



Volume XI Issue 3 Year 2026 | Page 841-852 | ISSN: 2527-9866

Received: 18-06-2026 | Revised: 20-07-2026 | Accepted: 08-08-2026

Sentiment Analysis of Indonesia's Economic Acceleration Program 2025 Using Support Vector Machine

Kamelia Lestari¹, Ermatita Ermatita²

^{1,2} Sriwijaya University, Palembang, South Sumatra, Indonesia,

e-mail: 09012622529004@student.unsri.ac.id¹, ermatita@unsri.ac.id²

*Correspondence: ermatita@unsri.ac.id

Abstract: Indonesia's Economic Acceleration Program 2025 is a government policy to accelerate national economic growth in response to the global economic slowdown. The implementation of this program has generated a variety of public responses on social media. This study aims to analyze public sentiment towards the Indonesian Economic Acceleration Program in the launch and implementation phases and identify differences in sentiment distribution in both phases. Data in the form of TikTok comments was collected through web scraping, then processed through preprocessing, lexicon-based sentiment labeling validated using manually labeled samples, TF-IDF feature representation, and classification using the Support Vector Machine. Model evaluation was carried out using 10-fold cross validation. The results of the study showed that SVM provided superior performance to the comparison model. In the launch phase, SVM achieved an accuracy of 81%, while in the implementation phase it achieved an accuracy of 79.5% with superior performance in all evaluation metrics. The distribution of sentiment in the launch phase was dominated by neutral sentiment by 67.1%, while in the implementation phase the proportion of negative sentiment increased to 46.38%. These results show that there is a difference in the distribution of sentiment between the launch and implementation phases, so that it can be an input in understanding the public's response to the Indonesian Economic Acceleration Program.

Keywords: Sentiment Analysis, Support Vector Machine, TF-IDF, Temporal Analysis, Public Policy

1. Introduction

The development of technology and communication today has brought significant changes in the political field, one of which is the use of social media that allows people to express opinions, criticisms, and support for a policy [1]. According to data from global digital research institutes, since January 2024 it shows that the number of social media users in the world has exceeded 5 billion [2]. Indonesia is one of the countries with the largest number of internet users in the world, which is more than 185 million users, of which around 139 million or 66.5% of them are active users of social media [3].

The high participation of the public in the digital space has resulted in various public opinions expressed in comments, uploads, reviews, and online discussions [4]. The collection of opinions describes public perceptions of various issues, services, and public policies [5]. Therefore, public policy not only affects social and economic conditions, but also directly affects public perception of the government [6]. Public perception is an important aspect in public policy evaluation, because it can reflect the level of acceptance, trust, and effectiveness of the policies that have been implemented [7].

On September 15, 2025, the Government of Indonesia through the Coordinating minister for Economic Affairs has launched a policy package known as the Indonesian Economic Acceleration Program or the 2025 Stimulus Package [8]. An acceleration program is a policy designed to accelerate the achievement of a goal in a shorter time than a regular program [9]. This program is present as a response to the global economic slowdown caused by geopolitical tensions, suppressing the openness of international trade, and the slowdown in cross-border trade which has an impact on declining export demand [10]. In addition, the occurrence of global inflation which shows an estimate of around 3.5% in 2024 and 2.9% in 2025 reflects continued price pressures in various countries [11], [12]. This condition has the potential to suppress international trade and weaken the economic recovery of various countries [13].

The Indonesian Economic Acceleration Program consists of eight policy aspects involving various economic sectors and is implemented gradually according to the readiness of each program [8]. Some policies will begin to be implemented in 2025, while other policies will be implemented in 2026 [8]. The division of implementation stages provides space for the government to evaluate the effectiveness of each policy before the entire program is implemented. In addition, the implementation that takes place gradually allows for a change in public perception in line with their experience in feeling the impact of policies. Therefore, sentiment analysis in the launch phase and implementation phase is important to understand the dynamics of changes in public opinion towards the Indonesian Economic Acceleration Program.

The implementation of the Indonesian Economic Acceleration Program has generated various public responses spread through social media. The difference in the level of digital literacy and the public's experience in receiving policy benefits has led to the emergence of diverse perceptions of the program [14], [15]. However, the policy evaluation carried out so far is still more focused on administrative indicators, such as budget absorption and the number of beneficiaries, while perception, satisfaction, and the level of public trust have not been the focus in policy evaluation [16]. In fact, these aspects are important regarding public acceptance and the sustainability of the implementation of a public policy.

Therefore, to deal with these challenges, an analytical approach is needed that is able to process public perception data on social media. The sentiment analysis approach with machine learning is a relevant solution because it is able to automatically process large amounts of text data so that the analysis process becomes more efficient [17]. This method can recognize complex and difficult to analyze patterns of public sentiment through manual observation [18]. With this ability, sentiment analysis can provide an overview of the tendency of public opinion towards a policy so that it can be considered in the policy evaluation process.

One of the algorithms *Machine Learning* widely used in sentiment analysis is the Support Vector Machine (SVM), as it has proven to be effective in classifying text data [19]. SVMs have a good ability to handle high-dimensional data, such as text data represented using TF-IDF [20]. This algorithm works by determining the optimal separator boundaries to differentiate each sentiment class [21]. In addition, SVM has good generalization ability and is relatively resistant to *Overfitting* in small to medium-sized datasets, which are commonly found in opinion data on social media [22]. Use of functions *kernel* also allows SVM to recognize more complex patterns in text data thereby improving classification results [20]. These characteristics make SVM suitable for classifying public sentiment towards the Indonesian Economic Acceleration Program based on comments on social media.

A number of previous studies have shown that the SVM algorithm performs well in text-based sentiment analysis. One of them is a study that analyzes public sentiment towards the Advanced Indonesia Cabinet using Twitter data by comparing SVM and Naive Bayes algorithms. The

results showed that the SVM algorithm obtained an accuracy of 89.60%, higher than Naive Bayes which obtained an accuracy of 85.74%, so that SVM was considered more effective in classifying text-based sentiments [23]. Another study also analyzed public sentiment on the issue of the DPR's right to inquiry related to the 2024 election using Twitter data with SVM algorithm and TF-IDF weighting. The results showed that the SVM model obtained 77% accuracy in classifying positive and negative sentiments. The findings show that SVM has adequate ability to identify public opinion on government policy issues that are developing on social media [24].

Another study analyzed public sentiment towards the Public Housing Savings (Tapera) policy using tweet data obtained from platform X. SVM and *Naive Bayes* by implementing SMOTE to address data imbalances. The results of the study show that public sentiment tends to be dominated by negative sentiment towards Tapera's policies. Moreover, the SVM algorithm produces better performance after SMOTE implementation with the highest accuracy of 93%, whereas *Naive Bayes* Achieved an accuracy of 89% [25].

Based on these studies, the SVM algorithm has shown good performance in classifying public sentiment towards various government policies. However, most studies still evaluate policy sentiment at a certain time, so it has not been able to explain how public perception develops as policy implementation takes place. As a result, the dynamics of changing public opinion at various stages of policy have not been widely researched. In addition, research on the 2025 Indonesian Economic Acceleration Program that compares public sentiment between the launch phase and the implementation phase is still very limited.

Therefore, this study analyzes public sentiment towards the 2025 Indonesian Economic Acceleration Program by comparing public perception in the launch phase and the policy implementation phase using the SVM algorithm. In contrast to previous studies that generally evaluated sentiment over a specific period of time, this study observed changes in public opinion trends at the time of program launch and during the implementation phase, thus providing an overview of the dynamics of public response to policies over time. The results of the research obtained are expected to provide benefits in social media-based sentiment analysis as one of the supporting approaches to evaluate government policies in a sustainable manner.

2.Methods

A. Data Description

The data used in this study is secondary data in the form of user comments on TikTok social media that discuss the 2025 Indonesian Economic Acceleration Program. The data was divided into two groups based on the time of data collection, namely the launch phase and the program implementation phase, to illustrate the distribution of public sentiment in each policy phase. Data on the launch phase was collected from uploads published on September 18, 2025 through several official accounts, namely IDX Channel Official, Liputan6, Ministry, and Official iNews, with a total of 4,000 comments. Meanwhile, data in the implementation phase was collected from official account uploads discussing the implementation of the 2025 Indonesian Economic Acceleration Program, including KompasTV, the Indonesian Ministry of Agriculture, and MetroTV, with a total of 4,000 comments.

Data collection in the implementation phase began after one of the programs in the 2025 Indonesian Economic Acceleration Program began to be implemented on September 20, 2025. Thus, the dataset in the implementation phase represents the community's response to the implementation of the policy in that period, while the dataset in the launch phase represents the community's response at the time the policy was announced. A comparison of the two phases was carried out to provide an overview of the difference in sentiment distribution in the two observation periods

B. Research Methods and Evaluation

1. Research Stages

This research was carried out through several main stages, namely data collection, *text preprocessing*, sampling, data labeling, feature extraction, modeling, and model evaluation. The overall research flow is shown in figure 1.

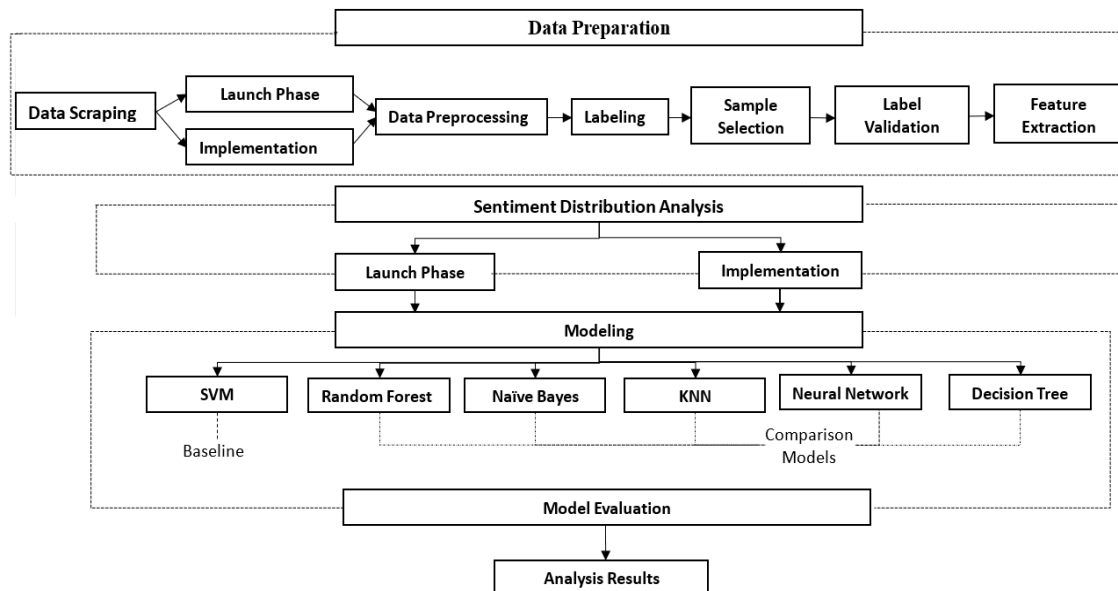


Figure 1. Research Stages

Figure 1 shows the research flow that began with data collection through *web scraping* of TikTok users' comments regarding the 2025 Indonesian Economic Acceleration Program. The dataset is separated into the launch phase and implementation phase, then *text preprocessing* is carried out which includes *cleaning*, *case folding*, *word normalization*, *tokenization*, *stopword removal*, and *stemming*. Furthermore, each comment is labeled using a *lexicon-based* approach. To evaluate the reliability of labeling, samples were determined using the Slovin formula with an error rate of 5%, then manual labeling was carried out by one annotator based on *sentiment polarity* guidelines. The results of manual labeling are compared with *lexicon-based* labeling to measure their suitability. After that, the data was represented using TF-IDF as a feature extraction technique, followed by the analysis of the sentiment distribution in both phases and the development of a classification model using SVM. Model performance was evaluated using *accuracy*, *precision*, *recall*, and *F1-score metrics*.

2. Data Preprocessing

The data that has been collected then goes through the *Text Preprocessing* which includes *cleaning*, *case folding*, *word normalization*, *tokenization*, *stopword removal*, and *Vote*. At the stage *Cleaning*, URL, *Mention*, *Hashtag*, emojis, symbols, excess punctuation, and numbers removed to reduce *Noise*. Especially in the implementation phase dataset, the figure is still maintained because it is related to one of the programs in the 2025 Indonesian Economic Acceleration Program, so it has the potential to affect sentiment analysis. Furthermore, the entire text is lowercase to maintain consistency of the feature representation at the TF-IDF extraction stage. Thus, capital letters and non-capitals are not considered as two distinct features by the model [26]. The next stage is the normalization of words by converting non-standard words into standard words using a dictionary containing about 15,500 word pairs arranged in an Excel file. After that it is done *Tokenization* Using the *Whitespace Based Tokenization*, which is appropriate because the text has gone through a process of cleanup and normalization so that each word separated by spaces can be considered a single token [27]. Furthermore, words that do not contribute to the meaning of the sentiment are removed using *Library Literature*,

then carried out *Vote Using Library* the same to change the adjective word into its basic form so that the variation of the word with the same meaning is represented as a single feature [28].

3. Data Labeling

In this study, sentiment labeling was carried out using the *lexicon-based* with three categories, namely positive, negative, and neutral. The use of three classes was chosen because comments on public policy not only contain support or rejection, but also contain a lot of questions and information that do not clearly show sentimental tendencies. To evaluate the reliability of labeling results *lexicon-based*, the number of samples is determined using the Slovin formula with an error rate of 5%, then selected using stratified random sampling so that each sentiment class is still represented. The launch phase consisted of 361 samples consisting of 101 negative comments, 243 neutral, and 17 positive, as well as 363 samples in the implementation phase consisting of 150 negative, 150 neutral, and 63 positive comments. The entire sample was then manually labeled by one independent annotator who was unaware of the labeling results *Lexicon based*. To maintain the consistency of the annotation process, the researcher developed labeling guidelines based on the concept *Sentiment polarity*, i.e. positive for comments that contain support or appreciation, negative for comments that contain criticism or rejection, and neutral for comments that are questionable or do not contain clear sentimental tendencies [29]. Table 1 presents the concept of polarity sentiment for manual sample labeling.

Table 1. Manual Labeling Concept

Label	Criteria
Positive	1. Stating a helpful or helpful policy 2. Provide support or approval 3. Saying Thank You 4. Containing hope 5. Showing satisfaction
Negatives	1. Criticizing policy implementation 2. Declaring the policy ineffective 3. Expressing disappointment 4. Contains sarcasm or sarcasm 5. Mention the incompatibility between theory and practice
Neutral	1. In the form of questions 2. Informative 3. Not showing support or rejection 4. Contains no emotional evaluation

4. Feature Extraction

At this stage, the text that has gone through the preprocessing process is converted into numerical form using the TF-IDF method. Numerical representation is performed on each training data within each fold cross validation before use, to classify the test data.

5. Model Support Vector Machine

In this study, SVM was used to classify comments into three sentiment classes, namely positive, negative, and neutral. Since SVM is essentially a binary classification algorithm, multiclass classification is carried out using the One vs Rest (OvR) strategy. The model was implemented using LinearSVC with parameter C with a value of 1.0 and class weight to overcome the imbalance of the amount of data in each sentiment class, while the other parameters used from the scikit-learn library.

$$f(x)=w.x+b \quad (1)$$

In Equation 1, w is the weight vector, x is the input feature vector, while b is the bias value [30].

6. Model Evaluation

Model evaluation was carried out using stratified 10-fold cross validation. In each fold, the TF-IDF feature formation process is used for evaluation. The accuracy, precision, recall, and F1-score values are calculated on each fold, then averaged as the final performance of the model. The calculation of evaluation metrics is based on the components of the Confusion Matrix, namely True Positive (TP), True Negative (TN), False Positive (FP), and False Negative (FN).

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (2)$$

Equation 2, TP is the number of correct predictions for the positive class, TN is the number of correct predictions for the negative class. Meanwhile, FP is a class that actually does not belong to the positive class, FN is a class that actually does not belong to the negative class.

$$Precision = \frac{TP}{TP + FP} \quad (3)$$

Equation 3 TP is the number of correct predictions for the positive class, while FP is the actual class that does not belong to the positive class. *Accuracy* used to measure the degree of accuracy of a model in predicting a sentiment class [31].

$$Recall = \frac{TP}{TP+FN} \quad (4)$$

Equation 4 TP is the sum of the correct predictions for the positive class, FN The actual class does not include the negative class. *Recall* is used to measure the model's ability to identify all the data that actually belongs to a sentiment class [31].

$$F1 - Score = 2x \frac{precision \times recall}{precision+recall} \quad (5)$$

Equation 5 *F1 Score* is the harmonic mean between *Accuracy* and *Recall* which is used to evaluate the balance between prediction accuracy and the model's ability to identify all relevant data [31].

3. Results and Discussion

A. Data Analysis Results

The *preprocessing* results showed that from each of the 4,000 comments collected in each phase, 3,623 net data were obtained in the launch phase and 3,855 net data in the implementation phase. The number of data eliminated in the launch phase was higher because 139 duplicate data and 238 data with *empty strings* were found, while in the implementation phase there were only 70 duplicate data and 75 data with *empty strings*. The results of data preprocessing are presented in table 2.

Table 2. Data Preprocessing Results

Phase	Remarks	Quantity
Launch	Raw Data	4000
	Duplicate	139
	Empty variable	0
	Empty String	238
Total Net Data		3.623
Implementation	Raw Data	4000
	Duplicate	70
	Empty Variables	0
	Empty string	75
Total Net Data		3.855

B. Label Validation

The results of the calculation using the confusion matrix showed that lexicon-based labeling had a degree of suitability as shown by an *accuracy* value of 80% in 361 sample data in the launch phase and 82% in 363 sample data in the implementation phase. These results showed that most of the labels produced with *lexicon-based* were consistent with the results of manual labeling.

C. Sentiment Distribution

The results of the sentiment distribution analysis showed a shift in public sentiment between the launch phase and the implementation phase. The results of the launch phase were dominated by neutral sentiment, while in the implementation phase it was dominated by negative sentiment. The results of this study are presented in figure 2.

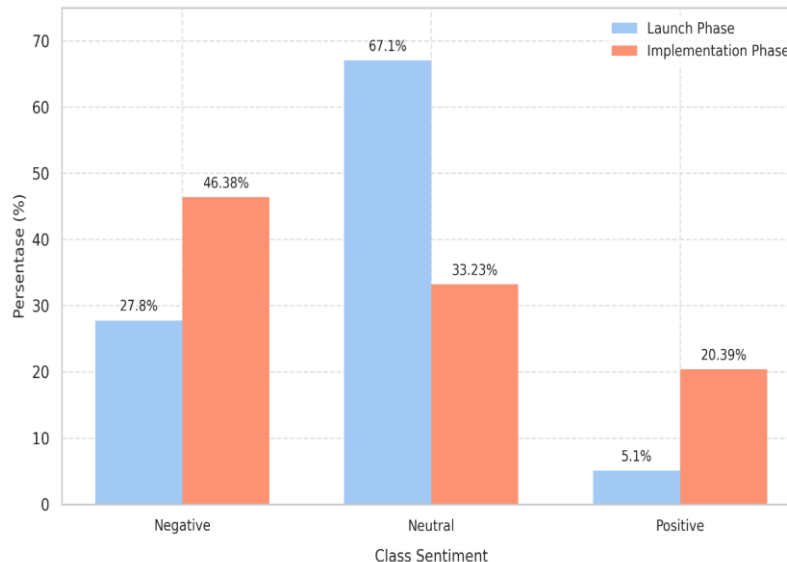


Figure 2. Results of the Second Phase of Sentiment Distribution

Figure 2 shows that in the launch phase, sentiment was dominated by neutral sentiment of 67.1%, followed by negative sentiment of 27.8%, and positive sentiment of 5.1%. The dominance of neutral sentiment shows that at the launch stage, most of the comments are still in the form of informative responses, such as opinions, expectations, and questions related to the announced program. On the other hand, in the implementation phase, there was a change in sentiment distribution, namely negative sentiment increased to 46.38% and positive sentiment increased to 20.39%, while neutral sentiment decreased to 33.23%. These results show that public comments in the implementation phase are more diverse than in the launch phase. An

increase in negative sentiment indicates that more comments contain negative criticism or judgment, while an increase in positive sentiment indicates an increase in comments containing support or appreciation for the program. These results are also supported by the word frequency analysis presented in Figure 3.

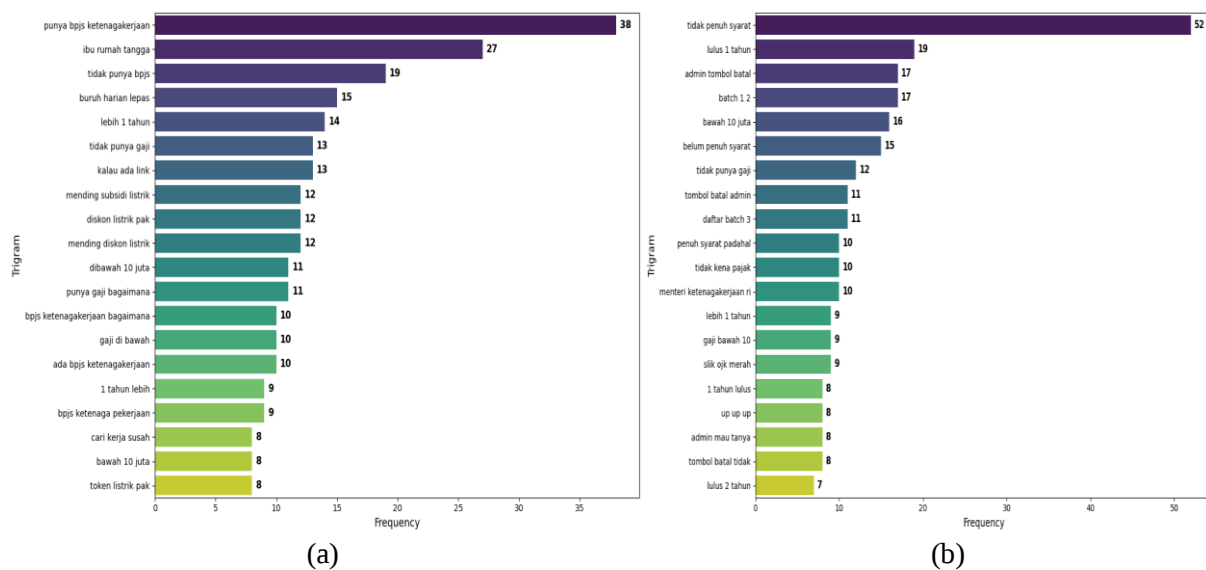


Figure 3. Word Frequency (a) the launch phase and (b) the implementation phase.

Figure 3(a) shows that in the launch phase the discussion was dominated by words related to the requirements of beneficiaries and the mechanism of the Indonesian Economic Acceleration Program, such as "Punya BPJS Ketenagakerjaan", "Ibu rumah tangga", and "tidak punya BPJS". In addition, the appearance of phrases such as "mending subsidi listrik" and "mending diskon listrik" indicates that there are comments that compare the program with other forms of assistance. This result is in line with the distribution of sentiment that was still dominated by the neutral class in the launch phase.

Figure 3(b) shows that in the implementation phase, the discussion shifted to the aspect of program implementation. The dominant words, such as "tidak penuh syarat", "slik OJK merah", and "dibawah 10 juta", indicate that the comments are more about administrative requirements and the mechanism for implementing the program. This shift in the topic of discussion is in line with the distribution of sentiments that became more diverse in the implementation phase, marked by an increase in negative and positive sentiments.

D. Model Performance Evaluation

The results of the evaluation showed that SVM performed in the launch phase with an accuracy of 81%. Meanwhile, in the implementation phase of the SVM model it also showed an accuracy of 79.5%. The results of the performance evaluation of all models are presented in detail in Table 3.

Table 3. Model Performance Results

Phase	Models	Accuracy	Accuracy	Recall	F1 Score
Launch Phase	SVM	81%	68,9%	69,3%	68,9%
	NN	79,3%	68,1%	63,8%	65,4%
	Random Forest	80%	74,6%	55,7%	58,5%
	Decision Tree	73,2%	58,6%	62,1%	59,9%
	Naive Bayes	74,9%	51,1%	44,5%	44,8%
	KNN	67,7%	60,4%	36,5%	33,3%
Implementation Phase	SVM	79,5%	77,9%	78,1%	77,9%
	NN	78,1%	77%	76,6%	76,7%

Random Forest	76,6%	75,8%	73,5%	74,2%
Decision Tree	72,8%	70,8%	71,2%	70,8%
Naive Bayes	61,1%	66,2%	54,3%	55,5%
KNN	46%	58,7%	45%	42,3%

Based on Table 3, in the launch phase SVM obtained the best performance with an accuracy of 81%, precision of 68.9%, *recall* of 69.3%, and an *F1-score* of 68.9%. In the implementation phase, SVM again showed the best performance with an accuracy of 79.5%, *precision* of 77.9%, *recall* of 78.1%, and an *F1-score* of 77.9%. These results show that SVM provides the most consistent classification performance compared to the comparative model in both phases of the study. The results of the evaluation were then strengthened through the *Confusion Matrix* presented in Figure 4.

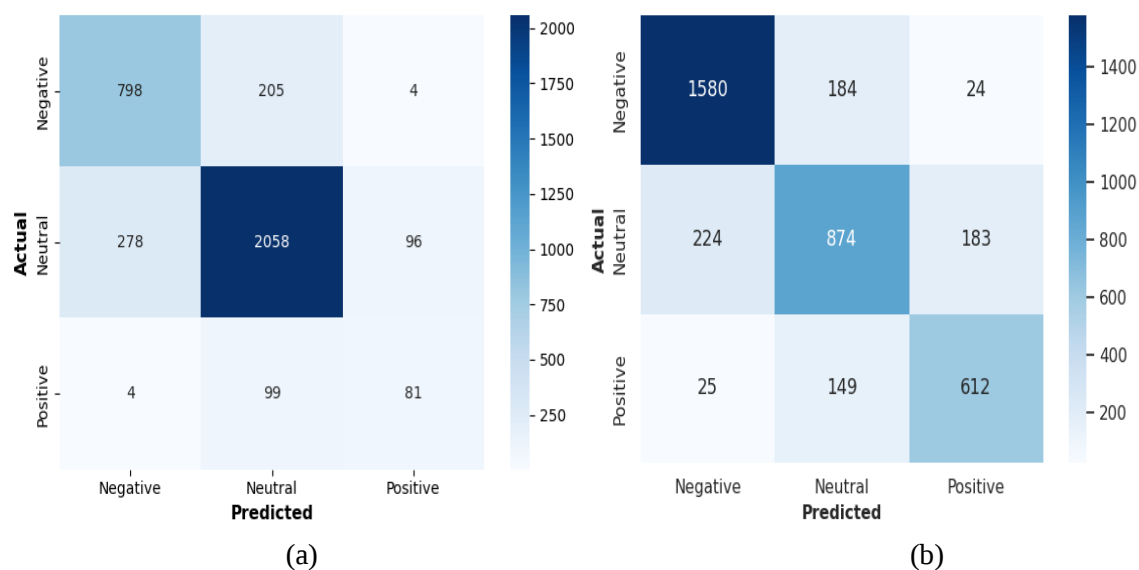


Figure 4. Confusion Matrix (a) the launch phase and (b) the implementation phase.

Figure 4(a) shows the results of the SVM model classification in the launch phase. Most of the data was correctly classified in all three sentiment classes, with the most correct predictions in the neutral class. Misclassification is still dominated by the exchange of predictions between negative and neutral classes, while errors between negative and positive classes are relatively few. These results show that the model is able to distinguish between the three classes of sentiment, although there are still characteristic similarities between the negative and neutral classes.

Figure 4(b) shows that in the implementation phase the model also managed to correctly classify most of the data in all three classes of sentiment. The pattern of misclassification remains dominated by the exchange of predictions between negative and neutral classes, followed by the exchange between neutral and positive classes, while the error between negative and positive classes is relatively lower. Overall, the model maintains good classification performance on the implementation data.

The results of this study are in line with the research of Purnomo and Septiana, which showed that SVM provides better classification performance than the comparative model in text-based sentiment analysis [32], [33]. In this study, SVM also showed the best performance in the launch and implementation phases based on all evaluation metrics, indicating its ability to classify sentiment consistently in both research phases. Meanwhile, the results of this study are different from those of Afandy and Parjito who reported SVM accuracy of 94.12%. These differences are influenced by data characteristics, social media platforms, class distribution, and

different research contexts [34]. This study uses TikTok comments that tend to be more diverse in informal language use and writing variations than Twitter data.

Overall, the results show that SVM performance is a model that is superior in both research phases from other comparative models. These results indicate that model evaluation is not enough based on accuracy alone, but also needs to consider precision, recall, and F1-score, especially on social media data with an unbalanced distribution of classes. In addition, the two-phase approach provides an overview of the differences in sentiment distribution at the stage of launching and implementing the Indonesian Economic Acceleration Program. However, these differences only describe patterns in the data analyzed and cannot be interpreted as a direct impact of policy implementation because it is still possible to be influenced by other factors, such as the timing of data collection, the context of the upload, and the characteristics of the audience.

4. Conclusions

Overall, the results of the study show that there is a difference in sentiment distribution in the two phases, namely neutral sentiment dominates in the launch phase, while in the implementation phase the sentiment distribution is dominated by negative sentiment. In terms of classification performance, SVM provides superior performance compared to other comparison models, so it is effectively used for multi-class sentiment classification on TikTok comment data represented using TF-IDF. However, the difference in sentiment distribution only describes the pattern in the analyzed data and cannot be interpreted as a direct impact of policy implementation. This study has several limitations, namely the use of data that only comes from TikTok, the difference in upload sources in the two phases, and sentiment labeling that still uses a *lexicon-based* approach with manual validation by one annotator. Therefore, further research is recommended to utilize data from various social media platforms, involve more than one annotator, and explore the representation of features and models based on *deep learning* to improve the performance of sentiment classification.

References

- [1] Y. Theocharis, S. Boulianne, K. Koc-Michalska, and B. Bimber, "Platform affordances and political participation: how social media reshape political engagement," *West European Politics*, vol. 46, no. 4, pp. 788–811, 2023, doi: 10.1080/01402382.2022.2087410.
- [2] Datareportal, "Digital 2024 Global Overview Report." Accessed: Jan. 31, 2026. [Online]. Available: <https://datareportal.com/reports/digital-2024-global-overview-report>
- [3] We Are Social and Meltwater, "Digital 2024: Indonesia." Accessed: Jan. 31, 2026. [Online]. Available: <https://datareportal.com/reports/digital-2024-indonesia>
- [4] S. Boulianne, "Twenty years of digital media effects on civic and political participation," *Communication research*, vol. 47, no. 7, pp. 947–966, 2020, doi: 10.1177/0093650218808186.
- [5] J. Ahn, H. K. Kim, L. A. Kahlor, L. Atkinson, and G.-Y. Noh, "The impact of emotion and government trust on individuals' risk information seeking and avoidance during the COVID-19 pandemic: A cross-country comparison," *Journal of Health Communication*, vol. 26, no. 10, pp. 728–741, 2021, doi: 10.1080/10810730.2021.1999348.
- [6] S. P. Bhanot and E. Linos, "Behavioral public administration: Past, present, and future," *Public Administration Review*, vol. 80, no. 1, pp. 168–171, 2020, doi: 10.1111/puar.13129.
- [7] M. Ylinen, "Incorporating agile practices in public sector IT management: A nudge toward adaptive governance," *Information Polity*, vol. 26, no. 3, pp. 251–271, 2021, doi: 10.3233/IP-200269.
- [8] kementerian Keuangan Republik Indonesia, "Pemerintah Luncurkan Program Paket Ekonomi 2025, Perkuat Penyerapan Tenaga Kerja dan Pertumbuhan Perekonomian," Kementerian Keuangan Republik Indonesia. Accessed: Oct. 06, 2025. [Online]. Available: <https://www.kemenkeu.go.id/informasi-publik/publikasi/berita-utama/Program-Paket-Ekonomi-2025>

- [9] M. Mazzucato and R. Kattel, "COVID-19 and public-sector capacity," *Oxford review of economic policy*, vol. 36, no. Supplement\1, pp. S256--S269, 2020, doi: 10.1093/oxrep/graa031.
- [10] D. Caldara and M. Iacoviello, "Measuring geopolitical risk," *American economic review*, vol. 112, no. 4, pp. 1194--1225, 2022, doi: 10.1257/AER.20191823.
- [11] M. S. Aiyar *et al.*, *Geo-economic fragmentation and the future of multilateralism*. International Monetary Fund, 2023.
- [12] M. I. Bimasena, I Kadek Yogi Prayoga, I Gusti Agung Putu Mahendra, and Purnama Sidik, "Benchmarking IndoBERT and Multilingual BERT for Indonesian Financial News Sentiment Classification," *JURNAL TEKNOLOGI DAN OPEN SOURCE*, vol. 9, no. 1, pp. 164--176, Jun. 2026, doi: 10.36378/jtos.v9i1.5583.
- [13] R. Baldwin and R. Freeman, "Risks and global supply chains: What we know and what we need to know," *Annual Review of Economics*, vol. 14, no. 1, pp. 153--180, 2022, doi: 10.1146/annurev-economics-051420-113737.
- [14] S. Diepeveen and M. Pinet, "User Perspectives on Digital Literacy as a Response to Misinformation," *Development Policy Review*, vol. 40, no. S2, p. e12671, 2022, doi: 10.1111/dpr.12671.
- [15] T. G. L. A. der Meer and Y. Jin, "Seeking formula for misinformation treatment in public health crises: The effects of corrective information type and source," *Health communication*, vol. 35, no. 5, pp. 560--575, 2020, doi: 10.1080/10410236.2019.1573295.
- [16] F. Molica and A. M. Santos, "Mapping uncharted territory: research gaps in EU cohesion policy from a policy-making perspective," *Regional Studies, Regional Science*, vol. 12, no. 1, pp. 517--531, 2025, doi: 10.1080/21681376.2025.2514503.
- [17] Tedyyana, Agus, Fajar Ratnawati, and Rezki Kurniati. "Rancangan sistem informasi penelitian dan pengabdian masyarakat Politeknik Negeri Bengkalis menggunakan metode UML (Unified Modeling Language)." *Sistemasi: Jurnal Sistem Informasi* 8.3 (2019): 413-423..
- [18] F. A. Acheampong, H. Nunoo-Mensah, and W. Chen, "Transformer models for text-based emotion detection: a review of BERT-based approaches," *Artificial Intelligence Review*, vol. 54, no. 8, pp. 5789--5829, 2021, doi: 10.1007/s10462-021-09958-2.
- [19] Q. Lu, X. Sun, Y. Long, Z. Gao, J. Feng, and T. Sun, "Sentiment Analysis: Comprehensive Reviews, Recent Advances, and Open Challenges," *IEEE Transactions on Neural Networks and Learning Systems*, vol. 35, no. 11, pp. 15092--15112, 2024, doi: 10.1109/TNNLS.2023.3294810.
- [20] A. Gasparetto, M. Marcuzzo, A. Zangari, and A. Albarelli, "A Survey on Text Classification Algorithms: From Text to Predictions," *Information*, vol. 13, no. 2, 2022, doi: 10.3390/info13020083.
- [21] V. Piccialli and M. Sciandrone, "Nonlinear optimization and support vector machines," *Annals of Operations Research*, vol. 314, no. 1, pp. 15--47, 2022, doi: 10.1007/s10479-022-04655-x.
- [22] A. Musthafa *et al.*, "Paket Ekonomi 2025: Strategi Pemerintah Mengakselerasi Pertumbuhan," *Applied Sciences (Switzerland)*, vol. 13, no. 1, pp. 193, 116377, 2025, doi: <https://doi.org/10.1016/j.eswa.2021.116377>.
- [23] S. N. Nugraha, R. Pebrianto, A. Latif, and M. R. Firdaus, "ANALISIS SENTIMEN TWITTER TERHADAP MENTERI INDONESIA DENGAN ALGORITMA SUPPORT VECTOR MACHINE DAN NAIVE BAYES," *E-Link: Jurnal Teknik Elektro dan Informatika*, vol. 17, no. 1, pp. 1--12, 2022, doi: 10.30587/e-link.v17i1.3965.
- [24] D. F. Sebastian, H. Sulistiani, and A. R. Isnain, "Sentiment analysis of public opinion on the right of inquiry in indonesia in 2024 using the support vector machine (svm) method," *Jurnal Teknik Informatika (Jutif)*, vol. 5, no. 4, pp. 1025--1034, 2024, doi: 10.52436/1.jutif.2024.5.4.1968.
- [25] A. M. Rizqiyah and I. K. D. Nuryana, "Analisis Sentimen Masyarakat terhadap Kebijakan Iuran Tabungan Perumahan Rakyat (TAPERA) pada Platform X Menggunakan Algoritma Naive Bayes Classifier dan Support Vector Machine," *Journal of Emerging Information Systems and Business Intelligence*, vol. 5, no. 3, pp. 298--306, 2024, doi: 10.26740/jeisbi.v5i3.64074@article{sebastian2024sentiment, title={Sentiment analysis of public opinion on the right of inquiry in indonesia in 2024 using the support vector machine (svm) method}, author={Sebastian, Dicky Fernanda and Sulistiani, Heni and Isnain, Auliya Rahman}, journal={Jurnal Teknik Informatika (Jutif)}, volume={5}, number={4}, pages={1025--1034}, year={2024} }.

- [26] M. Siino, I. Tinnirello, and M. La Cascia, "Is text preprocessing still worth the time? A comparative survey on the influence of popular preprocessing methods on Transformers and traditional classifiers," *Information Systems*, vol. 121, no. March 2023, p. 102342, 2024, doi: 10.1016/j.is.2023.102342.
- [27] H. Allam, L. Makubvure, B. Gyamfi, K. N. Graham, and K. Akinwolere, "Text classification: How machine learning is revolutionizing text categorization," *Information*, vol. 16, no. 2, p. 130, 2025, doi: 10.3390/info16020130.
- [28] A. G. Ghifari, G. Y. Ananada, and K. Purwandari, "A Comparative Sentiment Analysis of Public Opinion on Indonesia's National Football Coach Using CRNN and SVM," *Procedia Computer Science*, vol. 269, pp. 1485–1493, 2025, doi: 10.1016/j.procs.2025.06.181.
- [29] C. Zhang, S. Xu, Z. Li, and S. Hu, "Understanding Concerns, Sentiments, and Disparities Among Population Groups During the COVID-19 Pandemic Via Twitter Data Mining: Large-scale Cross-sectional Study," *J Med Internet Res*, vol. 23, no. 3, p. e26482, Mar. 2021, doi: 10.2196/26482.
- [30] H. Sain, H. Kuswanto, S. W. Purnami, and S. P. Rahayu, "Fuzzy support vector machine for classification of time series data: A simulation study," *Decision Science Letters*, vol. 12, no. 3, pp. 487–498, Jun. 2023, doi: 10.5267/j.dsl.2023.5.002.
- [31] J. Opitz, "A closer look at classification evaluation metrics and a critical reflection of common evaluation practice," *transactions of the association for computational linguistics*, vol. 12, pp. 820–836, 2024, doi: 10.1162/tacl_a_00675.
- [32] D. Purnomo, W. Firgiawan, and N. Nur, "Komparasi Algoritma Random Forest, Naïve Bayes, dan SVM pada Sentimen Kebijakan PPN 12%," *Jurnal Tekno Kompak*, vol. 19, no. 2, pp. 155–167, May 2025, doi: 10.33365/jtk.v19i2.122.
- [33] I. Septiana and D. Alita, "Perbandingan Random Forest dan SVM dalam Analisis Sentimen Quick Count Pemilu 2024," *Jurnal Informatika: Jurnal Pengembangan IT*, vol. 9, no. 3, pp. 224–233, 2024, doi: 10.30591/jpit.v9i3.6640.
- [34] Y. Afandy and Parjito, "Perbandingan SVM dan Random Forest pada Analisis Sentimen Kebijakan Tabungan Perumahan Rakyat Berdasarkan Data Media Sosial X," *SemanTIK: Teknik Informasi*, vol. 11, no. 1, 2025, doi: 10.55679/semantik.v11i1.93.