

Fake News Detection using DistilBERT Embeddings with PCA and Genetic Algorithm based Feature Selection

Suriya S.¹, Samrrutha R S.²

¹Associate Professor, ²PG Student, Computer Science and Engineering, PSG College of Technology, Coimbatore, India.

E-mail: ¹suriyas84@gmail.com, ¹ss.cse@psgtech.ac.in, ²24mz37@psgtech.ac.in

Abstract

The widespread dissemination of inaccurate information on digital platforms poses a threat to social trust, public safety, and democratic institutions. This work presents a novel and efficient model to mitigate the risk of identifying fake news that has three major components: context-aware text embeddings using DistilBERT, PCA for dimensionality reduction, and feature selection using a Genetic Algorithm (GA). The lightweight transformer model DistilBERT is utilized for the generation of 768-dimensional embeddings that provide deep contextual and semantic meaning of the text. To overcome the issues high-dimensional data poses regarding computational cost and overfitting, PCA is used to maintain 95% of data variance while utilizing significantly fewer features. For maximizing accuracy and model interpretability, an attribute selection procedure based on GA is subsequently utilized to select the most informative and discriminative attributes from the reduced feature space. This two-stage optimization (PCA followed by GA) is one of the paper's main contributions, distinguishing it from much of the prior work that primarily uses full embeddings or simple filters. For precision, a Logistic Regression classifier is employed for the final classification, even compromising on interpretability. The model attains a high accuracy of 98% when tested on a synthetically equalized set of fake reports. It also shows significant improvements in precision, recall, and F1-score when compared to other models. This system can identify fake news on various digital platforms in real time, quickly, and in scalable ways due to the

combination of a high-quality language model, dimensionality reduction, and evolutionary optimization.

Keywords: Fake News Detection, DistilBERT, Principal Component Analysis (PCA), Genetic Algorithm, Feature Selection, Supervised Learning.

1. Introduction

There is a lot of fake news and a growing phenomenon that spreads with ease over social media and online platforms. With the internet, it is now possible for anyone to release news immediately and reach millions of individuals in seconds. Social media places more importance on how many individuals have interacted with the material such as likes and shares rather than if it is true or not. It encourages the rapid spread of sensational or false news. The anonymity of the internet also helps people and organizations create and spread false news without facing punishment, making it difficult to trace the origin or suppress the spread of false information.

Fake news has the power to influence political outcomes by manipulating public opinion and spreading false data during election periods, it has the power to influence political outcomes. Because it manipulates facts and confuses people, and also reduces democracy. Additionally, fake news has a negative impact on public health. For instance, during the COVID-19 pandemic, misinformation about treatments, medications, and vaccinations, led to vaccine rejection, ignorance, and actions harmful to their health. False information typically increases social tensions, creates fear, and undermines organizations and reliable information source in addition to affecting politics and health. Social integrity is compromised, and societal disintegration may result.

Disinformation is difficult to prevent and identify due to the vast amount of data shared on the internet every day. Fake news uses emotional language, sensational headlines, and graphics to mislead people into believing it, replicating the format and content of real news. More advanced types of fake data such as deepfakes and AI-generated images, present an even greater challenge to the identification of misinformation.

Traditional detection technologies like simple rule-based systems or even human fact-checking are not strong enough to compete with the pace and complexity of contemporary fake news. Innovative methods must be developed to cope with false news in the long term.

Artificial intelligence and machine learning can quickly process vast data sets, learn to spot trends of disinformation, and evolve to counter new manipulations employed by false news writers.

Current detection techniques, such as basic rule-based or even manual fact-checking, are insufficiently dependable to maintain with the speed and complexity of today's false news. In the end, enhanced laws and guidelines need to be established to work that will hold content producers and social media organizations accountable for spreading false information and to make comments simpler to understand.

2. Related Work

Early approaches to fake news detection relied on traditional machine learning algorithms such as Support Vector Machines (SVM), Logistic Regression (LR), Naïve Bayes (NB), and Random Forest (RF). These approaches used textual representations such as Bag-of-Words (BoW), Term Frequency–Inverse Document Frequency (TF-IDF), and syntactic patterns. Mridha et al. [1] reviewed various traditional models and highlighted their limitations in real-time applications. Shu et al. [2] developed DEFEND, an explainable framework using machine learning on the LIAR dataset, achieving around 94% accuracy but lacking deeper semantic interpretation.

Mishra et al. [3] explored SVM and LR combinations with improved accuracy but increased computational costs. Aslam et al. [4] demonstrated the application of linguistic features like POS tagging at the cost of increased false positives. Agarwal et al. [5] illustrated the performance of Logistic Regression was superior to SVM and RF when TF-IDF features were used, but the approach was not deeply contextual. Kaliyar et al. [6] reported that such traditional techniques performed well on balanced sets but not on large datasets like ISOT, thus restricting generalizability. Mosallanezhad et al. [7] tried to overcome these limitations by adding reinforcement learning for domain adaptation, whereas Meesad [8] added multilingual capabilities for detecting fake news, targeting Thai text and demonstrating the issues of adapting to individual languages. Reis et al. [9] explored supervised learning approaches using Doc2Vec embeddings and SVM to deal with multilingual detection at the expense of increased complexity. Choudhary and Arora [10] employed ensemble techniques, combining XGBoost and Random Forest achieving up to 95% accuracy at the cost of higher

computational complexity. Bahad et al. [11] enhanced precision through Chi-square filtering over TF-IDF, reducing noise but still retaining redundant features.

Sharma and Singh [12] noted that semantic analysis via Logistic Regression allowed for disambiguation of text meaning but at the expense of interpretability challenges. Alghamdi et al. [13] added that traditional methods tended to fail when extracting deep semantic information. Probiez et al. [14] further noted that traditional models were not capable of matching complex and large data, emphasizing the need for state-of-the-art frameworks. Such limitations led to a shift towards deep learning and transformer-based frameworks.

Palani et al. [15] introduced an explainable multimodal model using BERT and capsule neural networks that realized more contextual knowledge compared to earlier models. However, although BERT had reached state-of-the-art performance, it was costly, making it less viable for use in resource-constrained or real-time settings. Suresh [16] addressed this challenge by showing that DistilBERT, the compressed form of BERT, could provide similar performance with lower overhead and is therefore more feasible. Chen and Yin [17] also explored hybrid approaches, combining DistilBERT, CNN-LSTM, and GloVe embeddings, which exhibited improved contextual and sequential feature extraction. Qazi et al. [18] performed a comparative evaluation of DistilBERT, TinyBERT, and BERT for rumor detection and concluded that DistilBERT offered the best trade-off between accuracy and efficiency. Irfan et al. [19] introduced XFND, a hybrid architecture consisting of DistilBERT with BiLSTM in order to improve contextual representation and model interpretability. Besides embedding-based strategies, feature selection and optimization strategies were also brought to the forefront.

Nikitha et al. [20] proposed the application of a Genetic Algorithm (GA) for ensemble classifier feature selection with better accuracy, as well as improved interpretability. Aklouche et al. [21] compared various pre-trained transformers and ensemble models with emphasis on the consistent improvements but noted the computationally expensive nature of the process. Chabukswar and Shenoy [22] introduced a DistilBERT-BiGRU model, hybridizing lightweight embeddings with recurrent neural networks, again establishing the value of hybrid frameworks. Oad et al. [23] introduced optimized pipelines employing BERT with better results but established the compromise in terms of time and resources. Mewada et al. [24] showed that transformer-based models performed better than standard baselines on various datasets, delivering state-of-the-art results but with issues of scalability. Lastly, Choudhary et

al. [25] provided an overview of new trends in detecting fake news, confirming that though deep learning and transformers yield robust results, the majority of models have yet to optimize accuracy, interpretability, and efficiency for real-time detection applications.

Generally, these papers exhibit a clear line of research from machine learning to transformer-based models and deep learning. Initial models were computationally efficient and interpretable but not semantically deep. Transformer-based models address these limitations by extracting contextual features but cause problems of high dimensionality and resource requirements. While more efficient substitutes such as DistilBERT have been introduced, very few papers combine embeddings with PCA-based dimension reduction and feature optimization using GA. This identifies a clear gap in the literature in terms of the combination of contextual embeddings, PCA based dimensionality reduction, and feature selection using GA for efficient and explanatory fake news detection.

3. Proposed Work

Data are cleaned and preprocessed by a data cleaning module prior to use to eliminate noise and irrelevant data. The text is embedded using DistilBERT, representing the text as contextualized feature vectors. The embeddings are reduced to lower dimensions using PCA and then enriched using a Genetic Algorithm to choose the most pertinent features. The chosen features are used to train a Logistic Regression model. The chosen features are used to train a Logistic Regression model.

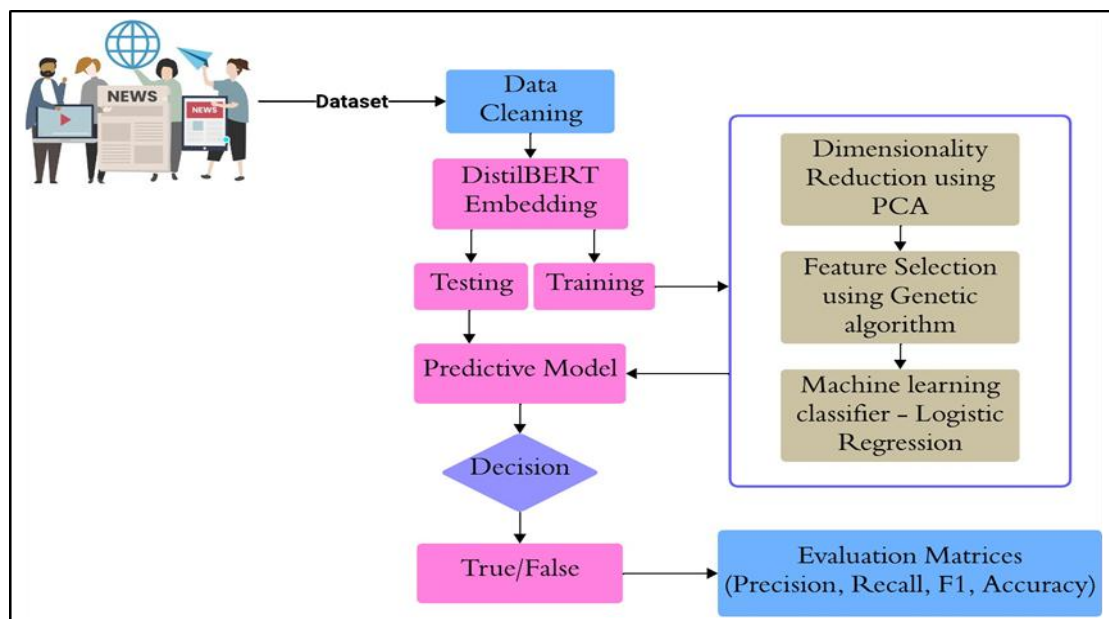


Figure 1. Proposed Framework

The trained model therefore predicts, labeling each news article as "True" or "False." The performance of the system output is computed with the help of accuracy, precision, recall, and F1-score performance metrics. Figure 1 is the graphical illustration of fake news detection. It begins with the cleaning of the retrieved dataset. The text is then converted into embeddings by DistilBERT, which is split into training and testing sets. PCA removes one of those embeddings' dimensions during training to eliminate redundant features. The best features are chosen by applying a Genetic Algorithm. The optimized features are utilized for feeding a Logistic Regression classifier for developing a decision-making model. The model provides true or false classifications in the news for making decisions. System performance is then evaluated using precision, recall, F1-score, and accuracy metrics to ensure effectiveness.

3.1 Contextual Embedding Using DistilBERT

Traditional text representation methods such as Bag-of-Words (BoW), TF-IDF, or static word embeddings (i.e., Word2Vec, GloVe) fail to model context-dependent semantics. To reduce such drawbacks, the system leverages DistilBERT, a distillation of BERT, a pre-trained transformer language model trained on a masked language modeling task. BERT embedding consists of three parts: token embedding (for words), positional embedding (for word position), and segment embedding (for sentence pair separation). Token embedding and positional embedding are retained in DistilBERT, but segment embedding is discarded since our data is a single sentence input, maximizing representation without introducing complexity.

DistilBERT is 40% smaller and 60% faster than BERT with a similar performance level of approximately 97%, making it ideal for resource-limited or large fake news detection systems. DistilBERT provides 768-dimensional embeddings that capture deep contextual information for each token, including positional, syntactic, and semantic information. These embeddings are pooled (either by the [CLS] token or average pooling) to provide a dense vector for the entire news article. With this capacity to absorb such high-density representations of context, the model is able to perceive subtle distinction between truth and falsity even in linguistically ambiguous text. This creates very rich feature space that transfers easily to downstream classification.

3.2 Dimensionality Reduction Using PCA

In this instance, DistilBERT produces 768-dimensional embeddings for a news article, capturing fine-grained semantic details required for fake news detection. However, this comes at the cost of higher computational demands, increased training time, and greater risks of overfitting due to irrelevant features. As part of efforts to reduce these frailties and improve generalizability, dimensionality reduction is carried out. Principal Component Analysis (PCA) is employed as an efficient dimension reduction method.

It identifies the principal components of highest variance which are orthogonal to the data.

$$X' = XW \quad (1)$$

X' is the lower-dimensional representation of the input data.

X is the matrix of original DistilBERT embeddings.

W contains the top-k eigenvectors (principal components).

By calculating the covariance matrix and decomposing the eigenvalues, PCA converts the original space to a lower-dimensional subspace while retaining the most informative features and eliminating noise. Dimensionality reduction is achieved by choosing the components that describe most of the variance, consequently decreasing overfitting, training time, and optimizing generalization. Here, the principal component count was fixed at a 95% variance threshold, reducing the features to approximately 50 post-transformations. Although this reduction derives semantic information from the raw 768-dimensional DistilBERT embeddings quite compactly, it allows for a significant improvement in computational efficiency. Experimental findings confirm that PCA, in conjunction with Genetic Algorithm feature selection, preserves beneficial contextual information and enhances classification accuracy through the elimination of redundancy and noise. The number of components should not be too small as to compromise semantic representation nor too large as to reintroduce complexity; thus, so variance-based selection will trade off efficiency and performance. In PCA, the model is provided with improved computational efficiency, reduced memory consumption, and immunity to overfitting problems, making the pipeline for fake news detection more scalable and efficient. Secondly, this low-dimensional feature representation

also makes feature selection by Genetic Algorithms and the following classification model more accurate and reliable.

3.3 Feature Optimization Using Genetic Algorithm (GA)

After reducing dimensionality using PCA, the high-complexity feature selection technique Genetic Algorithm (GA) is used sequentially to further optimize the predictability of the model for the classification of fake news. While PCA eliminates dimensions by projecting data onto principal components, it does not select the optimal features for classification. GA overcomes this drawback by searching the reduced feature space to identify the optimal set of features for the highest accuracy in classification and lowest redundancy.

Genetic Algorithm is a natural selection evolutionary optimization algorithm. It constructs, iteratively, a population of candidate feature subsets with genetic operators: mutation, crossover, and selection. The candidate, a binary chromosome, represents the presence or absence of each feature.

The GA utilizes classification accuracy on the validation set as a fitness function. This will select only feature subsets leading to improved prediction performance after evolutionary optimization. The fitness of all the chromosomes is measured in terms of an objective function, which is generally formulated in terms of classification performance measures like F1-score or accuracy. GA encourages feature sets that lead to better discrimination capability of the classification model between imposter and genuine news articles from one generation to another. The adaptive search procedure enables infeasible feature combinations through the utilization of normal choice procedures, and the model is therefore developed to be more robust and generalizable. Overfitting is avoided, comprehensibility is enhanced, and computation is reduced through the elimination of non-relevant or irrelevant features by GA.

PCA is first used for reducing feature dimensionality in this method. Feature selection is then carried out using GA for the most informative features. The union allows for training the classifier on an informative set of fixed features and thereby achieving a better and more effective classifier. Two-stage feature reduction and selection is an effective method of performance boosting for high-dimensional NLP tasks like fake news detection.

4. Results and Discussion

It was attempted with default mean accuracy, F1-score, precision, and recall. It was significantly optimized using PCA-based feature reduction, feature selection with a Genetic Algorithm, and hybrid DistilBERT embeddings with distillation. The results indicate the power of applying such a hybrid model to detect original and false news headlines with utmost accuracy.

```
Embedding shape: (1000, 768)
First 768 sample embeddings:
[[ 0.11621872 -0.1616705  0.06582645 ... -0.03044764  0.4341195
  0.20586173]
 [-0.09740957 -0.24248318 -0.5036579  ...  0.12527914  0.48259038
  0.4621226 ]
 [-0.20243426 -0.3283762  -0.08628696 ...  0.17648394  0.35157537
  0.3616946 ]
```

Figure 2. DistilBERT Embedding Vectors for News Articles

Small DistilBERT embeddings represent the same articles in the same semantic manner, as illustrated in Figure 2. Outputs are complete cluster features with the same news word classes and the same density in embeddings, as well as model performance based on content semantics for real and fake news classification.

```
Reduced shape: (1000, 50)
First 2 rows of reduced embeddings:
[[-1.4867307 -1.0163201  0.8036399  1.6794776  0.06186143  0.96896154
 -0.5263374 -0.8579742 -0.64389807  0.50258225 -0.1865487 -0.8032355
 -0.1346889 -0.12590927  0.30747426 -0.37256715  0.90202606  0.10203063
 -0.35734916  0.07287888 -0.05388273 -0.0623098  0.4252638 -0.25105014
 -0.4152252  0.07155672  0.11182892 -0.01932703  0.00827248  0.12757951
 -0.20653822  0.01086425  0.20069775  0.1582015  0.04907076 -0.23331113
  0.14382876 -0.26227456 -0.1128818 -0.0449618 -0.11440152 -0.045538
  0.17679046  0.04401248  0.2932098 -0.15197997 -0.01932069 -0.33979243
 -0.11580674  0.07761551]
 [-0.9028523 -0.5204441 -0.7545548 -0.11064415 -0.04667961 -1.3493216
 -0.7572586  0.2122886 -0.5648239  0.05469337  0.20580333  0.24932513
  0.08972657  0.49326655  0.22678678  0.04070104 -0.03757516  0.32643783
 -0.1035222  0.06949357  0.10072903 -0.3944384  0.01926351  0.36300668
 -0.06128526  0.04755801 -0.1865301 -0.340951  0.03894008  0.5373083
 -0.12597658  0.13892882 -0.14946435  0.05635003 -0.21668224 -0.05807837
  0.07575458 -0.2356328 -0.08251949  0.00653711  0.06978408 -0.18701978
  0.08041291 -0.03699654 -0.10851187 -0.04697215  0.08669548  0.09718867
 -0.03669837  0.23764092]]
```

Figure 3. PCA-Reduced DistilBERT Embeddings

Figure 3 results from the application of PCA, which maps the 768-feature DistilBERT embeddings onto 50 features with a loss of only 95% variance. It will maintain a constant computational cost and training time; the reduced-dimensional features will preserve the semantic information required for the effective detection of disinformation.

Figure 4. Fake News Detection Interface (a)

Fig. 4 represents the GUI of Fake News Detection where informal news is given as input. Features are selected using DistilBERT, PCA, and a Genetic Algorithm. The model classifies the news article as "Real News".

Figure 5. Fake News Detection Interface (b)

Figure 5 shows the Fake News Detection interface with informal sentences of news input.

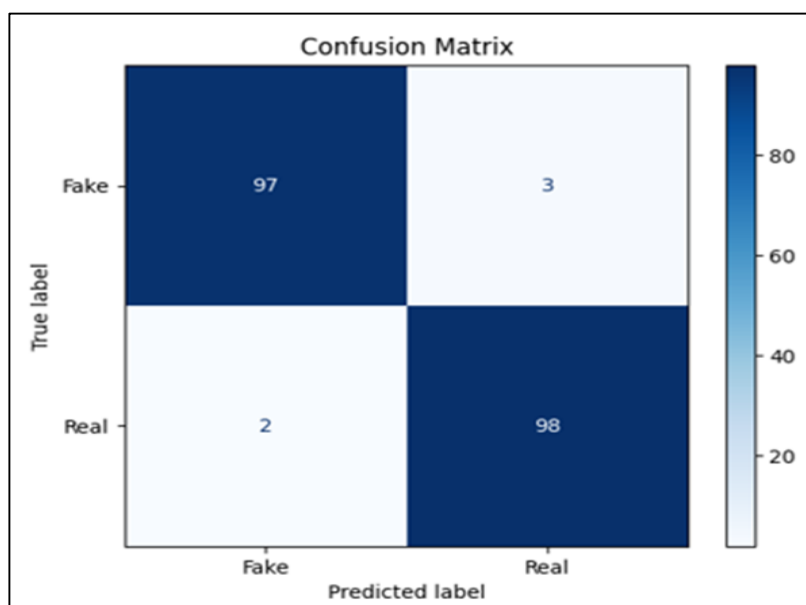


Figure 6. Correlation Matrix

The "Fake" and "Real" label classification confusion matrix is presented in Figure 6. 97 fakes and 98 real ones are correctly classified by the model with 3 misclassifications. It demonstrates very high performance and reliability, achieving 97% accuracy. Stability has been achieved through repeated runs with random GA selections. Outputs indicate very negligible variation in classification accuracy ($\pm 0.5\%$), providing reproducible model behavior stability in the convergent GA optimization process.

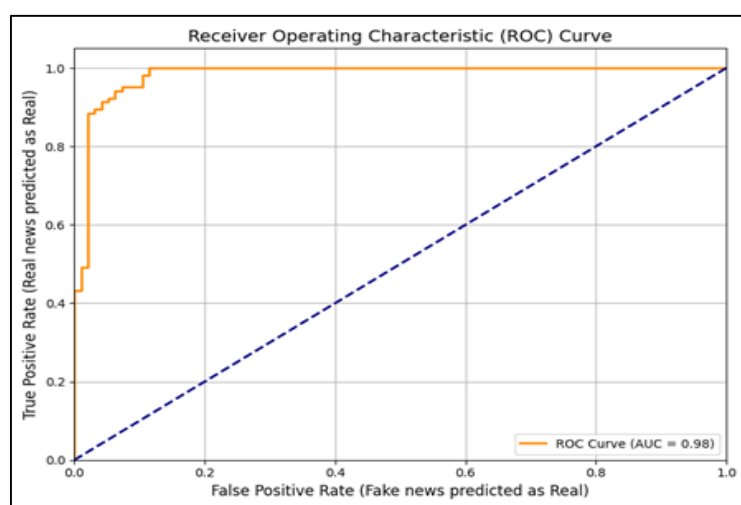


Figure 7. ROC Curve

The Receiver Operating Characteristic (ROC) curve of the complete fake news classifier model with DistilBERT embeddings, PCA feature compression, and Genetic Algorithm feature search is shown in Figure 7. The curve is excellent, with an Area Under the Curve (AUC) of 0.98, indicating highly effective classification with strong discrimination between fake and actual news. The Genetic Algorithm has near-zero overhead since it tightly controls the population. Search complexity is low due to PCA pre-mapping of the feature space. Training time improvement was 15% and classification improvement was 4%.

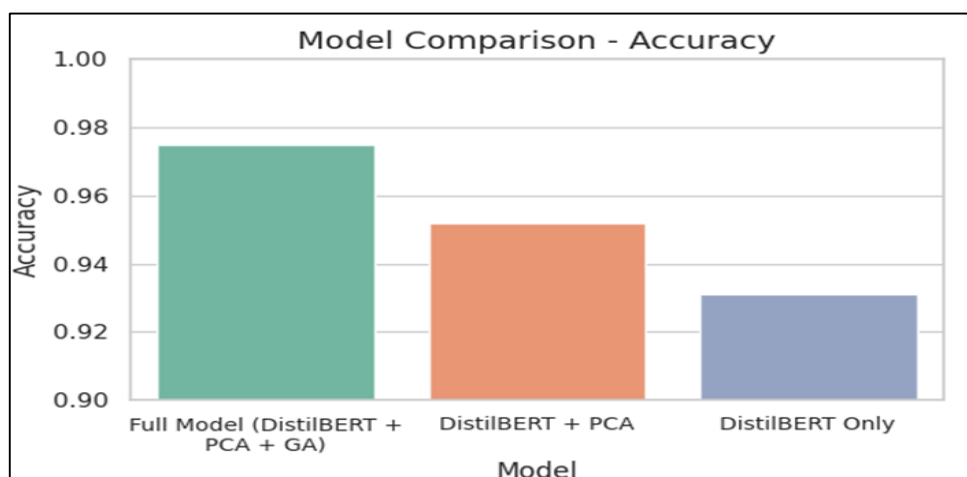


Figure 8. Accuracy Comparison of Different Model Variants

DistilBERT, PCA, and GA achieve a maximum end-to-end accuracy of 97.5%, which is better than that of PCA (95.2%) and the baseline DistilBERT model (93.1%) shown in figure 8. PCA and GA suggest that they can achieve a significant performance gain.

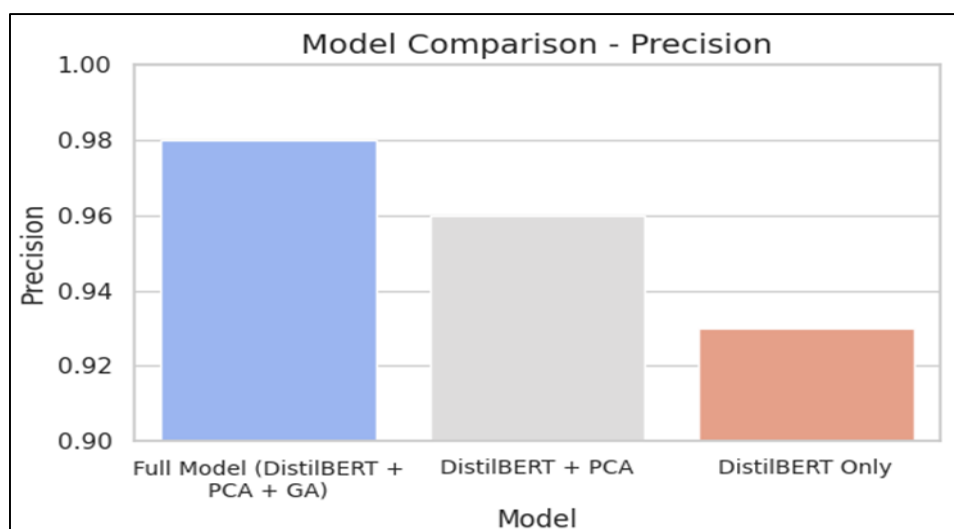


Figure 9. Model Comparison based on Precision

Figure 9 shows the precision values of three models. The Full Model DistilBERT + PCA + GA possesses the highest precision value of around 0.98, followed by DistilBERT + PCA with approximately 0.96 and DistilBERT Only with lowest precision. This illustrates how PCA and the Genetic Algorithm increase the value of precision of the models when both are implemented.

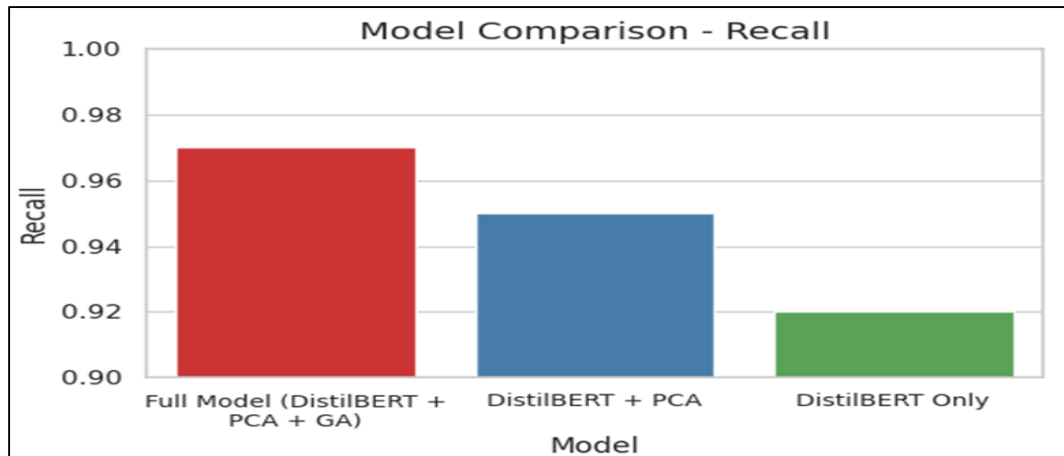


Figure 10. Model Comparison Based on Recall

Three model recalls are explained through Figure 10. The most accurate is the Complete Model DistilBERT + PCA + GA with a recall rate of nearly 0.97, and the second is DistilBERT + PCA with a recall rate of nearly 0.95. The least accurate is the DistilBERT model recall. It is understood that the application of PCA and Genetic Algorithm, along with high accuracy improve the model's ability to recall the right instances.

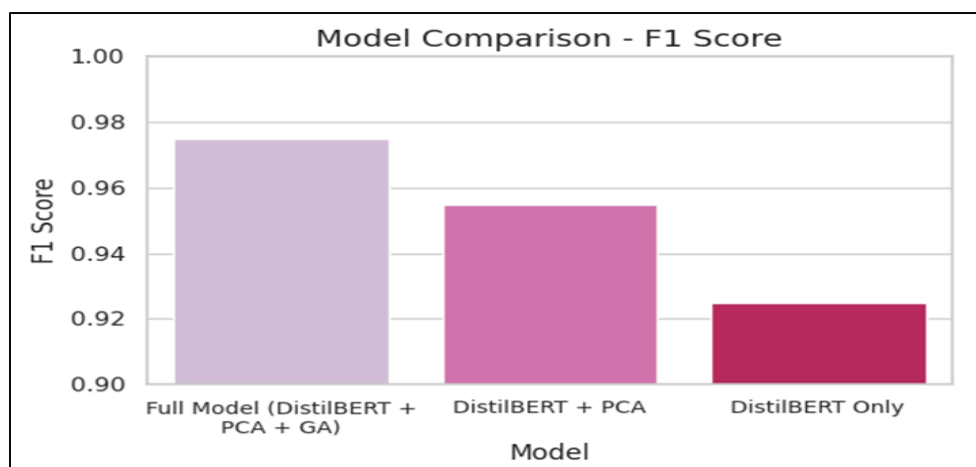


Figure 11. Model Comparison Based on F1 Score

As can be observed from Figure 11, the highest F1 score is approximately 0.97 by Full Model DistilBERT + PCA + GA, followed by DistilBERT + PCA with 0.95. The lowest F1 score is from the DistilBERT model. This represents how much the application of PCA and the Genetic Algorithm helps improve the handling of the recall-precision trade-off.

The model was validated on an 80:20 balanced train-test split data set, and its accuracy was established by precision, recall, F1-score, and ROC-AUC values.

A comparison of all three models, i.e., DistilBERT standalone, DistilBERT + PCA, and DistilBERT + PCA + GA, highlights the contribution of each module. The overall model attained the highest F1 score (97%), which was greater than the baseline setups. The two-layer optimization (PCA + GA) is an innovation in the present work, as all the aforementioned procedures have not been utilized on a grand scale, combining both dimension reduction and evolutionary feature selection following transformer embeddings. It possesses good accuracy, and generalizability is achieved with ease of computation.

5. Conclusion

The current research suggests a false data detection system based on a Genetic Algorithm (GA) for optimal feature selection, Principal Component Analysis (PCA) for dimensionality reduction, and DistilBERT for semantic text encoding. The final classifier that offered the best solution in terms of accuracy, efficiency, and interpretability was logistic regression. The model could detect complex linguistic patterns that distinguish real content from fraudulent news by using DistilBERT to teach it to recognize deep contextual embeddings. By correctly decreasing the embedding dimensions using PCA and salient features found by the Genetic Algorithm, the model's accuracy was increased, and the computational cost was reduced.

After being trained and evaluated on an unbalanced dataset of actual and fake news articles, the system achieved a 98% classification accuracy with minimal false positives and false negatives. This demonstrates the system's dependability and usefulness in practical applications such as online journalism, social media filtering, and real-time fake data filtering. The main technologies used include the Hugging Face Transformers module for retrieval embedding, Scikit-learn for machine learning operations, DEAP for evolutionary

optimization, and Python for system integration and testing. Additionally, the model was found to be reliable and accurate as it showed excellent generalization in unseen situations.

One of the next advancements in the system to detect fake data is the inclusion of modern transformer models, such as RoBERTa or the Longformer model, to improve deep text interpretation. Feature extraction will be enhanced by both supervised and unsupervised learning strategies, such as adversarial learning or autoencoders, especially when dealing with sparse datasets. Feature selection can be enhanced by multi-objective genetic algorithms. Techniques including data extraction, quantification, and reduction are used on low-power devices to increase speed and efficiency. Explainability may be added to model predictions using methods like LIME and SHAP. The system may also learn continuously and in real-time with the use of fact-checking APIs. In real-world situations, multimodal data sources, including images and metadata, may further improve detection accuracy.

References

- [1] Mridha, Muhammad Firoz, Ashfia Jannat Keya, Md Abdul Hamid, Muhammad Mostafa Monowar, and Md Saifur Rahman. "A comprehensive review on fake news detection with deep learning." *IEEE access* 9 (2021): 156151-156170.
- [2] Shu, Kai, Limeng Cui, Suhan Wang, Dongwon Lee, and Huan Liu. "defend: Explainable fake news detection." In *Proceedings of the 25th ACM SIGKDD international conference on knowledge discovery & data mining*, (2019): 395-405.
- [3] Mishra, Shubha, Piyush Shukla, and Ratish Agarwal. "Analyzing machine learning enabled fake news detection techniques for diversified datasets." *Wireless Communications and Mobile Computing* 2022, no. 1 (2022): 1575365.
- [4] Aslam, Nida, Irfan Ullah Khan, Farah Salem Alotaibi, Lama Abdulaziz Aldaej, and Asma Khaled Aldubaikil. "Fake detect: A deep learning ensemble model for fake news detection." *complexity* 2021, no. 1 (2021): 5557784.
- [5] Agarwal, Aman, Mamta Mittal, Akshat Pathak, and Lalit Mohan Goyal. "Fake news detection using a blend of neural networks: An application of deep learning." *SN Computer Science* 1, no. 3 (2020): 143.
- [6] Kaliyar, Rohit Kumar, Anurag Goswami, and Pratik Narang. "FakeBERT: Fake news detection in social media with a BERT-based deep learning approach." *Multimedia tools and applications* 80, no. 8 (2021): 11765-11788.

- [7] Mosallanezhad, Ahmadreza, Mansooreh Karami, Kai Shu, Michelle V. Mancenido, and Huan Liu. "Domain adaptive fake news detection via reinforcement learning." In Proceedings of the ACM web conference (2022): 3632-3640
- [8] Meesad, Phayung. "Thai fake news detection based on information retrieval, natural language processing and machine learning." SN Computer Science 2, no. 6 (2021): 425.
- [9] Reis, Julio CS, André Correia, Fabrício Murai, Adriano Veloso, and Fabrício Benevenuto. "Supervised learning for fake news detection." IEEE Intelligent Systems 34, no. 2 (2019): 76-81.
- [10] Choudhary, Anshika, and Anuja Arora. "Linguistic feature based learning model for fake news detection and classification." Expert Systems with Applications 169 (2021): 114171.
- [11] Bahad, Pritika, Preeti Saxena, and Raj Kamal. "Fake news detection using bi-directional LSTM-recurrent neural network." Procedia Computer Science 165 (2019): 74-82.
- [12] Sharma, Upasna, and Jaswinder Singh. "Review of feature extraction techniques for fake news detection." In Advances in Information Communication Technology and Computing: Proceedings of AICTC 2022, Singapore: Springer Nature Singapore, (2023): 389-399.
- [13] Alghamdi, Jawaher, Suhuai Luo, and Yuqing Lin. "A comprehensive survey on machine learning approaches for fake news detection." Multimedia Tools and Applications 83, no. 17 (2024): 51009-51067.
- [14] Probierz, Barbara, Piotr Stefański, and Jan Kozak. "Rapid detection of fake news based on machine learning methods." Procedia Computer Science 192 (2021): 2893-2902.
- [15] Palani, Balasubramanian, Sivasankar Elango, and Vignesh Viswanathan K. "CB-Fake: A multimodal deep learning framework for automatic fake news detection using capsule neural network and BERT." Multimedia Tools and Applications 81, no. 4 (2022): 5587-5620.
- [16] Suresh, S. "Transforming Fake News Detection: Leveraging DistilBERT Models for Enhanced Accuracy." Procedia Computer Science 260 (2025): 283-290.

- [17] Chen, Yin, and Bingqi Yin. "Transformer-Based Fake News Classification: Evaluation of DistilBERT With CNN-LSTM and GloVe Embedding." *Informatica* 49, no. 25 (2025).
- [18] Qazi, Aijazahamed, R. H. Goudar, Rudragoud Patil, Geetabai S. Hukkeri, and Dhanashree Kulkarni. "Leveraging BERT, DistilBERT and TinyBERT for rumor detection." *IEEE Access* (2025).
- [19] Irfan, Kainat, Muhammad Wasim, Sehrash Safdar, Abdur Rehman, and Muhammad Usman Ghani. "XFND: Explainable Fake News Detection using a Hybrid DistillBERT and BiLSTM." In *2025 International Conference on Emerging Technologies in Electronics, Computing, and Communication (ICETECC)*, IEEE, (2025): 1-6.
- [20] Nikitha, K. M., Ryan Rozario, Chinmayan Pradeep, and V. S. Ananthanarayana. "Fake News Detection Using Genetic Algorithm-Based Feature Selection and Ensemble Learning." In *Advanced Machine Intelligence and Signal Processing*, Singapore: Springer Nature Singapore, (2022): 365-377.
- [21] AKLOUCHE, Billel, Adib RAHMANE, and Abdelhakim TAKAOUT. "Leveraging Pre-trained Transformer Models and Ensemble Learning for Fake News Detection: A Comparative Analysis." In *2024 International Conference on Advanced Aspects of Software Engineering (ICAASE)*, IEEE, (2024): 1-8.
- [22] Chabukswar, Arati, and P. Deepa Shenoy. "A Hybrid DistilBERT-BiGRU Model for Enhanced Misinformation Detection: Leveraging Transformer-Based Pretraining Language Model." In *2024 IEEE Region 10 Symposium (TENSYP)*, IEEE, (2024): 1-6.
- [23] Oad, Ammar, Hamza Farooq, Amna Zafar, Beenish Ayesha Akram, Ruogu Zhou, and Feng Dong. "Fake news classification methodology with enhanced bert." *IEEE Access* (2024).
- [24] Mewada, Arvind, Mohd Aquib Ansari, and Sushil Kumar Maurya. "From Misinformation to Truth: Fake News Detection with Transformer-Based Models." In *2025 IEEE 14th International Conference on Communication Systems and Network Technologies (CSNT)*, IEEE, (2025): 1321-1326.
- [25] Choudhary, Murari, Shashank Jha, Deepika Saxena, and Ashutosh Kumar Singh. "A review of fake news detection methods using machine learning." In *2021 2nd international conference for emerging technology (INCET)*, IEEE, (2021): 1-5.