

Analysis of TF-IDF and TF-RF Feature Extraction on Product Review Sentiment

Keisha Priya Harmandini^{1)*}, Kemas Muslim L²⁾

¹⁾²⁾School of Computing, Telkom University, Indonesia

¹⁾harmandinikeisha@student.telkomuniversity.ac.id, ²⁾kemasmuslim@telkomuniversity.ac.id

Submitted : Jan 9, 2024 | Accepted : Jan 25, 2024 | Published : Apr 1, 2024

Abstract: Sentiment analysis of product reviews is critical in understanding customer views and satisfaction, especially in the context of e-commerce applications. A marketplace provides channels where users can submit reviews of the products they purchase. However, due to the large number of reviews in a marketplace, analyzing them is no longer feasible to be performed manually. This research proposes a machine learning implementation to perform sentiment analysis on product reviews. In this research, the product review dataset on Shopee marketplace is used for sentiment analysis by comparing TF-IDF and TF-RF feature extraction using the SVM algorithm with stages of dataset, labeling, feature extraction and accuracy results. The importance of the comparison between TF-IDF and TF-RF feature extraction in this research is related to the need to evaluate and determine which feature extraction method is most effective in increasing the accuracy of sentiment analysis. TF-IDF and TF-RF are two methods commonly used in text analysis, and a comparison of their performance can provide deep insight into the effectiveness of each in the context of product sentiment analysis. Thus, through this comparison, this research aims to determine the best approach that can provide the highest accuracy results, so that the results can serve as a guide for further research. Based on the evaluation, the highest accuracy value is achieved at 92.87% by using TF-IDF and SVM classifiers which outperformed previous research.

Keywords: sentiment analysis, shopee, svm, tf-idf, tf-rf

INTRODUCTION

In the dynamic context of e-commerce in Indonesia, the marketplace is the main foundation for the growth of this industry. With the rapid growth of this sector, the choice of e-commerce platform for research becomes crucial. Shopee is a well-known marketplace in Indonesia. Currently Shopee is ranked second as the e-commerce platform with the highest number of monthly site visitors in the first quarter of 2021 in Indonesia (Cahyaningtyas et al., 2021).

Sentiment analysis or also known as opinion mining, is one component of text mining. This field studies people's opinions, emotions, evaluations, behavior and reactions towards entities such as related products and services (Pratama, 2022). Sentiment analysis is a method used for identifying views or opinions on a subject (such as an individual, organization, or product) in a set of data (Nasukawa & Yi, 2003). According to (Muchammad Shiddieqy Hadna et al., 2016), opinion refers to the expression of a person's attitude regarding an issue that involves conflict. Opinions can take the form of mentions on social media, articles on news sites, or personal blogs.

Based on the problems above, this research will use the SVM algorithm as a classifier algorithm to classify products review in the Shopee application. In previous research (Putri & Lhaksmana, 2023) sentiment analysis was carried out by comparing the term weighting method using SVM classification. The data used in the research came from tweets with the hashtag #permendikbud30. Based on the test results, the highest F1-Score value was obtained for TF-RF with SVM kernel rbf function of 51%. There have been several previous studies, namely sentiment analysis using SVM with an accuracy of 80.90% (Hantoro et al., 2022). Then sentiment analysis using SVM with linear kernel with TF-IDF and unigram extraction features achieved an accuracy of 84.31% (Purbaya et al., 2023).

In this research, a comparison was made of feature extraction with the SVM algorithm through the stages of dataset, labeling, feature extraction and accuracy results. The final result of this research is the accuracy of feature extraction for product reviews on Shopee. The focus of this research is to determine the performance comparison between TF-IDF and TF-RF extraction features in sentiment analysis on product reviews on Shopee, as well as determining which extraction features are more suitable for sentiment analysis on product reviews on Shopee. The

*name of corresponding author



This is anCreative Commons License This work is licensed under a Creative Commons Attribution-NonCommercial 4.0 International License.

selection of feature extraction methods is a key element in this research, with the aim of determining a method that is more suitable and effective in increasing the accuracy of sentiment analysis.

TF-IDF is considered because this method has proven effective in extracting features that reflect the importance of a word in a document to the entire document collection. TF-IDF considers the frequency of word occurrences and simultaneously assigns weights based on how common or rare the word appears in all documents. This approach can provide a better representation of word significance in sentiment analysis. Meanwhile, TF-RF is considered to take advantage of the advantages of Random Forest in modeling complex and non-linear relationships between features. By utilizing an ensemble of decision trees, TF-RF is able to handle the complex characteristics of text data in product reviews, which other methods may not be able to accommodate well. Therefore, through a comparison between TF-IDF and TF-RF, this research aims to provide insight into the performance of each method in the context of product sentiment analysis. It is hoped that the research results will provide a deeper understanding of which feature extraction method is more appropriate and effective for use in analyzing sentiment in product reviews on the e-commerce platform. The limitation of the problem in this research is that the dataset used is product reviews on Shopee, totaling 11,318 data in English and divided into two classes of positive and negative sentiment.

LITERATURE REVIEW

Table 1. Performance comparison

Publication Year and Reference	Method	Dataset	Result
(Hantoro et al., 2022)	SVM Classification	Shopee review on Google Play	Accuracy value 80.90%
(Rangga et al., 2019)	K-Nearest Neighbor and SVM	Tweets containing "Donald Trump" or relating to Donald Trump	SVM is superior with a value of 89.70% without K-Fold Cross Validation and 88.76% with K-Fold Cross Validation
(Pang et al., 2002)	Naïve Bayes, Maximum Entropy, and SVM	Movie review	Highest accuracy value using SVM achieved 82.9%
(Purbaya et al., 2023)	SVM, TF-IDF, and N-gram	Keyword "keripik" on Shopee search	SVM using the TF-IDF method and N-Gram model with performance values of 88.4% accuracy, 87.3% precision, 88.4% recall, and 86.9% F1-Score
(Norindah Sari et al., 2023)	Supervised Delta TF-IDF and Unsupervised TF-IDF	Comments about COVID-19 symptoms on Twitter	The accuracy value with supervised Delta TF-IDF was 88.5% and with unsupervised TF-IDF the accuracy value was 87.9%
(Kaburuan et al., 2022)	Naïve Bayes	Review of one of the women's home clothing products or home dresses sold on Shopee	The accuracy obtained reached 90.03%
This research	SVM, TF-IDF, and TF-RF	Product review on Shopee	The accuracy obtained reached 92.87%

In gaining an in-depth understanding of this topic, the literature review becomes the foundation for exploring the main concepts that have been investigated and developed by previous researchers. Text Mining is a powerful technique in Natural Language Processing (NLP) that plays a crucial role in sentiment analysis. In the context of

Shopee Indonesia, sentiment analysis using Support Vector Machines (SVM) has been implemented. The dataset comprises 990 training samples and 110 test samples, and the study has reported an impressive accuracy rate of 80.90% (Hantoro et al., 2022). The SVM algorithm is a classification algorithm that is superior to the K-Nearest Neighbor (KNN) method with an accuracy level of 89.70% (Rangga et al., 2019). This comparison suggests that SVM, with its ability to efficiently handle high-dimensional data and capture complex decision boundaries, is a robust choice for sentiment analysis tasks. The reported accuracy level of 89.70% indicates the effectiveness of the SVM algorithm in classifying sentiments, showcasing its potential for applications in various domains, including e-commerce platforms like Shopee Indonesia. Other studies have performed sentiment classification in film reviews using various machine learning techniques. This research uses Naïve Bayes, Maximum Entropy and SVM techniques. Several approaches are also used for feature extraction, such as unigram, unigram+bigram, unigram+POS, adjective, and unigram+position. Experimental results show that SVM has the best performance with an accuracy level of 82.9% (Pang et al., 2002).

With the SVM algorithm using TF-IDF and N-gram models, the experimental results show that the best performance for SVM classification with Linear kernels can be achieved through the application of TF-IDF and unigram feature extraction. Using the lexicon approach for sentiment classification produces an accuracy of 84.31% for total positive reviews. In the unigram model, the acquisition and precision of the linear kernel reached 88.4% and 87.3%, while the recall value reached 88.4%. The F1-Score assessment matrix for Unigram reached 86.9%, while for Bigram and Trigram it was 78.5% and 77.4% respectively. Overall, the unigram model that uses linear kernels in the SVM algorithm shows significant potential for application in the development of various system focuses (Purbaya et al., 2023).

Comparison of the two types of weighting using the Random Forest algorithm. The results of classification performance with supervised weighting of Delta TF-IDF produce better accuracy of 88.5%, while with unsupervised weighting it produces accuracy of 87.9% (Norindah Sari et al., 2023).

In addition to the SVM and KNN algorithms, sentiment analysis research has explored the effectiveness of Naïve Bayes classification on product reviews within the Shopee marketplace. Sentiment analysis research has also been carried out on product reviews on Shopee using Naïve Bayes classification. Data was collected using the web crawling method from reviews written by users who purchased one of the women's home clothing products or house dresses sold on the Shopee marketplace. The accuracy obtained reached 90.03% with a total dataset of 2907 data (Kaburuan et al., 2022).

The literature underscores the importance of sentiment analysis in deciphering user opinions on e-commerce platforms, notably in Shopee Indonesia. Utilizing Support Vector Machines (SVM) for sentiment analysis in this context has consistently shown superior accuracy compared to K-Nearest Neighbor (KNN). This highlights SVM's effectiveness in capturing nuanced sentiments. The literature review can be seen in table 1.

METHOD

In this research, a system will be created that has the ability to perform sentiment classification on text containing product reviews and compare the extraction features of TF-IDF and TF-RF which can be seen in Figure 1.

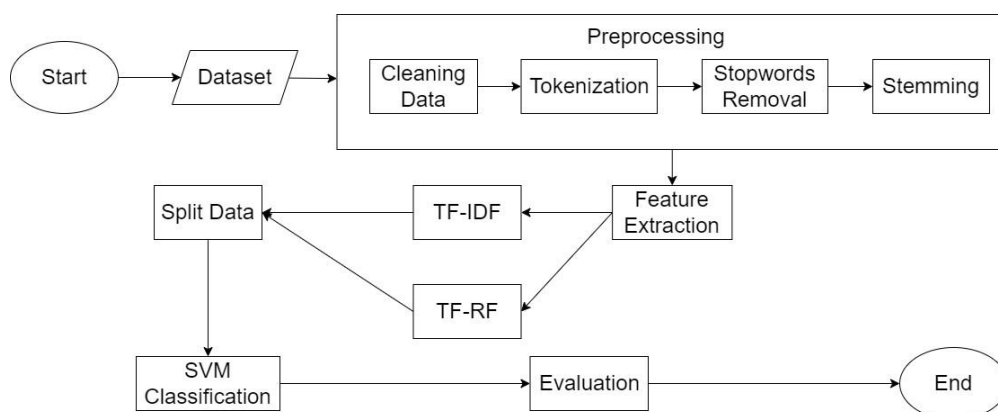


Fig. 1 System overview flow chart.

Dataset

The dataset used in this research is a dataset obtained from Kaggle containing product reviews on Shopee. Shopee is a popular e-commerce platform and is widely used by the public. The data used amounted to 11,318 data in English and has not been labeled so the next stage is data labeling. Data labeling uses (1) for positive sentiment and (-1) for negative sentiment.

Preprocessing

Preprocessing is a crucial step in data preparation for classification, involving cleaning to remove unnecessary elements like punctuation, emoticons, and symbols. This process also includes case folding, converting capital letters to lowercase, and removing numbers, ensuring the data only contains alphabetical characters. The objective is to maintain a clean dataset, devoid of irrelevant components. The data cleaning process can be seen in table 2.

Table 2. Example of data cleaning process

Before	After
Fast delivery and the product is as like the photo displayed.	fast delivery and the product is as like the photo displayed
Received item in order. Item is as mentioned and will recommend others to buy as it is value for money.	received item in order item is as mentioned and will recommend others to buy as it is value for money
The short was delivered faster then expected.	the short was delivered faster then expected

The second step is tokenization, tokenization is the process of dividing sentences in a set of text documents into chunks of words or characters according to system requirements. Separating sentences and words is done based on spaces in sentences or paragraphs. At the tokenization stage, special characters such as punctuation are also removed and all words are changed to lowercase. These pieces are referred to as tokens (Manning et al., 2008). The tokenization process can be seen in table 3.

The third step is stopwords removal, which is a procedure aimed at eliminating frequently occurring words that lack significant meaning. Some verbs, adjectives, and other adverbs are included in the stopwords list. The main purpose of stopword removal is to eliminate words that have no meaning, thereby increasing processing speed and performance. Stopwords include determiners, conjunctions, prepositions, and similar words (Riany et al., 2016). The stopwords removal process can be seen in table 4.

Table 3. Example of tokenization process

Before	After
fast delivery and the product is as like the photo displayed	['fast', 'delivery', 'and', 'the', 'product', 'is', 'as', 'like', 'the', 'photo', 'displayed']
received item in order item is as mentioned and will recommend others to buy as it is value for money	['received', 'item', 'in', 'order', 'item', 'is', 'as', 'mentioned', 'and', 'will', 'recommend', 'others', 'to', 'buy', 'as', 'it', 'is', 'value', 'for', 'money']
the short was delivered faster then expected	['the', 'short', 'was', 'delivered', 'faster', 'then', 'expected']

Table 4. Example of stopwords removal process

Before	After
['fast', 'delivery', 'and', 'the', 'product', 'is', 'as', 'like', 'the', 'photo', 'displayed']	['fast', 'delivery', 'product', 'like', 'photo', 'displayed']
['received', 'item', 'in', 'order', 'item', 'is', 'as', 'mentioned', 'and', 'will', 'recommend', 'others', 'to', 'buy', 'as', 'it', 'is', 'value', 'for', 'money']	['received', 'item', 'order', 'mentioned', 'will', 'recommend', 'others', 'buy', 'value', 'money']
['the', 'short', 'was', 'delivered', 'faster', 'then', 'expected']	['the', 'short', 'delivered', 'faster', 'then', 'expected']

The fourth step involves stemming, which is a procedure that links diverse morphological variations of words to their basic or common forms (Tala, 2003). It's important to note that the stemming algorithm differs between languages; for instance, English and Indonesian employ distinct algorithms due to their unique morphologies. In English texts, the process primarily involves removing suffixes. However, in Indonesian language texts, stemming is more intricate and complex, necessitating the removal of various affixes to obtain the base form of a word (Wahyudi et al., 2017). The stemming process can be seen in table 5.

Table 5. Example of stemming process

Before	After
['fast', 'delivery', 'product', 'like', 'photo', 'displayed']	['fast', 'delivery', 'product', 'like', 'photo', 'display']
['received', 'item', 'order', 'mentioned', 'will', 'recommend', 'others', 'buy', 'value', 'money']	['receive', 'item', 'order', 'mention', 'will', 'recommend', 'others', 'buy', 'value', 'money']
['the', 'short', 'delivered', 'faster', 'then', 'expected']	['the', 'short', 'delivery', 'faster', 'then', 'expect']

TF-IDF

TF-IDF is a method feature weighting is very popular and often used in text analysis. This method has a fairly high level of accuracy and recall (Ye et al., 2018). Term Frequency (TF) is process to count the number of words that appear in the dataset (Gumilang, 2018). The following equation can be used to calculate the Term Frequency (TF) weight.

$$TF(t, d) = \frac{\text{Total number of terms in the document } d}{\text{Number of occurrences of term } t \text{ in the document } d} \quad (1)$$

Inverse Document Frequency (IDF) is process to get the number of datasets that contain words the key you are looking for. This can be calculated using the following equation.

$$IDF(t, D) = \log \left(\frac{\text{Number of documents containing the term } t+1}{\text{Total number of documents in the corpus } |D|} \right) + 1 \quad (2)$$

So the TF-IDF equation is.

$$TFIDF(t, d, D) = TF(t, d) \times IDF(t, D) \quad (3)$$

TF-RF

The TF-RF method is a combination of Term Frequency (TF) and Relevance Frequency (RF) with the aim of improving weighting performance words compared to other methods. This method considers the relevance of a document by looking at the frequency of appearance of terms in the category related (Wu & Gu, 2014).

$$TFRF = TF(t, d) \times \log \left(2 \frac{a}{\max(1, C)} \right) \quad (4)$$

TF(t, d) is the frequency of appearance of term t in document d, a is the number of documents not in a certain category that contain the term, and C is the number of documents in a certain category that contain the term.

Split Data

The data split process is carried out by dividing the dataset into two parts, namely training data and test data. Training data functions as material that will be studied by the system, which will later be used to create a learning model based on a predetermined classification method. Meanwhile, the test data will utilize the learning model to classify text based on the data provided. In this research there were 11,318 data with 10,349 positive data and 969 negative data, Due to unbalanced data, the SMOTE algorithm is used to balance positive data with negative data. The data division will use ratio of 80:20.

SVM Classification

The SVM algorithm proves to be highly efficient for text classification tasks. Geometrically, a binary SVM classifier can be visualized as a hyperplane in the feature space, effectively separating points representing positive instances from those representing negative instances. During training, the classifier selects a unique hyperplane that maximally separates known positive instances from known negative instances, creating a margin, which is the distance from the hyperplane to the nearest points in both sets (Feldman & Sanger, 2007).

SVM hyperplanes are exclusively determined by a small subset of training instances known as support vectors, rendering the remaining training data irrelevant for the trained classifier. This characteristic sets SVM apart from other categorization algorithms. An advantageous feature of the SVM classifier is its theoretically justified approach to the overfitting problem, allowing it to perform well regardless of the feature space's dimensionality. Additionally, SVM requires no parameter adjustment, as there exists a theoretically motivated "default" choice of

*name of corresponding author



This is anCreative Commons License This work is licensed under a Creative Commons Attribution-NonCommercial 4.0 International License.

parameters that has demonstrated optimal performance through experimental validation (Feldman & Sanger, 2007). SVM can expressed by the equation.

$$(w \cdot x_i) + b = 0 \quad (5)$$

in data x_i which is included in class -1 can be formulated as in equation

$$(w \cdot x_i + b) \leq 1, y_i = -1 \quad (6)$$

while data x_i which is included in class +1 can be formulated as in equation

$$(w \cdot x_i + b) \geq 1, y_i = 1 \quad (7)$$

where w is weight, x is data or input and b is bias.

Evaluation

In this research, testing was carried out to evaluate the results of the algorithm used. One of the evaluation methods used is the confusion matrix. This method is very helpful in analyzing the quality of the classifier. After the confusion matrix is carried out, values such as accuracy, recall, precision, and f1-score can be calculated and displayed in percentage form. The following is an example of a confusion matrix table in table 6.

Table 6. Confusion Matrix

		Predicted Class	
		1	-1
Actual Class	1	True Positive	False Negative
	-1	False Positive	True Negative

Accuracy is a metric that gauges the extent to which a models's predictions align with the actual outcomes of the modeled reality (Sammut & Webb, 2011). The equation for calculating accuracy is written as follows.

$$Accuracy = \frac{TP+TN}{TP+FN+FP+TN} \quad (8)$$

Recall represents the ratio of positive documents accurately predicted as positive (Sammut & Webb, 2011) divided by the sum of True Positives and False Negatives. The equation for calculating Recall is written using the following formula.

$$Recall = \frac{TP}{TP+FN} \quad (9)$$

Precision is the proportion of documents predicted positive that are really positive (Sammut & Webb, 2011). The equation for calculating Precision is written as follows.

$$Precision = \frac{TP}{TP+FP} \quad (10)$$

F1-Score is a calculation to measure performance by combining precision and recall values (Sammut & Webb, 2011). The equation for calculating the F-1 Score is written as follows.

$$F1 - Score = 2 \frac{Precision \times Recall}{Precision + Recall} \quad (11)$$

RESULT

These research steps consist of dataset collection, data preprocessing, feature extraction, SVM classification and evaluation. In detail, the first stage carried out was the collection of datasets where the dataset was obtained from Kaggle which contained product reviews on Shopee as many as 11,318 English data. Then there is a preprocessing stage where at this stage data cleaning, tokenization, stopwords removal and stemming are carried out, then the preprocessed data is implemented extraction features with TF-IDF and TF-RF. The sentiment analysis results revealed an imbalanced distribution between negative and positive sentiments, with 969 instances classified as negative and 10,349 instances as positive. Recognizing the importance of balanced datasets for robust model training, the researchers opted for oversampling using the SMOTE algorithm. In this research, the evaluation was conducted following the modeling process, which involved implementing the SVM algorithm and employing both TF-IDF and TF-RF feature extraction techniques. The evaluation considered different ratios of training data to test

*name of corresponding author



This is anCreative Commons License This work is licensed under a Creative Commons Attribution-NonCommercial 4.0 International License.

data (60:40, 70:30, 80:20, and 90:10) in order to determine the most optimal accuracy outcomes. The evaluation metrics, including accuracy, precision, recall, and F1-score, are presented in Table 7 and Table 8.

Figure 2 visually represents the impact of the SMOTE algorithm on the dataset. This graphical representation serves to illustrate how the oversampling technique has effectively balanced the distribution of sentiments, thereby enhancing the model's ability to learn from both negative and positive instances. This preprocessing step contributes to the overall reliability and fairness of the sentiment analysis model, ensuring that it can effectively capture patterns from both sentiment categories during training and subsequent evaluations.

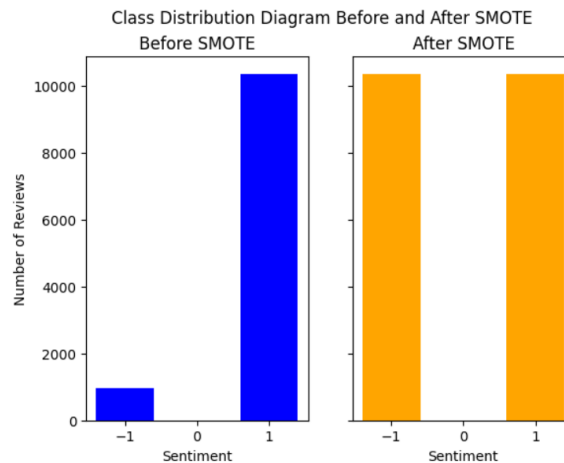


Fig. 2 Class distribution diagram before and after SMOTE.

Table 7. Performance evaluation result with TF-IDF

Ratio	Accuracy	Recall	Precision	F1-Score
60:40	92.82%	92.84%	93.02%	92.82%
70:30	92.67%	92.68%	92.90%	92.66%
80:20	92.87%	92.85%	93.16%	92.86%
90:10	92.80%	92.85%	93.10%	92.79%

Table 8. Performance evaluation result with TF-RF

Ratio	Accuracy	Recall	Precision	F1-Score
60:40	91.15%	91.17%	91.28%	91.15%
70:30	91.19%	91.20%	91.33%	91.19%
80:20	91.32%	91.31%	91.50%	91.32%
90:10	91.54%	91.59%	91.76%	91.54%

Based on the two tables above, it can be concluded that split data using a ratio of 80:20 with the TF-IDF extraction feature obtained the most optimal results with an accuracy of 92.87%.

DISCUSSIONS

Based on performance evaluation using the SVM algorithm with TF-IDF and TF-RF extraction features, split the data with four comparisons (60:40, 70:30, 80:20 and 90:10) to get the most optimal results for accuracy, recall, precision and f1-score from the results of split data testing using a ratio of 80:20 to get the most optimal results on the TF-IDF extraction feature with an accuracy of 92.87%, recall 92.85%, precision 93.16% and f1-Score 92.86% with a confusion matrix that can be seen in table 9.

Table 9. Confusion matrix SVM with TF-IDF

		Predicted Class	
		1	-1
Actual Class	1	1827	229
	-1	66	2018

Based on table 9, SVM classification and TF-IDF extraction features produce quite accurate results for product review sentiment on Shopee compared to using TF-RF extraction features. The resulting confusion matrix values are as follows.

1. True Positive (TP) totaled 1827, where the model classified 1827 data that were truly positive and also predicted as positive by the model.
2. True Negative (TN) totaled 2018, where the model classified 2018 data that were truly negative and also predicted as negative by the model.
3. False Positive (FP) totaled 66, where the model classified 66 data that were actually negative but were incorrectly predicted as positive by the model.
4. False Negative (FN) totaled 229, where the model classified 229 data that were actually positive but were incorrectly predicted as negative by the model.

CONCLUSION

In research on sentiment analysis of product reviews on Shopee, the use of the SVM algorithm with the TF-IDF extraction feature produces higher accuracy than the use of the TF-RF extraction feature, with a data split ratio of 80:20. The use of TF-IDF gives higher weight to unique and informative words. Although not very common, TF-RF may try to integrate the Random Forest method with term frequency. However, in this case the method does not provide as good performance as TF-IDF. With the TF-IDF extraction feature, an accuracy of 92.87%, recall of 92.85%, precision of 93.16% and f1-Score of 92.86% whose accuracy outperforms previous research and shows that the model is able to properly classify product review sentiment. Although high accuracy has been achieved, it is necessary to continue to develop the model by trying other classification algorithms or other extraction features. Developing research by analyzing sentiment by product category or by exploring more complex reviews can provide additional insights.

REFERENCES

- Cahyaningtyas, C., Nataliani, Y., & Widiarsari, I. R. (2021). Analisis sentimen pada rating aplikasi Shopee menggunakan metode Decision Tree berbasis SMOTE. *AITI: Jurnal Teknologi Informasi*, 18(Agustus), 173–184.
- Feldman, R., & Sanger, J. (2007). *The text mining handbook: advanced approaches in analyzing unstructured data*.
- Gumilang, Z. A. N. (2018). *IMPLEMENTASI NAÏVE BAYES CLASSIFIER DAN ASOSIASI UNTUK ANALISIS SENTIMEN DATA ULASAN APLIKASI E-COMMERCE SHOPEE PADA SITUS GOOGLE PLAY*.
- Hantoro, K., Handayani, D., & Setiawati, S. (2022). A Implementation of Text Mining In Sentiment Analysis of Shopee Indonesia Using SVM. *Bulletin of Information Technology (BIT)*, 3(2), 115–120. <https://doi.org/10.47065/bit.v3i1.282>
- Kaburuan, E. R., Sari, Y. S., & Agustina, I. (2022). Sentiment Analysis on Product Reviews from Shopee Marketplace using the Naïve Bayes Classifier. *Lontar Komputer : Jurnal Ilmiah Teknologi Informasi*, 13(3), 150. <https://doi.org/10.24843/lkjiti.2022.v13.i03.p02>
- Manning, C., Raghavan, P., & Schütze, H. (2008). *Introduction to Information Retrieval*.
- Muchammad Shiddieqy Hadna, N., Insap Santosa, P., & Wahyu Winarno, W. (2016). STUDI LITERATUR TENTANG PERBANDINGAN METODE UNTUK PROSES ANALISIS SENTIMEN DI TWITTER. In *Seminar Nasional Teknologi Informasi dan Komunikasi*.
- Nasukawa, T., & Yi, J. (2003). Sentiment analysis: Capturing favorability using natural language processing. *Proceedings of the 2nd International Conference on Knowledge Capture, K-CAP 2003*, 70–77. <https://doi.org/10.1145/945645.945658>
- Norindah Sari, S., Reza Faisal, M., Kartini, D., Budiman, I., & Hamonangan Saragih, T. (2023). *Perbandingan Ekstraksi Fitur dengan Pembobotan Supervised dan Unsupervised pada Algoritma Random Forest untuk Pemantauan Laporan Penderita COVID-19 di Twitter* (Vol. 11, Issue 1).
- Pang, B., Lee, L., & Vaithyanathan, S. (2002). *Thumbs up? Sentiment Classification using Machine Learning Techniques*. EMNLP. <http://www.cs.cornell.edu/people/pabo/movie-review-data/>.
- Pratama, M. Y. (2022). *ANALISA SENTIMEN TERHADAP PENGGUNAAN APLIKASI SHOPEE FOOD PADA TWITTER MENGGUNAKAN METODE NAÏVE BAYES DAN SUPPORT VECTOR MACHINE (SVM)*.
- Purbaya, M. E., Rakhmadani, D. P., Maliana Puspa Arum, & Luthfi Zian Nasifah. (2023). Implementation of n-gram Methodology to Analyze Sentiment Reviews for Indonesian Chips Purchases in Shopee E-Marketplace. *Jurnal RESTI (Rekayasa Sistem Dan Teknologi Informasi)*, 7(3), 609–617. <https://doi.org/10.29207/resti.v7i3.4726>
- Putri, M. R., & Lhaksana, K. M. (2023). *Analisis Sentimen Terhadap Tweet Pelecehan Seksual Dengan Perbandingan Metode Term Weighting Menggunakan Klasifikasi SVM Terhadap Tagar Permendikbud30*.

- Rangga, M., Nasution, A., & Hayaty, M. (2019). Perbandingan Akurasi dan Waktu Proses Algoritma K-NN dan SVM dalam Analisis Sentimen Twitter. *JURNAL INFORMATIKA*, 6(2), 212–218. <http://ejournal.bsi.ac.id/ejurnal/index.php/ji>
- Riany, J., Fajar, M., & Lukman, M. P. (2016). *Penerapan Deep Sentiment Analysis pada Angket Penilaian Terbuka Menggunakan K-Nearest Neighbor*.
- Sammut, C., & Webb, G. I. (2011). *Encyclopedia of Machine Learning*.
- Tala, F. Z. (2003). *A Study of Stemming Effects on Information Retrieval in Bahasa Indonesia*.
- Wahyudi, D., Susyanto, T., Nugroho, D., Studi Teknik Informatika, P., Sinar Nusantara Surakarta, S., & Studi Sistem Informasi, P. (2017). *IMPLEMENTASI DAN ANALISIS ALGORITMA STEMMING NAZIEF & ADRIANI DAN PORTER PADA DOKUMEN BERBAHASA INDONESIA*.
- Wu, H., & Gu, X. (2014). *Reducing Over-Weighting in Supervised Term Weighting for Sentiment Analysis*.
- Ye, J., Jing, X., & Li, J. (2018). Sentiment Analysis Using Modified LDA. *Lecture Notes in Electrical Engineering*, 473, 205–212. https://doi.org/10.1007/978-981-10-7521-6_25