

IDENTIFYING THE NATURE OF USERS IN SOCIAL MEDIA USING SVM ALGORITHM

J. Mohamed Aslam,
Research Scholar,
Department of Computer Science,
Rajah Serfoji Government College,
Thanjavur, Tamilnadu, India.

Dr. K. Mohan Kumar,
Research Guide & Head,
Department of Computer Science,
Rajah Serfoji Government College,
Thanjavur, Tamilnadu, India.

Abstract: Social network has become a very popular way for internet users to communicate and interact in online. Users spend plenty of time on famous social networks (e.g., Facebook, Twitter, SinaWeibo, etc.), for reading news, discussing events and posting messages. The main aim of this work is to find out the spammers and non-spammers (good and bad friends or followers). For this work more than 1000 messages collected from different friends. Then a labeled dataset of users are constructed and applied into SVM (Support Vector Machines) based spammer detection algorithm. This algorithm gives the true positive rate of spammers and non-spammers reaching 99.1% and 99.9% respectively.

Keywords: Support Vector Machines, Application Program Interface, Radial Basis Function, spammer, non-spammer

I. INTRODUCTION

Online social network, such as Facebook, Twitter, Weibo, etc., has become one of the entertainments for internet users to keep communications with their friends [1–3]. According to Statista report, the number of social network users has reached 1.61 billion until late 2013, and is estimated to be around 2.33 billion users in the globe, until the end of 2017. An along with great technical and commercial success, social network platform also provides a large amount of opportunities for broadcasting spammers, which spreads malicious messages and behavior[4]. According to Nexgate report, during the first half of 2013, the growth of social spam has been 35.5%. A social spam message is potentially seen by all the followers and recipients' friends., discussing possible solutions for spam detection and filtering. However, they are either ineffective or based on too much considered conditions (e.g., a lot of content and behavior feature, etc.) [5]. This system contains the following main features. They are (i) adopts the spammer feature to detect spammer, (ii) test the results over face book, (iii) can implement any biggest social network site in this world and (iv) Under the face book API, a specific dataset crawler is developed to extract any unauthorized users' public messages inside the face book platform.

II. PROBLEM DESCRIPTION

Online social networks are widely use these days for the purpose of communication. Users can share more type of information among friends. But there exist some social network users who misuse the features of these social networks and promote the spreading of malicious content. They do this by uploading the malicious post in other user page. These contents spread at a fast rate. There is no proper mechanism to detect these malicious posts immediately and remove it effectively. The organization and scale of these trolling campaigns, however, suggests that the practice has moved into a new phase, whereby corporations and governments seek to control the discourse surrounding popular events (both real and imagined) on social media. However, the Times discovered that the agency always routed its Internet traffic through proxy servers, thus

rendering useless the easy path to doing so via the originating IP addresses. The relationship between the work of humanists and forensic examiners cannot be denied: both groups seek to make accurate predictions about uncertainties related to textual data. There is often enough distinct information in even just a handful of sentences for a human reader to understand that they are from a common source. For instance, function words can be coupled with the most frequent punctuation, or other stable features, becoming more flexible and discriminative for use in a bag-of-words representation that disregards grammar, but preserves the underlying statistics of language. They are useful complements that can enhance the accuracy of attribution models. The presence of some labeled data often improves results dramatically. Accordingly, supervised approaches now dominate the field. It is not always possible to perform meaningful semantic analysis on sample sizes as small as tweets with any of today's topic modeling algorithms. Feasible if they are considering just a single testing tweet in an actual investigation. It also makes supervised learning methods unreliable, because a user might make references to the same person across her messages, creating a strong bias towards that particular feature in training; and any message with a reference to that same person would subsequently be misclassified as being from that user. There are still many other native features to be explored in social media, including the actual social graph of a suspect. Other meta-data such as timestamps, device records, and browser records could also be exploited, though they are often unreliable and, in some cases, easily forged. Therefore, a more generalized solution that incorporates knowledge of the problem, general and data-driven features and, possibly, network connectivity and the social network of suspects would be, certainly, more difficult to tamper with the existing system. This problem will be resolved while using in the proposed methodology.

III. PROPOSED METHODOLOGY

Based on dataset and feature collection described in the previous section, a supervised machine learning model is introduced for spammer's identification. Supervised learning is the machine learning task of inferring a function

from labeled training data that consists of a set of training examples..

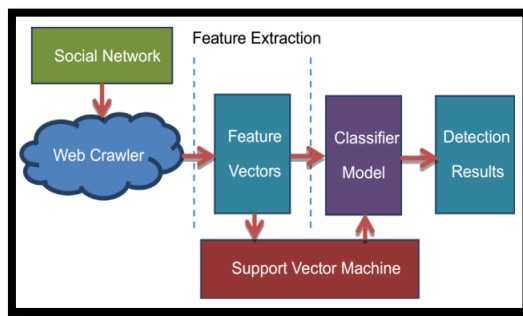


Figure 1: SVM Architecture

Figure-1 illustrates the basic concept of proposed spammer detection model. In this solution, training data is converted to a series of feature vectors that consist of a set of values for attributes. These vectors construct the input of supervised machine learning algorithm. After training, a classification model is applied to distinguish whether the specific user belongs to normal user or spammer.

SVM classifier: The spammer detection solution is based on a non-linear support vector machine (SVM) classifier with the Radial Basis Function (RBF) kernel. This could be achieved through the implementation provided by libSVM, an integrated software for supporting vector classification, regression and distribution estimation. The SVM with RBF kernel function has two such training parameters: C controls over fitting of the model; and gamma controls the degree of nonlinearity. In the experiment, apply a parameter selection tool provided by libSVM to select parameters automatically with a 5-fold cross-validation.

Ratio of spammer to non-spammer: To use complete training dataset for testing work and achieve F-measure value of spammer and non-spammer as 91.6% and 93.2% respectively. This might not be the optimized result. In order to achieve higher spammer detection accuracy, the ratio of spammer to non-spammer in the training dataset is changed as follows: 10:1, 8:1, 6:1 4:1, 2:1, 1:2, 1:4, 1:6, 1:8 and 1:10, with corresponding classification accuracy result illustrated in Fig. 6. It shows that F-measure value of both spammer and non-spammer grows simultaneously when the ratio of spammer decreases, and reaches the highest accuracy of about 99.5% and 99.9% when the ratio is set to 1:2. After that, the accuracy drops quickly while the ratio of non-spammer rises.

Software Tool: A new software tool is developed to resolve the above said conflicts. The tool developed with the following modules.

- Social network creation
- Source Data
- Classification Strategies
- Analysing character module

IV.RESULTS AND DISCUSSION

The following figures show the screen shots of the developed software.



Figure 2: Login page

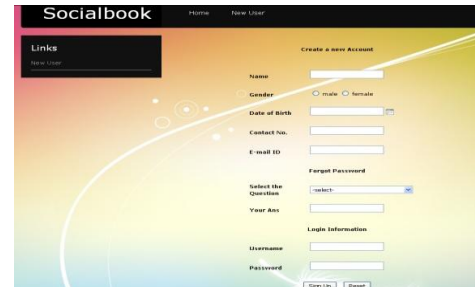


Figure 3: Details page



Figure 4: Profile page



Figure 5: Access page

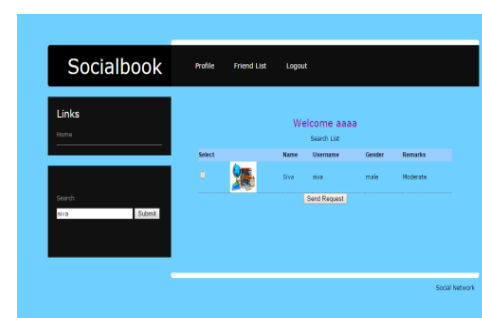


Figure 6: Analyzing of Good or Bad User

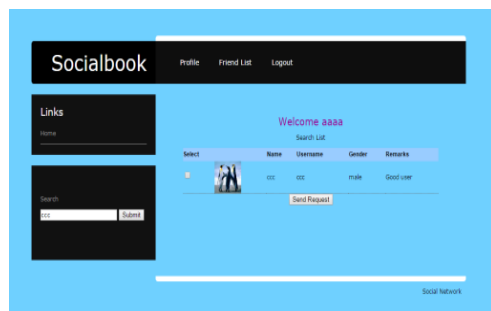


Figure 7: Identification of Good User

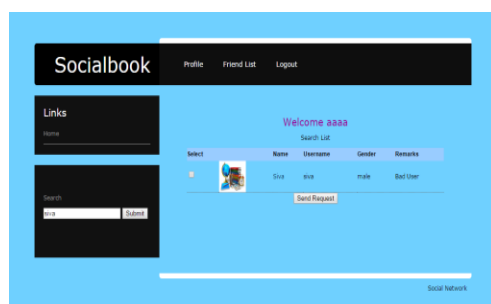


Figure 8: Identification of Bad User

The following table arrived by comparing the attributes such as Correctness, Efficiency, Flexibility, integrity, interoperability, maintainability, portability, reliability, reusability, testability, usability and availability with existing system and our proposed system.

S.NO	Attributes	Existing system	Proposed system
1	Correctness	70	90
2	Efficiency	80	99
3	Flexibility	70	98
4	Integrity/Security	80	99
5	Interoperability	70	95
6	Maintainability	70	99
7	Portability	81	97
8	Reliability	79	98
9	Reusability	89	96
10	Testability	80	98
11	Usability	89	99
12	Availability	80	99

Table 1: Software Quality Attributes Comparison

The graphical representation of the above table -1 is given below.

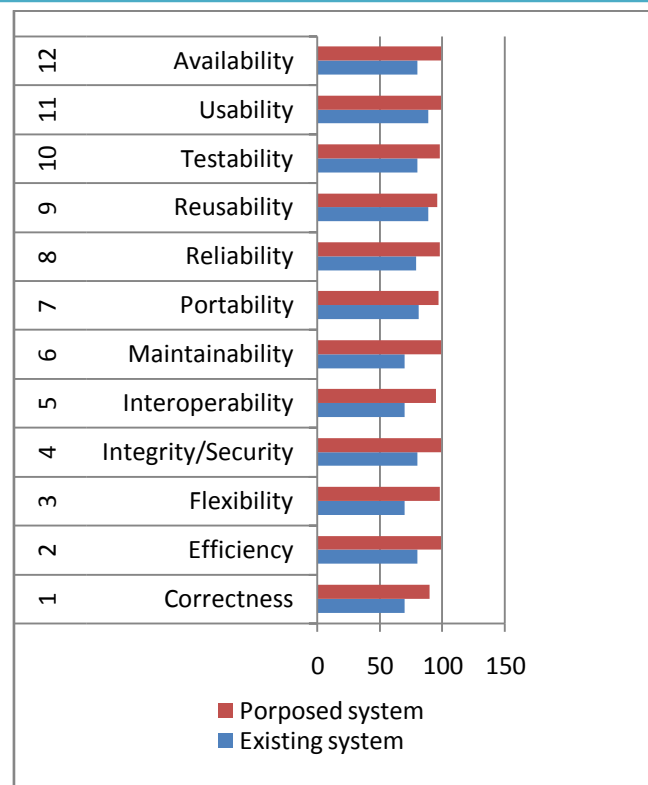


Figure 8: Software Quality Attributes by Graph.

The above figure shows that the proposed system improves the various attributes which are needed for good quality software.

V.CONCLUSION

The advances in digital and multimedia technology are significantly impacting human behaviors and social interactions. The main idea of their global research project is to develop an automatic system for detecting suspicious profiles in the social media, through which we can uncover suspicious behavior and interests of users as well. At present a system for detecting suspicious posts in social network using similarity approach in text analysis. Their approach is based on similarity with comparing social network seized posts with a suspicious predefined database.

VI. FUTURE WORK

For future work, plan to improve the system in term of execution time, developing automated classification and using other knowledge resources in order to improve the precision rates, the semantic of exchanged information will be used to identify more significant suspicious profiles.

REFERENCES

- [1]. Facebook, (<http://www.facebook.com/>).
- [2]. Welcome to Twitter, (<http://twitter.com/>).
- [3]. Weibo – SINA, (<http://english.sina.com/weibo/>).
- [4]. Statista, (<http://www.statista.com/>).
- [5]. Nexgate. 2013 State of Social Media Spam, (<http://nexgate.com/wp-content/uploads/2013/09/Nexgate-2013-State-of-Social-Media-Spam-Research-Report.pdf>), 2013.

- [6]. A. Abbasi and H. Chen. Applying authorship analysis to extremist group web forum messages. *IEEE Intelligent Systems*, 20(5):67–75, 2005.
- [7]. A. Abbasi and H. Chen. Writeprints: A stylometric approach to identity-level identification and similarity detection in cyberspace. *ACM Transactions on Information Systems*, 26(2):7, 2008.
- [8]. S. Afroz, A. Caliskan-Islam, A. Stoleran, R. Greenstadt, and D. McCoy. Doppelganger finder: Taking stylometry to the underground. In *IEEE Security and Privacy*, 2014.
- [9]. J. Albadarneh, B. Talafha, M. Al-Ayyoub, B. Zaqaibeh, M. Al-Smadi, Y. Jararweh, and E. Benkhelifa. Using big data analytics for authorship authentication of arabic tweets. In *IEEE/ACM Intl. Conference on Utility and Cloud Computing*, 2015.
- [10]. M. Almishari, D. Kaafar, E. Oguz, and G. Tsudik. Stylometric linkability of tweets. In *Workshop on Privacy in the Electronic Society*, 2014.
- [11]. A. Anderson, M. Corney, O. de Vel, and G. Mohay. Multi-topic email authorship attribution forensics. In *ACM Conference on Computer Security*, 2001.
- [12]. Yui Arakawa, Akihiro Kameda, Akiko Aizawa, and Takafumi Suzuki. Adding twitter-specific features to stylistic features for classifying tweets by user type and number of retweets. *Journal of the Association for Information Science and Technology*, 65(7):1416–1423, 2014.
- [13]. S. Argamon, M. Koppel, and G. Avneri. Style-based text categorization: What newssystem am I reading? In *AAAI Workshop on Text Categorization*, pages 1–4, 1998.
- [14]. S. Argamon, C. Whitelaw, P. Chase, S. R. Hota, N. Garg, and S. Levitan. Stylistic text classification using functional lexical features. *Journal of the American Society for Information Science and Technology*, 58(6):802–822, 2007.
- [15]. Juola & Associates. Computational analysis of authorship and identity for immigration. Whitesystem, 2016. 2002.
- [16]. H. Baayen, H. Van Halteren, and F. Tweedie : *Literary and Linguistic Computing*, 11(3):121–132, 1996.