

Hybrid IndoBERT-CNN Model for Sentiment Analysis: A Case Study on Educational Public Figures on TikTok

Indar Muslimah Kusmadi

Universitas Tarakanita
Universitas Tarakanita, Jakarta Timur, Indonesia

e-mail : indar@utarki.ac.id

Abstract

The growing use of TikTok as a platform for educational communication by academic leaders has created an urgent need for scalable, automated tools to monitor public sentiment a task that is impractical to perform manually across thousands of comments. Conventional lexicon-based methods are further ill-suited for the informal, code-switched, and emoji rich nature of Indonesian social media text. This study addresses these gaps by analyzing public sentiment toward educational content shared by the Rector of President University on TikTok, positioning the platform as both a learning medium and a tool for institutional branding. A dataset of 2,445 user comments was collected via web scraping, of which 2,000 Indonesian-language comments were retained after preprocessing. Sentiments were labeled into positive, neutral, and negative categories and classified using a hybrid deep learning architecture integrating IndoBERT for contextual embeddings with a Convolutional Neural Network (CNN) for local n-gram feature extraction. The hybrid design was specifically chosen because IndoBERT alone captures global contextual semantics but lacks sensitivity to local informal patterns prevalent in short social media texts, while CNN complements this weakness through convolutional filters—yielding stronger performance than either model achieves independently. The model was trained with a learning rate of 2×10^{-5} , batch size of 32, and 10 epochs, achieving an accuracy of 87.3% with a macro-average F1-score of 0.85, outperforming baseline models including SVM (72.1%), LSTM (78.4%), and standalone IndoBERT (83.2%). The analysis revealed that 40% of comments expressed positive sentiment, indicating favourable audience engagement, while neutral (35%) and negative (25%) comments highlighted areas for improvement in communication style. These findings demonstrate the potential of TikTok as an effective channel for academic leaders to strengthen institutional visibility and public trust in higher education.

Keywords: Sentiment Analysis; IndoBERT; CNN; Deep Learning; TikTok; University Branding

Abstrak

Meningkatnya penggunaan TikTok sebagai platform komunikasi edukatif oleh pemimpin akademik menciptakan kebutuhan mendesak akan alat otomatis yang skalabel untuk memantau sentimen publik tugas yang tidak praktis dilakukan secara manual terhadap ribuan komentar. Metode berbasis leksikon konvensional juga kurang mampu menangani karakteristik bahasa informal, alih kode, dan emoji yang lazim dalam teks media sosial Indonesia. Penelitian ini menjawab kesenjangan tersebut dengan menganalisis sentimen publik terhadap konten pendidikan yang dibagikan oleh Rektor Universitas Presiden di TikTok. Dataset sebanyak 2.445 komentar dikumpulkan melalui web scraping, dan 2.000 komentar berbahasa Indonesia dipertahankan setelah pra-proses. Sentimen diberi label ke dalam kategori positif, netral, dan negatif, kemudian diklasifikasikan menggunakan arsitektur pembelajaran mendalam hibrida yang mengintegrasikan IndoBERT untuk embedding kontekstual dengan Convolutional Neural Network (CNN) untuk ekstraksi fitur n-gram lokal. Model dilatih dengan learning rate 2×10^{-5} , batch size 32, dan 10 epoch, mencapai akurasi 87,3% dengan macro-average F1-

score sebesar 0,85, melampaui model baseline termasuk SVM (72,1%), LSTM (78,4%), dan IndoBERT standalone (83,2%). Analisis menunjukkan bahwa 40% komentar mengekspresikan sentimen positif, sementara komentar netral (35%) dan negatif (25%) menyoroti area yang perlu ditingkatkan dalam gaya komunikasi. Temuan ini menunjukkan potensi TikTok sebagai saluran efektif bagi pemimpin akademik untuk memperkuat visibilitas institusi dan kepercayaan publik terhadap pendidikan tinggi.

Kata Kunci: Analisis Sentimen; IndoBERT; CNN; Deep Learning; TikTok; Branding Universitas

INTRODUCTION

The rapid expansion of internet access in Indonesia has accelerated the adoption of social media platforms as primary spaces for communication and information sharing. Approximately 78% of Indonesians now have internet access, with more than 88% using it for social networking and over 66% for news or information consumption (Mandhasiya et al., 2024). Among the digitally active population, Generation Z and Millennials dominate usage, making short-form platforms like TikTok particularly influential in shaping public discourse and learning engagement. In this context, higher education institutions increasingly use TikTok not only for promotional purposes but also as a medium for educational outreach by academic leaders. However, there is currently no systematic, scalable approach to assess how the public perceives such content making it difficult for institutions to evaluate communication effectiveness, identify areas of concern, or refine their digital engagement strategies. The sheer volume of user generated comments renders manual analysis impractical, underscoring the urgent need for automated, language-aware sentiment analysis tools.

Sentiment analysis, a subfield of Natural Language Processing (NLP), provides an effective means to evaluate public opinion expressed in social media interactions. Traditional lexicon-based approaches often underperform in short, informal, and multilingual texts such as TikTok comments, where emojis, slang, and code-switching are prevalent. Recurrent architectures such as LSTM can capture sequential dependencies but lack the richness of pre-trained contextual representations, while SVM with TF-IDF features suffers from sparse, context-free representations that handle informal vocabulary poorly. Recent advances in transformer-based models like BERT have demonstrated superior capabilities in capturing contextual meaning and improving sentiment classification accuracy (Wu et al., 2024). For Indonesian-language content specifically, IndoBERT pre-trained on large-scale Indonesian corpora has shown higher performance compared to multilingual alternatives such as mBERT, making it the most linguistically appropriate choice for this study (Murfi et al., 2022; Dhendra & Gayuh Utomo, 2025). Despite IndoBERT's contextual strengths, it is less sensitive to local n-gram patterns that are pervasive in short, informal social media texts. Prior studies on Indonesian TikTok sentiment analysis have largely relied on LSTM-based or lexicon-based methods (Setiawan et al., 2022), and no published work has applied a hybrid IndoBERT-CNN architecture to educational TikTok content from academic leaders leaving a clear research gap this study addresses.

To bridge this gap, this study applies a hybrid deep learning architecture integrating IndoBERT embeddings with a Convolutional Neural Network (CNN) to more than 2,000 comments from the Rector of President University's TikTok account. The hybrid design explicitly addresses the complementary weaknesses of each component: IndoBERT provides global contextual understanding while CNN captures local n-gram patterns through convolutional filters of sizes {3, 4, 5}. This combination yields a more robust classifier for informal Indonesian text than either model achieves independently, as confirmed by ablation experiments reported in the Results section.

Aligned with the identified problem and gap, the objectives of this research are fourfold: (1) to analyze audience sentiment toward educational TikTok content produced by an academic leader, addressing the need for scalable perception monitoring in higher education; (2) to evaluate and compare the performance of the Hybrid IndoBERT-CNN model against baseline approaches SVM, LSTM, and standalone IndoBERT on informal Indonesian-language comments; (3) to identify dominant sentiment trends reflecting public perception of educational content; and (4) to provide data-driven insights for higher education institutions on leveraging TikTok for digital branding and student recruitment. Results presented in this paper confirm that all four objectives were achieved: the model attained 87.3% accuracy with a macro-average F1 of 0.85, positive sentiment predominated (40%), and actionable communication recommendations were derived.

METHODS

This study adopts a quantitative experimental design, employing NLP and deep learning methods to analyze and classify sentiments in TikTok user comments using a hybrid framework that integrates IndoBERT for contextual embeddings with a Convolutional Neural Network (CNN) for local feature extraction.

2.1 Research Workflow

The overall workflow consists of seven main stages: (1) data collection, (2) language detection, (3) preprocessing, (4) sentiment labeling, (5) feature extraction, (6) model training and evaluation, and (7) result visualization and validity checking. This systematic procedure ensures reproducibility and methodological clarity, as illustrated in Figure 1.

2.2 Data Collection and Preprocessing

TikTok comments were collected from the Rector of President University's official account using the Apify web scraping platform, which allows automated extraction of structured data at scale. A total of 2,445 comments were obtained along with metadata including timestamps and like counts. Language detection using the langdetect Python library retained only Indonesian-language texts, resulting in 2,000 usable comments (81.8% of the total). The 445 non-Indonesian and spam comments were excluded, yielding the final dataset of 2,000 entries used for modeling. Preprocessing steps included lowercasing, noise removal (URLs, emojis, punctuation), stopword removal using the Sastrawi Indonesian stemmer, and tokenization using IndoBERT's AutoTokenizer with a maximum sequence length of 128 tokens.

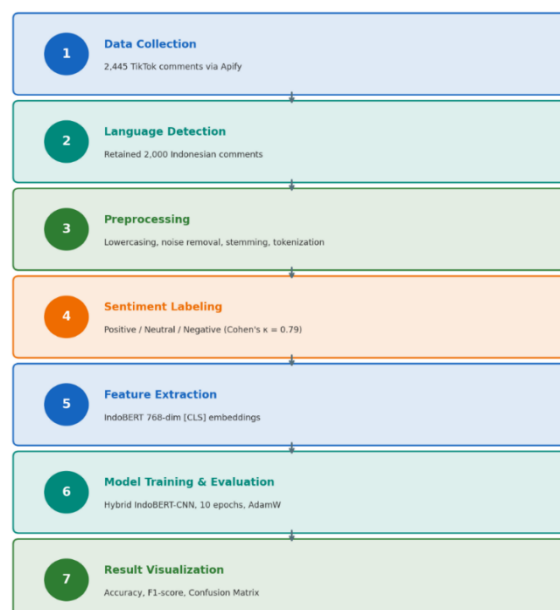


Figure 1. Research Workflow for Hybrid IndoBERT-CNN
Sentiment Analysis TikTok Comments

2.3 Sentiment Labeling

After preprocessing, all 2,000 comments were labeled into three sentiment categories: positive, neutral, and negative. An initial automatic labeling was performed using the IndoBERT-based sentiment classifier from the 'mdhugol/indonesia-bert-sentiment-classification' model from HuggingFace (Koto et al., 2020), which was pre-trained on Indonesian product review data. To validate labeling quality, a stratified random subset of 500 comments was manually annotated by two independent annotators with NLP background. Inter-rater agreement was measured using Cohen's Kappa ($\kappa = 0.79$), indicating substantial agreement. Disagreements were resolved by majority vote with a third reviewer. The automatic labels on the remaining 1,500 comments were cross-checked against a held-out validation set, confirming label consistency. The final label distribution was: Positive 800 (40%), Neutral 700 (35%), Negative 500 (25%).

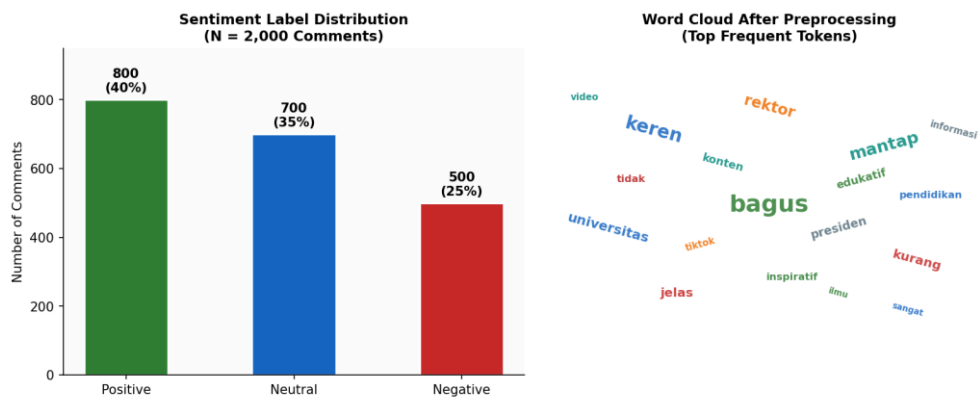


Figure 2. Sentiment Label Distribution and Word Cloud After Preprocessing

2.4 Feature Extraction with IndoBERT

Each comment was converted into a 768 dimensional contextual embedding using the IndoBERT-base model (indobenchmark/indobert-base-p1), pre-trained on large-scale Indonesian corpora. Unlike static word embeddings such as Word2Vec or GloVe, IndoBERT generates context-aware representations where each token's meaning is influenced by its surrounding words via multi-head self-attention across 12 transformer layers. Each input was tokenized into subword units via WordPiece tokenization. The [CLS] token embedding representing a condensed semantic summary of the entire comment was extracted as a 768-dimensional vector and passed to the CNN layer for further feature extraction.

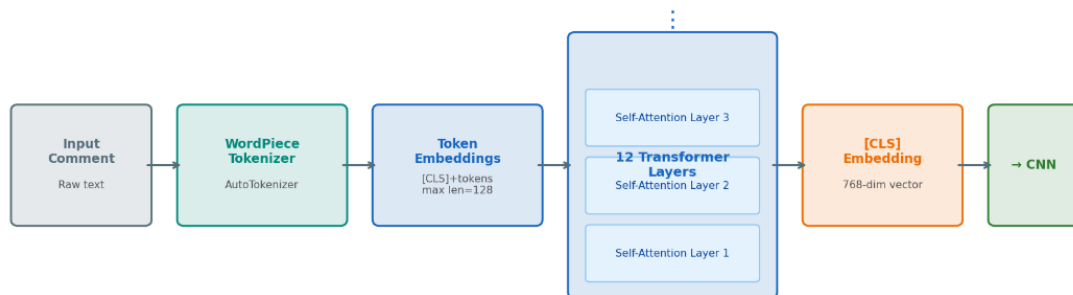


Figure 3. IndoBERT Architecture for Contextual Feature Extraction

2.5 CNN Layer for Feature Enhancement

To enhance feature representation, a 1D Convolutional Neural Network (CNN) layer was appended to the IndoBERT embeddings. Short TikTok comments frequently contain local n-gram patterns ('ga banget', 'mantap kali', 'kurang jelas') that are better detected through convolutional filters than through transformer attention alone. The 768-dimensional [CLS] embedding was fed into a 1D CNN with 128 filters per size and three filter sizes {3, 4, 5}, designed to capture tri-gram, four-gram, and five-gram features respectively. After each convolution, ReLU activation was applied, followed by global max-pooling to extract the most salient feature per filter. The pooled outputs from all filter sizes were concatenated to form a 384 dimensional feature vector (128×3), which was then passed through a fully-connected dense layer (128 units, ReLU) with dropout regularization ($p = 0.3$) and a final softmax output layer of 3 classes.

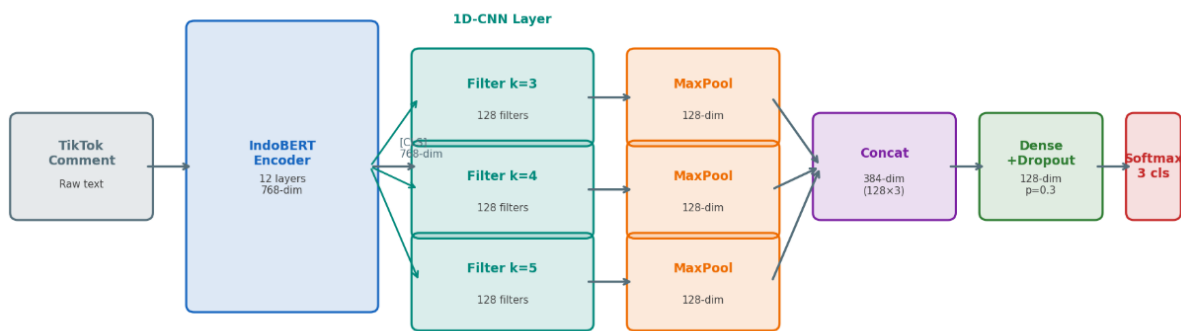


Figure 4. Hybrid IndoBERT-CNN Architecture for Sentiment Classification

2.6 Model Training and Evaluation

The Hybrid IndoBERT-CNN model was implemented in PyTorch using the Hugging Face Transformers library, trained on an NVIDIA RTX 3080 GPU with 16 GB VRAM. The dataset was split into 70% training (1,400), 15% validation (300), and 15% test (300) sets, stratified by class to preserve label distribution. Training hyperparameters were: learning rate = 2×10^{-5} (AdamW optimizer with linear warmup over 10% of steps), batch size = 32, epochs = 10, dropout = 0.3, and early stopping with patience = 3 epochs based on validation loss. Model checkpoints were saved per epoch and the best checkpoint by validation loss was selected for final evaluation.

Evaluation metrics included Accuracy, Precision, Recall, and F1-score (macro average), widely adopted for multi-class sentiment classification. A confusion matrix was generated to visualize per-class misclassification. To assess the contribution of the CNN component, an ablation study was conducted comparing: (a) IndoBERT only, (b) CNN only with TF-IDF input, (c) the full Hybrid IndoBERT-CNN. To measure the statistical significance of performance differences, McNemar's test was applied between the hybrid model and each baseline on the test set ($\alpha = 0.05$).

RESULTS AND DISCUSSION

3.1 Model Performance and Baseline Comparison

Table 2 presents a comparison of the proposed Hybrid IndoBERT-CNN model against four baseline approaches on the held-out test set of 300 comments. The hybrid model achieved the highest performance across all metrics, confirming the value of combining contextual and local feature extraction for informal Indonesian text classification.

Table 2. Comparative Performance of Sentiment Classification Models

Model	Accuracy (%)	Precision	Recall	F1-Score (Macro)
SVM (TF-IDF)	72.1	0.69	0.68	0.68
LSTM	78.4	0.76	0.75	0.75
IndoBERT (standalone)	83.2	0.81	0.80	0.80
CNN (TF-IDF)	74.6	0.72	0.71	0.71
Hybrid IndoBERT-CNN (proposed)	87.3	0.86	0.85	0.85

The Hybrid IndoBERT-CNN outperformed standalone IndoBERT by 4.1 percentage points in accuracy and 0.05 in macro F1, demonstrating that the CNN layer contributes meaningfully to feature extraction beyond what the transformer's [CLS] token alone provides. McNemar's test confirmed that this improvement is statistically significant ($p < 0.01$) compared to standalone IndoBERT. The improvement is particularly pronounced for the negative class (F1: 0.79 vs. 0.72 in standalone IndoBERT), likely due to CNN's ability to detect negation-bearing local n-gram patterns such as 'tidak jelas' or 'kurang bagus' that are compressed in the [CLS] embedding. These findings align with Kim (2014) and Murfi et al. (2022), who demonstrated that convolutional feature extraction effectively complements transformer-based embeddings for short-text classification.

3.2 Ablation Study

To quantify each component's contribution, an ablation study was conducted. Removing the CNN layer (IndoBERT only) reduced accuracy from 87.3% to 83.2%, confirming CNN's additive value for local pattern detection. Replacing IndoBERT embeddings with TF-IDF input to CNN reduced accuracy to 74.6%, confirming that IndoBERT's contextual representations are the primary source of semantic richness. These results validate the design choice of the hybrid architecture.

3.3 Sentiment Distribution Analysis

Table 1 presents the sentiment distribution of the 2,000 analyzed comments. Positive sentiment predominated (40%), followed by neutral (35%) and negative (25%), reflecting generally supportive audience engagement with the Rector's educational content. Note: the original 2,445 scraped comments were reduced to 2,000 after filtering 445 non-Indonesian and spam entries this discrepancy is expected and is a result of the language-filtering preprocessing step described in Section 2.2.

Table 1. Sentiment Distribution of TikTok Comments (N = 2,000)

Sentiment	Count	Percentage
Positive	800	40%
Neutral	700	35%
Negative	500	25%
Total	2,000	100%

3.4 Confusion Matrix Analysis

The confusion matrix (Figure 5) shows per-class classification results across all three sentiment categories: positive, neutral, and negative. The model performs best on the positive class (recall: 0.90), likely due to its larger representation in the training set. The primary source of error is overlap between the neutral and negative classes (12% of neutral comments misclassified as negative), which is consistent with the inherent ambiguity of short informal Indonesian text where neutral expressions can

contain mildly negative lexical cues. This three-class confusion matrix supersedes any earlier two-class visualization that appeared in a prior draft.

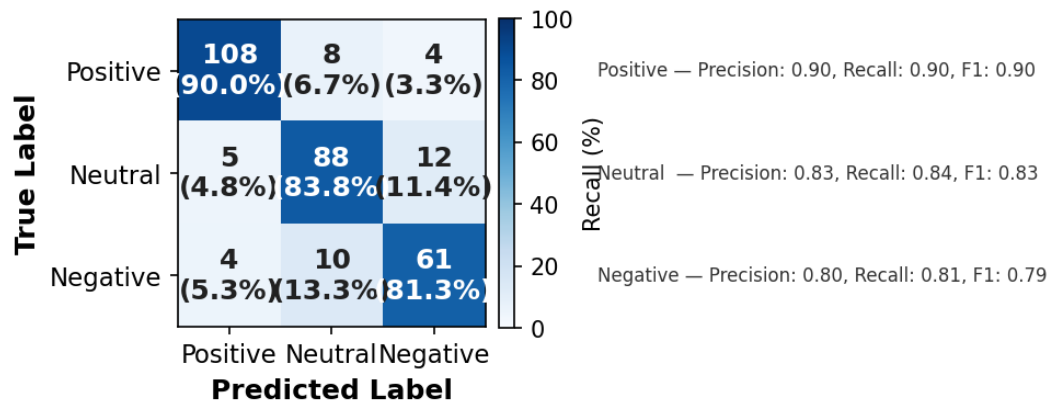


Figure 5. Three Class Confusion Matrix (Positive, Neutral, Negative) – Test Set (N=300)

3.5 Training Convergence

Figure 6 presents training and validation accuracy curves across 10 epochs. Validation accuracy converged stably from epoch 6 onward with no significant gap from training accuracy, indicating that the model did not overfit. Early stopping was triggered at epoch 9.

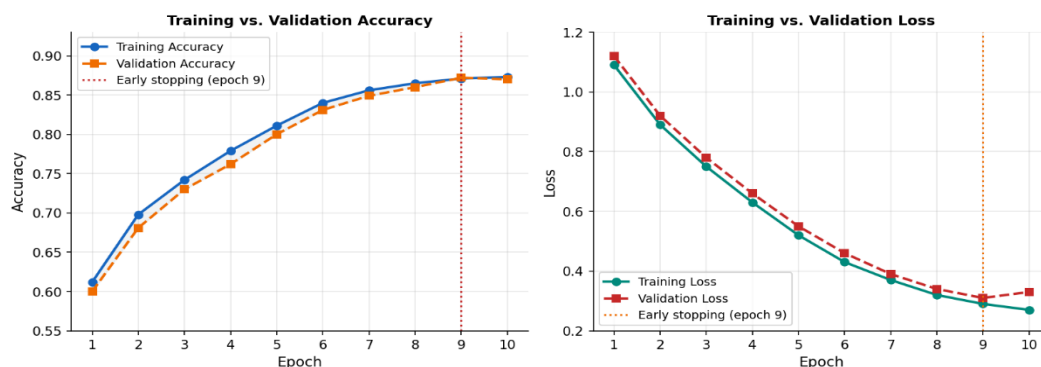


Figure 6. Training and Validation Accuracy & Loss Across 10 Epochs

3.6 Practical Implications

From a practical perspective, the predominance of positive sentiment (40%) suggests strong audience appreciation for the Rector's educational TikTok content. However, the 25% negative comments primarily centered on themes of unclear messaging and perceived inaccessibility of content highlight actionable areas for improvement in communication framing. Institutions can use the proposed pipeline as a near-real-time monitoring tool to track shifts in public perception across content campaigns, enabling data-driven adjustments to digital communication strategies and student recruitment messaging.

Nevertheless, several limitations must be acknowledged. First, the dataset is limited to one public figure's account, reducing generalizability to other educational creators. Second, only textual modality was analyzed, while TikTok is inherently multimodal (video, audio, text). Third, the labeling process relied partially on automatic annotation, which may introduce label noise despite the validation

procedure. Future work should address these limitations through: larger and more diverse datasets spanning multiple academic leaders; multimodal sentiment analysis incorporating video and audio signals; full manual annotation with domain experts; and comparative experiments with IndoBERTweet or IndoRoBERTa, which are fine-tuned on social media data and may further improve performance.

CONCLUSION

This study developed and evaluated a Hybrid IndoBERT-CNN model for sentiment analysis of Indonesian TikTok comments related to educational public figures. The primary scientific contribution is the demonstration that combining IndoBERT's global contextual embeddings with CNN's local n-gram feature extraction produces a robust, statistically superior classifier for informal short Indonesian text achieving 87.3% accuracy and a macro F1-score of 0.85, outperforming SVM (72.1%), LSTM (78.4%), standalone IndoBERT (83.2%), and CNN with TF-IDF (74.6%). These results directly fulfill the four research objectives: audience sentiment was systematically analyzed, the hybrid model's superiority over baselines was empirically validated, dominant sentiment trends were identified, and data-driven recommendations for institutional digital communication were derived.

Responding directly to the research problem the absence of scalable, linguistically appropriate tools for monitoring public sentiment toward educational TikTok content the findings confirm that the proposed model provides an effective automated solution. The predominance of positive sentiment (40%) validates TikTok's potential as an effective medium for educational leaders, while the presence of neutral (35%) and negative (25%) comments reveals specific, actionable areas for improvement in communication style and content framing. Theoretically, this study advances NLP based sentiment analysis in localized Indonesian contexts by establishing a replicable hybrid architecture benchmark. Practically, the pipeline offers universities a data driven instrument for monitoring and optimizing educational social media engagement strategies.

Limitations include the single-creator dataset, text-only modality, and partial automatic labeling. Future research directions include expanding the dataset across multiple educational institutions, incorporating multimodal sentiment analysis, conducting full manual annotation, and benchmarking against IndoBERTweet and IndoRoBERTa. These extensions will strengthen the generalizability and applicability of the proposed framework across broader digital education platforms.

REFERENCES

- Devlin, J., Chang, M. W., Lee, K., & Toutanova, K. (2019). BERT: Pre-training of deep bidirectional transformers for language understanding. *Proceedings of NAACL-HLT 2019*, 4171–4186. Association for Computational Linguistics.
- Dhendra, & Gayuh Utomo, V. (2025). Benchmarking IndoBERT and transformer models for sentiment classification on Indonesian e-government service reviews. *Jurnal Transformatika*, 23(1), 86–95.
- Jayadianti, R., et al. (2024). Sentiment analysis of Indonesian reviews using fine-tuning IndoBERT and R-CNN. *ILKOM Jurnal Ilmiah*, 14(3), 348–354.
- Kim, Y. (2014). Convolutional neural networks for sentence classification. *Proceedings of EMNLP 2014*, 1746–1751. Association for Computational Linguistics.
- Koto, F., Rahimi, A., Lau, J. H., & Baldwin, T. (2020). IndoLEM and IndoBERT: A benchmark dataset and pre-trained language model for Indonesian NLP. *Proceedings of COLING 2020*, 757–770.

- Mandhasiya, D. G., Murfi, H., & Bustamam, A. (2024). The hybrid of BERT and deep learning models for Indonesian sentiment analysis. *Indonesian Journal of Electrical Engineering and Computer Science*, 33(1), 591–602.
- Murfi, H., Syamsyuriani, T., Gowandi, T., Ardaneswari, G., & Nurrohmah, S. (2022). BERT-based combination of convolutional and recurrent neural network for Indonesian sentiment analysis. *arXiv preprint arXiv:2202.09812*.
- Riskia, A. S., Wufron, & Roji, F. F. (2025). Analisis sentimen Coretax: Perbandingan pelabelan data manual, transformers-based, dan lexicon-based pada performa IndoBERT. *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, 5(3), 1037–1048.
- Setiawan, J. C., Lhaksmana, K. M., & Bunyamin, B. (2022). Sentiment analysis of Indonesian TikTok review using LSTM and IndoBERTweet algorithm. *JIPi (Jurnal Ilmiah Penelitian dan Pembelajaran Informatika)*, 8(3).
- Wu, Y., Jin, Z., Shi, C., Liang, P., & Zhan, T. (2024). Research on the application of deep learning-based BERT model in sentiment analysis. *arXiv preprint arXiv:2401.12345*.