

ML and Opt Lect 19

- Backprop for MLPs
- Automatic differentiation & ML frameworks
- Autoencoders
- Example of using Pytorch to learn a 2-layer classifier for Fashion MNIST

Let's give the relevant backprop formulas for the MLP case

$$o^l = \sigma(a^l) = \sigma(\omega^l o^{l-1} + b^l) \quad \leftarrow \text{use this to relate } \nabla_{o^l} f \text{ to } \nabla_{\omega^l} f \text{ and } \nabla_{b^l} f$$
$$o^{l+1} = \sigma(a^{l+1}) = \sigma(\omega^{l+1} o^l + b^{l+1}) \quad \leftarrow \text{use this to relate } \nabla_{o^l} f \text{ to } \nabla_{o^{l+1}} f$$

We have, by the chain rule,

$$J_{o^{l+1}}(o^l) = \text{diag}(\sigma'(a^{l+1})) \cdot \omega^{l+1} \in \mathbb{R}^{n_{l+1} \times n_l}$$

$$J_{o^l}(b^l) = \text{diag}(\sigma'(a^l)) \in \mathbb{R}^{n_l \times n_l}$$

and $J_{o^l}(\omega^l) \in \mathbb{R}^{n_l \times (n_l \times n_{l-1})}$ satisfies

$$\begin{aligned} [J_{o^l}(\omega^l)]_{i,:} &= \left[\frac{\partial (o^l)_i}{\partial \omega^l} \right] \in \mathbb{R}^{n_l \times n_{l-1}} \\ &= \text{diag}(\sigma'(a^l)) e_i (o^{l-1})^T \end{aligned}$$

so we can show that for any vector $v \in \mathbb{R}^{n_d}$,

$$\mathcal{J}_{o_d}(\omega^d)^T v = \text{diag}(\sigma'(a^d)) v (o^{d-1})^T \in \mathbb{R}^{n_d \times n_{d-1}}$$

Putting these pieces together

Backprop for MLPs

$$\text{temp} \leftarrow \nabla_{o^L} f$$

$$\nabla_{b^L} f = \text{diag}(\sigma'(a^L)) \text{temp} + \lambda \nabla_{b^L} R_L$$

$$\nabla_{w^L} f = \text{diag}(\sigma'(a^L)) \text{temp} (o^{L-1})^T + \lambda \nabla_{w^L} R_L$$

for $l = L-1, \dots, 1$

$$\text{temp} \leftarrow (w^{l+1})^T \text{diag}(\sigma'(a^{l+1})) \text{temp}$$

$$\nabla_{b^l} f \leftarrow \text{diag}(\sigma'(a^l)) \text{temp} + \lambda \nabla_{b^l} R_l$$

$$\nabla_{w^l} f \leftarrow \text{diag}(\sigma'(a^l)) \text{temp} (o^{l-1})^T + \lambda \nabla_{w^l} R_l$$

Notes: (i) in practice we cache the preactivations a^l during the forward pass

(ii) cost of forward & backward passes similar

(iii) costs of passes dominated by lin algebra ops (mat vec prods)

BP is a core computational tool for working with NNs.

Become familiar with it — we will see that it gives rise to some important optimization phenomena when training NNs.

- We can implement BP for any computational DAG with (sub)differentiable neurons. In practice doing this by hand is tedious, error-prone, and requires careful optimizations for efficiency.
- Now we have available several high-quality ML frameworks whose main focus is to automate BP as much as possible: you specify how $f(x)$ is computed in a forward pass, then they deduce the backward pass using autodifferentiation. They also optimize the computations and support GPGPUs for much faster computation.

Common ML Frameworks:

- **Pytorch** (we use in this class, more popular now in research)

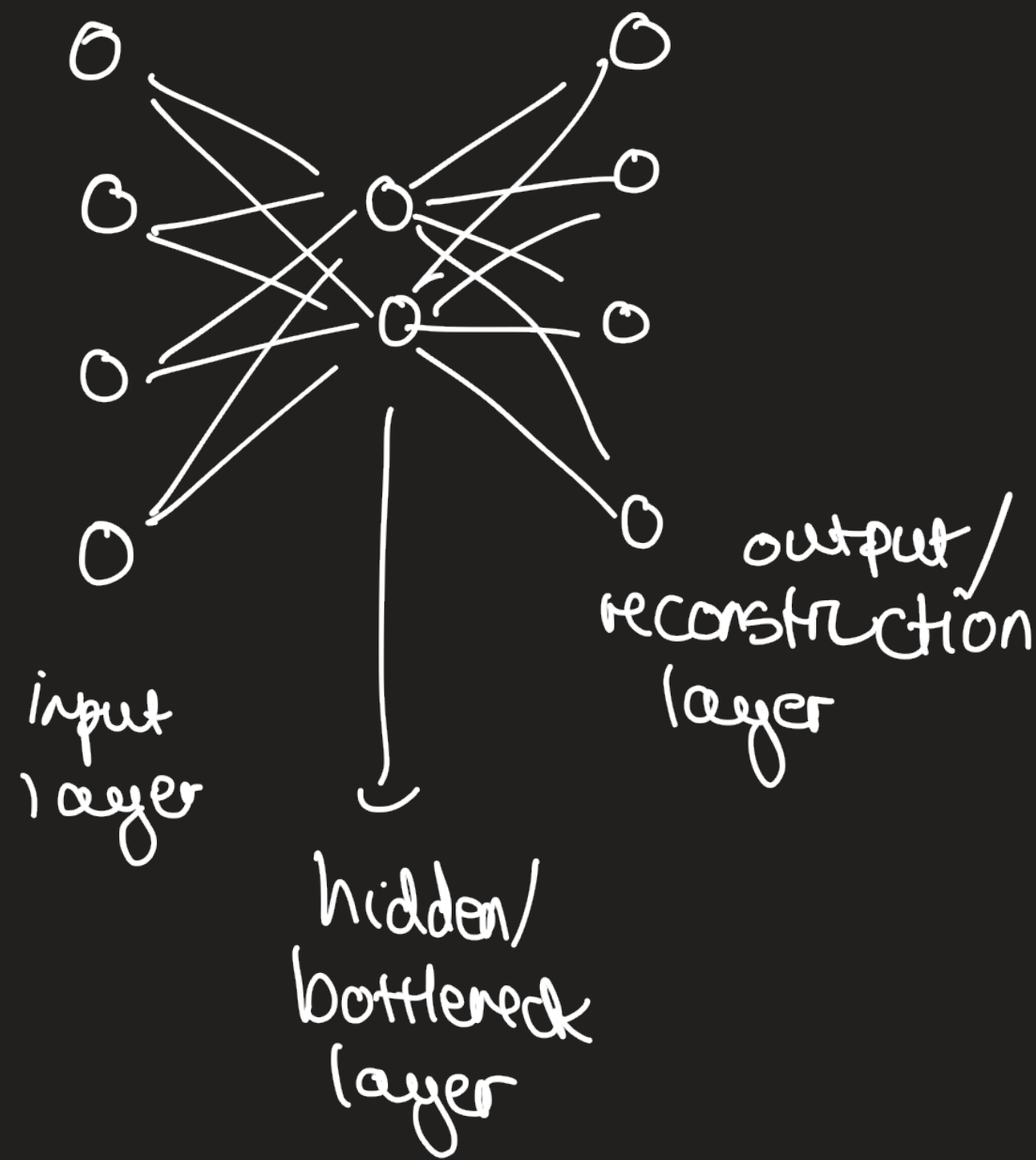
2016 Facebook Research

- **Tensorflow** (used the last time course was taught; more used in production)

2015 Google Brain

Autoencoders

- Nonlinear generalization of PCA & SVD
- Typically autoencoders are used to create useful features for a downstream task, in an unsupervised way



Goal: learn a (compressed) representation that does a good job of reconstructing the input

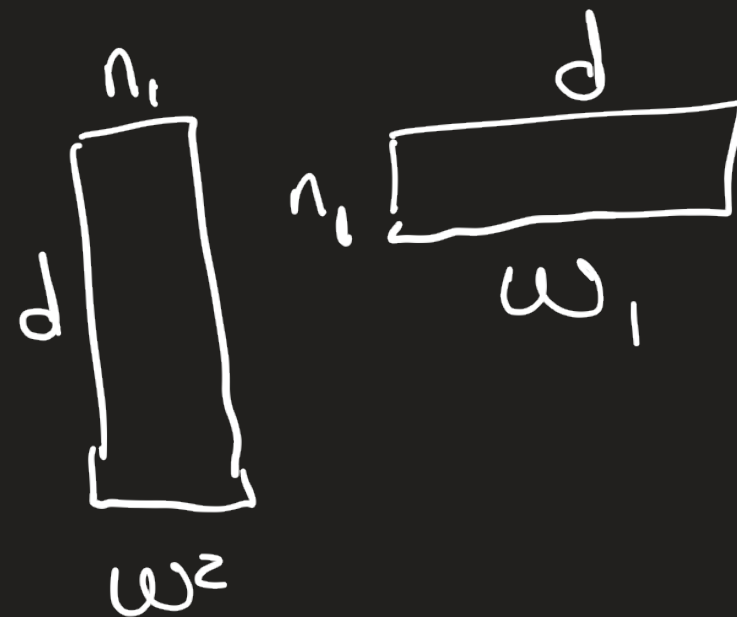
Ex ℓ_2 -loss $\ell(u, v) = \|u - v\|_2^2$
and $\sigma(x) = x$

$$\begin{aligned}\hat{y}(x) &= o^2 = \sigma(\omega^2 o' + b^2) \\ &= \sigma(\omega^2 \sigma(\omega' x + b') + b^2) \\ &= \omega^2 \omega' x + \omega^2 b' + b^2\end{aligned}$$

fix $b' = b^2 = 0$ then

$$\hat{y}(x) = \omega^2 \omega' x$$

where



so the autoencoder learns ω^2, ω'

to minimize

$$\begin{aligned}f(\omega', \omega^2) &= \frac{1}{n} \sum_{i=1}^n \|\hat{y}(x_i) - x_i\|_2^2 \\ &= \frac{1}{n} \sum_{i=1}^n \|\omega' \omega^2 x_i - x_i\|_2^2 = \frac{1}{n} \|\omega' \omega^2 X - X\|_F^2\end{aligned}$$

where $X = [x_1 | \dots | x_n]$

so fitting an autoencoder with $\ell(u, v) = \|u - v\|_2^2$
and $b^1 = b^2 = 0$ is equivalent to solving

$$\omega_*^2, \omega_*^1 = \underset{\substack{\omega^2 \in \mathbb{R}^{n \times d} \\ \omega^1 \in \mathbb{R}^{d \times n_1}}}{\operatorname{argmin}} \|\omega^2 \omega^1 X - X\|_F^2$$

this is the SVD!

Clearly autoencoders are a more flexible method of
dim reduction than SVD because one can take a
nonlinear σ and nonzero biases.

Deep autoencoders :

