
Convolutional Neural Networks (CNNs)

for image classification

- **convolutional layers**, stride, *à trous*
- **pooling** (max and average)
- **fully connected layers**
- **data augmentation**
- **class activation map (CAM)**

Towards Convolutional Neural Networks

Remember (slide from Topic 3):

..... **convolution** operation is defined

$$g[i, j] = \sum_{u=-k}^k \sum_{v=-k}^k h[u, v] \cdot f[i - u, j - v]$$

It is written as:

$$g = h * f = \sum_{u=-k}^k \sum_{v=-k}^k h[-u, -v] \cdot f[i + u, j + v]$$

- You should also remember that convolution is a **linear operation**
Thus, it can be written as $g = \mathbf{W}_h f$
- CNNs use **convolutions as very sparse linear transformations**.
- In the context of (large) images, such NN design is motivated by **efficiency** and **neighborhood processing** - **we will learn filters**

Early Work on CNNs

Fukushima (1980) – Neo-Cognitron

LeCun (1998) – Convolutional Networks (ConvNets)

- similarities to Neo-Cognitron
- success on character recognition

Other attempts at deeply layered Networks trained with backpropagation

- not much success (e.g. very slow, diffusing/vanishing gradient)

Lately - significant training improvements

- various tricks (batch normalization, drop-outs, residual links, etc.)

Convolutional Network: Motivation

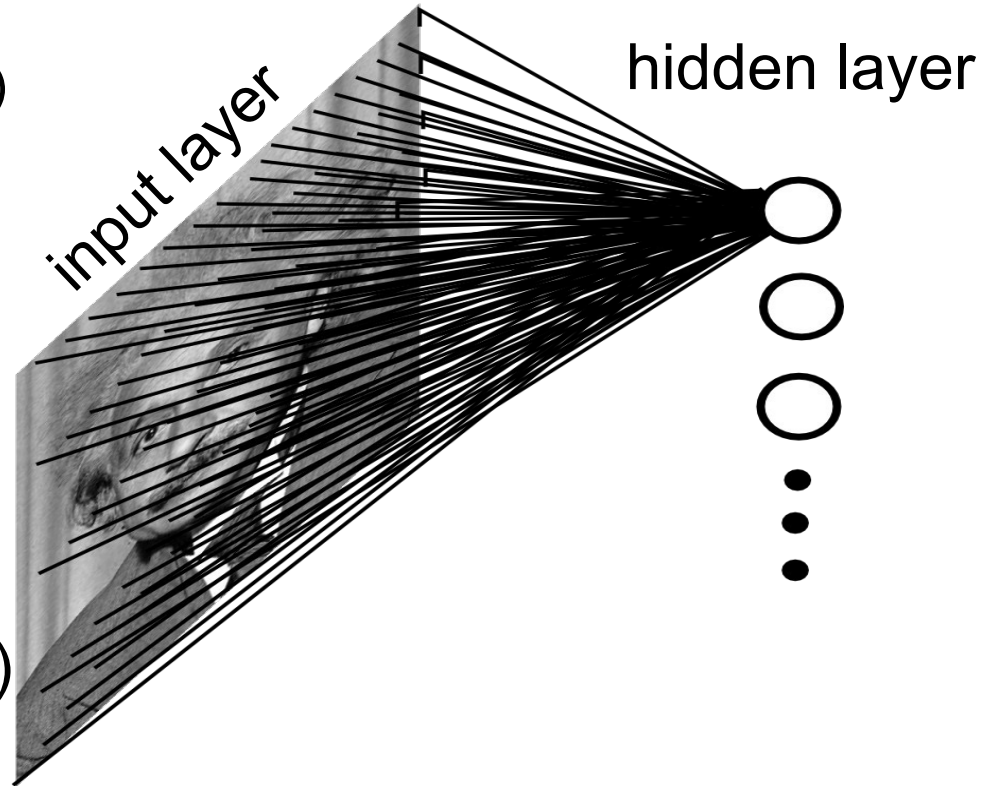
Consider a **fully connected network** (most weights $W[i,j] \neq 0$)

Example: 200 by 200 image,
 4×10^4 connections to one
hidden unit

For 10^5 hidden units $\rightarrow 4 \times 10^9$
connections

But distant pixels are unrelated
(correlations are mostly local)

**Do not waste resources by
connecting unrelated pixels**



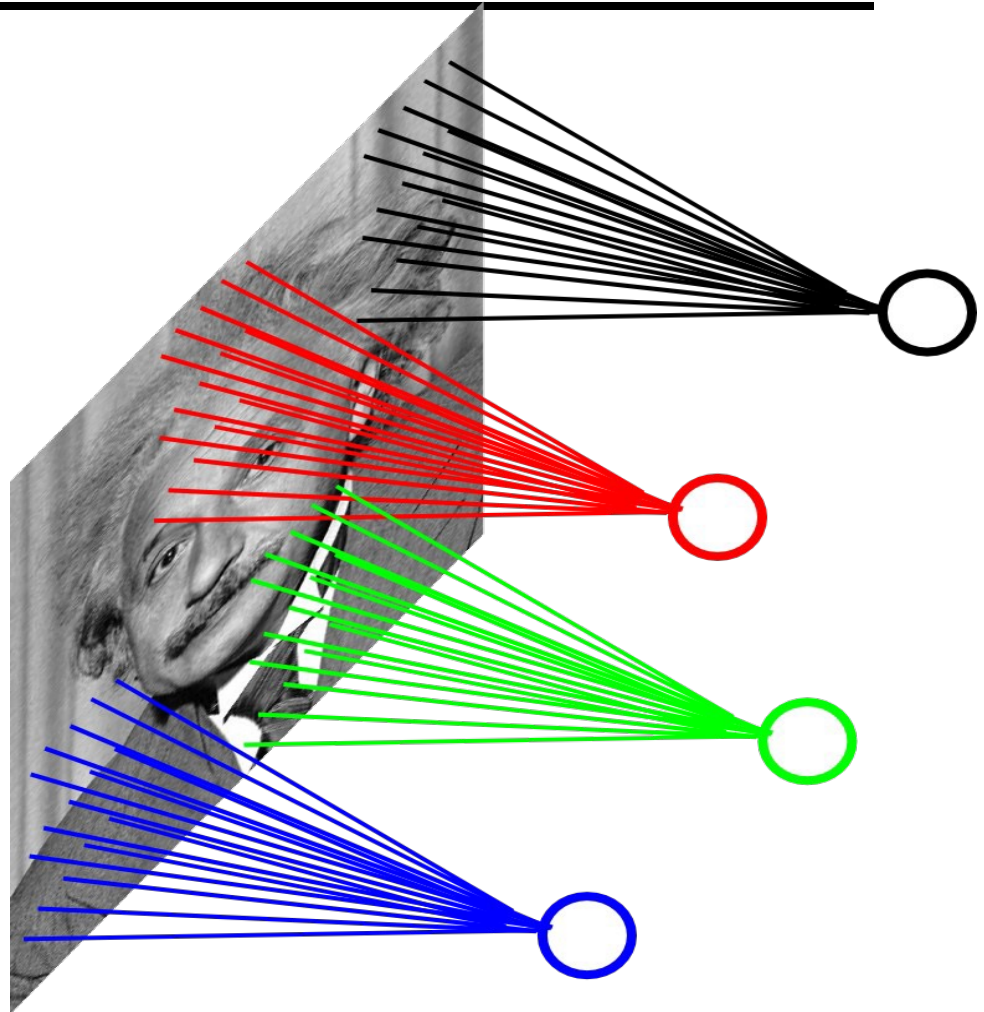
Convolutional Network: Motivation

Connect only pixels in a local patch, say 10×10

For 200 by 200 image, 10^2 connections to one hidden unit

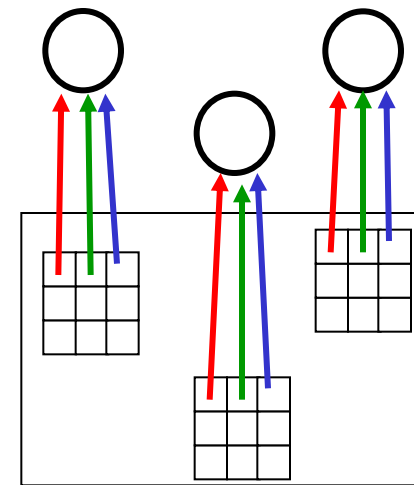
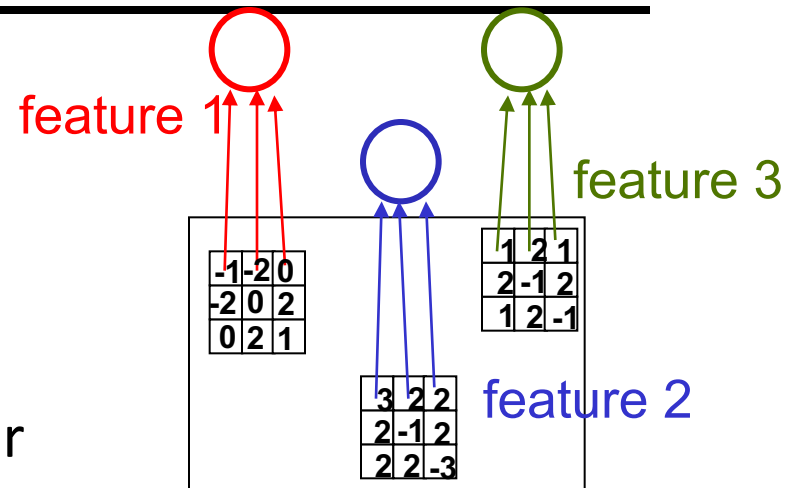
For 10^5 hidden units $\rightarrow 10^7$ connections

- contrast with 4×10^9 for fully connected layer
- factor of 400 decrease



Convolutional Network: Motivation

- Intuitively, each neuron learns a good feature (a filter) in one particular location
- If a feature is useful in one image location, it should be useful in all other locations
 - *stationarity*: statistics is similar at different locations
- Idea: make all neurons detect the **same feature at different positions**
 - i.e. **share parameters** (network weights) across different locations
 - greatly reduces the number of tunable parameters to learn



red connections have equal weight
green connections have equal weight
blue connections have equal weight

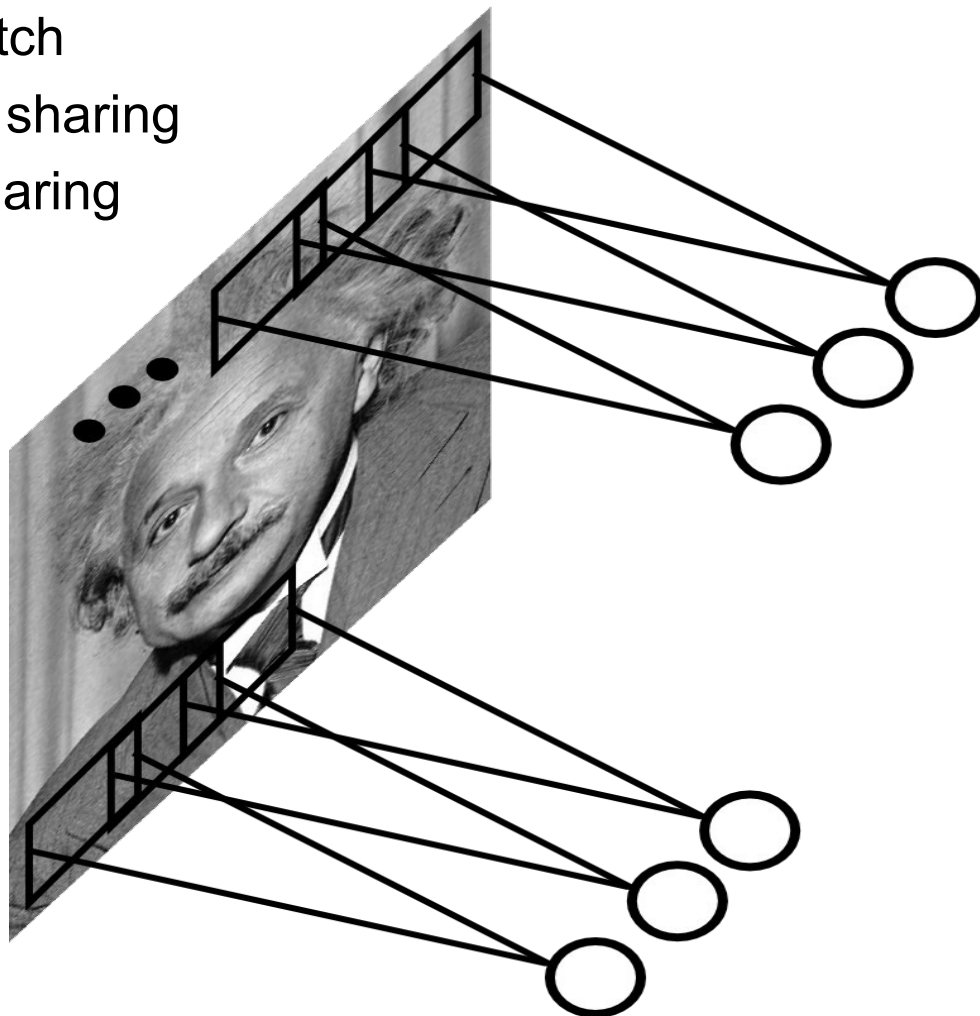
ConvNets: Weight Sharing

Much fewer parameters to learn

For 10^5 hidden units and 10×10 patch

- 10^7 parameters to learn without sharing
- 10^2 parameters to learn with sharing

Does not depend
on the number of
hidden units

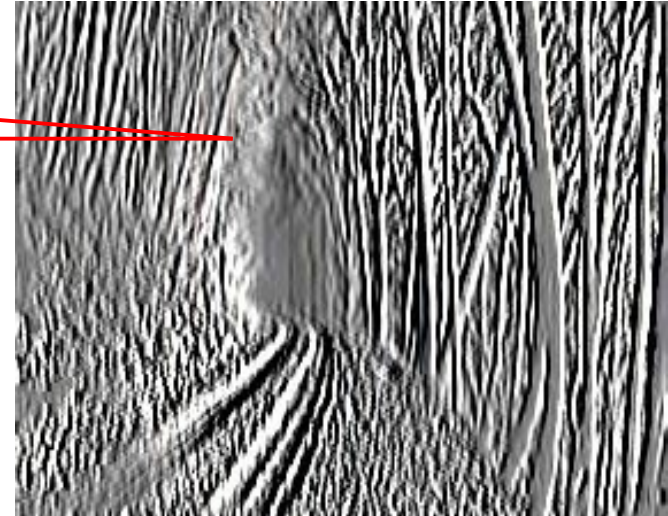


Filtering via Convolution Recap

Recall filtering with convolution for feature extraction



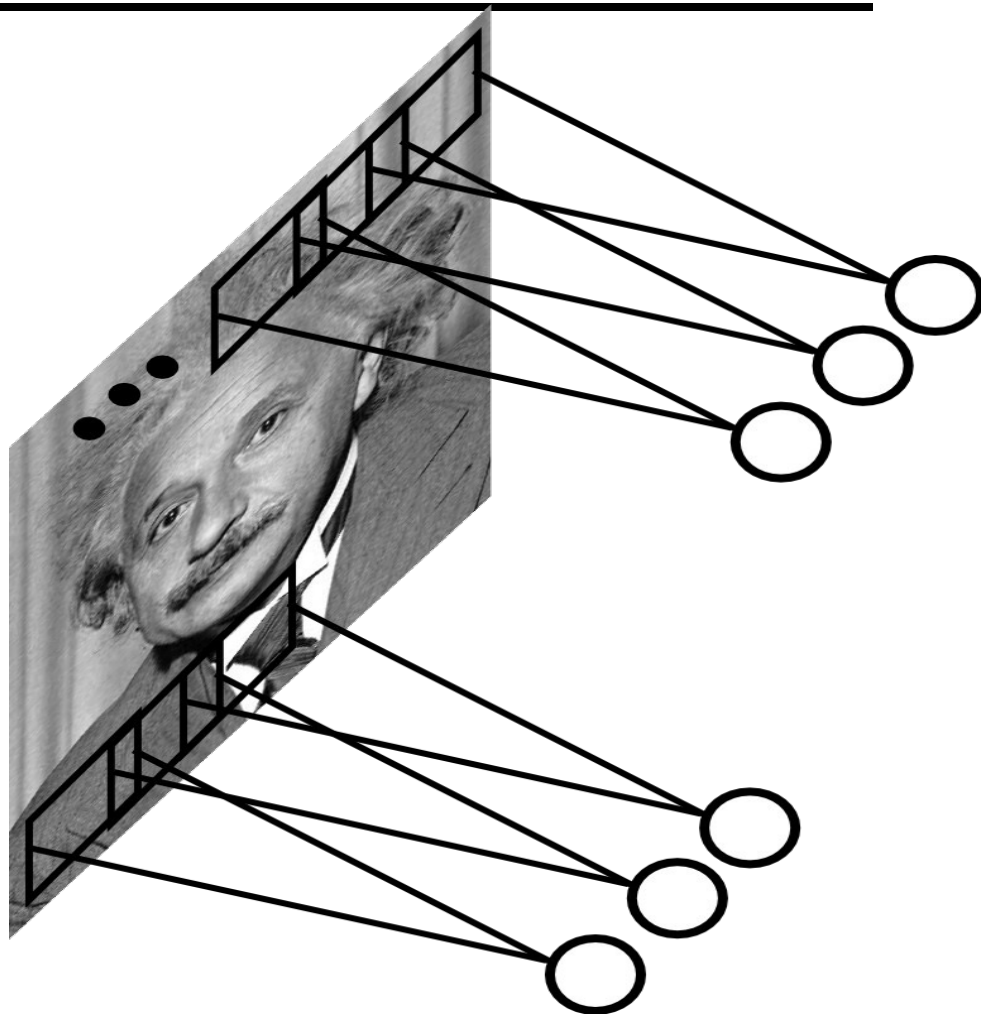
$$\begin{matrix} * & \begin{bmatrix} -1 & 0 & 1 \\ -1 & 0 & 1 \\ -1 & 0 & 1 \end{bmatrix} & = \end{matrix}$$



Convolutional Layer

Same as convolution with
some fixed filter

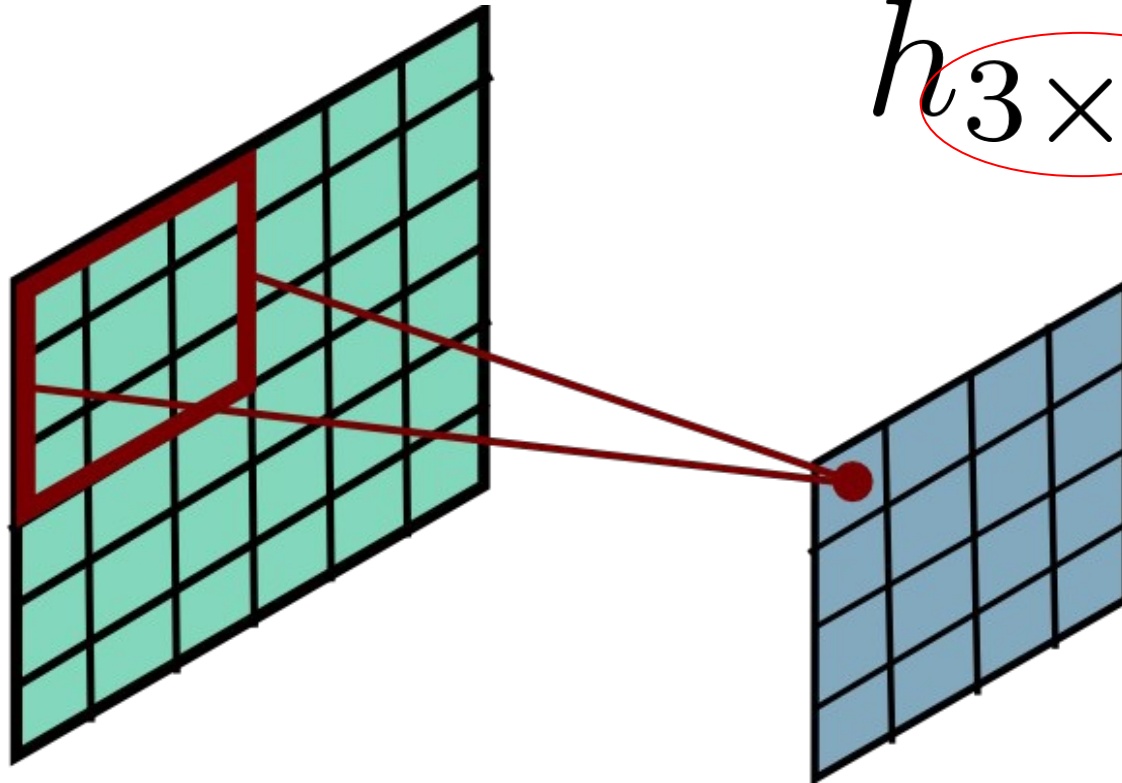
But here the filter
parameters
will be learned



Convolutional Layer

convolution kernel

h 3×3 size



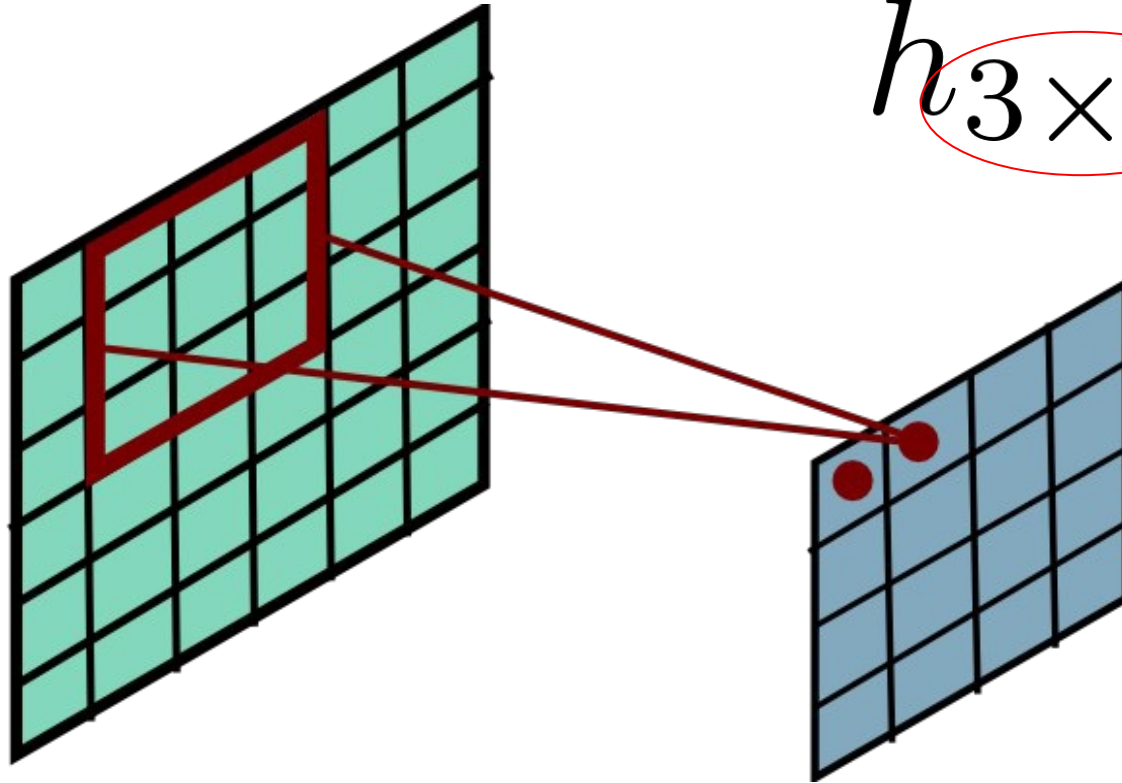
input

output

Convolutional Layer

convolution kernel

h 3×3 size



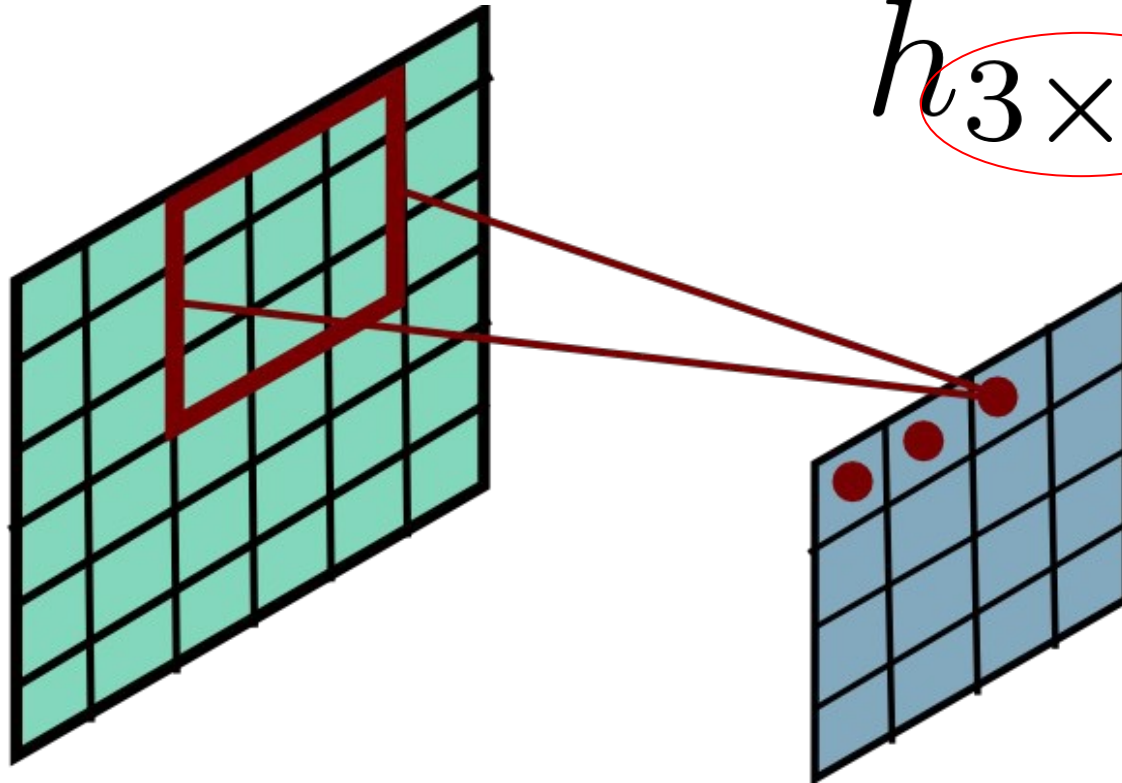
input

output

Convolutional Layer

convolution kernel

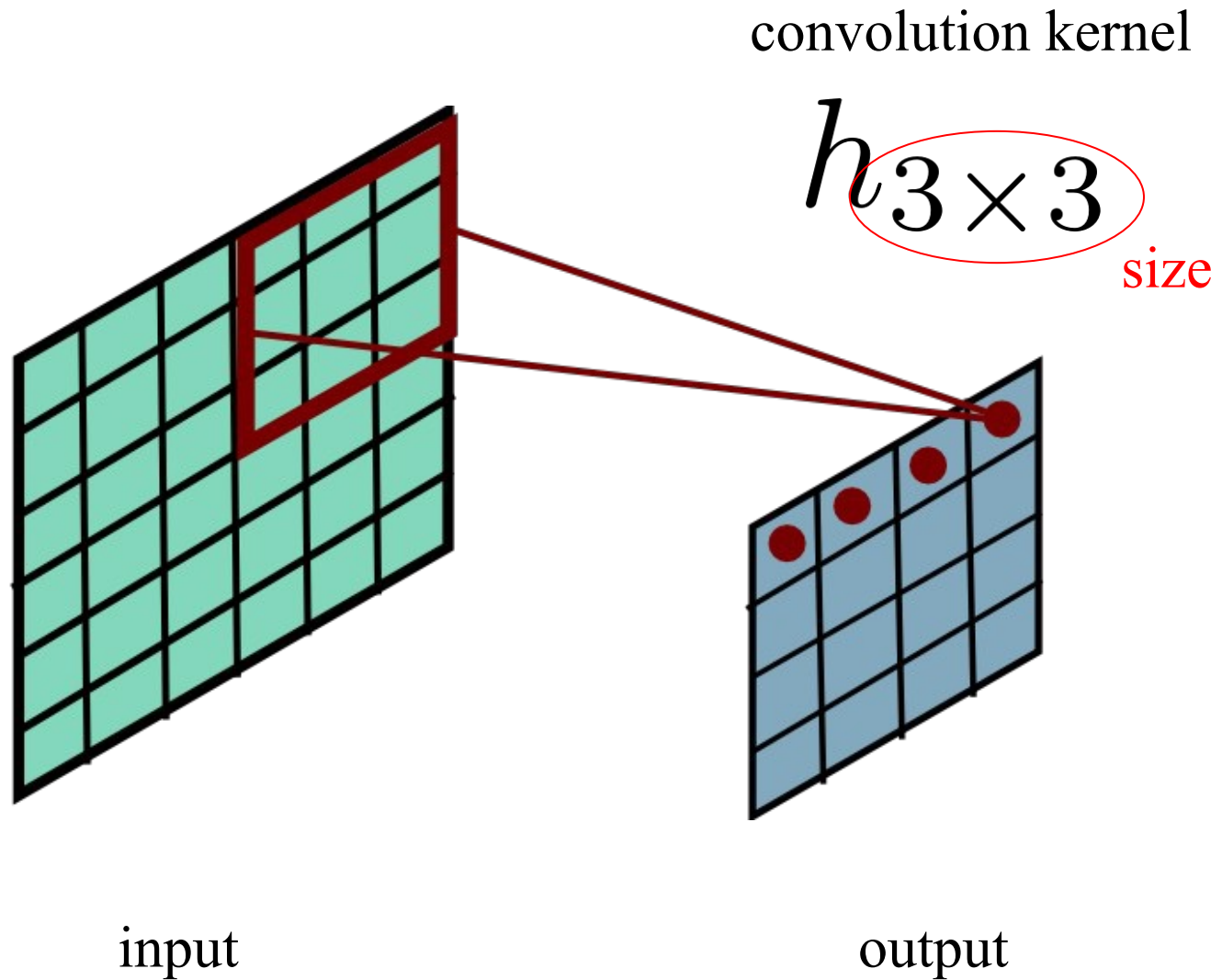
h 3×3 size



input

output

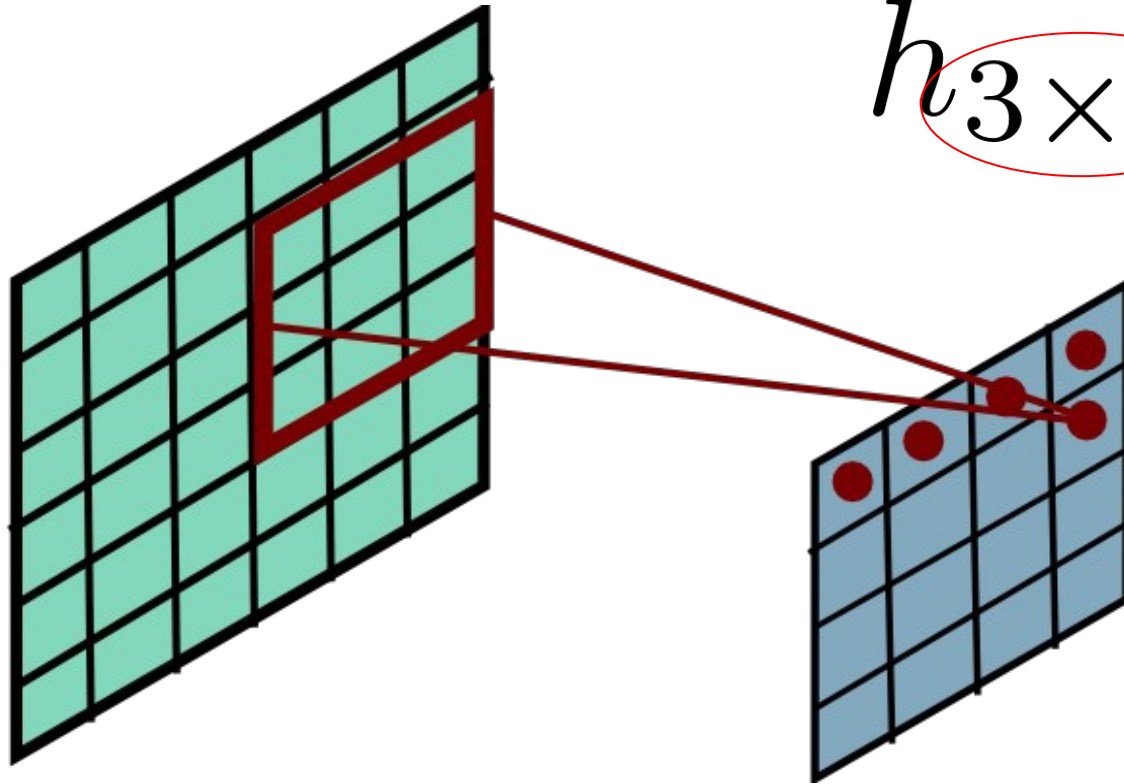
Convolutional Layer



Convolutional Layer

convolution kernel

h 3×3 size



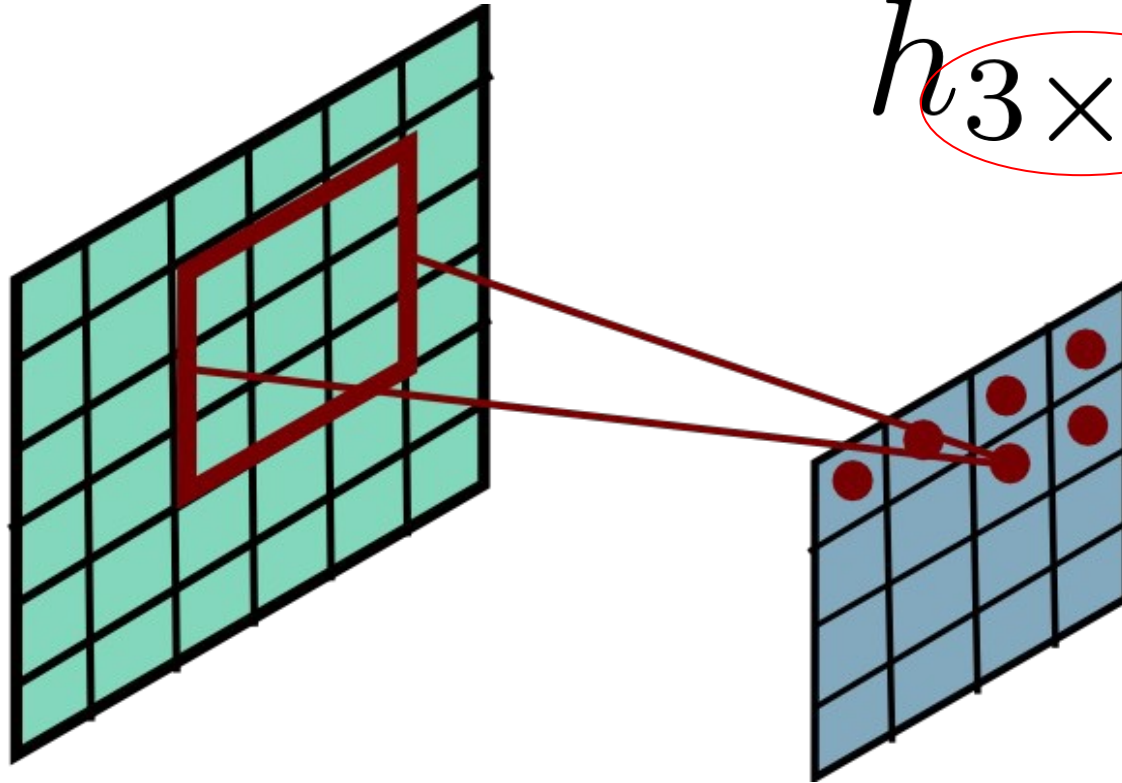
input

output

Convolutional Layer

convolution kernel

h 3×3 size



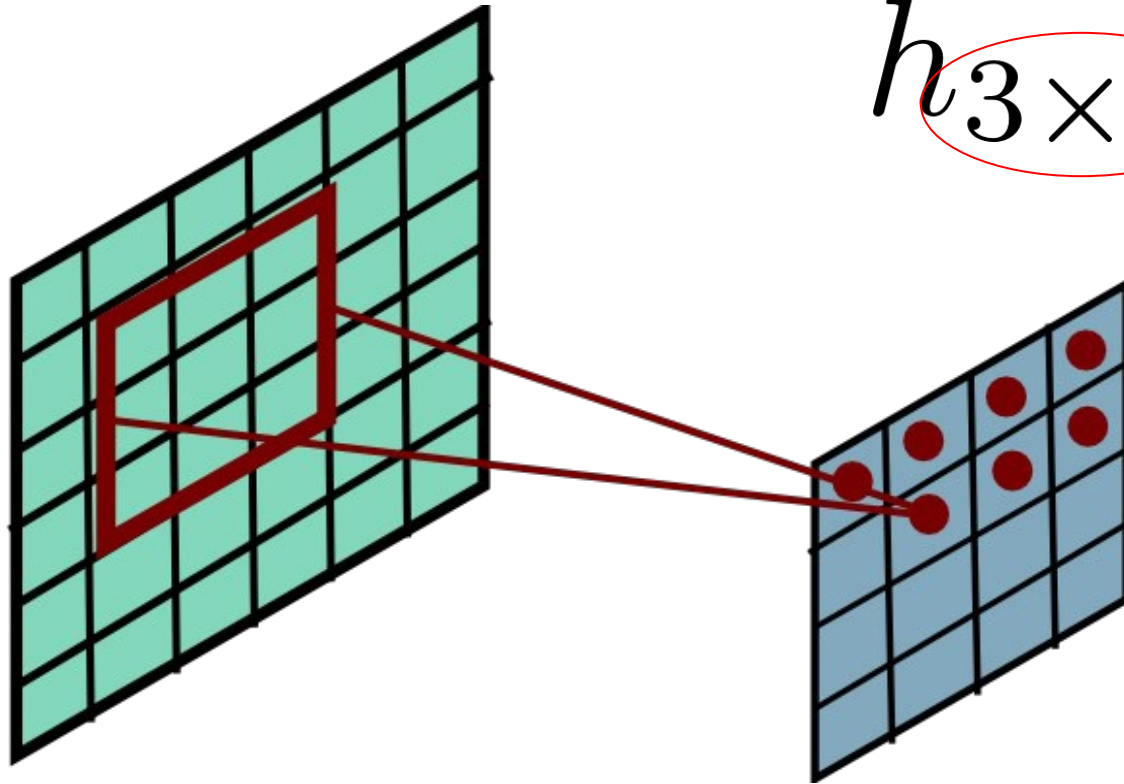
input

output

Convolutional Layer

convolution kernel

h 3×3 size



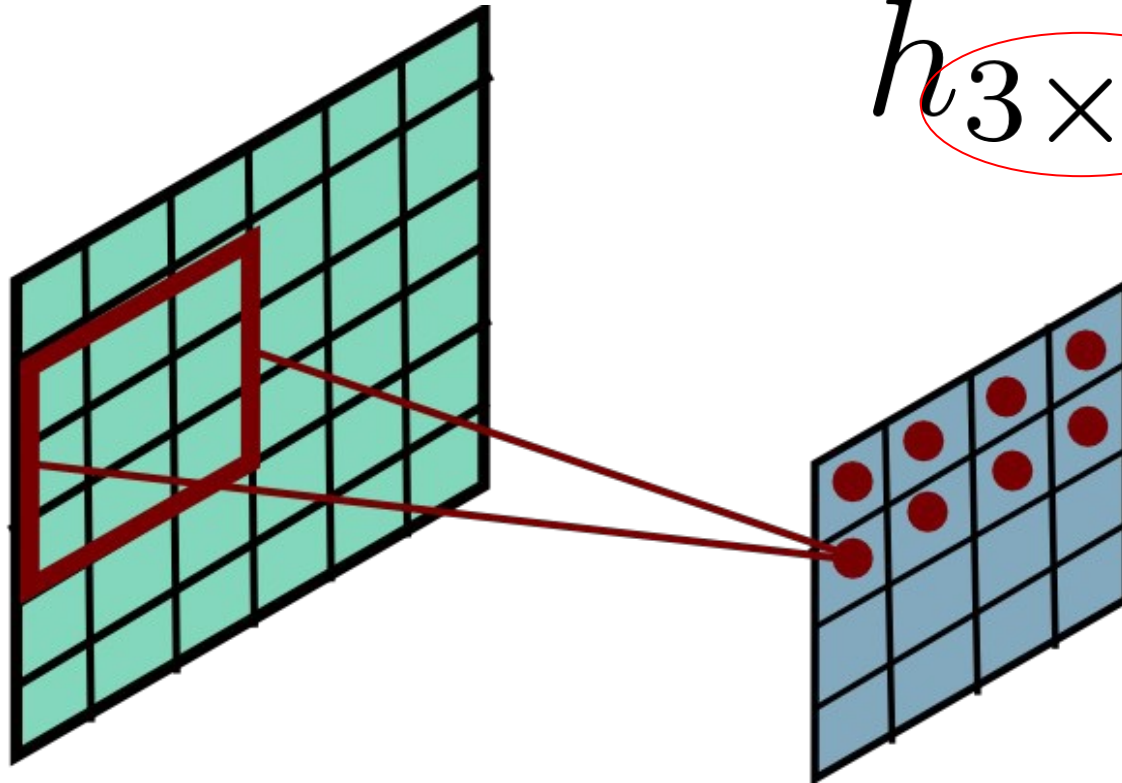
input

output

Convolutional Layer

convolution kernel

h 3×3 size



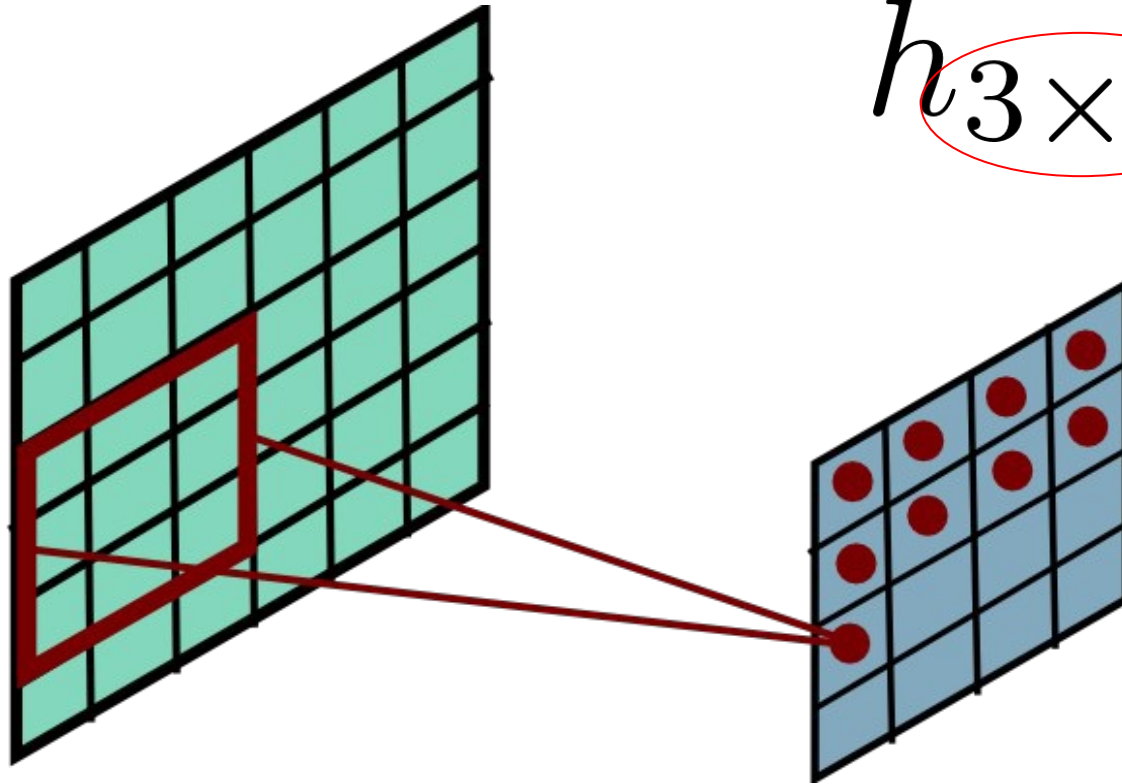
input

output

Convolutional Layer

convolution kernel

h 3×3 size



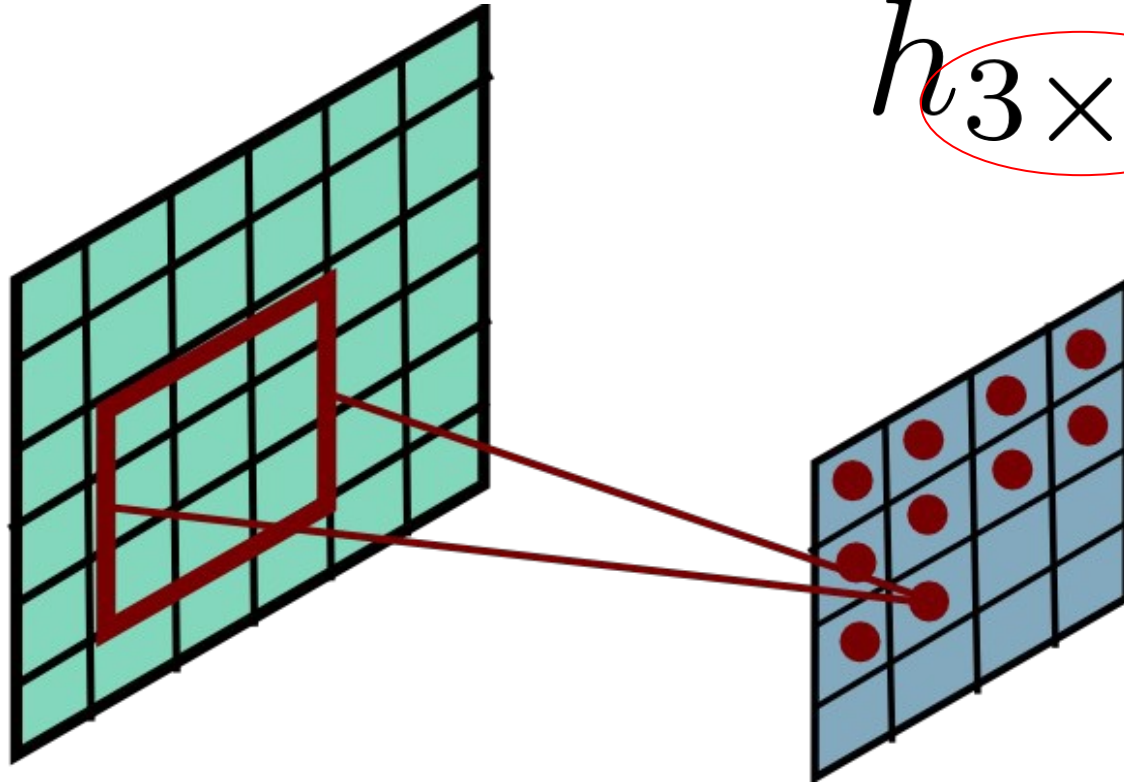
input

output

Convolutional Layer

convolution kernel

h 3×3 size



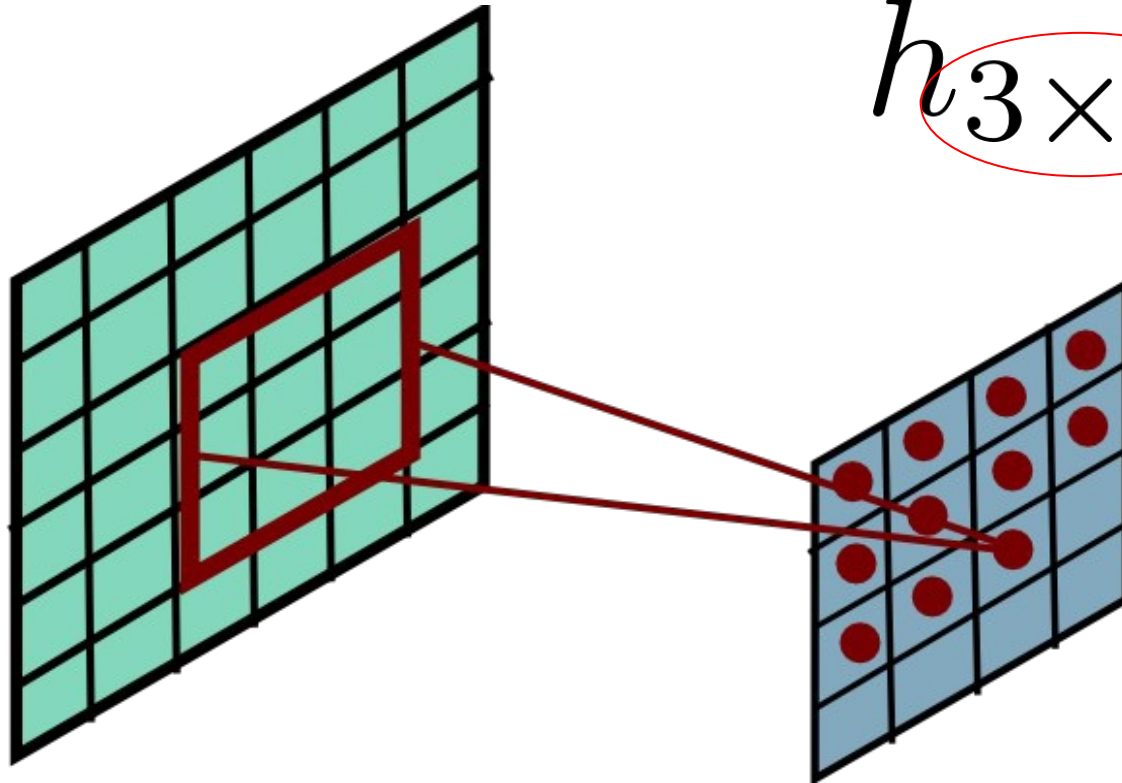
input

output

Convolutional Layer

convolution kernel

h 3×3 size



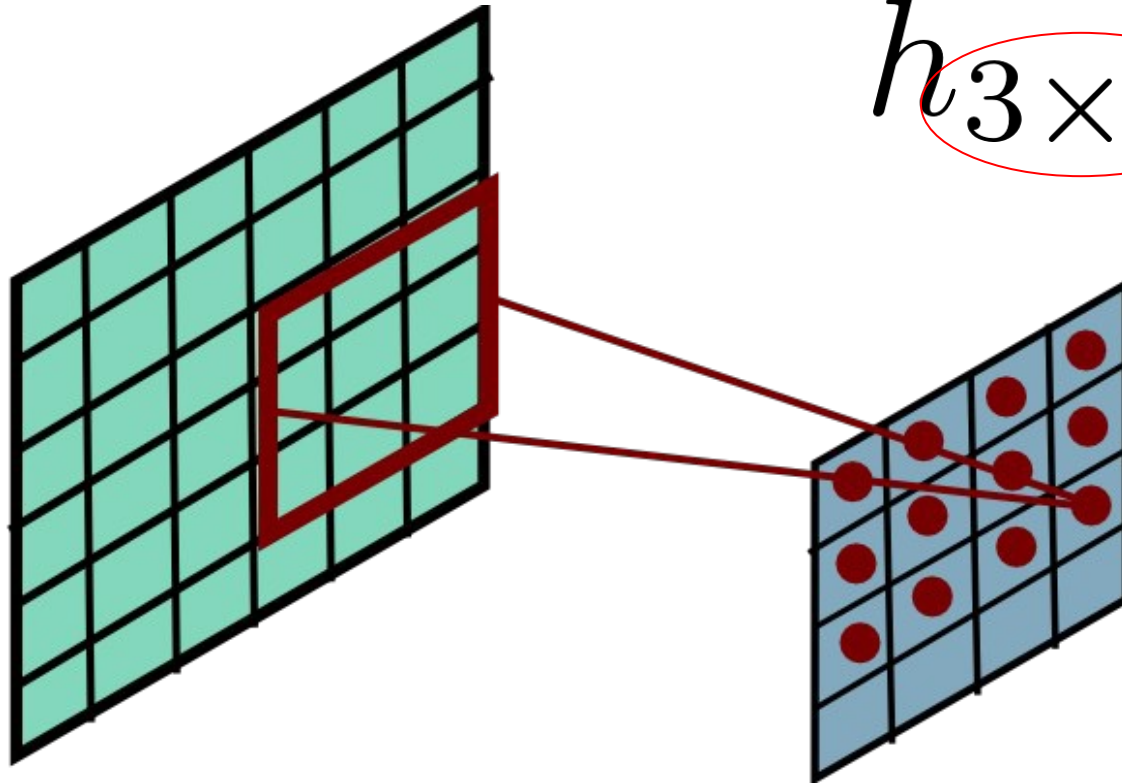
input

output

Convolutional Layer

convolution kernel

h 3×3 size



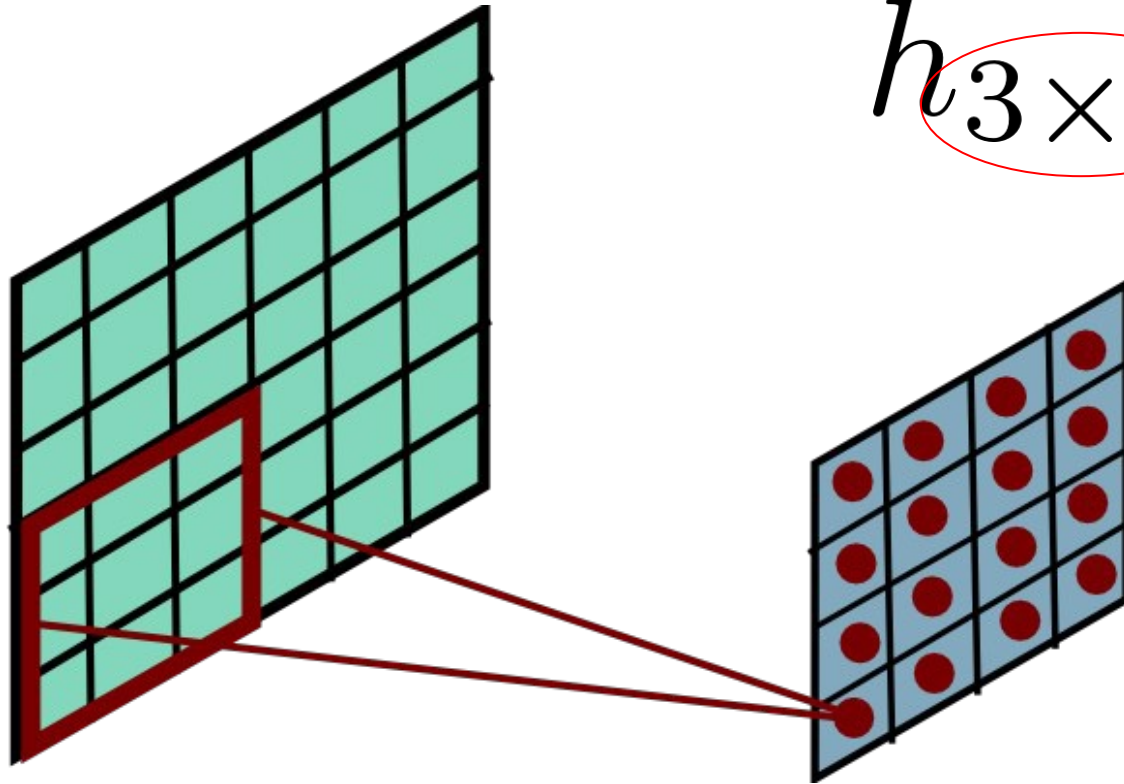
input

output

Convolutional Layer

convolution kernel

h 3×3 size

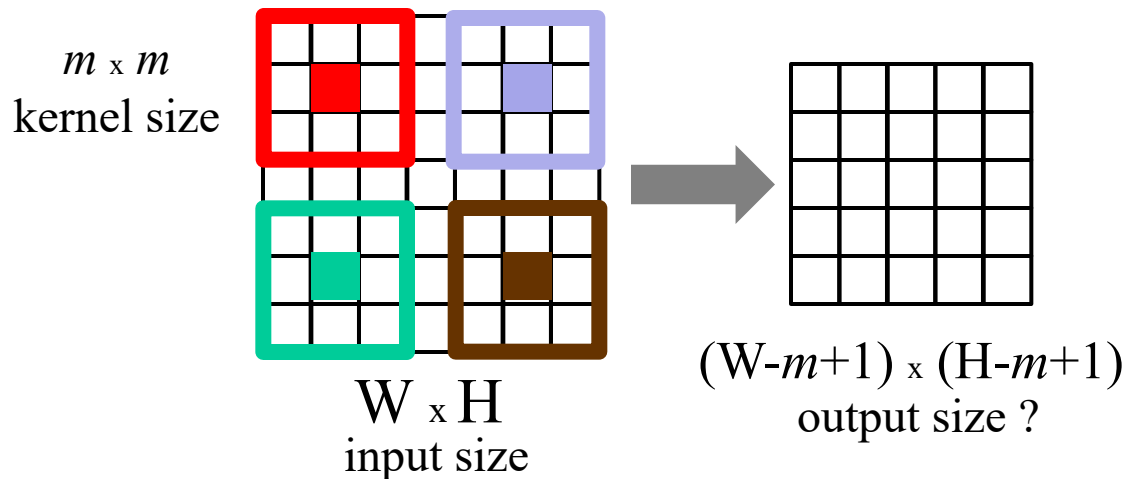


input

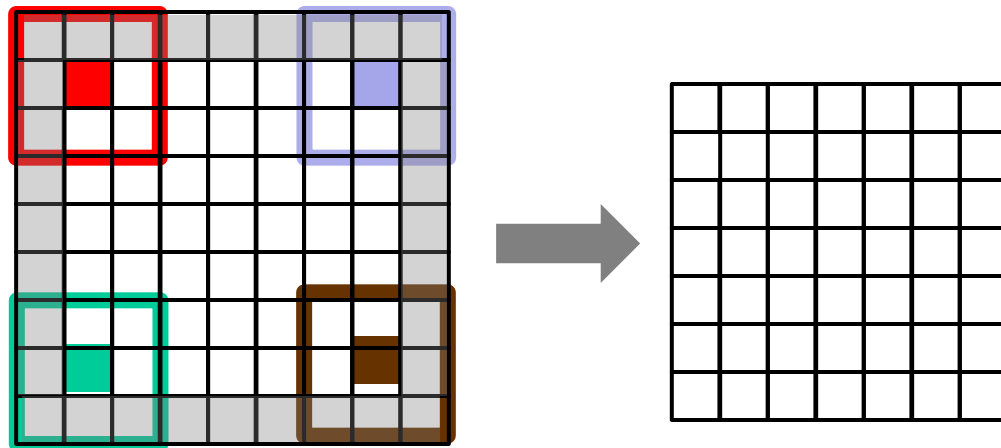
output

Convolutional Layer - Size Change

Output is usually slightly smaller because the borders of the image are left out



If want output to be the same size, zero-pad the input



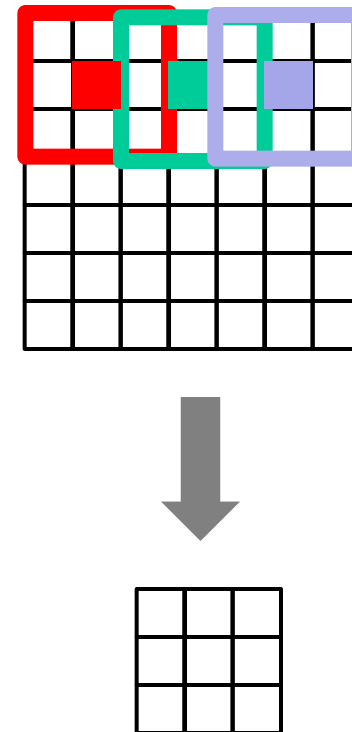
Convolutional Layer - Stride

Can apply convolution only to some pixels (say every second)

- output layer is smaller

Example

- stride = 2 means apply convolution every second pixel
- makes output image approximately **twice smaller** in each dimension
 - image not zero-padded in this example



strided convolution

minimizes information sharing/duplication

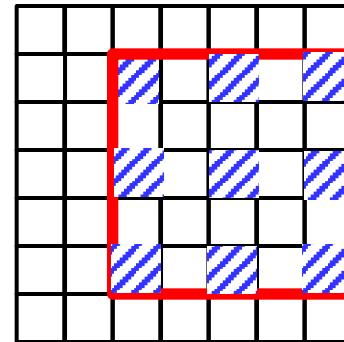
(overlap of kernel windows in the input)

but also reduces spatial resolution of the output

Convolutional Layer - Dilation

It may be helpful to **increase kernel size**
to **enlarge “receptive field”**
for each element of the output

But larger kernels could be expensive...



Use only **subset of points** within the kernel's window

atrous convolution (Fr. *à trous* – hole)

a.k.a. ***dilated convolution***

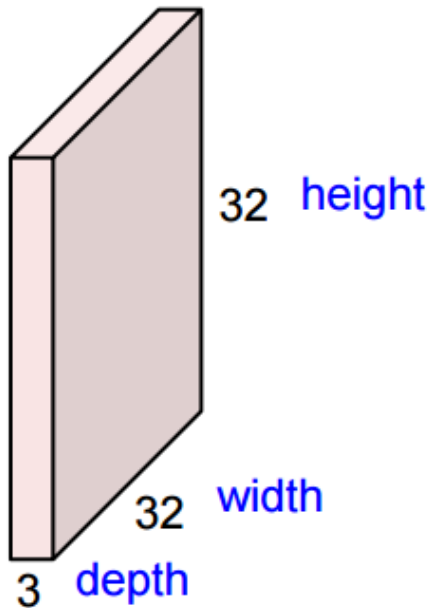
larger *receptive field* (5x5) for output elements
while effectively using smaller kernels (3x3)

It often makes sense to combine atrous convolution with stride

Convolutional Layer – Feature Depth

Input image is usually color, has 3 channels or depth 3

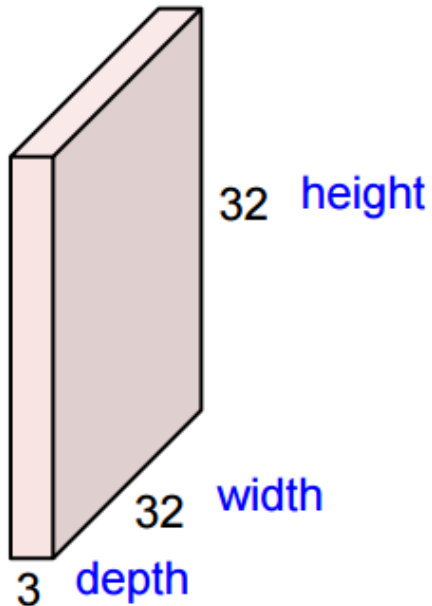
32x32x3 image



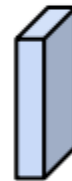
Convolutional Layer – Feature Depth

Convolve 3D image with 3D filter

32x32x3 image



5x5x3 filter



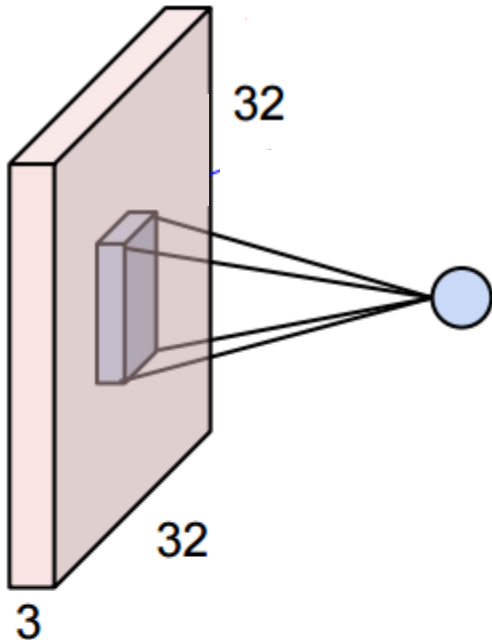
75 parameters

Convolutional Layer – Feature Depth

Each convolution step is a 75 dimensional dot product between the $5 \times 5 \times 3$ filter and a piece of image of size $5 \times 5 \times 3$

Can be expressed as $\mathbf{w}^t \mathbf{x}$, 75 parameters to learn ($\mathbf{w}$)

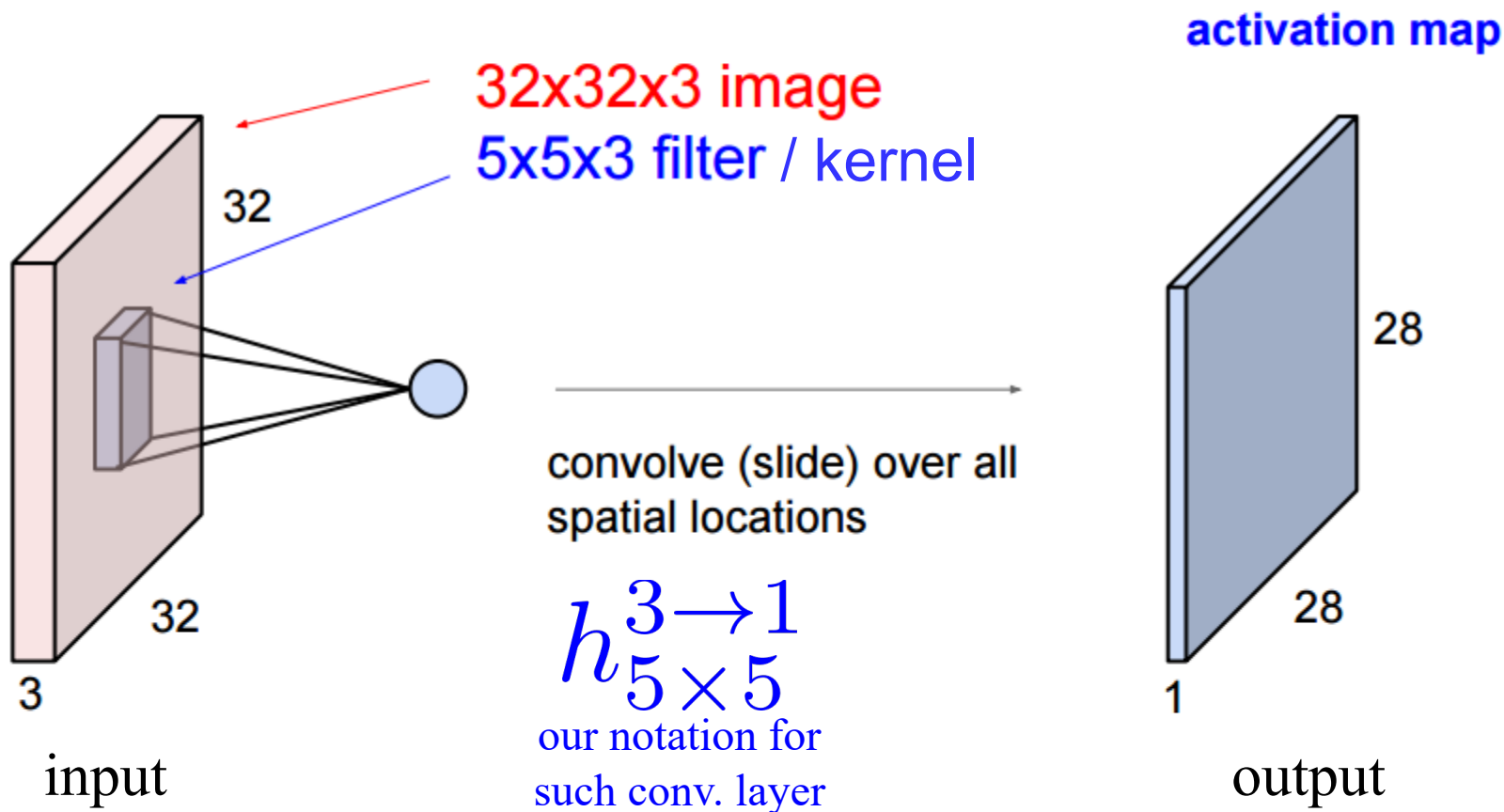
Can add bias $\mathbf{w}^t \mathbf{x} + b$, 76 parameters to learn ($\mathbf{w}, b$)



Convolutional Layer

Convolve 3D image with 3D filter

- result is a 28x28x1 activation map, no zero padding used



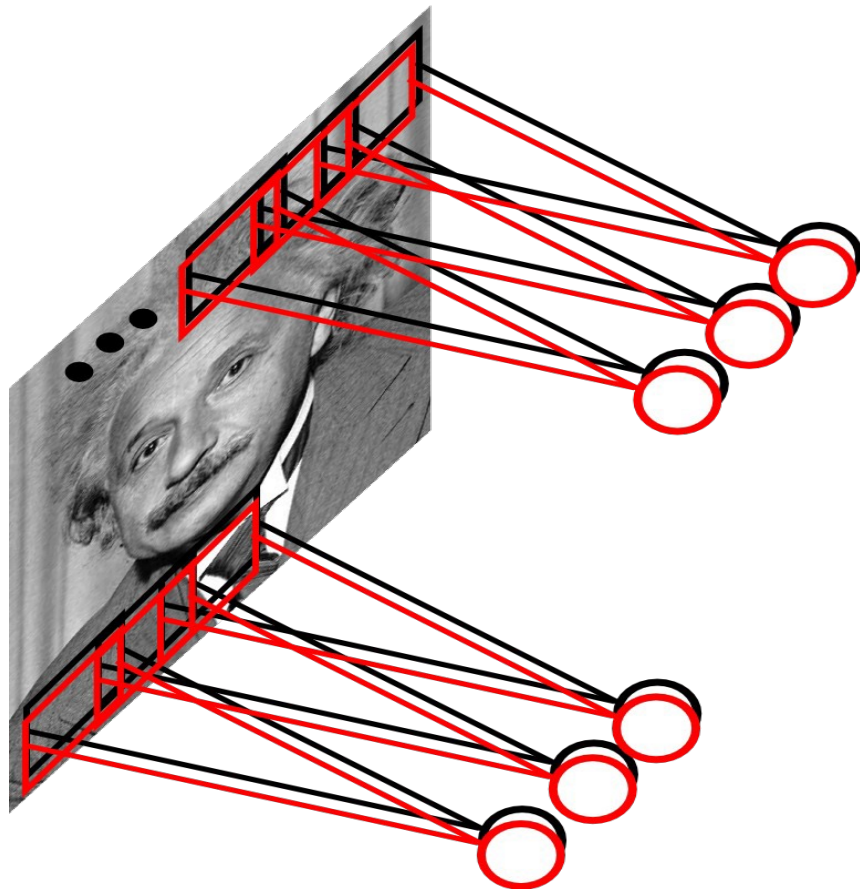
Convolutional Layer

One filter is responsible for
one feature type

Learn multiple filters

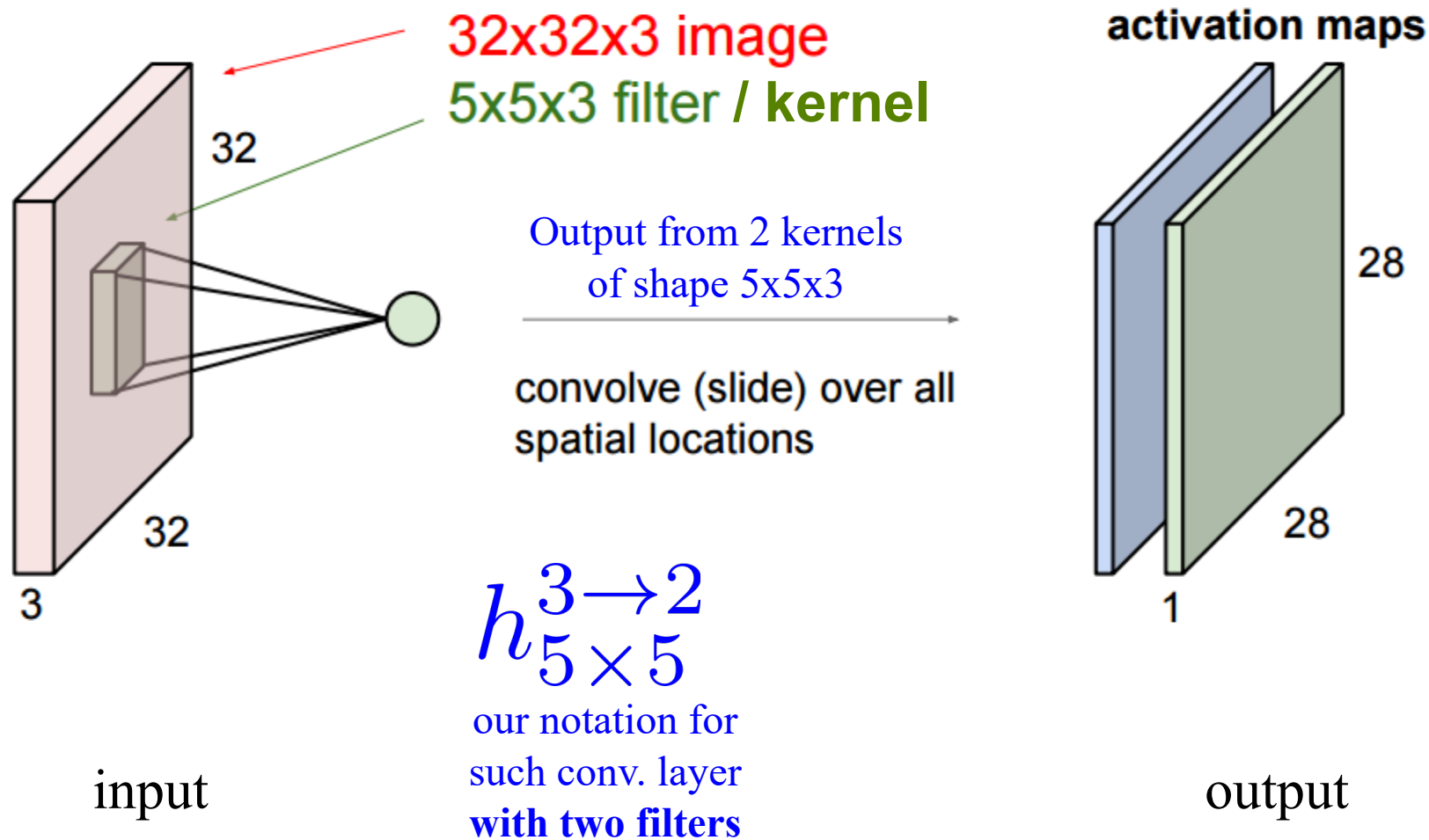
Example:

- 10x10 patch
- 100 filters
- only 10^4 parameters to learn



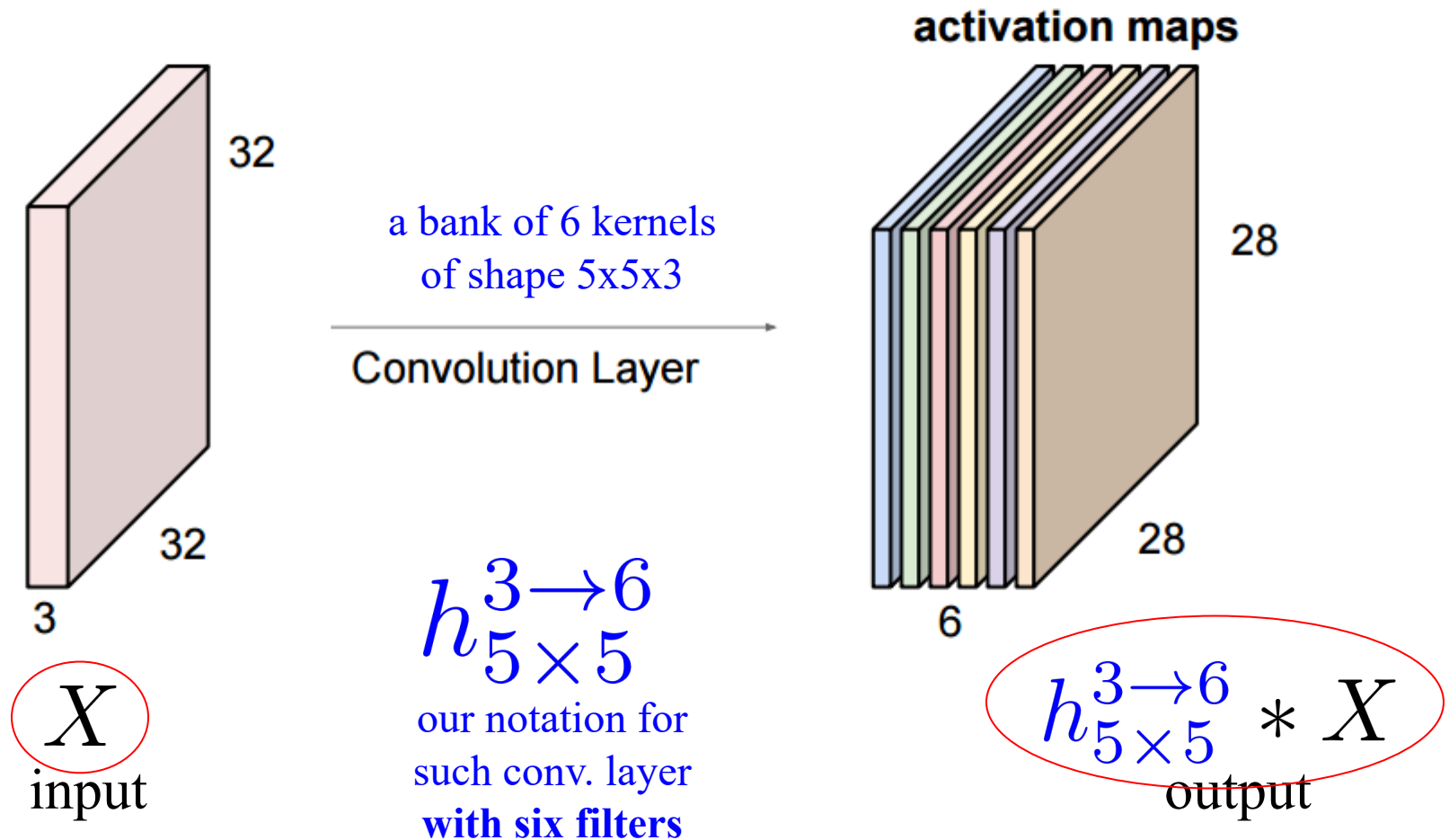
Convolutional Layer

Consider one **extra filter**



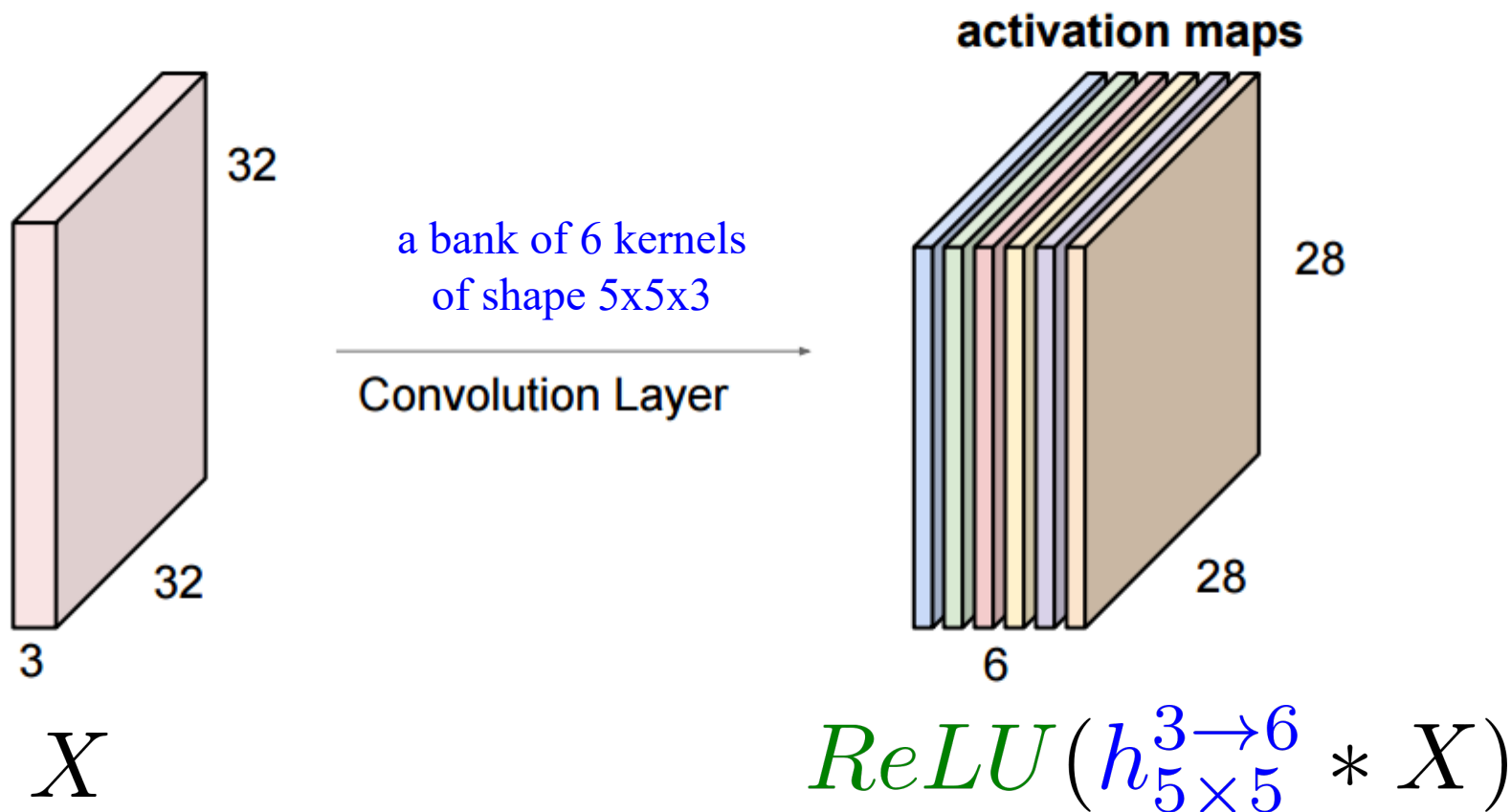
Convolutional Layer

- If have 6 filters (each of size 5x5x3) get 6 activation maps, 28x28 each
- Stack them to get new 28x28x6 “image”



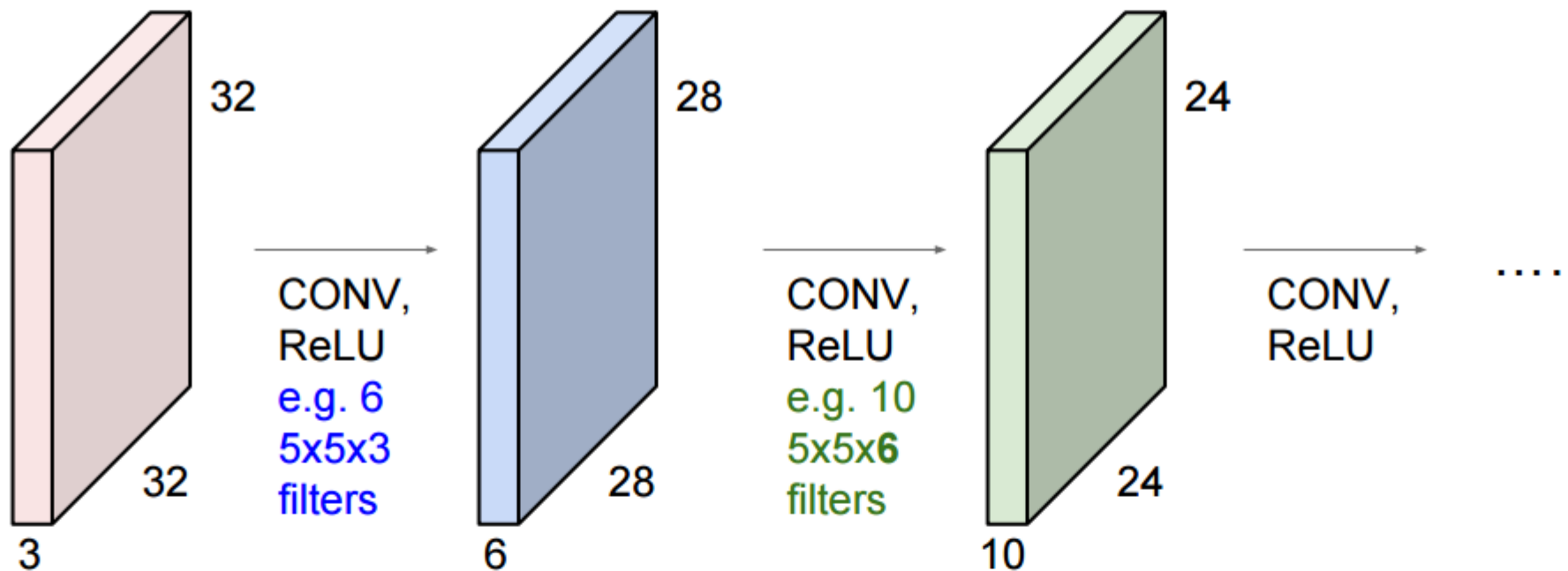
Convolutional Layer

Apply activation function (say **ReLU**) to the activation map



Several Convolution Layers

Construct a sequence of convolution layers interspersed with activation functions



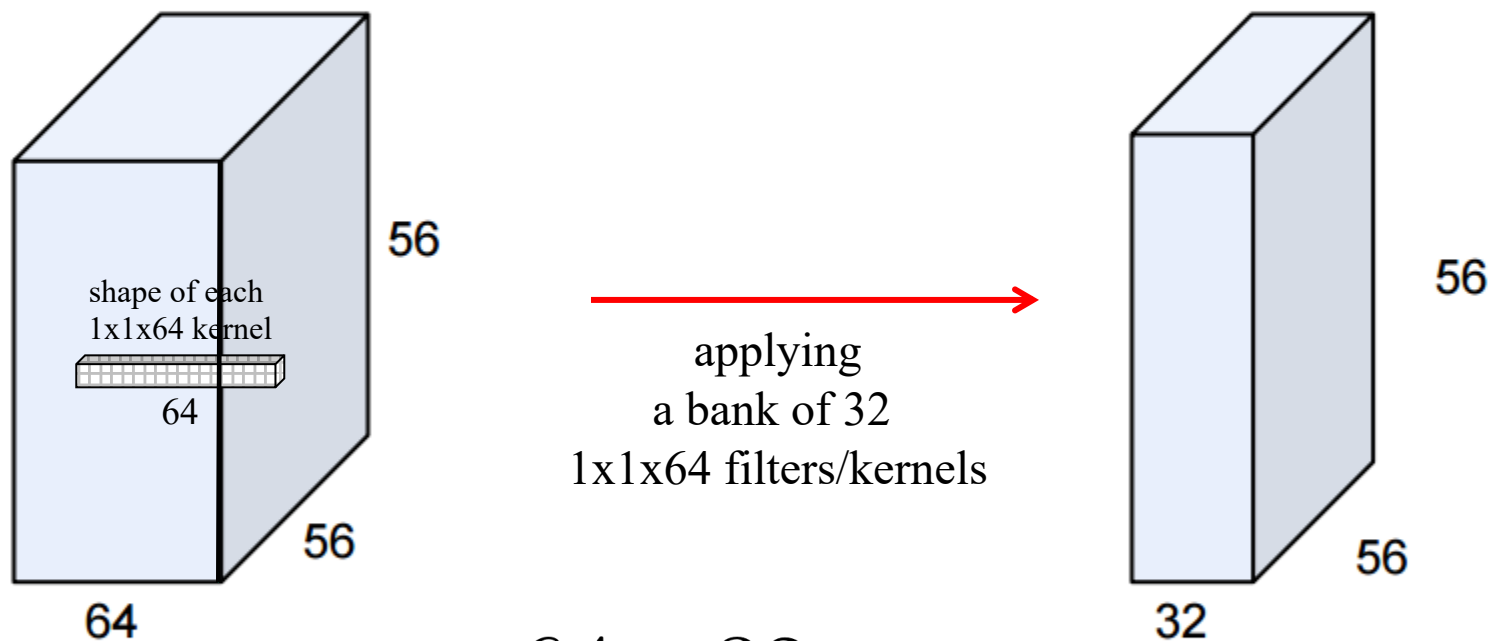
$$ReLU(h_{5 \times 5}^{3 \rightarrow 6} * X) \quad ReLU(h_{5 \times 5}^{6 \rightarrow 10} * X)$$

Convolutional Layer

1x1 convolutions make perfect sense

Example

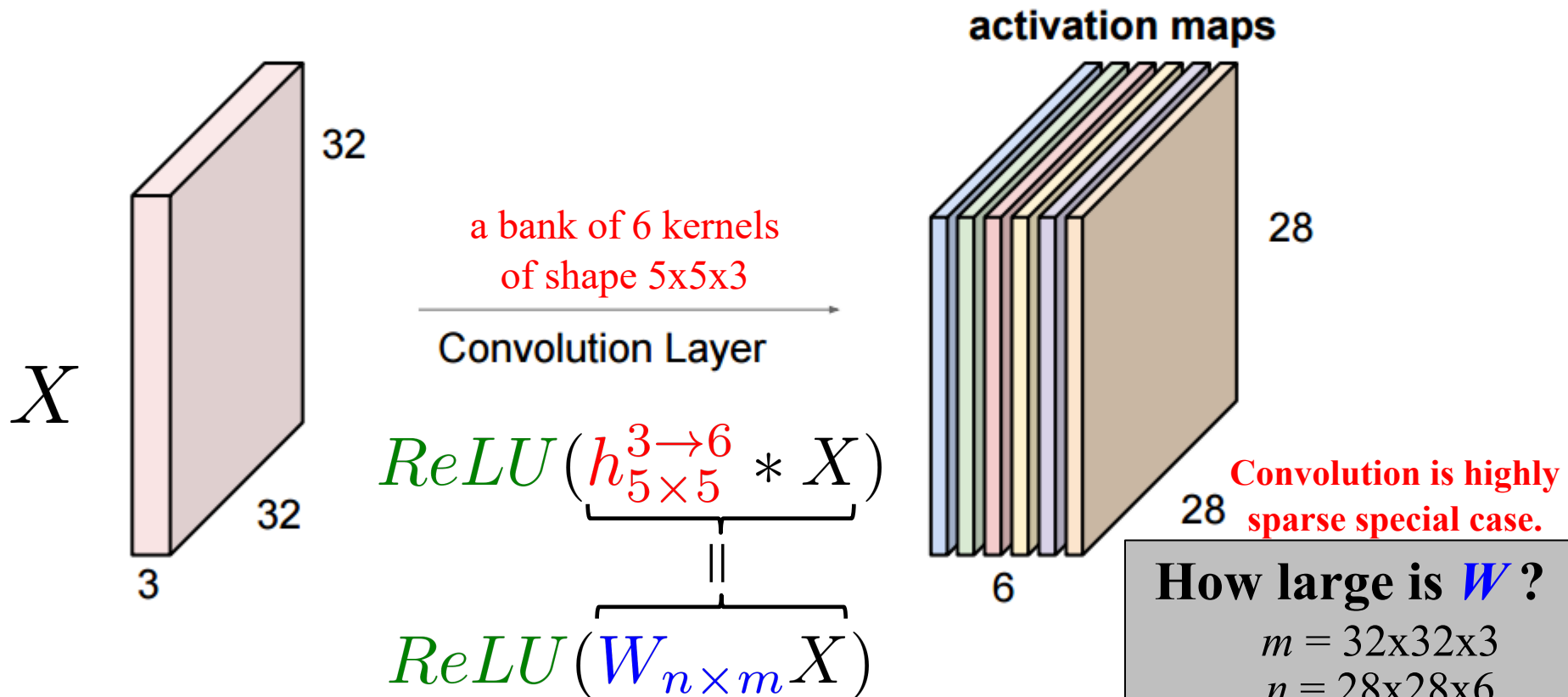
- Input image of size 56x56x64
- Convolve with 32 filters, each of size 1x1x64



$$h_{1 \times 1}^{64 \rightarrow 32} * X$$

Convolutional Layer vs Fully Connected

For example, assume that we applied **ReLU** to the activation maps



Convolution is highly sparse special case.

How large is W ?

$$m = 32 \times 32 \times 3$$

$$n = 28 \times 28 \times 6$$

1.4m parameters

vs. **450 parameters**

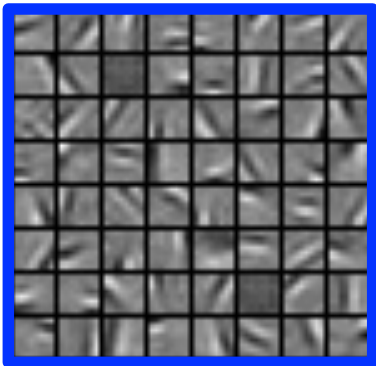
for 6 kernels 5x5x3

The **convolution** is a linear transform.
So, we can equivalently express it via
matrix multiplication for some matrix W .

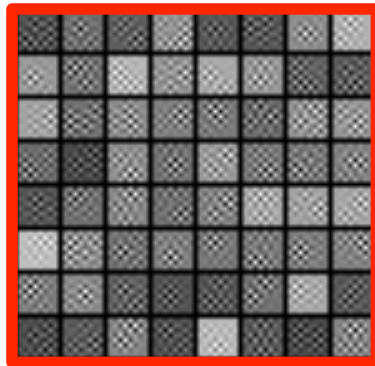
Check Learned Convolutions

- Good training: learned filters exhibit structure and are uncorrelated

GOOD

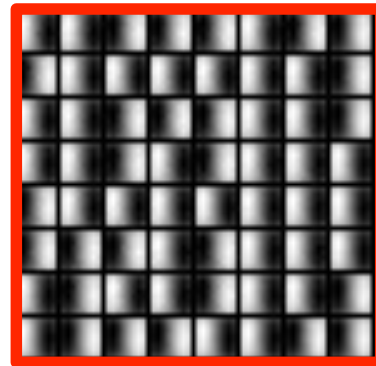


BAD



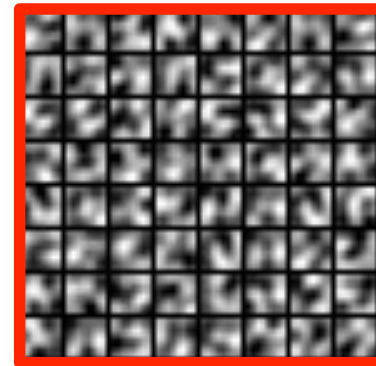
too noisy

BAD



too
correlated

BAD



lack
structure

Convolutional Layer Summary

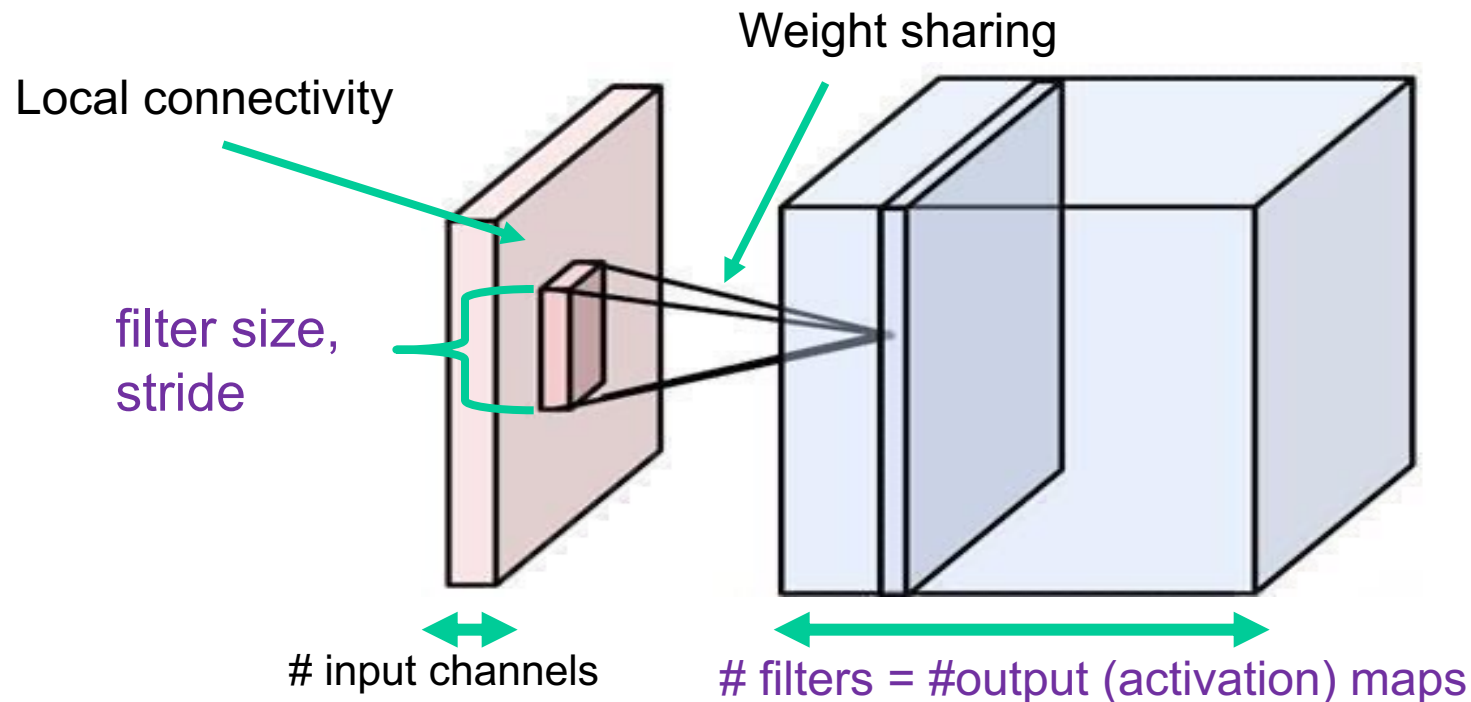
Local connectivity

Weight sharing

Handling multiple input/output channels

Retains location associations

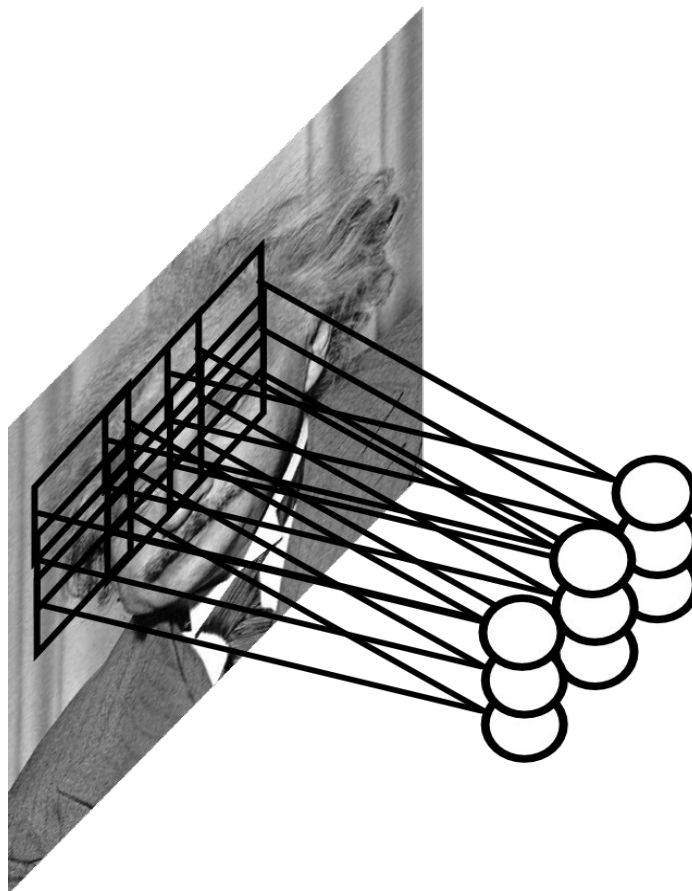
Transforms 3D tensor into 3D tensor (**tensor flow**)



Pooling Layer

Say a filter is an *eye* detector

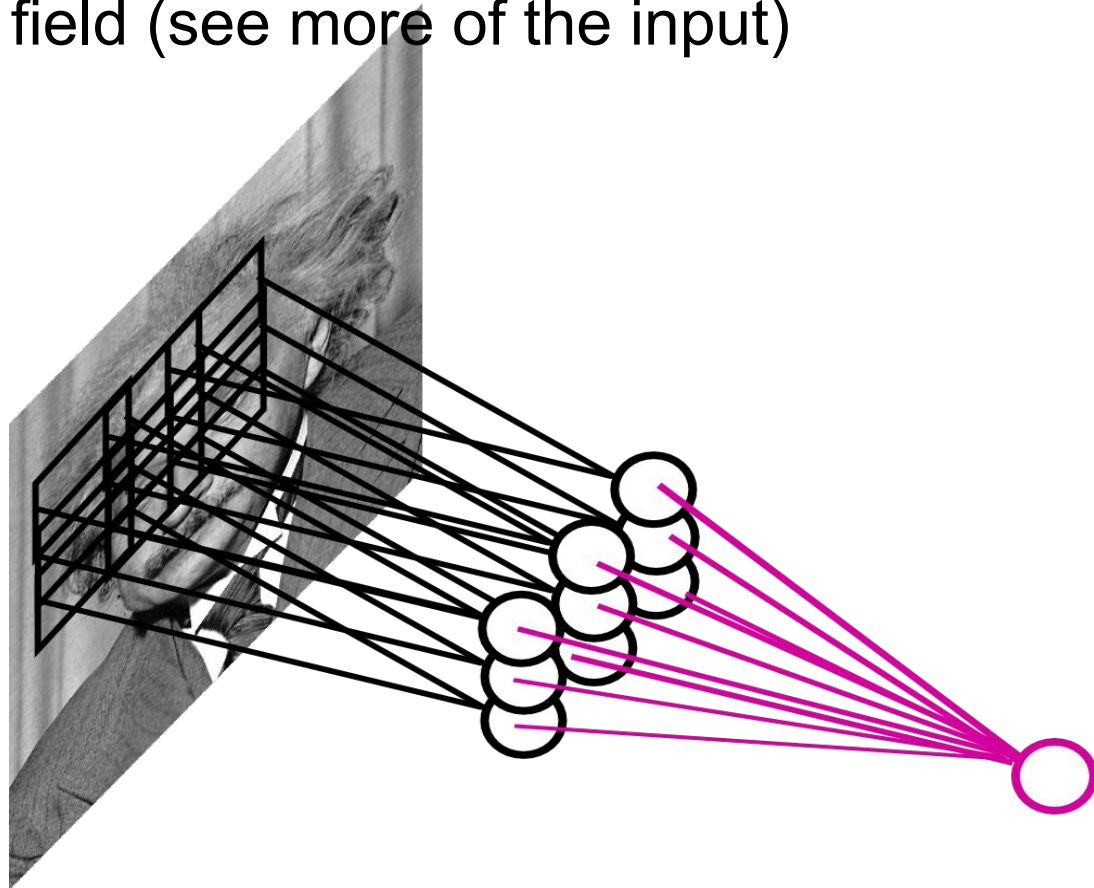
Want detection to be robust to precise *eye* location



Pooling Layer

Pool responses at different locations

- by taking **max**, **average**, etc.
- robustness to exact spatial location
- also larger receptive field (see more of the input)
- Usually pooling applied with stride > 1
- This reduces resolution of output map
- But we already lost resolution (precision) by pooling



Pooling Layer: Max Pooling Example

Single depth slice

1	1	2	4
5	6	7	8
3	2	1	0
1	2	3	4

max pool with 2x2 filters
and stride 2



6	8
3	4

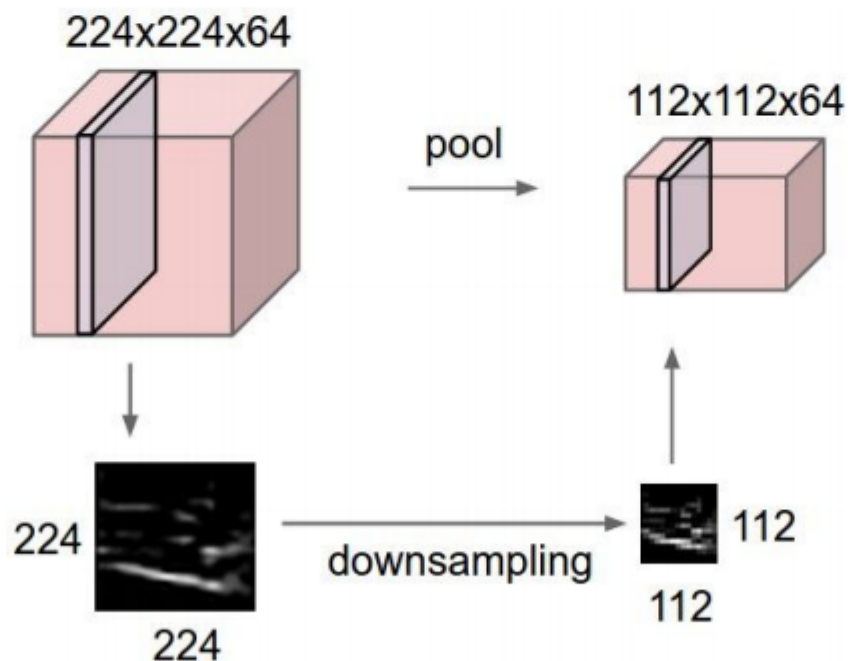
$$\text{Pool}_{2 \times 2}^{st2}(X)$$

our notation for
2 by 2 pooling layer
with stride 2

- pooling can be interpreted as ***downsampling***
- general forms of averaging can be used, e.g. $\left(\frac{1}{N} \sum_{i=1}^N X_i^p \right)^{\frac{1}{p}}$ Hölder mean
where $p=\infty$ implies max and $p=1$ arithmetic mean

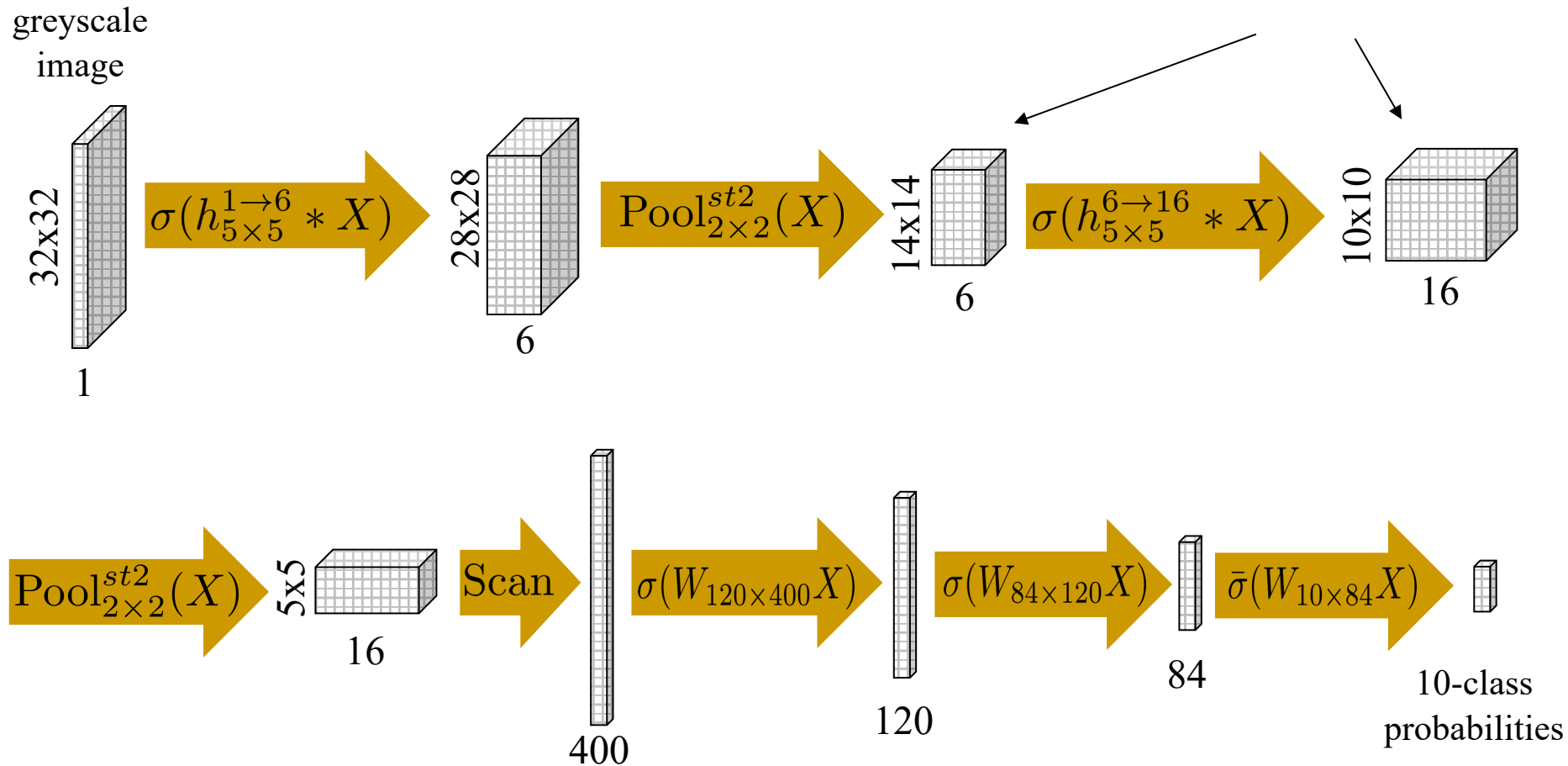
Pooling Layer

Pooling usually applied to each activation map separately



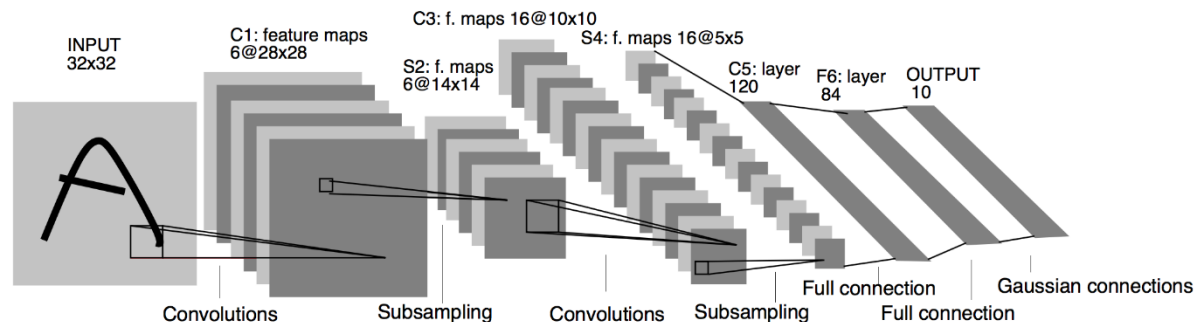
Basic CNN example (à la **LeNet** -1998)

NOTE: transformation of multi-dimensional arrays (**tensors**)



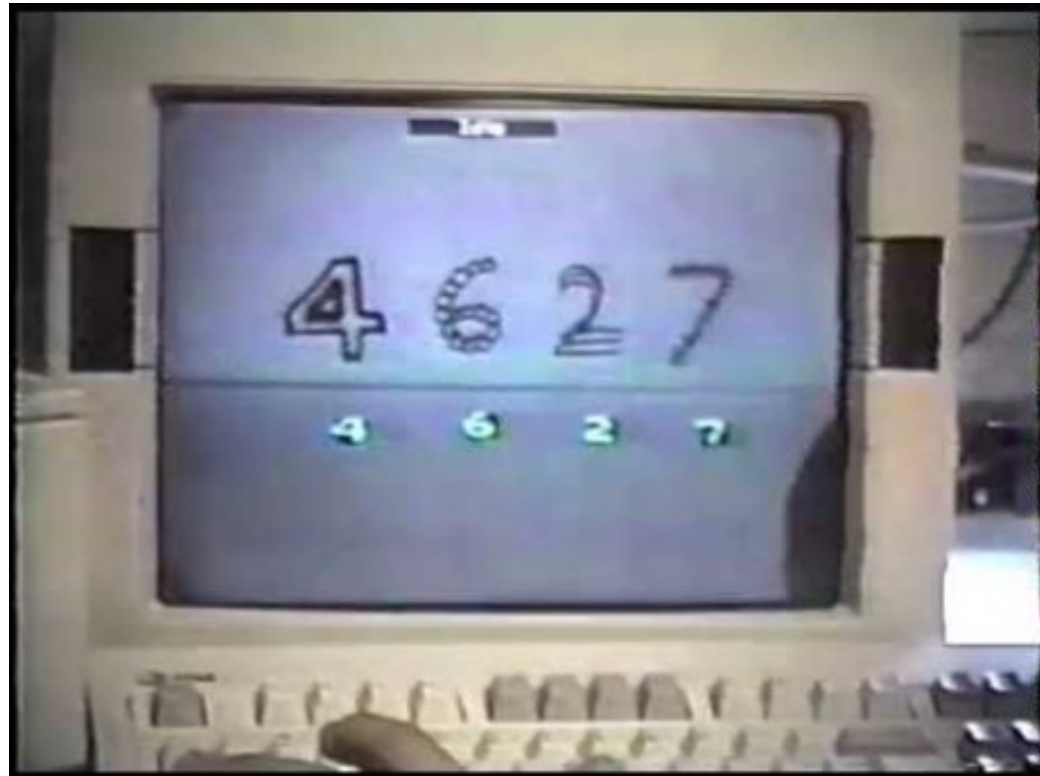
First CNN architectures for classification

- **first CNNs (1982-89)** *Neocognitron: A new algorithm for pattern recognition tolerant of deformations and shifts in position*
(a.k.a. *convNets*)
K. Fukushima, S. Miyake - Pattern Recognition 1982
Handwritten digit recognition with a back-propagation network
Y. LeCun et al - NIPS 1989
- **LeNet (1998)** *Handwritten digit recognition with a back-propagation network*
Y. LeCun, L. Bottou, Y. Bengio, P. Haffner - Proc.of IEEE 1998



First CNN architectures for classification

- **first CNNs (1982-89)** *Neocognitron: A new algorithm for pattern recognition tolerant of deformations and shifts in position*
(a.k.a. *convNets*)
K. Fukushima, S. Miyake - Pattern Recognition 1982
Handwritten digit recognition with a back-propagation network
Y. LeCun et al - NIPS 1989
- **LeNet (1998)**

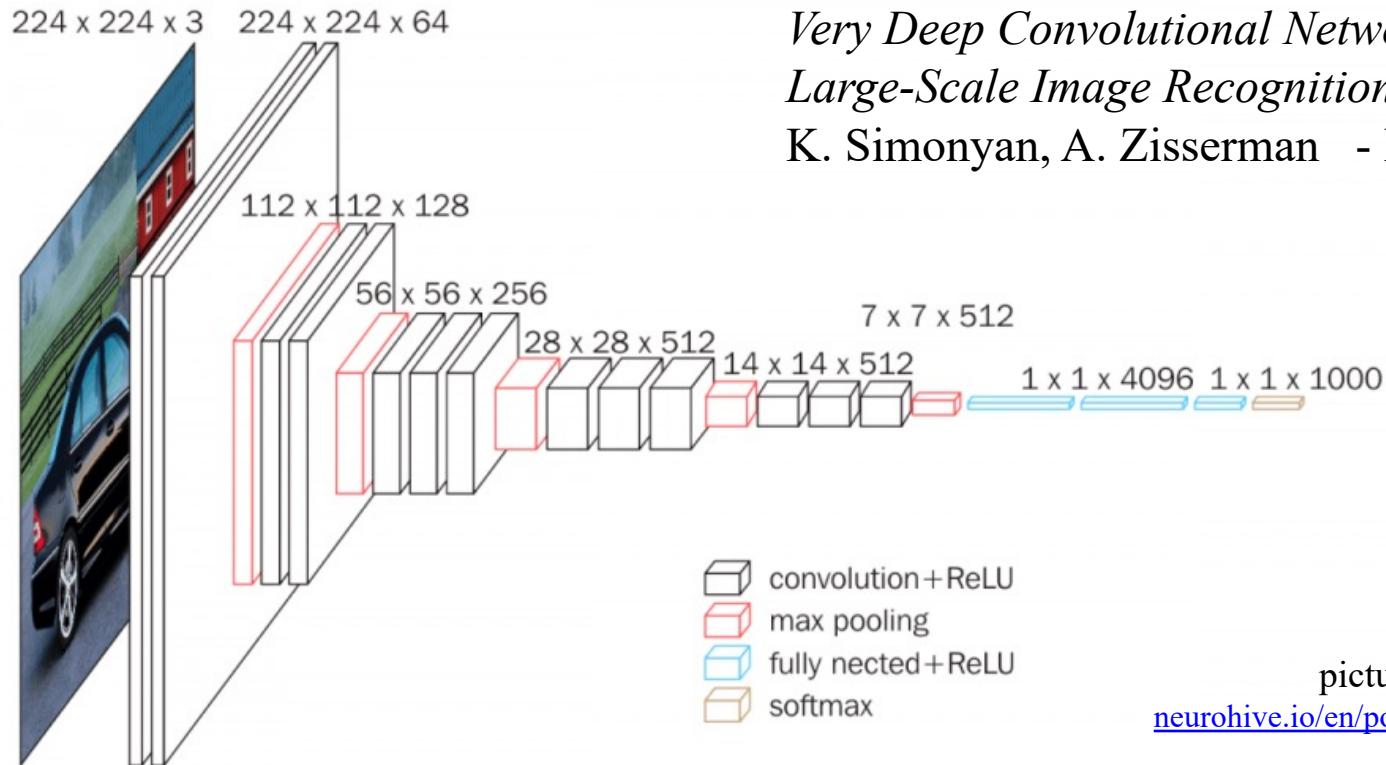


Deep CNN architectures for classification

- **AlexNet** (2012) *ImageNet classification with deep convolutional neural networks*
Alex Krizhevsky, Ilya Sutskever, Geoffrey Hinton - NIPS 2012.
- **VGG** (2014) *Very Deep Convolutional Networks for Large-Scale Image Recognition*
K. Simonyan, A. Zisserman - ICLR 2015
<http://www.robots.ox.ac.uk/~vgg/practicals/cnn/index.html>
- **ResNet** (2016) *Deep residual learning for image recognition*
K. He, X. Zhang, S. Ren, J. Sun. - CVPR 2016

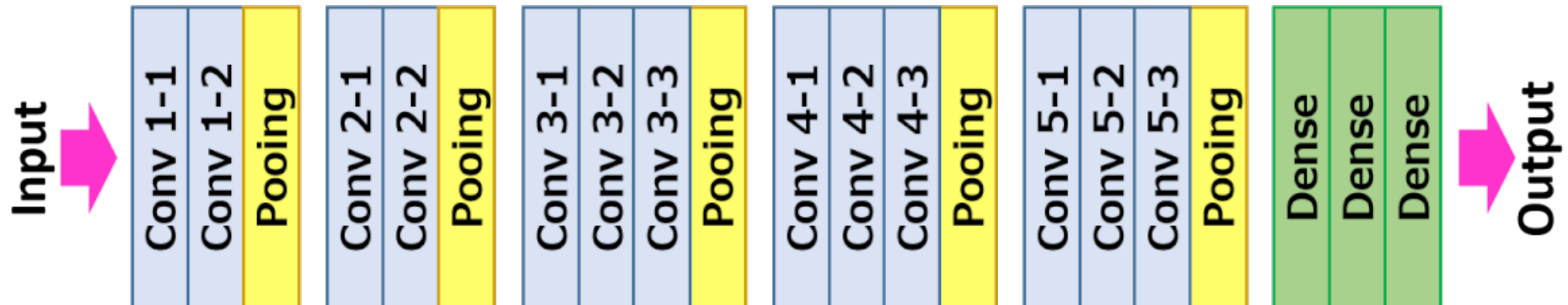
VGG -16

Very Deep Convolutional Networks for Large-Scale Image Recognition
K. Simonyan, A. Zisserman - ICLR 2015



picture credits

neurohive.io/en/popular-networks/vgg16/



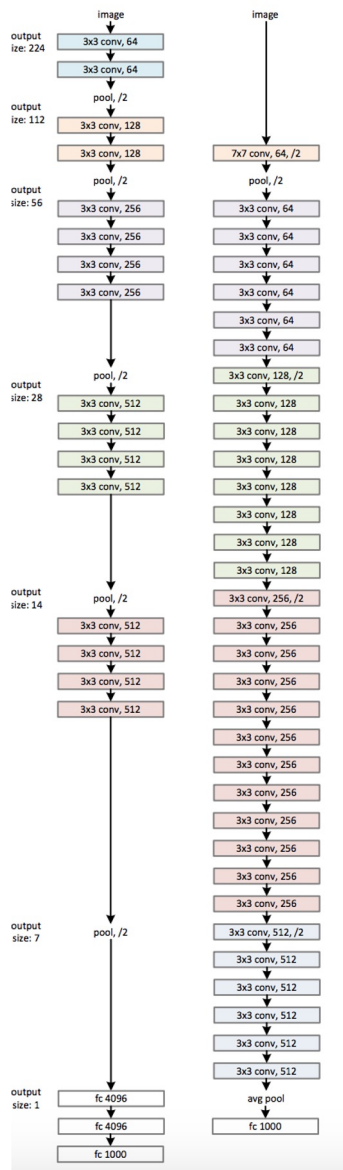
ResNet

very deep 😊

one of the
state of the art
on *image net*

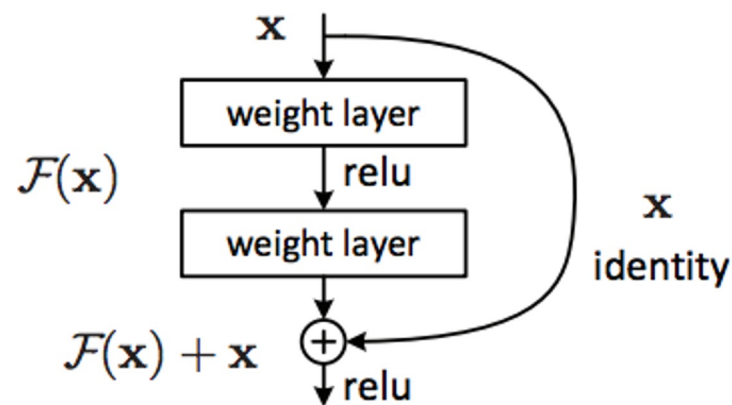
www.image-net.org

- very large dataset
of labeled images
>14,000,000



Deep residual learning for image recognition. K. He, X. Zhang, S. Ren, and J. Sun. CVPR 2016

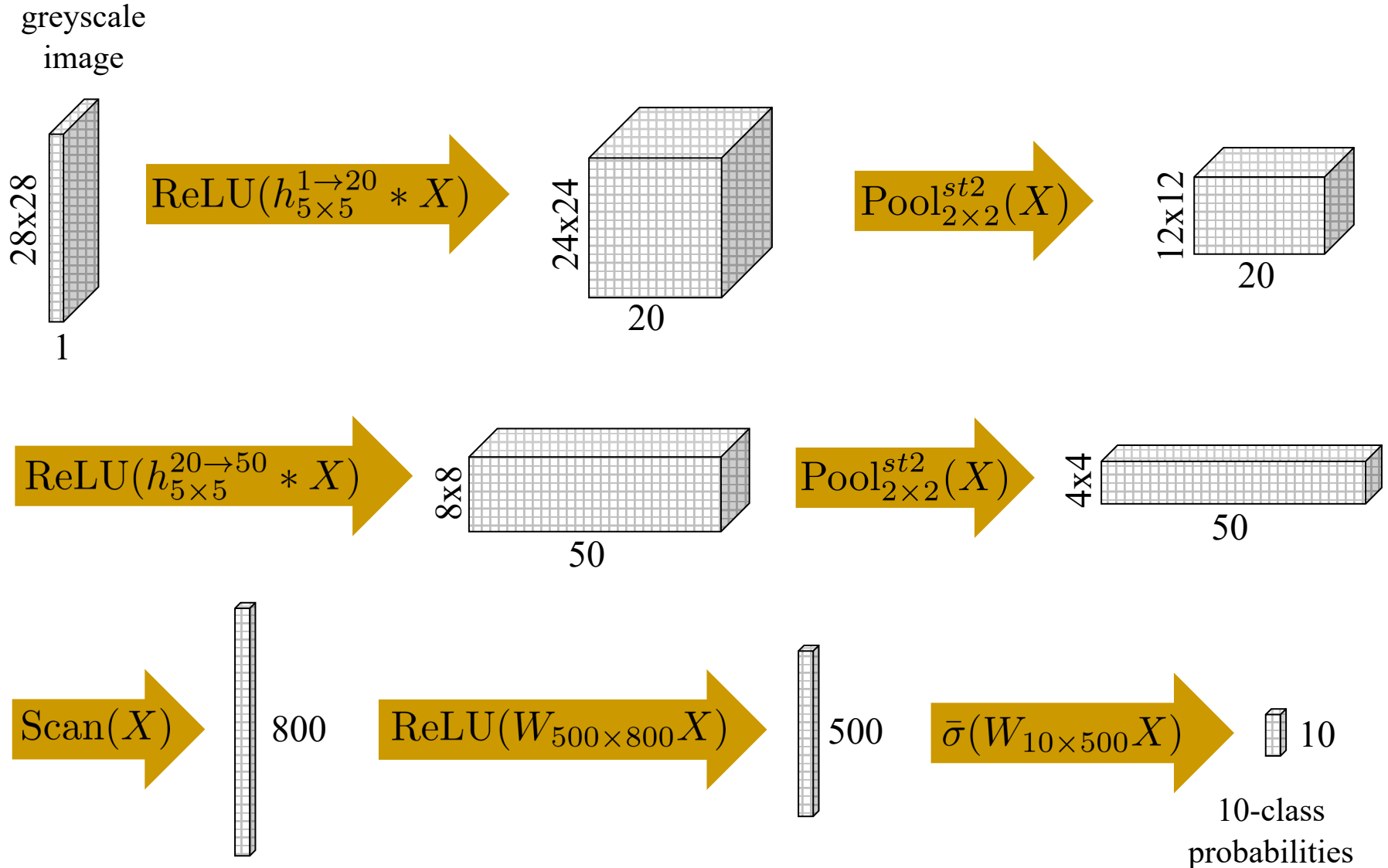
key technical trick

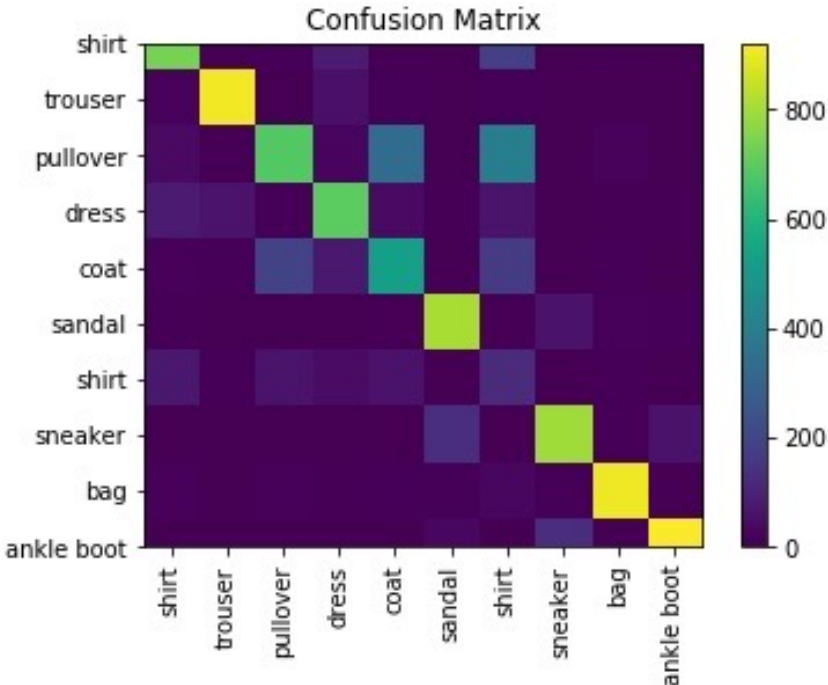


resnet block

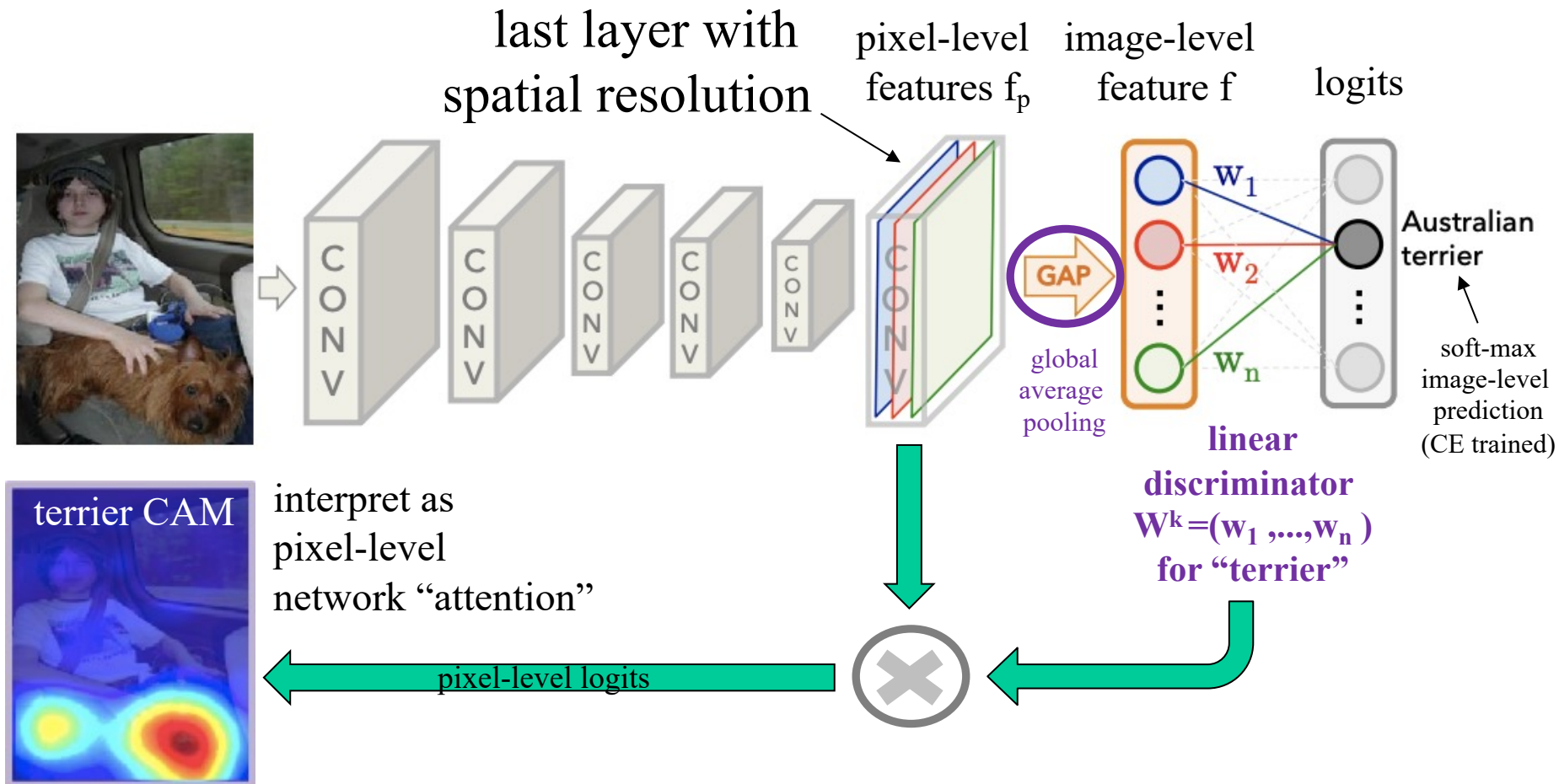
(residual link helps gradient descent)

FashionMNIST classification example





Class-activation Map (CAM)



CVPR 2016: "Learning Deep Features for Discriminative Localization"
B.Zhou, A.Khosla, A. Lapedriza, A.Oliva, A.Torralla

NOTE: motivates ideas for **object localization**, as well as **image-level supervision for semantic segmentation**

Practical Issues for NN training

- data augmentation
 - addresses limited training data
 - e.g. transform available labeled data (domain and range transformations, deformations, cropping, etc.)
- stochastic gradient descent (SGD)
 - batch size selection
 - batch normalization (subtract batch *mean* and divide by batch *st.d.*)
- hyper-parameter tuning (e.g. learning rate)
 - break test data into “*validation*” data + (real) “*testing*” data
 - real testing data is often hidden, and one must use validation data for internal testing purposes, as in assignment 5
- debugging