

INTELLIGENT DETECTION OF MALICIOUS SOCIAL MEDIA PROFILES USING MULTI-DIMENSIONAL ANALYTICS

¹R MANOJ KUMAR, ²KODTGELA ABHIGNA, ³LOMTE SRI KANTH, ⁴KUNTA NANDHINI, ⁵LAVANYA SHAHI, ⁶MAGAPU KUSUMA SATYA SAI HARI VINAY

¹ASSISTANT PROFESSOR, ^{2,3,4,5&6}UG SCHOLAR

DEPARTMENT OF CSE, NARSIMHA REDDY ENGINEERING COLLEGE (UGC- AUTONOMOUS) MAISAMMAGUDA (V), KOMPALLY, SECUNDERABAD,
TELANGANA-500100

ABSTRACT

The rapid growth of social media platforms has significantly increased online communication and information sharing. However, it has also led to the emergence of malicious profiles, fake accounts, bots, and fraudulent users that spread misinformation, perform phishing attacks, and manipulate online communities. Detecting such malicious profiles has become a critical challenge for maintaining trust and security in social media ecosystems. This research proposes an Intelligent Detection System for Malicious Social Media Profiles using Multi-Dimensional Analytics. The proposed framework analyzes multiple dimensions of user behavior, including profile attributes, content activity, network relationships, and temporal interaction patterns. By integrating these diverse data sources, the system can effectively identify suspicious behaviors that indicate malicious intent. Machine learning and data analytics techniques are employed to process large volumes of social media data and classify users as legitimate or malicious. The framework extracts key features such as posting frequency, follower-following relationships, content similarity, and engagement patterns. These features are then analyzed using advanced classification algorithms to detect abnormal patterns associated with fake or harmful accounts. The proposed system improves detection accuracy by combining behavioral analysis, content analysis, and network analysis in a unified multi-dimensional framework. Experimental results demonstrate that the approach can effectively identify malicious profiles while reducing false positives. This research contributes to enhancing trust, security, and reliability in social media platforms, helping to protect users from harmful online activities.

INTRODUCTION

Social media platforms have become an essential part of modern communication, enabling people to share information, interact with communities, and access digital services. Popular platforms such as Facebook, Twitter (X), and Instagram host billions of users worldwide. While these platforms provide numerous benefits, they also face significant security challenges due to the increasing presence of **malicious profiles**, fake accounts, spambots, and coordinated misinformation campaigns. Malicious social media profiles are often created to perform activities such as **spreading misinformation, phishing attacks, online fraud, identity theft, and spam distribution**. These accounts can manipulate public opinion, promote harmful content, and reduce the trust of genuine users in social networking platforms. Traditional detection methods mainly rely on **manual moderation or rule-based systems**, which are often insufficient to handle the massive volume of data generated by millions of users. To address these challenges, researchers have started using **advanced data analytics and artificial intelligence techniques** for automated malicious profile detection. Multi-dimensional analytics plays a crucial role in this process by analyzing various

aspects of user behavior and account characteristics. These dimensions may include **profile information, user activity patterns, network connections, content analysis, and temporal behavior**. By combining these multiple dimensions, intelligent detection systems can identify suspicious patterns and differentiate between legitimate users and malicious accounts more effectively. Machine learning and deep learning models can learn from large datasets and continuously improve their ability to detect emerging threats and sophisticated attack strategies. The proposed approach focuses on developing an **intelligent framework that integrates multi-dimensional analytics with machine learning techniques** to enhance the detection of malicious social media profiles. This framework aims to improve platform security, protect user privacy, and increase trust among social media users by enabling early identification and mitigation of malicious activities.

LITERATURE REVIEW

Introduction to Malicious Profile Detection

Social media platforms such as Facebook, Instagram, and Twitter have become essential communication channels, but they also face serious security issues due to the rise of **fake or malicious profiles**. These accounts are often used for activities such as phishing, misinformation spreading, identity theft, and spam campaigns. Automated detection systems are therefore required to identify suspicious accounts and protect genuine users. Machine learning-based detection systems analyze profile attributes, behavioral patterns, and interaction networks to distinguish legitimate users from malicious actors.

Early Rule-Based and Feature-Based Detection Methods

Early research on malicious profile detection relied on **rule-based systems and manual feature analysis**. These systems examined simple indicators such as incomplete profile information, abnormal activity levels, and suspicious friend connections. Although these approaches were easy to implement, they lacked scalability and struggled to detect sophisticated fake accounts created using automated bots or coordinated networks. As a result, researchers began exploring data-driven machine learning approaches for more accurate detection.

Machine Learning Approaches for Fake Profile Detection

Machine learning algorithms such as **Support Vector Machines (SVM), Decision Trees, Random Forest, and Naïve Bayes** have been widely used for detecting malicious social media accounts. These models analyze multiple attributes such as profile metadata, posting frequency, network connections, and user interactions to classify accounts as genuine or fake. Studies show that ensemble methods like Random Forest often provide higher accuracy due to their ability to combine multiple decision models and handle complex data patterns.

However, traditional machine learning methods require manual feature engineering and may struggle to detect evolving attack strategies. Malicious actors continuously adapt their techniques, making it necessary for detection models to learn from large datasets and dynamic user behaviors.

Deep Learning Techniques for Malicious Profile Detection

Recent research has focused on **deep learning models** such as Convolutional Neural Networks (CNN), Recurrent Neural Networks (RNN), and Graph Neural Networks (GNN) to improve detection accuracy. These models can automatically extract complex features from textual content, images, and user behavior patterns. Deep learning systems also analyze profile images, social interactions, and content patterns to detect suspicious activities such as bot behavior, coordinated misinformation campaigns, or impersonation attempts.

For example, AI-based systems can analyze behavioral anomalies such as rapid posting activity, repetitive interactions, or abnormal follower patterns. These indicators help identify automated bot accounts or malicious users attempting to manipulate platform engagement metrics.

Multi-Dimensional Analytics for Social Media Security

Modern detection frameworks use **multi-dimensional analytics**, which combines several data sources to improve detection performance. Instead of relying on a single feature type, these systems analyze multiple dimensions such as:

- **Profile attributes:** username, profile completeness, account age
- **Behavioral patterns:** posting frequency, login activity, interaction patterns
- **Content analysis:** text sentiment, spam keywords, misinformation signals
- **Network structure:** follower relationships and community structures

By integrating these dimensions, detection systems can identify coordinated fake account networks and complex malicious behaviors that single-feature models might miss.

Challenges and Research Gaps

Despite significant progress, several challenges remain in malicious profile detection research:

- Handling **large-scale social network data** efficiently
- Detecting **coordinated bot networks and social bots**
- Addressing **data imbalance between real and fake profiles**
- Adapting to **rapidly evolving attack techniques**

Researchers are increasingly exploring **hybrid models combining machine learning, deep learning, and graph-based analysis** to build more robust and scalable detection systems.

Summary

The literature indicates that malicious social media profile detection has evolved from simple rule-based approaches to advanced **AI-driven multi-dimensional analytics frameworks**. Modern systems combine behavioral analysis, social network structure, and content analysis using machine learning and deep learning techniques. These approaches significantly improve detection accuracy and help maintain trust and security within social media platforms.

SYSTEM ANALYSIS

EXISTING SYSTEM

- The paper [1] constitutes the DeepSybil model, intended to enhance security in online social networks by categorizing Sybil attacks. The model employs sophisticated methodologies, including the Gannet Optimisation Algorithm and Graph Attention Networks, to detect fraudulent accounts and destructive behaviors in OSN precisely. Data collection is essential in the model, where the SNFA dataset is utilized to distinguish between genuine and fraudulent accounts.
- The proposed methodology [2] classifies a dataset comprising trusted files, malware, and ransomware by employing multiple deep learning models, including LeNet, AlexNet, Standard-CNN, and VGG-16. To augment the comprehensibility of the model's predictions, the authors utilized the Gradient-weighted Class Activation Mapping (Grad-CAM) technique [10], which identifies and emphasizes the significant areas of the input images that play a crucial role in the model's decision-making process.
- This paper [8] presents a pioneering procedure for detecting anomalous social network behaviors in real time using CNN. The CNN model differentiates between benign and dangerous users by scrutinizing user-profiles and dynamic actions. The effort employs a thorough dataset to extract and classify characteristics, improving the precision of detecting abnormal behavior.
- The authors [11] utilized a hybrid strategy incorporating content filtering and collaborative filtering algorithms to compare spam tweets and similar malicious profiles. Additionally, a recommendation system was implemented to assess the quality of tweets and videos disseminated by system profiles. In addition, the authors employed a convolutional neural network (CNN) algorithm for spam identification. Recent research has categorized online social network (OSN) profiles into verified, genuine, and fake accounts. This classification is indispensable for the creation of precisely targeted detection strategies. For example, Ellaky et al. accompanied a systematic literature review that categorized accounts based on their characteristics and behaviors and highlighted the evolution of detection systems [12].
- User activity patterns, profile features, and social connections are frequently examined to identify fraudulent accounts. ML has fundamentally

influenced the detection of fraudulent accounts. Accounts have been classified using innumerable algorithms based on several features: supervised learning, unsupervised learning, and hybrid models. Numerous studies employ supervised learning techniques to classify accounts, including SVM, RF, and ANN [13].

- The researchers have investigated unsupervised methods to identify malware. The author proposed a procedure that uses clustering techniques to combine similar user profiles, thereby facilitating the identification of anomalies that suggest bot activity. Integrating multiple algorithms or techniques has enhanced the efficacy of hybrid approaches [14], [15]. For example, integrating temporal activity data with profile-centric features has improved detection accuracy.

DISADVANTAGES

- An existing system presents a pioneering procedure for detecting anomalous social network behaviors in real time using CNN. The CNN model differentiates between benign and dangerous users by scrutinizing user-profiles and dynamic actions. The effort employs a thorough dataset to extract and classify characteristics, improving the precision of detecting abnormal behavior.
- Conventional approaches to inauthentic account detection, including rule-based systems and rudimentary machine learning models, frequently need to be revised to confront these obstacles.

PROPOSED SYSTEM

- The paper aims to propose an innovative model that integrates visual, temporal, and network analysis to detect fraudulent accounts on OSNs precisely and reliably. To augment detection capabilities by implementing sophisticated feature extraction methodologies from snaps, user behavioral patterns, and social network frameworks. Assess the proposed model using the Cresci 2017 dataset to ascertain its efficacy in practical applications and to juxtapose its performance with conventional machine learning techniques. Create a precise and scalable model capable of real-time application across diverse social media platforms, effectively countering the menace of fraudulent accounts.
- Recursive Feature Elimination (RFE) is a technique for choosing the most essential features in a dataset by iteratively eliminating the least significant ones. This method employs recursion to eliminate the least critical characteristics, considering the proposed methodology's performance and enabling the model to prioritize the most informative features.

ADVANTAGES

- Visual and Temporal Analysis: The model incorporates pioneering techniques for analyzing profile and user activity patterns over time, capturing subtle visual clues and behavioral anomalies indicative of fake or bot accounts.

- Graph-Based Network Analysis: The model employs graph-based methodologies to identify anomalies in social connections and network structures, augmenting its capacity to detect fraudulent accounts.

IMPLEMENTATION

MODULES

SERVICE PROVIDER

In this module, the Service Provider has to login by using valid user name and password. After login successful he can do some operations such as Train & Test Datasets, View Trained and Tested Datasets Accuracy in Bar Chart, View Trained and Tested Datasets Accuracy Results, View Detected User Profile Status, View Detected User Profile Status Ratio, Download Predicted Data Sets, View Detected User Profile Status Ratio Results, View All Remote Users.

VIEW AND AUTHORIZE USERS

In this module, the admin can view the list of users who all registered. In this, the admin can view the user's details such as, user name, email, address and admin authorizes the users.

REMOTE USER

In this module, there are n numbers of users are present. User should register before doing any operations. Once user registers, their details will be stored to the database. After registration successful, he has to login by using authorized user name and password. Once Login is successful user will do some operations like REGISTER AND LOGIN, Detection Of User Profile Status, View Your Profile.

BOT ACCOUNT DETECTION (FAD)

Multimodal data, such as visual content, temporal activity patterns, and network interfaces, are scrutinized by the proposed method for distinguishing fake bot accounts on OSNs. The technique mentioned earlier, Fake Bot Account Detection (FAD), marks using innovative deep learning algorithms to evaluate the data. A sumptuous explanation of the proposed methodology, comprising three principal features: preprocessing, feature extraction, and classification. A cohesive classification model then monitors this. The proposed method capitalizes on the capabilities of state-of-the-art techniques to extract visual features, model intricate network structures, and capture sequential user behaviors. To extract pertinent features from profiles, user activity records, and social graphs, the proposed FAD methodology commences by pre-processing the Cresci 2017 dataset. Figure 2 contributes to the inclusive framework of the Proposed Model.

ALGORITHMS

DECISION TREE CLASSIFIERS

Decision tree classifiers are used successfully in many diverse areas. Their most important feature is the capability of capturing descriptive decision making knowledge from the supplied data.

Decision tree can be generated from training sets. The procedure

for such generation based on the set of objects (S), each belonging to one of the classes $C_1, C_2, \dots, C_k$ is as follows:

Step 1. If all the objects in S belong to the same class, for example C_i , the decision tree for S consists of a leaf labeled with this class

Step 2. Otherwise, let T be some test with possible outcomes $O_1, O_2, \dots, O_n$. Each object in S has one outcome for T so the test partitions S into subsets $S_1, S_2, \dots, S_n$ where each object in S_i has outcome O_i for T. T becomes the root of the decision tree and for each outcome O_i we build a subsidiary decision tree by invoking the same procedure recursively on the set S_i .

GRADIENT BOOSTING Gradient boosting is a [machine learning](#) technique used in [regression](#) and [classification](#) tasks, among others. It gives a prediction model in the form of an [ensemble](#) of weak prediction models, which are typically [decision trees](#).^{[1][2]} When a decision tree is the weak learner, the resulting algorithm is called gradient-boosted trees; it usually outperforms [random forest](#). A gradient-boosted trees model is built in a stage-wise fashion as in other [boosting](#) methods, but it generalizes the other methods by allowing optimization of an arbitrary [differentiable loss function](#).

K-NEAREST NEIGHBORS (KNN)

- Simple, but a very powerful classification algorithm
- Classifies based on a similarity measure
- Non-parametric
- Lazy learning
- Does not “learn” until the test example is given
- Whenever we have a new data to classify, we find its K-nearest neighbors from the training data

Example

- Training dataset consists of k-closest examples in feature space
- Feature space means, space with categorization variables (non-metric variables)
- Learning based on instances, and thus also works lazily because instance close to the input vector for test or prediction may take time to occur in the training dataset

LOGISTIC REGRESSION CLASSIFIERS

Logistic regression analysis studies the association between a categorical dependent variable and a set of independent (explanatory) variables. The name *logistic regression* is used

when the dependent variable has only two values, such as 0 and 1 or Yes and No. The name *multinomial logistic regression* is usually reserved for the case when the dependent variable has three or more unique values, such as Married, Single, Divorced, or Widowed. Although the type of data used for the dependent variable is different from that of multiple regression, the practical use of the procedure is similar. Logistic regression competes with discriminant analysis as a method for analyzing categorical-response variables. Many statisticians feel that logistic regression is more versatile and better suited for modeling most situations than is discriminant analysis. This is because logistic regression does not assume that the independent variables are normally distributed, as discriminant analysis does. This program computes binary logistic regression and multinomial logistic regression on both numeric and categorical independent variables. It reports on the regression equation as well as the goodness of fit, odds ratios, confidence limits, likelihood, and deviance. It performs a comprehensive residual analysis including diagnostic residual reports and plots. It can perform an independent variable subset selection search, looking for the best regression model with the fewest independent variables. It provides confidence intervals on predicted values and provides ROC curves to help determine the best cutoff point for classification. It allows you to validate your results by automatically classifying rows that are not used during the analysis.

NAÏVE BAYES

The naive bayes approach is a supervised learning method which is based on a simplistic hypothesis: it assumes that the presence (or absence) of a particular feature of a class is unrelated to the presence (or absence) of any other feature. Yet, despite this, it appears robust and efficient. Its performance is comparable to other supervised learning techniques. Various reasons have been advanced in the literature. In this tutorial, we highlight an explanation based on the representation bias. The naive bayes classifier is a linear classifier, as well as linear discriminant analysis, logistic regression or linear SVM (support vector machine). The difference lies on the method of estimating the parameters of the classifier (the learning bias). While the Naive Bayes classifier is widely used in the research world, it is not widespread among practitioners which want to obtain usable results. On the one hand, the researchers found especially it is very easy to program and implement it, its parameters are easy to estimate, learning is very fast even on very large databases, its accuracy is reasonably good in comparison to the other approaches. On the other hand, the final users do not obtain a model easy to interpret and deploy, they does not understand the interest of such a technique. Thus, we introduce in a new presentation of the results of the learning process. The classifier

is easier to understand, and its deployment is also made easier. In the first part of this tutorial, we present some theoretical aspects of the naive bayes classifier. Then, we implement the approach on a dataset with Tanagra. We compare the obtained results (the parameters of the model) to those obtained with other linear approaches such as the logistic regression, the linear discriminant analysis and the linear SVM. We note that the results are highly consistent. This largely explains the good performance of the method in comparison to others. In the second part, we use various tools on the same dataset ([Weka 3.6.0](#), [R 2.9.2](#), [Knime 2.1.1](#), [Orange 2.0b](#) and [RapidMiner 4.6.0](#)). We try above all to understand the obtained results.

RANDOM FOREST

Random forests or random decision forests are an ensemble learning method for classification, regression and other tasks that operates by constructing a multitude of decision trees at training time. For classification tasks, the output of the random forest is the class selected by most trees. For regression tasks, the mean or average prediction of the individual trees is returned. Random decision forests correct for decision trees' habit of overfitting to their training set. Random forests generally outperform decision trees, but their accuracy is lower than gradient boosted trees. However, data characteristics can affect their performance. The first algorithm for random decision forests was created in 1995 by Tin Kam Ho[1] using the random subspace method, which, in Ho's formulation, is a way to implement the "stochastic discrimination" approach to classification proposed by Eugene Kleinberg. An extension of the algorithm was developed by Leo Breiman and Adele Cutler, who registered "Random Forests" as a trademark in 2006 (as of 2019, owned by Minitab, Inc.). The extension combines Breiman's "bagging" idea and random selection of features, introduced first by Ho[1] and later independently by Amit Geman[13] in order to construct a collection of decision trees with controlled variance. Random forests are frequently used as "blackbox" models in businesses, as they generate reasonable predictions across a wide range of data while requiring little configuration.

SVM

In classification tasks a discriminant machine learning technique aims at finding, based on an *independent and identically distributed (iid)* training dataset, a discriminant function that can correctly predict labels for newly acquired instances. Unlike generative machine learning approaches, which require computations of conditional probability distributions, a discriminant classification function takes a data point x and assigns it to one of the different classes that are a part of the classification task. Less powerful than generative approaches, which are mostly used when prediction involves outlier detection, discriminant approaches require fewer computational resources and less training data, especially for a multidimensional feature space and when only posterior probabilities are needed. From a geometric perspective, learning a classifier is equivalent to

finding the equation for a multidimensional surface that best separates the different classes in the feature space. SVM is a discriminant technique, and, because it solves the convex optimization problem analytically, it always returns the same optimal hyperplane parameter—in contrast to *genetic algorithms (GAs)* or *perceptrons*, both of which are widely used for classification in machine learning. For perceptrons, solutions are highly dependent on the initialization and termination criteria. For a specific kernel that transforms the data from the input space to the feature space, training returns uniquely defined SVM model parameters for a given training set, whereas the perceptron and GA classifier models are different each time training is initialized. The aim of GAs and perceptrons is only to minimize error during training, which will translate into several hyperplanes' meeting this requirement.

CONCLUSION

The research presents a sophisticated model that combines a visual feature extractor, temporal activity extractor, and network structure analyzer to identify fake accounts accurately. The method was trained and assessed using the Cresci 2017 dataset. Each component focuses on distinct data parts: visual characteristics derived from profile photographs, temporal patterns extracted from activity logs, and relational structures obtained from social graphs. The visual feature extractor efficiently captures global interdependencies in visual data, offering a resilient feature representation of profile photos. The temporal activity extractor network is designed to analyze temporal sequences by learning from activity patterns across time. This ability is essential for detecting abnormal behaviors characteristic of fraudulent accounts. The network structure analyzer leverages the relational nature of social graphs to comprehend the interactions and interconnections among accounts. By integrating these potent methodologies, our model attains exceptional precision and resilience in identifying counterfeit accounts. Although the proposed methodology demonstrates encouraging outcomes, numerous opportunities exist for further investigation and enhancement. Applying sophisticated data augmentation methods to visual data enhances the variety of training examples, strengthening the visual transformer's resilience and generalization capabilities. Investigate advanced techniques for integrating multimodal data from visual, temporal, and network sources. Attention processes applied to different modalities can improve the integration of various groups of features. Enhance the model to ensure it can efficiently handle more extensive datasets and more intricate social graphs. Enhance the comprehensibility and clarity of the model's predictions by devising techniques to increase its

explainability and interpretability. It is essential to win consumers' trust and comprehend the fundamental factors contributing to identifying fraudulent accounts. Use transfer learning to exploit pre-trained models on enormous visual and temporal data datasets, diminishing the requirement for abundant labeled data and accelerating the training process. Emphasize the ability to detect and respond to bogus accounts immediately using real-time detection capabilities. This entails optimizing the model to achieve minimal delay in making inferences and seamlessly incorporating it into monitoring systems that operate in real-time. Verify the model's effectiveness on several datasets besides Cresci 2017 to confirm its suitability for social media platforms and situations. This involves doing cross-platform validation to assess the model's capacity to generalize. Discuss ethical considerations about privacy and bias in the detection of false accounts. Verify that the model adheres to data protection rules and lacks biases that may result in unjust treatment of specific user demographics. By following these future paths, we may improve false account detection models' efficiency, resilience, and relevance, hence promoting safer and more reliable social media platforms. Real-time detection necessitates models that progression and categorize incoming data streams with minimal latency, which can be computationally expensive when applied to large-scale OSNs. Resolving these complications necessitates the optimisation of the framework's computational architecture, maybe consuming distributed processing or low-latency models designed for streaming data. Scalability is a critical aspect, as OSNs can comprise of billions of people and activities. Consequent efforts ought to quintessence on amending the framework for big data contexts by exploiting effectual storage options, parallel processing methodologies, and dynamic model updates. Furthermore, the usage of adaptive learning techniques might facilitate the model's evolution in tandem with embryonic bot behaviours and innovative social media trends.

REFERENCES

- [1] B. Antony and S. Revathy, "Enhancing security in online social networks: Introducing the deepsybil model for Sybil attack detection," *Multimedia Tools Appl.*, vol. 83, no. 14, pp. 41911–41937, Oct. 2023, doi: 10.1007/s11042-023-16851-3.
- [2] G. Ciaramella, G. Iadarola, F. Martinelli, F. Mercaldo, and A. Santone, "Explainable ransomware detection with deep learning techniques," *J. Comput. Virol. Hacking Techn.*, vol. 20, no. 2, pp. 317–330, Sep. 2023, doi: 10.1007/s11416-023-00501-1.
- [3] G. Jethava and U. P. Rao, "Exploring security and trust mechanisms in online social networks: An extensive review," *Comput. Secur.*, vol. 140, May 01, 2024, Art. no. 103790, doi: 10.1016/j.cose.2024.103790.
- [4] R. J. Krishna, T. Gopalakrishnan, M. Divyapushpalakshmi, K. Amarendra, P. Dadheech, and S. Sengan, "Security and privacy concerns in social networks mathematically modified Metaheuristic-based approach," *J. Discrete Math. Sci. Cryptography*, vol. 27, nos. 2–A, pp. 371–382, Mar. 2024, doi: 10.47974/jdm-sc-1892.
- [5] F. El Mendili, M. Fattah, N. Berros, Y. Filaly, and Y. E. B. E. Idrissi, "Enhancing detection of malicious profiles and spam tweets with an automated honeypot framework powered by deep learning," *Int. J. Inf. Secur.*, vol. 23, no. 2, pp. 1359–1388, Apr. 2024, doi: 10.1007/s10207-023-00796-7.
- [6] R. R. Sekar, T. D. Rajkumar, and K. R. Anne, "Deep fake detection using an optimal deep learning model with multi head attention-based feature extraction scheme," *Vis. Comput.*, vol. 2024, pp. 1–18, Jul. 2024, doi: 10.1007/s00371-024-03567-0.
- [7] A. Shah, S. Varshney, and M. Mehrotra, "Detection of fake profiles on online social network platforms: Performance evaluation of artificial intelligence techniques," *Social Netw. Comput. Sci.*, vol. 5, no. 5, p. 489, Apr. 2024, doi: 10.1007/s42979-024-02839-9.
- [8] X. Wang, K. Wang, K. Chen, Z. Wang, and K. Zheng, "Unsupervised Twitter social bot detection using deep contrastive graph clustering," *Knowl.-Based Syst.*, vol. 293, Jun. 2024, Art. no. 111690, doi: 10.1016/j.knosys.2024.111690.
- [9] M. Aljabri, R. Zagrouba, A. Shaahid, F. Alnasser, A. Saleh, and D. M. Alomari, "Machine learning-based social media bot detection: A comprehensive literature review," *Social Netw. Anal. Mining*, vol. 13, no. 1, p. 20, Jan. 2023, doi: 10.1007/s13278-022-01020-5.
- [10] K. Antipova and H. Horban, "Positional encoding for transformers," in *Traditional and Innovative Scientific Research: Domestic and Foreign Experience*. Riga, Latvia: Baltija Publishing, 2024, doi: 10.30525/978-9934-26-436-8-1.
- [11] D. Punkamol and R. Marukatat, "Detection of account cloning in online social networks," in *Proc. 8th Int. Electr. Eng. Congr. (IEEECON)*, Mar. 2020, pp. 1–4, doi: 10.1109/IEEECON48109.2020.229558.
- [12] Z. Ellaky, F. Benabbou, and S. Ouahabi, "Systematic literature review of social media bots detection systems," *J. King Saud Univ.-Comput. Inf. Sci.*, vol. 35, no. 5, May 2023, Art. no. 101551, doi: 10.1016/j.jksuci.2023.04.004.
- [13] F. Ghorbanpour, M. Ramezani, M. A. Fazli, and H. R. Rabiee, "FNR: A similarity and transformer-based approach to detect multi-modal fake news in social media," *Social Netw.*

Anal. Mining, vol. 13, no. 1, p. 56, Mar. 2023, doi: 10.1007/s13278-023-01065-0.

[14] B. Goyal, N. S. Gill, P. Gulia, O. Prakash, I. Priyadarshini, R. Sharma, A. J. Obaid, and K. Yadav, "Detection of fake accounts on social media using multimodal data with deep learning," *IEEE Trans. Computat. Social Syst.*, early access, Aug. 7, 2024, doi: 10.1109/TCSS.2023.3296837.

[15] K. Hayawi, S. Saha, M. M. Masud, S. S. Mathew, and M. Kaosar, "Social media bot detection with deep learning methods: A systematic review," *Neural Comput. Appl.*, vol. 35, pp. 8903–8918, Mar. 2023, doi: 10.1007/s00521-023-08352-z.