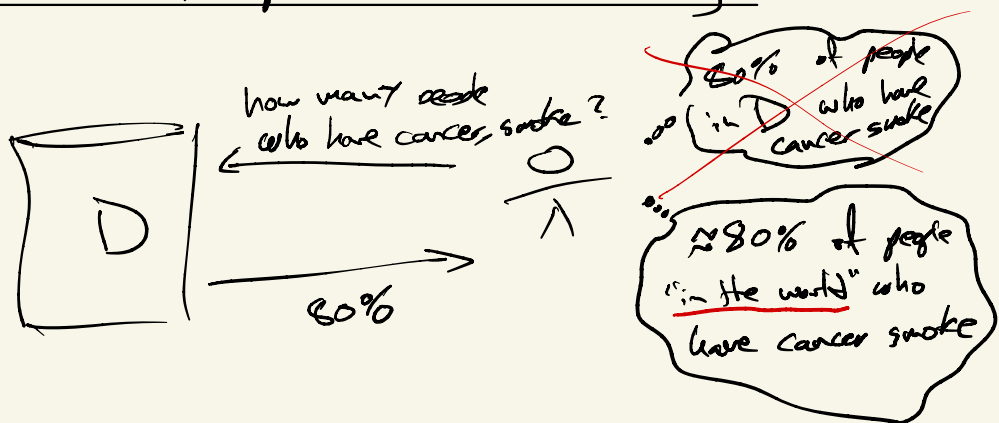



Differentially private learning



Generalization for simple stats

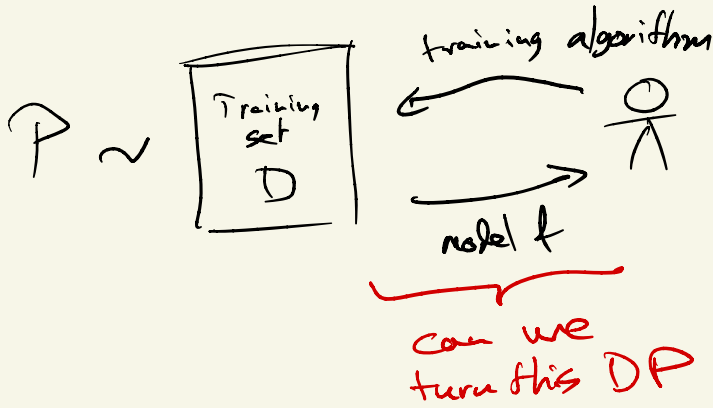
- D consists of n samples taken i.i.d from $\mathcal{P}$
 - we compute empirical mean $\bar{x} = \frac{1}{n} \sum x_i$
 - in expectation, $|\bar{x} - \mu| = O(1/\sqrt{n})$
- $\mathbb{E}_{\mathcal{P}^n}[\bar{x}]$ $\uparrow$ sampling error

Note: • sampling error: $O(1/\sqrt{n})$

• noise added to $\bar{x}$ to get DP: $O(1/\sqrt{n})$

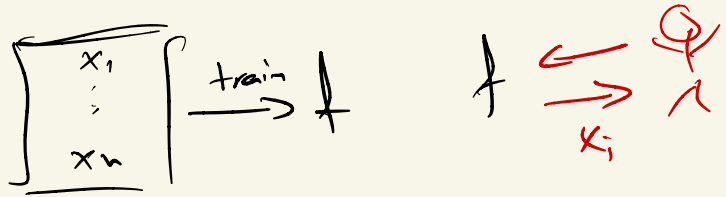
privacy comes for free

Machine Learning



Why private learning?

- ML models can "overfit" to training data
- ML models can "memorize" some data



Can we prevent this?

Non-private learning

Setting: • distribution P over inputs X with labels Y

• Goal: learn a model $f_\theta: X \rightarrow Y$
that "generalizes" ← parameters of the model

Defining generalization:

* loss function $\ell(\theta, x, y)$ measures the "error" of f_θ on input (x, y)

ex: $\ell(\theta, x, y) = \mathbb{1}\{f_\theta(x) \neq y\} \in \{0, 1\}$

* risk $\mathcal{L}(\theta)$ is expected loss over P :

$$\mathcal{L}(\theta) = \mathbb{E}_{(x, y) \sim P} [\ell(\theta, x, y)]$$

"test loss"
"test error"

we want a model with low risk!

* what do we get?

Training set $D = \{(x_1, y_1), \dots, (x_n, y_n)\}$
where $(x_i, y_i) \stackrel{\text{iid}}{\sim} P$

empirical risk $\hat{\mathcal{L}}_D(\theta) = \frac{1}{n} \sum_{i=1}^n \ell(\theta, x_i, y_i)$ "training loss"
or error

Empirical risk minimization: (ERM)

- find θ^* that minimize the loss on the training set

$$\theta^* = \underset{\theta \in \text{ALL MODELS}}{\operatorname{argmin}} \hat{L}_D(\theta)$$

ignore whether this is poly time...

ex: all 10-layer neural nets with all parameters and architecture...

- "Generalization gap": "overfitting"

$$|\text{test error} - \text{train error}|$$

$$= |L(\theta^*) - \hat{L}_D(\theta^*)|$$

Note: if we can fit the training set exactly ($\hat{L}_D(\theta^*) = 0$), then the generalization gap is just $L(\theta^*)$

- How to control the generalization gap?
regularization, data augmentation, etc.

How to turn ERM private?

- Output perturbation Mostly of theoretical interest

apply a standard training algorithm $T(D)$ to get a model f , and then "add noise" to f .

↳ not clear what this means?
e.g. adding noise to Θ

↳ we need to know the sensitivity of the ERM algorithm

↳ this is possible for "simple" models (linear/logistic regression)

- Objective perturbation mostly theoretical interest

instead of solving $\argmin_{\Theta} \hat{L}_D(\Theta)$

solve $\argmin_{\Theta} \hat{L}_D(\Theta) + \underbrace{b^T \Theta}_{\text{random}}$

↳ for some models, we can show this is DP!

↳ we are solving ERM exactly, just with a different objective

- Gradient perturbation
Specific training algorithm: gradient descent

Private Projected Gradient Descent

$$\Theta_0 \leftarrow \Omega \quad \leftarrow \text{set of parameters}$$

for $t = 0 \dots T-1$:

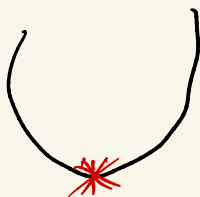
$$g \leftarrow \nabla_{\Theta_t} \hat{\mathcal{L}}_D(\Theta_t) = \frac{1}{n} \sum_{i=1}^n \nabla_{\Theta_t} \ell(\Theta_t, x_i, y_i)$$

$$\Theta_{t+1} \leftarrow \Theta_t - \eta \cdot (g + N(0, \sigma^2 \cdot I_{\text{dim}}))$$

$$(\Theta_{t+1} \leftarrow \text{Project } \Theta_{t+1} \text{ onto } \Omega)$$

return Θ_T

ERM



Output perturbation



Objective perturbation



Gradient perturbation

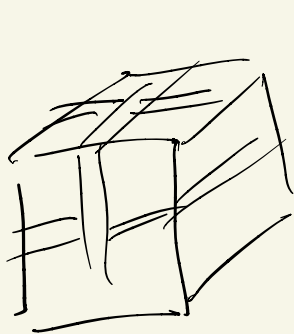


Utility: in convex setting, in worst-case
the training error of private ERM
is $O(\sqrt{d}/\epsilon)$ (larger than for
Standard ERM)
 $\leftarrow$ dimension of the model (# of learnable parameters)

Private Deep learning

train models where $d \gg n$
and no convexity
 $\nwarrow$ # of parameters $\nwarrow$ # of samples

Theory: this won't generalize $\Rightarrow$ Wrong
DP Theory: this can't be learned with privacy $\Rightarrow$ Wrong



DP-SGD

"a miracle"

Crucial step: making the batch gradient private

* Sample a batch of examples B

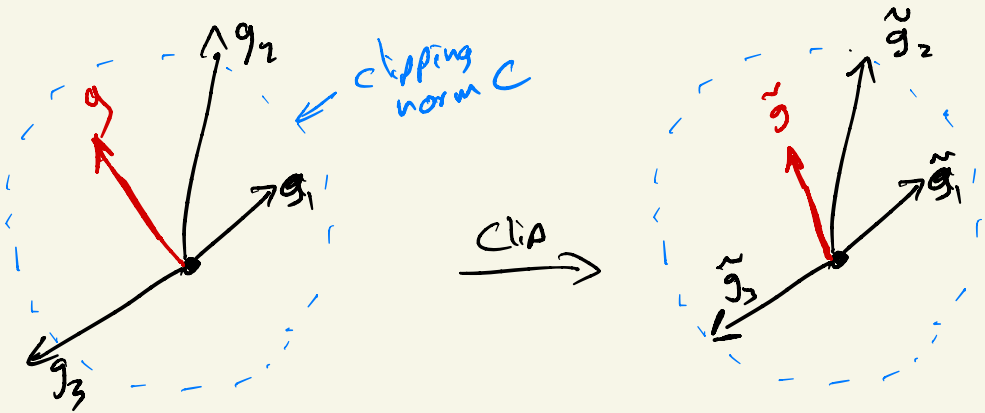
* calculate the mean gradient

$$g = \frac{1}{|B|} \sum_{i=1}^{|B|} \underbrace{\nabla_{\theta} \mathcal{L}(\theta, x_i, y_i)}_{g_i}$$

$$= \frac{1}{n} \sum_{i=1}^n g_i$$

how large can this vector be
(in ℓ_2 norm)?

* Clip the gradients to enforce
a maximal ℓ_2 norm



Observations: 1) this biases the gradient
2) implementing this in practice (efficiently) is non-trivial

3) how to choose C ?

- large: $\tilde{g}$ is less biased, but you will need to add more noise for DP
- small: $\tilde{g}$ is very biased, but less noise
- in practice: C is a hyperparameter

DP-SGD (Abadi et al. 2016)

$$\theta_0 \leftarrow \Omega$$

for $t = 0 \dots T-1$ do

sample a batch $S \subseteq [n]$ of size m

$$g_i \leftarrow \nabla_{\theta_t} \ell(\theta_t, x_i, y_i) \quad \forall i \in S$$

$$g_i \leftarrow C \cdot g_i / \max(C, \|g_i\|_2)$$

$$\tilde{g} \leftarrow \frac{1}{m} \sum_{i \in S} g_i \quad \leftarrow \text{these all have } \ell_2 \text{ norm} \leq C$$

$$\theta_{t+1} \leftarrow \theta_t + \mathcal{N}(\tilde{g} + N(0, \sigma^2 \cdot I_{\text{dim}}))$$

Analysis

1) "Naïve"

- the sensitivity of the gradient $\tilde{g}$ is $\frac{2C}{m}$
- we can add Gaussian noise

$$\sigma^2 = 2(2C)^2 \log\left(\frac{2.5}{\delta}\right) / m^2 \epsilon^2 \quad \text{to get } (\epsilon, \delta)\text{-DP}$$

for one gradient step

✗ Basic composition: $(\epsilon T, \delta T)$ -DP

✗ Advanced composition: $(\epsilon \sqrt{T \log(1/\delta)}, \delta T)$ -DP
↳ a lot better, but still not enough!

Amplification by subsampling

Lemma: let M be (ϵ, δ) -DP. Let $M'(D)$ be the algo that subsamples each element of D w.p p to build $\tilde{D}$ and then outputs $M(\tilde{D})$
Then M' is (ϵ', δ') -DP for $\epsilon' \approx p \cdot \epsilon$
 $\delta' = p \cdot \delta$

Let $p = \frac{m}{n}$ $\leftarrow$ batch size
 $\leftarrow$ # of training samples

Then DP-SGD is $O(\epsilon_p \sqrt{T \log(1/\delta)}, \epsilon_p T)$
DP

Even better composition?

Insight: advanced composition applies to aux
(ϵ, δ)-DP algorithm

• can we get a better analysis for
the subsampled Gaussian mechanism?

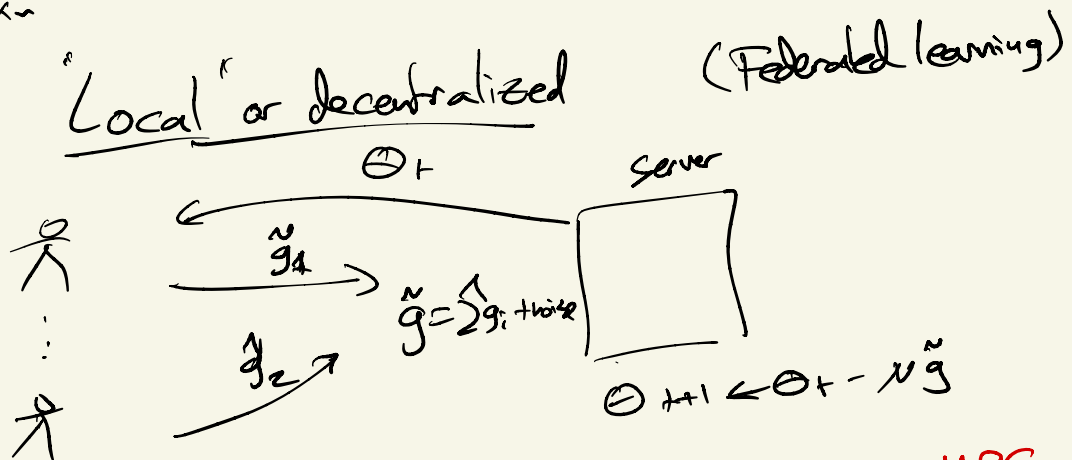
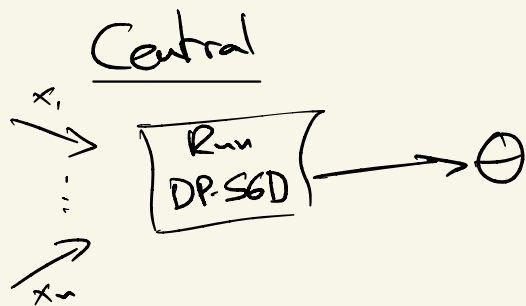
YES!

How? "Tedious" with (variants or generalizations
of DP: Rényi-DP, CDP, Gaussian-DP)

Result: $O(\epsilon_p \sqrt{T \log(1/\delta)}, \epsilon_p T)$
 $\uparrow$ save a $\log(1/\delta)$ term
 $\uparrow$ save a pT term

DP-SGD in practice

Local vs central mode of DP.



Q: . who computes the noisy sum?
 $\rightarrow$ MPC
 $\rightarrow$ Trusted server
 $\rightarrow$ multiple non-colluding servers

- how to do random subsampling?
 $\hookrightarrow$ what if clients aren't always online?
 $\hookrightarrow$ how to keep the subsample secret from the server(s)?