

Lecture:

Heuristic notions of privacy and an introduction to DP

In the last two lectures, we saw examples of reconstruction attacks that allow for (sometimes efficient) reconstruction of sensitive attributes from queries of bounded error. If the goal, however, was simply the "hiding" of the sensitive attribute of a single individual, whose quasi-identifiers are known, then there are simple mechanisms that provide "anonymity" in a heuristic sense.

1. K-Anonymity [Sweeney (2002) (see also Sweeney (1997))]:

A dataset $\mathcal{Q} = \{(q_i, s_i) : i \in [N]\}$ respects K-anonymity, if for any

$\uparrow$ sens. attribute
 $\downarrow$ quasi-identifiers

given $q \in \mathcal{Q}$, we have

$$\#\{i : q_i = q\} = 0 \text{ or } \#\{i : q_i = q\} \geq K.$$

Clearly, a dataset by itself may not respect K-anonymity, and one typically carries out "generalization" (or binning) or "suppression" (or replacing with '*') of attributes to give rise to "equivalence classes" that respect K-anonymity. generalized quasi-identifiers.

Example: Setting of hospital records.

Quasi-identifiers		Sensitive attribute
Gender	Postal Code	Disease
Male	600020	Heart disease
Male	600018	Lung infection
Female	600036	Osteoporosis
Female	600036	Cervical cancer
Female	600036	Osteoporosis

Upon generalization, this yields a dataset $\bar{D}$ that respects 2-anonymity:

(Gender, Postal code)		Disease
(Male, 600018) or (Male, 600020)		Heart disease Lung infection
(Female, 600036)		Osteoporosis Osteoporosis Cervical cancer

Eg. class (with arrows pointing to the rows)

However, K-anonymity has an obvious pitfall: What if all the sensitive attributes in a given eg. class are equal? Then, it is easy for an adversary to identify the sensitive attribute of a user, given knowledge of his/her quasi-identifiers.

To ameliorate this issue:

2. l-Diversity [Machanavajjhala, Kifer, Gehrke, Venkatasubramanian (2007)]:

A dataset $\mathcal{D} = \{(q_i, s_i) : i \in [N]\}$ respects l -diversity if each quasi-identifier has, associated with it, at least l **distinct** sensitive attributes.

HW: Does the dataset $\overline{\mathcal{D}}$ above respect l -diversity? If so, what is the largest such l ?

However, in very recent work [Rameshwar, Tandon (2025)], we argued that in the "worst-case", "composing" together / "linking" datasets could rapidly degrade anonymity in the sense of l -diversity.

... see also "t-closeness"

[Li, Li, Venkatasubramanian, "t-Closeness: Privacy beyond k-anonymity and l-diversity," (2007)]

However, we are interested in rigorous notions of privacy, which also "compose" gracefully, and are backed by formal results in probability/information theory.

Differential Privacy

The chief idea in defining a formal notion of privacy is to address the problem of tracing: an adversary must not be able to distinguish b/w whether a dataset contained/did not contain a sample from a given user, purely based on the outputs of a privacy-preserving mechanism.

Def 1: Datasets $\mathcal{D}, \mathcal{D}'$ with the same collection of users $[N]$ are said to be neighbouring if $\mathcal{D} = (x_1, x_2, \dots, x_N)$ and $\mathcal{D}' = (x'_1, x'_2, \dots, x'_N)$, with $x_j \neq x'_j$, for some $j \in [N]$ and $x_i = x'_i$, for all $i \neq j$.

Def 2 (DP): A randomized algorithm (or mechanism) $A: \mathcal{D}_n \rightarrow \mathcal{Y}$ is ϵ -differentially private (or ϵ -DP), for $\epsilon > 0$, if, for every pair $\underline{x}, \underline{x}'$ of neighbouring datasets, and every $S \subseteq \mathcal{Y}$, we have

$$\Pr[A(\underline{x}) \in S] \leq e^\epsilon \Pr[A(\underline{x}') \in S].$$

Remarks: (i) Note that, by the definition above, we must have, $\forall \underline{x}, \underline{x}'$ neighbouring & all $S \subseteq \mathcal{Y}$,

$$e^{-\epsilon} \Pr[A(\underline{x}') \in S] \leq \Pr[A(\underline{x}) \in S] \leq e^\epsilon \Pr[A(\underline{x}') \in S].$$

(ii) The definition above is inspired by the structure of the tracing attack we saw last class: we want log-likelihoods to be "small".

(iii) How does one set ϵ ? Typically, one would like ϵ to be "small" (≈ 0.1 or ≈ 1), but nothing stops us from setting $\epsilon \gg 1$. Clearly, in the latter case, the privacy guarantees (on small datasets) are unlikely to be meaningful.

[HW: Look up the values of ϵ used in the 2020 US Census releases, the 2024 Israel release of live birth statistics, and in real-world deployments by big tech firms. What do you think of these values?]

Achieving DP: A First Pass on Binary-Valued Data

Suppose that $D = (x_1, \dots, x_N)$ is such that $x_i \in \{0, 1\}$, $\forall i \in [N]$. A simple way of "privatizing" this dataset is to randomize each bit so that it is biased towards the true value.

This leads to RR [Randomized Response, Warner (1965)]:

→ For each $i \in [N]$, set $y_i = \begin{cases} x_i, & \text{w.p. } 1-p, \\ 1-x_i, & \text{w.p. } p \end{cases}$ for $p \in (0, 1/2)$.

→ Return $(y_1, \dots, y_N)$.

HW: ① Argue that $\sum_{i=1}^N y_i$ is a "good estimate" of the statistic $\sum_{i=1}^N x_i$, i.e.,

$$\left(\mathbb{E} \left[\left(\sum y_i - \sum x_i \right)^2 \right] \right)^{1/2} = O(\sqrt{n}).$$

② For what value of ϵ does $RR(p)$ satisfy ϵ -DP?

③* Get rid of the "bias" term in ①, by constructing a new estimator that is an unbiased estimate of $\sum x_i$. Can you use a Chernoff bound to show that this estimate is "close" to $\sum x_i$, with high prob.?
(at most α away)