

Back to the Drawing Board: Revisiting the Design of Optimal Location Privacy-preserving Mechanisms

Simon Oya
University of Vigo
simonoya@gts.uvigo.es

Carmela Troncoso
IMDEA Software Institute
carmela.troncoso@imdea.org

Fernando Pérez-González
University of Vigo
fperez@gts.uvigo.es

ABSTRACT

In the last years we have witnessed the appearance of a variety of strategies to design optimal location privacy-preserving mechanisms, in terms of maximizing the adversary's expected error with respect to the users' whereabouts. In this work, we take a closer look at the defenses created by these strategies and show that, even though they are indeed optimal in terms of adversary's correctness, not all of them offer the same protection when looking at other dimensions of privacy. To avoid "bad" choices, we argue that the search for optimal mechanisms must be guided by complementary criteria. We provide two example auxiliary metrics that help in this regard: the conditional entropy, that captures an information-theoretic aspect of the problem; and the worst-case quality loss, that ensures that the output of the mechanism always provides a minimum utility to the users. We describe a new mechanism that maximizes the conditional entropy and is optimal in terms of average adversary error, and compare its performance with previously proposed optimal mechanisms using two real datasets. Our empirical results confirm that no mechanism fares well on every privacy criteria simultaneously, making apparent the need for considering multiple privacy dimensions to have a good understanding of the privacy protection a mechanism provides.

CCS CONCEPTS

• Security and privacy → Privacy-preserving protocols; • Networks → Location based services;

KEYWORDS

Location Privacy; Mechanism Design; Mechanism Evaluation; Quantifying Privacy

1 INTRODUCTION

Location based services raise important privacy concerns regarding the private information that exposing accurate location to service providers reveals [13, 14, 16, 20, 30]. To protect users' privacy, the academic community has proposed a wide variety of location privacy-preserving mechanisms [3, 15, 17–19, 21, 23, 28, 29] that mostly work altering the users' actual location before exposing it

to the service provider. The privacy evaluation of these proposals typically does not consider a strategic adversary, fostering an arms race in which defenses and attacks succeed each other without ever providing clear location privacy guarantees. To counter this effect, recent efforts focus on cutting the arms race short by either embedding the adversarial knowledge on the design process [5, 24, 27], or providing guarantees independent of the adversary's prior [2, 5, 24].

In this paper, we focus on sporadic user-centric protection mechanisms based on randomization, which preserve privacy by reporting a noisy version of the real location to the service provider according to a probability distribution. These mechanisms are adequate for applications that require infrequent location exposure, and can be run locally by the user. In this scenario, approaches that embed the adversarial knowledge on the design process are based on a Bayesian modeling of the adversary [26], and find optimal noise-generating mechanisms via linear optimization in which a target privacy objective is sought in presence of utility constraints [27]. On the other hand, approaches that provide privacy guarantees independent of the adversary's prior are based on *geo-indistinguishability* [2], an adaptation of differential privacy [9] to two-dimensional spaces, used by a number of works [11, 12, 22]. Geo-indistinguishability can be achieved optimally in terms of utility using expensive linear programming [5], or suboptimally using efficient remapping techniques that increase the utility of the query [6]. Finally, the Bayesian and the geo-indistinguishability approaches have been combined by Shokri [24] to obtain mechanisms that guarantee geo-indistinguishability while achieving a good performance against the Bayesian adversary.

Following the recommendation by Shokri et al. [26], which has been taken as the standard by the community, all of these approaches use the adversary's *correctness*, i.e., how close the adversary's estimate is to the correct answer, to evaluate location privacy. Usually, the adversary's correctness is measured as her expected estimation error, where this error is modeled using some distance metric between the real location and the adversary's estimation [25].

In this paper, we aim at understanding the properties of the mechanisms output by these design strategies. We find that, when the target privacy notion is the adversary's expected estimation error, there are many optimal mechanisms that meet a desired quality loss constraint. While this may seem advantageous, we show that following such an optimization objective may result in the selection of naive mechanisms that obviously provide little privacy, e.g., alternating the exposure of the actual user location and a far away location. Indeed, this mechanism complies on average with the constraints of the problem. Yet, it results on little uncertainty for the adversary, effectively providing a false perception of privacy.

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than ACM must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.

CCS '17, October 30–November 3, 2017, Dallas, TX, USA

© 2017 Association for Computing Machinery.

ACM ISBN 978-1-4503-4946-8/17/10...\$15.00

<https://doi.org/10.1145/3133956.3134004>

To counter such effect we argue that, depending on the user's preferences, the search for an optimal location privacy-preserving mechanism needs to consider more criteria than the error, contradicting the belief established by Shokri et al. [26]. As examples of complementary metrics to guide the design of protection mechanisms we propose the use of information-theoretic metrics, e.g., the conditional entropy, or a worst-case bound for quality loss. We provide efficient methods to construct mechanisms with respect to these criteria, and demonstrate that the remapping method introduced in [6] to improve the utility of geo-indistinguishability-based methods is in fact a straightforward generic scheme to build an optimal mechanism in terms of the expected estimation error from *any* obfuscation mechanism. We evaluate the effectiveness of the different mechanisms according to different privacy criteria using two real location datasets concluding that, generally, mechanisms that are optimal for one criterion do not necessarily perform well on others.

To summarize, we make the following contributions:

- ✓ We provide a theoretical characterization of optimal location privacy-preserving mechanisms in terms of the mean adversarial error. We show that, for a given average quality loss, there is more than one optimal protection mechanism that maximizes the average privacy. This family of mechanisms forms a convex polytope in which different mechanisms provide different privacy guarantees.
- ✓ We demonstrate the limitations of evaluating defenses solely considering the correctness of the adversary [26], and advocate for the use of complementary criteria to guide the design of location privacy-preserving mechanisms where the privacy guarantees provided are better understood.
- ✓ We provide algorithms to efficiently design mechanisms based on criteria other than the adversary's error. Furthermore, we demonstrate that remapping, previously proposed as an enhancement to geo-indistinguishability, is not only beneficial to improve the utility of this technique but can be used as a generic method to turn any obfuscation mechanism into optimal in terms of average adversarial error.
- ✓ We evaluate prior and new location privacy-preserving mechanisms on two real location datasets. Our results confirm that it is difficult to find optimal mechanisms that fare well on all criteria. This demonstrates that previous approaches to design location privacy-preserving mechanisms, while having solid foundations, oversimplify the design problem and generate defenses that overestimate the level of privacy offered to the user.

This paper is organized as follows. In Section 2, we introduce our system model, and the quality loss and privacy metrics we consider in the paper. In Section 3 we study the consequences of choosing the average adversary error as the standard metric to evaluate location privacy, illustrating that mechanisms that are optimal by this criterion may provide little privacy. In Section 4 we propose to consider auxiliary metrics to avoid bad mechanism choices in the optimization. As examples, we study the use of the conditional entropy and the worst-case quality loss. We evaluate several mechanisms built according to these new criteria in Section 5, and offer our conclusions in Section 6.

2 SYSTEM MODEL

We now describe our system model, which is in agreement with the framework for location privacy proposed by Shokri et al. [26], and introduce the notation used throughout the paper, which is summarized in Table 1.

We consider a set of users that send queries with a geographical position of interest to a location based service to obtain a service (e.g., finding points of interest or nearby friends). The location of interest can be the current location of the user or some other location the user is interested in querying about. Users wish to obtain utility from the location based service, while keeping their whereabouts private from an adversary that can observe the locations in the queries, e.g., an eavesdropper of the user-server communication, or the service provider itself. In order to protect their locations, users employ a *location privacy-preserving mechanism* that perturbs their location prior to exposing it to the server. We consider a strategic adversary that knows the protection mechanism operation, and has some knowledge about the users movement patterns. Given the observed perturbed location and her knowledge, the adversary tries to infer the user real location.

We model the locations queried by the users as a discrete set of *points of interest* denoted by $\mathcal{X} \doteq \{x_1, x_2, \dots, x_N\}$. We refer to these locations as *real* or *input locations* since they are the actual locations that are input to the location privacy-preserving mechanism. We use $\pi(x)$ to denote the *prior* probability that a user in the population queries the service provider about location x ($\pi(x) \geq 0$ and $\sum_{x \in \mathcal{X}} \pi(x) = 1$). This prior can either represent the global behavior of all the users as in [6], or be tailored to a particular user, but we assume that it is known both by the user and the adversary and that it can be used to design the privacy-preserving mechanism. We also consider independence between queries, i.e., that the input locations x from the same or other users are samples from i.i.d. random variables given by π .

The set of possible locations reported by the location privacy-preserving mechanism is denoted by $\mathcal{Z}$. We assume that users can report any location in the world $\mathcal{Z} = \mathbb{R}^2$. We refer to these locations as *output locations*, as they are the outputs of the privacy-preserving mechanism. The mechanism itself is denoted by f and modeled as a set of (continuous) conditional probability distributions, where $f(z|x)$ denotes the probability density function (pdf) of reporting the output location $z \in \mathbb{R}^2$ when the real location of the user is $x \in \mathcal{X}$ (note that $f(z|x) \geq 0$ and $\int_{\mathbb{R}^2} f(z|x) dz = 1$ for all $x \in \mathcal{X}$). We represent discrete mechanisms, i.e., mechanisms with a discrete output domain, in $\mathbb{R}^2$ with the Dirac delta function δ . For example, the mechanism that maps any $x \in \mathcal{X}$ to two particular outputs $z_1, z_2 \in \mathbb{R}^2$ with the same probability would be $f(z|x) = 0.5\delta(z - z_1) + 0.5\delta(z - z_2)$. For integration purposes, $\delta(z - z')$ must be understood as a two-dimensional Gaussian pdf centered at z' whose variance is arbitrarily small.

When using a privacy-preserving mechanism f to obtain privacy, the user experiences a loss on the quality of service due to the fact that she reports a location that might not be the location of interest, and may even be far away from this one. We use $P(f, \pi)$ to denote the *privacy* of the user, and $Q(f, \pi)$ to denote her *quality loss*. We specify particular instantiations of these functions below.

Table 1: Summary of notation

Symbol	Meaning
x	Input location the user is interested in querying about.
$\mathcal{X}$	Set of valid input locations or points of interest.
z	Output location released by the mechanism, $z \in \mathbb{R}^2$.
$\hat{x}$	Adversary's estimation of the input location, $\hat{x} \in \mathbb{R}^2$.
$\pi(x)$	Prior probability that a user wants to query about x .
$f(z x)$	Privacy mechanism. Pdf of $z \in \mathbb{R}^2$ given $x \in \mathcal{X}$.
$f_Z(z)$	Pdf of z , i.e., $f_Z(z) = \sum_{x \in \mathcal{X}} \pi(x) \cdot f(z x)$.
$p(x z)$	Posterior probability of x given z .
$d_Q(x, z)$	Quality loss distance function between x and z .
$\bar{Q}$	Average quality loss metric, in (1).
Q^+	Worst-case quality loss metric, in (2).
$d_P(x, \hat{x})$	Privacy distance function between x and $\hat{x}$.
P_{AE}	Average error privacy metric, in (5).
P_{CE}	Conditional entropy privacy metric, in (9)
P_{GI}	Geo-Indistinguishability privacy metric, in (11)

2.1 Quality Loss Metrics

We consider two possible definitions of quality loss: the average loss, and the worst-case loss. To this end we introduce $d_Q(x, z)$, a function that quantifies how much quality of service is lost by a user reporting output location z when she is interested in input location x . Larger values of $d_Q(x, z)$ indicate a larger loss, and therefore a worse utility performance for the user. The canonical choice for this function is the Euclidean distance: $d_Q(x, z) = \|x - z\|_2$. Note that $d_Q(\cdot)$ does not need to be a metric in the mathematical sense: it could be any function that maps an input location and a released location to a loss value (e.g., a feeling-based utility metric as in [1, 4]).

Average Loss. The average loss measures how much quality a user loses on average, and can be written as:

$$\bar{Q}(f, \pi) = \sum_{x \in \mathcal{X}} \int_{\mathbb{R}^2} \pi(x) \cdot f(z|x) \cdot d_Q(x, z) dz. \quad (1)$$

This metric has been the typical choice of utility in the related literature [2, 5–7, 27] since it is very intuitive. This metric also has the advantage of being linear with the mechanism f , which is very useful towards reducing the computational cost of mechanism design algorithms. Moreover, it makes the analysis of optimal algorithms in terms of average loss tractable.

Worst-case Loss. Given a function that quantifies the point-wise loss as defined above, $d_Q(x, z)$, the worst-case loss is defined as:

$$Q^+(f, \pi) = \max_{\substack{x, z \\ \pi(x) > 0 \\ f(z|x) > 0}} d_Q(x, z). \quad (2)$$

The worst-case loss measures how much utility the user loses in the worst case possible. For example, if $d_Q(x, z)$ is the Euclidean distance and the user wants to query about x , a mechanism with $Q^+(f, \pi) \leq 2\text{km}$ ensures that the output z will not be further than 2km away from x . This property is very helpful for many

applications that target nearby-type of services, since if the reported location is very far from the desired location then the result of the query would be generally useless for the user.

2.2 Privacy Metrics

We present now three notions of privacy: the average adversary error, the conditional entropy of the posterior distribution, and geo-indistinguishability.

Average Error. The average error is the de-facto standard to measure location privacy since Shokri et al. [26] argued that incorrectness determines the privacy of users. Consider that the adversary knows the prior π and the mechanism f chosen by the user. With this information, she produces an estimate $\hat{x} \in \hat{\mathcal{X}}$ of the user's input location x . The choice of $\hat{x}$ depends on the computational power of the adversary. Since we assume that the user has the freedom to report any location in $\mathbb{R}^2$, we also assume an unbounded adversary that can estimate locations on the whole world $\hat{\mathcal{X}} = \mathbb{R}^2$. Upon observing z , the adversary can build a *posterior* probability mass function over the inputs, denoted as $p(x|z)$:

$$p(x|z) = \frac{\pi(x) \cdot f(z|x)}{\sum_{x' \in \mathcal{X}} \pi(x') \cdot f(z|x')}. \quad (3)$$

Let $d_P(x, \hat{x})$ be a function that quantifies the magnitude of the adversary's error when deciding that the input location was $\hat{x}$ when the input location is actually x . As in the case of the average loss $\bar{Q}$, this function $d_P(\cdot)$ does not necessarily need to be a metric (e.g., it can include the user sensitivity to an adversary learning semantic information such as in [1]). Given an output location z , the optimal decision for the adversary in terms of minimizing the average error is

$$\hat{x}(z) = \underset{\hat{x} \in \mathbb{R}^2}{\operatorname{argmin}} \left\{ \sum_{x \in \mathcal{X}} p(x|z) \cdot d_P(x, \hat{x}) \right\}. \quad (4)$$

The average adversary's error, or just average error, is defined as the mean error incurred by an adversary that chooses the estimation $\hat{x}$ optimally given each observed z . Let $f_Z(z) = \sum_{x \in \mathcal{X}} \pi(x) \cdot f(z|x)$ be the probability density function of z . Then, the average error is:

$$P_{AE}(f, \pi) = \int_{\mathbb{R}^2} f_Z(z) \sum_{x \in \mathcal{X}} p(x|z) \cdot d_P(x, \hat{x}(z)) dz \quad (5)$$

$$= \int_{\mathbb{R}^2} \min_{\hat{x} \in \mathbb{R}^2} \left\{ \sum_{x \in \mathcal{X}} \pi(x) \cdot f(z|x) \cdot d_P(x, \hat{x}) \right\} dz. \quad (6)$$

Note that mechanisms designed with P_{AE} inherently protect against a strategic adversary, since the metric embeds the adversary's estimation. This metric has been used as part of the design objective in previous works [26, 27], and as a way of comparing the performance in terms of privacy of mechanisms designed with other different privacy goals in mind [2, 5–7].

Conditional Entropy. The conditional entropy is an information-theoretic metric that can be used to measure the adversary's uncertainty about the user's real location when z is released. After observing z , the adversary builds the posterior $p(x|z)$ using (3). The uncertainty of the adversary regarding the value of x given z can

be measured as the entropy of this posterior:

$$H(x|z) \doteq - \sum_{x \in \mathcal{X}} p(x|z) \cdot \log(p(x|z)). \quad (7)$$

The conditional entropy measures the *average* entropy of the posterior after z is released. Formally,

$$P_{CE}(f, \pi) = \int_{\mathbb{R}^2} f_Z(z) \cdot H(x|z) dz, \quad (8)$$

where $f_Z(z)$ is the probability density function of z , and $H(x|z)$ is a function of z as defined in (7). Alternatively, using only the prior π and the mechanism f , the conditional entropy can be written as

$$P_{CE}(f, \pi) = - \sum_{x \in \mathcal{X}} \int_{\mathbb{R}^2} \pi(x) \cdot f(z|x) \cdot \log \left(\frac{\pi(x) \cdot f(z|x)}{\sum_{x' \in \mathcal{X}} \pi(x') \cdot f(z|x')} \right) dz. \quad (9)$$

Note that this metric does not depend on the geography of the problem, i.e., on the particular values of x or z . If we use the base-two logarithm in the formula, then P_{CE} can be interpreted as how many bits of information the adversary needs on average to completely identify x . This metric was disregarded as a possible privacy metric in [26] due to being uncorrelated with the average error. In this work, we challenge such conclusion showing that considering solely the correctness of the adversary may lead to the design of mechanisms that offer low privacy. We show in Section 4 how using the conditional entropy as a complementary privacy metric helps to avoid choosing those undesirable mechanisms.

Geo-Indistinguishability. Geo-indistinguishability is an extension of the concept of differential privacy, originally a notion of privacy in databases, to the location privacy scenario. It was originally proposed in [2] and other works have continued the research on this line [5–7]. Formally, ϵ -geo-indistinguishability requires the following condition to be fulfilled by a location privacy-preserving mechanism f ,

$$\int_A f(z|x) dz \leq e^{\epsilon \cdot d_P(x, x')} \cdot \int_A f(z|x') dz, \quad \forall x, x' \in \mathcal{X}, \forall A \subseteq \mathbb{R}^2. \quad (10)$$

This requirement ensures that given an area $A \subseteq \mathbb{R}^2$, the probability of reporting a point z in that area if the original location was x over any other location x' within some distance around x , is similar, and therefore x and x' have some degree of statistical indistinguishability. In this definition, $d_P(x, x')$ is a function that quantifies how indistinguishable x and x' are: smaller values of $d_P(x, x')$ indicate a higher indistinguishability, as the constraint becomes tighter. The privacy parameter in this definition is ϵ : larger values of ϵ indicate a looser constraint that allows $f(z|x)$ and $f(z|x')$ to be more different, and therefore x and x' become more distinguishable. Smaller values of ϵ force the probability density functions $f(z|x)$ and $f(z|x')$ to be closer, providing more privacy. Note that, if for a single input location x there is a positive probability of reporting the output in a region $A \subseteq \mathbb{R}^2$, $\int_A f(z|x) dz > 0$, then that must also be true for every other input location x' . Also, note that geo-indistinguishability is independent of the prior π .

The typical choice of $d_P(x, x')$ in geo-indistinguishability is the Euclidean distance [2, 5]. Many geo-indistinguishability mechanisms rely on the fact that $d_P(x, x')$ is a metric (specifically, in the

fact that it satisfies the triangular inequality $d_P(x, x') \leq d_P(x, z) + d_P(x', z)$) to prove that they meet the condition in (10).

Although geo-indistinguishability is generally considered a privacy guarantee and not itself a metric, we can adapt it to represent an equivalent concept to our generic metric $P(f, \pi)$. Given a mechanism that provides ϵ -geo-indistinguishability, it is straightforward to see that it is also ϵ' -geo-indistinguishable if $\epsilon' > \epsilon$. Since a smaller ϵ denotes more privacy, it makes sense to define the geo-indistinguishability level provided by a mechanism f according to the smallest ϵ it guarantees. Also, since we are defining $P(f, \pi)$ as a magnitude that grows with the protection of the users, we choose to define our measure of geo-indistinguishability, $P_{GI}(f)$, as the inverse of the smallest ϵ guaranteed by the mechanism. Given the mechanism f , we write

$$P_{GI}(f) = \inf_{\substack{x, x' \in \mathcal{X} \\ z \in \mathbb{R}^2}} d_P(x, x') \cdot \left| \log \frac{f(z|x)}{f(z|x')} \right|^{-1}, \quad (11)$$

where we assume by convention that $\log(\frac{0}{0}) = 0$ and that $d_P(x, x') = \|x - x'\|_2$ is the Euclidean distance. Larger values of P_{GI} indicate more privacy, and the mechanism guarantees $1/P_{GI}$ -geo-indistinguishability.

3 LIMITATIONS OF THE EXPECTED ADVERSARY ERROR BASED EVALUATION

The most standard way to assess the location privacy provided by two mechanisms has been the evaluation of the trade-off between their average adversary error P_{AE} and their average loss $\bar{Q}$. The use of the average error as yardstick for location privacy was proposed in [26] under the general notion of correctness, and its use as a way of comparing mechanisms was followed by many of the subsequent works [1, 2, 5–7, 27]. The choice of distance functions $d_P(\cdot)$ and $d_Q(\cdot)$ for both the average error and the average loss in these works is mostly the Euclidean distance [1, 2, 5, 6, 27] although some of them also consider the Hamming distance [5, 26, 27] or semantic distances for privacy [1, 7].

In this section, we show the problems that stem from this established 2-dimensional evaluation approach. We start by studying the properties of mechanisms that are optimal according to these two metrics. Then, we introduce a new mechanism that we call the *coin mechanism*, and use it as an example that brings to light the flaws of judging the privacy of a mechanism by its performance in terms of average error and average loss.

3.1 Study of the Established Mechanism Evaluation

We start our analysis by assuming that the choice of distance functions $d_P(\cdot)$ and $d_Q(\cdot)$ is the same for simplicity, which is a typical choice in related works (e.g., both are the Euclidean distance). We denote this by $d_P(\cdot) \equiv d_Q(\cdot)$. At the end of the section, we argue what happens when this is not the case. We also introduce two definitions. First, let $\mathcal{F}_Q$ be the set of all the mechanisms that achieve an average loss smaller or equal than Q . Formally,

$$\mathcal{F}_Q \doteq \{f \mid \bar{Q}(f, \pi) \leq Q\}. \quad (12)$$

Also, let $\mathcal{F}_Q^{\text{opt}} \subseteq \mathcal{F}_Q$ be the set of all mechanisms $f \in \mathcal{F}_Q$ that are optimal in terms of average adversary error, i.e.,

$$\mathcal{F}_Q^{\text{opt}} \doteq \{f \mid f \in \mathcal{F}_Q, P_{\text{AE}}(f, \pi) \geq P_{\text{AE}}(f', \pi) \ \forall f' \in \mathcal{F}_Q\}. \quad (13)$$

We call a mechanism inside $\mathcal{F}_Q^{\text{opt}}$ optimal, since it achieves as much privacy as possible among all the mechanisms with the same quality loss. We state the following lemma:

LEMMA 3.1. *The set of optimal mechanisms with respect to the average privacy P_{AE} and the average loss $\bar{Q}$ is a convex polytope.*

PROOF. Let the privacy achieved by any mechanism in $\mathcal{F}_Q^{\text{opt}}$ be $P_{\text{opt}}(Q)$. Then, we can define this set as

$$\mathcal{F}_Q^{\text{opt}} = \{f \mid P_{\text{AE}}(f, \pi) = P_{\text{opt}}(Q), \ \bar{Q}(f, \pi) \leq Q\}, \quad (14)$$

and since $P_{\text{AE}}(f, \pi)$ and $\bar{Q}(f, \pi)$ are linear operations with f , (14) can be written as an intersection of half-spaces, which forms a convex polytope. $\square$

Note that the proof also applies to the case where $d_P(\cdot) \neq d_Q(\cdot)$ (e.g., privacy as the average Hamming error of the adversary and quality loss as the average Manhattan distance). The same outcome can be derived for the conditional entropy and geo-indistinguishability, although we leave those results out of the scope of this work.

This lemma shows that there is a *family* of optimal mechanisms that lie inside a convex polytope, instead of just a single mechanism. All of them provide the same (maximal) privacy for the same quality loss constraint so, in principle, they are equally useful. In what follows, we show why this is not the case.

We start by introducing the concept of remapping. A remapping g is a function $g: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ that maps an output $z \in \mathbb{R}^2$ to another output $z' \in \mathbb{R}^2$ according to the probability density function $g(z'|z)$. It is well known that if we generate a mechanism $f' = f \circ g = \int_{\mathbb{R}^2} g(z'|z) \cdot f(z|x) dz$, then the privacy of f' in terms of average error, conditional entropy or geo-indistinguishability is not smaller than that of f . This is reasonable, as the remapping g is independent from x , and thus it does not reveal any information about it. The optimal Bayesian remapping is defined as follows:

Definition 3.2 (Optimal remapping). Given a mechanism f , its optimal remapping is the one that minimizes the average loss of the composition $f' = f \circ g$, i.e., $g(z'|z) = \delta(z' - r(z))$, where

$$r(z) = \operatorname{argmin}_{z' \in \mathbb{R}^2} \sum_{x \in \mathcal{X}} \pi(x) \cdot f(z|x) \cdot d_Q(x, z'). \quad (15)$$

This remapping assigns each location z to the location $r(z)$ in (15), and is used in [6] as a way of improving the utility of geo-indistinguishability mechanisms. Now, we show that it can also be used not only to reduce the quality loss of mechanisms but to achieve optimal mechanisms in terms of average error privacy:

THEOREM 3.3. *Let g be an optimal remapping for mechanism f , and let f' be the composition $f' = f \circ g$. If $d_P(\cdot) \equiv d_Q(\cdot)$, then f' is an optimal mechanism, i.e., $f' \in \mathcal{F}_{\bar{Q}(f', \pi)}^{\text{opt}}$.*

The proof is provided in the Appendix.

This theorem provides a straightforward way of building an optimal mechanism f' from any mechanism f . The idea is to reassign each output z of f to another symbol z' such that the average quality loss is minimized. Doing this for every output ensures that the quality loss cannot be further reduced, and since the distance function used to evaluate quality loss and privacy is the same, the best estimation the adversary can do of x is just to keep the released value. Note that the $\bar{Q}(f', \pi) \leq \bar{Q}(f, \pi)$. This means that, in order to find an optimal mechanism f' for a target quality loss $\bar{Q}(f', \pi) = Q$ using the remapping strategy, one has to adjust the loss of the mechanism f (e.g., by tuning its variance if it is a noise mechanism) until f' achieves the desired average loss Q .

It is straightforward to see that, if the optimal remapping for a mechanism f is just doing nothing, then it means f is optimal:

COROLLARY 3.4. *If the optimal remapping in (15) for a mechanism f is $g(z'|z) = \delta(z' - z)$, then f is optimal for its quality loss Q , i.e., $f \in \mathcal{F}_Q^{\text{opt}}$.*

This is a very convenient way of proving the optimality of a mechanism when $d_P(\cdot) \equiv d_Q(\cdot)$. Another way of seeing that such mechanism is optimal, is by realizing that with this choice of metrics, the privacy is upper bounded by the quality loss $P_{\text{AE}}(f, \pi) \leq \bar{Q}(f, \pi)$, and the upper bound is indeed achieved when an optimal mechanism is used. We note that the fact that $P_{\text{AE}}(f, \pi) = \bar{Q}(f, \pi)$ for optimal mechanisms is not new, as it was already mentioned in [2] about the mechanisms in [27].

3.2 The Coin Mechanism and the Flaws of the Traditional Approach

We now discuss the following mechanism, which we call *the coin mechanism*, and prove that it is optimal. Let z^* be the output location that minimizes the average quality loss of a mechanism that always reports that location regardless of the input x . Formally,

$$z^* \doteq \operatorname{argmin}_{z \in \mathbb{R}^2} \sum_{x \in \mathcal{X}} \pi(x) \cdot d_Q(x, z). \quad (16)$$

As an example, if we measure the point-to-point loss as the mean squared error $d_Q(x, z) = \|x - z\|_2^2$, then z^* will be given by the mean $z^* = \sum_{x \in \mathcal{X}} \pi(x) \cdot x$. If the loss is measured as the Euclidean distance, then z^* is the geometric median of π . Given a generic distance function $d_Q(\cdot)$, the optimal output location z^* can be computed by solving the optimization problem in (16).

Let Q^* be the average quality loss achieved by a mechanism that always reports z^* regardless of the input. We construct the following mechanism, which we denote f_{coin} . First, we fix a desired quality loss $Q \leq Q^*$ and compute $\alpha \doteq 1 - Q/Q^*$. Then, we build

$$f_{\text{coin}}(z|x) = \alpha \cdot \delta(z - x) + (1 - \alpha) \cdot \delta(z - z^*), \quad (17)$$

where z^* is in (16). This mechanism can be easily explained and implemented simulating a coin flip. We first set our desired quality loss $Q \leq Q^*$. Note that it would not make sense to fix Q to a value larger than Q^* since we would not achieve more privacy by doing so; a mechanism that always reports z^* and has an average loss of Q^* yields the highest privacy allowed by π . Then, we compute $\alpha = 1 - Q/Q^*$ and set it as the probability that our coin shows heads. Assume we are interested in querying about a location $x \in \mathcal{X}$, so we flip the coin. If the coin shows heads, then we report our desired

location $z = x$. If the coin hits tails, then we report z^* regardless of the value of x . It is easy to see that the average loss of (17) is indeed Q , by the linearity of this metric with f .

PROPOSITION 3.5. *The coin mechanism obtained for quality loss Q achieves the maximum average adversarial error possible given a constraint on the average quality loss, i.e., $f_{\text{coin}}(Q) \in \mathcal{F}_Q^{\text{opt}}$, if both are measured with the same distance function $d_P(\cdot) \equiv d_Q(\cdot)$.*

The proof is straightforward using the result in Corollary 3.4.

We now reason why, even though the coin mechanism is optimal by the standards that have been used to evaluate privacy in prior works (i.e., P_{AE} and $\bar{Q}$), this mechanism is hardly desirable for any user. When the coin shows heads, the adversary observes z . If $z \neq z^*$, the adversary knows for sure that the user was interested in querying about $x = z$ and therefore the user has no privacy at all. In this case, for privacy issues, there was no point in using the mechanism. When the coin shows tails, the user is mapped far away to z^* . The adversary observes z^* and has no idea where the user is, besides the prior π that was already known by her. In this case, the privacy of the user is maximal, but the quality loss is very large, since z^* is almost always very far away from the user. The quality loss is so large that the utility the user gets from this realization of the mechanism can be considered zero, so we can say that there was no point in using the mechanism in this case either. We have reached the issue we mentioned earlier: there is a mechanism, optimal by classic location privacy standards [26], that is useless both from the privacy and the quality loss point of view. This shows that there is a fundamental problem with the classic way that has been used to evaluate location privacy mechanisms.

3.3 The reach of this problem

One could think that the problem of this bi-dimensional evaluation approach lies on the fact that one cannot use the same metric to measure quality loss and privacy, e.g., the Euclidean distance. However, even with different metrics, mechanisms similar to the coin can be derived. For example, if privacy is the average mean squared error and quality loss is measured as the average Manhattan distance (i.e., the l_1 norm), a deterministic mechanism that consists on reporting the real location on most of the places and mapping to the other side of the Earth in some others is optimal, due to the fact that the MSE grows quadratically with the distance, while the l_1 (or any l_p norm) does not. In our evaluation, we show an example where a mechanism optimized for P_{AE} and $\bar{Q}$ with a different pair of distance functions $d_P(\cdot) \neq d_Q(\cdot)$ suffers from the coin issue. The problem does not arise from the particular distance functions $d_P(\cdot)$ and $d_Q(\cdot)$ one uses to evaluate the average error and loss, but from the fact that these metrics are *averages*, and as such they do not restrict the minimum privacy of a single use of the mechanism or the maximum quality loss of the mechanism, they just ensure that the average is good. We believe that, while evaluating the average behavior of a mechanism is not an erroneous notion per-se, it must be handled with care to avoid undesirable results, such as the coin mechanism.

As a concluding remark, we would like to note that we have shown this problem assuming that the outputs of the mechanism and the values estimated by the adversary are points in $\mathbb{R}^2$, for

notational simplicity and generality. An important fraction of previous works [1, 5, 7, 26, 27] assume a discrete model where the set of output values $\mathcal{Z}$ and estimated values $\hat{\mathcal{X}}$ are the centers of a grid over the map or points of interest such as $\mathcal{X}$. In these scenarios, one can derive a similar mechanism, where hitting tails means that the user reports the location out of the allowed ones that minimizes the average error. That mechanism can also be shown to be optimal in terms of average error and loss, although it is not a desirable mechanism for any user. For completeness, we also evaluate this scenario in our experiments. The same applies to the case where instead of having discrete input locations $\mathcal{X}$, users can report any point in $\mathbb{R}^2$ (for example, a tracking or a date finder application). The coin mechanism in (17) can be applied directly to this scenario, and it can be shown to be optimal (changing the summations over $\mathcal{X}$ to integrals). It is clear that using the traditional evaluation approach has flaws in all these scenarios and we must find a solution to this.

4 COMPLEMENTARY MECHANISM EVALUATION CRITERIA

So far we have seen that evaluating mechanisms based solely on the average error and quality loss does not reflect whether a mechanism is actually more beneficial than another one, due to the fact that some undesirable mechanisms are deemed optimal by this approach. In this section, we propose a solution to this evaluation procedure that consists in incorporating complementary evaluation criteria that add different perspectives to the performance of a mechanism in terms of privacy and quality loss.

We propose two metrics, that are not intended to be used as a replacement of the average error and average loss but in combination with them, adding new dimensions to the privacy vs. quality loss trade-off. The first metric we propose is the conditional entropy, a privacy metric that helps detecting inconsistent mechanisms such as the coin. The second one is the worst-case loss, a quality loss metric that provides a way of staying out of mechanisms that might yield no utility for the user at all. We comment on the implementation of mechanisms that take these metrics into consideration, and propose a mechanism that maximizes the conditional entropy while being optimal in terms of average error and quality loss. We finish the section describing other alternative privacy metrics.

4.1 The Conditional Entropy as a Complementary Metric

4.1.1 Usefulness of the Conditional Entropy. One of the problems of the coin mechanism can be seen from an information-theoretic point of view. The coin is a binary mechanism, in the sense that each input location can only be mapped to itself or to a fixed point in the map. From the adversary's perspective, this means that if the coin shows heads the adversary has no uncertainty at all about the user's input location, and if it shows tails the uncertainty is maximal. The conditional entropy can be used to detect these scenarios where the adversary has no uncertainty about x . Recalling (8), the conditional entropy can be written as

$$P_{\text{CE}}(f, \pi) = \int_{\mathbb{R}^2} f_Z(z) \cdot H(x|z) dz, \quad (18)$$

where $H(x|z) \doteq -\sum_{x \in \mathcal{X}} p(x|z) \cdot \log(p(x|z))$ is the entropy of the posterior after a location z is released. It is clear that (18) is an average over the entropy of all the posteriors. However, contrary to the average error, the conditional entropy is an average over functions $H(z|x)$ that are strictly concave with f . This means that in order to perform well in terms of the conditional entropy, a mechanism must spread its uncertainty among every posterior $p(x|z)$ instead of achieving maximal uncertainty with some outputs and zero uncertainty with others, as the coin does.

Another interesting property of the entropy is that it is not a geographical metric. The entropy of a posterior $H(x|z)$ does not depend on the coordinates of the input locations or the semantic information tied to them (e.g., if the location is a hospital or a club). The entropy only depends on how evenly the posterior is distributed among the input locations. This probabilistic aspect of privacy, defined as *uncertainty* in [26], cannot be captured by other privacy notions such as correctness (e.g., the average adversary error). Due to the geographic nature of the location privacy problem, we cannot judge a mechanism based solely on its entropy. However, using it as an additional dimension of privacy gives a more complete picture of the performance of a mechanism.

We would like to point out that this notion of uncertainty provided by the entropy was disregarded as a reasonable privacy metric in [26] based on the fact that, since it is not correlated with the adversary error, it does not capture how hard is for the adversary to estimate the real input location. We claim that it is indeed the fact that the entropy is not correlated with the adversary error which gives it a special value as a *complementary* metric of privacy. The same way that semantic location privacy metrics have been proposed together with geographic metrics [1, 7] to give different perspectives on the problem, the conditional entropy is a tool that gives valuable information about the protection provided by the mechanism not captured by the average error.

We would like to make two remarks regarding the entropy. First, the conditional entropy $P_{CE}(f, \pi)$ must be taken into account together with the mutual information $I(X; Z)$ to get a full picture of the information-theoretic properties of the mechanism. The conditional entropy represents the average amount of uncertainty the adversary has about the real location x after observing z . A small value of conditional entropy indicates low uncertainty, and therefore we might get the impression that a mechanism with such small value provides low privacy. However, it might have been possible that the entropy of the prior was already low, and therefore even if the mechanism was perfect from the privacy point of view (i.e. it did not reveal any information, $I(X; Z) = 0$), there is nothing any mechanism could have done to avoid having a low conditional entropy. We must therefore take into account the mutual information or, equivalently, the entropy of the prior π , when interpreting the value given by the conditional entropy.

The second remark is that the conditional entropy must not be tailored to a particular adversary with a possibly wrong knowledge of the prior π . In this work, we have assumed that the prior π models the choice of input locations by the users, and therefore the correct way of computing the entropy is by using π in the formulas above. This entropy must be regarded as the uncertainty that a very strong passive adversary with full knowledge of the behavior of the users would have when observing z .

4.1.2 Implementation of Mechanisms with large Conditional Entropy. We now look for a mechanism that is optimal in terms of the average error and average loss, i.e., a mechanism in $\mathcal{F}_Q^{\text{opt}}$, that also achieves as much conditional entropy as possible. This problem is equivalent to the rate-distortion problem [8] of finding a pdf $f(z|x)$ that minimizes the mutual information between x and z subject to a quality loss constraint, which can be solved iteratively by implementing the Blahut-Arimoto algorithm. For this, we must first restrict our output to a discrete alphabet $\mathcal{Z}$ for computational reasons. The more points we assign to this alphabet and the more evenly we cover the space where we want to compute the mechanism with them, the better its performance will be. Since both the input and output domains are discrete, the mechanism is determined by the probabilities of reporting z when the user is in x , that we denote by $p(z|x)$ here for clarity. We start with an initial mechanism, for example uniform mapping $p(z|x) = 1/|\mathcal{Z}|$. Then, we perform the following steps:

- (1) We compute the probability mass function of each the output:

$$P_Z(z) = \sum_{x \in \mathcal{X}} \pi(x) \cdot p(z|x), \quad \forall z \in \mathcal{Z}. \quad (19)$$

- (2) We update the mechanism as follows:

$$p(z|x) = P_Z(z) \cdot e^{-b \cdot d_Q(x,z)}, \quad \forall x \in \mathcal{X}, z \in \mathcal{Z}. \quad (20)$$

- (3) We normalize the mechanism:

$$p(z|x) = \frac{p(z|x)}{\sum_{z' \in \mathcal{Z}} p(z'|x)}, \quad \forall x \in \mathcal{X}, z \in \mathcal{Z}. \quad (21)$$

We skip this step for the outputs z with $P_Z(z) = 0$.

- (4) We repeat these steps until the change in the probabilities $p(z|x)$ is below some threshold.

The value of b in the second step needs to be tuned to change the quality loss of the mechanism $\bar{Q}(f, \pi)$ and cannot be pre-computed to achieve an exact value of average loss. Larger values of b yield mechanisms with less quality loss, and therefore less average error privacy and less conditional entropy. Finally, we obtain our mechanism $f(z|x)$ by applying the optimal remapping to the discrete mechanism defined in $\mathcal{X} \rightarrow \mathcal{Z}$ by the probabilities $p(z|x)$. This ensures that the resulting mechanism is optimal from the adversary error privacy point of view.

We make two remarks regarding this algorithm. The first one is about its computational cost. The operations in the three steps above are not expensive as they only include multiplications and additions. The number of elements we need to compute in order to build $p(z|x)$ is $N \doteq |\mathcal{X}| \cdot |\mathcal{Z}|$. The first step above consists of N products and additions. In the second step $e^{-b \cdot d_Q(x,z)}$ can be precomputed as b , $\mathcal{X}$ and $\mathcal{Z}$ do not change during the algorithm, so we only have to make N multiplications, and in the third step we compute $|\mathcal{X}|$ values of $\sum_{z' \in \mathcal{Z}} p(z'|x)$ and then perform N divisions. It is clear then that the cost grows with the sizes of $\mathcal{X}$ and $\mathcal{Z}$. However, the algorithm only needs to be computed once for all the users, which can be done in the cloud, and even if the prior π varies we can use a previously computed algorithm as initialization of the iteration above to get a fast update of the mechanism.

The second remark is that the mechanism produced by this algorithm also satisfies $2b$ -geo-indistinguishability (the proof is in the Appendix). This is a byproduct property that was not part

of the reasoning behind the algorithm and it does not imply that the conditional entropy and geo-indistinguishability are related. In fact, these are *fundamentally different* notions: the former is an average metric that only considers the probabilistic (and not the geographic) aspect of the problem, while the latter is a worst-case metric that also considers the geography of the problem. Also, if we truncate the optimal conditional entropy mechanism, we obtain a mechanism that is almost optimal in terms of conditional entropy but does not provide *any* level of geo-indistinguishability.

We evaluate this mechanism and others with respect to the conditional entropy and the traditional metrics in Section 5.

4.2 The Worst-Case Quality Loss as a Complementary Metric

4.2.1 Usefulness of the Worst-Case Quality Loss. After analyzing the privacy problems of the coin mechanism, we now turn to the utility point of view. The great drawback of the coin mechanism from the quality loss perspective is that if the coin shows tails then the server's response to the user's query will most likely be useless due to the great quality loss incurred by reporting z^* . We can think of many applications where, if the Euclidean distance between x and z is larger than a certain value, the user gets literally nothing from the server response. For example, if we are close to a point of interest x and we want to find a nearby hospital, querying about a location z in another city will likely return a useless response from the server. In that case, we could think of generating another output and query the server again because we did not get what we were hoping for. By doing so, the privacy properties of the mechanism change, and in the case of the coin it is equivalent to always revealing our true location.

A solution to this utility issue consists in imposing a worst-case quality loss constraint on the mechanism, i.e.,

$$Q^+(f, \pi) = \max_{\substack{x, z \\ \pi(x) > 0 \\ f(z|x) > 0}} d_Q(x, z) \leq Q_{\max}^+. \quad (22)$$

To put it simply, we want a mechanism that releases output locations within $Q_{\max}^+$ from the input location, i.e., a *bounded mechanism*. The upper bound $Q_{\max}^+$ would be tuned depending on the application in question, so that a user never gets a worthless result. When used together with the average error and the average loss, the worst-case loss metric reveals those mechanisms we might want to avoid using. It is easy to see that the coin mechanism, although optimal in terms of P_{AE} and $\bar{Q}$, gives a very large value of $Q^+(f_{\text{coin}}, \pi)$, which manifests its uselessness.

An interesting consequence of setting a maximum worst-case quality loss constraint when designing a mechanism is that it can simplify the computational cost of the protocol that implements or computes it. For example, take the case of the works in [5, 27], where authors assume a discrete set of output locations $\mathcal{Z}$ and propose to solve a linear program to find an optimal mechanism (in terms of average error and geo-indistinguishability, respectively). The constraint in (22) reduces the amount of variables that need to be computed in these programs (only a subset of $\mathcal{Z}$ are possible outputs for each input $x \in \mathcal{X}$), as well as the amount of constraints, which in turn decreases drastically the computational cost of the problem. In other implementations of mechanisms, where f is not

explicitly derived but computed by adding (continuous) noise and then computing a remapping using the posterior (c.f. [6]), having a worst-case quality loss constraint reduces the amount of inputs that need to be considered when computing the posterior, effectively reducing the computational cost of the algorithm.

Finally, we would like to note that this metric exposes a basic problem with geo-indistinguishability mechanisms. As mentioned before, when using a geo-indistinguishability mechanism, if a user with input location x has non-zero probability of reporting $z \in A \subseteq \mathbb{R}^2$, then when the input location is any other $x' \in \mathcal{X}$ she must assign a non-zero probability to reporting $z \in A$. This means that for any geo-indistinguishable mechanism f , the worst-case quality loss metric $Q^+(f, \pi)$ gives a huge value and the probability of getting a useless response from the server would be larger than zero. One could argue that, given the nature of the geo-indistinguishability guarantee, the probability of reporting a location z far from x is low and decreases exponentially with the distance between them, so we could disregard such an event from happening. However, if we really truncate the mechanism to ensure that the probability of going very far is zero, then the mechanism does not provide any geo-indistinguishability guarantee at all. It is then clear that geo-indistinguishability mechanisms are problematic from the quality loss point of view, and if a user gets zero utility from a realization of the mechanism she cannot re-use it immediately, otherwise the privacy guarantee is violated. We comment on a possible solution to this problem below.

4.2.2 Implementation of Mechanisms with Worst-Case Quality Loss Constraint. Now we set the task of designing a mechanism that achieves a good value of worst-case quality loss or, alternatively, that ensures that the worst-case quality loss is below some bound $Q^+(f, \pi) \leq Q_{\max}^+$. The straightforward approach, given a mechanism f , is to truncate the mechanism (for example, by generating samples of z until one of them ensures that $d_Q(x, z) \leq Q_{\max}^+$, and then releasing that z). This approach is reasonable, but one must take into account that the privacy properties of this new truncated mechanism f' are not the same as the original mechanism f , and therefore they must be re-evaluated.

Another issue that concerns the design of bounded mechanisms is that a deterministic remapping (15) might violate a Q^+ constraint (i.e., even if f guarantees the Q^+ constraint, a composition $f' = f \circ g$ might not guarantee it). Finding a bounded mechanism that achieves as much privacy as an unbounded one in $\mathcal{F}_Q^{\text{opt}}$ can be an impossible task, due to the fact that the polytope defined by $Q^+(f, \pi) \leq Q_{\max}^+$ might be disjoint with $\mathcal{F}_Q^{\text{opt}}$. However, we can lose some privacy with respect to an optimal unbounded mechanism in exchange for a better worst-case quality loss guarantee by enforcing the bounding constraint $Q^+(f, \pi) \leq Q_{\max}^+$.

4.3 Other Complementary Metrics

Now, we finally outline other metrics that can be used together with the average error and average quality loss to assess the privacy of mechanisms, and leave the development of mechanisms taking them into account as subject for future work.

Geo-indistinguishability (10) inherently ensures that an input location x is mapped to a nearby location with more probability than

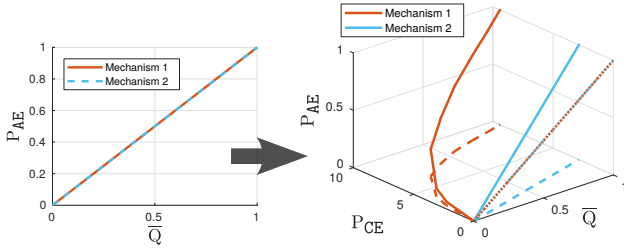


Figure 1: Two mechanisms that perform equally in the P_{AE} vs. $\bar{Q}$ plane, might behave very differently in practice. This is revealed by considering a multi-dimensional characterization of privacy.

to a far location, which solves the privacy issue we illustrated with the coin mechanism. However, this privacy notion is not compatible with a worst-case quality loss constraint by definition, due to the fact that $f(z|x) > 0$ implies $f(z|x') > 0$, $\forall x' \in \mathcal{X}$. A possible approach to solve this utility issue of geo-indistinguishability can be to relax its definition, allowing a small tolerance value $\Delta \ll 1$, i.e.,

$$\int_A f(z|x) dz \leq e^{\epsilon \cdot d_P(x, x')} \cdot \int_A f(z|x') dz + \Delta, \quad \begin{matrix} \forall x, x' \in \mathcal{X}, \\ \forall A \subseteq \mathbb{R}^2. \end{matrix} \quad (23)$$

Other interesting metrics to assess the privacy of mechanisms are those based on the worst-case output. For example, the worst-case output average error, defined as

$$P_{WC-AE}(f, \pi) = \min_{\substack{z \in \mathbb{R}^2 \\ f_Z(z) > 0}} \min_{\hat{x} \in \mathbb{R}^2} \left\{ \sum_{x \in \mathcal{X}} \pi(x) \cdot f(z|x) \cdot d_P(x, \hat{x}) \right\}, \quad (24)$$

measures the average error of the adversary's estimation in the most vulnerable output. When applied to the coin mechanism, this metric would reveal its privacy issue, since $P_{WC-AE}(f_{\text{coin}}, \pi) = 0$.

On the other hand, the worst-case output conditional entropy, defined as

$$P_{WC-CE}(f, \pi) = \min_{\substack{z \in \mathbb{R}^2 \\ f_Z(z) > 0}} \sum_{x \in \mathcal{X}} p(x|z) \cdot \log p(x|z), \quad (25)$$

reveals the uncertainty the adversary has after observing z in the worst case (for the user). If there is any output value z that leaks a lot of information about the real location x (as it happens with every $z \neq z^*$ in the coin mechanism), this metric highlights it.

The metrics introduced throughout this section add additional dimensions to the privacy and quality loss evaluation procedure, revealing features not captured by the standard 2-dimensional approach based on the average error and the average loss. An example of this new characterization of privacy is shown in Fig. 1 where we show the performance of two mechanisms as a 3-D plot of P_{AE} , P_{CE} and $\bar{Q}$, together with the projections in the P_{AE} - $\bar{Q}$ and P_{CE} - $\bar{Q}$ planes. In the next section, we show similar examples (albeit with 2-dimensional plots, for clarity) of particular location privacy preserving mechanisms.

5 EVALUATION

In this section, we assess the performance of different location privacy-preserving mechanisms with respect to different privacy notions. Our experiments confirm that relying on a single metric for evaluation can lead to an erroneous assessment of the privacy provided by a mechanism. We divide our evaluation into two parts. First, we consider the continuous scenario introduced in Section 2 and use real datasets to evaluate the performance of unbounded mechanisms, and of mechanisms that guarantee a maximum worst-case quality loss. Second, we consider a simpler scenario where the locations can only belong to a discrete set, and evaluate other defenses that have been proposed in the literature. All our experiments are performed using Matlab.¹

5.1 Continuous Scenario

For this part of the evaluation, we consider that users are interested in querying about Points of Interest (PoIs) in a discrete set but they can report any point in $\mathbb{R}^2$ to the server (see Section 2). We also consider that the adversary performs her estimation in $\mathbb{R}^2$. We build the set of PoIs using the Gowalla² and Brightkite³ real-world datasets. Following the approach of the finite domain evaluation in [6], we restrict the PoIs to a finite region of San Francisco area between the latitude coordinates (37.5395 and 37.7910) and longitude (-122.5153 and -122.3789). We choose the San Francisco area because it contains a big density of points of interest and a large number of user check-ins, which ensures that the data is rich and representative of what one would expect from users living in the area. On the other hand, considering a finite region allows us to evaluate mechanisms whose computational cost increases with the number of points of interest, such as the exponential and exponential posterior mechanisms. We transform the PoIs into Cartesian coordinates in kilometers using the Haversine formula with respect to the center of the region. We end up with $|\mathcal{X}| = 9701$ PoIs for Gowalla and $|\mathcal{X}| = 8898$ for Brightkite, distributed in an area of roughly $28\text{km} \times 12\text{km}$. As example, the distribution of PoIs for Gowalla is shown in Fig. 2. For each dataset, we compute the prior π by counting how many users check-in on each point of interest and normalizing the resulting histogram. The obtained priors are shown in Fig. 3. We see that, in both datasets, there is a single point of interest x_{top} that draws a lot of attention from the users ($\pi(x_{\text{top}}) \approx 0.04$ in Gowalla and $\pi(x_{\text{top}}) \approx 0.23$ in Brightkite).

We evaluate six location-privacy preserving mechanisms, measuring their performance in terms of the average adversary error (P_{AE}), conditional entropy (P_{CE}) and geo-indistinguishability (P_{GI}) for different values of average quality loss ($\bar{Q}$). We always use the Euclidean distance for the quality loss $d_Q(x, z) = \|x - z\|_2$, and therefore the optimal remapping in (15) is obtained by computing the geometric median of the posterior. We compute this median using Weiszfeld's iterative method. We first evaluate the mechanisms without any bounds on their worst-case quality loss, and then imposing such constraint.

The first three mechanisms we evaluate consist in adding noise in the continuous plane and then remapping them. We generate

¹<https://www.mathworks.com/products/matlab.html>

²<https://snap.stanford.edu/data/loc-gowalla.html>

³<https://snap.stanford.edu/data/loc-brightkite.html>

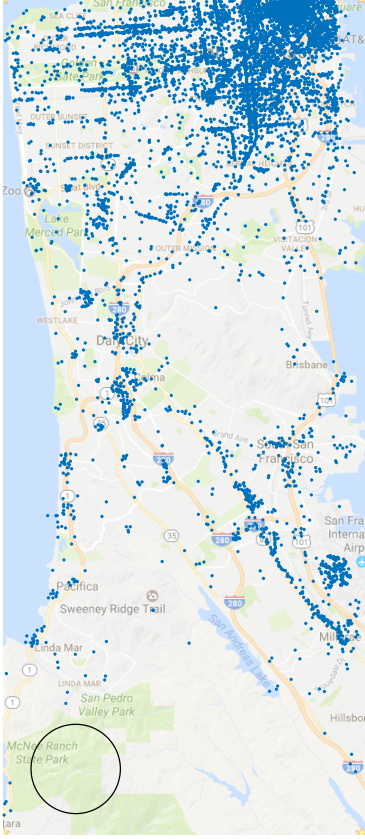


Figure 2: Points of interest in the San Francisco region taken from Gowalla dataset.

this noise in polar coordinates, sampling θ from a uniform distribution in $(0, 2\pi)$ and the radius r from a distribution specified below. Since for these algorithms we cannot find a closed form expression for $f(z|x)$, we evaluate them empirically. To this end we sample π to obtain x , we obtain z adding the noise and performing the remapping, and then we measure privacy according to each metric. We report averages over 5 000 repetitions. These mechanisms are:

- **[Lap] Planar Laplacian noise** plus remapping [6]. To generate the radius of the Laplace noise, we first sample p uniformly in the interval $(0, 1)$. Then, following [2], we set $r = \frac{1}{\epsilon} \left(W_{-1} \left(\frac{p-1}{e} \right) + 1 \right)$ where W_{-1} is the -1 branch of the Lambert W function. We test different values of ϵ from 0.4km^{-1} to 40km^{-1} , so that the average loss varies between 0.05 and 5km.
- **[Gau] Bi-dimensional Gaussian noise** plus remapping. To generate Gaussian noise, we sample the radius from a Rayleigh distribution, varying its mean from 0.05 to 5km.
- **[Cir] Uniform circular noise** plus remapping. In this case, we sample the radius $r \in (0, R)$ from $f(r) = r/R^2$, where R is the maximum radius of the circle, which we vary from 0.075km to 7.5km. This ensures an average loss that varies between 0.05 and 5km.

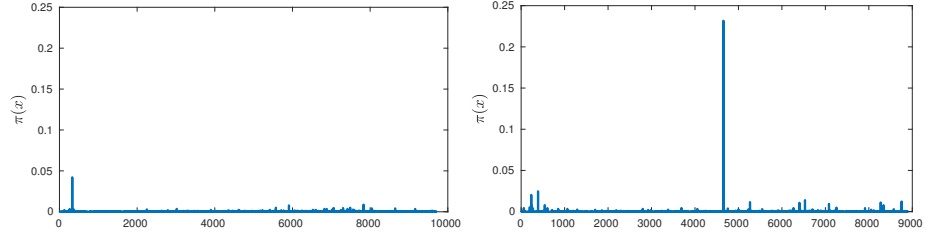


Figure 3: Priors π for Gowalla (left) and Brightkite (right) datasets.

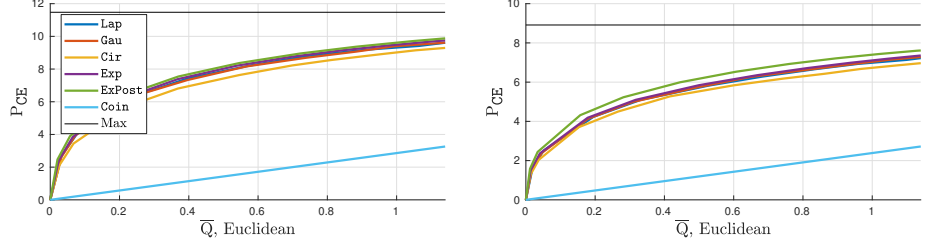


Figure 4: Conditional entropy vs. average quality loss for Gowalla (left) and Brightkite (right) datasets.

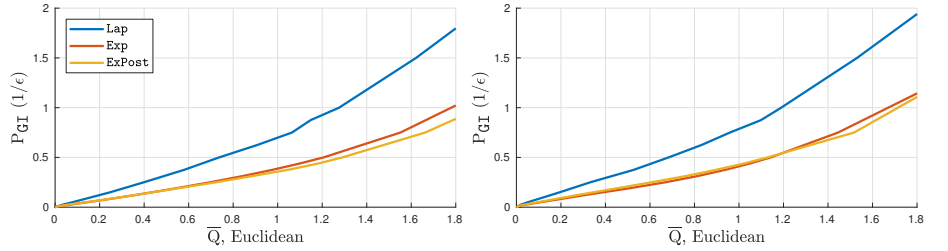


Figure 5: Geo-Ind Privacy P_{GI} vs. average quality loss for Gowalla (left) and Brightkite (right) datasets.

Second, we evaluate three mechanisms that output values in a discrete set, whose conditional probability density functions $f(z|x)$ can be computed arithmetically. This allows us to exactly determine their privacy and quality loss performance. These mechanisms are:

- **[Coin] The coin mechanism**, explained in Sect. 3.2. We vary its average loss $\bar{Q}$ from 0 to 2.
- **[Exp] The Exponential mechanism** plus optimal remapping. The exponential mechanism is a general differential privacy technique that can be applied to provide geo-indistinguishability [10]. We set $\mathcal{Z} = \mathcal{X}$ and set a parameter b , then compute the probability of mapping each input x to an output z as $p(z|x) = a \cdot e^{-b \cdot d_Q(x,z)}$, where a ensures that $\sum_{z \in \mathcal{Z}} p(z|x) = 1$. Then, we apply an optimal remapping to the outputs of this function and obtain $f(z|x)$. In the experiments, we vary b from 0.4km^{-1} and 40km^{-1} .
- **[ExPost] Exponential posterior mechanism**, proposed in Section 4.1.2. In our experiments we set the discrete output alphabet of this algorithm to $\mathcal{Z} = \mathcal{X}$.

5.1.1 Results for unbounded mechanisms (no Q^+ constraint).

When the worst-case quality loss is not constrained, the optimal remapping ensures that all mechanisms are optimal in terms of

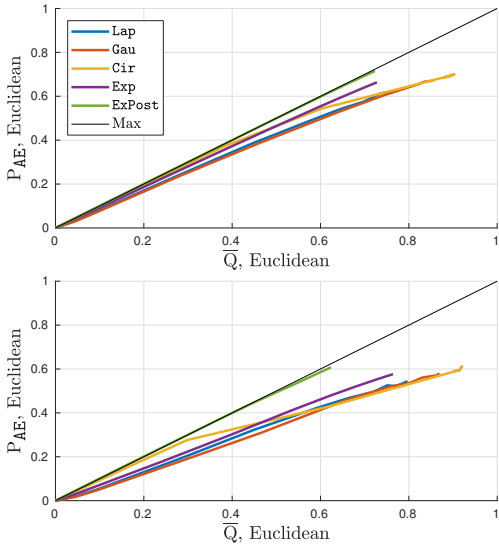


Figure 6: Average error vs. average quality loss for different bounded mechanisms.

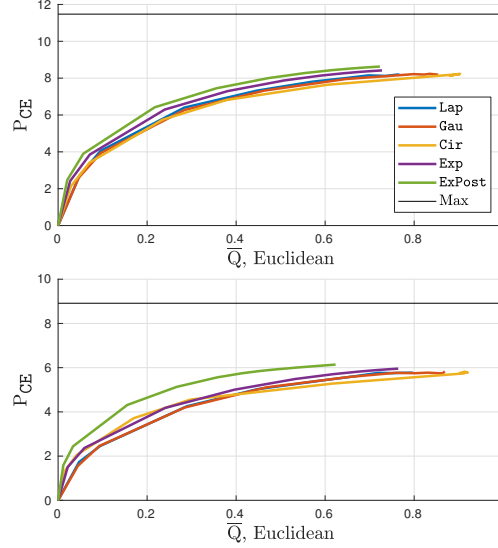


Figure 7: Conditional entropy vs. average quality loss for different bounded mechanisms.

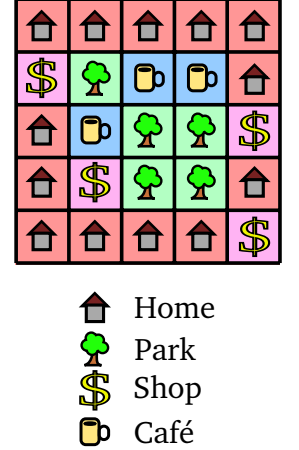


Figure 8: Semantic map of the discrete synthetic scenario.

average error, i.e., $P_{AE} = \bar{Q}$ (see Fig. 11 in the Appendix). This shows that the optimal remapping applied to *any* mechanism achieves an optimal performance, whether it was Laplacian noise or a binary selection of a location such as Coin, as we proved in Sect. 3.

Figure 4 shows the mechanisms' performance in terms of conditional entropy P_{CE} , where the horizontal black line represents the maximum entropy achievable, i.e., the entropy of the prior π . Unsurprisingly, ExPost outperforms the rest of the mechanisms, as it is optimized with respect to this metric. The relative improvement of ExPost with respect to the other algorithms is slightly better in Brightkite than in Gowalla. This is due to the fact that in Brightkite the most frequent PoI is more popular than in Gowalla (see Fig. 3), and thus performing well in this location is crucial to achieve a good overall privacy level in Brightkite. The iterative structure of ExPost allows this mechanism to refine its performance and be more effective than the rest of the mechanisms around this PoI. We note, however, that this refinement comes at the price of an increase in computational cost. Overall, all the mechanisms achieve a similar performance in terms of conditional entropy, except for the coin, that performs poorly. This reinforces the critique in Sect. 3.2: even though Coin is optimal in terms of the average adversary error, measuring its performance in terms of conditional entropy reveals its privacy flaws.

Figure 5 shows the mechanisms' performance in terms of geo-indistinguishability $P_{GI}(f)$ (we recall that $P_{GI}(f) = 1/\epsilon$), only for Lap, Exp and ExPost, as these are the only algorithms that guarantee this property. As already seen in [6], the Laplace noise outperforms the exponential mechanism, and ExPost performs similar to the latter.

5.1.2 Results for bounded mechanisms. We now impose a worst-case quality loss constraint of $Q_{\max}^+ = 1.5\text{km}$ to the mechanisms (as a reference, we show a circle of radius 1.5km in Fig. 2). To implement this constraint in the mechanisms, we truncate their

output at 1.5km and then apply the optimal remapping that respects the worst-case loss constraint. We do this by solving the problem in (15) with constraints. We do not evaluate the coin mechanism in this scenario, since it almost always violates the Q^+ constraint.

The results for the average adversary error as Euclidean distance are shown in Fig. 6. As expected, the mechanisms obtained after the remapping in this scenario are not necessarily optimal. We see that ExPost achieves a result that is close to the optimal mechanism in the unbounded case, while the other mechanisms achieve less average privacy. We conjecture this is due to the iterative nature of ExPost, that refines its performance, while the other mechanisms are not optimized regarding the worst-case loss constraint. Again, ExPost achieves a wider advantage in Brightkite for the same reason explained above.

Figure 7 shows the performance of the bounded mechanisms in terms of conditional entropy. The results are similar to those in the unbounded scenario, with ExPost outperforming the others with a slightly wider advantage in this case. As bounded mechanisms do not achieve geo-indistinguishability, we do not evaluate the performance with respect to this metric in this scenario.

5.2 Discrete scenario

We now consider a simple synthetic scenario and evaluate the optimal mechanisms obtained following the method by Shokri et. al [27]. In this work, the authors propose a linear program that finds a mechanism f inside the polytope of optimal mechanisms for P_{AE} given a constraint $\bar{Q}$, i.e., $f \in \mathcal{F}_Q^{\text{opt}}$. This approach is very versatile, as it can be computed for any pair of distance functions $d_P(\cdot)$ and $d_Q(\cdot)$. We set our synthetic scenario under the assumptions of that work: the input and output alphabets are discrete and identical $\mathcal{X} = \mathcal{Z}$, and the adversary can only estimate locations inside that same alphabet $\hat{\mathcal{X}} = \mathcal{X}$. For simplicity, we consider that the set of locations in $\mathcal{X}$ are the centers of the cells that make a 5×5

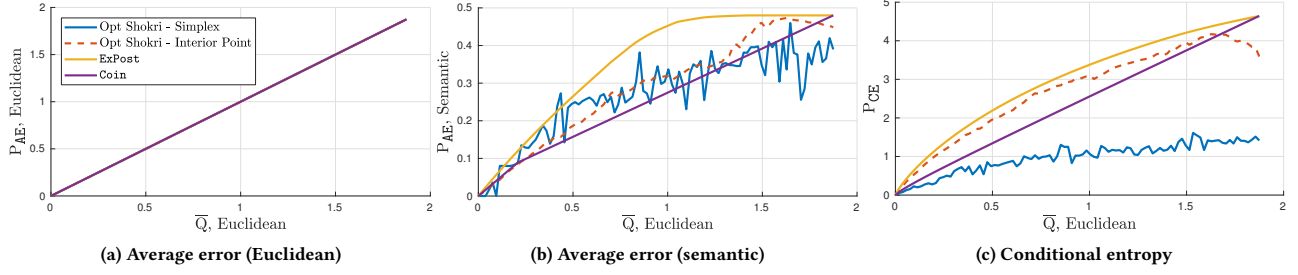


Figure 9: Performance of Shokri et. al’s algorithm optimized for the adversary error in terms of Euclidean distance, compared to the coin mechanism and exponential posterior mechanism.

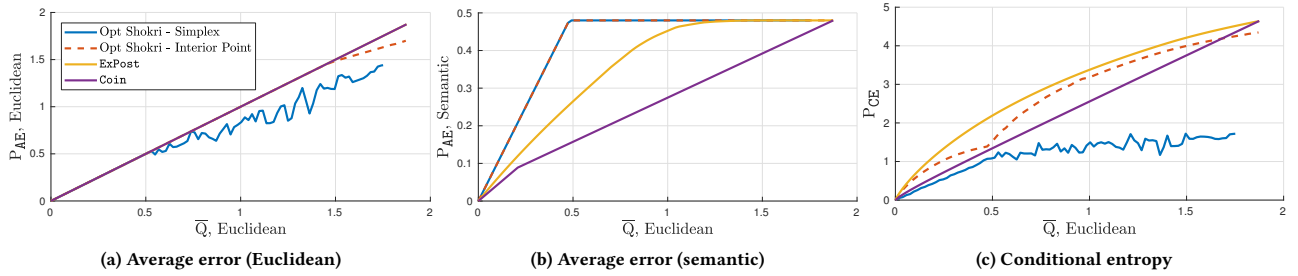


Figure 10: Performance of Shokri et. al’s algorithm optimized for the adversary error in terms of semantic distance, compared to the coin mechanism and the exponential posterior mechanism.

square grid and assign a tag to each location that can be “Home”, “Park”, “Shop” or “Café”, as depicted in Fig. 8. We consider that the prior is uniform $\pi(x) = 1/25$, $\forall x \in \mathcal{X}$. We measure the point-wise loss as the Euclidean distance $d_Q(x, z) = \|x - z\|_2$ and consider two point-wise metrics of privacy: the Euclidean distance and a semantic distance defined as the Hamming distance between tags, i.e., $d_P(x, z) = 0$ if $\text{Tag}(x) = \text{Tag}(z)$, and $d_P(x, z) = 1$ otherwise. This metric is similar to the semantic metric in [1]. The average error computed using this distance function represents the probability that an adversary guesses incorrectly the tag of x .

We evaluate ExPost and Coin together with the optimal mechanism proposed in [27]. For the latter, we solve the linear program to find optimal mechanisms in terms of maximizing P_{AE} using the Euclidean distance (Fig. 9) and the semantic distance we defined (Fig. 10). As expected, the optimal mechanisms (Shokri et. al) achieve the optimal privacy when evaluated using the adversary’s error for which they are optimized (Figs. 9a and 10b), but not when evaluated against a different metric (Figs. 9b and 10a). ExPost and Coin achieve maximum privacy in terms of Euclidean distance, as before, but not in terms of semantic distance. This example emphasizes that optimizing a mechanism with respect to a privacy metric may provide very bad performance with respect to other privacy criteria.

This experiment also shows another important idea: even though the solutions of the linear program both achieve approximately the same performance in terms of average error (optimal in Figs. 9a and 10b, suboptimal in Figs. 9b and 10a), they exhibit a radically different behavior in terms of conditional entropy. Indeed, using the mechanism computed with the simplex algorithm (a mechanism at

a vertex of $\mathcal{F}_Q^{\text{opt}}$), the adversary has much less uncertainty about x on average than if the user had implemented a mechanism from the interior of the polytope. This difference in entropy is also what allows us to tell apart a mechanism such as ExPost from Coin. Note that the mechanism computed by solving the linear program with the simplex algorithm performs even worse than the coin in terms of entropy, illustrating the dangers of optimizing privacy in only one dimension.

6 CONCLUSIONS

In this work, we have demonstrated the problems of using a single privacy metric as indicator of the performance of location privacy preserving mechanisms. We have proven that there is more than one optimal protection mechanism in terms of maximizing the average adversary error for a given average quality loss, and that the family of mechanisms that fulfill such condition behave differently in terms of other privacy metrics. Thus, optimizing defenses with only one privacy metric in mind may lead to mechanisms that offer poor protection in other dimensions of privacy. To avoid selecting underperforming mechanisms we propose the use of complementary criteria to guide the choice. We provide two example auxiliary metrics: the conditional entropy and the worst-case loss. We propose an optimal mechanism with respect to the former, and provide means to implement mechanisms according to the latter.

We evaluate the mechanisms, comparing them to previous work, on two real datasets. Our experiments confirm two important ideas: first, that we cannot find a mechanism that performs optimally with

respect to every privacy metric. Second, that even if a mechanism performs well in a particular metric it does not imply that it is necessarily beneficial for the user. Our findings reveal the need to take a step back in mechanism design to integrate privacy as a multi-dimensional notion, in order to avoid solutions that provide a false perception of privacy.

A APPENDIX

A.1 Proof of Theorem 3.3

In order to prove this result, first notice that, when $d_P(\cdot) \equiv d_Q(\cdot)$, the quality loss $\bar{Q}$ is an upper bound of privacy P_{AE} :

$$\begin{aligned} P_{AE}(f, \pi) &= \int_{\mathbb{R}^2} \min_{\hat{x} \in \mathbb{R}^2} \left\{ \sum_{x \in \mathcal{X}} \pi(x) \cdot f(z|x) \cdot d_P(x, \hat{x}) \right\} dz \\ &\leq \int_{\mathbb{R}^2} \left\{ \sum_{x \in \mathcal{X}} \pi(x) \cdot f(z|x) \cdot d_Q(x, z) \right\} dz = \bar{Q}(f, \pi), \quad (26) \end{aligned}$$

Now, assume that $f' = f \circ g$, and therefore

$$z = \operatorname{argmin}_{z' \in \mathbb{R}^2} \sum_{x \in \mathcal{X}} \pi(x) \cdot f'(z|x) \cdot d_Q(x, z'). \quad (27)$$

The optimal adversary estimation of x given z given in (4) can be written as

$$\hat{x}(z) = \operatorname{argmin}_{\hat{x} \in \mathbb{R}^2} \sum_{x \in \mathcal{X}} \pi(x) \cdot f'(z|x) \cdot d_P(x, \hat{x}). \quad (28)$$

We see that since $d_P(\cdot) \equiv d_Q(\cdot)$ the optimal adversary estimation is doing nothing, i.e., $\hat{x}(z) = z$. This implies that $P_{AE}(f', \pi) = \bar{Q}(f', \pi)$, and since we have achieved the upper bound on privacy given in (26), f' is optimal.

A.2 Geo-indistinguishability of the posterior exponential mechanism.

We recall that the geo-indistinguishability guarantee requires the following condition to be fulfilled (now written for discrete mechanisms, where $p(z|x)$ denotes the probability of reporting z when the original location is x):

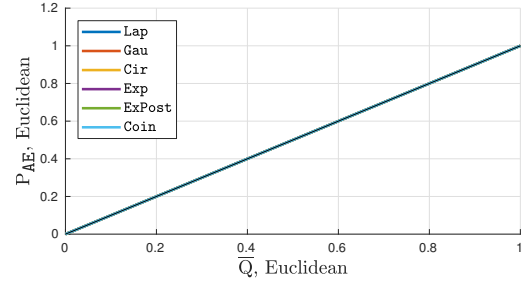
$$p(z|x) \leq e^{\epsilon \cdot d_P(x, x')} \cdot p(z|x'), \quad \forall x, x' \in \mathcal{X}, z \in \mathcal{Z}, \quad (29)$$

where $d_P(x, x')$ is the Euclidean distance.

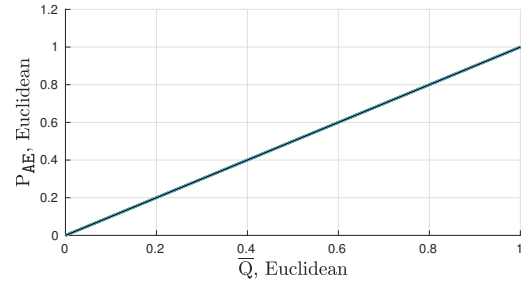
The last iteration of the ExPost algorithm in 4.1.2 returns a mechanism that can be written for a particular input x and output z as

$$p(z|x) = \begin{cases} \frac{P_Z(z) \cdot e^{-b \cdot d_Q(x, z)}}{\sum_{z' \in \mathcal{Z}} P_Z(z') \cdot e^{-b \cdot d_Q(x, z')}} & \text{if } P_Z(z) > 0, \\ 0, & \text{if } P_Z(z) = 0. \end{cases} \quad (30)$$

where $d_Q(x, z)$ is the Euclidean distance. In the second case, the geo-indistinguishability guarantee is trivially achieved since given any pair of input locations $x, x' \in \mathcal{X}$, $p(z|x) = p(z|x') = 0$. For the first case, we use the triangular inequality $d_Q(x, z) + d_Q(x', z) \geq$



(a) Gowalla



(b) Brightkite

Figure 11: Average error vs. average quality loss for different unbounded mechanisms.

$d_Q(x, x')$ to write

$$p(z|x) = \frac{P_Z(z) \cdot e^{-b \cdot d_Q(x, z)}}{\sum_{z' \in \mathcal{Z}} P_Z(z') \cdot e^{-b \cdot d_Q(x, z')}} \quad (31)$$

$$\leq \frac{P_Z(z) \cdot e^{b \cdot d_Q(x, x')} \cdot e^{-b \cdot d_Q(x', z)}}{\sum_{z' \in \mathcal{Z}} P_Z(z') \cdot e^{-b \cdot d_Q(x, z')}} \quad (32)$$

$$\leq \frac{P_Z(z) \cdot e^{b \cdot d_Q(x, x')} \cdot e^{-b \cdot d_Q(x', z)}}{\sum_{z' \in \mathcal{Z}} P_Z(z') \cdot e^{-b \cdot d_Q(x, x')} \cdot e^{-b \cdot d_Q(x', z')}} \quad (33)$$

$$= \frac{P_Z(z) \cdot e^{-b \cdot d_Q(x', z)}}{\sum_{z' \in \mathcal{Z}} P_Z(z') \cdot e^{-b \cdot d_Q(x', z')}} \cdot e^{2b \cdot d_Q(x, x')} \quad (34)$$

$$= e^{2b \cdot d_Q(x, x')} \cdot p(z|x'), \quad (35)$$

which satisfies the geo-indistinguishability for $\epsilon = 2b$ or $P_{GI} = 1/2b$, if $d_Q(\cdot)$ is the Euclidean distance. This concludes the proof.

A.3 Performance of the unbounded mechanisms in terms of the average error

When the average error (Euclidean) and the average quality loss (Euclidean) are used to evaluate the performance of the mechanisms described in Section 5, we achieve the trivial result $P_{AE} = \bar{Q}$. This is shown in Fig. 11 for completeness.

ACKNOWLEDGMENTS

This work is partially supported by EU H2020-ICT-10-2015 NEXTLEAP (GA n 688722), the Agencia Estatal de Investigación (Spain) and the European Regional Development Fund (ERDF) under projects

WINTER (TEC2016-76409-C2-2-R) and COMONSENS (TEC2015-69648-REDC), and by the Xunta de Galicia and the European Union (European Regional Development Fund - ERDF) under projects Agrupación Estratéxica Consolidada de Galicia accreditation 2016-2019 and Red Temática RedTEIC 2017-2018. Simon Oya is funded by the Spanish Ministry of Education, Culture and Sport under the FPU grant.

REFERENCES

- [1] Berker Ağır, Kévin Huguenin, Urs Hengartner, and Jean-Pierre Hubaux. 2016. On the Privacy Implications of Location Semantics. *Proceedings on Privacy Enhancing Technologies* 2016, 4 (2016), 165–183.
- [2] Miguel E Andrés, Nicolás E Bordenabe, Konstantinos Chatzikokolakis, and Catuscia Palamidessi. 2013. Geo-indistinguishability: Differential privacy for location-based systems. In *Proceedings of the 2013 ACM SIGSAC conference on Computer & communications security*. ACM, 901–914.
- [3] Alastair R. Beresford and Frank Stajano. 2003. Location Privacy in Pervasive Computing. *IEEE Pervasive Computing* 2, 1 (2003), 46–55.
- [4] Igor Bilogrevic, Kévin Huguenin, Stefan Mihaila, Reza Shokri, and Jean-Pierre Hubaux. 2015. Predicting users' motivations behind location check-ins and utility implications of privacy protection mechanisms. In *22nd Network and Distributed System Security Symposium (NDSS)*.
- [5] Nicolás E Bordenabe, Konstantinos Chatzikokolakis, and Catuscia Palamidessi. 2014. Optimal geo-indistinguishable mechanisms for location privacy. In *Proceedings of the 2014 ACM SIGSAC Conference on Computer and Communications Security*. ACM, 251–262.
- [6] Konstantinos Chatzikokolakis, Ehab Elsalamouny, and Catuscia Palamidessi. 2016. Practical Mechanisms for Location Privacy. (2016).
- [7] Konstantinos Chatzikokolakis, Catuscia Palamidessi, and Marco Stronati. 2015. Constructing elastic distinguishability metrics for location privacy. *Proceedings on Privacy Enhancing Technologies* 2015, 2 (2015), 156–170.
- [8] Thomas M Cover and Joy A Thomas. 2012. *Elements of information theory*. John Wiley & Sons.
- [9] Cynthia Dwork. 2006. Differential Privacy. In *Automata, Languages and Programming, 33rd International Colloquium, ICALP 2006 (Lecture Notes in Computer Science)*, Michele Bugliesi, Bart Preneel, Vladimiro Sassone, and Ingo Wegener (Eds.), Vol. 4052. Springer, 1–12.
- [10] Cynthia Dwork. 2008. Differential privacy: A survey of results. In *International Conference on Theory and Applications of Models of Computation*. Springer, 1–19.
- [11] Kassem Fawaz, Huan Feng, and Kang G. Shin. 2015. Anatomization and Protection of Mobile Apps' Location Privacy Threats. In *24th USENIX Security Symposium*, Jaeyoon Jung and Thorsten Holz (Eds.). USENIX Association, 753–768.
- [12] Kassem Fawaz and Kang G. Shin. 2014. Location Privacy Protection for Smartphone Users. In *ACM SIGSAC Conference on Computer and Communications Security*, Gail-Joon Ahn, Moti Yung, and Ninghui Li (Eds.). ACM, 239–250.
- [13] Julien Freudiger, Reza Shokri, and Jean-Pierre Hubaux. 2012. Evaluating the privacy risk of location-based services. In *Financial Cryptography and Data Security*. Springer, 31–46.
- [14] Sébastien Gambs, Marc-Olivier Killijian, and Miguel Núñez del Prado Cortez. 2011. Show Me How You Move and I Will Tell You Who You Are. *Transactions on Data Privacy* 4, 2 (2011), 103–126.
- [15] Bugra Gedik and Ling Liu. 2005. Location Privacy in Mobile Systems: A Personalized Anonymization Model. In *25th International Conference on Distributed Computing Systems (ICDCS)*. IEEE Computer Society, 620–629.
- [16] Philippe Golle and Kurt Partridge. 2009. On the Anonymity of Home/Work Location Pairs. In *International Conference on Pervasive Computing (LNCS)*, Hideyuki Tokuda, Michael Beigl, Adrian Friday, A. J. Bernheim Brush, and Yoshito Tobe (Eds.), Vol. 5538. Springer, 390–397.
- [17] Marco Gruteser and Dirk Grunwald. 2003. Anonymous Usage of Location-Based Services Through Spatial and Temporal Cloaking. In *International conference on Mobile systems, applications and services*. ACM, 31–42.
- [18] B. Hoh and M. Gruteser. 2005. Protecting Location Privacy Through Path Confusion. In *International Conference on Security and Privacy for Emerging Areas in Communications Networks*. 194–205. <https://doi.org/10.1109/SECURECOMM.2005.33>
- [19] H. Kido, Y. Yanagisawa, and T. Satoh. 2005. An anonymous communication technique using dummies for location-based services. In *Pervasive Services, 2005. ICPS '05. Proceedings. International Conference on*. 88–97.
- [20] John Krumm. 2007. Inference Attacks on Location Tracks. In *5th International Conference on Pervasive Computing (LNCS)*, Anthony LaMarca, Marc Langheinrich, and Khai N. Truong (Eds.), Vol. 4480. Springer, 127–143.
- [21] Hua Lu, Christian S. Jensen, and Man Lung Yiu. 2008. PAD: privacy-area aware, dummy-based location privacy in mobile services. In *ACM International Workshop on Data Engineering for Wireless and Mobile Access*. ACM, 16–23. <https://doi.org/10.1145/1626536.1626540>
- [22] Changsha Ma and Chang Wen Chen. 2014. Nearby Friend Discovery with Geo-indistinguishability to Stalkers. *Procedia Computer Science* 34 (2014), 352–359.
- [23] Joseph T. Meyerowitz and Romit Roy Choudhury. 2009. Hiding stars with fire-works: location privacy through camouflage. In *15th Annual International Conference on Mobile Computing and Networking (MOBICOM)*, Kang G. Shin, Yongguang Zhang, Rajive Bagrodia, and Ramesh Govindan (Eds.). ACM, 345–356.
- [24] Reza Shokri. 2015. Privacy Games: Optimal User-Centric Data Obfuscation. *PoPETS* 2015, 2 (2015), 299–315.
- [25] Reza Shokri, Julien Freudiger, Murtuza Jadhwal, and Jean-Pierre Hubaux. 2009. A distortion-based metric for location privacy. In *ACM Workshop on Privacy in the Electronic Society, WPES*, Ehab Al-Shaer and Stefano Paraboschi (Eds.). ACM, 21–30.
- [26] Reza Shokri, George Theodorakopoulos, Jean-Yves Le Boudec, and Jean-Pierre Hubaux. 2011. Quantifying location privacy. In *Security and privacy (sp), 2011 IEEE symposium on*. IEEE, 247–262.
- [27] Reza Shokri, George Theodorakopoulos, Carmela Troncoso, Jean-Pierre Hubaux, and Jean-Yves Le Boudec. 2012. Protecting location privacy: optimal strategy against localization attacks. In *Proceedings of the 2012 ACM conference on Computer and communications security*. ACM, 617–627.
- [28] Yu Wang, Dingbang Xu, Xiao He, Chao Zhang, Fan Li, and Bin Xu. 2012. L2P2: Location-aware location privacy protection for location-based services. In *INFOCOM, 2012 Proceedings IEEE*. 1996–2004. <https://doi.org/10.1109/INFOCOM.2012.6195577>
- [29] Tun-Hao You, Wen-Chih Peng, and Wang-Chien Lee. 2007. Protecting Moving Trajectories with Dummies. In *International Conference on Mobile Data Management*. 278–282.
- [30] Yu Zheng, Lizhu Zhang, Xing Xie, and Wei-Ying Ma. 2009. Mining Interesting Locations and Travel Sequences from GPS Trajectories. In *Proceedings of the 18th International Conference on World Wide Web*. ACM, 10.