

Enhancing Natural Disaster Monitoring: A Deep Learning Approach to Social Media Analysis Using Indonesian BERT Variants

Karlina Elreine Fitriani¹, Mohammad Reza Faisal¹, Muhammad Itqan Mazdadi¹, Fatma Indriani¹, Dodon Turiyanto Nugrahadi¹, and Septyan Eka Prastya²

¹ Department of Computer Science, Lambung Mangkurat University, Banjarbaru, Indonesia

² Department of Information Technology, Sari Mulia University, Banjarmasin, Indonesia

ABSTRACT

Social media has become a primary source of real-time information that can be leveraged by artificial intelligence to identify relevant messages, thereby enhancing disaster management. The rapid dissemination of disaster-related information through social media allows authorities to respond to emergencies more effectively. However, filtering and accurately categorizing these messages remains a challenge due to the vast amount of unstructured data that must be processed efficiently. This study compares the performance of IndoRoBERTa, IndoRoBERTa MLM, IndoDistilBERT, and IndoDistilBERT MLM in classifying social media messages about natural disasters into three categories: eyewitness, non-eyewitness, and don't know. Additionally, this study analyzes the impact of batch size on model performance to determine the optimal batch size for each type of disaster dataset. The dataset used in this study consists of 1000 messages per category related to natural disasters in the Indonesian language, ensuring sufficient data diversity. The results show that IndoDistilBERT achieved the highest accuracy of 81.22%, followed by IndoDistilBERT MLM at 80.83%, IndoRoBERTa at 79.17%, and IndoRoBERTa MLM at 78.72%. Compared to previous studies, this study demonstrates a significant improvement in classification accuracy and model efficiency, making it more reliable for real-world disaster monitoring. Pre-training with MLM enhances IndoRoBERTa's sensitivity and IndoDistilBERT's specificity, allowing both models to better understand context and optimize classification results. Additionally, this study identifies the optimal batch sizes for each disaster dataset: 32 for floods, 128 for earthquakes, and 256 for forest fires, contributing to improved model performance. These findings confirm that this approach significantly improves classification accuracy, supporting the development of machine learning-based early warning systems for disaster management. This study highlights the potential for further model optimization to enhance real-time disaster response and improve public safety measures more effectively and efficiently.

PAPER HISTORY

Received Des. 12, 2024

Revised Jan. 20, 2025

Accepted Feb. 10, 2025

Published Feb 22, 2025

KEYWORDS

Machine Learning;
Social Media;
Natural Disaster;
IndoRoBERTa;
IndoDistilBERT

CONTACT:

¹Mohammad Reza Faisal
reza.faisal@ulm.ac.id

1. INTRODUCTION

Indonesia is one of the countries highly vulnerable to various natural disasters, such as earthquakes, floods, landslides, and volcanic eruptions, due to its location along the Pacific Ring of Fire [1]. Therefore, accurate information is essential to anticipate further impacts.

Innovative big data techniques open up opportunities to understand and mitigate the impacts of natural disasters. One such approach is leveraging social media to monitor user responses, including information about the type, location, time of occurrence, and requests for assistance during disasters [2]. The social media platform X (formerly Twitter) has become a primary source of critical information, particularly real-time news that is easily accessible to the public [3]. This advantage makes social media a crucial tool for gathering up-to-date

information [4]. Additionally, user data from X holds significant potential for analyzing social patterns and responses to events. In natural disasters, this data can be utilized to monitor situations in real-time, identify eyewitnesses, and support emergency decision-making.

The use of social media messages for disaster management can be enhanced with artificial intelligence, which accelerates the identification of messages related to the event [5]. The artificial intelligence system will classify social media messages into three categories: eyewitness, noneyewitness, and don't know. Messages in the eyewitness category are those posted by individuals present at the location during the disaster. Messages in the noneyewitness category are about natural disasters posted by users not at the scene of the event. Meanwhile, messages in the don't know category contain the keyword "natural disaster" but are not related to the actual disaster

[6]. Through data processing from social media, various big data techniques can be applied to identify important information related to disaster situations.

One of the commonly used big data techniques is Bidirectional Encoder Representations from Transformers (BERT) for text classification [7]. BERT can deeply understand the context of language and provide accurate results with minimal adjustments for various Natural Language Processing (NLP) tasks, making it a revolutionary solution that outperforms traditional NLP models in a wide range of tasks [8].

Research [9] used Twitter social media messages to automatically detect fake news through BERT classification, achieving an accuracy of 61%. Subsequent research [10] used Twitter data to analyze customer sentiment towards Garuda Indonesia and compared the performance of the Multinomial Naive Bayes model with an accuracy of 64.8%, SVM with an accuracy of 71.6%, and BERT with an accuracy of 75.6%, showing that BERT had the best performance. The following study [11] used Twitter data to classify messages related to COVID-19, with results showing that the IndoBERT model achieved an accuracy of 89.50%, outperforming BERT with an accuracy of 81.50% in classifying Indonesian language messages. This indicates that BERT is less optimal in understanding the Indonesian language in depth. Therefore, this study explores other Indonesian variants of BERT to optimize the classification of Indonesian language messages, particularly in the domain of natural disasters.

BERT is a base model, while Robustly Optimized BERT Approach (RoBERTa) is an enhanced version with modifications to the training process to improve performance and language understanding in specific contexts [12]. IndoRoBERTa, the Indonesian language variant of RoBERTa, is optimized for efficient pre-training and text classification in Indonesian, with deep contextual understanding [13]. RoBERTa is often trained with more data and longer text sequences, making it more effective in certain natural language processing tasks compared to BERT [14]. IndoDistilBERT, the lightweight version of Distilled BERT (DistilBERT) for the Indonesian language, offers faster inference and competitive performance, making it an ideal choice for applications that require efficient computational resources [15].

Furthermore, a study by [16] compared the text classification results of tweets using Indonesian RoBERTa and IndoDistilBERT, achieving accuracies of 0.705882 and 0.501960, respectively, indicating that Indonesian RoBERTa outperforms IndoDistilBERT for classifying Indonesian language text.

The study [17] shows that Masked Language Modeling (MLM) is a key component in modern language models such as BERT. This technique has proven to perform better than models that do not use MLM. However, existing research has not specifically compared the performance of models with and without MLM. Therefore, this study aims to explore and compare the performance of IndoRoBERTa and IndoDistilBERT models with and

without MLM, in order to understand the extent to which MLM affects the effectiveness of these models.

Based on literature reviews and previous findings, the selection of IndoRoBERTa and IndoDistilBERT along with their MLM variants is driven by the need to handle Indonesian text more effectively. IndoRoBERTa excels in deep contextual understanding, while IndoDistilBERT is lighter and faster. The use of MLM in both models aims to enhance the ability to understand language patterns more accurately. The comparison of the performance of IndoRoBERTa and IndoDistilBERT in NLP tasks for classifying disaster messages from social media has rarely been explored. The novelty of this research lies in the application of machine learning methods with pre-trained MLM on IndoRoBERTa and IndoDistilBERT to enhance classification performance.

Based on the above description, this research focuses on comparing the performance of various NLP models, namely IndoRoBERTa, IndoRoBERTa with MLM, IndoDistilBERT, and IndoDistilBERT with MLM, in classifying disaster-related messages from Indonesian-language social media. This study uses a machine learning approach to train and optimize these classification models, as well as investigates the effect of batch size parameters on model performance to achieve optimal results. With this approach, the research contributes to: (1) Recommending the best model and batch size to be applied in disaster early warning systems; (2) Providing insights into the performance of disaster message classification using Indonesian-based language models; and (3) Serving as a reference for future research in the development of NLP models for sentiment analysis and disaster message classification.

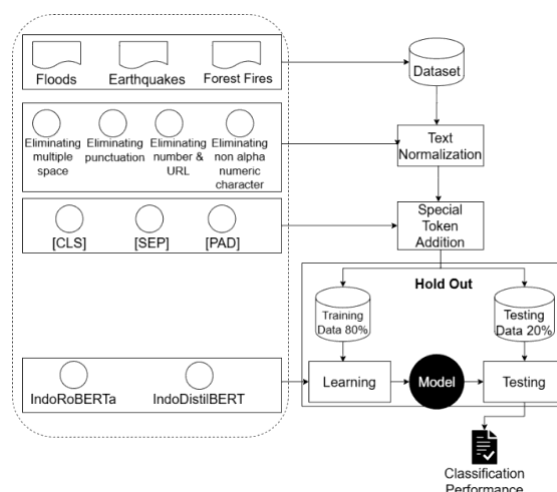


Fig. 1. The Research Procedure with the Pre-Trained Model

2. MATERIALS AND METHOD

This study involves the implementation of two approaches using pre-trained language models, IndoRoBERTa and

IndoDistilBERT, for text classification in disaster-related datasets, including floods, earthquakes, and forest fires. Additionally, the study explores the impact of training these models with a Masked Language Model (MLM) approach before classification. The research procedure with the pre-trained model can be seen in Fig. 1. The research procedure with the training model using MLM can be seen in Fig. 2.

A. Dataset

The dataset used in this research is derived from previous studies [6]. This research uses three datasets with the details provided in Table 1. Each dataset consists of 3000 Twitter messages. The data used consists of comments related to natural disasters, which were collected through a crawling process from the Twitter/X social media platform about natural disasters, specifically floods, earthquakes, and forest fires. This data was then manually labeled and divided into three classes: eyewitness, noneyewitness, and don'tknow. All tweets

Table 1. Dataset

Dataset	Class Label			Total
	eyewitness	non-eyewitness	don't know	
Floods	1000	1000	1000	3000
Eartquakes	1000	1000	1000	3000
Forest Fires	1000	1000	1000	3000

Text normalization in this study serves to clean the text by transforming it from an inconsistent form into a standard format, making it easier for the model to understand [18]. The steps of text normalization in this study include: (1) Eliminating Multiple Spaces to ensure cleaner and more consistent text; (2) Eliminating Punctuation to remove unnecessary marks; (3) Eliminating Numbers & URLs to avoid irrelevant data; (4) Eliminating Non-Alphanumeric Characters to create cleaner text [19].

C. Special Token Addition

The addition of special tokens in this study helps the model understand the structure and context of the data, resulting in more accurate outputs in text classification [20]. The tokens commonly used include the Classification Token (CLS), Separator Token (SEP), and Padding Token (PAD) [21], which are also used in this study. The explanations are as follows: (1) The CLS token is placed at the beginning of the input and represents the entire text; (2) The SEP token is a separator used between two parts of the text; (3) The PAD token is used to equalize the length of inputs within a batch.

D. Model Training MLM

Masked Language Modeling (MLM) is a pre-training method where some tokens in the input text are hidden, and the model is trained to predict the missing tokens based on the remaining context [22]. According to research [23] experiments with a masking rate of 40% are beneficial for training large models, as demonstrated by [22] however, it did not yield better results compared to the default 15% masking rate in their experiments. In this stage of the research, 15% of the words in the text are randomly masked to train the model to predict the missing words, learn patterns and context in Indonesian, and enhance its capabilities before fine-tuning. Of the 15% of selected tokens, 80% are replaced with the [MASK] token, 10% with random tokens, and 10% remain unchanged. Fig. 3 it shows the MLM training process with BERT. This masking strategy helps the model learn contextual relationships by relying on surrounding words to predict the missing tokens. As a result, the pre-trained model gains a better understanding of semantic structures, enhancing its performance on various NLP tasks.

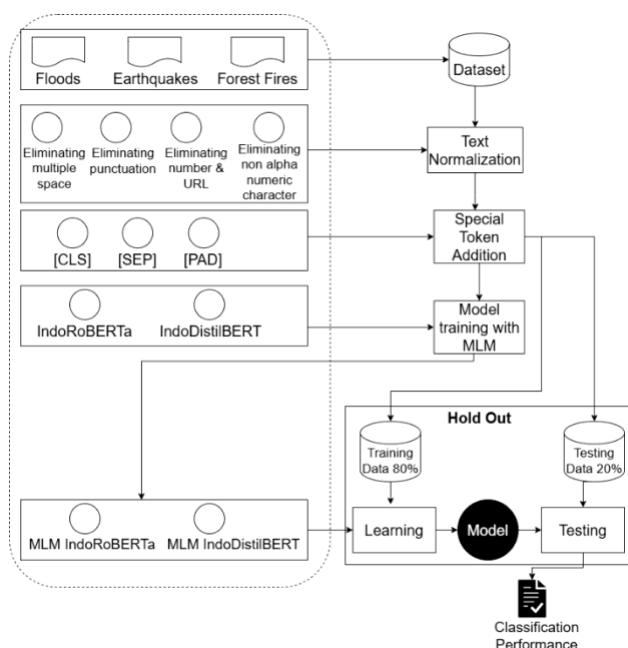


Fig. 2. The Research Procedure with the Training Model Using

are in Indonesian and have undergone manual labeling to ensure their alignment with each respective category. The selection of this dataset is relevant as it reflects real-time information that can be utilized for disaster early warning systems. The dataset can be downloaded from the following link:

<https://github.com/rezafaisal/NaturalDisasterOnTwitter>.

Each category contains 1000 messages, making these three datasets balanced, as each class label has the same number of instances.

B. Text Normalization

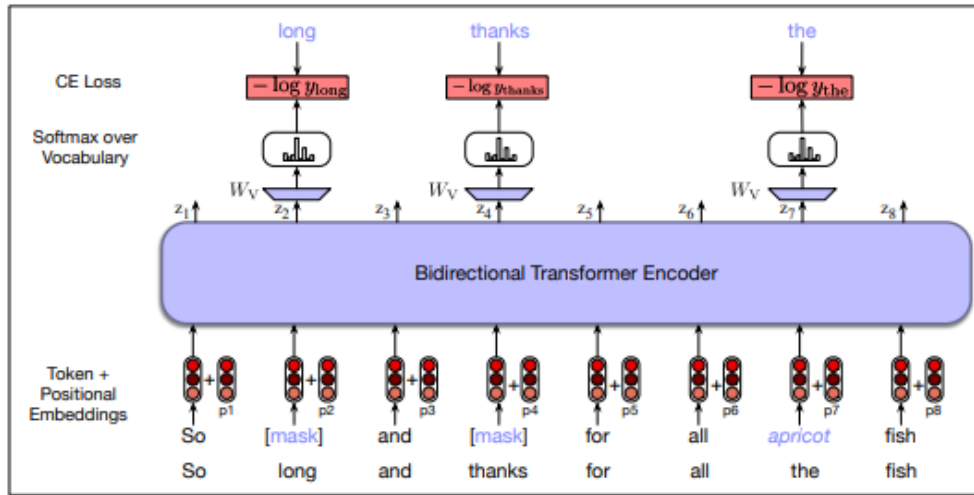


Fig. 3. MLM training process with BERT [24]

In Fig. 3 explains the training of a language model with the masking technique, where two tokens are masked, one is replaced with a random word, and its probability is calculated for the training loss. Models like BERT use subword tokens, even though the input is presented as whole words [24]. The equations in Masked Language Modeling (MLM) serve three main roles. Equation (1), the MLM loss function, calculates the expectation of the negative log probability of the masked token based on its context. Equation (2) computes the probability of the masked token using the softmax function to determine its vector representation. Equation (3) applies the cross-entropy loss function to optimize the model during training [25] [22].

$$L(C) = \mathbb{E}_{x \in C} \mathbb{E}_{\mathcal{M} \subset x} [\sum_{x_1 \in \mathcal{M}} \log p(x_i | \bar{x})] \quad (1)$$

$|\mathcal{M}| = m |x|$

Equation (1), C denotes the text corpus, x represents a sentence within the corpus, and $\mathcal{M} \subset x$ is the set of words in the sentence that are selected for masking. The notation $|\mathcal{M}| = m |x|$ indicates the number of tokens masked in sentence x , while $p(x_i | \bar{x})$ represents the model's probability of predicting the original token x_i based on the observed context $\bar{x}$.

$$p_0(x_i | \hat{x}) = \frac{\exp(e_{x_i}^T h_i)}{\sum_{x' \in \mathcal{V}} \exp(e_{x'}^T h_i)} \quad (2)$$

In this equation (2), x_i, h_i is the contextual representation of the masked token, and $\mathcal{V}$ represents the model's vocabulary.

$$\mathcal{L}_{MLM} = \mathbb{E}(-\sum_{i \in \mathcal{M}} \log p_0(x_i | \hat{x})) \quad (3)$$

Equation (3), this loss function measures the average negative log-likelihood of the masked token predictions, allowing the model to learn to predict missing words more accurately. The variable descriptions remain the same as

in Equations (1) and (2). The number of equations used in this study is limited because the MLM technique fundamentally requires only these core components without significant variations in additional equations.

E. Classification

This study employed a train-test split validation strategy with an 80:20 ratio, along with a validation split of the same ratio on the training data, where stratification was applied to maintain balanced class distribution in each subset. This approach was chosen to reduce bias and ensure the reliability of the model's performance metrics. The data prepared in the previous step was then further divided using the hold-out technique with the same ratio [26]. This ensured that each subset maintained a representative class distribution for a more reliable evaluation, reducing the risk of model bias toward majority classes. The training data will be used to build classification models using: (1) pre-trained IndoRoBERTa; (2) pre-trained IndoDistilBERT; (3) MLM IndoRoBERTa; and (4) MLM IndoDistilBERT.

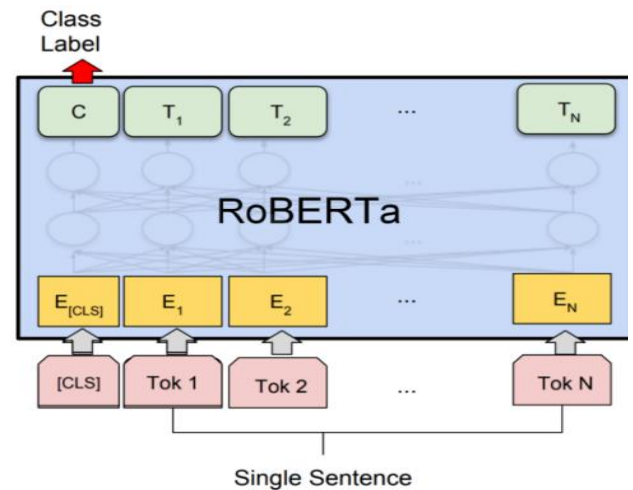


Fig. 4. RoBERTa Architecture [27]

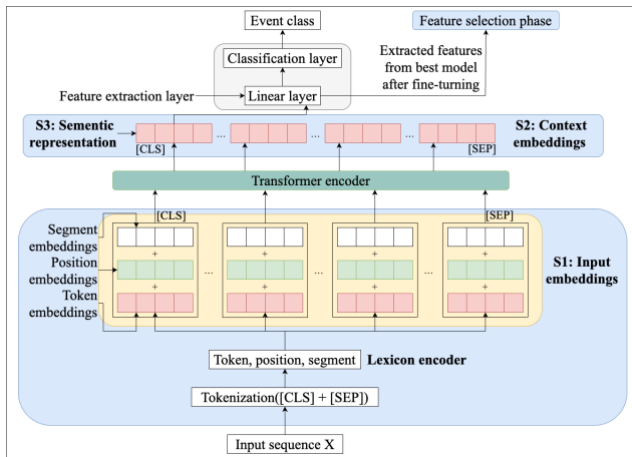


Fig. 5. DistilBERT feature extraction model [29]

IndoRoBERTa is a RoBERTa-based language model designed for Indonesian, trained using the Open Super-Large Crawled ALManACH corpus (OSCAR) dataset, which is 17.05 GB in size and includes 166 languages. By utilizing this large dataset, IndoRoBERTa is capable of understanding various text variations in Indonesian. With 124 million parameters, the model can comprehend different text variations in Indonesian and learn complex language patterns. OSCAR also provides strong generalization capabilities for various types of text [16].

In Fig. 4, RoBERTa uses the Byte-Pair Encoding (BPE) tokenizer to split the text into tokens, adding the [CLS] token at the beginning without the [SEP] token to separate sentences. The tokens are converted into embeddings by summing the token, segment, and position embeddings, and then processed through transformer layers. RoBERTa is more efficient for longer texts because it does not require the [SEP] token [14]. Additionally, IndoRoBERTa has been widely used for various natural language processing tasks, including text classification, sentiment analysis, and named entity recognition. In text classification, the model processes input text by generating contextual embeddings, which are then passed through a classifier to assign predefined labels based on learned patterns. IndoRoBERTa uses the <s> and </s> tokens instead of the [CLS] and [SEP] tokens, in accordance with the Byte-Pair Encoding (BPE) architecture [28] which differs from WordPiece-based models like IndoDistilBERT. IndoDistilBERT is a version of DistilBERT trained on Indonesian Wikipedia and newspaper data, using WordPiece tokenization and a vocabulary of 32000 tokens. This model reduces the number of parameters by 40% compared to BERT, while maintaining 95% of its performance and being 60% faster [16]. Training with Indonesian-language data enables this model to understand the variations of the Indonesian language. The proposed feature extraction model architecture based on DistilBERT is shown in Fig. 5. The feature extraction starts by converting the tweets into word embedding vectors (S1). The transformer encoder uses the self-attention mechanism to generate contextual

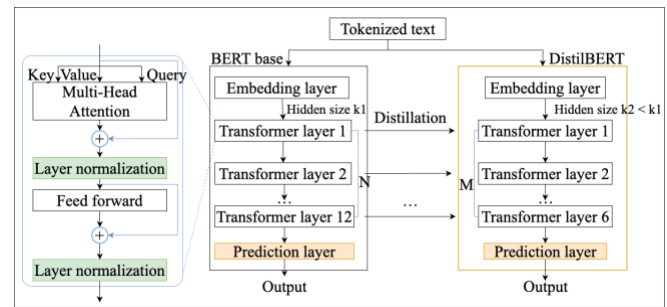


Fig. 6. DistilBERT model architecture and components [29]

embeddings (S2), which are then combined into a semantic vector (S3) before being processed by the fully connected layer to produce a vector of dimension d .

The [CLS] token at the beginning and [SEP] token at the end are used for token representation. DistilBERT simplifies BERT's parameters by 40% and speeds up inference by 60%. The [CLS] token representation is used for class prediction with the softmax function, as shown in Classifier Equation 4.

$$P_r(c|X) = \text{softmax}(W^T \cdot X) \quad (4)$$

Equation (4) shows the probability of tweets X is classified into classes c by using the Softmax function. In Equation (4), W^T is the transpose of the weight matrix W , which is used to multiply feature vectors X . The embedding vector (S3) is processed with a fully connected layer of dimension 128 for feature reduction before classification, where W is the weight matrix [29].

By applying the Softmax function, the model assigns probability scores to each class, ensuring the sum of all probabilities equals one. The fully connected layer with 128 dimensions refines these features before classification, enhancing the model's ability to differentiate disaster-related categories.

Table 2. Pre-Training Model Parameters

No	Parameter	IndoRoBERTa	IndoDistilBERT
1	Optimizer	Adam	Adam
2	Learning Rate	2e-5	2e-5
3	Batch Size	16	16
4	Max Length	128	128
5	Epoch	5	5
6	MLM Probability	15	15

The architecture and components of the DistilBERT model are shown in Fig. 6. DistilBERT has 6 layers (50%

of BERT) with 66 million parameters, created through the simplification of BERT. This process involves merging information from several BERT layers and transferring weights from deep layers to shallow layers. DistilBERT is more efficient with fewer parameters and higher speed, making it suitable for applications with computational limitations. This model uses one of the two BERT layers, making it smaller and faster without sacrificing performance [29].

F. Classification Performance

In this study, the classification performance stage evaluates the model's ability to classify disaster-related messages on floods, earthquakes, and forest/land fires into three categories: eyewitness, noneyewitness, and don't know. Evaluation is performed using a confusion matrix to calculate accuracy, specificity, and sensitivity, which are useful in assessing the model's predictions [30], [31]. Fig. 7 shows the Confusion Matrix.

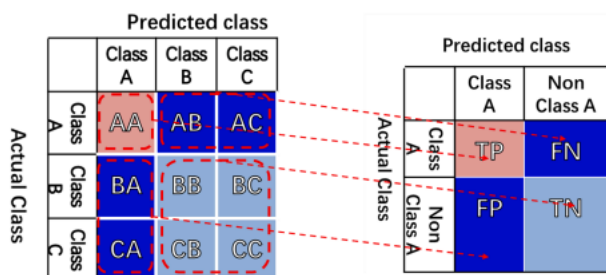


Fig. 7. Confusion Matrix [32]

In Fig. 7, there are four terms that represent the classification results in the confusion matrix, namely [32]: (1). True Positive (TP); (2). True Negative (TN); (3). False Positive (FP); and (4). False Negative (FN).

G. Experiment Settings

This section presents the parameters used in this study. The model's concept and architecture were implemented using Python 3.11.11 and TensorFlow 2.18.0, with hardware support from NVIDIA Tesla T4 and A100 GPUs, along with other additional libraries. This study conducted exploratory experiments on batch size to determine its impact on model performance. The tested batch sizes were 16, 32, 64, 128, and 256, while other factors such as learning rate (2e-5) and the number of epochs (5) were kept constant. This research helps identify the optimal batch size configuration for each disaster data set. The details of the parameters used in the pre-training of MLM IndoRoBERTa and IndoDistilBERT in this study are presented in detail in Table 2.

Table 3 presents in detail the parameters used in the IndoRoBERTa and IndoDistilBERT models in this study.

Table 3. Classification Model Parameters

No	MLM	Classification Model	Parameter
1	No	IndoRoBERTa	Optimizer = Adam Learning Rate = 2e-5 Batch Size = 16, 32, 64, 128, 256 Epoch = 10 Maxt Length = 128
2	Yes		
3	No		
4	Yes		
		IndoDistilBERT	

3. RESULTS

The result of adding special tokens involves the addition of specific tokens, namely the token <s> or [CLS] at the beginning of the sentence, the token </s> or [SEP] at the end of the sentence, and the token <pad> or [PAD] to complete sentences with fewer words than the maximum word count set. Then, each token is encoded based on its index in the vocabulary. The result of adding the special tokens is shown in Table 4.

Table 4. Results of Adding Special Tokens

IndoRoBERTa	IndoDistilBERT
['<s>', 'sak', 'ing', '[CLS]', 'sak', '##ing', 'Gdahsyat', 'nya', 'dahsyat', '##nya', 'kabut', 'Gkabut', 'Gasap', 'asap', 'kebakaran', 'hutan', 'Gkebakaran', 'Ghutan', 'dan', 'lahan', 'di', 'Gdan', 'Glahan', 'Gdi', 'beberapa', 'kabupaten', 'di', 'Gbeberapa', 'riau', 'sampai', 'gubernur', 'Gkabupaten', 'Gdi', 'Gr', '##nya', 'yang', 'iau', 'Gsampai', 'mendukung', 'jokowi', 'Ggubernur', 'nya', 'beberapa', 'hari', 'yang', 'Gyang', 'Gmendukung', 'lalu', 'mengatakan', 'kabut', 'Gj', 'okowi', 'Gbeberapa', 'asap', 'akibat', 'kebakaran', 'Ghari', 'Gyang', 'Gglalu', 'hutan', 'dan', 'lahan', 'di', 'Gmengatakan', 'Gkabut', 'riau', 'belum', 'apa', 'eh', 'Gasap', 'Gakibat', 'kini', 'dia', 'kenapa', 'is', 'Gkebakaran', 'Ghutan', '##pa', 'riau', '##mer', '##de', 'Gdan', 'Glahan', 'Gdi', '##ka', '[SEP]', '[PAD]',..., 'Gr', 'iau', 'Gbelum', '[PAD]']	['<s>', 'sak', 'ing', '[CLS]', 'sak', '##ing', 'Gdahsyat', 'nya', 'dahsyat', '##nya', 'kabut', 'Gkabut', 'Gasap', 'asap', 'kebakaran', 'hutan', 'Gkebakaran', 'Ghutan', 'dan', 'lahan', 'di', 'Gdan', 'Glahan', 'Gdi', 'beberapa', 'kabupaten', 'di', 'Gbeberapa', 'riau', 'sampai', 'gubernur', 'Gkabupaten', 'Gdi', 'Gr', '##nya', 'yang', 'iau', 'Gsampai', 'mendukung', 'jokowi', 'Ggubernur', 'nya', 'beberapa', 'hari', 'yang', 'Gyang', 'Gmendukung', 'lalu', 'mengatakan', 'kabut', 'Gj', 'okowi', 'Gbeberapa', 'asap', 'akibat', 'kebakaran', 'Ghari', 'Gyang', 'Gglalu', 'hutan', 'dan', 'lahan', 'di', 'Gmengatakan', 'Gkabut', 'riau', 'belum', 'apa', 'eh', 'Gasap', 'Gakibat', 'kini', 'dia', 'kenapa', 'is', 'Gkebakaran', 'Ghutan', '##pa', 'riau', '##mer', '##de', 'Gdan', 'Glahan', 'Gdi', '##ka', '[SEP]', '[PAD]',..., 'Gr', 'iau', 'Gbelum', '[PAD]']

IndoRoBERTa	IndoDistilBERT
1, 1, 1, 1, 1, 1, 1, 1, 1, 1,	1, 1, 1, 1, 1, 1, 1, 1, 1, 1,
1, 1, 1, 1, 1, 1, 1, 1, 1, 1,	1, 1, 1, 1, 1, 1, 1, 1, 1, 1,
1, 1, 1, 1,	1, 1, 1, 1, 1, 1, 1, 1, 1, 1,
1, 1, 1, 1, 1, 1, 1, 1, 1, 1,	1, 1, 1, 1, 1, 1, 1, 1, 1, 1,
1, 1, 1, 1, 1, 1, 1, 1, 1, 1,	1, 1, 1, 1, 1, 1, 1, 1, 1,
1, 1, 1, 1,	0,..., 0
1, 1, 1, 1, 1, 0,..., 0	

IndoRoBERTa	IndoDistilBERT
0, 20763, 316, 21499,	3, 4469, 1529, 16031,
327, 22559, 12439,	1538, 15934, 11353, 7797,
8026, 2499, 291, 3581,	3327, 1509, 4162, 1495,
277, 664, 3188, 277,	1788, 2198, 1495, 6704,
385, 3182, 865, 6082,	2016, 3840, 1538, 1510,
327, 292, 2534, 348,	3654, 15071, 1788, 1994,
15750, 664, 968, 292,	1510, 2579, 3270, 15934,
1526, 2123, 22559,	11353, 3028, 7797, 3327,
12439, 2224, 8026,	1509, 4162, 1495, 6704,
2499, 291, 3581, 277,	3093, 3397, 15269, 3244,
385, 3182, 1989, 1978,	1831, 17212, 1744, 2453,
225, 378, 2533, 839,	6704, 4668, 4365, 1914, 1,
21169, 2461, 1029, 385,	2,...,2
343, 318, 3138, 2, 1,...,1	

0, instructing the model to ignore them. This improves efficiency by preventing unnecessary computations on padding. As shown in [Table 6](#), IndoRoBERTa and IndoDistilBERT assign ones to word tokens and zeros to padding. This distinction ensures accurate processing of meaningful text while disregarding placeholders.

These results highlight the importance of choosing an optimal batch size to balance accuracy, specificity, and sensitivity in disaster classification. Further investigations could explore fine-tuning strategies to enhance model performance for specific disaster types.

Classification Model	Dataset	Batch Size	Classification Performance		
			Accuracy (%)	Specificity (%)	Sensitivity (%)
IndoRoBERTa	Floods	16	76.00	84.70	81.71
		32	77.00	86.41	81.71
		64	74.33	79.31	84.21
		128	76.83	81.97	84.83
		256	75.67	88.83	76.00
	Earthquakes	16	75.50	76.22	79.38
		32	71.17	80.34	67.93
		64	72.00	68.21	84.70
		128	73.83	77.01	78.53
		256	72.50	85.98	67.40

Classification Model	Dataset	Batch Size	Classification Performance		
			Accuracy (%)	Specificity (%)	Sensitivity (%)
IndoRoBERTa MLM	Forest Fires	16	82.17	97.46	99.23
		32	85.17	98.48	98.00
		64	84.67	98.99	98.11
		128	83.83	98.99	98.06
		256	85.17	98.49	98.14
	Floods	16	75.00	80.86	87.57
		32	76.33	68.65	95.19
		64	75.67	77.06	89.66
		128	76.00	76.40	86.03
		256	75.67	86.67	78.49
	Earthquakes	16	73.50	62.16	87.50
		32	75.00	74.71	80.95
		64	74.67	72.78	82.89
		128	75.00	78.38	76.17
		256	74.67	70.62	83.59
	Forest Fires	16	82.83	96.97	99.40
		32	84.17	98.49	99.39
		64	84.00	98.48	98.82
		128	83.00	98.48	97.97
		256	84.83	99.50	98.77
IndoDistilBERT	Floods	16	77.17	79.33	85.87
		32	77.83	82.07	83.25
		64	76.67	81.97	80.87
		128	75.17	77.84	79.03
		256	75.83	77.84	84.18
	Earthquakes	16	77.83	74.33	87.44
		32	77.33	85.26	73.30
		64	77.00	85.49	72.16
		128	78.50	75.42	85.71
		256	77.17	79.49	79.19
	Forest Fires	16	86.17	98.48	99.35
		32	86.33	98.98	98.01
		64	86.67	98.48	96.13
		128	87.00	99.49	87.00
		256	87.33	98.99	98.17
IndoDistilBERT MLM	Floods	16	77.83	85.41	82.95
		32	77.33	86.67	78.86
		64	75.67	79.88	83.33
		128	75.67	87.08	75.28
		256	77.33	81.36	83.43
	Earthquakes	16	76.33	82.68	77.72
		32	76.33	72.47	87.18
		64	77.83	82.26	76.65
		128	78.17	82.22	81.05
		256	76.83	81.18	75.92
	Forest Fires	16	85.33	97.98	96.34
		32	85.33	97.98	99.31
		64	82.83	98.98	97.52
		128	85.83	98.48	97.45
		256	86.50	98.99	96.00

4. DISCUSSION

This research was carried out with various batch sizes (16, 32, 64, 128, 256), while other factors such as learning rate ($2e-5$) and number of epochs (5) were kept constant. Based on Table 7 in the results section, it can be concluded that the best performance for each model and dataset with varying batch sizes is summarized in Table 8. Table 8 shows that IndoRoBERTa with a batch size of 256 on the wildfire dataset achieved the best results. While pre-training MLM on the IndoRoBERTa model did

not significantly improve overall performance, it enhanced sensitivity. Meanwhile, the IndoDistilBERT model demonstrated the highest accuracy on the wildfire dataset with a batch size of 256. Pre-training MLM on this model provided good performance, particularly improving specificity and performance on the flood dataset with a batch size of 32. Generally, larger batch sizes, such as 128 and 256, yielded more optimal results compared to smaller batch sizes. The wildfire dataset consistently delivered the best classification performance compared to the flood and earthquake datasets.

Table 8. Classification Performance Results with the Highest Batch Size for Each Dataset and Model

Classification Model	Dataset	Batch Size	Classification Performance		
			Accuracy (%)	Specificity (%)	Sensitivity (%)
IndoRoBERTa	Floods	128	76.83	81.97	84.83
	Earthquakes	16	75.50	76.22	79.38
	Forest Fires	256	85.17	98.49	98.14
IndoRoBERTa MLM	Floods	32	76.33	68.65	95.19
	Earthquakes	32	75.00	74.71	80.95
	Forest Fires	256	84.83	99.50	98.77
IndoDistilBERT	Floods	32	77.83	82.07	83.25
	Earthquakes	128	78.50	75.42	85.71
	Forest Fires	256	87.33	98.99	98.17
IndoDistilBERT MLM	Floods	16	77.83	85.41	82.95
	Earthquakes	128	78.17	82.22	81.05
	Forest Fires	256	86.50	98.99	96.00

Table 9. Best Model Classification Performance Results with Batch Size on Each Disaster Dataset

Classification Model	Dataset	Batch Size	Classification Performance		
			Accuracy (%)	Specificity (%)	Sensitivity (%)
IndoDistilBERT MLM	Floods	32	77.83	85.41	82.95
IndoDistilBERT	Earthquakes	128	78.50	75.42	85.71
	Forest Fires	256	87.33	98.99	98.17

Table 10. Model Performance Results

No	Classification Model	Classification Performance		
		Accuracy (%)	Specificity (%)	Sensitivity (%)
1	IndoRoBERTa	79.17	85.56	87.45
2	IndoRoBERTa MLM	78.72	80.95	91.64
3	IndoDistilBERT	81.22	85.49	89.04
4	IndoDistilBERT MLM	80.83	88.87	86.67

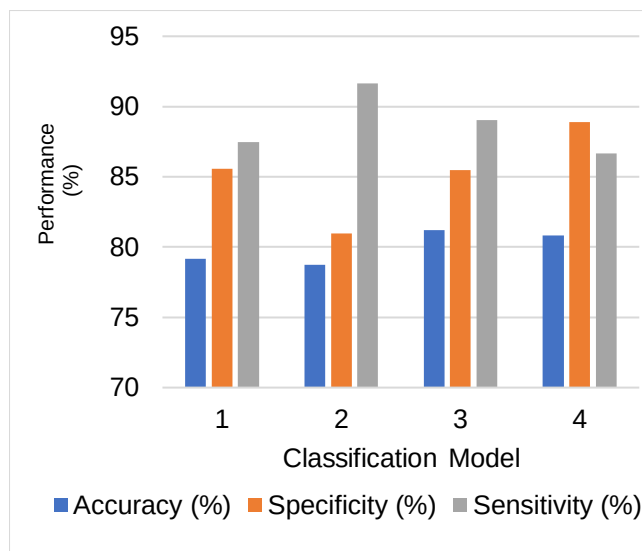


Fig. 8. Comparison of the Performance of Each Classification Model

of 128 and on the wildfire dataset with a batch size of 256. These results indicate that models without pre-training MLM can perform better with appropriate batch configurations and specific dataset types, with the wildfire dataset consistently delivering the highest results. Based on Table 10, the performance of each classification model shows significant differences, both with and without pre-training MLM. The IndoRoBERTa model without MLM performed well; however, after pre-training MLM, sensitivity improved, although there was a slight decrease in accuracy and specificity. Meanwhile, the IndoDistilBERT model without MLM demonstrated the best performance across all metrics. With pre-training MLM, specificity increased, though sensitivity slightly decreased, along with a minor drop in accuracy. These results indicate that pre-training MLM can enhance performance in specific aspects, such as sensitivity for IndoRoBERTa and specificity for IndoDistilBERT. The performance comparison of each classification model is shown in Fig. 8. Upon analyzing the results, it is evident

Table 11. Classification Performance Results from Previous Studies

Dataset	Previous Research Accuracy (%)	Model	Our Research Accuracy (%)
Floods	74	1D CNN [6]	77.83
	77.87	SVM [33]	
Earthquakes	64.13	SVM [34]	78.50
	74.33	1D CNN [6]	
Forest Fires	81.14	1D CNN [6]	87.33
	81.87	2D CNN [35]	

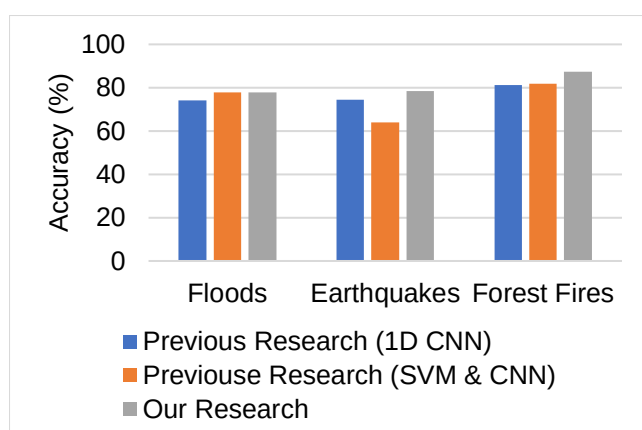


Fig. 9. Classification Performance Results from Previous Research

that Model 3, IndoDistilBERT, outperforms the other models. Table 12 presents the classification model performance from previous studies using the same dataset. Based on the results presented, it is evident that transformer-based models, particularly IndoDistilBERT, exhibit superior performance in classifying disaster-related messages. The improvements in accuracy compared to previous studies indicate that leveraging deep learning approaches can enhance disaster monitoring and response efforts. Future work can focus on optimizing model performance by incorporating additional features, such as temporal and geospatial data, to improve classification accuracy. Moreover, expanding the dataset with real-time social media posts could further validate the model's effectiveness in practical scenarios.

Fig. 9 illustrates the performance comparison between the classification models in this study and those from previous research. The accuracy values presented in the results section represent the highest accuracy achieved for each natural disaster case. As a comparison model, this study uses results from previous research, as shown in Table 11. The selection of comparison models in Table 11 as a baseline is made to ensure performance comparisons remain relevant to the models used in this

study. Based on the one-way ANOVA test results, an F-statistic of 0.7947 and a two-tailed p-value of 0.4941 were obtained, indicating that the difference between previous research and this study is not statistically significant ($p > 0.05$). The conclusion is that the accuracy improvement is not sufficiently significant. Further optimization is needed to achieve a significant difference. However, the accuracy achieved in this study surpasses the results obtained from classification using the 1D CNN model [6]. These findings indicate that IndoDistilBERT and IndoDistilBERT MLM have strong capabilities in text classification, thereby enhancing model performance.

Based on the results obtained, the error analysis shows that the model excels in detecting the 'eyewitness' category due to explicit disaster-related phrases. However, weaknesses are observed in misclassifications between the 'non-eyewitness' and 'don't know' categories, often caused by ambiguous content. A limitation of this study lies in the use of pre-training MLM, which did not always have a significant impact on improving the overall model performance. Although this technique is designed to help the model better understand the language context, the research findings show that the accuracy of models with pre-training MLM tends to be lower compared to those without pre-training. Additionally, pre-training MLM requires high resources, and its inconsistent results across different datasets highlight the need for further evaluation to optimize its performance.

This study contributes to the development of a social media-based disaster early warning system. The resulting model can enhance the efficiency of detecting emergency messages on social media, supporting decision-making in disaster mitigation.

5. CONCLUSION

This study shows that the IndoDistilBERT model outperforms IndoRoBERTa in classifying social media messages about natural disasters. The results suggest that lighter models like IndoDistilBERT can still deliver optimal performance despite having lower complexity. The analysis also reveals that pre-training MLM does not always significantly impact accuracy improvement in classifying natural disaster messages on social media, although this method can enhance sensitivity in IndoRoBERTa and specificity in IndoDistilBERT.

The IndoDistilBERT model demonstrated the best performance in classifying social media messages about natural disasters with an accuracy of 81.22%, specificity of 85.49%, and sensitivity of 89.04%. The IndoRoBERTa model achieved an accuracy of 79.17%, specificity of 85.56%, and sensitivity of 87.45%. Using pre-training MLM on IndoRoBERTa resulted in an accuracy of 78.72%, specificity of 80.95%, and sensitivity of 91.64%, while with IndoDistilBERT and pre-training MLM, accuracy reached 80.83%, specificity 88.87%, and sensitivity 86.67%. These results indicate that pre-training MLM does not outperform the models without additional pre-training in terms of overall performance.

This study also found that the best batch size for each

type of natural disaster dataset is 32 for the flood dataset with an accuracy of 77.83%, specificity of 85.41%, and sensitivity of 82.95%; 128 for the earthquake dataset with an accuracy of 78.50%, specificity of 75.42%, and sensitivity of 85.71%; and 256 for the wildfire dataset with an accuracy of 87.33%, specificity of 98.99%, and sensitivity of 98.17%, which contributed to improving the model's performance.

In future research, evaluations will focus on hyperparameter tuning techniques to optimize the performance of pre-trained models on MLM and other Indonesian BERT variants, including IndoRoBERTa and IndoDistilBERT. The primary goal is to improve model accuracy and overall performance in classifying disaster-related messages on social media.

6. ACKNOWLEDGMENT

This research was supported by funding from the DRTPM Research Program of Indonesia's Ministry of Education, Culture, Research, and Technology. Main Contract Number: 056/E5/PG.02.00.PL/2024, Derivative Contract Number: 1026/UN8.2/PG/2024.

REFERENCES

- [1] M. Fuady, R. Munadi, and M. A. K. Fuady, "Disaster mitigation in Indonesia: between plans and reality," *IOP Conf Ser Mater Sci Eng*, vol. 1087, no. 1, p. 012011, Feb. 2021, doi: 10.1088/1757-899x/1087/1/012011.
- [2] I. Budiman, M. R. Faisal, D. T. Nugraha, Muliadi, M. K. Delimayanti, and S. E. Prastya, "Harvesting Natural Disaster Reports from Social Media with 1D Convolutional Neural Network and Long Short-Term Memory," in *2023 8th International Conference on Informatics and Computing, ICIC 2023*, Institute of Electrical and Electronics Engineers Inc., 2023. doi: 10.1109/ICIC60109.2023.10382045.
- [3] D. Murthy, "Sociology of Twitter/X: Trends, Challenges, and Future Research Directions," 2024, doi: 10.1146/annurev-soc-031021.
- [4] H. Sa'diah *et al.*, "Gender Classification on Social Media Messages Using fastText-base Feature Extraction and Long Short-Term Memory," *Journal of Electronics, Electromedical Engineering, and Medical Informatics*, vol. 6, no. 3, pp. 243–252, Jul. 2024, doi: 10.35882/ijeemi.v6i3.407.
- [5] C. J. Powers *et al.*, "Using artificial intelligence to identify emergency messages on social media during a natural disaster: A deep learning approach," *International Journal of Information Management Data Insights*, vol. 3, no. 1, Apr. 2023, doi: 10.1016/j.jjimei.2023.100164.
- [6] M. R. Faisal, I. Budiman, F. Abadi, D. T. Nugraha, M. Haekal, and I. Sutedja, "Applying Features Based on Word Embedding Techniques to 1D CNN for Natural Disaster Messages Classification," in *2022 5th International Conference on Computer and Informatics Engineering, IC2IE 2022*, Institute of Electrical and Electronics Engineers Inc., 2022, pp. 192–197. doi: 10.1109/IC2IE56416.2022.9970188.
- [7] K. Nemkul, "Use of Bidirectional Encoder Representations from Transformers (BERT) and Robustly Optimized Bert Pretraining Approach (RoBERTa) for Nepali News Classification," *Tribhuvan University Journal*, vol. 39, no. 1, pp. 124–137, Jun. 2024, doi: 10.3126/tuj.v39i1.66679.
- [8] M. S. Sayeed, V. Mohan, and K. S. Muthu, "BERT: A Review of Applications in Sentiment Analysis," Jun. 01, 2023, *Ital Publication*. doi: 10.28991/HIJ-2023-04-02-015.
- [9] V. Nair, D. J. Pareek, and S. Bhatt, "A Knowledge-Based Deep Learning Approach for Automatic Fake News Detection using BERT on Twitter," in *Procedia Computer Science*, Elsevier B.V., 2024, pp. 1870–1882. doi: 10.1016/j.procs.2024.04.178.

- [10] B. Prasetyo, Ahmad Yusuf Al-Majid, and Suharjo, "A Comparative Analysis of MultinomialNB, SVM, and BERT on Garuda Indonesia Twitter Sentiment," *PIKSEL: Penelitian Ilmu Komputer Sistem Embedded and Logic*, vol. 12, no. 2, pp. 445-454, Sep. 2024, doi: 10.33558/piksel.v12i2.9966.
- [11] I. Budiman *et al.*, "Classification Performance Comparison of BERT and IndoBERT on Self-Report of COVID-19 Status on Social Media Porównanie wyników klasyfikacji BERT i IndoBERT w zakresie samodzielnego zgłaszania statusu COVID-19 w mediach społecznościowych," 2024.
- [12] C. H. Lin and U. Nuha, "Sentiment analysis of Indonesian datasets based on a hybrid deep-learning strategy," *J Big Data*, vol. 10, no. 1, Dec. 2023, doi: 10.1186/s40537-023-00782-9.
- [13] A. Conneau *et al.*, "Unsupervised Cross-lingual Representation Learning at Scale," Nov. 2019, [Online]. Available: <http://arxiv.org/abs/1911.02116>
- [14] Y. Liu *et al.*, "RoBERTa: A Robustly Optimized BERT Pretraining Approach," Jul. 2019, [Online]. Available: <http://arxiv.org/abs/1907.11692>
- [15] V. Sanh, L. Debut, J. Chaumond, and T. Wolf, "DistilBERT, a distilled version of BERT: smaller, faster, cheaper and lighter," Oct. 2019, [Online]. Available: <http://arxiv.org/abs/1910.01108>
- [16] G. Z. Nabillah and D. Suhartono, "Personality Classification Based on Textual Data using Indonesian Pre-Trained Language Model and Ensemble Majority Voting," *Revue d'Intelligence Artificielle*, vol. 37, no. 1, pp. 73-81, Feb. 2023, doi: 10.18280/ria.370110.
- [17] M. Joshi, D. Chen, Y. Liu, D. S. Weld, J. Zettlemoyer, and O. Levy, "SpanBERT: Improving Pre-training by Representing and Predicting Spans," Jul. 2019, [Online]. Available: <http://arxiv.org/abs/1907.10529>
- [18] A. Ahmad Aliero, B. Sulaimon Adebayo, H. Olanrewaju Aliyu, A. Gogo Tafida, B. Umar Kangiwa, and N. Muhammad Dankolo, "Systematic Review on Text Normalization Techniques and its Approach to Non-Standard Words," 2023. [Online]. Available: <https://www.researchgate.net/publication/374166354>
- [19] R. Lamsal, M. R. Read, and S. Karunasekera, "CrisisTransformers: Pre-trained language models and sentence encoders for crisis-related social media texts," *Knowl Based Syst*, vol. 296, Jul. 2024, doi: 10.1016/j.knsys.2024.111916.
- [20] L. Putelli, A. E. Gerevini, A. Lavelli, T. Mehmood, and I. Serina, "On the Behaviour of BERT's Attention for the Classification of Medical Reports," 2022. [Online]. Available: <http://ceur-ws.org>
- [21] M. Lukasik, B. Dadachev, G. Simões, and K. Papineni, "Text Segmentation by Cross Segment Attention," Apr. 2020, [Online]. Available: <http://arxiv.org/abs/2004.14535>
- [22] A. Wettig, T. Gao, Z. Zhong, and D. Chen, "Should You Mask 15% in Masked Language Modeling?," Feb. 2022, [Online]. Available: <http://arxiv.org/abs/2202.08005>
- [23] Y. Meng *et al.*, "Representation Deficiency in Masked Language Modeling," Feb. 2023, [Online]. Available: <http://arxiv.org/abs/2302.02060>
- [24] J. Daniel and J. H. Martin, "Speech and Language Processing," 2023.
- [25] Y. Meng *et al.*, "Representation Deficiency in Masked Language Modeling," Feb. 2023, [Online]. Available: <http://arxiv.org/abs/2302.02060>
- [26] A. Areshey and H. Mathkour, "Transfer Learning for Sentiment Classification Using Bidirectional Encoder Representations from Transformers (BERT) Model," *Sensors*, vol. 23, no. 11, Jun. 2023, doi: 10.3390/s23115232.
- [27] R. Khusuma, W. Maharani, and P. H. Gani, "Personality Detection On Twitter User With RoBERTa," *JURNAL MEDIA INFORMATIKA BUDIDARMA*, vol. 7, no. 1, p. 542, Feb. 2023, doi: 10.30865/mib.v7i1.5598.
- [28] B. Richardson and A. Wicaksana, "COMPARISON OF INDOBERT-LITE AND ROBERTA IN TEXT MINING FOR INDONESIAN LANGUAGE QUESTION ANSWERING APPLICATION," *International Journal of Innovative Computing, Information and Control*, vol. 18, no. 6, pp. 1719-1734, Dec. 2022, doi: 10.24507/ijic.18.06.1719.
- [29] H. Adel *et al.*, "Improving Crisis Events Detection Using DistilBERT with Hunger Games Search Algorithm," *Mathematics*, vol. 10, no. 3, Feb. 2022, doi: 10.3390/math10030447.
- [30] I. T. Ali, Y. Rahayu, and A. Setiawan, "J EEE MI i Web Application Based on Machine Learning for Diabetes Detection using Microstrip Resonator and Streamlit," vol. 6, no. 3, pp. 2656-8624, 2024, doi: 10.35882/ijeemi.v6i3.2.
- [31] I. Fatwasauri, "Desain 3 Dimension Baby Incubator Using SketchUp Application Based On Indonesia National Standards," *Indonesian Journal of Electronics, Electromedical Engineering, and Medical Informatics*, vol. 6, no. 3, Aug. 2024, doi: 10.35882/ijeemi.v6i3.6.
- [32] L. Haifeng, H. (Cynthia) Hou, and Z. Gou, "User-centric approach to optimizing thermal comfort in university classrooms: Utilizing computer vision and Q-XGBoost reinforcement learning," *Energy Build*, vol. 323, Nov. 2024, doi: 10.1016/j.enbuild.2024.114808.
- [33] M. K. Delimayanti, R. Sari, M. Laya, M. R. Faisal, Pahrul, and R. F. Naryanto, "The effect of pre-processing on the classification of twitter's flood disaster messages using support vector machine algorithm," in *Proceedings of ICAE 2020 - 3rd International Conference on Applied Engineering*, Institute of Electrical and Electronics Engineers Inc., Oct. 2020. doi: 10.1109/ICAEE50557.2020.9350387.
- [34] S. Monika Nooralifa, M. Reza Faisal, F. Abadi, R. Adi Nugroho, J. A. Yani Km, and K. Selatan, "IDENTIFIKASI OTOMATIS PESAN SAKSI MATA PADA MEDIA SOSIAL SAAT BENCANA GEMPA," vol. 08, no. 2, p. 129.
- [35] M. Reza Faisal, M. Itqan Mazdadi, R. Adi Nugroho, F. Abadi, I. A. Komputer FMIPA ULM Yani St KM, and S. Kalimantan, "EYE WITNESS MESSAGE IDENTIFICATION ON FOREST FIRES DISASTER USING CONVOLUTIONAL NEURAL NETWORK." [Online]. Available: <http://fb.me/6sFXlyEcj>

AUTHOR BIOGRAPHY



Karlina Elreine Fitriani is from Tanah Bumbu, South Kalimantan. Since 2021, she has been studying in the Department of Computer Science at Universitas Lambung Mangkurat. Her current research focuses on data science. Additionally, her final project involves research centered on text classification of disaster-related messages sourced from social media. Through this research, she aims to support disaster monitoring and the development of more effective early warning systems. She can be contacted at email: 2111016120006@mhs.ulm.ac.id.



Mohammad Reza Faisal was born in Banjarmasin. Following his graduation from high school, he pursued his undergraduate studies in the Informatics department at Pasundan University in 1995, and later majored in Physics at Bandung Institute of Technology in 1997. After completing his bachelor's program, he gained experience as a training trainer in the field of information technology and software development. Since 2008, he has been a lecturer in computer science at Universitas Lambung Mangkurat, while also pursuing his master's program in Informatics at Bandung Institute of Technology in 2010. In 2015, he furthered his education by pursuing a doctoral degree in Bioinformatics at Kanazawa University, Japan. To this day, he continues his work as a lecturer in Computer Science at Universitas Lambung Mangkurat. His research interests encompass Data Science, Software Engineering, and Bioinformatics. He can be contacted at email: reza.faisal@ulm.ac.id.



Muhammad Itqan Mazdadi is a lecturer in the Department of Computer Science, Lambung Mangkurat University. His study interest is centered on Data Science and Computer Networking. Before becoming a lecturer, he completed his undergraduate program in the Computer Science Department at Lambung Mangkurat University in 2013. He then completed his master's degree from the Department of Informatics at Islamic Indonesia University, Yogyakarta. Currently, he serves as the Secretary of the Computer Science Department at Lambung Mangkurat University. He is dedicated to fostering a collaborative and innovative learning environment that encourages students to explore and excel in the field of computer science. He can be contacted at email: mazdadi@ulm.ac.id.



Fatma Indriani is a lecturer in the Department of Computer Science at Lambung Mangkurat University. She completed her undergraduate program in the Informatics Department at Bandung Institute of Technology and master's degree at Monash University, Australia in 2012. Her latest education is a doctorate in bioinformatics at Kanazawa University, Japan, 2022. The research fields she focuses on are

Data Science and Bioinformatics. She can be contacted at email: f.indriani@ulm.ac.id.



Dodon Turianto Nugrahadi is a lecturer in Department of Computer Science, Lambung Mangkurat University. His research interest is centered on Data Science and Computer Networking. He completed his bachelor's degree in Informatics Engineering in the UK, Petra, Surabaya in 2004. After that, he pursued a master's degree in Information Engineering at Gajah Mada University, Yogyakarta in 2009. His current area of research revolves around Network, Data Science, Internet of Things (IoT), and network Quality of service (QoS). He can be contacted at email: dodonturianto@ulm.ac.id.

Septyan Eka Prastya is a lecturer in the Department of Information Technology, Sari Mulia University. He completed his undergraduate education in Computer Science, Lambung Mangkurat University in 2013. He also completed his master's education at the Islamic University of Indonesia in 2016. Currently, he is conducting research in the field of software engineering and machine learning. He can be contacted at email: septyan.e.prastya@unism.ac.id.

