

CHAPTER 1

INTRODUCTION

Research on cross-lingual fake news detection aims to address the challenges raised by language diversity in global information. The main challenge when using different languages is to ensure that fake news detection models can accurately understand and analyze news content without losing meaning or context. This is important as most fake news detection datasets only support one language at a time, making it difficult to apply in multilingual environments [1]. Also, the development of methods for detecting fake news in low-resource languages is limited by the lack of training data [2]. To address this issue, cross-lingual research utilizes methods such as cross-lingual embedding or transfer learning that enable cross-lingual comparison and analysis without translating the entire dataset. Research [2] shows cross-lingual knowledge transfer to be a focus of research as it can help overcome the limitations of training data in the target language. Then how research [3] further explains that cross-lingual embedding can help deal with cross-lingual challenges and transfer learning to align texts from different languages into the same vector space, thus enabling interlanguage comparison and analysis without translating the entire dataset. This approach emphasizes language diversity and the need for models that can adapt to various multilingual datasets.

In line with these challenges, the widespread use of online social media presents both opportunities and difficulties in terms of news spread. Social media platforms disseminate news from diverse sources, including real and fake news, significantly shaping our perspectives on rapidly circulating information. However, this information is not always reliable, and it becomes increasingly difficult to assess the accuracy of news content. As discussed in recent studies, social media platforms target users' emotions through news dissemination, often manipulating perceptions [4]. Consequently, distinguishing between real and fake news on these platforms has become a pressing issue in research. Due to the unique characteristics of social networks and the strategies used by fake news creators, existing detection algorithms may be ineffective, highlighting the need for more robust models that can account for not only the content but also the context of the information [5].

Various approaches have been explored to detect fake news, with deep learning models being used either independently or in hybrid configurations. One such study [6] used the CNN model for detecting fake news in Chinese, incorporating TF-IDF feature extraction and TextRank keyword extraction methods. Despite not achieving the highest accuracy, the CNN model performed admirably, reaching an accuracy of 94%. Furthermore, CNN models have been extensively tested in both standalone and hybrid settings, demonstrating robust performance in fake news detection. A hybrid approach [7], combining CNN and LSTM models with dimensionality reduction techniques like Principal Component Analysis

(PCA) and chi-square, achieved an impressive accuracy of 97.8%. This highlights the potential of hybrid deep learning models, though additional research is required further to refine their performance and adaptability to real-world scenarios.

Building on these insights, the goal of this research is to address the limitations associated with Indonesian-language datasets for fake news detection. Although Indonesian datasets may be relatively smaller compared to datasets in other languages, they provide a valuable and unique opportunity to focus on a language with a distinct linguistic structure and cultural context. By employing a hybrid learning approach, this study integrates multiple machine learning models with techniques such as transfer learning and MUSE embedding. By utilizing MUSE, models can transfer knowledge learned from one language to another, making it extremely valuable for low-resource languages such as Indonesian. However, cross-lingual fake news detection remains a challenge due to variations in dataset quality, linguistic differences, and limitations in knowledge transfer based on embeddings. These approaches facilitate cross-lingual analysis without necessitating the translation of the entire dataset. Transfer learning enables the model to leverage the knowledge from the source language (Indonesian) and apply it to other target languages, while MUSE embedding allows for aligning word representations across languages, ensuring better contextual understanding. This hybrid approach is expected to enhance the adaptability of fake news detection models across Indonesian and other languages, offering more accurate and robust solutions, especially for languages with limited datasets [8], [9].

To address this challenge, this research proposes a hybrid learning approach that integrates collaborative learning with MUSE embeddings to enhance cross-lingual fake news detection. By training the model with source and target languages within the framework of collaborative learning, this study aims to improve the model's adaptability and cross-lingual generalization with different linguistic structures. Furthermore, this research also explores how linguistic similarities, dataset quality, and embedding effectiveness impact the overall performance of cross-lingual classification.

Furthermore, the use of MUSE embedding as a key component in this study allows word representations from different languages to be in the same space. This greatly assists the model in understanding different language variations, thus strengthening its cross-lingual fake news detection capabilities. MUSE embedding provides a richer context for the model to learn from Indonesian source data and transfer that knowledge to other target languages. With this approach, this research not only contributes to fake news detection but also seeks to develop a more effective methodology to deal with the challenges of research with multiple languages.

Table 1.1: Cross-lingual Advantages and Disadvantages

Reference	Focus	Method	Advantages	Disadvantages
Han, S. (2022) [2]	Cross-lingual fake news detection in low-resource languages	Transfer learning using LSTM-CNN hybrid with multilingual embeddings	Effective in detecting fake news in low-resource languages. It also, Encourages transfer learning for better generalization across languages.	Requires pre-trained models from resource-rich languages. Performance may degrade for languages with significantly fewer resources.
Arif et al. (2022) [10]	Multi-class and cross-lingual fake news detection	Use of multilingual embeddings and joint learning for multi-class classification	Improves cross-lingual transfer by leveraging multilingual data. Then, it can handle multiple languages simultaneously, enhancing versatility.	Complexity increases with the number of languages. Also, data imbalance can affect model performance across different languages.
Mahendra et al. (2018) [11]	Cross-lingual word sense disambiguation	Supervised learning with cross-lingual datasets	Enhances word sense disambiguation by using bilingual corpora. It can also be adapted to a variety of languages with parallel corpora.	Dependence on high-quality bilingual datasets. and limited to word-level disambiguation and may not generalize to other NLP tasks.

Reference	Focus	Method	Advantages	Disadvantages
Van Thin et al. (2023) [12]	Aspect-based sentiment analysis in multiple languages	Zero-shot and joint training strategies with multilingual contextualized embeddings	Allows sentiment analysis in multiple languages without direct supervision for each language. And use of joint training can improve model efficiency.	Requires a high-quality multilingual contextual model. Also, performance may vary depending on the language pair and dataset quality.

1.1 Theoretical Framework

LSTM (Long Short-Term Memory) is used to process text data sequences in an efficient way for fake news detection. LSTM can remember information in the long term, which allows the model to handle further word dependencies in the text. This is important in fake news detection as the context in a news story often requires a deep understanding of the word order present in a sentence or paragraph [13]. LSTM works by using three main components: forget gate, input gate, and output gate—that effectively control the information that is passed or forgotten along the data sequence, allowing for a better understanding of patterns in the text. This is particularly useful in processing long texts often found in fake news, which can have interconnected claims and require long-term context to analyze [7].

CNN is a type of neural network architecture that is very effective in handling grid-shaped data such as images or spatial data. CNNs use convolutional layers to extract features from the input and pooling layers to reduce dimensionality and improve computational efficiency. CNN is applied to recognize patterns or features in data. In the context of fake news detection, CNN can be used to identify patterns or text features related to the truth or fake news [14]. The LSTM-CNN hybrid model is developed for text classification tasks by combining two different types of neural networks, namely LSTM (Long Short-Term Memory) and CNN (Convolutional Neural Network). The model utilizes the power of LSTM to capture long-term dependencies in text sequences, while CNN is used to extract locally relevant features, such as patterns or phrases that may indicate the sentiment or characteristics of the text. This combination is expected to improve the model's performance in text classification, be it for fake news detection or other classification tasks [15].

1.2 Conceptual Framework/Paradigm

In this study, the main variable relations are used in the context of cross-lingual fake news detection using the LSTM and CNN models and the combination of the two models, namely the LSTM-CNN hybrid. As it has been applied in research [16], the LSTM network is effective in modeling long-term dependencies in sequential data. That is because LSTM has a memory that “remembers” past information and makes judgments depending on what has been learned. LSTM captures contextual information from the text by maintaining memory over past inputs, allowing it to differentiate between true and false news based on learned sequential patterns. Then, how CNN works by extracting key local features from the sequential representations to identify relevant textual patterns that are crucial in distinguishing fake news from real news. CNN is particularly effective in capturing n-gram patterns and other local dependencies in text, making it a valuable addition to the model. This integration leverages the capabilities of LSTM’s ability to learn long-term dependencies and CNN’s effectiveness in extracting key features from textual representations. By combining these two models, the hybrid LSTM-CNN architecture enhances the model’s capability to capture complex linguistic patterns, ultimately improving the accuracy of fake news detection.

1.3 Statement of the Problem

Fake news is designed to mislead the public by presenting a narrative that closely resembles actual news and often capitalizes on users’ emotions, such as anger or fear, to increase its spread. Recognizing fake news is challenging, especially when relying purely on text content, as it often uses persuasive yet ambiguous language. While fake news datasets in Indonesian are limited and often underrepresentative, this research explores methods to overcome these limitations through a hybrid learning approach, which allows for more accurate and adaptive models. The limited availability of Indonesian language datasets, which are often classified as low-resource, provides a distinct advantage by encouraging the development of more specialized and efficient models tailored to the local context. Fake news in Indonesian often relates to local contexts, such as Indonesian politics or cultural issues, which are not always covered in other language datasets. Different languages have unique grammatical structures, syntax and writing styles. Models trained in one language may not be able to handle these differences when applied to another language. Fake news is often created to take advantage of country or language-specific social and political differences, so it is important for cross-lingual fake news detection models to be able to understand different contexts in order to effectively detect fake news. This research aims to address these challenges through a hybrid learning approach that is more adaptive and can improve accuracy in detecting fake news, especially in the context of Indonesian.

1.4 Objective and Hypotheses

The objective of this research is to develop a hybrid learning approach for fake news detection that effectively combines joint learning techniques with MUSE embedding to enhance cross-lingual capabilities. This hybrid approach is designed to address the challenges posed by the limitations of Indonesian datasets, which are particularly relevant in the context of low-resource languages. By utilizing transfer learning and advanced embedding techniques, this research aims to overcome the problem of insufficient data for languages with limited resources, offering a more efficient way to tackle fake news detection in such contexts. Additionally, this study intends to enable cross-lingual fake news detection without the need for complete translation of entire datasets, thus streamlining the process and making it more scalable to other languages with fewer resources.

Moreover, this research aims to improve the adaptability and accuracy of fake news detection models when applied to multilingual datasets, especially in the context of languages with significant structural differences, such as Indonesian and other target languages. The hypothesis driving this study is that the integration of MUSE embedding, along with a hybrid learning approach that incorporates transfer learning and multiple machine learning models, will result in a significant improvement in the accuracy of cross-lingual fake news detection. It is expected that this combined approach will outperform traditional single-language fake news detection methods, particularly in scenarios where low-resource languages such as Indonesian are involved. By leveraging the strengths of both transfer learning and cross-lingual embeddings, the research is anticipated to provide richer contextual understanding, better model adaptability, and improved performance across a wide range of multilingual datasets. This approach is expected to significantly advance the field of fake news detection, making it more inclusive and effective for diverse linguistic environments.

1.5 Assumption

This research assumes that the cross-lingual approach with MUSE embeddings can effectively transfer knowledge between languages so that the model can detect fake news with good performance even if it is trained on a dataset in a combined language, including the source language, and tested on another language as a target language. In addition, it is assumed that hybrid LSTM-CNN, LSTM, and CNN models can optimally process multiple languages with different difficulty levels, given the structural variability in the languages used.

1.6 Scope and Delimitation

This research is limited to cross-lingual fake news detection using a hybrid LSTM-CNN model, LSTM, and CNN with MUSE embeddings as a method of mapping words between languages. The dataset used includes Indonesian as the source dataset and Malay, German, Spanish, and English as the target datasets. This research only focuses on text as a source of information.