

Sentiment Analysis of YouTube Users on Blackpink Kpop Group Using IndoBERT

Received:
28 May 2024

Accepted:
18 July 2024

Published:
1 August 2024

^{1*}Slamet Riyadi, ²Lathifah Khansa Salsabila, ³Cahya Damarjati,
⁴Rohana Abdul Karim

¹⁻³Department of Information Technology, Faculty of Engineering,
Universitas Muhammadiyah Yogyakarta, Indonesia

⁴Faculty of Electrical and Electronic Engineering, Universiti Malaysia
Pahang Al-Sultan Abdullah, Malaysia

E-mail: ¹riyadi@umy.ac.id, ²lathifah.khansa.ft20@umy.ac.id,
³cahya.damarjati@umy.ac.id, ⁴rohanaak@umpsa.edu.my

*Corresponding Author

Abstract— Background: The Korean Pop (K-Pop) phenomenon has become an important part of popular culture worldwide, with Blackpink being one of the most influential groups. Analyzing sentiment toward Blackpink is urgent, given its growing popularity and wide influence among fans worldwide. In the present technological era, social media platforms such as YouTube have evolved into a space where artists and their fans may interact with each other. As a consequence, social media has become a powerful tool for assessing the emotional tone and sentiment conveyed by individuals. **Objective:** This research aims to explore the trend of public sentiment towards Blackpink and evaluate how well the IndoBERT model analyzes the sentiment of Indonesian texts. **Methods:** The objective of this study is to examine the pattern of public sentiment towards Blackpink and assess the proficiency of the IndoBERT model in analyzing the sentiment of Indonesian writings. **Results:** The findings demonstrated that the IndoBERT model had an exceptional level of precision, achieving a 98% accuracy rate. In addition, it obtained a fl, recall, and accuracy score of 95%. The remarkable results demonstrate the efficacy of the IndosBERT technique in evaluating the emotion of Indonesian-language literature towards Blackpink. **Conclusion:** This study enhances the knowledge of how fans and audiences react to K-pop material and establishes a foundation for future research and advancement. The impressive precision of the IndoBERT model showcases its capacity for sentiment analysis in Indonesian literature, making it a useful tool for future research endeavors.

Keywords— Sentiment Analysis; Blackpink; IndoBERT; Youtube; K-Pop

This is an open access article under the CC BY-SA License.



Corresponding Author:

Slamet Riyadi,
Faculty of Engineering, Department of Information Technology,
Universitas Muhammadiyah Yogyakarta, Yogyakarta, Indonesia.
Email: riyadi@umy.ac.id
Orchid ID: <https://orcid.org/0000-0003-1981-8876>



I. INTRODUCTION

K-pop, or Korean Pop, is one of the music genres from South Korea. K-pop is part of the "Korean Wave" or Hallyu. This term refers to the popularity of pop culture from the Land of Gingseng, including Korean television shows, music, and movies throughout Asia and other parts of the world [1], [2]. This data processing is essential for text summarization, slang or dialect sentence conversion, and classification [3]. Text classification automatically classifies data by grouping text into predefined categories [4], [5]. Text classification is usually used as sentiment analysis, and one of the most used methods is Natural Language Processing (NLP) [6], [7]. NLP allows measuring and analyzing the sentiment of a sentence and determining which parts are important[8], [9].

Research related to sentiment analysis on one of the K-pop groups has previously been conducted by Noviana et al [10] using the Naïve Bayes and Support Vector Machine (SVM) methods. This research produced an accuracy value of 81% for the SVM method, while the Naïve Bayes method produced an accuracy value of 79%. However, although Natural Language Processing (NLP) and deep learning have advanced, machine learning architecture still struggles to classify text. Context from earlier words cannot be maintained in machine learning algorithms. Because of this, the created context frequently ignores the sentence's word order[11].

In addition, the previous research discussed user sentiment about the K-pop concert NCT 127 [12]. After the data is collected through crawling and preprocessing, the results show that Twitter users tend to give a positive opinion about K-pop concerts. Research conducted by Rizkina et al. [13] examines the analysis of emotions in K-pop performances and the development of classification strategies utilizing the Naïve Bayes Classifier approach. The data acquired from Twitter exhibited a significant level of accuracy. This study investigates the positive and negative emotions sent by individuals on Twitter. The method used includes preprocessing data, data labeling, TF-IDF, and implementing the Naïve Bayes Algorithm; the result shows a value accuracy of 82%.

Language models must be trained frequently using a huge amount of data to predict words accurately. A pre-trained language model has been trained on a substantial amount of data [14]. One of the most widely used pre-trained models with transformer architecture is BERT (Bidirectional Encoder Representation from Transformers), which is thought to be an efficient way to comprehend complex texts because it is made to use bidirectional text representation[15], [16]. Bert's pre-trained models may be found in other languages, such as Indonesian. An Indonesian pre-trained model based on BERT is called IndoBERT [17], [18].

Several research projects have thoroughly examined and categorized the emotional content of written text in Indonesian utilizing the IndoBERT technique. Nevertheless, this study has not specifically investigated the emotional reaction, particularly focused on certain issues within the K-pop genre. This research distinguishes itself from previous polls by focusing on the sentiment especially directed at Blackpink. This K-pop girl group is generally acknowledged as the most influential worldwide. This research used the Indonesian Bidirectional Encoder Representations from Transformers (IndoBERT) methodology to analyze Blackpink's YouTube social media data sentiment. This research aims to ascertain the response of fans and viewers towards their work. This work employs the IndoBERT methodology, which utilizes state-of-the-art improvements in natural language processing to improve the accuracy of sentiment analysis [19], [20]. Consequently, this enables a deeper comprehension of the public's response to Blackpink. This method is unique since a dearth of research has focused on analyzing emotions towards Blackpink. Hence, the objective of this study is to examine the patterns in public sentiment towards Blackpink and evaluate the effectiveness of the IndoBERT model. The study's notable contribution lies in its concentration on a particular and well acknowledged topic within the K-pop business, distinguishing it from prior studies and emphasizing its distinctiveness and significance.

II. RESEARCH METHOD

The stages in this research can be seen in Fig 1.

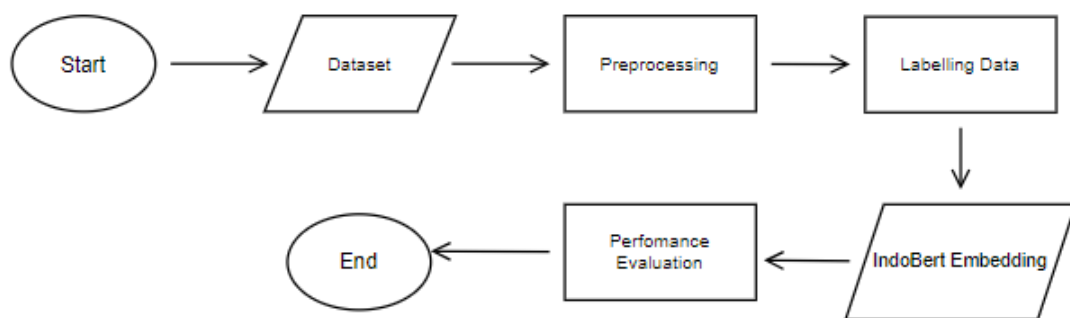


Fig 1. Research Diagram

A. Dataset

The study used user comment data from the YouTube social media platform. The data was collected by evaluating the audience's response to the Shut Down Music Video through review. For this study, a total of 3,971 data points were collected from September 2022 to February 2023. Next, the data will be divided into 80% for training data, 10% for validation data, and 10% for

testing data. The dataset distribution can be seen in Table 1. Dataset examples can be seen in Table 2.

Table 1. Dataset Distribution

Training	Validation	Test
3176	397	398

Table 2. Dataset Example

No	Time	Review
1	2023-01-06T02:08:24Z	liriknya kerenn bgt siiii wkwwkwk, apalagi pas di rap jennie
2	2022-09-20T20:38:51Z	sumpah ROSE si ratuuu Outfittt gk ada Obatt ijo ijo ðŸ˜šðŸ˜šðŸ˜š & btw Maaf nih Rose Pacar ku ðŸ˜šðŸ˜šðŸ˜šðŸ˜š
3	2022-09-16T09:06:36Z	Jennie cantik ya.....ðŸ™,,
4	2022-11-09T15:07:15Z	Kyknya pink venom kali ini mmg ngumpulin smua MV lama2 nya , property+ part2 nya
5	2022-09-18T23:20:21Z	Keren bgt yaa MV blackpink kali ini ngegabungin semua konsep MV sebelumnya di MV yg sekarang keren bgttttðŸŽ

B. Preprocessing

A crucial aspect of doing sentiment analysis is the preparation of the text. The process aids in cleansing and standardizing the text data, facilitating its analysis [21]. Preprocessing involves many processes aimed at eliminating unnecessary words or characters that are irrelevant to the categorization process. The process involves the elimination of characters and numbers, dividing the text into smaller units called tokens, reducing words to their base form (stemming), and eliminating often used words (stop words) [22]. Figure 2 depicts the sequential steps involved in the preprocessing procedure.

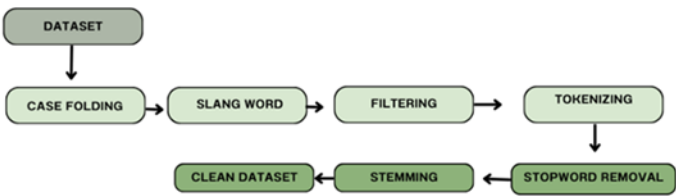


Fig 2. Preprocessing stage [11]

1. Case Folding: Case folding is the process of converting all uppercase characters in the dataset to lowercase. The objective of this procedure is to transform all characters in the dataset to lowercase with the aim of facilitating generalization [23].
2. Slangword: translates the slang/slang sentences used by replacing the sentences with standard sentences according to the dictionary [11].
3. Filtering: the filtering step is intended to clean or filter the review text. The text will be cleaned of unwanted elements.
4. Tokenizing is a procedure used to divide sentences into parts of words, punctuation marks, and other expressions that have meaning according to the language rules [23].
5. Stop word removal: is the removal of common and irrelevant words that are unlikely to convey much sentiment, such as "and" "the" "from" [21].
6. Stemming: To improve the performance of sentiment analysis, stemming is one of the preprocessing steps [24]. Stemming is a general method for processing and retrieving natural language information. The main goal is to reduce words to their basic or root forms [25].

C. Labelling Data

A lexicon-based labeling method was used in this study. This method involves words linked to scores that indicate positive, negative, or neutral characters[26]. Analyzing data sets in Indonesian, this research implements the Inset Lexicon dictionary. The dictionary contains words and their weights, which determine the polarity value [27], [28]. The polarity value is calculated by adding up all the word weights in the review text[29]. After each word in the text is labeled, the overall sentiment score is calculated by counting the number of words with positive and negative values [30]. The formula commonly used to calculate the sentiment score (StSc) is:

$$StSc = \frac{\text{number of positive words} - \text{number of negative words}}{\text{Total number of words}} \quad (1)$$

Reviews have positive sentiment if the polarity is greater than zero, and negative sentiment if the polarity is lower than zero will become neutral if the polarity is zero.

$$StSc = \begin{cases} 1 & \text{positive word} > \text{negative} \\ -1 & \text{positive word} < \text{negative} \\ 0 & \text{other} \end{cases} \quad (2)$$

D. IndoBERT Embedding

Word Embedding is the term used to describe the process of transforming words into vectors [31]. IndoBERT is a BERT model that has undergone training on an Indonesian language dataset using masked language modeling [32]. The IndoBERT Transformer was trained using over 220 million words extracted from three main texts. The 90 million words of the Indonesian web Corpus, 55 million words from news items such as Kompas, Tempo, and Liputan6, and 74 million words from the Indonesian Wikipedia[33]. Adding a special token to the start and finish of the sentence- the token [CLS] at the start and [SEP] at the end, is the first step in the embedding process [34]. These tokens are used to differentiate between the start and finish of sentences. Subsequently, the word chunk tokenization approach is used to transform sentences into individual words or tokens [31]. the purpose of this procedure is to provide language processing (NLP) models with the capability to comprehend the significance and context of various words in a given text or phrase [35]. Fig 3 contains the embedding flow using IndoBERT. Segment embedding allows BERT to understand context, identify sentence boundaries, and understand relationships in text. If the input consists of one sentence, it will form a zero-embedding segment [36]. Apart from that, positional embedding is used to embed the position and meaning of each word in a sentence. In other words, it is a way to show the order of words in a sentence [37].

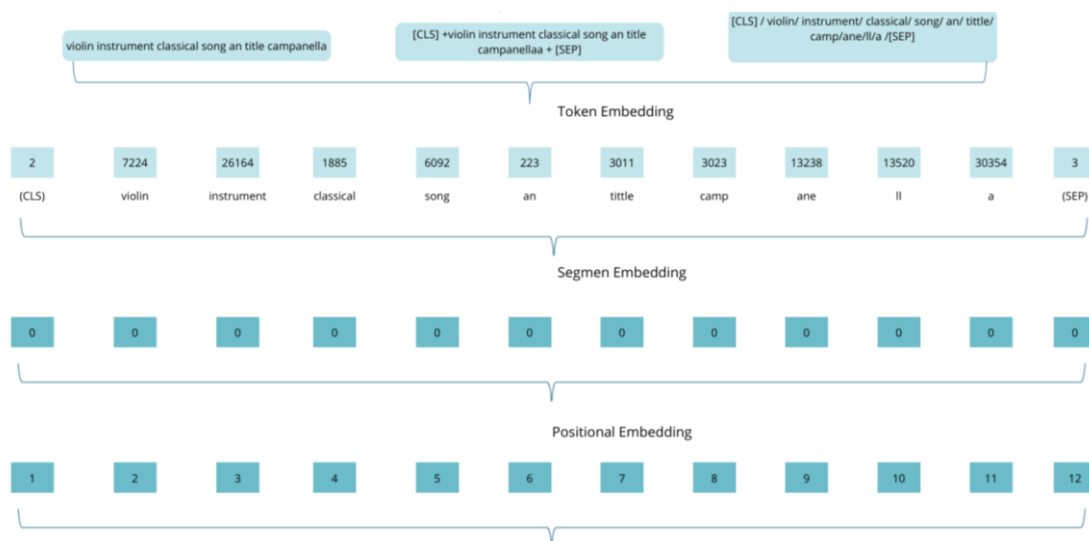


Fig 3. IndoBERT Embedding [31]

Performance Evaluation

Performance evaluation after collecting classification results, the model's effectiveness should be assessed. Accuracy and F1 measures are the two assessment measures used in this study[38], [39]. The True Positive (TP), False Positive (FP), True Negative (TN), and False Negative (FN) values must first be determined to compute this metric. The quantity of data

accurately categorized as positive is shown by the TP number, as is the amount of data correctly identified as negative by FP and as negative by FN. The extent to which a classification model can correctly predict each case in a dataset is seen by a metric known as accuracy. Accuracy is the ratio of accurate predictions to all other predictions—a formula to calculate the accuracy score.

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (3)$$

On the other hand, the F1 score is an evaluation metric that combines recall and precision, providing a deeper understanding of the performance of a classification model. Recall quantifies how well the model generates real, applicable, and positive outcomes. The percentage of actual positive results to all positive expectations is known as recall.

$$Precision = \frac{TP}{TP + FP} \quad (4)$$

$$Recall = \frac{TP}{TP + FN} \quad (5)$$

$$Fscore = 2 * \frac{Precision * Recall}{Precision + Recall} \quad (6)$$

III. RESULT AND DISCUSSION

This research used a dataset including 3,971 samples. The research indicated that the neutral attitude group accounted for the largest share, with 48.75% or 1,936 samples. The presence of positive emotion was found in 28.03% of the samples, amounting to 1,113 instances. By contrast, 23.22% of the sample showed negative reactions, totaling 922 incidents. This finding was obtained through the use of the Inset Lexicon dictionary in the labelling process. The results of this poll deserve attention as they emphasize unfavourable emotions towards Blackpink, and show a fair and equal perspective among society as a whole. The comprehensive labeling result is seen in Fig 4.

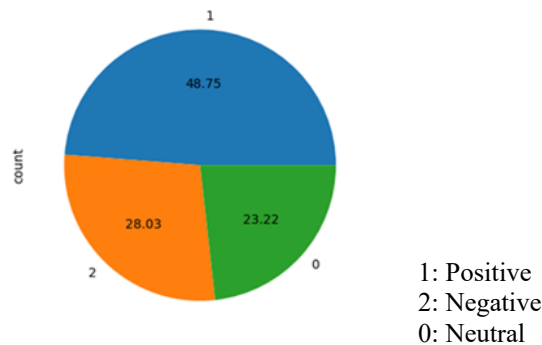


Fig 4. Labeling Result

The PyTorch framework used with the Python programming language is used in this research. In this study, the number of epochs is 10, the number of batches is 16, and the learning rate is $2e-5$ as parameters for Adam optimization. Fig 5 displays the training curve for the dataset. In general, the model demonstrates a consistent improvement in performance throughout the training process. Both the training loss and validation loss consistently drop, while the training accuracy and validation accuracy steadily climb until they reach their peak in the final epoch. Fig 6 shows the IndoBERT confusion matrix with the best performance in this research. This matrix shows the model's ability to classify data correctly, as shown in the figure. The model correctly predicted 882 events as positive, 1,529 as neutral, and 738 as negative. In 10 cases, the model incorrectly predicted neutral sentiment as negative. There was 1 case of neutral sentiment incorrectly predicted as positive, 10 cases of the model incorrectly predicted positive sentiment as negative, and six positive events incorrectly predicted as neutral.

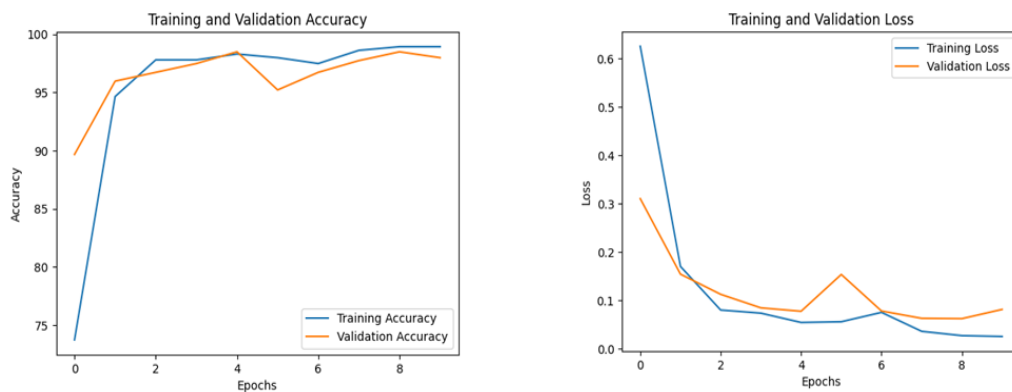


Fig 5. Accuracy and Loss Training Curve

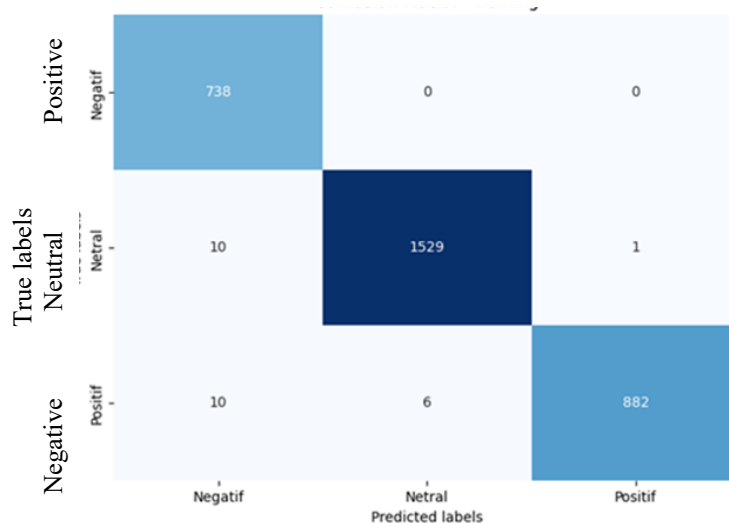


Fig 6. Confusion Matrix

All training models can learn at each epoch, continuously improving results. The validation results show that the average loss is 0.0085. With an accuracy of 98% on validation data, the model could correctly classify approximately 98% of validation data. Precision measures how well the model can recover all positive cases. All classes have a precision above 0.95, which shows that the model has few false positives. All classes have a recall above 0.95, which shows that the model has few false negatives. All classes have high f1-scores, indicating that the model has a good balance between precision and recall. The evaluation results show that the model performs excellently in classifying validation data with high accuracy, precision, recall, and f1-score levels for each class. The results of the performance evaluation are shown in Table 3.

The study results indicate that the IndoBERT model has a remarkable accuracy rate of 98%, with precision, recall, and f1-scores above 0.95. This demonstrates the efficacy of the model in comprehending the emotional content of Indonesian text pertaining to Blackpink. The findings of this study align with previous studies that have shown the exceptional efficacy of the IndoBERT model in assessing sentiment in Indonesian literature. Previous research using the Naïve Bayes and Support Vector Machine (SVM) methods achieved accuracy rates of 81% for SVM and 79% for Naïve Bayes, highlighting the challenges faced in text classification utilizing both approaches [10]. Previous studies used the Support Vector Machine (SVM) technique to examine the sentiment of Twitter users towards the NCT 127 K-pop concert [12]. Furthermore, other research endeavors employed the Naïve Bayes Classifier approach to assess the sentiment of K-pop concerts, resulting in significant accuracy rates [13]. The consistent use of sentiment analysis methodologies across many domains enhances the credibility of the present results obtained using the IndoBERT model.

Table 3. Performance Evaluation

	Precision	Recall	F1-score
Negative	0.95	1.00	0.97
Netral	0.98	0.98	0.98
Positif	1.00	0.96	0.98
Accuracy	0.98		

IV. CONCLUSION

This study focuses on the sentiment analysis of Indonesian literature via the IndoBERT technique. A total of 3,971 data points were obtained from the gathering of comments on Blackpink's "Shut Down" music video, covering the period from September 2022 to February 2023. The sentiment analysis resulted in the classification of 1,936 comments as neutral, 1,113

comments as positive, and 922 comments as negative. The model evaluation demonstrated outstanding performance, with an accuracy score of 98% and F1 score, recall, and precision all above 95%. The results illustrate the effectiveness of the IndoBERT approach for analyzing the emotion of Indonesian texts, establishing it as a dependable methodology in this domain.

Nevertheless, there are limitations to this study that should be acknowledged in further research. An inherent constraint is the concentration on a solitary dataset derived from YouTube comments, which may not comprehensively depict the range of emotions conveyed across various platforms or situations. Further investigation might examine the utilization of IndoBERT on a wider array of datasets and contemplate the creation of sophisticated BERT variations specifically designed for analyzing Indonesian text. Furthermore, staying updated with advancements in natural language processing and integrating novel techniques might significantly improve the precision and dependability of sentiment analysis in Indonesian texts. Continuous development is essential for ensuring the relevance and effectiveness of sentiment analysis technologies.

Author Contributions: *Slamet Riyadi*: Conceptualization, Methodology, Writing - Original Draft, Writing - Review & Editing, Supervision. *Lathifah Khansa Salsabila*: Software, Investigation, Data Curation, Writing - Original Draft. *Cahya Damarjati*: Investigation, Data Curation. *Rohana Abdul Karim*: Review.

Funding: This work was supported by Universitas Muhammadiyah Yogyakarta under Hibah Penelitian Internal 2023/2024 Grant number 50/R-LRI/XII/2023.

Acknowledgments: The author would like to thank Universitas Muhammadiyah Yogyakarta for the research and financial support for the publication.

Conflicts of Interest: The authors declare no conflict of interest.

Data Availability: The data cannot be openly shared for the protection of study participant privacy.

Informed Consent: There were no human subjects.

Animal Subjects: There were no animal subjects.

ORCID:

Slamet Riyadi: <https://orcid.org/0000-0003-1981-8876>

Lathifah Khansa Salsabila: <https://orcid.org/0009-0009-0480-8078>

Cahya Damarjati: <https://orcid.org/0000-0003-4389-6321>

Rohana Abdul Karim: <https://orcid.org/0000-0002-4469-912X>

REFERENCES

- [1] C. Liu, "The Research on the Influence of KPOP (Korean Popular Music) Culture on Fans," *Communications in Humanities Research*, vol. 4, no. 1, pp. 63–68, May 2023, **doi:** 10.54254/2753-7064/4/20220177.

- [2] S. M. Justine Miguel and J. Xavier Chavez, "The Korean Wave: A Quantitative Study On K-Pop's Aesthetic Presence in The Philippines Multimedia Industry," 2023.
- [3] M. Novo-Lourés, R. Pavón, R. Laza, D. Ruano-Ordas, and J. R. Méndez, "Using Natural Language Preprocessing Architecture (NLPA) for Big Data Text Sources," *Sci Program*, vol. 2020, no. 1, p. 2390941, Jan. 2020, doi: 10.1155/2020/2390941.
- [4] A. Gasparetto, M. Marcuzzo, A. Zangari, and A. Albarelli, "A Survey on Text Classification Algorithms: From Text to Predictions," *Information 2022, Vol. 13, Page 83*, vol. 13, no. 2, p. 83, Feb. 2022, doi: 10.3390/INFO13020083.
- [5] S. Minaee, N. Kalchbrenner, E. Cambria, N. Nikzad, M. Chenaghlu, and J. Gao, "Deep Learning--based Text Classification," *ACM Computing Surveys (CSUR)*, vol. 54, no. 3, Apr. 2021, doi: 10.1145/3439726.
- [6] S. Rohajawati *et al.*, "Unveiling Insights: A Knowledge Discovery Approach to Comparing Topic Modeling Techniques in Digital Health Research," *INTENSIF: Jurnal Ilmiah Penelitian dan Penerapan Teknologi Sistem Informasi*, vol. 8, no. 1, pp. 111–125, Feb. 2024, doi: 10.29407/INTENSIF.V8I1.22058.
- [7] A. Fadlil, S. Sunardi, and R. Ramdhani, "Similarity Identification Based on Word Trigrams Using Exact String Matching Algorithms," *INTENSIF: Jurnal Ilmiah Penelitian dan Penerapan Teknologi Sistem Informasi*, vol. 6, no. 2, pp. 253–270, Aug. 2022, doi: 10.29407/INTENSIF.V6I2.18141.
- [8] I. Lauriola, A. Lavelli, and F. Aioli, "An introduction to Deep Learning in Natural Language Processing: Models, techniques, and tools," *Neurocomputing*, vol. 470, pp. 443–456, Jan. 2022, doi: 10.1016/j.neucom.2021.05.103.
- [9] E. Lindrawati, E. Utami, and A. Yaqin, "Comparison of Modified Nazief&Adriani and Modified Enhanced Confix Stripping algorithms for Madurese Language Stemming," *INTENSIF: Jurnal Ilmiah Penelitian dan Penerapan Teknologi Sistem Informasi*, vol. 7, no. 2, pp. 276–289, Aug. 2023, doi: 10.29407/INTENSIF.V7I2.20103.
- [10] R. Noviana and I. Rasal B A Jurusan, "PENERAPAN ALGORITMA NAIVE BAYES DAN SVM UNTUK ANALISIS SENTIMEN BOY BAND BTS PADA MEDIA SOSIAL TWITTER," *Jurnal Teknik dan Science*, vol. 2, no. 2, pp. 51–60, Jun. 2023, doi: 10.56127/JTS.V2I2.791.
- [11] G. Z. Nabiilah, S. Y. Prasetyo, Z. N. Izdiyar, and A. S. Girsang, "BERT base model for toxic comment analysis on Indonesian social media," in *Procedia Computer Science*, Elsevier B.V., 2022, pp. 714–721. doi: 10.1016/j.procs.2022.12.188.
- [12] Dessy Angelina, U. Hayati, and G. Dwilestari, "Penerapan Metode Support Vector Machine Pada Sentimen Analisis Pengguna Twitter Terhadap Konser K-Pop," *Kopertip : Jurnal Ilmiah Manajemen Informatika dan Komputer*, vol. 7, no. 1, pp. 14–23, Feb. 2023, doi: 10.32485/kopertip.v7i1.251.
- [13] N. Q. Rizkina and F. N. Hasan, "Analisis Sentimen Komentar Netizen Terhadap Pembubaran Konser NCT 127 Menggunakan Metode Naive Bayes," *Journal of Information System Research (JOSH)*, vol. 4, no. 4, pp. 1136–1144, Jul. 2023, doi: 10.47065/JOSH.V4I4.3803.
- [14] C. P. Chai, "Comparison of text preprocessing methods," *Nat Lang Eng*, vol. 29, no. 3, pp. 509–553, May 2023, doi: 10.1017/S1351324922000213.
- [15] J. Devlin, M.-W. Chang, K. Lee, K. T. Google, and A. I. Language, "BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding," *Proceedings of the 2019 Conference of the North*, pp. 4171–4186, 2019, doi: 10.18653/V1/N19-1423.
- [16] M. Li *et al.*, "TrOCR: Transformer-Based Optical Character Recognition with Pre-trained Models," *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 37, no. 11, pp. 13094–13102, Jun. 2023, doi: 10.1609/AAAI.V37I11.26538.
- [17] A. Rahmawati, A. Alamsyah, and A. Romadhony, "Hoax News Detection Analysis using IndoBERT Deep Learning Methodology," *2022 10th International Conference on*

- Information and Communication Technology, ICoICT 2022*, pp. 368–373, 2022, **doi:** 10.1109/ICOICT55009.2022.9914902.
- [18] S. Saadah, K. M. Auditama, A. A. Fattahila, F. I. Amorokhman, A. Aditsania, and A. A. Rohmawati, “Implementation of BERT, IndoBERT, and CNN-LSTM in Classifying Public Opinion about COVID-19 Vaccine in Indonesia,” *Jurnal RESTI (Rekayasa Sistem dan Teknologi Informasi)*, vol. 6, no. 4, pp. 648–655, Aug. 2022, **doi:** 10.29207/RESTI.V6I4.4215.
- [19] J. H. Joloudari *et al.*, “BERT-deep CNN: state of the art for sentiment analysis of COVID-19 tweets,” *Soc Netw Anal Min*, vol. 13, no. 1, pp. 1–14, Dec. 2023, **doi:** 10.1007/S13278-023-01102-Y/METRICS.
- [20] S. Srivastava, M. K. Sarkar, and C. Chakraborty, “Machine Learning Approaches for COVID-19 Sentiment Analysis: Unveiling the Power of BERT,” *2024 IEEE 14th Annual Computing and Communication Workshop and Conference, CCWC 2024*, pp. 92–97, 2024, **doi:** 10.1109/CCWC60891.2024.10427866.
- [21] L. Igual and S. Seguí, “Basics of Natural Language Processing,” pp. 195–210, 2024, **doi:** 10.1007/978-3-031-48956-3_10.
- [22] P. M. Lavanya and E. Sasikala, “Deep learning techniques on text classification using Natural language processing (NLP) in social healthcare network: A comprehensive survey,” *2021 3rd International Conference on Signal Processing and Communication, ICSPSC 2021*, pp. 603–609, May 2021, **doi:** 10.1109/ICSPSC51351.2021.9451752.
- [23] W. Bourequat and H. Mourad, “Sentiment Analysis Approach for Analyzing iPhone Release using Support Vector Machine,” *International Journal of Advances in Data and Information Systems*, vol. 2, no. 1, pp. 36–44, Apr. 2021, **doi:** 10.25008/IJADIS.V2I1.1216.
- [24] S. Al-Saqqa, A. Awajan, and S. Ghoul, “Stemming Effects on Sentiment Analysis using Large Arabic Multi-Domain Resources,” *2019 6th International Conference on Social Networks Analysis, Management and Security, SNAMS 2019*, pp. 211–216, Oct. 2019, **doi:** 10.1109/SNAMS.2019.8931812.
- [25] L. Albraheem and H. S. Al-Khalifa, “Exploring the problems of sentiment analysis in informal Arabic,” *ACM International Conference Proceeding Series*, pp. 415–418, 2012, **doi:** 10.1145/2428736.2428813.
- [26] M. Danubianu Stefan, A. Barila, M. Danubianu, and B. Gradinaru, “Romanian-Lexicon-Based Sentiment Analysis for Assessing Teachers’ Activity,” *IJCSNS International Journal of Computer Science and Network Security*, vol. 22, no. 10, p. 43, 2022, **doi:** 10.22937/IJCSNS.2022.22.10.7.
- [27] D. Fimoza, A. Amalia, and T. Henny Febriana Harumy, “Sentiment Analysis for Movie Review in Bahasa Indonesia Using BERT,” *2021 International Conference on Data Science, Artificial Intelligence, and Business Analytics, DATABIA 2021 - Proceedings*, pp. 27–34, 2021, **doi:** 10.1109/DATABIA53375.2021.9650096.
- [28] V. Bonta, N. Kumareash, and N. Janardhan, “A Comprehensive Study on Lexicon Based Approaches for Sentiment Analysis,” *Asian Journal of Computer Science and Technology*, vol. 8, no. S2, pp. 1–6, Jan. 2019, **doi:** 10.51983/AJCST-2019.8.S2.2037.
- [29] S. Anggina, N. Y. Setiawan, and F. A. Bachtiar, “Analisis Ulasan Pelanggan Menggunakan Multinomial Naïve Bayes Classifier dengan Lexicon-Based dan TF-IDF Pada Formaggio Coffee and Resto,” *is The Best Accounting Information Systems and Information Technology Business Enterprise this is link for OJS us*, vol. 7, no. 1, pp. 76–90, Sep. 2022, **doi:** 10.34010/aisthebest.v7i1.7072.
- [30] C. S. G. Khoo and S. B. Johnkhan, “Lexicon-based sentiment analysis: Comparative evaluation of six sentiment lexicons,” <https://doi.org/10.1177/0165551517703514>, vol. 44, no. 4, pp. 491–511, Apr. 2017, **doi:** 10.1177/0165551517703514.

- [31] G. Z. Nabiilah, I. N. Alam, E. S. Purwanto, and M. F. Hidayat, "Indonesian multilabel classification using IndoBERT embedding and MBERT classification," *International Journal of Electrical and Computer Engineering*, vol. 14, no. 1, pp. 1071–1078, Feb. 2024, **doi:** 10.11591/ijece.v14i1.pp1071-1078.
- [32] B. Wilie *et al.*, "IndoNLU: Benchmark and Resources for Evaluating Indonesian Natural Language Understanding," 2020. Accessed: Jul. 14, 2024. [Online]. Available: <https://aclanthology.org/2020.aacl-main.85>
- [33] F. Koto, A. Rahimi, J. H. Lau, and T. Baldwin, "IndoLEM and IndoBERT: A Benchmark Dataset and Pre-trained Language Model for Indonesian NLP," Nov. 2020, [Online]. Available: <http://arxiv.org/abs/2011.00677>
- [34] H. Chen *et al.*, "Pre-trained image processing transformer," *Proceedings of the IEEE Computer Society Conference on Computer Vision and Pattern Recognition*, pp. 12294–12305, 2021, **doi:** 10.1109/CVPR46437.2021.01212.
- [35] R. Patil, S. Boit, V. Gudivada, and J. Nandigam, "A Survey of Text Representation and Embedding Techniques in NLP," *IEEE Access*, vol. 11, pp. 36120–36146, 2023, **doi:** 10.1109/ACCESS.2023.3266377.
- [36] G. Zain Nabiilah, I. Nur Alam, E. Setyo Purwanto, and M. Fadlan Hidayat, "Indonesian multilabel classification using IndoBERT embedding and MBERT classification," *International Journal of Electrical and Computer Engineering (IJECE)*, vol. 14, no. 1, pp. 1071–1078, 2024, **doi:** 10.11591/ijece.v14i1.pp1071-1078.
- [37] A. Zhao and Y. Yu, "Knowledge-enabled BERT for aspect-based sentiment analysis," *Knowl Based Syst*, vol. 227, p. 107220, Sep. 2021, **doi:** 10.1016/J.KNOSYS.2021.107220.
- [38] S. Sucipto, D. D. Prasetya, and T. Widiyaningtyas, "Educational Data Mining: Multiple Choice Question Classification in Vocational School," *Matrik: Jurnal Manajemen, Teknik Informatika, dan Rekayasa Komputer*, vol. 23, no. 2, pp. 367–376, 2024, **doi:** 10.30812/matrik.v23i2.3499.
- [39] H. Hairani and T. Widiyaningtyas, "Augmented Rice Plant Disease Detection with Convolutional Neural Networks," *INTENSIF: Jurnal Ilmiah Penelitian dan Penerapan Teknologi Sistem Informasi*, vol. 8, no. 1, pp. 27–39, Feb. 2024, **doi:** 10.29407/INTENSIF.V8I1.21168.