

Fake News Detection

NAGAVELLI YOGENDER NATH¹, GATTU RAMYA², R. PRASANTH REDDY³, K. MANI RAJU⁴

¹Assistant Professor, Department of Computer Science & Engineering (AI&ML), Sumathi Reddy Institute of Technology for Women, Hyderabad.

²Assistant Professor, Department of Computer Science & Engineering, Vignan Institute of Management and Technology for Women, Hyderabad.

³Assistant Professor, Department of Computer Science & Engineering, Malla Reddy (MR) deemed to be University, Hyderabad.

⁴Assistant Professor, Department of Computer Science & Engineering, Malla Reddy (MR) deemed to be University, Hyderabad.

Abstract- In today's world, social media platforms are important means of information diffusion, and people trust them without questioning their authenticity. Social media is a major factor in propagating fake news. Thus, to mitigate the consequences of fake news, we create an NLP model to differentiate fake and real news. Here machine learning algorithms has been used for enhancing fake news detection performance with NLP. Models trained using max entropy classifier, where news content is scanned for sentences that could indicate the news is fake based on existing NLP libraries. TF-IDF weighting is used to score certain pieces of text, so that detection is fast on any updates or incoming messages due to its fast computation time and high recall rate (low mistake rate). Here the proposed project's purpose is to detect fake and misleading news from social media networks.

Indexed Terms- Fake news, Recurrent Neural Network, TF-IDF, NLP.

I. INTRODUCTION

Information that has been altered to resemble news media content but has a different management structure and purpose is known as fake news. It's constantly blowing up on social media. Online blogs, journals, newspapers, and forums all support the expansion of the online publication sector. It's challenging to identify reliable news sources. The proliferation of false information demands the creation of efficient analysis tools that can disclose the veracity of information. People who regularly utilise social media are significantly impacted by news inaccuracy, both positively and negatively. A calamity must be

detected as soon as possible in order to prevent it [1]. Consequently, scientists are focussing on creating methods and algorithms that can effectively identify false information. The mainstream media's editing techniques and standards to secure the dependability and credibility of content. People who are more engaged in political discussions and stock prices are the ones who are drawn to fake news, which also affects their mental health by causing stress, anxiety, and sadness. Instead of focussing on individual articles, one should focus on the original stories that were first published by the authorised publishers in order to stop the spread of fake news [2].

Social media platforms' ongoing expansion has led to a high degree of information source diversification, which is further bolstered by user-friendly accessibility and publisher affordability. Social media improves the quality of our interpersonal interactions by making it simpler and easier to connect and communicate with others. On the other hand, the standard of interpersonal connections is under danger [3]. People have made social media programming crucial to their lives because of the internet's absorption into daily life. As a result of technological advancements, people can now easily communicate with anyone, anywhere in the world. Therefore, when people use social media excessively and let technology control communication, it has a big effect on interpersonal relationships [4].

Online blogs, journals, newspapers, and forums all support the expansion of the online publication sector. People who regularly utilise social media are significantly impacted by news inaccuracy, both positively and negatively. One of the biggest issues facing the modern, digitally connected world is fake

news. The victims may suffer severely as a result of the dissemination of false information and fake news, which can cause confusion and rumours. People might not know how fake news might impact their life or how to react when it does. Fake news detection is therefore crucial to preventing the spread of false information and fake news. The study's suggested solution seeks to be implemented in actual social media platforms and remove the negative experience of users receiving false information from unreliable sources [5].

II. LITERATURE REVIEW

According to [6], the dissemination of false information is currently harming society. We sometimes think a lot of phoney news is true because of how widely it spreads. We consequently encounter problems and deny ourselves a great deal of positive and practical news. We must endeavour to automatically identify fake news in order to save people's lives from these different issues. The process of detecting fake news is quite difficult. In this research, we describe our deep learning-based method for multi-class fake news identification. Using information from the task organisers, we employed a Long Short Term Memory (LSTM) model for multi-class bogus news identification. On the training data, our top-performing model had an accuracy of 0.98. When it came to identifying fake news, our trained model responded accurately.

According to [7], web-based media is a platform for freely expressing one's opinions and thoughts and has made communication easier than it was in the past. Additionally, this increases the possibility that someone will purposefully release a fake phrase. The problem of people being exposed to false information and unquestioningly accepting it also arises from the reliable access to a variety of online statistics sources. Because of this, it's important for us to comprehend and promote such content through online media. Without knowing the news's source, it is extremely difficult to distinguish between verified facts and fabrications in the modern era of information generated by internet media. In this study, we suggest several methods to confirm whether or not the gathered news is fraudulent. Natural Language Processing (NLP) is the method utilised for this.

Numerous other techniques, such as text categorisation and classification modelling, are also employed, and the outcomes have been analysed. Data was gathered from multiple sources, and different methods, such as SVM and Naive Bayes LSTM, were employed to confirm whether the news was accurate.

In [8], two distinct dimensionality reduction techniques—Principle Component Analysis (PCA) and Chi-Square—are combined with a hybrid neural network architecture that combines the capabilities of CNN and LSTM to address the aforementioned problem. Before sending the feature vectors to the classifier, this paper suggested using dimensionality reduction techniques to lower their dimensionality. A dataset with four different stances—agree, disagree, discuss, and unrelated—was obtained from the Fake News Challenges (FNC) website in order to construct the logic. In order to detect bogus news, the nonlinear characteristics are put into PCA and chi-square, which yields more contextual features. Finding out how a news story feels about its headline is the driving force behind this study. Results are improved by 4 and F1 score with the suggested model. The experimental findings demonstrate that PCA performs better than state-of-the-art techniques and chi-square, with 97.8

Due to the vast volume of information shared on digital media in recent years, fake news has proliferated on the Internet [9]. This has prompted us to investigate which of the three Deep Learning models—Long Short-Term Memory (LSTM), Neural Network with Keras (NN-Keras), Neural Network with TensorFlow (NN-TF), and Naive Bayes (NB) and Support Vector Machine (SVM)—performs the best. Using two distinct English-language news datasets, we looked at five models. Four metrics were used to assess the models' performance: accuracy, precision, recall, and F1-score. The findings obtained demonstrated that deep learning models outperformed typical machine learning models in terms of accuracy. The LSTM model has performed better than every other model that was looked at. Its accuracy was 94.21 percent on average.

In order to draw readers, sway beliefs, and boost internet click revenue, a lot of fake news has been produced online in recent years [10]. However, assessing and determining the reliability of news may

be a difficult and time-consuming task. The creation of this fake information has become a global problem, and its effects are well-known for causing confusion over facts and incorrect decision-making. This is merely due to the fact that the majority of research conducted thus far on the identification of such news, particularly in real time, is not very effective. Therefore, reviewing the key works that have addressed this issue is our goal in this work. The study's findings indicate that two main strategies—linguistic and network—have been proposed. We shall attempt to cite a collection of automatic web and social network false news detecting systems in this post.

The research on fake news identification is reviewed in [11], which divides detection techniques into two categories: machine learning-based and knowledge-based. The two types of machine learning-based techniques discussed in this study are conventional techniques and neural network techniques. For conventional methods, it offers Na⁺ve Bayes and Support Vector Machines (SVMs). For neural network approaches, in addition to recurrent neural networks (RNNs) and convolutional neural networks (CNNs). The provided approaches are also discussed in the study. A method for device mastery, namely surveyed reading, is employed in [12] to obtain false information. In particular, this work employed a database of non-fiction stories to teach the device mastery version, which makes use of the Python module Scikit-learn. We used a bag of words, the term frequency, the inverse document frequency, and the bi diagram frequency to extract records from the database using textual content illustration patterns. Next, we investigated type strategies, specifically the feasible type and the linear department of the name and content material, to see if it changed into a typical/no-click on feed in a fake/real sequence. Our analysis concludes that line segregation functions well using the TF-IDF version within the content material segmentation procedure [13]. When compared to the TF-IDF and the term bag of words, the Bi-gram frequency version provided a significantly lower accuracy in theme separation.

Fake news detection is an important but challenging topic in Natural Language Processing (NLP), as stated in [14]. Social networking sites' rapid rise has greatly expanded the amount of information available while

also increasing the spread of false information. The influence of fake news has increased as a result, occasionally affecting offline society and posing a threat to public safety. Given the large amount of Web content, automatic fake news identification is a useful natural language processing problem that can help all online content creators reduce the time and effort required to identify and stop the spread of fake news. Here, we go over the challenges of identifying fake news and related problems. We thoroughly examine the issue formulations, datasets, and NLP solutions developed for this task, along with their advantages and disadvantages. Based on our results, we outline future research directions, such as more precise, thorough, equitable, and useful detection algorithms. According to [15], the fundamentals of data mining include analysing the data from several perspectives, classifying it, and then summarising it. Data analysis is made possible by data mining software. The fundamentals of data pre-processing, classification, clustering, and the WEKA tool are presented in this study. Weka is a tool for data mining. We outline how to use the WEKA tool for these technologies in this document. It offers the ability to categorise the data using a variety of techniques.

Regular expression (RE), a language for specifying text search strings, has been one of the unsung successes of standardisation in computer science in [16]. This practical language is used in all computer languages, word processors, and text processing tools such as the Unix tools grep or Emacs. With regard to text normalisation, edit distance provides a means of quantifying both of these intuitions regarding string similarity. Formally speaking, the minimum number of editing operations (such as insertion, deletion, and substitution) required to change one string into another is known as the minimum edit distance between two strings.

A global platform for studying computational methods of learning is Machine Learning [17]. The journal publishes articles that present significant findings on a wide range of learning techniques applied to a variety of learning problems. In 2021, machine learning achieved a number of incredible feats, and important research articles led to technological advancements that are utilised by billions of people. This field's research is developing at a rapid rate, helping you stay

up to date. The most significant current scientific research publications are compiled here. They cover topics such as applications research, difficulties and techniques research, and research methodological issues [18]. Papers that make assertions regarding learning issues or approaches back them up with actual research, theoretical analysis, or analogies to psychological phenomena. Applications papers demonstrate how to use learning strategies to address significant application issues. The way machine learning research is carried out is improved by research methodology publications.

III. METHODOLOGY

(1) Train machine learning model is the first of the project's two main components. (2) Model deployment. Data collection is the first step in the pipeline, which is followed by data exploration, data preparation and cleaning, machine learning model training, and model deployment.

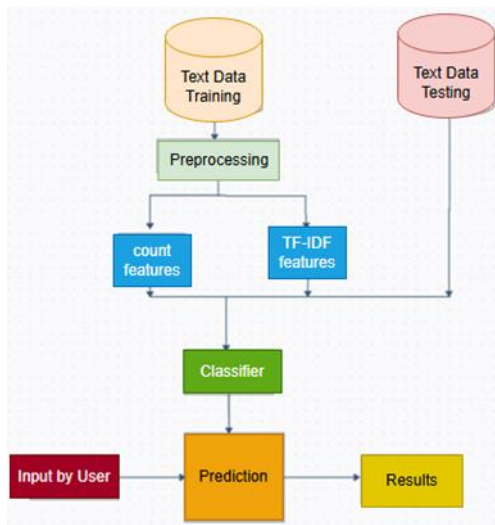


Fig. 1: Flowchart

A. Data collection and cleaning

Numerous sources, such as UCI Machine Learning and Kaggle, are available for gathering news data. During the data cleaning phase, all empty, repetitive, and troublesome rows are eliminated from the dataset in addition to the ID and URL columns. This is due to the fact that empty or repetitive items in any column only make up around 3% of the total rows in the combined dataset. The problematic rows are also eliminated because some of the news titles or content

contain tabs ("t") or newlines ("n"), which cause the text in the (.csv) file to migrate into a new row or column and disrupt the dataset.

B. Data Exploration

Understanding the features and patterns of the data is the goal of data exploration, which comes after data gathering [19]. Data exploration can also lessen the likelihood of having very unbalanced data, which can significantly affect the model that is later trained. After the dataset has been gathered and cleaned, data exploration is done to show the distribution or ratio of fake news to real news, word counts, or even the creation of a word cloud that shows the words that appear most frequently.

C. Data Preparation

The process of preparing and converting data into a machine-understandable context before feeding it into the machine learning model for training is known as data preparation [20]. In natural language processing, regular expressions are an essential tool for defining text search strings because they excel at matching patterns. Among its primary tasks are stemming, sentence segmentation, word segmentation and normalization, and more [21]. Words in a cascade are tokenized using a straightforward regular expression replacement. When the index is constructed, stop words—terms that are commonly employed in sentences—will be eliminated.

Preprocessing news titles and contents using text processing techniques like regular expression to eliminate punctuation and special characters, tokenization to divide the text into words, lemmatization to return the words to their root word, and stop words removal to eliminate common and meaningless words are all part of the project's data preparation [22]. To convert the keywords into machine-understandable numerical vectors, vectorization must be performed after the keywords have been obtained.

The N-gram model, which is a contiguous sequence of n items from a sample of text, is a commonly used technique in language processing. By counting in the corpus and normalizing to a phrase or any other word sequence depending on the words that came before it, it determines the likelihood [23]. "Term frequency in

information retrieval" is what the TF-IDF stands for. This graphic illustrates the importance of a word in a document. It helps account for the fact that some words appear infrequently and increases the proportionate number of times a word appears. Another method for turning words to number vectors is one-hot encoding, since machines can only recognize numbers.

D. Model Training

In this project we use MaxEnt Classifier as the training model.

Maximum entropy (maxent) classifier has been a popular text classifier, by parameterizing the model to achieve maximum categorical entropy, with the constraint that the resulting probability on the training data with the model being equal to the real distribution

E. Model Deployment

This project uses Flask to create a website in order to deploy the model. A daily news column is part of the website's design. News headlines have been used to categorise all of the news. We develop a page to identify authentic or fraudulent news. The user interface that was developed has the ability to gather and disseminate news, regardless of its veracity.

IV. RESULTS AND DISCUSSION

Results for every step of the machine learning process are displayed in this section.

A. Training Model

The maxent classifier is used in this project to train the model. The provided data is used in Trigram vectorizer, TF-IDF weighting, and explicit POS tagging.

1) POS Tagging: Adding a prefix to each word with its type (Noun, Verb, Adjective, . . .). e.g: I went to school =_i PRP-I VBD-went TO-to NN-school Also, after lemmatization, it will be 'VB-go NN-school', This highlights the sentence's semantics and makes it clear what its goal is. This will assist the classifier in distinguishing between various sentence kinds.

```
Model Classifier on non-tagged text
GridSearchCV(cv=ShuffleSplit(n_splits=5, random_state=12345, test_size=0.2, train_size=None),
             estimator=LogisticRegression(),
             param_grid={'C': [0.001, 0.01, 0.1, 1, 10, 100, 1000],
                         'penalty': ['l1', 'l2']})

Cv-scores
Accuracy: nan (s/nan) for params: {'C': 0.001, 'penalty': 'l1'}
Accuracy: 0.885 (s/-0.816) for params: {'C': 0.001, 'penalty': 'l2'}
Accuracy: nan (s/nan) for params: {'C': 0.01, 'penalty': 'l1'}
Accuracy: 0.927 (s/-0.828) for params: {'C': 0.01, 'penalty': 'l2'}
Accuracy: nan (s/nan) for params: {'C': 0.1, 'penalty': 'l1'}
Accuracy: 0.938 (s/-0.827) for params: {'C': 0.1, 'penalty': 'l2'}
Accuracy: nan (s/nan) for params: {'C': 1, 'penalty': 'l1'}
Accuracy: 0.938 (s/-0.828) for params: {'C': 1, 'penalty': 'l2'}
Accuracy: nan (s/nan) for params: {'C': 10, 'penalty': 'l1'}
Accuracy: 0.931 (s/-0.834) for params: {'C': 10, 'penalty': 'l2'}
Accuracy: nan (s/nan) for params: {'C': 100, 'penalty': 'l1'}
Accuracy: 0.931 (s/-0.831) for params: {'C': 100, 'penalty': 'l2'}
Accuracy: nan (s/nan) for params: {'C': 1000, 'penalty': 'l1'}
Accuracy: 0.932 (s/-0.824) for params: {'C': 1000, 'penalty': 'l2'}

Best Estimator Params
LogisticRegression(C=1000)
```

Fig. 2: POS tagging

2) TF-IDF Weighting: We calculate the TFIDF score of each term in a piece of text. The text will be tokenized into sentences and each sentence is then considered a text item. We will also apply those on the cleaned text and the concatenated POS tagged text.

```
Model on original text + TF-IDF
GridSearchCV(cv=ShuffleSplit(n_splits=5, random_state=12345, test_size=0.2, train_size=None),
             estimator=LogisticRegression(),
             param_grid={'C': [0.001, 0.01, 0.1, 1, 10, 100, 1000],
                         'penalty': ['l1', 'l2']})

Cv-scores
Accuracy: nan (s/nan) for params: {'C': 0.001, 'penalty': 'l1'}
Accuracy: 0.944 (s/-0.147) for params: {'C': 0.001, 'penalty': 'l2'}
Accuracy: nan (s/nan) for params: {'C': 0.01, 'penalty': 'l1'}
Accuracy: 0.768 (s/-0.623) for params: {'C': 0.01, 'penalty': 'l2'}
Accuracy: nan (s/nan) for params: {'C': 0.1, 'penalty': 'l1'}
Accuracy: 0.808 (s/-0.633) for params: {'C': 0.1, 'penalty': 'l2'}
Accuracy: nan (s/nan) for params: {'C': 1, 'penalty': 'l1'}
Accuracy: 0.986 (s/-0.084) for params: {'C': 1, 'penalty': 'l2'}
Accuracy: nan (s/nan) for params: {'C': 10, 'penalty': 'l1'}
Accuracy: 0.931 (s/-0.832) for params: {'C': 10, 'penalty': 'l2'}
Accuracy: nan (s/nan) for params: {'C': 100, 'penalty': 'l1'}
Accuracy: 0.939 (s/-0.825) for params: {'C': 100, 'penalty': 'l2'}
Accuracy: nan (s/nan) for params: {'C': 1000, 'penalty': 'l1'}
Accuracy: 0.963 (s/-0.829) for params: {'C': 1000, 'penalty': 'l2'}

Best Estimator Params
LogisticRegression(C=1000)
```

Fig. 3: TF-IDF Weighting

3) Trigram vectorizer: We use the Trigram vectorizer, which vectorizes triplets of words rather than each word separately.

```
Model on original text + TF-IDF
GridSearchCV(cv=ShuffleSplit(n_splits=5, random_state=12345, test_size=0.2, train_size=None),
             estimator=LogisticRegression(),
             param_grid={'C': [0.001, 0.01, 0.1, 1, 10, 100, 1000],
                         'penalty': ['l1', 'l2']})

Cv-scores
Accuracy: nan (s/nan) for params: {'C': 0.001, 'penalty': 'l1'}
Accuracy: 0.944 (s/-0.147) for params: {'C': 0.001, 'penalty': 'l2'}
Accuracy: nan (s/nan) for params: {'C': 0.01, 'penalty': 'l1'}
Accuracy: 0.768 (s/-0.623) for params: {'C': 0.01, 'penalty': 'l2'}
Accuracy: nan (s/nan) for params: {'C': 0.1, 'penalty': 'l1'}
Accuracy: 0.808 (s/-0.633) for params: {'C': 0.1, 'penalty': 'l2'}
Accuracy: nan (s/nan) for params: {'C': 1, 'penalty': 'l1'}
Accuracy: 0.986 (s/-0.084) for params: {'C': 1, 'penalty': 'l2'}
Accuracy: nan (s/nan) for params: {'C': 10, 'penalty': 'l1'}
Accuracy: 0.931 (s/-0.832) for params: {'C': 10, 'penalty': 'l2'}
Accuracy: nan (s/nan) for params: {'C': 100, 'penalty': 'l1'}
Accuracy: 0.939 (s/-0.825) for params: {'C': 100, 'penalty': 'l2'}
Accuracy: nan (s/nan) for params: {'C': 1000, 'penalty': 'l1'}
Accuracy: 0.963 (s/-0.829) for params: {'C': 1000, 'penalty': 'l2'}

Best Estimator Params
LogisticRegression(C=1000)
```

Fig. 4: Trigram vectorizer

A. Model Deployment

There are two phases in the deployment. First phase of the website shows daily news. All the real time news has been shown in the website. The website has been developed by using flask.

We categorize whether the news is authentic or fraudulent in the second stage. Additionally, the user interface allows the user to write content or copy news from other sources. The text processing pipeline is also visible to the user. It displays the preprocessed, post-tagged, and original text. In order for the user to comprehend how the text is processed.



Fig. 5: News Home page

After applying the machine learning algorithms, the news will be predicted as real or fake. The prediction results will be shown after searching the URL link based on the backend classifiers.



Fig. 6: Real News Detection

Fig. 7: Fake News Detection

V. CONCLUSION

In this project, we developed a platform to identify false information on social media. The machine learning model we trained for this project can identify false information with an accuracy of up to 94.3%. Because of their quick computation time and high recall rate, the models have been trained with news titles and content that are appropriate for usage in social media apps where users would respond quickly to any updates or incoming messages. Additionally, we may add latest news to our website by pulling it from the News API. Models can be further enhanced in subsequent research by adjusting the parameters to attain even greater recall and accuracy. To further enhance the models in the future, more research can also be done on the news photos, videos, and text. We have investigated machine learning methods to

identify authentic or fraudulent news. We have defined two components in our project: data parsing and categorisation. We looked into a number of tools and libraries for data parsing, but using Python libraries was the most straightforward approach. We were able to extract the data and save it in a structured format by employing them. The accuracy of this system's news predictions ranges from 80 to 90 percent. Millions of people who use social media will benefit from this. Human reliability is the input. We are unable to obtain the outcome till and unless a URL is provided for the detection. Therefore, automating this system may not work in the future. As a result, the input determines whether the news is predicted to be authentic or fraudulent. There are certain restrictions on our project, such the fact that the data sets are kept on the local system in spreadsheet format. As a result, unless and unless it is shared in some other way, no developer will have access to this collected data. Another drawback is that handling vast amounts of data is challenging. This restriction may be strengthened in the future.

REFERENCES

- [1] Bhaskar Majumdar, Md. Rafiuzzaman Bhuiyan, Md. Arif Hasan, Md. Sanzidul Islam, Sheak Rashed Haider Noori, "Multi Class Fake News Detection using LSTM Approach", International Conference on System Modeling Advancement in Research Trends (SMART) Dec 2021.
- [2] Hunt Allcott and Matthew Gentzkow. 2017. Social media and fake news in the 2016 election. *Journal of economic perspectives* 31, 2 (2017), 211–236.
- [3] Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, et al. 2020. Language models are few-shot learners. *Advances in neural information processing systems* 33 (2020), 1877–1901.
- [4] Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2018. Bert: Pre-training of deep bidirectional transformers for language understanding. *arXiv preprint arXiv:1810.04805* (2018).
- [5] Weizhi Xu, Junfei Wu, Qiang Liu, Shu Wu, and Liang Wang. 2022. Evidence-aware fake news detection with graph neural networks. In *Proceedings of the ACM web conference 2022*. 2501–2510.
- [6] Xuan Zhang and Wei Gao. 2023. Towards llm-based fact verification on news claims with a hierarchical step-by-step prompting method. *arXiv preprint arXiv:2310.00305* (2023).
- [7] Naveen Sai Bommina, Nandipati Sai Akash, Uppu Lokesh, Dr. Hussain Syed, Dr. Syed Umar, "A Hybrid Optimization Framework for Enhancing IoT Security via AI-based Anomaly Detection", *International Journal on Recent and Innovation Trends in Computing and Communication*, (2023) ISSN: 2321-8169 Volume: 11 Issue: 3.
- [8] Mr. Dikshendra Daulat Sarpate, and Dr. B.G Nagaraja, "CONVOLUTION NEURAL NETWORK-BASED SPEECH EMOTION RECOGNITION USING MFCCS", *International Journal of Communication Networks and Information Security*, 2023/12/10
- [9] Habeeb, M. S., & Babu, T. R. (2022). Network intrusion detection system: a survey on artificial intelligence-based techniques. *Expert Systems*, 39(9), e13066
- [10] Uppu Lokesh, Naveen Sai Bommina, Nandipati Sai Akash, Dr. Hussain Syed, Dr. Syed Umar. (2021). Deep Reinforcement Learning with Genetic Algorithm Tuning for Intrusion Detection in IoT Systems. *International Journal of Communication Networks and Information Security (IJCNIS)*, 13(3), 582–595.
- [11] Mr DD Sarpate, "Design of Dual Band Microstrip for satellite Applications", 2nd International Conference on Recent Innovations in Engineering & Technology 2020.
- [12] Uppu Lokesh, Naveen Sai Bommina, Nandipati Sai Akash, Dr. Hussain Syed, Dr. Syed Umar, "Designing Energy-Efficient and Secure IoT Architectures Using Evolutionary Optimization Algorithms", *International Journal of Applied Engineering & Technology*, Vol. 4 No.2, September, 2022.
- [13] Usman, M., Zubair, M., Hussein, H. S., Wajid, M., Farrag, M., Ali, S. J., ... & Habeeb, M. S. (2021). Empirical mode decomposition for analysis and filtering of speech signals. *IEEE Canadian Journal of Electrical and Computer Engineering*, 44(3), 343–349.

- [14] Naveen Sai Bommina , Nandipati Sai Akash, Uppu Lokesh , Dr. Hussain Syed , Dr. Syed Umar, "Multi-Objective Genetic Algorithms for Secure Routing and Data Privacy in IoT Networks", International Journal of Communication Networks and Information Security (IJCNIS), (2020), 12(3), 632–643.
- [15] K Sankar, Divya Rohatgi, S Balakrishna Reddy, "COX Regressive Winsorized Correlated Convolutional Deep Belief Boltzmann Network for Covid-19 Prediction with Big Data", Grenze International Journal of Engineering & Technology (GIJET), Grenze ID: 01.GIJET.9.1.547, © Grenze Scientific Society, 2023.
- [16] Nandipati Sai Akash, Naveen Sai Bommina, Uppu Lokesh, Hussain Syed, Syed Umar, "Optimized Block Chain-Enabled Security Mechanism for IoT Using Ant Colony Optimization", International Journal on Recent and Innovation Trends in Computing and Communication, (2023), 11(10), 1226–1233.
- [17] Singh, A., Gupta, M., Raj, A., Gupta, S. K., & Habeeb, M. S. (2020, December). TWDM-PON: The Enhanced PON for Triple Play Services. In 2020 5th IEEE International Conference on Recent Advances and Innovations in Engineering (ICRAIE) (pp. 1-5). IEEE.
- [18] Naveen Sai Bommina , Nandipati Sai Akash, Uppu Lokesh , Dr. Hussain Syed , Dr. Syed Umar, "Privacy-Preserving Federated Learning for IoT Devices with Secure Model Optimization", International Journal of Communication Networks and Information Security (IJCNIS), (2021), 13(2), 396–405.
- [19] RS Supriya Khaitan, Divya Rohatgi, Sana Nalband, Tejali Mhatre, Shweta Patil, "Enhancing Essay Grading Efficiency and Consistency through Two-Layer LSTM Models and Attention Mechanisms", Journal of Information Systems Engineering and Management 10 (2), 191-202.
- [20] Naveen Sai Bommina, Uppu Lokesh, Nandipati Sai Akash, Dr. Hussain Syed, Dr. Syed Umar, "Optimized AI Models for Real-Time Cyberattack Detection in Smart Homes and Cities", International Journal of Applied Engineering & Technology, Vol. 4 No.1, June, 2022.
- [21] M. Mukhedkar, D. Rohatgi, V.A. Vuyyuru, K.V.S.S. Ramakrishna, Y.A. Baker El-Ebiary, V.A. Asir Daniel, "Feline wolf net: A hybrid lion-grey wolf optimization deep learning model for ovarian cancer detection", Int. J. Adv. Comput. Sci. Appl., 14 (9) (2023)
- [22] Umar, Syed, Bommina Naveen Sai, Nagineni Sai Lasya, Doppalapudi Asutosh, and Lohitha Rani. "Machine Learning based Sentiment Analysis of Product Reviews Using DeepEmbedding." Journal of Optoelectronics Laser 41, no. 6(2022): 108-113.
- [23] R. Gnanakumaran, Divya Rohatgi, A K Sampath, Nidhi Nagar, D. Amuthaguka, Raj Kumar Gupta, "Robust Extreme Learning Machine based Sentiment Analysis and Classification", 2023 5th International Conference on Smart Systems and Inventive Technology (ICSSIT), (2023), DOI: 10.1109/ICSSIT55814.2023.10061017.