

Fake Profile Identification in E-Learning Platform using Machine Learning

Kalai Selvi T.¹, Abirami S.², Aravind S.³, Ariharan M.⁴

¹Associate professor, ^{2,3,4}Student, Department of CSE, Erode Sengunthar Engineering College, Perundurai, Erode.

Email: ¹tkalaiselvi1281@gmail.com, ²abirami12345abi@gmail.com, ³aravindkdu579@gmail.com, ⁴ariharanariharan343@gmail.com

Abstract

In E-Learning Platforms like LinkedIn, identifying fake profiles poses significant challenges, affecting trust and engagement. Fake Profiles can mislead users, dilute the quality of interactions and undermine the credibility of the Platform. This survey explores the application of machine learning techniques for detecting fake profiles on e-learning platforms. The study examines various features that can indicate suspicious behaviours, such as anomalous login patterns, unusual interaction metrics, and inconsistent profile information. Through a comprehensive review of existing research, the study identifies the key challenges in implementing machine learning solutions for fake profile detection, such as data privacy concerns, feature engineering, and scalability. Additionally, the survey highlights the potential of using machine learning models and ensemble techniques to enhance detection accuracy. By consolidating insights from prior studies, this survey aims to provide a foundation for future research and development in safeguarding e-learning platforms against fake profiles, thereby enhancing a secure and trustworthy digital learning environment.

Keywords: Fake Profiles, E-Learning Platforms, Machine Learning, User Behaviour, Fraud Detection.

1. Introduction

Fake profiles on e-learning platforms pose significant risks to the integrity and security of these systems. [2] The creation of such fraudulent accounts may be driven by various

motivations, including but not limited to cheating on assessments, manipulating course participation metrics, exploiting system vulnerabilities, or harvesting course materials for resale. In some instances, fake profiles may also be used to disseminate malicious content or engage in deceptive practices such as impersonation or phishing, undermining the trust and authenticity that learners and educators place on these platforms.

The existence of fake profiles not only undermines the credibility of e-learning systems but also hinders meaningful data collection, analysis, and educational outcomes. For instance, fake users skew the insights derived from platform usage data, such as course engagement, completion rates, or learner performance. [8],[9] In academic and certification-driven environments, where credentials are highly valued, the detection of fake profiles is critical for maintaining the legitimacy of qualifications and the reputation of educational institutions.

Detecting fake profiles in e-learning platforms is a multifaceted challenge. Unlike traditional fraudulent activities where suspicious behaviour can be directly linked to financial transactions, fake profiles often mimic legitimate user behaviours, making detection more difficult. Furthermore, the scalability of modern e-learning platforms, which provide to millions of users, complicates the process of manual identification of fake accounts. [6] The reliance on self-reported information, such as names, email addresses, and demographic data, makes it easier for malicious actors to fabricate or obscure their identity. Automated detection mechanisms are essential in identifying fake profiles in a scalable and efficient manner.

Several characteristics are common among fake profiles, including incomplete or erroneous information, unusual usage patterns, and inconsistent interactions with course materials. [10], [11] However, these attributes alone are insufficient to conclusively determine fraudulent behaviour, as legitimate users may also exhibit similar traits, especially in large and diverse learning environments. This has led researchers to explore more sophisticated approaches, such as machine learning, to accurately differentiate between authentic and fake users.

2. Machine Learning Role

In the case of fraudulent profiles on e-learning platforms, machine learning plays a very critical role by using its ability to analyse and process large amounts of data very effectively. It helps in developing predictive models that can identify anomalies and patterns which indicate

fake profiles. Applying supervised learning techniques, it is possible to train algorithms using annotated datasets to identify profiles as real or fake from behavioural information, profile features, login behaviour, and engagement statistics. Additionally, unsupervised techniques such as clustering may be employed to identify unusual patterns in user behaviour such as numerous accounts with similar characteristics and questionable timelines of interactions. Machine learning algorithms such as decision trees, support vector machines, and neural networks enable the automation of the detection process, which minimizes the need for human intervention and increases scalability. The system can continuously improve its detection accuracy through feedback loops with the incorporation of machine-learning models within the platform, thus making it a robust and adaptive solution for maintaining the integrity and reliability of e-learning platforms.

3. Literature Survey

Bhrugumalla L. V. S Aditya et al. [1] in their study proposed a Heterogeneous Social Media Analysis for efficient Deep Learning Fake-Profile Identification. The methodology is structured around utilizing deep-transfer learning and advanced feature optimization techniques to enhance the detection of fake profiles across multiple social media platforms. First, a variety of datasets are collected from different sources, such as user posts, likes, comments, multimedia content, and login behaviors. Each data type is processed using specific methods: Audio, 1D signal, for example, passes through Fourier, Cosine, Convolutional, Gabor, and Wavelet Transforms. Images and videos pass through similar kinds of 2D transforms. Word2Vec transforms textual data, creating a feed to a binary Convolutional Neural Network, bCNN, that classifies whether profiles are real or fake.

Chen Lin, John et al. [2] in their research, “Shilling Black-Box Recommender System by Learning to Generate Fake User Profiles” suggests Leg-UP, a novel attack model targeted at generating undetectable fake user profiles to compromise recommender systems (RS). The methodology begins with sampling real user behavior patterns, which are used as templates to train a generative adversarial network (GAN). The generator in the GAN generates fake profiles by imitating real user ratings, outputting discrete values that mimic actual user interactions. For the transferability of the attack across different RS models, the parameters of the generator are optimized on a surrogate RS model.

Pradeep Kumar Roy, et al. [3] in their research “Fake Profile Detection on Social Networking Websites: A Comprehensive Review” includes fake account detection on social networking websites. It is preceded by a comprehensive literature review to identify recent advancements and challenges followed by data collection from popular social media platforms to gather information on registered users and label fake accounts. Then, the extracted features, such as profile information, behavioral patterns, network structure, and content analysis, are used to train and test machine learning models for the identification of fake accounts.

Faiza masood, et al. [4] in their research “Spammer Detection and Fake User Identification on Social Networks” proposed a method of spammers detection and fake user identification on social networks especially on Twitter based upon a critical review and categorization of existing Hidden Markov Model (HMM) spam detection techniques. The current study categorizes these approaches based upon their focus areas, namely, the identification of fake content, spammers through URLs, spammers in trending topics, and fake users. With user features including all profile characteristics, content features in regard to the nature of posts, and graph feature network connections in addition to patterns in interaction structures, this methodology makes an evaluation between these methods with an approach to eventually provide an even more complete taxonomy which makes the efficient identification of spammers easy.

Greeshma Lingam, et al. [5] in the research “Particle Swarm Optimization on Deep Reinforcement Learning for Detecting Social Spam Bots and Spam-Influential Users in Twitter Network” presented a comprehensive literature review of recent developments and challenges in the detection of spammers and fake users, which develops a taxonomy of Twitter spam detection approaches based on classifying techniques by the capability to detect fake content, spam based on URL, spam in trending topics, and fake users. The techniques presented are then analyzed and compared based on various features, such as user features, content features, graph features, structure features, and time features, to identify their relative strengths and weaknesses.

Abdelouahab Amira, et al. [6] in his research “Detection and Analysis of Fake News User's Communities in Social Media” proposed the methodology for detecting and analyzing fake news user communities in social media focuses on identifying organized groups involved in spreading fake news without relying on prior knowledge of the news content or user profiles. The process starts with the construction of a spatial-temporal similarity graph, a new graph

structure that connects social accounts participating in early stages of fake news campaigns based on their timing and activity patterns.

Xuxin Zhang, et al. [7] in his research proposed an “Attacking Recommender Systems with Plausible Profile”. Recommender systems (RS) have become an essential component of web services due to their excellent performance. Despite their great success, RS have proved to be vulnerable to data poisoning attacks, which inject well-crafted fake profiles into RS, so that the target items can be maliciously recommended. The study begins with the fact that it is easy to detect current poisoning attacks in RS; the features of the created fake profiles cannot be at odds with those of the normal profiles all the time.

Shigang hu, et al. [8] proposed the methodology of CISER is the integration of reviewer credibility analysis, user interest profiling, and fine-grained sentiment analysis to improve product recommendations. The process begins with a candidate feature extraction module, which uses context and sentiment confidence to identify important product features from the review corpus. Then, the reviewer credibility analysis assigns scores based on reviewer expertise, trustworthiness, and influence to filter out untrustworthy reviews. Finally, the user interest mining module analyzes review aesthetics and patterns to understand user preferences

Nahid R. Abid Althaqafi, et al. [9] in his research described “The impact of the Weighted Features on the Accuracy of X-Platforms User Credibility Detection Using Supervised ML”. The methodology for the X-Platform User Credibility Detection is a multi-step approach, which designs the supervised machine learning system to detect user credibility on the X-Platform, extract features from the user profiles, apply the feature weighting methods, namely Principal Component Analysis (PCA), correlation-coefficient algorithms, and tree-based approaches such as Extra Trees Classifier, to measure the importance of each of the features, train the model using the weighted features to detect user credibility and evaluate the performance of the different feature-weighting methods on various categories of datasets to get the best detection accuracy from any of them, making use of a publicly available dataset called ArPFN.

Moona Kanwal, et al. [10] in his research on “Machine Learning Approach to Classification of Online User's By Exploiting information Seeking Behavior” focused on classifying the online users using their information-seeking behavior which includes search, sharing, and verifying the information. The study begins by gathering user feedback by creating

a questionnaire which captures the diversity across gender, occupation, and age. Feature engineering is used to draw out relevant behavioral traits from the data. These are then fed into a K-Means clustering algorithm, which clusters users based on intent-based profiles.

Giuseppe Sansonetti, et al. [11] proposed the method of detecting the unreliable user in social media, mainly on the platforms of Twitter, is to carry out a detailed review and classification of the existing spam detection techniques. It starts by discussing several approaches that could be used to identify spammers and fake users and then develops a taxonomy based on the capabilities of techniques for detection, including identification of fake content, spam related to URLs, spam in trending topics, and fake user accounts. Each technique is compared along multiple dimensions, including user features (profile attributes), content features (characteristics of the posts), graph features (network interaction patterns), structural features (overall account structure), and temporal features (activity timelines).

Hichem Felouat et al. [12] on his study titled as, “eKYC-DF: A Large-Scale Deepfake Dataset for Developing and Evaluating EKYC Systems” addressed the development of the eKYC-DF dataset uses methodology targeting the challenges arising from the deepfake technology towards eKYC systems. The process starts with the collection of a large-scale dataset containing more than 228,000 high-quality fake videos that represent diverse demographics in terms of age, gender, and ethnicity.

Ana Aguilera, et al. [13] in their work proposed an approach to expand an existing credibility model on Twitter by adding bot detection capability, compute the credibility of each tweet according to three: text, account/user and social impact measures using alternative filters to analyze text and user account attributes, take account bot verification in the computing user credibility, build T-CREO based implementation of the extended model for real-time credibility computation, train and test machine learning algorithms over supervised learning.

Jan Eloff et al. [14] proposed a methodology for detecting fake identities on SMP based on the application of machine learning techniques that intend to distinguish between accounts forged by humans and those churned out by bots. The research starts with a dataset of users' accounts, which shall comprise verified human accounts besides identified fake accounts. Critical features are inferred from account characteristics already available within the user's profile; these include the friend count, follower count, and friend-to-follower ratio, which had been well proven in other bot detection studies.

Muhammad Amir Mehmood, et al. [15] for the purpose of trend promoters detection in Twitter it involves several developing a dataset of labelled users containing users categorized into trend promoters and normal users designing four discerning features, including number of total tweets, duplicate tweets, overlapping ngram, and peak-to-mean ratio, for trend promoters classification examining features used for spam and bot accounts classification to filter three efficacious features for trend promoters identification utilizing these seven features to develop the Push-To-Trend framework, which achieves an accuracy of applying the framework to identify and analyze trend promoters from the Urdu tweets repository and analyzing the results to reveal the proportion of trend promoters generating tweets related to hashtags.

The GitHub repository "Fake Profile Detection using ML" by harshitkgupta et al [16] focuses on identifying fake profiles in social networks through machine learning models, including Support Vector Machine (SVM), Neural Network, and Random Forest. Implemented in Python with libraries like NumPy, Pandas, and scikit-learn, the study provides both Jupyter Notebooks and Python scripts for ease of use. The repository includes detailed instructions for running the code and offers results in HTML and PDF formats. This research is a valuable resource for researchers and practitioners aiming to enhance online security through automated fake profile detection. The Figures 1to 3 illustrates the ROC of the three models neural network, random forest, and SVM in fake profile identification.

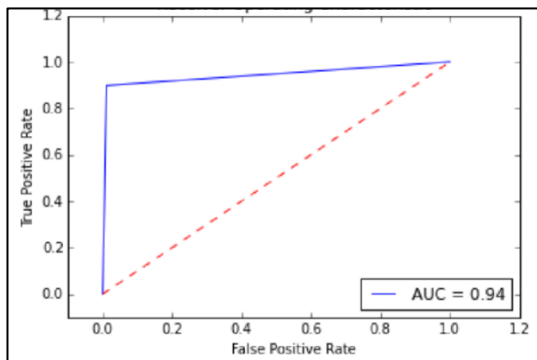


Figure 1. Neural Network

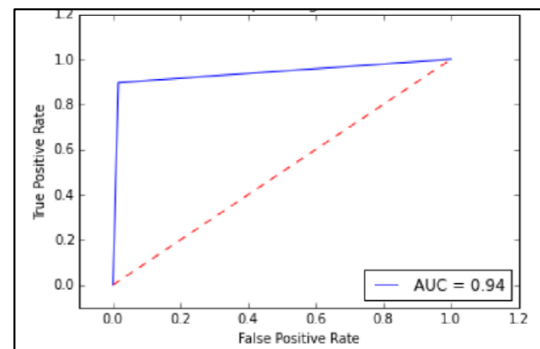


Figure 2. Random Forest

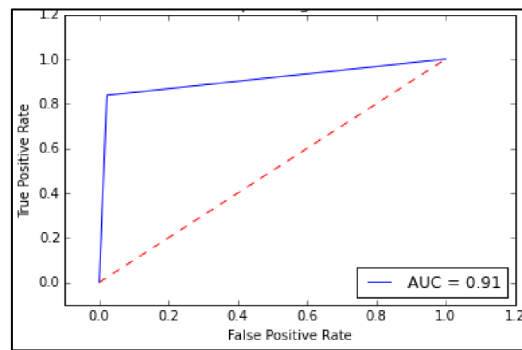


Figure 3. Support Vector Machine

ROC curves for the models Neural Network, Random Forest, and SVM are provided in comparison charts along with the AUC scores for each of them. It can be clearly seen that the ensemble model has a better performance in differentiating between classes as compared to the SVM model, which had a slightly lesser AUC of 0.91, while the higher AUC score of 0.94 for both Neural Network and Random Forest was obtained. The ensemble approach strengthens the model's classification ability through the strengths of different algorithms, and the curve is steeper and closer to the ideal top-left corner of the graph [16].

4. Comparative Study

The Table 1 below show the summarization of the literature review.

Table 1. Comparative Study

S. No	Citation	Algorithm Used	Dataset Used	Merits	Demerits	Application
1	[1]	Convolutional neural network (CNN)	Heterogeneous Social Media Data Source: Twitter API, Facebook API	Improve accuracy (8.3 %) to detecting fake profile. Enhanced precision (5.9%). Better Recalling (6.5% increase).	Complex for Require significant Computational resources. May Require continuous updates to stay ahead of evolving fake Profile Strain.	Social Media

2	[2]	Generative Adversarial Networks(GANs)	Synthetic Dataset Source: Amazon Reviews Dataset	Leg Up methodology offers improve attract transferability and invisibility, can be used to simulate real users by directly outputting discrete Rating	Its reliance on a generative adversarial network may introduce complexity and potential biases.	Recommender System
3	[3]	Support Vector Machines(SVM)	Social Networking Data Source:Stanford Network Analysis Platform	To Provide a comprehensive survey of recent develop in fake account. Generalized framework for fake profile detection.	Only focus the Summarizing a fake account detection. That may not contribute to the development of new detection methods.	Social Media
4	[4]	Logistic Regression	Social Network Data Source: DARPA Twitter Bot Dataset	It is used to offers an taxonomy for classifying approaches and compare technique for various feature, serving as a valuable researchers.	Solution provided are not novel	Social Networks
5	[5]	Deep Q-Learning algorithm used for detecting social bots.	Twitter Data Source: Botometer, Twitter Network Data	Provides an effective approach for detecting social spam bots, identifying spam users and reduce spam content spreading	Particle significant computational resources and detection. It reliance on deep reinforcement learning.	Twitter
6	[6]	Community Detection Algorithms like Girvan-Newman	Social Media Data Source: BuzzFeedNews, PolitiFact Fake News Dataset	Detecting an organized groups participating in fake news and labeling, identification of statistical structure of a community.	May require large amounts of data computational resources and its reliance on a fake news classifier may introduce the limitations.	Social Media
7	[7]	Generative Adversarial Networks(GANs)	Synthetic Profile Data Source: MovieLens	Provides a new approach for generating the fake profile and compare to the attack	Used for only malicious purposes, such as generating fake profile to manipulate	Recommender

			Dataset, Yelp Dataset	performance to maintaining high plausibility.	recommender systems	
8	[8]	Sentiment analysis using NLP models like LSTM	Online Product Review Data Source: Amazon Review Data, Yelp Dataset	The High Performance and effectively combines to achieve the mean average precision.	Multi module approach and Implementation. A optimal Performance for each module	Product Recommendation
9	[9]	Supervised machine learning system to detect user credibility.	Cross-Platform User Data Source: Twitter User Dataset (Credibility Corpus)	Improved Accuracy and lead to better performance in the part of classification problem and effective different data set.	Relies on the complexity in the implementation and require parameter optimal performance.	X-Platform
10	[10]	Clustering algorithms like K-means	User Behaviour Data Source: Social Media	Improve to provide understanding of user class intent. Accurate classification.	Involves the multiple steps to clustering and interaction of the data subject bias limitations.	Online User
11	[11]	Deep learning models like LSTMs	Social Media Data Source: Twitter Bot Challenge Dataset	Provides comprehensive analysis and effective detection. Identifies malicious user from both an automatic and human point of view	Relies on the Quality and their limitations. That introduce complexity the implementation and require tuning of parameters	Social Media
12	[12]	Convolutional Neural Networks (CNNs)	Deepfake Dataset Source: eKyc-DF Dataset	High quality for fake news and diverse terms. Develop the evaluating eKYC System. Evaluation of toolkit and trained models	Limited Generalizability and deepfake of an the dataset concerns	EKYC

13	[13]	Decision Trees for bot detection	Twitter Data Source: Credibility Dataset from Twitter, Botometer Data	Improve Credibility Analysis. High Accuracy(77%). Real Time Analysis	Bot Does not Provide any other Language(bot as English and Spanish)	Twitter
14	[14]	Support Vector Machines (SVMs) or ensemble methods like XGBoost.	Synthetic Data Source: Human-Bot Interaction Dataset, Twitter API	This method builds on existing research on detecting fake account created by bots or computer. Simple and Accessible features.	Not effective in detecting fake identities. Still this method is in early stages and not suited for further validation	Bots such as Social bots, Spam bots, Chatbots, Transaction bots, Gaming bots, Scraper bots, Impersonator bots, Informational bots, Malicious bots, Click bots.
15	[15]	Network-based propagation algorithms such as Label Propagation.	Hashtag Data Source: Twitter	High Accuracy(Push trend frame work 0.09%). for large datasets.	Designed only for specific dataset	Hashtags platform such as Linkedin,youtube Instagram, Twitter etc

The table presents a list of machine learning algorithms used for fake profile detection across various datasets and applications. CNNs enhance accuracy and recall, but their complexity and computational demands pose challenges. GANs simulate real users, but risk biases. Support Vector Machines and Logistic Regression are simple but effective. Deep Q-Learning and community detection algorithms detect social bots but require large datasets. Clustering algorithms and deep learning models face challenges like data bias and the need for parameter tuning. To mitigate these challenges, the study suggests using ensemble techniques to improve the differentiation between fake and genuine profiles.

5. Proposed Work

Ensemble model methodology is integrating various machine learning algorithms in an ensemble for improving the precision and reliability of the identification of fraudulent profiles within the e-learning platform. With the merits of different classifiers, such as decision trees, support vector machines, and neural networks, ensemble techniques like Random Forest,

Gradient Boosting, and AdaBoost can detect subtle patterns and anomalies in user behavior with remarkable accuracy. These models will examine aspects such as unusual login patterns, irregular interaction metrics, and contradictory profile details for suspect profiles. Ensemble techniques overcome challenges like data imbalance and feature relevance because the aggregated predictions of the individual models help reduce the chance of overfitting while increasing generalization. In addition, their ability to scale with the size of data typical for an e-learning platform is essential. This methodology provides a reliable and highly efficient solution to safeguard the integrity of the platform and enhance trust and secure interactions between users. The Figure.4 illustrates the general workflow of the proposed.

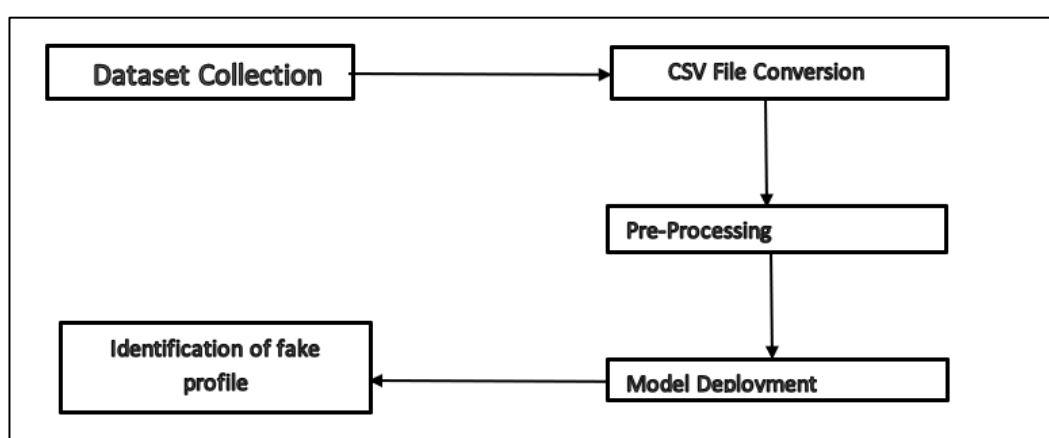


Figure 4. Workflow of Proposed Work

6. Discussion

The survey on fake-profile detection across social networks and e-learning platforms underlines the importance of machine learning and deep learning for overcoming spurious-profile-related challenges. The study reflect the increasingly sophisticated methodologies, which integrate heterogeneous social media data, content analysis, and network properties to provide enhanced detection accuracy. A very significant advance in this direction is the hybrid integration of deep reinforcement learning and particle swarm optimization to detect spam bots. Hybrid approaches promise much to specifically counter certain threats, like spam-influential users on social media platforms such as Twitter. In addition, weaknesses in advanced recommender systems have been unveiled by profile-based attacks, emphasizing the need for more robust systems. The exploration of fake news classification and community dynamics through complex network-inspired models emphasizes the efficacy of utilizing network theory for higher accuracy in fake information detection. Contributions like eKYC-DF datasets and

tools like CrediBot are instrumental in addressing deepfake threats and bot behavior, respectively, while ensuring regulatory compliance and credibility analysis. The application of supervised machine learning for feature weights and the scalability of deep learning-based detection algorithms really highlight their utility in practical applications. However, there remains much that is necessary for real-time adaptability in interoperability across various platforms. Therefore, research directions are focused on developing resilient and scalable solutions to protect digital platforms from fake profiles and improve defenses within an increasingly interconnected environment. These findings provide a good basis for future research, focusing on more secure and reliable e-platforms.

6.1 Future Work

Future work will involve collecting datasets from e-learning platforms to detect fake profiles, aiming to enhance security and authenticity. Pre-processing techniques, such as handling missing values, normalizing features, and removing anomalies, will be applied to prepare the data. Feature extraction will focus on identifying key attributes indicative of fake profiles. Advanced methods, including Natural Language Processing (NLP) for analyzing text-based data and ensemble techniques like Random Forests and Gradient Boosting Machines, will be explored to improve model performance and accuracy.

7. Conclusion

In summary, the issue of detecting fraudulent profiles on e-learning platforms like LinkedIn is very crucial in order to preserve user confidence, participation, and the overall integrity of the platform. It is possible to develop a reliable automated mechanism to distinguish between authentic users and deceptive accounts by applying machine learning methodologies to examine user activity, profile thoroughness, and connection dynamics. Training algorithms on labelled datasets will help the model to recognize anomalies in user behaviour, such as unusual connection requests and inconsistencies in profile information, thus improving the accuracy of fraud detection. Effective implementation is also necessary to improve the quality of user interaction on the platform while providing a safer and more reliable learning environment for all participants. In conclusion, the survey on the existing methods toward fake profiles identification in e-learning environment, provides a good basis for future research, focusing on more secure and reliable e-platforms.

References

- [1] Aditya, Bhrugumalla LVS, and Sachi Nandan Mohanty. "Heterogenous Social Media Analysis for Efficient Deep Learning Fake-Profile Identification." *IEEE Access* (2023).
- [2] Lin, Chen, Si Chen, Meifang Zeng, Sheng Zhang, Min Gao, and Hui Li. "Shilling black-box recommender systems by learning to generate fake user profiles." *IEEE Transactions on Neural Networks and Learning Systems* 35, no. 1 (2022): 1305-1319.
- [3] Roy, Pradeep Kumar, and Shivam Chahar. "Fake profile detection on social networking websites: a comprehensive review." *IEEE Transactions on Artificial Intelligence* 1, no. 3 (2020): 271-285.
- [4] Masood, Faiza, Ahmad Almogren, Assad Abbas, Hasan Ali Khattak, Ikram Ud Din, Mohsen Guizani, and Mansour Zuair. "Spammer detection and fake user identification on social networks." *IEEE Access* 7 (2019): 68140-68152.
- [5] Lingam, Greeshma, Rashmi Ranjan Rout, Durvasula VLN Somayajulu, and Soumya K. Ghosh. "Particle swarm optimization on deep reinforcement learning for detecting social spam bots and spam-influential users in twitter network." *IEEE Systems Journal* 15, no. 2 (2020): 2281-2292.
- [6] Amira, Abdelouahab, Abdelouahid Derhab, Samir Hadjar, Mustapha Merazka, Md Golam Rabiul Alam, and Mohammad Mehedi Hassan. "Detection and analysis of fake news users' communities in social media." *IEEE Transactions on Computational Social Systems* (2023).
- [7] Zhang, Xuxin, Jian Chen, Rui Zhang, Chen Wang, and Ling Liu. "Attacking recommender systems with plausible profile." *IEEE Transactions on Information Forensics and Security* 16 (2021): 4788-4800.
- [8] Hu, Shigang, Akshi Kumar, Fadi Al-Turjman, Shivam Gupta, and Simran Seth. "Reviewer credibility and sentiment analysis based user profile modelling for online product recommendation." *Ieee Access* 8 (2020): 26172-26189.
- [9]-Althaqafi, Nahid R., Hessah A. Alsalamah, and Walaa N. Ismail. "The Impact of the Weighted Features on the Accuracy of X-Platform's User Credibility Detection Using Supervised Machine Learning." *IEEE Access* (2024).

- [10] Kanwal, Moona, Najeed A. Khan, Najma Ismat, Aftab A. Khan, and Muzammil Ahmad Khan. "Machine Learning Approach to Classification of Online Users by Exploiting Information Seeking Behavior." *IEEE Access* (2024).
- [11] Sansonetti, Giuseppe, Fabio Gasparetti, Giuseppe D'aniello, and Alessandro Micarelli. "Unreliable users detection in social media: Deep learning techniques for automatic detection." *IEEE Access* 8 (2020): 213154-213167.
- [12] Felouat, Hichem, Huy H. Nguyen, Trung-Nghia Le, Junichi Yamagishi, and Isao Echizen. "eKYC-DF: A Large-Scale Deepfake Dataset for Developing and Evaluating eKYC Systems." *IEEE Access* (2024).
- [13] Aguilera, Ana, Pamela Quinteros, Irvin Dongo, and Yudith Cardinale. "CrediBot: Applying Bot Detection for Credibility Analysis on Twitter." *IEEE Access* (2023).
- [14] Van Der Walt, Estée, and Jan Eloff. "Using machine learning to detect fake identities: bots vs humans." *IEEE access* 6 (2018): 6540-6549.
- [15] Kausar, Soufia, Bilal Tahir, and Muhammad Amir Mehmood. "Push-to-trend: A novel framework to detect trend promoters in trending hashtags." *IEEE Access* 10 (2022): 113005-113017.
- [16] Fake-Profile-Detection-using-ML/pdf at master · harshitkgupta/Fake-Profile-Detection-using-ML · GitHub