

Leveraging Prior Knowledge Asymmetries in the Design of Location Privacy-Preserving Mechanisms

Nazanin Takbiri	Virat Shejwalkar	Amir Houmansadr	Dennis L. Goeckel	Hossein Pishro-Nik
Electrical and	Information and	Information and	Electrical and	Electrical and
Computer Engineering	Computer Sciences	Computer Sciences	Computer Engineering	Computer Engineering
UMass-Amherst	UMass-Amherst	UMass-Amherst	UMass-Amherst	UMass-Amherst
ntakbiri@umass.edu	vshejwalkar@cs.umass.edu	amir@cs.umass.edu	goeckel@ecs.umass.edu	pishro@ecs.umass.edu

Abstract—The prevalence of mobile devices and Location-Based Services (LBS) necessitates the study of Location Privacy-Preserving Mechanisms (LPPM). However, LPPMs reduce the utility of LBSes due to the noise they add to users' locations. Here, we consider the remapping technique, which presumes the adversary has a perfect statistical model for the user location. We consider this assumption and show that under practical assumptions on the adversary's knowledge, the remapping technique leaks privacy not only about the true location data, but also about the statistical model. Finally, we introduce a novel method termed *Randomized Remapping* to provide a trade-off between leakage of the users' location and leakage of the users' model for a given utility.

Index Terms—Location-Based Service (LBS), information leakage, obfuscation, remapping technique, Location Privacy Preserving Mechanisms (LPPMs), utility-privacy trade-off.

I. INTRODUCTION

Contemporary mobile devices offer a wide spectrum of location-based services (LBS), such as ride sharing and navigation. LBSes collect large amounts of users' location data to tailor the service provided to each user's specific needs. To address the significant threat to user privacy due to location data sharing, location privacy-preserving mechanisms (LPPMs) have been introduced. LPPMs preserve privacy by sharing obfuscated versions of true location data, but at the cost of utility degradation [2], [3]. In [4], an adversary is assumed who has perfect knowledge of the distribution of a user's location data, and a *remapping technique* is proposed to improve the utility-privacy trade-offs of LPPMs. The remapping technique shares a location that is the adversary's best estimate of the true location, and consequently improves utility as the estimate is closer to the true location than the obfuscated location.

We model remapping as a general utility improvement technique for releasing not just location data but any type of data, e.g., IoT application data. We consider Gaussian distributed private data whose privacy is protected by the addition of Gaussian noise. Therefore, our analysis of remapping generalizes to various domains, e.g., sensor networks [5] and

distributed consensus [6], which consider Gaussian distributed data as a promising substitute for the real data, and it can be adapted to users' check-ins modeled as a multi-center Gaussian model [7].

We consider a friend (e.g., an IoT application) without any prior statistical information about user behavior and an adversary with statistical information about the user's behavior. This may occur, for example, when each intended recipient is either naive or only looking at a single datum or a small set of data from the user, whereas the adversary is sophisticated and has access to the user's data across a large time period. In such a case, the adversary can use their statistical advantage to obtain a better estimate of the user's data than the friend. Remapping recognizes this asymmetry of knowledge and reveals a more accurate version of the data that the adversary would have been able to obtain anyway. Thus, the remapping technique does not incur privacy loss, but improves utility for the user. Not surprisingly, this approach has garnered a growing amount of interest in the privacy community [8]–[12], hence motivating a fundamental analysis.

Note that the remapping technique implicitly assumes that the sophisticated adversary has perfect knowledge of the statistical model of the user data. However, this is not the case in practice, as the adversary's knowledge is imperfect due to multiple reasons, e.g., the adversary not having an infinite history, the user not reporting some of her data, or the reported data being noisy. We explore the remapping technique from the lens of this practical setting where the adversary has an *imperfect knowledge* of the statistical model of the user's data.

Contributions: After introducing our framework in Section II, we provide the first information-theoretic look to explain the operation of the remapping technique in Section III. As acknowledged briefly in [4], a risk of remapping is that it relies on accurate knowledge of the adversary's statistical model. In Section IV, we show that privacy leaks in two ways when erroneous assumptions are made about the adversary's model: (i) the adversary obtains a more accurate version of the data than they would have had without remapping, and (ii) the adversary is able to improve their knowledge of the statistics of the users' data beyond what they would have been able to do without remapping. In Section V, we present a random remapping algorithm, where data points

This work was supported by the National Science Foundation under grants CCF-1421957 and CNS-1739462.

Due to space limitation, detailed derivations, additional discussion, and more references are provided in the long version of the paper [1].

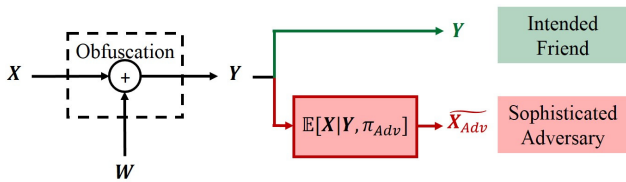


Fig. 1: System Model: Case where additive obfuscation (without remapping) is applied to the user's location. The (naive) intended friend does not have a prior distribution for X and hence employs Y for the user's locations. A sophisticated adversary, who possesses a prior distribution for X , can use this prior to obtain a better estimate of the user's location.

are independently remapped at random. The proposed random remapping provides a hyperparameter which allows the user to tune the trade-off between different types of privacy leakage at a given user utility.

.eps

II. SYSTEM MODEL AND METRICS

Consider a system where a user generates a location X which should be protected from a potential adversary. To preserve the privacy of the user's true location, the obfuscated location is obtained by adding noise W to X . In other words, the reported noisy version of location (Y) is obtained as $Y = X + W$. As shown in Figure 1, there exists an "intended" friend (e.g., an LBS) who does not have prior statistical knowledge about the user behavior, and a "sophisticated" adversary who has knowledge about the prior behavior of the user (π_{Adv}). The adversary observes the noisy reported location Y and uses it to find the estimate $\tilde{X}_{Adv}$, which denotes the estimate of the adversary given their observed location (Y) and their knowledge of the prior about the user (π_{Adv}) as $\tilde{X}_{Adv} = \mathbb{E}[X|Y, \pi_{Adv}]$. As a result, there exist asymmetries in knowledge and/or sophistication between the intended friend and the adversary. The remapping technique, which is introduced by Chatzikokolakis et al. [4], exploits these asymmetries to publish a more accurate version of the location that the sophisticated adversary would have been able to obtain anyway. As shown in Figure 2, each reported location is remapped into the best possible location according to the perfect prior information of the adversary.

Location Data Model: We adopt a Gaussian model. User traces are assumed to be independent and identically distributed (i.i.d.) Gaussian series, and each data location is drawn from a normal distribution with mean μ and variance σ_s^2 , in other words, we assume $X(k) = \mu + S$, where $S \sim \mathcal{N}(0, \sigma_s^2)$, thus $X(k) \sim \mathcal{N}(\mu, \sigma_s^2)$. We also assume there exists some underlying prior for the distribution of the mean (μ); we also take this to be Gaussian, and hence assume $\mu \sim \mathcal{N}(0, \sigma_\mu^2)$.

Obfuscation Mechanism: The obfuscated location is obtained by passing the data location through an additive white Gaussian noise (AWGN) channel. Hence, Y , the reported location of the user, is the sum of the true location, X , and the noise,

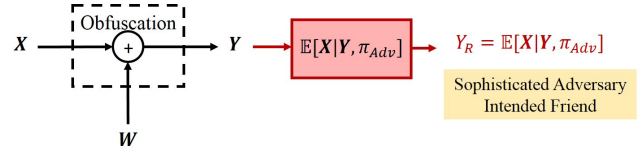


Fig. 2: Remapping: X is the user's true location, W is the amount of noise added through the obfuscation process, Y is the noisy reported location after applying obfuscation, and Y_R is the remapped location which is the best possible estimate of the adversary according to perfect prior knowledge about the user.

W , where W is drawn from a zero-mean normal distribution with variance equal to σ_w^2 . Thus, we have

$$Y = X + W \sim \mathcal{N}(\mu, \sigma_s^2 + \sigma_w^2).$$

Sophisticated Adversary Model: The adversary logs the user's locations over time to generate a prior about the behavior of the user and performs an inference attack to estimate the best possible location given this generated prior. Note that the remapping literature [4] has considered a *perfect prior* for the adversary. In reality, the adversary, however strong, does not have an infinite time history of user's data or have exact knowledge of the user's whereabouts, so she cannot build the perfect prior. In this paper, different adversarial settings have been considered: in Section III, we assume an adversary with a perfect prior, and in Section IV, we assume an adversary with an imperfect prior. It is critical to note that the adversary knows the mechanism of the obfuscation, but she does not know the exact value of the noise which will be added during obfuscation and does not have any auxiliary information or side information about the user's location.

Remapping Mechanism: In the absence of remapping, and given the *perfect prior* for the adversary, the adversary can estimate X using the reported noisy version of the location (Y) as:

$$Y_R = \mathbb{E}[X|Y, \mu] = \frac{\sigma_w^2}{\sigma_s^2 + \sigma_w^2} \mu + \frac{\sigma_s^2}{\sigma_s^2 + \sigma_w^2} Y, \quad (1)$$

where Y_R is the estimate of the adversary given the observed location (Y) and perfect knowledge of the prior (π_{Adv}). Remapping simply notes that, since the adversary obtains Y_R anyway (as shown in Figure 2), we might as well provide it to the applications to improve the utility [4].

Metrics: We use mean squared error (MSE) as a metric to quantify both utility degradation and privacy. In this paper, " $\mathcal{U}$ " denotes the MSE of the intended application/friend which quantifies utility degradation. In addition, " $\mathcal{P}$ " denotes the MSE of the adversary about the true location, and " $\hat{\mathcal{P}}$ " denotes the MSE of the adversary about the statistical model. Note that both " $\mathcal{P}$ " and " $\hat{\mathcal{P}}$ " quantify the level of privacy. We use MSE as metric in this paper due to the intuitive results and insights it provides about the shortcomings of the state-of-the-art remapping technique. However, we can also employ

the mutual information as a metric to quantify both utility and privacy. Note that the level of privacy can be quantified as $H(X|Y) = H(X) - I(X;Y)$, where $H(\cdot)$ is the entropy and $I(\cdot)$ is the mutual information. Thus, our goal is to minimize the information leakage to maximize the level of privacy at the highest possible utility. We have obtained similar results using mutual information as metric; these results can be found in the long version of the paper [1].

Discussion 1: The reason behind employing Gaussian models is related to the motivation for the paper. Recently, remapping was proposed to improve the utility of LPPMs without compromising privacy [4], which is an emerging technique of high importance in the field of privacy. Although, the leakage of private information has been noted, there has been no quantification of this important effect. Also the leakage of the statistical model (distribution), which is another major contribution of this paper, has been completely overlooked in the literature. In this paper, we have used Gaussian models not only because it make sense in some application [5]–[7], but also because these Gaussian models allow for a straightforward analysis to arrive at the answer to this compelling question.

III. CASE 1: ADVERSARY WITH PERFECT KNOWLEDGE

In this section, we assume the adversary knows the exact statistical distribution of the user data, which for our model means the exact value of the mean (μ).

A. Without Remapping

Without remapping, the user's intended friend, who does not have any knowledge of the statistical model for the user's data, observes only the noisy location (Y). The utility degradation is:

$$\mathcal{U}_{NR}^{(I)} = \mathbb{E} \left[\left(\tilde{X}_{App} - X \right)^2 \right] = \mathbb{E} \left[(Y - X)^2 \right] = \sigma_w^2. \quad (2)$$

In comparison to the user's friend, the sophisticated adversary obtains $\tilde{X}_{Adv} = \mathbb{E}[X|Y, \mu]$. Thus, $\mathcal{P}_{NR}^{(I)}$ which quantifies the level of privacy is calculated as:

$$\mathcal{P}_{NR}^{(I)} = \mathbb{E} \left[\left(\tilde{X}_{Adv} - X \right)^2 \right] = \mathbb{E} \left[(Y_R - X)^2 \right] = \frac{\sigma_w^2 \sigma_s^2}{\sigma_s^2 + \sigma_w^2}. \quad (3)$$

B. With Remapping

In this case, both the adversary and the user's friend observe the same reported location, $\tilde{X}_{Adv} = \tilde{X}_{App} = Y_R = \mathbb{E}[X|Y, \mu]$. Now, the MSE of the adversary and the MSE of the application are equal:

$$\mathcal{U}_R^{(I)} = \mathcal{P}_R^{(I)} = \mathbb{E} \left[(Y_R - X)^2 \right] = \frac{\sigma_w^2 \sigma_s^2}{\sigma_s^2 + \sigma_w^2}. \quad (4)$$

Since the intended friend/application is oblivious to the prior statistical knowledge about the user behavior, the MSE of the adversary is always smaller than or equal to the MSE of the application ($\mathcal{P} \leq \mathcal{U}$). Thus, we can conclude that the remapping technique provides the best utility among techniques satisfying the same level of privacy under the assumption that the adversary has perfect knowledge of the statistical model for the user data.

IV. CASE 2: ADVERSARY WITH IMPERFECT KNOWLEDGE

Here, we assume the adversary has a noisy version of the prior information, as might be obtained from a learning set of limited length. Specifically, the adversary has $\check{\mu} = \mu + E$, where E has a zero-mean normal distribution with variance equal to σ_e^2 , as would be the case if $\check{\mu}$ were the minimum mean square estimate (MMSE) based on prior observations with additive Gaussian obfuscation. Note that the distribution of E is not known to the PPM designers since the friend/application algorithm does not have knowledge on how the adversary has obtained her information and how accurately (e.g. amount of history observed). We consider not only the leakage of the true location (X) but also the leakage of the distribution of the true location (μ), which is a serious issue as such leakage would improve future estimates of the adversary.

A. Without Remapping

If remapping is not employed, the user's intended friend observes the reported location (Y). Thus, the utility is quantified as:

$$\mathcal{U}_{NR}^{(II)} = \mathbb{E} \left[\left(\tilde{X}_{App} - X \right)^2 \right] = \mathbb{E} \left[(Y - X)^2 \right] = \sigma_w^2. \quad (5)$$

In contrast, the sophisticated adversary uses both $Y = \mu + S + W$ and $\check{\mu} = \mu + E$ to improve knowledge not only about the true location (X) but also about the distribution of the true location (μ). Now, $\hat{\mathcal{P}}_{NR}^{(II)}$ which quantifies the MSE of the adversary about the distribution of the true location (μ) is:

$$\hat{\mathcal{P}}_{NR}^{(II)} = \mathbb{E} \left[(\tilde{\mu}_{Adv} - \mu)^2 \right] = \frac{\sigma_e^2 \sigma_\mu^2 (\sigma_s^2 + \sigma_w^2)}{(\sigma_\mu^2 + \sigma_e^2) (\sigma_s^2 + \sigma_w^2) + \sigma_e^2 \sigma_\mu^2}. \quad (6)$$

and $\mathcal{P}_{NR}^{(II)}$ which quantifies the MSE of the adversary about the true location (X) is:

$$\begin{aligned} \mathcal{P}_{NR}^{(II)} &= \mathbb{E} \left[\left(\tilde{X}_{Adv} - X \right)^2 \right] \\ &= \frac{\sigma_s^2 \sigma_w^2}{\sigma_s^2 + \sigma_w^2} + \frac{\sigma_w^4 \sigma_e^2 \sigma_\mu^2}{(\sigma_s^2 + \sigma_w^2) ((\sigma_\mu^2 + \sigma_e^2) (\sigma_s^2 + \sigma_w^2) + \sigma_e^2 \sigma_\mu^2)}. \end{aligned}$$

B. With Remapping

The user's friend observes the remapped location, so the utility of the system is:

$$\mathcal{U}_R^{(II)} = \mathbb{E} \left[\left(\tilde{X}_{App} - X \right)^2 \right] = \mathbb{E} \left[(Y_R - X)^2 \right] = \frac{\sigma_s^2 \sigma_w^2}{\sigma_s^2 + \sigma_w^2}. \quad (7)$$

However, the adversary observes not only Y_R , but also $\check{\mu}$, and uses both of these observations to estimate $\tilde{\mu}_{Adv}$ and $\tilde{X}_{Adv}$. Now, the MSE of the adversary about the distribution of the true location (μ) is:

$$\hat{\mathcal{P}}_R^{(II)} = \mathbb{E} \left[(\tilde{\mu}_{Adv} - \mu)^2 \right] = \frac{\sigma_s^4 \sigma_e^2 \sigma_\mu^2}{\sigma_s^4 (\sigma_\mu^2 + \sigma_e^2) + \sigma_e^2 \sigma_\mu^2 (\sigma_s^2 + \sigma_w^2)}, \quad (8)$$

and the MSE of the adversary about the true location (X) is:

$$\mathcal{P}_R^{(II)} = \mathbb{E} \left[\left(\tilde{X}_{Adv} - X \right)^2 \right] = \frac{\sigma_s^2 \sigma_w^2}{\sigma_s^2 + \sigma_w^2}. \quad (9)$$

Numerical results demonstrating what can be learned from these expressions will be presented in Section V.

C. Discussion: Leakage of the Statistical Model

From (8), we can conclude that increasing the obfuscation noise, somewhat surprisingly, increases the leakage about the distribution of the true location (μ) when remapping is employed. Note that $Y_R = \mathbb{E}[X|Y, \check{\mu}]$ depends on two parameters: 1) $\check{\mu} = \mu + E$ and 2) $Y = X + W$; thus, if we increase the obfuscation noise by increasing σ_w^2 , Y_R relies less on Y and more on $\check{\mu}$. Now in the extreme case, where σ_w^2 goes to infinity, the observed location (Y) is useless and, as a result, $Y_R = \mathbb{E}[X|Y, \check{\mu}] = \mu$. Note that protecting the statistical model is important since leaking the statistical model gives the adversary the opportunity to get a better estimate of the user's future locations. Hence, remapping technique leaks complete information about the statistical model (μ) as σ_w^2 goes to infinity.

V. RANDOMIZED REMAPPING AND NUMERICAL RESULTS

As derived in Section IV, the remapping technique can leak information about the distribution of the true location (μ) if the adversary does not have the perfect prior about the user. Here, we introduce a new technique called randomized remapping to improve privacy. This technique gives us an opportunity to provide a trade-off between leakage of the users' location and leakage of the users' model for a given utility. In this randomized remapping technique, each location data is replaced with its remapped location (Y_R) with probability p_H . Here, Z denotes the output of our randomized remapping technique and can be expressed as:

$$Z = \begin{cases} Y_R, & \text{with probability } p_H, \\ Y, & \text{with probability of } 1 - p_H. \end{cases}$$

The user's friend observes Z ; thus, according to the law of total expectation, the MSE of the application is

$$\mathcal{U}_{Rand}^{(II)} = \mathbb{E}[(Z - X)^2] = p_H \mathcal{U}_R^{(II)} + (1 - p_H) \mathcal{U}_{NR}^{(II)}. \quad (10)$$

However, the adversary observes both Z and $\check{\mu} = \mu + E$ to estimate the true location (X) and distribution of the true location (μ). Equations (11) and (12) present the analysis for a genie-aided adversary, where a genie tells the adversary for each published data point whether remapping was applied to that data point or not. In this case, we can calculate $\hat{\mathcal{P}}_{Rand}^{(III)}$ which indicates the MSE of the adversary about the distribution of the true location (μ) as:

$$\hat{\mathcal{P}}_{Rand}^{(III)} = \mathbb{E}[(\tilde{\mu}_{Adv} - \mu)^2] = p_H \hat{\mathcal{P}}_R^{(II)} + (1 - p_H) \hat{\mathcal{P}}_{NR}^{(II)}. \quad (11)$$

Figure 3a shows the MSE of the adversary about the statistical model ($\hat{\mathcal{P}}_{Rand}^{(III)}$) versus the MSE of the intended application/friend ($\mathcal{U}_{Rand}^{(II)}$). We can also calculate $\mathcal{P}_{Rand}^{(III)}$ which

indicates the MSE of the adversary about the true location (X) as:

$$\mathcal{P}_{Rand}^{(III)} = \mathbb{E}[(\tilde{X}_{Adv} - X)^2] = p_H \mathcal{P}_R^{(II)} + (1 - p_H) \mathcal{P}_{NR}^{(II)}. \quad (12)$$

Figure 3b shows the MSE of the adversary about the true location ($\mathcal{P}_{Rand}^{(III)}$) versus the MSE of the intended application/friend ($\mathcal{U}_{Rand}^{(II)}$).

From Figures 3a and 3b, we can conclude that remapping technique leaks less information about the user's true location (as is shown in Figure 3b), but leaks significant information about the statistical model (as is shown in Figure 3a). As a result, we proposed randomized remapping to enable a trade-off between leakage of the users' location and leakage of the users' model for a given utility that cannot be obtained with standard remapping. Randomized remapping provides a hyperparameter, p_H , which allows us to tune this trade-off. Figure 3c shows that for a given utility, randomized remapping provides a range of possible location leakages and statistical model leakages. Figure 3d shows the trade-off between leakage of the users' location and leakage of the users' model for the case MSE of the application is equal to 0.6. Thus, PPM designers can choose p_H depending on the application and tune the trade-off between leakage of the users' location and leakage of the users' model.

Discussion 2: Note that the actual randomized remapping where the adversary lacks knowledge of whether $Z = Y_R$ or $Z = Y$ (i.e. whether remapping was performed on a given published data point or not), would perform much better against the true adversary rather than the genie-aided adversary considered here. In particular, the friend who is only interested in a single data point or does not perform complicated statistical modeling does not care about whether remapping was employed, and thus cares little about whether $Z = Y_R$ or $Z = Y$; he will use the published data value at face value regardless. However, for the adversary not knowing whether $Z = Y_R$ or $Z = Y$, they are presented with a complicated mixture of Gaussians model, which significant impairs their ability to update their statistical model. The derivation of the adversary's optimal detector and its performance characterization is challenging and therefore relegated to future work.

VI. CONCLUSIONS

Remapping, a state-of-the-art utility improvement technique for location privacy preserving mechanisms, has been shown to leak no privacy if the adversary has *perfect* knowledge of a user's statistical model. To understand the impact of remapping under real-world scenarios, we perform an information-theoretic analysis of remapping under practical assumptions on the adversary's knowledge. We show that if the adversary has *imperfect* knowledge of the statistical model, the standard remapping technique provides the best utility, but leaks significant information about the statistical model. To remedy such, we proposed a preliminary version of randomized remapping

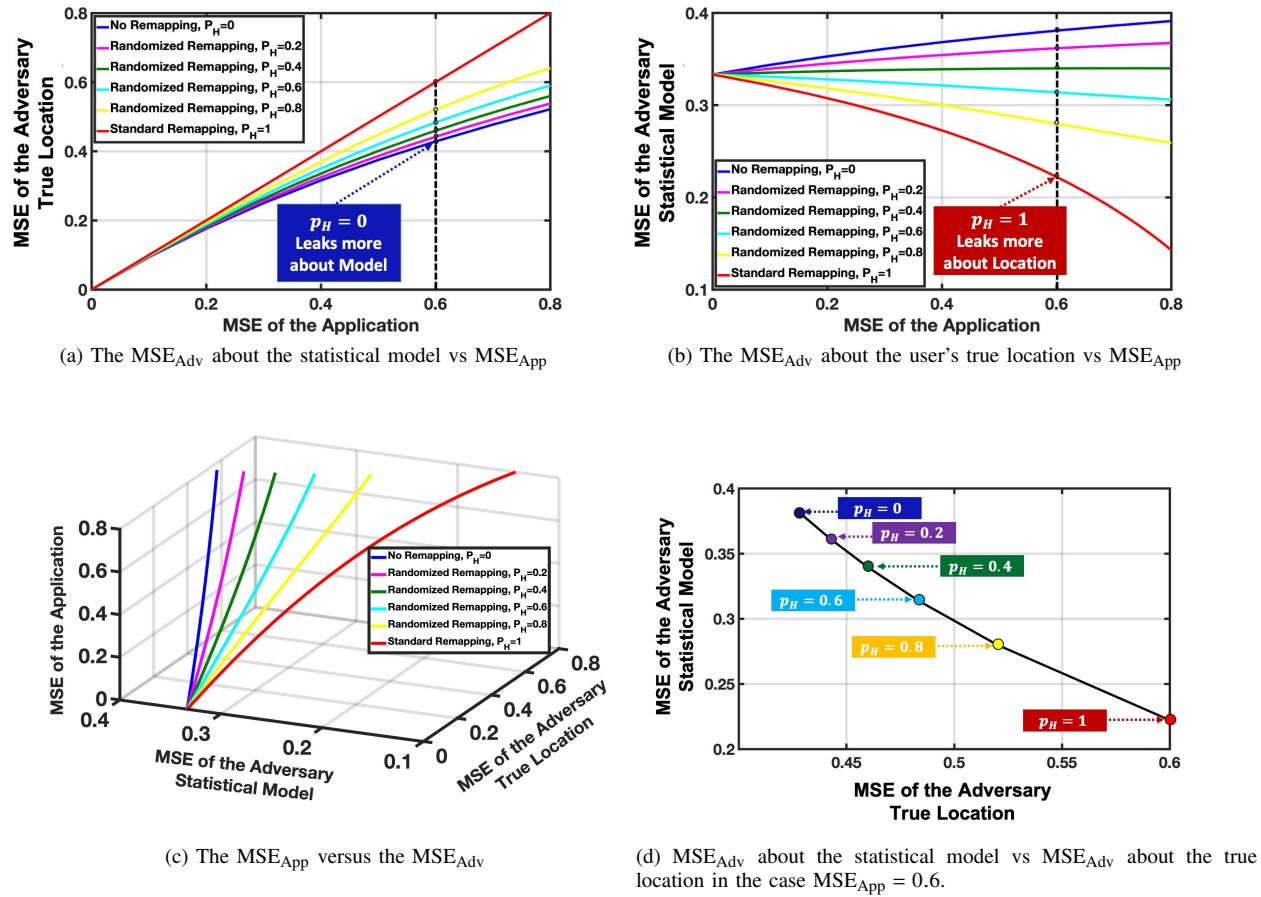


Fig. 3: The MSE of the adversary versus the MSE of the application for three cases. Case 1: remapping technique is not employed ($p_H = 0$), Case 2: a randomized remapping technique is employed with $p_H = 0.2, 0.4, 0.6$, and 0.8 , and Case 3: standard remapping [4] is employed ($p_H = 1$). Here, we assume $\sigma_\mu^2 = \sigma_e^2 = \sigma_s^2 = 1$ and σ_w^2 is swept from 0 to 3 with steps of 0.1.

which provides a hyperparameter which allows us to tune the trade-off between leakage of the users' location and leakage of the users' model for a given utility based on our need. We have taken the first necessary steps to open up new avenues for further research in this domain, and designing an optimal randomized remapping technique is an interesting topic for future research.

REFERENCES

- [1] N. Takbiri, A. Houmansadr, D. L. Goeckel, and H. Pishro-Nik, "Remapping," <http://www.ecs.umass.edu/ece/pishro/Papers/remapping.pdf>, January.
- [2] R. Shokri, G. Theodorakopoulos, C. Troncoso, J.-P. Hubaux, and J.-Y. Le Boudec, "Protecting location privacy: optimal strategy against localization attacks," in *Proceedings of the 2012 ACM conference on Computer and communications security*. ACM, 2012, pp. 617–627.
- [3] M. Duckham and L. Kulik, "A formal model of obfuscation and negotiation for location privacy," in *International conference on pervasive computing*. Springer, 2005, pp. 152–170.
- [4] K. Chatzikokolakis, E. ElSalamouny, and C. Palamidessi, "Efficient utility improvement for location privacy," *Proceedings on Privacy Enhancing Technology*, vol. 2017, no. 4, pp. 210–231, 2017.
- [5] A. Vempaty, O. Ozdemir, K. Agrawal, H. Chen, and P. K. Varshney, "Localization in wireless sensor networks: Byzantines and mitigation techniques," *IEEE Transactions on Signal Processing*, vol. 61, no. 6, pp. 1495–1508, 2012.
- [6] D. Wagner, "Resilient aggregation in sensor networks," in *SASN*, vol. 4. Citeseer, 2004, pp. 78–87.
- [7] C. Cheng, H. Yang, I. King, and M. R. Lyu, "Fused matrix factorization with geographical and social influence in location-based social networks," in *AAAI*, 2012.
- [8] A. Palia and R. Tandon, "Optimizing noise level for perturbing geo-location data," *Future of Information and Communication Conference*, pp. 63–73, 2018.
- [9] S. Oya, C. Troncoso, and F. Pérez-González, "Is geo-indistinguishability what you are looking for?" in *Proceedings of the 2017 Workshop on Privacy in the Electronic Society*. Dallas, Texas, USA: ACM, 2017, pp. 137–140.
- [10] R. Mendes and J. P. Vilela, "On the effect of update frequency on geo-indistinguishability of mobility traces," in *Proceedings of the 11th ACM Conference on Security and Privacy in Wireless and Mobile Networks*. Stockholm, Sweden: ACM, 2018, pp. 271–276.
- [11] Y. Kawamoto and T. Murakami, "Differentially private obfuscation mechanisms for hiding probability distributions," *CoRR*, vol. abs/1812.00939, 2018. [Online]. Available: <http://arxiv.org/abs/1812.00939>
- [12] S. Oya, C. Troncoso, and F. Pérez-González, "A tabula rasa approach to sporadic location privacy," *CoRR*, vol. abs/1809.04415, 2018. [Online]. Available: <http://arxiv.org/abs/1809.04415>