

Knowledge Discovery in Tweets for the Prevention of Inference Attacks

D. Sai Eswari, Afreen Rafiq and R. Deepthi

Department of Computer Science and Engineering, Malla Reddy Engineering College for Women,
Hyderabad – 500014, Telangana, India; saieswari3@gmail.com, afreenrafiq2@gmail.com,
deepthideepu66@gmail.com,

Abstract

Objectives: The objective of our proposed study is to eliminate the inference attacks through knowledge discovery process without compromising the accuracy of the systems. **Methods/Analysis:** In this paper, we have explored an enhanced privacy adoption scheme for twitter users. A supervised learning model that prevents the inference attacks using hit analytics and metadata knowledge derivation systems. By integrating the extracted feature vectors, we have eliminated the inference attacks caused by third party applications and elegantly classified the malignant and benign URLs. **Findings:** The proposed model is experimented in twitter dataset. The tweet URLs are collected and segmented into different URLs and similar URLs. The performance analysis is in terms of frequency of URLs usage, User's assurance level via hit analytics and accuracy. It is evident from the results that higher usage of twitter network is employed and as the rate of users increases, the proposed classifier detects the benign and malicious URLs at significant level of 83%. **Novelty/Improvement:** A recent development made in the digital era has seen a phenomenal growth of social users. The social users make use of social networks to share their opinions about the facts, events etc. Our proposed study has incorporated on twitter users where trained classifier is developed from the behavior analysis of tweets users.

Keywords: Click Analytics and Third Party Applications, Digital Era, Public Knowledge, Social Networks, Twitter, URL Services

1. Introduction

Due to the developments made in the Online Social Networks, the significance of data analysis relied upon the attributes used by individual and group users. Due to security reasons, some of the attributes are alone revealed in their user's account profiles. In order to find the knowledge about the user's experience, some hidden attributes has to be studied. Most of the prior works has focused on Twitter¹ which is one of the social networks that depicts widespread knowledge of the adopted service and user's tendency to keep their sensitive data. Twitter is the most renowned networks that used by different people. A survey reported that 350 million messages have been posted by 150 million active users. With the use of third party applications like Windows, MAC, iOS etc, the URL shortening services has been performed. In some cases, the

length of URL will be restricted for achieving better services².

URL shortening service is one of the most focused studied for deriving the public knowledge. The best instances of the URL shortenings are bit.ly and goo.gl. When any user clicks the URL, then the knowledge about that URL services is extracted. The visitor information can also be derived from the analysis of visitor's history³. The target of the URL shortening is to remove the limit setup in message characters over different web services. The phenomenal growth in web technologies throws several threats which are a daunting task to be discovered. Henceforth, none of the study has revealed the user's assurance and consciousness about the benign or malignant URLs.

Privacy control⁴ over the web services is quite ineffective against the inference attacks. The different security

*Author for correspondence

operations are provided for different web technologies. However, as they do not consider the actual content of what is being shared, these simple policies provide no help with limiting the unanticipated impact of global inference⁵. The paper is focused on preventing the inference attacks by analyzing the metadata on Twitter data⁶. The rest of the paper is organized as follows: Section 2 describes the related work; Section 3 describes the proposed work; Section 4 describes the performance evaluation and concludes in Section 5.

2. Related Work

This section depicts the prior works carried out in knowledge derivation using social networks. In accord to web services and policies, the sensitive information should be shared and protected for the authorized users. Hence, different security operations were suggested by research communities. The author in⁷ studied about the friend-follower ratio, URL ratio and messenger ratio about the spam messages. They discussed about how the spammers steal the sensitive data from tweets. Honey-Profiles are used for validating the tweets from three kinds of social networking. They described about the features used for deriving the knowledge. Similarly, the study was extended to the directed graph processes where vertices represent the user accounts and edge represents the types of associations among the users⁸. In order to detect the attacks, the in-degree and out-degree of the user's profiles are depicted. It acts as machine learning model to obtain the user's knowledge about the security controls. They also detected the spam accounts using machine learning models⁹.

The author in¹⁰ discussed about the undirected graph model for predicting the behaviors of spammer. To clearly depict the feature of the message, distance and connectivity metric has been explored. Their objective was to predict the inherent features of the messaging systems and also classifies the whether the message is benign or spam. The author in¹¹ discussed the detection model based on user's interaction process, in specific to analysis on tweets used, number of tweets, and unique URL numbers. The system inspects every message and evaluates the feature values before rendering the message to the intended recipients and makes immediate decision on whether or not the message under inspection are dropped. The author in¹² studied about the detection of malicious websites under content based features and host based features of URL

systems. They suggested a prediction model for online algorithms to devise the features of the systems.

The author in¹³ studied about the honeyclient systems. When the page is loaded, the honeyclient systems execute code over the file system to detect the unexpected behaviors. The serious demerits about honeyclient system are the higher consumption of time and cost which indirectly affects the scalability of the systems¹⁴. To enhance the system's scalability, an efficient filtering system is required. Prophiler is the static technique used for investigating on malicious messages which performs like filtering process¹⁵. This technique revealed that most malicious tweets arouse from impressive words and hash tags.

From the previous analysis, it is inferred that the research issues to be concentrated are:

- Overwhelmed growth of unstructured textual data,
- Stealing of sensitive information via browsing history analysis,
- Lack of privacy controls in social networks,
- Lack of figuring out the inference attackers, and
- Lack of triggering accuracy of user's location.

3. Prevention of Inference Attacks using Knowledge Discovery Process

This section depicts the workflow of proposed methodology using real time dataset, Twitter using an enhanced public hit analytics and metadata knowledge derivation methods. The four major processes of proposed model are explained as follows:

3.1 Data Collection

The initial step of our proposed model is the acquisition of data. It is performed in two steps: 1. Acquisition of tweets with its URLs, and 2. Creeping for redirected URLs. Since, Twitter Streaming process is employed for obtaining the context information of the public opinions. A creeping thread is applied to receive the tweets from URLs.

3.2 Feature Extraction

Feature extraction process is the second step that permits to find significant attributes used for predicting

the inference attacks. It contains three parts, namely: 1. Aggregating similar domains, 2. Hit analytics, and 3. metadata knowledge derivation.

Aggregating similar domains: When the tweets are gathered in tweet window, it is allowed for monitoring the different URL or similar URL. The similar URLs are thus substituted by the name of its domain. Example,

Hit Analytics: It checks how many times the users visited the twitter networks. It also checks the actions performed by new visitors. During an incessant monitoring process, the twitter users and their followers are also monitored. The inquiries should incessantly monitor both static and dynamic modification done to the URLs. Correspondingly, the inquiries over user's history make to portray the behavior of users.

Metadata knowledge derivation: The basic feature of the inference attacks is studied for the URLs, goo.gl, and bit.ly. The attack system tries to differentiate shortened URLs hit on virtual users from shortened URLs hit by real Twitter users. Thus, the extracted features are preserved into real-valued vector form. By analyzing this vector form, the benign URLs are moved into whitelists and malignant URLs are moved into blacklists.

3.3 Training

The training process acts as supervised learning systems. It contains two processes, namely, retrieval status of tweets accounts and training classifier. The training classifier contains two classes, URL from blacklist accounts as malicious and URL from white lists accounts as benign. The training vectors are periodically updated by its feature vectors.

3.4 Classification

The target of the classification process is to classify the benign and malicious URLs from the obtained feature vectors. The trained classifier executes on the URLs, if it returns higher number of malicious feature vectors then it is classified as malignant class else benign class. The proposed process is explained in Figure 1.

4. Experimental Analysis

This section depicts the simulation analysis carried out in the Twitter Datasets. With the help of click analytics process, we have analyzed the URL shortened services on goo.gl and bit.ly. These two URL have received 1 million

of queries per day. Table 1 depicts the sample twitter dataset (Table 1).

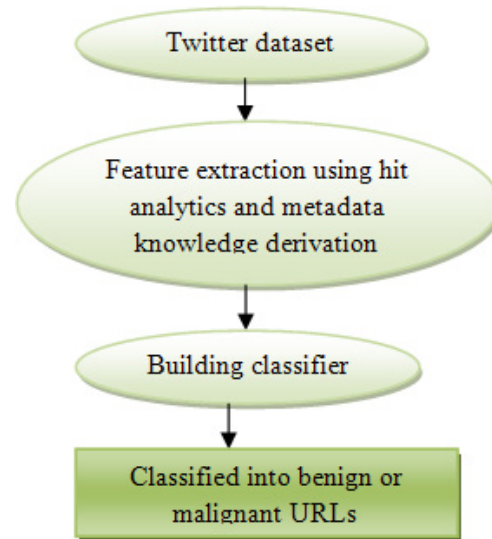


Figure 1. Proposed workflow.

Table 1. Data collection

No. of followers	Goo.gl	Bit.ly
100-1K	5	6
1k- 10k	9	5
10k-100k	35	12
100k-1M	19	23
Total	68	46

The following parameters were analyzed for prediction of inference attacks in terms of:

Usage of Frequency: Sample analysis is carried out to differentiate the usage of shortened URLs under different domains like social networking sites, emails, webs and others. The Figure 2 presents the importance of shortening URLs services.

User assurance level via hit analytics: This analysis deals on the user's assurance level over the privacy operations. In order to use the secured URLs, several consecutive measures have been adopted by URL shortening services. The URLs are ranked by the users and the highest ranked score is taken for building user's assurance level. The Figure 3 presents the confidence analysis in terms of website popularity, well-defined security measures and spantime services using click analytics process.

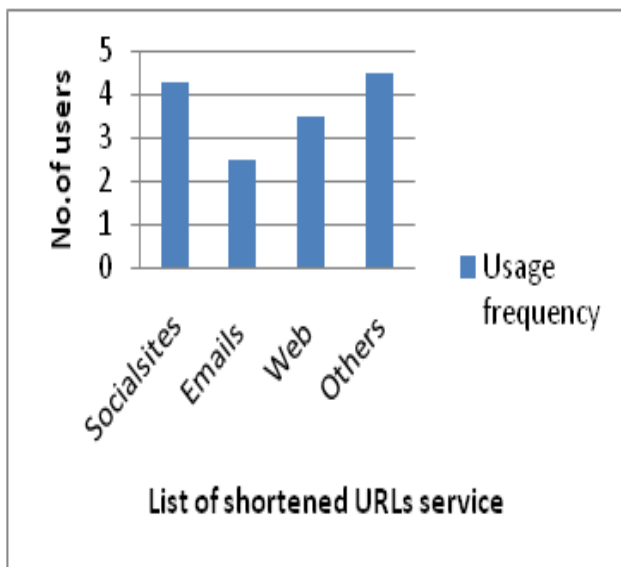


Figure 2. Frequency usage.

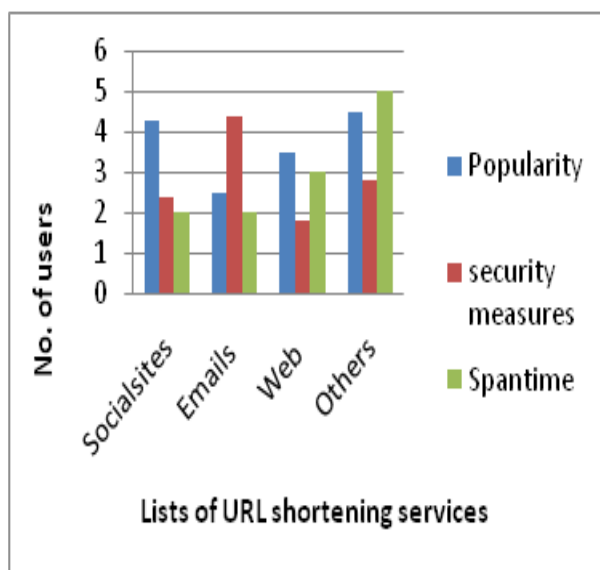


Figure 3. Assurance analysis through hit analytics process.

Accuracy: By estimating the accuracy, we have found the prevention of inference attacks via metadata knowledge derivation. In real time scenarios, it is quite common that similar URLs are used by the users. Hence, the user's behavior is analyzed for injecting malicious URLs. It is computed as:

$$\sum_{a,b \in \text{set of pairs in tweets}} = \frac{J(a,b)}{|\text{set of pairs in tweets}|}$$

Where $J(a,b)$ dictates the jaccard index between two set of users (Figure 4).

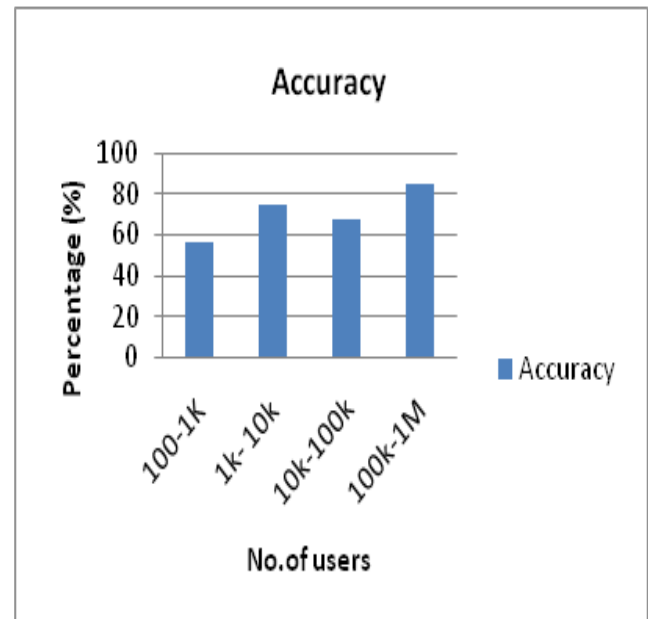


Figure 4. Accuracy rate.

5. Conclusion

In the recent days, the usage of short URLs has been employed in several real-time applications. The exploration of new threats is introduced by different sources. According to the usage of social users, new threats like malware, phishing, spam etc are defined from the behaviors of visiting users. Though prior works have been researched in shortening URLs service, the issue of inference attacks can't be solved. In this study, we have proposed a supervised learning model that prevents the inference attacks using click analytics and metadata knowledge derivation systems. Hence, the aim of this paper is to resolve the gap persists in social networking sites in specific to URL shortening services. We have analyzed the proposed learning model in terms of usage frequency, user's confidence and accuracy rate using real-time dataset, Twitter. It is evident from the results that our proposed model yields better accuracy rate.

6. References

1. Jonghyuk Song et. al. Inference attack on browsing history of twitter users using public click analytics and twitter metadata, IEEE Transactions on Dependable and Secure Computing. 2014.
2. Neumann A. Analyzing security implications of URL shortening services, Diploma Thesis, RWTH Aachen University, 2011.

3. Maggi F, Frossi A, Zanero S, Stringhini G, Stone-Gross B, Kruegel C, Vigna G. Two years of short URLs internet measurement: Security threats and countermeasures, Intl. World Wide Web (WWW) Conference, Rio de Janeiro, 2013. Crossref.
4. Neumann A, Barnickel J, Meyer U. Security and privacy implications of URL shortening services, Web 2.0 Security and Privacy 2011 Conference, Oakland, USA, May 2011.
5. Antoniadis D, Athanasopoulos E, Polakis I, Ioannidis S, Karagiannis T, Kontaxis G, Markatos EP. Web: The web of short URLs, 2011 Intl. World Wide Web (WWW) Conference, Hyderabad, India, March 2011. Crossref.
6. Iversion AI. Spamhaus and URL shortening services, Spam Ressource. Available at: <http://www.spamresource.com/2011/03/spamhaus-url-shortening-services.html>.
7. Lab MX. Increase in usage of URL shorteners in spam campaigns, Available at: Crossref.
8. Lauretti D. Facebook is blocking links from Google's URL shortening service, Examiner, March 2013.
9. Lab MX. Shortened URLs: The real dangers behind and how to avoid troubles. Available at: Crossref.
10. Hoffman S. Cligs URL shortening service hacked, users redirected, CRN Technology News for Solution Providers and the IT Channel, June 2009, p.1.
11. Rajab M, Ballard L, Lutz N, Mavrommatis P, Provos N. CAMP: Content-agnostic malware protection, 20th Annual Network and Distributed System Security Symposium, CA: USA, February 24, 2013.
12. Merritt M. Family online safety guide, 4th ed. Norton Symantec Press, December 2012.
13. Consoi CA. Dealing with image spam, Virus Bulletin. Dec. 2006; 1-3.
14. Fry R. Malware defense and automation: Fully integrated defense operation, RSA Conference, February 24-27, 2014.