

A Bidirectional GRU Approach with Hyperparameter Optimization for Sentiment Classification in Game Reviews

Alif Januantara Prima¹, Nur Alamsyah^{*2}, Titan Parama Yoga³, Budiman⁴, Imannudin Akbar⁵,
Acep Hendra⁶

^{1,2,3,4,5,6}Universitas Informatika Dan Bisnis Indonesia, Indonesia

E-mail: ¹alif.jp21@student.unibi.ac.id, ^{*2}nuralamsyah@unibi.ac.id, ³titanparama@unibi.ac.id

⁴budiman@unibi.ac.id, ⁵imannudin@unibi.ac.id, ⁶acephendra@unibi.ac.id

Abstract

Sentiment analysis is essential for understanding user perspectives, particularly in contexts such as game reviews, where feedback directly influences product perception and user engagement. This study proposes a comparative approach employing Gated Recurrent Unit (GRU), hyperparameter-tuned GRU, and Bidirectional GRU models to classify sentiments within a game review dataset. The experimental pipeline begins with standard preprocessing and tokenization, followed by vectorization and supervised model training. Hyperparameter tuning is conducted using Keras Tuner to determine the optimal configuration of embedding dimensions, GRU units, dropout rates, and learning rates. The best-performing model—a Bidirectional GRU with optimized parameters—achieves a validation accuracy of 85.37% and demonstrates superior performance across key metrics, including precision, recall, and F1-score. Despite limitations posed by a relatively small and imbalanced dataset, the Bidirectional GRU exhibits strong generalization capability. Future work will explore class balancing techniques and the incorporation of pretrained word embeddings to further enhance model performance.

Keywords—Sentiment Analysis, Game Reviews, Bidirectional GRU, Hyperparameter Optimization, Deep Learning

Submission: May 19, 2025

Accepted: June 08, 2025

Published: July 10, 2025

DOI Article: <https://doi.org/10.25134/ilkom.v19i2.399>

1. INTRODUCTION

Sentiment analysis is a fundamental task in Natural Language Processing (NLP) that seeks to extract subjective opinions from text and classify them as positive, negative, or neutral [1]. Within the domain of game reviews, sentiment analysis offers valuable insights into user satisfaction, gameplay experience, and overall market reception [2][3]. As digital gaming platforms continue to expand, the volume of user-generated textual feedback grows accordingly presenting both a significant opportunity and a challenge for extracting meaningful patterns through automated analysis [4].

One of the main challenges in sentiment classification of game reviews lies in the nature of the data itself: user reviews are often short, informal, and structurally diverse [5]. Moreover, these datasets are frequently imbalanced, with a disproportionate number of positive reviews compared to negative ones [6]. These characteristics necessitate the use of models capable of capturing contextual meaning,

managing noisy and brief text, and maintaining robust performance despite limited and skewed data distributions [7].

Early approaches to sentiment analysis primarily relied on traditional machine learning algorithms such as Support Vector Machines (SVM), Naive Bayes, and logistic regression [8]. These models typically required extensive feature engineering, utilizing techniques like term frequency-inverse document frequency (TF-IDF), bag-of-words, and n-gram modeling [9]. While effective in certain applications, such methods are inherently limited in their ability to capture sequential dependencies and semantic nuances within language [10].

Advancements in deep learning have shifted the focus toward recurrent neural networks (RNNs), particularly Long Short-Term Memory (LSTM) and Gated Recurrent Unit (GRU) architectures. GRU models, known for being more computationally efficient than LSTMs, have demonstrated strong performance in various sequence classification tasks [11]. Further enhancements, such as the Bidirectional GRU, enable the model to process text in both

forward and backward directions, thereby improving its ability to capture contextual information—a crucial factor in sentiment analysis [12].

Recent studies have validated the effectiveness of GRU and its variants in sentiment classification. For example, Kaur et al. (2024) [13] implemented a GRU-based architecture for Twitter sentiment analysis and reported performance gains over standard RNNs. Likewise, Liu et al. (2024) [14] showed that Bidirectional GRU outperformed unidirectional models in classifying short-text reviews, demonstrating its advantage in handling the brevity and variability typical of user-generated content.

Hyperparameter tuning has been widely recognized as a critical component in optimizing deep learning models. Alamsyah et al. (2024) [15] demonstrated that random search can outperform traditional grid search in improving model performance. More recently, frameworks such as Keras Tuner and Optuna have introduced automated tools that streamline the tuning process for key parameters, including learning rate, number of hidden units, embedding dimensions, and dropout rates—often yielding substantial performance gains [16]. However, to the best of our knowledge, no prior study has conducted a systematic comparison of GRU, hyperparameter-tuned GRU, and Bidirectional GRU within a single experimental framework applied to game review sentiment datasets. This gap in the literature forms the basis and motivation for the present study.

The primary contribution of this research is a comprehensive evaluation of GRU-based architectures for sentiment classification on game review data, with particular focus on the effects of hyperparameter optimization and bidirectional structure. Unlike previous studies that investigated GRU or LSTM in isolation, this work systematically compares three configurations: a baseline GRU model, a GRU model optimized using Keras Tuner, and a Bidirectional GRU model employing the best-tuned parameters. By applying automated hyperparameter tuning, the study demonstrates measurable improvements in classification metrics such as accuracy, precision, and F1-score. Moreover, the use of Bidirectional GRU enables the model to capture richer contextual information from both directions of the input sequence, enhancing its generalization capability—even in the presence of limited and imbalanced data. These findings offer a practical foundation for future sentiment analysis research involving short-text reviews and underscore the

importance of integrating architectural enhancements with parameter optimization in deep learning applications.

2. RESEARCH METHODOLOGY

The research method employed in this study consists of several sequential stages, beginning with data collection and ending with model evaluation. The overall framework is illustrated in Figure 1, which outlines the complete pipeline for sentiment classification using Bidirectional GRU and its hyperparameter-optimized variant. The process begins with the collection of game review data, followed by a series of preprocessing steps including text cleaning, handling missing values, and converting labels into numeric format. The cleaned text is then transformed into sequences of integers through tokenization and subsequently padded to ensure uniform input length.

Subsequently, the dataset is divided into training and testing subsets, which serve as the basis for training two model configurations: a baseline Bidirectional GRU and a tuned Bidirectional GRU optimized through hyperparameter search. Both models are evaluated using key performance metrics, including accuracy, precision, recall, F1-score, and a confusion matrix. The objective of this methodological pipeline is to assess the extent to which hyperparameter tuning and bidirectional architecture contribute to improved accuracy, generalization, and robustness in sentiment classification of game review texts.

2.1. Data Collected

The dataset used in this study was sourced from the Kaggle platform, which contains labeled game review texts for sentiment classification tasks [17]. The dataset contains 2,000 labeled game reviews, with a class distribution of approximately 65% positive and 35% negative sentiments, indicating an imbalanced dataset. The average review length is around 23 words per entry, with substantial variation in style, length, and sentence structure. These characteristics make the dataset well-suited for evaluating model performance on short and informal user-generated text, which is typical in game review platforms.

Each data entry consists of a short user-written review, its cleaned version, and a sentiment label indicating whether the review is positive or negative. The dataset is particularly suitable for evaluating the performance of

sequence-based models due to the informal and varied nature of user-generated text. The features available in the dataset are described in Table 1.

Table 1. Description of Dataset Features

Feature Name	Data Type	Description
ReviewID	String	A unique identifier for each review entry.
Review	Text	The original raw review text as written by the user.

Feature Name	Data Type	Description
Processed_Review	Text	The cleaned and normalized version of the review, used as model input.
Label	Categorical	The sentiment category of the review, labeled as either Positive or Negative.

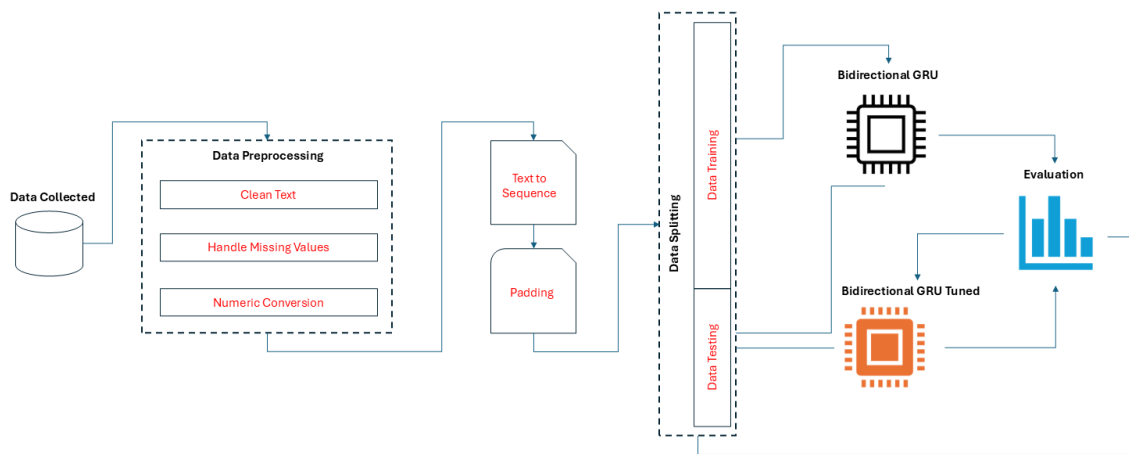


Figure 1. Pipeline For Sentiment Classification

2.2 Data Preprocessing

The data preprocessing stage is essential to prepare the raw textual data into a clean and structured format suitable for deep learning models [18]. As illustrated in Figure 1, this stage involves three main steps. The first step, clean text, transforms all characters in the review text to lowercase, removes punctuation, numbers, special characters, and excessive whitespace. This step helps eliminate noise that may hinder the model's learning process. The second step, handle missing values, involves removing any records that contain null or missing values in either the review content or the sentiment label.

This ensures the reliability of the dataset and avoids errors during the tokenization and training stages. The third step, numeric conversion, converts the sentiment labels from categorical format ("Positive" or "Negative") into binary numeric values (1 for positive, 0 for negative). This transformation allows the classification task to be modeled as a binary classification problem, aligning with the structure of the neural network's output layer. Overall, these preprocessing steps ensure that the

dataset is properly formatted for tokenization and subsequent training of the Bidirectional GRU models.

2.3 Text To Sequence And Padding

After the data has been cleaned and prepared, the next step is to convert the textual input into a numerical format that can be processed by neural networks [19]. This process begins with Text to Sequence, where each cleaned review is transformed into a sequence of integers using a tokenizer. The tokenizer assigns a unique integer to each word based on its frequency in the corpus, resulting in sequences of varying lengths depending on the number of words in each review.

Since neural network models require input data to be of uniform length, the second step, Padding, is applied to ensure consistency. Padding adds zeros to the end of shorter sequences (post-padding) so that all input sequences match the length of the longest sequence in the dataset. This step is crucial for batch processing during model training, as it allows efficient computation and avoids dimension mismatch errors. Together, the text-

to-sequence conversion and padding steps produce a well-structured, fixed-length input matrix that can be directly fed into the Bidirectional GRU models for training [20].

2.4 Data Splitting

To evaluate the performance and generalization ability of the model, the dataset is divided into two subsets: training data and testing data. This process, known as *data splitting*, is a standard procedure in machine learning and deep learning workflows to prevent overfitting and to assess how well the model performs on unseen data.

In this study, the dataset is split using a ratio of 80:20, where 80% of the data is allocated for training and the remaining 20% for testing. To ensure better generalization, especially given the relatively small and imbalanced nature of the dataset, we adopted a 5-fold cross-validation strategy instead of a simple train-test split. This approach divides the data into five equal parts, using four parts for training and one for testing in each iteration, allowing every sample to be used for both training and validation. This method improves robustness and reduces the risk of overfitting. The training data is used to train the Bidirectional GRU and Bidirectional GRU Tuned models by allowing them to learn patterns and features from labeled examples. Meanwhile, the testing data is kept separate and is only used during the final evaluation phase to measure the model's accuracy, precision, recall, and F1-score on data that the model has never seen before. This approach ensures that the evaluation metrics reflect the model's true predictive capability and not just its ability to memorize the training data.

2.5 Bidirectional GRU

The core model used in this study is the Bidirectional Gated Recurrent Unit (Bidirectional GRU), an enhanced variant of the GRU network that can capture contextual information from both past (previous words) and future (next words) in a sequence. Unlike unidirectional GRU that only processes sequences in a forward direction, Bidirectional GRU employs two GRU layers: one that reads the input from left to right, and another that reads it from right to left. Their outputs are concatenated, enabling the model to learn richer representations of the input data.

Each GRU unit uses gating mechanisms—reset gate and update gate—to control the flow of information. The mathematical formulation of a standard GRU cell is as follows:

$$z_t = \sigma(W_z \cdot x_t + U_z \cdot h_{t-1}) \quad \{(update\ gate)\} \quad (1)$$

$$r_t = \sigma(W_r \cdot x_t + U_r \cdot h_{t-1}) \quad \{(reset\ gate)\} \quad (2)$$

$$\tilde{h}_t = \tanh(W_h \cdot x_t + U_h \cdot (r_t \odot h_{t-1})) \quad \{(candidate\ hidden\ state)\} \quad (3)$$

$$h_t = (1 - z_t) \odot h_{t-1} + z_t \odot \tilde{h}_t \quad \{(final\ hidden\ state)\} \quad (4)$$

Where x_t is the input at time step t , h_{t-1} is the previous hidden state, z_t is the update gate, r_t is the reset gate, $\tilde{h}_t$ is the candidate activation, h_t is the new hidden state, $\odot$ denotes element-wise multiplication, σ is the sigmoid activation function and $\tanh$ is the hyperbolic tangent activation.

2.6 Bidirectional GRU Tuned

To improve the predictive performance of the baseline Bidirectional GRU model, hyperparameter tuning was carried out using the Keras Tuner framework. This tuning process systematically explores the model's hyperparameter space to identify the best configuration that minimizes validation loss while maximizing generalization. The tuning approach adopted in this study uses random search across several key hyperparameters, including the embedding dimension, the number of GRU units, dropout rates, the number of neurons in the dense layer, and the learning rate used during optimization.

The best-performing configuration obtained from this process is summarized in Table 2.

Table 2. Best Hyperparameter Configuration for Bidirectional GRU

Hyperparameter	Value
Embedding Dimension	64
GRU Units	32
Dropout Rate	0.3
Dense Layer Units	128
Learning Rate	0.01

This tuned configuration enables the model to learn more effectively by balancing

complexity and regularization. Mathematically, the hyperparameter tuning process can be represented as an optimization objective:

$$\theta^* = \arg \min_{\theta \in \Theta} \mathcal{L}_{vl}(f(x; \theta)) \quad (5)$$

where θ represents the set of hyperparameters, Θ denotes the hyperparameter search space, $f(x; \theta)$ is the Bidirectional GRU model parameterized by θ , and $\mathcal{L}_{vl}$ is the validation loss, measured using binary cross-entropy. This optimization process allows the model to achieve a validation accuracy of 85.37%, outperforming both the standard GRU and untuned Bidirectional GRU models. These results confirm that hyperparameter tuning plays a crucial role in enhancing the model's ability to generalize to unseen data.

2.6 Evaluation

The performance of each model was assessed using a set of well-established evaluation metrics, namely accuracy, precision, recall, and F1-score. These metrics were calculated specifically for the positive sentiment class, as it represented the majority class within the dataset and served as a key focus of the analysis. Accuracy reflects the overall correctness of the model's predictions, while precision measures the proportion of positive predictions that were actually correct. Recall, on the other hand, indicates the model's ability to identify all actual positive instances. F1-score provides a harmonic balance between precision and recall, especially useful in cases where the class distribution is imbalanced.

To complement these metrics, confusion matrices were also generated to visualize the classification results in greater detail. Each confusion matrix illustrates the number of correctly and incorrectly predicted instances across both positive and negative classes. This enables a deeper understanding of how the model handles misclassifications, particularly in detecting minority class instances.

Based on the evaluation results, the Bidirectional GRU model with hyperparameter tuning consistently outperformed the baseline models across all evaluation criteria. Its improved performance confirms the positive impact of tuning and bidirectional context in enhancing model generalization and reliability in sentiment classification tasks.

3. RESULT

The experimental phase of this study began with the implementation of a baseline

GRU model, serving as the reference architecture for sentiment classification on game reviews. This model was trained using default settings, without any hyperparameter tuning. Although it achieved a respectable accuracy of 78%, it exhibited limitations in generalization—particularly when handling imbalanced class distributions.

To address this, the second experiment applied hyperparameter optimization to the GRU architecture. Using Keras Tuner, key parameters such as embedding size, number of GRU units, dropout rate, and learning rate were fine-tuned to enhance model performance. While the optimized GRU maintained the same overall accuracy (78%) as the baseline, it demonstrated improved training stability and consistency across evaluation metrics.

The third and final experiment focused on combining bidirectional sequence modeling with hyperparameter optimization. A Bidirectional GRU model was trained using the best-tuned configuration. This architecture enabled the model to capture both past and future context in the input sequence, leading to improved representational power. As a result, the Bidirectional GRU Tuned model achieved a higher accuracy of 83%, making it the best-performing approach in this study.

A comparison of evaluation results is presented in Table 3, while the accuracy comparison is visually illustrated in Figure 2.

Table 3 Comparison Of Accuracy

Model	Accuracy	Precision	Recall	F1-Score
GRU (Baseline)	0.78	0.78	1.00	0.88
GRU (Tuned)	0.78	0.78	1.00	0.88
Bidirectional GRU (Tuned)	0.83	0.82	1.00	0.90

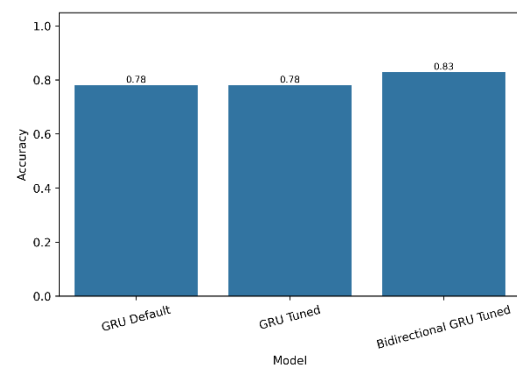


Figure 2. Accuracy Comparison

4. DISCUSSION

The experimental results of this study reveal a clear progression in model performance as architectural enhancements and hyperparameter optimization were introduced. The baseline GRU model, despite being simple and efficient, exhibited limited capability in distinguishing between sentiment classes, particularly when handling imbalanced data. While it achieved a recall of 1.00—indicating all positive reviews were correctly identified—its relatively lower precision and F1-score reflected a tendency to overpredict the majority class.

Introducing hyperparameter tuning to the GRU model resulted in more stable training behavior and consistent metric scores. However, the improvement was marginal, suggesting that architectural limitations still constrained the model's representational power. The tuned GRU model maintained the same accuracy and recall as the baseline, but showed no significant increase in overall classification effectiveness.

The most significant improvement was observed when implementing the Bidirectional GRU with optimized hyperparameters. This configuration not only enhanced the accuracy to 83%, but also improved the precision and F1-score, indicating a better balance between sensitivity and specificity. The bidirectional architecture enabled the model to capture richer context by analyzing input sequences in both forward and backward directions—an advantage particularly beneficial for short and informal reviews typical of user-generated game content.

Furthermore, the use of hyperparameter tuning in combination with bidirectional modeling proved to be a key factor in achieving better generalization, as evidenced by the improved performance across all evaluation metrics. These findings confirm that careful tuning and architectural design can significantly influence the effectiveness of deep learning models in sentiment analysis tasks.

5. CONCLUSION

This study has demonstrated a structured approach to improving sentiment classification performance on game review data using GRU-based models. Starting from a baseline GRU architecture, we observed that while the model could achieve high recall, it lacked precision and failed to generalize well under class imbalance. Incorporating hyperparameter tuning helped stabilize training and marginally improved model performance,

but the architectural limitations of unidirectional GRU remained a bottleneck.

The best results were obtained by implementing a Bidirectional GRU with tuned hyperparameters. This model achieved the highest accuracy and F1-score, indicating a more balanced and reliable classification capability. The bidirectional structure allowed the model to capture contextual information more effectively, while hyperparameter optimization contributed to improved learning efficiency and generalization. The overall findings affirm that model architecture and parameter tuning play a critical role in enhancing the effectiveness of deep learning approaches for sentiment analysis.

For future work, several directions can be explored to further strengthen the model. These include addressing class imbalance through techniques such as oversampling or cost-sensitive training, experimenting with pretrained word embeddings such as FastText or BERT-based representations, and incorporating attention mechanisms to highlight key parts of the review text. Additionally, expanding the dataset and applying domain adaptation methods may help improve model robustness across different genres of user-generated content. Such enhancements will contribute to more accurate and context-aware sentiment analysis systems, particularly in domains where informal and brief textual expressions are prevalent.

6. SUGGESTION

Based on the findings of this study, several recommendations are proposed to guide future research in sentiment analysis of short-form review data. A primary area for enhancement is the management of class imbalance, which could be addressed through techniques such as oversampling, applying class-weighting strategies, or implementing data augmentation to increase the representation of minority classes. While the Bidirectional GRU model, optimized through hyperparameter tuning, demonstrated strong performance, future studies may benefit from incorporating attention mechanisms to enable the model to more effectively identify and prioritize contextually important words within a sentence.

Another promising avenue involves the integration of pretrained word embeddings—such as FastText, Word2Vec, or IndoBERT—which can enrich semantic understanding and improve classification accuracy, particularly for texts written in Bahasa Indonesia. Additionally, expanding the dataset to encompass a broader range of review types and user expressions could

enhance the model's generalizability. Finally, deploying the proposed model in real-world applications—such as automated feedback systems or sentiment monitoring tools in game development—would not only validate its practical utility but also provide opportunities for further refinement and adaptation.

REFERENSI

- [1] M. Ranjan, S. Tiwari, A. M. Sattar, and N. S. Tatkar, "A New Approach for Carrying Out Sentiment Analysis of Social Media Comments Using Natural Language Processing," *Eng. Proc.*, vol. 59, no. 1, p. 181, 2024.
- [2] X. Tong, I. Willcock, and Y. Sun, "Unravelling Player's Insights: A Comparative Analysis of Topic Modelling Techniques on Game Reviews and Video Game Developers' Perspectives," *IEEE Trans. Games*, 2024.
- [3] T. Guzvinecz and J. Sz\Hucs, "Textual Analysis of Virtual Reality Game Reviews.," *Infocommunications J.*, vol. 16, 2024.
- [4] W. Erpurini, A. G. Putrada, N. Alamsyah, S. F. Pane, and M. N. Fauzan, "Confirmatory Factor Analysis for The Impact of Students' Social Medial on University Digital Marketing," in *2023 International Conference on Computer Science, Information Technology and Engineering (ICCoSITE)*, IEEE, 2023, pp. 615–620.
- [5] L. Yue, W. Chen, X. Li, W. Zuo, and M. Yin, "A survey of sentiment analysis in social media," *Knowl. Inf. Syst.*, vol. 60, pp. 617–663, 2019.
- [6] L. Wang, M. Han, X. Li, N. Zhang, and H. Cheng, "Review of classification methods on unbalanced data sets," *Ieee Access*, vol. 9, pp. 64606–64628, 2021.
- [7] N. Alamsyah, A. P. Kurniati, and others, "Airfare Fluctuation Analysis with Event and Sentiment Features by Stacking Ensemble Model," in *2024 Ninth International Conference on Informatics and Computing (ICIC)*, IEEE, 2024, pp. 1–6.
- [8] T. V. Subbamma, N. Mantravadi, A. R. Shalom, N. Tajne, G. Sreenivasulu, and V. S. Goud, "Natural Language Processing for Sentiment Analysis Techniques and Applications," *Metall. Mater. Eng.*, vol. 31, no. 3, pp. 194–200, 2025.
- [9] A. G. Putrada, N. Alamsyah, I. D. Oktaviani, and M. N. Fauzan, "A hybrid genetic algorithm-random forest regression method for optimum driver selection in online food delivery," *J. Ilm. Tek. Elektro Komput. Dan Inform. JITEKI*, vol. 9, no. 4, pp. 1060–1079, 2023.
- [10] A. A. Kumar, "Semantic memory: A review of methods, models, and current challenges," *Psychon. Bull. Rev.*, vol. 28, no. 1, pp. 40–80, 2021.
- [11] N. Alamsyah, T. P. Yoga, B. Budiman, and others, "IMPROVING TRAFFIC DENSITY PREDICTION USING LSTM WITH PARAMETRIC ReLU (PReLU) ACTIVATION," *JITK J. Ilmu Pengetah. Dan Teknol. Komput.*, vol. 9, no. 2, pp. 154–160, 2024.
- [12] J. V. Tembhurne and T. Diwan, "Sentiment analysis in textual, visual and multimodal inputs using recurrent neural networks," *Multimed. Tools Appl.*, vol. 80, no. 5, pp. 6871–6910, 2021.
- [13] A. Kaur, V. Kukreja, and D. Kundra, "Echoes of Emotion: Enhancing Sentiment Analysis with CNN-GRU Synergy," in *2024 International Conference on Big Data Analytics in Bioinformatics (DABCon)*, IEEE, 2024, pp. 1–7.
- [14] J. Liu, Y. Sun, Y. Zhang, and C. Lu, "Research on Online Review Information Classification Based on Multimodal Deep Learning," *Appl. Sci.*, vol. 14, no. 9, p. 3801, 2024.
- [15] N. Alamsyah, B. Budiman, T. P. Yoga, and R. Y. R. Alamsyah, "COMPARISON LINEAR REGRESSION AND RANDOM FOREST MODELS FOR PREDICTION OF UNDERGROUND DROUGHT LEVELS IN FOREST FIRES," *J. Techno Nusa Mandiri*, vol. 21, no. 2, pp. 81–86, 2024.
- [16] J. B. L. Bernardo, A. Taparugssanagorn, H. Miyazaki, B. M. Pati, and U. Thapa, "Robust Human Activity Recognition for Intelligent Transportation Systems Using Smartphone Sensors: A Position-Independent Approach," *Appl. Sci.*, vol. 14, no. 22, p. 10461, 2024.
- [17] trilism, "EXCEL SENTIMENT FOR GAME REVIEW." 2025.
- [18] N. Alamsyah, A. P. Kurniati, and others, "A novel airfare dataset to predict travel agent profits based on dynamic pricing," in *2023 11th International Conference on Information and Communication Technology (ICoICT)*, IEEE, 2023, pp. 575–581.

- [19] N. Alamsyah, A. P. Kurniati, and others, "Event Detection Optimization Through Stacking Ensemble and BERT Fine-tuning For Dynamic Pricing of Airline Tickets," *IEEE Access*, 2024.
- [20] I. S. A. Khaleq, "Enhancing Clickbait Detection Through Deep Learning and Language-specific Analysis in English and Kurdish," PhD Thesis, Polytechnic University, 2023.