

RoBERTa-GCN: A Novel Approach for Combating Fake News in Bangla Using Advanced Language Processing and Graph Convolutional Networks

MEJBAH AHAMMAD¹, AL SANI², KHALILUR RAHMAN³, MD TANVIR ISLAM⁴, MD MOSTAFIZUR RAHMAN MASUD⁵, MD. MEHEDI HASSAN⁶, (*Member, IEEE*), MOHAMMAD ABU TAREQ RONY⁷, SHAH MD NAZMUL ALAM⁸, MD. SADDAM HOSSAIN MUKTA^{9*}

¹Department of Computer Science and Engineering, American International University Bangladesh, Dhaka, Bangladesh

²Department of Industrial and Production Engineering, Jashore University of Science and Technology, Jashore, Bangladesh

³Department of Computer Science and Engineering, Sylhet Engineering College, Sylhet, Bangladesh

⁴Department of Software, Sungkyunkwan University, Gyeonggi-do, South Korea

⁵School of Computing, Universiti Teknologi Malaysia, Johor Bahru, Malaysia

⁶Computer Science and Engineering Discipline, Khulna University, Khulna, Bangladesh

⁷Department of Statistics, Noakhali Science and Technology University, Noakhali, Bangladesh

⁸Department of Computer Science and Engineering, Arizona State University

⁹LUT School of Engineering Science, LUT University, Lappeenranta, Finland

Corresponding author: Md. Saddam Hossain Mukta (Saddam.Mukta@lut.fi)

ABSTRACT In this era of widespread information, combating fake news in less commonly represented languages like Bengali is a significant challenge. Fake news is a critical issue in Bangla, a language that a vast population uses but lacks adequate natural language processing tools. To address this, our research introduces RoBERTa-GCN, a cutting-edge model combining RoBERTa with a graph neural network (GCN) to accurately identify fake news in Bangladesh. The dataset we utilized comprises articles from 22 prominent Bangladeshi news portals covering diverse subjects such as politics, sports, economy, and entertainment. This comprehensive dataset enables the model to learn and adapt to the intricacies of the Bangla language and its news ecosystem, facilitating effective fake news detection across various content categories. Our approach integrates the RoBERTa model, adapted for Bangla, with GCN's expertise in processing relational data, forming an effective means to differentiate between authentic and fake news. This study's key achievement is the creation and application of the RoBERTa-GCN model to the Bangla language, an area not thoroughly explored in previous research. The findings show that RoBERTa-GCN surpasses existing methods, achieving impressive accuracy rates of 98.60%, highlighting its capability as a robust model for preserving news integrity in the digital era, especially for the Bangla-speaking population.

INDEX TERMS Fake news detection, Graph neural network, NLP, Bangla language, Deep learning, Machine learning, Encoder

I. INTRODUCTION

IN Recent years, the emergence of online news platforms, including social media, news blogs, and online newspapers, has prompted individuals to actively seek and consume news due to their advantages, such as rapid information dissemination, convenient accessibility, and cost-effectiveness [1]. Meanwhile, easy access to social media platforms enables the rapid and widespread spread of false information, including fake news, within the digital realm [2].

“Fake news” refers to articles that might delude or deceive readers by disseminating fabricated information [3]. For the sake of traffic, some self-media and internet users spread a significant amount of unverified news that was subsequently reprinted and unthinkingly followed to expand and undermine the credibility and authority of mainstream media. However, it additionally resulted in economic, political, and other hidden risks [4], [5].

The propagation of false news serves to create fear, pro-

mote racist ideology, and provoke acts of bullying and violence against innocent individuals [6]. According to intelligence studies, social media rumors have a particularly long-lasting influence on less clever people, preventing them from making optimal judgments [7].

Fake news is widely recognized as a significant threat to global trade, media, and democracy, carrying severe ramifications [8]. Around 25% of Americans visited a fraudulent news website six weeks before the 2016 US election. This occurrence has been hypothesized as a potential factor that impacted the election's outcome [9]. A fake news report claimed the explosion resulted in injuries to US President Barack Obama, leading to a significant loss of \$130 billion in the stock market [10].

In addition to having a higher rate of social media penetration, Bangladesh is a South Asian nation plagued by poverty and rumors. As of the beginning of 2023, Bangladesh's internet user population stood at 66.94 million, as reported by Datareportal [11], [12]. Bangladesh has encountered many quantifiable incidences of misinformation in recent years. Ten individuals were injured, and rioters fatally beat five in July 2019 due to widespread rumors regarding the anticipated human sacrifice during the Padma Bridge construction project [13]. An additional case of misinformation resulting in violence occurred in 2012 when a local Buddhist youth's Facebook post featuring an image of a burned Quran, the holy book of Islam, incited an enraged mob to torch 12 Buddhist temples, pillage, and burn over 50 houses in a Buddhist community situated in Ramu, Cox's Bazar [14]. During the COVID-19 pandemic, vaccination skepticism in Bangladesh was encouraged by allegations that coronavirus vaccines "contain a microchip" that makes it possible for Western governments to spy on people. Such unjustified conspiracy concepts pose a huge obstacle to the country's public health programs [15]. Another false news stated that just 200,000 vaccinations would be imported for the first installment, with all of this being delivered to the government leaders and members of the ruling party in Bangladesh [16].

To counter the spread of misinformation, dedicated platforms like PolitiFact [17], FactCheck [18], and JaChai [19] manually review and update potential fake news articles in online media, providing detailed explanations and citing logical and factual reasons to debunk the false claims. Despite their efforts, computational tools have recently been employed to fight the threat of false news [3]. One effort to identify fake news involves using multi-perspective speaker profiles [20]. Conversely, another approach introduces an automated fact-checking methodology leveraging third-party sources to verify the precision of news articles [21]. Additionally, a methodology has been established to detect incongruities between Bangla news headlines and body content [22].

Distinguishing fake news in Bangla is significantly more challenging than in other languages because of its structure. Approximately 265 million people speak Bangla as their native language, making it the 7th most widely spoken language

globally [23]. The majority of the readers are unable to recognize fake news with their own knowledge. There are several studies available in English. As far as we know, limited resources or computational methods are currently available to address the menace of fake news generated in Bangla, which poses a potential threat to this sizable community. Therefore, to address these challenges, we propose a custom method of detecting Bangla fake news utilizing the power of RoBERTa and GCN. In addition, our research through extensive data analysis and exploring various methods for fake news detection marked by the following key contributions:

- 1) We propose a custom model named RoBERTa-GCN for Bangla fake news detection. To the best of our knowledge, we are the first to integrate RoBERTa and GCN for Bangla fake news detection.
- 2) The combination of RoBERTa's contextual embeddings with GCN's structural learning is an innovative approach that enhances Bangla fake news detection by leveraging the relationships within news content, an aspect often overlooked by other models.
- 3) We benchmark several SOTA models against Ban-FakeNews and compare them with our proposed RoBERTa-GCN, which outperforms all existing models by achieving 98.60% accuracy.

The structure of the subsequent sections of this document is meticulously designed to facilitate a coherent presentation of our research findings. Section II is dedicated to a critical examination of pertinent literature, situating our work within the broader scholarly context. Section III, we explain overall data collection and analysis procedures. In Section IV, we delineate our proposed methodology, elaborating on the innovative approaches and techniques employed in our investigation. Section V presents a rigorous analysis of the data obtained through the application of advanced methodological procedures, offering a comparative evaluation of the results. Finally, Section VI synthesizes the key insights and discoveries emanating from our study, underscoring their significance and implications for the field.

II. RELATED WORKS

The dissemination of false information on online media platforms has led to a need for clarification among various individuals, giving rise to numerous challenges. Several researchers have employed diverse techniques to control the spread of fake news.

Rai *et al.* [35] used Long Short-Term Memory (LSTM) and bidirectional encoder representation from transformers (BERT) to classify fake news. Their results showed higher accuracy compared to the baseline models on the PolitiFact and GossipCop datasets. Additionally, Rai *et al.* employed a vanilla BERT model combined with an LSTM layer to investigate performance improvements further. Nasir *et al.* [24] showed how to use a hybrid deep learning (DL) model that combines convolutional and recurrent neural networks (CNN-RNN) to sort fake news. This model did better than non-hybrid methods on the ISOT and FA-KES datasets and

TABLE 1. Summarizing various methods in the literature for detecting fake news, detailing the datasets used, implemented techniques, and their overall performances.

Ref.	Year	Dataset	Technique	Accuracy (%)
[24]	2021	FA-KES(804 news articles)	Hybrid CNN-RNN	60.00
[25]	2021	300 articles	Random Forest	82.00
[26]	2020	726 news articles	Random Forest	85.00
[27]	2023	English articles, English tweets, Urdu articles, Urdu tweets, Arabic Articles, Arabic Tweets	BERT	90.00
[3]	2020	22 most popular and mainstream trusted news portals (50,000 news articles)	Support Vector Machine	91.00
[28]	2023	English ISOT(44,493 news articles), Dravidian-language fake news(25,347 news articles)	XLM-RoBERT	93.31
[29]	2022	Facebook (676 individual posts)	DistilBERT, XLM-RoBERTa, BERT	97.00
[30]	2023	IBFND	BERT	97.22
[31]	2021	George McIntire(6,335 news articles), Kaggle(20,761 news articles), Gossipcop(21,641 news articles), Politifact(948 news articles)	BerConvoNet	97.45
[32]	2022	Bengali newsletter-based research database (49000 news articles gathered from various Bangalese news sources)	Support Vector Machine	98.08
[33]	2023	Twitter and YouTube (10,700 posts and news articles)	CT-BERT+BiGRU	98.46
[34]	2022	Buzzfeed(1700 news), GossipCop(17,520 news), ISOT(44,900 news), Politifact(700 news)	LSTM-LF	98.50

shows promise for future tests that will use it on other datasets as well. Alghamdi *et al.* [33] explored the impact of freezing and unfreezing parameters in a neural network architecture used for BERT and CT-BERT. They found that models like BiGRU and CT-BERT achieved superior performance on the COVID-19 fake news dataset.

Sudhakar and Kaliyamurthi [36] employed advanced machine learning classifiers, specifically Logistic Regression(LR) and Naive Bayes(NB), where the G-power experiment ensures accuracy and T-test analysis displays that LR works better than NB for independent samples. Choudhary and Arora [37] applied a sequential neural model incorporating an LSTM-based word embedding model. The feature-based sequential model performed better than ML-based models and LSTM-based word embeddings in less time. Choudhary *et al.* [31] showed that BerConvoNet, using BERT in the news embedding block with various kernel sizes, effectively distinguishes real from fake news across benchmark datasets, outperforming other state-of-the-art (SOTA) models. Khullar and Singh [38] emphasized the need for a distributed client-server architecture to protect data, proposing a model named F-FNC, utilizing LSTM, CNN-LSTM, Bi-LSTM, and CNN-BiLSTM algorithms, the model outperformed other distributed DL classifiers based on different metrics. Raja *et al.* (raja2023fake) used transfer learning with mBERT and XLM-R models to effectively detect fake news in the Dravidian Fake dataset, achieving high accuracy in multilingual settings. Similarly, Kaliyar *et al.* (kaliyar2020fndnet) developed FNDNet, a CNN-based model that identifies biased features to classify fake news, showing improved performance on benchmark datasets for social media.

To notice fake news, Xia *et al.* [39] utilized a multi-head

attention model to find major emergencies in a think tank and public opinion stage and a hybrid CNN, Bi-LSTM, and attention mechanism model, which helped to increase metrics values like loss, accuracy, f1-score, and recall. To distinguish fake news from SOTA models by adding variational Bi-LSTM, autoencoder, and semantic topic-related features, Hosseini *et al.* [40] used LDAVAE, an integrated LDA, supervised Bi-LSTM VAE probabilistic model where the model ablation measured the accuracy and performance of neural embeddings and topic features. Palani and Elango [41] proposed a hybrid BERT-BiLSTM-CNN (BBC-FND) model structure consisting of three primary layers: an embedding layer, a feature representation layer, and a classification layer that then captures applicable information, patterns, and global meaning to detect the forecast news validity. Okunoye and Ibor [42] used genetic search to select neural architecture, utilizing neural architecture and DL techniques to classify bogus news. Malhotra and Malik [43] evaluated SVM, LR, CatBoost, XGBoost, multinomial, NB, and RF ensemble models, where the deep auto-ViML model and passive-aggressive classifier worked well. SVM was the fastest at 0.245 ms. Ghamdi *et al.* [27] utilized natural language processing (NLP), BERT, and human programming to detect misinformation in Twitter and website content. A labeled dataset in English, Arabic, and Urdu was used to highlight the efficiency of their models in detecting false material. Mohawesh *et al.* [44] proposed a semantic technique combining pre-trained word2vec vectors with a capsule neural network to enhance multilingual fake news detection. This approach surpassed existing methods by incorporating relational variables extracted from the text. Dixit *et al.* [34] showed the four-step false news detection models: data pre-processing, feature reduction, feature extraction, and LSTM-

LF classification. It worked better than previous ways of finding fake news on Buzzfeed, GossipCop, ISOT, and Politifact datasets. Jain et al. [45] utilized a resilient DL model to detect false information assertions using associated embedding, attention methods, and pertinent metadata to show that it worked effectively across real datasets. In contrast to CNN, Palani et al. [46] used the BERT model to pull out written parts that kept their meaning and came up with the Capsule network (CapsNet) algorithm, which helped combine useful data to make a lot of data representations for accurate fake news detection. Ravish et al. [47] suggested using a multi-layered Principal Component Analysis (PCA) feature selection method along with Multiclass Support Vector Machines (MSVMs) for classification. This improved model accuracy across ten datasets, especially when more characteristics were needed to validate the feature extraction methods. Hos-sain et al. [3] studied and tested cutting-edge models such as BERT, CNN, Bi-LSTM, LR, and RF. These models used common language variables and neural networks to make low-resource language research more advanced. To detect fake news classification, Hasib et al. [48] used DT, SGD, BERT, CNN, and ANN models, where BERT worked better than any other models and achieved maximum accuracy in balanced datasets and good performance through cross-validation with varied K values.

Pranto et al. [29] experimented with ML models, among which BERT was the most prominent automatic detection method for Bengali false news classification. When BERT was applied to Bengali Facebook posts related to the COVID-19 dataset, 10 topics were identified and grouped into three groups: system, belief, and social. Roman et al. [30] enhanced the Bangla Fake News Dataset by addressing the imbalance, using web scraping and Google Translate, achieving the highest F1-score via the BERT model, surpassing the BanFakeNews dataset's performance. Ali et al. [32] conducted a feasibility study on detecting Bangla false news from social media, employing diverse feature extraction methods and ML algorithms, ultimately achieving optimal accuracy using LSTM in their proposed system. Anjum et al. [25] used a Random Forest(RF) classifier to distinguish bogus news from a combination of false and legitimate sources, with an accuracy of 82%. It included more than 300 articles. Islam et al. [26] employed a data mining approach to classify Bengali fake news in South Asia, analyzing 726 articles, and found RF to archive an accuracy rate of 85%. Bhattacharjee et al. [49] introduced BanglaBERT and BanglishBERT and achieved results in NLU for Bangla, along with new datasets. As highlighted in the literature review and Table 1, numerous models have shown potential for fake news detection, but most are tailored for English. This paper addresses this gap by proposing a specialized model for detecting fake news in the Bangla language.

III. DATASET: BANFAKENEWS

In this paper, we utilized a public dataset named "*Ban-FakeNews*", which was specially introduced for developing

models for classifying fake and real news [3]. "*BanFake-News*" was collected from 22 most prominent online news sites in Bangladesh, which was used to acquire false news data for our research. The dataset is publicly available on Kaggle.¹ Ensuring the dataset's diversity involved selecting news stories from various fields, such as sports, politics, economics, and the environment. Here is a synopsis of all the news stories that were part of our data collection:

- kalerkantho.com
- jagonews24.com
- banglanews24.com
- banglatribune.com
- jugantor.com
- dhakatimes24.com
- ittefaq.com.bd
- somoynews.tv
- dailynayadiganta.com
- bangla.bdnews24.com
- prothomalo.com
- bd24live.com
- risingbd.com
- dailyjanakantha.com
- bd-pratidin.com
- channelionline.com
- samakal.com
- independent24.com
- rtnn.net
- bangla.thereport24.com
- mzamin.com
- bhorerkagoj.net

Table 2 presents the framework of a dataset specifically created for identifying fake news in the Bengali language. The dataset consists of columns representing the article ID, domain, publication date, category, title, text, and a label indicating the authenticity of the piece. Every column has a distinct function for analysis. The bar chart in Figure 1 represents a dataset breakdown for noticing fake news in Bangla. Each bar corresponds to a different news domain, and the length of the bar reflects the number of articles from that domain included in the dataset. Here's an analysis based on the visualization:

TABLE 2. Features of the fake news dataset used in this study.

Columns	Descriptions
articleID	Unique identifier for each article
domain	The domain in which the publication was made
date	Date of publication of the article
category	Category or genre of the article
headline	Headline or title of the article
content	Main content or body of the article
label	Label indicating the authenticity of the article

Domain Diversity: The dataset comprises articles from diverse domains. This diversity is beneficial for building a robust fake news detection mechanism, allowing the model to learn from various writing styles and content types.

Volume of Articles: The domains contribute a varying number of articles, with kalerkantho.com and jagonews24.com providing the largest amount. In contrast, bhorerkagoj.net and mzamin.com contribute the least.

Data Representation: The dataset reflects the diversity of domains and the proportionality of fake and real news within each domain to prevent overfitting to specific domain characteristics that are not related to news authenticity.

¹BanFakeNews dataset available on: <https://www.kaggle.com/datasets/cryptexcode/banfakenews>. Accessed: 05, September, 2024.

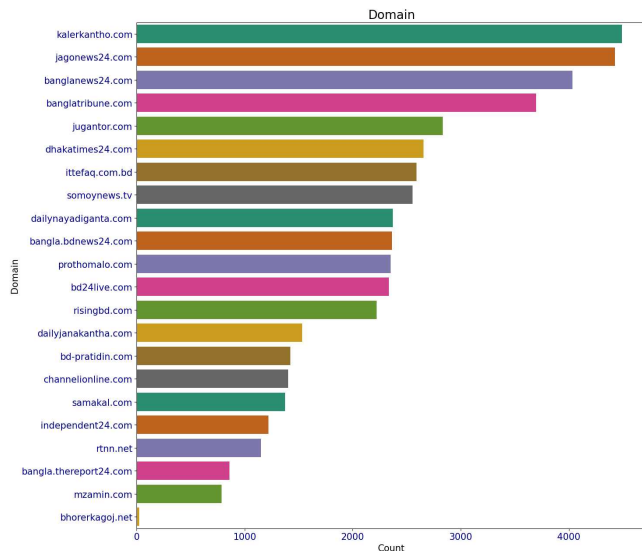


FIGURE 1. Distribution of Bangla fake news by domain, showing the count of fake news articles from various news websites.

The utilization of the chart guarantees that the dataset employed to train a model for detecting false news is equitably distributed and accurately reflects the actual dissemination of news across various Bangla domains. Furthermore, Figure 2 illustrates the category-wise distribution of the dataset, showing an even spread across five categories: politics, sports, economy, entertainment, and technology. Each category represents roughly one-fifth of the total, with entertainment slightly leading. This balanced distribution suggests that fake news is not concentrated in a single category but is spread across various topics. Such a balance in the dataset is valuable for developing a fake news detection model that remains unbiased and effective across multiple categories. Similarly, Figure 3 presents a polar bar chart highlighting the frequency of top words associated with fake news in Bengali. The chart features several spikes, each representing the frequency of a different word. Notably, one word stands out with a significantly higher frequency, indicating its common usage in the context of fake news. This visual tool helps in identifying key terms prevalent in fake news articles, which can be critical for enhancing the performance of machine learning algorithms in text classification. Figure 4 provides a comparative visualization of word clouds of the dataset. These clouds showcase the most frequently used words within each news type, with the size of the words reflecting their frequency.

A. DATA PREPROCESSING

The processing of Bangla fake news data involves specialized steps tailored to the nuances of the Bengali language. It includes text normalization, the removal of irrelevant characters, and contextual analysis. Additionally, algorithms adapt to linguistic patterns for accurate classification. Significantly contributing to preserving information integrity within the Bengali-speaking community, this procedure guarantees de-

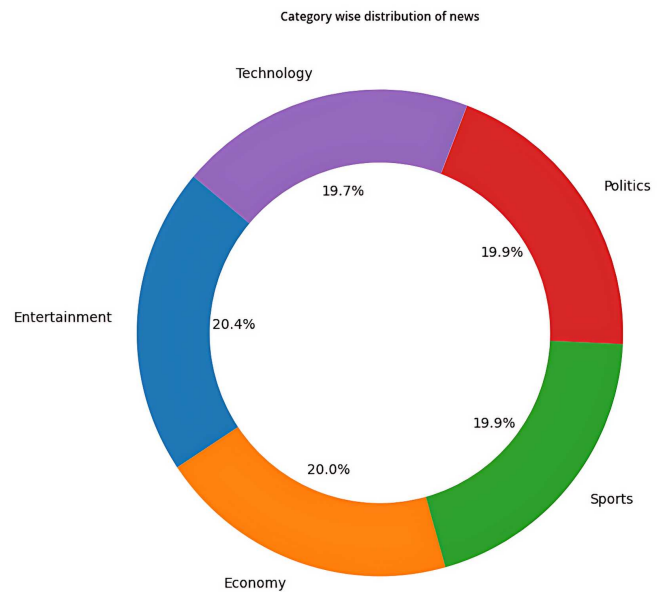


FIGURE 2. Category Wise distribution of News in Donut Chart for Bangla fake news detection.

pendable identification and filtration of misinformation. In this paper, we used the BanFakeNews [3] dataset, which we preprocess as shown in Figure 5 and further described as:

1) Text Cleaning

The Bangla Fake News Text Cleaning Process is a meticulous and crucial step in data preparation, especially tailored for the unique characteristics of the Bengali language. This procedure involves several key activities: removing non-textual elements, correcting typographical errors, and standardizing various dialectal forms prevalent in the Bangla script. Distinguishing characters, punctuation, and numerical information that fail to enhance the semantic value of the text is a particular emphasis. The process also includes normalizing colloquial expressions and idioms, ensuring the text is analytically relevant and linguistically accurate. This thorough cleaning is essential for effectively analyzing and detecting fake news in Bangla text datasets.

- **Text Normalization:** Converting text to a uniform format by lowercasing characters, standardizing dates and numbers, and resolving inconsistencies.
- **Noise Removal:** Eliminating irrelevant characters and symbols like HTML tags, URLs, emojis, special characters, and extra whitespaces.
- **Dealing with Punctuation:** Removing or utilizing punctuation marks to improve sentence structure parsing.
- **Handling Numerals:** Convert numbers to words or remove them based on their analysis and relevance.
- **Tokenization:** Breaking the text into tokens such as words, numbers, or symbols, considering the unique script and grammar of the Bangla language.

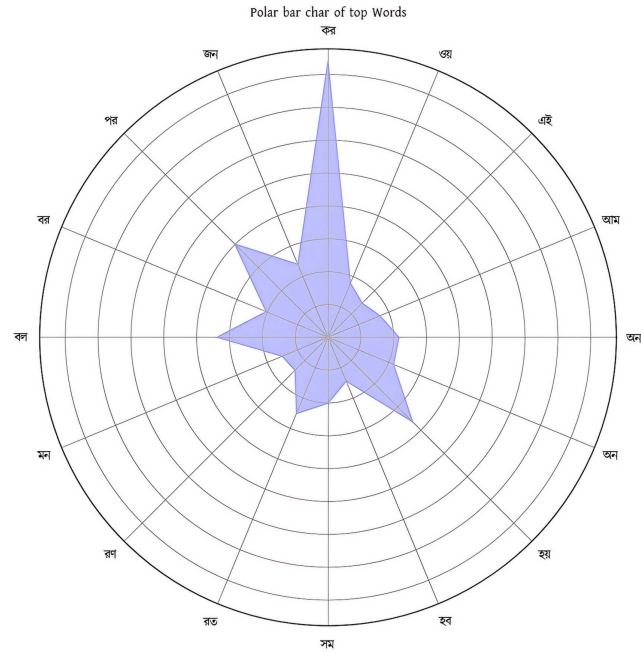


FIGURE 3. Polar Bar chart visualizing frequency of top words used in Bangla language fake news articles

- **Removing Stop Words:** Discard frequently occurring words that have little meaning, such as prepositions, conjunctions, and certain verbs.
- **Stemming and Lemmatization:** Reducing words to their root form through stemming or using vocabulary and morphological analysis for lemmatization, which is complex due to Bangla's morphology.
- **Handling Slang and Abbreviations:** Standardize or translate slang terms and abbreviations to maintain clarity in the dataset.

2) Feature Extraction

The cleaning procedure is crucial for producing Bangla text data to identify false news. The procedure involves making essential modifications to the text to enhance the efficiency of feature extraction methods, including Bag of Words (BOW) [50], TF-IDF [51], Word Embedding [52], and N-Grams [53]. This stage rigorously removes discrepancies, standardizes textual deviations, and eradicates non-linguistic components unique to the Bengali script. The thoroughness of the cleaning process is crucial, as it directly impacts the quality of the recovered features. Consequently, the accuracy of recognizing fake news within the linguistically varied Bangla dataset is affected.

3) Semantic Analysis

Semantic analysis encompasses a multifaceted procedure that interprets the latent meanings, sentiments, and veracity present in Bengali text to identify false news. This phase involves sentiment analysis to gauge the emotional tone, theme identification to understand the central subjects, contextual analysis to grasp situational meanings, and entity

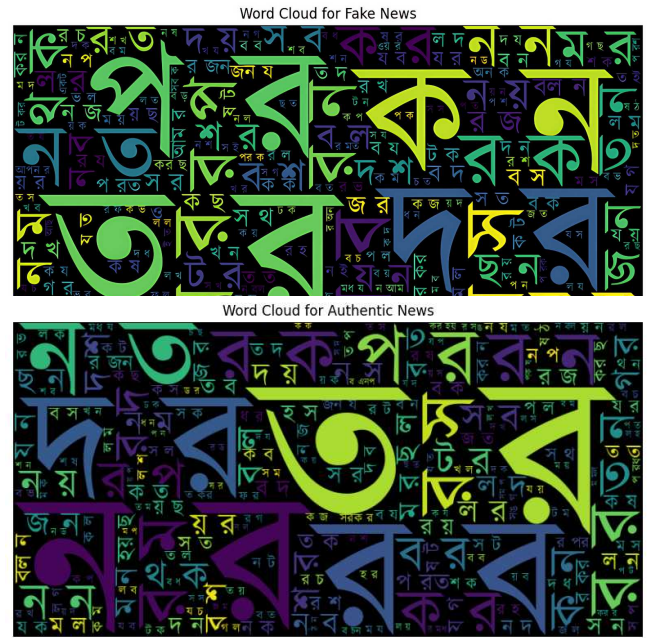


FIGURE 4. Word cloud of the dataset we used for RoBERTa-GCN.

recognition to pinpoint and categorize names, places, and organizations. These steps are tailored to comprehend the complexities of Bangla syntax and semantics, allowing for an in-depth examination of news content to identify and flag misinformation accurately.

4) Content Authentication

In the fight against fake news within the Bangla-speaking community, content authentication stands as a pivotal line of defense. This intricate process employs sophisticated fact-checking algorithms tailored to the Bangla language, scrutinizing the veracity of claims and narratives. Source credibility analysis evaluates the trustworthiness of information origins, while cross-referencing mechanisms compare data across multiple databases and records. These measures, essential in the digital age, are designed to safeguard the truth and combat the spread of misinformation in Bangla news, thereby preserving the integrity of the informational ecosystem for millions of Bengali speakers.

5) Syntactic Analysis

Syntactic analysis in Bangla fake news detection is a structured approach to dissecting the grammatical intricacies of text. It involves constructing parsing tree structures that map the grammatical hierarchy and relationships between words. Grammar analysis examines the rules specific to Bangla syntax, ensuring sentences are constructed correctly. Sentence structure analysis then delves deeper, evaluating the arrangement of words and phrases to discern patterns that might indicate misinformation. This rigorous syntactic scrutiny is vital for automated systems to effectively identify and flag fake news in the richly complex Bangla language.

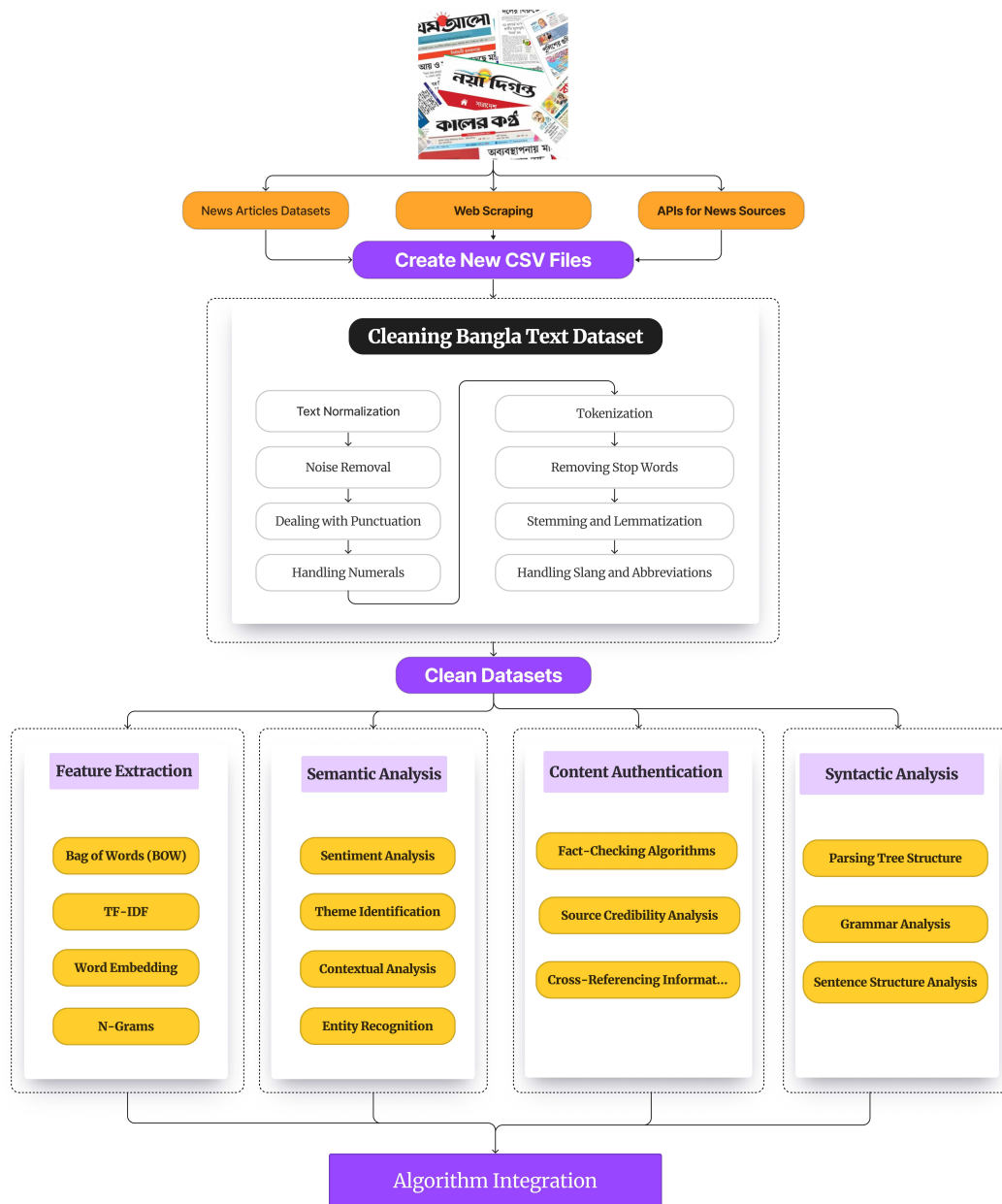


FIGURE 5. Workflow for preparing bangla text data for analysis, featuring cleaning, feature extraction, and various analytical methods.

6) Pseudo Algorithm for Bangla Fake News Data Preprocessing

The Algorithm 1 describes a process for preprocessing the dataset for training our proposed RoBERTa-GCN. It begins by fetching articles from Bangla news sources. These articles are preprocessed, which includes normalization, tokenization, noise removal, etc., to clean the text. Next, features are extracted using NLP techniques like BOW and TF-IDF. A pre-trained model is then loaded to analyze these features. Furthermore the pseudocode shows the pipeline for assessing each article and determines whether the content is real or

fake. If an article is predicted to be fake, it is flagged, and relevant stakeholders are notified. This automated process streamlines the task of identifying misinformation in Bangla news articles. Table 3 outlines the symbols and their meanings used in Algorithm 1.

IV. PROPOSED: ROBERTA-GCN

In NLP and misinformation detection, we present ‘RoBERTa-GCN,’ an innovative Convolutional (CNN) model to detect false information in the Bangla language. This model is a fusion of RoBERTa’s [54] enhanced lan-

TABLE 3. Definition of the symbols used in Algorithm 1.

Symbol	Definition
$\mathcal{S}$	Set of news articles
$\mathcal{L}$	Labeling function
$\mathcal{P}$	Articles preprocessing
$\mathcal{F}$	Feature extraction
$\mathcal{M}$	Instantiate a pre-trained model
ψ	Classification function
$\mathcal{F}\uparrow\downarrow\}$	Flag articles procedure

language processing abilities, particularly optimized for dynamic contexts and the intricate relational data handling capabilities of Graph Convolutional Network (GCN) [55]. Tailored to classify news articles into ‘Authentic’ or ‘Fake’ categories, RoBERTa-GCN is meticulously trained on a diverse dataset of Bangla news, capturing the language’s unique characteristics and cultural nuances. It features a dynamic GCN component, crucial for adapting to the ever-evolving language use and patterns in fake news distribution, ensuring accuracy and robustness. This schematic outlines a Bangla fake news detection model that ingests news articles, segregating headlines and contents. It employs encoders like RoBERTa-GCN for contextual understanding and topic encoders for thematic insights. The model constructs a dynamic graph to encapsulate inter-feature relationships. Features are pooled to synthesize the data into a compact

Algorithm 1 Pseudocode for preprocessing the “BanFake-News” dataset for training our proposed RoBERTa-GCN.

Require: $\mathcal{S}$
Ensure: $\mathcal{L} : \mathcal{S} \rightarrow \{0, 1\}$, where 1 corresponds to FAKE and 0 to REAL
1: **Procedure** Preprocess_and_Classify($\mathcal{S}$):
2: $\mathcal{S}' \leftarrow \{s' \mid \forall s \in \mathcal{S}, s' = \mathcal{P}(s)\}$ {Preprocess each article}
3: $\mathcal{V} \leftarrow \{v \mid \forall s' \in \mathcal{S}', v = \mathcal{F}(s')\}$ {Extract features for each article}
4: $\mathcal{M} \leftarrow \mathcal{M}()$ {Load the machine learning model for prediction}
5: $\mathcal{L} \leftarrow \{(s, \psi(\mathcal{M}, v)) \mid \forall s \in \mathcal{S}, \forall v \in \mathcal{V}, s' = \mathcal{P}(s)\}$ {Label each article as FAKE or REAL}
6: **for each** $s \in \mathcal{S}$ **do**
7: **if** $\mathcal{L}(s) = 1$ **then**
8: $\mathcal{F}\uparrow\downarrow\}(s, 1)$ {Flag the article as FAKE}
9: **end if**
10: **end for**
11: **return** $\mathcal{L}$
12: **Function** $\mathcal{P}(s)$:
13: $\mathcal{P} : \mathcal{S} \rightarrow \mathcal{S}'$
14: **return** Normalize(Tokenize(RemoveNoise(s))) {Return a preprocessed article}
15: **Function** $\mathcal{F}(s')$:
16: $\mathcal{F} : \mathcal{S}' \rightarrow \mathbb{R}^n$ {Mapping to an n -dimensional feature space}
17: **return** ExtractFeatures(s') {Return the n -dimensional feature vector}
18: **Function** $\mathcal{M}()$:
19: $\mathcal{M} : \emptyset \rightarrow \text{Model}$
20: **return** RetrievePretrainedModel() {Return the pre-trained fake news detection model}
21: **Function** $\psi(\mathcal{M}, v)$:
22: $\psi : \text{Model} \times \mathbb{R}^n \rightarrow \{0, 1\}$
23: **return** $\mathcal{M}(v)$ {Apply model $\mathcal{M}$ to feature vector v and return label}
24: **Procedure** $\mathcal{F}\uparrow\downarrow\}(s, \text{label})$:
25: $\mathcal{F}\uparrow\downarrow\} : \mathcal{S} \times \{0, 1\} \rightarrow \text{Void}$
26: **if** label = 1 **then**
27: MarkArticleAsFake(s) {Mark the article s as FAKE}
28: NotifyStakeholders(s) {Notify stakeholders of the FAKE article}
29: **end if**

TABLE 4. Definition of the symbols used in Algorithm 2.

Symbol	Definition
x	Bangla news article
F_{RoBERTa}	pre-trained RoBERTa model
F_{GCN}	Graph Convolutional Network model
$\mathcal{P}$	Article preprocessing
$\mathcal{E}$	Article embedding via RoBERTa
$\mathcal{G}$	Graph construction via embedding
$\mathcal{C}$	Classification function
$\mathcal{E}val$	Evaluation function

representation, which is then processed through several fully connected neural network layers. These layers reduce the information to a binary classification that, using a sigmoid function, distinguishes between real news and propaganda or fake news. This sophisticated architecture Figure 6 ensures nuanced detection tailored for the Bangla language.

As shown in Algorithm 2, our proposed RoBERTa-GCN begins by preprocessing the news article and then generates embeddings using a pre-trained RoBERTa [54] model. These embeddings are then used to construct a graph, where nodes represent different aspects of the article, and edges define their relationships. The GCN [55] is then employed to perform graph convolution on this structure, enhancing the node features by aggregating information from neighboring nodes. The resulting graph is then classified using a classification function, and the final output is evaluated based on accuracy, precision, recall, and F1-score. The notation used in Algorithm 2 is described in Table 4.

To represent both the headline and the content as a single

Algorithm 2 RoBERTa-GCN for news classification

Require: $x, F_{\text{RoBERTa}}, F_{\text{GCN}}$
Ensure: y - classification label (0 for Real, 1 for Fake)
1: **Procedure** RoBERTa-GCN($x, F_{\text{RoBERTa}}, F_{\text{GCN}}$):
2: $x' \leftarrow \mathcal{P}(x)$
3: $E \leftarrow \mathcal{E}(x', F_{\text{RoBERTa}})$
4: $G(V, E) \leftarrow \mathcal{G}(E)$
5: $G'(V', E') \leftarrow \text{GraphConvolution}(G, F_{\text{GCN}})$
6: $y_{\text{pred}} \leftarrow \mathcal{C}(G')$
7: $\mathcal{E}val(y_{\text{pred}}, y_{\text{true}})$
8: **return** $\mathcal{E}val$
9: **end Procedure**
10: **Function** $\mathcal{P}(x)$:
11: $\mathcal{P} : x \rightarrow x'$
12: **return** Preprocess(x)
13: **Function** $\mathcal{E}(x', F_{\text{RoBERTa}})$:
14: $\mathcal{E} : x' \rightarrow E$
15: **return** $F_{\text{RoBERTa}}(x')$
16: **Function** $\mathcal{G}(E)$:
17: $V \leftarrow \text{CreateNodes}(E)$
18: $E \leftarrow \text{CreateEdges}(V)$
19: **return** $G(V, E)$
20: **Function** GraphConvolution(G, F_{GCN}):
21: **for each** $v \in V$ **do**
22: $V' \leftarrow F_{\text{GCN}}(v, G)$
23: **return** $G'(V', E)$
24: **Function** $\mathcal{C}(G')$:
25: $\mathcal{C} : G' \rightarrow \{0, 1\}$
26: **return** Classify(G')
27: **Function** $\mathcal{E}val(y_{\text{pred}}, y_{\text{true}})$:
28: **return** (accuracy, precision, recall, F_1)

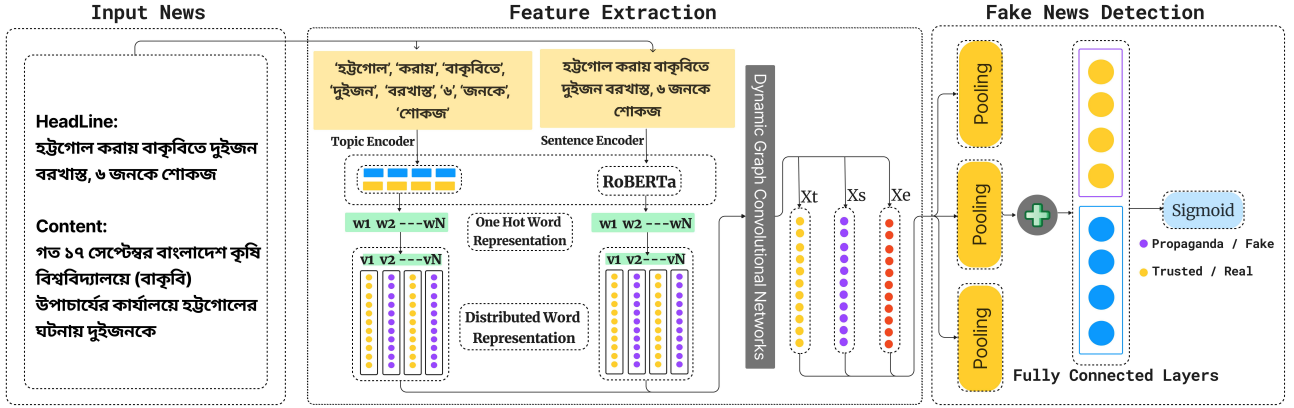


FIGURE 6. A convoluted working flow demonstrates a Bangla fake news detection model.

vector, embeddings from RoBERTa [54] for each word or subword token are first obtained. Then, a pooling operation (e.g., mean pooling or max pooling) is applied across these embeddings. This pooling process reduces the embeddings to a single vector that captures the essence of the entire article, providing a holistic representation that serves as input for the GCN. The graph is constructed by treating the unified representation vector as a feature of a single node or multiple nodes (if further segmenting the article). Edges are defined based on the relationships between these nodes, which could be determined through similarity measures, contextual relationships, or predefined rules capturing the structural and semantic connections within the text. This model is designed specifically to tackle the challenges of detecting fake news in Bangla, leveraging both language understanding and graph-based learning.

A. REPRESENTATION OF ROBERTA-GCN

As discussed earlier, RoBERTa-GCN is the final result of integrating the RoBERTa [54] and GCN [55] models. In this section, we present the mathematical representations that detail the integration of these models, forming our proposed RoBERTa-GCN for Bangla fake news classification.

1) RoBERTa Embedding Layer (E)

Given a Bangla news article represented as a sequence of tokens $x = \{x_1, x_2, \dots, x_n\}$, the RoBERTa [54] embedding layer transforms these tokens into dense vector representations:

$$E(x) = \{E(x_1), E(x_2), \dots, E(x_n)\} \quad (1)$$

where $E(x_i) \in \mathbb{R}^d$ is the d -dimensional embedding of token x_i obtained through RoBERTa.

2) GCN Layer

Consider a graph $G = (V, E)$ constructed from the dataset, with vertices V representing news articles and edges E encapsulating the relationships between them. The GCN [55]

layer utilizes graph convolution to gather data from neighboring nodes. The operation for the l -th layer of GCN is defined as follows:

$$H^{(l+1)} = \sigma(\tilde{A}H^{(l)}W^{(l)}) \quad (2)$$

In the given context, $\tilde{A}$ denotes the normalized adjacency matrix, $H^{(l)}$ represents the node feature matrix at layer l , and $W^{(l)}$ signifies the weight matrix for layer l .

3) Dynamic Adaptation Mechanism (D)

The dynamic component adjusts the GCN weights over time or other changing factors, represented by:

$$W^{(l)}(t) = g(W^{(l)}, t) \quad (3)$$

where g is a function that modifies the weights $W^{(l)}$ at each layer l in response to the temporal dimension t .

4) Classification Layer (C)

By passing the output of the final GCN layer $H^{(L)}$ through a classification layer, the probability that an individual article is of the 'Authentic' or 'Fake' class is determined:

$$P(y|x) = \text{sigmoid}\left(C\left(H^{(L)}\right)\right) \quad (4)$$

A sigmoid activation function is used in the RoBERTa-GCN architecture's classification layer to turn the output of the last graph convolutional layer into a probability distribution over binary class labels. Mathematically, this is articulated as $P(y|x) = \sigma(W_c h + b_c)$, where σ denotes the sigmoid function, W_c represents the weights of the classification layer, h is the output vector from the last GCN layer, and b_c is the bias term. The output of the sigmoid function, denoted by $\sigma(z) = \frac{1}{1+e^{-z}}$, is a scalar between 0 and 1, which interprets the likelihood of the input article x being classified as 'Authentic' (label close to 1) or 'Fake' (label close to 0).

5) Performance Optimization and Regularization (O)

The training objective includes a regularization component to mitigate overfitting.

$$L_{\text{total}} = L_{\text{original}} + \lambda \sum_{l=1}^L \|W^{(l)}\|_F^2 \quad (5)$$

Where L_{total} represents the total loss function, L_{original} is the original loss function, such as cross-entropy for classification. λ is the regularization coefficient. $\sum_{l=1}^L \|W^{(l)}\|_F^2$ denotes the sum of the squared Frobenius norms of the weight matrices $W^{(l)}$ at each layer l , which is a common regularization term to prevent overfitting.

6) Integration of RoBERTa-GCN

The complete "RoBERTa-GCN" model, incorporating all the components, can be formally represented as:

$$\begin{aligned} \text{RoBERTa-GCN}(x, t) &= P(y|x) \\ &= \text{softmax} \left(C(GCN(E(x), G(V, E), t)) \right) \end{aligned} \quad (6)$$

This equation shows the whole process that the model goes through, from using RoBERTa to embed Bangla news articles to classifying them using a dynamically adapted GCN. The training process also includes regularization and performance optimization. The evaluation metrics offer a quantifiable measure of the model's effectiveness in accurately categorizing news reports as 'Authentic' or 'Fake.'

The problem at hand is to construct a predictive model, denoted as RoBERTa-GCN(x, t), which discerns the veracity of Bangla news articles. The input x represents the textual data of an article, and t encapsulates temporal features. The model leverages the RoBERTa embeddings $E(x)$ to transform x into a vector space, subsequently processed by a (GCN). The GCN captures the inter-article relational dynamics over time, formalized as $GCN(E(x), G(V, E), t)$. The final output $y \in \{0, 1\}$, where 0 denotes 'Authentic' and 1 signifies 'Fake', is generated when a softmax function is applied to the output of the GCN, resulting in the probability distribution $P(y|x)$. The effectiveness of the model is assessed using conventional binary classification metrics.

7) Input News

The input consists of a headline and content represented as sequences of words:

$$\text{Headline} = (h_1, h_2, \dots, h_m) \quad (7)$$

$$\text{Content} = (c_1, c_2, \dots, c_n) \quad (8)$$

where h_i and c_j are tokens indexed by a function ϕ .

8) Feature Extraction

a: Topic Encoder

The topic encoder maps topics to one-hot encoded vectors in $\mathbb{R}^k$:

$$\text{TE}(t_i) = \mathbf{e}_i \quad \text{where} \quad \mathbf{e}_{ij} = \begin{cases} 1 & \text{if } i = j \\ 0 & \text{otherwise} \end{cases} \quad (9)$$

b: Sentence Encoder (RoBERTa)

The sentence encoding function by RoBERTa transforms the sentence into embedding vectors:

$$\text{SE}(c) = (\mathbf{e}_{c_1}, \mathbf{e}_{c_2}, \dots, \mathbf{e}_{c_n}) \quad (10)$$

c: Word Representation

Each word from the vocabulary is transformed into a dense vector via an embedding matrix:

$$\mathbf{e}_{w_i} = E \mathbf{v}_{w_i} \quad (11)$$

9) Dynamic s (GCN)

At each time step t , the GCN updates node features based on the graph structure:

$$H^{(l+1)}(t) = \sigma \left(\left(I - D(t)^{-\frac{1}{2}} A(t) D(t)^{-\frac{1}{2}} \right) H^{(l)}(t) W^{(l)} \right) \quad (12)$$

10) Fake News Detection

a: Pooling Layer

A pooling function reduces the dimensionality of the feature matrix:

$$h_{\text{pool}} = \text{pool} \left(\bigoplus_{t=1}^T H^{(L)}(t) \right) \quad (13)$$

b: Fully Connected Layers

The fully connected layers are defined recursively:

$$\mathbf{z}_l = \sigma(W_l \mathbf{z}_{l-1} + \mathbf{b}_l), \quad \mathbf{z}_0 = h_{\text{pool}} \quad (14)$$

c: Sigmoid Output

The final classification is determined by the sigmoid function:

$$P(\text{Fake}|\mathbf{z}_{L_{\text{fc}}}) = \sigma(\mathbf{w}_{\text{out}}^\top \mathbf{z}_{L_{\text{fc}}} + b_{\text{out}}) \quad (15)$$

B. EXPERIMENTAL SETUP

The dataset used in our study is first described, with an emphasis on its background and the methods used to compile it. After that, we'll go over the data pretreatment processes that were used to better feed our analytical models. After that, we go into depth about our research to find the optimal number of subjects and the process of hyperparameter tweaking in the modeling technique. Finally, we detail the measures used to evaluate our all-encompassing method for identifying Bangla-language fake news. This research constructs machine learning models using widely-used Python modules such as sklearn, keras, and tensorflow. The research tests

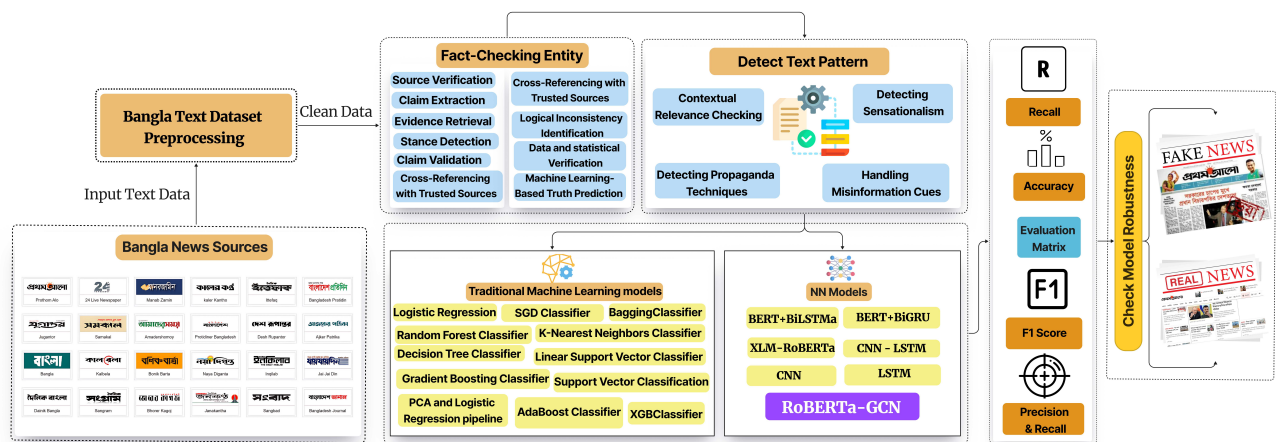


FIGURE 7. Overview of a comprehensive Bangla news detection framework. The framework includes preprocessing of Bangla text datasets, fact-checking entities, detecting text patterns, various machine learning models, and evaluation metrics to ensure robust fake news detection.

are carried out in a Google Colab setting using a powerful GPU backend with 52 GB of system RAM. The experimental setup utilizes an Intel processor as the system model. The investigations included using the Python 3 programming language to train and test the efficacy of the models. The evaluation measures included in the experiments include accuracy, precision, recall, F1-score, and time complexity, which are utilized to evaluate the performance of machine learning models.

1) Modeling Approach

Figure 7 outlines a comprehensive approach for detecting fake news in Bangla. It begins with preprocessing a dataset of Bangla news articles, and then a fact-checking entity runs various checks like source verification, evidence retrieval, and stance detection. Both traditional ML models and Neural Network (NN) models are used for analysis. The process also includes checking for sensationalism and propaganda techniques. Finally, an evaluation matrix with metrics like recall, accuracy, and the F1-score assesses the performance of the models. This structured framework is essential for developing effective tools to combat misinformation in Bangla news.

2) Hyperparameter Tuning

Hyperparameter tuning is a critical phase in ML where the ideal hyperparameters for a specific model are determined. This process entails experimenting with various hyperparameter combinations and assessing the model's effectiveness using a test set. The goal of hyperparameter tuning is to identify the hyperparameters that yield the highest performance on the test set. This step is vital as it can enhance the effectiveness of the ML methods. Table 5 shows the set of hyperparameters we utilized in our experiments.

Fact-Checking Entity. The Fact-Checking Entity for Bangla fake news operates through a systematic methodol-

ogy to ensure the credibility of information. It initiates source verification to authenticate the origin. Claim extraction isolates the main assertions for scrutiny. Evidence retrieval then gathers data supporting or refuting the claims. Stance detection assesses the evidence's position relative to the claim, followed by claim validation against established facts. Cross-referencing with trusted sources further corroborates findings, while logical inconsistency identification spots contradictions. Data and statistical verification quantitatively analyze the plausibility. Finally, machine learning-based truth prediction algorithms synthesize these steps, predicting the likelihood of the information being factual.

Detect Text Pattern: The Detect Text Pattern step in Bangla fake news analysis is a sophisticated process designed to identify distinctive linguistic patterns that are commonly associated with misinformation. This method involves computational analysis of the text structure, seeking out anomalies or irregularities in the use of language. It encompasses the identification of stylistic features, narrative frameworks, and discourse constructs that diverge from standard journalistic practices. By leveraging natural language processing tailored to the Bangla language, this step is pivotal in automating the recognition of fabricated content, streamlining the fact-checking process in the expansive and diverse Bangla-speaking media landscape.

Model Tuning. In addressing Bangla fake news, a suite of traditional machine learning models is deployed to discern patterns indicative of misinformation. These models include K-Nearest Neighbors(KNN) [56], RF [57], AdaBoost [58], Support Vector Machines [59], Stochastic Gradient Descent [60], XGBoost [61], Decision Trees [62], Bagging [63], Gradient Boosting(GB) [64], and pipelines integrating with . Each algorithm brings unique strengths in handling the complexities of Bangla text data, from probabilistic output to ensemble methods like RF and GB that capitalize on collective decision-making to improve prediction accuracy.

TABLE 5. The set of Hyperparameter settings used across different models for optimization.

Technique	Hyperparameters Values
Logistic Regression	solver='liblinear', penalty='l2', C=1.0, max_iter=100, fit_intercept=True, class_weight=None, random_state=None
Random Forest	n_estimators=100, max_depth=None, criterion='gini', min_samples_split=2, min_samples_leaf=1, max_features='auto'
SVC	C=1.0, kernel='rbf', degree=3, gamma='scale', coef0=0.0, shrinking=True, probability=False, tol=1e-3, cache_size=200
XGBC	n_estimators=100, max_depth=3, learning_rate=0.1, verbosity=1, objective='binary:logistic', booster='gbtree'
Decision Tree	criterion='gini', splitter='best', max_depth=None, min_samples_split=2, min_samples_leaf=1, min_impurity_decrease=0.0
GBC	loss='deviance', learning_rate=0.1, n_estimators=100, subsample=1.0, criterion='friedman_mse'
KNC	n_neighbors=5, weights='uniform', algorithm='auto', leaf_size=30, p=2, metric='minkowski', metric_params=None
AdaBoost	base_estimator=None, n_estimators=50, learning_rate=1.0, algorithm='SAMME.R'
SGD	loss='hinge', penalty='l2', alpha=0.0001, l1_ratio=0.15, fit_intercept=True, max_iter=1000, tol=1e-3, shuffle=True, verbose=0
LinearSVC	penalty='l2', loss='squared_hinge', dual=True, tol=1e-4, C=1.0, multi_class='ovr', fit_intercept=True, max_iter=1000

These techniques form a robust analytical framework essential for automating the detection of fake news in Bangla. Advanced NN models provide a sophisticated approach to text analysis. Utilizing architectures like BERT with Bi-LSTM or BiGRU layers enhances contextual understanding. XLM-RoBERTa [65] offers cross-lingual capabilities, which are crucial for diverse datasets. LSTM [66] networks capture long-term dependencies within text, while CNNs excel in pattern recognition. Combining CNN with RNN layers leverages both local feature extraction and sequence prediction. These NN models are fine-tuned on Bangla text data, capturing the language's unique syntactic and semantic nuances, thus elevating the accuracy and reliability of fake news classification.

Model Performance Evaluation. In the combat against Bangla fake news, evaluating the performance of detection models is paramount. Key metrics such as precision and recall, F1-score, and overall accuracy serve as the benchmarks. Precision measures the model's ability to correctly identify fake news, while recall assesses its capability to detect as many fake instances as possible. The F1-score provides a harmonic balance between precision and recall, which is crucial for models where equilibrium is needed. Ensuring model robustness involves testing with diverse data, including verified real and deceptive news articles, to validate the model's resilience and its effectiveness in the nuanced Bangla linguistic context.

C. METRICS EVALUATION

We adhere to a systematic approach that rigorously evaluates the model's performance using a comprehensive set of statistical metrics, including accuracy, precision, recall, and the F1-score. These metrics are meticulously calibrated to address the syntactic and semantic intricacies of the Bengali language. A comprehensive evaluation assesses the model's ability to distinguish between factual and fraudulent content accurately, ensuring its effectiveness in real-world scenarios where the accuracy of information is critical. The ultimate objective is to refine these models to achieve near-perfect precision, thereby enhancing the integrity of the Bangla digital news landscape.

1) Coherence UCI

The Coherence Uniform Coherence Index (UCI) evaluation for Bangla fake news detection is an advanced metric assessing the logical consistency of text within the dataset. By evaluating how well the narrative within the news articles holds together, it specifically addresses the particular challenge that Bangla's rich linguistic structure poses. This step is pivotal in discerning the authenticity of news by examining the flow and connection of ideas, an aspect where fabricated stories often falter. The evaluation ensures that the Bangla text data maintains a coherent and believable storyline, which is a hallmark of genuine news, thereby enhancing the detection process's effectiveness against the spread of misinformation. The UCI is formulated as follows:

$$UCI = \frac{1}{N} \sum_{i=1}^N \left(\frac{\text{Number of Coherent Segments}}{\text{Total Number of Segments}} \right)_i \quad (16)$$

UCI represents the Uniform Coherence Index for the dataset. N is the number of documents in the dataset. i refers to an individual document in the dataset. Number of Coherent Segments in document i is the count of segments that are determined to be coherent. Total Number of Segments in document i is the count of all segments in the document.

The Pointwise Mutual Information (PMI) evaluation in Bangla fake news detection quantifies the association between words within the text, which is critical for understanding contextual nuances in the Bengali language. This statistical measure helps identify unusual word pairings often found in fabricated news, enhancing the accuracy of detecting misinformation in Bangla text datasets. The mathematical formula for PMI is as follows:

$$\text{PMI}(x, y) = \log \left(\frac{P(x, y)}{P(x) \times P(y)} \right) \quad (17)$$

$\text{PMI}(x, y)$ is the PMI between events x and y . $P(x, y)$ is the joint probability of x and y occurring together. $P(x)$ and $P(y)$ are the individual probabilities of x and y occurring. $\log$ is the logarithm function.

2) Coherence NPMI

In the Bangla fake news analysis, the coherence Normalized NPMI evaluation is a pivotal step. It assesses the contextual congruence of words within Bangla text, which is crucial for detecting fabricated stories. This metric finely gauges the naturalness and logical flow of news content, enhancing the robustness of misinformation identification. The mathematical formula for NPMI is as follows:

$$\text{NPMI}(x, y) = \frac{\text{PMI}(x, y)}{-\log P(x, y)} \quad (18)$$

$\text{NPMI}(x, y)$ is the NPMI between events x and y . $\text{PMI}(x, y)$ is the pointwise mutual information, calculated as $\log \left(\frac{P(x, y)}{P(x) \times P(y)} \right)$. $P(x, y)$ is the joint probability of x and y occurring together. $\log$ represents the natural logarithm function.

3) Coherence CV

The COHERENCE CV (Coherence Vector) evaluation in Bangla fake news analysis is a sophisticated tool that leverages linguistic pattern recognition to ascertain thematic consistency. This method critically appraises the congruity of textual elements, effectively discerning anomalies indicative of misinformation, thus bolstering the integrity of news content in the Bengali language. The Coherence Vector (CV) involves several steps in its formulation:

- 1) Representing words as vectors using techniques like Word2Vec [67], GloVe [68], or FastText [69].
- 2) Forming sentence or document vectors by aggregating word vectors.
- 3) Measuring coherence through cosine similarity between an adjacent sentence or segment vectors.

The cosine similarity is calculated as:

$$\text{Cosine Similarity}(\vec{A}, \vec{B}) = \frac{\vec{A} \cdot \vec{B}}{|\vec{A}| |\vec{B}|} \quad (19)$$

$\vec{A}$ and $\vec{B}$ are vector representations of sentences or text segments. $\cdot$ denotes the dot product of the vectors. $|\vec{A}|$ and $|\vec{B}|$ are the magnitudes of the vectors.

4) Topic Diversity

The topic diversity evaluation in Bangla fake news analysis employs advanced text analytics to scrutinize thematic variations and breadth. This metric is instrumental in detecting aberrant topic distributions, a hallmark of disinformation. By leveraging language-specific models, it significantly refines the process of discerning authentic news from deceptive narratives in Bengali content. The Topic Diversity (TD) can be quantified using the entropy formula:

$$TD = - \sum_{i=1}^n p(i) \log p(i) \quad (20)$$

Here, n is the total number of unique topics in the text. $p(i)$ is the probability of the occurrence of topic i . $\log$ is the logarithm function, typically the natural logarithm.

5) Evaluation Metrics

For the binary classification task in Bangla Fake News Detection, which discerns between "Fake News" and "Real News", the total number of samples T evaluated is calculated using the following metric:

$$T = \sum_{i=1}^M (TP_i + FN_i) + \sum_{i=1}^M (FP_i + TN_i) \quad (21)$$

- **Clean Accuracy (Cln):** Clean Accuracy, often used to assess the performance of a model on a dataset free from adversarial examples or noise, is particularly crucial in contexts such as fake news detection, where precision is paramount. For a Bangla Fake News Detection system with two categories (usually "Fake News" and "Real News"). Given weights w_1 and w_2 for 'Real News' and 'Fake News' respectively, and a confidence threshold θ , the Clean Accuracy (Cln) is defined as:

$$Cln = \frac{w_1 \cdot TP_1 + w_2 \cdot TN_2}{W} \quad (22)$$

where W is the weighted sum of all true positive and true negative classifications that exceed the confidence threshold θ , computed as:

$$W = \sum_{i=1}^M w_i \cdot (TP_i + TN_i) \quad \text{for all } i \text{ where } P(\hat{y}_i | x_i) \geq \theta \quad (23)$$

Here, TP_1 is the number of True Positives for 'Real News', TN_2 is the number of True Negatives for 'Fake News', and $P(\hat{y}_i | x_i)$ denotes the predicted probability that sample x_i is classified as label $\hat{y}_i$. The sum W normalizes the accuracy score, considering the confidence of the predictions and the assigned weights to each class, which corrects for class imbalance.

where $M = 2$ represents the two categories of news.

- **Robust Accuracy (Boa):** Robust Accuracy (often denoted as Boa for "Best Overall Accuracy") is a metric designed to measure the performance of a classification model in the presence of adversarial examples or when the model is subjected to certain perturbations or noise. In the context of Bangla Fake News Detection, where there are two categories (fake news and Authentic news), the robust accuracy metric would ideally account for the model's resilience against such adversarial conditions. The following provides the formulation of Boa:

$$\begin{aligned} Boa = \frac{1}{N} \sum_{i=1}^N \left(w_{adv} \cdot \mathbb{I}[\hat{y}_i^{adv} = y_i] \cdot r_i \right. \\ \left. + w_{cln} \cdot \mathbb{I}[\hat{y}_i^{cln} = y_i] \cdot (1 - r_i) \right) \\ - \lambda \cdot \text{MCR} \end{aligned} \quad (24)$$

where, N is the total number of samples. $\hat{y}_i^{adv}$ and $\hat{y}_i^{cln}$ are the predicted labels under adversarial and clean

TABLE 6. Performance comparison of traditional Machine Learning algorithms based on various performance metrics such as Accuracy (ACC), Precision (PRE), Recall (REC), Matthews Correlation Coefficient (MCC), Hinge Loss (HL), Sensitivity (SEN), Specificity (SPE), Positive Predictive Value (PPV), Negative Predictive Value (NPV) without employing any classification thresholds.

Models	ACC	PRE	REC	MCC	HL	SEN	SPE	PPV	NPV
Linear Regression (LR)	0.879	0.920	0.971	0.475	0.037	0.921	0.323	0.910	0.871
Random Forest (RF)	0.917	0.837	0.866	0.580	0.041	0.805	0.438	0.879	0.926
Support Vector Machine (SVC)	0.853	<u>0.941</u>	0.905	<u>0.695</u>	0.029	0.934	0.249	0.961	0.845
XGBClassifier (XGBC)	0.890	0.921	0.807	0.515	0.058	0.780	0.964	0.809	0.868
Decision Tree (DT)	0.917	0.958	0.982	0.545	0.071	0.393	0.883	<u>0.953</u>	0.792
Gradient Boosting (GBC)	0.884	0.889	0.941	0.747	0.021	0.378	0.923	0.940	0.880
K-Nearest Neighbors (KNN)	0.954	0.923	<u>0.972</u>	0.632	0.304	0.892	0.332	0.871	0.814
AdaBoost	0.836	0.859	0.879	0.485	<u>0.287</u>	0.987	0.931	0.846	<u>0.915</u>
Stochastic Gradient Decent (SGD)	0.769	0.897	0.729	0.428	0.039	0.181	0.845	0.778	0.894
Linear SVC	<u>0.937</u>	0.889	0.918	0.325	0.031	<u>0.952</u>	<u>0.935</u>	0.912	0.884

conditions, respectively, y_i is the true label, w_{adv} and w_{cln} are weights for adversarial and clean condition performances. r_i is the resilience factor for each sample, $\mathbb{I}$ is the indicator function. λ is a regularization parameter, MCR is the Misclassification Rate.

- **Probability of Attack Success:** In the Probability of Attack Success (Succ) in Bangla Fake News Detection, where the system categorizes news into two classes ('Fake News' and 'Authentic News'), we can consider a scenario where the system is under adversarial attack. The aim would be to measure how often such attacks successfully deceive the system. This metric is particularly relevant in the context of robust machine learning, where models must resist adversarial examples designed to cause misclassification. The complex equation for Succ is defined as follows:

$$\text{Succ} = \frac{1}{N} \sum_{i=1}^N \left(s_i \cdot \mathbb{I}[\hat{y}_i^{adv} \neq y_i] \cdot p_i^{atk} + (1 - s_i) \cdot \mathbb{I}[\hat{y}_i^{cln} \neq y_i] \cdot (1 - p_i^{atk}) \right) \quad (25)$$

where, N is the total number of samples. s_i is the sensitivity factor for each sample. $\hat{y}_i^{adv}$ and $\hat{y}_i^{cln}$ are the predicted labels under adversarial and clean conditions, respectively. y_i is the true label. p_i^{atk} represents the probability of an attack on each sample. $\mathbb{I}[\cdot]$ is the indicator function.

V. EXPERIMENTAL RESULTS

The following sections discuss our experimental results and compare different traditional Machine learning and Deep learning models with our proposed RoBERTa-GCN.

A. RESULTS OF MACHINE LEARNING MODELS

The performance evaluation of various machine learning models for fake news classification, as shown in Tables 6 and 7, reveals distinct strengths across different metrics. In the analysis without classification thresholds as presented in Table 6, K-Nearest Neighbors emerged as the top performer with the highest accuracy of 0.954, MCC of 0.632, and specificity of 0.871. The decision Tree excelled in the recall

at 0.982 and precision at 0.958, highlighting its capability to identify true positives, while Random Forest (RF) led to a negative predictive value of 0.926.

Our extensive evaluations of the same models for different classification thresholds are presented in Table 7. At the 0.05 threshold, Random Forest showed superior results with the highest accuracy of 0.992, MCC of 0.908, and specificity of 0.965. XGBClassifier led the recall at 0.998 and exhibited the lowest hinge loss of 0.018. At the 0.1 threshold, K-Nearest Neighbors stood out with the highest accuracy of 0.966, recall of 0.996, and sensitivity of 0.996, making it the most reliable for detecting true positives. The Decision Tree maintained its strength in MCC at 0.915 and specificity at 0.942, indicating robust classification capabilities. Support Vector Machine models also performed well, particularly in precision and positive predictive value. Overall, K-Nearest Neighbors, Random Forest, and Decision Tree models consistently ranked among the top performers, with specific accuracy, precision, and recall strengths, making them suitable for different aspects of fake news classification tasks depending on the specific metric requirements.

B. RESULTS OF DEEP LEARNING MODELS

A comparative analysis of the performance metrics for various deep learning models, including BERT + BiLSTM, BERT + BiGRU, XLM-RoBERTa, LSTM, CNN, CNN-LSTM, and the proposed RoBERTa-GCN model is presented in Table 8. Among the models, the proposed RoBERTa-GCN achieves the highest accuracy, with a score of 0.986, demonstrating its superior ability to classify Bangla news articles as authentic or fake correctly. It also outperforms the other models in precision, achieving 0.972, Matthews correlation coefficient with a score of 0.957, sensitivity at 0.986, and negative predictive value reaching 0.978. These results indicate its robustness and reliability in detecting fake news. Additionally, RoBERTa-GCN exhibits the lowest Hamming loss at 0.017 and specificity at 0.985, further confirming its effectiveness in minimizing classification errors.

Compared to traditional models like BERT + BiLSTM and CNN-LSTM, which achieved accuracies of 0.944 and 0.913, respectively, RoBERTa-GCN's performance is significantly higher. This result highlights the advantage of integrating

TABLE 7. Comparative of various machine learning models based on different performance metrics for the classification thresholds of 0.05 and 0.1.

Models	Fake ML 0.05								Fake ML 0.1							
	ACC	PRE	REC	MCC	HL	SPE	PPV	NPV	ACC	PRE	REC	MCC	HL	SPE	PPV	NPV
LR	0.970	0.973	<u>0.997</u>	0.583	0.031	0.785	0.943	0.875	0.900	0.874	0.935	0.785	0.045	0.786	0.914	0.895
RF	0.992	<u>0.992</u>	0.941	0.908	0.085	0.965	0.893	0.936	0.894	<u>0.965</u>	0.892	0.698	<u>0.023</u>	0.887	0.897	0.912
SVC	<u>0.986</u>	0.988	0.998	0.829	0.014	0.895	0.885	0.882	0.908	0.895	0.910	0.723	0.012	0.935	0.793	0.874
XGBC	0.992	0.917	0.998	<u>0.906</u>	<u>0.018</u>	0.912	0.895	0.910	0.914	0.905	0.941	0.866	0.085	0.787	0.698	<u>0.932</u>
DT	0.979	0.994	0.985	0.787	0.208	<u>0.935</u>	0.784	0.897	0.936	0.963	0.932	<u>0.915</u>	0.036	0.894	<u>0.940</u>	0.913
GBC	0.965	0.967	0.996	0.477	0.036	0.893	0.936	<u>0.940</u>	0.943	0.918	0.894	0.943	0.047	0.942	0.895	0.923
KNN	0.969	0.971	0.997	0.564	0.031	0.782	0.890	0.935	0.893	0.936	0.967	0.896	0.065	0.914	0.918	0.914
AdaBoost	0.961	0.970	0.991	0.460	0.039	0.938	0.901	0.941	0.915	0.841	0.944	0.875	0.074	0.874	0.932	0.895
SGD	0.964	0.944	0.895	0.451	0.036	0.923	0.913	0.785	0.966	0.969	0.996	0.506	0.063	<u>0.941</u>	0.892	0.923
LinearSVC	0.971	0.975	0.896	0.602	0.029	0.789	<u>0.941</u>	0.931	<u>0.945</u>	0.892	0.914	0.651	0.071	0.936	0.969	0.935

TABLE 8. Performance metrics comparison for traditional DL algorithms with RoBERTa-GCN

Models	ACC	PRE	REC	MCC	HL	SEN	SPE	PPV	NPV
BERT + BiLSTMa	<u>0.944</u>	0.941	0.961	0.918	0.024	0.917	0.932	0.935	0.918
BERT + BiGRU	0.936	0.952	0.916	0.880	0.081	0.735	<u>0.978</u>	0.914	0.904
XLM-RoBERTa	0.953	0.948	0.924	0.941	0.125	<u>0.983</u>	0.214	0.951	<u>0.958</u>
LSTM	0.890	<u>0.954</u>	0.934	0.935	0.061	0.940	0.964	<u>0.969</u>	0.898
CNN	0.917	0.938	0.932	<u>0.947</u>	0.042	0.823	0.883	0.974	0.897
CNN-LSTM	0.913	0.789	0.941	0.937	<u>0.018</u>	0.878	0.923	0.957	0.934
RoBERTa-GCN	0.986	0.972	<u>0.913</u>	0.957	0.017	0.986	0.985	<u>0.969</u>	0.978

TABLE 9. Comparison of our proposed RoBERTa-GCN with different model learning techniques and their performances based on accuracy (ACC) metrics.

Models	Year	Methods	ACC
SVM [3]	2020	Machine Learning	0.910
Random Forest [26]	2020	Machine Learning	0.850
Hybrid CNN-RNN [24]	2021	Deep Learning	0.600
BerConvoNet [31]	2021	Deep Learning	0.974
BERT [27]	2022	Machine Learning	0.900
BanglaBERT [49]	2022	Deep Learning	0.834
XLM-RoBERT [29]	2023	Deep Learning	0.933
FNDNet [70]	2023	Deep Learning	0.983
CT-BERT+BiGRU [33]	2023	Deep Learning	<u>0.984</u>
RoBERTa-GCN (Our study)	2024	Deep Learning	0.986

RoBERTa's contextual language understanding with the relational learning capabilities of GCN, making RoBERTa-GCN the most effective model in this study for Bangla fake news detection.

C. COMPARATIVE ANALYSIS

Figure 8 presents a comprehensive evaluation of various machine learning models against multiple performance metrics. The performance scores are plotted to facilitate a visual comparison, highlighting the efficacy of each model across the metrics. Adding RoBERTa-GCN makes it possible to compare it to more advanced models that use both transformer architectures and GCN. These models are hybrid models. Moreover, Table 9 compares the performance accuracy of various learning techniques published between 2020 and 2023 for fake news detection. Early approaches, such as the Hybrid CNN-RNN [24] model in 2021, achieved a moderate accuracy of 0.600, indicating room for improvement in deep learning models at that time. As the field advanced, significant gains were seen with models like BERT [27] and

BanglaBERT [49], which achieved 0.90 and 0.834 accuracy, respectively in 2022, and traditional machine learning techniques like SVM [3] and RF [26], which reported 0.910 and 0.85 accuracy, respectively in 2020. Transfer learning also proved effective, with XLM-RoBERTa [29] achieving 0.9331 accuracy in 2023. Among deep learning models, BerConvoNet [31] and FNDNet [70] showed strong performance with accuracies of 0.974 and 0.983, respectively, while the CT-BERT+BiGRU [33] model reached 0.984 in 2023. The key finding from this table is that our proposed RoBERTa-GCN model outperforms all other models with the highest accuracy of 0.986. This highlights the superior effectiveness of integrating RoBERTa's language understanding with GCN's relational learning capabilities, marking a significant advancement in the field of Bangla fake news detection. Overall, the table underscores the steady progression of model performance over time, culminating in the SOTA results achieved by RoBERTa-GCN.

The ROC curve in Figure 9 displays the performance of a binary classifier, which could be the RoBERTa-GCN model, by plotting the True Positive Rate (TPR) against the False Positive Rate (FPR). The ROC curve shows an excellent diagnostic ability of the model, as indicated by the AUC of 0.96, suggesting that the model can distinguish between classes with high accuracy. The curve closely approaches the upper left corner of the graph, which represents an ideal classifier yielding a high TPR and a low FPR. This is indicative of a highly effective model in terms of both sensitivity (identifying true positives) and specificity (correctly avoiding false positives). The model's performance, as illustrated by this ROC curve, is therefore exemplary, with the curve nearly reaching the perfect classification point at the top left.

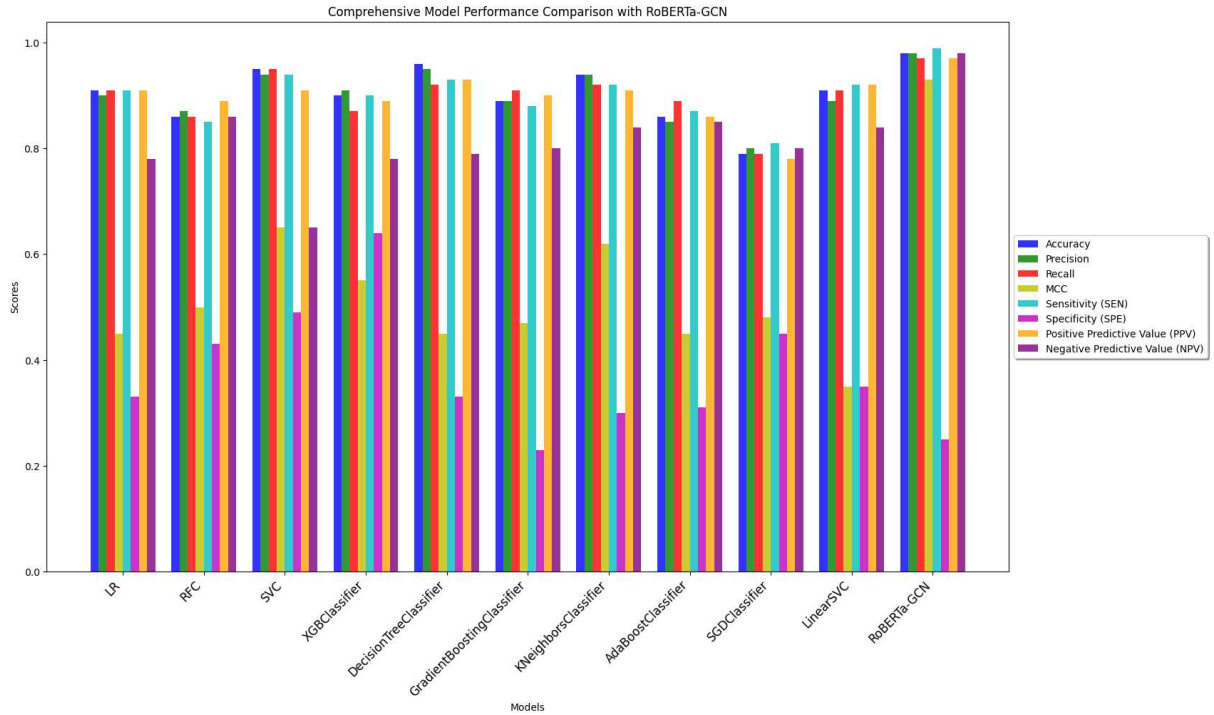


FIGURE 8. Comparison of proposed RoBERTa-GCN with other models based on different performance metrics.

D. EXPLAINABILITY OF ROBERTA-GCN

Figure 10 depicts the output of a Local Interpretable Model-agnostic Explanations (LIME) analysis applied to our proposed RoBERTa-GCN model. It helps to explain the predictions of our RoBERTa-GCN model by highlighting the features, in this case, words that contribute most to the model's prediction. Based on the Figure 10 we can get some insights as follows:

- 1) The first text has a prediction probability of 0.92 for being authentic and 0.08 for being fake. Certain words are highlighted, which likely influenced the model's decision to classify it as authentic.
- 2) The second text has the same prediction probabilities

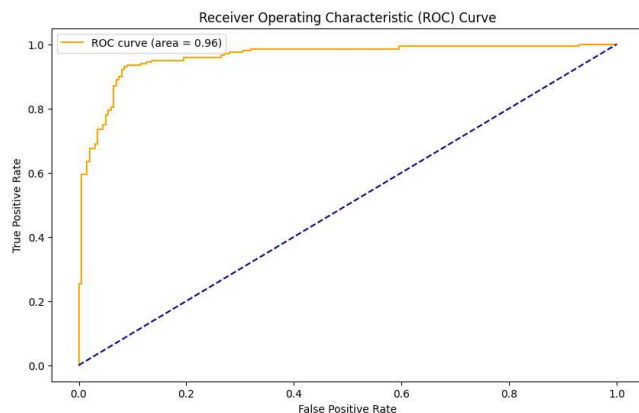


FIGURE 9. The obtained ROC curve of our proposed RoBERTa-GCN model.

as the first one, with different words highlighted. This suggests that those words have a significant impact on the model's prediction of authenticity.

- 3) The third text is predicted with absolute certainty (probability 1.00) to be fake, and no probability is assigned to it being authentic. Highlighted words in the text are the features that the model found most indicative of it being fake.

The right side of the image, where the texts with highlighted words are shown, illustrates how LIME provides local interpretability. For each prediction, it points out the specific words that have the highest weight in the model's decision-making process. This is crucial for understanding why the model makes certain decisions and can be used to improve the model by, for example, adjusting feature weightings or by providing more training data to reduce biases.

Moreover, the LIME analysis reveals that our proposed RoBERTa-GCN model does not rely solely on superficial cues or common phrases but rather on a nuanced understanding of the text's context, enabled by the integration of RoBERTa's language understanding and GCN's relational reasoning. This interpretability is crucial for validating the model's reliability, especially in sensitive applications like fake news detection, where understanding the rationale behind a classification is as important as the classification itself.

VI. CONCLUSION

This paper represents a pivotal stride in combating the escalating problem of misinformation in the digital age, particularly within the context of the Bengali language. As the



FIGURE 10. LIME analysis for Bangla fake news detection, highlighting words influencing model predictions. Blue indicates words supporting authenticity; orange indicates words supporting fake news.

digital ecosystem becomes increasingly saturated with deceptive content, the need for sophisticated, language-specific solutions has never been more pressing. Addressing this need, the research introduces the RoBERTa-GCN model, a fusion of RoBERTa's advanced natural language processing capabilities and GCN's adeptness at managing relational data. This model is meticulously tailored to grasp the intricacies of the Bangla language, encompassing its rich linguistic nuances and cultural contexts. Trained on a comprehensive dataset collected from diverse Bangladeshi news sources, the RoBERTa-GCN model is rigorously optimized to distinguish authentic news from fabricated narratives effectively. Its performance, benchmarked against several baseline models, demonstrates a marked improvement, affirming its potential as a critical tool in the arsenal against misinformation.

The significance of the RoBERTa-GCN model transcends its technical achievements; it addresses a crucial gap in the field of fake news detection for languages that have traditionally been underrepresented in global NLP research. By leveraging state-of-the-art NLP techniques and adapting to the unique challenges posed by the Bangla language, the model not only sets a new standard for fake news detection but also underscores the importance of incorporating linguistic and cultural idiosyncrasies into AI-driven solutions. The model's evaluation, employing metrics like accuracy, precision, recall, and F1-score, attests to its robustness and reliability in real-world settings. Furthermore, the use of advanced coherence measures like UCI, PMI, and NPMI ensures that the model's performance is both consistent and dependable, making it a formidable tool against the spread of misinformation. The study's all-around approach, which

includes a new model design and a deep understanding of regional linguistic traits, opens the door for more research in this area and shows how important customized AI solutions are for keeping information safe in a world that is becoming more and more connected.

A. LIMITATIONS AND FUTURE SCOPES

Although RoBERTa-GCN is good at identifying Bangla fake news, it has certain shortcomings that might make it less applicable to other languages and training datasets. A potential future direction for improving the model's usefulness is to investigate ways to include other South Asian languages. In addition, there is potential for modifying RoBERTa-GCN to identify more nuanced types of disinformation, such as biased journalism or satire. To further explore the model's efficacy across other social media platforms and include real-time data analysis, more studies might be conducted to expand its practical applicability.

ACKNOWLEDGEMENT

We would like to acknowledge Dr. Sami Azam and Dr. Asif Karim from the Faculty of Science and Technology, Charles Darwin University, Casuarina, NT 0909, Australia, for their assistance in securing funding and their thorough review of this paper.

REFERENCES

- [1] M. H. Goldani, R. Safabakhsh, and S. Momtazi, "Convolutional neural network with margin loss for fake news detection," *Information Processing & Management*, vol. 58, no. 1, p. 102418, 2021.

- [2] M. Narra, M. Umer, S. Sadiq, H. Karamti, A. Mohamed, I. Ashraf, et al., "Selective feature sets based fake news detection for covid-19 to manage infodemic," *IEEE Access*, vol. 10, pp. 98724–98736, 2022.
- [3] M. Z. Hossain, M. A. Rahman, M. S. Islam, and S. Kar, "Banfake-news: A dataset for detecting fake news in bangla," *arXiv preprint arXiv:2004.08789*, 2020.
- [4] D. Wang, W. Zhang, W. Wu, and X. Guo, "Soft-label for multi-domain fake news detection," *IEEE Access*, 2023.
- [5] T. Jiang, J. P. Li, A. U. Haq, A. Saboor, and A. Ali, "A novel stacking approach for accurate detection of fake news," *IEEE Access*, vol. 9, pp. 22626–22639, 2021.
- [6] M. G. Hussain, M. R. Hasan, M. Rahman, J. Protim, and S. Al Hasan, "Detection of bangla fake news using mnb and svm classifier," in 2020 International Conference on Computing, Electronics & Communications Engineering (iCCECE), pp. 81–85, IEEE, 2020.
- [7] A. Roets et al., "'fake news': Incorrect, but hard to correct. the role of cognitive ability on the impact of false information on social impressions," *Intelligence*, vol. 65, pp. 107–110, 2017.
- [8] H. Saleh, A. Alharbi, and S. H. Alsamhi, "Opcnn-fake: Optimized convolutional neural network for fake news detection," *IEEE Access*, vol. 9, pp. 129471–129489, 2021.
- [9] E. Grave, P. Bojanowski, P. Gupta, A. Joulin, and T. Mikolov, "Learning word vectors for 157 languages," *arXiv preprint arXiv:1802.06893*, 2018.
- [10] S. Vosoughi, D. Roy, and S. Aral, "The spread of true and false news online," *science*, vol. 359, no. 6380, pp. 1146–1151, 2018.
- [11] S. KEMP, "Digital 2023: Bangladesh — datareportal – global digital insights," <https://datareportal.com/reports/digital-2023-bangladesh>, Feb.13, 2023. (Accessed on 12/21/2023).
- [12] M. Hossain, "Nothing but fb. why bangladeshis never took to twitter, threads and the like | the business standard," <https://www.tbsnews.net/features/panorama/nothing-fb-why-bangladeshis-never-took-twitter-threads-and-667186>, July.18, 2023. (Accessed on 11/15/2023).
- [13] S. Report, "Mobs beat 2 dead for 'kidnapping' | daily star," <https://www.thedailystar.net/frontpage/news/mobs-beat-2-dead-kidnapping-1774471>, July .21, 2019. (Accessed on 12/07/2023).
- [14] I. Ahmed and J. A. Manik, "A hazy picture appears | the daily star," <https://www.thedailystar.net/news-detail-252212>, Oct .3, 2012. (Accessed on 12/02/2023).
- [15] R. Rafe, "Misinformation mars bangladesh vaccination drive – dw – 01/27/2021," <https://www.dw.com/en/covid-bangladesh-vaccination-drive-marred-by-misinformation/a-56360529>, Jan .27, 2021. (Accessed on 01/05/2024).
- [16] T. Report, "Busting the top 3 fake news of the week | the business standard," <https://www.tbsnews.net/thoughts/busting-top-3-fake-news-week-173236>, Dec .18, 2020. (Accessed on 01/10/2024).
- [17] "PolitiFact," <https://www.politifact.com/>. (Accessed on 01/25/2024).
- [18] "Factcheck.org - a project of the annenberg public policy center," <https://www.factcheck.org/>. (Accessed on 01/09/2024).
- [19] "Verification - fact check bangladesh. fact-check bangladesh," <https://www.jachai.org/>. (Accessed on 01/20/2024).
- [20] Y. Long, Q. Lu, R. Xiang, M. Li, and C.-R. Huang, "Fake news detection through multi-perspective speaker profiles," in Proceedings of the eighth international joint conference on natural language processing (volume 2: Short papers), pp. 252–256, 2017.
- [21] G. Karadzhov, P. Nakov, L. Márquez, A. Barrón-Cedeño, and I. Koychev, "Fully automated fact checking using external sources," *arXiv preprint arXiv:1710.00341*, 2017.
- [22] M. A. Haque Palash, A. Khan, K. Islam, M. A. Al Nasim, and R. M. Bin Shahjahan, "Incongruity detection between bangla news headline and body content through graph neural network," in The Fourth Industrial Revolution and Beyond: Select Proceedings of IC4IR+, pp. 375–387, Springer, 2023.
- [23] R. Kawsar, "Bangla ranked at 7th among 100 most spoken languages worldwide," <https://www.dhakatribune.com/world/201648/bangla-ranked-at-7th-among-100-most-spoken>, Feb .17, 2020. (Accessed on 01/01/2024).
- [24] J. A. Nasir, O. S. Khan, and I. Varlamis, "Fake news detection: A hybrid cnn-rnn based deep learning approach," *International Journal of Information Management Data Insights*, vol. 1, no. 1, p. 100007, 2021.
- [25] A. Anjum, M. Keya, A. K. M. Masum, and S. R. H. Noori, "Fake and authentic news detection using social data strivings," in 2021 12th International Conference on Computing Communication and Networking Technologies (ICCCNT), pp. 1–5, IEEE, 2021.
- [26] F. Islam, M. M. Alam, S. S. Hossain, A. Motaleb, S. Yeasmin, M. Hasan, and R. M. Rahman, "Bengali fake news detection," in 2020 IEEE 10th International Conference on Intelligent Systems (IS), pp. 281–287, IEEE, 2020.
- [27] M. A. Al Ghamdi, M. S. Bhatti, A. Saeed, Z. Gillani, and S. H. Almotiri, "A fusion of bert, machine learning and manual approach for fake news detection," *Multimedia Tools and Applications*, pp. 1–18, 2023.
- [28] E. Raja, B. Soni, and S. K. Borgohain, "Fake news detection in dravidian languages using transfer learning with adaptive finetuning," *Engineering Applications of Artificial Intelligence*, vol. 126, p. 106877, 2023.
- [29] P. B. Pranto, S. Z.-U.-H. Navid, P. Dey, G. Uddin, and A. Iqbal, "Are you misinformed? a study of covid-related fake news in bengali on facebook," *arXiv preprint arXiv:2203.11669*, 2022.
- [30] S. Rohman, J. Ferdous, S. M. R. Ullah, and M. A. Rahman, "Ibfnf: An improved dataset for bangla fake news detection and comparative analysis of performance of baseline models," in 2023 International Conference on Next-Generation Computing, IoT and Machine Learning (NCIM), pp. 1–6, IEEE, 2023.
- [31] M. Choudhary, S. S. Chouhan, E. S. Pilli, and S. K. Vipparthi, "Berconvonet: A deep learning framework for fake news classification," *Applied Soft Computing*, vol. 110, p. 107614, 2021.
- [32] M. A. Ali, M. L. Matubber, V. Sharma, and B. Balamurugan, "An improved and efficient technique for detecting bengali fake news using machine learning algorithms," in 2022 6th International Conference On Computing, Communication, Control And Automation (ICCUBEA), pp. 1–4, IEEE, 2022.
- [33] J. Alghamdi, Y. Lin, and S. Luo, "Towards covid-19 fake news detection using transformer-based models," *Knowledge-Based Systems*, vol. 274, p. 110642, 2023.
- [34] D. K. Dixit, A. Bhagat, and D. Dangi, "Automating fake news detection using ppca and levy flight-based lstm," *Soft Computing*, vol. 26, no. 22, pp. 12545–12557, 2022.
- [35] N. Rai, D. Kumar, N. Kaushik, C. Raj, and A. Ali, "Fake news classification using transformer based enhanced lstm and bert," *International Journal of Cognitive Computing in Engineering*, vol. 3, pp. 98–105, 2022.
- [36] M. Sudhakar and K. Kaliyammurthi, "Effective prediction of fake news using two machine learning algorithms," *Measurement: Sensors*, vol. 24, p. 100495, 2022.
- [37] A. Choudhary and A. Arora, "Linguistic feature based learning model for fake news detection and classification," *Expert Systems with Applications*, vol. 169, p. 114171, 2021.
- [38] V. Khullar and H. P. Singh, "f-fnc: Privacy concerned efficient federated approach for fake news classification," *Information Sciences*, vol. 639, p. 119017, 2023.
- [39] H. Xia, Y. Wang, J. Z. Zhang, L. J. Zheng, M. M. Kamal, and V. Arya, "Covid-19 fake news detection: A hybrid cnn-bilstm-am model," *Technological Forecasting and Social Change*, vol. 195, p. 122746, 2023.
- [40] M. Hosseini, A. J. Sabet, S. He, and D. Aguiar, "Interpretable fake news detection with topic and deep variational models," *Online Social Networks and Media*, vol. 36, p. 100249, 2023.
- [41] B. Palani and S. Elango, "Bbc-fnd: An ensemble of deep learning framework for textual fake news detection," *Computers and Electrical Engineering*, vol. 110, p. 108866, 2023.
- [42] O. B. Okunoye and A. E. Ibor, "Hybrid fake news detection technique with genetic search and deep learning," *Computers and Electrical Engineering*, vol. 103, p. 108344, 2022.
- [43] P. Malhotra and S. K. Malik, "Fake news detection using ensemble techniques," *Multimedia Tools and Applications*, 2023.
- [44] R. Mohawesh, S. Maqsood, and Q. Althebyan, "Multilingual deep learning framework for fake news detection using capsule neural network," *Journal of Intelligent Information Systems*, pp. 1–17, 2023.
- [45] V. Jain, R. K. Kaliyar, A. Goswami, P. Narang, and Y. Sharma, "Aenet: an attention-enabled neural architecture for fake news detection using contextual features," *Neural Computing and Applications*, vol. 34, no. 1, pp. 771–782, 2022.
- [46] B. Palani, S. Elango, and V. Viswanathan K, "Cb-fake: A multimodal deep learning framework for automatic fake news detection using capsule neural network and bert," *Multimedia Tools and Applications*, vol. 81, no. 4, pp. 5587–5620, 2022.
- [47] R. Katarya, D. Dahiya, S. Checker, et al., "Fake news detection system using featured-based optimized msvm classification," *IEEE Access*, vol. 10, pp. 113184–113199, 2022.

- [48] K. M. Hasib, N. A. Towhid, K. O. Faruk, J. Al Mahmud, and M. Mridha, "Strategies for enhancing the performance of news article classification in bangla: Handling imbalance and interpretation," *Engineering Applications of Artificial Intelligence*, vol. 125, p. 106688, 2023.
- [49] A. Bhattacharjee, T. Hasan, W. U. Ahmad, K. Samin, M. S. Islam, A. Iqbal, M. S. Rahman, and R. Shahriyar, "Banglabert: Language model pretraining and benchmarks for low-resource language understanding evaluation in bangla," *arXiv preprint arXiv:2101.00204*, 2021.
- [50] W. A. Qader, M. M. Ameen, and B. I. Ahmed, "An overview of bag of words; importance, implementation, applications, and challenges," in 2019 international engineering conference (IEC), pp. 200–204, IEEE, 2019.
- [51] J. Ramos et al., "Using tf-idf to determine word relevance in document queries," in *Proceedings of the first instructional conference on machine learning*, vol. 242, pp. 29–48, Citeseer, 2003.
- [52] Z. Yin and Y. Shen, "On the dimensionality of word embedding," *Advances in neural information processing systems*, vol. 31, 2018.
- [53] G. Sidorov, F. Velasquez, E. Stamatatos, A. Gelbukh, and L. Chanona-Hernández, "Syntactic n-grams as machine learning features for natural language processing," *Expert Systems with Applications*, vol. 41, no. 3, pp. 853–860, 2014.
- [54] Y. Liu, M. Ott, N. Goyal, J. Du, M. Joshi, D. Chen, O. Levy, M. Lewis, L. Zettlemoyer, and V. Stoyanov, "Roberta: A robustly optimized bert pretraining approach," *arXiv preprint arXiv:1907.11692*, 2019.
- [55] T. N. Kipf and M. Welling, "Semi-supervised classification with graph convolutional networks," *arXiv preprint arXiv:1609.02907*, 2016.
- [56] L. E. Peterson, "K-nearest neighbor," *Scholarpedia*, vol. 4, no. 2, p. 1883.
- [57] L. Breiman, "Random forests," *Machine learning*, vol. 45, pp. 5–32, 2001.
- [58] Y. Freund and R. E. Schapire, "A decision-theoretic generalization of on-line learning and an application to boosting," *Journal of computer and system sciences*, vol. 55, no. 1, pp. 119–139, 1997.
- [59] C. Cortes and V. Vapnik, "Support-vector networks," *Machine learning*, vol. 20, pp. 273–297, 1995.
- [60] J. Bottou, "Stochastic gradient descent tricks," in *Neural Networks: Tricks of the Trade: Second Edition*, pp. 421–436, Springer, 2012.
- [61] T. Chen and C. Guestrin, "Xgboost: A scalable tree boosting system," in *Proceedings of the 22nd acm sigkdd international conference on knowledge discovery and data mining*, pp. 785–794, 2016.
- [62] W.-Y. Loh, "Classification and regression trees," *Wiley interdisciplinary reviews: data mining and knowledge discovery*, vol. 1, no. 1, pp. 14–23, 2011.
- [63] L. Breiman, "Bagging predictors," *Machine learning*, vol. 24, pp. 123–140, 1996.
- [64] J. H. Friedman, "Greedy function approximation: a gradient boosting machine," *Annals of statistics*, pp. 1189–1232, 2001.
- [65] A. Conneau, K. Khandelwal, N. Goyal, V. Chaudhary, G. Wenzek, F. Guzmán, E. Grave, M. Ott, L. Zettlemoyer, and V. Stoyanov, "Un-supervised cross-lingual representation learning at scale," *arXiv preprint arXiv:1911.02116*, 2019.
- [66] M. Sundermeyer, R. Schlüter, and H. Ney, "Lstm neural networks for language modeling," in *Interspeech*, vol. 2012, pp. 194–197, 2012.
- [67] K. W. Church, "Word2vec," *Natural Language Engineering*, vol. 23, no. 1, pp. 155–162, 2017.
- [68] J. Pennington, R. Socher, and C. D. Manning, "Glove: Global vectors for word representation," in *Proceedings of the 2014 conference on empirical methods in natural language processing (EMNLP)*, pp. 1532–1543, 2014.
- [69] A. Joulin, E. Grave, P. Bojanowski, M. Douze, H. Jégou, and T. Mikolov, "Fasttext. zip: Compressing text classification models," *arXiv preprint arXiv:1612.03651*, 2016.
- [70] R. K. Kaliyar, A. Goswami, P. Narang, and S. Sinha, "Fndnet—a deep convolutional neural network for fake news detection," *Cognitive Systems Research*, vol. 61, pp. 32–44, 2020.



Award from the International Conference ICETET-SIP-22.

MEJBAH AHAMMAD received his BSc in computer science and engineering with the major of computer vision and his MSc in computer science with the major of artificial intelligence from American International University, Bangladesh. He is currently the CEO of Bytes of Intelligence, New York, USA, and works as a deep learning and AI-specialized consultant and instructor at aiQuest Intelligence, Dhaka, Bangladesh. He has authored over 11 research papers. He has the Best Research



gories to optimize manufacturing processes, enhance production efficiency, and innovate smart industrial systems. His academic and research career demonstrates a deep fascination with the intersection of engineering and cutting-edge technologies, emphasizing an interdisciplinary approach. AL SANI is eager to contribute to the advancement of intelligent industrial systems through groundbreaking research and innovative solutions.

AL SANI received his B.Sc. degree in Industrial and Production Engineering from Jashore University of Science and Technology (JUST), Jashore-7408, Bangladesh. His current research interests include artificial intelligence (AI), machine learning (ML), and deep learning (DL), specifically image processing, computer vision, and natural language processing. Additionally, he is passionate about exploring the integration of industrial engineering principles with AI technologies



computer vision, medical image analysis and natural language processing. He is passionate about applying sophisticated computational tools to real-world problems. He is also actively involved in the academic and research areas.

KHALILUR RAHMAN received the B.Sc. degree in Computer Science and Engineering from Sylhet Engineering College, Sylhet-3100, Bangladesh. His research interest in Deep learning, Machine learning, and Artificial intelligence. He has worked on various projects in the field of application of machine learning and deep learning. His current research interests include artificial intelligence (AI), machine learning (ML), and deep learning (DL), specifically image processing,



passionate about applying advanced computational techniques to solve real-world challenges and has been actively involved in academic and research communities.

MD TANVIR ISLAM is a graduate student pursuing an M.Sc. in Computer Science and Engineering at Sungkyunkwan University (SKKU), South Korea. Specializing in AI-driven image processing, his research has led to publications in top-tier journals and conferences such as "*Alexandria Engineering Journal*" and "*ACM Multimedia*". He has experience as a Research Assistant at "*VIS2KNOW Lab*", focusing on computer vision, deep learning, and machine learning. He is



MD MOSTAFIZUR RAHMAN MASUD received his Bachelor of Science degree in Software Engineering from Universiti Teknologi Malaysia, Johor Bahru, Malaysia. His academic interests focus on the intersection of machine learning, deep learning, natural language processing (NLP), and computer vision. During his studies, he contributed to several research projects exploring AI-driven solutions to real-world challenges, with a particular emphasis on improving machine perception and language understanding. He has actively participated in technical seminars and hackathons to sharpen his skills. His research aims to apply AI technologies in areas like automation, healthcare, and human-computer interaction. He aspires to contribute to advancements that bridge the gap between computational models and practical applications in various industries.



MOHAMMAD ABU TAREQ RONY obtained his Bachelor of Science degree in Statistics from Noakhali Science and Technology University, located in Noakhali, Bangladesh. Additionally, he possesses expertise in devising advanced analytics strategies using data. His professional experience is diverse and includes 3 years of research in areas such as Artificial Intelligence. He is currently working as a Research Data Scientist at aiQuest Intelligence, Dhaka, Bangladesh. He has published articles in refereed journals and conference proceedings such as IEEE Access, Elsevier Data in Brief, MDPI Children, IEEE, and Springer International Conferences.

Moreover, He actively engages in partnerships with international researchers, recognizing that research plays an indispensable role in fostering innovation. Overall, Rony is a hardworking individual who has taught himself various skills such as Data Analysis, Statistics, ML, and DL.



MD. MEHEDI HASSAN (MEMBER, IEEE) is a dedicated and accomplished researcher, is completed the Master of Science (M.Sc.) degree in computer science and engineering at Khulna University, Khulna, Bangladesh in 2024. He is a member of the prestigious Institute of Electrical and Electronics Engineers (IEEE). Mehedi completed his BSc degree in Computer Science and Engineering from North Western University, Khulna in 2022, where he excelled in his studies and demonstrated a strong aptitude for research. As the founder and CEO of The Virtual BD IT Firm and VRD Research Laboratory, Bangladesh, Mehedi has established himself as a highly respected leader in the fields of biomedical engineering, data science, and expert systems. Mehedi's research interests are broad and include important human diseases, such as oncology, cancer, and hepatitis, as well as human behavior analysis and mental health. He is highly skilled in association rule mining, predictive analysis, machine learning, and data analysis, with a particular focus on the biomedical sciences. As a young researcher, Mehedi has published 60 articles in various international top journals and conferences, which is a remarkable achievement. His work has been well-received by the research community and has significantly contributed to the advancement of knowledge in his field. Overall, Mehedi is a highly motivated and skilled researcher with a strong commitment to improving human health and well-being through cutting-edge scientific research. His accomplishments to date are impressive, and his potential for future contributions to his field is very promising. Additionally, he serves as an academic editor of PLOS One, PLOS Digital Health, Discover Applied Sciences (Springer) and a reviewer for 67 prestigious journals. Md. Mehedi Hassan has filed more than 3 patents out of which 3 are granted to his name.



SHAH MD NAZMUL ALAM is doing his Masters in Computer Science at Arizona State University, Tempe. From 2021 to the present, he was a research assistant with the organization Bytes of Intelligence in New York, USA. His research interests include healthcare, agriculture, and business with machine learning and deep Learning.

...