

An brief tour of Differential Privacy

Your guide:

Avrim Blum

Computer Science Dept

Itinerary

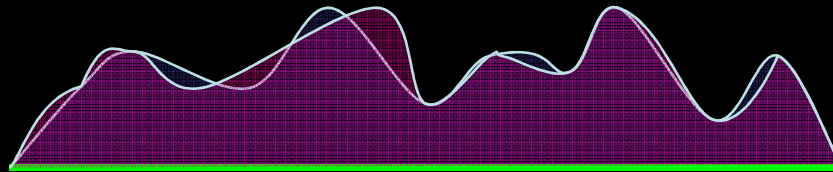
- Stop 1: A motivating example. Why seemingly similar notions from crypto aren't sufficient.
- Stop 2: Definition of differential privacy and a basic mechanism for preserving it.
- Stop 3: Privacy/utility tradeoffs: ask a silly (sensitive) question, get a silly answer.
- Stop 4: Other kinds of mechanisms, releasing sanitized databases, more privacy/utility tradeoffs, and discussion.

A preliminary story

- A classic cool result from theoretical crypto:
 - Say you want to figure out the average salary of people in the room, without revealing anything about your own salary other than what is inherent in the answer.
- Turns out you can actually do this. In fact, any function at all. "secure multiparty computation".
 - It's really cool. Want to try?
- Anyone have to go to the bathroom?
 - What happens if we do it again?

Differential Privacy [Dwork et al.]

- “Lets you go to the bathroom in peace”
 - What we want is a protocol that has a probability distribution over outputs



such that if person i changed their input from x_i to any other allowed x_i' , the relative probabilities of any output do not change by much.

- So, for instance, can pretend your input was any other allowed value you want.
- Can view as model of “plausible deniability”.
 - Even if no bad intent, who knows what prior info people have?

Differential Privacy: Definition

It's a property of a protocol A which you run on some dataset X producing some output $A(X)$.

- A is ϵ -differentially private if for any two neighbor datasets X, X' (differ in just one element $x_i \rightarrow x'_i$),



for all outcomes v ,

$$e^{-\epsilon} \leq \Pr(A(X)=v)/\Pr(A(X')=v) \leq e^{\epsilon}$$

$\approx 1-\epsilon$

probability over
randomness in A

$\approx 1+\epsilon$

Differential Privacy: Definition

It's a property of a protocol A which you run on some dataset X producing some output $A(X)$.

- A is ϵ -differentially private if for any two neighbor datasets X, X' (differ in just one element $x_i \rightarrow x_i'$),

View as model of plausible deniability

(pretend after the fact that my input was really x_i')

for all outcomes v ,

$$e^{-\epsilon} \leq \Pr(A(X)=v)/\Pr(A(X')=v) \leq e^{\epsilon}$$

$\approx 1-\epsilon$

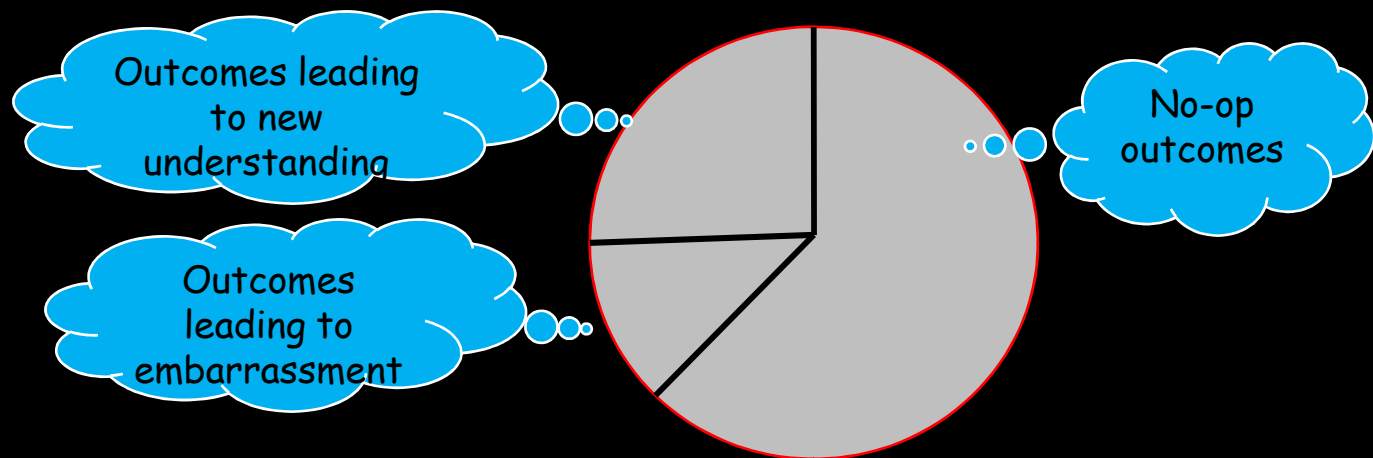
probability over
randomness in A

$\approx 1+\epsilon$

Differential Privacy: Definition

It's a property of a protocol A which you run on some dataset X producing some output $A(X)$.

- A is ϵ -differentially private if for any two neighbor datasets X, X' (differ in just one element $x_i \rightarrow x_i'$),



for all outcomes v ,

$$e^{-\epsilon} \leq \Pr(A(X)=v)/\Pr(A(X')=v) \leq e^{\epsilon}$$

$\approx 1-\epsilon$

probability over
randomness in A

$\approx 1+\epsilon$

Differential Privacy: Definition

It's a property of a protocol A which you run on some dataset X producing some output $A(X)$.

- A is ϵ -differentially private if for any two neighbor datasets X, X' (differ in just one element $x_i \rightarrow x_i'$),

What if you participate in two protocols A and B ?

$$e^{-2\epsilon} \leq \Pr(A(X)=v \ \& \ B(X)=w) / \Pr(A(X')=v \ \& \ B(X')=w) \leq e^{2\epsilon}$$

for all outcomes v ,

$$e^{-\epsilon} \leq \Pr(A(X)=v) / \Pr(A(X')=v) \leq e^{\epsilon}$$

$\approx 1-\epsilon$

probability over
randomness in A

$\approx 1+\epsilon$

So,
combination
is 2ϵ -DP.

Differential Privacy: Definition

It's a property of a protocol A which you run on some dataset X producing some output $A(X)$.

- A is ϵ -differentially private if for any two neighbor datasets X, X' (differ in just one element $x_i \rightarrow x_i'$),

OK, great. How can we achieve it? What kind of ϵ can we get with reasonable utility?

Silly algorithm: $A(X)=0$ no matter what. Or $A(X)=\text{unif}[0,b]$

for all outcomes v ,

$$e^{-\epsilon} \leq \Pr(A(X)=v)/\Pr(A(X')=v) \leq e^{\epsilon}$$

$\approx 1-\epsilon$

probability over
randomness in A

$\approx 1+\epsilon$

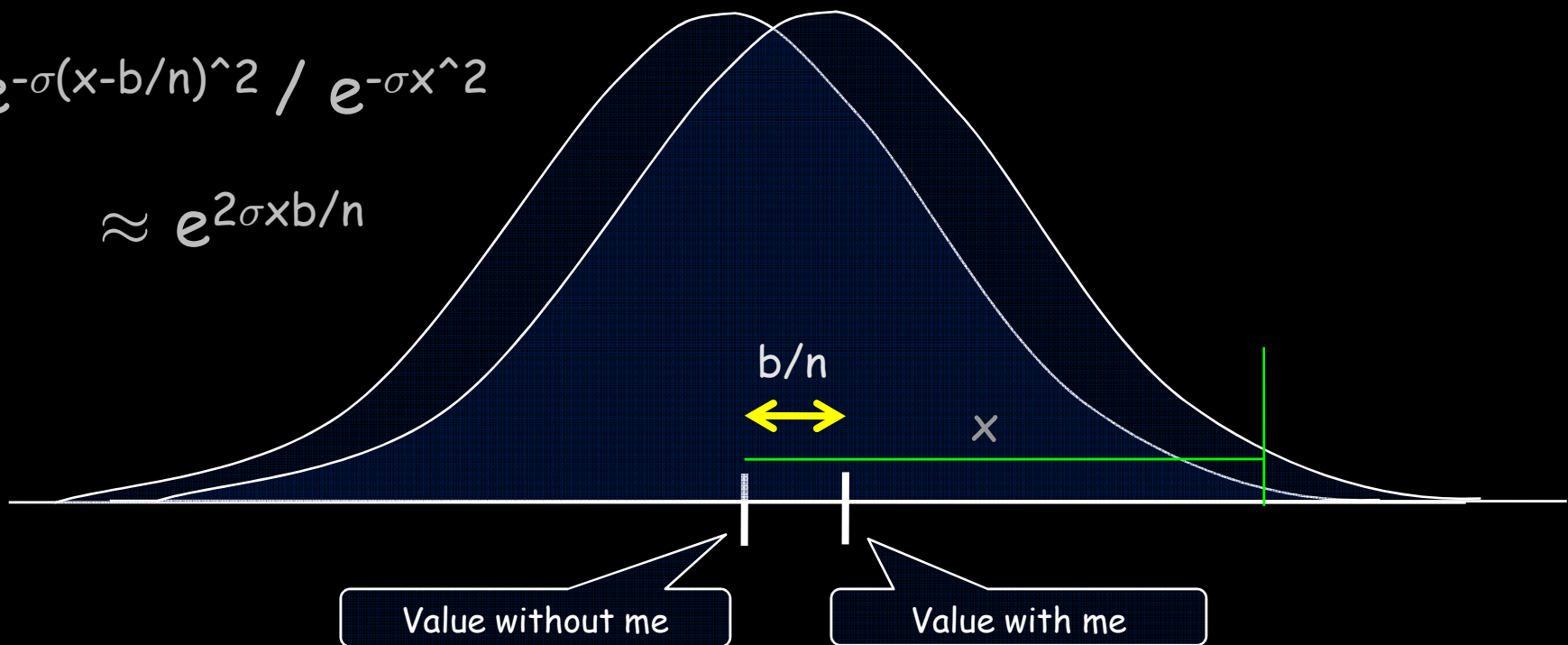
Differential Privacy via output perturbation

Say have n inputs in range $[0, b]$. Want to release average while preserving privacy.

- Natural idea: take output and perturb with noise.
- First thought: add Gaussian noise.

$$e^{-\sigma(x-b/n)^2} / e^{-\sigma x^2}$$

$$\approx e^{2\sigma x b/n}$$



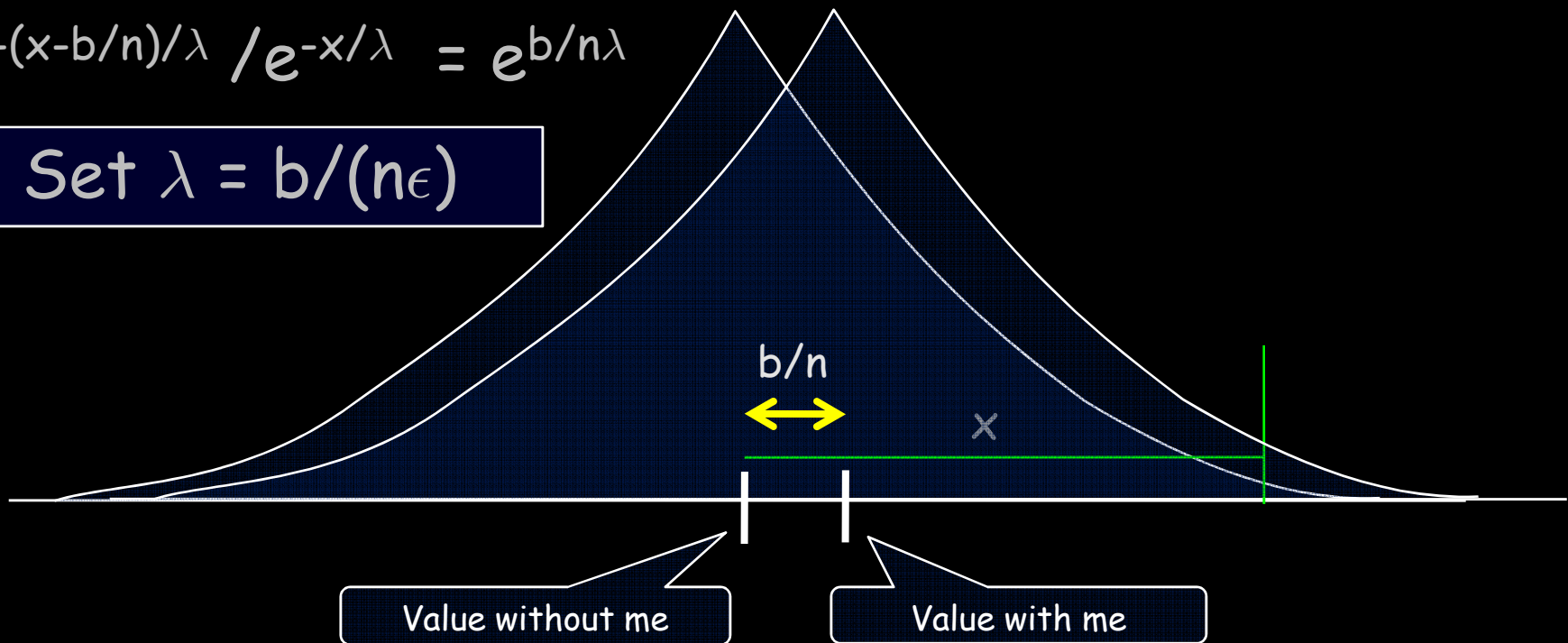
Differential Privacy via output perturbation

Say have n inputs in range $[0, b]$. Want to release average while preserving privacy.

- Natural idea: take output and perturb with noise.
- Better: Laplace (or geometric) distrib $p(x) \propto e^{-|x|/\lambda}$

$$e^{-(x-b/n)/\lambda} / e^{-x/\lambda} = e^{b/n\lambda}$$

Set $\lambda = b/(n\epsilon)$

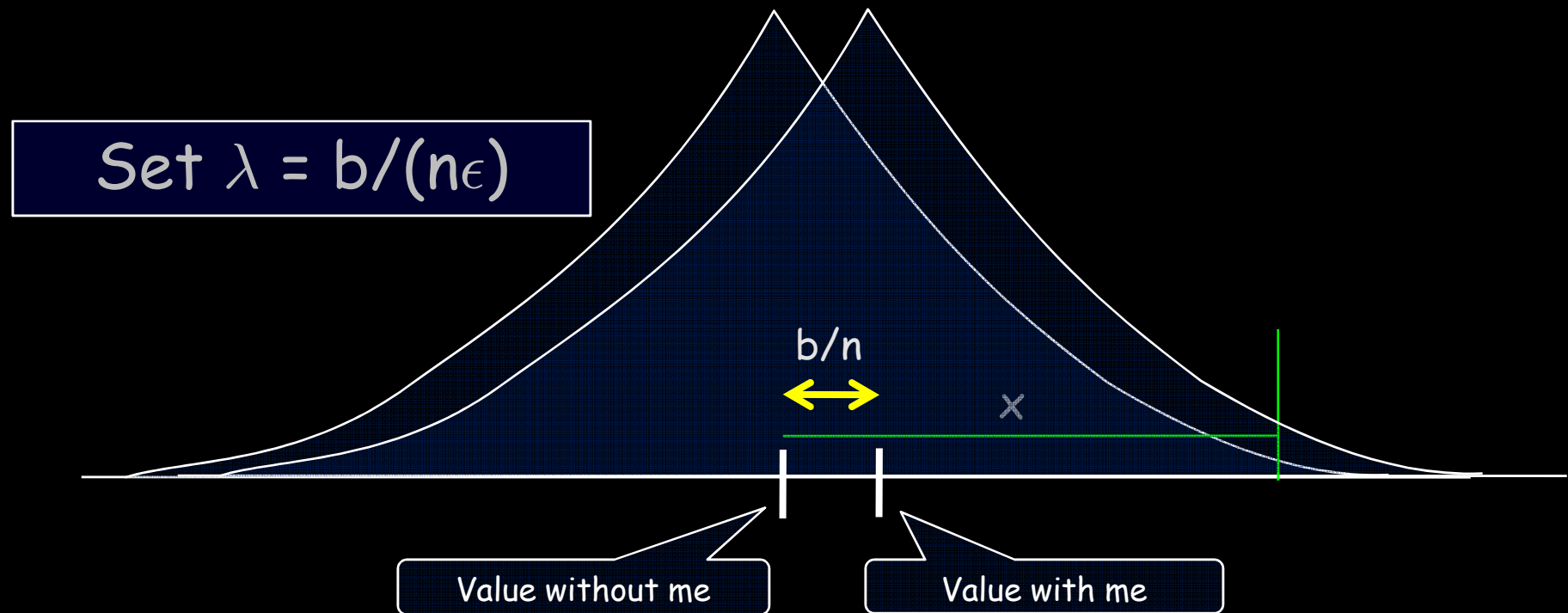


"Laplace mechanism"

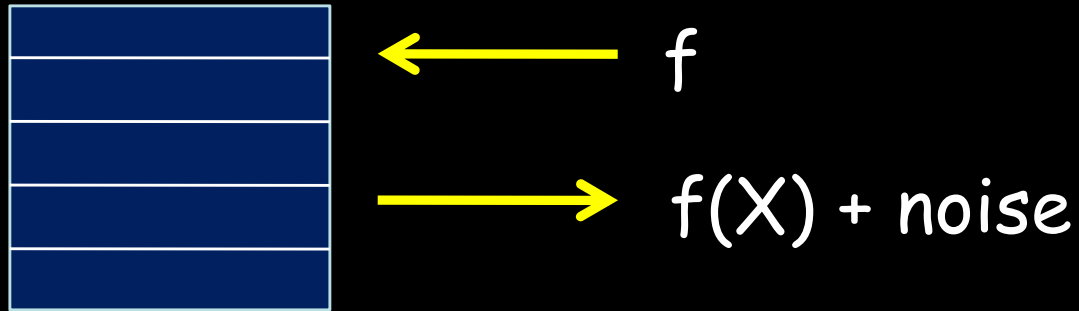
So, add noise roughly $1/\epsilon \times$ (effect any individual can have on outcome) gives desired ratio $e^\epsilon \approx (1+\epsilon)$.

If want answer within $\pm \alpha b$, need $n \geq 1/(\epsilon \alpha)$.

Utility/privacy/database-size tradeoff



Laplace mechanism more generally



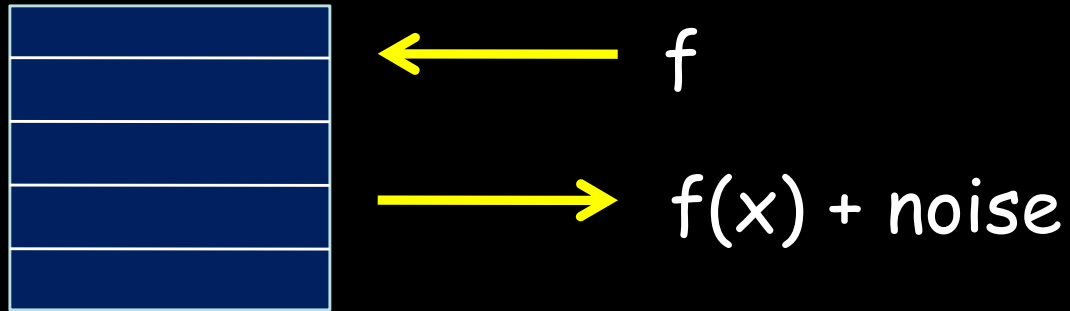
- E.g., f = standard deviation of income
- E.g., f = result of some fancy computation.

Global Sensitivity of f :

$$GS_f = \max_{\text{neighbors } X, X'} |f(X) - f(X')|$$

- Just add noise $\text{Lap}(GS_f / \epsilon)$.

What can we do with this?

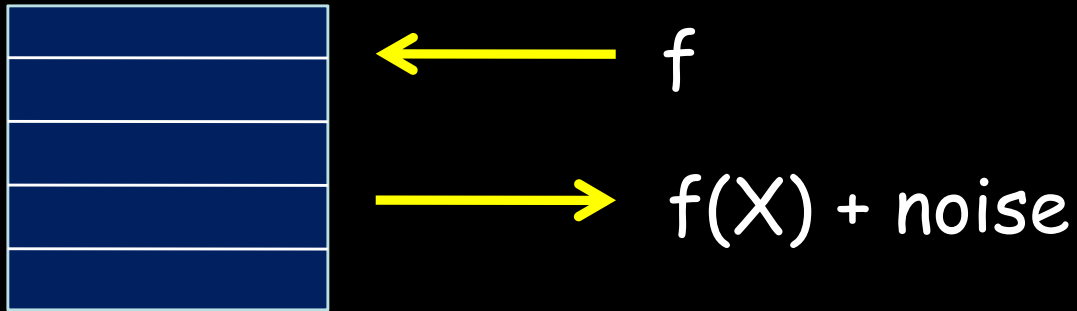


- Interface to ask questions
- Run learning algorithms by breaking down interaction into series of queries.
- *But*, each answer leaks some privacy:
 - If k questions and want total privacy loss of ϵ , better answer each with ϵ/k .

Remainder of talk

- Local sensitivity / Smooth sensitivity [Nissim-Raskhodnikova-Smith '07]
- Objective perturbation [Chaudhuri-Monteleoni-Sarwate '08]
- Sample and Aggregate [NRS '07]
- Exponential Mechanism [McSherry-Talwar '07]
- What can you say about publishing a sanitized database?

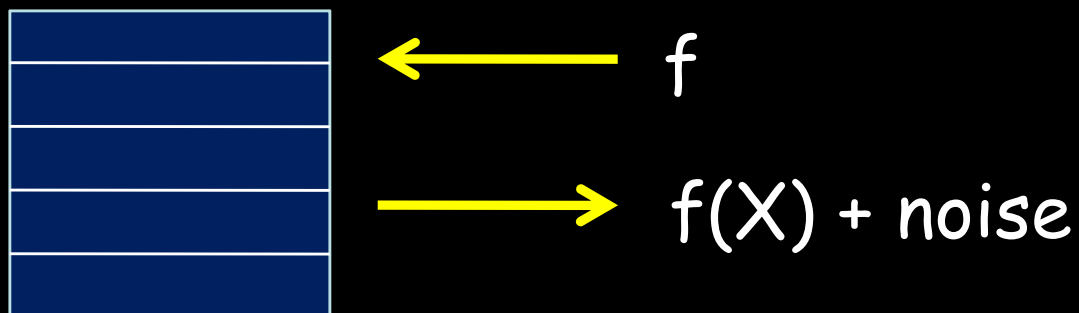
Local Sensitivity



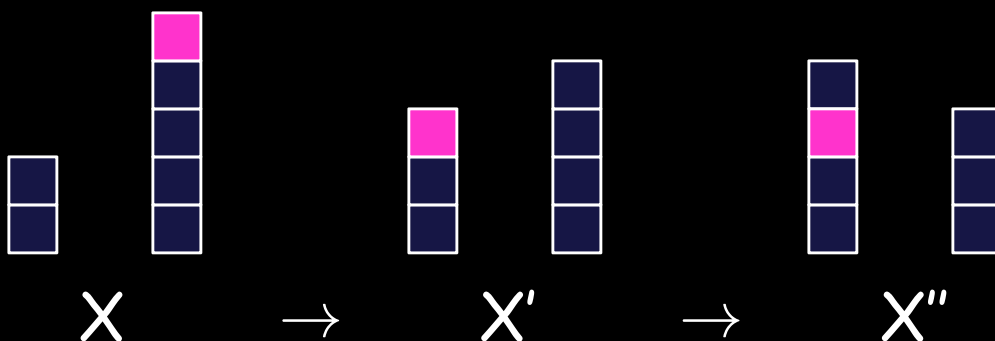
- Consider f = median income
 - On some databases, f could be *very* sensitive. E.g., 3 people at salary=0, 3 people at salary= b , and you.
 - But on many databases, it's not.
 - If f is not very sensitive on the actual input X , does that mean we don't need to add much noise?

$$LS_f(X) = \max_{\text{nbrs } X'} |f(X) - f(X')|$$

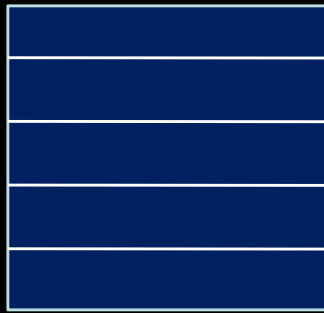
Local Sensitivity



- Consider f = median income
 - If f is not very sensitive on the actual input X , does that mean we don't need to add much noise?
- Be careful: what if sensitivity itself is sensitive?



Smooth Sensitivity



← f

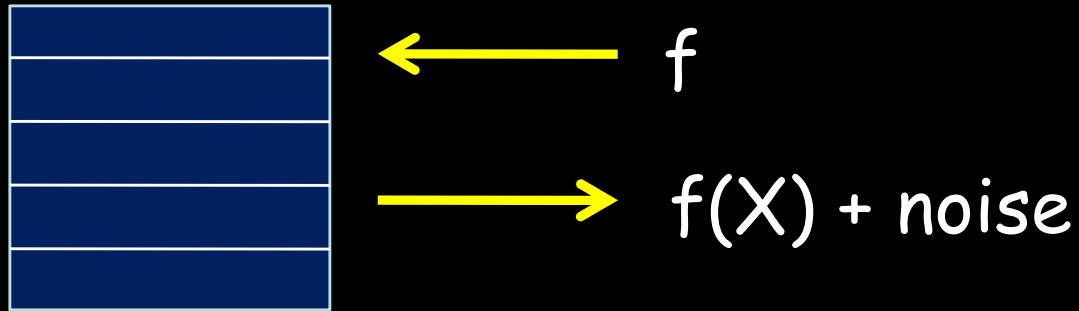
→ $f(X) + \text{noise}$

- [NRS07] prove can instead use (roughly) the following smooth bound instead:

$$\text{Max}_Y [LS_f(Y) \cdot e^{-\epsilon d(X,Y)}]$$

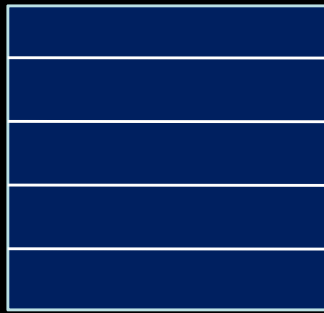
- With Anupam Datta, Jeremiah Blocki, Or Sheffet: looking at how to efficiently compute for various graph quantities in networks.

Aside...



- What should differential privacy mean in a networks context?
 - You can plausibly claim your actual set of neighbors was anything you want
 - Too stringent? How about "subject to being a non-celebrity?"
 - Too risky? Impact of your presence in network on other parts of network structure?

Smooth Sensitivity



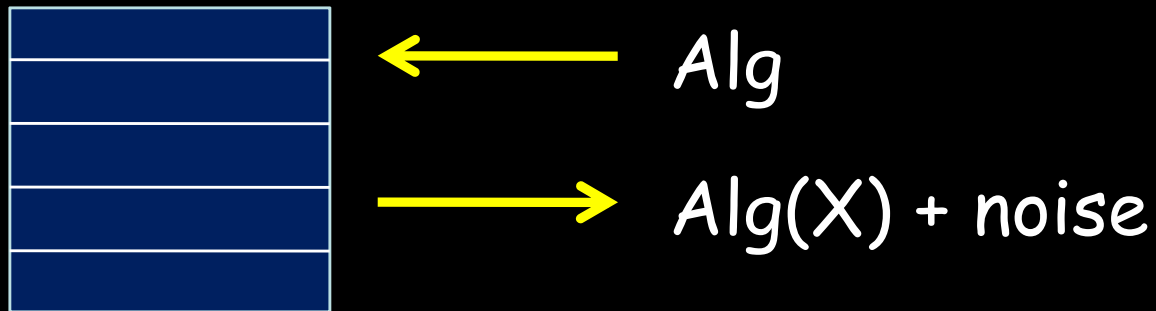
Alg



$\text{Alg}(X) + \text{noise}$

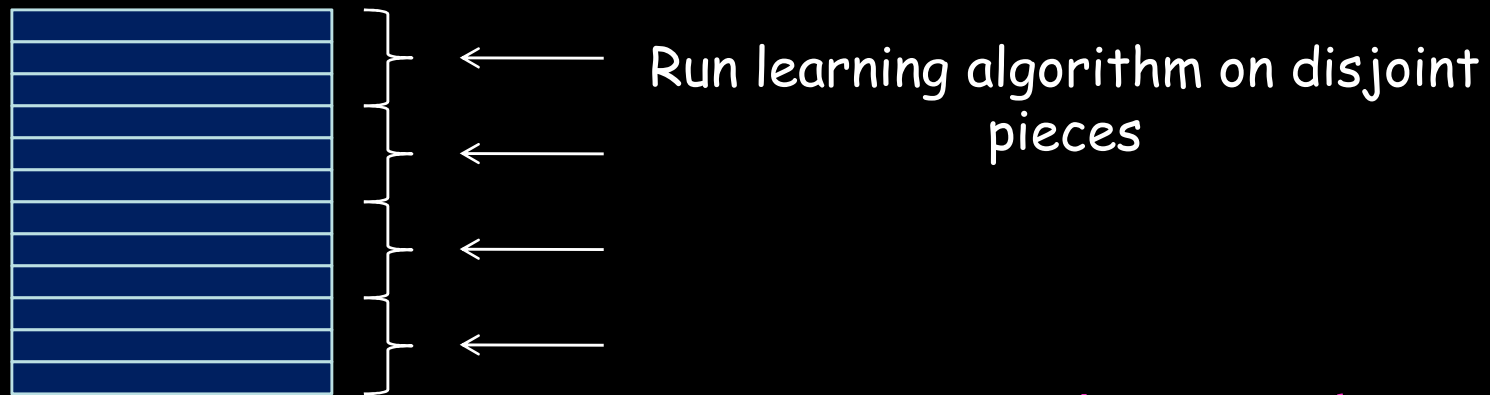
- In principle, could apply sensitivity idea to any learning algorithm (say) that you'd like to run on your data.
- But might be hard to figure out or might give really bad bounds.

Sample-and-aggregate (also [NRS07])

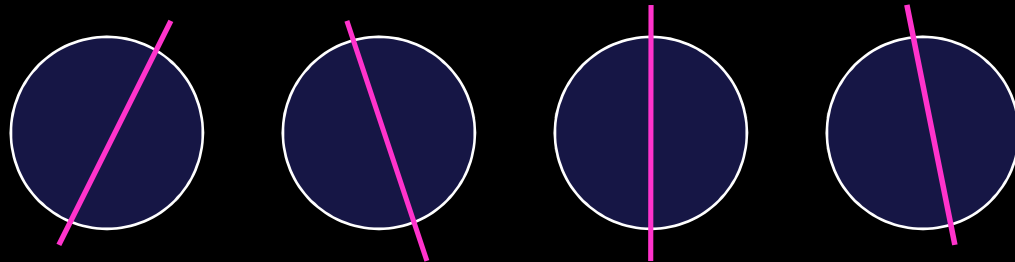


- Say you have some learning algorithm and hard to tell how sensitive it would be to changing a single input.
- Some way to run it privately anyway?

Sample-and-aggregate (also [NRS07])

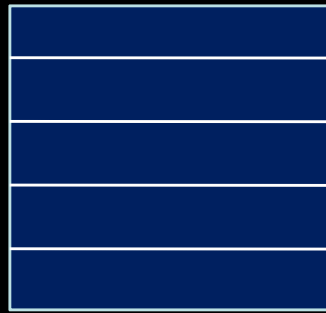


- *Get outputs*



- *Then average these outputs.*

Objective perturbation [CMS08]

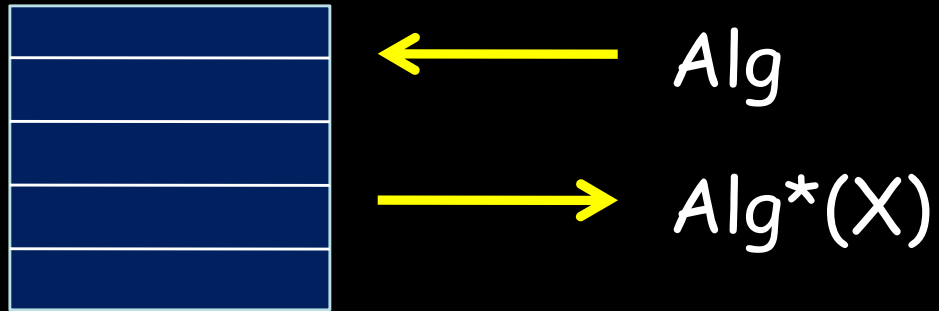


← $\text{Alg}^* = \text{Alg} + \text{noise}$

→ $\text{Alg}^*(X)$

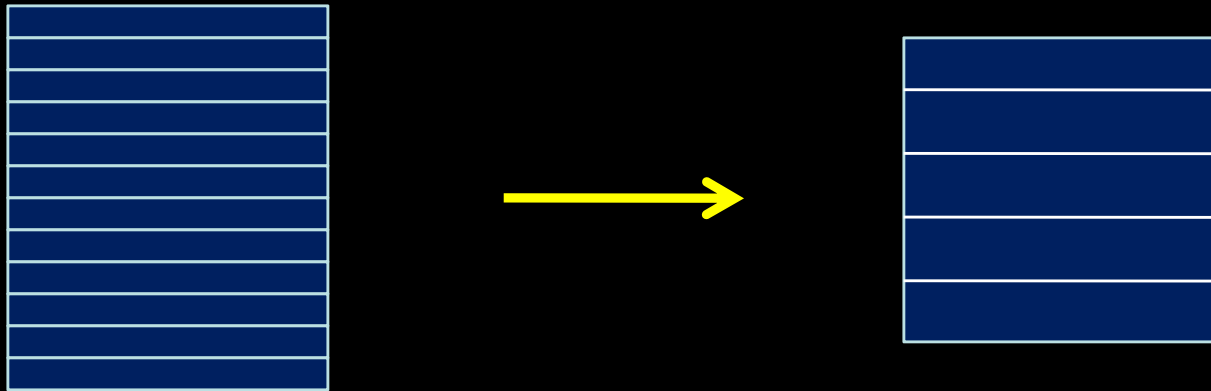
- Idea: add noise to the objective function used by the learning algorithm.
- Natural for algorithms like SVMs that have regularization term.
- [CMS] show how to do this, if use a smooth loss function.
- Also show nice experimental results.

Exponential Mechanism [MT07]



- What about running some generic optimization algorithm?
- [[skip for now]]

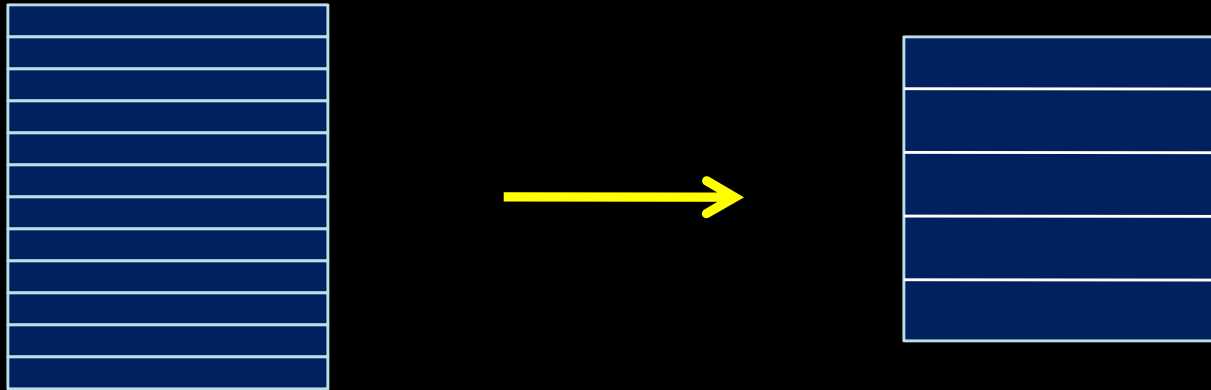
What about outputting sanitized databases?



- So far, just question-answering. Each answer leaks some privacy - at some point, have to shut down.
- What about outputting a sanitized database that people could then examine as they wish?

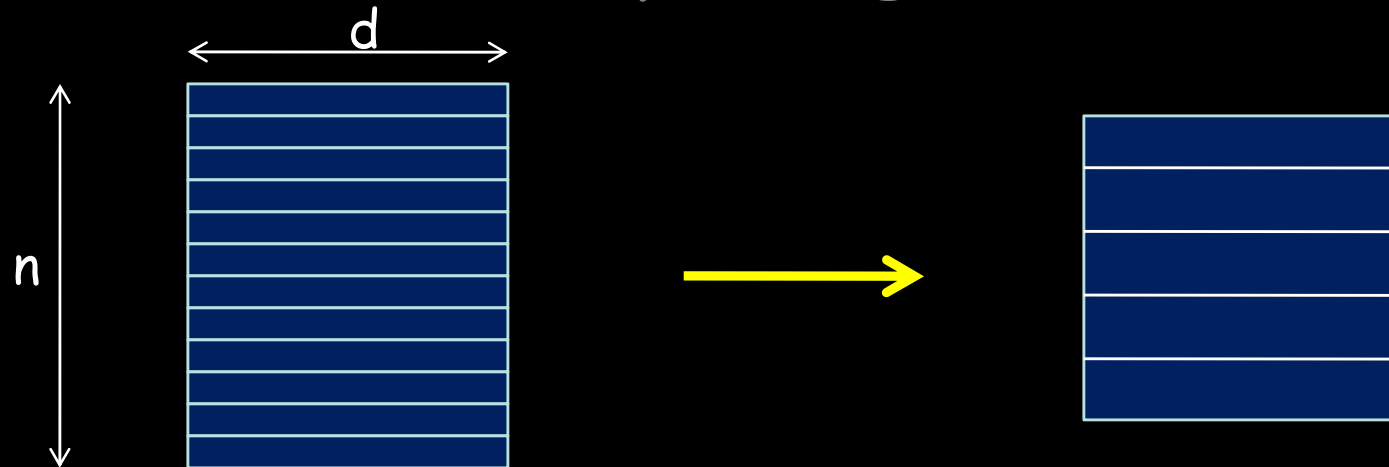
And is related to the original database...

What about outputting sanitized databases?



- Could ask a few questions (using previous mechs) and then engineer a database that roughly agrees on these answers.
- But really, we want a database that matches on questions we haven't asked yet.
- Do you need to leak privacy in proportion to number of questions asked?

What about outputting sanitized databases?

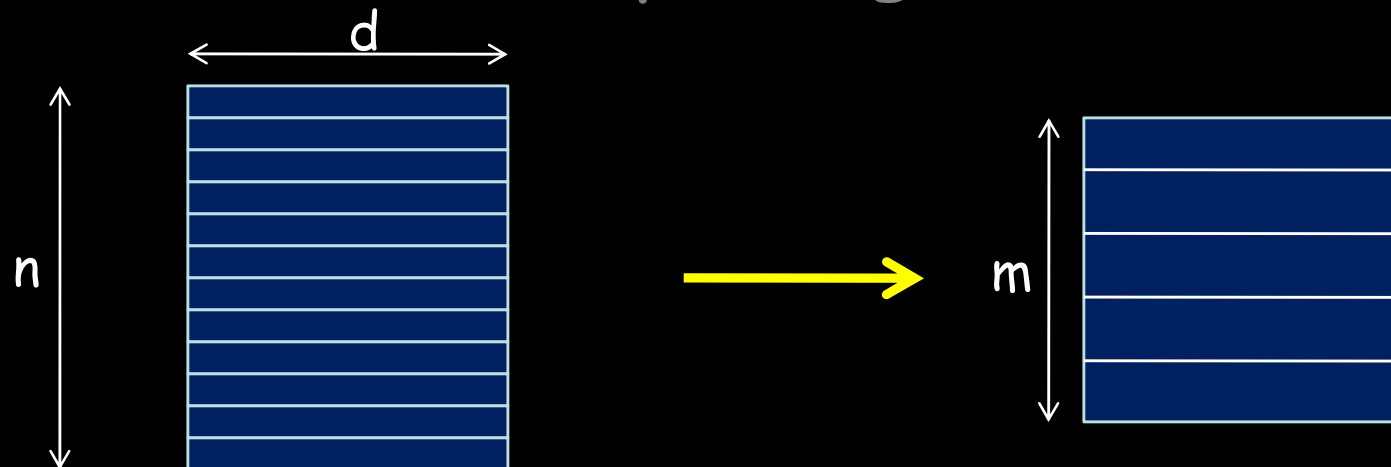


Actually, no you don't [B-Ligett-Roth] (At least not for count-queries)

- Fix a class C of quantities to preserve. E.g., fraction of entries with $x[i_1]=1, x[i_2]=0 \dots x[i_k]=1$.
- Want ϵ -privacy and preserve all $q \in C$ up to $\pm \alpha$.
- [BLR] show: in principle, can do with database of size only $n = O(d \log |C|)$.

Allowing exponentially-many questions!

What about outputting sanitized databases?

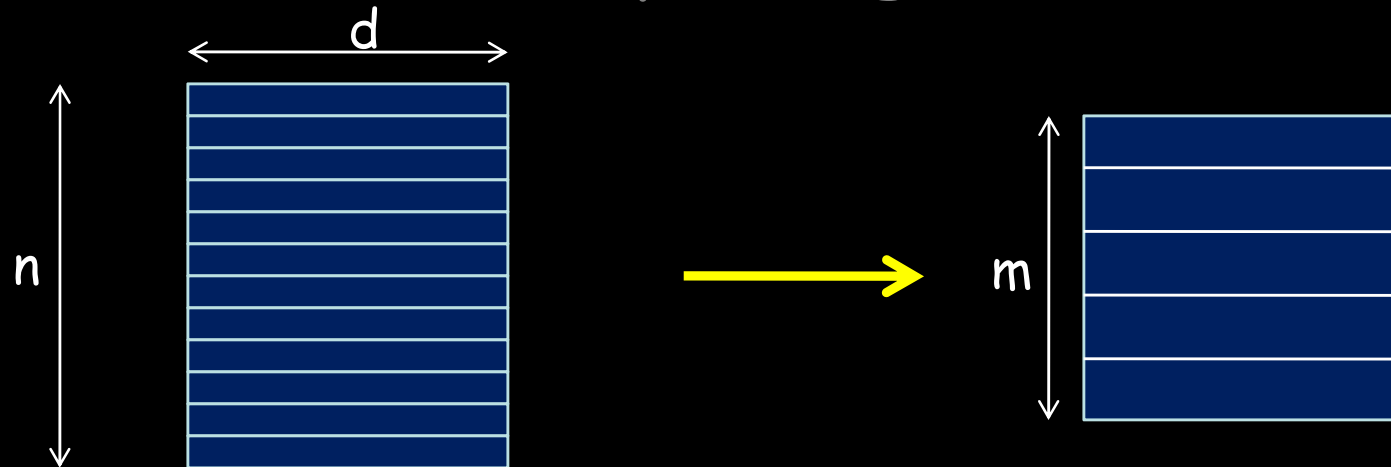


$$\Pr(D) \propto e^{-O(\epsilon n \text{ penalty}(D))}$$

Idea:

- Standard results from learning/statistics say that there exist small databases that apx preserve all quantities in C . $m = O(\log |C|)$ is sufficient.
- Put explicit distribution on them, using exponential mechanism of [McSherry-Talwar]
- Solve to get $n \approx (d \log C)/(\epsilon \alpha^3)$

What about outputting sanitized databases?



$$\Pr(D) \propto e^{-O(\epsilon n \text{ penalty}(D))}$$

Idea:

- Standard results from learning/statistics say that there exist small databases that apx preserve all quantities in C . $m = O(\log |C|)$ is sufficient.
- Put explicit distribution on them, using exponential mechanism of [McSherry-Talwar]
- But, seems extremely hard to get efficient alg.

Differential Privacy summary & discussion

Positives:

- Clear semantic definition. Any event (anything an adversary might do to you) has nearly same prob if you join or don't join, lie or tell the truth.
- Nice composability properties.
- Variety of mechanisms developed for question answering in this framework.
- *Some* work on sanitized database release.

Differential Privacy summary & discussion

Negatives / open issues

- It's a pessimistic/paranoid quantity, so may be more restrictive than needed.
- " ϵ " is not zero. Privacy losses add up with most mechanisms (but see, e.g., [RR10],[HR10])
- Doesn't address group information.
- Notion of "neighboring database" might need to be different in network settings.
- ...