

Lecture 23: Lower Bounds for Interior Point

Instructor: Gavin Brown

Scribe: Gavin Brown

Disclaimer: This document is intended as an informal supplement to in-class note-taking. It has not been given the level of scrutiny expected in polished lecture notes, let alone that reserved for peer-reviewed publications.

Today we will see another approach for lower bounds in DP learning. Unlike packing and fingerprinting, what we'll see today is heavily tailored to a particular task: returning any value that's in the "interior" of a dataset. We'll prove (most of) a surprising lower bound for this task. Although we won't go much beyond this, today's material provides a good starting point for understanding DP PAC learning.

1 Setup and Theorem

Let X be a totally ordered domain, e.g., $X \subseteq [0, 1]$ or $X = \{0, 1, 2, \dots, |X| - 1\}$. Usually we've used the notation $[k] = \{1, 2, \dots, k\}$, but today we'll write $[k] = \{0, 1, \dots, k - 1\}$. I'll try to include reminders.

Definition 1.1. An algorithm $A : X^n \rightarrow X$ solves the *interior point problem* on X with error probability β if for all $X \in X^n$

$$\Pr \left[\min_{i \in [n]} x_i \leq A(X) \leq \max_{i \in [n]} x_i \right] \geq 1 - \beta.$$

This is a very simple task (do not confuse it with "interior point" methods in optimization). Without privacy, returning any data point works. The median is a robust solution. Note that Definition 1.1 does not require $A(X)$ to be a member of the dataset. Under DP, we of course do not expect this to be the case.

Bun, Nissim, Stemmer, and Vadhan [2015] proved the following theorem.

Theorem 1.2 (BNSV). Fix $\varepsilon \in (0, 1/4)$ and $\delta(n) \leq 1/(50n^2)$. For every $n \in \mathbb{N}$, solving the interior point problem on X with error less than $1/4$ and $(\varepsilon, \delta(n))$ -DP requires $n = \Omega(\log^*(|X|))$.

Here $\log^*$ is the *iterated logarithm*, the minimum integer k such that

$$\underbrace{\log(\log(\dots \log(z)))}_{k \text{ times}} \leq 1.$$

This is an extremely slow-growing function: for $z \in (16, 65536]$, $\log^*(z) = 4$, and for $z \in (65536, 2^{65536}]$, $\log^*(z) = 5$.

Before discussing the proof, let's discuss the theorem some more.

- On the one hand, it is extremely weak: if you are willing to believe all of our data points can be represented using a computer the size of the known universe (which contains perhaps 2^{266} atoms), you don't need many examples!
- On the other hand, it is extremely strong: without additional assumptions (about, e.g., an underlying data distribution), many tasks for real-valued statistics are *impossible* under DP.

- The lower bound has quite a different flavor from the other techniques we've seen. It's an interesting direction to understand how these approaches relate to each other.
- There is an interesting line of work on upper bounds for this and related problems.
- BNSV studied the interior point problem in the context of private *PAC learning*. PAC learning is the central formal setting in learning theory. This lower bound sits in a larger body of literature on DP-PAC learning, the highlight of which is the equivalence between (i) DP-PAC learnability and (ii) online learnability. We won't cover this in any great detail, but it's one of the most notable results in learning theory in the last several years.

Quantitatively, the theorem requires $\delta(n) \leq 1/(50n^2)$. Different proof techniques (as covered in Mark Bun's dissertation) apply to larger δ , all the way up to the optimal $\delta \approx 1/n$. Why is this optimal? Consider the following $(0, \delta)$ -DP algorithm, which is similar to the $\text{LeakyRR}_{\epsilon, \delta}$ mechanism from Lecture 11.

Algorithm 1 Name-and-Shame

Input: Data $x_1, \dots, x_n$; $\delta > 0$

```

1:  $S \leftarrow \emptyset$ 
2: for  $i = 1$  to  $n$  do
3:   with probability  $\delta$ , add  $x_i$  to  $S$ 
4: end for
5: return  $S$ 

```

If Algorithm 1 returns a non-empty S then we have an interior point. Since each $x_i \in S$ independently with probability δ , we have

$$\Pr[S \neq \emptyset] = 1 - (1 - \delta)^n \approx \delta n,$$

when δn is not too large. In particular, if $\delta \approx 1/n$ we can expect to succeed with constant probability.

2 Main Lemma and Plan for Proof

Our presentation will closely follow that of BNSV, except we will draw some pictures and skip some steps. Our main lemma will require some notation. Let

$$P_n := \frac{e^\epsilon}{e^\epsilon + 1} + (e^\epsilon + 1) \sum_{j=1}^n \delta(j) \leq 1 - \beta$$

$$b(n) := 1/\delta(n).$$

We'll pretend that $b(n)$ is an integer: it will be the numerical base we use to represent data points later. Note that we snuck in an assumption about the parameters, by asking that $P_n \leq 1 - \beta$. This is where the restriction on $\delta(n)$ comes in, in particular if $P_n \leq 1 - \beta$ for all n then we certainly need $\sum_j \delta(j)$ to converge.

We also define, recursively,

$$S(1) = 2 \quad \text{and} \quad S(n+1) = b(n)^{S(n)}.$$

Here is the main lemma we will prove.

Lemma 2.1. *For every $n \in \mathbb{N}$, there exists a distribution D_n over datasets in $[S(n)]^n = \{0, 1, \dots, S(n) - 1\}^n$ such that for every $(\varepsilon, \delta(n))$ -DP algorithm A ,*

$$\Pr_{X \sim D_n} [\min X \leq A(X) \leq \max X] \leq P_n.$$

The distribution we construct will *not* be a product distribution. The samples will be correlated in very specific way.

Our proof will use induction and have three main steps. First is the **(i) base case**; the proof for D_1 is simple. For induction, we assume the interior point problem is hard for datasets from D_n and prove that implies D_{n+1} is hard. We have **(ii) construction**: given D_n , how do we define D_{n+1} ? Finally, we have the **(iii) reduction**: we have to show that D_{n+1} is hard.

A bit more before we dive in. When we draw a dataset from D_n , we get n items, each of which is an integer in $[S(n)] = \{0, 1, \dots, S(n) - 1\}$, where $S(n)$ is defined above. So for $n = 1$, D_1 is a distribution over single-item datasets, the data point lives in $[S(1)] = [2] = \{0, 1\}$.

Now, D_{n+1} gives $n + 1$ items, each in $[S(n + 1)]$, where $S(n + 1) = b(n)^{S(n)}$ by definition. So $n + 1$ points

$$(y_0, y_1, \dots, y_n) \sim D_{n+1},$$

and you should think of items in this dataset as having base $b(n)$:

$$y_i = \underbrace{\boxed{y_{i,0}} \mid \boxed{y_{i,1}} \mid \boxed{y_{i,2}} \mid \dots \mid \boxed{y_{i,S(n)-1}}}_{S(n) \text{ digits}}.$$

Thus the *length* of data points drawn from D_{n+1} , when written in base $b(n)$, is equal to the *domain size* of data points from D_n . This means the domain size gets exponentiated every time we move from n to $n + 1$. This tower of exponents is exactly the inverse of the theorem's iterated logarithm.

3 Proving the Main Lemma

3.1 Base Case

Take $D_1 = \text{Uniform}(\{0, 1\})$. Let A be $(\varepsilon, \delta(1))$ -DP. A gets a single bit x and solves the interior point problem if $A(x) = x$. We need to show that

$$\Pr_{x \sim D_1} [A(x) = x] \leq P_1 = \frac{e^\varepsilon}{e^\varepsilon + 1} + (e^\varepsilon + 1)\delta(1).$$

This is easy. Let $\neg x$ denote the negation of $x \in \{0, 1\}$. We apply the DP definition:

$$\begin{aligned} \Pr_{x \sim D_1} [A(x) = x] &= \frac{1}{2} \sum_{x \in \{0,1\}} \Pr[A(x) = x] \\ &\leq \frac{1}{2} \sum_{x \in \{0,1\}} e^\varepsilon \Pr[A(\neg x) = x] + \delta(1) \\ &= \frac{1}{2} \sum_{x \in \{0,1\}} e^\varepsilon (1 - \Pr[A(\neg x) = \neg x]) + \delta(1). \end{aligned}$$

Since $\sum_{x \in \{0,1\}} \Pr[A(\neg x) = \neg x] = \sum_{x \in \{0,1\}}$, we have

$$\Pr_{x \sim D_1} [A(x) = x] \leq e^\varepsilon - e^\varepsilon \Pr[A(x) = x] + \delta(1).$$

Rearranging, we are done:

$$\Pr_{x \sim D_1} [A(x) = x] \leq \frac{e^\varepsilon + \delta(1)}{e^\varepsilon + 1} \leq \frac{e^\varepsilon}{e^\varepsilon + 1} + (e^\varepsilon + 1)\delta(1) = P_1.$$

3.2 Induction Step: The Construction

Given a sample $(x_1, \dots, x_n) \sim D_n$, we construct a sample from D_{n+1} as follows. First, draw $y_0 \sim \text{Uniform}([S(n+1)])$, i.e.

$$y_0 = \boxed{y_{0,0}} \boxed{y_{0,1}} \boxed{y_{0,2}} \cdots \boxed{y_{0,S(n)-1}},$$

where each block (or base $b(n)$ “digit”) is drawn uniformly and independently. Then for each $i = 1, \dots, n$ independently, we draw

$$y_i = \underbrace{\boxed{y_{0,0}} \boxed{y_{0,1}} \cdots \boxed{y_{0,x_i}}}_{\text{match } y_0 \text{'s digits}} \underbrace{\boxed{y_{i,x_i+1}} \cdots \boxed{y_{i,S(n)-1}}}_{\text{uniform and indep.}}$$

So y_i matches y_0 on its x_i most significant base- $b(n)$ digits and is independent on the remainder. If we relabel the points so they are ordered, $x_1 \leq x_2 \leq \dots \leq x_n$, then the overall picture is

$$\begin{array}{l} y_0 = \boxed{y_{0,0}} \boxed{y_{0,1}} \boxed{y_{0,2}} \boxed{y_{0,3}} \boxed{y_{0,4}} \boxed{y_{0,5}} \boxed{y_{0,6}} \\ y_1 = \boxed{y_{0,0}} \boxed{y_{0,1}} \boxed{y_{1,2}} \boxed{y_{1,3}} \boxed{y_{1,4}} \boxed{y_{1,5}} \boxed{y_{1,6}} \\ \quad \underbrace{\hspace{1.5cm}}_{x_1 \text{ digits}} \\ y_2 = \boxed{y_{0,0}} \boxed{y_{0,1}} \boxed{y_{0,2}} \boxed{y_{2,3}} \boxed{y_{2,4}} \boxed{y_{2,5}} \boxed{y_{2,6}} \\ \quad \underbrace{\hspace{1.5cm}}_{x_2 \text{ digits}} \\ \vdots \\ y_n = \boxed{y_{0,0}} \boxed{y_{0,1}} \boxed{y_{0,2}} \boxed{y_{0,3}} \boxed{y_{0,4}} \boxed{y_{0,5}} \boxed{y_{n,6}} \\ \quad \underbrace{\hspace{1.5cm}}_{x_n \text{ digits}} \end{array}$$

and all the black blocks match across their columns, while the blue blocks are all independent and uniform. Remember that each block here is a number in $[S(n)] = \{0, 1, \dots, S(n) - 1\}$.

3.3 Induction Step: The Reduction

Now for the meat of the proof. Assume that any $(\varepsilon, \delta(n))$ -DP algorithm, when run on inputs from D_n , has success probability at most P_n . Now consider an algorithm A which is $(\varepsilon, \delta(n+1))$ -DP and operates on datasets in $[S(n+1)]^n$. We will show it cannot solve the interior point problem on samples from D_{n+1} with probability greater than P_{n+1} . To do this, we will use A to construct an algorithm $\hat{A}$ which works for data on D_n .

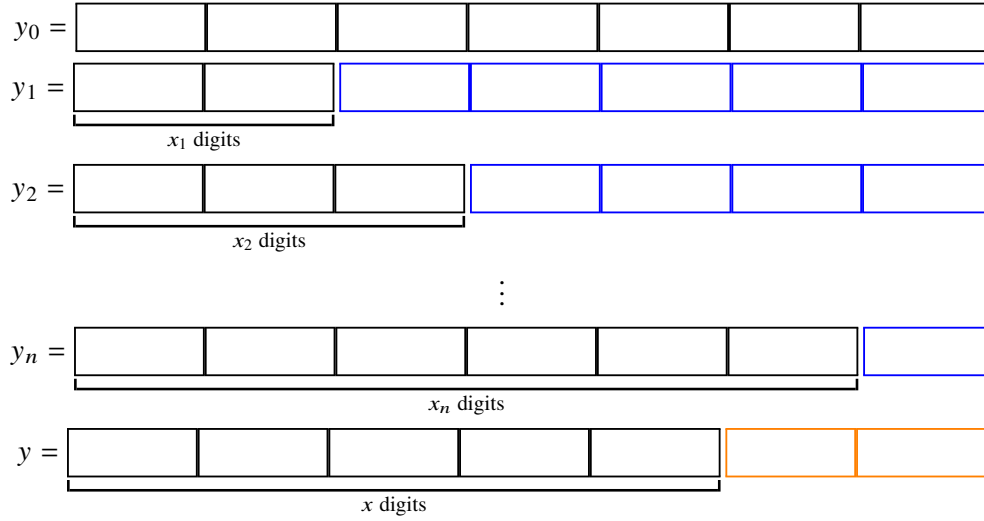
Algorithm 2 $\hat{A}$

Input: Data $X = (x_1, \dots, x_n) \in [S(n)]^n$

- 1: Sample $Y = (y_0, y_1, \dots, y_n) \in [S(n+1)]^{n+1}$ as in Section 3.2, starting from $(x_1, \dots, x_n)$
 - 2: Let $y \leftarrow A(Y)$
 - 3: Let $x \in [S(n)]$ be the length of the longest prefix of y (in base $b(n)$) which matches y_0
 - 4: **return** x
-

Exercise 3.1. Prove that if A is $(\varepsilon, \delta(n+1))$ -DP then $\hat{A}$ is $(\varepsilon, \delta(n))$ -DP.

We will show that, when the call to A in Algorithm 2 returns an interior point of $(y_0, y_1, \dots, y_n)$, then almost always the x returned by $\hat{A}$ is an interior point for $(x_1, \dots, x_n)$. Here's a figure of such an outcome.

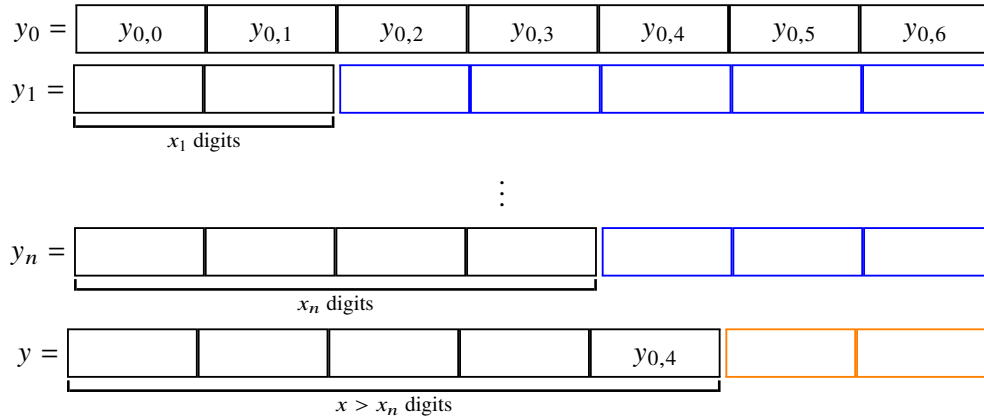


Black blocks match with y_0 , blue blocks are generated randomly, and orange blocks are the output of A which does not match y_0 . In general, y will match y_0 on some number of blocks, it could be as few as 0 or as many as $S(n)$.

The first part of the argument is not mathematically hard, but it may take a bit of thought.

Exercise 3.2. Prove that if y is an interior point of $(y_0, y_1, \dots, y_n)$, then $\min\{x_i\} \leq x$.

The second part of the argument uses privacy: if $x \geq \max\{x_i\}$, then A has successfully guess an integer which appeared only once in the data. Consider:



It seems A 's output y contains the number $y_{0,4} \in [b(n)]$. But this can't be very likely, since $b(n)$ is large and A 's output cannot depend too heavily on any one input. We formalize this.

Lemma 3.3. *We have $x > \max x_i$ with probability at most $\frac{e^\varepsilon}{b(n)} + \delta(n+1)$.*

Proof. Fix a dataset $X = (x_1, \dots, x_n) \in [S(n)]^n$; let $w = \max x_i$. Fix also $\hat{A}$'s choice of $Y \in [S(n+1)]^{n+1}$, except for the entry $y_{0,w+1} \in [b(n)]$. As discussed above, if $x > \max x_i$ then $y_{w+1} = y_{0,w+1}$. We will show this unlikely to happen (regardless of the fixed values above).

For $z \in [b(n)]$ let Y_z denote Y with digit $y_{0,w+1} \leftarrow z$. Taking randomness over the coins of A and the uniform choice of $Z \in_R [b(n)]$, we have

$$\begin{aligned} \Pr[A(Y_Z)_{w+1} = Z] &= \frac{1}{b(n)} \sum_{z \in [b(n)]} \Pr[A(Y_z)_{w+1} = z] \\ &= \frac{1}{b(n)^2} \sum_{\substack{z \in [b(n)] \\ z' \in [b(n)]}} \Pr[A(Y_z)_{w+1} = z] \\ &\leq \frac{1}{b(n)^2} \sum_{\substack{z \in [b(n)] \\ z' \in [b(n)]}} e^\varepsilon \Pr[A(Y_{z'})_{w+1} = z] + \delta(n+1). \end{aligned}$$

Note that the second equality is essentially a no-op: we merely introduced another variable. But now we can write Z, Z' as independent random variables:

$$\Pr[A(Y_Z)_{w+1} = Z] \leq e^\varepsilon \Pr[A(Y_{Z'})_{w+1} = Z] + \delta(n+1).$$

No matter the value of $A(Y_{Z'})_{w+1}$, the probability Z matches it is $\frac{1}{b(n)}$. This completes the proof. $\square$

Combining the exercise and lemma, we have

$$\Pr_{X \sim D_n} [\min X \leq \hat{A}(X) \leq \max X] \geq \Pr_{X \sim D_n} [\min X \leq \hat{A}(X) \leq \max X] - \frac{e^\varepsilon}{e^\varepsilon + 1} - \delta(n+1).$$

Since our inductive hypothesis says

$$\Pr_{X \sim D_n} [\min X \leq \hat{A}(X) \leq \max X] \leq P_n,$$

we have

$$\begin{aligned} \Pr_{X \sim D_n} [\min X \leq \hat{A}(X) \leq \max X] &\leq P_n + \frac{e^\varepsilon}{e^\varepsilon + 1} + \delta(n+1) \\ &\leq P_{n+1}, \end{aligned}$$

as promised.

4 Conclusions and Further Reading

What have we proved? For each integer n , we constructed a distribution over datasets in $[S(n)]^n$ such that any algorithm has poor performance on it. In other words, if you want to succeed with a data domain of size $S(n)$ or larger, you'll need more than n samples. Recall that $S(n)$ grows extremely quickly with n :

$$S(1) = 2, \quad S(n+1) = (1/\delta(n))^{S(n)}.$$

With a little algebra, if we set $|\mathcal{X}| = S(n)$ we'll get our lower bound of $n \gtrsim \log^*(|\mathcal{X}|)$ in the main theorem.

This first proof appeared in Bun et al. [2015], but Mark Bun's PhD dissertation contains more discussion and multiple proofs for the hardness of DP interior point [Bun, 2016]. One of the proofs uses Ramsey theory in a manner similar to its use in the DP-PAC/online learnability equivalence [Alon et al., 2022]. I'll also mention that the recursive lower bound construction parallels the development of upper bounds for the DP interior point problem. We'll talk a bit more about how these can be used in the next couple lectures.

References

Noga Alon, Mark Bun, Roi Livni, Maryanthe Malliaris, and Shay Moran. Private and online learnability are equivalent. *ACM Journal of the ACM (JACM)*, 69(4):1–34, 2022.

Mark Bun, Kobbi Nissim, Uri Stemmer, and Salil Vadhan. Differentially private release and learning of threshold functions. In *2015 IEEE 56th Annual Symposium on Foundations of Computer Science*, pages 634–649. IEEE, 2015.

Mark Mar Bun. *New Separations in the Complexity of Differential Privacy*. PhD thesis, 2016.