

A Review on Spam Detection in Social Media Networks

R.S. Gaikwad¹, Dr. B.L.Gunjal²

¹ME Student, Department of Computer Engineering, Amrutvahini COE, Sangamner, Maharashtra, India

²Associate Professor, Dept. of Computer Engg. & HOD Dept. of IT, Amrutvahini COE, Sangamner, Maharashtra, India

ABSTRACT

Day by day us Online Social Networking is increasing in large amount such as Twitter, Facebook and Myspace. People are sharing their views, post comments, making new friends on OSN. OSN is also useful in advertisement and promotion of products, services, and topics. But there are some threats in OSN. Spammers or illegitimate users affect legitimate users and also the whole network. Sometimes legitimate users have to compromise their accounts. So detecting spammers and campaign promoters is a challenging task. So in this paper we tried to study different types of spamming attacks and campaigns in Online Social networks and related work in spam detection. There may be different approaches for detecting spam such as algorithms, features, and datasets. These things can affect the spam detection Performance.

Keyword: - Spam detection, Clustering, Data stream, Online Social Networks.

1. Introduction

As social media is now a day's increasingly popular source of public information, companies, organizations and individuals are actively using social media platforms to promote their products, services, ideas and ideologies. These things can be harmful to normal users because come times it try to influence user's behaviors/opinions/decisions. Spammers post unwanted messages containing typical words of a trending topic and URLs, usually obfuscated by URL shorteners, which can lead users to completely unrelated websites. Spammer's intention may be driven by several goals, such as to spread advertise to make sales, distribute pornography, viruses, phishing, or simple just to compromise system status. These types of spam can contribute to de-value real time search services. Spam which is referred as unsolicited messages containing malicious links that directs victims to external sites containing malware downloads, phishing, drug sales, or scams. But how do we know that a message and campaigns on OSN are bad or are spam or non-spam. Because sometimes campaigns may be conducted by health organization to promote good habits and try to stay away from bad habits. Ex. government health agency may conduct an anti-smoking campaign on OSN. So detecting spam or non spam is challenging task in OSN. Other challenges like data, feature, and model which used in spam detection may affect detection accuracy. Also data may be stored as whole or it may be real time streaming data which may need different approach and online learning algorithms. So in this paper we tried to study different approaches in spam detection in OSN, the datasets which were used and features or attributes of users and spammers.

2. Literature Survey

C. Chen, J. Zhang, Y. Xie, Y Xiang, W Zhou [1] Carried out work performance evaluation of different machine learning based streaming spam detection methods which is on feature, data, model aspects. They used dataset of 600 million tweets to label 6.5 million spam tweets Trend Micro's Web Reputation system is used and 12 features to represent tweets. Spam detection carried out as problem of binary classification and machine learning algorithms are used to solve it. They evaluated the different factors impact to the spam detection performance, included training data size, feature discretization, data sampling, time-related data, spam to non-spam ratio, and machine learning algorithms. Finally they conclude that it is important to do feature Feature discretization. Increasing training data size is not beneficial and when tweets sampled continuously instead of random selection then classifier performs well.

F. Benevenuto, G. Magno, T. Rodrigues, and V Almeida [2] proposed spammer detection on twitter. They considered the problem of spammers detecting on Twitter. The large dataset of Twitter collected which includes 1.9 billion links, 54 million users, and 1.8 billion tweets. With tweets from 2009 related to three famous trending topics, a large labeled collection of users constructed manually and classified into non-spammers and spammers. they then identified a number of characteristics related to tweet content attributes and user social behavior attributes, which were used to detect spammers. 39 content attributes and 23 user behavior attributes related to tweets were used. They then used these attributes for machine learning process for classifying users as either non-spammers or spammers. A non-linear Support Vector Machine (SVM) classifier used with the Radial Basis Function (RBF) kernel that allow SVM models to perform separations with very complex boundaries. They investigated the feasibility of applying a supervised machine learning method to identify spammers.

Z. Miller, B. Dickinson, W. Deitrick, W. Hu, A. Wang [3] Proposed data stream clustering approach for Twitter spammer detection. They treated it as anomaly detection problem instead of Classification problem. For this study data set of 3239 user accounts and a sample tweet from each account included. Three feature categories were used: content, user information, and tweet text. User information from previous studies and introduced 95 one-gram features were analyzed. They proposed two algorithms for stream clustering, Den-Stream and Stream-KM++ were modified to make easy spam identification. These algorithms clustered outlier as spammer and normal user. They combined Den-Stream and Stream-KM++ for maintaining low-false positive rate and to improve performance of spam detection.

K. Thomas, C. Grier, V. Paxson, D. Song [4] Work examines the tools, techniques, and support infrastructure for abuse of online social networks. Dataset of nearly 1.8 billion tweets sent by 32.9 million Twitter accounts over collected. In this approach they examined accounts suspended by Twitter for abusive behavior. They manually verify a sample of suspended accounts and find the vast majority were suspended for spamming, providing us with a rich source of ground truth for measuring spam. To identify different classes of spammers they used blacklist. SURBL and all its affiliated lists (e.g. SpamCop, Joewein); Google Safebrowsing, both malware and phishing; and URIBL are three blacklist used. This study provides tools techniques used for spamming and highlighted those features.

K. Thomas, C. Grier, J. Ma, V. Paxson, D. Song [5] Proposed Monarch, which determines whether the URLs direct to spam by using a real-time system that crawls URLs as they are submitted to web services. They investigate the distinctions between twitter and email spam, which includes misuse of public web hosting and redirector services. Also evaluated scalability of Monarch to protect twitter services. They fed posted URLs for classification in the system. System then visit URL and collect its data which includes hosting infra, page behavior, page content. System then converts this raw data into real valued features and provides these to classifier training and live decision making. Then final decision and warnings are generated. They used spark map-reduce framework for classification using Distributed LR with L1-regularization and Stochastic gradient descent for LR (LRsgd) algorithms. Monarch also used blacklist of URLs for effectiveness. Finally they given a difference between email and twitter spam which are the the abuse of generic redirectors and public web hosting, the persistence of features over time, and overlap of spam features.

X. Jin, C. X. Lin, J. Luo, J. Han [6] proposed *SocialSpamGuard*, a data mining based online social media spam detection system for social n/w security. For real-time detection they integrate a GAD clustering algorithm to cluster in large scale with learning algorithm. To show spam activity they introduced content features and image features and social network features. The major goal was detect spam post from infected users and spammers. They fist collected historical media data and both content and social network futures were extracted. Then they build a classification model and detect spam active learning. Then make prediction and alarms to client using online active learning.

Q Cao, M. Sirivianos, X. Yang, T. Pregueiro [7] Proposed tool in the hands of OSN operators, which called as *SybilRank*. By using social rank properties it rank users according to their perceived likelihood of being fake (Sybils). They used hadoop prototype to demonstrate efficient computation. SybilRank detect users that are suspected to be Sybils after three stages. At 1st stage Propagation trust via $w = O(\log n)$ power iterations. In 2nd Stage, SybilRank ranks nodes based on their degree-normalized trust. In the last stage, SybilRank assigns portions of

fake nodes in the intervals of the ranked list. Experiments were performed on social graph dataset from different social sites.

S. Ghosh, B. Viswanath, F. Kooti [8] Examined link farming in twitter network and proposed a method to depress the spirit of the activity. They studied users who establish links to spammers and reasons to their behavior. Around 41,352 suspended accounts were examined users who connects them. They Proposed Collusionrank a ranking system that penalizes users for connecting to spammers. To fight link farming in Twitter a Pagerank-like approach is used. the initial scores towards a set of bad nodes negatively biased. The Collusionrank score of a node is computed based on the score of its followings. Thus user pushed down in ranking who follow spammers by getting negative score of higher magnitude.

H. Costa, F. Benevenuto, L. Merschmann [9] In this study they investigated tip spam detection task on Apontador which famous Brazilian LBSN system. To differentiate spam from non-spam they indentified number of attributes based on labeled collection of tip provided by Apontador and also crawled information of user. Attributes were extracted from the tip content and behavior of user. They used supervised learning algorithm to classify spam or non-spam based on attributes. Dataset obtained from Apontador, two sets of data, 1. Labeled tips spam or non spam and 2. Crawled data to enhance features. obtain information of places, users and the social graph of more than 137,000 users with their links of follower and followees.

E. Tan, L. Guo, S. Chen, X. Zhang, and Y. Zhao [10] Proposed unsupervised method UNIK UNsupervised social network spam detection. By identifying previously SD2 limitations new approach was proposed. Instead of directly identify spammers, it removing non-spammers from the network, by using both the social graph and the user-link graph. They collected dataset from from a large commercial social blog site over 10 months. They first applied supervised learning approach to to detect spammers and used result for evaluation. They investigated that when when there is an increasing level of attacks SD2's performance degrades. UNIK overcome these limitations of SD2. UNIK system first identify non-spammers by applying the SD2 algorithm to the social graph. UNIK then make a whitelist of non spammers. It then use this whitelist with user link graph to detect spammers.

C. Yang, R. C. Harkreader, G. Gu [11] Proposed a fist neighbor based detection features to identify twitter spammers. They first done empirical analysis of the evasion tactics used by Twitter spammers on a large dataset containing around 500,000 Twitter accounts and more than 14 million tweets. Then they designed some features to identify twitter spammers. They used 24 features for detection. For dataset they developed Twitter crawler that taps into Twitter's Streaming API. They focused on harmful links to phishing and malware sites. Also analyzed evasive tactics that spammers are using to evade present machine learning detection schemes. Profile-based, Content based are the two evasive tactics. They designed 10 new detection features including three Neighbor-based features, three Graph-based features, three Automation-based features and one Timing-based feature. Then Evaluated robustness of these features.Used ML-based algorithm.

S. Lee, and J. Kim [12] They proposed a suspicious URL detection system for Twitter called WARNINGBIRD. They focused on correlated redirect chains of URLs in a number of tweets instead of landing pages of individual URLs in each tweet. They also used shared resources to identify suspicious URLs. They developed statistical classifier with features identified from correlated URLs and tweet context information. They Introduced 12 features for classifying suspicious URLs on Twitter. Dataset the collection of tweets with URLs and crawling for URL redirections were used obtained using Twitter API. They collected 27,895,714 tweet samples with URLs. For training two subcomponents: retrieval of account statuses and the training classifier used. Finally conclude that WARNINGBIRD is robust when protecting against conditional redirection.

C. Yang, R. Harkreader, J. Zhang [13] they first empirically analyzed how criminal accounts mix into and survive in the whole Twitter. They targeted criminal accounts defined by Twitter Rule. They also analyzed inner social relationships in the criminal community composed of criminal accounts and outer social relationships. Through analysis of criminal accounts social relationships they design an inference algorithm to catch more criminal accounts. The dataset for experiments contains 485,721 Twitter accounts with 14,401,157 tweets and 5,805,351 URLs. inner social relationships analyzed by visualizing its relationship graph and revealing its relationship characteristics. Outer social relationships by extracting and characterizing criminal supporters, they shown typical characteristics of these supporters. Criminal account Inference Algorithm(CIA) is used. CIA is based on the

following two observations: (1) criminal accounts tend to be socially connected; (2) criminal accounts usually share similar topic/keywords/URLs to attract victims, thus having strong semantic coordination among them.

H. Li, A. Mukherjee, B. Liu, R. Kornfield and S. Emery. [14] To discover campaigns, accounts of their promoters and how they organized and implemented so it can discover dynamics of internet marketing. They proposed a system in context of twitter. Main aim was to identify user accounts which are involved in promotion. For that streaming tweets from twitter and particular keyword of topic given. They solved problem as a relational classification problem and solve it using typed Markov Random Fields (T-MRF). Used three real-life datasets from the health science domain related to smoking. Two datasets were about two anti-smoking campaigns conducted by the Centers for Disease Control and Prevention (CDC), a government health agency in the USA, and one dataset is about electronic cigarettes (or e-cigarettes) promotions on Twitter. As there are different interactions or dependencies among the nodes they used two types of entities and three types of nodes 1 user: promoter or a non-promoter, 2 URL: promoted or organic, 3 burst: planned or normal burst. For each dataset they manually labeled 800 users. They performed 5 random runs. For each run, they randomly select 400 users for training and the dependencies.

J. Song, S. Lee and J. Kim [15] Proposed a system which uses relation features to detect message is spam or not. They reported that account features can easily be made-up by spammers cannot be collected until a number of malicious activities have been done by spammers. So using relation features, such as the distance and connectivity between a message receiver and a message sender, to decide whether the current message is spam or not, because they are easy to collect and difficult for spammers to control. Distance is the length of the shortest path and the connectivity is measured by using min-cut and random walk and used in classification. The 11 features that are used in classification. This method has two problems the message will be filtered out if new account user send message to his friend and message from infected friend.

M. Egele, G. Stringhini, C. Kruegel, and G. Vigna. [16] Proposed a system to detect compromised user accounts of Twitter and Facebook. It identifies accounts which experience a sudden change in behavior using a composition of statistical modeling and anomaly detection approach. They look for group of accounts which experiences similar changes in short period. They developed a tool COMPA which run on large-scale dataset of more than 1.4 billion publicly-available Twitter messages, as well as on a dataset of 106 million Facebook messages. Identified features to train model. Two types of model used, Mandatory model and Optional model. They evaluated messages for each model and combine results into a overall anomaly score of message.

Alex Hai Wang [17] Proposed a spam detection system is to identify suspicious users on Twitter. Also to explore the “follower” and “friend” relationships among Twitter a directed social graph model is proposed. According to twitter spam they identified content-based features and graph-based features for spam detection. They collected dataset from Twitter of around 500K tweets, 25K users, and 49M follower/friend relationships. To differentiate the suspicious behaviors from normal ones Bayesian classification algorithm is applied. This system can achieve 89% precision.

G. Stringhini, C. Kruegel, G. Vigna. [18] Proposed a system to detect spam on social network and how they operate. They created a set of honeynet accounts on three major social networks, Facebook, Twitter, Myspace. Total 900 and 300 accounts created on each site. They identified characteristics that used to detect spammers in a social network. To differentiate between real users and spam bots, they manually verify all the profiles that contacted them. Through this process, they noticed that spam bots share some common character, and formalized them in features that they then used for automated spam detection. Four categories of bots are Displayer, Bragger, Poster, Whisperer. To classify spammers and legitimate users they used machine learning techniques and six features.

S. Yardi, D. Romero, G. Schoenebeck, and D. Boyd [19] they examined spam around a one-time Twitter meme “robotpickuplines”. They demonstrate that the existence of structural network differences between spam accounts and legitimate users. Also given some challenges to differentiate spammers and non spammers. They used a simple algorithm to identify spam in #robotpickuplines with some of these properties: (1) keyword detection; (2) searches for URLs; and (3) username pattern matches. They used manually labeled 300 tweets as spam or non spam from # robotpickuplines and use that simple algorithm. They address some issues age of account, frequency of tweets, friends, spammers clustered, spammers location according to network in their study.

B. Wang, A. Zubiaga, M. Liakata, R. Procter. [20] Proposed a system for tweets spam detection using tweet-inherent features. For this task they identified four different feature sets and five classification algorithms using large hand-labeled dataset. They used two datasets for experiments (1) Social Honeypot Dataset Their dataset which consists of 22,223 spammers and 19,276 legitimate users along with their most recent tweets and (2) 1KS-10KN Dataset which contains 1,000 spammers and 10,000 legitimate users. Some preprocessing tasks are also applied to datasets i.e. decoding HTML entities, and expanding contractions with apostrophes to standard spellings. They used features or attributes which includes 11 user features, 17 content attributes and N-gram models, Sentiment features. For classification and evaluation they used 5 classification algorithms implemented using scikit-learn: Decision Tree, and Random Forests, Bernoulli Naive Bayes, K-Nearest Neighbour (KNN), Support Vector Machines (SVM).

C. Grier, K. Thomas, V. Paxson, M. Zhang. [21] In this study, to gain an insight into the underlying techniques used to attract users they group spam URLs into campaigns and identify trends that uniquely distinguish phishing, malware, and spam. Also they found that the blacklists are slow in identification of new threats. They also downloaded click through rates using bit.ly API. Proposed a technique to identify two types of spamming accounts, one created primarily for spamming and second those which are compromised by spammers. Datasets of over 200 million tweets from the stream and crawled 25 million URLs. They also collected the complete history for over 120,000 users with public accounts, half of which have sent spam identified by their blacklists; the history was an additional 150 million tweets sent by these users. They used clustering algorithm to cluster campaigns.

K. Lee, J. Caverlee, S. Webb. [22] To uncover social spammers in online social systems they proposed and evaluated a honeypot-based approach and also to preserve community value. (1) Development of social honeypots to harvest misleading spam profiles from social networking communities; and (2) Statistical analysis of the properties of these spam profiles for creating spam classifiers for filtering out existing and new spammers are the two key components of this system. The observations of implementation based on MySpace and Twitter. To identify previously unknown spammers they developed a machine learning based classifiers based on some profile features. They used two datasets: (1) MySpace Dataset: 1.5 million of public profiles collected and (2) Twitter Dataset: consists of 215,345 user profiles, 4,040,415 tweets. To monitor spammers' behaviors and log their information a social honeypots as information system resources was used. To detect social spam behavior they deploy social honeypot consisting of a legitimate profile and an associated bot. evidence of the spam candidate collected by social honeypot's bot. They extracted observable features from collected spam profiles. These features with a classifier SVM used for classification.

H. Gao, J. Hu, C. Wilson [23] This study aims to quantify and characterize spam campaigns from online online social networks. Their detection approach focuses on two techniques that group together wall posts that share either the same (possibly obfuscated) destination URL, or strong textual similarity. They model these post as node in a larg graph. They modeled each wall post in the form *<description, URL>* pair that is then used in detection phase. They reduced the problem of identifying spam campaigns to a problem of identifying connected subgraphs inside the similarity graph. They developed an Post Similarity Graph Clustering algorithm for that. They used a dataset of 187 million wall posts written to the profile walls of 3.5 million Facebook users. They identified over 200,000 malicious wall posts attributable to 57,000 malicious accounts using threshold-based techniques for spam detection.

C. Castillo, M. Mendoza, B. Poblete. [24] Proposed the automatic methods to assess the credibility of a given set of tweets. Aim was to classify trending topics in to creditable or non creditable. They used features from user's posting and reposting actions. By studying activities they extracted discussion topic. Then they manually labeled topics whether it is from event news or informal conversation. After that relevant features from each labeled topic were extracted and used them to build a classifier. Features were from four categories Message, User, Topic, and Propagation. A supervised learning was use in classifier to detect news, events.

3. CONCLUSIONS

So from this study we conclude that there are different approaches in spam detection. Some approaches used features to detect spam with the Machine learning algorithms. Futures involved were content features, user based features, URL-based features, features based on social graph. Content based features are number of hashtags per number of words on each tweet, number of URLs per words, number of words of each tweet, number of characters of each tweet, number of URLs on each tweet, number of hashtags on each tweet, number of numeric

characters that appear on the text, number of users mentioned on each tweet, number of times the tweet has been retweeted. User-based features include the number of followers, the number of accounts followed, and the follower ratio, account-age, no-user favorites, no-lists, and no-tweets. Features based on the social graph, such as , betweenness centrality, and bidirectional links ratio, local clustering coefficient. Datasets were gathered from OSN in different times. Twitter Streaming API is available for data collection. Also data collected from Facebook and Myspace. But Data may be streaming or may be historical stored data. It can affect spam detection. Data may be collected by honeypot accounts. For detection of spam or non-spam classification and clustering algorithms were used. Generally used machine learning algorithms are random forest, C4.5 decision tree, Bayes network, Naive Bayes, k-nearest neighbor, and support vector machine. Also some researchers developed their own algorithm based on features. Some researchers used a blacklist of URLs in spam detection. Some of its examples are SURBL, Google Safebrowsing, URIBL. Another one is Trend Micro's Web Reputation which maintains a large dataset of URL reputation records.

So finally there is impact of different factors on the spam detection performance, which contain spam to nonspam ratio, feature discretization, training data size, data sampling, time-related data, and machine learning algorithms. Three aspects data, attributes or features and approach are important in spam detection.

4. ACKNOWLEDGEMENT

I am very much thankful of Dr. B.L. Gunjal for their guidance and consistent encouragement in this paper work.

6. REFERENCES

- [1] C. Chen, J. Zhang, Y. Xie, Y. Xiang, W.i Zhou, M. Hassan, AlElaiwi, and M. Alrubaian, "A Performance Evaluation of Machine Learning-Based Streaming Spam Tweets Detection", in *IEEE TRANSACTIONS ON COMPUTATIONAL SOCIAL SYSTEMS*, pp 1-12, 2016.
- [2] F. Benevenuto, G. Magno, T. Rodrigues, and V. Almeida, "Detecting spammer on Twitter," presented at the 7th Annu. Collab. Electron. Messaging Anti-Abuse Spam Conf., Redmond, WA, USA, Jul. 2010.
- [3] Z. Miller, B. Dickinson, W. Deitrick, W. Hu, and A. H. Wang, "Twitter spammer detection using data stream clustering," *Inf. Sci.*, vol. 260, pp. 64–73, Mar. 2014.
- [4] K. Thomas, C. Grier, D. Song, and V. Paxson, "Suspended accounts in retrospect: An analysis of Twitter spam," in *Proc. ACM SIGCOMM Conf. Internet Meas.*, pp. 243–258, 2011.
- [5] K. Thomas, C. Grier, J. Ma, V. Paxson, and D. Song, "Design and evaluation of a real-time URL spam filtering service," in *Proc. IEEE Symp. Sec. Privacy*, pp. 447–462, 2011.
- [6] X. Jin, C. X. Lin, J. Luo, and J. Han, "Social spamguard: A data miningbased spam detection system for social media networks," *PVLDB*, vol. 4, no. 12, pp. 1458–1461, 2011.
- [7] Q. Cao, M. Sirivianos, X. Yang, and T. Pregueiro, "Aiding the detection of fake accounts in large scale social online services," in *Proc. Symp. Netw. Syst. Des. Implement.(NSDI)*, pp.197–210, 2012.
- [8] S. Ghosh et al., "Understanding and combating link farming in the Twitter social network," in *Proc. 21st Int. Conf. World Wide Web*, pp. 61–70, 2012.
- [9] H. Costa, F. Benevenuto, and L. H. C. Merschmann, "Detecting tip spam in location-based social networks," in *Proc. 28th Annu. ACM Symp. Appl. Comput.*, pp. 724–729, 2013.
- [10] E. Tan, L. Guo, X. Zhang, and Y. Zhao, "Unik: Unsupervised social network spam detection," in *Proc. 22nd ACM Int. Conf. Inf. Knowl. Manage.*, San Francisco, CA, USA, pp. 479–488, Oct. 2013.
- [11] C. Yang, R. Harkreader, and G. Gu, "Empirical evaluation and new design for fighting evolving Twitter spammers," *IEEE Trans. Inf. Forensics Sec.*, vol. 8, no. 8, pp. 1280–1293, Aug. 2013.
- [12] S. Lee and J. Kim, "Warningbird: A near real-time detection system for suspicious URLs in Twitter stream," *IEEE Trans. Dependable Secure Comput.*, vol. 10, no. 3, pp. 183–195, May/Jun. 2013.
- [13] C. Yang, R. Harkreader, J. Zhang, S. Shin, and G. Gu, "Analyzing spammers' social networks for fun and profit: A case study of cyber criminal ecosystem on Twitter," in *Proc. 21st Int. Conf. World Wide Web*, pp. 71–80, 2012.
- [14] X. Zhang, S. Zhu, and W. Liang, "Detecting spam and promoting campaigns in the Twitter social network," in *Proc. IEEE 12th Int. Conf. Data Mining (ICDM)*, pp. 1194–1199, 2012.
- [15] J. Song, S. Lee, and J. Kim, "Spam filtering in Twitter using sender receiver relationship," in *Proc. 14th Int. Conf. Recent Adv. Intrusion Detect.*, pp. 301–317, 2011.
- [16] M. Egele, G. Stringhini, C. Kruegel, and G. Vigna, "COMPA: Detecting compromised accounts on social networks," presented at the 20th. Annu. Netw. Distrib. Syst. Sec. Symp., San Diego, CA, USA, Feb.. 24–27, 2013.
- [17] A. H. Wang, "Don't follow me: Spam detection in Twitter," in *Proc. Int. Conf. Sec. Cryptogr. (SECRYPT)*,

pp. 1–10, 2010,

[18] G. Stringhini, C. Kruegel, and G. Vigna, “Detecting spammers on social networks,” in *Proc. 26th Annu. Comput. Sec. Appl. Conf.*, pp. 1–9, 2010.

[19] S. Yardi, D. Romero, G. Schoenebeck, and D. Boyd, “Detecting spam in a Twitter network,” *First Monday*, vol. 15, nos. 1–4, Jan. 2010.

[20] B. Wang, A. Zubiaga, M. Liakata, and R. Procter, “Making the most of tweet-inherent features for social spam detection on Twitter,” arXiv preprint arXiv:1503.07405, 2015.

[21] C. Grier, K. Thomas, V. Paxson, and M. Zhang, “@SPAM: The underground on 140 characters or less,” in *Proc. 17th ACM Conf. Comput. Commun. Sec.*, pp. 27–37, 2010.

[22] K. Lee, J. Caverlee, and S. Webb, “Uncovering social spammers: social honeypots + machine learning,” in *Proc. 33rd Int. ACM SIGIR Conf. Res. Develop. Inf. Retrieval.*, pp. 435–442, 2010.

[23] H. Gao, J. Hu, C. Wilson, Z. Li, Y. Chen, and B. Y. Zhao, “Detecting and characterizing social spam campaigns,” in *Proc. 10th ACM SIGCOMM Conf. Internet Meas.*, pp. 35–47, 2010.

[24] C. Castillo, M. Mendoza, and B. Poblete, “Information credibility on Twitter,” in *Proc. 20th Int. Conf. World Wide Web*, pp. 675–684, 2011.

BIOGRAPHIES

	Mr. R. S. Gaikwad is Pursuing Master in Engineering from Amrutvahini College of Engineering Sangamner. Received BE degree from Savitribai Phule Pune university. His interests Areas are Data Mining, Big Data, Software Engineering.
	Dr. B. L. Gunjal is Associate Professor in computer Engineering and Head of Information Technology Department in Amrutvahini College of Engineering, Sangamner. Received PhD degree in Computer Engineering from Savitribai Phule Pune university. Also received “BEST Teacher Award” by Savitribai Phule Pune university. Her Research interests in Image Processing, Data Mining.