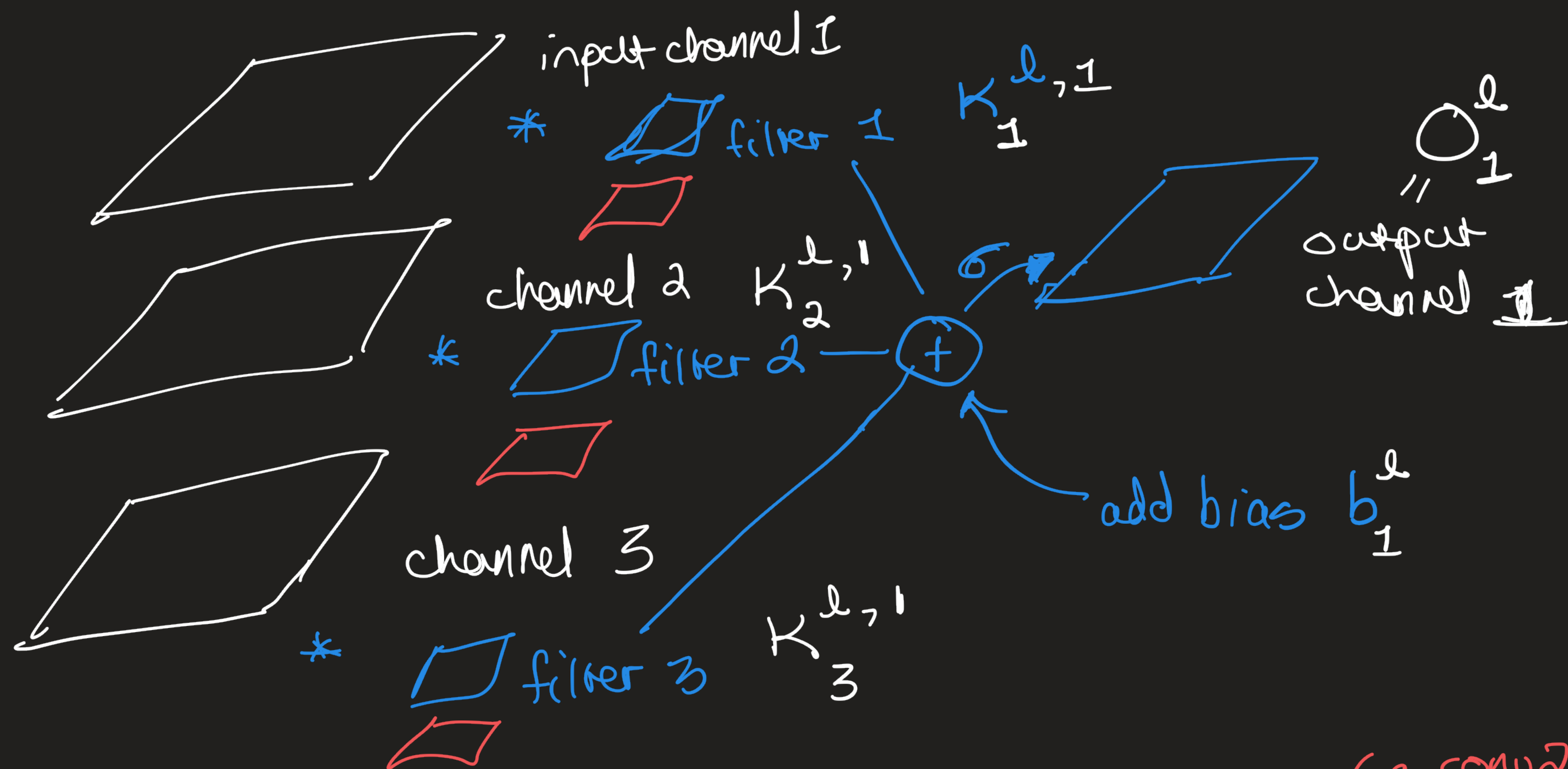


Today:

CNNs

- working w/ multi-channel images
- backprop for CNNs
- issues w/ backprops: vanishing & exploding gradients
- remedies:
 - pretraining
 - (2014) dropout
 - (2015) batch normalization
 - modify architecture of the network
- common CNN architectures:
 - VGG (repeating structures)
 - GoogLeNet (inception architecture)
 - ResNets (residual blocks)

Multi-channel inputs to convolutional layers?



ex: consider 3 channel input w/ 32 channel output (a conv2D layer w/ 32 filters)
and each filter is $3 \times 3 \Rightarrow$ you have

$$\underbrace{3 \times 3}_{\text{filter size}} \times \underbrace{32}_{\text{\# outputs}} \times \underbrace{3}_{\text{inputs}} + 32^{\text{\# outputs (bias)}} = 896 \text{ parameters for this layer}$$

Say we have n^l channels on convolutional layer l and for each channel we have kernels

$K_j^{l,c}$ where $j = 1, \dots, n_{l-1}$ and $c = 1, \dots, n_l$

and we write O_j^l as the j th output channel on layer l

A_j^l as the preactivation of the j th channel on layer l

$$\underline{A_c^l} = \left[\sum_{t=1}^{n_{l-1}} O_t^{l-1} * K_t^{l,c} + b_c^l \right]$$

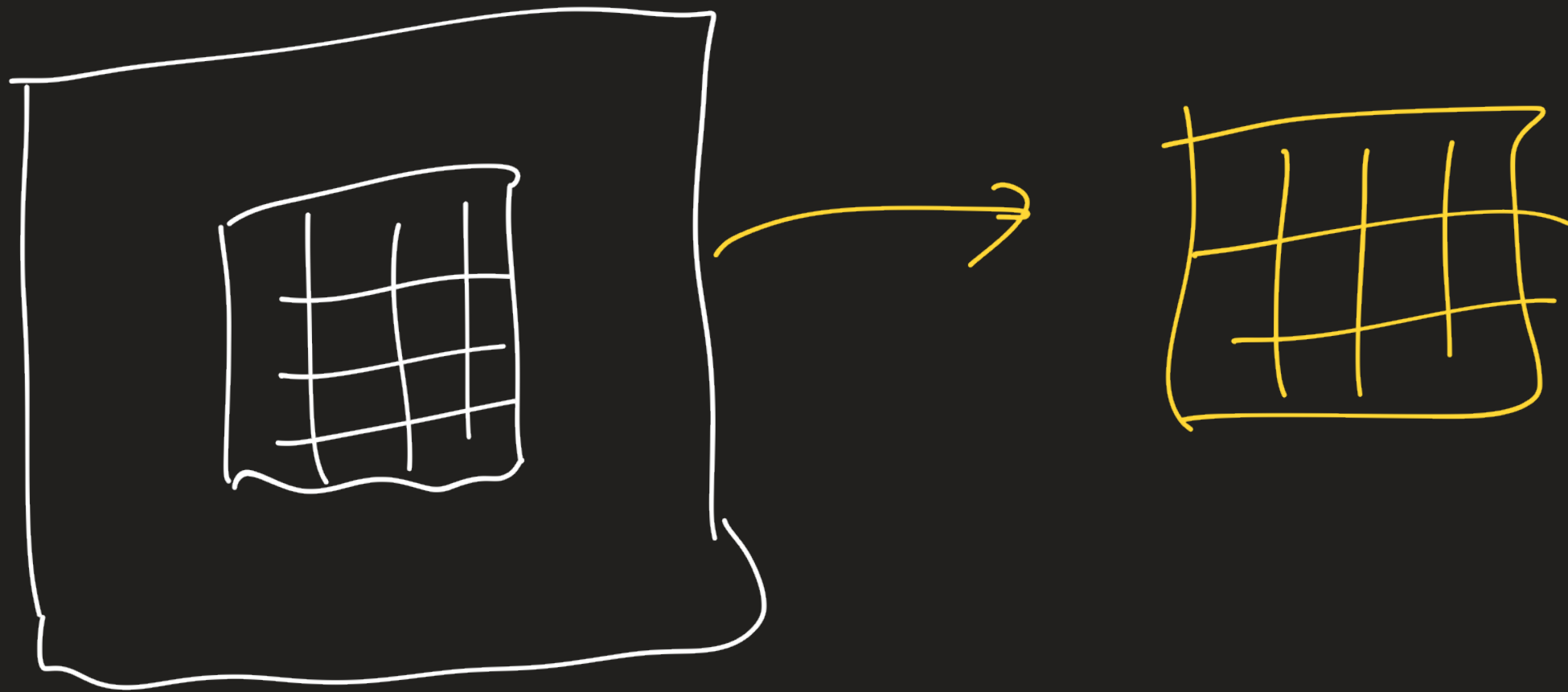
↑
bias for the j th channel on layer l

$$O_c^l = \sigma(A_c^l) \quad \text{for each output channel } c \text{ on layer } l$$

Backprop for CNNs

Now we want to learn $\frac{\partial f}{\partial K_t^{e,c}}$, $\frac{\partial f}{\partial b_c^l}$

Difficulty:



Vanishing & Exploding Gradients

Phenomenon that arises in deep neural networks: $L-1$

$$\|\nabla_{w^L} f\|_2 \rightarrow \begin{cases} 0 & \text{vanishing gradients w.r.t. } L-1 \\ \infty & \text{exploding gradients w.r.t. } L-1 \end{cases}$$

Why? Because of the chain rule:

$$\nabla_{o^L} f = \left[w^{L+1} \right]^T \cdot \left[\nabla_{o^{L+1}} f \odot \sigma'(a^{L+1}) \right]$$

and

$$\nabla_{w^L} f = \text{diag}(\sigma'(a^L)) (\nabla_{o^L} f) \cdot [o^{L-1}]^T$$

$$\nabla_{b^L} f = \nabla_{o^L} f \odot \sigma'(a^L)$$

Two considerations:

- 1) The saturation of our activation function, given by $\sigma'(a^{l+1})$ and $\sigma'(a^l)$



- 2) The norm of our weight matrix W^{l+1}
- if, on average, $\|W^{l+1}x\|_2 \geq \|x\|_2$
then we will get exploding gradients
 - if, on average, $\|W^{l+1}x\|_2 \leq \|x\|_2$
then we will get vanishing gradients

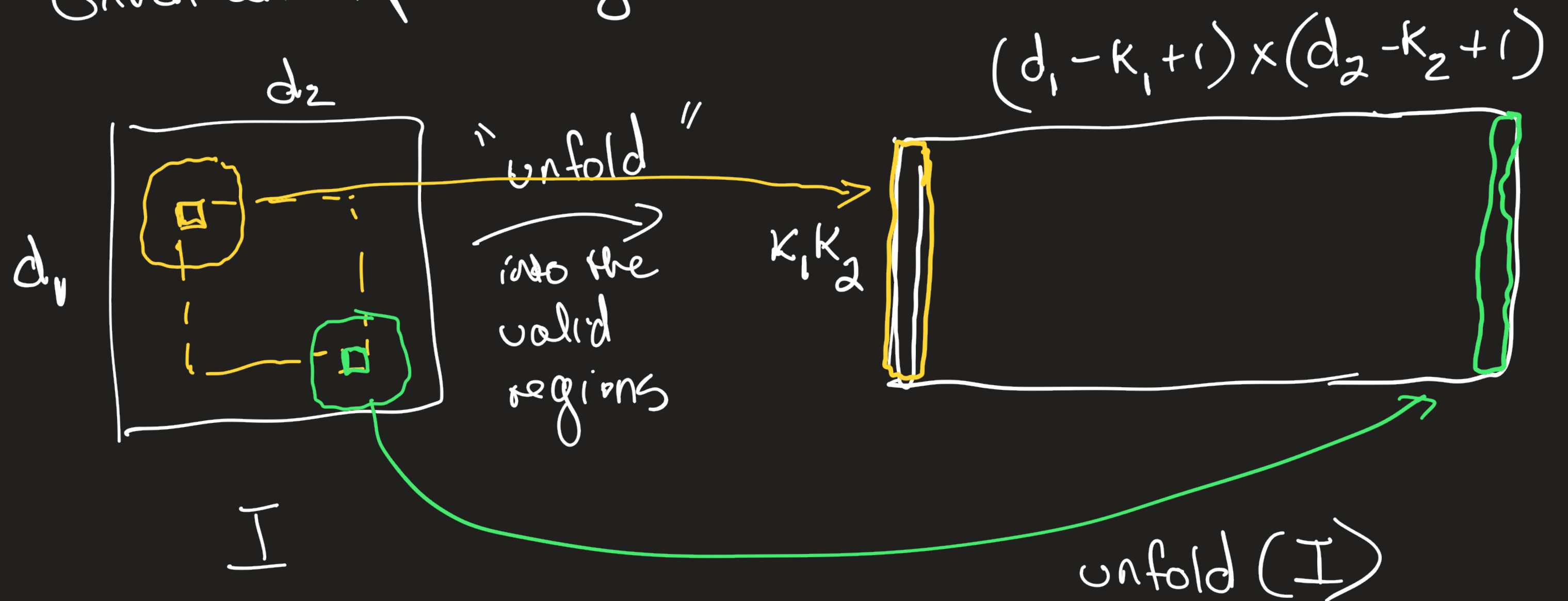
Vanishing & Exploding gradients are difficult to handle w/ first order optimization algorithms.

- the scales of the gradients can vary wildly b/w layers

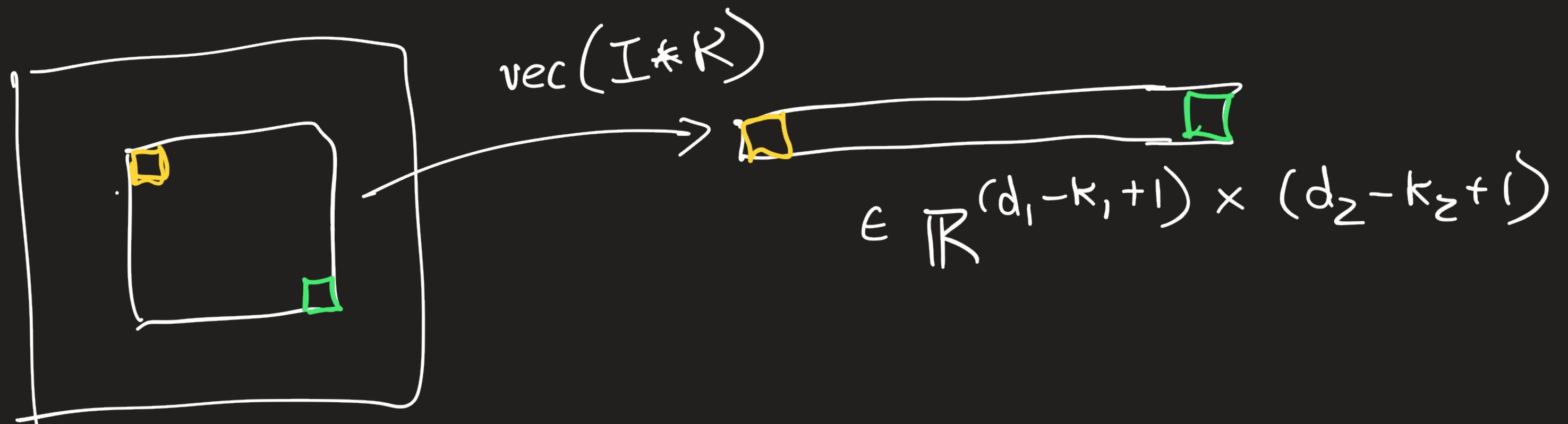
- Next class: remedies for this problem

Basic idea is to note that we can represent our images however we want (and the operations on them)

Given an input image I note



$$\text{vec}(I \star K) \in \mathbb{R}^{(d_1 - k_1 + 1) \times (d_2 - k_2 + 1)}$$



Question: What is the relationship b/w $\text{vec}(I \star K)$ and $\text{unfold}(I)$?

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

Diagram illustrating the vectorization of the convolution operation $I * K$. The left side shows a vector of size $\mathbb{R}^{(d_1 - k_1 + 1) \times (d_2 - k_2 + 1)}$ with a yellow circle at the start and a green circle at the end. The right side shows a vector of size $\mathbb{R}^{k_1 \times k_2}$ multiplied by a matrix of size $\mathbb{R}^{k_1 \times k_2 \times [(d_1 - k_1 + 1), (d_2 - k_2 + 1)]}$. The matrix is represented by a rectangle with a yellow vertical strip on the left and a green vertical strip on the right.

Consequence:

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

$$\Rightarrow I * K = \text{unvec}(\text{vec}(K) \cdot \text{unfold}(I))$$

this is a vector, and we know how to backprop w.r.t. param vectors