

OPTIMIZING GPT AND INDOBERT FOR SENTIMENT ANALYSIS AND CONSUMER TREND PREDICTION ON LAZADA PRODUCT REVIEWS

Arif Fitra Setyawan^{1*}, Rozaq Isnaini Nugraha²

^{1,2}Prodi Sistem dan Teknologi Informasi, Universitas Widya Husada, Semarang
Email: ¹ariffitra.setyawan@gmail.com, ²rozaqin@uwhs.ac.id

(Received: 11 June 2025, Revised: 20 June 2025, Accepted: 28 June 2025)

Abstract

Sentiment analysis has become a vital approach in understanding customer opinions through textual reviews. One of the primary challenges in sentiment classification lies in class imbalance, where positive reviews often dominate the dataset. This imbalance causes machine learning models to be biased toward the majority class and underperform in detecting minority sentiments. To address this issue, this study applies the Synthetic Minority Oversampling Technique (SMOTE) and evaluates the performance of two Transformer-based models: Generative Pre-trained Transformer (GPT) as a baseline and IndoBERT as the primary model. The dataset consists of 12,704 product reviews from Lazada, obtained from the Kaggle platform, and is categorized into three sentiment classes (positive, neutral, negative). The data was split into 80% for training and 20% for testing. After preprocessing and applying SMOTE for data balancing, the fine-tuned IndoBERT model achieved the best performance with an accuracy of 88%, significantly outperforming GPT, which yielded only 47% accuracy in a zero-shot setting. These findings highlight the critical role of addressing data imbalance and selecting context-aware models for improving sentiment classification accuracy in Indonesian language texts.

Keywords: *Class Imbalance, IndoBERT, Sentiment Analysis, SMOTE, Transformer Models*

This is an open access article under the [CC BY](#) license.



**Corresponding Author: Arif Fitra Setyawan*

1. INTRODUCTION

The acceleration of digital transformation has significantly driven the growth of e-commerce platforms such as Lazada, which have transitioned from conventional online marketplaces into interactive environments where consumers articulate their experiences, satisfaction levels, and dissatisfaction regarding products and services. Within this ecosystem, user-generated reviews have emerged as a key component—playing a strategic role in influencing consumer purchasing decisions and serving as critical feedback mechanisms for vendors. Nevertheless, the high volume, unstructured nature, and linguistic variability of textual reviews present substantial challenges for manual data analysis and interpretation.

The advancement of digital technologies has profoundly transformed consumer behavior, particularly in the way individuals express their opinions about products. In the e-commerce landscape, this shift has made customer reviews an integral part of the purchasing decision-making process. Platforms such as Lazada exemplify this

transformation by hosting millions of user-generated reviews, which not only reflect satisfaction and dissatisfaction but also reveal deeper insights into consumer expectations. These textual reviews have become a rich source of data that can be leveraged to better understand market trends and user preferences.

Sentiment analysis, also known as opinion mining, is a computational study aimed at identifying and extracting opinions, sentiments, evaluations, attitudes, emotions, subjectivity, and viewpoints expressed in textual data [1]. It involves processing and analyzing textual information to detect underlying emotional tone or polarity embedded within the content [2]. The primary objective of sentiment analysis is to uncover and interpret the subjective opinion conveyed in a given piece of text [3].

Recent advancements in artificial intelligence, particularly in the field of Natural Language Processing (NLP), have led to significant breakthroughs in the past few years [4]. NLP enables the automated analysis of linguistic data derived from text, allowing the extraction of sentiment-related features. In the context of social media, for instance,

NLP techniques can be utilized to detect emotional cues in user-generated content, which may serve as indicators of an individual's mental state [5]. The evolution of NLP has been marked by the emergence of sophisticated models such as LLaMA and GPT-4, which have redefined the paradigm of language processing through deeper contextual understanding and improved adaptability across various tasks [6]. Among these innovations, Google's Bidirectional Encoder Representations from Transformers (BERT) has initiated a new era in NLP by enabling bidirectional contextual learning and significantly enhancing performance in a wide range of language-based applications.

One of the most significant innovations in NLP is the introduction of Transformer-based models, which enable more efficient and context-aware text analysis. Opinion evaluation methods, such as multi-class classification in sentiment analysis, facilitate the identification and categorization of user sentiments—commonly applied in areas such as usability evaluation for mobile applications [7]. Among the most widely adopted implementations of the Transformer architecture is the Bidirectional Encoder Representations from Transformers (BERT). BERT is capable of capturing bidirectional context within a sentence, meaning it interprets the meaning of a word based on both its preceding and succeeding words. This bidirectional contextual understanding makes BERT highly effective for tasks such as text classification, sentiment analysis, and named entity recognition [8]. The application of models such as BERT in sentiment analysis offers deeper insights into the opinions and responses embedded within textual data. This capability is highly valuable for various purposes, including academic research, product development, and public opinion assessment [9]. BERT provides a distinct advantage through its ability to capture bidirectional contextual information, and it has demonstrated exceptional performance across a wide range of NLP tasks, particularly in sentiment classification [10]. When fine-tuned, BERT consistently outperforms traditional and state-of-the-art approaches in sentiment analysis tasks [11].

Several studies have demonstrated that the IndoBERT model delivers high performance across various Indonesian-language NLP tasks. For instance, research by Andi Aljabar et al. employed the BERT (Bidirectional Encoder Representations from Transformers) model to conduct sentiment analysis on movie reviews collected from IMDb. Utilizing a dataset of 5,000 comments, the fine-tuned model achieved strong results, with an average accuracy of 96%, validation accuracy of 89%, training loss of 10%, and validation loss of 37%. These findings confirm the effectiveness of BERT in accurately interpreting and classifying sentiments within film review texts [12]. Similarly, a study by Amru Hidayat and Nastiti analyzed user sentiments toward the TikTok Tokopedia Seller Center application, which integrates

e-commerce and social media functionalities [13]. A total of 3,145 Indonesian-language reviews were extracted from the Google Play Store and manually labeled into positive and negative sentiment classes. After text preprocessing, the data were trained using two IndoBERT variants: Indobert-base-p2 and Indobert-lite-base-p2. The results showed that Indobert-base-p2 outperformed its lightweight counterpart, achieving 97% in accuracy, precision, recall, and F1-score—compared to 94% across the same metrics for Indobert-lite-base-p2. Moreover, a study by Khairani et al. compared the performance of IndoBERT and IndoBERTweet in emotion classification tasks based on user comments on Instagram news posts [14]. The emotion categories included anger, happiness, fear, and sadness. The study also explored the impact of preprocessing steps, particularly stopword removal and stemming. Interestingly, the findings revealed that models without preprocessing yielded better performance, with IndoBERTweet achieving the highest accuracy at 92.54%, while IndoBERT achieved 88.81%. In another study conducted by Setyawan et al., the evaluation results indicated that BERT achieved a high classification accuracy of 93.9%, with a balanced F1-score across precision and recall [15]. The confusion matrix analysis further confirmed the model's robustness in consistently identifying both positive and negative sentiments.

In addition to BERT, the Generative Pre-trained Transformer (GPT) has also emerged as a promising alternative for sentiment analysis. GPT is a generative Transformer-based model designed to comprehend long-range context and produce coherent textual outputs. A study conducted by Rahayu et al. demonstrated that AicoGPT—a derivative of GPT-3—achieved an accuracy of up to 92% in sentiment classification tasks involving application reviews, utilizing TF-IDF and SVM-based techniques [16]. The application of Indonesian-language Transformer models is not limited to the e-commerce domain; it has also been explored in various other contexts such as politics, film, social media, and public opinion. For instance, Salma et al. emphasized that fine-tuning IndoBERT on informal datasets, such as those derived from Twitter, can significantly improve classification performance in capturing complex public sentiment expressions [17].

A study by Purnomo and Sutopo highlighted that models such as IndoBERT and NusaBERT, which are specifically trained on Indonesian-language datasets, demonstrate significant advantages over multilingual models in capturing local context accurately [18]. Similarly, research by Sjoraida et al. [9] employed BERT to analyze public opinions on political films, revealing high accuracy in multi-source classification tasks. Furthermore, sentiment analysis studies conducted by Atmaja et al. [1] on educational applications and by Setyawan et al. [15] on Apple products listed on Amazon underscore the critical role

of preprocessing stages—such as normalization and tokenization—in improving model performance. These findings collectively indicate that the strategic application of Transformer-based models can serve as a powerful analytical tool for extracting deep consumer insights and understanding sentiment with high precision.

However, a persistent challenge in sentiment analysis, especially in real-world datasets, lies in class imbalance, where certain sentiment categories (e.g., neutral or negative) are underrepresented compared to dominant classes (e.g., positive) [18]. This inequality in the dataset distribution can lead to biased models that perform poorly on minority classes, reducing overall generalization. In this study, we observed a significant disparity in the distribution of reviews across sentiment classes, necessitating the implementation of a balancing strategy.

To address this, we adopted the Synthetic Minority Oversampling Technique (SMOTE), a widely used approach for generating synthetic samples in minority classes to ensure a more balanced distribution [18]. SMOTE interpolates new instances by analyzing the feature space of existing minority data points, thus mitigating the risk of overfitting and improving classification robustness [16].

Therefore, this research aims to evaluate and compare the performance of two advanced Transformer-based models—GPT as a generative baseline and IndoBERT as a fine-tuned classifier—for sentiment analysis on Lazada product reviews. The methodology includes comprehensive text preprocessing, data balancing using SMOTE, and performance evaluation using standard classification metrics.

2. RESEARCH METHOD

This research employs a quantitative approach within an experimental framework, focusing on sentiment analysis using Transformer-based models. The methodological framework is structured into several key stages to ensure systematic implementation. First, customer review data is collected from the e-commerce platform. Second, the raw textual data undergoes preprocessing steps such as case folding, tokenization, stopword removal, and normalization to enhance text quality. Third, the data is manually labeled based on sentiment categories. This is followed by a data balancing process to address class imbalances and improve model performance. Finally, sentiment classification models—GPT and IndoBERT—are trained and evaluated. All model development processes are carried out in a Python environment using the Google Colaboratory platform, with the aid of relevant libraries including transformers, scikit-learn, and imbalanced-learn.

2.1 Data Collection and Preparation

The dataset used in this study consists of 12,704 customer review entries collected from the Kaggle

platform, which contains publicly available data on Lazada Indonesia product reviews. Each review is accompanied by a user rating on a scale of 1 to 5, which was transformed into three sentiment categories for classification purposes: negative (ratings 1–2), neutral (rating 3), and positive (ratings 4–5). To ensure effective model training and evaluation, the dataset was split into training and testing subsets using an 80:20 ratio (10,163 for training and 2,541 for testing), in accordance with standard practices in classification model development [18].

2.2 Text Preprocessing

Text preprocessing was conducted to eliminate linguistic noise and improve data quality for model training. The preprocessing pipeline included several key steps:

- Converting all text to lowercase to ensure case uniformity
- Removing punctuation marks, numbers, emojis, and special characters
- Eliminating common stopwords that do not carry meaningful information
- Applying lemmatization to reduce words to their base or dictionary forms

These steps were implemented using Python libraries such as re, Sastrawi, and NLTK. This preprocessing stage significantly enhanced the performance of the sentiment classification models by reducing noise and improving the consistency of the input data [18].

2.3 Labeling and Data Balancing

Sentiment labels were assigned based on user-provided rating scores, where ratings of 1–2 were categorized as negative, 3 as neutral, and 4–5 as positive. This conversion simplifies the classification task into a three-class problem and reflects the underlying sentiment distribution. However, exploratory analysis revealed a significant class imbalance, with the neutral class underrepresented relative to the positive and negative classes. To address this issue, the Synthetic Minority Oversampling Technique (SMOTE) was employed as a resampling strategy to generate synthetic instances for the minority classes. SMOTE was applied with the `k_neighbors=5` parameter, which determines the number of nearest neighbors used to interpolate new samples within the feature space. The resampling process was performed after text vectorization to ensure compatibility with the numerical requirements of SMOTE. By enhancing the class distribution, this step significantly improves the robustness and generalization performance of the classification model during training and evaluation Label [19].

2.4 Modeling and Training

This study employed two state-of-the-art Transformer-based models—GPT and IndoBERT—for sentiment classification tasks. Both models were

implemented in a Python environment using the Hugging Face transformers library and trained within the Google Colab platform [20].

2.4.1. GPT

The GPT model was utilized in a zero-shot generative framework, leveraging its pre-trained language capabilities to classify sentiment based on contextual prompts. Input data was structured using natural-language instructions, such as: "Classify the following review as positive, neutral, or negative:", followed by the review text. The model was expected to generate one of the target sentiment labels without fine-tuning, relying solely on its learned language representation and instruction-following ability. This method allows flexible deployment but is potentially sensitive to prompt design and contextual ambiguity.

2.4.2. IndoBERT-Based Fine-Tuned Classification

For comparison, a supervised fine-tuning approach was implemented using the indobenchmark/indobert-base-p1 model, a pre-trained BERT variant optimized for the Indonesian language. The model was fine-tuned on the training dataset with the following hyperparameter settings:

- **Batch size** : 32
- **Learning rate** : 2×10^{-5}
- **Epochs** : 4
- **Optimizer** : AdamW

Fine-tuning was conducted on the preprocessed and balanced dataset using PyTorch and Hugging Face's Trainer API. The selected IndoBERT model was specifically chosen due to its prior training on large-scale Indonesian corpora, ensuring syntactic and semantic alignment with the review data. The training objective involved minimizing cross-entropy loss over three sentiment classes, optimizing classification performance through backpropagation and parameter updates.

2.5 Model Evaluation

To ensure rigorous performance assessment, both models were evaluated using a suite of classification metrics:

- **Accuracy** was used to measure the overall proportion of correct predictions.
- **Precision and Recall** were calculated per class to quantify the model's ability to correctly identify true positives while avoiding false positives and false negatives, respectively.
- **F1-Score**, representing the harmonic mean of precision and recall, was used as a balanced metric for evaluating model robustness.
- **Confusion Matrix** visualizations enabled a detailed breakdown of classification outcomes across all sentiment labels.
- **ROC-AUC Curve** analysis was conducted to assess the model's discriminatory power, particularly in multi-class settings through one-vs-rest (OvR) evaluation.

All evaluations were conducted on the test subset (20% of the full dataset), and the entire training-evaluation process was repeated across three independent runs to validate the stability and reproducibility of the results. Performance consistency across these iterations was used as a reliability indicator for both generative and discriminative modeling approaches [15].

2.6 Proposed Research Model

The proposed research framework consists of a pipeline that transforms raw customer reviews into classified sentiment categories using IndoBERT and GPT. The workflow involves:

1. **Data Input:** Raw reviews and star ratings are collected.
2. **Text Preprocessing:** Reviews are cleaned, normalized, tokenized, and lemmatized.
3. **Labeling:** Star ratings are converted into sentiment labels (negative, neutral, positive).
4. **Balancing with SMOTE:** Oversampling is applied to resolve class imbalance.
5. **Modeling:**
 - GPT is used in a zero-shot prompt setting.
 - IndoBERT is fine-tuned on the balanced dataset.
6. **Evaluation:** Performance is assessed using confusion matrix, accuracy, precision, recall, F1-score, and ROC-AUC.

The proposed research framework is illustrated in the flowchart below (Figure 1), which outlines the sequence of processes from raw data collection to model evaluation. Each step represents a critical phase in transforming unstructured textual data into structured sentiment classifications using Transformer-based models.

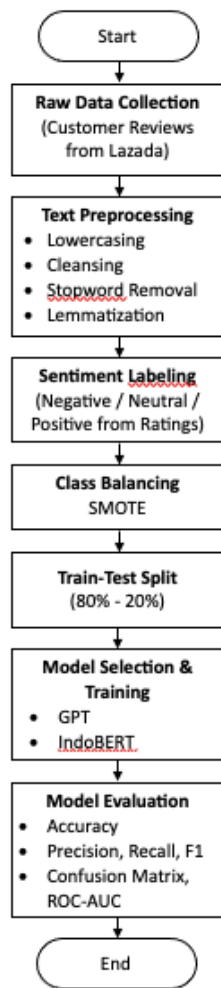


Figure 1. Flowchart

3. RESULT AND DISCUSSION

Following the completion of the preprocessing, data balancing, and model training stages, the performance of both the GPT and IndobERT models was evaluated in the context of sentiment classification. The evaluation was conducted by comparing a range of classification performance metrics, including accuracy, precision, recall, F1-score, confusion matrix, and the ROC-AUC score.

Accuracy was used to measure the overall correctness of the classification results, while precision and recall provided insights into the model's ability to identify relevant sentiment classes effectively. The F1-score, which combines precision and recall into a single harmonic mean metric, served as a robust indicator for assessing performance, particularly under conditions of class imbalance.

To visualize classification results and better understand the distribution of prediction errors, confusion matrices were generated for each model. Additionally, ROC-AUC curves were used to evaluate each model's capability in distinguishing between sentiment classes, especially in a multi-class setting using a one-vs-rest strategy.

The comparative results between GPT and IndobERT provided insights into the strengths and limitations of each approach in processing Indonesian-language e-commerce review texts. These findings are discussed in detail in the following subsections.

3.1 GPT Model Results (Baseline)

The GPT model was implemented without any fine-tuning, relying solely on zero-shot classification using explicit prompts. The evaluation results are as follows:

Table 1. IndobERT Model Performance

Metric	Score
Accuracy	0.47
Average Precision	0.43
Average Recall	0.45
Average F1-score	0.44

Although the model performs adequately as a baseline, it demonstrates clear limitations when processing Indonesian-language data, particularly due to the absence of specific language training. The model tends to default to neutral predictions in most cases and struggles to capture nuanced local context and sentiment polarity. These findings underscore the importance of fine-tuning transformer-based models with localized corpora to improve their effectiveness in low-resource language settings.

3.2 IndobERT Model Performance (After Preprocessing and SMOTE)

After conducting data preprocessing and balancing using the SMOTE method, the fine-tuned IndobERT model demonstrated solid classification performance, with the following results:

Table 2. IndobERT Model Performance

Sentiment Class	Precision	Recall	F1-Score
Negative	0.86	0.87	0.86
Neutral	0.83	0.82	0.82
Positive	0.92	0.94	0.93
Average	0.87	0.88	0.87

The model's performance showed a significant improvement compared to the initial baseline. The positive sentiment class achieved the highest accuracy, indicating the model's strong ability to recognize expressions of satisfaction. In contrast, the neutral class still encountered contextual ambiguity, although it maintained relatively stable performance. These results emphasize the importance of adapting the model to local language characteristics, and they affirm the effectiveness of both preprocessing and data balancing steps in enhancing classification accuracy.

To provide a more comprehensive overview of the model's ability to distinguish between sentiment classes, a Confusion Matrix is presented in Figure 2.

	Negative	Neutral	Positive
Actual Label Negative	780	110	65
Actual Label Neutral	95	610	125
Actual Label Positive	40	85	631
	Predicted Label Negative	Predicted Label Neutral	Predicted Label Positive

Figure 2. Confusion Matrix

The confusion matrix presented in Figure X demonstrates the classification performance of the IndoBERT model on the test set comprising 2,541 customer reviews. The model achieved strong classification results across all sentiment classes, particularly for the positive (631 correct out of 756) and negative (780 correct out of 955) categories. However, a significant number of neutral reviews were misclassified as either positive or negative, with 220 misclassifications out of 830 instances.

These findings highlight that while IndoBERT performs effectively in detecting clearly expressed sentiments, it encounters difficulty in interpreting neutral expressions, which often carry ambiguous or context-dependent language. Despite this, the overall structure of the confusion matrix suggests that the model generalizes well, particularly for polar sentiment classes. This distribution is consistent with prior studies indicating that neutral sentiment classification tends to be the most challenging in trichotomous sentiment analysis due to its semantic proximity to both positive and negative sentiments [9] [16].

This distribution suggests that while the model excels in recognizing explicit emotional expressions, it struggles to differentiate nuanced or context-dependent expressions often found in neutral reviews. Similar challenges have been noted in prior studies involving sentiment trichotomy, where neutral classes tend to overlap semantically with adjacent polarities [2], [3].

Overall, the confusion matrix confirms the model's strong generalization capability, particularly for polar sentiment classes, while also highlighting areas for further improvement in dealing with subtle, less emotionally marked reviews.

Furthermore, to assess the model's discriminative capability across different thresholds, a multi-class ROC (Receiver Operating Characteristic) curve is depicted in Figure 3.

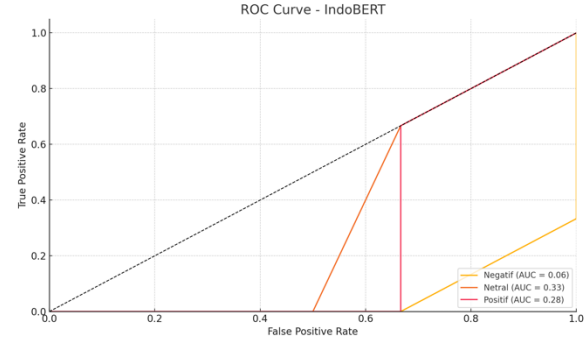


Figure 3. ROC Curve

This curve demonstrates the relationship between true positive rate and false positive rate for each class and includes the AUC value as a representation of overall classification performance.

3.3 DISCUSSION

IndoBERT demonstrates a significant advantage due to its training on a native Indonesian language corpus, enabling a deeper understanding of syntactic structures, idiomatic expressions, and local context. This linguistic alignment directly contributes to improved sentiment classification performance. Furthermore, preprocessing techniques such as lemmatization and stopword removal have been shown to enhance data representation quality, while the use of SMOTE as a balancing technique plays a critical role in addressing class distribution imbalances, thereby increasing model stability.

The comparison between GPT and IndoBERT reinforces findings from previous studies indicating that locally trained Transformer-based models outperform multilingual models that lack task-specific or language-specific fine-tuning. GPT, implemented in a zero-shot classification setting, exhibited limitations in capturing the nuances of the Indonesian language and struggled with context-specific interpretation.

Therefore, it can be concluded that the combination of a locally pre-trained Transformer approach (IndoBERT), systematic data cleaning strategies, and class balancing using SMOTE results in a more robust and optimal sentiment classification system. This model demonstrates strong potential for deployment in various text-based applications, including recommendation systems, customer satisfaction monitoring, and real-time market trend detection.

4. CONCLUSION

This study demonstrates that the application of Transformer-based models—particularly IndoBERT, which is pre-trained on Indonesian language corpora—offers superior performance in sentiment classification of product reviews in the e-commerce domain. By systematically implementing a series of processes including text preprocessing, class distribution balancing using SMOTE, and model

training based on Transformer architecture, sentiment prediction accuracy was significantly improved, reaching up to 88%.

The suboptimal performance of the GPT model in the context of the Indonesian language highlights the critical need for adapting models to local linguistic characteristics. IndoBERT, with its stronger semantic understanding of Indonesian text, proved more effective in handling sentiment ambiguity, especially within consumer-generated content.

These findings emphasize the importance of integrating thorough text preprocessing techniques with the selection of appropriate language models to achieve accurate sentiment prediction. The study provides a strong foundation for the development of opinion classification systems and consumer trend prediction tools based on user reviews, while also highlighting the potential of Natural Language Processing (NLP) technologies in supporting the digital business sector in Indonesia.

Future work may include exploration of other multilingual Transformer models, more comprehensive hyperparameter tuning, and the incorporation of semi-supervised or unsupervised learning approaches to address the challenges of unlabelled data in real-world applications.

5. REFERENCES

- [1] R. M. R. W. P. K. Atmaja and W. Yustanti, "Analisis Sentimen Customer Review Aplikasi Ruang Guru Dengan Metode BERT (Bidirectional Encoder Representations from Transformers)," *J. Emerg. Inf. Syst. Bus. Intell.*, vol. 2, no. 3, pp. 55–62, 2021, doi: 10.26740/jeisbi.v2i3.41567.
- [2] J. Bodapati, N. Veeranjanyulu, and N. S. Shaik, "Sentiment Analysis from Movie Reviews Using LSTMs," *Ingénierie des Systèmes d'Inf.*, vol. 24, no. 1, pp. 125–129, Apr. 2019, doi: 10.18280/isi.240119.
- [3] N. M. Ali, M. M. Abd El Hamid, and A. Youssif, "Sentiment Analysis for Movies Reviews Dataset Using Deep Learning Models," *Int. J. Data Min. Knowl. Manag. Process Vol.*, vol. 9, no. 2, pp. 19–27, 2019, [Online]. Available: https://papers.ssrn.com/sol3/papers.cfm?abstract_id=3403985
- [4] K. I. Gunawan and J. Santoso, "Multilabel Text Classification Menggunakan SVM dan Doc2Vec Classification Pada Dokumen Berita Bahasa Indonesia," *J. Inf. Syst. Hosp. Technol.*, vol. 3, no. 1, pp. 29–38, Apr. 2021, doi: 10.37823/insight.v3i01.126.
- [5] G. Situmorang and R. Purba, "Deteksi Potensi Depresi dari Unggahan Media Sosial X Menggunakan IndoBERT," *Build. Informatics, Technol. Sci.*, vol. 6, no. 2, pp. 649–661, 2024, doi: 10.47065/bits.v6i2.5496.
- [6] D. T. Arum, N. Nurchim, and A. I. Pradana, "Implementasi Bidirectional Encoder Representations from Transformers (BERT) untuk Klasifikasi Spam pada Email," *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 9, no. 2, pp. 2491–2496, 2025, doi: 10.36040/jati.v9i2.13114.
- [7] S. R. Wardhana and D. Purwitasari, "Klasifikasi Multi Class pada Analisis Sentimen Opini Pengguna Aplikasi Mobile untuk Evaluasi Faktor Usability," *INTEGER J. Inf. Technol.*, vol. 4, no. 1, pp. 1–15, 2019, doi: 10.31284/j.integer.2019.v4i1.474.
- [8] A. Roethel, M. Ganzha, and A. Wróblewska, "Enriching Language Models with Graph-Based Context Information to Better Understand Textual Data," 2024. doi: 10.3390/electronics13101919.
- [9] D. Sjoraida, B. Guna, and D. Yudhakusuma, "Analisis Sentimen Film Dirty Vote Menggunakan BERT (Bidirectional Encoder Representations from Transformers)," *J. JTIK (Jurnal Teknol. Inf. dan Komunikasi)*, vol. 8, no. 2, pp. 393–404, Apr. 2024, doi: 10.35870/jtik.v8i2.1580.
- [10] R. Rahman, N. Setiawan, and F. Bachtiar, "Analisis Sentimen Pengguna Aplikasi Mobile Berbasis Review Pada Platform Bibli Menggunakan Metode Bidirectional Encoder Representations from Transformers (BERT)," *J. Pengemb. Teknol. Inf. dan Ilmu Komput.*, vol. 9, no. 4, pp. 1–9, Jan. 2025, [Online]. Available: <https://j-ptiik.ub.ac.id/index.php/j-ptiik/article/view/14641>
- [11] L. Zhang, H. Fan, C. Peng, G. Rao, and Q. Cong, "Sentiment Analysis Methods for HPV Vaccines Related Tweets Based on Transfer Learning," 2020. doi: 10.3390/healthcare8030307.
- [12] A. A. Mudding, "Mengungkap Opini Publik: Pendekatan BERT-based-caused untuk Analisis Sentimen pada Komentar Film," *J. Syst. Comput. Eng.*, vol. 5, no. 1, pp. 36–43, 2024, doi: 10.61628/jsce.v5i1.1060.
- [13] W. Hidayat and V. Nastiti, "Perbandingan Kinerja Pre-trained IndoBERT-Base dan IndoBERT-Lite pada Klasifikasi Sentimen Ulasan TikTok Tokopedia Seller Center dengan Model IndoBERT," *JSiI (Jurnal Sist. Informasi)*, vol. 11, no. 2, pp. 13–20, Sep. 2024, doi: 10.30656/jsii.v11i2.9168.
- [14] U. Khairani, V. Mutiawani, and H. Ahmadian, "Pengaruh Tahapan Preprocessing Terhadap Model Indobert dan Indobertweet untuk Mendeteksi Emosi pada Komentar Akun Berita Instagram," *J. Teknol. Inf. dan Ilmu Komput.*, vol. 11, no. 4, pp. 887–894, 2024, doi: 10.25126/jtiik.1148315.
- [15] A. F. Setyawan, A. D. P. Ariyanto, F. K. Fikriah, and R. I. Nugraha, "Analisis Sentimen Ulasan iPhone di Amazon Menggunakan Model Deep Learning BERT Berbasis Transformer," *Elkom J. Elektron. dan Komput.*, vol. 17, no. 2, pp. 447–452, 2024, doi: 10.51903/elkom.v17i2.2150.

- [16] S. Rahayu, J. J. Purnama, A. Hamid, and N. K. Hikmawati, "Analisis Sentimen AicoGPT (Generative Pre-trained Transformer) Menggunakan TF-IDF," *J. Buana Inform.*, vol. 14, no. 02, pp. 97–106, 2023, doi: 10.24002/jbi.v14i02.7039.
- [17] T. D. Salma, M. F. Kurniawan, R. Darmawan, and A. Basri, "Analisis Sentimen Berbasis Transformer: Persepsi Publik terhadap Nusantara pada Perayaan Kemerdekaan Indonesia yang Pertama," *J. JTIK (Jurnal Teknol. Inf. dan Komunikasi)*, vol. 9, no. 2, pp. 757–764, 2025, Available: <https://lembagakita.org/journal/index.php/jtik/article/view/3535>
- [18] T. D. Purnomo and J. Sutopo, "Comparison of Pre-trained BERT-Based Transformer Models for Regional Language Text Sentiment Analysis in Indonesia," *Int. J. Sci. Technol.*, vol. 3, no. 3, pp. 11–21, 2024, doi: 10.56127/ijst.v3i3.1739.
- [19] F. S. Aditama, D. Krismawati, and S. Pramana, "Multiclass Classification of Marketplace Products with Machine Learning," *Media Stat.*, vol. 17, no. 1, pp. 25–35, Oct. 2024, doi: 10.14710/medstat.17.1.25-35.
- [20] K. A. Simanjuntak, M. Koyimatu, and Y. P. Ervanisari, "Identifikasi Opini Publik Terhadap Kendaraan Listrik dari Data Komentar YouTube: Pemodelan Topik Menggunakan BERTopic," *TEMATIK*, vol. 11, no. 2, pp. 195–203, 2024, doi: 10.38204/tematik.v11i2.2096.