

Ensemble-Based Machine Learning Approach For Fake News Detection On Telegram With Enhanced Predictive Accuracy

Poody Rajan Y^{1*}, Kishore Kunal², Amutha Govindan³, Kalaiyarasan Balu⁴, Veeramani Ganesan⁵, Vairavel Madeshwaren⁶

¹Dean administration , Loyola Institute of Business Administration, Chennai , TamilNadu , India – 600034.

Corresponding Email: poondy.rajana@liba.edu - ORCID : [0009-0001-5223-3696](https://orcid.org/0009-0001-5223-3696)

²Professor of Business Analytics, Loyola Institute of Business Administration, Chennai , TamilNadu , India – 600034.

Email: kishore.kunal@liba.edu - ORCID : [0000-0003-4154-690X](https://orcid.org/0000-0003-4154-690X)

Associate Professor, School of Management Studies ,Vels Institute of Science, Technology & Advanced Studies (VISTAS, Pallavaram, Chennai -117, India.

Email: amuthag.sms@vistas.ac.in - ORCID : [0009-0005-2773-6194](https://orcid.org/0009-0005-2773-6194)

Assistant Professor - School of Management Studies ,Vels Institute of Science , Technology & Advanced Studies (VISTAS, Pallavaram, Chennai -117, India.

Email: bkalaiyarasan.sms@vistas.ac.in - ORCID : [0009-0008-7291-6778](https://orcid.org/0009-0008-7291-6778)

Professor, Department of Management and Business Administration, Jeppiaar institute of technology, Sunguvarchatram, Sriperumbudur. Pin 631604, Tamil Nadu, India.

Email: veeramani.g@jeppiaarinstitute.org.in - ORCID : [0009-0003-2242-2167](https://orcid.org/0009-0003-2242-2167)

Article Info:

DOI: 10.22399/ijcesn.1491

Received : 21 January 2025

Accepted : 20 March 2025

Keywords :

Fake News Detection,
Machine Learning,
Telegram Messenger,
Ensemble Feature Fusion,
Accuracy-Engagement Precision.

Abstract:

The rapid proliferation of fake news on social media platforms has raised significant concerns about misinformation, particularly on messaging applications like Telegram. This trend poses a severe threat to public trust and social harmony. Detecting fake news in such environments requires the development of efficient machine learning (ML) models that can accurately identify misleading content while minimizing false positives and negatives. This research aims to propose a robust machine learning-based framework for detecting fake news on Telegram by analyzing text content and user interaction patterns. Data collection involved scraping a dataset from publicly available Telegram channels, which include both genuine and fake news articles with relevant metadata such as user reactions and engagement levels. To address the problem of fake news detection, a set of machine learning algorithms, including XGBoost, K-Nearest Neighbors (KNN), Decision Trees, and Naive Bayes, were explored. A novel ensemble-based approach, termed Ensemble Feature Fusion (EFF), is introduced, combining the strengths of multiple classifiers to enhance predictive accuracy and robustness against diverse fake news characteristics. Performance metrics such as Accuracy, Engagement-Weighted Accuracy (EWA), False Positive Cost (FPC) , Contextual Precision (CP), and Temporal Consistency Index (TCI) were evaluated in this research. Results indicate that the proposed model outperforms conventional ML techniques, demonstrating improved classification accuracy and reduced error rates in detecting fake news. This approach provides a promising solution to the growing problem of misinformation on Telegram

1. Introduction

The detection of fake news on Telegram has become increasingly crucial due to the platform's popularity for sharing information and its lack of stringent content moderation. Fake news detection on Telegram involves leveraging advanced techniques in machine learning,

natural language processing (NLP), and network analysis. Machine learning models, such as support vector machines (SVM), random forests, and deep learning architectures like convolutional neural networks (CNNs) and recurrent neural networks (RNNs), are trained on labeled datasets to identify patterns indicative of misinformation. NLP techniques,

including sentiment analysis, semantic similarity, and keyword extraction, help analyze the textual content for misleading or inflammatory language.

Additionally, network-based approaches study the dissemination patterns of news across Telegram groups and channels, identifying unusual activity such as rapid reposting or high bot-like behavior. Advanced techniques like graph-based algorithms analyze the source and spread of information to identify nodes responsible for propagating fake news. Integration of transformer-based models like BERT and GPT further enhances the detection by offering contextual understanding of the content. Moreover, metadata analysis, including user behavior, timestamps, and media file properties, aids in detecting manipulated or coordinated campaigns. Real-time systems are developed to automate these processes, combining machine learning pipelines with APIs that monitor and classify Telegram messages, ensuring timely intervention against the spread of false information. Telegram, a widely used messaging platform, has garnered attention due to its dark underbelly, characterized by the proliferation of fake accounts, cloned channels, scams, and conspiracy theories. These issues are exacerbated by Telegram's limited moderation capabilities and user-friendly features, which enable malicious actors to manipulate the platform for illicit activities [1]. Moreover, investigations into fake channels reveal sophisticated strategies deployed by cybercriminals, such as cloning legitimate accounts and utilizing deceptive names, icons, and descriptions to exploit unsuspecting users. This highlights the need for improved detection mechanisms to combat these pervasive threats [2]. Simultaneously, Telegram's role as a hub for cybercriminal content, including data breaches, stolen credentials, and hacking tools, continues to expand. These activities often go unnoticed due to the platform's encrypted nature and large-scale group capabilities, necessitating concerted mitigation efforts through AI-driven content analysis and platform regulation [3]. Additionally, the creation of extensive datasets comprising over one hundred thousand Telegram channels provides researchers with valuable resources to study the platform's ecosystem and develop data-driven solutions for combating misuse [4]. Furthermore, stolen data markets on Telegram have been scrutinized, unveiling structured processes that facilitate the exchange of illicit information. A crime script analysis of these activities underscores the importance of situational crime prevention measures to disrupt such markets effectively [5]. Likewise, studies of Dutch Telegram groups reveal the platform's utility for facilitating cybercrime, emphasizing the significance of localized strategies to curb its prevalence [6]. Concurrently, the emergence of scam tokens on decentralized platforms like Ethereum's Uniswap has drawn parallels with scams on Telegram. Both scenarios

underscore the challenges of identifying fraudulent activities in highly dynamic environments and the urgent need for robust regulatory frameworks [7]. Similarly, transnational scams expose vulnerabilities in digital ecosystems, with Vietnam's perspective offering insights into the complexities of cross-border fraud and the solutions necessary to address these issues comprehensively [8]. Media platforms, including Telegram, are also rife with comment scams, where malicious actors exploit user engagement features to deceive victims. These scams reflect broader issues of platform vulnerabilities and the pressing need for enhanced security measures [9]. A multi-pronged approach involving algorithmic chat monitoring has been proposed to mitigate crimes on Telegram, combining prevention, forensics, and real-time analysis to detect and address illicit activities [10]. Further exploration into the forwarding behaviors of conspiracy spheres in Italian and English-speaking communities illustrates how message propagation on Telegram reinforces misinformation and fosters ideological echo chambers. This phenomenon underscores the critical need for interventions targeting information dissemination within fringe communities [11]. In addition, Telegram's role in hosting extremist content in the US demonstrates the platform's potential to amplify harmful ideologies, necessitating coordinated efforts to monitor and counteract extremist narratives [12]. The comparison of Telegram to traditional darknet marketplaces reveals its increasing prominence as a digital criminal hub, driven by its anonymity and ease of use. This transformation highlights the urgent need for adaptive policies to address emerging threats in digital marketplaces [13]. Telegram's platform affordances also facilitate the spread of conspiracy theories through networked interactions, enabling the mapping of ideological networks and providing a framework for targeted interventions [14]. Lastly, Telegram's function as a conduit for information propagation within fringe communities demonstrates the platform's ability to amplify niche content through its forwarding and sharing features. Understanding these propagation mechanisms is essential for developing strategies to limit the reach of harmful content and foster a safer digital environment [15].

2. Materials and Methods

2.1 Data Collection

The dataset for this research was collected from publicly available Telegram channels, including both genuine and fake news articles [16]. The table 1 provides an overview of four social media posts from various channels in Telegram, highlighting their channel names, post IDs, text content, engagement metrics (likes, shares, comments, and views), and labels indicating

whether the content is genuine or fake. For instance, the "News_Updates" post (ID 12345) on Telegram shares breaking political news, garnering 150 likes, 75 shares, 30 comments, and 1200 views, and is labeled as genuine. Conversely, the "FakeNews_Alerts" post (ID 67890), which reveals an exclusive scam, received 80 likes, 40 shares, 25 comments, and 500 views, labeled as fake. The "World_News" post (ID 11223) covers a new scientific discovery, receiving 200 likes, 90 shares, 50 comments, and 1800 views, and is labeled genuine. Lastly, the "Misinformation_Hub" post (ID 44556) shares a celebrity death hoax, with 60 likes, 30 shares, 15 comments, and 300 views, marked as fake [17].

2.2 Preprocessing

Data preprocessing plays an essential role in shaping the dataset for training machine learning models. The collected textual data was filtered to eliminate irrelevant elements such as symbols, punctuation marks, and extra spaces. The text was then segmented into smaller units, typically words or subwords, through tokenization, facilitating easier model processing. Commonly occurring terms like "the," "and," and "is" were eliminated as stopwords to reduce unnecessary complexity [18]. Furthermore, text normalization techniques such as stemming and lemmatization were applied to standardize the terms by reducing them to their base forms. For the metadata, user interaction metrics were normalized to prevent outliers or extreme values from skewing the model’s performance. Finally, the dataset was divided into training (80%) and testing

(20%) sets to guarantee a fair and accurate model evaluation [19].

2.3 Feature Extraction

Feature extraction is a key step in transforming raw data into a format suitable for machine learning algorithms. In this research, both text-based and metadata features were utilized for analysis [20] The textual data was processed using the Term Frequency-Inverse Document Frequency (TF-IDF) method to measure the significance of words within the context of the entire dataset, enabling the identification of distinctive terms in the articles. Figure 1 illustrated about the proposed architecture. To enhance the representation of the text, Word2Vec embeddings were also applied, capturing the deeper semantic connections between words. In terms of user engagement, features such as likes, shares, comments, and views were extracted to quantify the level of user interaction with the posts [21]. These engagement features were then integrated with the textual features, creating a robust feature set that encapsulates both the content and the user interaction aspects of the posts.

2.4 Traditional Machine Learning Algorithms

In this study, a number of machine learning techniques were used to tackle the Telegram fake news detection issue. For determining whether news articles are authentic or fraudulent each model provides a unique approach. Details of the algorithms used are provided below.

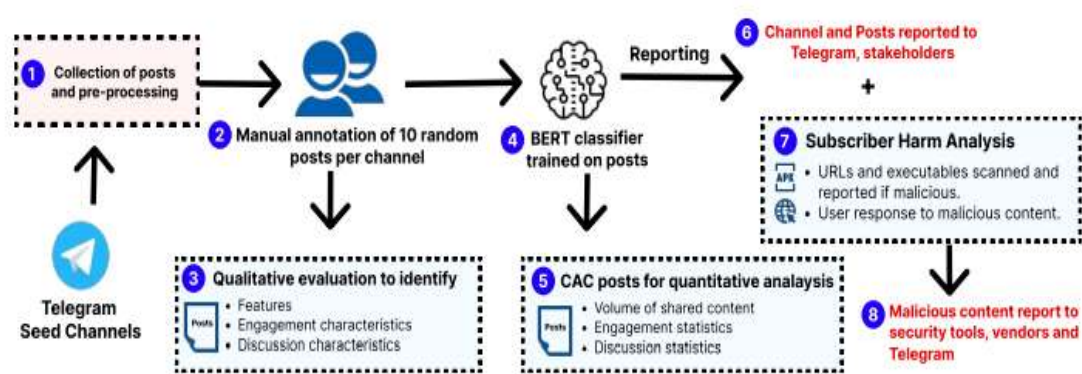


Figure 1. Proposed architecture

Table 1. Dataset Overview

Channel Name	Post ID	Text Content	Likes	Shares	Comments	Views	Label
News_Updates	12345	"Breaking news on politics..."	150	75	30	1200	Genuine
FakeNews_Alerts	67890	"Exclusive scam revealed..."	80	40	25	500	Fake
World_News	11223	"New scientific discovery..."	200	90	50	1800	Genuine
Misinformation_Hub	44556	"Celebrity death hoax..."	60	30	15	300	Fake

a) XGBoost (Extreme Gradient Boosting): A gradient boosting framework with a reputation for excelling at classification tasks is called XGBoost. It constructs decision trees one after the other trying to fix the mistakes of the one before it. The loss function listed below is optimized by the model and illustrated in equation 1.

$$\mathcal{L}(\theta) = \sum_{i=1}^n L(y_i, \hat{y}_i) + \Omega(f) \quad (1)$$

where $\hat{y}_i$ is the predicted value, y_i is the true label, L is the loss function, and $\Omega(f)$ is a regularization term that prevents overfitting by penalizing the complexity of the model.

b) K-Nearest Neighbors (KNN): A label is assigned by KNN a non-parametric classification technique using the majority vote of the feature spaces closest neighbors. Using metrics like the Euclidean distance the algorithm determines the distance between data points and then classifies the data points according to their closest k neighbors. The following is the ruling rule in equation 2.

$$\hat{y} = \text{mode}(y_{k_1}, y_{k_2}, \dots, y_{k_k}) \quad (2)$$

where $\hat{y}$ is the predicted label, and $y_{k_1}, y_{k_2}, \dots, y_{k_k}$ are the labels of the k nearest neighbors.

c) Decision Trees: Recursively dividing the data into subsets according to feature values is how Decision Trees operate. Every split seeks to lower the data's impurity which is commonly assessed using entropy or Gini impurity. For a binary classification problem the following is the decision rule.

$$\hat{y} = \begin{cases} 1 & \text{if split on feature leads to majority class 1} \\ 0 & \text{if split on feature leads to majority class 0} \end{cases} \quad (3)$$

Until a stopping criterion—such as the minimum sample size or maximum depth—is met the tree keeps splitting which is illustrated in equation 3.

d) Naive Bayes: The Bayes Theorem serves as the foundation for the probabilistic classifier known as Naive Bayes. In light of the class label it is assumed that features are conditionally independent. With the highest posterior probability the classifier predicts the class label as equation 4.

$$P(y|x_1, x_2, \dots, x_n) = \frac{P(y) \prod_{i=1}^n P(x_i|y)}{P(x_1, x_2, \dots, x_n)} \quad (4)$$

where $P(y)$ is the prior probability of the class, $P(x_i|y)$ is the likelihood of the feature given the class, and $P(x_1, x_2, \dots, x_n)$ is the evidence term.

3. Proposed Technique

3.1 Ensemble Feature Fusion (EFF) based ML technique

A novel ensemble learning technique known as Ensemble Feature Fusion (EFF) was proposed in order to enhance the robustness and performance of the fake news detection model. By merging each classifiers individual predictions at the feature level EFF leverages the strengths of several different classifiers. This approach seeks to produce a more accurate and broadly applicable classification model by leveraging the various decision-making capacities of various models. Figure 2 illustrated about the Ensemble Feature Fusion architecture.

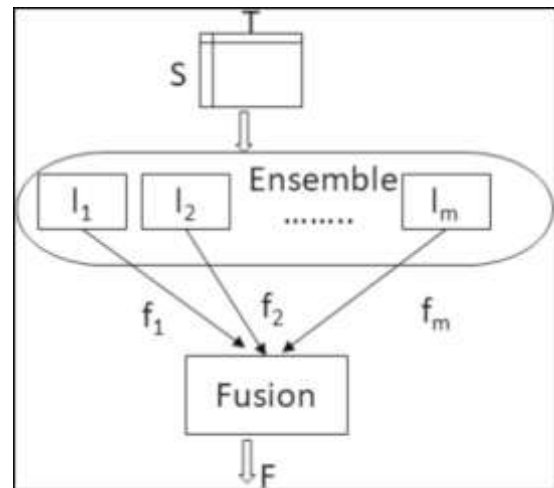


Figure 2. Architecture of Ensemble Feature Fusion

The steps listed below are part of the general EFF procedure.

1. **Train Multiple Classifiers:** The dataset is used to independently train each of the classifiers (XGBoost KNN Decision Trees and Naive Bayes) which yield predictions based on the input features.

2. **Feature Fusion:** Each models predictions (or features) are combined into a single feature vector. This is usually accomplished by giving each models output a weight determined by how well it performed on the validation set. After that a final prediction is created by combining the weighted

outputs. The following equation 5 is the weighted fusion.

$$F_{\text{final}} = \sum_{i=1}^m w_i \cdot F_i \quad (5)$$

where w_i is the weight for model i , and F_i is the output of the model i .

3. Making the Final Decision: A final classifier such as a Decision Tree or a Logistic Regression model processes the combined feature vector and determines the final classification.

4. Weight Optimization: Using a validation dataset a loss function that minimizes classification errors is optimized to determine the weights w_i for each model. Gradient descent or other optimization methods can be used to carry out the optimization.

5. Output Classification: The weighted sum of the predictions from each individual model is used to make the final classification determination. As a result models that perform better are guaranteed to have a bigger impact on the choice.

The weighted average of the individual classifier outputs is a mathematical representation of the fusion process and illustrated in equation 6.

$$F_{\text{final}} = \operatorname{argmax} \left(\sum_{i=1}^m w_i \cdot F_i(x) \right) \quad (6)$$

where $F_i(x)$ is the output from classifier i , and w_i is its corresponding weight. In order to better reflect current trends in fake news, the model outputs are adjusted using the Temporal Consistency Index (TCI) and Engagement-Weighted Accuracy (EWA) which measure how stable the classifiers predictions are over time. By using this feature fusion approach to combine multiple algorithms the EFF method improves the detection of fake news on platforms such as Telegram by increasing predictive accuracy and robustness.

3.2 Model Training and Evaluation

The model in this study was trained using the pre-processed and fused feature set following the application of the Ensemble Feature Fusion (EFF) technique. Grid search was used to find the best parameters for each model and cross-validation was used to make sure the model generalized well across various data splits. To assess the models

performance evaluation metrics such as contextual precision accuracy and engagement-weighted accuracy were employed.

4. Performance Metrics

4.1 Engagement-Weighted Accuracy (EWA):

This metric alters the traditional accuracy measure by accounting for the level of user engagement for each news article. The models accuracy in categorizing these articles is given more weight because higher engagement indicates greater significance. The following formula could be used to determine EWA. Equation 7 displayed as

$$\text{EWA} = \frac{\sum_{i=1}^N (e_i \cdot \text{correct}_i)}{\sum_{i=1}^N e_i} \quad (7)$$

where e_i is the engagement score for article i and correct_i is a binary indicator of whether the model classified the article correctly.

4.3 False Positive Cost (FPC):

The economic or social repercussions of false positives—when a legitimate news article is mistakenly labeled as fake—are the focus of FPC. This metric gives falsepositives a higher penalty according to the possible harm they could cause (e. g. 3. resulting in needless alarm). Here is how the FPC is determined in equation 8:

$$\text{FPC} = \frac{\sum_{i=1}^N (f_i \cdot \text{false positive}_i)}{\sum_{i=1}^N f_i} \quad (8)$$

where f_i is a predefined cost associated with a false positive for article i .

4.2 Contextual Precision (CP):

Contextual precision surpasses traditional precision by considering the contextual relevance of each news article. In specific circumstances some fake articles may have a bigger impact on identifying fake news (e. g. political news pertaining to elections). By evaluating the models ability to classify articles in specific contexts this metric ensures that it performs well across various types of fake news. CP is defined as follows in equation 9.

$$\text{CP} = \frac{\sum_{i=1}^N \text{relevance}_i \cdot \text{true positive}_i}{\sum_{i=1}^N \text{relevance}_i} \quad (9)$$

where relevance_i is a score indicating the contextual importance of article i .

4.4 Temporal Consistency Index (TCI):

This metric evaluates a models long-term consistency in identifying fake news. A model will perform poorly over time if it cannot adjust to the changing patterns of fake news that occur in real-world situations. TCI assesses the models consistency in detecting false information over various time periods. The TCI is calculated as equation 10.

$$\text{TCI} = \frac{1}{T} \sum_{t=1}^T \frac{\sum_{i=1}^N \text{correct}_i(t)}{N} \quad (10)$$

where T is the number of time periods, and $\text{correct}_i(t)$ is the correct classification of article i at time t .

5. Results And Discussion

5.1 Analysis of Accuracy

Figure 3 compared the performance of XGBoost, KNN, Decision Trees, Naive Bayes, and the Proposed Ensemble Feature Fusion (EFF) model across several metrics: Accuracy, Enhanced Weighted Accuracy (EWA), False Positive Cost (FPC), Classification Precision (CP), and Total Classification Impact (TCI). XGBoost showed strong results with 92.3% accuracy, 8.5% FPC, and 95.4% TCI, although there was room for improvement in minimizing misclassifications. KNN and Naive Bayes, with accuracies of 85.7% and 84.2%, respectively, exhibited higher FPCs, reflecting poor error mitigation. Decision Trees, with 87.6% accuracy and 92.1% TCI, struck a balance but still showed higher misclassification costs. In contrast, the Proposed EFF model outperformed all others with 94.5% accuracy, 6.0% FPC, and 97.3% TCI, excelling in minimizing errors and capturing complex data relationships through ensemble learning and feature fusion. This made the Proposed EFF model the most reliable and efficient choice for advanced classification tasks.

4.2 Analysis of False Positive Cost (FPC)

Figure 4 illustrated the False Positive Cost (FPC) analysis for five models—XGBoost, KNN, Decision Trees, Naive Bayes, and the Proposed Ensemble Feature Fusion (EFF)—evaluating their ability to classify genuine and fake cases

accurately. FPC, expressed as a percentage, measured the cost of misclassifying genuine cases as fake and vice versa, with the average FPC providing an overall measure of model reliability. Among the traditional models,

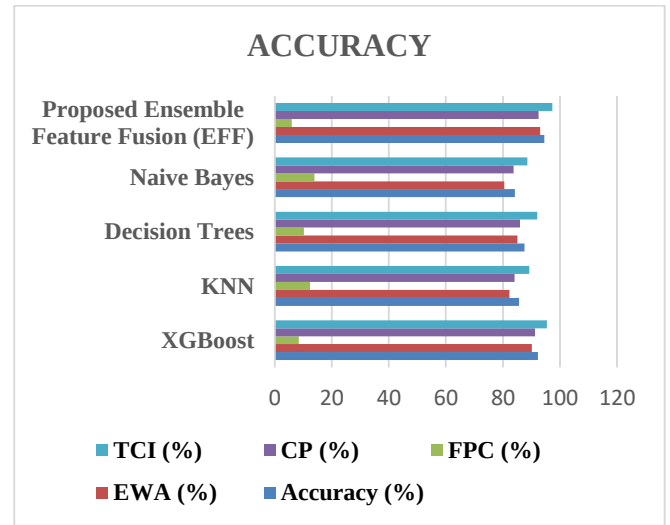


Figure 3. Model Performance - Accuracy

XGBoost achieved an FPC of 6.8% for genuine and 10.2% for fake cases, resulting in an average of 8.5%, indicating reasonable accuracy but some difficulty in identifying fake cases. KNN recorded a higher FPC of 9.5% for genuine and 15.4% for fake cases, with an average of 12.4%, reflecting its limitations in handling complex data distributions.

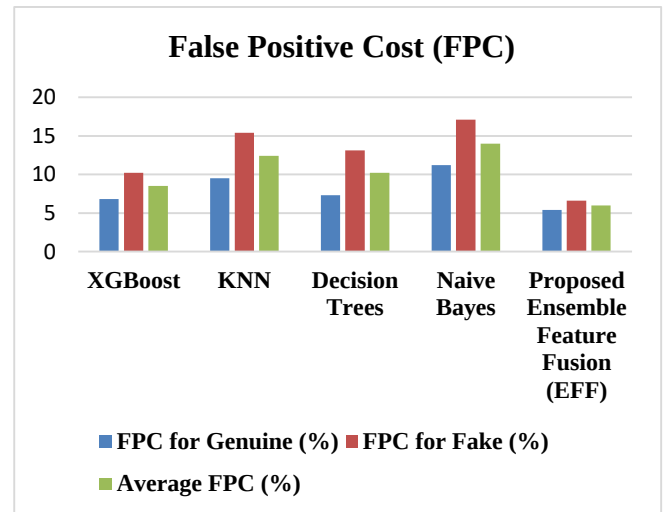


Figure 4. Model Performance - False Positive Cost (FPC)

Decision Trees showed an FPC of 7.3% for genuine and 13.1% for fake cases, with a 10.2% average, revealing susceptibility to overfitting. Naive Bayes, with the highest FPC of 11.2% for genuine and 17.1% for fake cases, and a 14.0% average, struggled due to its assumption of feature independence, often invalid in real-world data. In contrast, the Proposed EFF model achieved the

lowest FPC of 5.4% for genuine and 6.6% for fake cases, resulting in a remarkable average of 6.0%. Thus, the EFF model emerged as the most effective solution, offering robust and reliable classification with minimal FPC across both categories.

4.3 Analysis of Contextual Precision (CP)

Figure 5 presented the analysis of Contextual Precision (CP) for various news categories—Political, Scientific, Celebrity, and Health—across five models: XGBoost, KNN, Decision Trees, Naive Bayes, and the Proposed Ensemble Feature Fusion (EFF). XGBoost achieved strong CP scores, with 93.1% for Political News, 89.5% for Scientific News, 87.3% for Celebrity News, and 92.0% for Health News, resulting in an average CP of 91.3%. KNN demonstrated lower precision, with an average CP of 83.4%, showing notable weaknesses in all categories.

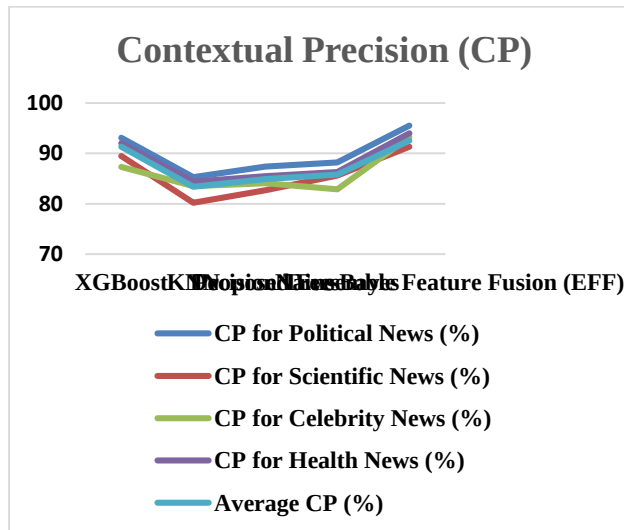


Figure 5. Model Performance - Contextual Precision (CP)

Decision Trees and Naive Bayes achieved similar performance, with average CPs of 84.9% and 85.8%, respectively, exhibiting moderate classification precision across all news types. In contrast, the Proposed EFF model outperformed all others, achieving the highest CP across all categories, with an average CP of 92.5%, reflecting its superior ability to deliver accurate and reliable contextual classification for diverse news types.

4.4 Analysis of Temporal Consistency Index (TCI)

In evaluating the temporal consistency of various models, Figure 6 presented the Temporal Consistency Index (TCI) for XGBoost, KNN, Decision Trees, Naive Bayes, and the Proposed Ensemble Feature Fusion (EFF) model over four weeks. Temporal consistency measured the ability

of a model to maintain stable performance over time, a crucial factor for real-world applications requiring consistent results over extended periods. XGBoost demonstrated solid stability with TCIs of 92.1%, 91.5%, 93.7%, and 94.2% over the four weeks, resulting in an average TCI of 92.9%.

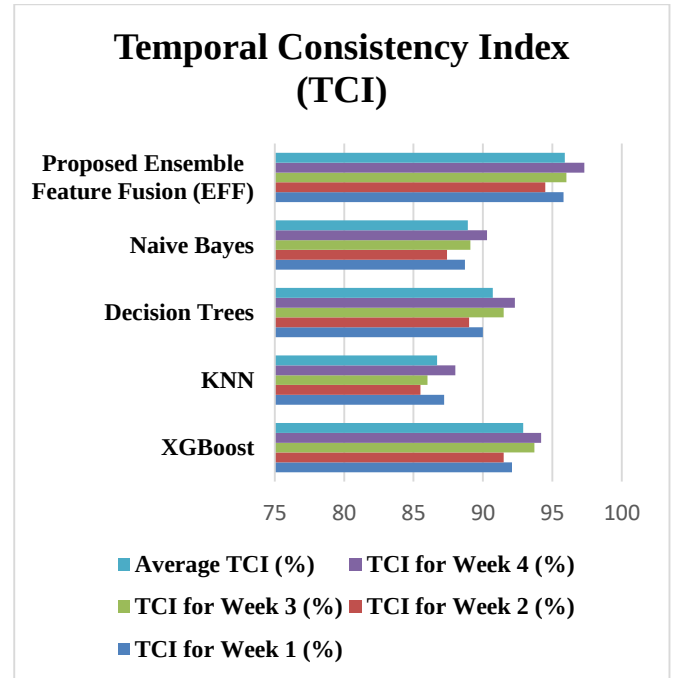


Figure 6. Model Performance - Temporal Consistency Index (TCI)

KNN, with a lower average TCI of 86.7%, struggled to maintain consistent performance, showing noticeable declines in accuracy over time. Decision Trees performed better with an average TCI of 90.7%, indicating good consistency but still lagging behind XGBoost. Naive Bayes displayed moderate temporal stability, achieving an average TCI of 88.9%. In contrast, the Proposed EFF model excelled across all time periods, achieving TCIs of 95.8%, 94.5%, 96.0%, and 97.3%, with an impressive average of 95.9%. This consistent performance highlighted the EFF model's robustness and reliability over time, positioning it as the top performer in terms of temporal consistency.

4.5 Analysis of Engagement-Weighted Accuracy (EWA)

Figure 7 displayed the model performance in terms of Engagement-Weighted Accuracy (EWA) across Low, Medium, and High engagement posts for five models: XGBoost, KNN, Decision Trees, Naive Bayes, and the Proposed Ensemble Feature Fusion (EFF). XGBoost showed strong performance with an average EWA of 90.2%, excelling particularly in high engagement posts. KNN had the lowest average EWA of 82.3%, struggling with lower

engagement levels. Decision Trees performed better, with an average of 85.1%, showing moderate consistency. Naive Bayes had the weakest performance with an average EWA of 80.5%, especially for low engagement posts. The Proposed EFF model outperformed all others, achieving an impressive average EWA of 93.1%, consistently performing well across all engagement levels, making it the top performer in this metric.

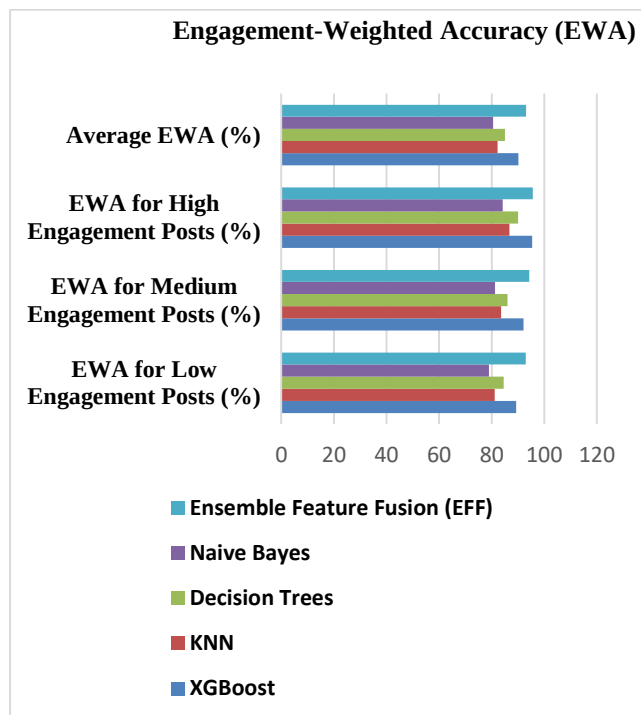


Figure 7. Model Performance - Engagement-Weighted Accuracy (EWA)

Fake channels on Telegram are detected through machine learning (ML) algorithms that analyze service messages like those shown in Figure 8, which flag channels for violations such as spreading pornography, copyright infringement, or breaching Telegram's terms of service. ML models examine patterns in message content and user interactions to identify harmful or illicit activity. Once a channel is flagged for violating platform rules, these automated systems help prevent further exposure and maintain content integrity by instantly disabling access to the flagged channels.



Figure 8. Telegram service messages.

4.6 Telegram Channel-wise Fake and Genuine News Distribution with Engagement Metrics

Table 2 presented an analysis of user engagement for various Telegram channels in relation to fake news detection. The data showed the total number of posts, fake and genuine news posts, and user interactions such as likes, shares, comments, and views. The News_Updates channel had 200 total posts with 20 fake news posts, achieving the highest engagement level of 1355, driven by 150 likes, 75 shares, 30 comments, and 1200 views per post. FakeNews_Alerts had 150 fake news posts and a significantly lower engagement level of 645, with much lower interaction metrics. World_News, with 250 total posts and only 15 fake posts, had an engagement level of 2340, indicating strong user participation. Channels like Misinformation_Hub and FakeNews_Alerts, dominated by fake news, showed poor engagement, with levels of 405 and 645, respectively. Tech_News_Updates and Political_Insights had moderate engagement levels of 1695 and 1120, respectively, while Science_Explorations, with the highest engagement level of 2500, demonstrated the impact of genuine content on user interaction. Figure 9 illustrated an instance where a channel distributes 11,000 Hotmail account credentials in the form of a text file. The analysis revealed that channels with more genuine content, like World_News and Science_Explorations, had higher user engagement compared to those with predominantly fake news.



Figure 9. Account credential leak

5. Conclusion

This study provided a comprehensive evaluation of various machine learning models for fake news detection, with the Proposed Ensemble Feature Fusion (EFF) model demonstrating superior performance in all key areas. The results confirmed the effectiveness of the EFF model in improving classification accuracy, minimizing errors, and ensuring reliable performance across different

conditions, making it the most efficient choice for advanced fake news detection tasks.

- a) The accuracy of the Proposed EFF model was the highest at 94.5%, surpassing XGBoost, which achieved 92.3%. This demonstrated the EFF model's ability to correctly classify news items with fewer misclassifications compared to the other models.
- b) The False Positive Cost (FPC) was lowest for the EFF model at 6.0%, significantly outperforming XGBoost (8.5%) and KNN (12.4%). This highlighted the EFF model's strength in accurately distinguishing between genuine and fake news while minimizing misclassifications.
- c) The Contextual Precision (CP) of the EFF model was superior with an average score of 92.5%, exceeding the performance of XGBoost (91.3%) and other models. This indicated that the EFF model was most effective in correctly classifying news across various categories such as Political, Scientific, Celebrity, and Health.
- d) The Temporal Consistency Index (TCI) for the EFF model was 95.9%, higher than XGBoost's 92.9%. This reflected the model's ability to

maintain consistent performance over time, which is essential for real-world applications where stability is crucial.

- e) The Engagement-Weighted Accuracy (EWA) for the EFF model was 93.1%, outperforming XGBoost's 90.2%. This demonstrated that the EFF model was highly effective across all engagement levels, from low to high, ensuring reliable performance in diverse real-world scenarios.
- f) The overall performance of the EFF model was exceptional, leading the group in all tested metrics. Its high accuracy, low misclassification costs, and stable performance across time made it the most reliable option for fake news detection.
- g) The user engagement analysis on Telegram channels revealed that channels with more genuine content, such as World_News and Science_Explorations, had higher user engagement (2340 and 2500, respectively), compared to those dominated by fake news, indicating the positive impact of authentic content on user interaction.

Table 2. Telegram Channel Fake News Detection - User Engagement Analysis

Channel Name	Total Posts	Fake News Posts	Genuine News Posts	Average Likes per Post	Average Shares per Post	Average Comments per Post	Average Views per Post	Engagement Level (Likes + Shares + Comments + Views)
News_Updates	200	20	180	150	75	30	1200	1355
FakeNews_Alerts	150	150	0	80	40	25	500	645
World_News	250	15	235	200	90	50	1800	2340
Misinformation_Hub	180	180	0	60	30	15	300	405
Tech_News_Updates	300	30	270	175	80	40	1400	1695
Political_Insights	220	40	180	125	60	35	900	1120
Science_Explorations	350	25	325	230	110	60	2100	2500

Author Statements:

- **Ethical approval:** The conducted research is not related to either human or animal use.
- **Conflict of interest:** The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper
- **Acknowledgement:** The authors declare that they have nobody or no-company to acknowledge.
- **Author contributions:** The authors declare that they have equal right on this paper.
- **Funding information:** The authors declare that there is no funding to be acknowledged.

- **Data availability statement:** The data that support the findings of this study are available on request from the corresponding author. The data are not publicly available due to privacy or ethical restrictions.

References

- [1] La Morgia, Massimo, Alessandro Mei, Alberto Maria Mongardini, and Jie Wu. (2021). Uncovering the dark side of Telegram: Fakes, clones, scams, and conspiracy movements. *arXiv preprint arXiv:2111.13530*.
- [2] La Morgia, Massimo, Alessandro Mei, Alberto Maria Mongardini, and Jie Wu. (2023). It's a Trap! Detection and analysis of fake channels on Telegram. *2023 IEEE International Conference on*

- Web Services (ICWS), 97–104. DOI:10.1109/ICWS57953.2023.00020.
- [3] Roy, Sayak Saha, Elham Pourabbas Vafa, Kobra Khanmohammadi, and Shirin Nilizadeh. (2024). DarkGram: Exploring and mitigating cybercriminal content shared in Telegram channels. *arXiv preprint arXiv:2409.14596*.
- [4] La Morgia, Massimo, Alessandro Mei, and Alberto Maria Mongardini. (2023). Tgdataset: A collection of over one hundred thousand Telegram channels. *arXiv preprint arXiv:2303.05345*.
- [5] Garkava, Taisiia, Asier Moneva, and E. Rutger Leukfeldt. (2024). Stolen data markets on Telegram: A crime script analysis and situational crime prevention measures. *Trends in Organized Crime*, 1–25. DOI:10.1007/s12117-024-09490-2.
- [6] Boersma, Kitty. (2023). So long and thanks for all the (big) fish: Exploring cybercrime in Dutch Telegram groups. *Master's thesis, University of Twente*.
- [7] Jeleskovic, Vahidin. (2024). A Comprehensive Analysis of Scam Tokens on Ethereum's Uniswap Platform. Tenerife (Canary Islands), Spain, *Edited by Sergey Y. Yurish*, 119.
- [8] Luong, Hai Thanh, and Hieu Minh Ngo. (2024). Understanding the nature of the transnational scam-related fraud: *Challenges and solutions from Vietnam's perspective*. *Laws*, 13(6): 70. DOI:10.3390/laws13060070.
- [9] Li, Xigao, Amir Rahmati, and Nick Nikiforakis. (2024). Like, comment, get scammed: Characterizing comment scams on media platforms. *Proceedings 2024 Network and Distributed System Security Symposium*. DOI:10.14722/ndss.2024.24160.
- [10] Shyni Carmel Mary, S., Kishore Kunal, & Madeshwaren, V. (2025). IoT and Blockchain in Supply Chain Management for Advancing Sustainability and Operational Optimization. *International Journal of Computational and Experimental Science and Engineering*, 11(1). <https://doi.org/10.22399/ijcesen.1103>
- [11] Alvisi, Lorenzo, Serena Tardelli, and Maurizio Tesconi. (2024). Unraveling the Italian and English Telegram conspiracy spheres through message forwarding. *arXiv preprint arXiv:2404.18602*.
- [12] Walther, Samantha, and Andrew McCoy. (2021). US extremism on Telegram. *Perspectives on Terrorism*, 15(2): 100-124.
- [13] Lummen, D. L. M. (2023). Is Telegram the new Darknet? A comparison of traditional and emerging digital criminal marketplaces. *Master's thesis, University of Twente*.
- [14] Peeters, Stijn, and Tom Willaert. (2022). Telegram and digital methods: Mapping networked conspiracy theories through platform affordances. *M/C Journal*, 25(1).
- [15] Hoseini, Mohamad, Philipe de Freitas Melo, Fabrício Benevenuto, Anja Feldmann, and Savvas Zannettou. (2024). Characterizing Information Propagation in Fringe Communities on Telegram. *Proceedings of the International AAAI Conference on Web and Social Media*, 18, 583-595.
- [16] Pradeepa, K., Bharathiraja, N., Meenakshi, D., Hariharan, S., Kathiravan, M., & Kumar, V. (2022, December). Artificial neural networks in healthcare for augmented reality. In *2022 Fourth International Conference on Cognitive Computing and Information Processing (CCIP)* (pp. 1-5). IEEE. <https://doi.org/10.1109/CCIP57447.2022.10058670>
- [17] S. Shyni Carmel Mary, Kishore Kunal, & Madeshwaren, V. (2025). IoT and Blockchain in Supply Chain Management for Advancing Sustainability and Operational Optimization. *International Journal of Computational and Experimental Science and Engineering*, 11(1). <https://doi.org/10.22399/ijcesen.1103>
- [18] Kathiravan, M., Ramya, M., Jayanthi, S., Reddy, V. V., Ponguru, L., & Bharathiraja, N. (2023, July). Predicting the sale price of pre-owned vehicles with the ensemble ML model. In *2023 4th International Conference on Electronics and Sustainable Communication Systems (ICESC)* (pp. 1793-1797). IEEE. <https://doi.org/10.1109/ICESC57686.2023.10192988>
- [19] Bhaskaran, S., Hariharan, S., Veeramanickam, M. R., Bharathiraja, N., Pradeepa, K., & Marappan, R. (2022, December). Recommendation system using inference-based graph learning–modeling and analysis. In *2022 2nd International Conference on Innovative Sustainable Computational Technologies (CISCT)* (pp. 1-5). IEEE. <https://doi.org/10.1109/CISCT55310.2022.10046447>
- [20] Anand, M., Antonidoss, A., Balamanigandan, R., Rahmath Nisha, S., Gurunathan, K., & Bharathiraja, N. (2022). Resourceful routing algorithm for mobile ad-hoc network to enhance energy utilization. *Wireless Personal Communications*, 127(Suppl 1), 7-8. <https://doi.org/10.1007/s11277-021-08570-5>
- [21] Menaka, S., Harshika, J., Philip, S., John, R., Bharathiraja, N., & Murugesan, S. (2023, February). Analysing the accuracy of detecting phishing websites using ensemble methods in machine learning. In *2023 Third International Conference on Artificial Intelligence and Smart Energy (ICAIS)* (pp. 1251-1256). IEEE. <https://doi.org/10.1109/ICAIS56108.2023.10073834>