

Membership Inference



Membership Inference Attacks

Recall: data reconstruction attacks
"blatantly non-private" algorithms
↳ show that something is not private

Defeating reconstruction attacks $\nRightarrow$ privacy

A quick detour through crypto

Key recovery

vs

IND-CPA ~~seems~~
in distinguishability

↳ clearly insecure

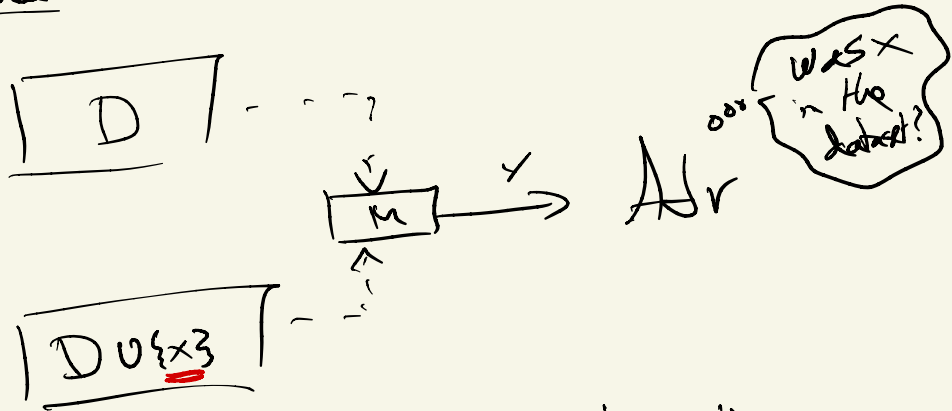
↳ easy to create schemes where key recovery is impossible but are clearly insecure

E.g. $\text{Enc}(k, m) := m$

or $:= \text{DES}(k_{1..n/2}, m)$

Differential privacy $\approx$ hardness of inferring membership

Experiment

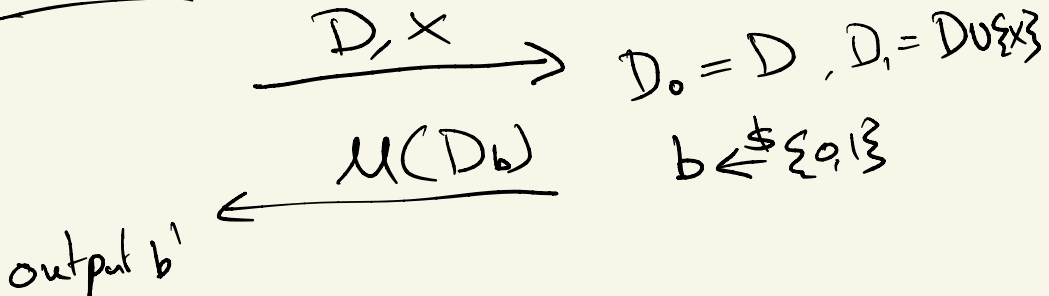


- "Membership inference attack"
- "Tracing attacks"

A formal security game for MI attacks

Adversary

Challenger



Adv wins if $b' = b$

Quantifying Adv's success

$$1) \text{Adv}_{\text{MI}}(\lambda) = 2 \left| \Pr[b' = b] - 1/2 \right|$$

in crypto this is fine because we want
 $\text{Adv}(\lambda) = \text{negl}(\lambda)$

With DP guarantees, we can never get
 $\text{negl}(\lambda)$ advantage (if we also want utility)

2) Confusion matrix:

$$\text{TPR} = \Pr[b' = 1 \mid b = 1]$$

$$\text{FPR} = \Pr[b' = 1 \mid b = 0]$$

$$\text{TNR} = \Pr[b' = 0 \mid b = 0]$$

$$\text{FNR} = \Pr[b' = 0 \mid b = 1]$$

Connection between DP and MI:

Theorem [KOV'15]. Let $\epsilon > 0$, $\delta \in [0, 1]$

A mechanism $\mathcal{M}$ is (ϵ, δ) -DP iff

$$e^{\epsilon} \cdot \text{FNR} \geq \text{FNR} - \delta \quad \text{and}$$

$$e^{\epsilon} \cdot \text{FPR} \geq \text{FPR} - \delta$$

for any adversary in the MI game

Observations:

DP protects against something much
stronger than reconstruction
(namely MI)

A mechanism M is (ϵ, δ) -DP iff

$$e^\epsilon \cdot \text{FNR} \geq \text{TPR} - \delta \quad \text{and} \quad (*)$$

$$e^\epsilon \cdot \text{FPR} \geq \text{TPR} - \delta$$

Proof 1) (ϵ, δ) DP $\Rightarrow (*)$

For simplicity, assume that Adv is deterministic
let S be the set of outputs of M for
which Adv outputs $b' = 0$.

$$\text{By DP: } \underbrace{P_r[M(D_0) \in S]}_{\text{FNR}} \leq e^\epsilon \underbrace{P_r[M(D_1) \in S]}_{\text{FNR}}$$

by symmetry, we can get the same for TPR / FPR

2) $(*) \Rightarrow (\epsilon, \delta)$ -DP

Assume that $\forall S \subseteq \mathcal{Y}$, if the attacker
outputs $b' = 0$ whenever M 's output is in S ,
then $(*)$ holds

$$\Rightarrow \Pr[u(D_0) \in S] \leq e^\epsilon \cdot \Pr[u(D_1) \in S] + \delta$$

and if we assume that all sets $S \subseteq Y$ where
 Adv outputs $b' = 1$ whenever u 's output
 is in S ,

$$\Pr[u(D_1) \in S] \leq e^\epsilon \Pr[u(D_0) \in S] + \delta$$

$\Rightarrow$ This is exactly (ϵ, δ) -DP

Corollary A mechanism is (ϵ, δ) -DP
 iff for all adversaries in the MI game
 we have:

$$e^\epsilon \geq \max\left(\frac{\text{TPR} - \delta}{\text{FNR}}, \frac{\text{TPR} - \delta}{\text{FPR}}\right)$$

example: if there exists an attacker
 with a TPR = 90% at a FPR of 10%
 then the mechanism is not ϵ -DP for
 $\epsilon \leq \ln 90 \approx 4.5$

Interpreting / applying MI :

1) Actual attack (goal of Adv : infer whether someone participated in some "sensitive" statistical experiment)

this is a rather weak attack

⇒ what we assume: Adv knows all the data about x

⇒ what we get: Adv learns one bit

2) lower-bounding / auditing DP

⇒ if we can find a MI attack, we know the system is not DP (for some ϵ budgets)

Analogy : Chosen plaintext attack

Adv chall
no, m $k \notin$
→ $\text{Enc}(m)$
←
↳ guesses b

↳ in practice, we don't really care about an attacker learning b

↳ why we care : if such an attack exists, the scheme is not secure

From MI to lower bounds on DP

Recall: lower bounds from reconstruction attacks
if Dinur Nissim attacks work $\Rightarrow$ scheme is not DP

We know: you can't answer n counting queries with $o(1)$ noise
 $\hookrightarrow$ check your HW!

Now: a similar (but tighter) result for the Gaussian mechanism

what we know: if we add noise $\Sigma = \tilde{O}(K/n)$ to each query, we can answer k counting queries with (ϵ, δ) -DP (Gaussian mechanism)

Theorem: let $\mathcal{U}(D) = \frac{1}{n} \sum_{i=1}^n \vec{x}_i + N(0, \Sigma^2 I_{nk})$
 $\vec{x}_1, \dots, \vec{x}_n \in \{0, 1\}^{n \times k}$

Assume $\sigma = o(\sqrt{K/n})$, and $\delta = o(1)$

Then $\mathcal{U}$ is not (ϵ, δ) -DP for any $\epsilon = o(1)$

Proof: we'll build a MI attack against $\mathcal{M}$

- Adv gets to pick D and $\vec{x}$

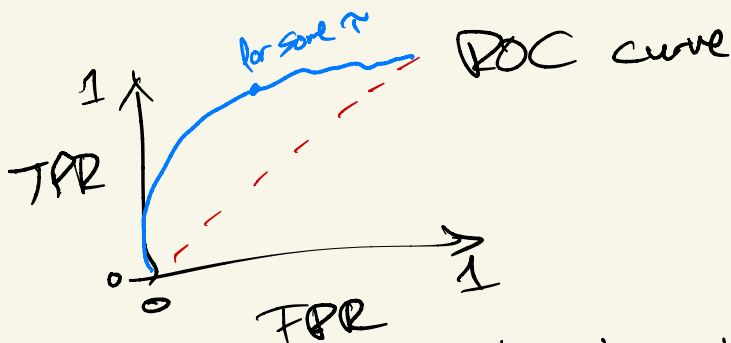
$$D_0 = \{\vec{0}, \dots, \vec{0}\} \quad \vec{x} = [1, 1, \dots, 1]$$

$$D_1 = \{\vec{0}, \dots, \vec{x}\}$$

- Adv sends $D, \vec{x}$ to the challenger and gets back $\vec{y} = \mathcal{M}(D_b)$

$$D_0 = D \quad D_1 = D \cup \vec{x}$$

- Adv computes $z = \langle \vec{y}, \vec{x} \rangle$ and outputs $b' = 1$ iff $z \geq \tau$



Intuition: • if $b=0$ (so $\vec{x}$ is not in the input)
then $\vec{y} = \mathcal{M}(D_0) = \mathcal{N}(0, \sigma^2 I)$

$$\mathbb{E}[\langle \vec{y}, \vec{x} \rangle] = 0$$

• if $b=1$, $\vec{y} = \mathcal{M}(D_1) = \begin{bmatrix} 1 \\ \vdots \\ 1 \end{bmatrix} + \mathcal{N}(0, \sigma^2 I)$

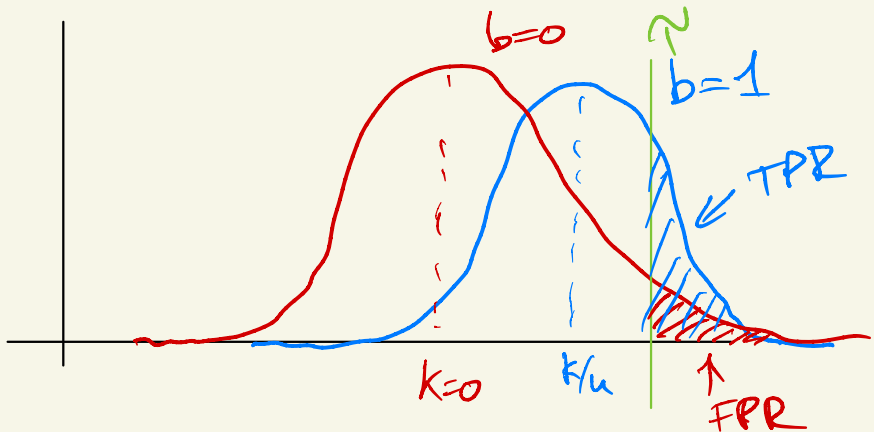
$$\mathbb{E}[\langle \vec{y}, \vec{x} \rangle] = \frac{(\sum x_i^2)}{L}$$

More formally:

$$\text{if } b=0: \mathbf{z} = \begin{bmatrix} N(0, \sigma^2) \\ \vdots \\ N(0, \sigma^2) \end{bmatrix} \cdot [1, \dots, 1] = \sum_{i=1}^K N(0, \sigma^2) \\ = N(0, K\sigma^2)$$

$$\begin{bmatrix} \mu_1 + N \\ \vdots \\ \mu_n + N \end{bmatrix} \cdot [1, \dots, 1]$$

$$\text{if } b=1: \mathbf{z} = \langle \vec{\gamma}, \vec{x} \rangle = N(0, K\sigma^2) + \langle \vec{\frac{\mathbf{K}}{n}}, \vec{x} \rangle \\ = N\left(\frac{K}{n}, K\sigma^2\right)$$



let's analyze the attack:

$$\text{set } \tau = k/u$$

$$\bullet \text{ TPR} = \Pr[z \geq \tau \mid b=1] = 50\%$$

by symmetry
of Gaussian

• upper bound FPR

$$\text{let } X \sim \mathcal{N}(0, \sigma^2) \text{ and } t > 0 \text{ then} \\ \Pr[X > t] \leq \exp\left(-\frac{t^2}{2\sigma^2}\right)$$

$$\begin{aligned} \text{FPR} &= \Pr[z \geq \tau \mid b=0] \\ &= \Pr[\mathcal{N}(0, k\sigma^2) \geq k/u] \\ &\leq \exp\left(-\frac{(k/u)^2}{2k\sigma^2}\right) \end{aligned}$$

$\sigma = o\left(\frac{\sqrt{k}}{n}\right)$

$$\exp\left(\underbrace{-\frac{k}{2n^2\sigma^2}}_{o(1)}\right) = o(1)$$

$$\begin{aligned} \exp\left(-\frac{k}{n^2 \sigma^2}\right) & \quad \sigma^2 = o\left(\frac{\sqrt{k}}{n}\right) \\ & = \exp\left(-\frac{k}{n^2 \frac{n^4}{k}}\right) \quad \text{e.g.} = \frac{\sqrt{k}}{n^2} \\ & = \exp(-u^2) = o(1) \end{aligned}$$

$$\text{TPR} = 50\%$$

$$\text{FPR} = o(1)$$

$$\begin{aligned} \text{by Corollary:} \quad \epsilon & \geq \ln\left(\frac{\overbrace{\text{TPR} - \delta}^{\text{constant}}}{\underbrace{\text{FPR}}_{\text{less than constant}}}\right) \\ & = \omega(1) \end{aligned}$$

What we showed: to answer k counting queries with the Gaussian mech. and (ϵ, δ) -DP for $\epsilon = o(1)$ and $\delta = o(1)$ we need $\sigma = \Omega\left(\frac{\sqrt{k}}{n}\right)$ and thus incur error $\Omega\left(\frac{\sqrt{k}}{n}\right)$ per query ∇

Membership inference for machine learning

Adversary

Challenger



train a model f
on a dataset Δ

guess whether
 x was in
the training data

Deep neural networks

• $f_{\theta}: X \rightarrow [0, 1]^C$ ← # of classes

↑ weights ↑ input

• Denote $f_{\theta}(x)_y$ the probability of class y

• A feed-forward network:

$$f(x) = \underbrace{\sigma(f_{x_0} \dots f_{x_{k-1}} \circ f_1(x))}_{z_{\theta}(x)}$$

↑ softmax ↑ k layers

$$z_{\theta}: X \rightarrow \mathbb{R}^C$$

$$\nabla L(z) = \left[\frac{e^{z_1}}{\sum_i e^{z_i}}, \dots, \frac{e^{z_c}}{\sum_i e^{z_i}} \right] e^{\{0,1\}^c}$$

$\nearrow$ softmax
 $\nearrow$ logits

Training a neural network : gradient descent $\nwarrow$ loss, eg. cross entropy

$$\Theta_{t+1} \leftarrow \Theta_t - \eta \sum_{(x,y) \in \mathcal{B}} \nabla_{\Theta} \ell(\Theta_t, x, y)$$

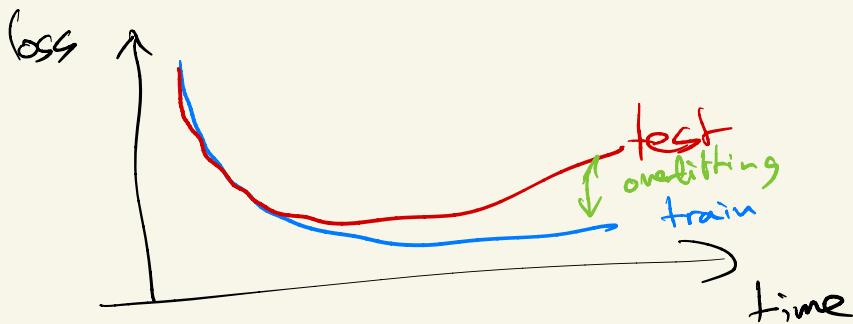
$\nwarrow$ batch

Overfitting

$$\hat{L}_D(\Theta) \leq L(\Theta)$$

$\nwarrow$ empirical loss
 $\nwarrow$ test loss

$\nearrow$ training set



How do to MI?

$$b \leftarrow \{0, 1\}^S$$

$$D_0 = D$$

$$D_1 = D \cup \{x\}$$

$$\theta \xrightarrow{2} \theta \leftarrow \text{Train}(D_b)$$

How to test
whether θ
comes from D_0 or D_1 ?

1) "White-box attack": look at the parameters θ and figure out if some info on x is encoded in there somewhere

2) "Black-box attack": we'll just look at model outputs $\theta(x)$, in particular

the loss $\ell(\theta, x, y)$

if $x \in D$
we expect
this loss to be $\approx \hat{L}_D(\theta, x, y)$

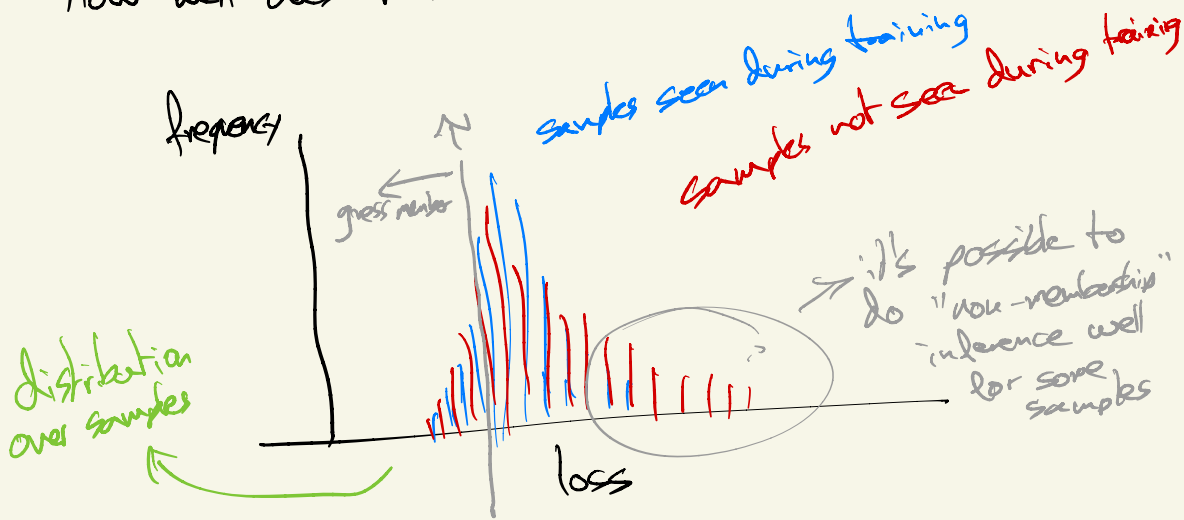
if $x \notin D$
we expect this
loss to be $\gg \hat{L}_D(\theta, x, y)$

A first RI attacks: Global loss thresholding

$$A_{\text{global}}(\theta, x, y, \tau) = \begin{cases} 1 & \text{if } \ell(\theta, x, y) \leq \tau \\ 0 & \text{otherwise} \end{cases}$$

↑
query access
is enough

How well does this work? (spoiler: not very well)



Evaluation

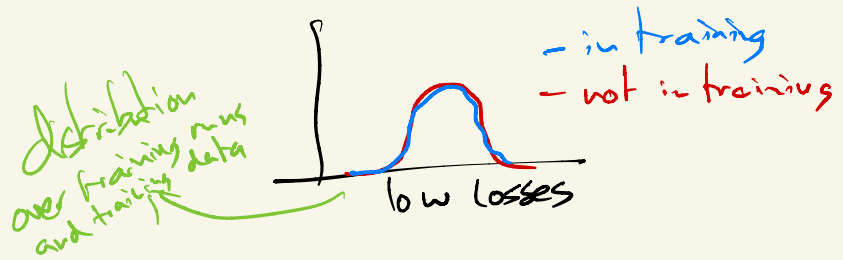
- Accuracy = $P_{\theta}[b' = b] \approx 70\%$
over many $x's \sim P$

- TPR to FPR ratio ≈ 1 (for small FPR)

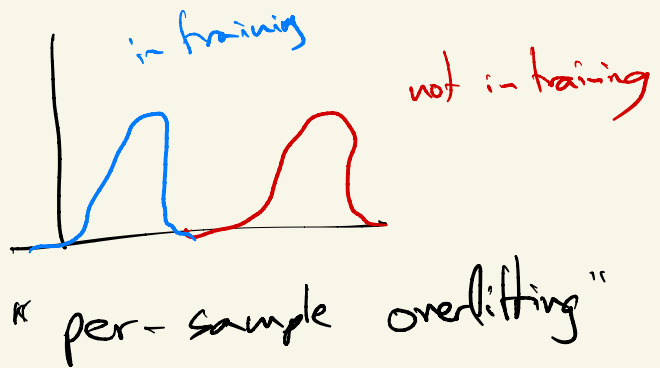
What we want : confident membership inference for some samples

Difficulty calibration

- inliers (a prototypical "cat" or "dog")



- outliers (a cat that looks like a dog)



The people most at risk from privacy inference might be the ones who used privacy the most

A stronger attack: per sample likelihood ratios

Likelihood ratio test

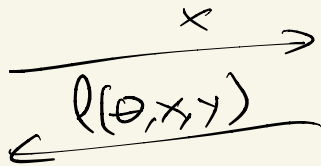
- Suppose you have a null hypothesis H_0 and an alternative hypothesis H_1 and you want to decide which is more likely given observed data $x \in D$

- likelihood ratio test:
$$\Lambda = \frac{\Pr[\text{data} | H_0]}{\Pr[\text{data} | H_1]}$$

- Neyman-Pearson: the test " $\Lambda \leq \tau$ " is the test that yields the highest TPR at any fixed FPR

How do we do a likelihood ratio test for a neural network?

Adv



Chall

trains θ on D_b

we want to compute

$$P_r \left[l(\theta, x, y) \mid \theta \text{ was trained on } D_b \right]$$

training randomness

we don't know how to compute this probability ...

We'll estimate it by "shadow models"

- train N models on datasets that either contain or don't contain x
- $D \leftarrow P^n$ let $D_0 = D$, $D_1 = D \otimes x^3$
 $\downarrow$ part $\downarrow$ $P \rightarrow N$
 θ θ

$$Loss_{IN} = \left\{ \text{loss}(\theta^{IN}, x, y) \mid \begin{matrix} D \leftarrow P \\ \theta^{IN} \leftarrow \text{Train}(D, \text{loss}) \end{matrix} \right\}$$

$$Loss_{OUT} = \left\{ \text{loss}(\theta^{OUT}, x, y) \mid \begin{matrix} D \leftarrow P \\ \theta^{OUT} \leftarrow \text{Train}(D) \end{matrix} \right\}$$

Last step: approximate $P_0[l(\theta, x, y) \mid H_0]$
 by $N(\mu_{IN}; \Sigma_{IN}^2)$
 $N(\mu_{OUT}, \Sigma_{OUT}^2)$

then: $\frac{Pr[l(\theta, x, y) \mid H_0]}{Pr[l(\theta, x, y) \mid H_1]}$ is the
 ratio of two Gaussians

Results

TPR

at FPR=1%



Aside:

Idea: 1) Collect a dataset D

2) calculate the TPR at FPR=1%
1% = all samples

3) remove all samples with
 $TPR > 1\% \Rightarrow D'$

4) train a model on D'

5) Q: how vulnerable are the
most vulnerable samples in D' ?

A: $\Rightarrow 10\%$ TPR at FPR=1%

What are MI attacks for machine learning useful for?

Auditing Privacy (lower bounds on DP)

Idea: Use Corollary 1 ($e^\epsilon \geq \frac{\text{TPR} - \delta}{\text{FPR}}$)
to lower bound privacy

Why?

1) Debugging DP-SGD implementations

Ex: • I claim $\epsilon = 0.2$ for my implementation

• our MI attack $\rightarrow \frac{\text{TPR}}{\text{FPR}} \approx e^{2.7} \approx 15$

by corollary 1, $\epsilon \geq 2.7$ 

2) Understanding the "Empirical privacy" of DP-SGD

Recall the analysis of DPSGD:

for one step of gradient descent, the analysis is tight if one gradient has norm $\geq C$ and is orthogonal to all other gradients.

$$\text{ex: } D = \{\vec{0}, \dots, \vec{0}\} \quad \vec{x} = \begin{bmatrix} 0 \\ \vdots \\ c \end{bmatrix}$$

The full analysis is tight if:

1) each step is tight

2) the attacker gets to see all intermediate model updates $\theta_1, \dots, \theta_{T-1}$

intuitively, the privacy of DP-SGD should be better than what the analysis says

We can estimate this empirically using MI attacks ! (DPSGD auditing)

(Roughly) the results: it seems that the DPSGD analysis is extremely conservative
(e.g. provable $\epsilon \leq 20$, the best MI attack we have might reach $\epsilon \geq 1$)

Q: - are our attacks bad?
- is the DPSGD analysis bad?