance between the two models was relatively small, Random Forest consistently demonstrated slightly higher accuracy, precision, recall, and F1-score compared to SVM, and also achieved superior results in its best-performing fold.

At the fold level, both models reached their optimal performance in Fold 3, where SVM attained 92.76% accuracy, while RF surpassed it with 93.80%. This suggests the greater stability of Random Forest and its potential for achieving optimal classification performance. The confusion matrix analysis further reinforces these findings: SVM performed well in identifying positive and negative sentiments, with high recall values (96.59% and 97.67%, respectively), but showed limitations in distinguishing neutral sentiment (recall 78.91%), often misclassifying it as negative or positive. In contrast, Random Forest demonstrated more balanced performance across all sentiment categories, with stronger recognition of neutral sentiment (recall 84.96%) while maintaining high recall for positive and negative classes.

From a linguistic perspective, the predominance of neutral sentiment (57.3%) indicates that most public responses were descriptive or cautious in nature, reflecting uncertainty or reservation toward Danantara. Negative sentiment (27.8%) highlights skepticism and criticism of the policy, whereas the relatively small proportion of positive sentiment (14.9%) suggests limited public support or appreciation. This interpretation was validated by a language expert with a background in Indonesian linguistics, ensuring that the sentiment labeling and subsequent analysis aligned with the contextual meaning of the data.

Overall, these results demonstrate that while SVM excels in identifying strong sentiment polarities (positive vs. negative), Random Forest provides more balanced and reliable performance across all sentiment categories, particularly for the more ambiguous neutral class. Consequently, both models offer complementary insights: SVM

may be more beneficial when the emphasis is on detecting clear sentiment polarities, whereas Random Forest is preferable for obtaining a comprehensive understanding of sentiment distribution. Importantly, the integration of such models into real-time public opinion monitoring systems could provide valuable input for policymakers, enabling more effective communication strategies and timely responses to societal concerns regarding national strategic policies such as Danantara.

Compared with previous sentiment analysis studies on Indonesian policy discourse reporting accuracy between 75% and 85%, the performance achieved in this study (>90%) suggests that preprocessing optimization and balanced training data significantly improve classification stability. However, differences in dataset characteristics and labeling strategies limit direct comparison, indicating the need for standardized benchmarks in Indonesian sentiment analysis research.

Despite promising results, several limitations should be acknowledged. First, the dataset size remains relatively moderate and limited to a specific time period following the policy announcement. Second, sentiment labeling relied partly on automated models, which may introduce semantic bias despite expert validation. Third, this study focuses on classical machine learning approaches and does not evaluate deep learning architectures such as BERT fine-tuning, which may yield different performance characteristics.

## CONCLUSION

This study analyzed public sentiment toward Danantara on the social media platform X using Support Vector Machine (SVM) and Random Forest (RF), based on a dataset of 1.434 Indonesian-language tweets. After preprocessing, automatic sentiment labeling using RoBERTa, feature extraction with TF-IDF, and class balancing using SMOTE, both models demonstrated strong classification performance, achieving accuracy levels above 90% under 5-fold stratified cross-validation. Sentiment labeling was further validated by a language expert with a background in Indonesian linguistics to ensure contextual consistency.

## REFERENCES

- Abdul, R. (2024). Analisis Sentimen Masyarakat Terhadap Undang-Undang Perlindungan Data Pribadi Pada Aplikasi X Dengan Metode Support Vector Machine. *Jurnal METHODIKA*, 10(1), 6–10. doi: <https://doi.org/10.46880/mtk.v10i1.2335>.
- Arsi, P. A., & Retno, W. (2021). Analisis Sentimen Wacana Pemindahan Ibu Kota Indonesia Menggunakan Algoritma

The experimental results revealed that while SVM performed well in detecting clear polarity between positive and negative sentiments, Random Forest consistently achieved slightly higher scores across all evaluation metrics. On average, RF achieved an accuracy of 91.47%, precision of 91.67%, recall of 91.47%, and an F1-score of 91.44%, surpassing SVM's accuracy of 91.06%. In the best-performing fold, RF reached 93.8% accuracy compared to SVM's 92.7%. Confusion matrix analysis further confirmed that RF exhibited superior performance in identifying neutral sentiment, which SVM frequently misclassified.

Beyond the technical evaluation, the sentiment distribution provides key insights for policymakers. The predominance of neutral sentiment indicates that much of the public remains cautious or undecided regarding Danantara, while the significant proportion of negative sentiment reflects skepticism and criticism that merit further attention. Meanwhile, the relatively limited positive sentiment suggests that public appreciation of the policy remains low.

In conclusion, Random Forest proved to be the more effective and reliable algorithm for sentiment analysis of Indonesian social media data in this study, particularly in capturing neutral sentiment alongside positive and negative categories. These findings highlight the potential of sentiment analysis as a valuable tool for informing government communication strategies and offering evidence-based insights into public perceptions of national strategic policies such as Danantara.

From a practical perspective, the findings demonstrate that sentiment analysis can serve as a complementary analytical tool for monitoring public perception toward newly established national policies. Policymakers may utilize such insights to design adaptive communication strategies and address public concerns more effectively.

Future research should explore larger longitudinal datasets, incorporate deep learning models, and examine topic-level sentiment dynamics to better capture evolving public discourse surrounding national economic initiatives.

- Support Vector Machine (SVM). *Jurnal Teknologi Informasi dan Ilmu Komputer*, 8(1), 147–156. doi: <https://doi.org/10.25126/jtiik.0813944>.
- Astuti, F. D., & Lenti, F. N. (2023). Perbandingan Algoritma Klasifikasi KNN, Gaussian Naive Bayes, dan Random Forest pada dataset LCMS tanaman keladi tikus yang diseimbangkan dengan metode Synthetic Minority Over-Sampling Technique. *Prosiding SINTaKS*, 2(1), 89–98. doi: <http://dx.doi.org/10.35842/sintaks.v2i1.27>.
- Ayu, C. D., Febiani, F., Ardhani, M. L., Syahwa, N., & Nuraya, A. S. (2025). Keterkaitan Danantara dengan Stabilitas Keuangan Makro di Indonesia. *Indonesian Research Journal on Education*, 5(2), 1026–1031. doi: <https://doi.org/10.31004/irje.v5i2.2418>.
- Dayani, A. D., Yuhandri, & Widi Nurcahyo, G. (2024). Analisis Sentimen Terhadap Opini Publik pada Sosial Media Twitter Menggunakan Metode SVM. *Jurnal KomtekInfo*, 11(1), 1–10. doi: [10.35134/komtekinform.v11i1.439](https://doi.org/10.35134/komtekinform.v11i1.439).
- Dongoran, D., & Sari, I. P. (2025). Implementasi Klasifikasi Data Tracer Study ... *Hello World Jurnal Ilmu Komputer*, 4(1), 12–24. doi: [10.56211/helloworld.v4i1.619](https://doi.org/10.56211/helloworld.v4i1.619).
- Ernawati, S., & Wati, R. (2024). Evaluasi Performa Kernel SVM ... *Jurnal Buana Informatika*, 15(1), 40–49. doi: <https://doi.org/10.24002/jbi.v15i1.7925>.
- Fachreza, R., & Handoko, W. T. (2024). Analisis Sentimen Masyarakat Terhadap Kinerja Trans Semarang ... *Progresif: Jurnal Ilmiah Komputer*, 20(2), 724–734. doi: [10.35889/progresif.v20i2.2093](https://doi.org/10.35889/progresif.v20i2.2093).
- Firdaus, A. A. F., Anton, & Imam, R. (2023). Analisis Sentimen Pada Proyeksi Pemilihan Presiden 2024 ... *Decode: Jurnal Pendidikan Teknologi Informasi*, 3(2), 236–245. doi: [10.51454/decode.v3i2.172](https://doi.org/10.51454/decode.v3i2.172).
- Firmansyach, W. A., Hayati, U., & Wijaya, Y. A. (2023). Analisa Terjadinya Overfitting dan Underfitting ... *Jurnal Mahasiswa Teknik Informatika*, 7(1), 262–269. doi: <https://doi.org/10.36040/jati.v7i1.6329>.
- Hadi, N., & Sugiarto, D. (2025). Analisis Sentimen Pembangunan IKN ... *Jurnal Informatika: Jurnal Pengembangan IT*, 10(1), 37–49. doi: [10.30591/jpit.v10i1.7106](https://doi.org/10.30591/jpit.v10i1.7106).
- Halim, S. F. N., & Azmi, U. (2023). Analisis Perbandingan Klasifikasi dan Penerapan SMOTE ... *Jurnal Sains dan Seni ITS*, 12(2), 127–134. doi: [10.12962/j23373520.v12i2.111833](https://doi.org/10.12962/j23373520.v12i2.111833).
- Herdayani, T. C., & Zailani, A. U. (2022). Sentiment Analysis Terkait Pemindahan Ibu Kota Indonesia ... *JTSI*, 3(2), 154–165. doi: <https://doi.org/10.35957/jtsi.v3i2.2920>.
- Hidayat, J. J., Setyowati, C., Amin, M. D. I., Bimasakti, K., & Werdana, A. P. (2025). Deep Learning-based Sentiment Analysis ... *Jurnal Media Computer Science*, 4(2), 277–292. doi: [10.37676/jmcs.v4i2](https://doi.org/10.37676/jmcs.v4i2).
- Indrayanto, C. G., Ratnawati, D. E., & Rahayudi, B. (2023). Analisis Sentimen Data Ulasan MyPertamina ... *JPTIHK*, 7(3), 1131–1139. doi: <https://j-ptiik.ub.ac.id/index.php/j-ptiik/article/view/12390>.
- Meriana, V., Wirayanti, N. K. A. W., Laksono, J. R. B., Binanto, I., & Sianipar, N. F. (2023). Perbandingan Algoritma KNN ... *ICCCI 2021*, 1(1), 1–3. doi: [10.35842/sintaks.v2i1.27](https://doi.org/10.35842/sintaks.v2i1.27).
- Permana, R. A., & Pratama, A. R. (2025). Analisis sentimen kenaikan harga BBM ... *Scientific Student Journal for Information, Technology and Science*, 6(1), 39–49.
- Purnama, S. J., & Chotib. (2022). Analisis Kebijakan Publik Pemindahan Ibu Kota Negara. *Jurnal Ekonomi dan Kebijakan Publik*, 13(2), 155–168. doi: [10.22212/jekp.v13i2.3486](https://doi.org/10.22212/jekp.v13i2.3486).
- Putra, K. T., Anggraini, S., Sutriani, L., & Impron, A. (2025). Analisis Sentimen Masyarakat Kalimantan Tengah ... *Journal Scientific of Mandalika*, 6(5), 1115–1123. doi: <https://doi.org/10.36312/10.36312/vol6iss5pp1115-1123>.
- Rahayu, S. P., Afuan, L., & Yunindar, G. A. (2025). Implementation of Text Mining ... *JUTIF*, 6(1), 359–368. doi: [10.52436/1.jutif.2025.6.1.4429](https://doi.org/10.52436/1.jutif.2025.6.1.4429).
- Rihastuti, S., & Rosyidi, A. (2025). Analisis Sentimen Pengguna Tiktok ... *JCS-TECH*, 5(1), 19–23.
- Setiawan, A., & Suryono, R. R. (2024). Analisis Sentimen IKN Menggunakan SVM dan Naïve Bayes. *Edumatic*, 8(1), 183–192. doi: [10.29408/edumatic.v8i1.25667](https://doi.org/10.29408/edumatic.v8i1.25667).
- Setiawan, A., Febrio, R. W., Purnama, M. A., & Efrizoni, L. (2024). Komparasi algoritma KNN, SVM, dan Decision Tree ... *Jurnal Informatika & Rekayasa Elektronika*, 7(1), 107–114. doi: <https://doi.org/10.36595/jjire.v7i1.1161>.

- Umam, F., Haryoko, A., & Ulum, M. (2024). Analisis Sentimen Pengguna Twitter Terhadap Mata Uang Kripto ... doi: <https://doi.org/10.55719/curtina.v5i2.1649>.
- Wijiyanto, W., Pradana, A. I., Sopingi, S., & Atina, V. (2024). Teknik K-Fold Cross Validation untuk Mengevaluasi Kinerja Mahasiswa. *Jurnal Algoritma*, 21(1), 239–248. doi: 10.33364/algoritma/v.21-1.1618.