

Original Research

Sentiment Analysis of the Joko Widodo Diploma Controversy Using NLP and Multi-Layer Perceptron

Herdiansyah ^{a*}, Nazori Suhandi ^b, Dwi Asa Verano ^c

^{a,b,c} Department of Informatics, Faculty of Computer Science and Science, Universitas Indo Global Mandiri, Palembang, Indonesia

*Corresponding Author: 12021110061@students.uigm.ac.id

ABSTRACT

The controversy regarding the authenticity of former Indonesian President Joko Widodo's diploma has become a public issue widely discussed on social media, particularly on YouTube. YouTube comments contain diverse and unstructured public opinions expressed in informal language, making manual analysis difficult. Therefore, this study aims to analyze public sentiment toward the Joko Widodo diploma controversy using Natural Language Processing (NLP) with the Multi-Layer Perceptron (MLP) algorithm. A total of 34,965 YouTube comments were collected and processed through cleaning, case folding, normalization, tokenization, stopword removal, and stemming. The processed text was transformed into numerical representations using Term Frequency–Inverse Document Frequency (TF-IDF). The MLP model then classified the comments into three sentiment classes: positive, negative, and neutral. Model performance was evaluated using a confusion matrix and accuracy, precision, recall, and F1-score metrics. The results show that the model achieved 91.56% accuracy on the testing data, with strong precision, recall, and F1-score across the three sentiment classes. These findings indicate that the proposed NLP-MLP approach can effectively classify public sentiment and support automated monitoring of public opinion on political issues discussed through YouTube comments.

Copyright © 2026 by Author(s). This is an open-access article distributed under the terms of the Creative Commons Attribution 4.0 International License (CC BY 4.0).

How to Cite:

Herdiansyah, N. Suhandi, and D. A. Verano, "Sentiment Analysis of the Joko Widodo Diploma Controversy Using NLP and Multi-Layer Perceptron," *J. Kecerdasan Buatan dan Teknol. Inf.*, vol. 5, no. 3, pp. 677–685, 2026.
Doi: <https://doi.org/10.69916/jkbt.v5i3.591>

KEYWORDS

Multi-Layer Perceptron
Natural Language Processing
Sentiment Analysis
TF-IDF
YouTube

ARTICLE INFO

Received: July 18, 2026

Accepted: August 28, 2026

Available Online: September 18, 2026

1 INTRODUCTION

Social media has transformed the way people communicate, interact, and share information in the digital era. It enables individuals to express opinions, collaborate, and establish virtual relationships with others [1]. Among various social media platforms, YouTube is an important source of public opinion because its comment section allows users to express support, criticism, and responses to uploaded content [2]. However, YouTube comments are generally unstructured and contain informal language, abbreviations, emojis, and sarcasm, making them challenging to analyze without computational techniques [3]. Recent international studies also show that YouTube comments can be effectively analyzed using machine learning and deep learning approaches, although model performance varies according to the dataset, topic, and representation method [4].

One issue that has attracted significant public attention in Indonesia is the controversy surrounding the alleged fake diploma of former President Joko Widodo. Although official clarification has been provided by the relevant university and government institutions, public discussions continue to develop across social media [5]. The wide range of opinions expressed by users reflects differing perspectives toward public figures, educational institutions, and government institutions [6].

Sentiment analysis is an effective approach for identifying public opinion from textual data. As a branch of text mining, sentiment analysis extracts information from text to determine whether opinions are positive, negative, or neutral [7]. Natural Language Processing (NLP) enables computers to process and analyze unstructured human language [8]. Previous studies have applied conventional machine learning methods, such as Support Vector Machine, to sentiment analysis of public opinions, while other research has investigated the Multi-Layer Perceptron with TF-IDF features [9]. Recent international research has also examined YouTube sentiment using deep learning models for conflict-related topics, conventional machine learning models for YouTube comments, and transformer-based models for video comments [10]. These studies demonstrate the effectiveness of automated sentiment analysis, but they focus on different topics, datasets, or model architectures and do not specifically examine the Joko Widodo diploma controversy using MLP on Indonesian YouTube comments [11]. This study therefore applies the Multi-Layer Perceptron (MLP), an Artificial Neural Network (ANN) algorithm with multiple hidden layers, to classify sentiment in YouTube comments related to the Joko Widodo diploma controversy [12]. The research specifically combines NLP preprocessing and TF-IDF feature representation with MLP classification to identify positive, negative, and neutral public sentiment [13]. This approach addresses the identified research gap by focusing on a specific Indonesian political controversy and evaluating MLP performance on a large YouTube comment dataset [14]. The findings are expected to provide a data-driven understanding of public perceptions and demonstrate the applicability of NLP and MLP for political sentiment analysis in Indonesia [15].

2 RESEARCH METHODS

2.1 Research Stages

The research workflow used in this study consists of several sequential stages, starting from problem identification and ending with the evaluation and discussion of the results, as illustrated in Figure 1.

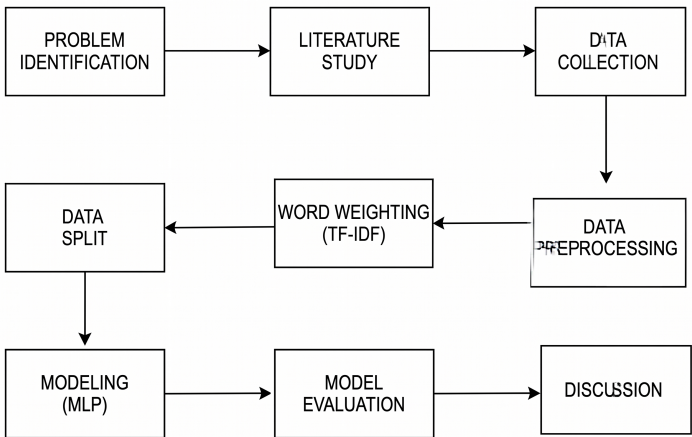


Figure 1. Research stages diagram

2.2 Problem Identification

The data used in this study were collected from public comments posted on YouTube videos discussing the controversy surrounding Joko Widodo’s diploma. The selected videos were obtained from several YouTube channels that frequently discuss political and public issues, including ILC, Rakyat Bersuara iNews, Indonesia Kita Garuda TV, Catatan Demokrasi tvOne, and Dua Arah Kompas TV. The comments were collected using a web crawling process implemented in Python through the Google Colab environment. The data collection process used the keyword “Joko Widodo diploma controversy in Indonesia” to identify relevant YouTube videos and comments. A total of 34,965 comments were collected. The dataset contained comment text and supporting metadata, including video ID, author, and publication date. Only the comment text was used as the primary variable for sentiment classification, while the remaining attributes were retained as supporting information. After collection, the dataset was cleaned by removing duplicate comments, spam, advertisements, and irrelevant symbols or characters. The resulting dataset was then prepared for sentiment labeling and subsequent preprocessing stages.

2.3 Data Collection

Data collection is a crucial initial stage of this research. The data were obtained from public comments posted on the YouTube platform, particularly on videos discussing the controversy surrounding Joko Widodo's diploma in Indonesia. YouTube was selected as the data source due to the high level of user interaction within its comment sections, which can directly represent public opinion.

In this study, the data were collected from several YouTube channels that frequently discuss political issues and public affairs, namely ILC, Rakyat Bersuara iNews, Indonesia Kita Garuda TV, Catatan Demokrasi tvOne, and Dua Arah Kompas TV. The research dataset was obtained through a YouTube comment crawling process on videos related to the Joko Widodo diploma controversy using the Python programming language in Google Colab. The collected data were stored in Microsoft Excel format, undergoing initial cleaning to remove duplicate comments, spam, and irrelevant symbols. The cleaned dataset was then manually labeled before being split into training and testing subsets.

2.4 Sentiment Labeling

Sentiment labeling was performed to classify each YouTube comment into one of three categories: positive, negative, or neutral. The labeling process was conducted manually based on the sentiment expressed in the content of each comment. The manual labeling was performed before model training so that each comment had a sentiment label that could be used as the target variable during classification.

Comments expressing approval, support, appreciation, or favorable opinions were categorized as positive. Comments containing criticism, rejection, dissatisfaction, or unfavorable opinions were categorized as negative. Comments that primarily conveyed factual information, questions, or opinions without a clear positive or negative orientation were categorized as neutral [16]. The labeled dataset was subsequently divided into training and testing data using an 80:20 ratio (27,972 training comments and 6,993 testing comments).

2.5 Data Preprocessing

The collected comments underwent several preprocessing stages to transform unstructured text into a suitable representation for classification. The preprocessing stages consisted of cleaning, case folding, normalization, tokenization, stopword removal, and stemming.

Cleaning removed punctuation, numbers, emojis, and irrelevant characters. Case folding converted all text into lowercase letters. Normalization standardized non-standard words, abbreviations, and informal social media expressions. Tokenization divided the text into individual words. Stopword removal eliminated common words less relevant to sentiment classification, while stemming reduced words to their root forms. These preprocessing stages were applied to reduce noise and produce a consistent textual representation before feature extraction.

2.6 Word Weighting Using TF-IDF

After preprocessing, the textual data were transformed into numerical feature vectors using the Term Frequency–Inverse Document Frequency (TF-IDF) method. TF measures the frequency of a term within a document, whereas IDF provides higher weights to terms that occur less frequently across the document collection. The combination of TF and IDF emphasizes informative terms while reducing the influence of common terms [17]. The resulting TF-IDF vectors served as the input features for the MLP classification model.

2.7 Data Splitting

The dataset was divided using an 80:20 ratio. Of the 34,965 opinion data instances collected, 27,972 (80%) were used as the training dataset to enable the model to learn patterns across positive, negative, and neutral sentiments. The remaining 6,993 data instances (20%) were used as the testing dataset to evaluate classification performance on unseen data. This 80:20 split provided an appropriate balance between training and testing processes, minimizing overfitting risks [18].

2.8 Modeling Using Multi-Layer Perceptron

The Multi-Layer Perceptron (MLP) was selected as the classification algorithm because it is capable of learning complex non-linear relationships between numerical input features and sentiment classes. In this study, the MLP receives TF-IDF feature vectors as input and predicts one of three sentiment classes: positive, neutral, or negative.

a. Multi-Layer Perceptron Model Architecture

The proposed MLP architecture consists of an input layer, two hidden layers, and an output layer. The first hidden layer contains 128 neurons, while the second hidden layer contains 64 neurons. The output layer consists of three neurons corresponding to the three sentiment classes and uses the Softmax activation function to generate class probabilities. The model architecture is illustrated in Figure 2.

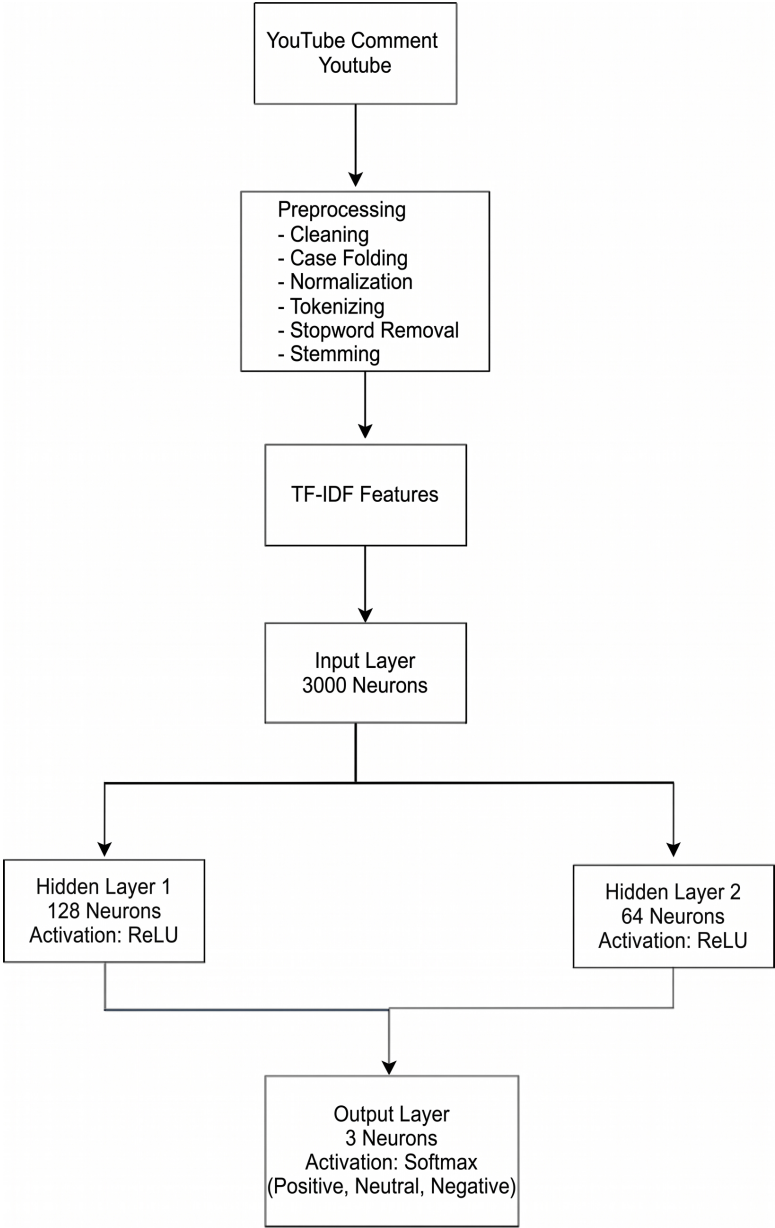


Figure 2. Multi-Layer Perceptron model architecture

The MLP training process used the Adam optimizer and Binary Cross-Entropy loss function. The learning rate, number of epochs, and batch size were configured during model training to ensure convergence, resulting in the final hyperparameter configuration yielding the optimal testing accuracy.

b. MLP Training Process

The model was trained using the Backpropagation algorithm. During forward propagation, TF-IDF feature vectors passed through the input layer and two hidden layers to generate predicted probabilities. Predictions were compared with actual sentiment labels using the loss function, and resulting errors were backpropagated to update weights and biases via the Adam optimizer. This process repeated across epochs until training was completed.

c. Algorithm Evaluation Procedure

The trained MLP model was evaluated using the testing dataset. A confusion matrix compared actual sentiment labels with predicted labels, measuring performance across accuracy, precision, recall, and F1-score metrics.

d. Model Evaluation Method

Model performance was evaluated using a Confusion Matrix to assess the model’s accuracy and reliability in classifying user comments. The Confusion Matrix compares actual sentiment classes with predicted classes, allowing classification errors to be clearly identified across positive, negative, and neutral categories.

3 RESULTS AND DISCUSSION

3.1 Data Collection Results

This study utilized secondary data consisting of user comments collected from YouTube regarding the Joko Widodo diploma controversy. Data were gathered from December 1, 2025, to December 20, 2025, using the YouTube Data API in Python. A total of 34,965 authentic comments were retrieved across news channels. The primary attribute used for sentiment analysis was the ‘comment’ text, while ‘video_id’, ‘author’, and ‘published’ timestamps served as supporting metadata.

3.2 Sentiment Labeling Results

Manual sentiment labeling categorized the dataset into three distinct classes: positive, negative, and neutral. The distribution of sentiment labels across the dataset is shown in Figure 3.

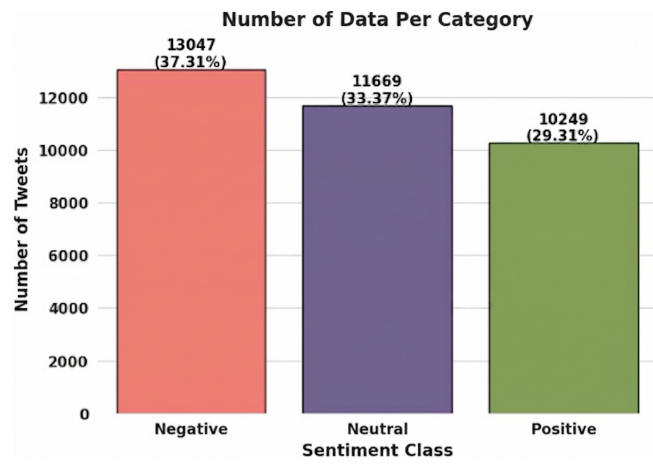


Figure 3. Sentiment labeling distribution results

3.3 Feature Extraction using TF-IDF

Following preprocessing (cleaning, case folding, normalization, tokenization, stopwords removal, and stemming), textual data were transformed into numerical TF-IDF feature matrices. The TF-IDF method effectively assigned higher weights to domain-specific informative words while suppressing non-discriminative terms across the corpus.

3.4 Data Splitting Execution

The dataset of 34,965 records was partitioned into 27,972 training samples (80%) and 6,993 testing samples (20%). The training set allowed the model to capture feature distributions across all three sentiment classes, while the unseen testing set provided an objective benchmark for generalization.

3.5 Model Evaluation

The MLP model achieved an accuracy of 91.56% on the testing dataset, compared with 98.54% on the training dataset. Although a minor performance drop occurred on unseen testing data, the high testing accuracy confirms that the model maintained strong generalization capabilities.

a. Confusion Matrix for the Training Data

The confusion matrix for the training data is presented in Figure 4.

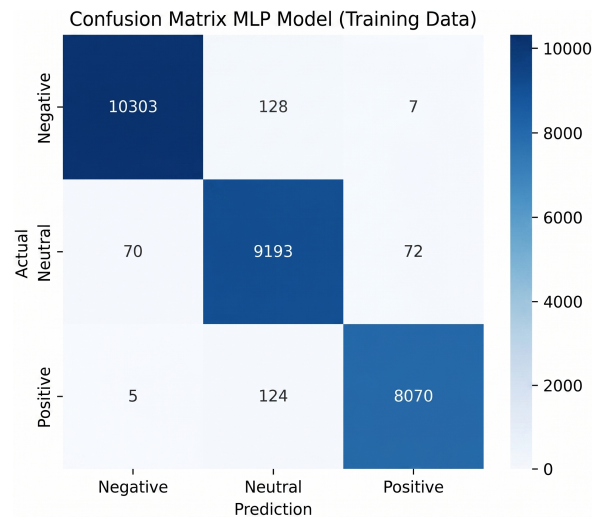


Figure 4. Confusion matrix for the training dataset

Based on the training confusion matrix, the model correctly classified 10,303 negative comments, 9,193 neutral comments, and 8,070 positive comments. Misclassifications remained minimal: for the Negative class, 128 were misclassified as Neutral and 7 as Positive; for Neutral, 70 were misclassified as Negative and 72 as Positive; for Positive, 124 were misclassified as Neutral and 5 as Negative.

b. Confusion Matrix for the Testing Data

The confusion matrix for the testing dataset is presented in Figure 5.

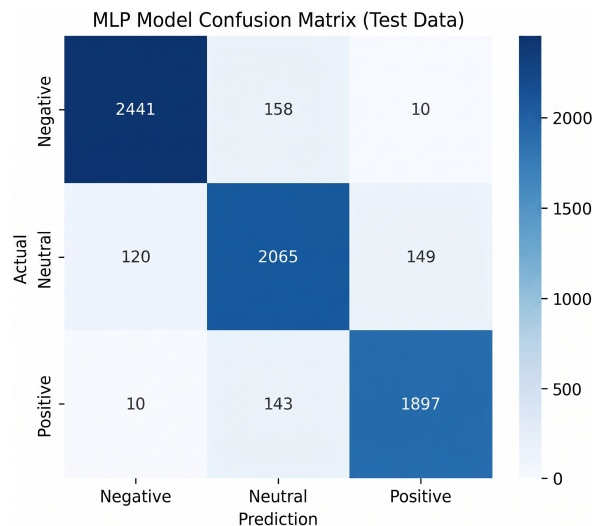


Figure 5. Confusion matrix for the testing dataset

The testing confusion matrix shows that the model correctly classified 2,441 Negative comments, 2,065 Neutral comments, and 1,897 Positive comments. For the Negative class, 158 comments were misclassified as Neutral and 10 as Positive. For Neutral, 120 were misclassified as Negative and 149 as Positive. For Positive, 143 were misclassified as Neutral and 10 as Negative.

The performance metrics for each sentiment class on the training data are calculated as follows:

Positive Class Calculation (Training):

$$\text{Precision}_{\text{positive}} = \frac{8070}{8070 + 79} = \frac{8070}{8149} = 0.9903 \text{ (99.03\%)} \quad (1)$$

$$\text{Recall}_{\text{positive}} = \frac{8070}{8070 + 129} = \frac{8070}{8199} = 0.9842 \text{ (98.42\%)} \quad (2)$$

$$\text{F1-Score}_{\text{positive}} = 2 \times \frac{0.9842 \times 0.9903}{0.9842 + 0.9903} = 0.9871 \text{ (98.71\%)} \quad (3)$$

Neutral Class Calculation (Training):

$$\text{Precision}_{\text{neutral}} = \frac{9193}{9193 + 252} = \frac{9193}{9445} = 0.9733 \text{ (97.33\%)} \quad (4)$$

$$\text{Recall}_{\text{neutral}} = \frac{9193}{9193 + 142} = \frac{9193}{9335} = 0.9847 \text{ (98.47\%)} \quad (5)$$

$$\text{F1-Score}_{\text{neutral}} = 2 \times \frac{0.9847 \times 0.9733}{0.9847 + 0.9733} = 0.9789 \text{ (97.89\%)} \quad (6)$$

Negative Class Calculation (Training):

$$\text{Precision}_{\text{negative}} = \frac{10303}{10303 + 75} = \frac{10303}{10375} = 0.9930 \text{ (99.30\%)} \quad (7)$$

$$\text{Recall}_{\text{negative}} = \frac{10303}{10303 + 135} = \frac{10303}{10465} = 0.9845 \text{ (98.45\%)} \quad (8)$$

$$\text{F1-Score}_{\text{negative}} = 2 \times \frac{0.9845 \times 0.9930}{0.9845 + 0.9930} = 0.9887 \text{ (98.87\%)} \quad (9)$$

Overall Accuracy Calculation (Training):

$$\text{Accuracy} = \frac{10303 + 9193 + 8070}{27970} = \frac{27556}{27970} = 0.9854 \text{ (98.54\%)} \quad (10)$$

The summary of model performance evaluation across training and testing datasets is presented in Table 1.

Table 1. Model Performance Evaluation Results (Training / Testing)

Sentiment Class	Precision (%)	Recall (%)	F1-Score (%)	Data Subset
Negative	99.30 / 94.94	98.45 / 93.56	98.87 / 94.23	(Train / Test)
Neutral	97.33 / 87.27	98.47 / 88.27	97.89 / 87.85	(Train / Test)
Positive	99.03 / 92.26	98.42 / 92.06	98.71 / 92.17	(Train / Test)
Overall Accuracy		98.54 %		Training Data
Overall Accuracy		91.56 %		Testing Data

3.6 Discussion

The results demonstrate that combining NLP preprocessing, TF-IDF feature extraction, and MLP classification effectively categorizes public sentiment in a large YouTube comment dataset. The model achieved a testing accuracy of 91.56%, with testing F1-scores reaching 94.23% for Negative, 92.17% for Positive, and 87.85% for Neutral classes.

Compared with prior literature, the proposed model exhibits strong competitive performance. Arasy et al. [19] applied MLP with TF-IDF to classify political tweets (300 instances), obtaining an accuracy of 66.67%. The higher accuracy achieved in this study (91.56%) is likely attributed to the substantially larger dataset size (34,965 comments), domain-specific preprocessing, and optimized neural network hidden layers. When compared with deep sequence models on YouTube comments, such as CNN-Bi-LSTM with Word2Vec (95.73% accuracy) [16] or Bagging Classifiers (95.8% accuracy) [20], the proposed MLP model yields slightly lower accuracy but offers significantly reduced computational complexity and faster inference times.

The lower performance observed in the Neutral class (F1-score 87.85%) represents an important finding. Neutral comments in political discussions frequently consist of factual queries, news quotes, or subtle expressions lacking explicit polarity keywords. Because TF-IDF measures term frequency rather than deep semantic context, lexical overlap occurs between neutral statements and weakly positive or negative comments. Nevertheless, the overall testing accuracy of 91.56% confirms that the NLP-MLP architecture is robust for automated social media public opinion monitoring.

4 CONCLUSION

This study implemented Natural Language Processing (NLP) and the Multi-Layer Perceptron (MLP) algorithm to analyze public sentiment regarding the Joko Widodo diploma controversy on YouTube. Using 34,965 user comments processed through standard text preprocessing and TF-IDF feature extraction, the MLP model achieved an overall testing accuracy of 91.56%. The model exhibited outstanding classification capabilities for

polar sentiments (Negative F1-score of 94.23% and Positive F1-score of 92.17%), while showing slightly reduced precision on ambiguous neutral statements (Neutral F1-score of 87.85%).

These findings demonstrate that the NLP-MLP framework is a viable, automated tool for monitoring public opinion on digital political discourse. Future work can extend this research by incorporating contextual language models (e.g., IndoBERT) and word embeddings to improve sentiment disambiguation in neutral and sarcastic comments.

ACKNOWLEDGMENTS

The authors express their gratitude to Universitas Indo Global Mandiri for providing institutional support and computational resources throughout this research project.

DECLARATIONS

Author Contributions:

Herdiansyah: Conceptualization, Methodology, Software (Prototyping), Investigation, Writing – Original Draft.

Nazori Suhandi: Validation, Formal Analysis, Data Curation, Writing – Review & Editing.

Dwi Asa Verano: Project Administration, Supervision, Resources, Writing – Review & Editing.

Funding:

The authors received no financial support for the research, authorship, and publication of this article.

Conflicts of Interest:

The authors declare that they have no competing financial or personal conflicts of interest regarding this publication.

Data Availability:

The data supporting the findings of this study (including preprocessed text and model parameters) are available from the corresponding author upon reasonable request.

AI Use Statement:

During manuscript preparation, the authors utilized ChatGPT solely for language refinement and readability improvement. All scientific content, data analyses, experimental models, and conclusions were originally developed and verified by the authors, who assume full responsibility for the published content.

Additional Information:

No additional information is available for this paper.

REFERENCES

- [1] I. Huda, "Implementasi Natural Language Processing (NLP)," *J. Nas. Tek. Elektro Dan Teknol. Inf.*, pp. 15–28, 2021.
- [2] M. R. Anggana, "Peran Natural Language Processing (NLP) dalam Membantu Manusia," *Kata.Ai*, vol. 11, pp. 62–68, 2022.
- [3] A. Liyih, S. Anagaw, M. Yibeyin, and Y. Tehone, "Sentiment analysis of the Hamas–Israel war on YouTube comments using deep learning," *Sci. Rep.*, vol. 14, pp. 1–9, 2024, doi: 10.1038/s41598-024-63367-3.
- [4] A. P. Widayari, "Sentiment Analysis of YouTube Comments on the 2025 DPR RI Demonstration Using Machine Learning," *J. Comput. Sci.*, vol. 10, no. 2, 2026.
- [5] A. Angdresey, L. Sitanayah, I. Lucky, and H. Tangka, "Sentiment Analysis for Political Debates on YouTube Comments using BERT Labeling, Random Oversampling, and Multinomial Naïve Bayes," *IEEE Access*, 2025.
- [6] N. Mahfudza and M. Ikhsan, "Sentiment Analysis of Youtube Comments on Indonesian Presidential Candidates in 2024 using Naïve Bayes Classifier Method," *JURIKOM (Jurnal Ris. Komput.)*, vol. 12, no. 2, pp. 140–148, 2025, doi: 10.30865/jurikom.v12i2.8538.

- [7] I. Aditya and A. Putra, "Analisis Komentar Youtube Terhadap Polemik Ijazah Presiden Ke 7 Indonesia Menggunakan Support Vector Machine," *Bull. Comput. Sci. Res.*, vol. 6, no. 1, pp. 445–457, 2025, doi: 10.47065/bulletincsr.v6i1.883.
- [8] U. Bert, N. Sholihah, F. F. Abdulloh, and M. Rahardi, "Sentiment Analysis on KPU Performance Post-2024 Election via YouTube Comments," *J. Inf. Technol.*, vol. 8, no. 4, pp. 2222–2232, 2024.
- [9] C. Zuriana, D. Iskandar, and M. Idham, "OPINI PUBLIK TERHADAP DEBAT CAPRES 2024: ANALISIS SENTIMEN DALAM KOMENTAR LIVE YOUTUBE KPU RI," *J. Media Inf.*, vol. 13, no. 2, pp. 472–485, 2024.
- [10] A. Info, "SENTIMENT ANALYSIS OF COMMENTS ON INDONESIAN POLITICAL SPEECH VIDEOS ON YOUTUBE USING MACHINE LEARNING," *J. Technol.*, vol. 7, no. 2, pp. 34–44, 2025.
- [11] A. A. Bastian and A. Perdana, "Sentiment Analysis of Public Comments on YouTube Regarding the Inaugural Speech of the 8th President of Indonesia Using VADER and BERT Methods," *Int. J. Polit. Sci.*, vol. 5, no. 2, pp. 126–140, 2025.
- [12] A. B. T. Y. Prawira and F. A. Tyas, "SENTIMENT ANALYSIS OF THE NATIONAL MANDATE PARTY (PAN) IN YOUTUBE COMMENTS USING THE TF-IDF AND COSINE SIMILARITY APPROACHES," *Metadata J.*, vol. 8, no. 1, pp. 111–125, 2026, doi: 10.47652/metadata.v8i1.937.
- [13] H. A. Amaliaputri, A. P. Wijaya, and D. Salim, "Sentiment Analysis of YouTube Comments on Indonesian 2024 Presidential Candidate Talk Show," *Procedia Comput. Sci.*, vol. 56, no. 7, pp. 2697–2703, 2026.
- [14] F. A. Aziz and L. S. Harahap, "Sentiment Analysis Regarding the Indonesian House of Representatives Rejecting the Constitutional Court Decision from Social Media Using Naive Bayes," *J. Soc. Media Res.*, vol. 10, no. 1, pp. 31–37, 2025.
- [15] B. Valentino and I. M. K. Karo, "Pemahaman Dasar Matriks Dan Aplikasinya Dalam Kehidupan Sehari-Hari," *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 9, no. 3, pp. 5381–5385, 2025, doi: 10.36040/jati.v9i3.14183.
- [16] P. Agusia, M. Uli, A. Manurung, V. Calista, and V. C. Mawardi, "Pemanfaatan Word Cloud Pada Analisis Sentimen Dalam Menggali Persepsi Publik," *J. Serba Inf.*, pp. 25–30, 2024.
- [17] D. Septiani and I. Isabela, "Analisis term frequency inverse document frequency (tf-idf) dalam temu kembali informasi pada dokumen teks," *J. Teknol. Inf.*, vol. 25, pp. 81–88, 2023.
- [18] F. Refindha, A. Harianto, Z. Alawi, and I. Aristia, "PENGARUH KOMPOSISI SPLIT DATA PADA AKURASI KLASIFIKASI PENDERITA DIABETES MENGGUNAKAN MACHINE LEARNING," *J. Data Sci.*, vol. 8, no. 1, pp. 36–44, 2025.
- [19] I. A. Arasy, S. Agustian, and L. Handayani, "Klasifikasi Sentimen Menggunakan Metode Multilayer Perceptron dengan Fitur TF-IDF," *MALCOM: Indones. J. Mach. Learn. Comput. Sci.*, vol. 5, no. 2, pp. 908–919, 2025.
- [20] S. Shevira, I. M. A. D. Suarjaya, and P. W. Buana, "Pengaruh Kombinasi dan Urutan Pre-Processing pada Tweets Bahasa Indonesia," *JITTER J. Ilm. Teknol. Dan Komput.*, vol. 3, no. 2, p. 1074, 2022, doi: 10.24843/jtrti.2022.v03.i02.p06.