

Privacy Leakage via De-anonymization and Aggregation in Heterogeneous Social Networks

Huaxin Li, Qingrong Chen, Haojin Zhu, *Senior Member, IEEE*, Di Ma, *Member, IEEE*, Hong Wen, *Member, IEEE*, and Xuemin (Sherman) Shen, *Fellow, IEEE*

Abstract—Though representing a promising approach for personalization, targeting, and recommendation, aggregation of user profiles from multiple social networks will inevitably incur a serious privacy leakage issue. In this paper, we propose a Novel Heterogeneous De-anonymization Scheme (NHDS) aiming at de-anonymizing heterogeneous social networks. NHDS firstly leverages the network graph structure to significantly reduce the size of candidate set, then exploits user profile information to identify the correct mapping users with a high confidence. Performance evaluation on real-world social network datasets shows that NHDS significantly outperforms the prior schemes. Finally, we perform an empirical study on privacy leakage arising from cross-network aggregation based on four real-world social network datasets. Our findings show that 39.9% more information is disclosed through de-anonymization and the de-anonymized ratio is 84%. The detailed privacy leakage of user demographics and interests is also examined, which demonstrates the practicality of the identified privacy leakage issue.

Index Terms—Data privacy, Social networks security, De-anonymization, Heterogeneous social networks.

1 INTRODUCTION

Social networks (online social networks, mobile social networks, vehicular social networks, etc.), or social media, have been extremely popular in current days. The latest statistics show that the number of active traditional social media users has exceeded 2.7 billion [1]. Along with overwhelming popularity of social networks, people enjoy abundant functionalities and services of a variety of social networks, including sharing status updates, posting photos, communicating with others, and making friends.

Due to the different functionalities of different social networks, a user tends to sign in multiple social networks for different purposes. According to the report conducted by Pew Research Center in 2015, 52% of online adults use two or more social media sites such as Facebook, Twitter, MySpace, or LinkedIn [2]. Aggregating user profiles from different social networks reveals various aspects of users. It is interesting that cross-network information represents a double-edged sword. On one hand, once the user's multiple accounts of different social networks are identified or mapped, these accounts' profiles, preferences, and activities can be collected to benefit personalization, targeting, and recommendation. The latest research pointed out that, the ads delivered by Google, one of the major ad networks, are personalized based on both users' demographic and interest profiles [3]. On the other hand, the adversary can exploit cross-network aggregation

to collect the information of various aspects of the target users, which will incur a serious privacy leakage issue [4]. This issue can not only exist in traditional social networks but also exist in new emerging social networks, like vehicular social networks. For example, Twittermobile car is able to send and receive Twitter messages, which contain the information including drivers' status, vehicle profiles, and real-time traffic notifications; RoadSpeak is a voice chatting system used by daily driving commuters or a group of people who are on a commuter bus or train [5]. These vehicular social-based applications exploit traditional online social networking services, like Facebook and Twitter, and thus are also under threat of de-anonymization attack.

In this study, we take an initial step towards investigating the following two questions: i) How can we design a practical and effective cross-network aggregation scheme for heterogeneous social networks? The proposed cross-network aggregation scheme is expected to link the target user's various accounts on different social media platforms and collect the user's profile in different aspects. ii) To what degree the cross-domain aggregations can reveal the different attributes of a user (e.g. interest, demographics).

One of the fundamental challenges of bridging the different social identities of the users on different social media is that the users tend to use varying usernames (screen names) or have unequal profiles (e.g. fields such as homepage, birthday, etc.) due to the increasing privacy concerns. The process of identifying user from a social network (e.g., anonymized network) based on another social network (e.g., auxiliary network) is called 'de-anonymization'. Recently, there is an increasing interest to study how to 'de-anonymize' or 're-identify' users across social networks, which mainly falls to the following two categories: profile based de-anonymization and structured based de-anonymization, which either suffer from high false positive or assume the social networks are aligned.

In this study, to answer the above questions, we first present a Novel Heterogeneous De-anonymization Scheme for heterogeneous social networks, which is coined as NHDS. Different from

- Huaxin Li, Qingrong Chen, and Haojin Zhu are with Shanghai Key Laboratory of Scalable Computing and Systems, Department of Computer Science and Engineering, Shanghai Jiao Tong University, Shanghai, 200240, China (E-mail: lihuaxin003@sjtu.edu.cn, chenqingrong@sjtu.edu.cn, zhu-hj@cs.sjtu.edu.cn).
- Di Ma is with the College of Engineering and Computer Science, University of Michigan-CDearborn, Dearborn, MI, 48128, USA (E-mail: dmadma@umich.edu).
- Hong Wen is with the National Key Laboratory of Science and Technology on Communication, University of Electronic Science and Technology of China, Chengdu, 611731, China. (E-mail: sunlike@uestc.edu.cn).
- Xuemin (Sherman) Shen is with the Department of Electrical and Computer Engineering, University of Waterloo, Waterloo, Ontario, Canada. (E-mail: sshen@uwaterloo.ca).

any previous works which either focus on profile based or structure based approach, NHDS aims to integrate the merits of two kinds of approaches. The motivation is that a real-world attacker is able to leverage as much information as she can to help de-anonymize in practice. Since both user profiles and network graph topology can be collected through web crawlers, platforms' APIs, or public datasets, a novel approach that leverages the merits of these two strategies is expected to achieve a higher performance. In particular, it firstly leverages the social network structure to significantly reduce the size of node candidate set. Then, it exploits user profile matching to further identify the correct mapping nodes with a high confidence. The seed nodes that act as the anchor points to align two or more heterogeneous social networks will be identified automatically.

To further investigate the privacy leakage caused by the cross-network aggregation, we apply the proposed NHDS algorithm to a large dataset involving four real-world heterogeneous online social networks, i.e., Livejournal, Flickr, Last.fm, and Myspace. We perform the de-anonymization algorithm and measure the privacy leakage arising from cross-network aggregation. The results are quite surprising in that, with the proposed de-anonymization algorithm, cross-network aggregations can reveal 39.9% uncovered attributes of users (e.g. interest, demographics).

To the best of our knowledge, the proposed work is the first empirical study which evaluates the impact of cross-network de-anonymization and aggregation on privacy leaking on the real-world datasets. From the privacy protection perspective, our study also reveals the potential risks to the community about user de-anonymization and information aggregation, and calls for the following research efforts on privacy-preserving personal recommendation. The major contributions of this paper can be summarized as follows:

- We propose the Novel Heterogeneous De-anonymization Scheme to de-anonymize users across heterogeneous social networks. The proposed scheme jointly exploits publicly available network structure information and user profile, which is expected to significantly increase the detection accuracy.
- We conduct extensive experiments on real-world heterogeneous social network datasets to demonstrate the effectiveness of our proposed scheme. The comparative results show that NHDS achieves high detection accuracy and maintains a considerable recall compared with the baseline.
- To understand the consequence of real-world de-anonymization attack or cross-domain aggregation, we investigate and quantify the information leakage through network aggregation based on de-anonymized social networks. The results show that 39.9% information is disclosed and the de-anonymized ratio (defined in Section. 6.3) is 84%, which raises a serious privacy concern.

The rest of paper is organized as follows. Section II discusses the related research works. Section III models preliminary concepts and formulates the problem. The proposed approach is presented in Section IV. Then, Section V evaluates the results based on a set of real-world social networks. Section VI investigates consequent privacy leakage via de-anonymization, and VII concludes this paper.

2 RELATED WORK

2.1 Structure based de-anonymization

De-anonymizing social networks is a hot research topic in recent years. Structure based de-anonymization works are based on the assumption that the different social networks of the same group users should show the similar network topology, which can be exploited for user identification [6], [7], [8]. The observation of this kind of approaches is that a user tends to build connections with similar users they are interested in or acquainted with in different social networks. Narayanan and Shmatikov performed the de-anonymization attack to large-scale directed social networks. They designed a de-anonymization algorithm by identifying some seeds and propagating based on structure similarity [9]. In [10], Nilizadeh et al. extended Narayanan and Shmatikov's attack by proposing a community-enhanced de-anonymizing scheme of social networks. Then, Lai [11] proposed to detect communities in social networks via user's interests and de-anonymize users in communities. Ji et al. also designed an Adaptive De-Anonymization framework for the scenario that the anonymized and auxiliary graphs have partial overlap [12]. Some papers modeled mobility traces as graphs and presented different attacks for de-anonymizing using online social networks as side channel [13], [14]. However, in heterogeneous social networks, this assumption may not always hold due to the fact that the users of different social networks may not always be overlapping. The diversity of usage patterns on different social networks will further render the inconsistency of the network structures of the different social networks. In our proposed method, we also exploit semantic publicly available information, such as user profile, to help de-anonymize users.

Besides, Ji et al. [15], [16] conducted the comprehensive quantification on the de-anonymizability of 24 real-world social networks with seed information in general scenarios. Later, in [17], a uniform and open-source secure graph data sharing/publishing system was proposed. Li et al. proposed a graph-based framework for privacy preserving data publish, which is a systematic abstraction of existing anonymity approaches and privacy criteria [18]. Qian et al. leveraged background knowledge graph to improve the de-anonymization performance [19]. But this work mainly focuses on de-anonymizing a graph anonymized from original graph and inferring some private attributes. Fu et al. proposed a graph node similarity measurement in consideration with both graph structure and descriptive information, and a deanonymization algorithm based on the measurement [20]. Zhang et al. targeted Twitter users in a metropolitan area by exploiting the strong geographic locality within communications on Twitter [42]. In our work, we try to de-anonymize heterogeneous social networks by considering both semantic information and structure information, and evaluate the privacy leakage after de-anonymization.

2.2 Profile based user matching

Public information and semantic information on social media or social network sites provide the evidence to match users of different social networks. Iofciu et al. used tags to identify users across social tagging systems such as Delicious, StumbleUpon and Flickr [21]. Olga et al. extracted features and developed supervised machine learning models which can perform entity matching between two profiles for a user by similar name and de-anonymizing a user's identity. [22] Goga et al. identified accounts

on different social network sites that all belong to the same user by exploiting only innocuous activity, such as location profiles, timing profiles, language profiles, that inherently comes with posted content [23]. Vosecky et al. identified users between Facebook and StudiVZ by exploiting various profile attributes [24]. Zafarani et al. [25] conducted an in-depth investigation of this problem by defining sophisticated features to model the behavior patterns of users in selecting usernames. Korayem et al. extracted four kinds of features, i.e. temporal activity similarity features, text similarity features, geographic similarity features, social connection similarity features, and apply machine learning techniques to find correct mapping [26]. Wondracek et al. introduced a technique that narrows down user identity by examining social-network group membership stolen from browsing history [27]. Zhang et al. [28] connected social networks users by considering both local and global consistency among multiple networks, but they treat both two consistencies as features and train an energy-based learning model. In [29] and [30], the first privacy-preserving personal profile matching scheme for mobile social networks was proposed by Li et al. In this scheme, an initiating user can be identified from a group of users the one whose profile best matches with his/her, with limited risk of privacy exposure. Later, two novel fine-grained private profile matching protocols were designed in [31], [32]. Different from these works, our proposed approach uses social structure to narrow down the candidate sets in order to achieve higher accuracy.

3 MOTIVATION AND MODELING

3.1 Motivation

As introduced above, the Novel Heterogeneous De-anonymization Scheme (NHDS) integrates both network structure and public information to de-anonymize social network users. On one hand, network structure and topology can be leveraged to de-anonymize social networks users with considerable accuracy as reported in state-of-the-art [9], [10]. However, this kind of approaches requires enough overlapping to ensure the accuracy by revisiting nodes and correcting early false mappings [9], thus can be less effective when the two networks are heterogeneous and not well aligned. On the other hand, profile information is useful to identify a person's different social networks accounts [24], [26], [42], but it may also cause many false positives in a large scale de-anonymization due to the same or similar profile attributes of different persons (e.g. same or similar nicknames of different persons). Since a social network contains a huge number of users, the possibility that different persons have the same or similar online profiles can be nonnegligible.

The first step of NHDS is that social network structure is leveraged to significantly reduce the size of node candidate sets. Then, user profile information is exploited to further identify the correct mapping nodes. Our motivation of this design is based on the observation that the graph structures (e.g., community, neighborhood) of social networks are likely to be similar for the same user groups [10], [11]. So community and neighborhood are employed to generate the candidates set and guide mapping process in a general view. Then in a specific view, a person is likely to have very similar profiles and generate same evidence on different social networks. Thus public available profile information is used to decide node mappings with a high confidence. Before presenting the scheme, we formulate the graph model, profile

information model, and attack model in remaining part of this section.

3.2 Graph Structure Modeling

Social network structure is usually represented as a graph, where each user is a node in the graph, and connections between a pair of users are represented as edges. Let $G = (V, E)$ represent a social network graph where V is a set of users and $E \subseteq V \times V$ is a set of directed/undirected links between users. $e(v_1, v_2)$ means that v_1 and v_2 are in friend relationship or follower/followee relationship where $e \in E, v_1, v_2 \in V$. As the important structures in the social network graph, *neighbor* and *community* are formally defined as follows:

Definition 1. A neighbor set R of a user $v_i \in V$ is a set of users $R_i = [v_j]_{j=1\dots}, \forall v_j : \exists e(v_j, v_i)$, who are directly in friend relationship or follow relationship with v_i . We further define a function α to form the neighbor set R_i of v_i , i.e. $\alpha(v_i) = R_i$.

Definition 2. A community C in a social network graph is a disjoint partition of vertices in $G(V, E)$. Formally, we denote communities in a graph as $\mathcal{C} = \{C_1, C_2, \dots, C_k\}$, where $C_i \neq \emptyset$ and $C_i \cap C_j = \emptyset$ if $i \neq j$ for $1 \leq i, j \leq k$. $\forall C_i \in \mathcal{C}, V_{C_i} \subset V$ and $E_{C_i} \subset V_{C_i} \times V_{C_i}$. In this paper, communities are defined by Infomap algorithm [33], which uses the probability flow of random walks on a network and decomposes the network into modules by compressing a description of the probability flow [33].

3.3 Profile Information Modeling

As platforms aiming at attracting attention, boosting self-presentation, promoting and sustaining social capital, social networks must allow part of user profile information to be available to public. The amount of publicly available profile information on social networks varies from each other, according to the platform defined and/or user defined privacy settings. For instance, Twitter allows users to follow other users without permissions and most profile are in public, while on Facebook, one user needs to send request to another for becoming 'friends', then the profile becomes visible. Meanwhile, a Facebook user might/might not choose to show his/her gender, status, hometown, on this platform. To exploit profile information across heterogeneous social networks, we firstly give a uniform definition:

Definition 3. Let $X_i = [x_{ik}]_{k=1\dots d}$ denote a set of attributes associated with the user $v_i \in V$ (for instance, username (screen name), location, self-description, etc), where d is the number of types of attributes and x_{ik} records the content of the k th attribute of user v_i . If a user v_i 's j th attribute is not available on the social network (e.g., a user chooses not to show her hometown on Twitter), $x_{ij} = \text{null}$. Note that a user may have several vectors modeling different profiles in social networks where he/she has an account. Common attributes between two vectors will be used to compute the profile similarity.

Since heterogeneous social network platforms contain different kinds of profile information, and some of them contains semantic or syntactic meaning, mapping two users' accounts from two heterogeneous social networks is similar to an ontology matching problem. In general, ontology matching determines an alignment for a pair of ontologies O_1 and O_2 . Each ontology consists of a set of discrete attributes which are usually represented

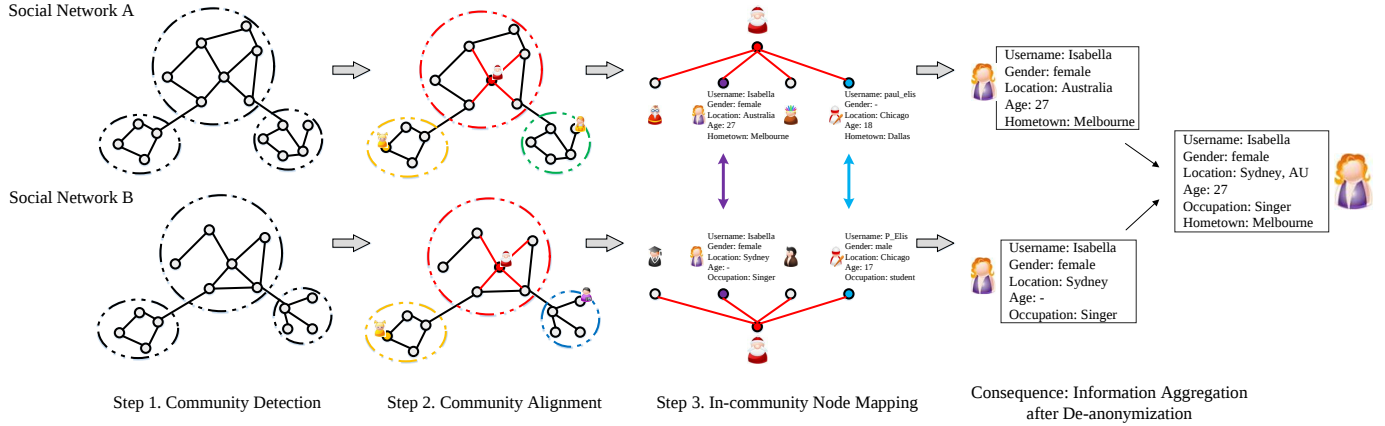


Fig. 1: Overview of our scheme

in the form of tables, classes, properties, and determines as output the relations. In our problem, profile matching can be defined as follows:

Definition 4. Given two profiles, $p_A = \{x_1, \dots, x_A\}$ and $p_U = \{x_1, \dots, x_U\}$. If $\text{type}(x_i) = \text{type}(x_j)$ for two attributes $x_i \in p_A$ and $x_j \in p_U$, the similarity between the two attributes is defined as:

$$\text{sim}_a = \text{matchScore}(x_i, x_j) \quad (1)$$

Here what the similarity is depends on which attribute is considered, it can be value similarity for ages or genders, string similarity for screen names, text similarity for descriptions or tweets, semantics similarity for locations or hometowns, etc., as presented in Section 4.4. Then the similarity of profiles is computed by:

$$\text{sim}_p = \frac{\sum_{r=1}^t w_r(\text{sim}_a)_r}{t} \quad (2)$$

where w_r is the weight given to attributes, and t is the number of attribute pairs of the same type between two profiles.

In Section. 4.4, we will show how to match users according to various profile attributes.

3.4 Attack Model

We assume two heterogeneous social network graphs G_A and G_U . G_A is denoted as the anonymous network graph and G_U is the auxiliary network graph. Note that the graph here is not necessary to be the whole graph of a social network, the de-anonymization attack can be conducted on partial graphs (i.e. subgraphs) collected by the attacker. That is to say, the attacker is able to obtain a subgraph $G = (V, E)$ and profile attributes X_i corresponding to $v_i \in V$ through published datasets or crawling sites. The goal of the attacker is to learn more information of the users across different networks by mapping users in G_A to users in G_U . To achieve this goal, the attacker needs to identify user accounts from two different social networks that belong to a same person in large scale and with a high confidence. This problem can be formally defined as follows.

Problem 1. Given (1) two different social network graphs $G_A = (V_A, E_A)$ and $G_U = (V_U, E_U)$, (2) sets of attributes X_i and X_j of $v_i \in V_A$ and $v_j \in V_U$ respectively, finding user mappings

$v_i \leftrightarrow v_j, v_i \in V_A, v_j \in V_U$ that belong to the same real persons accurately by iteratively computing:

$$\text{argmax}_{v_i \in \text{Cand.}_A, v_j \in \text{Cand.}_B} S(X_i, X_j) \quad (3)$$

where S is a function to compute similarity between X_i and X_j , as shown in Equ. 2. Cand._A and Cand._U are two candidate sets for potential correct mappings generated by community and neighbor structure in G_A and G_U , respectively.

4 NOVEL HETEROGENEOUS DE-ANONYMIZATION SCHEME

In this section, we introduce our proposed Novel Heterogeneous De-anonymization Scheme (NHDS).

4.1 Scheme Overview

Figure. 1 illustrates our proposed scheme which has three main steps: (1) Communities Detection: communities in both networks are detected according to graph structure, (2) Communities Alignment: seeds are automatically identified based on profiles, and communities that contain the same pairs of seeds are aligned, (3) In-community node mapping: in each pair of aligned communities, nodes with high similarity score, which is computed by profile similarity in Equ. 2, is accepted as a mapping, and mapping process is propagated to the neighbors. Algorithm. 1 presents the whole procedure, and the details and time complexity are introduced in the following sub-sections.

Algorithm 1 Algorithm of proposed scheme

Input : $G_A < V_A, E_A >, G_U < V_U, E_U >$, threshold θ

Output: Mappings of users μ'

//Communities detection

$\mathcal{C}_A = \text{Infomap}(G_A)$

$\mathcal{C}_U = \text{Infomap}(G_U)$

//Communities alignment

$\mu = \text{SelectSeeds}(V_A, V_B)$

$\text{CommPairs} = \text{AlignCommunities}(\mathcal{C}_A, \mathcal{C}_U, \mu)$

//In-community node mapping

$\mu' = \text{InCommunityMapping}(\text{CommPairs}, \mu, \theta)$

return μ'

4.2 Communities Detection

The goal of first step is to partition social network graphs G_A and G_U into two sets of communities $\mathcal{C}_A = \{c_1, \dots, c_m\}$ and $\mathcal{C}_B = \{c_1, \dots, c_n\}$. We devise the communities detection algorithm based on Infomap algorithm [33], which has a low time complexity, to partition disjoint, non-overlapping communities $\mathcal{C}_A$ and $\mathcal{C}_U$ for two graphs, respectively. In brief, Infomap finds the shortest multilevel description of the random walker therefore giving us the best hierarchical clustering of the network - the optimal number of levels and modular partition at each level - with respect to the dynamics on the network. So another merit of using Infomap algorithm is that it generates $\mathcal{C}_A, \mathcal{C}_U$ with different scales at different levels so that we can choose communities with similar scale for aligning. The algorithm for communities detection and division is denoted as the *Infomap*($\cdot$) function in Algorithm. 1, and the time complexity is $O(|E|)$.

4.3 Communities Alignment

For aligning communities $C_i \in \mathcal{C}_A$ and $C_j \in \mathcal{C}_U$, [10] proposes to treat each community as a node in a graph, then propagate the communities mapping process from some ‘community seeds’ using an improved version of [9]. However, in practice, we find that communities can be more easily aligned given the publicly available profile information. We choose the username (or screen name) to identify seeds for two reasons. Firstly, the username (screen name) must be available on every social network’s website, so the attacker has enough chances to obtain or crawl them. Secondly, as shown in [9], [26], the possibility that two accounts with the same usernames (screen names) do not belong to a user is less than 5%. Thus we design the following algorithm, which aligns $\mathcal{C}_A$ and $\mathcal{C}_U$ according to the number of same usernames in communities according to the algorithm described as the following two steps.

- 1) The first step is to find all user pairs with same usernames $\mu = \{\dots, (u_i, u_j)_k, \dots\}$ where $u_i \in V_A$ and $u_j \in V_U$. Greedy searching will cause a high complexity of $O(|V_A||V_U|)$. Instead, this process can be implemented by a hash table so that the time complexity can be reduced to $O(|V_A| + |V_U|)$. This procedure is denoted as *SelectSeeds*($\cdot$) function in Algorithm. 1.
- 2) In the second step, an initial confidence score $cs_{i,j}$ (that indicates whether two communities should be aligned) for each pair of communities (C_i, C_j) , where $C_i \in \mathcal{C}_A, 1 \leq i \leq m, C_j \in \mathcal{C}_U, 1 \leq j \leq n$, is set as 0. For each pair $(u_p, u_q) \in \mu$, $cs_{i,j}$ is added by one, given $u_p \in C_i$ and $u_q \in C_j$. Then, all confidence scores cs for communities pairs are examined, if $cs_{i,j}$ exceeds a threshold θ_{cs} , C_i and C_j are aligned. The time complexity of this step is $O(|\mu|) < O(|V_A| + |V_U|)$. This procedure is denoted as *AlignCommunities*($\cdot$) function in Algorithm. 1.

The overall complexity of communities alignment algorithm is $O(|V_A| + |V_U|)$, as described above. Our overall evaluations show that our communities division and alignment only slightly reduce the recall rate.

4.4 In-community Node Mapping

Algorithm. 2 describes our in-community node mapping algorithm. Within each pair of aligned communities, a propagating

and mapping algorithm (in *Propagation*($\cdot$)) is performed locally. This algorithm takes two graphs of communities $G_{c1} = (V_{c1}, E_{c1})$, $G_{c2} = (V_{c2}, E_{c2})$ and the set of seeds in these communities $\mu_{c1,c2}$ selected in the previous step as input. It iteratively finds new mappings in the neighbor sets of seeds in $\mu_{c1,c2}$, and extends mapping process based on graph structure. At each iteration, the algorithm computes similarity scores within two neighbor sets $R_u = \alpha(u)$ and $R_v = \alpha(v)$, which are generated by two already mapped users u and v . It picks a user r_u in R_u and computes similarity score with users in R_v and find out a r_v with the highest score. The similarity score is computed by the *MatchScore*($\cdot$) function in Algorithm. 2, which will be discussed later. If the score exceeds a pre-defined threshold θ , r_u and r_v are accepted as a successful mapping. If an already mapped node is mapped to another node with a higher similarity score, the previous mapping is replaced by the new mapping with higher score. The process is halted if no new mapping is explored, and unvisited nodes in propagation process of both communities are gathered to compare to find remaining mappings. The time complexity of propagation is $O(\min\{|V_{c1}|, |V_{c2}|\}d_{c1}d_{c2})$, where d_{c1} and d_{c2} are bounds on the degree of the nodes in V_{c1} and V_{c2} , respectively.

Algorithm 2 Algorithm of the InCommunityMapping($\cdot$)

Input : community pairs *CommPairs*, seeds μ , threshold θ

Output: μ' with more mappings of users

for $(C_a, C_u)_j \in \text{CommPairs}$ **do**

$\mu_j = \text{Propagation}(C_a, C_u)$

end

return $\mu' = \bigcup_{j=1, \dots, \text{len}(\text{CommPairs})} \mu_j$

Procedure *Propagation*(R_1, R_2)

$\mu_j \subset \mu$ //the seeds set of $(C_a, C_u)_j$

while *exists* $\langle v_1, v_2 \rangle \in \mu_j$ *is unvisited* **do**

$R_1 = \alpha(v_1), R_2 = \alpha(v_2)$

for r_1 *in* R_1 **do**

for r_2 *in* R_2 **do**

$\text{scores}[r_1].\text{add}(\text{MatchScore}(r_1, r_2))$

end

if $\text{MAX}(\text{scores}[r_1]) > \theta$ **then**

$r_{\text{max}} = \text{user with MAX}(\text{scores}[r_1])$

 add $\langle r_1, r_{\text{max}} \rangle$ into μ_j and mark *unvisited*

end

end

 Mark $\langle v_1, v_2 \rangle$ *visited*

end

return μ_j

To compute the similarity score between two users (nodes), profile matching are exploited. Due to the variety of profile attributes, *MatchScore*($\cdot$) implements an “if-then” rule to compute similarity scores of different kinds of profile attributes in different methods. After that, an overall similarity score of the two users is given by assigning weights empirically to score of different attributes as the form of Equ. 2. Different techniques are leveraged to compute similarity score of different kinds of attributes.

4.4.1 Value Matching

Direct value matching can be used to match non-string literal attributes, such as gender, birth date, and personal website (or

link). This kind of attributes usually have fixed formats on social networks, which indicates high confidence of rejecting (if genders of two users are different) or accepting (if birth date or personal websites of two users are same) a potential mapping.

4.4.2 Syntactic Matching

Syntactic matching is applied to attributes that are usually shown as strings (e.g. username and person name). These attributes on different social networks often have editing differences, such as difference among “Jones, David”, “David Jones”, and “D. Jones”. So string matching metric can be used for syntactically matching these attributes. In order to avoid the influence of abbreviation or acronym, Monge-Elkan algorithm, a recursive string matching algorithm, is applied [34]. The basic idea of this method is to break input string into tokens. Then the best matching token are compared to get the score as follows.

$$MongeElkan(A, B) = \frac{1}{|A|} \sum_{i=1}^{|A|} \max_{j=1}^{|B|} \{dist(A_i, B_j)\} \quad (4)$$

where A and B are two strings, and $dist()$ refers to a secondary distance function used to compute similarity between tokens of A and B . In a lot of functions computing edit-distance, *Jaro-Winkler similarity* is chosen as the secondary distance function in our problem, due to its remarkable performance in previous research on name-matching tasks [35]. Monge-Elkan algorithm returns 1 if two string are fully matched or one abbreviates the other; return 0 if there is no match between the two strings.

4.4.3 Keywords Matching

Keywords matching is a tool to handle attributes that are in the form of texts (e.g. self-description of users or ‘aboutme’). Since the contents a person posts and the words the person uses are likely to be similar, the similarity of texts, which is usually measured by some keyword-based matching methods, are used to determine whether two texts come from a same person. Given two texts, we first compute TF-IDF weights, which normalize each word count by the number of people that used it (in the comparison set, say, neighbour set R mentioned above), to get two weights vectors d and e , respectively. Then we compute the cosine similarity between d and e as the text similarity score:

$$Cosine(d, e) = \frac{d \cdot e}{\|d\| \|e\|} = \frac{\sum_{i=1}^n d_i e_i}{\sqrt{\sum_{i=1}^n d_i^2} \sqrt{\sum_{i=1}^n e_i^2}} \quad (5)$$

4.4.4 Semantic Matching

Previous three matching methodologies are based on similarity of string or text. However, some of attributes, such as locations and hometown, might be semantically equal when the term or words are totally different. For example, location attribute on Flickr is referred as the country where a user is, where the location attribute on Myspace can be filled in freely by the user. So a user who lives in Michigan State of America might have location of “America” on Flickr and “Michigan” on Myspace. If we only consider string similarity, “America” and “Michigan” have a large edit distance, and might be considered to be from different users literally. But they are semantically related and possible to be from the same user. In order to solve this problem, we use GeoNames database [36] to check whether two locations are semantically related. If they are actually a same location, we give a high score;

if they are in inclusion relationship (say, a city in a state or a state in a country), a medium score is given; if they are two independent locations, a low score is assigned.

With these four methods for computing similarity of attributes, $MatchScore(\cdot)$ is able to compute the similarity score of users, and return it to $Propagation(\cdot)$.

5 EVALUATIONS OF PROPOSED SCHEME

In this section, we evaluate our proposed NHDS scheme by conducting experiments on a set of real-world social networks data.

5.1 Datasets

The datasets of four real-world heterogeneous online social networks, i.e., Livejournal, Flickr, Last.fm, and Myspace, are obtained from [28]. The datasets include node information, edge information, and profile information of a subset of users of these social networks.

- *LiveJournal* is a social networking site and blogging platform that allows users to find each other through journaling and interest-based communities. The dataset consists of 3,017,286 users and 19,360,690 friend relationship.
- *Flickr* is an image hosting online community for sharing, storing, and organizing photos. The dataset consists of 215,495 users and 9,114,557 friend relationship.
- *Last.fm* is the world’s largest online music catalogue and has been recognized as a popular social network for music enthusiasts. Last.fm builds detailed profiles of users’ musical tastes and preferences. The dataset consists of 136,420 users and 1,685,524 friend relationship.
- *MySpace* is a social networking website offering an interactive, user-submitted network of friends, personal profiles, blogs, groups, photos, music, and videos. The dataset consists of 854,498 individuals and 6,489,736 friend relationship.

We build undirected social network graphs according to ‘friend’ or ‘follow’ relationship in these social networks. The statistics of the graphs are shown in Table. 1. These social networks not only provide different services and have different utility, but also have different graph properties. For example, Flickr has an average degree of 85.59 while the average degree of Livejournal is only 15.19. The heterogeneous structure increases the difficulty of de-anonymization.

In order to evaluate the results, we obtain the ground truth data from [28], [37], which contain pair-wise matched user id of two social networks. The data were originally collected by Perito et al [37] through Google Profiles service by allowing users to integrate different social network services.

TABLE 1: Statistics of social networks

Network	Nodes	Edges	Av. Degree
Livejournal	3,017,286	19,360,690	12.83
Flickr	215,495	9,114,557	85.59
Last.fm	136,420	1,685,524	24.71
Myspace	854,498	6,489,736	15.19

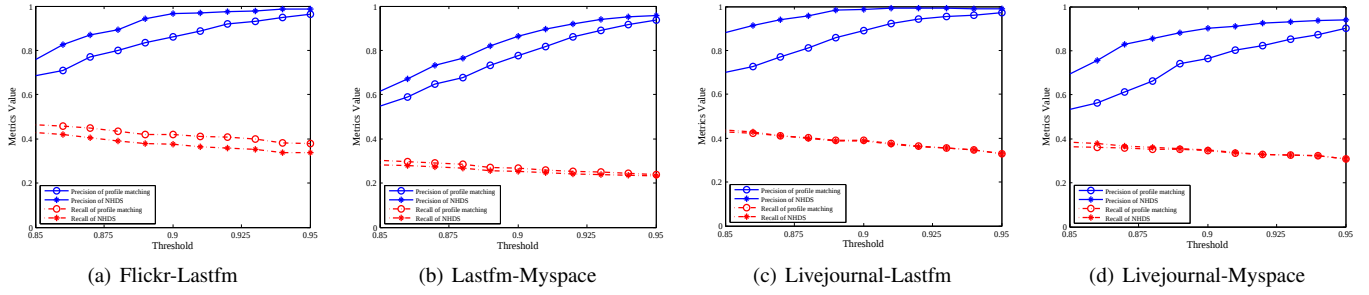


Fig. 2: Performance comparisons between profile-based matching and proposed NHDS

5.2 Experiments

Since our scheme is a combination of graph structure and publicly available profile information, we evaluate the results by comparing our approach with approaches that only exploit profile information, and approaches that are only based on graph structure, respectively. To quantitatively evaluate the algorithm, we consider the two widely-used metrics:

- **Precision:** In all mappings returned by the de-anonymization algorithm, the percentage of correct mapping. Since our goal is to find out correct mappings rather than find out incorrect mappings, the concept of precision here is same as the concept of accuracy.
- **Recall:** The percentage of correct mapping retrieved by algorithm in all mappings in ground truth.

The codes of experiments are written in Python and the programs are ran on a server with Intel[®] Xeon[®] 2.4GHz 14-core CPU and 64GB memory.

5.2.1 Evaluations of NHDS

As introduced above, the θ_{cs} and θ are the only two setting parameters of our scheme, so we set $\theta_{cs} = 1$ and vary θ to evaluate the results in the following parts. Due to the users' awareness of privacy protection and social networks' privacy settings, few common profile attributes are available for all users across multiple social network graphs. In order to evaluate our proposed NHDS scheme, we select username (screen name), which is widely available for all users of all social networks in our datasets, as profile information, and perform overall evaluations firstly. Then, comparative experiments by considering limited graph structures and more profile attributes are conducted on part of users to show that the performance under different situation. The direct profile-based matching represented by [24], i.e. computing profile similarity between each user in one social network and all users in the other social network and find the most similar one, is used as the baseline.

Fig. 2 shows how our NHDS scheme outperforms the profile-based matching by tuning threshold θ . The precision of our approach is obviously higher than the baseline by slightly sacrificing the recall. It reflects that graph structures (community/neighborhood) are useful to filter incorrect matchings, thus increasing the precision. As introduced in Section 4.4, the threshold θ is a similarity criterion that accepts a pair of nodes as a mapping in our algorithm, i.e. if the similarity score between two nodes exceeds the θ , the two nodes are accepted as a potential mapping. So, the higher θ is set, the more similar the accepted nodes are, and thus fewer potential mappings will be returned.

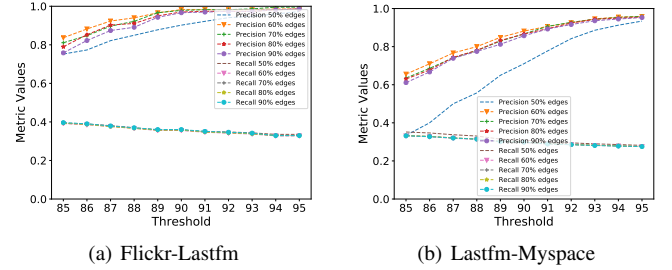


Fig. 3: Performance after randomly removing edges. In the legend, for instances, “90% edges” means 90% edges are remaining in the graph, which also means 10% edges are removed

So θ actually reflects the trade-off between precision and recall. An attacker can choose θ according to his/her requirement of this trade-off in practice. When the threshold is set to 0.9, the precision of matching users can be more than 90% with a recall of 40% for Flickr-Lastfm and Livejournal-Lastfm. The results show that large scale accurate de-anonymization can be performed.

Impacts of Graph Structure. Sometimes it is difficult to access a full view of a social network graph due to access limitations and privacy policies. Especially from the attacker's perspective, data that the attacker can collect are usually processed data where proper noise has been added. So it is interesting to investigate whether the attack is still feasible when limiting an attacker to a restricted view on the graph. We simulate different portions of graph to be analyzed by randomly removing different percentages (i.e., 10%, 20%, 30%, 40%, and 50%) of edges from original data. As shown in Fig. 3, removing certain portions (e.g., 10%-40%) of edges will not obviously decrease the precision. After removing 50% of edges, precision noticeably drops because the graph structure is significantly affected and cannot act as the guide to reduce false positive. Meanwhile, the recall is almost not affected. These results show that, though randomly removing of edges disturbs the graph structure, the attacker is still likely to leverage profile attributes to precisely identify correct mappings. The comparison between graph based approach and our approach in Section 5.2.2 also shows involving profile attributes can help identify mappings when graph structure is not well aligned.

Impacts of Multiple Profile Attributes. Then, we consider more profile attributes, including username, nickname, status, gender, location, and aboutme, to de-anonymize users between Flickr and Myspace, because the two social networks have the most available common profile attributes. We set the weights as $w(\text{username}) = 0.2, w(\text{nickname}) = 0.2, w(\text{status}) = 0.1, w(\text{location}) = 0.2, w(\text{gender}) = 0.2, w(\text{aboutme}) =$

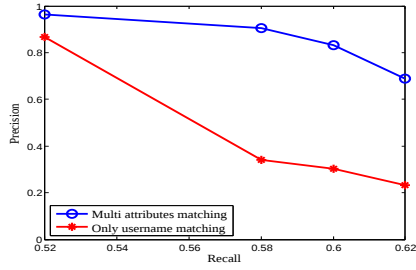


Fig. 4: Comparison of multi-profiles and username matching

0.1 and compute similarity scores as proposed in Section 4.4. The performance is compared with the scenario that only username is used to calculate similarity scores. Fig. 4 shows that when more profile attributes are considered between Flickr and Myspace, the precision is improved when the recall is kept the same. Similar conclusions can also be drawn on other social networks. For instances, 46% precision by considering name, location, links, aboutme v.s. 40% precision by considering username only when recall are both 29% for de-anonymization between Livejournal and Lastfm.

5.2.2 Comparisons with profile-based approaches

Previous studies exploit various user profile information to connect individuals between social networks, including usernames [24], tags [21], activities [23], group membership [27], and multiple kinds of profile attributes [25], [26], [28], [39], [41]. In order to reflect properties of our approach, NHDS is further compared with some existing profile-based approaches, including direct profile matching [24], MNA [39], SiGMA [41], and COSNET [28]. Fig. 6 shows the precision and recall of these approaches for instances. We can clearly observe that the precision and the recall is always a trade-off. This is because, if targeting at a higher precision, an algorithm must have strict criteria of identifying a mapping to ensure its correctness, thus resulting in lower recall; if targeting at a higher recall, i.e., more mappings are returned, the algorithm has to loose criteria to accept more possible mappings, as well as more potential false positive. Compared with other algorithms, NHDS provides the highest precision though sacrificing the recall. We believe that given a low precision, the high recall is nonsense - the attacker still cannot be confident about whether obtained mappings are correct. So our approach has its benefits for providing the most accurate result.

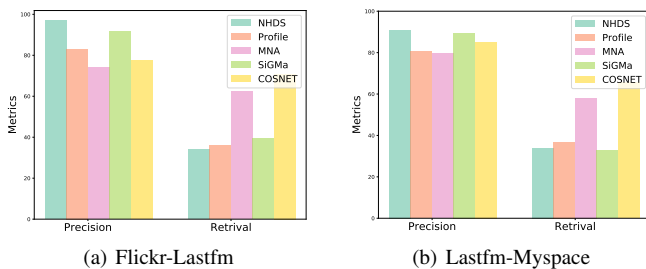


Fig. 6: Comparisons of profile based approaches

5.2.3 Comparisons with graph-based approaches

As mentioned above, numerous graph-based de-anonymizing algorithms have been proposed. However, only a few of them are

suitable to real-world heterogeneous social networks for various reasons. Some techniques are constrained by their restrict requirements of social networks of the same size (or same number of nodes) [6], [8], sybil users [38] or high computation capability for large scale networks [13], while others have only been evaluated between noisy graph and its original graph [10], [12], [19]. For reference, we test the well-known graph-based NS algorithm proposed by Narayanan and Shmatikov [9] and percolation-based de-anonymization algorithm [39] using the open-source evaluation system proposed in [17]. As a result, only few correct mappings are reported by the two previous algorithms on the heterogeneous social network datasets feeding more than 100 seeds. One possible reason of the unsatisfactory performance is pure graph-based approaches require enough overlaps of the network graphs to propagate and correct false mappings at the beginning of the mapping [9]. But it is usually difficult to obtain datasets with ideal overlaps from two heterogeneous social networks, which limits the performance of graph-based approaches in practice. According to [9], 30.8% of the mappings were re-identified correctly between a Twitter dataset (Av. degree of 37.7) and a Flickr dataset (Av. degree of 32.2), which is far away from our results on more heterogeneous networks datasets. The results show that introducing profile attributes of nodes obviously increases the successful rate of de-anonymizing.

6 REAL-WORLD PRIVACY LEAKAGE EVALUATION

We have shown that our approach is able to de-anonymize users accurately in a large scale. As shown in the last step in Figure. 1, aggregated profiles from de-anonymization reveal more privacy of the user. To understand the severity of de-anonymization attack and provide reference of social network privacy protection, it is worth to investigate to what extent the user's privacy will be exposed. To quantify the user privacy leakage after de-anonymizing, we retrieve de-anonymized users returned by our approach as described in Section. 5.2.1 ($\theta = 0.9$), and quantify real-world information leakage through de-anonymization.

6.1 Profile appearance on different social networks

Through the profile aggregation of accounts from different social networks, more previously unknown information of the same person can be obtained. For the ease of presentation, we firstly introduce and define two types of the previously unknown information which can be gained after aggregation as follows.

6.1.1 Platform preserved information

Different social networks contain different profile information, which can be exposed in public, according to the platform settings and utility. Fig. 5 illustrates an example of user profile setting pages of Flickr and Lastfm, where users can choose to fill in this page to demonstrate their profiles to the public. We can see that the Flickr allows users to show more profile information, including the singleness, occupation, hometown, current city, than Lastfm, as highlighted in red in Fig. 5(a). As a result, these profile information can be exposed through Flickr, while being preserved by Lastfm.

Similarly, Table. 2 lists profile attributes revealed on social networks contained in our datasets. For example, Flickr provides users' gender and location on the platforms, while gender and occupation of users are available in Myspace. Also for the location

(a) Flickr's profile setting page

(b) Lastfm's profile setting page

Fig. 5: An example to illustrate platform preserved information

TABLE 2: User profile revealed on social networks

	gender	age	links	status	interests	location	hometown	education	occupation	aboutme	orientation, income, etc.
Livejournal			✓		✓	✓		✓		✓	
Flickr	✓		✓	✓		✓	✓		✓	✓	
Lastfm	✓	✓	✓			✓				✓	
Myspace	✓	✓	✓	✓	✓		✓	✓	✓	✓	✓

information in Flickr and occupation information in Myspace, only one of the two platforms contains this information, while the other one preserves it. So, we denote this kind of information as *platform preserved information*. Given pairs of users de-anonymized between two social networks, the attackers might know more kinds of profile attributes of users, say the location information of Myspace users or occupation information of Flickr users. The aggregations of platform preserved information is coined as *platform preserved information aggregation*, which aggregate attributes that one platform is preserved through de-anonymizing from another platform.

6.1.2 User preserved information

On the other hand, though two social networks contain common profile information settings (e.g., the gender information appeared on Flickr and Myspace in our example), some profile information can still be preserved from the public by the users. That's because a person may choose not to fill in and show all his/her profile information online. Fig. 7 shows online profiles of two real-world Flickr users (important information is blurred in order to preserve users' privacy). Compared with the user B, the user A chooses to show his/her contact email, hometown, and current city to public while hides his/her occupation. We denote this kind of information as *user preserved information*. The attackers may have the chances to uncover information that the users preserved on a social network through de-anonymizing and aggregation. For example, the attacker can still collect the Flickr user A's occupation from Myspace or the user B's address from other social networks. We call it as *User preserved profile uncovering*.

In summary, the platform preserved information is the information that attackers can not obtain from this social network platform but can be collected from other social networks, due to the platform settings. And the user preserved information is

(a) Flickr profile of user A

(b) Flickr profile of user B

Fig. 7: An example to illustrate user preserved information

the information that users choose not to show on the platform though they have the options to show. We will quantify these two information leakage in the following parts.

6.2 Platform preserved information aggregation

Firstly, we evaluate the information leakage from *platform preserved profile aggregation*, i.e., aggregation of different kinds of attributes between two social networks. We compute the percentage of attributes exposure of de-anonymized users from two heterogeneous social networks. Table. 3 shows the attributes that are available on one platform but not available on the other platform, due to the platform settings, and corresponding exposed ratio. The exposed ratio of an attribute is calculated by the percentage of users who show this attribute online within all the de-anonymized users. From Table. 3, we can observe various personal information, e.g., orientation, income, ethnicity, are possible to be exposed from Myspace, while almost half users in Flickr reveal their locations and occupation. So the attackers are possible to obtain more complete user profiles and their locations after de-anonymizing. Similarly, a de-anonymized Livejournal user's gender and age are likely to be collected through its matched Lastfm account. So

de-anonymization across heterogeneous social networks makes it possible to learn users information from more aspects, and construct the more comprehensive profiles of users.

TABLE 3: Platform preserved information leakage

De-anonymization	Source	Attributes	Exposed Ratio
Flickr-Myspace	Flickr	occupation	45.52%
		location	51.36%
		age	49.89%
	Myspace	interests	25.27%
		orientation	21.09%
		bodytype	15.60%
		ethnicity	17.14%
		religion	22.41%
		children	19.56%
		smokedrink	16.70%
		education	18.90%
		occupation	20.87%
		income	5.71%
		schools	28.13%
		email	3.86%
Livejournal-Lastfm	Livejournal	schools	36.48%
		interests	62.66%
		birthdate	74.24%
		age	79.16%
	Lastfm	gender	100.00%

6.3 User preserved information uncovering

Then, we evaluate the information leakage of *uncovering user preserved attributes* - attributes that are available on the settings of both two social networks but users might not fill in on both sites. For the convenience of explanation, we first denote some concepts. A common attribute that is possible to be shown on two platforms can be classified into three types according to the users' settings:

- *Known attribute* means the attribute that users show on both platforms. So the attackers have already known the attribute of the users without de-anonymization from an auxiliary network.
- *Unknown attribute* means the attribute that users did not show on both platforms. So the attackers can not obtain the attribute even though two accounts are matched.
- *De-anonymized attribute* means attribute that users show on one site but not on the other. So the attacker can collect the attribute of users through de-anonymizing.

For example, as shown in Table. 2, both Flickr and Myspace provide settings of gender, hometown, and occupation to users. If a user has his gender on the both platforms, shows hometown on Flickr but not on Myspace, and doesn't show occupation on both sites, gender is a *known attribute*, hometown is a *de-anonymized attribute*, and occupation is a *unknown attribute*, as defined above.

Then, in order to evaluate how many users' attributes can be obtained after de-anonymization, we further define different portions of users for a specific attribute:

- *Known portion* means percentage of users who show the attribute on both platforms.
- *Unknown portion* means percentage of users who did not show the attribute on both platforms.
- *De-anonymized portion* means percentage of users who show the attribute on one platform but not on the other.

Fig. 8 shows the comparison of de-anonymized portions and known portions of attributes in different set of de-anonymization

experiments (Livejournal-Lastfm, Livejournal-Myspace, Flickr-Myspace, Lastfm-Myspace from top group of bars to bottom group of bars). Some attributes, such as *links*, *hometown*, *interests*, contain high de-anonymized portions, which exceed 45%. And the average percentage of de-anonymized portion (Equation. 6), which represents the information gained from de-anonymization, is 39.9%. It indicates notable information leakage through de-anonymization.

$$Info_Gain = \frac{de-anonymized\ portion + known\ portion}{known\ portion} - 1 \quad (6)$$

Furthermore, we evaluate how much previously invisible information the attackers can obtain from the de-anonymization attack. Since the known portion of attribute is already visible for the attacker, we define the de-anonymized ratio as:

$$Radio = \frac{de-anonymized\ portion}{1 - known\ portion} = \frac{de-anonymized\ portion}{de-anonymized\ portion + unknown\ portion} \quad (7)$$

Fig. 9 shows the de-anonymized ratio of attributes in different de-anonymization experiments (Livejournal-Lastfm, Livejournal-Myspace, Flickr-Myspace, Lastfm-Myspace from top group of bars to bottom group of bars). The average ratio is up to 0.84. The result shows that the attackers can obtain a great portion of previously unknown profile information through de-anonymization.

Our quantified evaluation shows that the privacy leakage through de-anonymization attack across real-world social networks is severe. Privacy preserving strategies and mechanisms for both protecting privacy of users and maintaining social network utility are still open research problem.

7 CONCLUSION

In this paper, we propose a practical Novel Heterogeneous De-anonymization Scheme (NHDS) for de-anonymizing real-world heterogeneous social networks, and evaluate and quantify the following privacy leakage. NHDS is a de-anonymizing scheme that exploits the network graph structure to significantly reduce the size of candidate set, and use user profile information to identify users with a high confidence. The performance evaluations of NHDS based on a dataset of four real-world social networks show that it achieves a high precision with a slight sacrifice of recall. We further quantify privacy leakage through de-anonymization. Evaluations show that notable portions of user information is disclosed. Privacy preserving in social networks is still an open challenge.

ACKNOWLEDGMENTS

This work is supported by National High-Tech R&D (863) Program (no. 2015AA01A707)

REFERENCES

- [1] D. Chaffey, "Global social media research summary 2016", 2016. <http://www.smartinsights.com/social-media-marketing/social-media-strategy/new-global-social-media-research/>
- [2] M. Duggan, N. Ellison, C. Lampe, A. Lenhart and M. Madden, "Social Media Site Usage 2014, Pew Research Center", 2015. <http://www.pewinternet.org/2015/01/09/social-media-update-2014/>.
- [3] W. Meng, R. Ding, S. P. Chung, S. Han, and W. Lee, "The Price of Free: Privacy Leakage in Personalized Mobile In-App Ads", In *NDSS*, 2016.

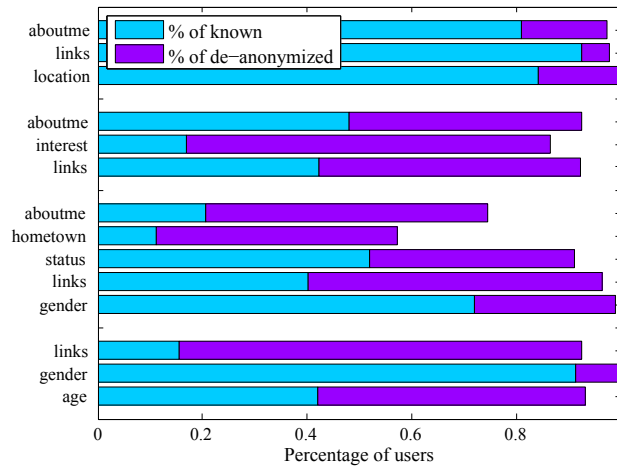


Fig. 8: Information gain through de-anonymization

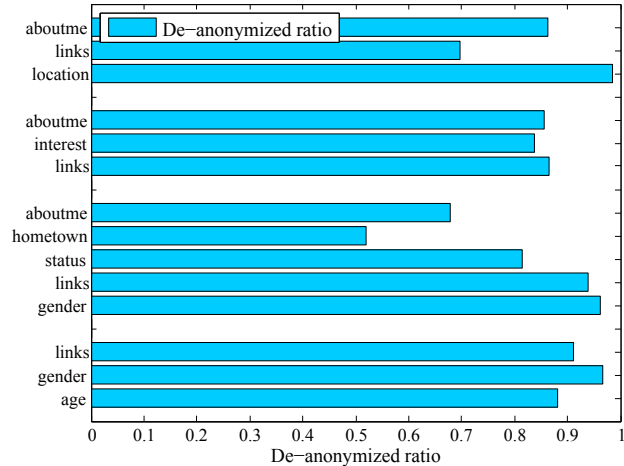


Fig. 9: De-anonymized ratio of attributes

- [4] H. Li, H. Zhu, S. Du, X. Liang, and X. Shen, "Privacy leakage of location sharing in mobile social networks: Attacks and defense", In *IEEE Transactions on Dependable and Secure Computing* vol: PP, Issue: 99, pp: 1-1, 2016.
- [5] A. M. Vegni, V. Loscri, "A survey on vehicular social networks," *IEEE Communications Surveys & Tutorials*, 17(4), 2397-2419, 2015.
- [6] N. Korula and S. Lattanzi, "An efficient reconciliation algorithm for social networks", *Proceedings of the VLDB Endowment*, 7(5), 377-388, 2014.
- [7] Z. Zhang, Q. Gu, T. Yue, and S. Su, "Identifying the same person across two similar social networks in a unified way: Globally and locally" In *Information Sciences*, 394, 53-67, 2017
- [8] P. Pedarsani, D. R. Figueiredo, and M. Grossglauser, "A bayesian method for matching two similar graphs without seeds". In *51st Annual Allerton Conference on Communication, Control, and Computing (Allerton)*, (pp. 1598-1607), 2013.
- [9] Narayanan A, Shmatikov V, "De-anonymizing social networks", In *30th IEEE Symposium on Security and Privacy*, (pp. 173-187), 2009.
- [10] Nilizadeh S, Kapadia A, Ahn Y Y, "Community-enhanced de-anonymization of online social networks", In *Proceedings of the 2014 ACM SIGSAC Conference on Computer and Communications Security*, (pp. 537-548), 2014.
- [11] S. Lai, H. Li, H. Zhu, N. Ruan, "De-anonymizing Social Networks: Using User Interest as a side-channel", In *2015 IEEE/CIC International Conference on Communications in China (ICCC)*, (pp. 1-5), 2015.
- [12] S. Ji, W. Li, M. Srivatsa, J. He and R. Beyah, "Structure based Data De-anonymization of Social Networks and Mobility Traces", *Springer Information Security*, 237-254, 2014.
- [13] M. Srivatsa and M. Hicks, "De-anonymizing mobility traces: Using social networks as a side-channel", In *Proceedings of the 2012 ACM conference on Computer and communications security*, (pp. 628-637), 2012.
- [14] S. Ji, W. Li, M. Srivatsa, J. He, and R. Beyah, "General graph data de-anonymization: From mobility traces to social networks", In *ACM Transactions on Information and System Security (TISSEC)*, 18(4), 12, 2016.
- [15] S. Ji, W. Li, N. Gong, P. Mittal, and R. Beyah, "On your social network de-anonymizability: Quantification and large scale evaluation with seed knowledge", *NDSS*, 2015.
- [16] S. Ji, W. Li, N. Gong, P. Mittal, and R. Beyah, "Seed-Based De-Anonymizability Quantification of Social Networks", In *IEEE Transactions on Information Forensics and Security*, 11(7), 1398-1411, 2016.
- [17] S. Ji, W. Li, P. Mittal, X. Hu, and R. Beyah, "SecGraph: A Uniform and Open-source Evaluation System for Graph Data Anonymization and De-anonymization", In *USENIX Security*, (pp. 303-318), 2015.
- [18] X.-Y. Li, C. Zhang, T. Jung, J. Qian, and L. Chen, "Graph-based privacy-preserving data publication", In *The 35th Annual IEEE International Conference on Computer Communications (INFOCOM)*, (pp. 1-9), 2016.
- [19] J. Qian, X.-Y. Li, C. Zhang, and L. Chen, "De-anonymizing Social Networks and Inferring Private Attributes Using Knowledge Graphs", In *The 35th Annual IEEE International Conference on Computer Communications (INFOCOM)*, (pp. 1-9), 2016.
- [20] H. Fu, A. Zhang, X. Xie, "Effective Social Graph De-anonymization Based on Graph Structure and Descriptive Information". In *ACM Transactions on Intelligent Systems and Technology* 6(4): 49:1-49:29, 2015.
- [21] T. Iofciu, P. Fankhauser, F. Abel, and K. Bischoff, "Identifying Users Across Social Tagging Systems", *ICWSM'11*, 2011.
- [22] O. Peled, M. Fire, L. Rokach, and Y. Elovici, "Matching entities across online social networks", *Neurocomputing*, 210, 91-106, 2016.
- [23] O. Goga, H. Lei, S. Parthasarathi, G. Friedland, R. Sommer and R. Teixeira, "Exploiting innocuous activity for correlating users across sites", In *Proceedings of the 22nd international conference on World Wide Web (WWW)*, (pp. 447-458), 2013.
- [24] J. Vosecky, D. Hong, and Y. Shen "User identification across multiple social networks", *IEEE Networked Digital Technologies (NDT)*, (pp. 360-365), 2009.
- [25] R. Zafarani and H. Liu, "Connecting users across social media sites: A behavioral-modeling approach", In *Proceedings of the 19th ACM SIGKDD international conference on Knowledge discovery and data mining*, (pp. 41-49), 2013.
- [26] M. Korayem and D. Crandall, "De-Anonymizing Users Across Heterogeneous Social Computing Platforms", *ICWSM'13*, 2013.
- [27] Wondracek G, Holz T, Kirda E, et al. "A practical attack to de-anonymize social network users", In *IEEE Symposium on Security and Privacy*, (pp. 223-238), 2010.
- [28] Y. Zhang, J. Tang, Z. Yang, J. Pei, and S. Yu, "COSNET: Connecting Heterogeneous Social Networks with Local and Global Consistency", *Proceedings of the 21th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, (pp. 1485-1494), 2015.
- [29] M. Li, N. Cao, S. Yu, and W. Lou, "FindU: Privacy-Preserving Personal Profile Matching in Mobile Social Network", In *Proceedings of IEEE INFOCOM*, (pp. 2435-2443), 2011.
- [30] M. Li, S. Yu, N. Cao, and W. Lou, "Privacy-Preserving Distributed Profile Matching in Proximity-based Mobile Social Networks", In *IEEE Transactions on Wireless Communications*, vol.12, no.5, 2013.
- [31] R. Zhang, Y. Zhang, J. Sun, G. Yan, "Fine-grained private matching for proximity-based mobile social networking", In *Proceedings of IEEE INFOCOM*, (pp. 1969-1977), 2012.
- [32] R. Zhang, Y. Zhang, J. Sun, G. Yan, "Privacy-Preserving Profile Matching for Proximity-Based Mobile Social Networking", In *IEEE Journal on Selected Areas in Communications*, vol. 31, no. 9, pp. 656-668, 2013.
- [33] M. Rosvall and C. T. Bergstrom, "Maps of random walks on complex networks reveal community structure", In *Proceedings of the National Academy of Sciences*, 105(4):1118-1123, 2008.
- [34] A. E. Monge, and E. Charles, "The Field Matching Problem: Algorithms and Applications", *Proceedings of the ACM SIGKDD international conference on Knowledge discovery and data mining*, (pp. 267-270), 1996.
- [35] W. Cohen, P. Ravikumar and S. Fienberg, "A comparison of string metrics for matching names and records." In *Kdd Workshop on Data Cleaning and Object Consolidation*, Vol. 3, 73-78, 2003.
- [36] GeoNames. <http://geonames.org/>. Retr. June 17, 2009.
- [37] D. Perito, C. Castelluccia, M. A. Kaafar, and P. Manils, "How unique and traceable are usernames?" In *International Symposium on Privacy Enhancing Technologies Symposium*, (pp. 1-17), 2011.

- [38] L. Backstrom, C. Dwork and J. Kleinberg, "Wherefore art thou r3579x?: anonymized social networks, hidden patterns, and structural steganography", In *Proceedings of the 16th international conference on World Wide Web (WWW)*, (pp. 181-190), 2007.
- [39] L. Yartseva and M. Grossglauser. "On the performance of percolation graph matching", In *Proceedings of the first ACM conference on Online social networks*, (pp. 119-130), 2013.
- [40] X. Kong, J. Zhang, and S. Y. Philip. "Inferring anchor links across multiple heterogeneous social networks", In *Proceedings of Conference on Information and Knowledge Management*, pp. 179188, 2013.
- [41] S. Lacoste-Julien, K. Palla, A. Davies, G. Kasneci, T. Graepel, and Z. Ghahramani. "Sigma: Simple greedy matching for aligning large knowledge bases", In *Proceedings of the 19th ACM SIGKDD international conference on Knowledge discovery and data mining*, pp. 572580, 2013.
- [42] J. Zhang, J. Sun, R. Zhang, and Y. Zhang, "Your actions tell where you are: Uncovering Twitter users in a metropolitan area," In *Proceedings of IEEE Conference on Communications and Network Security (CNS)*, pp. 424-432, 2015.



Huaxin Li is a graduate student working towards his M.Sc. degree in Department of Computer Science and Engineering, Shanghai Jiao Tong University. He received the B.Sc. degree in Department of Computer Science and Engineering, Shanghai Jiao Tong University, China, in 2011. His research interests include social networks privacy, smartphone security, network security and privacy, and machine learning.



Qingrong Chen is currently pursuing the bachelor's degree with Shanghai Jiao Tong University. His research interests are social network security, Bitcoin security, and network privacy.



Haojin Zhu (IEEE M'09-SM'16) received his B.Sc. degree (2002) from Wuhan University (China), his M.Sc.(2005) degree from Shanghai Jiao Tong University (China), both in computer science and the Ph.D. in Electrical and Computer Engineering from the University of Waterloo (Canada), in 2009. Since 2017, he has been a full professor with Computer Science department in Shanghai Jiao Tong University. His current research interests include network security and privacy enhancing technologies. He published 35 international journal papers, including JSAC, TDSC, TPDS, TMC, TWC, TVT, and 60 international conference papers, including ACM CCS, ACM MOBICOM, ACM MOBIHOC, IEEE INFOCOM, IEEE ICDCS. He received a number of awards including: IEEE ComSoc Asia-Pacific Outstanding Young Researcher Award (2014), Top 100 Most Cited Chinese Papers Published in International Journals (2014), Supervisor of Shanghai Excellent Master Thesis Award (2014), Distinguished Member of the IEEE INFOCOM Technical Program Committee (2015), Outstanding Youth Post Expert Award for Shanghai Jiao Tong University (2014), SMC Young Research Award of Shanghai Jiao Tong University (2011). He was a co-recipient of best paper awards of IEEE ICC (2007) and Chinacom (2008) as well as IEEE GLOBECOM Best Paper Nomination (2014). He received Young Scholar Award of Changjiang Scholar Program by Ministry of Education of P.R. China in 2016.



Di Ma is an Associate Professor in the Computer and Information Science Department at the University of Michigan-Dearborn, where she leads the Security and Forensics Research Lab (SAFE). She is broadly interested in the general area of security, privacy, and applied cryptography. Her research spans a wide range of topics, including smartphone and mobile device security, RFID and sensor security, vehicular network and vehicle security, computation over authenticated/encrypted data, fine-grained access control, secure storage systems, and so on. Her research is supported by NSF, NHTSA, AFOSR, Intel, Ford, and Research in Motion. She received the PhD degree from the University of California, Irvine, in 2009. She was with IBM Almaden Research Center in 2008 and the Institute for Infocomm Research, Singapore in 2000-2005. She won the Tan Kah Kee Young Inventor Award in 2004.



Hong Wen received the M.Sc. degree from Sichuan Union University of Sichuan, P. R. China, in 1997. She pursued her Ph.D. degree in Communication and Computer Engineering Dept. at the Southwest Jiaotong University (Chengdu, P. R. China). Then she worked as an associate professor in the National Key Laboratory of Science and Technology on Communications at UESTC, P. R. China. From January 2008 to August 2009, she was a visiting scholar and postdoctoral fellow in the ECE Dept. at University of Waterloo. Now she holds the professor position at UESTC, P. R. China. Her major interests focus on wireless communication system security.



Xuemin (Sherman) Shen (IEEE M'97-SM'02-F09) received the B.Sc.(1982) degree from Dalian Maritime University (China) and the M.Sc. (1987) and Ph.D. degrees (1990) from Rutgers University, New Jersey (USA), all in electrical engineering. He is a Professor and University Research Chair, Department of Electrical and Computer Engineering, University of Waterloo, Canada. He is also the Associate Chair for Graduate Studies. Dr. Shen's research focuses on resource management in interconnected wireless/wired networks, wireless network security, social networks, smart grid, and vehicular ad hoc and sensor networks. He is an elected member of IEEE ComSoc Board of Governor, and the Chair of Distinguished Lecturers Selection Committee. Dr. Shen served as the Technical Program Committee Chair/Co-Chair for IEEE Globecom'16, Infocom'14, IEEE VTC'10 Fall, and Globecom'07, the Symposia Chair for IEEE ICC'10, the Tutorial Chair for IEEE VTC'11 Spring and IEEE ICC'08, the General Co-Chair for ACM Mobihoc'15, Chinacom'07 and QShine'06, the Chair for IEEE Communications Society Technical Committee on Wireless Communications, and P2P Communications and Networking. He also serves/served as the Editor-in-Chief for IEEE Network, Peer-to-Peer Networking and Application, and IET Communications; a Founding Area Editor for IEEE Transactions on Wireless Communications; an Associate Editor for IEEE Transactions on Vehicular Technology, Computer Networks, and ACM/Wireless Networks, etc.; and the Guest Editor for IEEE JSAC, IEEE Wireless Communications, IEEE Communications Magazine, and ACM Mobile Networks and Applications, etc. Dr. Shen received the Excellent Graduate Supervision Award in 2006, and the Outstanding Performance Award in 2004, 2007, 2010, and 2014 from the University of Waterloo, the Premier's Research Excellence Award (PREA) in 2003 from the Province of Ontario, Canada, and the Distinguished Performance Award in 2002 and 2007 from the Faculty of Engineering, University of Waterloo. Dr. Shen is a registered Professional Engineer of Ontario, Canada, an IEEE Fellow, an Engineering Institute of Canada Fellow, a Canadian Academy of Engineering Fellow, a Royal Society of Canada Fellow, and a Distinguished Lecturer of IEEE Vehicular Technology Society and Communications Society.