

Measuring Privacy Leakage in Term of Shannon Entropy

Ricky Aditya^{1,*}, Boris Skoric²

¹*Department of Mathematics, Faculty of Science and Technology,
Sanata Dharma University, Yogyakarta, Indonesia*

²*Security Group, Eindhoven University of Technology, Eindhoven,
The Netherlands*

**Corresponding Author: y_ricky_aditya@yahoo.com*

(Received 17-05-2019; Revised 28-06-2019; Accepted 31-07-2019)

Abstract

Differential privacy is a privacy scheme in which a database is modified such that each user's personal data are protected without affecting significantly the characteristics of the whole data. Example of such mechanism is Randomized Aggregatable Privacy-Preserving Ordinal Response (RAPPOR). Later it is found that the interpretations of privacy, accuracy and utility parameters in differential privacy are not totally clear. Therefore in this article an alternative definition of privacy aspect are proposed, where they are measured in term of Shannon entropy. Here Shannon entropy can be interpreted as number of binary questions an aggregator needs to ask in order to learn information from a modified database. Then privacy leakage of a differentially private mechanism is defined as mutual information between original distribution of an attribute in a database and its modified version. Furthermore, some simulations using the MATLAB software for special cases in RAPPOR are also presented to show that this alternative definition does make sense.

Keywords: differential privacy, RAPPOR, Shannon entropy, mutual information, privacy leakage.

1 Introduction

In digitalized era when many things can be done online, privacy becomes a more serious issue, especially if our personal data have to be submitted online for some reasons. Even with their published privacy policies (something that most users never read it properly), there are some room for privacy violations. Here we will not talk about the hackers or any outsiders, because the ones who violate privacy might come from the authorized parties.

The most annoying case is when some parties use their authorities to leak someone's private data but there is no laws or rules which can conclude it as a privacy violation and therefore they cannot be punished. For example, our medical record data which are recorded in a hospital's database. Our data, together with other persons' data, might be used by other parties who want to learn something from the database, let us say a medicine company or a medical research center. We never know if they really just access the database for gaining only the necessary information, or they may search for our personal data.

A basic and simplest way to prevent this is by hiding the names of data owners, i.e. making the data to be anonymous. Unfortunately, this may be not enough to protect our private data. They can still access any other data, such as height, weight, age, gender, etc. Consider some persons with a very rare attribute, for examples : too tall, too short, too fat, too thin, and many more. By looking at one specific attribute or two, they can uniquely determine them and as consequence, can leak their private information. They, of course, violate those persons' privacy but we cannot say that they break any laws or rules in the privacy policies. Suppose that someone is famous as the tallest guy in his/her city. Roughly saying, as long as they do not ask the hospital who the tallest guy in this database is, and the hospital do not inform it either, no laws or rules are broken.

Based on this kind of issues, many data security researchers try to create a new privacy protocol to protect any private information. One of them is called as differential privacy. The idea is to modify the original database such that each user's personal data

are protected but characteristics of the whole database do not change significantly. Therefore other parties are still able to learn any information about the whole database but they are unable to learn any personal information.

As a very simple example, there are five persons : A, B, C, D and E. The fact is A and B are smokers, while the others not. After modification, the smokers become C and E. Here the fact that A and B are smokers is hidden, but it preserves the fact that two of those five persons are smokers. Note that other parties know that the database has been modified, so they cannot judge C and E as smokers. Therefore if they just want to know the proportion of smokers in the database, they will not get it wrong but they will not know who the real smokers are.

In practical case, of course, we will work on much larger database with various attributes. We do not have to preserve the exact proportion of any attributes, but we need to keep it with a small margin of errors. The concept of differential privacy will be discussed in the next section, together with some specific mechanisms which can be used.

2 Differential Privacy and RAPPOR

The idea of differential privacy came first in Dwork et.al. [1] in 2006. In their work, an idea to protect privacy by adding noise to the data is introduced. At that time, it had not been named as differential privacy, the name came later after some subsequent research. After few years working thoroughly on this area, a more comprehensive concept of differential privacy are later published in Dwork and Roth [2]. Concepts and definitions in this section are based on [1] and [2].

2.1 Differential Privacy

Now we go to the definition of differential privacy. Let a database is represented in a table in which the rows represent the users and the columns represent the attributes. Sometimes the parties who have authorized access to the database only need to take some samples of users and not all of them. We do not always know what they want to look for, but we can assume that they have full authorities to do so.

We say that two sub-databases are neighboring to each other if one is obtained by adding or deleting one row from the other. If the database is not modified, it is possible to learn about one specific user by learning two databases : one database that containing him/her and the database that is obtained by eliminating him/her from the previous one. Therefore, in order to protect that user's privacy the modification mechanism needs to eliminate this possibility. This leads to a definition of differentially private mechanism.

Definition 2.1. Let A be a mechanism to modify a database D with D' as output. The mechanism A is said to be (ϵ, δ) -differentially private, where both ϵ and δ are non-negative numbers, if for any neighboring sub-databases x_1 and x_2 of D , and for any subset S of D' , it satisfies :

$$\Pr[A(x_1) \in S] \leq e^\epsilon \cdot \Pr[A(x_2) \in S] + \delta \quad (1)$$

Equation (1) can be interpreted as the outputs from two neighboring databases has only very small and insignificant difference such that (almost) nothing can be learned about the user who differs them. If the numbers ϵ dan δ be smaller, then the differences become more insignificant and the privacy becomes stronger. In some specific cases, the parameter δ in (1) is set to be 0 and then the mechanism is said to be ϵ - differentially private.

Now we talk about the accuracy of a differential privacy mechanism. In this context we are concerned about information from a database which can be used to answer predicate counting queries. The class of those queries is called as concept class, usually denoted by C . Set of any possible values of a database is called as data universe, usually denoted by X . Output of a predicate counting query c on a database x , denoted by $c(x)$, is the proportion of elements in x which satisfy that predicate. For example, proportion of smokers or proportion of patients with heart problem in a medical record database. Then we have this definition of accuracy.

Definition 2.2. For any $c \in C$, a mechanism A on database x is said to be α -accurate for c if $|c(x) - c(A(x))| < \alpha$. Moreover, A is said to be (α, γ) -accurate for a concept class C if A is α -accurate for $(1-\gamma)$ fractions of queries in C .

Above definition can be interpreted as even though each personal data has been modified, but the proportion of users who satisfy a predicate does not change significantly. We need α to be smaller for a better accuracy. Considering that it is very difficult to create a mechanism that can be accurate for all queries in a concept class, then parameter γ is introduced. If γ is smaller, then more queries can be answered accurately. Furthermore, an utility parameter of a mechanism can also be defined based on its accuracy parameter.

Definition 2.3. *Let C be a concept class and X is a data universe. A modification mechanism A is said to have (α, β, γ) -utility with respect to C and X if for a database x it holds that $\Pr[A \text{ is } (\alpha, \gamma)\text{-accurate}] \leq \beta$.*

There are several kind of mechanisms which can be used to modify database which satisfy differential privacy principles. In this section we will introduce the Randomized Aggregatable Privacy-Preserving Ordinal Response (RAPPOR) mechanism. The next sub-section will discuss more about RAPPOR.

2.2 Randomized Aggregatable Privacy-Preserving Ordinal Response (RAPPOR)

Let a database consists of several attributes which each of them can be divided into several categories. For example, we can categorize people according to their genders (male/female), age range ($\leq 20, 21 - 40, 41 - 60, \geq 60$) and many more. Then for each attribute, each user is represented as the category he/she belongs to. To represent in which category a user belongs to, we can also use as a binary vector with exactly one 1 and 0 otherwise, where position of the 1 denotes the category he/she belongs to.

These binary vector representations will then be modified randomly based on a probability distribution and sent to other parties. Thus they will receive an already modified database. To learn about distribution of categories for each attribute, they have to take the aggregate values of each category. Because of this, later we will call them as data aggregator. The data aggregator does not know the actual distribution of categories, but he/she may know the probability distribution that is used to modify the database.

However, this knowledge should not be enough to leak actual information of the entire database.

There are several kind of RAPPOR mechanisms, as presented in Wang et.al. [3]. In this article we will discuss two kind of RAPPOR mechanisms, which are

1. RAPPOR with direct representation

This kind of data aggregation mechanism works as follows. Let there are m categories in an attribute and a user i belongs to category C_i' . After modification, user i belongs to category C_i' . Probability that user i still in his/her actual category ($C_i = C_i'$) is γ and for each category j where $j \neq C_i'$, probability of user i belongs to category j after modification is $\gamma/(m - 1)$.

2. RAPPOR with unary representation.

In this mechanism, category of a user i is represented as a binary vector $X_i = (X_{i1}, X_{i2}, \dots, X_{im})$ where $X_{ij} = 1$ if user i belongs to category j and otherwise $X_{ij} = 0$. Then this binary vector will be modified by adding noise independently on each bit. Here a bit 0 can be flipped to 1 or vice versa. For each bit, probability of binary flip from 0 to 1 is β_0 and probability of binary flip from 1 to 0 is β_1 . If $\beta_0 = \beta_1$, then it is called as symmetric scheme. The modified vector is then denoted as Z_i and this will be sent to the aggregator. Note that after modification, it is possible to have more than one 1s or no 1s at all.

In next sections, we will not discuss the privacy and accuracy aspects using Definition 2.1. and Definition 2.2., but we will use a different approach instead, that is, by using concepts from information theory and we will see how it could work.

3 Re-defining Privacy Leakage in Term of Shannon Entropy

There are some open problems from the concepts of differential privacy explained in the previous section. For example, in a differentially private scheme, we want to determine the values of ϵ and δ such that its privacy can be considered as good enough and the values of α , β and γ such that it has good accuracy and/or utility. We are also

interested in the practical interpretation of those parameters in a specific mechanism and how changes of one or two parameters affect the others.

It is difficult to answer those questions since we do not have a well-defined measurements of some parameters in differential privacy. Thus we might need another way of measuring the strength of privacy and accuracy. In Wang et.al. [4], an idea that linked differential privacy and mutual-information privacy was introduced. Therefore it should be possible to learn differential privacy using information theoretic approach. In this section we will use similar idea to re-define some aspects of differential privacy in the language of information theory.

3.1 Shannon Entropy and Mutual Information

Intuitively, stronger privacy will imply worse accuracy and vice versa. As a consequence, we cannot have both aspects at each highest level and we should try to find a solution for “optimizing” both privacy and accuracy. Therefore their measurements have to be “sensibly comparable”. In this section an alternative definition for privacy aspect in differential privacy based on information theory point of view will be introduced. Some basic definitions in information theory, based on Cover and Thomas [5], will be revisited first.

Definition 3.1. Let X be a random variable with probability distribution P and probability mass function $p_x = \Pr[X = x]$. Shannon entropy of X , denoted by $H(X)$, is defined as :

$$H(X) = E[-\log(p_x)] = \sum_{x \in X} p_x \cdot \log\left(\frac{1}{p_x}\right). \quad (2)$$

Binary entropy function h of an event with probability p is defined as :

$$h(p) = p \cdot \log\left(\frac{1}{p}\right) + (1-p) \cdot \log\left(\frac{1}{1-p}\right). \quad (3)$$

In some books, Shannon entropy is often called just by the word “entropy”. There are many interpretations of Shannon entropy. One of them is the number of binary (yes/no) questions which need to be asked in order to learn an output if the probability

distribution is known. This interpretation might be not totally accurate, but it is sensible enough to define privacy aspect. If the aggregator needs to ask too many questions in order to learn about an individual data, then we can say that the privacy is strong enough.

After being modified, a database might still give some partial information about its actual data. By learning an already modified database, an aggregator might be able to leak some actual information without knowing the original one. This “leakage” can be represented as mutual information between an original database and its modified version. The following is the definition of mutual information.

Definition 3.2. *Let X and Y be two random probability distributions. Mutual information between X and Y , denoted by $I(X;Y)$, can be computed using these equivalent formulas:*

$$\begin{aligned} I(X;Y) &= H(X) - H(X|Y) = H(Y) - H(Y|X) \\ &= H(X) + H(Y) - H(XY) = H(XY) - H(X|Y) - H(Y|X) \end{aligned} \quad (4)$$

Moreover, if X and Y have joint probability distribution p_{xy} and their respective marginal probability mass functions are $p_x = \sum_y p_{xy}$ and $p_y = \sum_x p_{xy}$, then :

$$I(X;Y) = \sum_{x,y} p_{xy} \cdot \log \left(\frac{p_{xy}}{p_x \cdot p_y} \right) \quad (5)$$

Based on Shannon entropy and mutual information as in Definition 3.1. and 3.2., we can create new definitions of privacy aspect of differential privacy. These will be discussed in next sub-section.

3.2 Alternative Definitions of Privacy Leakage

If we go to our implementation, then we can directly get an idea to define a privacy leakage. Suppose that C is the actual distribution of an attribute in a database and after modification the distribution becomes C' . The privacy leakage can be defined as mutual information between C and C' , i.e. $I(C;C')$, which can be interpreted as the number of binary questions asked by an aggregator to learn information about actual distribution C that can be answered by his/her knowledge on modified distribution C' . By this

interpretation, a stronger privacy scheme should have smaller value of mutual information between actual distribution and modified distribution.

Meanwhile, defining accuracy and utility are trickier. If we want to make everything well-defined, then we need a sensible interpretation about utility in term of the number of binary questions. This is for making a comparable measure between privacy and accuracy-utility. If one entropy in privacy has different interpretation with one entropy in accuracy/utility, then these measures are incomparable and we will not get what we expect in the beginning. This can be an open problem for any possible further research. In next section we are going to do some simulations using the MATLAB software to justify whether our definitions of privacy leakage and utility really make sense or not.

4 Simulation using the MATLAB Software

Alternative definition of privacy leakage introduced in previous section look make sense, but sometimes we need to justify them using some simulations in real and practical cases. Here we do simulations on privacy leakage first. Since computation of a big enough database would take long enough time to compute, we start with some special cases in small database which their computations do not take much time to complete.

4.1 Case I : RAPPOR with direct representation

Recall the mechanism of RAPPOR with direct representation introduced in Section 2. In this mechanism, a user which belongs to a category will have probability $1 - \gamma$ to stay in his/her actual category and probability $\gamma/(m - 1)$ to move into each of other categories, where m denotes the number of categories. Therefore if we know the actual distribution C , we can compute the entropy of conditional probability distribution $C'|C$ as below

$$\begin{aligned} H(C'|C) &= H\left(\left(1-\gamma, \frac{\gamma}{m-1}, \dots, \frac{\gamma}{m-1}\right)\right) = (1-\gamma) \cdot \log \frac{1}{1-\gamma} + (m-1) \cdot \left[\frac{\gamma}{m-1} \cdot \log \left(\frac{m-1}{\gamma} \right) \right] \\ &= (1-\gamma) \cdot \log \frac{1}{1-\gamma} + \gamma \cdot \left[\log(m-1) + \log \frac{1}{\gamma} \right] = (1-\gamma) \cdot \log \frac{1}{1-\gamma} + \gamma \cdot \log \frac{1}{\gamma} + \gamma \cdot \log(m-1) \end{aligned}$$

$$= h(\gamma) + \gamma \cdot \log(m-1) \quad (6)$$

Now we want to compute entropy of modified distribution C' . Let the actual distribution be $(p_1, p_2, \dots, p_m)$. We will determine the probability that a user would end up in category j , no matter what his/her actual category is. Let denote that probability as r_j . If a user is originally in category j (with probability p_j), then his/her probability to stay in category j is $(1 - \gamma) \cdot p_j$. If he/she is originally in another category i (with probability $p_i, i \neq j$), then his/her probability to move to category j is $(\gamma \cdot p_i)/(m - 1)$. Taking sum of these disjoint cases, we get a formula of r_j that is

$$r_j = (1 - \gamma) \cdot p_j + \sum_{\substack{i=1 \\ i \neq j}}^m \left(\frac{\gamma}{m-1} \right) \cdot p_i = (1 - \gamma) \cdot p_j + \left(\frac{\gamma}{m-1} \right) \cdot (1 - p_j) = \frac{\gamma}{m-1} + \left(1 - \frac{\gamma m}{m-1} \right) \cdot p_j \quad (7)$$

Thus distribution of C' is $(r_1, r_2, \dots, r_m)$, and mutual information between C and C' is

$$I(C'; C) = H(C') - H(C'|C) = \sum_{i=1}^m r_j \cdot \log \frac{1}{r_j} - h(\gamma) - \gamma \cdot \log(m-1) \quad (8)$$

where r_j is as defined in (7). This is the measurement of privacy leakage in this mechanism with probability parameter γ .

After obtaining formula (8), we can try to do a computation of it. To simplify the computation, we try on a special case where the actual distribution is uniform, i.e. the users are distributed uniformly into m categories with probability $1/m$ each. In this case we will have :

$$\begin{aligned} r_j &= \frac{\gamma}{m-1} + \left(1 - \frac{\gamma m}{m-1} \right) \cdot p_j = \frac{\gamma}{m-1} + \left(1 - \frac{\gamma m}{m-1} \right) \cdot \frac{1}{m} = \frac{\gamma}{m-1} + \frac{1}{m} - \frac{\gamma}{m-1} = \frac{1}{m} \\ I(C'; C) &= \sum_{i=1}^m r_j \cdot \log \frac{1}{r_j} - h(\gamma) - \gamma \cdot \log(m-1) \\ &= \sum_{i=1}^m \frac{1}{m} \cdot \log m - h(\gamma) - \gamma \cdot \log(m-1) = \log m - h(\gamma) - \gamma \cdot \log(m-1) \end{aligned} \quad (9)$$

Now we compute formula (9) of variable γ . We consider several cases with different number of categories 2, 3, 4 and 5 categories. Results of these computations are shown in Figure 1.

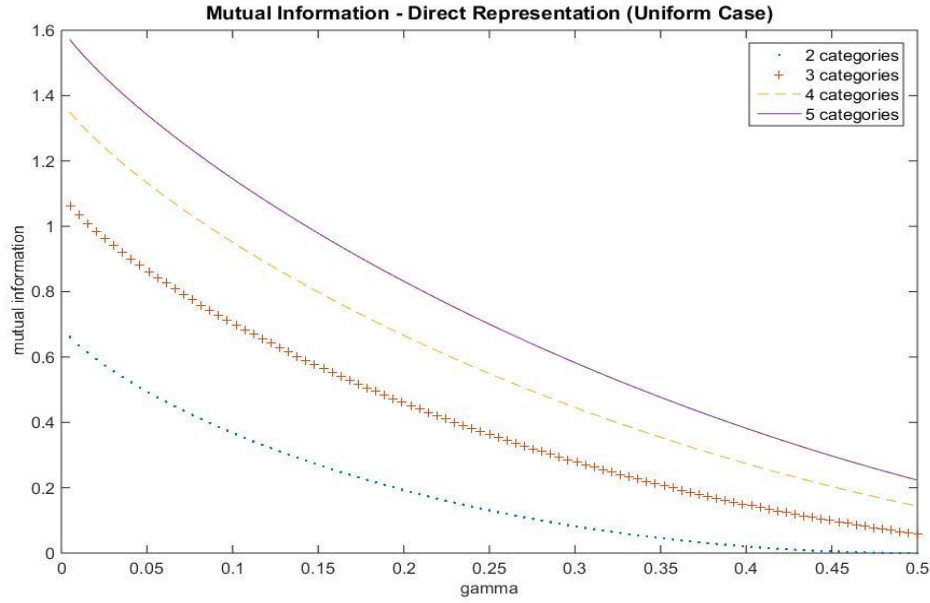


Figure 1. Privacy leakage of RAPPOR with direct representation. Graphs of mutual information $I(C'; C)$ as a function of the noise parameter γ are plotted for $m = 2, 3, 4$ and 5 , where the actual distribution C is uniform.

We can see the behavior of those graph. When γ is closer to 0.5 , the mutual information is closer to 0 and therefore get stronger privacy. We also see that if there are more categories, the value of mutual information is also bigger. However, we have not been able to compare multiple cases with different number of categories. Look at the fact that formula (9) depends on the value of m and if m is bigger, then $I(C'; C)$ shall be bigger too. This leads to a possible kind of “normalization”, which makes the value of mutual information fall in interval $[0,1]$. If $\gamma = 0$, the “normalized” mutual information should be equal to 1 , which means that the aggregator is fully able to learn any information in the database since he/she receives the original one. Unfortunately we are yet to find a formulation about the normalization factor.

4.2 Case II : RAPPOR with unary representation

Now we move on to another case of RAPPOR with unary representation. To simplify the case, we will consider the symmetric case when $\beta_0 = \beta_1$ (to avoid many subscripts, later they are both written as β). Each user can only belong to one category and

therefore his/her binary vector representation C contains exactly one 1 in his/her category's position and 0 otherwise. Given an arbitrary binary m -vector z , we will compute the probability that the binary vector representation will be modified into z . This shall depend on how many 1s are contained in z , i.e. the Hamming weight of z , usually denoted by $w(z)$.

If the 1 in original vector C is not flipped to 0, then from $m - 1$ 0s in C , there are $w(z) - 1$ of them which are flipped into 1 and the other $m - w(z)$ 0s are not flipped. Its probability will be $\beta^{w(z)-1} (1 - \beta)^{m-w(z)+1}$. In other side, if the 1 in C is flipped to 0, then from $m-1$ 0s in C , there are $w(z)$ of them which are flipped into 1 and the other $m - w(z) - 1$ 0s are not flipped. Its probability will be $\beta^{w(z)+1} (1 - \beta)^{m-w(z)-1}$. As a result, the probability distribution Z of a modified database, knowing that a user belongs to category j , can be written as :

$$\Pr[Z = z | C = j] = \left(\frac{1 - \beta}{\beta} \right)^{2z_j} \beta^{w(z)+1} (1 - \beta)^{m-w(z)-1} \quad (10)$$

It looks like a tricky task to compute the entropy $H(Z|C)$ since we have to take sum from any possible binary vectors z with various Hamming weights and positions of their 1s. However, by using some binomial properties in Rosen [6], we can obtain a pretty simple result below :

$$H(Z|C) = m \cdot h(\beta) = m \left(\beta \log \frac{1}{\beta} + (1 - \beta) \log \frac{1}{1 - \beta} \right) \quad (11)$$

Defining $\Pr[Z = z]$, for any m -binary vector, is a lot more difficult. We have to consider any possible original position of the single 1, multiply it by its actual probability and take sum of them. Based on (10) we can compute that probability mass function, denoted by $Q(z)$, as :

$$\begin{aligned} Q(z) &= \Pr[Z = z] = \sum_{j=1}^m p_j \cdot \Pr[Z = z | C = j] \\ &= \sum_{j=1}^m p_j \left(\frac{1 - \beta}{\beta} \right)^{2z_j} \beta^{w(z)+1} (1 - \beta)^{m-w(z)-1} \end{aligned}$$

$$= \left(\frac{\beta}{1-\beta} \right)^{w(z)} \beta (1-\beta)^{m-1} \sum_{j=1}^m p_j \left(\frac{1-\beta}{\beta} \right)^{2z_j} \quad (12)$$

Note that values of z_j are either 0 or 1, so we can divide the last sigma form in (12) into two cases when $z_j = 0$ and when $z_j = 1$. If we define $p(z) = \sum_j p_j \cdot z_j$, then (12) can be simplified into

$$\begin{aligned} Q(z) &= \left(\frac{\beta}{1-\beta} \right)^{w(z)} \beta (1-\beta)^{m-1} \left[\sum_{\substack{j=1 \\ z_j=1}}^m p_j \left(\frac{1-\beta}{\beta} \right)^2 + \sum_{\substack{j=1 \\ z_j=0}}^m p_j \left(\frac{1-\beta}{\beta} \right)^0 \right] \\ &= \left(\frac{\beta}{1-\beta} \right)^{w(z)} \beta (1-\beta)^{m-1} \left[p(z) \left(\frac{1-\beta}{\beta} \right)^2 + 1 - p(z) \right] \\ &= \left(\frac{\beta}{1-\beta} \right)^{w(z)} \beta (1-\beta)^{m-1} \left[1 + p(z) \frac{1-2\beta}{\beta^2} \right] \end{aligned} \quad (13)$$

Since entropy calculation involves logarithm and nothing can be simplified from logarithm of a sum, it will be difficult to simplify the $H(Z)$ term in this case.

Therefore we try to do a “brute force” for calculating mutual information $I(Z; C) = H(Z) - H(Z|C)$ by directly using the last row of (13) in calculating $H(Z)$ term. Again we set the actual distribution of categories to be uniform and we compute multiple cases of different number of categories : 2, 3, 4 and 5 categories. Results of these computations are shown in Figure 2.

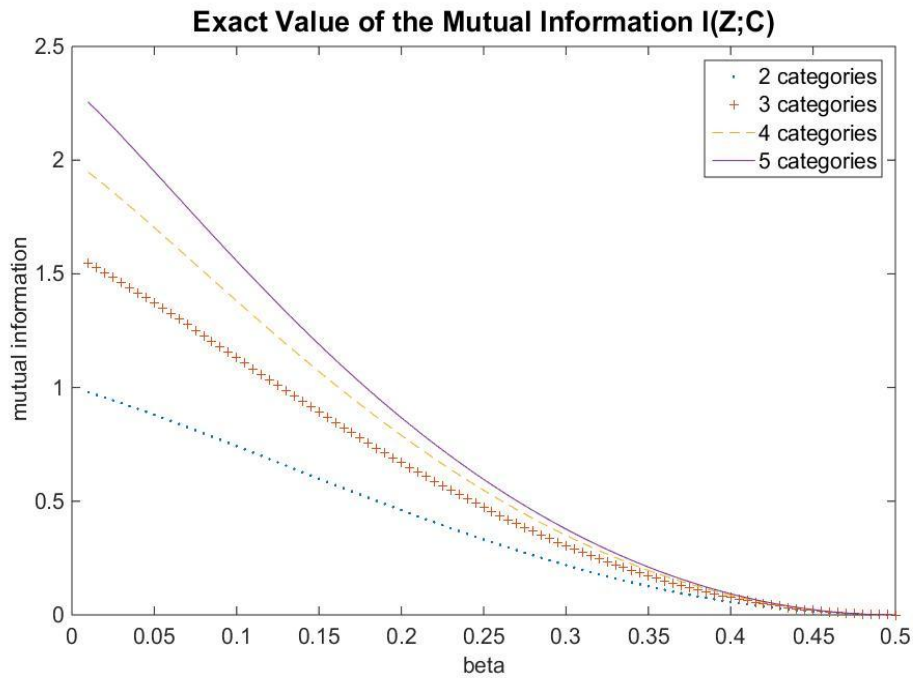


Figure 2. Privacy leakage of RAPPOR with unary representation. Graphs of mutual information $I(Z'; C)$ as a function of the noise parameter β are plotted for $m = 2, 3, 4$ and 5 , where the actual distribution C is uniform.

We see similar behavior with previous case. For any number of categories, their graphs are monotonically decreasing on interval $[0, 0.5]$. The difference is that all graphs tend to 0 when $\beta = 0.5$. We can interpret this as the aggregator is unable to learn anything when $\beta = 0.5$, i.e. the binary flip is totally random. More categories also imply bigger value of mutual information, but they are also yet to be normalized. Also note that if the range of β is extended to $[0, 1]$, then those graph will be monotonically increasing. Let us imagine if $\beta = 1$, then all binary vectors will be completely flipped (0 to 1 or vice versa) and the aggregator can easily determine the original ones. We can also intuitively conclude that cases when $\beta = t$ and $\beta = 1 - t$ are practically similar.

Apart from those two presented cases, we have tried to do computation for other mechanisms, but some of them have very complicated formula and be very difficult to compute. Some computations even need several days to be completed. Computation for a big enough number of categories is also yet to be done. There are two possible

solutions simplifying the computation, or determining upper/lower bound of the privacy leakage which is easier to compute.

5 Conclusions

From what are discussed in this article, we have several points of conclusions and feedbacks for any possible further research, which are :

1. By interpreting entropies as number of binary questions which are need to asked for learning information on a database, it is possible to re-define privacy and accuracy-utility aspects of a differential privacy scheme in term of entropies. In this article the former has been done.
2. Privacy leakage of a differentially private mechanism can be defined as the mutual information between actual distribution of categories of an attribute in a database and its modified version. This definition does make sense and some simulations with MATLAB have been done to justify it.
3. Defining accuracy and utility of a differentially private mechanism in term of entropies is a trickier task to do. One entropy in the definition of utility should have similar and comparable interpretation with one entropy in privacy leakage. This may still be very open problem.
4. Definitions of privacy leakage here is still lack of “normalization”. To make it totally comparable between any attributes with various number of categories, we might need to normalize them into a specific range (possibly $[0,1]$) and their normalization factors are yet to be determined. These could be some open problems to solve in further research.

Acknowledgements

This is a part of our work in the “Bridging the Gap between Theory and Practice in Data Privacy” (BRIDGE) project at Technische Universiteit Eindhoven in 2017, funded by the Netherlands Organization for Scientific Research (NWO) and National Science Foundation (NSF). This project should lead to a PhD degree in which initially the first

author was the PhD candidate. However, after several months the first author decided not to continue working on this project.

References

- [1] C. Dwork, F. McSherry, K. Nissim, and A. Smith, “Calibrating noise to sensitivity in private data analysis,” *Theory of Cryptography Conference*, New York, USA, 265–284, 2006.
- [2] C. Dwork and A. Roth, “The algorithmic foundations of differential privacy,” *Foundations and Trends in Theoretical Computer Science*, **9** (3–4), 211–407, 2014.
- [3] T. Wang, J. Blocki, N. Li, and S. Jha, “Locally differentially private protocols for frequency estimation,” *USENIX Security Symposium, Vancouver, British Columbia, Canada*, 729–745, 2017.
- [4] W. Wang, L. Ying, and J. Zhang, “On the relation between identifiability, differential privacy and mutual information privacy,” *IEEE Transactions on Information Theory*, **62** (9), 5018–5029, 2016.
- [5] T. M. Cover and J. A. Thomas, *Elements of Information Theory*, Second Ed, John Wiley & Sons Publication, Hoboken, USA, 2006.
- [6] K. H. Rosen, *Discrete Mathematics and Its Applications*, Seventh Ed, Mc Graw-Hill Education, New York, USA, 2011.