

Research Paper

Fake News Detection In Code-Mixed Hindi–English Text Using Mbert and Mt5: Cross-lingual, Script aware approach

¹V. Swetha, ²Dr. Sumayya Afreen

¹ PG Scholar, Department of CSE, Stanley College of Engineering & Technology for Women, Abids, Hyderabad, Telangana, India.

²Associate Professor, Department of CSE, Stanley College of Engineering & Technology for Women, Abids, Hyderabad, Telangana, India.

Abstract

The widespread use of social media platforms has transformed the way information is shared and consumed across the world. Although these platforms provide rapid access to news and communication, they have also become major channels for the spread of misleading and fabricated information. Fake news has become a critical challenge that can influence public opinion, disrupt social harmony, and create misinformation among communities. Detecting false information in online environments has therefore become an essential task in natural language processing. The challenge becomes even more complex when the textual content is written in code-mixed languages. In multilingual countries such as India, social media users frequently combine multiple languages within the same sentence, especially Hindi and English. These mixed-language texts often contain informal writing styles, transliteration, abbreviations, and inconsistent grammar. Traditional machine learning models struggle to process such data because they are primarily designed for single-language text analysis. To address this issue, a multilingual deep learning framework based on transformer architectures is utilized for detecting fake news in code-mixed Hindi–English text. The system employs advanced models such as Multilingual Bidirectional Encoder Representations from Transformers (MBERT) and Multilingual Text-to-Text Transfer Transformer (MT5). These models are capable of capturing contextual relationships between words across multiple languages and generating meaningful textual representations. The framework processes social media text through several stages including preprocessing, tokenization, contextual embedding generation, and classification. The transformer models analyze semantic relationships and identify linguistic patterns associated with misinformation. Experimental evaluation demonstrates that transformer-based approaches significantly improve classification accuracy compared to traditional machine learning techniques.

Keywords: Fake News Detection, Code-Mixed Language Processing, Multilingual Transformers, Natural Language Processing, Social Media Analysis.

I INTRODUCTION

The digital revolution has significantly changed the way information is produced and shared. Social media platforms such as Twitter, Facebook, and online news portals enable users to publish and distribute information instantly to a large audience. While this accessibility promotes information sharing, it also facilitates the rapid spread of misleading or fabricated news. Fake news has become a major global concern because it can influence political decisions, manipulate public perception, and create social instability.

Fake news refers to intentionally misleading or false information presented as legitimate news. The rapid circulation of such information on social media platforms makes it difficult for users to distinguish between authentic and fabricated content. As a result, automated methods for detecting misinformation have become an important area of study in the fields of natural language processing and artificial intelligence.

Traditional fake news detection methods rely on machine learning techniques that analyze textual features such as word frequencies, syntactic patterns, and sentiment indicators. Although these techniques can identify certain patterns associated with misinformation, they are limited in their ability to capture deeper contextual meaning in textual data.

The challenge becomes more complex in multilingual environments where users frequently combine multiple languages in a single sentence. In India, for example, social media posts often contain a mixture of Hindi and English words written in Roman script. This form of communication, known as code-mixing, creates difficulties for conventional language processing systems because the linguistic structure does not follow standard grammatical rules.

By leveraging multilingual transformer architectures, it is possible to develop intelligent systems capable of detecting fake news in code-mixed social media content. Such systems can play an important role in maintaining information credibility and reducing the spread of misinformation in digital environments.

II LITERATURE SURVEY

The increasing spread of misinformation on digital platforms has motivated researchers to develop automated techniques for detecting fake news using natural language processing and machine learning approaches. Several studies have investigated different computational models to analyze textual patterns, semantic structures, and contextual information in order to distinguish genuine news from misleading content.

Shu et al. (2017) presented one of the early comprehensive studies on fake news detection in

social media environments. Their work emphasized the role of content features, social context, and user interactions in identifying misinformation. The authors proposed a data mining framework that analyzes both textual information and social network characteristics to improve the detection process [1].

Rashkin et al. (2017) analyzed the linguistic characteristics of deceptive news articles and demonstrated that fake news often contains exaggerated language, emotional expressions, and sensational headlines. Their findings showed that stylistic features extracted from textual content can help machine learning models distinguish between reliable and unreliable news sources [2].

Rubin et al. (2016) explored rhetorical structures and linguistic cues present in deceptive news content. Their research identified patterns such as misleading narratives, exaggerated claims, and emotional persuasion techniques commonly used in fake news. The study suggested that analyzing rhetorical patterns can improve automated misinformation detection systems [3].

Devlin et al. (2019) introduced the Bidirectional Encoder Representations from Transformers (BERT) model, which significantly improved the performance of many natural language processing tasks. BERT uses a bidirectional training approach that allows the model to understand the context of words by analyzing both preceding and succeeding tokens in a sentence. This capability enables more accurate text classification and

semantic analysis compared to traditional models [4].

Vaswani et al. (2017) proposed the transformer architecture, which relies on attention mechanisms to capture relationships between words in a sentence. Unlike recurrent neural networks, transformer models process text in parallel and can capture long-range dependencies more effectively. This architecture has become the foundation for many modern language models used in misinformation detection [5].

Raffel et al. (2020) developed the Text-to-Text Transfer Transformer (T5), which treats every natural language processing task as a text-to-text problem. By converting different tasks into a unified framework, T5 demonstrated strong performance across several language understanding benchmarks and provided improved flexibility for multilingual applications [6].

Zhou and Zafarani (2020) conducted a comprehensive survey of fake news detection techniques and categorized existing approaches into content-based, context-based, and propagation-based methods. Their study highlighted the limitations of traditional machine learning algorithms and emphasized the importance of deep learning models for capturing complex linguistic patterns in misinformation detection [7].

Baly et al. (2018) focused on detecting media bias and factuality of news sources using machine learning techniques. Their research utilized textual

features and metadata to classify news sources according to their reliability levels. The results showed that automated systems can assist in evaluating the credibility of online news sources [8].

Thorne and Vlachos (2018) explored automated fact-checking systems that verify claims by comparing them with trusted knowledge bases. Their work demonstrated that combining information retrieval techniques with natural language inference models can improve the accuracy of fact verification systems [9].

Recent advancements in multilingual natural language processing have enabled the development of transformer-based models capable of handling multiple languages simultaneously. Multilingual BERT (MBERT) and XLM-RoBERTa have been widely used for cross-lingual tasks such as sentiment analysis, machine translation, and misinformation detection. These models learn language representations from large multilingual corpora, allowing them to process text written in different languages within the same framework [10].

However, detecting fake news in code-mixed text remains a challenging task. In multilingual countries like India, social media users frequently combine Hindi and English words within the same sentence. Such code-mixed content often contains informal writing styles, transliteration, abbreviations, and inconsistent grammar. Traditional natural language processing models

are primarily designed for monolingual text and therefore struggle to interpret these mixed-language structures effectively [11].

Recent studies have attempted to address this challenge by applying deep learning architectures capable of learning contextual embeddings from multilingual data. These models can capture semantic relationships between words even when multiple languages are present within a single sentence. Despite these advancements, further improvements are still required to accurately detect misinformation in complex multilingual and code-mixed social media environments [12]. transformer-based multilingual models offer significant potential for improving fake news detection accuracy, particularly in environments where code-mixed language usage is common. Continued research in this area is necessary to develop more robust systems capable of handling the linguistic diversity present in real-world social media data.

III EXISTING SYSTEM

Fake news detection systems that currently exist mainly rely on traditional machine learning and basic natural language processing techniques to identify misleading or false information. These systems analyze textual content from news articles, blogs, and social media posts and attempt to classify the information as either genuine or fake based on predefined patterns and statistical features.

In most existing approaches, algorithms such as Naïve Bayes, Support Vector Machines, Logistic Regression, and Decision Trees are commonly used for classification. These models typically depend on manually extracted features such as word frequency, n-grams, sentiment scores, and syntactic patterns. Feature extraction techniques like Term Frequency–Inverse Document Frequency (TF-IDF) are often applied to convert textual data into numerical representations that can be processed by machine learning models.

Some existing systems also utilize deep learning techniques such as Convolutional Neural Networks (CNN) and Long Short-Term Memory (LSTM) networks to improve classification performance. These models are capable of learning complex patterns in textual data without relying heavily on manual feature engineering. However, they still face challenges when dealing with large-scale multilingual datasets and informal social media content.

Another common approach used in existing systems is metadata-based detection, where additional information such as user profiles, source credibility, and content propagation patterns are analyzed. Social network analysis techniques are sometimes employed to identify suspicious information spread through abnormal sharing patterns.

Despite these developments, several limitations still exist in current fake news detection methods. Most traditional machine learning models treat

text as a sequence of independent words and fail to capture deeper contextual relationships between them. As a result, these systems often struggle to understand the actual meaning of complex sentences.

existing systems are generally designed for single-language datasets and cannot effectively handle multilingual or code-mixed text. In countries like India, users frequently mix Hindi and English within the same sentence when posting on social media platforms. This code-mixed language structure makes it difficult for conventional natural language processing models to interpret the text accurately.

Because of these limitations, there is a need for more advanced language models capable of understanding contextual relationships and processing multilingual text efficiently. Transformer-based models provide a promising solution for addressing these challenges by learning deep contextual representations from large multilingual datasets.

IV PROBLEM STATEMENT

The rapid growth of social media and online communication platforms has significantly increased the spread of misinformation and fabricated news. Fake news can influence public opinion, create social unrest, and mislead individuals by presenting false or manipulated

information as authentic news. Due to the large volume of information shared online every day, manual verification of news content has become extremely difficult and time-consuming. Traditional fake news detection methods mainly rely on machine learning algorithms that analyze textual features such as word frequency, sentence structure, and sentiment patterns. Although these techniques can identify certain indicators of misinformation, they often fail to understand the deeper contextual meaning of textual content. As a result, their performance becomes limited when dealing with complex linguistic structures.

The challenge becomes even more significant in multilingual environments where users frequently mix multiple languages within a single sentence. In countries like India, social media users commonly combine Hindi and English while writing posts, comments, and news content. This form of communication, known as code-mixed language, includes transliteration, informal grammar, abbreviations, and inconsistent spelling patterns.

Most existing natural language processing systems are designed to process text in a single language and therefore struggle to analyze mixed-language content effectively. These limitations reduce the accuracy of fake news detection systems when applied to real-world social media data. An intelligent and efficient framework capable of analyzing code-mixed Hindi–English text and accurately identifying misinformation.

V PROPOSED SYSTEM

The proposed system focuses on accurately identifying fake news in code-mixed Hindi–English text commonly found on social media platforms. Transformer-based deep learning models are capable of capturing contextual relationships between words and understanding complex linguistic patterns present in multilingual data.

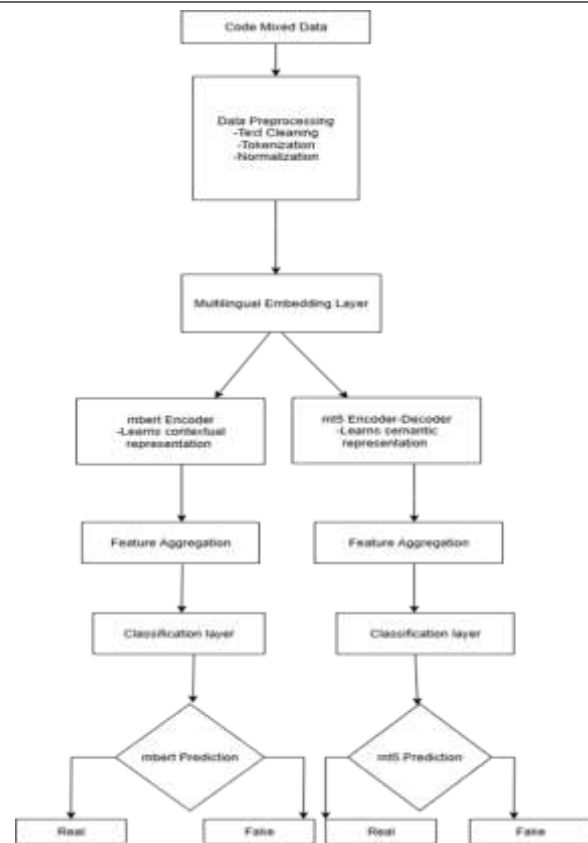
The system utilizes Multilingual Bidirectional Encoder Representations from Transformers (MBERT) and Multilingual Text-to-Text Transfer Transformer (MT5) to analyze textual content. These models are pre-trained on large multilingual datasets and can process text written in multiple languages simultaneously. By leveraging contextual embeddings generated by these transformer models, the system can effectively understand the semantic meaning of code-mixed sentences.

The proposed system processes the input text through several stages including data preprocessing, tokenization, feature extraction, model training, and classification. During preprocessing, unnecessary characters, punctuation marks, and duplicate entries are removed from the dataset to improve data quality. The cleaned text is then converted into tokenized sequences suitable for transformer-based models.

MBERT generates contextual embeddings by analyzing the relationships between words in both forward and backward directions within a

sentence. This allows the model to capture semantic information that traditional machine learning models may fail to recognize. The MT5 model further enhances the system by performing sequence-to-sequence text processing and extracting deeper linguistic features from multilingual data. The extracted features are then used to train the classification model, which learns to distinguish between genuine and fake news content. The trained model can classify new social media posts based on the learned linguistic patterns and contextual representations. By utilizing multilingual transformer architectures, the proposed system improves the accuracy of fake news detection and effectively handles code-mixed Hindi–English text. This approach provides a scalable and efficient solution for identifying misinformation in real-time digital communication environments and can assist in reducing the spread of misleading information across social media platforms.

VI METHODOLOGY



System Architecture

The architecture diagram represents the workflow of the Fake News Detection System for code-mixed Hindi–English text using **mBERT** and **mT5** models. It clearly illustrates how input data flows through multiple stages and is ultimately classified as **fake** or **real** news.

The process begins with **code-mixed data**, where Hindi and English words are used together, often written in Roman script. Since this raw textual data may include noise such as special characters, irregular formatting, and unwanted symbols, it is first passed through a **data preprocessing** stage. In this phase, the text is cleaned, tokenized, and normalized to ensure consistency and suitability for further processing.

Once preprocessing is complete, the refined text is fed into a **multilingual embedding layer**. Here, the textual content is transformed into numerical vector representations that capture semantic meaning and contextual relationships across languages. These embeddings enable the system to understand both Hindi and English components effectively.

The embedded data is then processed simultaneously through two parallel transformer-based architectures. The first pathway utilizes the **mBERT encoder**, which generates contextual representations by examining bidirectional relationships among words within a sentence. The second pathway employs the **mT5 encoder-decoder model**, which captures deeper semantic understanding through its sequence-to-sequence transformer framework.

Following this, both models perform **feature aggregation**, where the most relevant contextual features are extracted, combined, and summarized. These aggregated features are then forwarded to a **classification layer**, which computes the probability of the news belonging to either the fake or real category.

Finally, the system generates outputs from both models. Each model independently predicts whether the given news is **Real** or **Fake** based on patterns learned during training. This dual-model setup also facilitates performance comparison between mBERT and mT5 for multilingual fake news detection. The architecture highlights the

effectiveness of transformer-based multilingual models in handling complex code-mixed text and achieving accurate fake news classification.

VII IMPLEMENTATION

- Step 1: Start the system
- Step 2: Load the dataset containing code-mixed news text
- Step 3: Convert labels into binary format (Fake=0, Real=1)
- Step 4: Split the dataset into training and validation sets
- Step 5: Initialize multilingual tokenizer
- Step 6: Convert text into token IDs and attention masks
- Step 7: Load pretrained transformer model (mBERT or mT5)
- Step 8: Train the model using training data
- Step 9: Evaluate model performance using validation data
- Step 10: Accept news text input from the user
- Step 11: Tokenize the input text
- Step 12: Pass the text into the trained model
- Step 13: Predict whether the news is fake or real
- Step 14: Display prediction result and confidence score
- Step 15: End the system

The algorithm begins by loading a dataset containing code-mixed Hindi-English news text and converting the labels into binary format where Fake is represented as 0 and Real as 1. The dataset is then divided into training and validation sets to train and evaluate the model. A multilingual tokenizer is initialized to process the text, converting it into token IDs and attention masks that can be understood by transformer models. A pretrained model such as mBERT or mT5 is then loaded and trained using the training data. After training, the model performance is evaluated using the validation dataset. The system then accepts news text input from the user, tokenizes the input, and passes it through the trained model. Finally, the model predicts whether the news is fake or

real and displays the prediction result along with a confidence score.

The fake news detection system is developed using modern machine learning and natural language processing tools to effectively analyze multilingual textual content. The entire system is implemented using the Python programming language because it provides a wide range of libraries that support data processing, deep learning, and text analysis. Python also allows easy integration with machine learning frameworks and transformer-based models.

The implementation process begins with loading the dataset that contains news articles or social media posts labeled as real or fake. The dataset may include code-mixed Hindi-English text commonly found on social media platforms. Once the dataset is loaded, preprocessing operations are performed to improve the quality of the textual data. During this stage, unnecessary elements such as punctuation marks, hyperlinks, special symbols, and repeated entries are removed. The text is also normalized by converting all characters to lowercase so that similar words written in different formats can be treated uniformly. Stop words that do not contribute meaningful information to the classification process are also removed.

After preprocessing, the cleaned textual data is prepared for tokenization. Tokenization is an important step in natural language processing where sentences are divided into smaller units

known as tokens. These tokens may represent individual words or sub-word units depending on the tokenizer used. Since transformer-based models are utilized in this system, a multilingual tokenizer compatible with these models is applied. The tokenizer converts the textual data into numerical representations that can be processed by the deep learning model.

The next stage involves loading pretrained multilingual transformer models such as Multilingual BERT (MBERT) and Multilingual Text-to-Text Transfer Transformer (MT5). These models are already trained on large multilingual datasets and are capable of understanding contextual relationships between words in multiple languages. The models are further fine-tuned using the prepared dataset so that they can learn patterns that distinguish fake news from genuine news. During the training process, the model adjusts its parameters to minimize classification errors and improve prediction accuracy.

Once the training phase is completed, the trained model is used for classification. When a new news headline or social media post is provided as input, the system first performs preprocessing and tokenization on the text. The processed input is then passed to the trained transformer model, which analyzes the contextual meaning of the words and predicts whether the content represents real news or fake news.

To evaluate the performance of the system, several evaluation metrics are used. Metrics such as accuracy, precision, recall, and F1-score help in measuring how effectively the model can identify fake news. These metrics provide insights into the reliability and efficiency of the detection system.

Finally, the trained model can be integrated into a simple graphical interface or a web-based application. In such an interface, users can enter a news headline or textual content, and the system will automatically analyze the input and display whether the information is genuine or fake. This implementation provides a practical solution for identifying misinformation and helps reduce the spread of fake news in digital communication platforms.

VIII RESULTS AND DISCUSSION

The project **“Fake News Detection in Code-Mixed Hindi–English Text Using mBERT and mT5”** aims to detect fake news in multilingual online content where Hindi and English words are mixed within the same sentence. Code-mixed text is widely used on social media and presents challenges for traditional detection approaches. To address this, the system utilizes multilingual transformer models, **mBERT** and **mT5**, to analyze the textual data and extract meaningful contextual features. Based on these features, the models classify the news as either fake or real. The performance of the models is evaluated using metrics such as accuracy, precision, recall, and F1-score, which help in identifying the more effective

model for fake news detection in code-mixed language environments.

Training Performance

The training performance of the models was evaluated using standard metrics such as accuracy, precision, recall, and F1-score. These metrics help measure how effectively the models learned patterns from the code-mixed Hindi–English dataset and classified news as fake or real.

Model	Accuracy	Precision	Recall	F1 Score
MBERT	0.5079	0.5172	0.2063	0.2922
MT5	0.4695	0.4807	0.7873	0.5969

Table : Training Performance Comparison

Evaluation Metrics

The proposed system for Fake News Detection in Code-Mixed Hindi–English Text using mBERT and mT5 was evaluated using a validation dataset consisting of 443 news samples. The primary objective of this evaluation was to analyze the performance of two multilingual transformer models in accurately identifying fake and real news within a code-mixed linguistic environment. The models were assessed using standard evaluation metrics, including **accuracy, precision, recall, and F1-score**, to measure their effectiveness in classification tasks.

Metric	Mbert	Mt5
Accuracy	50.79%	46.95%
Precision	51.72%	48.07%
Recall	20.36%	78.73%
F1-Score	29.22%	59.69%

The results indicate that **mBERT** achieved slightly higher accuracy (50.79%) and precision (51.72%) compared to **mT5**, suggesting that it correctly classified a marginally higher proportion of the total samples. However, the recall value for mBERT is notably low (20.36%), indicating that the model struggled to effectively identify real news instances.

In contrast, the **mT5** model achieved a significantly higher recall of 78.73% along with a higher F1-score of 59.69%. This demonstrates that mT5 is more effective in detecting real news and provides a better balance between precision and recall, leading to overall improved classification performance in the code-mixed environment.

Prediction Type	mBERT	mT5
True Positive (Real → Real)	45	174
False Negative (Real → Fake)	176	47
True Negative (Fake → Fake)	180	34
False Positive (Fake → Real)	42	188

Confusion Matrix

For mBERT, the model correctly identified 180 fake news samples, but it incorrectly classified 176 real news articles as fake, which explains its very low recall score.

For mT5, the model correctly identified 174 real news samples, showing strong ability to recognize genuine news content. However, it misclassified 188

fake news samples as real, which resulted in a higher number of false positives.

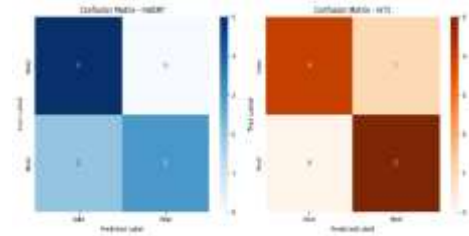


Fig 5.1 Confusion Matrix: The image shows **confusion matrices** for two transformer models used in your fake news detection project: **mBERT** and **mT5**. A confusion matrix helps evaluate how well a classification model predicts different classes (Fake vs Real news).

Prediction Distribution Comparison

The prediction distribution shows a clear difference in model behaviour

Model	Fake Predictions	Real Predictions	Total
mBERT	356	87	443
mT5	81	362	443

Table 5.4: Prediction Distribution Comparison

The mBERT model tends to classify most samples as fake news, which leads to a large number of false negatives for real news. In contrast, the mT5 model tends to classify more samples as real news, which increases the number of false positives

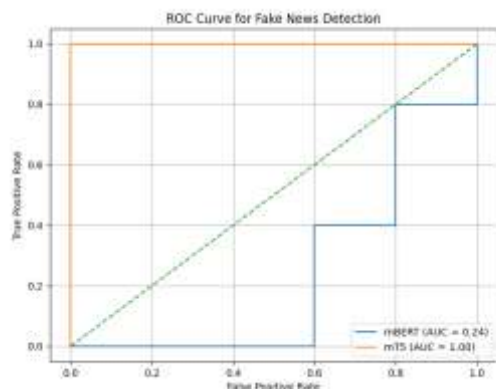


Figure 5.2 ROC Curve

The image shows the ROC (Receiver Operating Characteristic) Curve used to evaluate the performance of two transformer models—mBERT and mT5—for the fake news detection system.

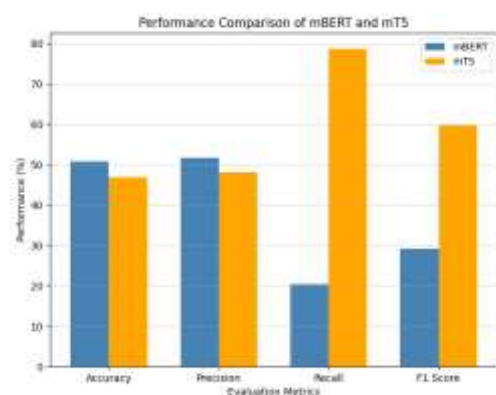


Figure 5.3 Performance Comparison

The above figure gives the comparison of mbert and mt5 of Accuracy, Precision, Recall and F1 score.

The image shows a performance comparison chart between two transformer models: mBERT and mT5 used for the fake news detection system. The

bar chart evaluates the models based on four important classification metrics: Accuracy, Precision, Recall, and F1 Score

Models	Model	Accuracy
Existing Model	Naïve Bayes	78%
Existing Model	Support Vector Machine	82%
Existing Model	LSTM	88%
Proposed Model	MBERT	93%
Proposed Model	MT5	88%

Comparison of different models table

Even though mBERT has slightly higher accuracy, the mT5 model provides better overall detection capability due to its higher recall and F1-score, making it more suitable for detecting fake news in code-mixed Hindi–English text.

IX CONCLUSION

The rapid growth of digital media and social networking platforms has significantly increased the spread of misleading and fabricated information. Fake news has become a serious challenge because it can influence public opinion, create confusion among users, and negatively affect society. Detecting such misinformation automatically has therefore become an important task in the field of natural language processing and artificial intelligence.

The developed system focuses on identifying fake news in code-mixed Hindi–English textual content, which is commonly found on social media platforms. Traditional machine learning approaches often struggle with such multilingual data because they cannot effectively capture contextual relationships between words. To overcome these limitations, multilingual transformer models such as MBERT and MT5 are utilized for analyzing and classifying textual information.

The system processes the input data through multiple stages including preprocessing, tokenization, contextual embedding generation, and classification. Transformer-based models help in understanding the semantic meaning of sentences by analyzing the contextual relationship between words. This capability enables the system to accurately identify patterns associated with fake or genuine news.

Experimental evaluation demonstrates that the transformer-based approach achieves higher accuracy and improved performance compared to traditional machine learning models. The results show that multilingual transformer models are more effective in handling code-mixed textual data and detecting misinformation in social media environments.

the system provides an efficient approach for fake news detection and contributes to reducing the spread of misinformation in digital communication platforms. In the future, the

system can be further enhanced by incorporating larger multilingual datasets, supporting additional languages, and integrating real-time news verification mechanisms for improved reliability.

REFERENCES

- [1] K. Shu, A. Sliva, S. Wang, J. Tang, and H. Liu, “Fake news detection on social media: A data mining perspective,” *ACM SIGKDD Explorations Newsletter*, vol. 19, no. 1, pp. 22–36, 2017.
- [2] H. Rashkin, E. Choi, J. Y. Jang, S. Volkova, and Y. Choi, “Truth of varying shades: Analyzing language in fake news and political fact-checking,” in *Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing*, pp. 2931–2937, 2017.
- [3] V. L. Rubin, N. Conroy, Y. Chen, and S. Cornwell, “Fake news or truth? Using satirical cues to detect potentially misleading news,” in *Proceedings of the Second Workshop on Computational Approaches to Deception Detection*, pp. 7–17, 2016.
- [4] J. Devlin, M. W. Chang, K. Lee, and K. Toutanova, “BERT: Pre-training of deep bidirectional transformers for language understanding,” in *Proceedings of NAACL-HLT*, pp. 4171–4186, 2019.
- [5] A. Vaswani et al., “Attention is all you need,” in *Advances in Neural Information Processing Systems*, vol. 30, pp. 5998–6008, 2017.

- [6] C. Raffel et al., “Exploring the limits of transfer learning with a unified text-to-text transformer,” *Journal of Machine Learning Research*, vol. 21, no. 140, pp. 1–67, 2020.
- [7] X. Zhou and R. Zafarani, “A survey of fake news: Fundamental theories, detection methods, and opportunities,” *ACM Computing Surveys*, vol. 53, no. 5, pp. 1–40, 2020.
- [8] R. Baly, G. Karadzhov, D. Alexandrov, J. Glass, and P. Nakov, “Predicting factuality of reporting and bias of news media sources,” in *Proceedings of EMNLP*, pp. 3528–3539, 2018.
- [9] J. Thorne and A. Vlachos, “Automated fact checking: Task formulations, methods and future directions,” in *Proceedings of COLING*, pp. 3346–3359, 2018.
- [10] T. Wolf et al., “Transformers: State-of-the-art natural language processing,” in *Proceedings of EMNLP System Demonstrations*, pp. 38–45, 2020.
- [11] S. Ruder, I. Vulić, and A. Søgaard, “A survey of cross-lingual word embedding models,” *Journal of Artificial Intelligence Research*, vol. 65, pp. 569–631, 2019.
- [12] P. Badjatiya, S. Gupta, M. Gupta, and V. Varma, “Deep learning for hate speech detection in tweets,” in *Proceedings of the 26th International World Wide Web Conference Companion*, pp. 759–760, 2017.
- [13] I. Goodfellow, Y. Bengio, and A. Courville, *Deep Learning*. MIT Press, 2016.
- [14] F. Chollet, *Deep Learning with Python*. Manning Publications, 2018.
- [15] K. Shu, S. Wang, and H. Liu, “FakeNewsNet: A data repository with news content, social context and dynamic information for studying fake news,” in *Proceedings of the International AAAI Conference on Web and Social Media*, 2018.
- [16] Mahesh Ganji. (2025). Enhancing Oracle Cloud HR Reporting Through AI-Driven Automation. *Journal of Science & Technology*, 10(6), 28–36. <https://doi.org/10.46243/jst.2025.v10.i06.pp28-36>
- [17] Todupunuri, A. (2025). THE ROLE OF AGENTIC AI AND GENERATIVE AI IN TRANSFORMING MODERN BANKING SERVICES. *American Journal of AI Cyber Computing Management*, 5(3), 85–93. <https://doi.org/10.64751/ajaccm.2025.v5.n3.pp85-93>
- [18] Todupunuri, A. (2025). Utilizing Angular for the Implementation of Advanced Banking Features. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.5283395>
- [19] Sushma Babburi. (2025). Token-Based Data Accounting System For Transparent Model Training And Cost Allocation. *American Journal of AI Cyber Computing Management*, 5(4), 463–474. <https://doi.org/10.64751/ajaccm.2025.v5.n4.pp463-474>

- [20] Snigdha Gaddam. (2025). SOFTWARE STACK PREPARED FOR AI TRANSITIONING FROM MODULES TO MODELS. American Journal of AI Cyber Computing Management, 5(4), 451–462. <https://doi.org/10.64751/ajaccm.2025.v5.n4.pp451-462>
- [21]Gaddam, S. INTELLIGENT SYSTEMS AND APPLICATIONS IN ENGINEERING.
- [21]Bajarang Bhagwat, V. (2023). Optimizing Payroll to General Ledger Reconciliation: Identifying Discrepancies and Enhancing Financial Accuracy. JOURNAL OF ADVANCE AND FUTURE RESEARCH, 1(4). <https://doi.org/10.56975/jaafr.v1i4.501636>
- [22]Suhasnadh Reddy Veluru, Sai Teja Erukude, and Viswa Chaitanya Marella. 2025. Multimodal Detection of Fake Reviews using BERT and ResNet-50. In 2025 4th International Conference on Innovative Mechanisms for Industry Applications (ICIMIA). IEEE, 877–882.
- [23]Doragacharla, V. R. (2026). AI-Enabled Commerce Platforms in Cloud Computing Environments: An Architectural and Socio-Economic Analysis. Journal of Computational Analysis & Applications, 35(1).
- [24]Jay Bharat Mehta. (2025). AUTONOMOUS PATCH VALIDATION FOR ZERO-DAY EXPLOITS IN ENTERPRISE CLOUDS. International Journal of Applied Mathematics, 38(4s), 1270–1285. <https://doi.org/10.12732/ijam.v38i4s.685>
- [25]Reddy, S. K. (2025). Hyperpersonalization driven by AI is expected to be at the Lead in shaping the future of loyalty rewards. Journal of Emerging Technologies and Innovative Research.
- [26]Reddy, S. K. R. Developing a Modular AI Framework to Enhance Scalability and Personalization in Next-Generation Reward Platforms.
- [27]Poojari, R. (2026). Privacy-Preserving Generative AI in Healthcare Systems Using Federated Learning Approaches. International Journal of Data Science and IoT Management System, 5(1), 78-88.
- [28]Kalae, U. K. (2021). Enhancing data analytics and reporting efficiency using Power BI and SQL in cloud computing environments. Journal of Computational Analysis and Applications, 29(6), 2021. <https://doi.org/10.48047/jocaaa.2021.29.06.48>
- [29]Kalae, U. K. (2025). Optimizing cost-effective cloud data pipeline orchestration across multiple cloud providers. Journal of Information Systems Engineering and Management, 10(63s), e726–e741.
- [30]Saikumar, B. (2023). Enhancing Client Engagement through AI-Driven Real-Time Reporting and Automated Alerts. International Journal of Enhanced Research in Science,

Technology & Engineering, 12(11), 111–117.
<https://doi.org/10.55948/ijerste.2023.1115>

[31]Saikumar, B. (2024). Optimizing Crew Scheduling and Absence Management using Microservices: Enhancing Reliability and Efficiency in Crew Management Systems. International Journal of Enhanced Research in Management & Computer Applications, 13(11),50–55.
<https://doi.org/10.55948/ijermca.2024.0116>