



Performance Comparison of FastText and Bi-LSTM for Multilabel Sentiment Analysis on Indonesian Social Media Data: A Case Study of the Asset Confiscation Bill

Theresia Hendrawati¹, Ni Luh Wiwik Sri Rahayu G^{2*}, Made Sena Dwi Tenaya³

^{1,2*,3} Institut Bisnis dan Teknologi Indonesia, Denpasar, Indonesia
[*wiwik@instiki.ac.id](mailto:wiwik@instiki.ac.id)

ARTICLE INFO

Article history:
Received 21 June 2026
Revised 19 July 2026
Accepted 24 July 2026

Keywords:
FastText;
Bi-LSTM;
Multilabel Sentiment
Analysis;
Natural Language
Processing;
Social Media Analytics

ABSTRACT

Public opinions expressed on social media are often multidimensional, making conventional single-label sentiment analysis insufficient for capturing the complexity of public discourse. Multilabel sentiment analysis enables a text instance to be associated with multiple categories simultaneously, providing a more comprehensive representation of public perceptions. This study presents a comparative analysis of FastText and Bidirectional Long Short-Term Memory (Bi-LSTM) for multilabel sentiment classification using Indonesian social media data related to the Asset Confiscation Bill (RUU Perampasan Aset). Data were collected from X (formerly Twitter) and YouTube and annotated into twelve predefined labels encompassing legal, political, economic, administrative, and governance dimensions. The proposed framework consisted of data collection, text preprocessing, multilabel annotation, model development, and performance evaluation using Hamming Score. Three experimental scenarios were conducted for each model to evaluate the impact of Label Attention and architectural optimization. Experimental results demonstrated that FastText consistently outperformed Bi-LSTM across all scenarios. FastText combined with Label Attention and Label Normalization achieved the highest Hamming Score of 0.987450, while Bi-LSTM attained its best performance of 0.967833 using Dense Layer Optimization and Label Attention. The findings indicate that FastText is more effective in handling noisy Indonesian social media texts due to its subword embedding capability. This study contributes empirical evidence regarding the effectiveness of lightweight embedding models for multilabel sentiment analysis and provides insights for future applications in public policy monitoring and social media analytics.

© 2026 Authors. Licensed under [CC BY-SA 4.0](https://creativecommons.org/licenses/by-sa/4.0/)

1. Introduction

The rapid growth of social media has transformed communication patterns and public participation in modern society. Social media platforms facilitate faster and more effective information exchange while enabling users to express opinions on various political, economic, and legal issues [1], [2]. Indonesia has experienced substantial growth in internet usage over the last decade, resulting in a significant increase in user-generated textual content across multiple digital platforms.

The abundance of textual data presents considerable challenges due to its unstructured nature and large volume. Consequently, Natural Language Processing (NLP) has emerged as an effective approach for extracting meaningful information from textual datasets [3]. NLP technologies have been widely implemented in various applications, including

Performance Comparison of FastText and Bi-LSTM for Multilabel Sentiment Analysis on Indonesian Social Media Data: A Case Study of the Asset Confiscation Bill	Theresia Hendrawati, et al
---	----------------------------

machine translation, question-answering systems, text summarization, and sentiment analysis [4].

Sentiment analysis is one of the most prominent NLP tasks and aims to identify emotions, opinions, and attitudes expressed within textual data [5]. Traditional sentiment analysis generally assigns a single label, such as positive, negative, or neutral, to a text instance. However, public discourse regarding government policies often contains multiple perspectives simultaneously, making single-label classification insufficient to capture the complexity of public opinion [6].

Multilabel sentiment analysis addresses this limitation by allowing a single text instance to be associated with multiple labels concurrently [7]. This approach has demonstrated promising performance in several domains, including news categorization, healthcare, and hate speech detection [8]. Moreover, multilabel classification is particularly relevant for policy-related discussions, where legal, social, economic, and political dimensions are frequently intertwined [9].

Despite its advantages, multilabel sentiment analysis remains challenging in social media environments due to the prevalence of noisy text, abbreviations, slang, and out-of-vocabulary (OOV) terms [10]. Therefore, selecting an appropriate text representation method is critical to achieving robust classification performance [11].

FastText has emerged as a lightweight embedding technique capable of handling OOV words through character n-gram representations [12]. Unlike conventional embedding approaches, FastText leverages subword information to generate more robust vector representations, making it suitable for processing noisy social media data [13]. Previous studies have reported that FastText offers competitive performance while maintaining computational efficiency across various sentiment analysis tasks [11].

Another widely adopted approach is Bidirectional Long Short-Term Memory (Bi-LSTM), an extension of Long Short-Term Memory networks that processes textual sequences in both forward and backward directions [14]. This capability enables Bi-LSTM to capture contextual dependencies more effectively than unidirectional architectures and has contributed to its success in numerous NLP applications [15].

Several studies have investigated the application of FastText and recurrent neural networks for sentiment classification. Fajar Pangestu *et al.* reported that a combination of FastText and LSTM achieved an accuracy of 89.5% in sentiment analysis of the Merdeka Curriculum on social media [13]. Similarly, Setiawan demonstrated that the integration of IndoBERT Lite and Bi-LSTM-CNN achieved a Hamming Score of 72.61% for multilabel hate speech detection in Indonesian Twitter data [7]. These findings indicate that both lightweight embedding models and sequence-based architectures exhibit significant potential for Indonesian NLP applications.

One public policy issue that has recently attracted considerable public attention in Indonesia is the Asset Confiscation Bill (RUU Perampasan Aset). The bill aims to strengthen the government's ability to recover assets derived from criminal activities through a Non-Conviction Based Asset Forfeiture mechanism [16]. Although the proposed legislation has received support from anti-corruption advocates, concerns have emerged regarding legal

Performance Comparison of FastText and Bi-LSTM for Multilabel Sentiment Analysis on Indonesian Social Media Data: A Case Study of the Asset Confiscation Bill	Theresia Hendrawati, et al
---	----------------------------

certainty, transparency, and potential violations of individual rights [17]. Consequently, discussions related to the bill have generated highly diverse and multidimensional public opinions across social media platforms [18].

To better understand these multidimensional perceptions, this study adopts a multilabel sentiment analysis framework consisting of twelve predefined categories, including legal, political, economic, transparency, and law enforcement dimensions [9]. Unlike conventional sentiment studies that focus solely on polarity classification, the proposed framework aims to identify multiple aspects of public discourse simultaneously [7].

This study compares the performance of FastText and Bi-LSTM using Indonesian social media data collected from X (formerly Twitter) and YouTube. Three experimental scenarios were designed to evaluate the impact of Label Attention mechanisms and architectural optimization strategies. Hamming Score was employed as the primary evaluation metric due to its suitability for multilabel classification problems [8].

Experimental results demonstrate that FastText consistently outperformed Bi-LSTM across all experimental settings. Specifically, FastText with Label Attention achieved the highest Hamming Score of 0.987450, whereas Bi-LSTM obtained its best score of 0.967833. These findings indicate that subword-based representations are highly effective for Indonesian social media texts characterized by lexical variations and OOV terms [10], [12]. The contributions of this study are threefold. First, it presents a comparative evaluation of FastText and Bi-LSTM for multilabel sentiment analysis on Indonesian social media data. Second, it introduces a twelve-label framework for analyzing public policy discourse. Third, it provides empirical evidence regarding the effectiveness of FastText in handling noisy textual environments commonly encountered in social media applications. The findings are expected to contribute to future NLP research and support the development of intelligent public opinion monitoring systems in Indonesia [19-23].

2. Method

This study employed a comparative experimental approach to evaluate the performance of FastText and Bidirectional Long Short-Term Memory (Bi-LSTM) in multilabel sentiment analysis of Indonesian social media texts. The overall research workflow consisted of five stages: data collection, text preprocessing, multilabel annotation, model development, and performance evaluation. The proposed methodology is illustrated in Fig. 1.

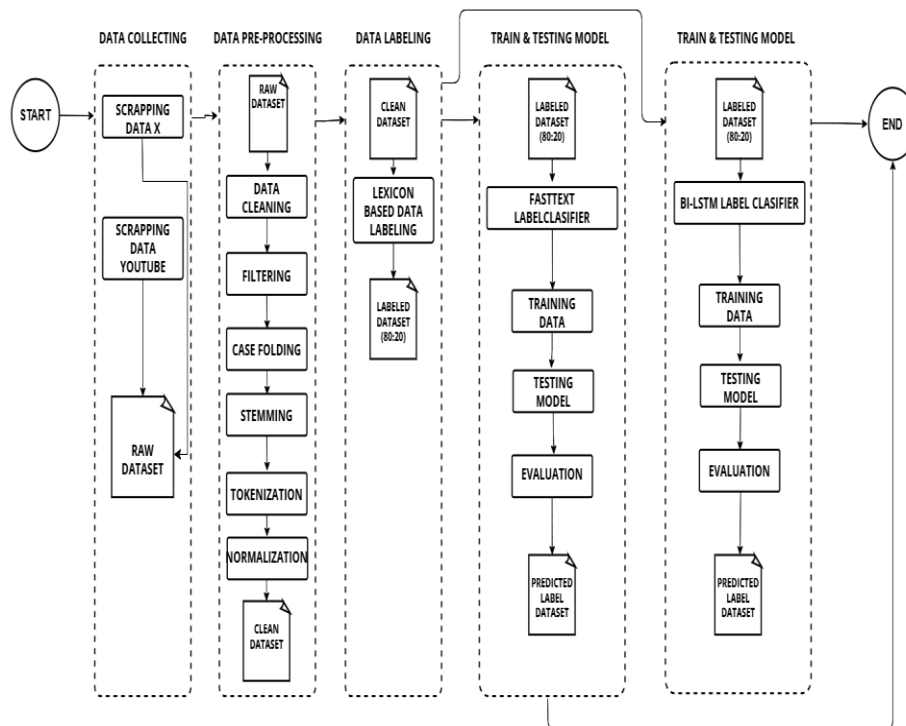


Figure 1. Research Stage

A. Data Collection

The dataset was collected from two major social media platforms, namely X (formerly Twitter) and YouTube, due to their extensive use for expressing public opinions regarding government policies. Web scraping techniques were employed to automatically retrieve textual data associated with discussions related to the Asset Confiscation Bill (RUU Perampasan Aset). Web scraping provides an efficient mechanism for extracting large volumes of unstructured textual information from online sources while minimizing manual intervention [20]. Data collection was conducted using predefined keywords associated with the Asset Confiscation Bill and related legal issues. The collected dataset comprised Indonesian-language social media posts and comments representing diverse perspectives on legal, economic, political, and social aspects of the policy. Duplicate entries, advertisements, and irrelevant content were removed during the initial filtering stage to ensure data quality.

B. Text Preprocessing

Text preprocessing was performed to transform raw textual data into a standardized format suitable for machine learning models. The preprocessing pipeline consisted of five stages: 1) Case Folding, 2) Cleansing, 3) Tokenization, 4) Stopword Removal, 5) Stemming [24-28].

Table 1. Contoh Hasil Pre-processing Data

Performance Comparison of FastText and Bi-LSTM for Multilabel Sentiment Analysis on Indonesian Social Media Data: A Case Study of the Asset Confiscation Bill	Theresia Hendrawati, et al
---	----------------------------

conversation_id_str	username	full_text
1.85574E+18	TheThanos_100	gini amat ya buzzer lowbat 16 persen halu tarik banget ironi
1.85522E+18	pemburupendusta	zelda sukma airin datang lagi andika perkasa hendrar prihadi padahal ibu puan pesan untuk para kader agar solid dukung total ruu rampas aset
1.85522E+18	MichelAdam717	rencana ubah ruu rampas aset jadi ruu pulih aset itu seperti berantas korupsi omong tipu tipu pieter zukifli november pengamat ganti diksi ruu rampas aset hilang esensi dinasti maling jokowi
1.85522E+18	Simon_Deja	ruu rampas aset jaman jokowi hanya jadi banner besar tanpa eksekusi dengan dominasi parlemen di dpr ri sekarang situasi sama di prabowo apa akan jadi banner usang atau eksekusi prabowo omong omong
1.85522E+18	dafern59	jhon sitorus prabowo saling dukung sama anggota komunitas karena merah tidak rendah hati gengsi dan sombong sudah ajak ganjar bareng prabowo tolak ibu ruu rampas aset tolak dan seterusnya ambil sikap oposisi jadi ingat kata pak jokowi saya gigit pakai cara saya semua sumber daya go

Case folding was applied to convert all characters into lowercase to maintain consistency across textual inputs. Cleansing was performed to remove URLs, hashtags, emojis, punctuation marks, and other non-textual elements commonly found in social media content. Subsequently, tokenization segmented textual data into individual tokens, followed by stopword removal to eliminate common Indonesian words that do not contribute significantly to sentiment representation. Finally, stemming was applied to reduce words to their root forms using Indonesian language stemming techniques [3], [4]. These preprocessing steps were essential for reducing textual noise and improving the overall effectiveness of multilabel classification models.

C. Multilabel Annotation

Unlike traditional sentiment analysis, which assigns a single label to each text instance, this study employed a multilabel annotation scheme consisting of twelve predefined categories. The labels were derived from previous studies, legal documents, and expert discussions concerning the Asset Confiscation Bill [7], [9].

1. Individual Rights: Opinions regarding the protection of individual rights in the Asset Forfeiture Bill.
2. Economy: The policy's impact on the economy and the investment climate.
3. Administrative Aspects: Evaluation of procedures and bureaucracy in the implementation of this bill.

Performance Comparison of FastText and Bi-LSTM for Multilabel Sentiment Analysis on Indonesian Social Media Data: A Case Study of the Asset Confiscation Bill	Theresia Hendrawati, et al
---	----------------------------

4. Legal Aspects: Potential conflicts with existing laws and regulations.
5. Transparency and Oversight: Issues regarding oversight and the potential for policy abuse.
6. Effectiveness of Law Enforcement: Assessment of the policy's ability to be implemented in law enforcement.
7. Corruption and Abuse of Authority: Focus on the prevention of corruption-related crimes.
8. State Assets: Evaluation of asset forfeiture affecting state management.
9. Judicial Process: The policy's impact on the rights of defendants in legal proceedings.
10. Politics: The influence of political dynamics on this policy.
11. Progressive: A policy perspective that represents a step forward in the legal system.
12. Regressive: Potential negative impacts on human rights or legal violations.

Each text instance could be assigned multiple labels simultaneously depending on the context expressed within the content. This annotation strategy enabled the proposed framework to capture the multidimensional characteristics of public opinion more effectively than conventional polarity-based sentiment classification approaches.

D. FastText Model

FastText was employed as the first classification model due to its ability to represent words using character n-grams. Unlike conventional embedding techniques, FastText utilizes subword information, enabling the model to generate embeddings for out-of-vocabulary terms frequently encountered in social media texts [29-33].

The FastText architecture used in this study consisted of:

- Input Layer
- Embedding Layer
- Average Pooling
- Dense Layer
- Output Layer

Three experimental scenarios were evaluated:

1. Baseline FastText.
2. FastText with Label Attention.
3. FastText with Label Attention and Label Normalization.

The incorporation of Label Attention was intended to improve the model's capability to identify relevant label relationships within multilabel contexts. Furthermore, label normalization was applied to mitigate label imbalance issues commonly observed in multilabel datasets [11], [13].

E. Bi-LSTM Model

The second model evaluated in this study was Bidirectional Long Short-Term Memory (Bi-LSTM). Bi-LSTM extends conventional LSTM architectures by processing sequences in

both forward and backward directions, thereby capturing contextual dependencies more comprehensively [14].

The Bi-LSTM architecture consisted of:

- Embedding Layer
- Bidirectional LSTM Layer
- Dense Layer
- Sigmoid Activation Layer

Three experimental configurations were investigated:

1. Baseline Bi-LSTM.
2. Bi-LSTM with Dense Layer Optimization.
3. Bi-LSTM with Dense Layer Optimization and Label Attention.

Dense Layer optimization was performed to determine the most suitable number of neurons for the multilabel classification task. Meanwhile, Label Attention was introduced to enhance the model's ability to focus on salient features associated with each label category [14], [15].

F. Experimental Design

To ensure a fair comparison between FastText and Bi-LSTM, both models were trained and evaluated using the same dataset and preprocessing pipeline. The experimental design is summarized in Table 2.

Table 2. Experimental Scenario

Scenario	FastText	Bi-LSTM
1	Baseline	Baseline
2	Label Attention	Dense Optimization
3	Label Attention + Normalization	Dense Optimization + Label Attention

The experiments were conducted under identical computational settings to minimize performance bias and ensure reproducibility.

G. Performance Evaluation

Model performance was evaluated using Hamming Score, which is widely adopted for multilabel classification tasks because it measures the proportion of correctly predicted labels across all instances [8]. The Hamming Score is defined as:

$$\text{Hamming Score} = (\text{TP} + \text{TN}) / (\text{TP} + \text{TN} + \text{FP} + \text{FN}) \quad (1)$$

where:

- TP denotes True Positive,
- TN denotes True Negative,
- FP denotes False Positive, and
- FN denotes False Negative.

A higher Hamming Score indicates better multilabel classification performance. Additionally, execution time was recorded to assess computational efficiency, enabling a

comprehensive comparison between FastText and Bi-LSTM in terms of both predictive accuracy and processing capability.

3. Result and Discussions

FastText Performance Evaluation

Three experimental scenarios were conducted to evaluate the effectiveness of FastText in multilabel sentiment classification. The first scenario utilized the baseline FastText architecture without additional optimization techniques. The second scenario incorporated a Label Attention mechanism to improve the model's ability to identify relationships among labels. Finally, the third scenario combined Label Attention with label normalization to address potential label imbalance issues.

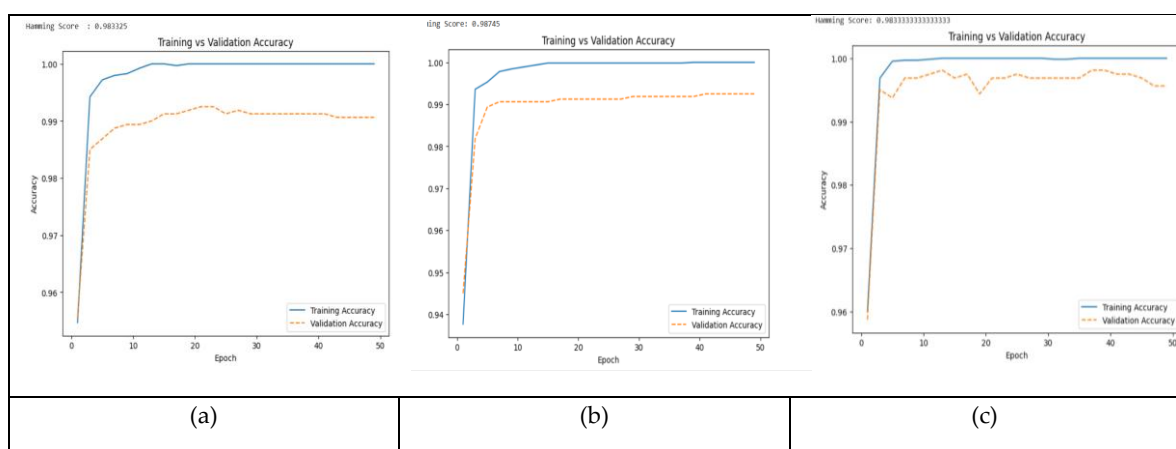


Figure 2. FastText Accuracy Graphs for Scenarios 1(a), 2(b), and 3(c)

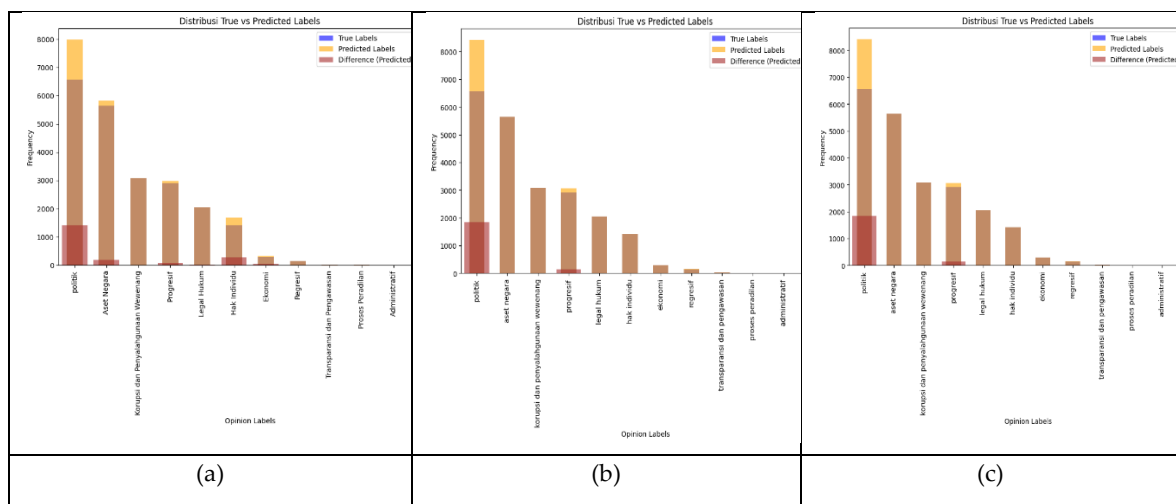


Figure 3. FastText Label Distribution Plots for Scenarios 1(a), 2(b), and 3(c)

The results demonstrate that the integration of Label Attention significantly improved the model's ability to identify relevant labels within a single text instance. Furthermore, the addition of label normalization contributed to a more balanced prediction across all twelve categories. FastText exhibited strong performance in handling noisy Indonesian social media texts. The model successfully captured semantic relationships despite the presence of abbreviations, typographical errors, and non-standard vocabulary commonly found in social media platforms. This capability can be attributed to FastText's character n-gram representation, which enables the generation of meaningful embeddings for words that were not explicitly encountered during training. Another notable finding is the computational efficiency of FastText. The training and inference processes required relatively low computational resources compared to deep learning architectures, making FastText a practical alternative for large-scale sentiment monitoring systems.

Bi-LSTM Performance Evaluation

The performance of Bi-LSTM was evaluated using three experimental configurations. The baseline model consisted of an embedding layer followed by a Bidirectional LSTM and a dense output layer. Subsequent experiments introduced Dense Layer optimization and Label Attention mechanisms to improve classification performance.

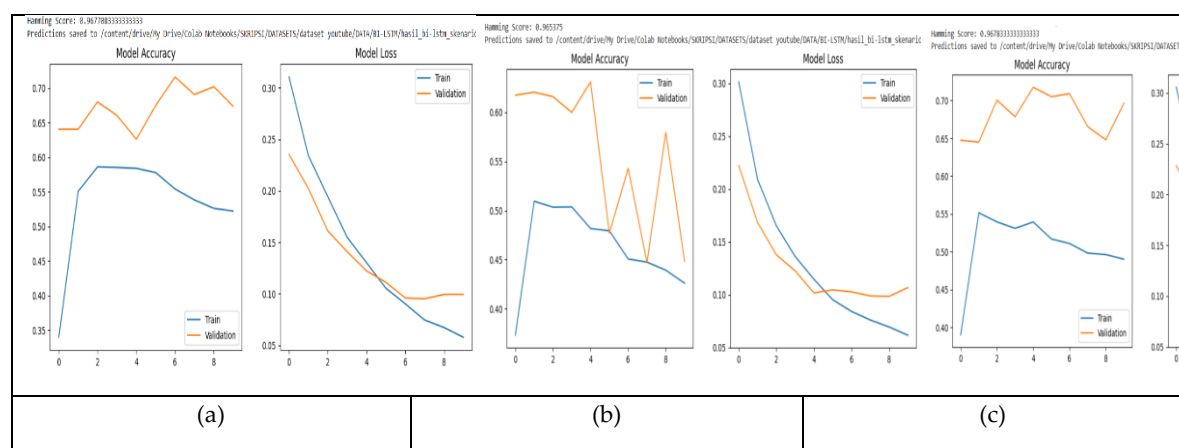
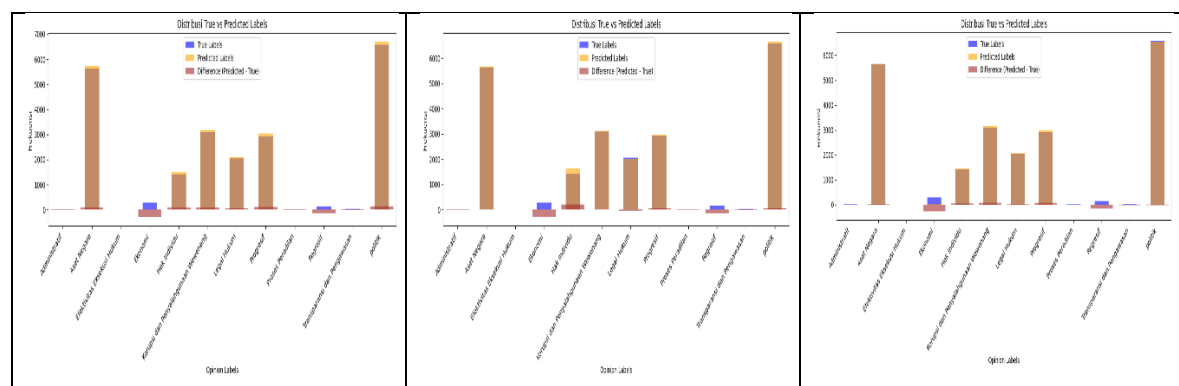


Figure 4. Bi-LSTM Accuracy Plots for Scenarios 1(a), 2(b), and 3(c)



(a)	(b)	(c)
-----	-----	-----

Figure 5. Bi-LSTM Label Distribution Plots for Scenarios 1(a), 2(b), and 3(c)

The findings indicate that Bi-LSTM benefited from architectural optimization. Specifically, Dense Layer adjustments improved the model's capability to map latent representations to the multilabel output space. The inclusion of Label Attention further enhanced performance by enabling the model to focus on salient features associated with each label category. Although Bi-LSTM demonstrated competitive performance, its Hamming Score remained lower than that of FastText across all experimental scenarios. Additionally, Bi-LSTM required substantially longer training times due to the sequential nature of recurrent neural network computations. During the experiments, Bi-LSTM showed greater sensitivity to noisy textual inputs. Informal language patterns, abbreviations, and spelling inconsistencies occasionally affected the quality of contextual representations generated by the model. Consequently, Bi-LSTM exhibited slightly reduced robustness when processing social media texts characterized by high lexical variability.

Model Evaluation

The results of the performance evaluation for each model are presented in Table 4.1, which includes several evaluation metrics for comparing the performance of the two models FastText and Bi-LSTM to identify the most optimal model. The following is a summary of the performance evaluation results for the multilabel classification models, as shown in Table 3.

Table 3. Model Evaluation

Model	Performance Evaluation			
	Dataset	Embedding Dimension	Testing Time (s)	Accuracy (Hamming Score)
Fasttext (Skenario 1)	Dataset 80:20	100	120 s	0.983325
Fasttext (Skenario 2)	Dataset 80:20	100	195 s	0.987450
Fasttext (Skenario 3)	Dataset 80:20	100	220 s	0.983333
Bi-LSTM (Skenario 1)	Dataset 80:20	128	504 s	0.967708
Bi-LSTM (Skenario 2)	Dataset 80:20	128	810 s	0.965375

Performance Comparison of FastText and Bi-LSTM for Multilabel Sentiment Analysis on Indonesian Social Media Data: A Case Study of the Asset Confiscation Bill				Theresia Hendrawati, et al
Bi-LSTM (Skenario 3)	Dataset 80:20	128	712 s	0.967833

Based on the analysis in the table above, it can be concluded that in the first scenario, the FastText and Bi-LSTM models were trained using a basic architectural mechanism. FastText demonstrated excellent performance with a Hamming Score of 0.983325 and a testing time of 120 seconds. Meanwhile, Bi-LSTM produced a Hamming Score of 0.967708 with a testing time of 504 seconds. This indicates that FastText is faster and more accurate than Bi-LSTM in the baseline scenario without additional labeling.

In the second scenario, applying the Label Attention mechanism to FastText resulted in an increase in the Hamming Score to 0.987450, although the testing time increased to 195 seconds. Meanwhile, Bi-LSTM recorded a slight decline in performance with a Hamming Score of 0.965375 and a testing time of 810 seconds. In this scenario, the number of neurons in the Dense Layer was adjusted to match the original labels, resulting in tighter label classification. This caused the accuracy to decrease because a greater number of predicted labels did not match the original labels. In the third scenario, additional processing—including label normalization and spelling adjustments (diff)—was applied to the FastText model. The results showed that FastText remained superior with a Hamming Score of 0.983333, although this was a slight increase from Scenario 1 and the testing time was longer at around 220 seconds; however, this demonstrates that the Label Attention mechanism can improve the model's accuracy in terms of predictions. Meanwhile, Bi-LSTM recorded a Hamming Score of 0.967833 with the longest testing time of 712 seconds; the training and testing process using the Label Attention mechanism resulted in improved label accuracy. This indicates that although Bi-LSTM is capable of handling normalized data well, it remains less efficient and accurate than FastText due to its weakness in understanding out-of-vocabulary (OOV) words, whereas FastText is far superior in this regard. This results in FastText's training and validation accuracy scores being higher and more stable. Overall, the FastText model consistently demonstrated the best performance across all scenarios, both in terms of accuracy and testing time efficiency. FastText achieved its highest performance in the second scenario with a Hamming Score of 0.987450, indicating that the Label Attention mechanism has a significant impact on the effectiveness of multilabel classification. Meanwhile, Bi-LSTM experienced score fluctuations that tended to be stable but required significantly longer testing times in every scenario.

Thus, it can be concluded that the FastText model is the most optimal method for the multilabel classification task of public opinion regarding the Asset Forfeiture Bill, both in terms of accuracy and processing speed.

4. Conclusion

This study presented a comparative evaluation of FastText and Bidirectional Long Short-Term Memory (Bi-LSTM) for multilabel sentiment analysis using Indonesian social media data related to the Asset Confiscation Bill. A multilabel framework consisting of twelve

Performance Comparison of FastText and Bi-LSTM for Multilabel Sentiment Analysis on Indonesian Social Media Data: A Case Study of the Asset Confiscation Bill	Theresia Hendrawati, et al
---	----------------------------

categories was developed to capture the multidimensional nature of public discourse, including legal, political, economic, transparency, and governance aspects. Three experimental scenarios were conducted for each model to assess the impact of Label Attention and architectural optimization on classification performance. The experimental results demonstrated that FastText consistently achieved superior performance across all scenarios. The highest Hamming Score of 0.987450 was obtained using FastText with Label Attention and Label Normalization, whereas Bi-LSTM achieved its best score of 0.967833 with Dense Layer Optimization and Label Attention. The findings indicate that FastText is more robust when handling noisy Indonesian social media texts due to its ability to generate subword representations and effectively process out-of-vocabulary terms. Furthermore, FastText exhibited lower computational complexity and faster processing times, making it a practical solution for large-scale and real-time sentiment monitoring systems. Overall, this study confirms that lightweight embedding models can provide highly competitive performance for multilabel sentiment analysis tasks in Indonesian Natural Language Processing applications. Future work may explore the integration of transformer-based architectures such as IndoBERT and multilingual language models, as well as the application of the proposed framework to other public policy domains and larger multilingual datasets.

Acknowledgment

Thank you to the DRPM of the Indonesian Institute of Business and Technology for the internal research grant under the INSTIKI Research and Development Program (IRDP)

References

- [1] Abdurrohim, I., and A. Rahman, "Penerapan Natural Language Processing untuk Analisis Sentimen terhadap Kebijakan Pemerintah," 2024.
- [2] Agatha Olivia Victoria, "Ahli: RUU Perampasan Aset Pastikan Pelaku Tak Nikmati Hasil Korupsi," 2024.
- [3] Akbar, I., M. Faisal, *et al.*, "Penerapan Long Short-Term Memory untuk Klasifikasi Multi-Label Terjemahan Al-Qur'an dalam Bahasa Indonesia," vol. 7, 2024.
- [4] Amanda, R., "Pemanfaatan Jejaring Sosial Twitter Sebagai Media Penyebaran Informasi Publik (Studi pada Akun Twitter @djplkemenhub151)," 2021.
- [5] Amien, M., "Sejarah dan Perkembangan Teknik Natural Language Processing (NLP) Bahasa Indonesia: Tinjauan tentang Sejarah, Perkembangan Teknologi, dan Aplikasi NLP dalam Bahasa Indonesia," *ELANG: Journal of Interdisciplinary Research*, vol. 1, no. 1, pp. 99–105, 2023, doi:10.48550/arXiv.2304.02746.
- [6] Bojanowski, P., E. Grave, A. Joulin, and T. Mikolov, "Enriching Word Vectors with Subword Information," *Transactions of the Association for Computational Linguistics*, vol. 5, pp. 135–146, 2017, doi:10.1162/tacl_a_00051.

Performance Comparison of FastText and Bi-LSTM for Multilabel Sentiment Analysis on Indonesian Social Media Data: A Case Study of the Asset Confiscation Bill	Theresia Hendrawati, et al
---	----------------------------

- [7] Cambria, E., and B. White, "Jumping NLP Curves: A Review of Natural Language Processing Research," *IEEE Computational Intelligence Magazine*, vol. 9, no. 2, pp. 48–57, 2014, doi:10.1109/MCI.2014.2307227.
- [8] Damayanti, A., I. D. Delima, *et al.*, "Pemanfaatan Media Sosial Sebagai Media Informasi dan Publikasi (Studi Deskriptif Kualitatif pada Akun Instagram @rumahkimkotatangerang)," 2023.
- [9] Danil, E., and N. Mulyati, "Peran Kejaksaan dalam Perampasan Aset dalam Pemberantasan Tindak Pidana Korupsi serta Kendala yang Dihadapi dalam Pelaksanaannya," vol. 6, no. 1, 2023.
- [10] Devlin, J., M.-W. Chang, K. Lee, and K. Toutanova, "BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding," in *Proc. NAACL-HLT*, 2019, pp. 4171–4186, doi:10.18653/v1/N19-1423.
- [11] F. Pangestu, *et al.*, "Analisis Sentimen pada Media Sosial X terhadap Implementasi Kurikulum Merdeka Menggunakan Metode FastText dan Long Short-Term Memory (LSTM)," 2024.
- [12] Hochreiter, S., and J. Schmidhuber, "Long Short-Term Memory," *Neural Computation*, vol. 9, no. 8, pp. 1735–1780, 1997, doi:10.1162/neco.1997.9.8.1735.
- [13] Howard, J., and S. Ruder, "Universal Language Model Fine-Tuning for Text Classification," in *Proc. ACL*, 2018, pp. 328–339, doi:10.18653/v1/P18-1031.
- [14] Jurafsky, D., and J. H. Martin, *Speech and Language Processing*, 3rd ed. Stanford, CA, USA: Stanford University, 2024.
- [15] Joulin, A., E. Grave, P. Bojanowski, and T. Mikolov, "Bag of Tricks for Efficient Text Classification," in *Proc. EACL*, 2017, pp. 427–431, doi:10.18653/v1/E17-2068.
- [16] Kantor Staf Kepresidenan, "Percepatan Pembahasan RUU Perampasan Aset untuk Penguatan Sistem Hukum Nasional," 2024.
- [17] Khder, M., "Web Scraping or Web Data Extraction Techniques: A Review," *International Journal of Advances in Soft Computing and Its Applications*, vol. 13, no. 1, pp. 69–83, 2021.
- [18] K. Sari, *et al.*, "Klasifikasi Teks Multilabel pada Artikel Berita Menggunakan Long Short-Term Memory dengan Word2Vec," 2020.
- [19] Liu, B., *Sentiment Analysis and Opinion Mining*. San Rafael, CA, USA: Morgan & Claypool Publishers, 2012.
- [20] Mikolov, T., K. Chen, G. Corrado, and J. Dean, "Efficient Estimation of Word Representations in Vector Space," *arXiv preprint*, arXiv:1301.3781, 2013.
- [21] Minaee, S., N. Kalchbrenner, E. Cambria, *et al.*, "Deep Learning Based Text Classification: A Comprehensive Review," *ACM Computing Surveys*, vol. 54, no. 3, pp. 1–40, 2021, doi:10.1145/3439726.

Performance Comparison of FastText and Bi-LSTM for Multilabel Sentiment Analysis on Indonesian Social Media Data: A Case Study of the Asset Confiscation Bill	Theresia Hendrawati, et al
---	----------------------------

- [22] Mustasaruddin, *et al.*, "Klasifikasi Sentiment Review Aplikasi MyPertamina Menggunakan Word Embedding FastText dan Support Vector Machine," 2023.
- [23] Reimers, N., and I. Gurevych, "Sentence-BERT: Sentence Embeddings Using Siamese BERT-Networks," in *Proc. EMNLP-IJCNLP*, 2019, pp. 3982–3992, doi:10.18653/v1/D19-1410.
- [24] Schuster, M., and K. K. Paliwal, "Bidirectional Recurrent Neural Networks," *IEEE Transactions on Signal Processing*, vol. 45, no. 11, pp. 2673–2681, 1997, doi:10.1109/78.650093.
- [25] Setiawan, "Analisis Sentimen Multilabel terhadap Ujaran Kebencian Menggunakan Gabungan Metode IndoBERT Lite dan Bidirectional LSTM-CNN dengan Grid Search Hyperparameter Optimization," 2022.
- [26] SEPRIADI NANDA, "Analisis Sentimen Review Aplikasi MyPertamina Menggunakan Word Embedding FastText dan Algoritma K-Nearest Neighbor," 2023.
- [27] Tsoumakas, G., and I. Katakis, "Multi-Label Classification: An Overview," *International Journal of Data Warehousing and Mining*, vol. 3, no. 3, pp. 1–13, 2007, doi:10.4018/jdwm.2007070101.
- [28] Umer, *et al.*, "FastText-Based Text Classification for Multilabel Learning Applications," 2023.
- [29] Vaswani, A., N. Shazeer, N. Parmar, *et al.*, "Attention Is All You Need," in *Advances in Neural Information Processing Systems*, vol. 30, 2017.
- [30] Veronika Aritonang, *et al.*, "Implementasi Bidirectional Long Short-Term Memory untuk Analisis Sentimen Bahasa Indonesia," 2022.
- [31] Wolf, T., L. Debut, V. Sanh, *et al.*, "Transformers: State-of-the-Art Natural Language Processing," in *Proc. EMNLP*, 2020, pp. 38–45, doi:10.18653/v1/2020.emnlp-demos.6.
- [32] Young, T., D. Hazarika, S. Poria, and E. Cambria, "Recent Trends in Deep Learning Based Natural Language Processing," *IEEE Computational Intelligence Magazine*, vol. 13, no. 3, pp. 55–75, 2018, doi:10.1109/MCI.2018.2840738.