

ML and Optimization Lecture 21

- Convolutional Neural Networks continued

- issues w/ deep NNs:

- overfitting

- vanishing & exploding gradients

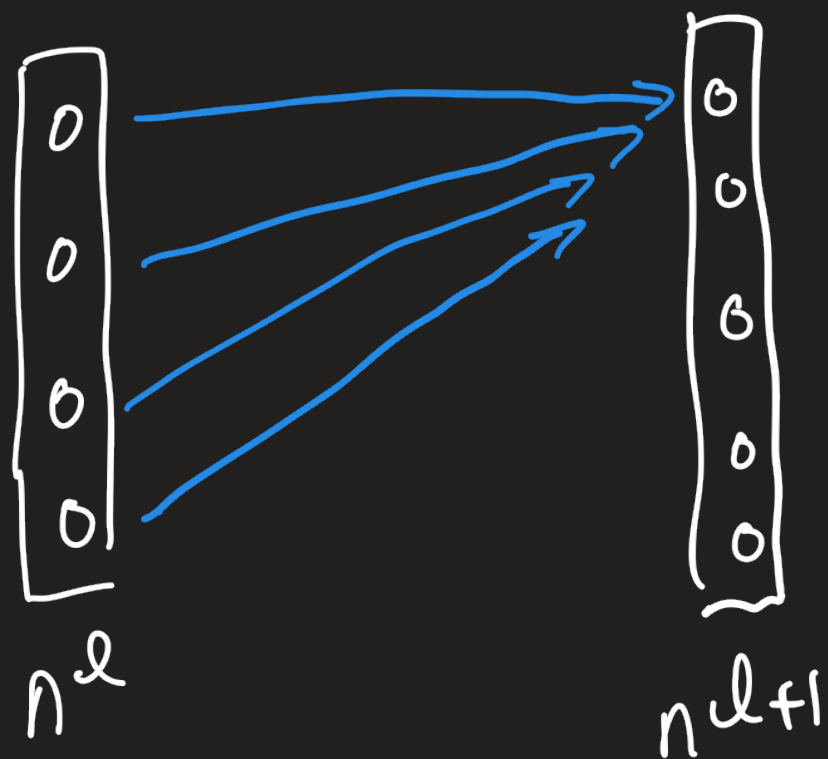
remedies:

- dropout (2014)

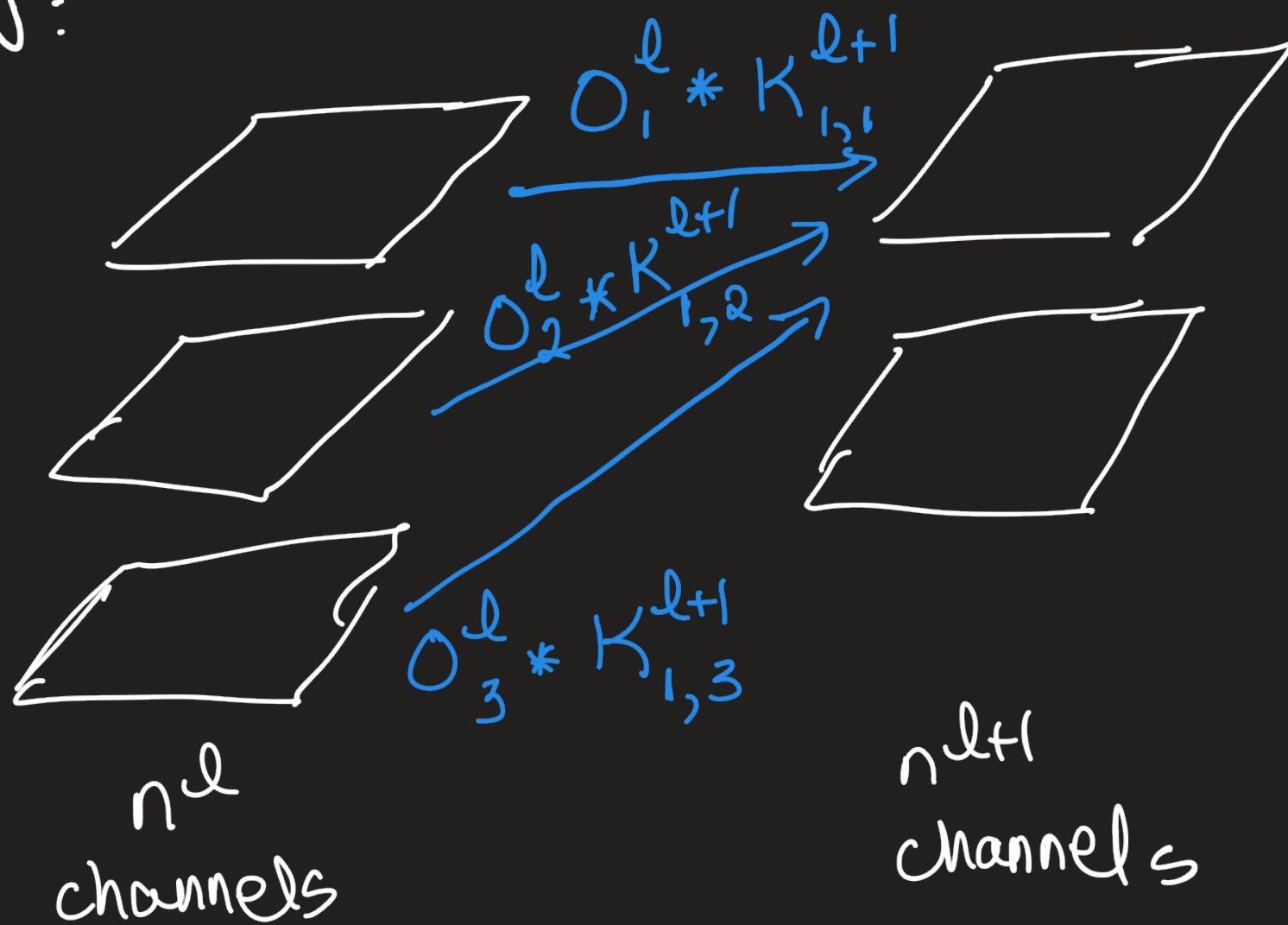
- (2015) batch normalization

Multi-channel CNNs

MLP:



CNN:



$$A_i^{l+1} = O_1^l * K_{1,i}^{l+1} + O_2^l * K_{2,i}^{l+1} + O_3^l * K_{3,i}^{l+1} + b_i^{l+1}$$
$$O_i^{l+1} = \sigma(A_i^{l+1})$$

In general if layer l has n_l channels, then

$$A_i^{l+1} = \left(\sum_{j=1}^{n_l} O_j^l * K_{i,j}^{l+1} \right) + b_i^{l+1}$$

$$O_i^{l+1} = \sigma(A_i^{l+1})$$

for $i=1, \dots, n_{l+1}$

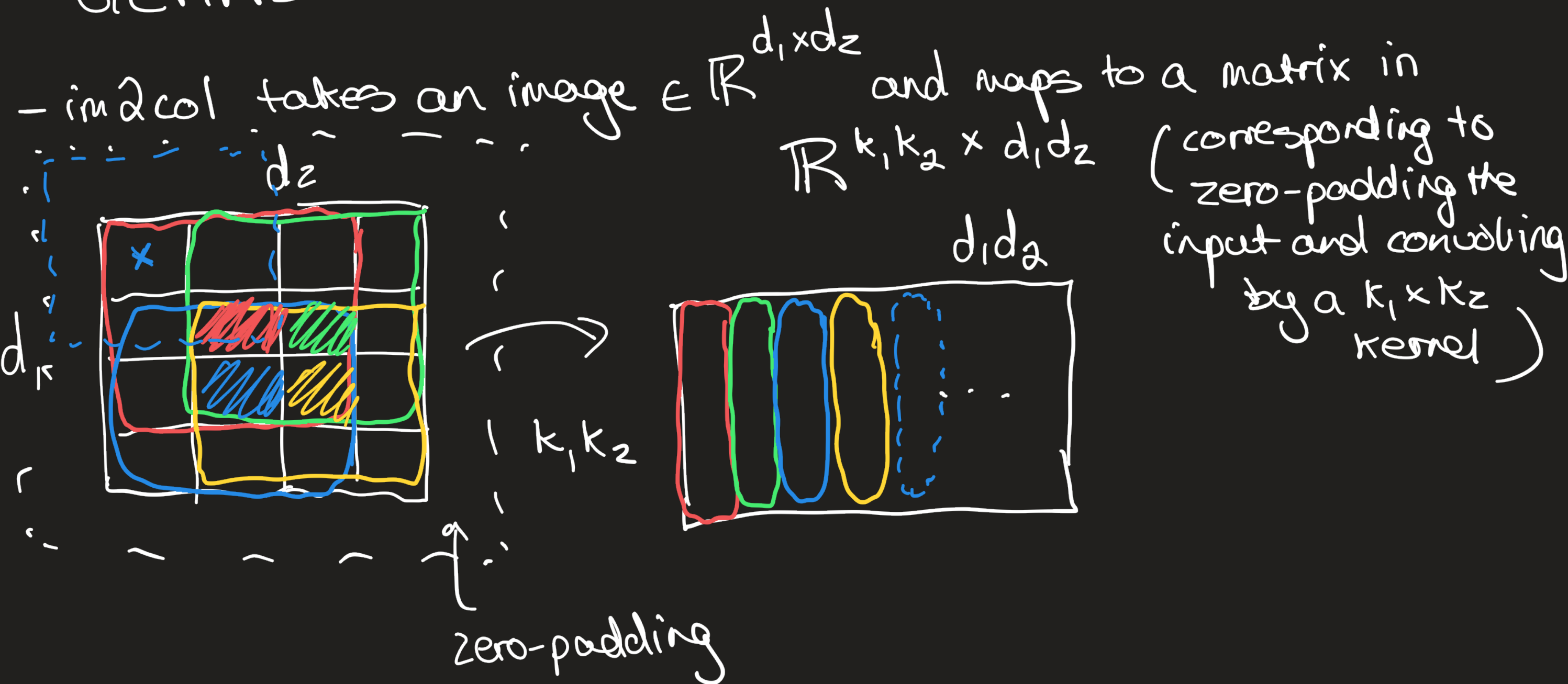
What about backprop?

- convolutions

Convolutions in practice (GEMM)

cf. Manas Sahni "Anatomy of a High-Speed Convolution"

reduce convolutions to the im2col operation and GEMMs



Then note that

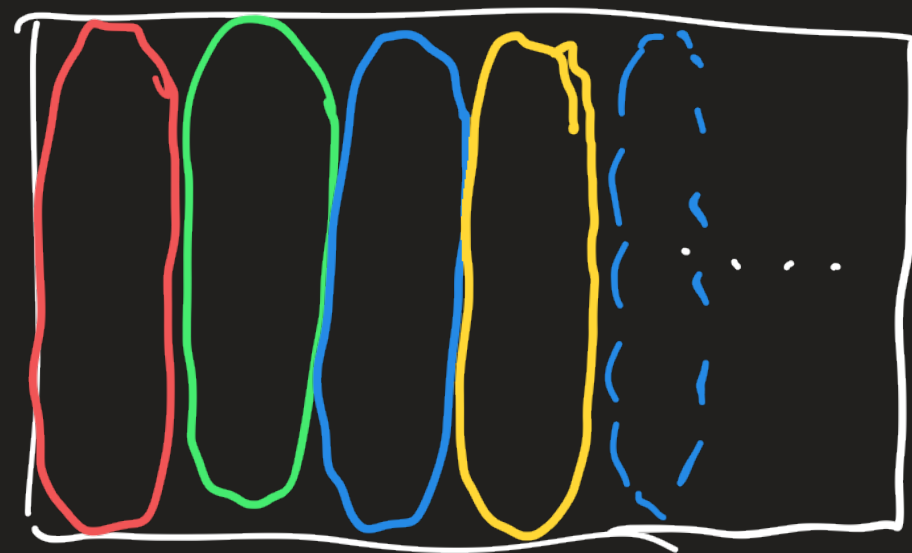
$$\text{vec}(I * K) = \text{vec}(K) \cdot \text{indcol}(I)$$

1	2	3
4	5	6
7	8	9

K



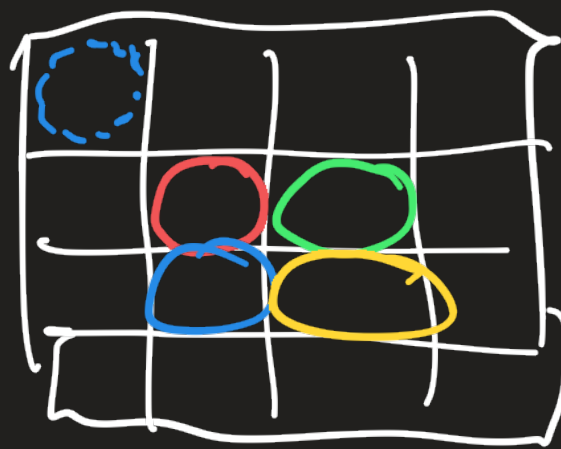
$\text{vec}(K)$



$\text{vec}(K) \cdot \text{indcol}(I)$



unvec



Issues with deep NNs (not just CNNs):

- overfitting (too much capacity for the amount of training data)
- vanishing & exploding gradients $\Rightarrow$ slow learning
- hyperparameter selection
 - optimal kernel sizes
 - # of channels per layer
 - # of layers
 - pooling types & locations
 - stride sizes
 - learning rates (function of minibatch size)
 - weight decay amount
 - more ...

Vanishing & Exploding Gradients

Phenomenon that

$$\|\nabla_{\mathbf{w}^l} f\|_2 \xrightarrow{l \rightarrow 1} \begin{cases} 0 & \text{vanishing gradients} \\ \infty & \text{exploding gradients} \end{cases}$$

Why? Because of the chain rule. E.g. for MLPs

$$\nabla_{\mathbf{O}^l} f = \text{diag}(\sigma'(a^{l+1})) (\mathbf{W}^{l+1})^T \nabla_{\mathbf{O}^{l+1}} f$$

Two considerations)

1) The saturation of our activation function



If preactivations are far from origin, $\sigma'(a^{l+1})$ close to zero, so

$$\|\nabla_{\mathbf{O}^l} f\|_2 \ll \|\nabla_{\mathbf{O}^{l+1}} f\|_2$$

2) The norm of our weight matrix W^{l+1}

— if on average $\|W^{l+1}x\|_2 \geq \|x\|_2$

then we get $\|\nabla_{O^l} f\|_2 \geq \|\nabla_{O^{l+1}} f\|_2$

— if on average $\|W^{l+1}x\|_2 \leq \|x\|_2$

then we get

$$\|\nabla_{O^l} f\|_2 \leq \|\nabla_{O^{l+1}} f\|_2$$

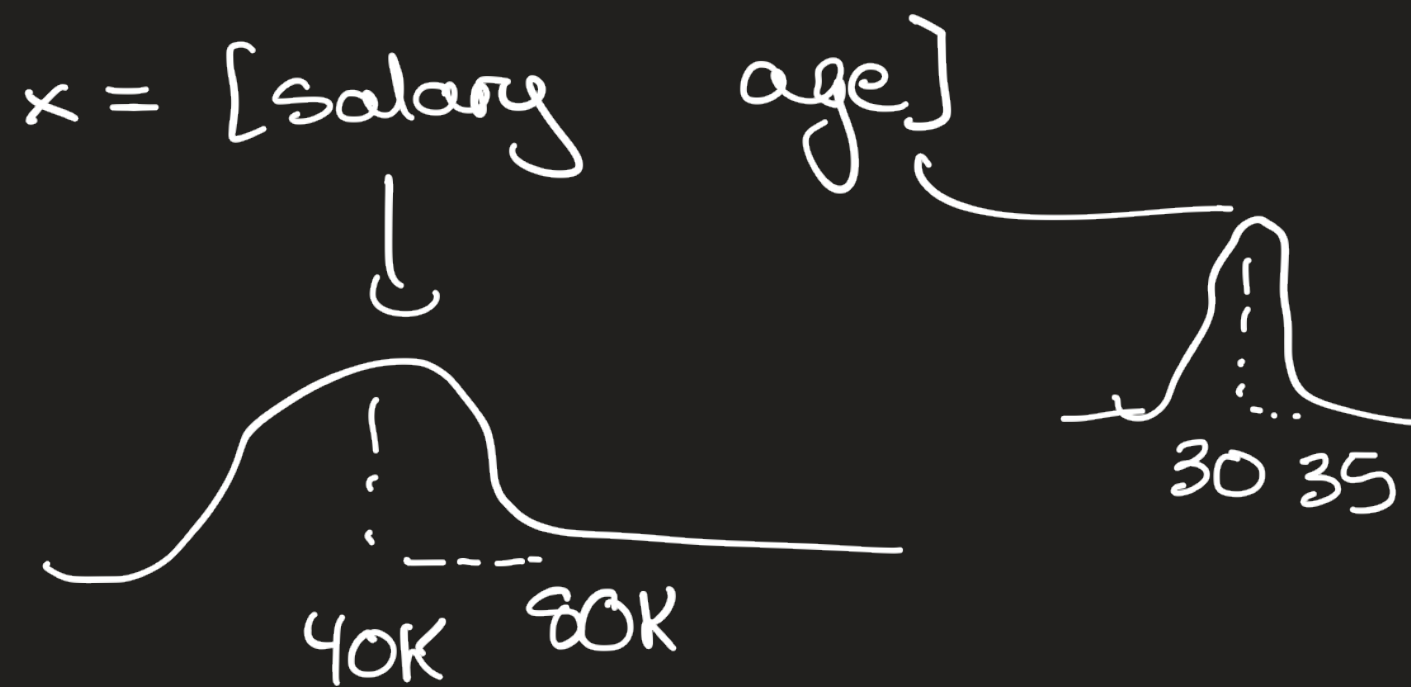
Consequence: naively training deep NN architectures either fails or gives poor performance.

Remedy: Batch Normalization (2015)

Idea: the feature distributions at each layer vary wildly, so normalize them so they look like standard gaussians



In convex ML we saw the advantage of normalizing our input data so the features are on the same scale
ex: training to predict credit default probability

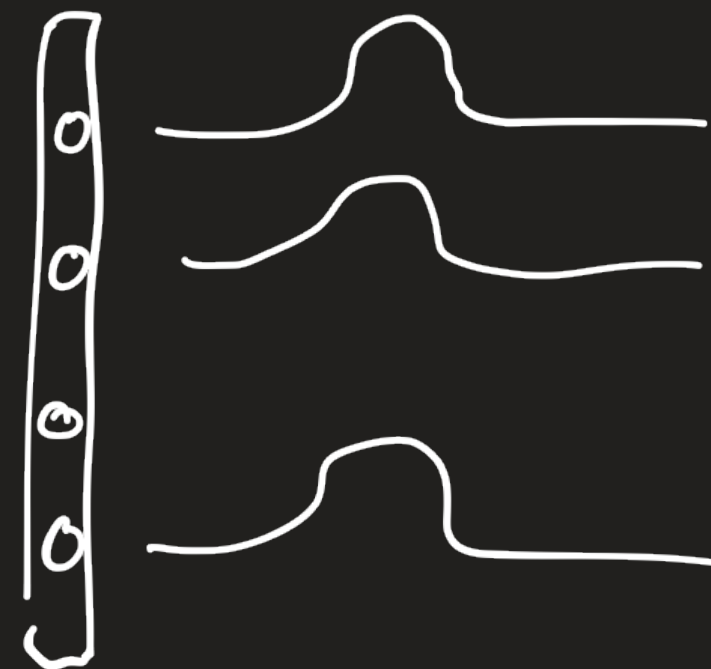
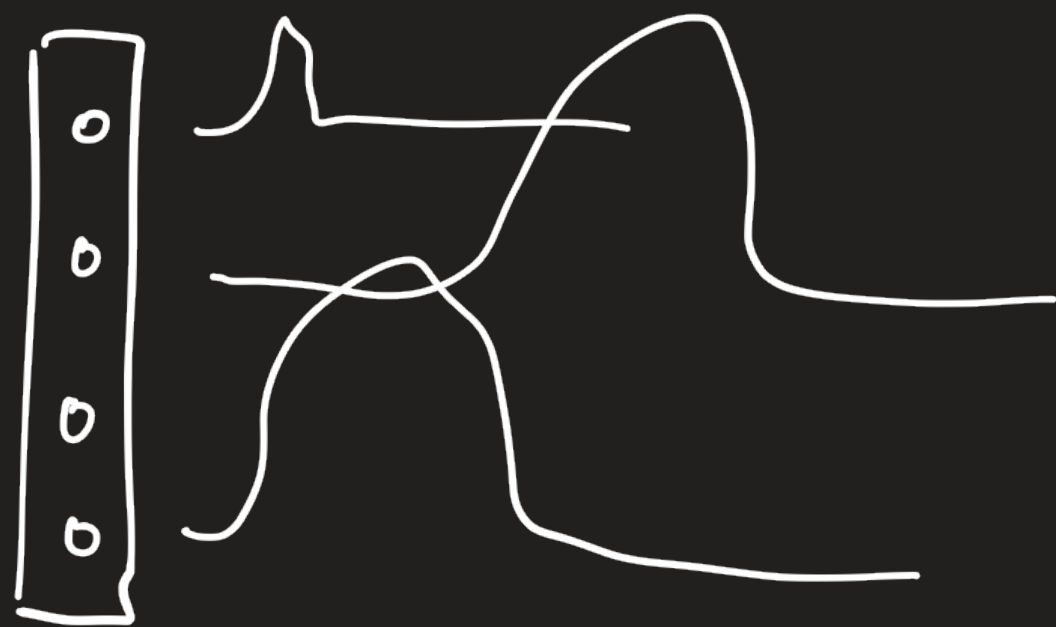


to get the features on the same scale we normalized them

$$x \mapsto \left[\frac{\text{salary} - \mu_{\text{salary}}}{\sigma_{\text{salary}}}, \frac{\text{age} - \mu_{\text{age}}}{\sigma_{\text{age}}} \right]$$

The figure shows two normalized normal distribution curves. The first curve, corresponding to the salary feature, has a peak at 0 and a dashed vertical line at 2. The second curve, corresponding to the age feature, also has a peak at 0 and a dashed vertical line at 2.

Idea of BN: add normalization layer after a regular layer to rescale and center the feature distributions



BN
layer

BN layer

Between layer l and $l-1$ add a BN layer:

$$a^l = \text{BN}_{\gamma^l, \beta^l}(\omega^l o^{l-1})$$

$$o^l = \sigma(a^l)$$

$$\text{BN}_{\gamma^l, \beta^l}(o^{l-1}) = \gamma^l \odot \frac{o^{l-1} - \mu_B}{\sqrt{\sigma_B^2 + \epsilon}} + \beta^l$$

μ_B = minibatch mean
$$= \frac{1}{m} \sum_{i=1}^m (o^{l-1})_i$$

σ_B^2 = minibatch variance =
$$\frac{1}{m} \sum_{i=1}^m ((o^{l-1})_i - \mu_B)^2$$

where $\gamma^l \in \mathbb{R}^{n_{l-1}}$ is a scale vector and $\beta^l \in \mathbb{R}^{n_{l-1}}$ is a shift vector

are useful to ensure we are in the nonlinear region of the activation function

Forward & Backward pass during training is straightforward b/c μ_B and σ_B^2 are computed using minibatch means (`model.train()` in pytorch)

What about inference time (`model.eval`)

1) After fitting, run all our training data through and compute

$$\mu^l = \frac{1}{n_{\text{train}}} \sum_{i=1}^{n_{\text{train}}} (o^{l-1})_i \in \mathbb{R}^{n_{l-1}}$$

$$\sigma^l = \frac{1}{n_{\text{train}}} \sum_{i=1}^{n_{\text{train}}} ((o^{l-1})_i - \mu^l)^2 \in \mathbb{R}^{n_{l-1}}$$

and use these

2) or maintain moving average estimates of population mean & variance

$$\mu^l \leftarrow \beta_1 \mu^l + (1 - \beta_1) \mu_B^l \quad \text{for each minibatch}$$

$$\sigma^l \leftarrow \beta_2 \sigma^l + (1 - \beta_2) \sigma_B^l \quad \text{for each minibatch}$$

and use μ^l and σ^l at inference time

