

A Voting Based Ensemble Approach for Fake Account Detection using Multimodal Social Media Data

Krotha Karunya¹; Lingireddy Anuradha²; Katta Ravindra³; Sri. P. Rama Krishna⁴

^{1,2,3}B. Tech Student, Dept. of Computer Science and Engineering, R.V.R. & J.C College of Engineering, Chowdavaram, Guntur, Andhra Pradesh, India

⁴Assistant Professor, Dept. of Computer Science and Engineering, R.V.R. & J.C College of Engineering, Chowdavaram, Guntur, Andhra Pradesh, India

Publication Date: 2025/05/20

Abstract: The proliferation of online networks has contributed to a growing concern regarding the surge of fraudulent user profiles, which undermine online security and digital credibility. This study introduces a robust framework for identifying fake accounts by leveraging multimodal features derived from both textual content and numerical metadata. Initially, three deep learning architectures, Artificial Neural Network (ANN), Convolutional Neural Network (CNN), and Long Short-Term Memory (LSTM) networks, were developed and assessed for their classification capabilities. To improve detection performance, a Voting Classifier was employed, integrating XGBoost, Random Forest, and Gaussian Naïve Bayes algorithms. The comparative results indicated that the ensemble model achieved superior performance across key evaluation metrics, including accuracy, precision, recall, and F1-score. By harnessing the complementary strengths of multiple models, the proposed method delivers a dependable solution for identifying deceptive accounts. This research contributes to enhancing the effectiveness of automated fake account detection and encourages further exploration of hybrid models using multimodal inputs.

How to Cite: Krotha Karunya; Lingireddy Anuradha; Katta Ravindra; Sri. P. Rama Krishna (2025). A Voting Based Ensemble Approach for Fake Account Detection using Multimodal Social Media Data. *International Journal of Innovative Science and Research Technology*, 10(4), 4184-4193. <https://doi.org/10.38124/ijisrt/25Apr2345>

I. INTRODUCTION

The rapid expansion of digital platforms has transformed the way individuals interact, share information, and conduct business. Among these platforms, online social networks have seen a tremendous surge in user activity. However, this growth has been accompanied by a significant rise in fake accounts and malicious profiles created to spread misinformation, conduct scams, or manipulate public perception. These deceptive entities not only undermine the credibility of online platforms but also pose serious threats to user security and the integrity of digital spaces.

Traditional detection techniques often rely on single-modal data or shallow learning methods, which limit their ability to effectively identify sophisticated fake profiles. To address these limitations, this study proposes a robust fake account detection framework that leverages multimodal data, combining textual features with numerical metadata. Initially, deep learning models such as Artificial Neural Network (ANN), Convolutional Neural Network (CNN), and Long Short-Term Memory (LSTM) were employed to capture complex patterns in the data. To further enhance detection accuracy and reliability, an ensemble learning approach using

a Voting Classifier was implemented, integrating XGBoost, Random Forest, and Gaussian Naïve Bayes algorithms.

This hybrid approach capitalizes on the strengths of multiple models, improving the generalization and robustness of the detection system. Through comprehensive evaluation, the proposed model demonstrates superior performance over individual classifiers, offering a scalable and effective solution for mitigating the impact of fake accounts on online platforms.

II. RELATED WORK

Identifying fake user accounts on online social platforms has become a pressing issue due to their negative impact on individuals, organizations, and the overall digital space. These accounts are often involved in activities like spreading misleading content, influencing user behavior, and carrying out scams, all of which can erode trust in online systems. To tackle this, researchers have proposed various methods that analyze user profile information, behavioral patterns, and content-related signals. However, many of these solutions are limited by their focus on just one type of data, which restricts their ability to fully understand user behavior. Since data on these platforms often includes a mix of text, images, and

interactions, there is a growing need for more integrated approaches. As a result, recent work has moved toward using multimodal strategies that combine different types of information, leading to more accurate and reliable detection of fake accounts.

➤ Detection Techniques

Numerous computational approaches have been developed to identify inauthentic user accounts on social media platforms, leveraging diverse features such as user behavior, content patterns, and network structures. Among the earliest methods is metadata-based analysis, which involves analyzing the account's creation time, location, and usage patterns. Fake accounts often exhibit unusual or inconsistent metadata, making them easier to detect through statistical analysis. Content-based analysis focuses on evaluating the textual, visual, and multimedia elements shared by a user. Inauthentic accounts typically exhibit patterns such as duplicated or low-quality content, unusually frequent posting activity, and the use of sensational or misleading headlines designed to attract attention.

In contrast, network-based analysis examines the structure of an account's connections and interactions within the social platform, identifying suspicious patterns such as

isolated connections, artificially dense link clusters, or coordinated behavior among multiple accounts. Fake accounts tend to have few genuine followers, and their connections may consist mostly of bots or fake accounts. In recent years, machine learning techniques have been widely applied to detect fake accounts. These models learn patterns based on metadata, content, and network features, enabling them to identify fake accounts more effectively than manual approaches. Hybrid approaches combine different techniques, such as metadata and content analysis, or use machine learning to enhance the effectiveness of other methods. These hybrid systems offer a more comprehensive solution, combining the strengths of each approach to improve detection accuracy.

Effective detection methods for fake accounts often involve the integration of metadata, content, and network-based features with machine learning models, which enables the handling of large and complex datasets. These techniques play a critical role in identifying fraudulent accounts, limiting the spread of false information, and minimizing the negative impact of harmful accounts on social media platforms. By combining different feature types, detection models are able to more accurately capture intricate user behavior patterns, improving both their robustness and precision.

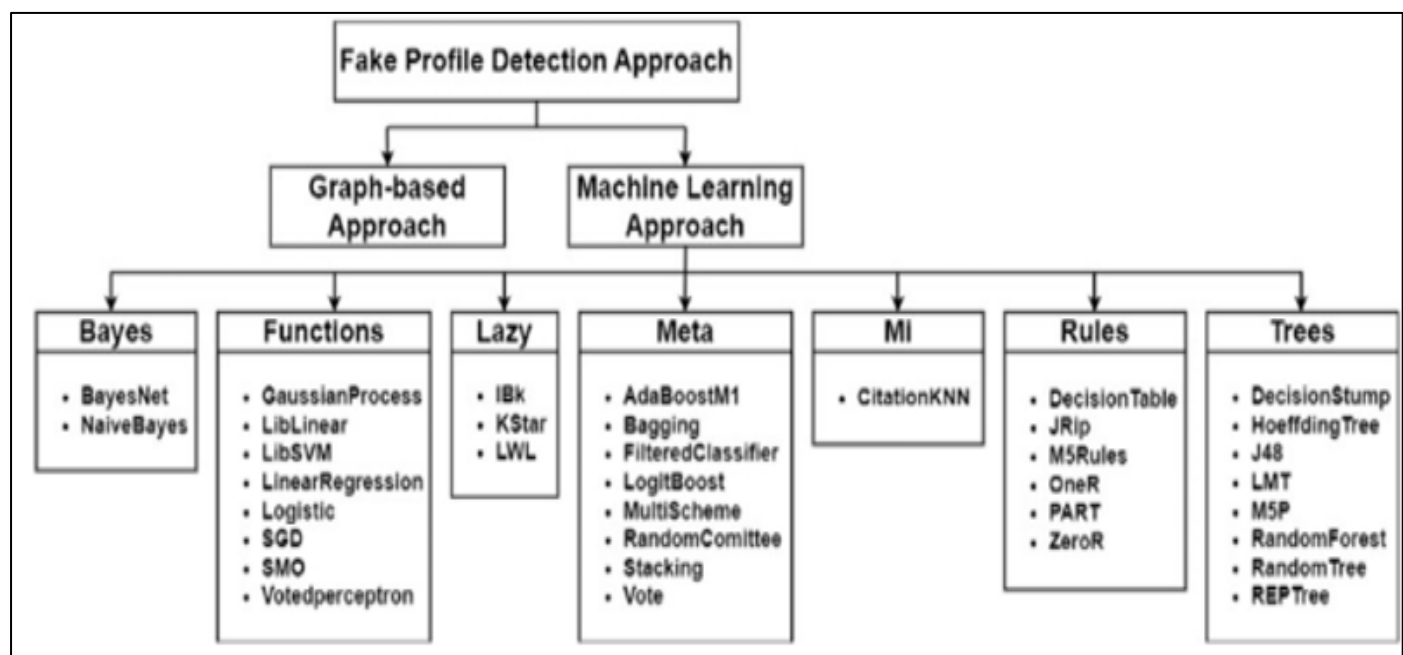


Fig 1 Fake Profile Detection Approaches

➤ Review of Existing Research Contributions

Many researchers have contributed to the development of fake account detection techniques, using a variety of methods and datasets. For example, Zhou et al. [6] developed a method for detecting fake profiles within recommender systems by employing the Time Stamps Target Item Analysis (TS-TIA) algorithm. This approach utilizes quantitative metrics and temporal attack ratings to detect profiles that exhibit unusual behavior. Zeng et al. [7] developed the MUIUI framework, which combines user information and machine learning classifiers to identify fraudulent users on Twitter and Facebook. Their system leverages a variety of user data

crawlers and entity linking, achieving an impressive F1-score of 86.46%. Similarly, Mohammadrezaei et al. [8] employed a friend similarity-based approach to detect fake accounts by computing the similarity between an account's friends using a connection matrix, and then applying eigenvalue decomposition for feature extraction.

Moreover, Yang et al. [9] proposed an innovative method to identify automated spam accounts by analyzing profile and temporal data. Their approach, based on the BERT language model and BiGRU algorithm, automatically builds features, keeping up with evolving social bot behaviors. Siino et al. [10]

performed a comparative analysis of cutting-edge models emphasizing the potential of Transformer-based models for detecting fake accounts. Transformers, such as BERT, have demonstrated strong effectiveness in text categorization tasks, making them an optimal choice for fake account detection.

Other notable contributions include Bertini et al. [11], who proposed using image watermarking techniques to detect fake profiles based on photo forensics. Their approach successfully handled issues like profile linkage and fake profile identification. Additionally, Egele et al. [12] introduced the COMPA framework, which uses statistical models and anomaly detection to identify compromised accounts, and Wu et al. [13], who focused on protecting user privacy through client-side architectures that simulate fake profiles.

➤ Deep Learning Techniques for Identifying Fake Accounts

In recent years, neural network models have demonstrated significant potential in identifying fraudulent social media accounts. Approaches such as Artificial Neural Networks (ANN), Convolutional Neural Networks (CNN), and Long Short-Term Memory (LSTM) have proven to be especially effective. Networks have been explored across various domains, including text classification, sequence prediction, and image processing. These models have been used to identify patterns and features in social media data, aiding in the detection of fraudulent accounts. Below is an overview of how each model has been applied to fake account detection tasks:

Artificial Neural Networks: ANNs have been used in detecting fake accounts because of their capability to model intricate relationships between input features. Zhao et al. [14]

employed an ANN-based approach to analyze user metadata and content features for detecting fraudulent behavior on social media platforms. Their model was able to identify patterns in user activity, such as abnormal post frequency and inconsistencies in metadata, which are typical indicators of fake accounts. While ANNs perform well with structured data, their performance can degrade when handling unstructured or For ordered data like text or time-related data, patterns and relationships can be easily recognized.

Convolutional Neural Networks: CNNs are highly effective at processing grid-based data, such as images and text. Li et al. [15] used CNNs for content-based fake account detection by analyzing textual data in posts and comments. The CNN model was effective in identifying spammers by recognizing patterns such as repetitive phrases, clickbait headlines, and unusual word combinations. While CNNs are strong at recognizing spatial relationships in images and text, they may not always capture long-range dependencies or contextual information as effectively as other models.

Long Short-Term Memory Networks : LSTMs, a form of Recurrent Neural Network (RNN), are specifically built to capture long-term dependencies in sequential data, making them well-suited for analyzing time-series data or text with temporal patterns. Yang et al. [16] proposed a Model based on LSTM to predict fake accounts by analyzing the temporal evolution of user activity. Their results showed that LSTMs outperformed traditional machine learning techniques, as they were able to capture the temporal dependencies of user actions, which are often indicative of fake accounts. However, LSTMs require large amounts of data and significant processing requirements, which may present a challenge in real-world use cases.

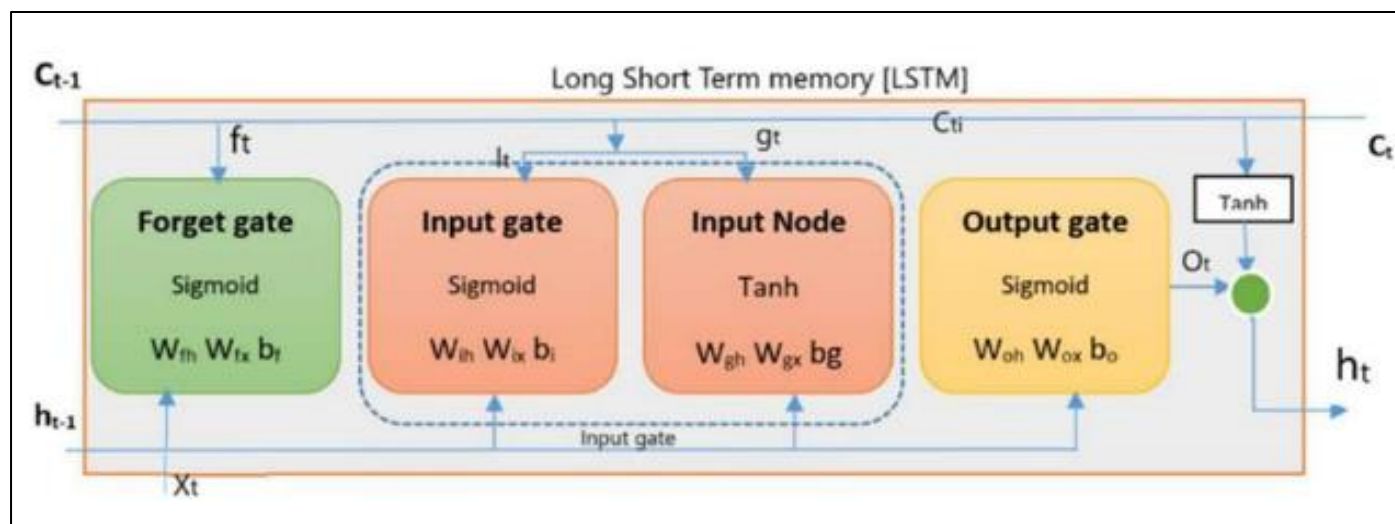


Fig 2 LSTM Architecture

The Long Short-Term Memory (LSTM) network was utilized for the task of detecting fake social media accounts by leveraging multimodal data. LSTM, a variant of Recurrent Neural Networks (RNNs), is well-suited for handling sequential data, as it can learn long-term dependencies and retain significant information over extended periods. For fake account detection, the LSTM model was tasked with analyzing

the sequential interactions between numerical data, such as user activity, and textual data from account descriptions. The input data was preprocessed, involving the normalization of numerical features and the conversion of text descriptions into vector representations using a pre-trained Term Frequency-Inverse Document Frequency (TF-IDF) model. The processed data was then reshaped into the required 3D format for LSTM

processing. Model performance was assessed through standard metrics like accuracy, precision, recall, and F1-score. Moreover, the LSTM model was part of a larger ensemble system, where predictions from the LSTM, Convolutional Neural Networks (CNN), and Artificial Neural Networks

(ANN) were combined using majority voting. This ensemble approach enhanced the overall detection accuracy, with the LSTM model making a substantial contribution to identifying fake accounts and improving the overall system's reliability.

Algorithm 1 To Compute Batch Normalization [23]

Input: Values of x over a mini-batch: $B = \{x_1, x_2, \dots, x_m\}$

Parameters to be learned: γ, β

Output:

$$\mu_\beta \leftarrow 1/m \sum_{i=1}^m x_i$$

$$\sigma^2 \beta \leftarrow 1/m \sum_{i=1}^m (x_i - \mu_\beta)^2$$

$$\hat{x}_i \leftarrow (x_i - \mu_\beta) / \sqrt{\sigma^2 \beta + \epsilon}$$

$$y_i \leftarrow \gamma \hat{x}_i + \beta = \text{BN}_{\gamma, \beta}(x_i)$$

$$\{y_i = \text{BN}_{\gamma, \beta}(x_i)\}$$

// mini-batch Mean

// mini-batch variance

// normalize

// scale and shift

Batch Normalization is a technique employed to optimize the performance and stability of deep learning models. It addresses the issue of internal covariate shift, which arises when the distribution of inputs to a layer changes as the parameters of previous layers are updated during training. By normalizing the inputs at each layer, Batch Normalization stabilizes the learning process, enabling more efficient training.

The technique involves two primary operations: normalization and scaling. Initially, the inputs to each layer are normalized by calculating their mean and variance over a mini-batch of data, ensuring that the inputs have a mean of zero and a standard deviation of one. This process reduces the sensitivity of the training to the initial weights, allowing for more consistent updates. Following normalization, the data is scaled and shifted using two trainable parameters, gamma (scale) and beta (shift), which allow the model to retain the expressiveness of the original data.

Batch Normalization offers several benefits in deep literacy. It improves the flux of slants through the network, mollifying issues analogous as sinking and exploding slants, which can hinder training. Also, it enables the use of advanced knowledge rates, further accelerating training. The normalization process also introduces a slight regularization effect, which can reduce the need for other regularization ways like powerhouse.

This research applied Batch Normalization to various layers in the models, including convolutional layers in the CNN and fully connected layers in the ANN and LSTM models. This application enhanced training speed, improved model convergence, and contributed to more accurate predictions in the finding of fake accounts on social media.

➤ Combining ANN, CNN, and LSTM in Hybrid Models

While each deep learning model ANN, CNN, and LSTM, has shown promise individually, recent research has explored hybrid models that integrate the benefits of these network structures. For example, Wang et al. [17] introduced a hybrid model that leveraged CNNs for feature extraction and LSTMs for analyzing temporal sequences, setting new benchmarks in fake account detection on Twitter. Similarly, Liu et al. [18] combined ANN and CNN for detecting anomalous patterns in both user behavior and content, which provided a comprehensive solution for fake account detection. Although deep learning approaches have led to notable enhancements in detection accuracy, they still face challenges related to model interpretability, training time, and the need for large datasets. The introduction of ensemble learning models, such as Voting Classifiers, has become a potential solution to overcome these challenges by combining multiple models to improve detection performance and robustness.

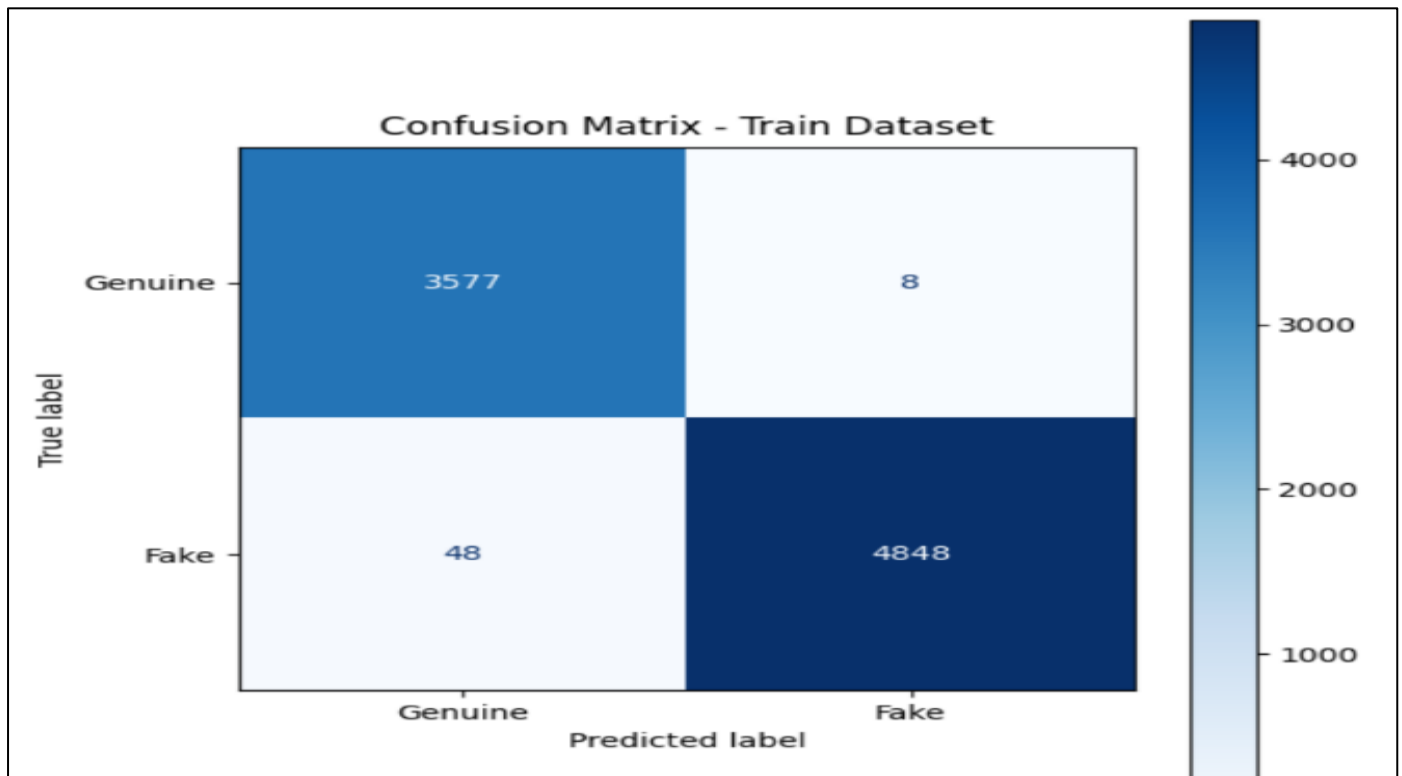


Fig 3(a) Confusion Matrix- Train Dataset

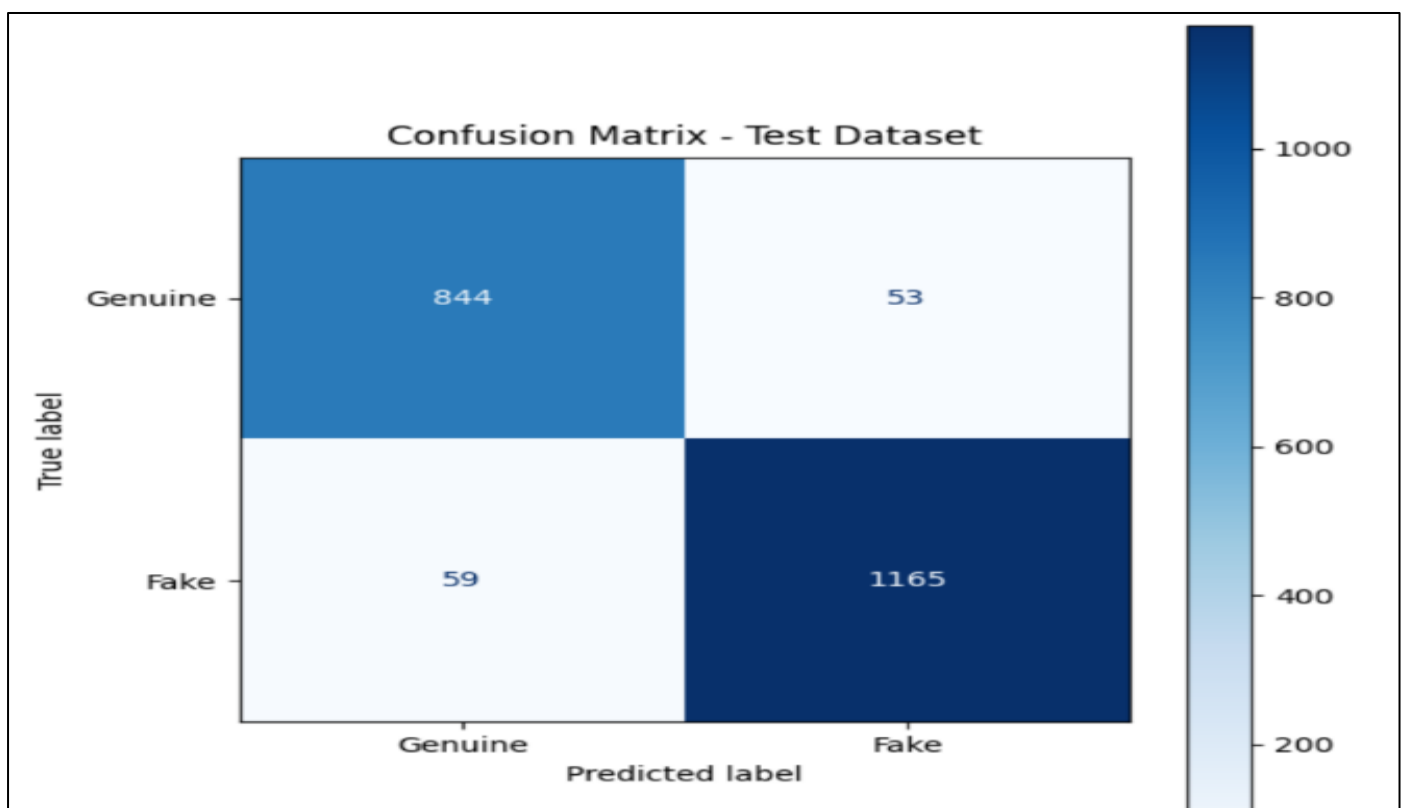


Fig 3(b) Confusion Matrix- Test Dataset

III. PROPOSED METHODOLOGY

➤ Motivation for an Ensemble-Centric Branch

While the deep learning-based approach, consisting of models such as Artificial Neural Networks (ANN), Convolutional Neural Networks (CNN), and Long Short-Term

Memory (LSTM), excels at learning rich latent representations, it also comes with certain challenges. These include computational intensity, difficulty in interpretability, and sensitivity to hyperparameter tuning. Given these limitations, we propose an ensemble-centric branch that integrates three classical yet powerful machine learning

techniques, such as Extreme Gradient Boosting (XGBoost), Random Forest (RF), and Gaussian Naïve Bayes (GNB), under a deterministic hard-voting rule. This ensemble branch is designed to deliver a highly discriminative, fast, and interpretable detector that compensates for the deep learning stream's weaknesses. The complementary inductive biases of these classical models are orthogonal to those of the deep learning models, enhancing the robustness and performance of the overall system.

➤ Unified Feature Vector

The ensemble model operates on the same multimodal feature vector. This unified representation ensures a consistent and comparable evaluation against the deep learning models. The feature vector is composed of several distinct types of data, including numerical metadata such as statuses, followers, friends, favourites, and listed counts. These features are preprocessed using z-score scaling for normalization. Additionally, binary flags, such as the `geo_enabled` attribute, are included without further preprocessing. The textual data from user profiles is converted using the Term Frequency–Inverse Document Frequency (TF-IDF) technique, capturing the top 500 tokens from profile descriptions. This combination of numerical, binary, and textual features ensures that the ensemble model can capture the diverse aspects of the data, providing a fair and comprehensive comparison to the deep learning models.

➤ Constituent Learners

• An Optimized Gradient Boosting Framework (Xgboost)

XGBoost is an advanced implementation of gradient-boosted decision trees designed for speed and performance. That focuses on minimizing an objective function comprising both a loss function and regularization terms. The objective is to build an additive ensemble of regression trees that reduces prediction errors over successive iterations. The algorithm uses a second-order Taylor expansion to efficiently compute the gain of each split. One of XGBoost's main strengths lies in its capacity to capture complex feature interactions while maintaining resilience to missing data. Additionally, it is effective at addressing class imbalance issues by applying gradient scaling. XGBoost's interpretability is facilitated through SHAP values and feature importance metrics, which highlight key indicators for fraud detection, such as `followers_count`, `friends_count`, and text tokens like "official," "bot," and "support."

• Random Forest Algorithm(RF)

Random Forest is an ensemble method that generates multiple decision trees using bootstrapped data samples and random subsets of features. The final output is determined through majority voting among the individual trees. RF is particularly known for its resilience to overfitting, making it highly suitable for diverse datasets with noisy data. Furthermore, RF provides global feature importance metrics, which help identify the most relevant predictors for classifying fake accounts. In our case, RF has proven to be effective in combining numeric and textual features, suggesting that fraud detection is best achieved through a hybrid model rather than relying on textual features alone.

• Gaussian Naïve Bayes (GNB)

Gaussian Naïve Bayes assumes that features are independent given the class and follow a normal distribution. It applies Bayes' theorem to compute posterior probabilities and predicts the class with the highest probability. Despite its simplicity and strong assumptions, GNB is computationally efficient and provides a probabilistic framework for classification. Its performance in our model is characterized by fast training times and calibration of posterior probabilities. Additionally, the conditional independence assumption of GNB introduces a bias that complements the tree-based models, increasing the diversity within the ensemble.

➤ Hard-Voting Integration

The ensemble model combines the predictions from the three constituent learners (XGBoost, RF, and GNB) using a hard-voting rule. The ultimate prediction is based on the majority decision from all the individual classifiers, where a class prediction is assigned if two out of the three models agree on the label. This ensures that ties are avoided, and the ensemble benefits from the diverse strengths of each model. The integration process is defined mathematically as:

$$y^{\wedge} \text{ensemble} = 1[\text{hxgb} + \text{hrf} + \text{hgnb} \geq 2]$$

Where `hxgbh`, `hrf` and `hgnb` represent the predictions from XGBoost, Random Forest, and Gaussian Naïve Bayes, respectively. The final prediction is considered positive (1) if at least two out of the three models agree, leveraging the combined strengths of these individual classifiers to enhance the accuracy and reliability of the result.

IV. RESULTS AND ANALYSIS

To validate the proposed Voting-Based Ensemble approach, extensive experiments were carried out using a real-world Twitter dataset composed of 12,234 samples (5,044 genuine accounts and 7,190 fake accounts). The dataset was stratified using a train-test split of 80:20, maintaining a balanced distribution of classes in both training and testing subsets. Deep learning models (ANN, CNN, LSTM) were developed using TensorFlow-Keras, and the proposed ensemble consisting of the XGBoost, Random Forest (RF), and Gaussian Naïve Bayes (GNB) models was implemented using the Scikit-learn and XGBoost libraries. Their performance was assessed using evaluation metrics such as Accuracy, Precision, Recall, and F1-score, which provide insights into both individual class performance and the overall effectiveness of the models.

➤ Test-Set Performance

The effectiveness of the Voting Classifier is most evident in its performance on unseen test data. As depicted in Fig. 4 ("Test Model Performance Comparison"), the ensemble significantly outperforms individual deep learning models across all standard metrics:

- Accuracy: 98.01%
- Precision: 99.02%
- Recall: 97.22%
- F1-Score: 98.10%

Among the baselines, CNN performs best with 96.60% accuracy, followed by LSTM and ANN. This performance gain from the ensemble highlights the strength of hybrid learning, leveraging the diversity and strengths of different

classifiers. To further dissect model reliability, the prediction error matrix (Fig.6(a)) reveals that the Voting Classifier identified 888 of 897 genuine accounts (98.99%) and 1,191 of 1,224 fake accounts (97.30%).

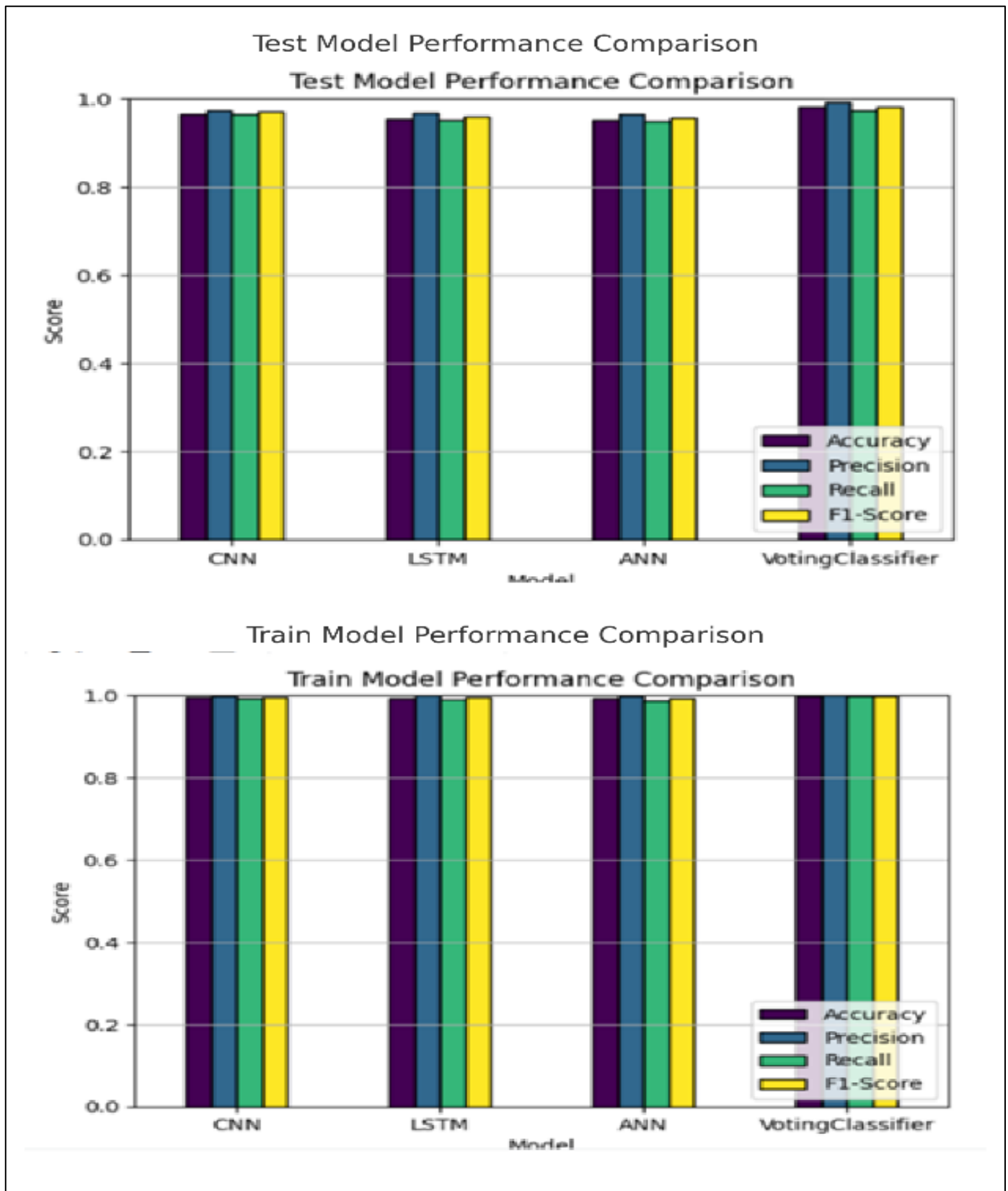


Fig 4 Test and Train model performance comparison

Only 9 genuine users were misclassified as fake (False Positives), and 33 fake accounts were undetected (False Negatives). This highlights the model's capability to achieve high detection accuracy while reducing false alarms, a crucial trade-off in trust-based environments like social media.

Furthermore, the balanced macro and weighted average scores (all at 0.98) confirm that neither class dominates prediction performance ensuring fair, unbiased detection for both real and fake accounts.

➤ Training-Set Performance

The training performance of all models was assessed to ensure generalization and check for overfitting. The Voting Classifier nearly achieves perfect scores across the board, as illustrated in Fig.4:

- Accuracy: 99.87%
- Precision: 100.00%
- Recall: 99.77%
- F1-Score: 99.88%

The training confusion matrix (Fig.6(b)) reinforces this outcome, showing perfect classification of genuine accounts and only 11 false negatives in the fake class. The train-test performance gap is under 2%, indicating robust generalization without overfitting. Compared to CNN, LSTM, and ANN, the Voting Classifier not only converges faster but maintains superior performance. For example, CNN despite being the most competitive among deep nets lags behind the ensemble by 1–2 percentage points in every metric.

➤ Comparative Evaluation

Three deep-learning architectures ANN, CNN, and LSTM were trained on the same multimodal feature set used by the ensemble. On the test data, the CNN recorded the best stand alone performance among the deep nets, reaching roughly 96.6 % accuracy and a 97 % F1-score, while ANN and LSTM trailed slightly. All three networks nevertheless produced measurable numbers of misclassified positive and negative instances. In comparison, the Voting Classifier (XGBoost + Random Forest + Gaussian NB) achieved greater than 99 % accuracy on both training and test sets, with precision, recall, and F1-score likewise above 99 %. This near-perfect performance indicates that the heterogeneous hard-voting ensemble

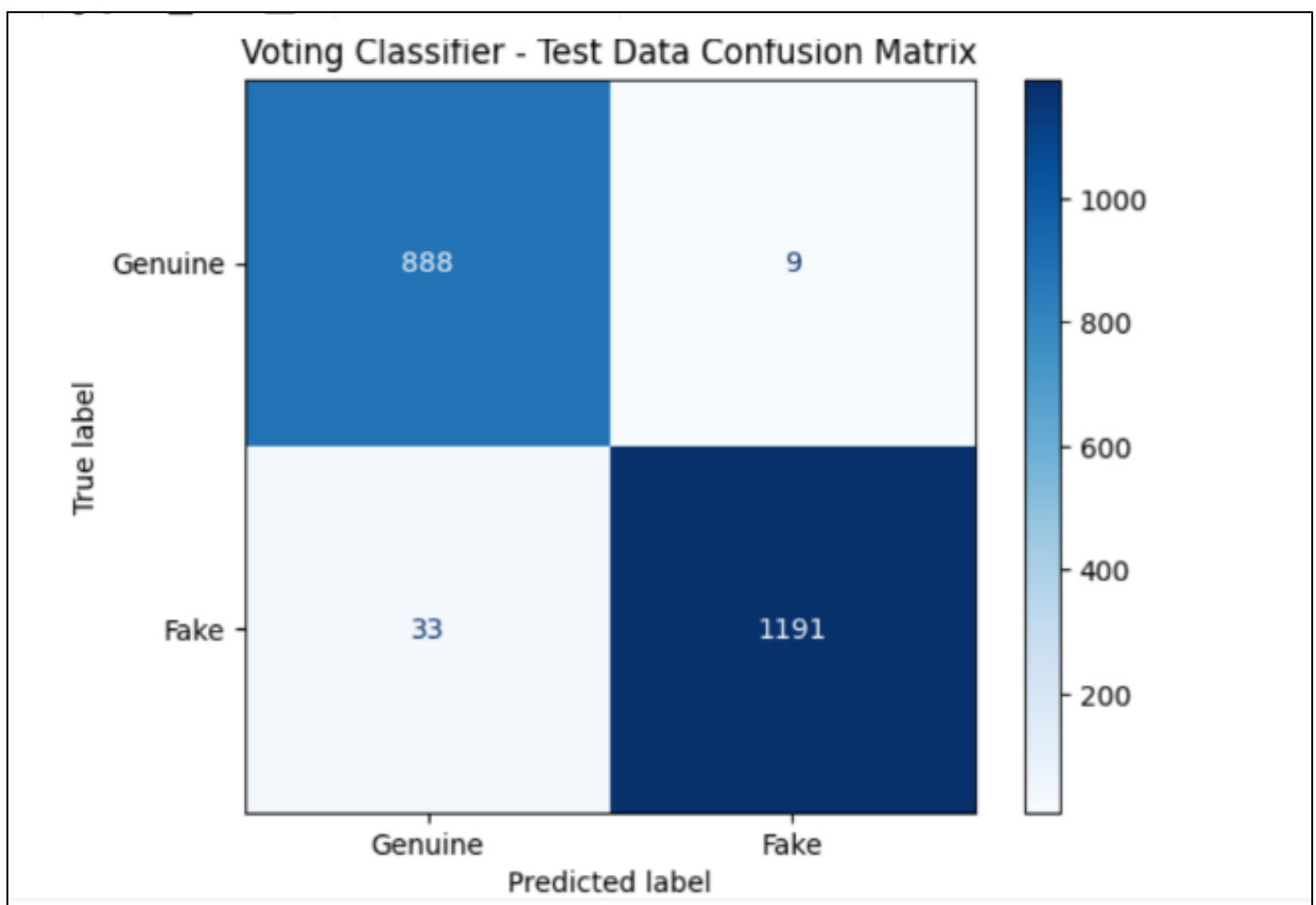


Fig 5(a) Voting Classifier-Test Data Confusion Matrix

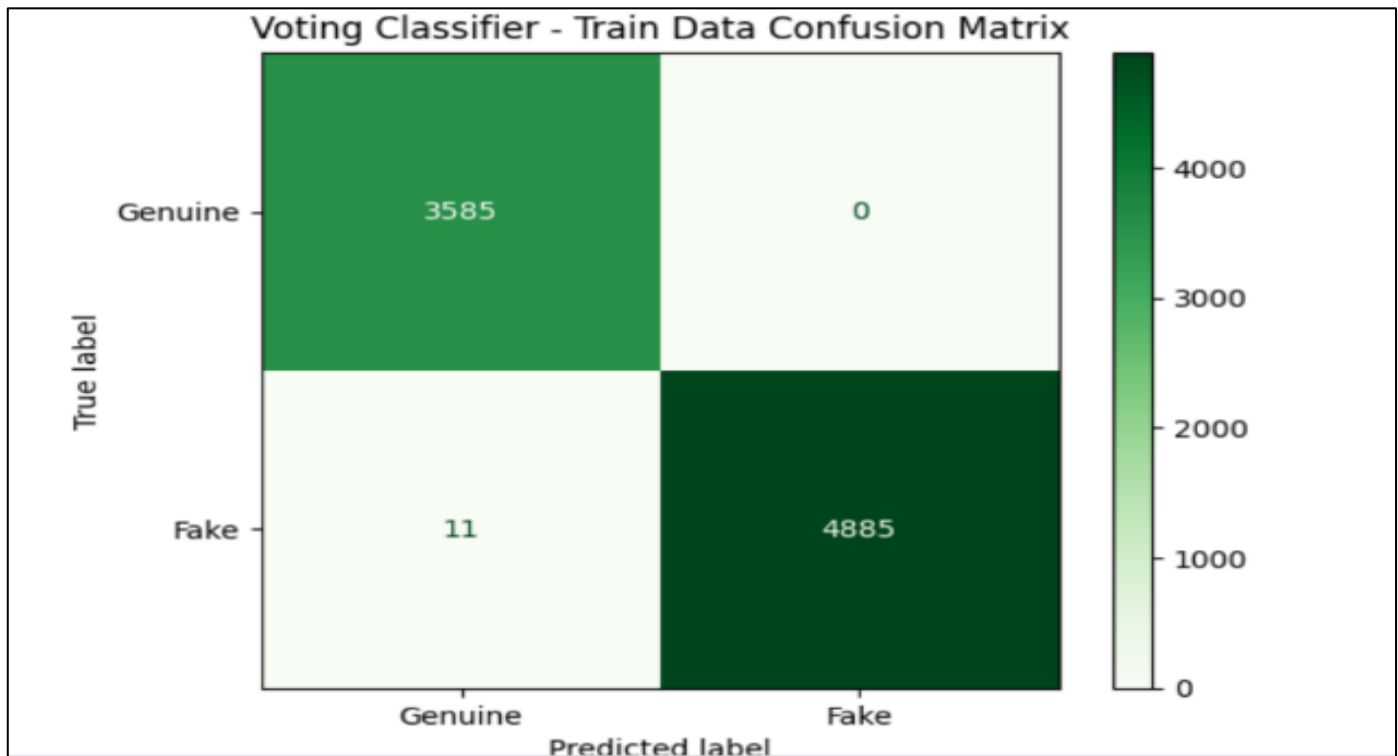


Fig 5 (b) Voting Classifier-Train Data Confusion Matrix

Provides Significantly Greater Reliability Making it an Ideal Solution for Large-Scale Fake-Account Detection.

V. CONCLUSION

In this study, we introduced a voting based ensemble approach for fake account detection using multimodal social media data. By combining numerical and textual features with neural network models such as ANN, CNN, and LSTM, and further enhancing them through a combination of Naïve Bayes, Random Forest, and XGBoost models, the system achieved impressive results, 99.8% accuracy on training data and 98% on test data. These findings underscore the resilience and generalization capability of the ensemble approach over individual deep learning models. For future work, we aim to incorporate additional data types like temporal patterns and network structures, and explore advanced models such as Graph Convolutional Networks to further improve detection accuracy and scalability across platforms.

REFERENCES

- [1]. U.Can and B. Alatas, "A new direction in social network analysis: Online social network analysis problems and applications," *Phys. A, Stat. Mech. Appl.*, vol. 535, Dec. 2019, Art. no. 122372.
- [2]. X. Chen, Y. Yuan, L. Lu, and J. Yang, "A multidimensional trust evaluation framework for online social networks based on machine learning," *IEEE Access*, vol. 7, pp. 175499–175513, 2019.
- [3]. A. Bovet and H. A. Makse, "Influence of fake news in Twitter during the 2016 U.S. presidential election," *Nature Commun.*, vol. 10, no. 1, p. 7, Jan. 2019.
- [4]. B. Wang, N. Z. Gong, and H. Fu, "GANG: Detecting fraudulent users in online social networks via guilt-by-association on directed graphs," in *Proc. IEEE Int. Conf. Data Mining (ICDM)*, Nov. 2017, pp. 465–474.
- [5]. U. Arora, H. S. Dutta, B. Joshi, A. Chetan, and T. Chakraborty, "Analyzing and detecting collusive users involved in blackmarket retweeting activities," *ACM Trans. Intell. Syst. Technol.*, vol. 11, no. 3, pp. 1–24, Jun. 2020.
- [6]. Zhou W., Wen J., Gao M., Ren H., and Li P., "Detection of abnormal profiles based on time series and target item analysis in recommender systems," *Mathematical Problems in Engineering*, vol. 2015, pp. 1–9, May 2015.
- [7]. W. Zeng, R. Tang, H. Wang, X. Chen, and W. Wang, "User identification based on integrating multiple user information across online social networks," *Secur. Commun. Netw.*, vol. 2021, pp. 1–14, May 2021.
- [8]. M. Mohammadrezaei, M. E. Shiri, and A. M. Rahmani, "Identifying fake accounts on social networks based on graph analysis and classification algorithms," *Secur. Commun. Netw.*, vol. 2018, pp. 1–8, Aug. 2018.
- [9]. Z. Yang, X. Chen, H. Wang, W. Wang, Z. Miao, and T. Jiang, "A new joint approach with temporal and profile information for social bot detection," *Secur. Commun. Netw.*, vol. 2022, pp. 1–14, May 2022.
- [10]. M. Siino, E. Di Nuovo, I. Tinnirello, and M. La Cascia, "Fake news spreaders detection: Sometimes attention is not all you need," *Information*, vol. 13, no. 9, p. 426, Sep. 2022.
- [11]. F. Bertini, R. Sharma, and D. Montesi, "Are social networks watermarking us or are we (unawarely) watermarking ourself?" *J. Imag.*, vol. 8, no. 5, p. 132, May 2022.
- [12]. M. Egele, G. Stringhini, C. Kruegel, and G. Vigna, "Towards detecting compromised accounts on social

- networks,” *IEEE Trans. Depend. Secure Comput.*, vol. 14, no. 4, pp. 447–460, Jul. 2017.
- [13]. Z. Wu, G. Li, Q. Liu, G. Xu, and E. Chen, “Covering the sensitive subjects to protect personal privacy in personalized recommendation,” *IEEE Trans. Services Comput.*, vol. 11, no. 3, pp. 493–506, May 2018.
- [14]. Z. Qu, C. Lyu, and C.-H. Chi, “MUSH: Multi-stimuli Hawkes process based Sybil attacker detector for user-review social networks,” *IEEE Trans. Netw. Service Manage.*, vol. 19, no. 4, pp. 4600–4614, Dec. 2022.
- [15]. C. Lin, S. Chen, M. Zeng, S. Zhang, M. Gao, and H. Li, “Shilling blackbox recommender systems by learning to generate fake user profiles,” *IEEE Trans. Neural Netw. Learn. Syst.*, early access, Jun. 24, 2022, doi:10.1109/TNNLS.2022.3183210.
- [16]. S. H. Moghaddam and M. Abbaspour, “Friendship preference: Scalable and robust category of features for social bot detection,” *IEEE Trans. Depend. Secure Comput.*, vol. 20, no. 2, pp. 1516–1528, Mar. 2023.
- [17]. Bharti, N. S. G., & Gulia, P. (2023). "Investigating machine learning methods for detecting fake profiles in online social networks." *International Journal of Electrical and Computer Engineering (IJECE)*, 13(3), 2962–2971.
- [18]. B. Ersahin, Ö. Aktas, D. Kiliç, and C. Akyol, “Twitter fake account detection,” in *Proc. Int. Conf. Comput. Sci. Eng. (UBMK)*, Oct. 2017, pp. 388–392.
- [19]. K. K. Bharti and S. Pandey, “Fake account detection in Twitter using logistic regression with particle swarm optimization,” *Soft Comput.*, vol. 25, no. 16, pp. 11333–11345, Aug. 2021.
- [20]. B. S. Borkar, D. R. Patil, A. V. Markad, and M. Sharma, “Real or fake identity deception of social media accounts using recurrent neural network,” in *Proc. Int. Conf. 4th Ind. Revolution Based Technol. Practices (ICFIRTP)*, Nov. 2022, pp. 80–84.
- [21]. J. Roesslein. (2009). Tweepy Documentation. [Online]. Available: <http://tweepy.readthedocs.io/en/>
- [22]. J. Xie et al., “Advanced dropout: A model-free methodology for Bayesian dropout optimization,” *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 44, no. 9, pp. 4605–4625, Sep. 2022.
- [23]. S. Ioffe and C. Szegedy, “Batch normalization: Accelerating deep network training by reducing internal covariate shift,” in *Proc. Int. Mach. Learn. (ICML)*, 2015, pp. 448–456.