

RoBERTa-BiLSTM-Conv1D Deep Learning Model for Detecting Persuasive Content in News

Arya Putra Kurniawan¹, Daniel Siahaan^{2*}, Brian Rizqi Paradisiaca Darnoto³

^{1,2}Institut Teknologi Sepuluh Nopember Keputih, Sukolilo, Surabaya, Jawa Timur

³Universitas Jember Krajan Timur, Sumbersari, Kec. Sumbersari, Kabupaten Jember, Jawa Timur

*Correspondence: daniel@if.its.ac.id

Abstract: The use of persuasive language is one of the defining features of native advertisements. Therefore, detecting persuasive content in news is essential, since native ads often appear disguised as legitimate news articles, it is crucial to identify and filter such content to maintain objectivity and improve the user experience. This study aims to detect news with persuasive content i.e. persuasive news in English language using a natural language processing (NLP) approach. The proposed method incorporates text summarization methods, pre-trained word embeddings, and deep learning models. An additional Conv1D layer has been added to improve the model's performance. The model were trained on an Indonesian news dataset translated into English using Google Translate API. Experimental results show that our proposed RoBERTa-BiLSTM-Conv1D model, outperformed other models, achieving 92% accuracy in identifying persuasive news in English. These findings indicate that the persuasive content detection model can be used for application in mainstream media environments to detect native ads in English language. In the future, the model can incorporate Indonesian and English news as training data to develop a cross-lingual native ads detection model.

Keywords: Deep learning, persuasive content, text summarization, word embedding

1. Introduction

Online media, which has emerged alongside the development of the digital era, has replaced the role of print media [1]. The leading cause of this phenomenon is the ease of accessibility of online media, which can be easily accessed through electronic devices, compared to print media, which can only be accessed by buying in specific locations. In addition to being used to convey objective news, online media can convey subjective advertising. Native Ads leverages online media by presenting ads in the form of news articles to inform specific brands, goods, and services [2]. Native ads content is an advertisement with the same typeface and writing style as a news article, so it is difficult to distinguish it from editorial news articles [3].

One of the characteristics of *native ads* is the use of persuasive language [4]. The information in the *native ads* news article uses persuasive language to attract consumers' attention to the ads it contains [5]. One of the persuasive methods used to attract the reader's attention is to describe the company in a positive light [6]. This practice is considered unfair or deceptive because some companies use persuasive news to attract consumers' attention in dishonest ways [7].

Persuasive content in the news needs to be detected because one of the characteristics of native ads is the use of persuasive language. By detecting persuasive content, native ads can be identified. Native ads differ from advertorials, which are advertisements disguised as news; Native Ads Also disguise ads as news. As a result, readers would apply the native ads (containing the reporter's opinion) as news (containing only facts). Detecting native ads would certainly help news portal administrators ensure that the news displayed is more objective and unbiased.

Currently, English is recognized as a foreign language that is essential for everyone to master to compete effectively on a global scale [8]. Therefore, various online news outlets in Indonesia have offered English versions of their content. The persuasive news must be filtered to ensure the presented news is objective. Currently, no established method exists for identifying persuasive news in English; this study aims to bridge this gap by proposing a method to detect persuasive news, particularly in English. To do this, this study uses the natural language processing method.

Natural Language Processing (NLP) models and examines human language through computational methods [9]. Natural language processing entails designing computational systems and methods to address real-world challenges in interpreting human language [10]. The main areas focus on key challenges in language processing, including language modeling, which involves measuring relationships between naturally occurring words; morphological analysis, which breaks down words into meaningful units and determines their actual parts of speech; syntactic analysis or parsing, which constructs sentence structures that can lead to understanding meaning; and semantic analysis, which aims to extract the meaning of words, phrases, and larger textual elements. Natural language processing can detect persuasive news automatically without human intervention.

Several studies have investigated methods for detecting persuasive content in news articles. Ho et al. [11] applied a hidden advertorial detection method to Chinese social media posts. Seven main features are observed in each social media post; these features are grouped into three categories: background descriptions, product-related information, and building closeness with the audience. The seven features and BERT were used as models to check the effectiveness of this study. The result is that most features can improve the model's ability to detect advertorials in social media posts.

Liu et al. [12] attempted to address the issue of identifying persuasive strategies in online news articles across multiple languages, specifically French, English, German, Polish, Italian, and Russian. Each paragraph of an online news article is translated into five additional languages to enhance the data. The mDeBERTa-v3 model is used to complete the detection task for the persuasion technique. Based on the results obtained, English has the best detection results compared to other languages. In the future, this research should aim to enhance the model's ability to operate on low-resource languages.

Other research by Darnoto et al. [13] detects persuasive content in online news articles that have previously been summarized. Text summarization of news articles is done to identify the main arguments and important points by shortening sentences that do not contain persuasive content. At the same time, the classification model is used to identify sentences that contain persuasive elements. To measure the performance of text summarization in this problem, a comparison was made between the model that used the text summarization method and the model that did not. The text summarization process uses the Latent Semantic Analysis (LSA) and TextRank methods. The models used for classification are CNN and BiLSTM. The results show that the text summarization method performs better, especially in the TextRank-BERT-BiLSTM model. To improve the model's performance, integrating additional features such as article author, publication date, or number of shares on social media is advised.

Each previous experiment on detecting persuasive news has utilized a word embedding model, a deep learning model, or a combination of both models. Word embedding has proven useful across various natural language processing applications [14]. Pre-trained word embeddings are fixed-length vector representations that encode general semantic meanings and linguistic structures found in natural language [15]. Pre-trained word embeddings are typically favored because they capture valuable prior knowledge about word semantics, which can be effectively transferred to specific applications and fine-tuned for various downstream tasks [16].

Deep learning has been widely used in NLP for text classification [17]. There are various architectural variants of deep learning, primarily defined by the types of layers, neural components, and connections they incorporate. One of the widely used deep learning methods for text classification is the Convolutional Neural Network (CNN). CNN processes various sequence

segments in parallel using different convolutional kernels [18]. As a result, they are widely employed in natural language processing tasks such as text classification.

This study proposes a persuasive news detection model utilizing text summarization techniques based on deep learning and word embedding features. Text summarization aims to summarize news documents and present text data that contains the main ideas of news texts, along with persuasive content. Pre-trained word embedding is used for vector representation to improve the model's classification capabilities. The deep learning approach is selected because, from various research fields, deep learning has surpassed traditional machine learning methods in performance [19]. We conducted several testing scenarios for the model, from which the best-performing scenario would be fine-tuned. The model was tuned by adding a layer to the deep learning architecture. We expected the added layer would improve the model's capability to recognize persuasive news. learning architecture. We expected the added layer would improve the model's capability to recognize persuasive news.

2. Research Method

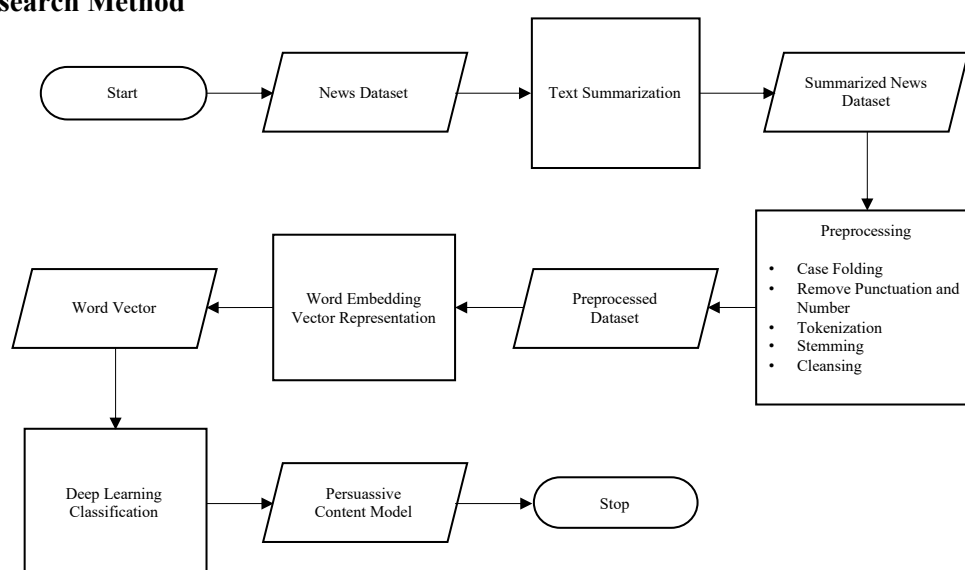


Figure 1 Deep learning architecture for detecting persuasive news.

The Architecture of Persuasive News Detection

Error! Reference source not found. shows the deep learning architecture used in this study to detect persuasive news using the NLP approach. The first step is to prepare a dataset for detecting persuasive news. The dataset contains persuasive news and regular news. Each news needs to be summarized. With the news data summarized, this study processes the data using various preprocessing techniques completed prerocessed news is then converted into word vectors using word embedding models. The word vectors are then used in the classification model using a deep learning approach. After determining which scenario model performs best, the next process is fine-tuning the model.z

2.1 Dataset

This study uses dataset published by Darnoto et al. [20]. The dataset consists of 1708 news articles, with 854 labeled as regular news and 854 labeled as persuasive news. The news in the dataset is still written in Indonesian. To create a persuasive news classification model in English, it is necessary to conduct a new translation of each news article from Indonesian to English. After translating the news into English, the news is summarized.

2.2. Text Summarization

A text summarization process is carried out to summarize the news in the dataset, simplifying the extracted features. This study used the TextRank algorithm [21] for the online news synthesis. TextRank summarizes text using a graph structure, where a word or phrase is represented as a

node, and the relationships between nodes reflect the similarity of meaning through the weights assigned to their edges. A comparison between news that has not been summarized and that has

Table 1 Comparison of Initial News and Summarized News

Raw News	Data Summarized News Data
JAKARTA - Before the digital era, having basic reading and writing skills may be enough. But now, everyone who is connected to the internet is also required to have digital skills. In order to accelerate Indonesia's digital transformation program and encourage the improvement of information and communication technology knowledge of the Indonesian people, efforts and strategies are needed to maximize digital literacy. In this regard, the Ministry of Communication and Information and Cyber-creation will launch the Digital Literacy Curriculum and Module, tomorrow, Friday (16/4/2021). Digital literacy aims to educate the public to think critically about the use of information and communication technology in daily life. To achieve this goal, the Ministry of Communication and Information Technology together with GNLD Siberkreasi network partners have prepared a Digital Literacy Roadmap 2021-2024 to increase public digital participation, encourage the development of community science in the fields of Information and Communication Technology (ICT) and digital, and encourage the level of digital transformation proficiency in the use of new technologies. The Digital Literacy Roadmap 2021-2024 is formulated in four frameworks in compiling the digital literacy curriculum, namely: Digital Skills, Digital Safety, Digital Safety, and Digital Culture. As well as three frameworks in compiling programs for three components of society: Digital Society, Digital Economy, and Digital Government. This framework was then downgraded to the Digital Literacy program which aims to make people proficient in using technology and digital media at the basic, intermediate, and advanced levels, and not forgetting digital literacy classes for an inclusive society.	To support Indonesia's digital transformation, the Ministry of Communication and Information, along with GNLD Siberkreasi, is launching a Digital Literacy Curriculum and Module on April 16, 2021. This initiative aims to enhance public understanding and critical thinking around information and communication technology (ICT). The Digital Literacy Roadmap 2021-2024 focuses on four main areas: Digital Skills, Digital Safety, Digital Ethics, and Digital Culture, and applies across three societal sectors: Digital Society, Digital Economy, and Digital Government.

Table 2 Preprocessing Steps

Preprocessing Step	Input	Output
Case Folding	To support Indonesia's digital transformation, the Ministry of Communication and Information, along with GNLD Siberkreasi, is launching a Digital Literacy Curriculum and Module on April 16, 2021. This initiative aims to enhance public understanding and critical thinking around information and communication technology (ICT).	to support indonesia's digital transformation, the ministry of communication and information, along with gnld siberkreasi, is launching a digital literacy curriculum and module on april 16, 2021. this initiative aims to enhance public understanding and critical thinking around information and communication technology (ict).
Remove Punctuation and Number	to support indonesia's digital transformation, the ministry of communication and information, along with gnld siberkreasi, is launching a digital literacy curriculum and module on april 16, 2021. this initiative aims to enhance public understanding and critical thinking around information and communication technology (ict).	to support indonesias digital transformation the ministry of communication and information along with gnld siberkreasi is launching a digital literacy curriculum and module on april this initiative aims to enhance public understanding and critical thinking around information and communication technology ict
Stopword Removal	to support indonesias digital transformation the ministry of communication and information along with gnld siberkreasi is launching a digital literacy curriculum and module on april this initiative aims to enhance public understanding and critical thinking around information and communication technology ict	support indonesias digital transformation ministry communication information gnld siberkreasi launching digital literacy curriculum module april initiative aims enhance public understanding critical thinking around information communication technology ict
Tokenization	support indonesias digital transformation ministry communication information gnld siberkreasi launching digital literacy curriculum module april initiative aims enhance public understanding critical thinking around information communication technology ict	["support", "indonesias", "digital", "transformation", "ministry", "communication", "information", "gnld", "siberkreasi", "launching", "digital", "literacy", "curriculum", "module", "april", "initiative", "aims", "enhance", "public", "understanding", "critical", "thinking", "around", "information", "communication", "technology", "ict"]
Lemmatization	["support", "indonesias", "digital", "transformation", "ministry", "communication", "information", "gnld", "siberkreasi", "launching", "digital", "literacy", "curriculum", "module", "april", "initiative", "aims", "enhance", "public", "understanding", "critical", "thinking", "around", "information", "communication", "technology", "ict"]	["support", "indonesia", "digital", "transformation", "ministry", "communication", "information", "gnld", "siberkreasi", "launch", "digital", "literacy", "curriculum", "module", "april", "initiative", "aim", "enhance", "public", "understanding", "critical", "thinking", "around", "information", "communication", "technology", "ict"]

been summarized is presented in Table 1. It can be seen from the example that there is a significant reduction in the news text used to create a classification model. The results of the The TextRank news summary focuses on maintaining the main points of the news without compromising its essence, ensuring that important features are not lost

2.3 Preprocessing

The next process preprocesses the synthesized news, using the NLTK library. This process would clean the data before it could be represented as word vectors. Table 2 outlines the five stages of preprocessing applied to the news data. The first processing stage is case folding, which is carried out to convert documents into lowercase letters so that differences are not detected due to the use of capital letters or lowercase letters. Then, removing punctuation and numbers because punctuation and numbers have no meaning or significance in the text analysis process. Then, stopwords removal is performed, which removes common words that frequently appear in sentences. Then, tokenization is carried out on the text, dividing it into pieces of words that compose it. The goal is for the word pieces to be considered separate entities and have a value in a matrix. Finally, the lemmatization process is carried out, which reduces words to their basic form.

2.4 Word Embedding

Algorithm 1 BERT Parameter

1	model = BERTModel(BERT)
2	tokens = tokenize_text(processed_text)
3	tokens.add(CLS,ESP)
4	token_ids = convert_tokens_to_ids(tokens)
5	token_ids.generate(attention_mask)
6	word_embeddings = model(token_ids)

Algorithm 2 BiLSTM Parameter

1	model = Sequential()
2	model.add(Bidirectional(LSTM(units=100,return_sequences=True),input_shape=(sequence_length, features)))
3	model.add(GlobalMaxPooling1D())
4	model.add(Dropout(rate=0.2))
5	model.add(Dense(units=2, activation='sigmoid'))
6	model.compile(optimizer='adam', loss='binary_crossentropy', metrics=['accuracy'])
7	model.fit(X_train, y_train, epochs=50, batch_size=128, validation_data=(X_val, y_val))

Before creating the classification model, the data must be transformed into word vectors, as computers can only process data in numerical form. The transformation was carried out using word embedding method. Word embedding methods can be classified into three main categories: conventional, distributional, and contextual models. This study uses contextual models with parameters from previous research [13], which utilizes two-word embeddings, i.e. Bidirectional Encoder Representations from Transformers (BERT) and Robustly Optimized BERT Pre-training Approach (RoBERTa). BERT is a word embedding model that utilizes transfer learning and is built on the transformer architecture [22]. Algorithm 1 describes the steps to using BERT. The process begins by initializing the pre-trained BERT model, followed by the tokenization of text data using the BERT tokenizer. This process involves adding a special BERT token, converting it into an ID token, and creating an attention mask. Then, embed the word BERT by inserting the completed, processed token into the BERT model. BERT provides a representation of the embeddings for each token and an embedding that represents the entire input. The BERT model consists of 124 million parameters, 12 layers, and a hidden size of 768.

RoBERTa is a variation of BERT that uses dynamic masking, which involves generating a new masking pattern each time a sequence is fed into the model [23]. This approach becomes essential when pre-training over more steps or using larger datasets. The RoBERTa model includes 12 layers, 522 million parameters, a hidden size of 768, and 12 attention heads. Both models perform well in text classification using Recurrent Neural Network (RNN) [20]. BERT

and RoBERTa have also been used with satisfactory results in fake news classification research [21].

2.5 Classification

Deep learning is a machine learning concept based on an artificial neural network [24]. Deep learning models often outperform shallow machine learning methods and conventional data analysis techniques. Deep learning is particularly beneficial in areas with extensive, high-dimensional data, making deep neural networks superior to machine learning methods in the majority of tasks that involve processing text, images, videos, and audio [25]. The models used in this study are BiLSTM and CNN.

Bidirectional Long Short-Term Memory (BiLSTM) is a *recurrent neural network* that processes data in two directions and works on two hidden layers [26]. BiLSTM has proven to excel at NLP tasks [27]. Algorithm 2 describes the layers and parameters used in BiLSTM. After the model initiation, a BiLSTM layer is added to process the sequence data, with a parameter of 100 units. After that, the GlobalMaxPooling1D layer is added to reduce the *output dimensions*. Layer Dropout is added for regularization at a rate of 0.2. Then, there is the Dense layer, which is used for binary classification with a sigmoid activation function. The model is compiled with the Adam optimizer, the 'binary_crossentropy' loss function, and the 'accuracy' metric. Finally, the models were trained on data and labels using 50 epochs and a batch size of 128.

$$\text{Accuracy} = \frac{\text{TP} + \text{TN}}{\text{TP} + \text{FP} + \text{TN} + \text{FN}} \quad (1)$$

$$\text{Precision} = \frac{\text{TP}}{\text{TP} + \text{FP}} \quad (2)$$

$$\text{Recall} = \frac{\text{TP}}{\text{TP} + \text{FN}} \quad (3)$$

$$\text{F1 Score} = 2 \times \frac{\text{Recall} \times \text{Precision}}{\text{Recall} + \text{Precision}} \quad (4)$$

$$\text{TPR} = \frac{\text{TP}}{\text{TP} + \text{FN}} \quad (5)$$

CNN is an artificial neural network that processes information in a feed-forward manner [28]. The CNN classification model has the same four layers as the BiLSTM model, but CNN uses an additional Conv1D layer as a replacement for the LSTM layer. Unlike LSTM models that rely on sequential contexts, Conv1D focuses more on recognizing local features. The parameter used in the CNN model is generally the same as the BiLSTM model. The Conv1D layer in this model uses 32 filters, a kernel size of 3, and a relu activation mode.

2.6 Evaluation Metrics

We employed four performance evaluation metrics to compare the performance of various methods, including accuracy, precision, recall, and F1 score. Accuracy measures how much of the prediction is correct (TP+TN) from the total number of samples (TP+FP+TN+FN). Precision measures how much of the positive class prediction (TP) is correct from the positive class predictions (TP + FP) sum. Recall measures how much the prediction of the correct positive class (TP) exceeds the total number of all positive classes (TP+FN). The F1-score measures the harmonic mean of recall and precision, providing a balanced measure between recall and precision. Equations 1 to 6 display the formula for the performance evaluation metrics.

We also utilized the AUC-ROC curve, a graph that displays the model's performance at all classification thresholds, with the true positive rate (TPR) on the y-axis and the false positive rate (FPR) on the x-axis. If the resulting curve is close to the baseline or a 45-degree diagonal line in the ROC space, then the results are less than satisfactory. Conversely, the result is considered satisfactory if the curve is close to the upper left corner or coordinate point (0.1) in the ROC space.

3. Results and Discussion

For detecting persuasive news in English-translated Indonesian news, this study employs the NLP approach, utilizing two-word embedding methods — BERT and RoBERTa — and two deep learning classifier models: CNN and BiLSTM. Based on the selected models, we devised four different scenarios, which are (1) CNN-BERT, (2) BiLSTM-BERT, (3) CNN-RoBERTa, and (4) BiLSTM-RoBERTa.

3.1 Model Evaluation

Table 3 presents a comparison of performance evaluation metrics across various scenarios in this study. It shows that CNN models performed less satisfactorily than BiLSTM models. From a word embedding perspective, RoBERTa outperforms BERT. In conclusion, RoBERTa-BiLSTM was the best-performing model out of the four scenarios. To maximize the performance of the model, we perform parameter tuning for the best-performance deep learning model, which is BiLSTM. We also examined whether adjusting the parameters would improve results or if the reverse is true.

Table 3 Performance Comparison Between Different Scenarios

Scenario	Accuracy	Precision	Recall	F1-Score	AUC
CNN-BERT	0.9	0.9	0.9	0.9	0.9
BiLSTM-BERT	0.89	0.89	0.89	0.89	0.89
CNN-RoBERTa	0.86	0.86	0.86	0.86	0.86
BiLSTM-RoBERTa	0.91	0.91	0.91	0.91	0.91

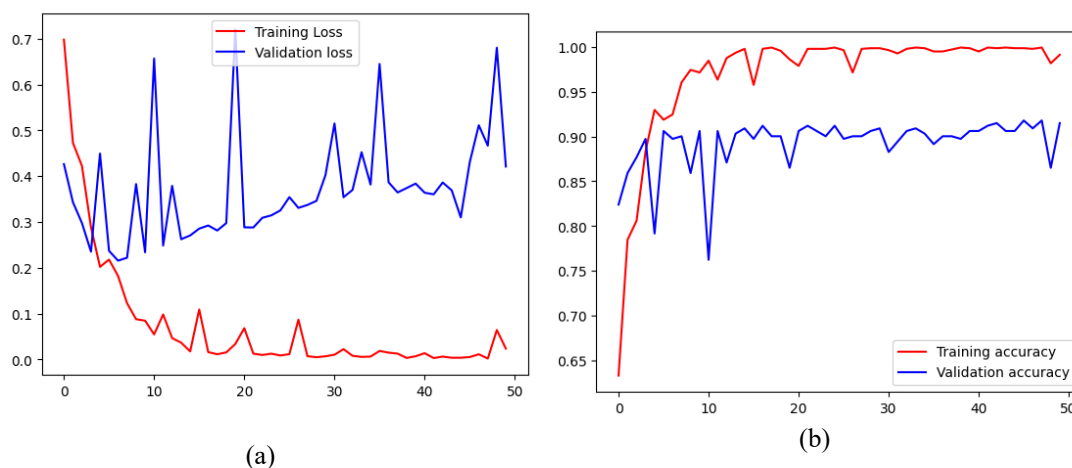


Figure 2 BiLSTM-RoBERTa-Conv1D Training and Validation (a) Loss (b) Accuracy

For the parameter tuning, this study adds an attention mechanism and Conv1D layer to the BiLSTM model [29]. The Conv1D layer used 64 filters, a kernel size of 7, and a real activation mode. **Table 4** shows different performances of BiLSTM models after parameter tuning. The results show that the performance of the BiLSTM-RoBERTa model has decreased due to parameter tuning. Out of three parameter tuning scenarios by adding the Conv1D layer and or adding an attention mechanism, only when adding the Conv1D layer does the BiLSTM – RoBERTa model increase performance from 91% to 92%. **Figure 2** (a) shows that the model validation accuracy is low at epoch 10. It also shows that the validation loss is high on epoch 10,

20, 37, and 50. The training loss of the BiLSTM–Roberta model is lower than its validation loss, suggesting that the model is overfitting the training data rather than learning general patterns, which makes it ineffective on unseen data. Similarly, the higher training accuracy than validation accuracy in **Figure 2 (b)** indicates that the data is overfitting. Specifically, at epoch 10, the training accuracy is 98%, while the validation accuracy is 86%. **Figure 3** shows the ROC curve of the BiLSTM–RoBERTa–Conv1D model, with the resulting curve lying near the (0, 1) point in the ROC space, indicating strong model performance. On the other hand, the BiLSTM–BERT–Conv1D–Attention Mechanism model achieved an increase in performance of almost 2%, from 89% to 91%, by incorporating the Conv1D layer and attention parameter. **Figure 4 (b)** shows that the model validation accuracy remains stable at approximately 91% until epoch 50. The BiLSTM–BERT validation loss is also stable compared to the BiLSTM – RoBERTa model's performance volatility. Although the maximum accuracy of

the BiLSTM–BERT model is lower than that of the BiLSTM–RoBERTa model, it demonstrates that the model performs more stably throughout the epochs. In the end, similar to the BiLSTM–RoBERTa–Conv1D model, the BiLSTM–BERT–Conv1D–Attention Mechanism model also suffers from overfitting data because there are differences between the training and validation loss, as shown in **Figure 4 (a) and (b)**.

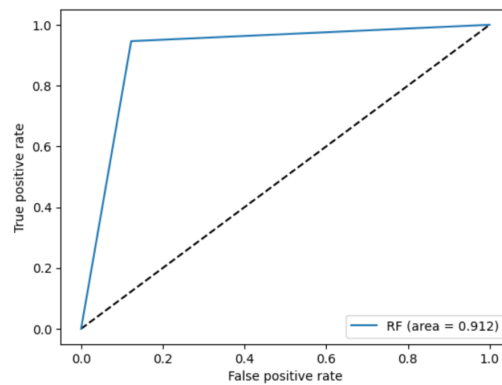


Figure 3 BiLSTM–RoBERTa ROC Curve

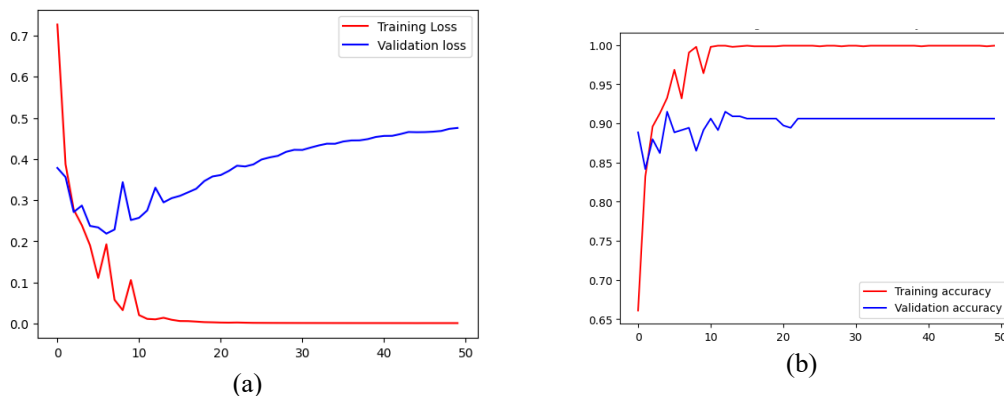


Figure 4 BiLSTM–BERT–Conv1D–Attention Mechanism Training and Validation (a) Loss (b) Accuracy

3.2 Discussion

RoBERTa–BiLSTM–Conv1D model proposed in this research performed better compared to other model for detecting persuasive content in news. The Conv1D layer functions as a local feature extractor that enhances the detection of important phrases or n-gram patterns that may be poorly captured by BiLSTM, thereby efficiently enhancing the representation of RoBERTa without disturbing the existing contextual structure. As previously evaluated, the proposed

model's accuracy of 92% has left an 8% chance for the model to predict data with the wrong label. Error! Reference source not found. illustrates an example of regular news being falsely predicted as persuasive news and vice versa. As is shown in the falsely predicted persuasive news, our model could be failed to predict the correct label mainly because of the sentence "deepen understanding of Indonesia's coffee production from cultivation to consumption—and highlights its diverse varieties," as this sentence is highly persuasive. Primarily, the sentence "highlights its diverse varieties" can pique the model's understanding of this sentence, thereby enticing the reader's curiosity. As the article mentions, the sentence is meant for the American consumer. For the persuasive news falsely predicted as regular news, we concurred that our model failed to predict the correct label due to the lack of direct sentences that persuaded the reader. As evident, the main idea of the article is the promotion of the Human Initiative's "InshaAllah Qurban 2020" program. This article lacks sentences with a direct, persuasive intent to encourage the reader to donate to the program. The suggestion for donating to the program is implicitly made in the sentence "reaffirmed its commitment to connecting those who can give with those in need, both now and in future qurbani events," which directs readers to donate in the future, although it is not explicitly mentioned. We concurred that our model's inability to capture this meaning is because of a failure to detect the article's implicit meaning. This problem can be alleviated by diversifying the train data by adding more implicit content and experimenting with different word embedding to see if there is a pre-trained embedding that can detect implicit content more stably.

The accuracy comparison between our proposed model (BiLSTM-RoBERTa-Conv1D) and models from previous works is shown in Table 6. Table 6 highlights previous studies most relevant to our work, selected based on their methodologies and subject matter. Our study is distinct in that it exclusively uses an English dataset consisting of persuasive news articles, an area that has not been explored in earlier research. We introduced a novel approach that combines text summarization with deep learning for classifying persuasive news articles. Additionally, we employed a fine-tuning method to enhance our model's performance. Despite its originality, the effectiveness of our method is evident; Table 6 presents its accuracy, which surpasses that of existing research in news article classification. The bold numbers in the table represent results on the English news dataset.

Table 5 Example of Data Predicted as False Persuasive News and False Regular News

False Predicted Persuasive News	False Predicted Regular News
The Indonesian Trade Attaché in Washington, D.C. held a screening of the documentary film <i>Legacy of Java</i> at the Indonesian Embassy on August 23, 2024, to promote Indonesian coffee in the U.S. market. Trade Attaché Ranitya Kusumadewi emphasized that the event helps introduce the unique flavors and high quality of Indonesian coffee to American consumers, who average three cups of coffee per day. The film aims to deepen understanding of Indonesia's coffee production—from cultivation to consumption—and highlights its diverse varieties. Directed by Budi Kurniawan, <i>Legacy of Java</i> builds on his earlier documentary <i>Aroma of Heaven</i> and explores the connection between sustainability and coffee cultivation in Java. The event also included coffee tastings and discussions on sustainability by Dua DC Coffee and The Klasik Beans Cooperative from Garut, West Java. Kusumadewi stated that the Embassy will regularly host coffee-related events to promote Indonesian specialty coffee and build networks among enthusiasts. In 2023, the U.S. was Indonesia's largest coffee export destination, with shipments worth \$215.96 million, making Indonesia the 10th-largest coffee supplier to the U.S.	Human Initiative expressed gratitude to donors who entrusted their qurbani through the InshaAllah Qurban 2020 program, which concluded on August 3, 2020, with a closing ceremony. According to President Tomy Hendrajati, around 236,760 beneficiaries across 19 Indonesian provinces and 8 countries received qurbani meat. The domestic distribution reached cities and rural areas, including Jakarta, Bandung, Surabaya, Yogyakarta, and Medan, among others. Internationally, qurbani was delivered to countries such as Myanmar, Syria, Palestine, Uganda, Somalia, Tanzania, Kenya, and Ethiopia. The distribution process followed Covid-19 health protocols and prioritized remote areas and healthcare workers, many of whom could not celebrate Eid al-Adha with their families. Human Initiative reaffirmed its commitment to connecting those who are able to give with those in need, both now and in future qurbani events.

As shown in previous section, the difference of test accuracy compared to the training accuracy is not far, proving that the persuasive news detection model can be used to detect persuasive content on English news portals. The sharp difference between training loss and accuracy loss of RoBERTa-BiLSTM-Conv1D model exhibits the overfitting problem. In future works, investigating solutions to overfitting can help assess their potential to enhance the performance of detection models. Data augmentation methods can also be applied to the dataset to enhance the data's quality and diversity, particularly by incorporating implicit news articles.

Table 6 Comparison of Accuracy Between Our Proposed Model and Alternative Model

Scenario	Accuracy
RoBERTa _{base} FakeNewsNet Dataset [30]	0.61
RoBERTa Indonesia Fake News [31]	0.83
BiLSTM Arabic Fake News Dataset [32]	0.85
BiLSTM-RoBERTa for Detection of Persuasive Content [13]	0.87

4. Conclusions

This study introduces a new architecture for detecting English language persuasive news, comparing different combinations of word embeddings and deep learning models. Specifically, the study evaluates the effectiveness of two word embedding techniques (BERT and RoBERTa) and two deep learning models (CNN and BiLSTM). The results from various combinations of these approaches were assessed for their performance in persuasive news detection. In addition, we parameter-tuned the best model based on the assessment. Parameter tuning is done by adding an attention mechanism or an additional layer.

Findings revealed that models using the BiLSTM approach consistently outperformed the CNN-based model. Notably, the combination of RoBERTa and BiLSTM achieved a high accuracy of 91% in identifying persuasive news. We fine-tuned the model's performance by adding the Conv1D layer. Finally, we improved the model's accuracy to 92% for our proposed BiLSTM-RoBERTa-Conv1D model, surpassing the accuracy of previous studies relevant to news articles. We also found that although BERT word embedding performs less accurately than RoBERTa, it has performed more stably throughout the epoch. In this research, we encountered an overfitting problem with our BiLSTM model. The study confirms that deep learning models can effectively distinguish between persuasive content and factual news. The model could also be applied in a web portal to detect news with persuasive content, allowing readers to have a better user experience while reading the news

References

- [1] M. Y. Saragih and A. I. Harahap, "The Challenges of Print Media Journalism in the Digital Era," *Budapest International Research and Critics Institute (BIRCI-Journal) : Humanities and Social Sciences*, vol. 3, no. 1, 2020, <https://doi.org/10.33258/birci.v3i1.805>
- [2] B. Eyada and A. Milla, "Native Advertising: Challenges and Perspectives," *Journal of Design Sciences and Applied Arts*, vol. 1, no. 1, 2020, <https://doi.org/10.21608/jdsaa.2020.70451>
- [3] M. Dahlén and M. Edenius, "When is advertising advertising? Comparing responses to non-traditional and traditional advertising media," *Journal of Current Issues and Research in Advertising*, vol. 29, no. 1, 2007, <https://doi.org/10.1080/10641734.2007.10505206>
- [4] B. R. P. Darnoto, D. Siahaan, and D. Purwitasari, "Deep Learning for Native Advertisement Detection in Electronic News: A Comparative Study," in *Proceedings - 11th*

- Electrical Power, Electronics, Communications, Control, and Informatics Seminar, EECCIS 2022*, 2022, <https://doi.org/10.1109/EECCIS54468.2022.9902953>
- [5] M. D. Molina, S. S. Sundar, T. Le, and D. Lee, "'Fake News' Is Not Simply False Information: A Concept Explication and Taxonomy of Online Content," *American Behavioral Scientist*, vol. 65, no. 2, 2021, <https://doi.org/10.1177/0002764219878224>
- [6] I. D. Romanova and I. V. Smirnova, "Persuasive techniques in advertising," *Training, Language and Culture*, vol. 3, no. 2, 2019, <https://doi.org/10.29366/2019tlc.3.2.4>
- [7] C. J. Hoofnagle and E. Meleshinsky, "Native Advertising and Endorsement: Schema, Source-Based Misleadingness, and Omission of Material Facts," *Technol Sci*, 2015, <https://techscience.org/a/2015121503/>
- [8] Dzul kifli Isadaud, M. Dzikrul Fikri, and Muhammad Imam Bukhari, "The Urgency Of English In The Curriculum In Indonesia To Prepare Human Resources For Global Competitiveness," *DIAJAR: Jurnal Pendidikan dan Pembelajaran*, vol. 1, no. 1, 2022, <https://doi.org/10.54259/diajar.v1i1.177>
- [9] D. Khurana, A. Koli, K. Khatter, and S. Singh, "Natural language processing: state of the art, current trends and challenges," *Multimed Tools Appl*, vol. 82, no. 3, 2023, <https://doi.org/10.1007/s11042-022-13428-4>
- [10] D. W. Otter, J. R. Medina, and J. K. Kalita, "A Survey of the Usages of Deep Learning for Natural Language Processing," *IEEE Trans Neural Netw Learn Syst*, vol. 32, no. 2, 2021, <https://doi.org/10.1109/TNNLS.2020.2979670>
- [11] M. C. Ho, C. Y. Chuang, Y. C. Hsu, and Y. Y. Chang, "Hidden Advertorial Detection on Social Media in Chinese," in *ROCLING 2021 - Proceedings of the 33rd Conference on Computational Linguistics and Speech Processing*, The Association for Computational Linguistics and Chinese Language Processing (ACLCLP), 2021, pp. 243–251, <https://aclanthology.org/2021.rocling-1.31/>
- [12] G. Liu, Y. R. Fung, and H. Ji, "NLUBot101 at SemEval-2023 Task 3: An Augmented Multilingual NLI Approach Towards Online News Persuasion Techniques Detection," in *17th International Workshop on Semantic Evaluation, SemEval 2023 - Proceedings of the Workshop*, 2023, <https://doi.org/10.18653/v1/2023.semeval-1.227>
- [13] B. R. P. Darnoto, D. Siahaan, and D. Purwitasari, "Automated Detection of Persuasive Content in Electronic News," *Informatics*, vol. 10, no. 4, Dec. 2023, <https://doi.org/10.3390/informatics10040086>
- [14] A. Moreo, A. Esuli, and F. Sebastiani, "Word-class embeddings for multiclass text classification," *Data Min Knowl Discov*, vol. 35, no. 3, 2021, <https://doi.org/10.1007/s10618-020-00735-3>
- [15] D. S. Asudani, N. K. Nagwani, and P. Singh, "Impact of word embedding models on text analytics in deep learning environment: a review," *Artif Intell Rev*, vol. 56, no. 9, 2023, <https://doi.org/10.1007/s10462-023-10419-1>
- [16] D. Erhan, A. Courville, Y. Bengio, and P. Vincent, "Why does unsupervised pre-training help deep learning?," in *Journal of Machine Learning Research*, 2010, <https://jmlr.org/papers/v11/erhan10a.html>
- [17] A. Mohammed and R. Kora, "An effective ensemble deep learning framework for text classification," *Journal of King Saud University - Computer and Information Sciences*, vol. 34, no. 10, 2022, <https://doi.org/10.1016/j.jksuci.2021.11.001>
- [18] Q. Li *et al.*, "A Survey on Text Classification: From Traditional to Deep Learning," *ACM Trans Intell Syst Technol*, vol. 13, no. 2, 2022, <https://doi.org/10.1145/3495162>
- [19] M. M. Taye, "Understanding of Machine Learning with Deep Learning: Architectures, Workflow, Applications and Future Directions," *Computers*, vol. 12, no. 5, 2023, <https://doi.org/10.3390/computers12050091>
- [20] B. R. P. Darnoto, D. Siahaan, and D. Purwitasari, "Electronic News Dataset for Native Advertisement Detection," *Scientific Data*, vol. 12, no. 1, Dec. 2025, <https://doi.org/10.1038/s41597-024-04341-6>

- [21] R. Mihalcea and P. Tarau, "TextRank: Bringing order into texts," in *Proceedings of the 2004 Conference on Empirical Methods in Natural Language Processing, EMNLP 2004 - A meeting of SIGDAT, a Special Interest Group of the ACL held in conjunction with ACL 2004*, 2004, <https://aclanthology.org/W04-3252/>
 - [22] A. Vaswani *et al.*, "Attention is all you need," in *Advances in Neural Information Processing Systems*, 2017. <https://doi.org/10.48550/arXiv.1706.03762>
 - [23] Y. Liu *et al.*, "RoBERTa: A Robustly Optimized BERT Pretraining Approach," 2019., <https://github.com/pytorch/fairseq>
 - [24] C. Janiesch, P. Zschech, and K. Heinrich, "Machine learning and deep learning," *Electronic Markets*, vol. 31, no. 3, pp. 685–695, Sep. 2021, <https://doi.org/10.1007/s12525-021-00475-2>
 - [25] Y. LeCun, G. Hinton, and Y. Bengio, "Deep Learning," *Nature*, vol. 521, 2015, <https://doi.org/10.1038/nature14539>
 - [26] M. Rhanoui, M. Mikram, S. Yousfi, and S. Barzali, "A CNN-BiLSTM Model for Document-Level Sentiment Analysis," *Mach Learn Knowl Extr*, vol. 1, no. 3, pp. 832–847, Sep. 2019, <https://doi.org/10.3390/make1030048>
 - [27] W. Yin, K. Kann, M. Yu, and H. Schütze, "Comparative Study of CNN and RNN for Natural Language Processing," *arXiv.org*, Feb. 2017, <http://arxiv.org/abs/1702.01923>
 - [28] M. A. Saleem, N. Senan, F. Wahid, M. Aamir, A. Samad, and M. Khan, "Comparative Analysis of Recent Architecture of Convolutional Neural Network," *Math Probl Eng*, vol. 2022, 2022, <https://doi.org/10.1155/2022/7313612>
 - [29] G. Liu and J. Guo, "Bidirectional LSTM with attention mechanism and convolutional layer for text classification," *Neurocomputing*, vol. 337, 2019, <https://doi.org/10.1016/j.neucom.2019.01.078>
 - [30] J. Alghamdi, Y. Lin, and S. Luo, "A Comparative Study of Machine Learning and Deep Learning Techniques for Fake News Detection," *Information (Switzerland)*, vol. 13, no. 12, 2022, <https://doi.org/10.3390/info13120576>
 - [31] S. F. N. Azizah, H. D. Cahyono, S. W. Sihwi, and W. Widiarto, "Performance Analysis of Transformer Based Models (BERT, ALBERT, and RoBERTa) in Fake News Detection," in *2023 6th International Conference on Information and Communications Technology, ICOIACT 2023*, 2023. <https://doi.org/10.1109/ICOIACT59844.2023.10455849>
 - [32] K. M. Fouad, S. F. Sabbeh, and W. Medhat, "Arabic fake news detection using deep learning," *Computers, Materials and Continua*, vol. 71, no. 2, 2022, <https://doi.org/10.32604/cmc.2022.021449>
 - [33] A. Praseed, J. Rodrigues, and P. S. Thilagam, "Hindi fake news detection using transformer ensembles," *Eng Appl Artif Intell*, vol. 119, Mar. 2023, <https://doi.org/10.1016/j.engappai.2022.105731>
-