

Identification of Fake ACCOUNT on Social Media using various Machine Learning Algorithm

Nita Balasaheb Kale, Prof. Anita Mahajan, Dr. Pankaj Agarkar,

PG Student: Department of Computer Engineering, Ajeenkya DY Patil School of Engineering Pune

Prof.: Department of Computer Engineering, Ajeenkya DY Patil School of Engineering Pune

HOD: Department of Computer Engineering, Ajeenkya DY Patil School of Engineering Pune

ABSTRACT: Currently, online social media platforms serve as the most prevalent and rapid means of exchanging information in the realm of technology. Individuals from diverse backgrounds predominantly devote a significant amount of their time to engaging with social networking sites. A vast quantity of information is generated and disseminated globally via social networks. These incentives have led to unauthorized individuals carrying out hostile actions against users of the social site. The creation of false accounts on social media is seen as more detrimental than any other kind of cybercrime. The detection of this infraction must occur prior to informing the customer about the creation of the fraudulent identity. Several algorithms and methodologies have been proposed for detecting fraudulent identities, mostly leveraging the huge volume of unprocessed data generated by social networks. This study aims to identify fraudulent identities on social media accounts using a dataset from Twitter. Diverse machine learning methods have been used to assess the suggested outcomes using natural language processing approaches. SVM, Fuzzy Random Forest, and Naïve Bayes have been used for classification purposes. The empirical study demonstrates the efficacy of the system and its superior accuracy compared to other machine learning algorithms and preexisting systems.

Keywords: Machine Learning, Naïve Bayes, Fuzzy Logic, Random Forest, twitter dataset, fake identity, bots, Natural Language processing, classification.

1.INTRODUCTION

Social media platforms have become an integral part of our lives, including all aspects of our daily social activities. People in specific social positions, such as those in the media, tend to prioritize their actions based on what they do in real life. Individuals possess behaviors, perceptions, actions, and habits inside these systems. (The cultural aspects of social media) As the importance of such media increases, it becomes increasingly crucial to study and evaluate the idea of Comprise. The study addresses the problem by specifically examining regional cultural characteristics with intellectual rigor. Google app inconsistencies are prevalent across the most popular social networking platforms. The prominence of social networking has rapidly escalated in recent times, making it crucial for advertising campaigns and celebrities seeking to enhance their online presence by augmenting their follower and fan base. However, counterfeit profiles, seemingly created on behalf of institutions or people, have the potential to undermine their reputation and diminish their count of followers and interactions. They continue to have false alerts and excessive vagueness with others. The presence of counterfeit accounts across many platforms has detrimental consequences that undermine the marketing and promotional capabilities of social media for businesses, while also establishing a foundation for online abuse. Within the realm of the internet, individuals possess particular inquiries about the safeguarding of their personal information.

Several social media platforms include Twitter, Google+, Youtube, Instagram, Flickr, Facebook, and Snapchat. A total of 823 million persons used social media on their cellphones daily, marking a significant increase over the previous fiscal quarter's 654 million users. Social networking platforms like Facebook now lack the capability to provide real-time updates for fake accounts. As a result, it is difficult for somewhat tech-savvy users to distinguish between genuine and fraudulent accounts. When dealing with large amounts of data, it is important to handle additional significant challenges such as data gathering, managing data streams, and delivering immediate user responses. These challenges are crucial for achieving dependable profile recognition performance. Prior research on fraudulent accounts focuses on conducting experiments to assess the effectiveness of preventative measures against deceptive user behavior patterns. A research project on social manipulation on Facebook utilizing the Google Maps API to analyze quantifiable information regarding the gender distribution of friends, access to data documents, clustering algorithms for mutual acquaintances, details about education and work, location knowledge of Facebook users, and shared interests. The security measures used to protect consumers from attackers include a comprehensive understanding of data, privacy regulations, strategies to bolster safety, and educational programs aimed at increasing awareness.

This effort will also address the problem posed by the increasing volume of data that is being generated at a rapid rate and covers a wide variety of topics. Figure 1 below illustrates the activities involved in detecting both genuine and counterfeit profiles. The vast amount of data gathered by social networks may be effectively used to distinguish between fake and genuine accounts, hence suggesting questions only from verified profiles. This is crucial for customers who are not experts in the field, as well as teens and children who may not be familiar with the safety measures. The study should focus on assisting users in promptly distinguishing between authentic and counterfeit profiles, as well as providing guidance on whether profiles should be reported and advising users on establishing connections with trustworthy partners.

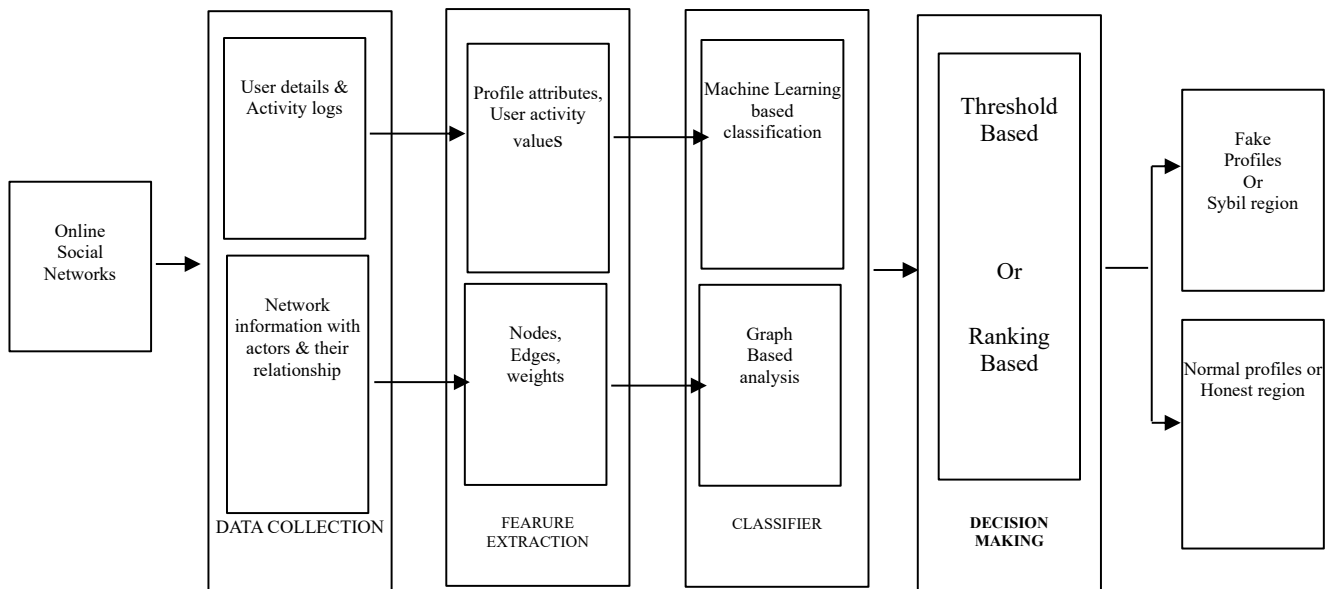


Figure 1 : General process flow in detection of the real or fake profiles on social media.

2.LITERATURE SURVEY

Drouin, Michelle, et al. [1] method primarily emphasizes the analysis of all public online activities. Individuals encounter significant levels of insincerity across various online platforms, such as social networking sites, online social websites, private chat rooms, and dating websites. They assert that people exhibit less honesty than expected on these platforms and are able to identify specific types of deception across multiple sites. When comparing several factors that predict integrity in various situations, it was shown that expectations about others' dishonest conduct were more significant than particular characteristics. This solution addresses these issues by using machine learning techniques and eliminating such identities..

L. M. Jupe et. al. [2] possible rationales for identification manipulation, methodology for verification, and manipulation of annexation. This study examined individuals who were instructed to deceive and tell the truth about their intrusion in three different scenarios: voluntary deception, coercive deception, and truth. The Verifiability Technique was used to assess the recordings, specifically outlining the material that has to be confirmed in the defendant's statements. Although there was speculation that individuals who lie would provide less verifiable specifics in their spoken remarks compared to those who tell the truth, the formal verification approach was unable to distinguish between factual and deceptive claims.

Y. Li, O. et. al. [3] the proposed system developments focus on two aspects: computationally, machine evolutionary techniques are utilized in a hierarchical manner that is both efficient and easily parallelizable. Furthermore, the method architecture is highly versatile and can be applied to various behavioral cluster-based problems and implementations with minimal modifications required.

T. Tuna et. al. [4] the system examines the similarities and differences between spammers and hackers, and then explores the possibility of speedier access and/or contextual attention to identify hackers more effectively, leading to improved results.

Commence the examination of the phenomenon of boundary-spanning creeping, wherein an individual may assume the role of a long-term reader within a friend's social circle due to the presence of weak boundaries, while also actively engaging in a specific social circle. Develop an analysis of connections and employ website crawling techniques to gather Twitter posts, with the aim of streamlining spam detection and enhancing accuracy. Additionally, endeavor to identify the genuine individuals responsible for creating bot identities.

P. Galan-Garcia et. al. [5] the current issue of user identity anonymity creates situations where individuals with false or unrelated identities regularly post articles, comments, or multimedia content in an attempt to mock or target certain individuals who may be unaware of the attack. These actions may significantly impact the offenders' reality, leading to situations where cyber attacks escalate into severe effects in real life. The current technique aims to identify and link counterfeit Twitter accounts that are used for deceitful activities to genuine attributes, even within the same network. This is achieved by evaluating the textual content of the comment threads generated by both profiles. The approach, under supervision, has been used to ascertain these profiles..

K. et. al. [6] Self-presentation refers to the way individuals present themselves to others, including their appearance, behavior, and communication style. Examinations of the primary five facets of the element in psychopathology. The curriculum delineates two approaches: (1) engaging in online deception to manipulate others and (2) seeking meaningful and long-lasting online connections. In addition to its primary objective, the program endeavored to construct scales for quantifying these structures, despite the absence of any already applicable metrics. Initially, focus on acquiring the knowledge and skills to effectively exploit the diverse range of human personality traits and mental disorders in order to predict deceitful behavior and online communication. The proposed strategy is expected to achieve an approximate accuracy rate of 85-88%, with a small percentage of unfavorable outcomes.

Thi Thu Thuong Le et. al. [7] Power disaggregation, also known as energy disaggregation, Conducted to validate the efficacy of the suggested system as an optimal alternative for minimizing our power use. Several algorithms are associated with this domain of machine learning. The categorization results from all those methods are not as robust as anticipated. In this research, we propose a novel method for developing an energy segmentation classifier using deep learning techniques. We use the Gated Recurrent Unit (GRU), which is based on the Recurrent Neural Network (RNN), to train the model using the UK DALE dataset in this domain. In addition, we compare our approach to the necessary energy propagation recurrent neural network (RNN).

Yong Zhang et. al. [8] A architecture called Comprehensive Attention with Recurrent Neural Networks (CA-RNN) is developed. It is capable of capturing previous, subsequent, and local characteristics of every sequential position. The Bidirectional Recurrent Neural Networks (BRNN) are used to include information from potential future events, while a convolutional layer is utilized to capture spatial characteristics. The traditional RNN is enhanced by two emergent modifications, namely LSTM and gated recurrent unit, to greatly enhance the effectiveness of the new architectural design. Another notable feature of the new model is its fully automated design, which can be prepared from start to finish without any human intervention.

Peddintiet. al. [9] developed a classifier that simplifies the four-class classification procedure into two binary classifications. The first classification determines if each identity is anonymous or particularly non-anonymous, while the second classification determines whether each account is recognized or quasi-identifiable. The classification of Twitter account information as 'anonymous,' 'identified', or 'unknown' is determined by combining the values of two classifiers. The majority of binary classifiers use Random Forest (RF) as a classification technique, with 100 tree branches. The choice of the classifier and the number of trees relies on the outcomes of the validation set and the bag failure. These classifiers are cost-sensitive meta classifiers that impose a greater penalty for misclassifying instances as anonymous or identifiable. The dataset used in this context originated from the social media platform, Twitter.

Philogene Kyle Dimpas et. al. [10] The detection of clickbait in both Filipino and foreign languages is achieved by the use of a recurrent long-term neural network with recurrent memory. This effort has collected Filipino and English headlines and determines if they are clickbait. The model used a BiLSTM-based neural network. The model utilizes Word2Vec to provide word representation and embedding for corpora. The experimental results demonstrated an accuracy rate of 91.5% when using the usual method..

Rajesh Purohit Bharat Sampatrao Borkar [11] examined the distinction between false detection and genuine social media identification by using Random Forest and Deep Convolutional Neural Network. This approach effectively mitigates the issue of false identity by bots in the classification process. Its main objective is to identify and detect instances of fraudulent human identification, since there has been little research conducted on this specific kind of deception. The search process employs two distinct algorithms, namely Random Forest (RF) and Neural Recurrent Network, to analyze the description..

B.Pandu Ranga Raju, B.Vijaya Lakshmi, C.V. Lakshmi Narayana [12] present a very effective and adaptable method for identifying fraudulent URLs, equipped with a comprehensive set of features that accurately capture the many attributes of phishing websites and their hosting platforms, including elements that are hard to counterfeit. The software leverages the Random Forests technique to achieve a combination of high detection capacity and low error levels. The experimental findings demonstrate that the software has the capability to enable the blacklist provider in generating automatic blacklists.

Roobaee Alroobaee [13] In order to optimize classification accuracy, we employ a diverse range of Machine Learning (ML) algorithms, including 'Support Vector Machines (SVM)', 'Random Forest (RF)', 'Multilayer Perceptron (MP)', 'Decision trees (DT)', 'Naïve Bayes (NB)', k-Nearest Neighbors (KNN), and the Deep Learning technique 'Long Short-Term Memory (LSTM)'. The findings indicate that the use of machine learning algorithms yielded distinct outcomes for each data set, namely age and gender. Deep learning algorithms have shown their utility in handling massive data sets.

3.RESEARCH METHODOLOGY

Social media platforms exert effect on several aspects such as science, information dissemination, public mobilization, employment prospects, and corporate operations. Scientists analyzed various social media platforms to assess their impact on individuals. Users may facilitate communication by creating a welcoming atmosphere for students to engage in academic activities. As a result, instructors worldwide can get acquainted with these platforms, enabling them to establish online classroom pages, distribute homework, facilitate discussions, and more. Many social enterprises may use these social media platforms to recruit talented and engaging people, therefore streamlining long-term research efforts. This initiative seeks to incorporate a mechanism for the automated categorization of fraudulent accounts with the purpose of safeguarding individuals' social identity. By utilizing this automated recognition strategy, websites can more effectively manage the vast number of personalities that can be handled automatically. For instance, if we want to assess inaccurate attributes according to their duration, date of publication of tweets, language, and geographical location. These points are dependent on each variation and component, and we will now outline the comprehensive study process, including the use of learning algorithms.

3.1 Tokenization

Tokenization is the process of dividing textual data into individual units such as words, phrases, symbols, or punctuation marks. The purpose of tokenization is to recognize the phrases inside an utterance. The sequence of tokens serves as data for further analysis, similar to the processes of sorting or sequential information mining. Tokenization is a key component of text summarization in both language studies and scientific inquiry. At the beginning, linguistic knowledge is a fundamental building component. Comprehensive expertise Data extraction strategies include the methodologies used for gathering information. Record tokenization is a precondition for the operation of a parser. This speech may be considered inconsequential, since the text is already encoded in comprehensible formats for mobile devices. However, there are still some unresolved difficulties, such as the use of exclamation points. Certain characters, such as brackets and hyphens, often need validation.

3.2 Stop Word Removal

Skip phrases are often used more extensively than conventional terms such as 'and,' 'are,' 'this,' etc. We may seem insignificant in the realm of document categorization. However, they must be eliminated. Moreover, the process of creating these recordings of stop phrases is intricate and inconsistent among written sources. This procedure simultaneously reduces the complexity of the text and improves the effectiveness of the approach. The textual content report gives information on phrases that are not crucial for text analysis applications.

3.3 Stemming and Lemmatization

The objective of both stemmed and character segmentation is to reduce the many types of inflection and derived variants to a standardized base form. Stemming typically comprises a simple heuristic procedure that truncates the endings of words in order to achieve a desired outcome more often than not. This process commonly entails deleting derived pronunciation. Stemming is the process of effectively analyzing words by studying their language and structure. It involves removing the intonation endings and determining the root or vocabulary form of a phrase, which is also known as technology. Natural Language Processing Preprocessing Text preprocessing is a crucial component of any NLP technique, and its importance is derived from the aforementioned procedures. This suggested system utilizes three distinct machine learning algorithms: Support Vector Machine (SVM), Fuzzy Random Forest (FRF), and Naïve Bayes (NB).

3.4 Support Vector Machine (SVM)

An SVM categorizes information by identifying the exceptional hyperplane that separates all elements of one category from those of another categorization. In the context of Support Vector Machines (SVM), the optimal hyperplane is determined by the longest distance between the two distinct groups. An SVM classifies data by identifying the exceptional hyperplane that separates certain characteristics of awareness in one group from those in the other. The support vectors are the data points that are closest to the separating hyperplane.

3.5 Naive Bayes

The Naive Bayes algorithm is a learning technique that calculates the likelihood of a certain class or category contributing to an item with specified attributes. To put it simply, the classifier is a deterministic one. The Naive Bayes approach is referred to as "naive" because it makes the assumption that the frequency of a certain attribute is independent of the occurrence of specific conditions. For instance, if we want to assess inaccurate characteristics according to their duration, date of publication, social media posts, language, and geographical location. These elements are interdependent and contingent upon several conditions. However, in my view, all these qualities contribute to the likelihood of automatically eliminating the wrong profile.

4. PROPOSED SYSTEM DESIGN

Several current methodologies have been developed to identify malicious activity or bots, as described in references [5] and [6]. However, these systems still encounter challenges such as a high false alarm rate and poor classification accuracy. Figure 2 illustrates the proposed method, which utilizes datasets from several social network websites to identify fake accounts throughout the full dataset.

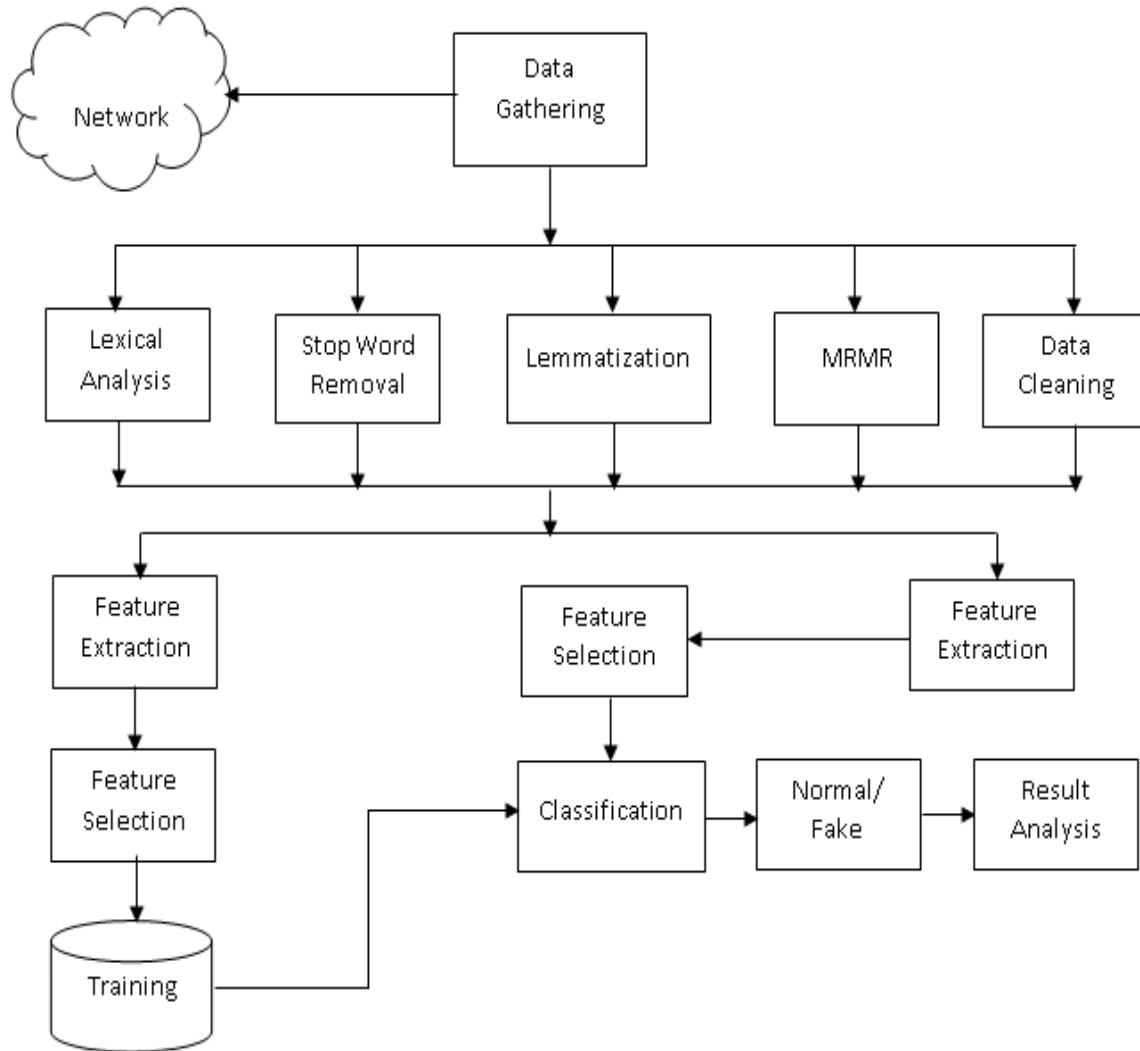


Figure 2: Proposed System Architecture Design

At first, the system retrieves data from the Twitter account using the Twitter API, which pulls the data from recent tweet comments read by users. The primary issue with social media programs is in their inability to effectively identify and distinguish between automated bots and fraudulent user accounts. We conducted this study by integrating NLP and Machine Learning algorithms to address the shortcomings present in current systems.

Initially, we collect data from many social media platforms. Once the data is obtained, it is then saved in both data repositories and data set files. The data has been gathered from several online sources such as Twitter, which may result in its lack of organization at times. Preprocessing this kind of data necessitates the use of a particular sampling strategy and data filtering tools. The systematic sampling approach has been used for data segregation, while the bloom filter has been utilized to eradicate misclassified cases. The electrical analysis should include both sentence recognition and tokenization. Tokenizing words involves storing them in a string array, which allows for convenient string verification. Stop word removal and lemmatization are two further NLP methods that fall within the domain of natural language processing. Following the cleaning procedure, the system employs several feature extraction approaches. Another approach used to provide similar capability is co-occurrence correlation, in addition to TF-IDF. The election has been conducted using several quality levels, and these attributes are then sent to the classification system. The same procedure has been used for both data training and testing, respectively. We used three distinct machine learning techniques, namely Fuzzy Random Forest, Naive Bayes, and support vector machine, in a sequential manner.

5. ALGORITHM DESIGN

5.1 Training using updated NB

Input: Training dataset Train Data[], Various activation functions[], Threshold Th

Output: Extracted Features Feature_set[] for completed trained module.

Step 1: Set input block of data d[], activation function, epoch size,

Step 2 : Features.pkl $\leftarrow$ ExtractFeatures(d[])

Step 3 : Feature_set[] $\leftarrow$ optimized(Features.pkl)

Step 4 : Return Feature_set[]

5.2 Testing using updated NB

Input: Training dataset TestDBLits [], Train dataset TrainDBLits[] and Threshold Th.

Output: Resultset < class_name, Similarity_Weight > all set which weight is greater than Th.

Step 1: For each testing records as given below equation

$$testFeature(k) = \sum_{m=1}^n (featureSetA[i] \dots A[n] \leftarrow TestDBLits)$$

Step 2 : Create feature vector from testFeature(m) using below function.

$$ExtractedFeatureSet(x)[t \dots n] = \sum_{x=1}^n (t \leftarrow (testFeature_{(k)}))$$

ExtractedFeatureSetx[t] holds the extracted feature of each instance for testing dataset.

Step 3: For each train instances as using below function

$$trainFeature(l) = \sum_{m=1}^n (featureSetA[i] \dots A[n] \leftarrow TrainDBLits)$$

Step 4 : Generate new feature vector from trainFeature(m) using below function

$$ExtractedFeatureSet(y)[t \dots n] = \sum_{x=1}^n (t \leftarrow (trainFeature_{(l)}))$$

Extracted_FeatureSet_Y[t] holds the extracted feature of each instance for training dataset.

Step 5 : Now evaluate each test records with entire training dataset

$$weight = calcSim(featureSetx) \parallel \sum_{x=1}^n FeatureSety[y]$$

Step 6 : Return Weight

6 RESULTS AND DISCUSSION

The implementation was carried out in a Windows open-source environment, using the Java Platform owing to its open-source availability. The data from the Twitter online application has been extracted using the freely accessible Twitter API. We generate diverse data segments to evaluate the classification precision of the system using three distinct machine learning techniques. Three distinct cross-validation procedures have been used for data splitting, namely 10-fold, 15-fold, and 20-fold.

6.1 Experiments using Naive Bayes

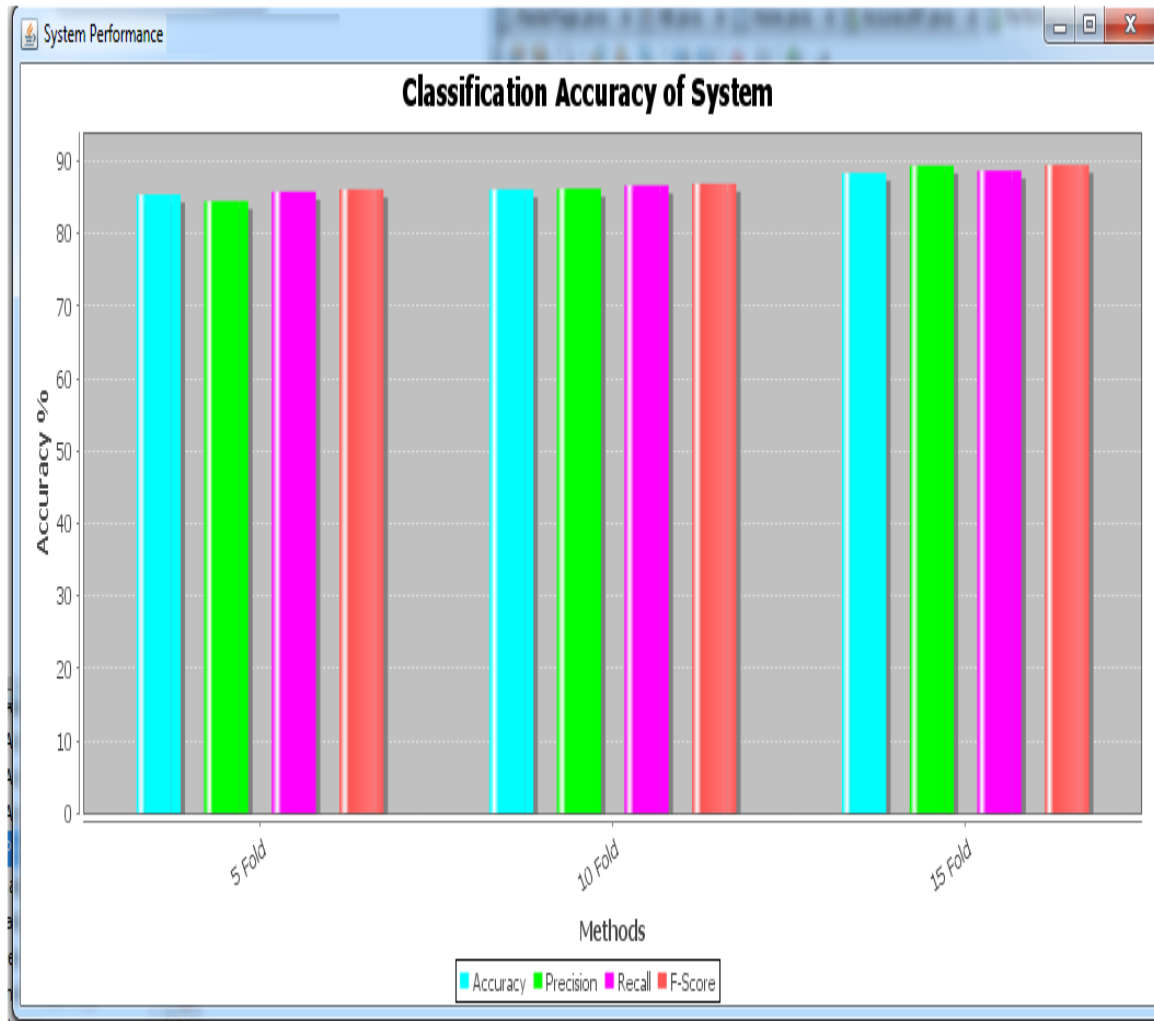


Figure 3 : Classification accuracy using Naïve Bayes with different cross validation

The picture above, labeled as picture 3, displays the classification accuracy of Naïve Bayes when applied to the twitter dataset. Similar tests were conducted using other cross-validation techniques, and the corresponding results are included in Table 1. Based on this data, we can deduce that 15-fold cross-validation yields the maximum classification accuracy of 89.40% for the Naïve Bayes model.

Table 1 : Classification accuracy with confusion matrix for Naïve Bayes

(NB)	5-Fold	10-Fold	15-Fold
Accuracy	85.40	86.10	89.40
Precision	84.40	86.50	89.40
Recall	85.70	86.65	89.70
F-Score	86.50	86.90	89.50

6.2 Experiment using Fuzzy Random Forest

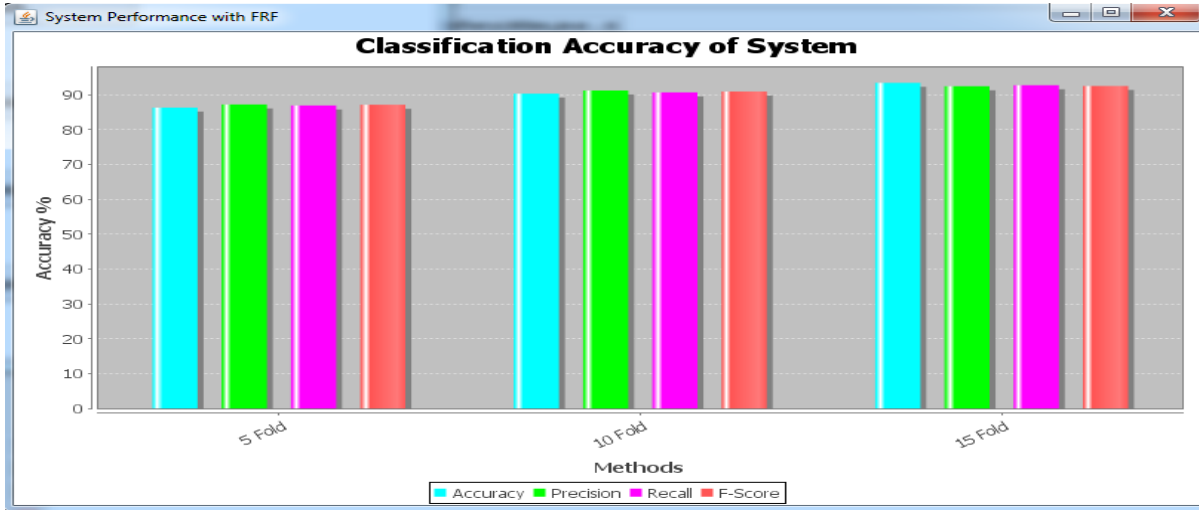


Figure 4: Classification accuracy using FRF with different cross validation

The picture above, labeled as picture 4, displays the classification accuracy of FRF using the twitter dataset. Similar tests were conducted with different cross-validations, and the corresponding results are included in Table 2. Based on this data, we can infer that 15-fold cross-validation yields the maximum classification accuracy of 90.60% for FRF.

Table 2 : Classification accuracy with confusion matrix for FRF

(FRF)	5-Fold	10-Fold	15-Fold
Accuracy	86.30	90.30	90.60
Precision	87.20	91.00	91.20
Recall	86.90	90.25	91.00
F-Score	87.10	90.80	91.60

6.3 Experiment using Support Vector Machine

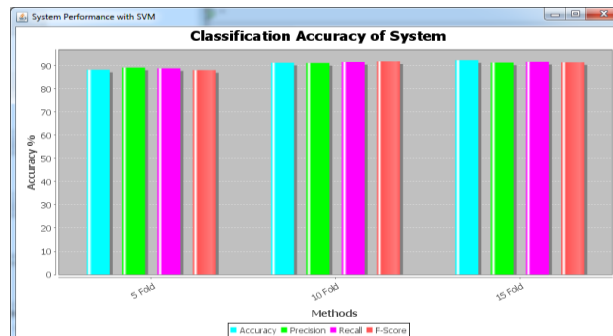


Figure 5 : Classification accuracy using SVM with different cross validation

The picture above, labeled as picture 5, displays the classification accuracy of Support Vector Machines (SVM) when applied to the twitter dataset. Similar tests were conducted using different cross-validation techniques, and the corresponding results are included in Table 3. Based on this data, we can deduce that 15-fold cross-validation yields the maximum classification accuracy of 91.40% for SVM.

Table 3 : Classification accuracy with confusion matrix for SVM

(SVM)	5-Fold	10-Fold	15-Fold
Accuracy	88.80	91.10	92.40
Precision	89.00	91.15	91.40
Recall	89.30	91.20	91.70
F-Score	88.70	91.60	91.50

The Weka tool utilizes the all classification method, specifically in the 3.8 Weka environment, for the implementation of machine learning algorithms. In due course, a cooperative endeavor and the comprehensive examination of many platforms are used to analyze the interrelatedness of social media and the substantial volume of data inside the social domain. The current study focuses on the practical uses of different types of assaults, with less attention given to this topic in previous studies. Data science is a rapidly developing discipline that is projected to become the greatest area of study in the next decade. The primary social data consists of large-scale data streams sent over the Internet. The use of big data analytics and analysis is beneficial across several domains.

7.CONCLUSION

This research article provides an overview of the most often used approaches for detecting Sybil or fraudulent accounts and media networks, and offers a critical analysis of the current level of research in this area. The various methodologies are compared based on their synthesized networking device and data set parameters. We have also focused on the newly offered solutions and analyzed their merits and downsides. Analysis of the success of these systems is the main topic of comparison. The issues that can be easily addressed are documented in the realm of fraudulent profile identification in digital social networks. Currently, despite the existence of several methods, there is still no all-encompassing solution for swiftly and easily identifying phony identities in social networks while ensuring the recognition of user information. Based on our experimental study, we can infer that Support Vector Machines (SVM) outperforms Naive Bayes (NB) and Random Forest (FRF) in all cross-validations. Machine efficiency in the future may be improved by using methods such as deep learning, which involves the use of specialized activation functions.

REFERENCES

1. Drouin, Michelle, Daniel Miller, Shaun MJ Wehle, and Elisa Hernandez. **Why do people lie online? “Because everyone lies on the internet,** Computers in Human Behavior 64 (2016): 134-142.
2. Jupe, Louise Marie, Aldert Vrij, Galit Nahari, Sharon Leal, and Samantha Ann Mann. **The lies we live: Using the verifiability approach to detect lying about occupation.,** Journal of Articles in Support of the Null Hypothesis 13, no. 1 (2016): 1-13.
3. Li, Yixuan, Oscar Martinez, Xing Chen, Yi Li, and John E. Hopcroft. **In a world that counts: Clustering and detecting fake social engagement at scale.,** In Proceedings of the 25th International Conference on World Wide Web, pp. 111-120. International World Wide Web Conferences Steering Committee, 2016.
4. Tuna, Tayfun, Esra Akbas, Ahmet Aksoy, Muhammed Abdullah Canbaz, Umit Karabiyik, Bilal Gonen, and Ramazan Aygun., **User characterization for online social networks. ,** Social Network Analysis and Mining 6, no. 1 (2016): 104.
5. Galan-Garcia, Patxi, Jose Gaviria de la Puerta, Carlos Laorden Gomez, Igor Santos, and Pablo García Bringas. **Supervised machine learning for the detection of troll profiles in twitter social network: Application to a real case of cyberbullying.** Logic Journal of the IGPL 24, no. 1 (2016): 42-53.
6. Stanton, Kasey, Stephanie Ellickson-Larew, and David Watson. **Development and validation of a measure of online deception and intimacy.,** Personality and Individual Differences 88 (2016): 187-196.
7. Kim, Jihyun, and Howon Kim., **Classification performance using gated recurrent unit recurrent neural network on energy disaggregation.,** In 2016 International Conference on Machine Learning and Cybernetics (ICMLC), vol. 1, pp. 105-110. IEEE, 2016.
8. Zhang, Yong, Meng Joo Er, Rajasekar Venkatesan, Ning Wang, and Mahardhika Pratama. **Sentiment classification using comprehensive attention recurrent models.,** In 2016 International joint conference on neural networks (IJCNN), pp. 1562-1569. IEEE, 2016.

9. Peddinti, Sai Teja, Keith W. Ross, and Justin Cappos. **Mining Anonymity: Identifying Sensitive Accounts on Twitter.**, arXiv preprint arXiv:1702.00164 (2017).
10. Dimpas, Philogene Kyle, Royce Vincent Po, and Mary Jane Sabellano. **Filipino and english clickbait detection using a long short-term memory recurrent neural network.**, In 2017 International Conference on Asian Language Processing (IALP), pp. 276-280. IEEE, 2017.
11. Rajesh Purohit Bharat Sampatrao Borkar , **Identification of Fake vs. Real Identities on Social Media using Random Forest and Deep Convolutional Neural Network**, in International Journal of Engineering and Advanced Technology, Issue-1 7347-7351 IJEAT 2019
12. B.Pandu Ranga Raju, B.Vijaya Lakshmi, C.V.Lakshmi Narayana, **Detection of Multi-Class Website URLs Using Machine Learning Algorithms**, *International Journal of Advanced Trends in Computer Science and Engineering*, pp. 1704-1712 ,Volume 9, No.2, 2020.
13. Roobaea Alroobaea, **An Empirical combination of Machine Learning models to Enhance author profiling performance**, *International Journal of Advanced Trends in Computer Science and Engineering*, pp. 2130-2137, Volume 9, No.2, 2020.