

Enhanced sentiment analysis and emotion detection in movie reviews using support vector machine algorithm

Aditiya Hermawan¹, Rico Yusuf¹, Benny Daniawan², Junaedi²

¹Department of Informatics Engineering, Faculty of Science and Technology, University of Buddhi Dharma, Tangerang, Indonesia

²Department of Information System, Faculty of Science and Technology, University of Buddhi Dharma, Tangerang, Indonesia

Article Info

Article history:

Received May 31, 2024

Revised Oct 9, 2024

Accepted Oct 18, 2024

Keywords:

Emotion detection

Movie review

Sentiment analysis

Support vector machine

Text mining

ABSTRACT

Films evoke diverse responses and reactions from audiences, captured through their reviews. These reviews serve as platforms for audiences to express opinions, evaluations, and emotions about films, reflecting the personal experiences and unique perceptions of the viewers. Given the vast volume of reviews and the distinctiveness of each perspective, automated analysis is essential for efficiently extracting valuable insights. This study employs the support vector machine (SVM) algorithm for classifying movie reviews into positive and negative categories. The dataset includes 50,000 IMDb movie reviews, split evenly between positive and negative sentiments. Each review is analyzed using the National Research Council Canada (NRC) emotion lexicon (NRCLex) to assign scores for emotions such as anger, disgust, fear, joy, sadness, and surprise. Subsequently, these reviews are further analyzed using term frequency-inverse document frequency (TF-IDF) for classification. The proposed algorithm achieves 90% accuracy, indicating its effectiveness in classifying sentiments in movie reviews. The study's findings confirm the potential of the SVM algorithm for broader applications in sentiment analysis and natural language processing. Additionally, integrating emotion detection enhances understanding of nuanced emotional content, providing a comprehensive approach to sentiment classification in large datasets.

This is an open access article under the [CC BY-SA](https://creativecommons.org/licenses/by-sa/4.0/) license.



Corresponding Author:

Aditiya Hermawan

Department of Informatics Engineering, University of Buddhi Dharma

Tangerang, Indonesia

Email: aditiya.hermawan@ubd.ac.id

1. INTRODUCTION

In the film industry, audiences express their sentiments about films through ratings and reviews. These reviews can be positive, neutral, or negative and are often posted on various audience review websites [1]. The increasing use of the Internet has made it easier for viewers to share their opinions and emotions about films, reaching a broad audience through online platforms and social media [2]. As a result, for their films to be in demand, companies must improve their ability to respond to audiences' sentimental reactions [3]-[5]. However, despite the abundance of these reviews, efficiently extracting valuable insights from them poses a significant challenge. Traditional methods for assessing emotional reactions, such as psychological signals detected through electrocardiograms and skin temperature analysis, are costly and time-consuming [6]. This presents a gap in efficiently and accurately analyzing the vast amount of sentiment data generated by movie reviews.

To address this challenge, a more efficient and scalable solution to identify emotions lies in text mining [7]-[9]. Text mining is a powerful technique and technology employed to systematically analyze large

volumes of textual data, such as movie reviews. It enables the automatic or semi-automatic discovery of recurring patterns, trends, and rules within the text, providing valuable insights into the sentiments and preferences of the audience [10], [11]. Text mining is a versatile approach that enables the analysis of audience opinions by extracting their sentiments, perceptions, emotions, and opinions. Basic emotional values expressed in audience reviews, using Paul Eckman's renowned research basic emotion is identified as six universal emotions shared across human cultures: anger, disgust, fear, joy, sadness, and surprise [12]. This study employs a lexicon-based method to extract these emotions from audience reviews.

Building on this approach, the National Research Council Canada (NRC) emotion lexicon (NRCLex) is an MIT-approved study on emotion recognition [13]-[15], and focuses on the association of specific words with various emotions. In this comprehensive linguistic resource, words are meticulously tagged with emotion labels corresponding to their associated basic emotions. By utilizing NRCLex, this research can effectively classify words within audience reviews into appropriate basic emotion categories. This comprehensive approach enhances the sentiment analysis and opinion mining conducted within this study, providing a deeper understanding of the emotional nuances expressed by the audience [16]-[18].

This study proposes further integration of the support vector machine (SVM) algorithm with emotion detection using NRCLex to enhance sentiment analysis. The SVM algorithm, widely recognized for its robust classification capabilities, is particularly well-suited for tasks in text mining, making it an ideal choice for classifying movie reviews into positive and negative sentiments. SVM has consistently proven its efficacy in text analysis, where its primary objective is to establish a linear boundary that effectively separates sentiments into distinct categories [19]. By employing SVM, this study aims to more effectively classify movie reviews into positive and negative sentiments. The novelty lies in combining the emotion scores from NRCLex with term frequency-inverse document frequency (TF-IDF) features to enhance the classification accuracy of SVM.

In this context, SVM has been pivotal in previous research for maximizing sentiment discernment within various textual datasets, including movie reviews. Its ability to navigate the complexities of textual data and discern intricate patterns has made it a valuable tool for precisely classifying reviews based on their sentiment expressions [20]. Prior studies have demonstrated SVM's excellence in identifying distinguishing features that differentiate positive sentiments from negative ones, contributing to accurate sentiment analysis [21].

Moreover, as a testament to its effectiveness in text mining and sentiment analysis, numerous studies and research papers have highlighted the prowess of SVM in classifying textual data [22]-[24]. These studies demonstrate SVM's strong capability to identify distinct patterns within complex text data, making it highly suitable for analyzing sentiment in various contexts. Researchers have effectively leveraged SVM as the primary classification tool, optimizing its application for tasks involving the analysis of words within reviews and assigning corresponding sentiment labels. This widespread use underscores SVM's robustness in accurately discerning positive and negative sentiments, thereby solidifying its reputation as a reliable model for text classification.

2. METHOD

The dataset used in this study is derived from a collection of 50,000 reviews sourced from internet movie database (IMDB) [25]. These reviews were meticulously curated, ensuring that no more than 30 reviews were included per movie to maintain diversity and representativeness. The construction of this dataset adheres to specific criteria, specifically to ensure a balanced distribution of positive and negative reviews. This balance is crucial as it prevents random guessing from achieving high accuracy, resulting in a more robust evaluation. In this context, a review is classified as negative if it scores less than 4 out of 10, while a review is considered positive if it has a score greater than or equal to 7 out of 10. Reviews falling in between these thresholds and classified as neutral are excluded from the dataset, thereby emphasizing the importance of analyzing strong and unequivocal sentiments. Table 1 provides a sample of the review dataset used in this study.

Table 1. Sample of review dataset

Review	Label
One of the other reviewers has mentioned that after watching just 1 Oz episode you'll be hooked. They are right, as this is exactly what happened with me.	Positive
Basically, there's a family where a little boy (Jake) thinks there's a zombie in his closet & his parents are fighting all the time. This movie is slower than a soap opera and suddenly, Jake decides to become Rambo and kill the zombieOK, first of all when you're going to make a film you must decide if its a thriller or a drama! As a drama the movie is watchable.	Negative

Figure 1 shows an overview of the system. The primary objective of this process is to transform the unstructured reviews into features, enabling them to be effectively trained by the SVM algorithm using the 'Label' as the target. Figure 2 illustrates the text preprocessing steps. Initially, the 'Reviews' column is transformed into tokens, dividing the text into smaller, more manageable units. These tokens are then uniformly converted to lowercase to ensure accurate matching with associated words in the basic emotion lexicon, a fundamental analysis component. Following this, stopwords are removed, which is crucial in enhancing token-to-lexicon matching efficiency. The data is streamlined by eliminating these unnecessary words while preserving its essential content [26]. Additionally, lemmatization reduces words to their base or root form. This step is essential for normalizing the text and ensuring that different forms of a word are treated as a single item, further improving the accuracy and efficiency of the analysis.

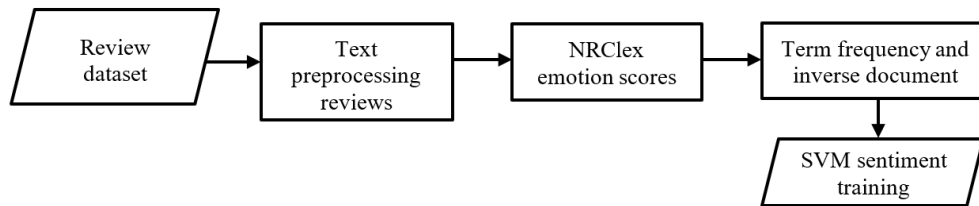


Figure 1. General view of the system

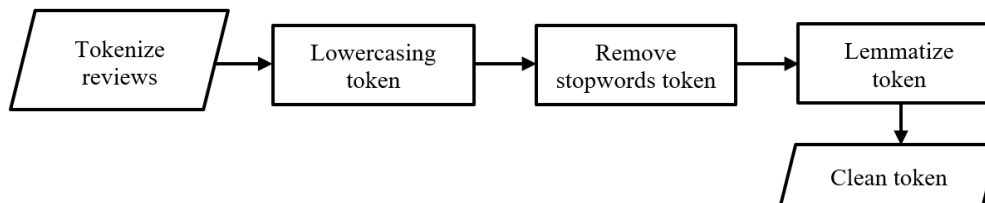


Figure 2. Text preprocessing

The primary purpose of text preprocessing is to make the textual data more readable and suitable for classification tasks [27], [28]. Tokenizing the text facilitates its analysis and classification, enabling the system to interpret and categorize the content effectively. Furthermore, eliminating stopwords reduces noise and irrelevant information, improving the data's sentiment analysis and emotion scoring efficiency. This meticulous preprocessing stage sets the foundation for the subsequent steps in the analysis, resulting in a set of clean tokens ready for scoring by NRClex.

Figure 3 shows NRClex emotion Scores. To generate emotion scores, a list of clean tokens is provided as input, and a comparison is made to determine whether these tokens match words in the six basic emotions lexicon. If a token is found to correspond with an entry in the lexicon, the NRClex 'word affect' method assigns a corresponding value. Conversely, if a token does not match any emotion in the lexicon, it receives no value.

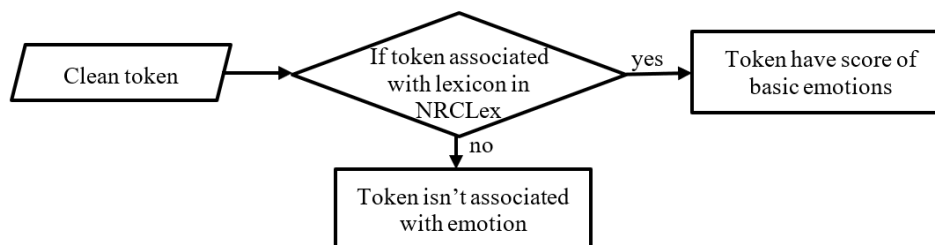


Figure 3. NRClex emotion scores

Figure 4 shows TF-IDF. Once the reviews have been assigned emotional values at this stage, the next step involves the TF-IDF process. TF-IDF is utilized to transform the reviews into a vector feature representation. This transformation allows each review to be represented as a numerical vector, with each

dimension in the vector corresponding to a unique token in the reviews [28]. TF-IDF evaluates the importance of each token within a review by considering its frequency (term frequency) and significance across the entire dataset (inverse document frequency). This process effectively captures the relative importance of each token within the context of the entire corpus of reviews.

By converting reviews into TF-IDF vectors, the data is now appropriately formatted for classification by the SVM algorithm [29]. SVM leverages these TF-IDF vectors to determine the relationship between the tokens and other tokens in the reviews, facilitating the classification of reviews based on their sentiment. This approach enables the SVM algorithm to discern patterns and trends within the TF-IDF feature space and accurately predict sentiment.

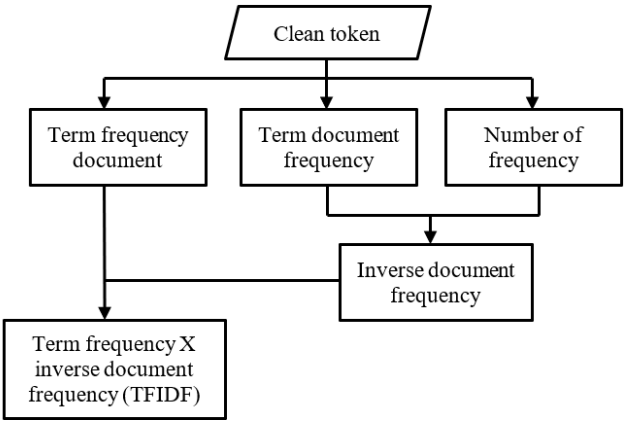


Figure 4. TF-IDF

Figure 5 shows the SVM sentiment training process. The SVM algorithm has been carefully modeled using a dataset meticulously annotated with sentiment labels and emotional values [30]. To facilitate robust training and evaluation of the SVM model, this dataset is strategically into two subsets: training and testing data. Within this partition, the testing data is a particularly intriguing subset as it intentionally excludes sentiment labels, challenging the SVM model to predict sentiment. To enable this prediction, a strategic selection of features is employed, encompassing both emotion scores and TF-IDF features. The emotion score feature encapsulates the emotional nuances within the reviews, providing an insightful layer of context. In contrast, the TF-IDF feature meticulously captures the distinctive linguistic patterns within each review, transforming them into numerical representations. Together, these features create intricate patterns within the textual data [31]. The primary objective behind incorporating these features is to empower the SVM model to effectively generalize and recognize patterns across a wide range of textual expressions [32]. This combination of emotion scores and TF-IDF features equips the SVM model to predict the sentiment of testing data that lacks predefined sentiment labels.

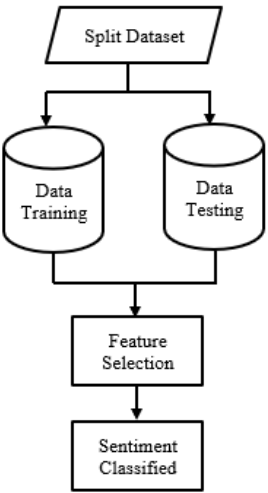


Figure 1. SVM sentiment training

This approach underscores the versatility of the SVM algorithm in sentiment analysis, as it successfully leverages multiple features to make accurate predictions and extract valuable insights from textual reviews. Combining emotion scores and TF-IDF features, the SVM model can recognize and generalize patterns across diverse textual data, capturing emotional nuances and distinctive linguistic patterns for more precise sentiment classification. As illustrated in Figure 5, the process begins with splitting the dataset into training and testing subsets, ensuring robust evaluation by training on one subset and testing on another.

During feature selection, emotion scores provide context by encapsulating the emotional content within reviews, while TF-IDF features capture the unique linguistic patterns of the text. This dual-feature approach enhances the model's ability to discern subtle differences between positive and negative sentiments. The primary goal is to enable the SVM model to effectively generalize and recognize patterns, thereby accurately predicting the sentiment of reviews in the testing data. This comprehensive strategy demonstrates the SVM's robustness and efficacy in sentiment analysis, making it a powerful tool for extracting meaningful insights from large volumes of textual data.

3. RESULTS AND DISCUSSION

3.1. Feature of emotion score and term frequency-inverse document frequency

In the context of results and discussion, this study investigates the performance of the SVM model with a particular focus on feature selection. The selected feature set encompasses two pivotal components. Firstly, the 'Basic Emotion Scores' feature is examined, quantifying the emotional content within the reviews by assessing the presence of six fundamental emotions: anger, disgust, fear, joy, sadness, and surprise. Table 2 shows the frequency of each emotion's occurrence and provides an 'Average' score, offering a normalized view of emotional intensity.

No	Score of basic emotion	Count	Average
1	Anger	163920	0.13
2	Disgust	126564	0.10
3	Fear	236738	0.20
4	Joy	243628	0.26
5	Sadness	196789	0.16
6	Surprise	155945	0.14
	Total	1123584	1.00

Complementing this, the 'TF-IDF' feature is incorporated, delving into the linguistic patterns within the textual data. By identifying words associated with the emotion lexicon, the feature selection using TF-IDF equips the model to predict sentiment with increased accuracy. It bridges the emotional context the Basic Emotion Scores represent with the linguistic nuances uncovered through TF-IDF. Consequently, the SVM model demonstrates its efficacy in precisely classifying reviews, providing comprehensive insights into the sentiment encapsulated within textual content. Subsequent sections will delve into a detailed discussion of the findings and their implications, shedding light on the effectiveness of these features in enhancing sentiment analysis and understanding emotional nuances in textual data.

3.2. Evaluation of support vector machine

To evaluate our SVM model's performance, we employed a rigorous testing approach using a review dataset of 50,000 reviews evenly split between positive and negative sentiments, ensuring balanced representation. The dataset was partitioned into training data, comprising 90% of the dataset, and testing data, accounting for the remaining 10%. To rigorously assess the model's accuracy and robustness, we opted for the 10-fold cross-validation method. This method involves creating a confusion matrix in ten iterations, each providing comprehensive insights into the model's performance.

The dataset is thoughtfully divided into ten equally balanced parts or folds in each iteration. The SVM model is trained on nine folds during training, while the remaining fold serves as the testing data. This process is repeated ten times, ensuring each dataset section becomes the testing data once. After completing all ten iterations, the results are aggregated to construct the confusion matrix. This serves as a cornerstone for evaluating the SVM model's performance and shedding light on its ability to classify reviews correctly [33].

Key performance evaluation metrics are meticulously computed based on the true positive (TP), true negative (TN), false positive (FP), and false negative (FN) values obtained from all iterations. Accuracy

provides an overall percentage of correct predictions, gauging the model's effectiveness, and is calculated using the formula: $(TP+TN)/(TP+TN+FP+FN)$. Precision measures the percentage of correct positive predictions out of all predicted positive results, serving as a vital indicator of the model's reliability, calculated as $TP/(TP+FP)$. Recall calculates the percentage of correct positive predictions compared to all positive data, offering insights into the model's ability to capture true positives using the $TP/(TP+FN)$ formula. The F1-score provides a balanced measure considering precision and recall, calculated as $2 \times (\text{Precision} \times \text{Recall}) / (\text{Precision} + \text{Recall})$. These performance metrics collectively offer a comprehensive assessment of the SVM model's capabilities and effectiveness in correctly classifying reviews, providing valuable insights into its overall performance.

In Table 3, performance metrics, including accuracy, precision, recall, and F1-score, comprehensively assess the SVM model's performance in classifying reviews. Accuracy measures overall correctness; precision assesses the accuracy of positive predictions; recall evaluates the model's ability to capture TP, while the F1-score combines precision and recall for a balanced assessment. Throughout all iterations, the SVM model consistently achieves an accuracy score surpassing 90%, indicating its proficiency in correctly classifying reviews, particularly distinguishing between positive and negative sentiments. These data points highlight the SVM model's robustness in sentiment classification, accurately capturing sentiments within movie reviews. Using features such as basic emotion scores and TF-IDF, in conjunction with the 10-fold cross-validation method, our SVM model excels in sentiment analysis, showcasing its potential for broader applications in natural language processing and text mining.

Table 3. SVM evaluation of 10-fold cross validation

Iteration	TP	FP	FN	TN	Accuracy
1	2265	276	203	2256	0.90
2	2357	254	203	2186	0.91
3	2157	273	206	2364	0.90
4	2297	282	215	2206	0.90
5	2333	277	238	2152	0.90
6	2326	244	215	2215	0.91
7	2271	283	204	2242	0.90
8	2296	261	207	2236	0.91
9	2318	237	219	2226	0.91
10	2241	235	229	2295	0.91

This research's extensive evaluation demonstrates that the SVM model consistently maintains an accuracy, precision, recall, and F1-score of approximately 0.90 or 90%. These performance metrics underscore the SVM model's suitability for sentiment classification tasks on review datasets. Leveraging these features and methods, our SVM model excels not only in sentiment analysis but also exhibits promising potential for diverse applications in natural language processing and text mining, as seen in the provided data in Table 4.

Table 4. Summary SVM evaluation of 10-fold cross validation

Actual	Predicted		Precision	Recall	F1-score	Support
	Positive	Negative				
Positive	22861	2622	0.90	0.91	0.91	25483
Negative	2139	22378	0.91	0.90	0.90	24517
Accuracy				0.90		
Macro avg				0.90		50000
Weighted avg				0.90		

3.3. Discussion

The findings of this study underscore the robustness and efficiency of the SVM algorithm in performing sentiment classification on movie reviews. By transforming the reviews into TF-IDF vectors, the data was made suitable for classification, enabling the SVM algorithm to achieve a high accuracy of 90%. This result highlights the algorithm's capability to discern subtle patterns and trends within the textual data, leading to precise sentiment predictions. The integration of emotion detection through the NRClex further enriched the sentiment analysis by providing a deeper understanding of the emotional content within the reviews. Combining these emotion scores with TF-IDF features created a comprehensive feature set that improved the algorithm's ability to generalize across different reviews.

Previous research conducted by Shrivastava *et al.* [34] proposed a sequence-based CNN method with word embeddings to detect emotions. The study results showed that the proposed CNN model outperformed LSTM and random forest classifiers by achieving an accuracy of 80.41%. However, in that study, the primary determination of emotions was done using Labelbox, and annotations that could not be scored automatically were done manually by three annotators. This approach allows for variations in the interpretation of emotions by each annotator, which in turn can affect the accuracy level of the model. Unlike our approach, we utilize a SVM combined with an NRCLEX to detect and classify emotions based on emotion scores and TF-IDF features. This combination allows the system to detect emotion patterns more consistently and reduces variability caused by subjective judgments from manual annotators.

The successful implementation of this approach suggests potential applications beyond movie reviews, such as customer feedback analysis, social media monitoring, and public opinion analysis. However, there are areas for future improvement and exploration. Incorporating more advanced machine learning techniques, such as deep learning models, could enhance the accuracy and depth of sentiment analysis. Additionally, expanding the dataset to include reviews from various genres, languages, and cultural contexts could provide a more comprehensive understanding of how sentiments and emotions vary across different demographics. The high accuracy and nuanced understanding of emotional content achieved in this research provide a strong foundation for further advancements in the field of natural language processing and sentiment analysis.

4. CONCLUSION

In conclusion, this research has successfully developed a website for sentiment analysis, incorporating opinion mining, text mining, and data mining techniques. The website efficiently extracts basic emotions, such as anger, disgust, fear, joy, sadness, and surprise, from user reviews and accurately determines sentiment polarity, distinguishing between positive and negative sentiments. Implementing the SVM algorithm for sentiment classification demonstrated exceptional performance, achieving an accuracy rate of 90% through rigorous testing using the 10 fold-cross validation method. The evaluation of the SVM algorithm revealed outstanding metrics, including an F1-score of 90%, a recall rate of 90%, and a precision rate of 90%. Notably, the study identified joy as the predominant emotion in positive sentiment reviews, while fear emerged as the dominant emotion in negative sentiment reviews. These findings underscore the research's effectiveness in emotion extraction and sentiment classification within user reviews. The SVM algorithm's high accuracy highlights its suitability for sentiment analysis tasks.

In summary, this research contributes valuable insights and a functional tool for analyzing user sentiments in online reviews, providing a foundation for further exploration in natural language processing and sentiment analysis. For future research, there are several avenues to explore and enhance the study further. One potential direction is to incorporate a more extensive and diverse dataset of film reviews, encompassing a more comprehensive range of genres, languages, and cultural contexts. This expanded dataset could provide a richer and more comprehensive understanding of how emotions and sentiments vary across film genres and audience demographics. Additionally, integrating advanced machine learning techniques, such as deep learning models and neural networks, could improve sentiment analysis accuracy and better capture the nuanced emotions expressed in movie reviews. Integrating other data sources, such as social media discussions and user-generated content related to films, could provide additional context and insights into audience sentiments and preferences. Overall, these future research directions aim to enhance the study's methodology, expand its scope, and improve its practical applications, contributing to the ongoing advancement of sentiment analysis within the film industry.

ACKNOWLEDGEMENTS

The authors would like to express their sincere gratitude to Buddhi Dharma University, Tangerang, Indonesia, for its invaluable support throughout the course of this research.

REFERENCES

- [1] M. Alzate, M. Arce-Urriza, and J. Cebollada, "Mining the text of online consumer reviews to analyze brand image and brand positioning," *Journal of Retailing and Consumer Services*, vol. 67, pp. 1–29, Jul. 2022, doi: 10.1016/j.jretconser.2022.102989.
- [2] M. Sykora, S. Elayan, I. R. Hodgkinson, T. W. Jackson, and A. West, "The power of emotions: Leveraging user generated content for customer experience management," *Journal of Business Research*, vol. 144, pp. 997–1006, May 2022, doi: 10.1016/j.jbusres.2022.02.048.
- [3] J. Serrano-Guerrero, M. Bani-Doumi, F. P. Romero, and J. A. Olivas, "Understanding what patients think about hospitals: A deep learning approach for detecting emotions in patient opinions," *Artificial Intelligence in Medicine*, vol. 128, pp. 1–11, Jun. 2022,

- doi: 10.1016/j.artmed.2022.102298.
- [4] B. Wan, P. Wu, C. K. Yeo, and G. Li, "Emotion-cognitive reasoning integrated BERT for sentiment analysis of online public opinions on emergencies," *Information Processing and Management*, vol. 61, no. 2, pp. 1–20, Mar. 2024, doi: 10.1016/j.ipm.2023.103609.
 - [5] J. Rogmann, J. Beckmann, R. Gaschler, and H. Landmann, "Media sentiment emotions and consumer energy prices," *Energy Economics*, vol. 130, pp. 1–15, Feb. 2024, doi: 10.1016/j.eneco.2023.107278.
 - [6] S. Gannouni, A. Aledaily, K. Belwafi, and H. Aboalsamh, "Electroencephalography based emotion detection using ensemble classification and asymmetric brain activity," *Journal of Affective Disorders*, vol. 319, pp. 416–427, Dec. 2022, doi: 10.1016/j.jad.2022.09.054.
 - [7] J. Hartmann, J. Huppertz, C. Schamp, and M. Heitmann, "Comparing automated text classification methods," *International Journal of Research in Marketing*, vol. 36, no. 1, pp. 20–38, Mar. 2019, doi: 10.1016/j.ijresmar.2018.09.009.
 - [8] J. C. Meléndez, E. Satorres, M. Reyes-Olmedo, I. Delhom, E. Real, and Y. Lora, "Emotion recognition changes in a confinement situation due to COVID-19," *Journal of Environmental Psychology*, vol. 72, pp. 1–6, Dec. 2020, doi: 10.1016/j.jenvp.2020.101518.
 - [9] Q. Wang, T. Su, R. Y. K. Lau, and H. Xie, "DeepEmotionNet: Emotion mining for corporate performance analysis and prediction," *Information Processing and Management*, vol. 60, no. 3, May 2023, doi: 10.1016/j.ipm.2022.103151.
 - [10] D. Li, B. Hu, Q. Chen, and S. He, "Towards Faithful Explanations for Text Classification with Robustness Improvement and Explanation Guided Training," in *Proceedings of the Annual Meeting of the Association for Computational Linguistics*, 2023, pp. 1–14. doi: 10.18653/v1/2023.trustnlp-1.1.
 - [11] F. Müller-Hansen *et al.*, "Attention, sentiments and emotions towards emerging climate technologies on Twitter," *Global Environmental Change*, vol. 83, pp. 1–11, Dec. 2023, doi: 10.1016/j.gloenvcha.2023.102765.
 - [12] P. Ekman, "What Scientists Who Study Emotion Agree About," *Perspectives on Psychological Science*, vol. 11, no. 1, pp. 31–34, Jan. 2016, doi: 10.1177/1745691615596992.
 - [13] S. M. Mohammad, "Practical and Ethical Considerations in the Effective use of Emotion and Sentiment Lexicons," 2020, doi: 10.48550/arXiv.2011.03492.
 - [14] S. M. Mohammad, "Ethics Sheet for Automatic Emotion Recognition and Sentiment Analysis," *Computational Linguistics*, vol. 48, no. 2, pp. 239–278, 2022, doi: 10.1162/coli_a_00433.
 - [15] S. M. Mohammad and P. D. Turney, "Crowdsourcing a Word-Emotion Association Lexicon," *Computational Intelligence*, vol. 29, no. 3, pp. 436–465, 2013, doi: 10.48550/arXiv.1308.6297.
 - [16] M. Rohse, R. Day, and D. Llewellyn, "Towards an emotional energy geography: Attending to emotions and affects in a former coal mining community in South Wales, UK," *Geoforum*, vol. 110, pp. 136–146, Mar. 2020, doi: 10.1016/j.geoforum.2020.02.006.
 - [17] A. Khunteta and P. Singh, "Emotion Cause Pair Extraction by Multi Task Learning on Enhanced English Dataset," *Procedia Computer Science*, vol. 218, pp. 766–777, 2022, doi: 10.1016/j.procs.2023.01.057.
 - [18] H. Huang, M. Fan, and C. A. Chou, "Graph-based learning of nonlinear physiological interactions for classification of emotions," *Pattern Recognition*, vol. 143, Nov. 2023, doi: 10.1016/j.patcog.2023.109794.
 - [19] J. Alcaraz, M. Labbé, and M. Landete, "Support Vector Machine with feature selection: A multiobjective approach," *Expert Systems with Applications*, vol. 204, pp. 1–14, Oct. 2022, doi: 10.1016/j.eswa.2022.117485.
 - [20] M. R. A. Rashid, K. F. Hasan, R. Hasan, A. Das, M. Sultana, and M. Hasan, "A comprehensive dataset for sentiment and emotion classification from Bangladesh e-commerce reviews," *Data in Brief*, vol. 53, pp. 1–14, Apr. 2024, doi: 10.1016/j.dib.2024.110052.
 - [21] X. Zheng and R. Schweickert, "Differentiating dreaming and waking reports with automatic text analysis and Support Vector Machines," *Consciousness and Cognition*, vol. 107, Jan. 2023, doi: 10.1016/j.concog.2022.103439.
 - [22] X. Luo, "Efficient English text classification using selected Machine Learning Techniques," *Alexandria Engineering Journal*, vol. 60, no. 3, pp. 3401–3409, Jun. 2021, doi: 10.1016/j.aej.2021.02.009.
 - [23] X. Shu and Y. Ye, "Knowledge Discovery: Methods from data mining and machine learning," *Social Science Research*, vol. 110, pp. 1–16, Feb. 2023, doi: 10.1016/j.ssresearch.2022.102817.
 - [24] A. F. Pathan and C. Prakash, "Cross-Domain Aspect Detection and Categorization using Machine Learning for Aspect-based Opinion Mining," *International Journal of Information Management Data Insights*, vol. 2, no. 2, pp. 1–18, Nov. 2022, doi: 10.1016/j.jjimei.2022.100099.
 - [25] A. L. Maas, R. E. Daly, P. T. Pham, D. Huang, A. Y. Ng, and C. Potts, "Learning word vectors for sentiment analysis," *ACL-HLT 2011 - Proceedings of the 49th Annual Meeting of the Association for Computational Linguistics: Human Language Technologies*, vol. 1, pp. 142–150, 2011.
 - [26] A. Petukhova and N. Fachada, "TextCL: A Python package for NLP preprocessing tasks," *SoftwareX*, vol. 19, pp. 1–6, Jul. 2022, doi: 10.1016/j.softx.2022.101122.
 - [27] D. Ramachandran and R. Parvathi, "Analysis of Twitter Specific Preprocessing Technique for Tweets," *Procedia Computer Science*, vol. 165, pp. 245–251, 2019, doi: 10.1016/j.procs.2020.01.083.
 - [28] N. H. Gabriela, R. Siautama, C. I. A. Amadea, and D. Suhartono, "Extractive Hotel Review Summarization based on TF/IDF and Adjective-Noun Pairing by Considering Annual Sentiment Trends," *Procedia Computer Science*, vol. 179, pp. 558–565, 2021, doi: 10.1016/j.procs.2021.01.040.
 - [29] F. Jáněz-Martino, R. Alaiz-Rodríguez, V. González-Castro, E. Fidalgo, and E. Alegre, "Classifying spam emails using agglomerative hierarchical clustering and a topic-based approach," *Applied Soft Computing*, vol. 139, pp. 1–16, May 2023, doi: 10.1016/j.asoc.2023.110226.
 - [30] P. H. Prastyo, I. Ardiyanto, and R. Hidayat, "A Combination of Query Expansion Ranking and GA-SVM for Improving Indonesian Sentiment Classification Performance," *Procedia CIRP*, vol. 189, pp. 108–115, 2021, doi: 10.1016/j.procs.2021.05.074.
 - [31] A. Orea-Giner, L. Fuentes-Moraleda, T. Villacé-Molinero, A. Muñoz-Mazón, and J. Calero-Sanz, "Does the Implementation of Robots in Hotels Influence the Overall TripAdvisor Rating? A Text Mining Analysis from the Industry 5.0 Approach," *Tourism Management*, vol. 93, pp. 1–10, Dec. 2022, doi: 10.1016/j.tourman.2022.104586.
 - [32] T. H. J. Hidayat, Y. Ruldeviyani, A. R. Aditama, G. R. Madya, A. W. Nugraha, and M. W. Adisaputra, "Sentiment analysis of twitter data related to Rinca Island development using Doc2Vec and SVM and logistic regression as classifier," *Procedia Computer Science*, vol. 197, pp. 660–667, 2021, doi: 10.1016/j.procs.2021.12.187.
 - [33] D. Valero-Carreras, J. Alcaraz, and M. Landete, "Comparing two SVM models through different metrics based on the confusion matrix," *Computers and Operations Research*, vol. 152, pp. 1–12, Apr. 2023, doi: 10.1016/j.cor.2022.106131.
 - [34] K. Shrivastava, S. Kumar, and D. K. Jain, "An effective approach for emotion detection in multimedia text data using sequence based convolutional neural network," *Multimedia Tools and Applications*, vol. 78, no. 20, pp. 29607–29639, Oct. 2019, doi: 10.1007/s11042-019-07813-9.

BIOGRAPHIES OF AUTHORS

Aditiya Hermawan    received a Master's degree in Computer Science from Budi Luhur University and is currently a Lecturer in the Informatics Engineering Study Program, Faculty of Science and Technology, Universitas Buddhi Dharma, Tangerang, Indonesia. In his role, he teaches various information technology courses and actively engages in research focused on new technologies and their applications. His research interests include data mining, machine learning, focusing on developing algorithms for data-driven predictions and decisions; database design, and blockchain technology. He can be contacted at email: aditiya.hermawan@ubd.ac.id.



Rico Yusuf    received his Bachelor's degree in Information Systems from Universitas Buddhi Dharma, Banten, Indonesia. He is actively involved in research projects focused on data mining and machine learning. As a member of the Data Science Research Group at Universitas Buddhi Dharma, he contributes to developing innovative solutions in information systems. His research interests include data mining, machine learning, and the optimization of information systems. He can be contacted at email: rico.yusuf@ubd.ac.id.



Benny Daniawan    pursued a Master's degree (S2) in Information Systems at Budi Luhur University Jakarta in 2014 and graduated in 2016 with a thesis titled "Evaluation of Lecturer Performance using AHP and TOPSIS Methods." The author is continuing doctoral studies in Computer Science at Satya Wacana Christian University Salatiga. Since 2017, the author has been working as a lecturer in the Information Systems Study Program at Buddhi Dharma University Tangerang. His research interests include multicriteria decision making (MCDM) and the technology acceptance model (TAM). He can be contacted at email: b3n2y.miracle@gmail.com.



Junaedi    completed his studies in Computer Science and obtained a Master's degree (S2) from Budi Luhur University and currently he is a lecturer in the Information Systems Studies Program, Faculty of Science and Technology, Buddhi Dharma University, Tangerang, Indonesia. In his roles, he taught various mathematics of information systems and is currently actively engaged in research focused on the latest technologies. His research interests include the internet of things, decision support systems, machine learning, and blockchain technology. He can be contacted at email: junaedi@ubd.ac.id.