

$$P(Y=+1|x^{(i)}) = \frac{P(x^{(i)}|Y=+1)P(Y=+1)}{P(x^{(i)})} \quad \text{Generative}$$

↓  
Posterior

$$P(Y=+1|x^{(i)}) = \frac{1}{1 + \exp(-x\theta)} \quad \rightarrow \text{Sigmoid}$$

directly

$x\theta = s = u \rightsquigarrow \text{LCF}$

$$P(Y=+1|x^{(i)}) = \frac{\exp(s_1)}{\exp(s_1) + \exp(s_2)} \rightsquigarrow \hat{y}_p \int_0^1$$

$$P(Y=2|x^{(i)}) = \frac{\exp(s_2)}{\exp(s_1) + \exp(s_2)} \rightsquigarrow \hat{y}_p \int_0^1$$

Softmax ↗ Minimize

$$-\text{avg log likelihood} = \text{cross entropy} = - \sum_{i=1}^N y_a^{(i)} \log \hat{y}_p^{(i)} = E(\theta)$$

$$\theta^{(t+1)} \leftarrow \theta^{(t)} - \alpha \nabla E_{\theta}(\theta)$$

$$\theta^{(t+1)} \leftarrow \theta^{(t)} - \alpha X^T (\hat{y}_p - y_a)$$

$$P(Y=x) = \frac{\exp(s_1)}{\exp(s_1) + \exp(s_2)} = \frac{\frac{\exp(s_1)}{\exp(s_1)}}{1 + \exp(s_2 - s_1)} = \frac{1}{1 + \exp(s_2 - s_1)} = \frac{1}{1 + \exp(-x\theta)}$$

$$- \left[ (s_2 - s_1) = x\theta_{\text{cat}} - x\theta_{\text{dog}} = x(\overbrace{\theta_{\text{cat}}}^{\theta} - \theta_{\text{dog}}) \right] = -x\theta$$

# CONVOLUTIONAL NEURAL NETWORK

Mahdi Roozbahani

Georgia Tech

**Great visualization tool:**

**<https://poloclub.github.io/cnn-explainer/>**

Slides are adopted on Ming Li (University of Waterloo – Deep learning part)  
with some modifications

# James Webb Space Telescope



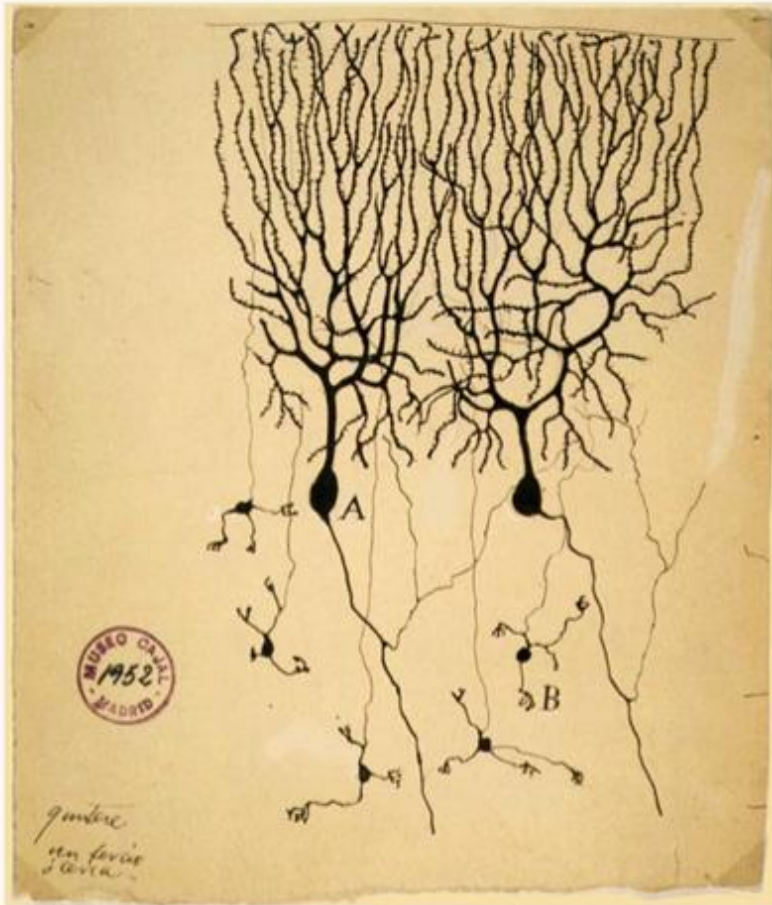
# The landscape of “mountains” and “valleys”



# A visual grouping of five galaxies



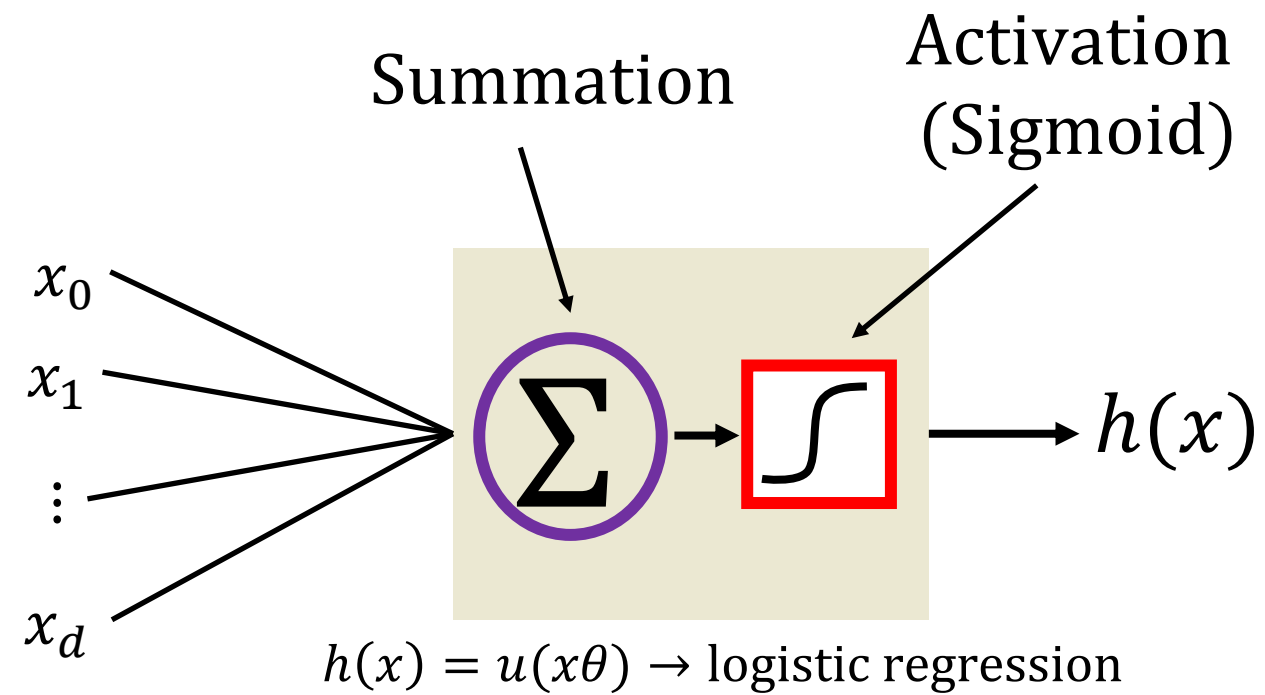
# Inspiration from Biological Neurons



The first drawing of a brain cells by Santiago Ramón y Cajal in 1899

**Neurons:** core components of brain and the nervous system consisting of

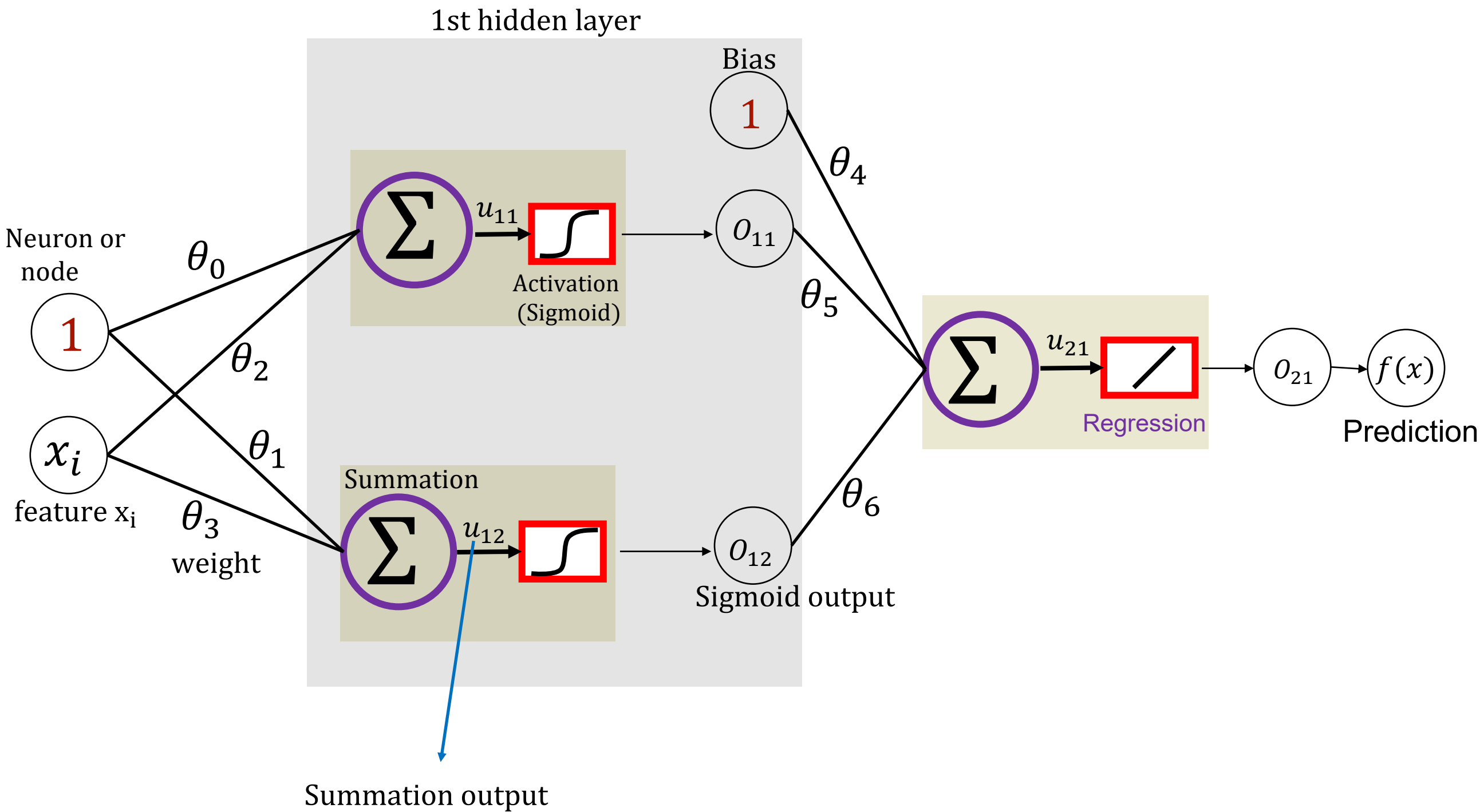
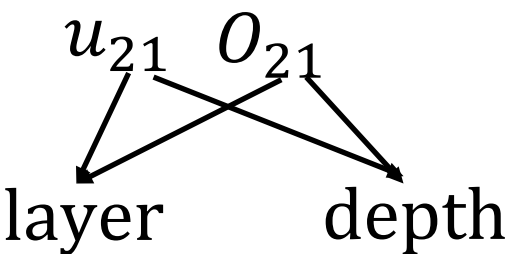
1. Dendrites that collect information from other neurons
2. An axon that generates outgoing spikes



$$\text{output} = \text{activation}(x\theta + b)$$

| Name of the neuron           | Activation function: $\text{activation}(z)$ |
|------------------------------|---------------------------------------------|
| Linear unit                  | $x\theta$                                   |
| Threshold/sign unit          | $\text{sign}(x\theta)$                      |
| Sigmoid unit                 | $\frac{1}{1 + \exp(x\theta)}$               |
| Rectified linear unit (ReLU) | $\max(0, x\theta)$                          |
| Tanh unit                    | $\tanh(x\theta)$                            |

# NN Regression

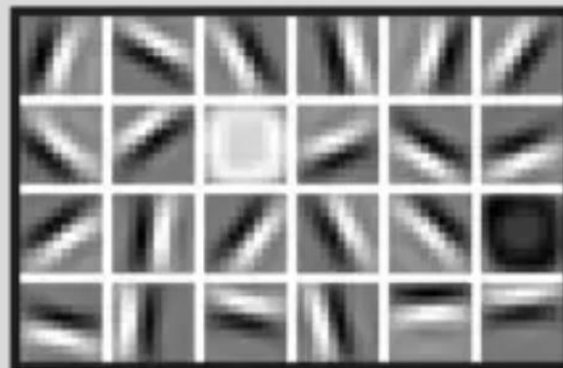


# FACIAL RECOGNITION

Deep-learning neural networks use layers of increasingly complex rules to categorize complicated shapes such as faces.



Layer 1: The computer identifies pixels of light and dark.



Layer 2: The computer learns to identify edges and simple shapes.



Layer 3: The computer learns to identify more complex shapes and objects.



Layer 4: The computer learns which shapes and objects can be used to define a human face.

$$\gamma_a \quad \gamma_D$$

$$\begin{bmatrix} 1 & 0 \end{bmatrix} \quad \begin{bmatrix} 0.7 & 0.3 \end{bmatrix}$$

$$X = \begin{bmatrix} x_1 & x_2 & \dots & x_d \end{bmatrix}$$

$$\theta_1 \quad \theta_d$$

$$\gamma_a = \begin{bmatrix} 1 \\ 0 \\ \vdots \end{bmatrix}$$

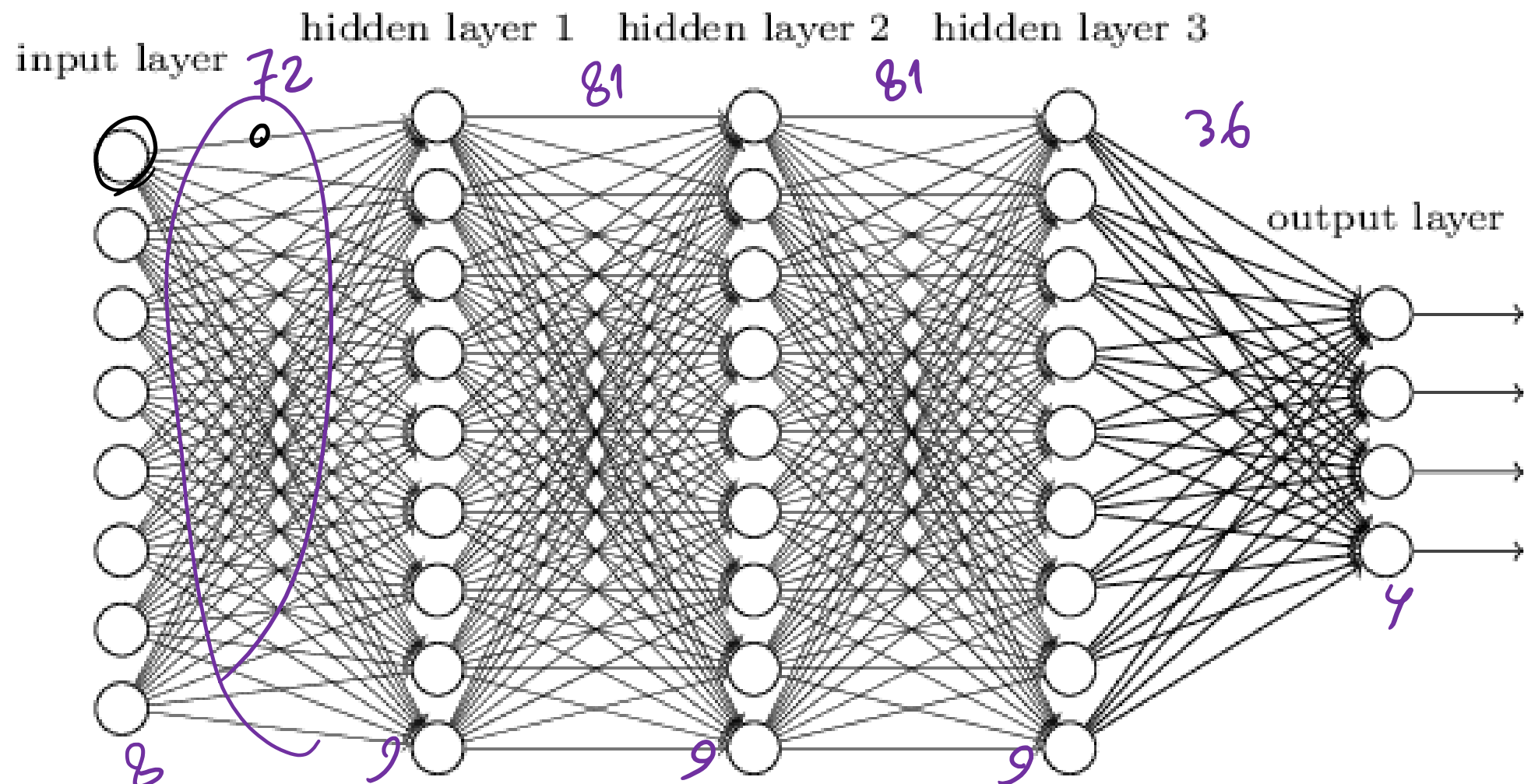
$$CE = - \sum \gamma_a \log \hat{\gamma}_p$$

$$\hat{\gamma}_p = \frac{\exp(s_1)}{\exp(s_1) + \exp(s_2)}$$

$$s_1 = X\theta = \theta_1 x_1 + \dots + \theta_d x_d$$

# Smaller Network: CNN

- We know it is good to learn a small model.
- From this fully connected model, do we really need all the edges?
- Can some of these be shared?



ANN

|     |     |     |
|-----|-----|-----|
| 1/9 | 1/9 | 1/9 |
| 1/9 | 1/9 | 1/9 |
| 1/9 | 1/9 | 1/9 |

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



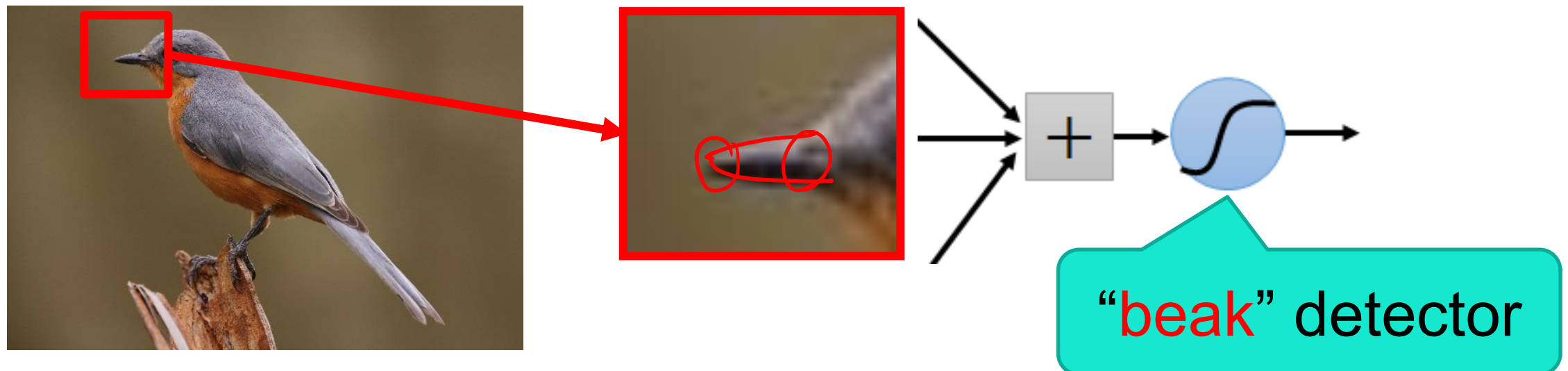
|   |   |   |   |
|---|---|---|---|
| 0 | 0 | 0 | 0 |
| 0 | 0 | 0 | 0 |
| 0 | 0 | 0 | 0 |
| 0 | 0 | 0 | 0 |

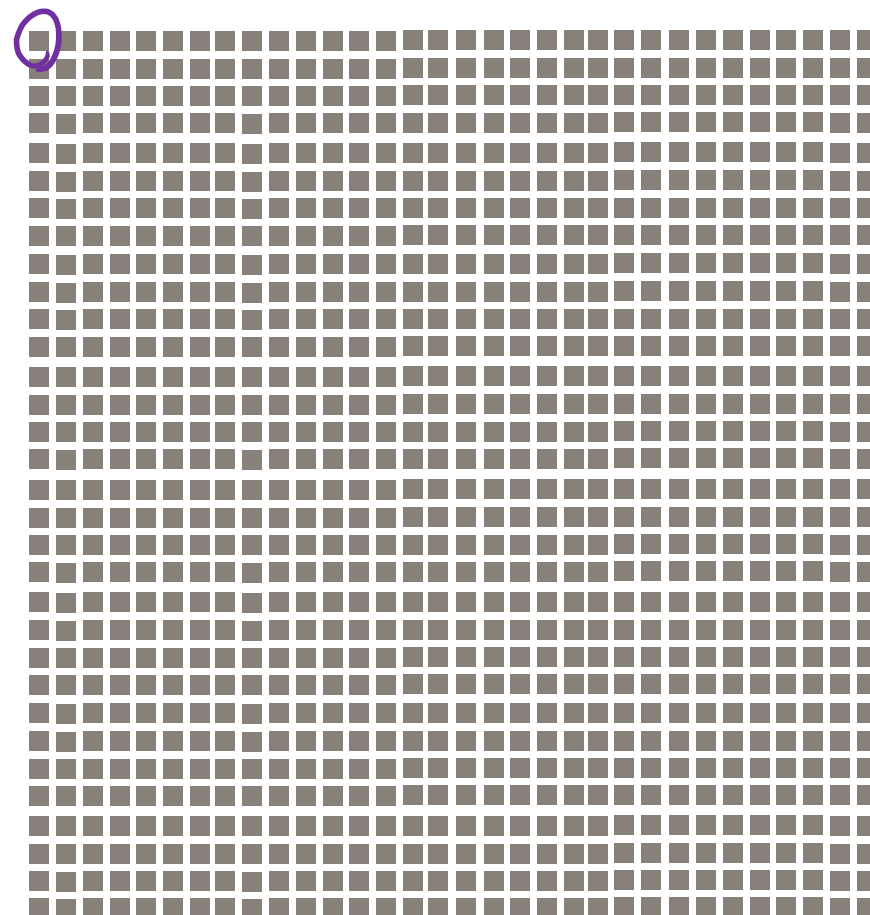
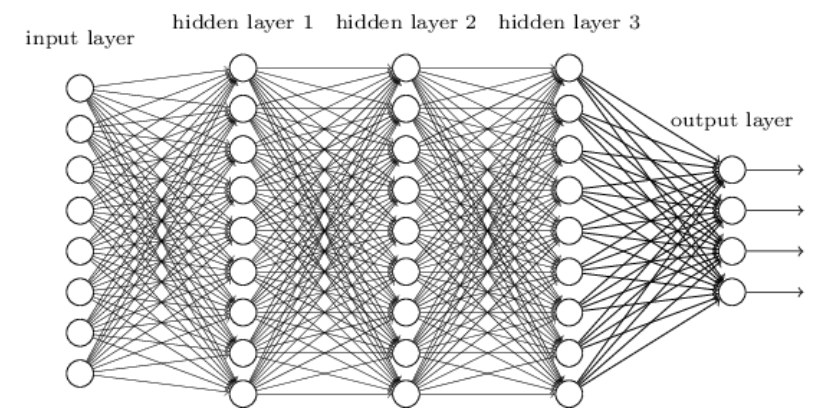
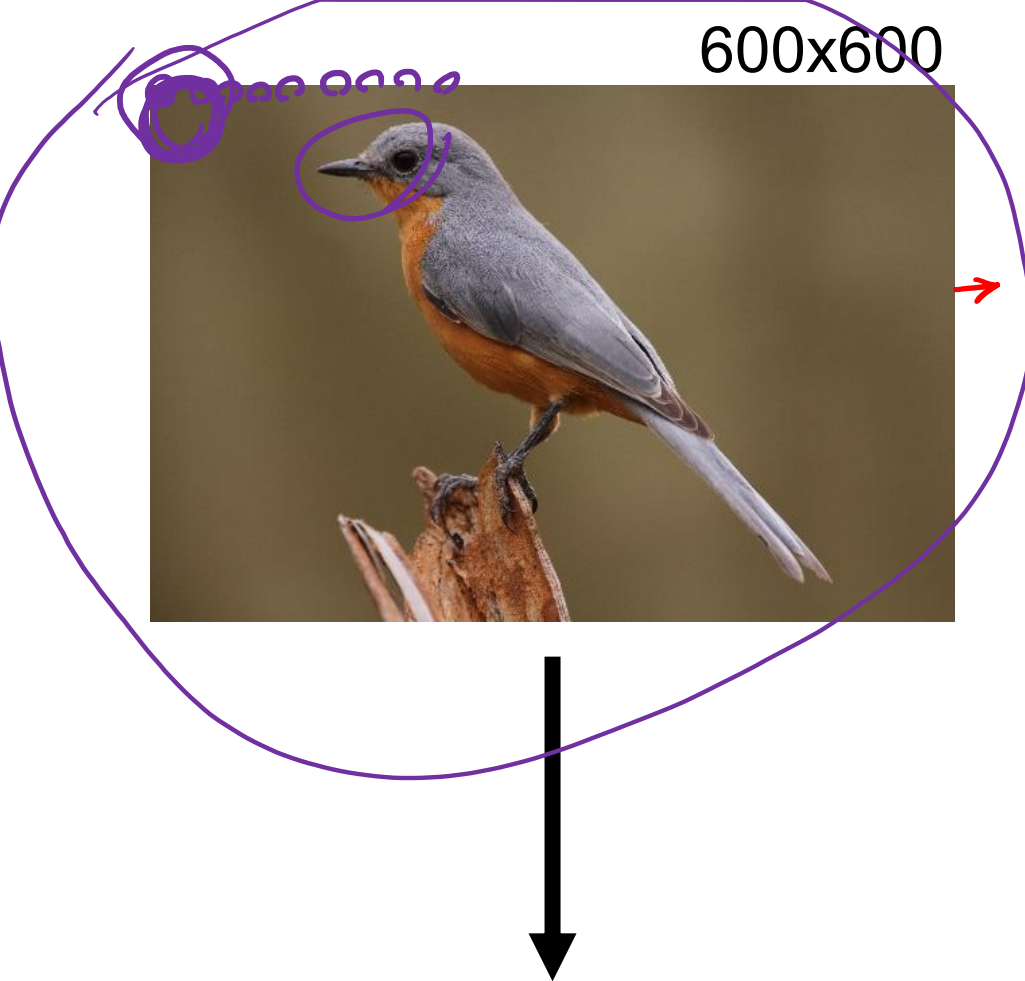


# Consider learning an image:

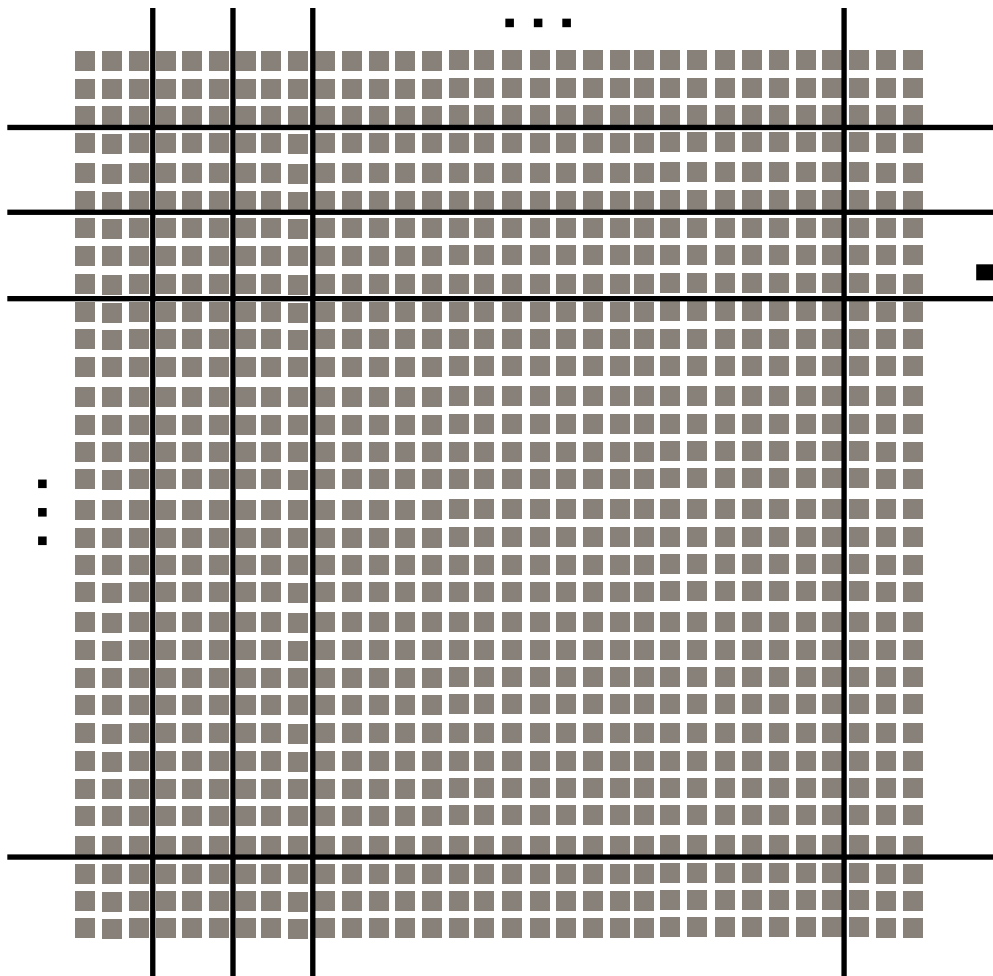
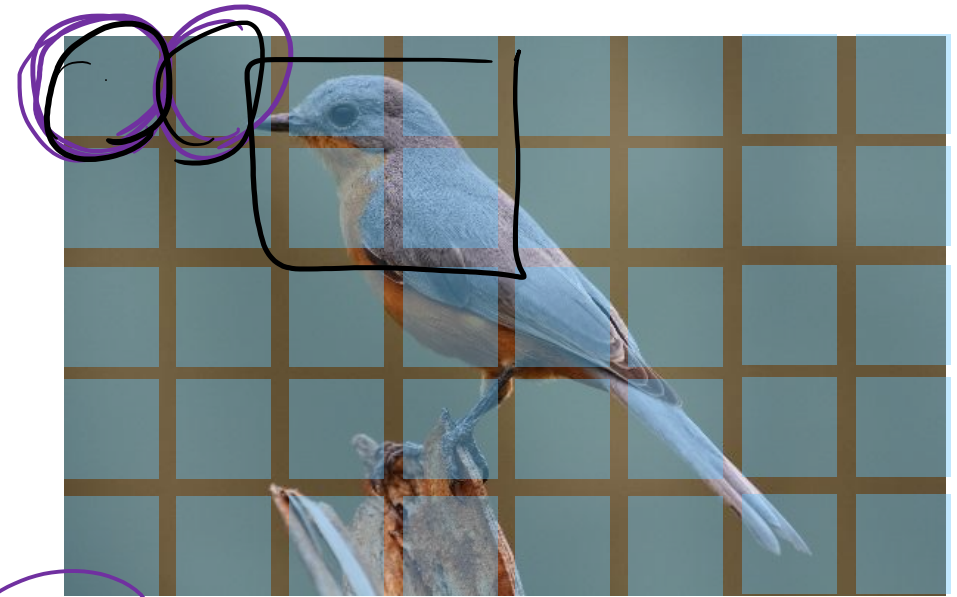
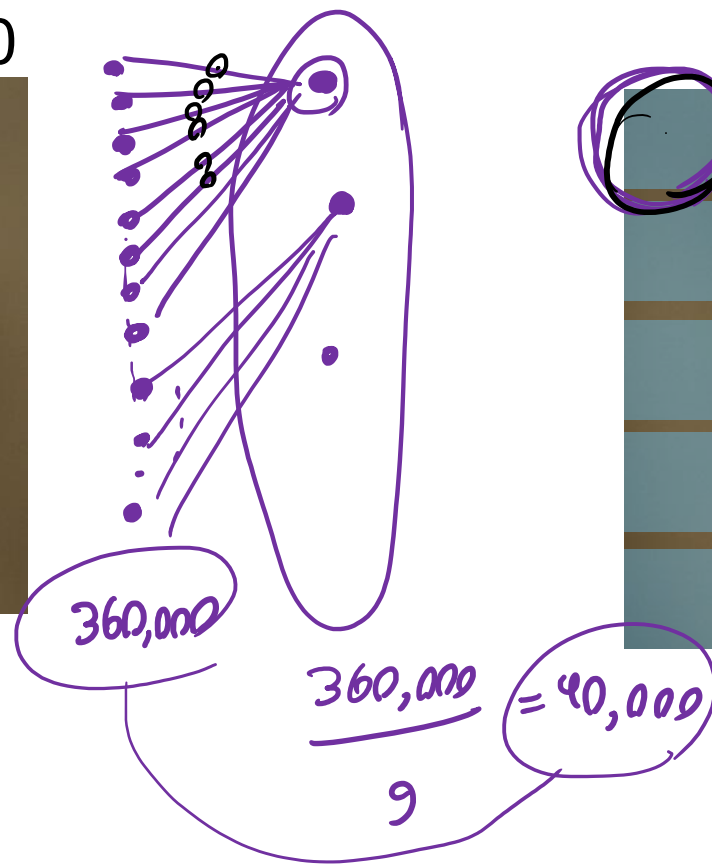
- Some patterns are much smaller than the whole image

Can represent a small region with fewer parameters





600x600



3x3



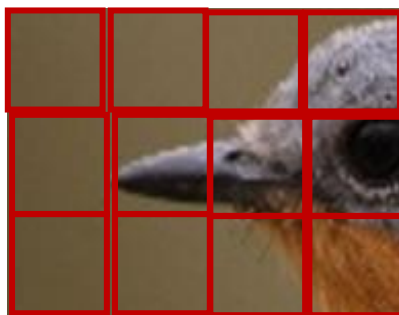
3x3



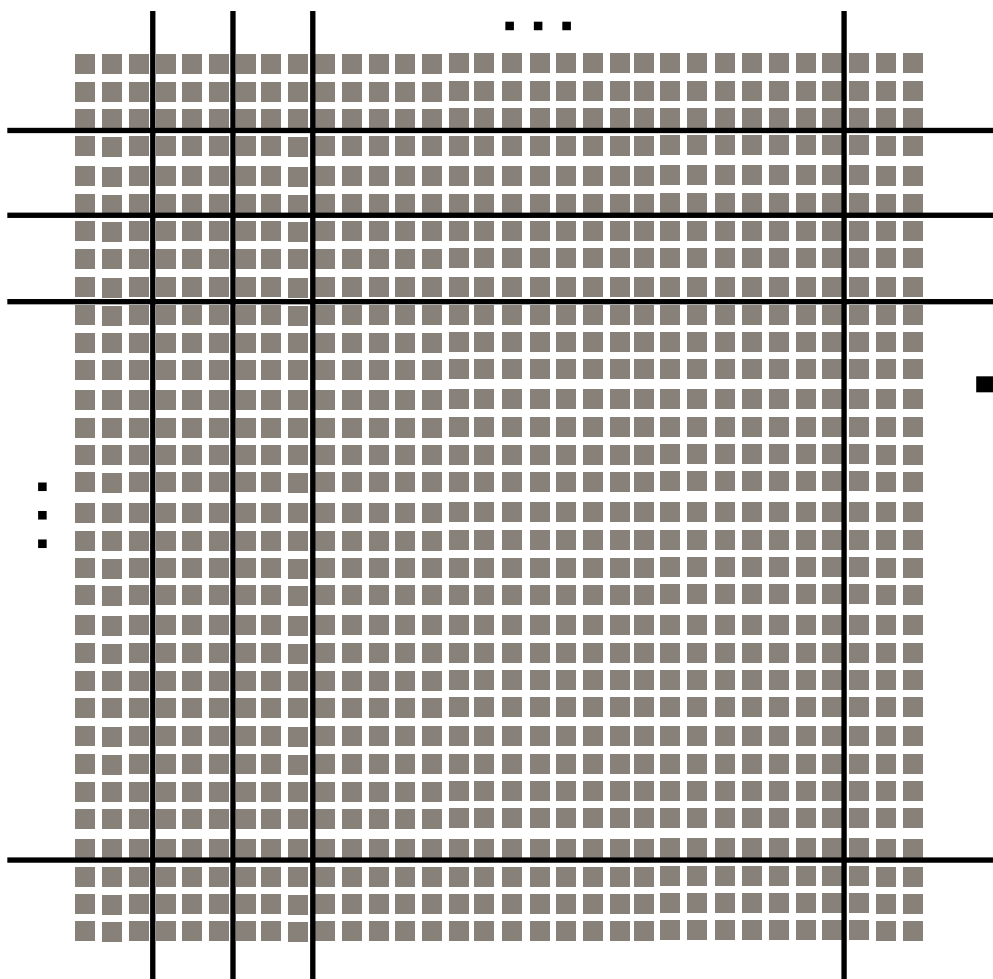
3x3

...

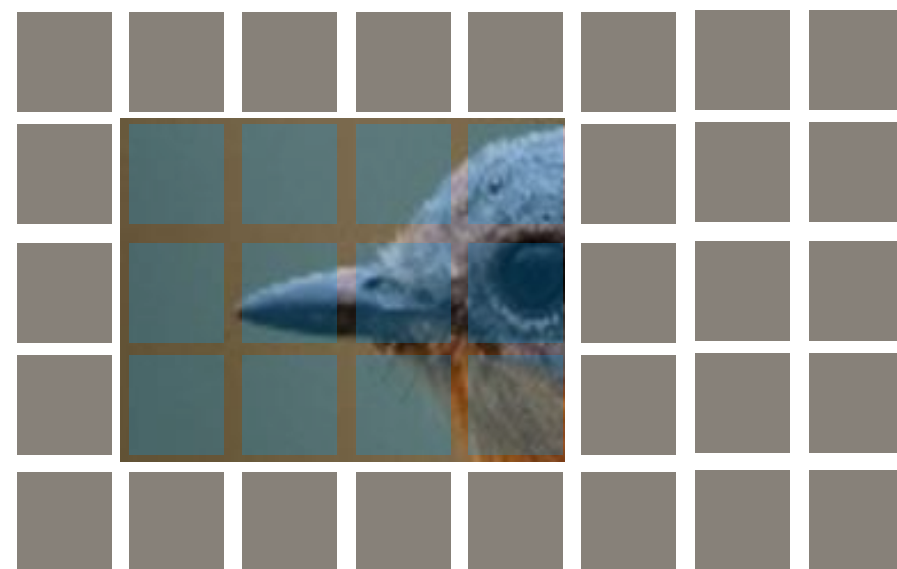




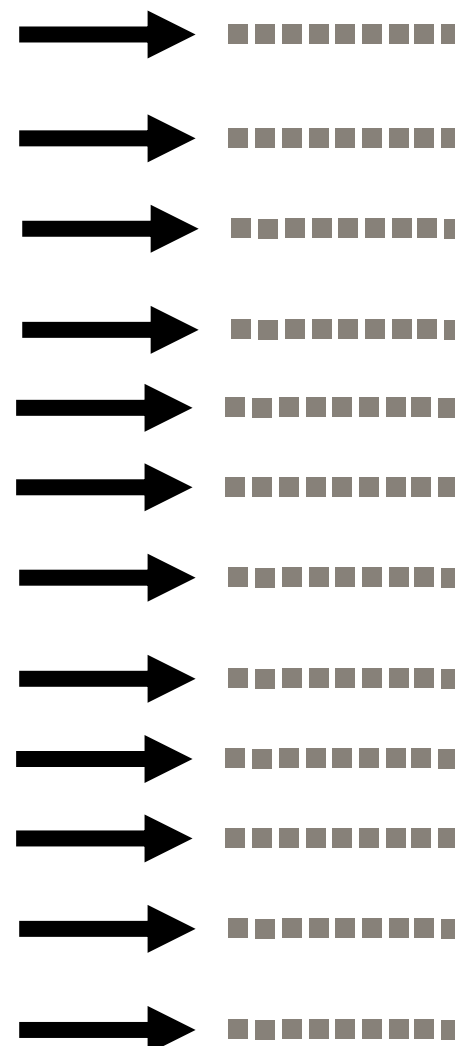
9x12



...

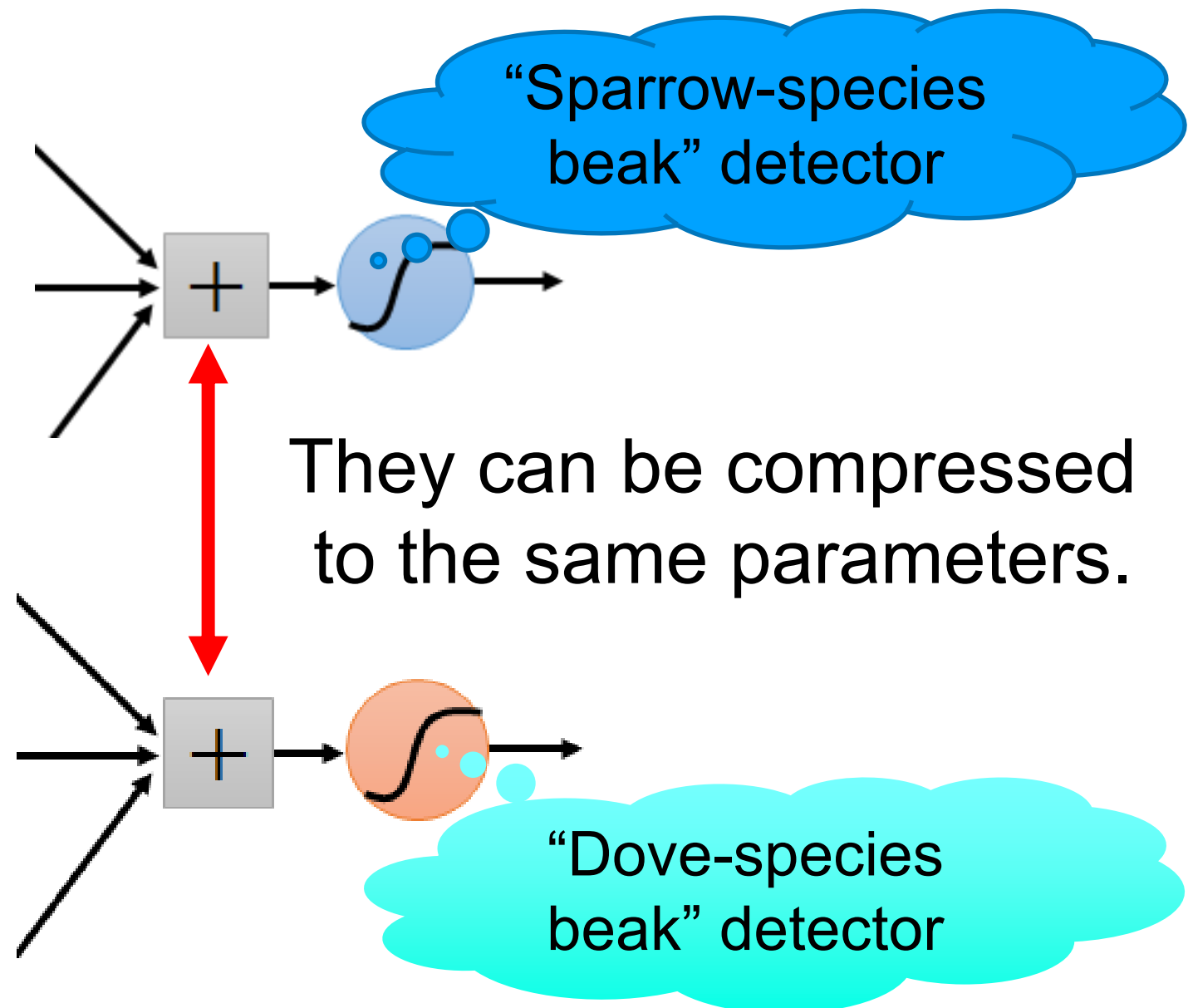
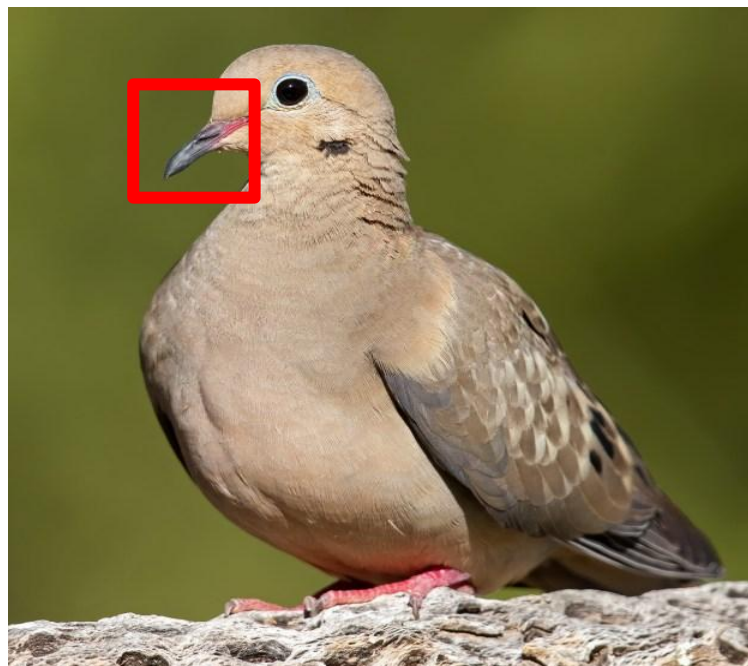
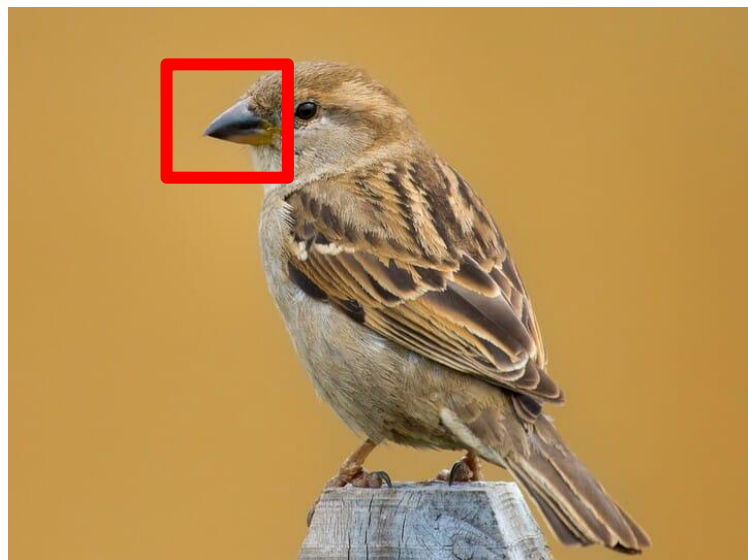


...



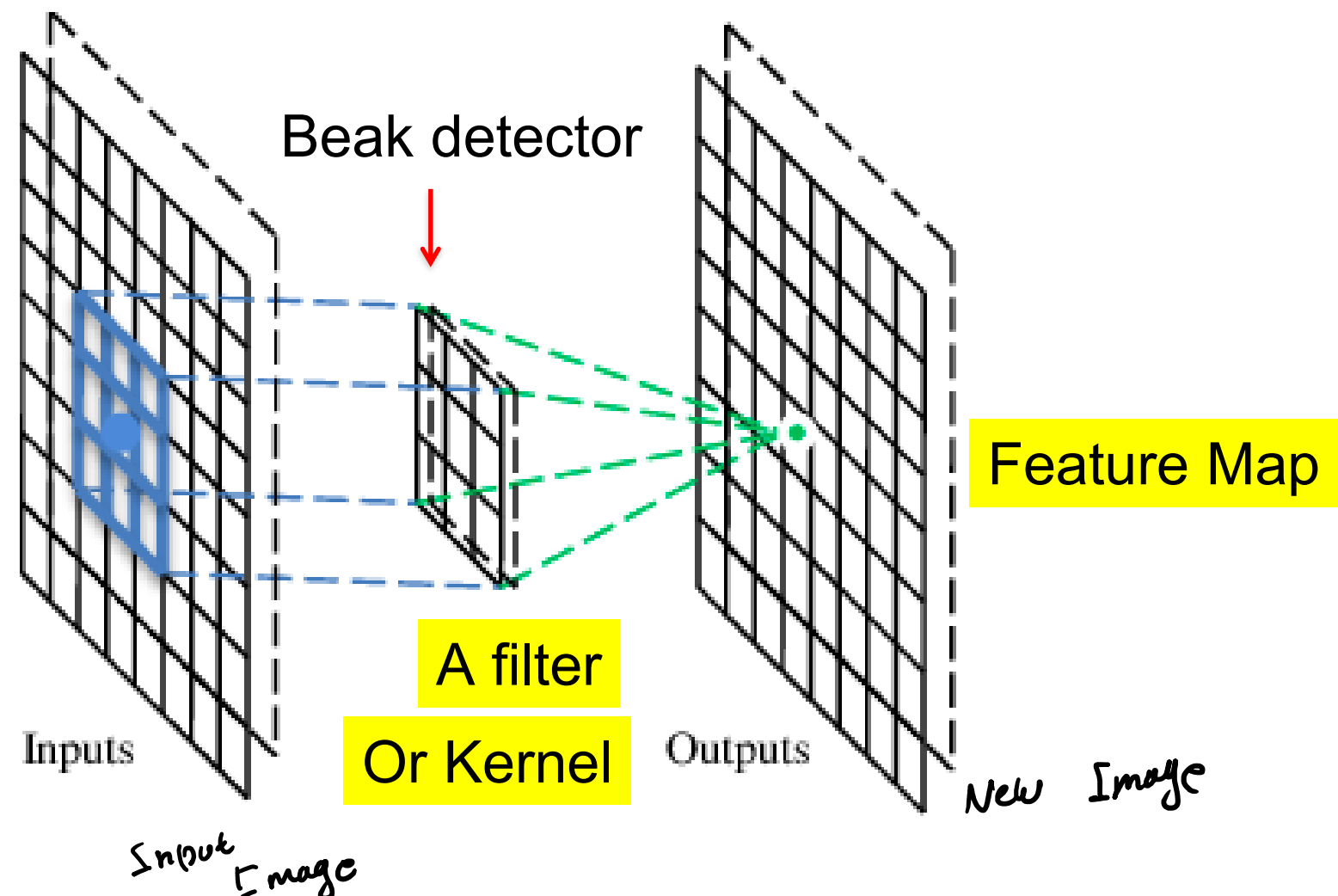
Same pattern appears in different places:  
They can be compressed!

What about training a lot of such “small” detectors  
and each detector must “move around”.



# A convolutional layer

A CNN is a neural network with some convolutional layers (and some other layers). A convolutional layer has a number of filters that does convolutional operation.



# Convolution

**These are the network parameters to be learned.**

|   |   |   |   |   |   |
|---|---|---|---|---|---|
| 1 | 0 | 0 | 0 | 0 | 1 |
| 0 | 1 | 0 | 0 | 1 | 0 |
| 0 | 0 | 1 | 1 | 0 | 0 |
| 1 | 0 | 0 | 0 | 1 | 0 |
| 0 | 1 | 0 | 0 | 1 | 0 |
| 0 | 0 | 1 | 0 | 1 | 0 |

6 x 6 image

|    |    |    |
|----|----|----|
| 1  | -1 | -1 |
| -1 | 1  | -1 |
| -1 | -1 | 1  |

Filter 1

|    |   |    |
|----|---|----|
| -1 | 1 | -1 |
| -1 | 1 | -1 |
| -1 | 1 | -1 |

Filter 2

⋮ ⋮

Each filter detects a small pattern (3 x 3).

# Convolution

$$\Theta_1 x_1 + \dots + \Theta_9 x_9 = LC$$

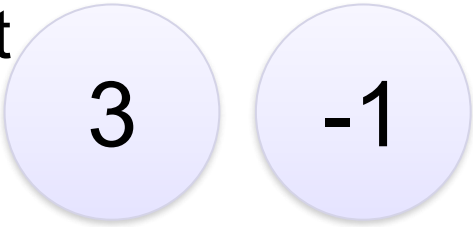
| $\theta_1$ | $\theta_2$ | $\theta_3$ |
|------------|------------|------------|
| 1          | -1         | -1         |
| -1         | 1          | -1         |
| -1         | -1         | 1          |

Filter 1

stride=1

| $x_1$ | $x_2$ | $x_3$ |   |   |   |
|-------|-------|-------|---|---|---|
| 1     | 0     | 0     | 0 | 0 | 1 |
| 0     | 1     | 0     | 0 | 1 | 0 |
| 0     | 0     | 1     | 1 | 0 | 0 |
| 1     | 0     | 0     | 0 | 1 | 0 |
| 0     | 1     | 0     | 0 | 1 | 0 |
| 0     | 0     | 1     | 0 | 1 | 0 |

Dot product  
→



6 x 6 image

# Convolution

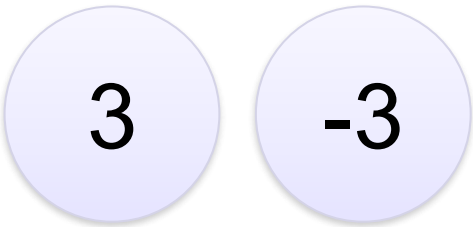
If stride=2

|   |   |   |   |   |   |
|---|---|---|---|---|---|
| 1 | 0 | 0 | 0 | 0 | 1 |
| 0 | 1 | 0 | 0 | 1 | 0 |
| 0 | 0 | 1 | 1 | 0 | 0 |
| 1 | 0 | 0 | 0 | 1 | 0 |
| 0 | 1 | 0 | 0 | 1 | 0 |
| 0 | 0 | 1 | 0 | 1 | 0 |

6 x 6 image

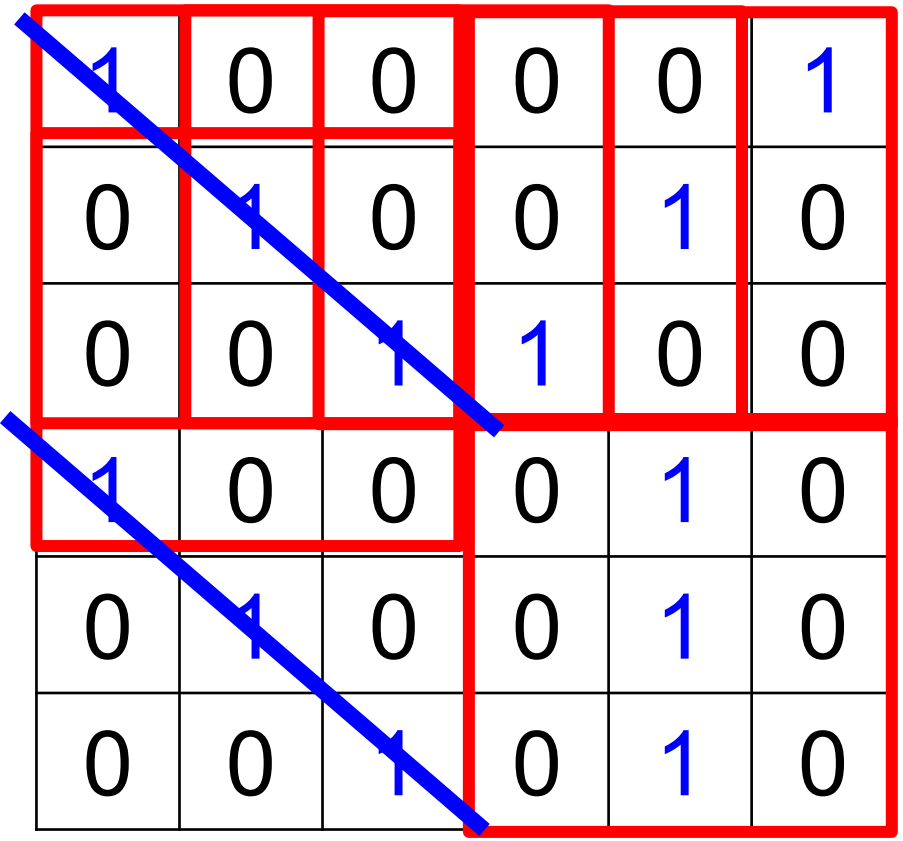
|    |    |    |
|----|----|----|
| 1  | -1 | -1 |
| -1 | 1  | -1 |
| -1 | -1 | 1  |

Filter 1

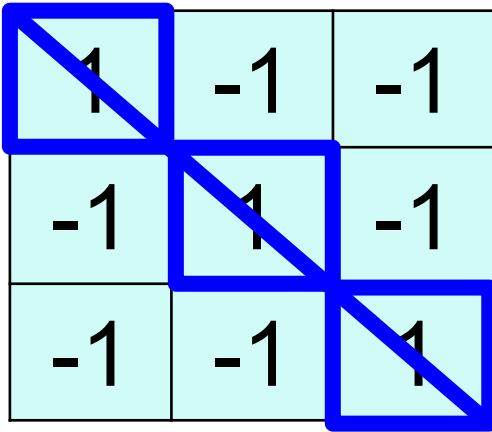


# Convolution

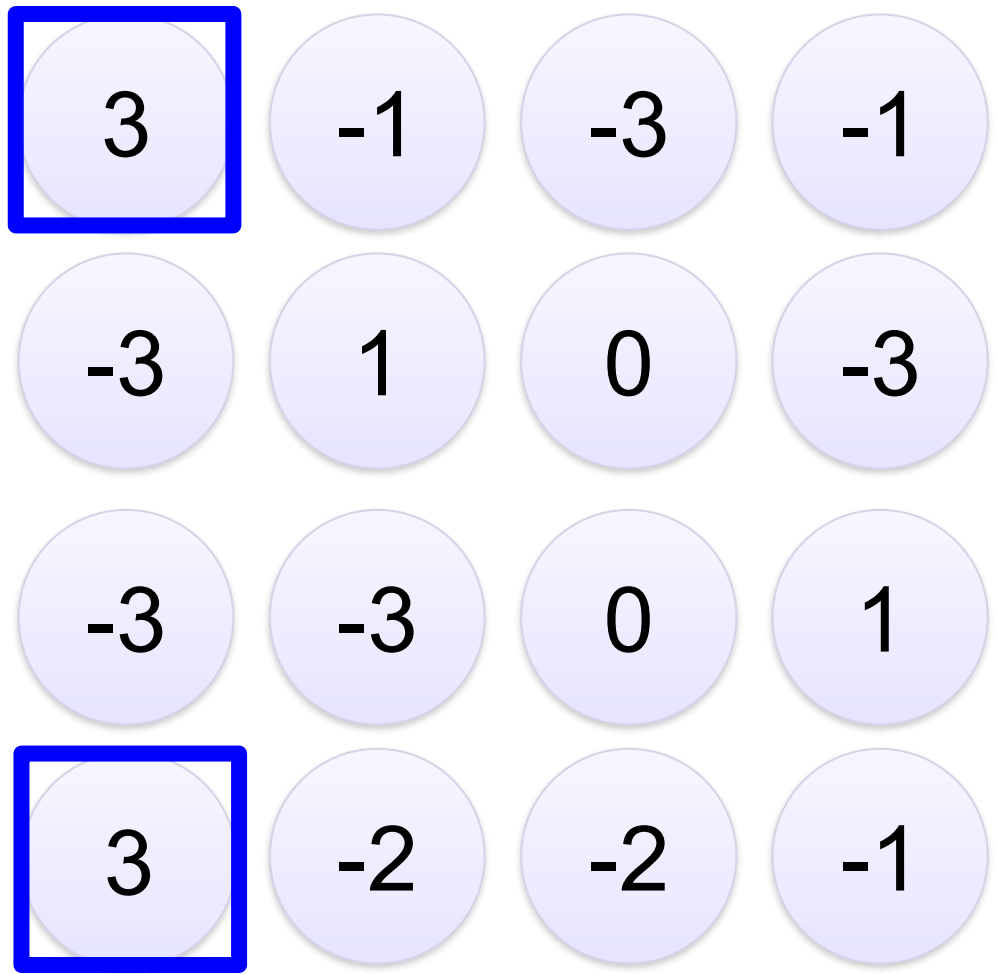
stride=1



6 x 6 image



Filter 1



# Convolution

stride=1

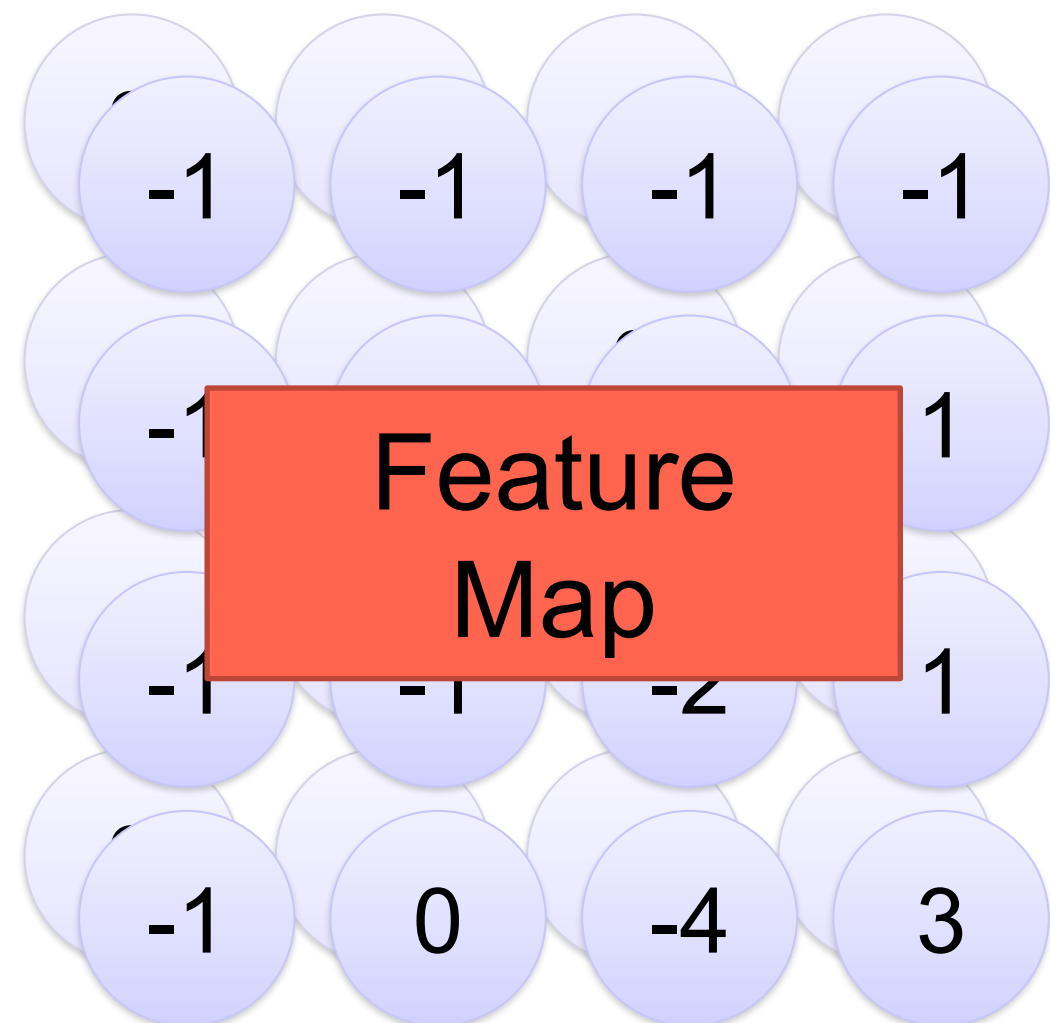
|   |   |   |   |   |   |
|---|---|---|---|---|---|
| 1 | 0 | 0 | 0 | 0 | 1 |
| 0 | 1 | 0 | 0 | 1 | 0 |
| 0 | 0 | 1 | 1 | 0 | 0 |
| 1 | 0 | 0 | 0 | 1 | 0 |
| 0 | 1 | 0 | 0 | 1 | 0 |
| 0 | 0 | 1 | 0 | 1 | 0 |

6 x 6 image

|    |   |    |
|----|---|----|
| -1 | 1 | -1 |
| -1 | 1 | -1 |
| -1 | 1 | -1 |

Filter 2

Repeat this for each filter



Two 4 x 4 images  
Forming 2 x 4 x 4 matrix

# Color image: RGB 3 channels

$$\text{Sub} = [1 \ 0 \ 0 \ 0 \ 1 \ 0 \ 0 \ 0 \ 1 \ 1 \ \dots]_{1 \times 27}$$

$$\text{F1} = [1 \ -1 \ -1 \ -1 \ \dots]_{1 \times 27}$$

Sub.  $\text{F1} \in \mathbb{R}$

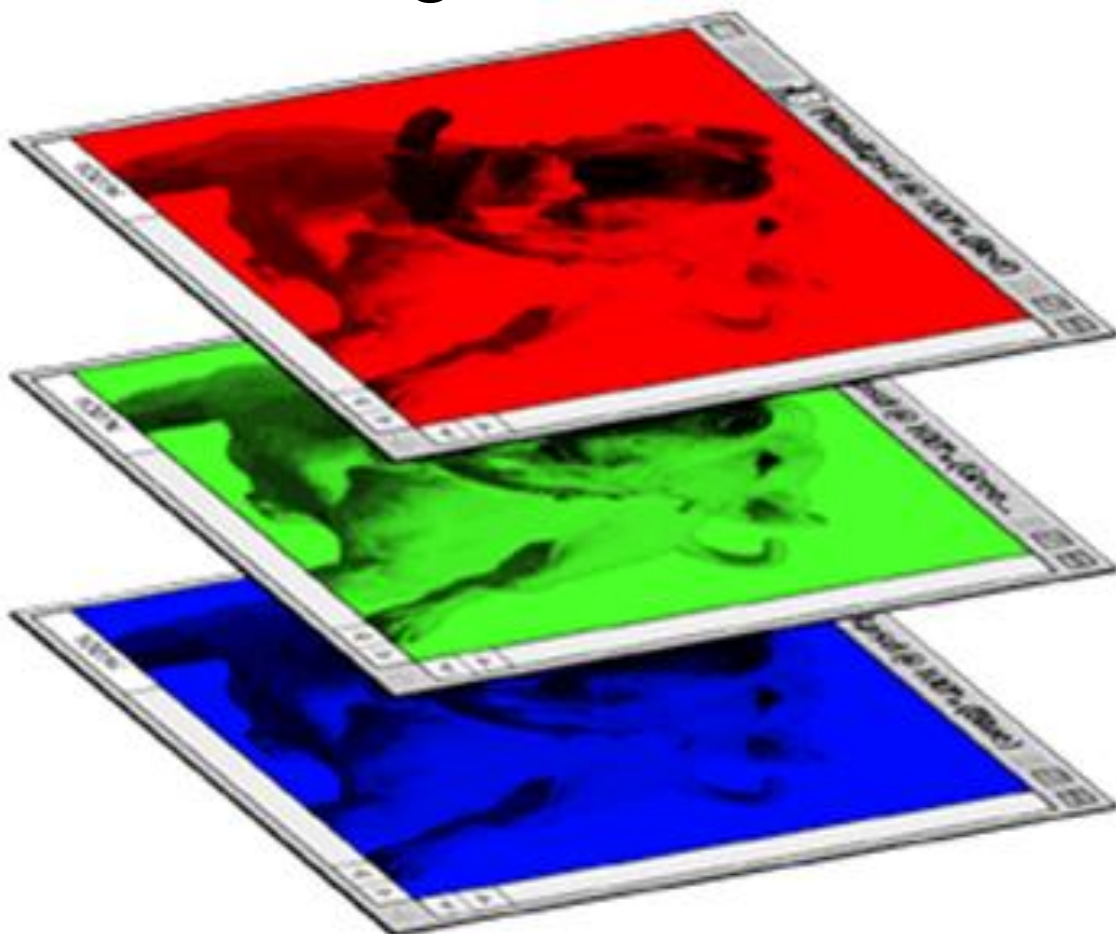
|    |    |    |
|----|----|----|
| 1  | -1 | -1 |
| -1 | 1  | -1 |
| -1 | -1 | 1  |

Filter 1

|    |   |    |
|----|---|----|
| -1 | 1 | -1 |
| -1 | 1 | -1 |
| -1 | 1 | -1 |

Filter 2

Color image



|   |   |   |   |   |   |
|---|---|---|---|---|---|
| 1 | 0 | 0 | 0 | 0 | 1 |
| 1 | 0 | 0 | 0 | 1 | 0 |
| 0 | 1 | 0 | 0 | 1 | 0 |
| 0 | 0 | 1 | 1 | 0 | 0 |
| 1 | 0 | 0 | 0 | 1 | 0 |
| 0 | 1 | 0 | 0 | 1 | 0 |
| 0 | 0 | 1 | 0 | 1 | 0 |

# Convolution v.s. Fully Connected

|   |   |   |   |   |   |
|---|---|---|---|---|---|
| 1 | 0 | 0 | 0 | 0 | 1 |
| 0 | 1 | 0 | 0 | 1 | 0 |
| 0 | 0 | 1 | 1 | 0 | 0 |
| 1 | 0 | 0 | 0 | 1 | 0 |
| 0 | 1 | 0 | 0 | 1 | 0 |
| 0 | 0 | 1 | 0 | 1 | 0 |

image

|    |    |    |
|----|----|----|
| 1  | -1 | -1 |
| -1 | 1  | -1 |
| -1 | -1 | 1  |

|    |   |    |
|----|---|----|
| -1 | 1 | -1 |
| -1 | 1 | -1 |
| -1 | 1 | -1 |

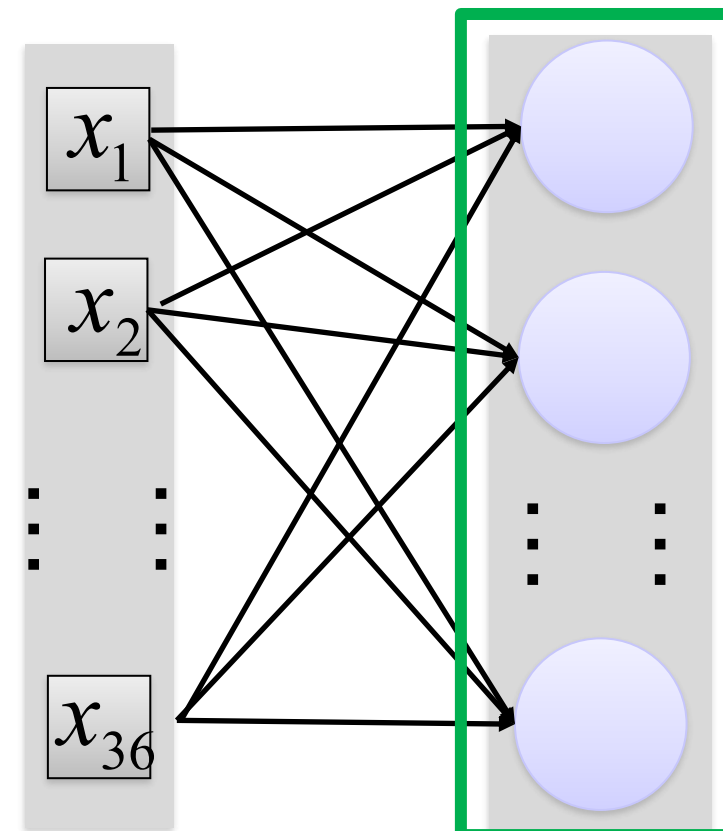


convolution

|    |    |    |    |
|----|----|----|----|
| -1 | -1 | -1 | -1 |
| -1 | -1 | -2 | 1  |
| -1 | -1 | -2 | 1  |
| -1 | 0  | -4 | 3  |

Fully-  
connected

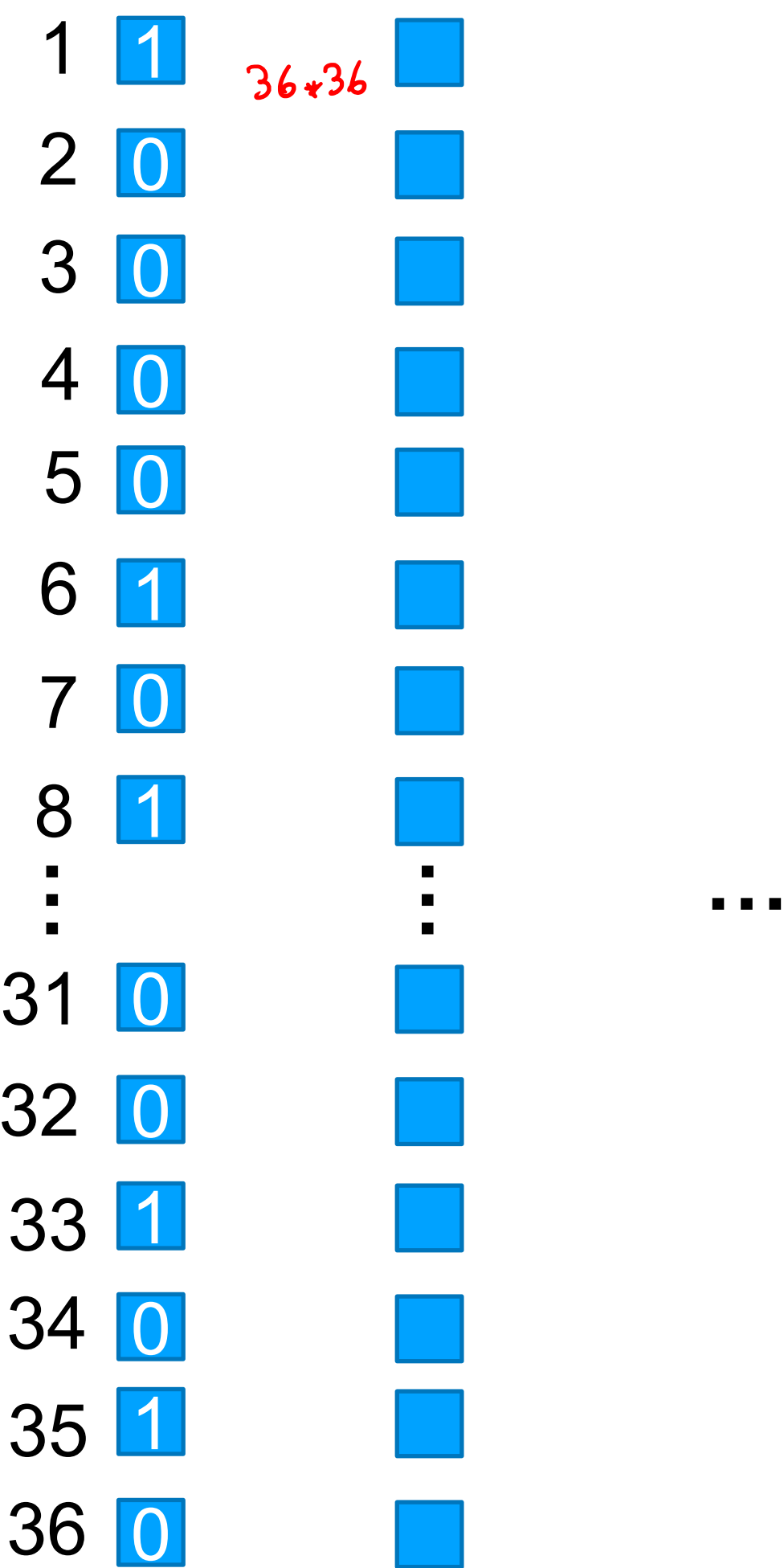
|   |   |   |   |   |   |
|---|---|---|---|---|---|
| 1 | 0 | 0 | 0 | 0 | 1 |
| 0 | 1 | 0 | 0 | 1 | 0 |
| 0 | 0 | 1 | 1 | 0 | 0 |
| 1 | 0 | 0 | 0 | 1 | 0 |
| 0 | 1 | 0 | 0 | 1 | 0 |
| 0 | 0 | 1 | 0 | 1 | 0 |



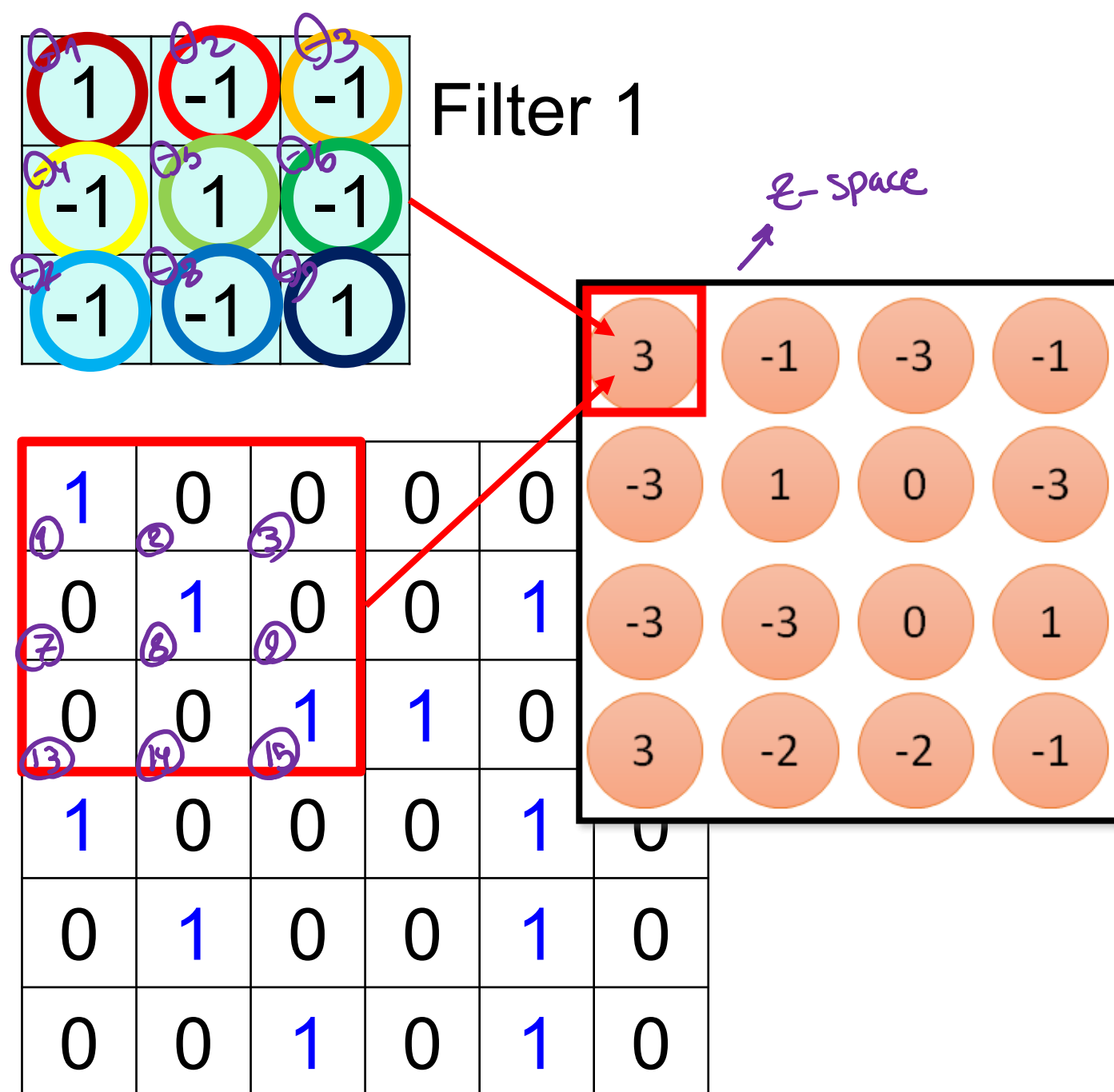
Conventional  
Fully Connected  
layers  
(FC layers)

|   |   |   |   |   |   |
|---|---|---|---|---|---|
| 1 | 0 | 0 | 0 | 0 | 1 |
| 0 | 1 | 0 | 0 | 1 | 0 |
| 0 | 0 | 1 | 1 | 0 | 0 |
| 1 | 0 | 0 | 0 | 1 | 0 |
| 0 | 1 | 0 | 0 | 1 | 0 |
| 0 | 0 | 1 | 0 | 1 | 0 |

6 x 6 image



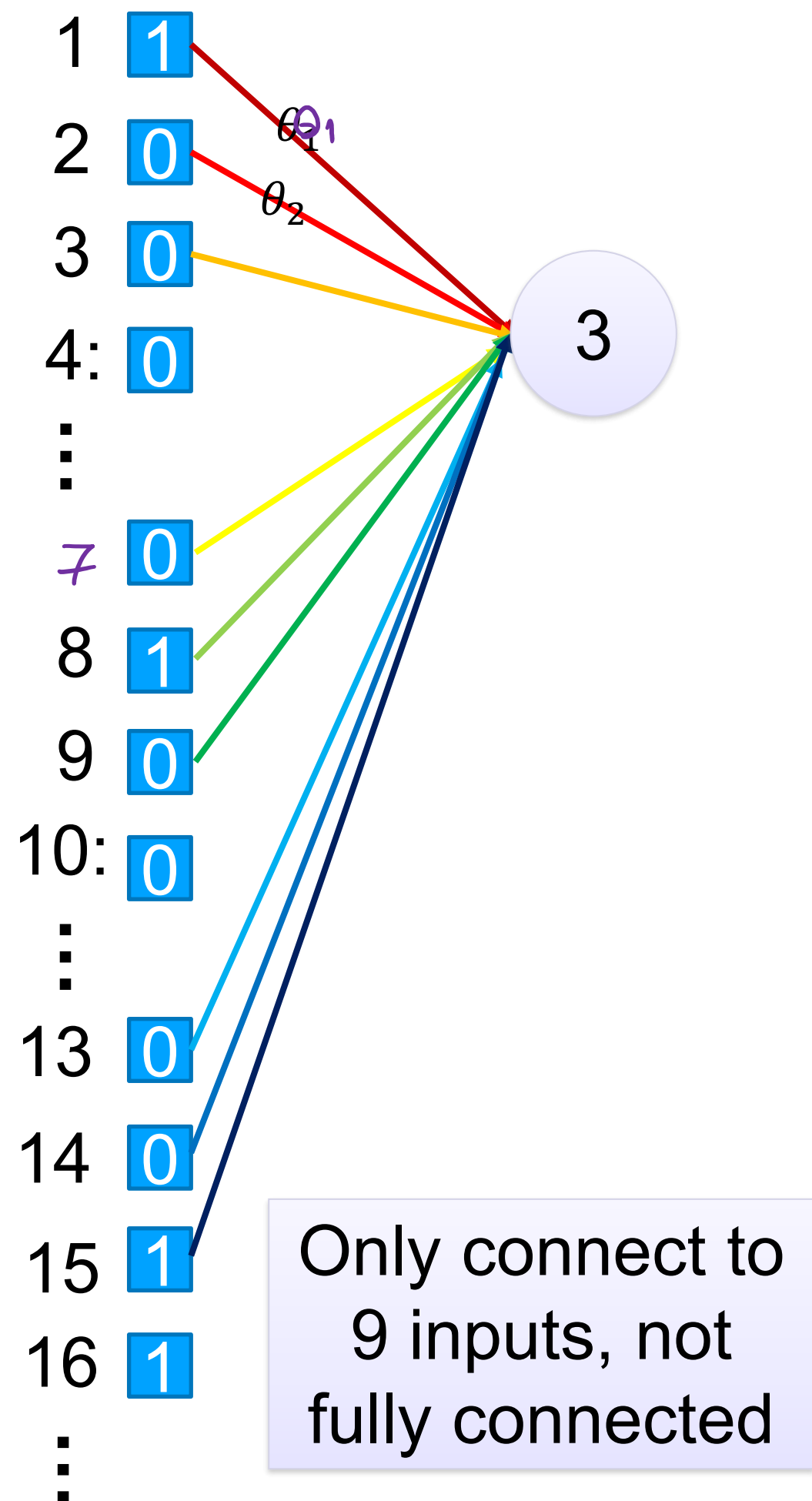
features 1<sup>st</sup> hidden layer

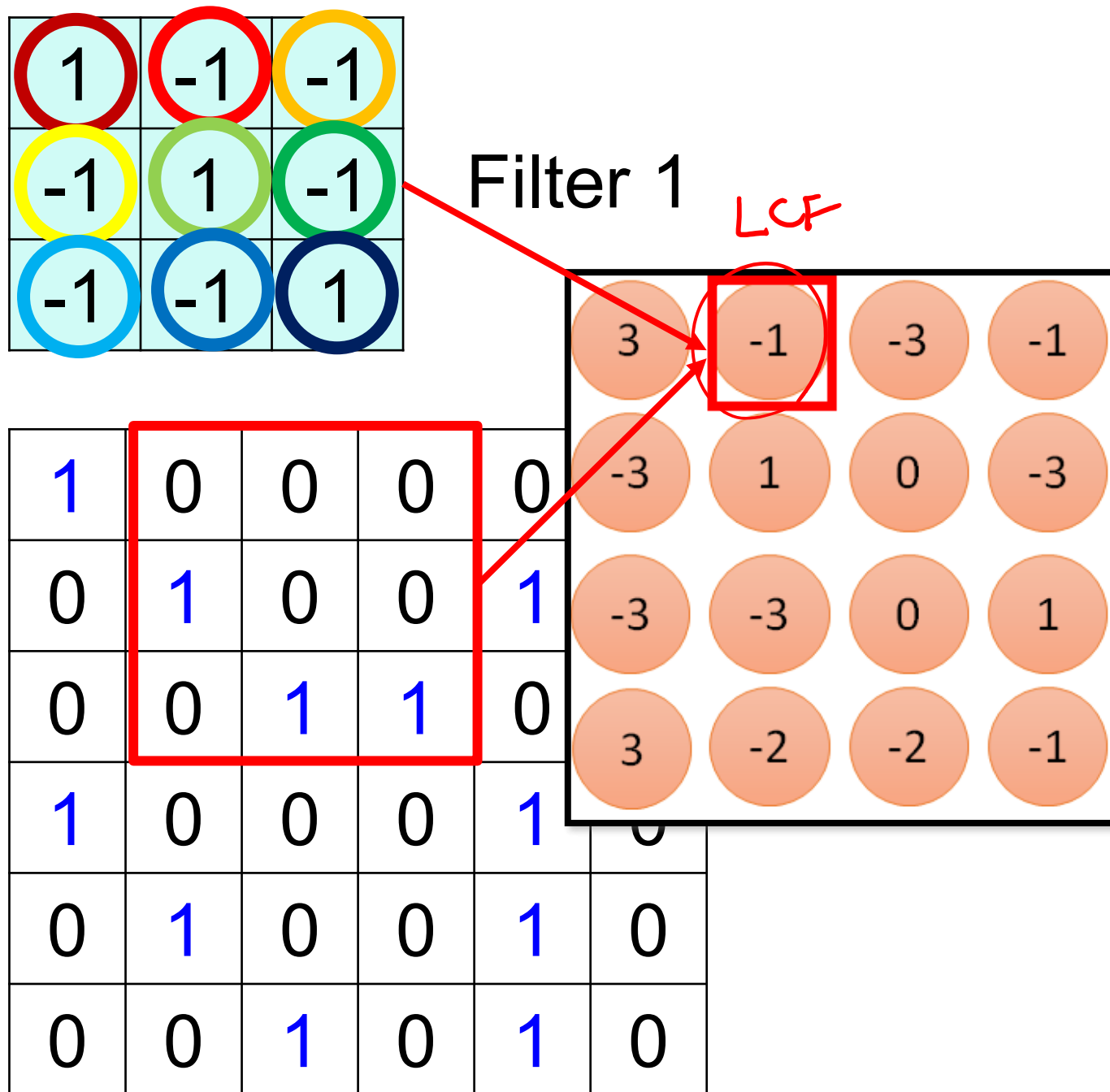


$x$   
space

6 x 6 image

fewer parameters!

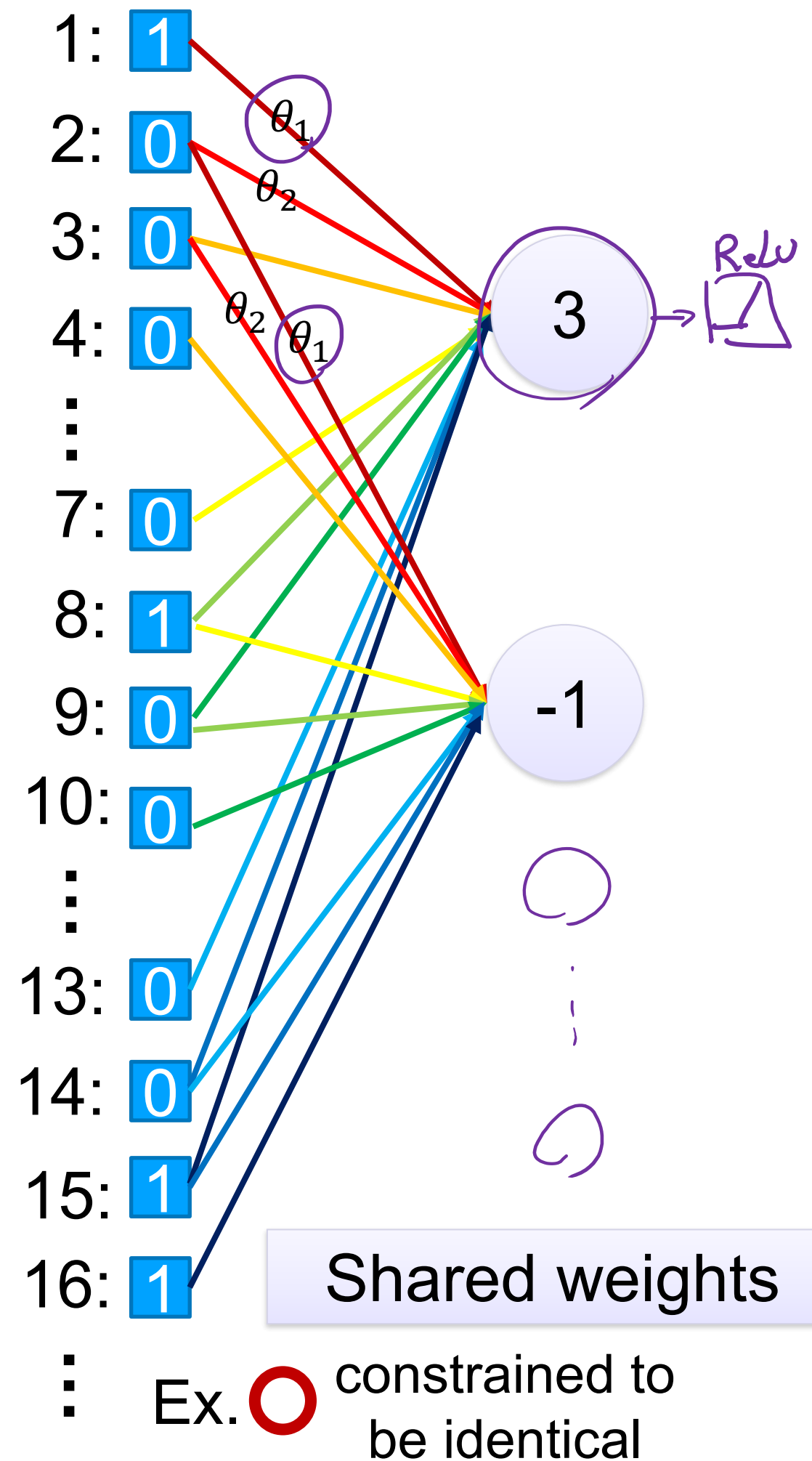




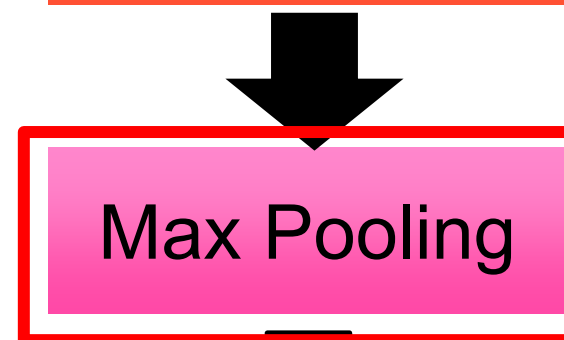
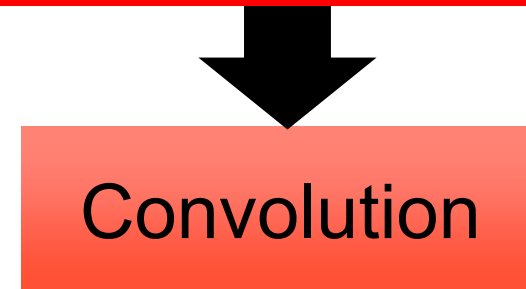
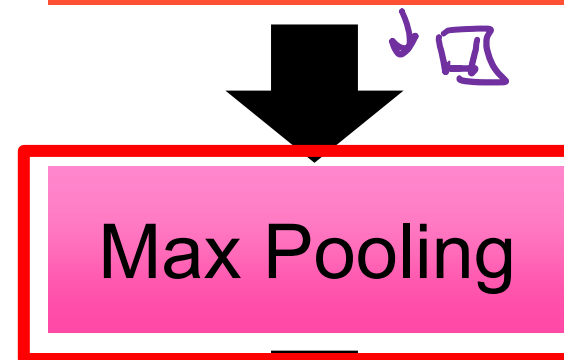
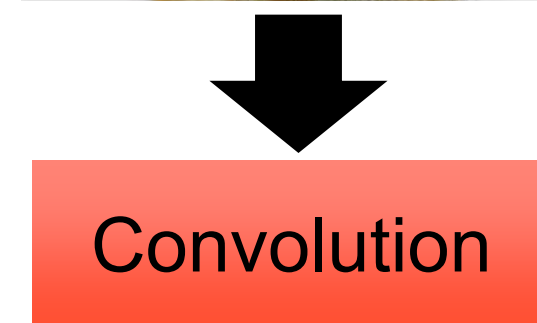
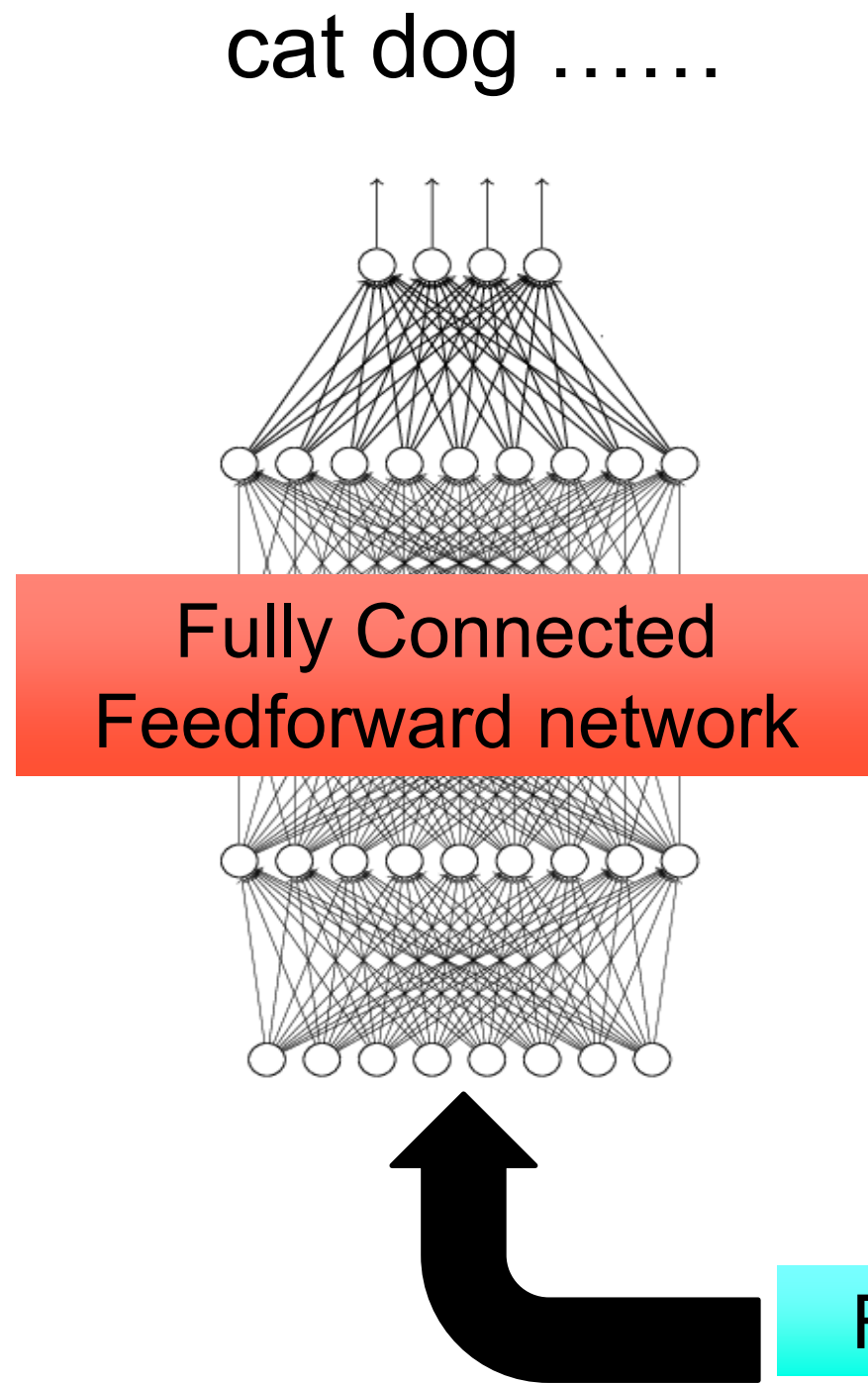
6 x 6 image

Fewer parameters

Even fewer parameters



# The whole CNN



Can repeat many times

Flattened

# Max Pooling

|    |    |    |
|----|----|----|
| 1  | -1 | -1 |
| -1 | 1  | -1 |
| -1 | -1 | 1  |

Filter 1

|    |   |    |
|----|---|----|
| -1 | 1 | -1 |
| -1 | 1 | -1 |
| -1 | 1 | -1 |

Filter 2

|    |    |    |    |
|----|----|----|----|
| 3  | -1 | -3 | -1 |
| -3 | 1  | 0  | -3 |
| -3 | -3 | 0  | 1  |
| 3  | -2 | -2 | -1 |

|    |    |    |    |
|----|----|----|----|
| -1 | -1 | -1 | -1 |
| -1 | -1 | -2 | 1  |
| -1 | -1 | -2 | 1  |
| -1 | 0  | -4 | 3  |

# Why Pooling

- Subsampling pixels will not change the object  
bird

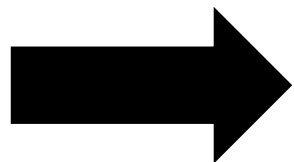


Subsampling



bird

We can subsample the pixels to make image smaller



fewer parameters to characterize the image

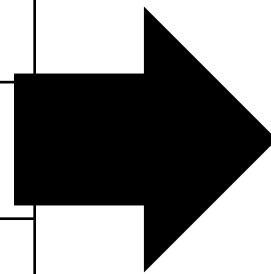
# A CNN compresses a fully connected network in three ways:

- Reducing number of connections
- Shared weights on the edges
- Max pooling further reduces the complexity

# Max Pooling

|   |   |   |   |   |   |
|---|---|---|---|---|---|
| 1 | 0 | 0 | 0 | 0 | 1 |
| 0 | 1 | 0 | 0 | 1 | 0 |
| 0 | 0 | 1 | 1 | 0 | 0 |
| 1 | 0 | 0 | 0 | 1 | 0 |
| 0 | 1 | 0 | 0 | 1 | 0 |
| 0 | 0 | 1 | 0 | 1 | 0 |

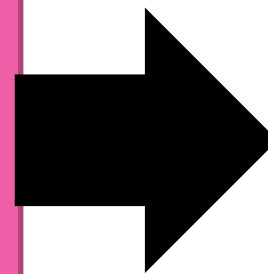
6 x 6 image



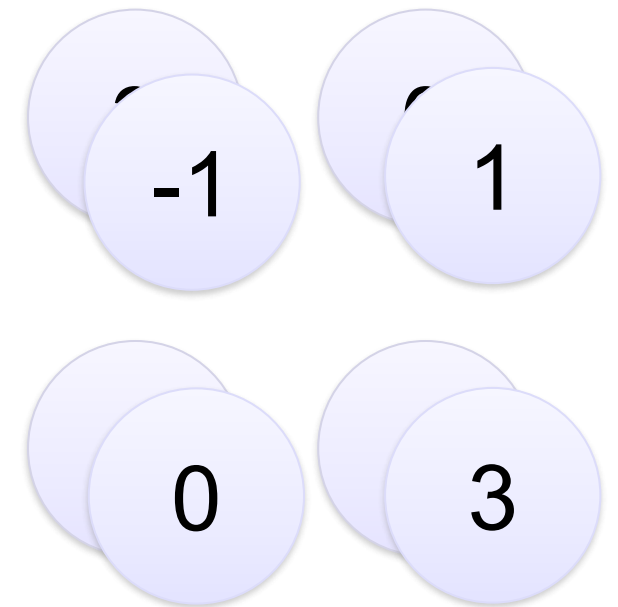
Conv



Max  
Pooling



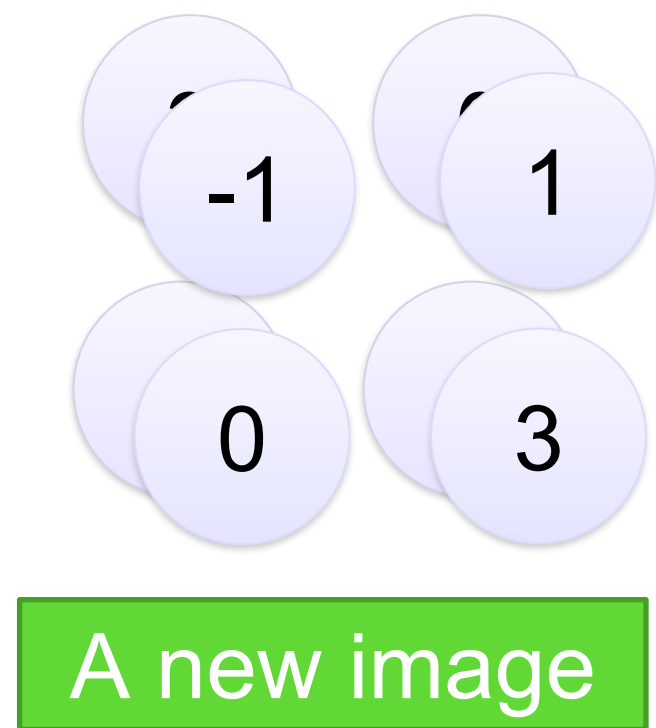
New image  
but smaller



2 x 2 image

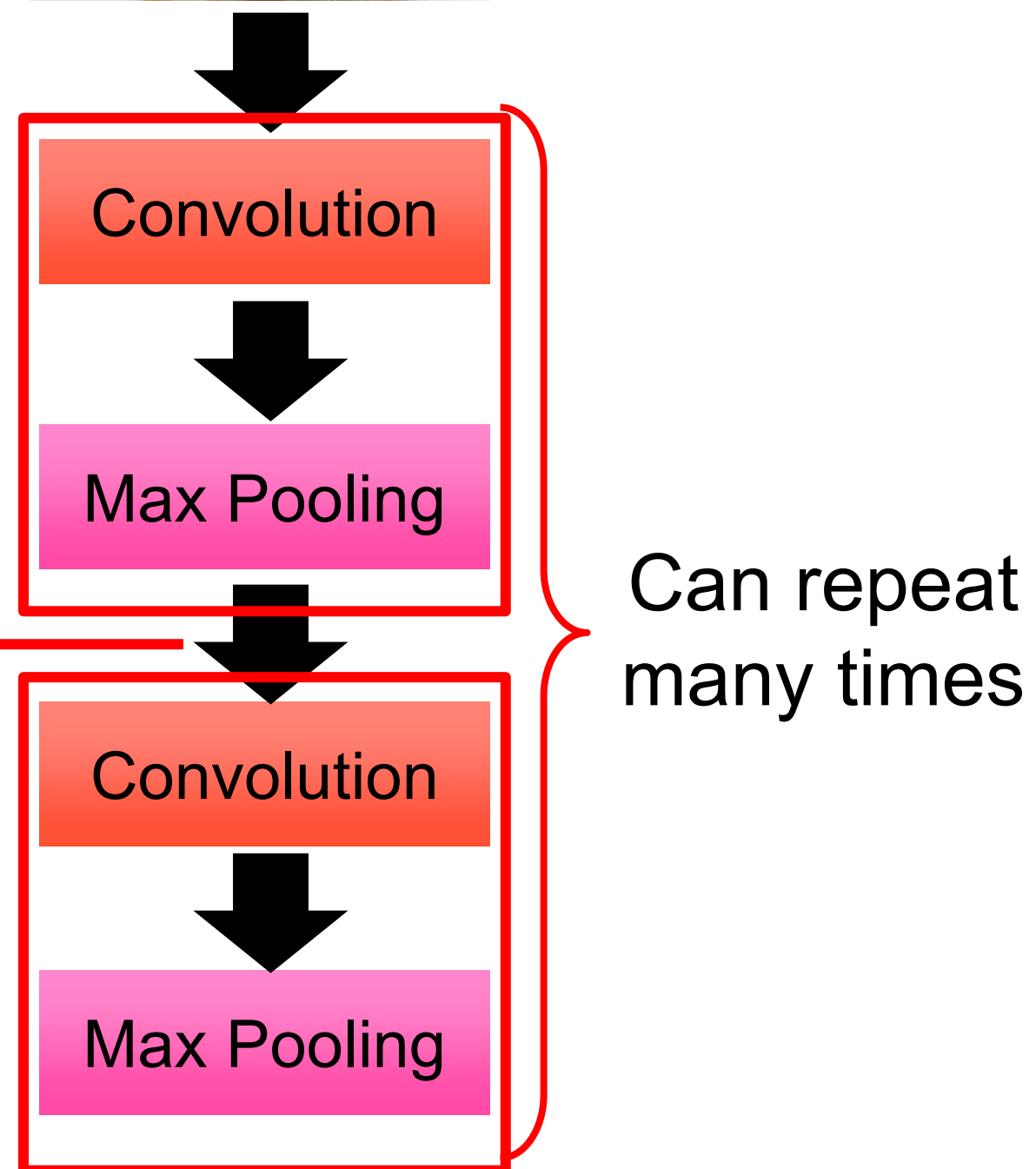
Each filter  
is a channel

# The whole CNN



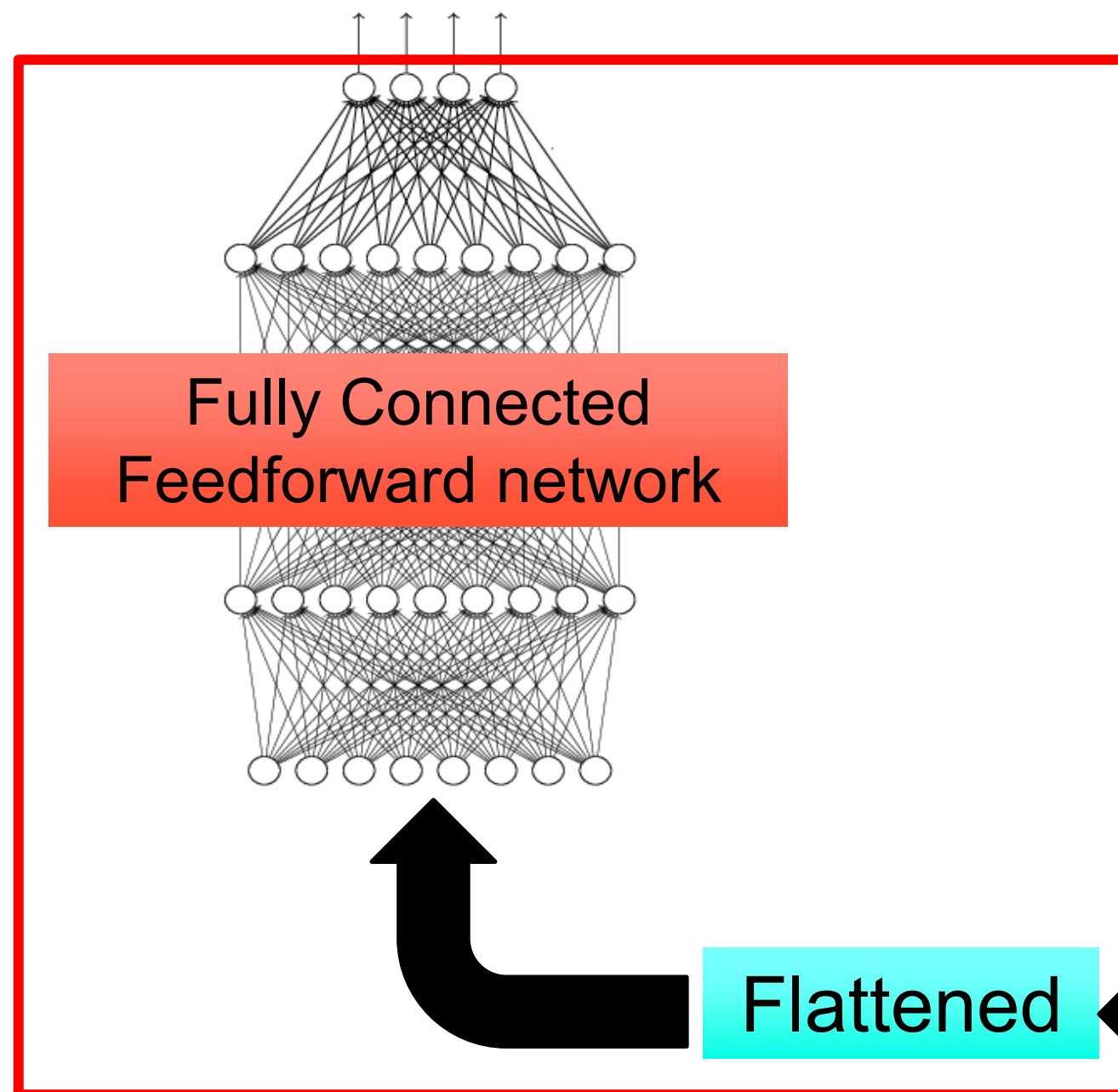
Smaller than the original image

The number of channels is the number of filters



# The whole CNN

cat dog .....



Convolution

Max Pooling

A new image

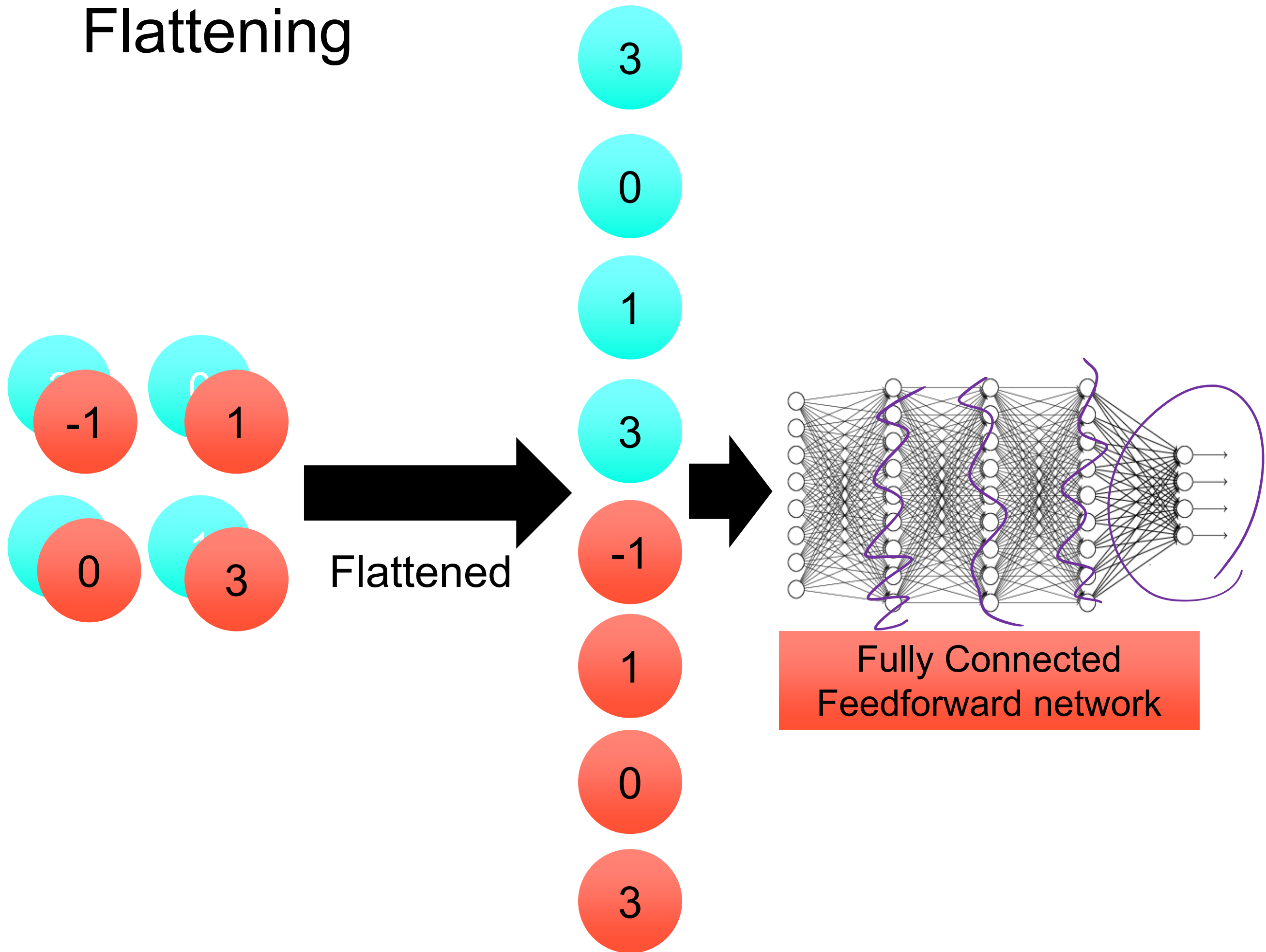
Convolution

Max Pooling

A new image

Flattened

# Flattening



# CNN in Keras

Only modified the *network structure* and *input format* (vector -> 3-D tensor)

```
model2.add( Convolution2D( 25,3,3,  
                           input_shape=(28,28,1)) )
```

|    |    |    |   |    |
|----|----|----|---|----|
| 1  | -1 | -1 | 1 | -1 |
| -1 | 1  | -1 | 1 | -1 |
| -1 | -1 | -1 | 1 | -1 |

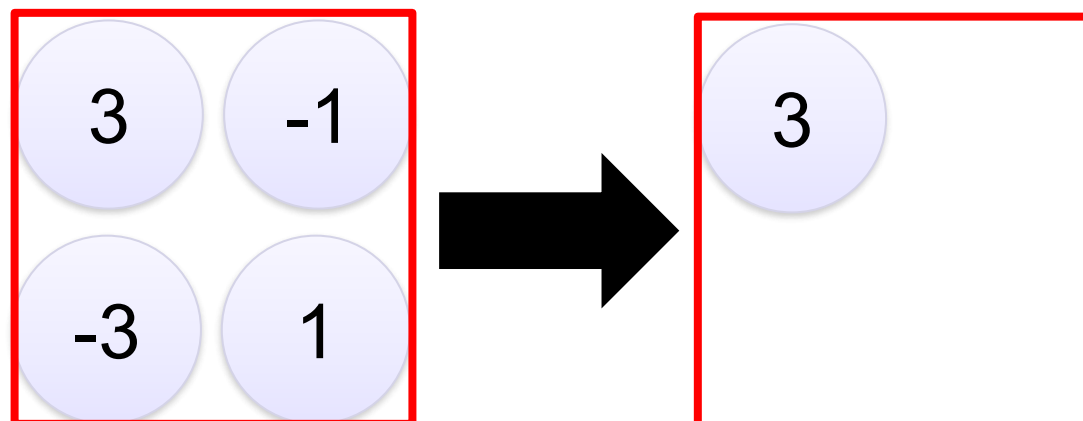
There are  
**25 3x3**  
filters.

Input\_shape = ( 28 , 28 , 1)

28 x 28 pixels

1: black/white, 3: RGB

```
model2.add(MaxPooling2D((2,2)))
```



input

Convolution

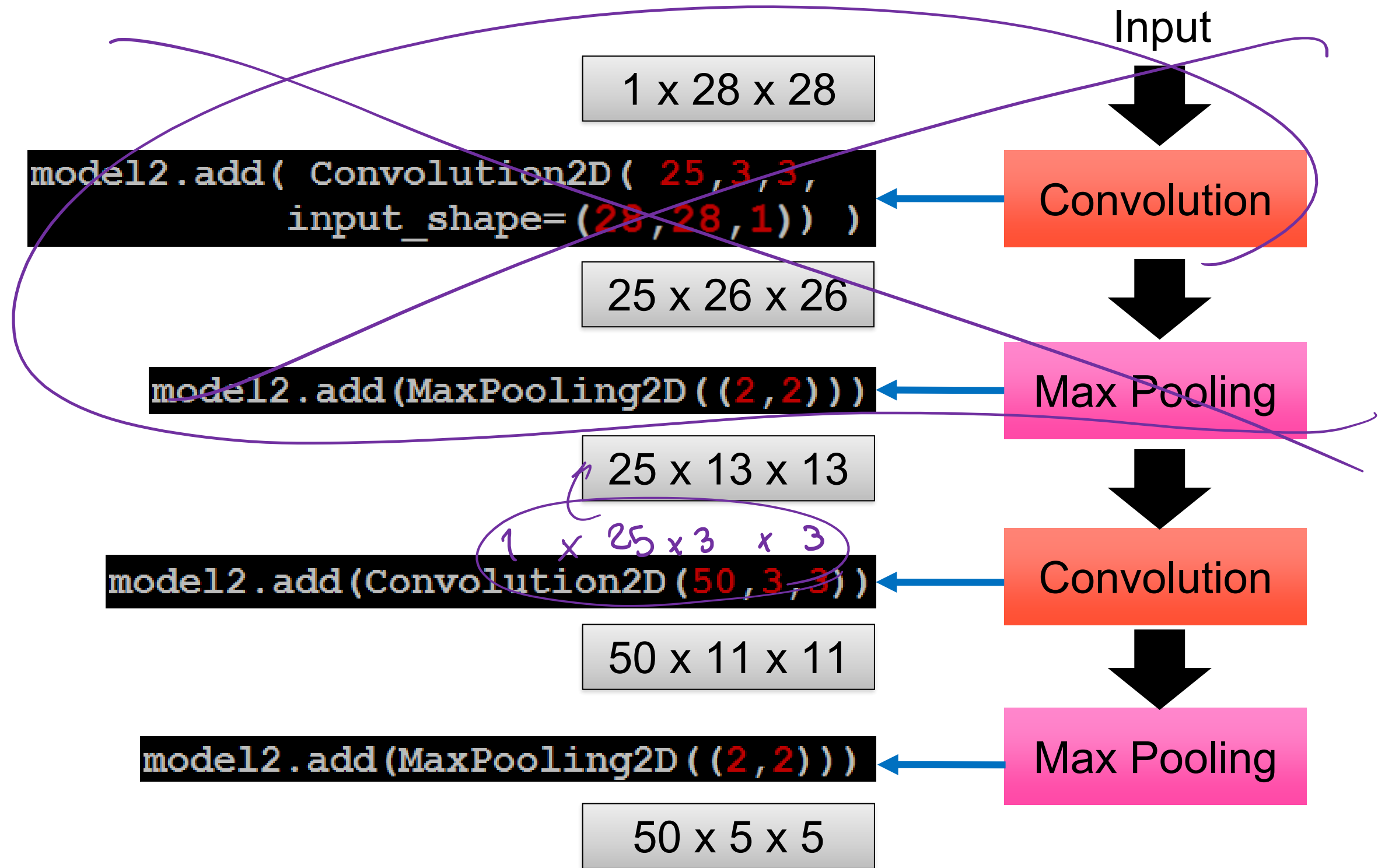
Max Pooling

Convolution

Max Pooling

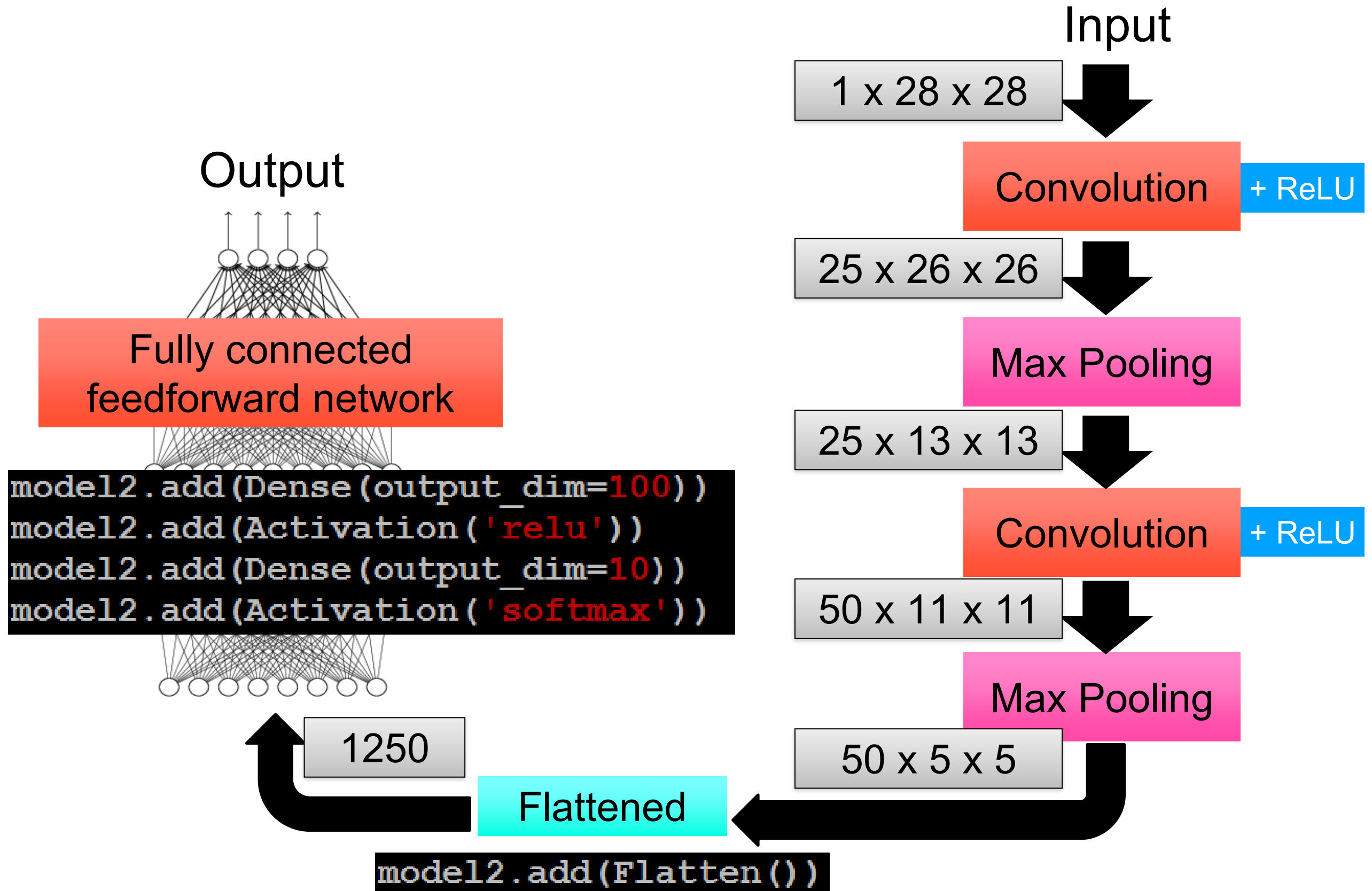
# CNN in Keras

Only modified the *network structure* and *input format* (vector -> 3-D array)



# CNN in Keras

Only modified the *network structure* and *input format* (vector -> 3-D array)



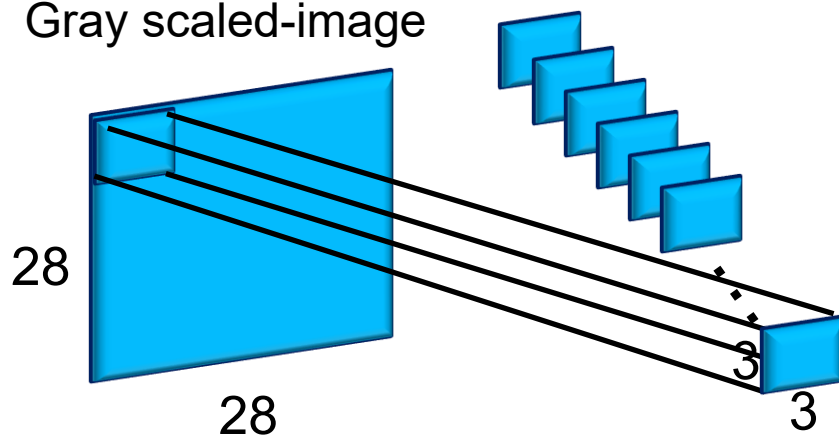
# Number of Parameters

$25 \times 3 \times 3 + 25$  parameters

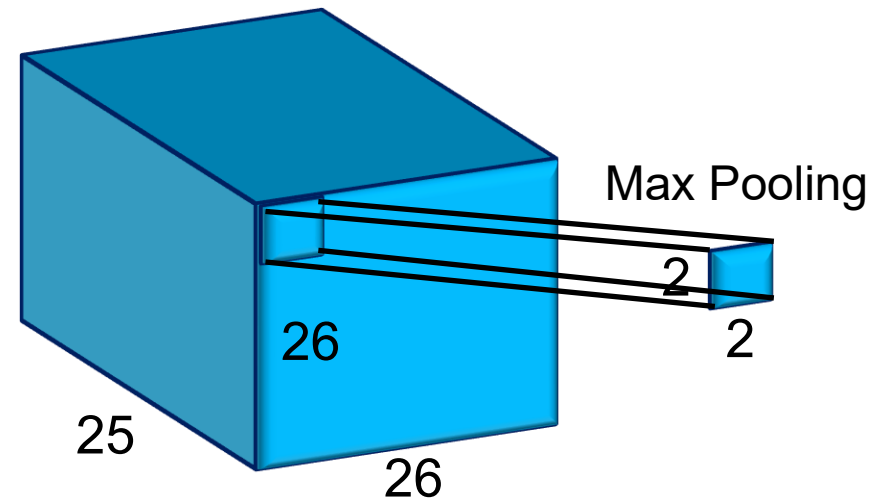
25 filters - Conv1

25:  $3 \times 3$

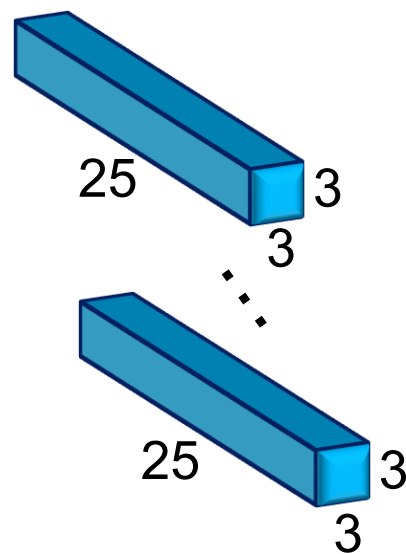
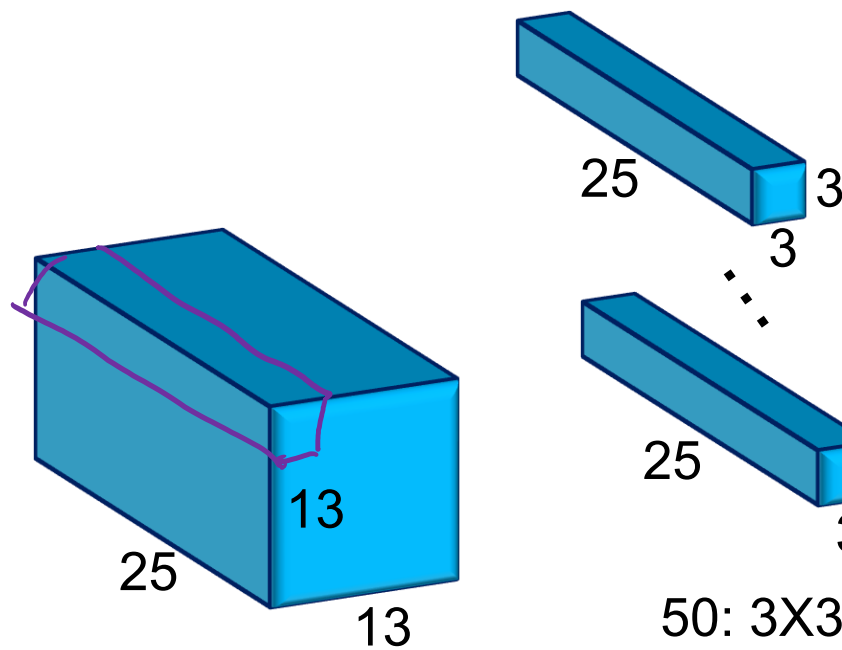
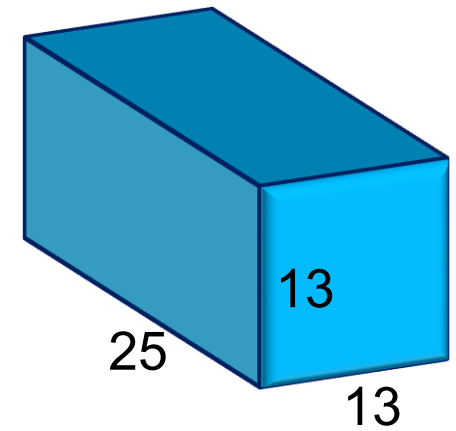
Gray scaled-image



Convolved result for 25 filters



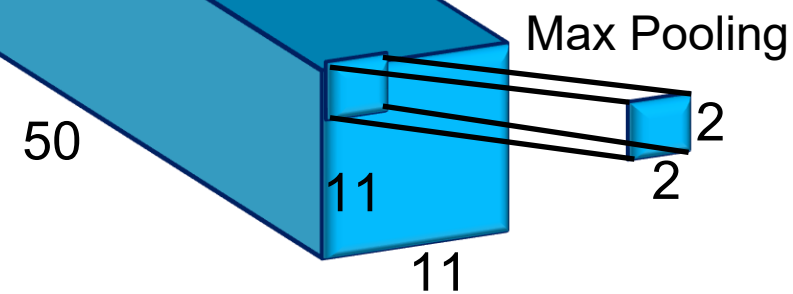
Result after Max Pooling



50:  $3 \times 3 \times 25$

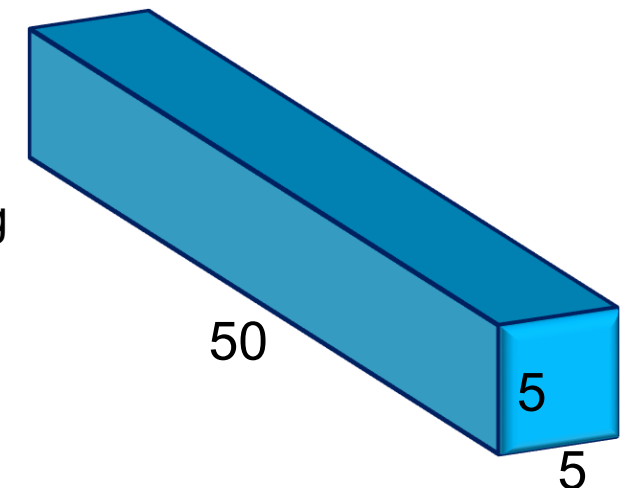
50 filters - Conv2

$50 \times 3 \times 3 \times 25 + 50$  parameters



Convolved result for 50 filters

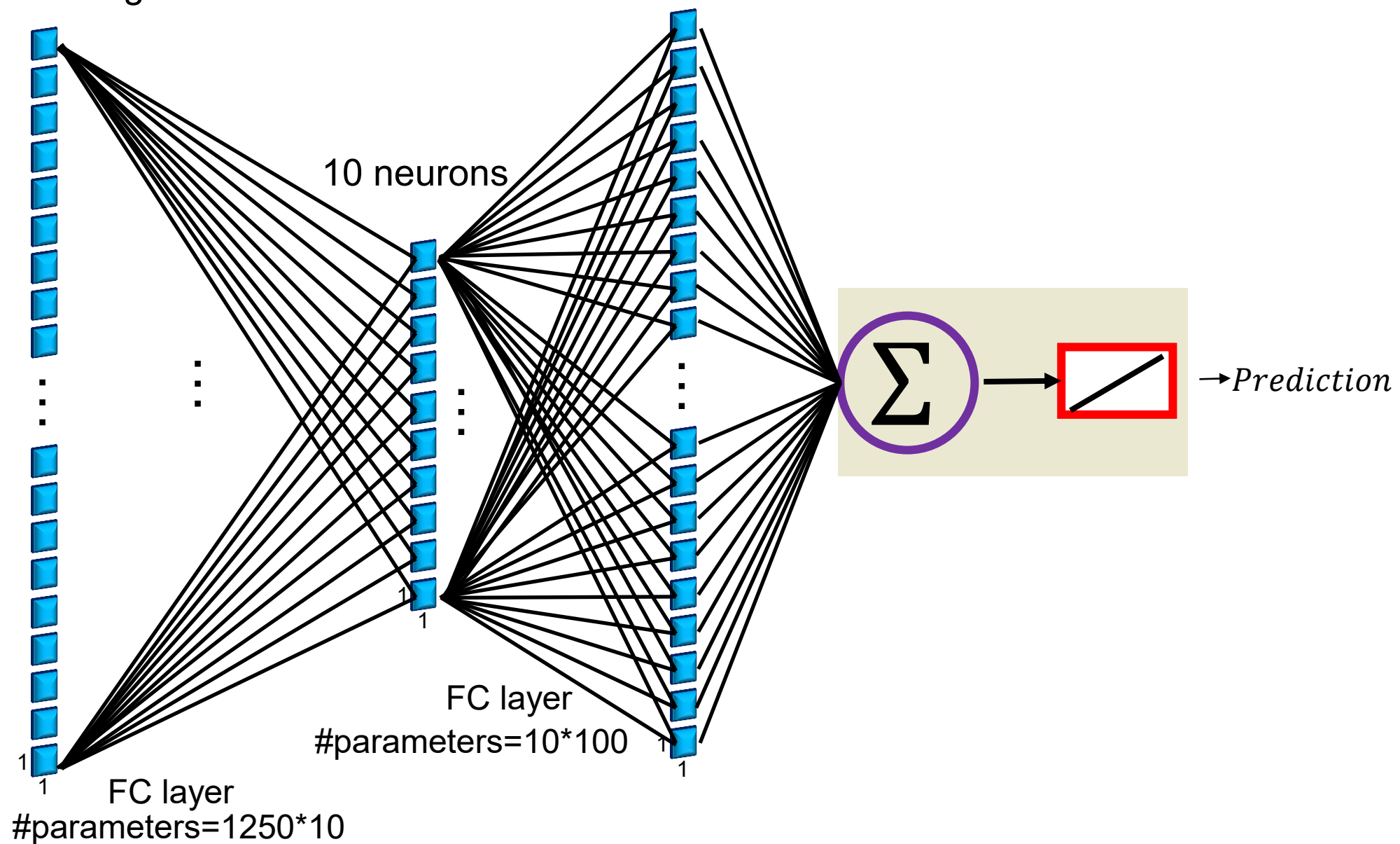
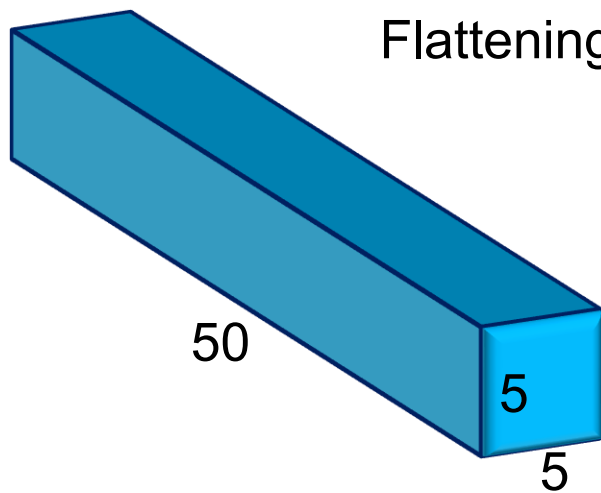
Result after Max Pooling



# neurons after flattening:  $50 \times 5 \times 5 = 1250$

100 neurons

Flattening



## 10 CNN Architecture