



Week 10

السنة الخامسة – هندسة المعلوماتية / الذكاء الصناعي

مقرر التعلم التلقائي

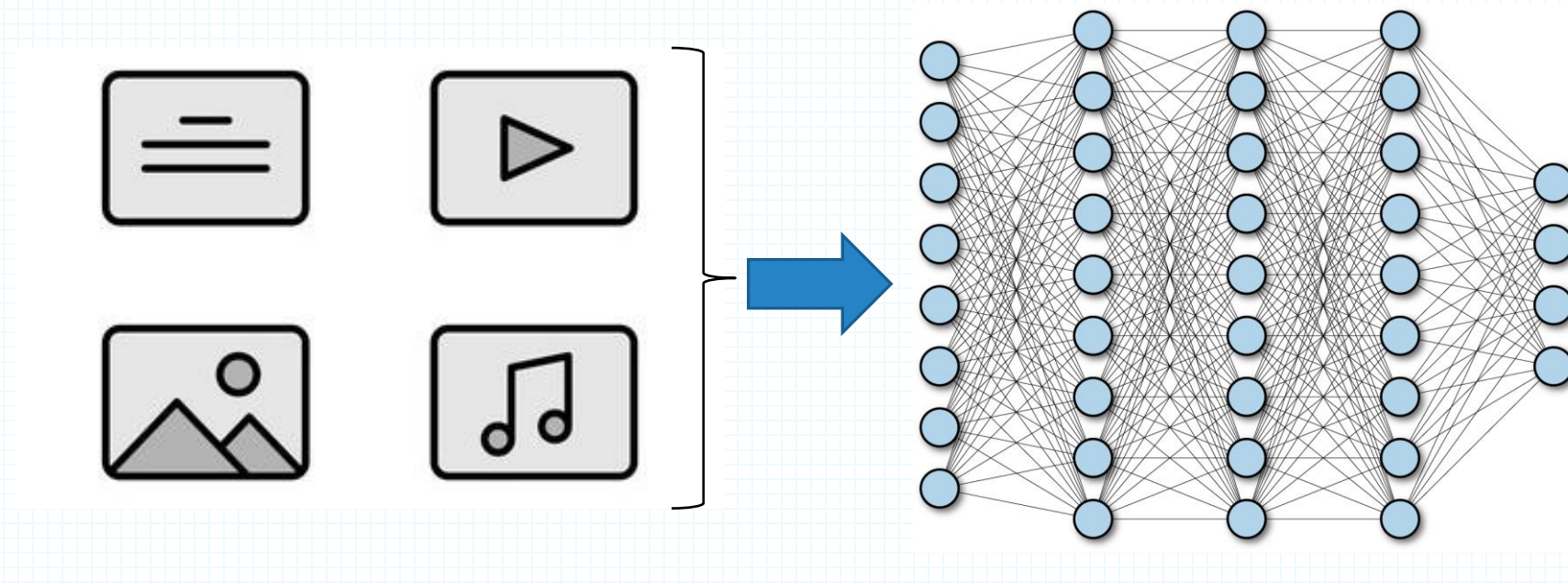
# Well Known DL Architecture (till 2016) CNN, RNN, LSTM

Dr. Riad Sonbol

[Access Course Materials](#) 

# Large networks

- What kind of neural networks can be used for large or variable length input vectors (e.g., time series)?
- Common networks:
  - CNN, RNN, etc



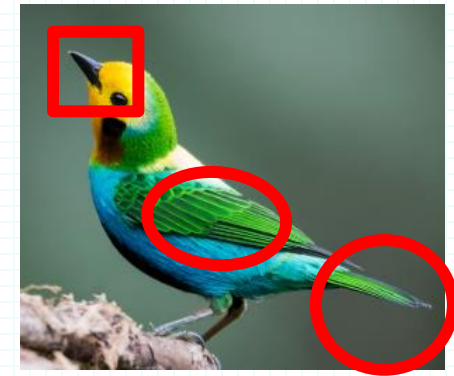
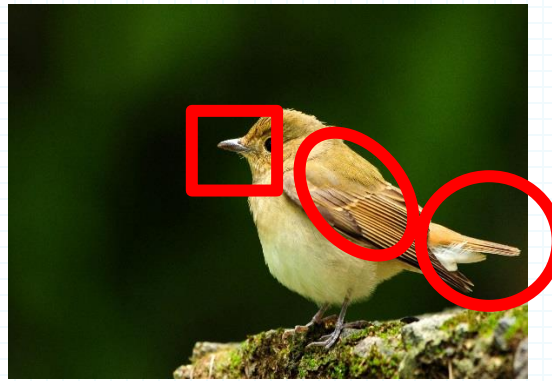
# الشبكة العصبونية التلافيفية

## Convolutional Neural Network (CNN)

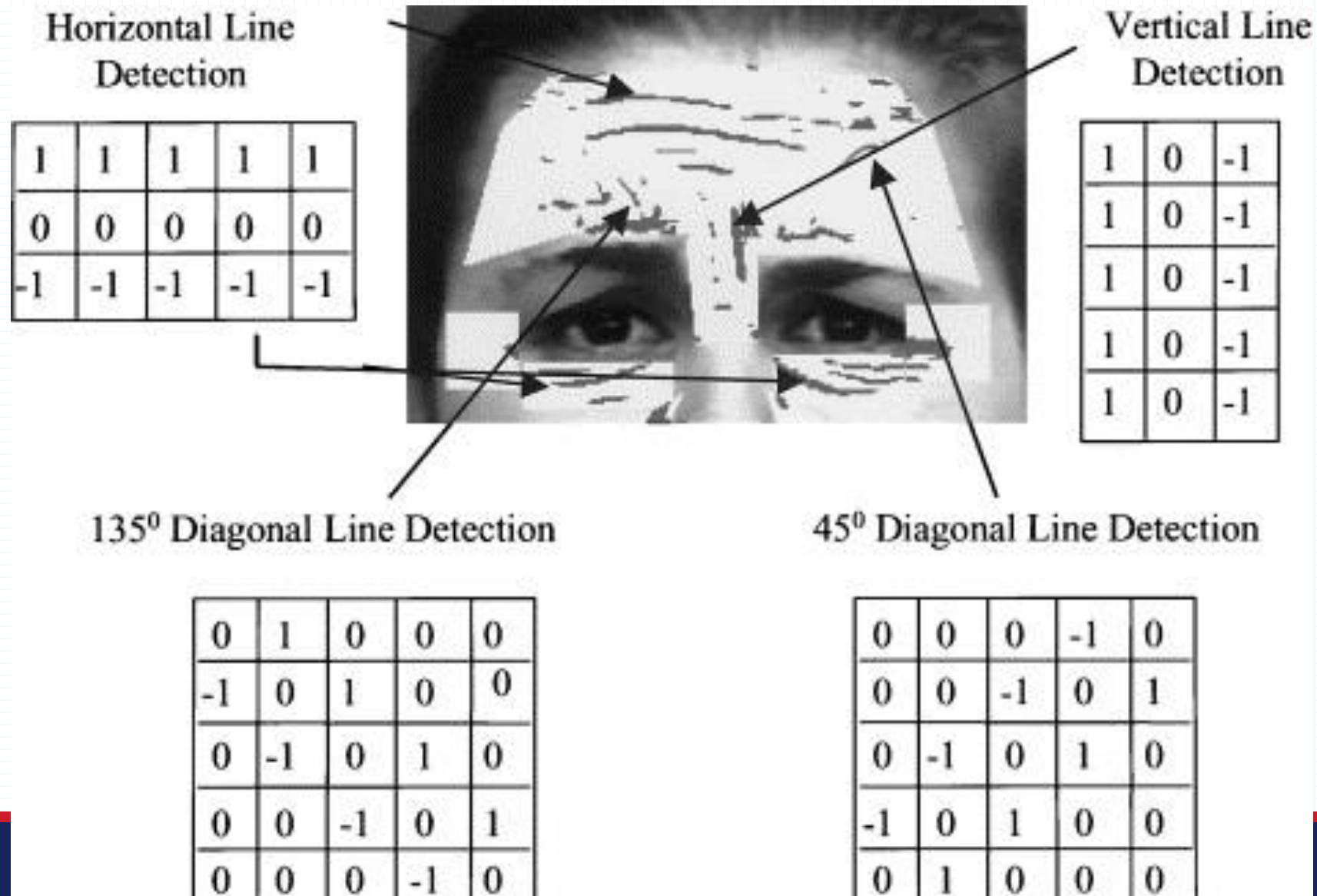
# Convolutions for feature extraction

---

The same patterns appear in different regions.



# Convolutions for feature extraction



# Convolutions for feature extraction

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

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

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

Filter

# Convolutions for feature extraction

---

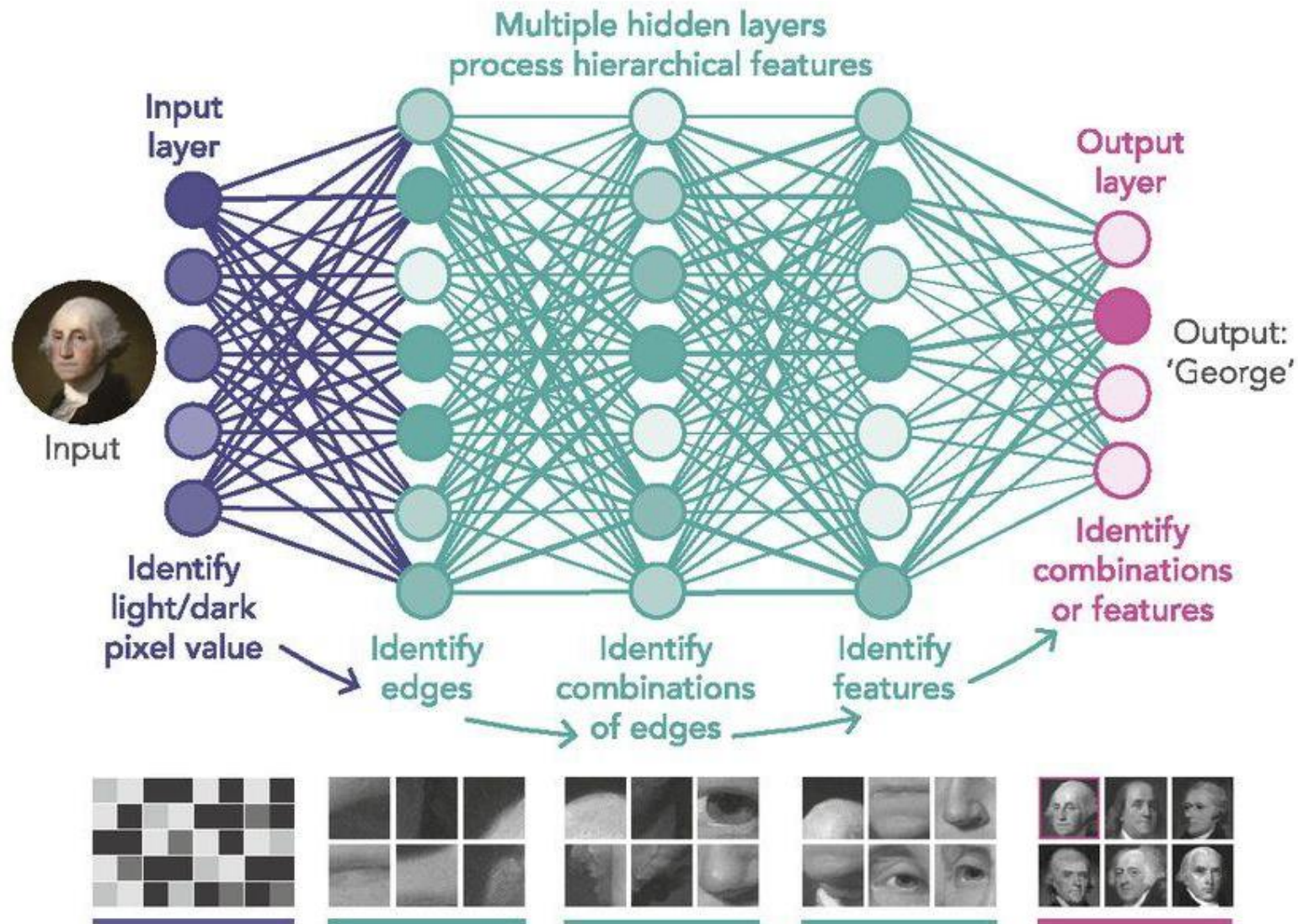
- In neural networks:
  - A **convolution** denotes the linear combination of a **subset of units** based on a **specific pattern of weights**.

$$a_j = \sum_i w_{ji} z_i$$

- Convolutions are often combined with an activation function to produce a feature

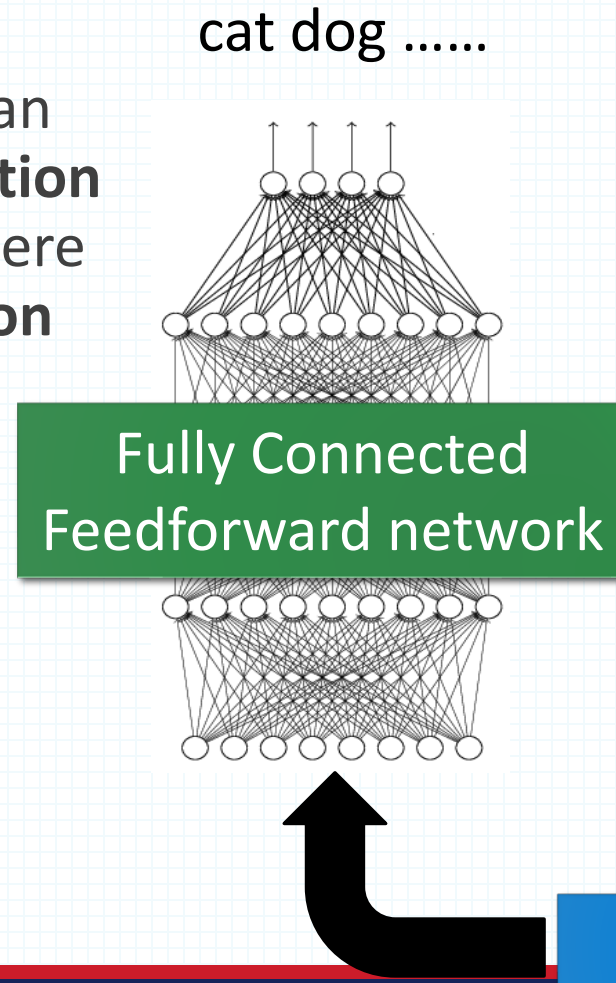
$$z_j = h(a_j) = h\left(\sum_i w_{ji} z_i\right)$$

# Deeper Neural Networks!



# Convolution Neural Network

- A **convolutional neural network** refers to any network that includes an **alternation of convolution and pooling layers**, where **some of the convolution weights are shared**.
- Architecture:



Convolution

Max Pooling

Convolution

Max Pooling

Can repeat  
many times

Flatten

# CNN – Convolution

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

Filter 1

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

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

# CNN – Convolution

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

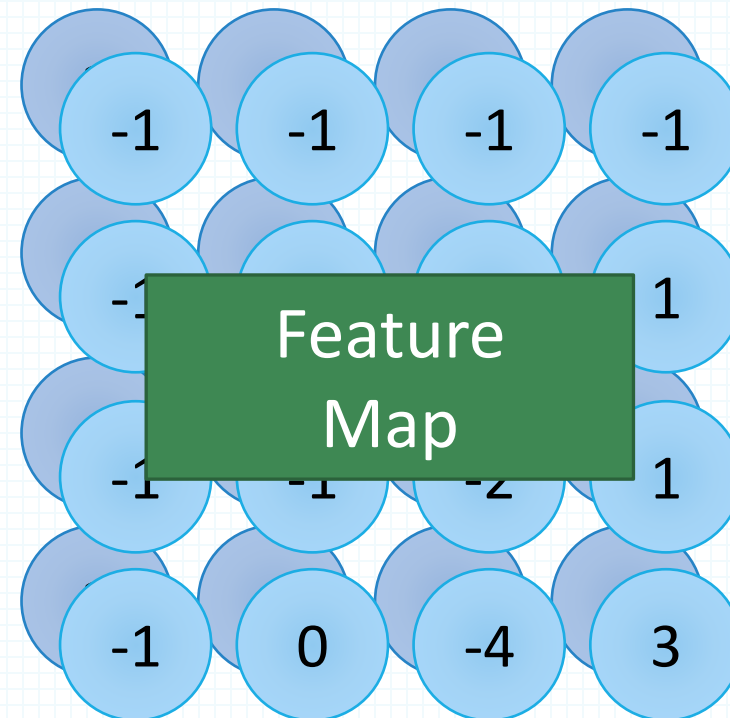
Filter 2

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

Do the same process for every filter



4 x 4 image

# CNN – Convolution

Those 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

Matrix

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

Filter 2

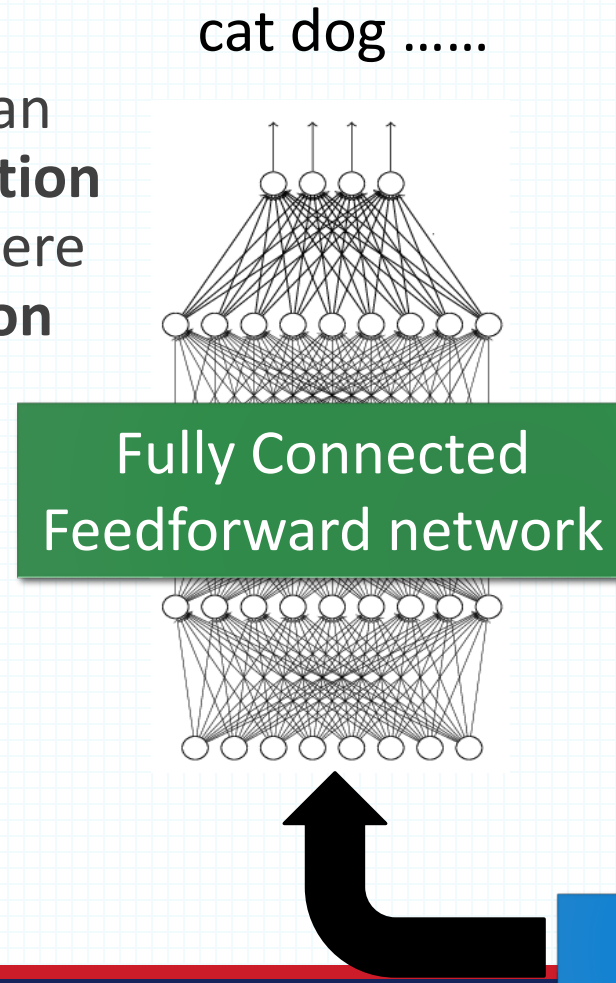
Matrix

⋮

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

# Convolution Neural Network

- A **convolutional neural network** refers to any network that includes an **alternation of convolution and pooling layers**, where **some of the convolution weights are shared**.
- Architecture:



Convolution

Max Pooling

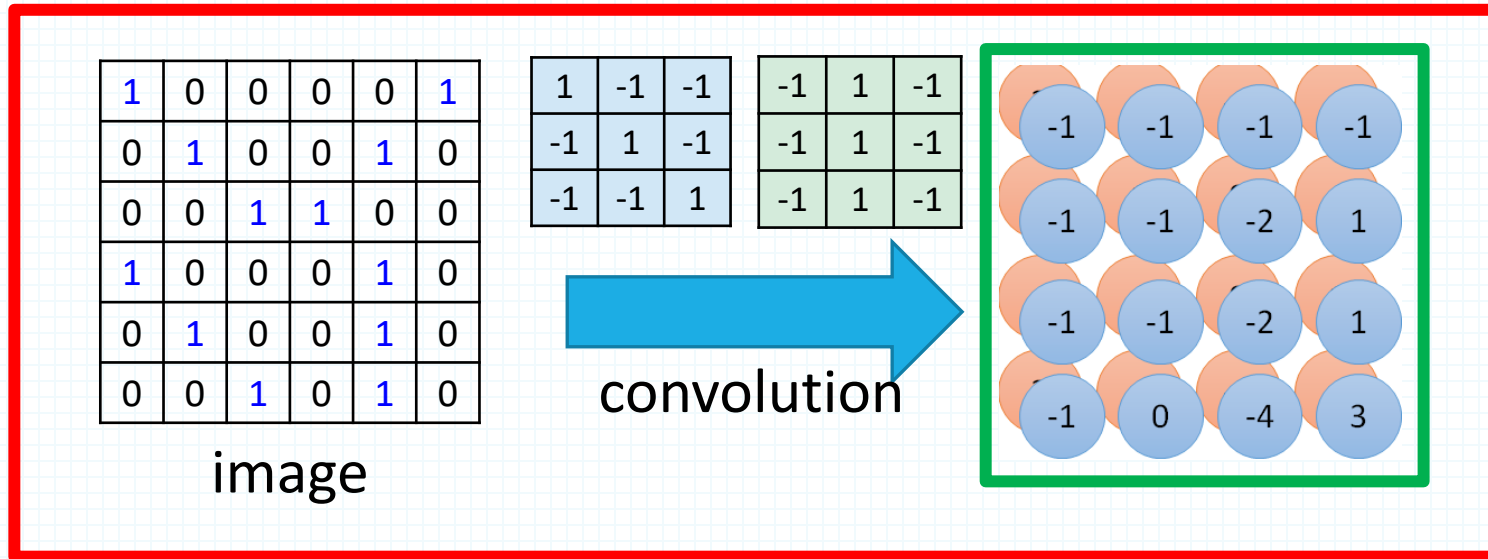
Convolution

Max Pooling

Flatten

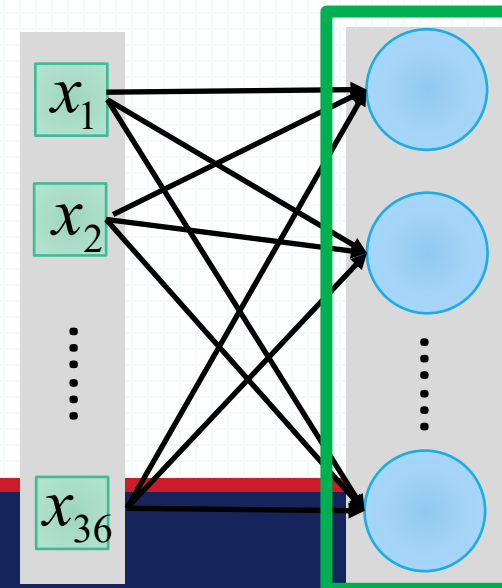
Can repeat  
many times

# Convolution v.s. Fully Connected



Fully-  
connected

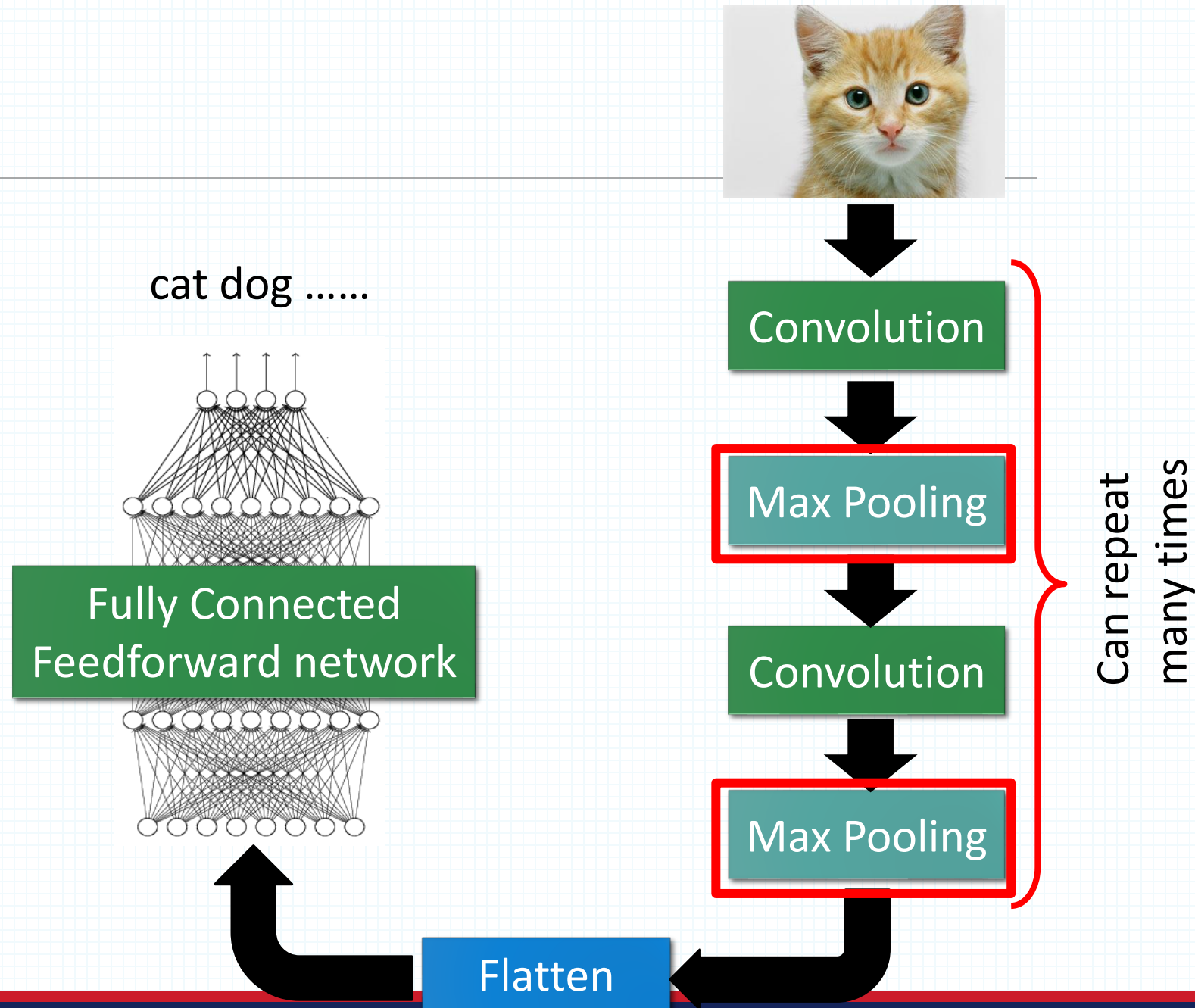
|   |   |   |   |   |   |
|---|---|---|---|---|---|
| 1 | 0 | 0 | 0 | 0 | 1 |
| 0 | 1 | 0 | 0 | 1 | 0 |
| 0 | 0 | 1 | 1 | 0 | 0 |
| 1 | 0 | 0 | 0 | 1 | 0 |
| 0 | 1 | 0 | 0 | 1 | 0 |
| 0 | 0 | 1 | 0 | 1 | 0 |



- Sparse interactions: Fewer connections
- Parameter sharing: Fewer weights
- Handle inputs of varying length

# Pooling

- Pooling: commutative mathematical operation that combines several units
- Examples:
  - max, sum, product, average, Euclidean norm, etc.
- Commutative property (order does not matter):
  - $\max a, b = \max(b, a)$



# CNN – Max Pooling

---

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

Filter 1

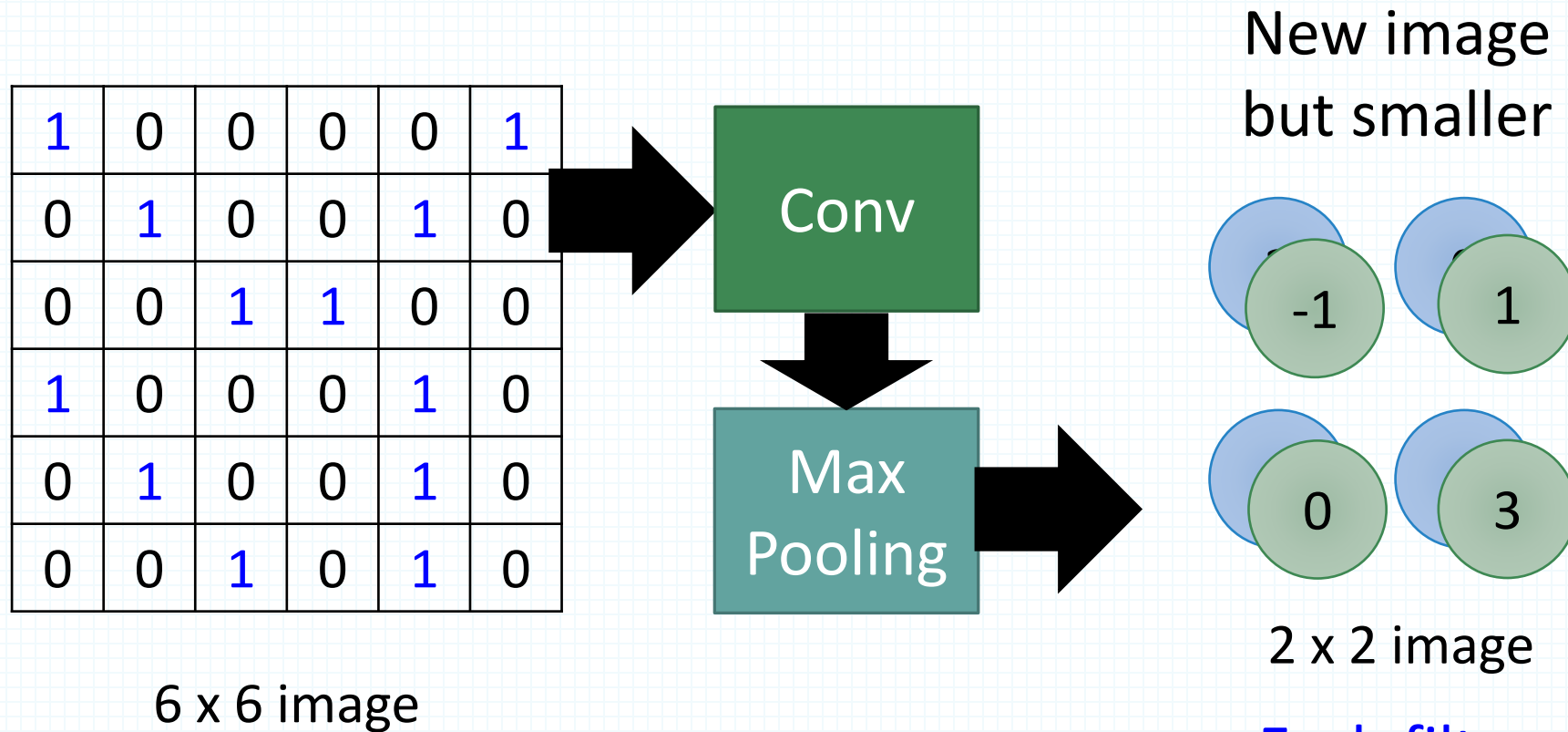
|    |   |    |
|----|---|----|
| -1 | 1 | -1 |
| -1 | 1 | -1 |
| -1 | 1 | -1 |

Filter 2

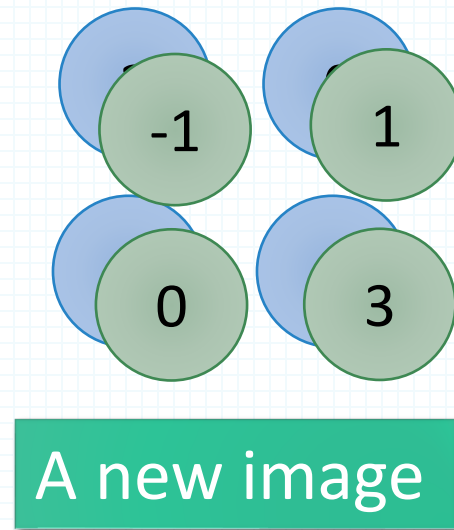
|    |    |    |    |
|----|----|----|----|
| 3  | -1 | -3 | -1 |
| -3 | 1  | 0  | -3 |
| -3 | -3 | 0  | 1  |
| 3  | -2 | -2 | -1 |

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

# CNN – Max Pooling

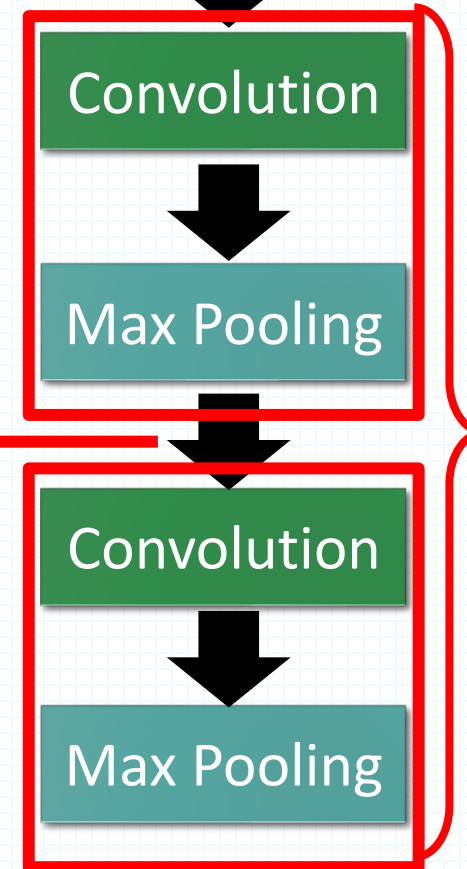


# The whole CNN



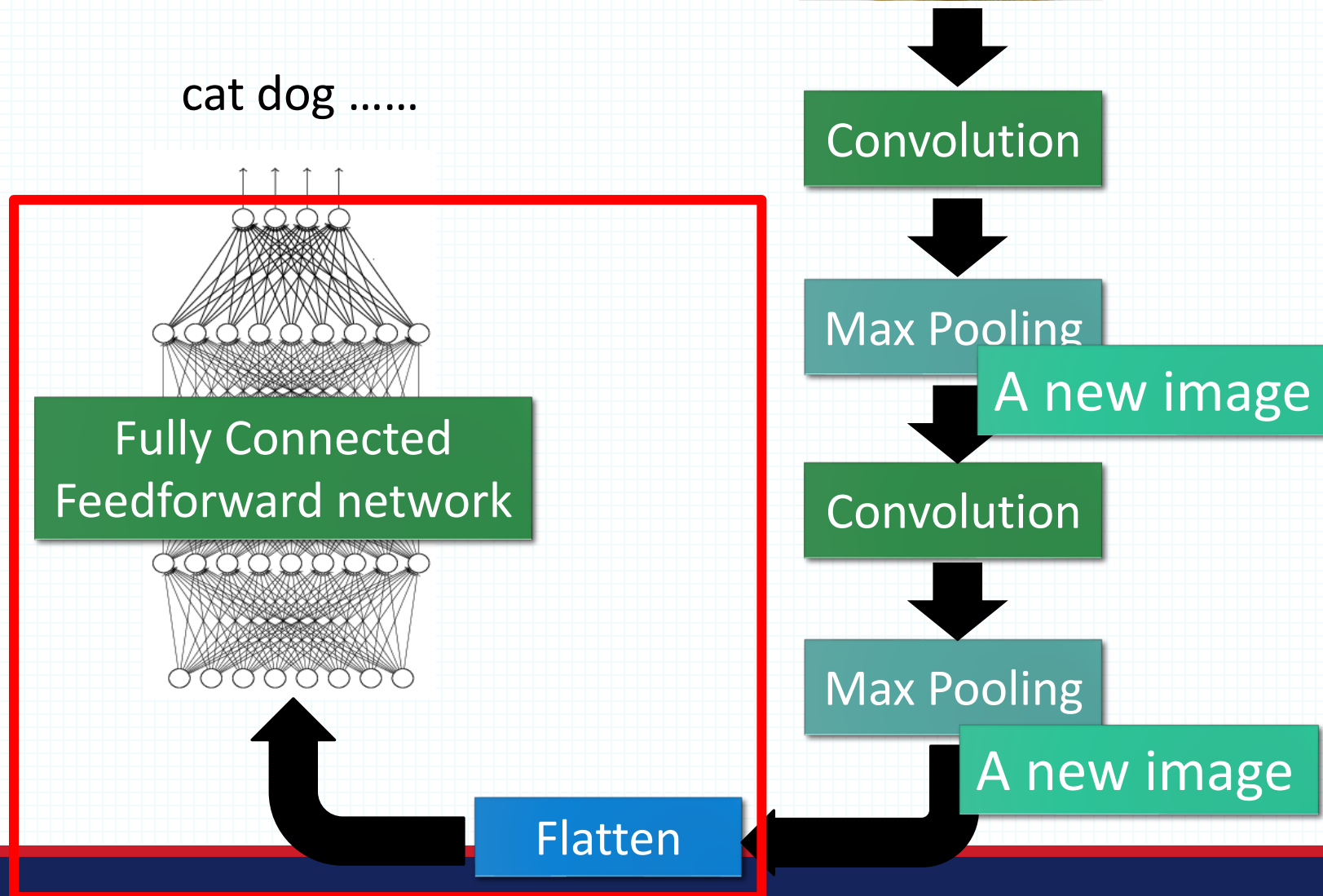
Smaller than the original image

The number of the channel is the number of filters

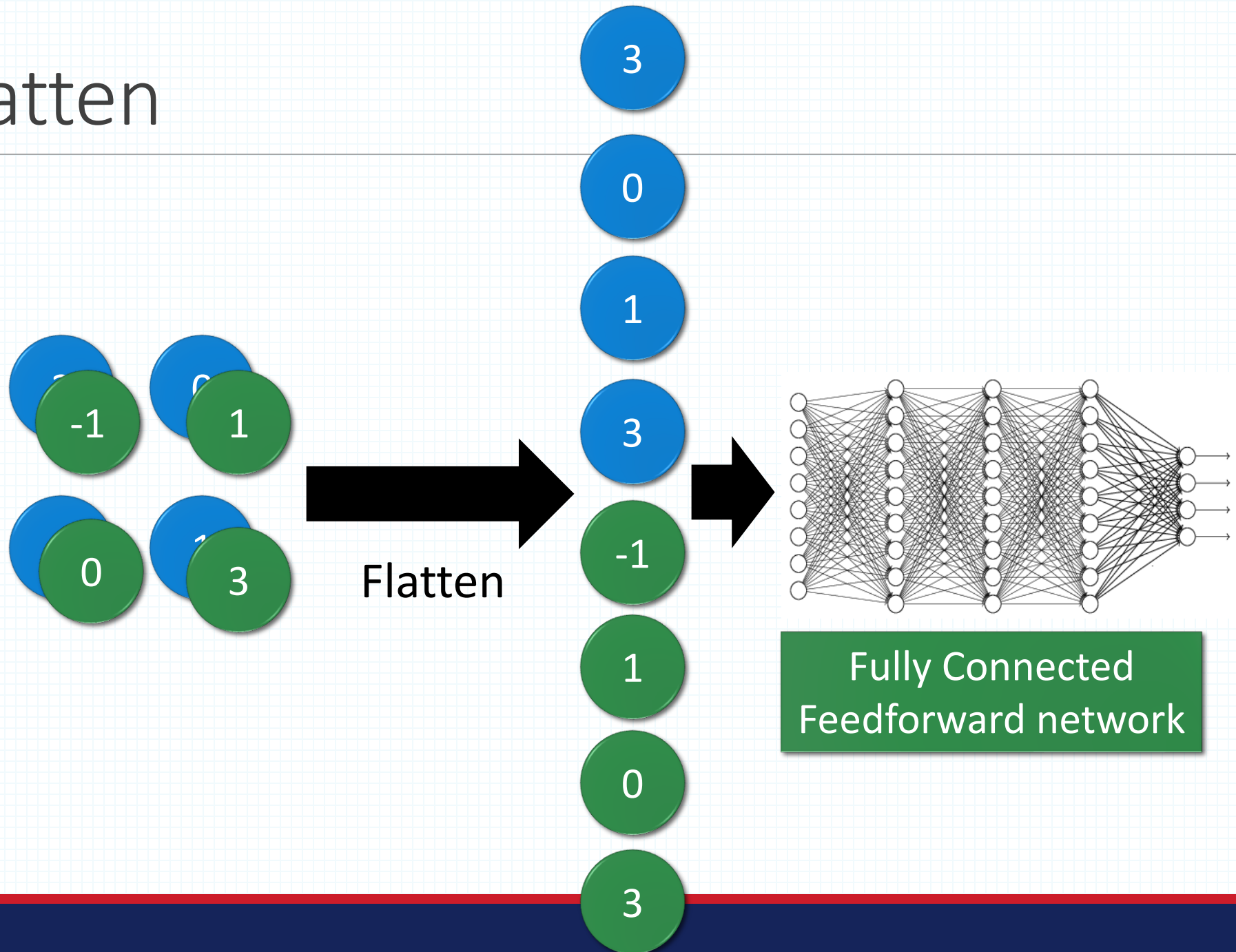


Can repeat many times

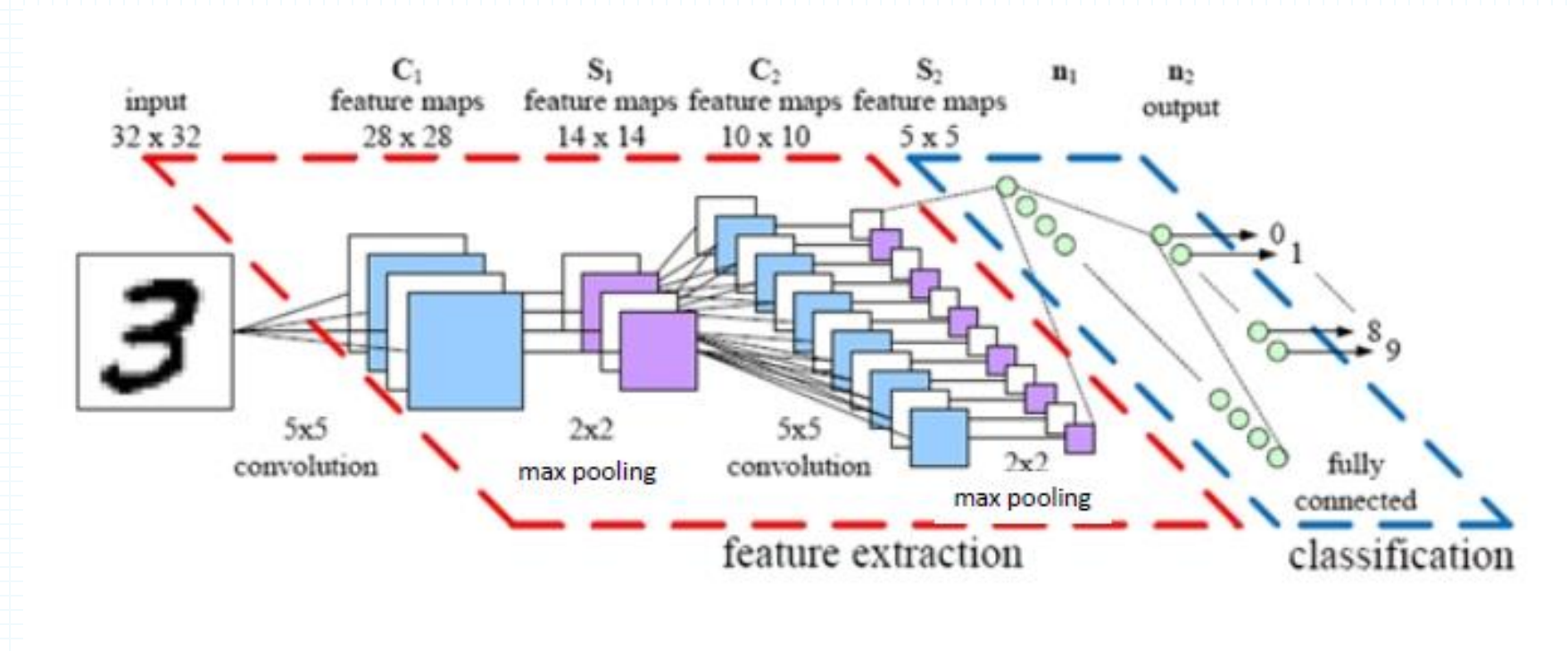
# The whole CNN



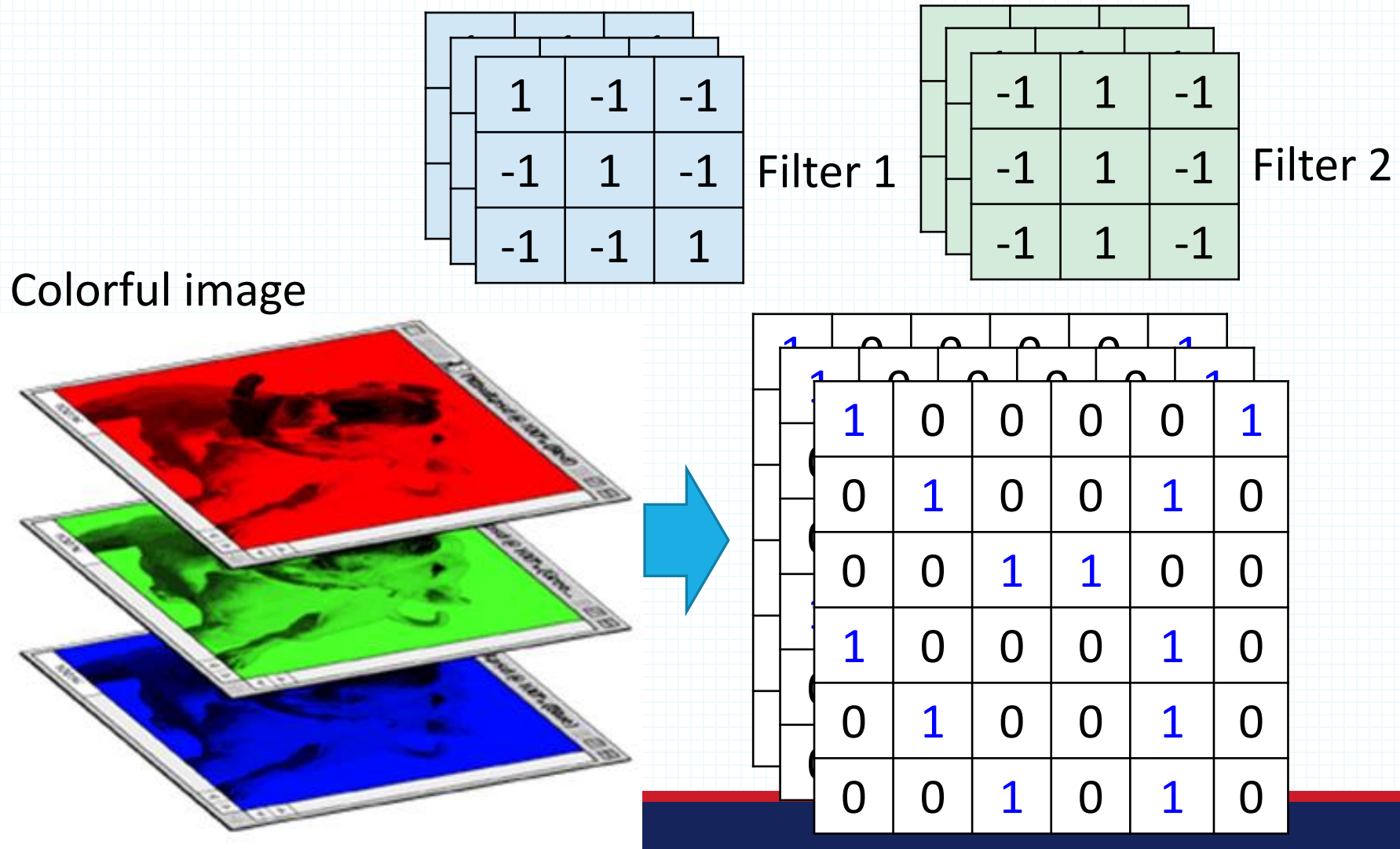
# Flatten



# Example



# CNN – Colorful image



# CNN – Zero Padding

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

Filter 1

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

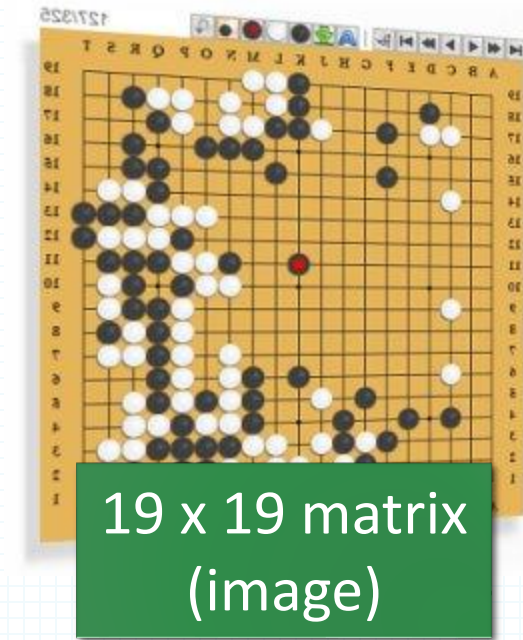
6 x 6 image

You will get another 6 x 6 images in this way



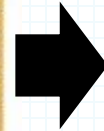
Zero padding

# More Application: Playing Go

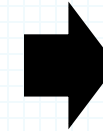


19 x 19 matrix  
(image)

Black: 1  
white: -1  
none: 0



Network



Next move  
(19 x 19  
positions)

19 x 19 vector

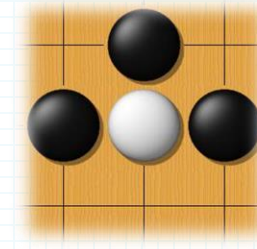
Fully-connected feedforward  
network can be used

But CNN performs much better.

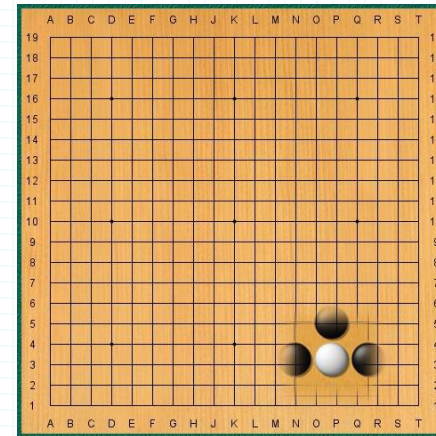
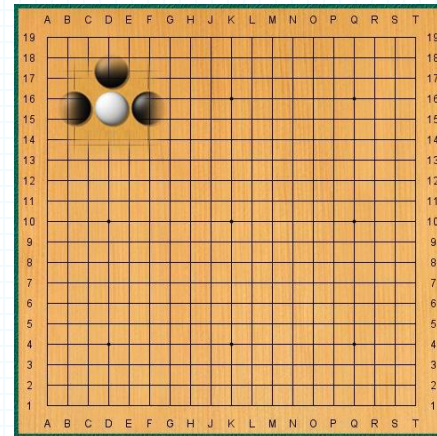
# Why CNN for playing Go?

---

- Some patterns are much smaller than the whole image



- The same patterns appear in different regions.  
Alpha Go uses 5 x 5 for first layer



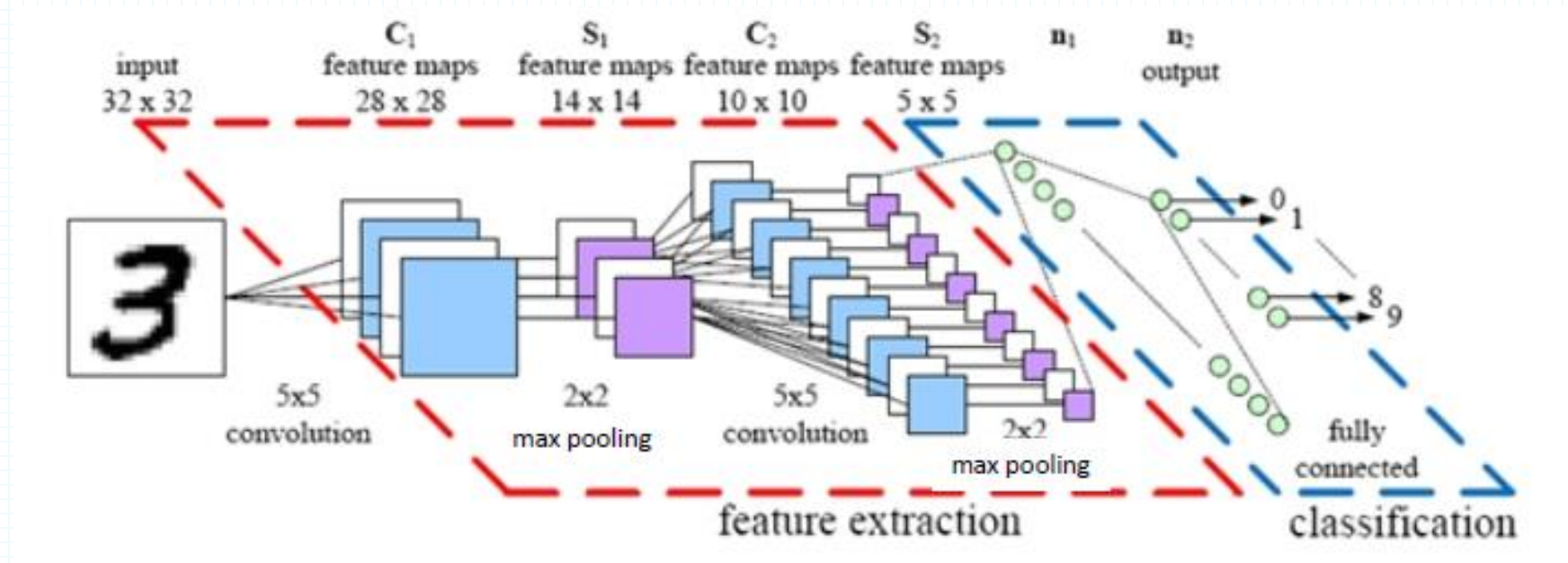
# Parameters

---

- **# of filters:** integer indicating the # of filters applied to each window.
- **kernel size:** tuple (width, height) indicating the size of the window.
- **Stride:** tuple (horizontal, vertical) indicating the horizontal and vertical shift between each window.
- **Padding:** “valid” or “same”. Valid indicates no input padding. Same indicates that the input is padded with a border of zeros to ensure that the output has the same size as the input.

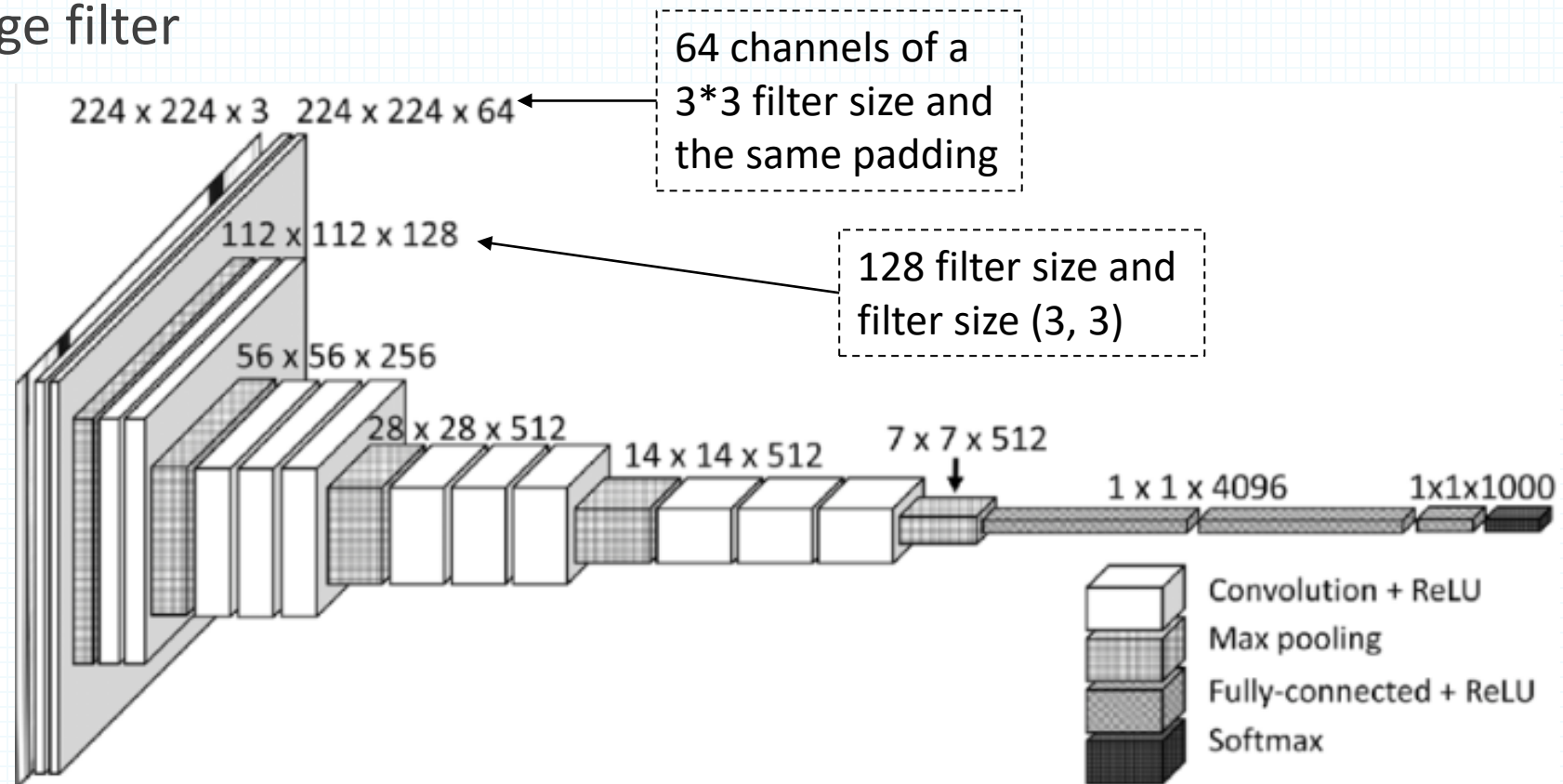
# Training

- Convolutional neural networks are trained in the same way as other neural networks
  - backpropagation
- Weight sharing:
  - Combine gradients of shared weights into a single gradient



# Architecture design: VGG

- What is the preferred filter size?
- VGG (Visual Geometry Group at Oxford, 2014): stack of small filters is often preferred to single large filter
  - Fewer parameters
  - Deeper network



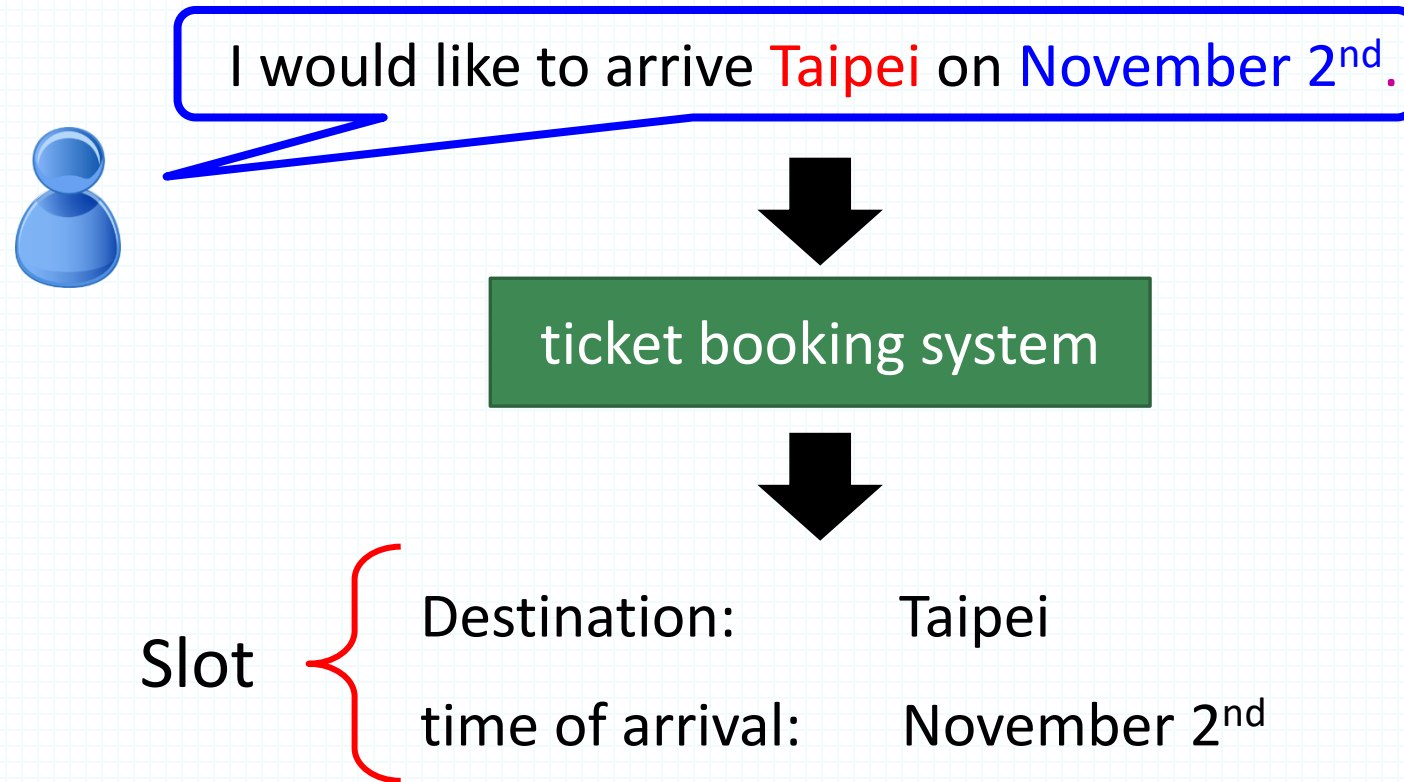
# الشبكات العصبونية التكرارية

## Recurrent Neural Networks

# Example Application

---

- Slot Filling



# Example Application

---

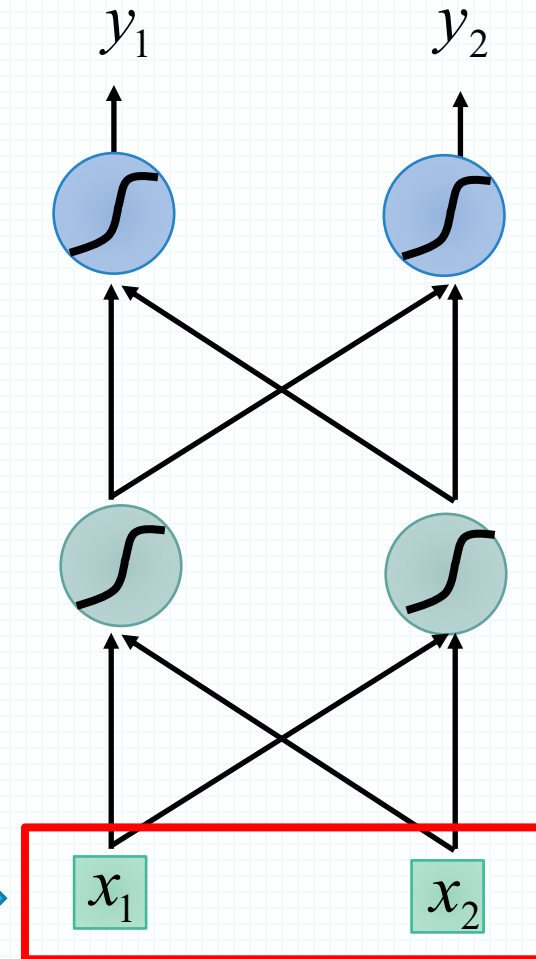
Solving slot filling by  
Feedforward network?

Input: a word

(Each word is represented  
as a vector)

How can we represent each word?

Taipei



# 1-of-N encoding

---

How to represent each word as a vector?

**1-of-N Encoding**    lexicon = {apple, bag, cat, dog, elephant}

The vector is lexicon size.

apple = [ 1   0   0   0   0 ]

Each dimension corresponds  
to a word in the lexicon

bag    = [ 0   1   0   0   0 ]

cat    = [ 0   0   1   0   0 ]

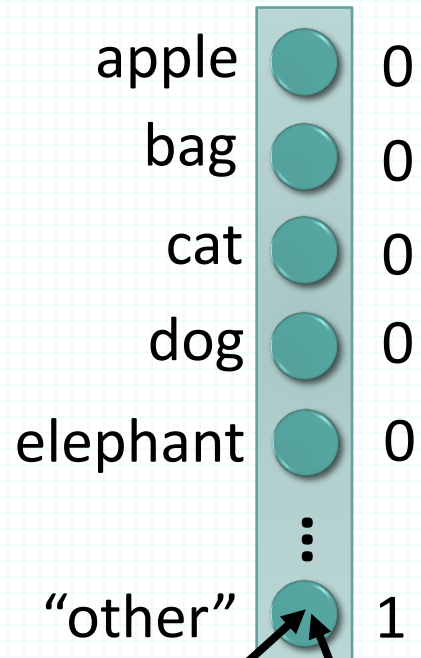
The dimension for the word  
is 1, and others are 0

dog    = [ 0   0   0   1   0 ]

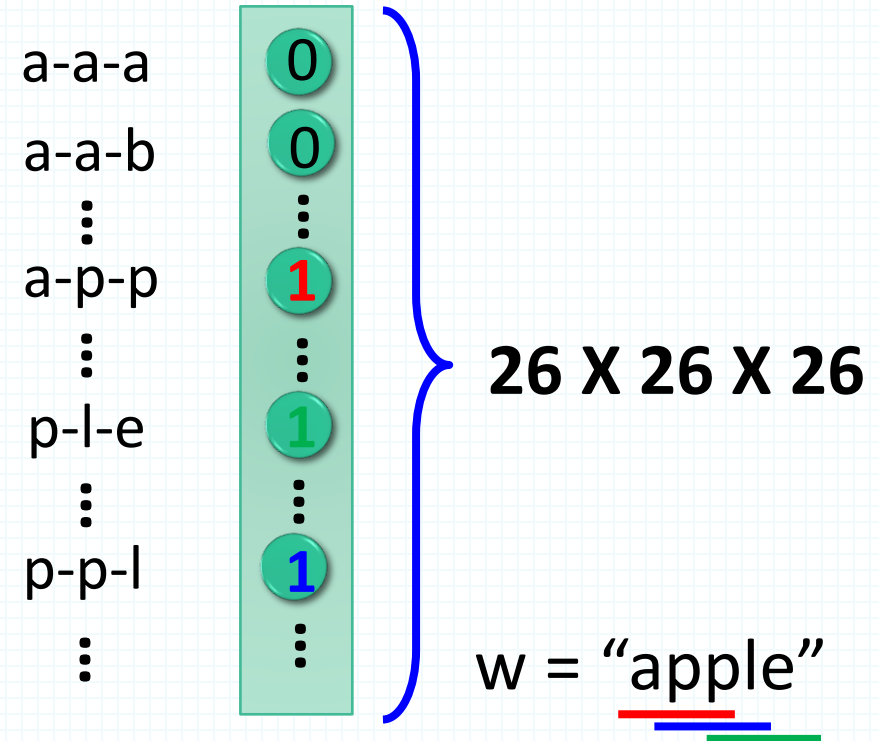
elephant = [ 0   0   0   0   1 ]

# Beyond 1-of-N encoding

## Dimension for "Other"



## Word hashing



# Example Application

Solving slot filling by  
Feedforward network?

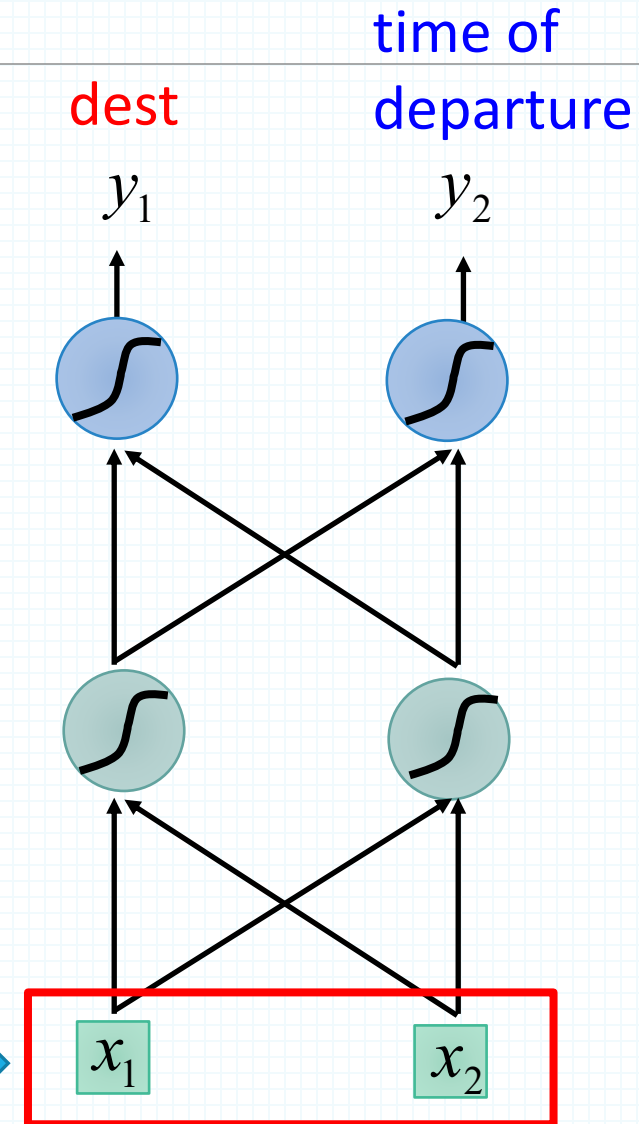
Input: a word

(Each word is represented  
as a vector)

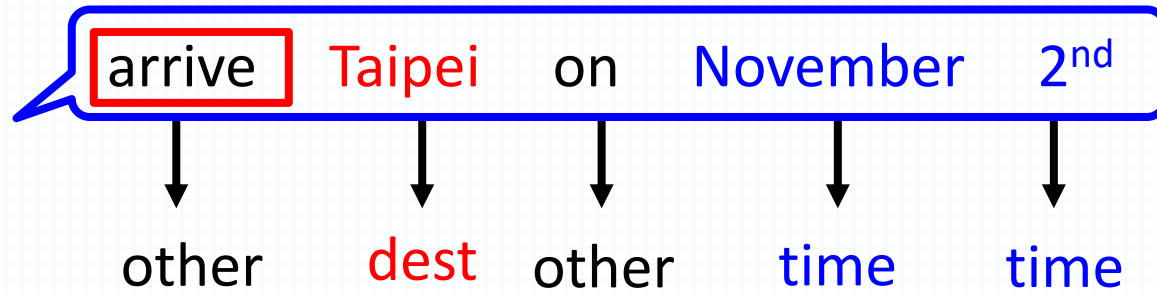
Output:

Probability distribution that  
the input word belonging to  
the slots

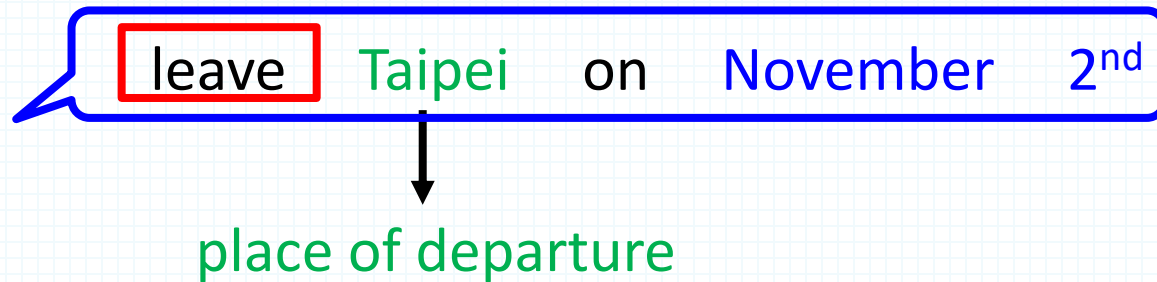
Taipei →



# Example Application

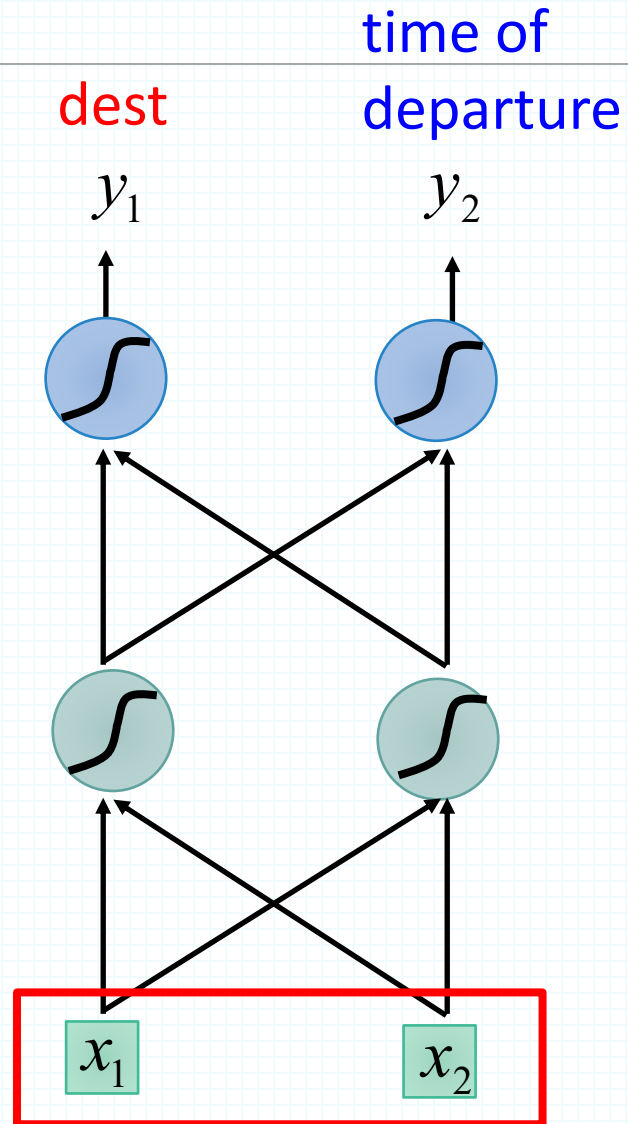


Problem?



Neural network  
needs memory!

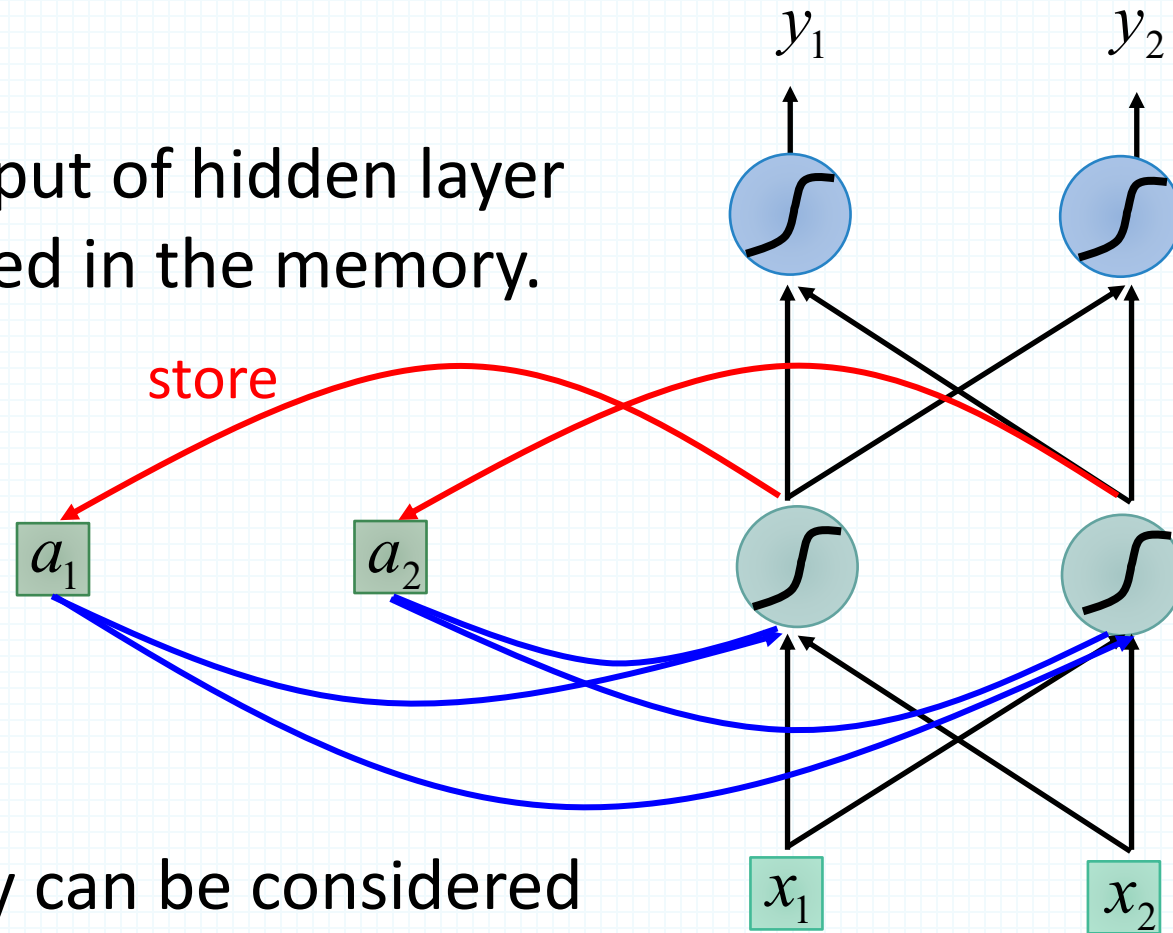
Taipei



# Recurrent Neural Network (RNN)

---

The output of hidden layer are stored in the memory.



Memory can be considered as another input.

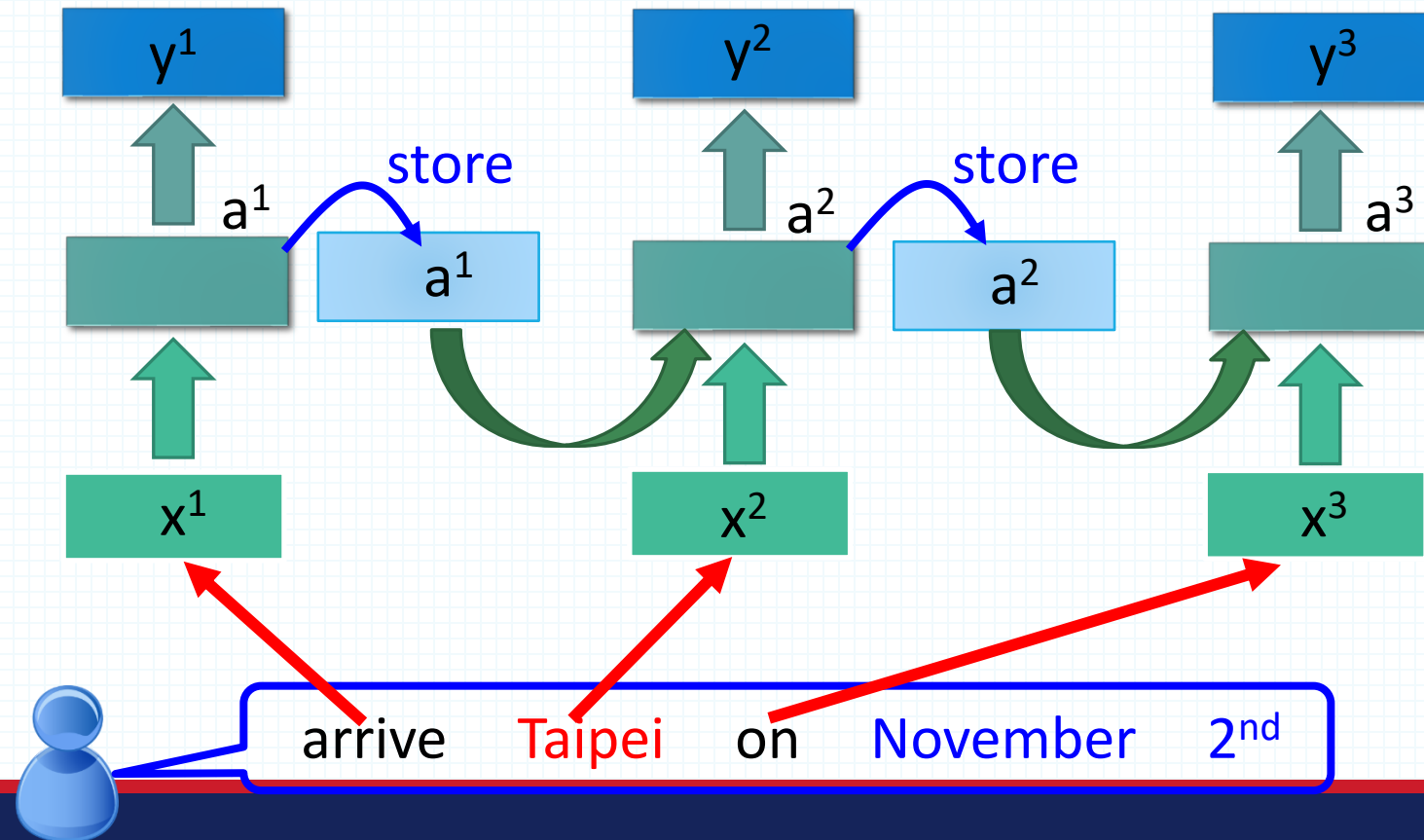
# RNN

The same network is used again and again.

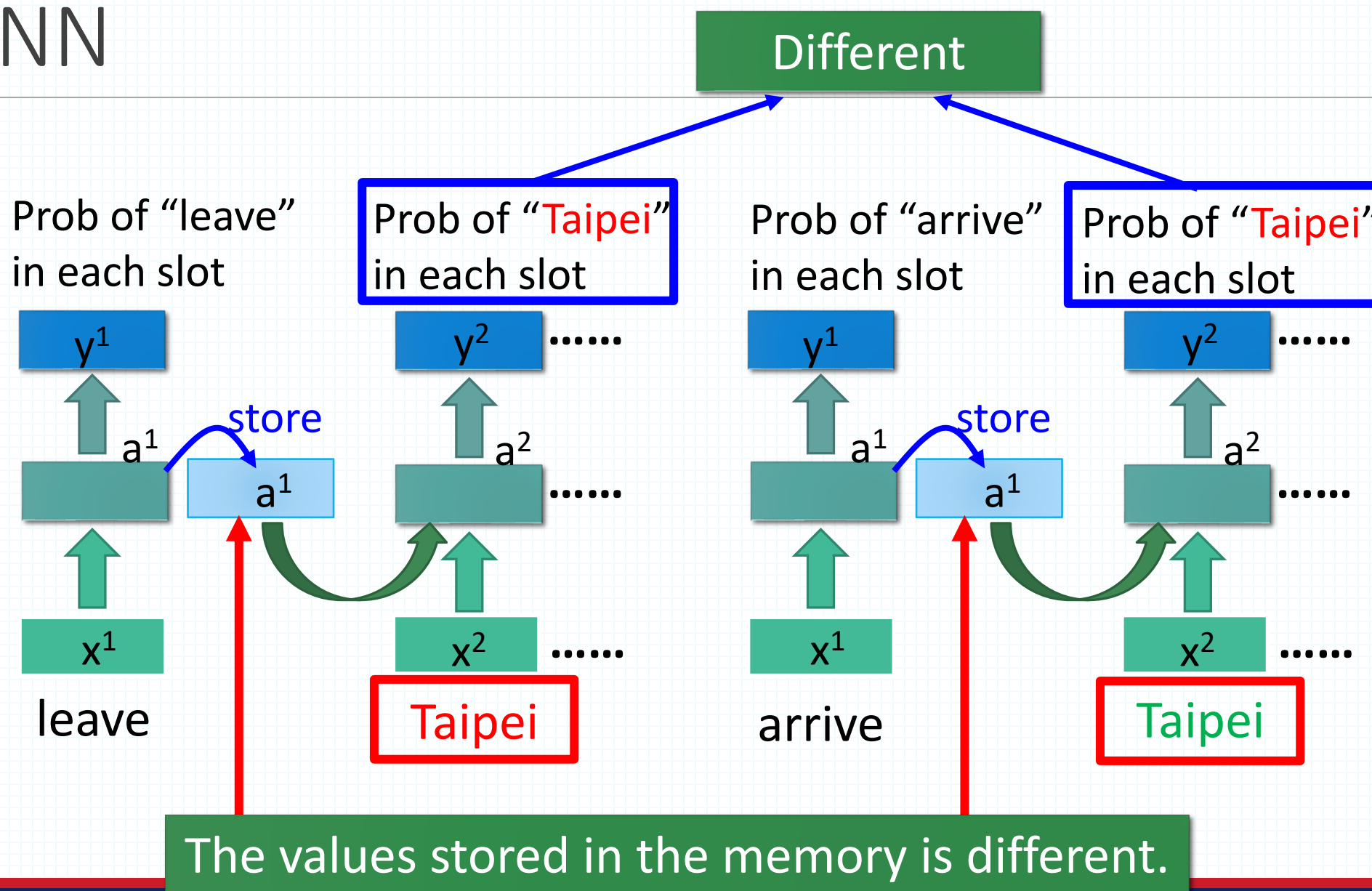
Probability of  
“arrive” in each slot

Probability of  
“**Taipei**” in each slot

Probability of  
“on” in each slot

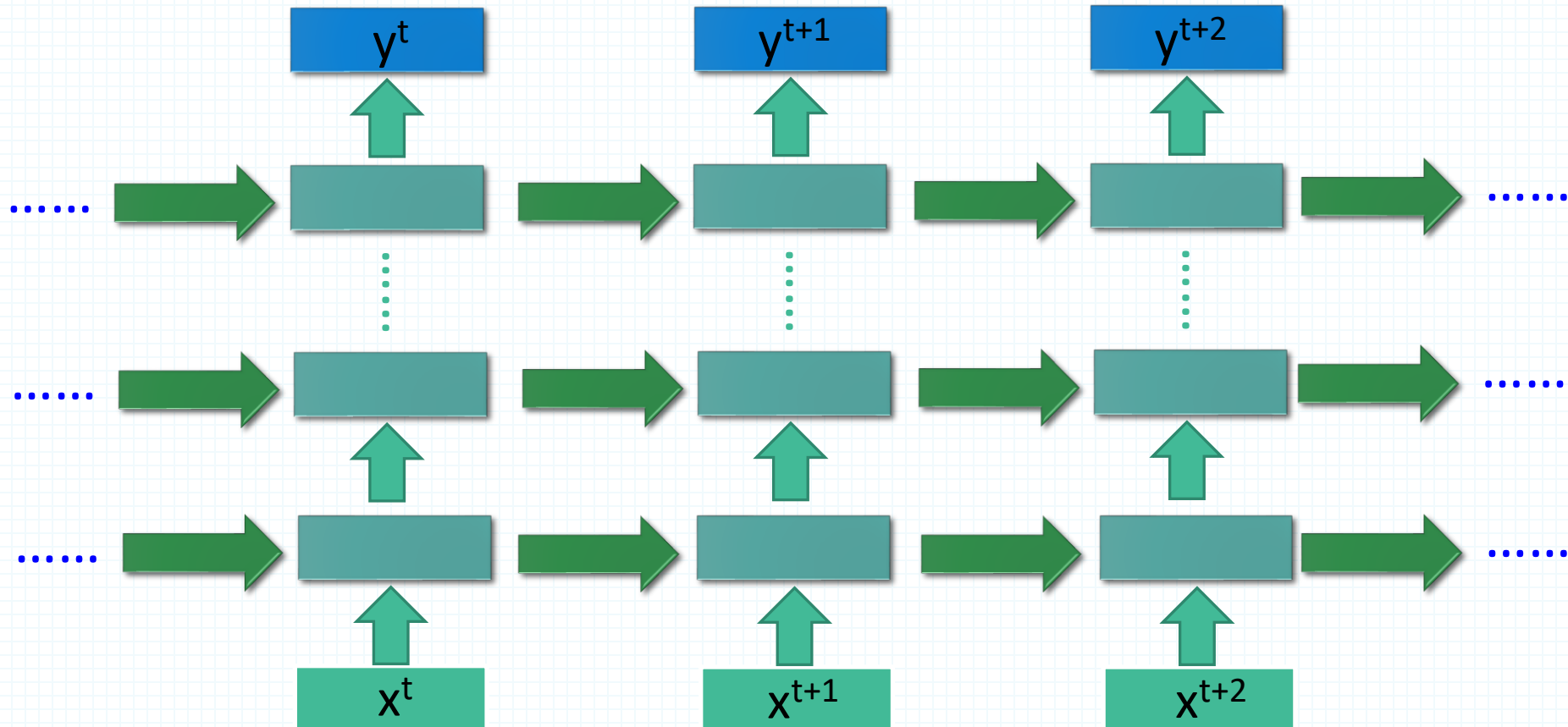


# RNN



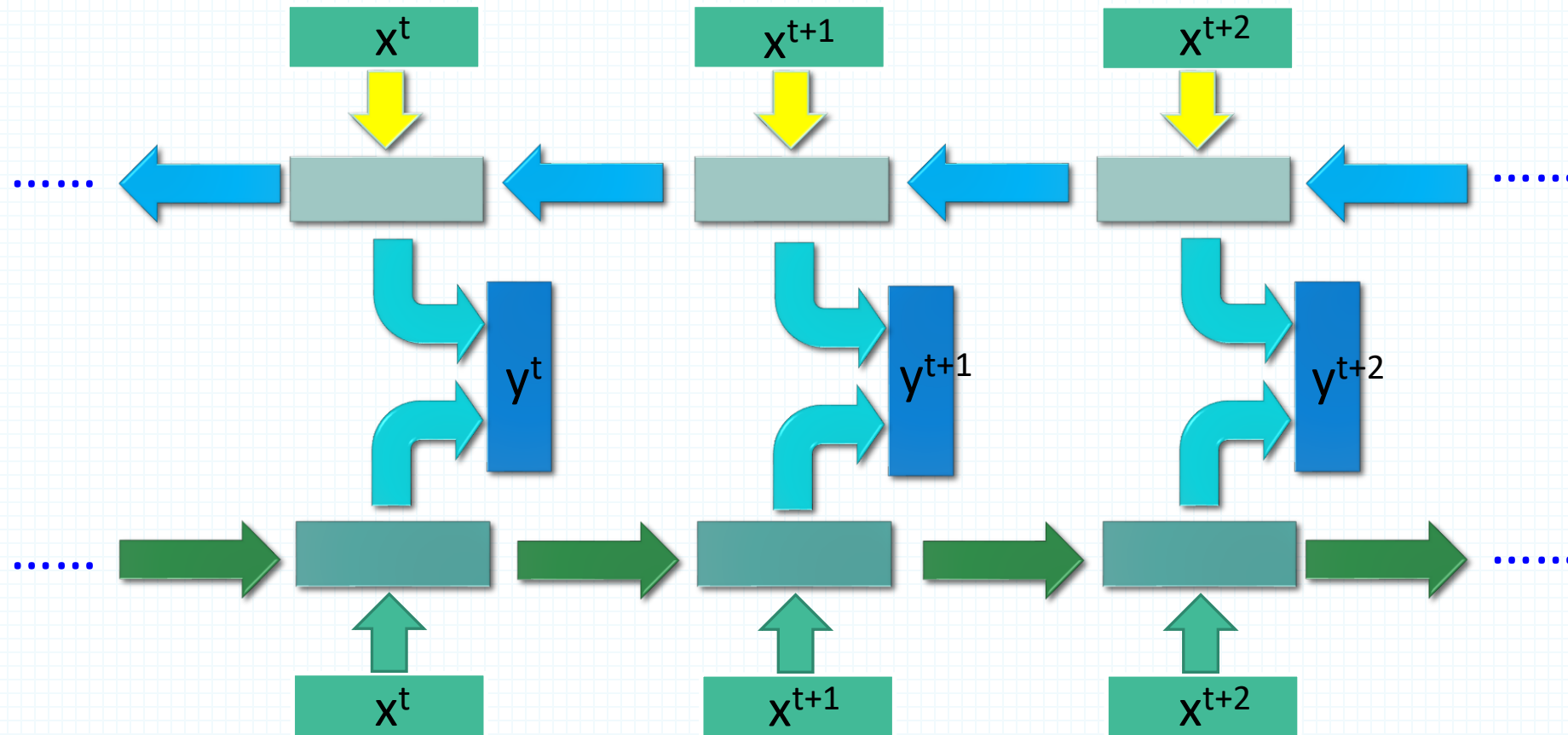
# Of course it can be deep ...

---



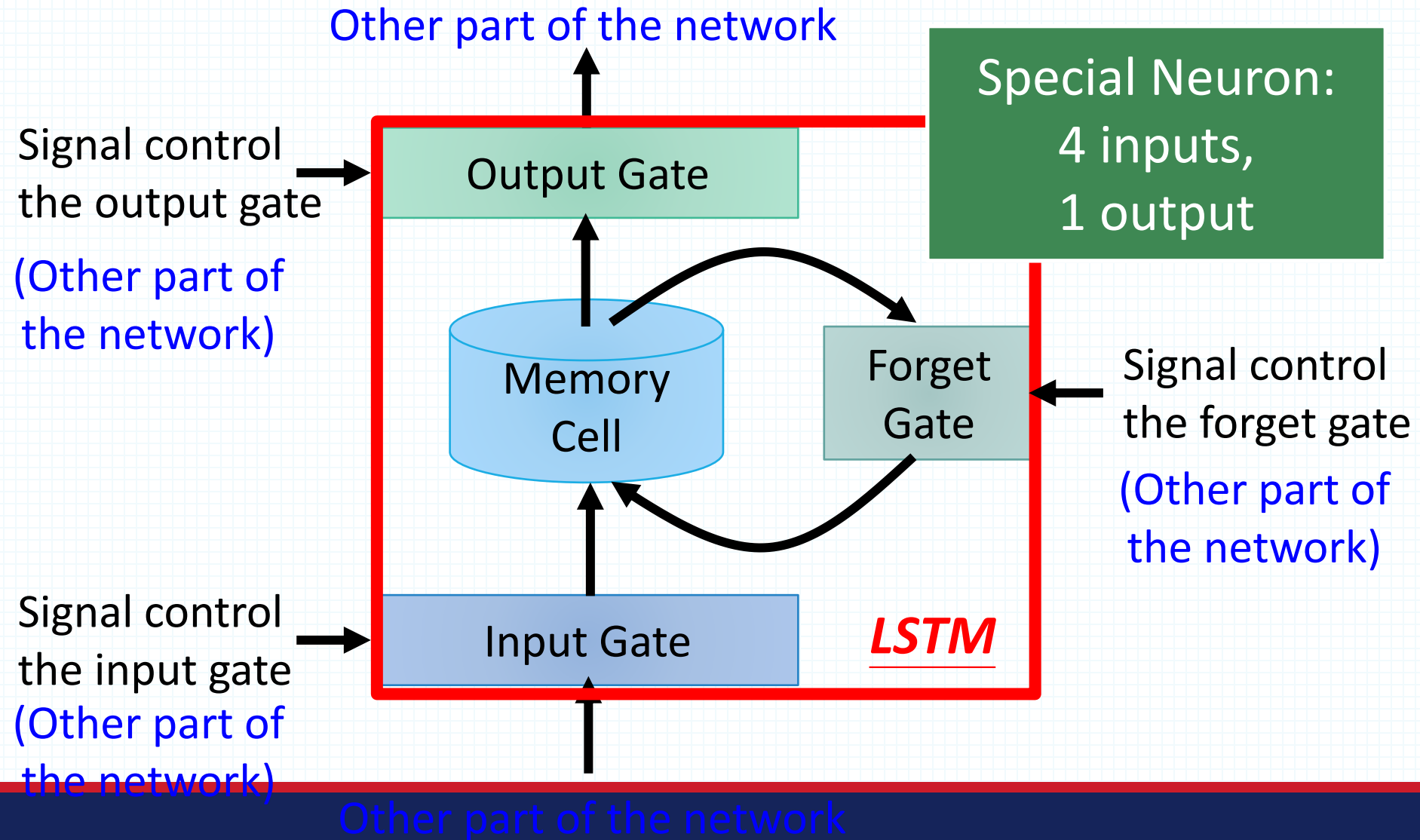
# Bidirectional RNN

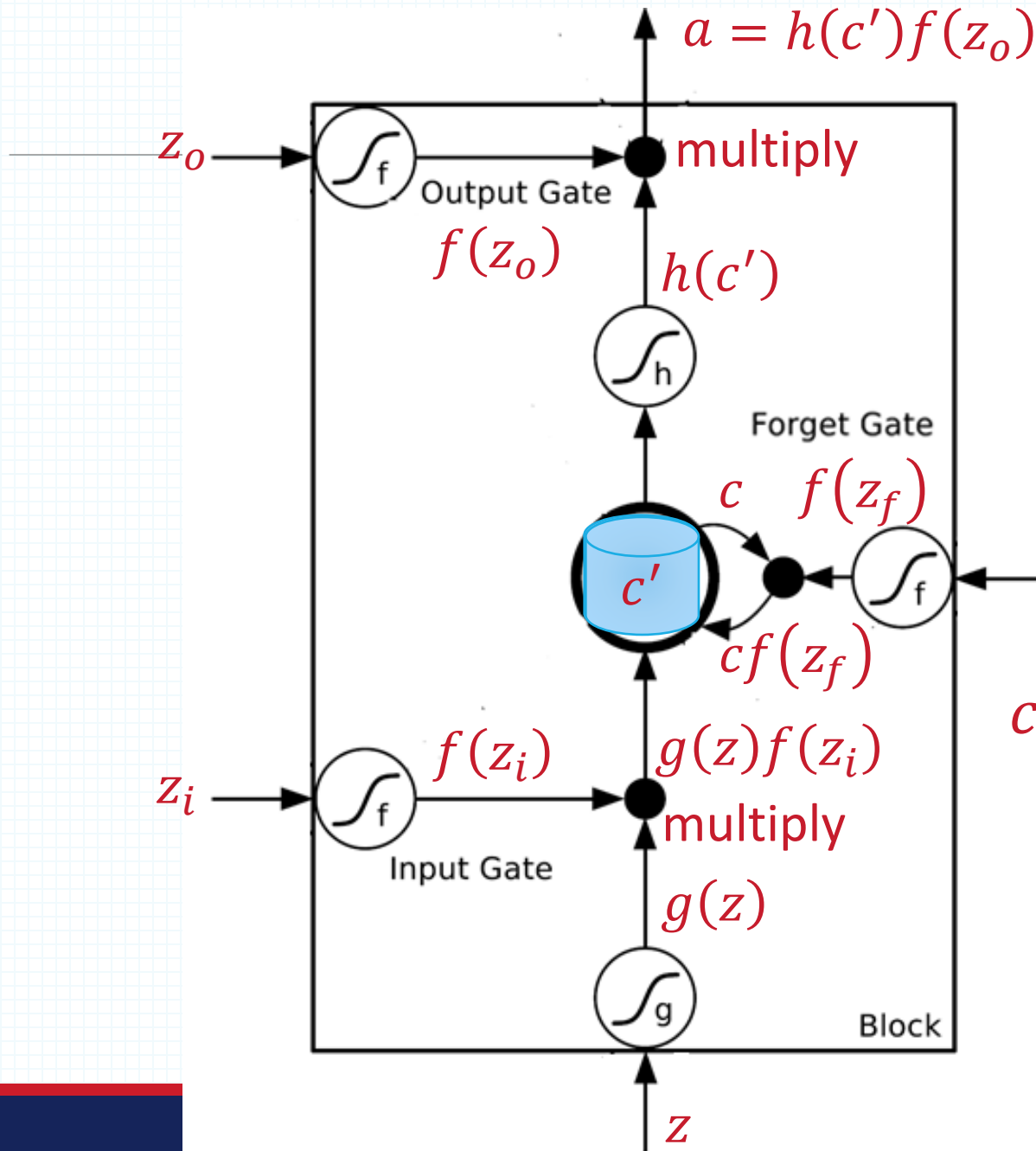
We can combine past and future evidence in separate chains



# Long Short-term Memory (LSTM)

# Long Short-term Memory (LSTM)





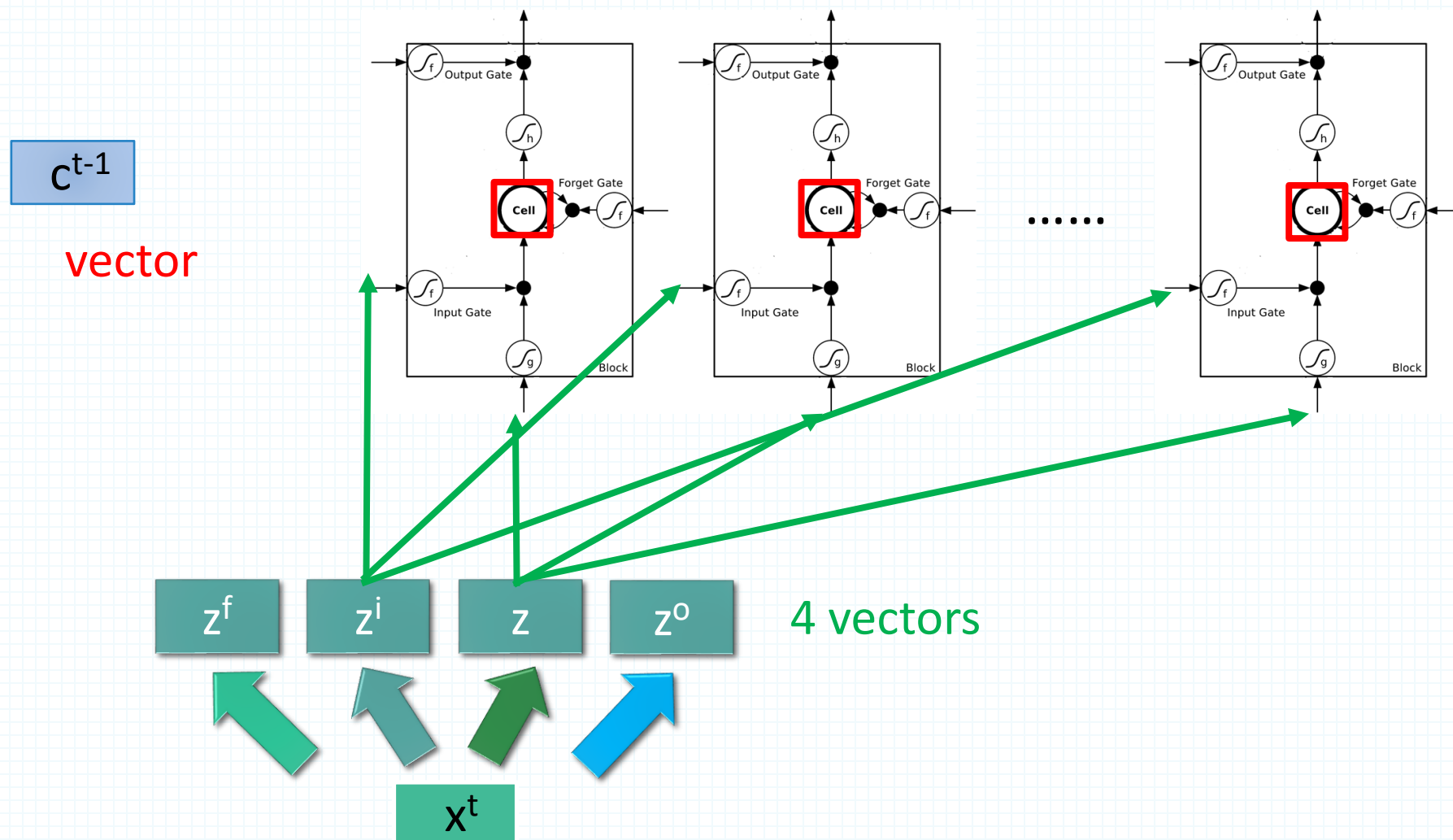
Activation function  $f$  is usually a sigmoid function

Between 0 and 1

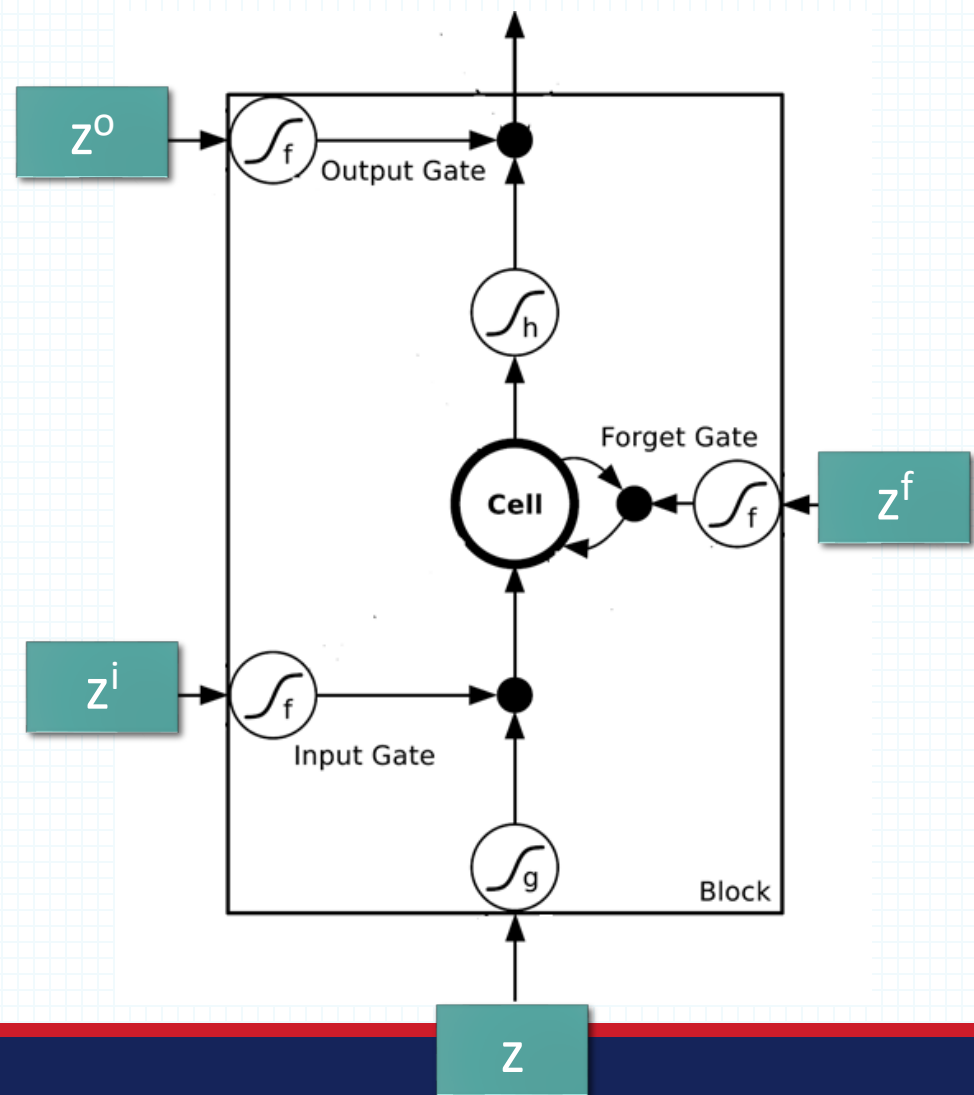
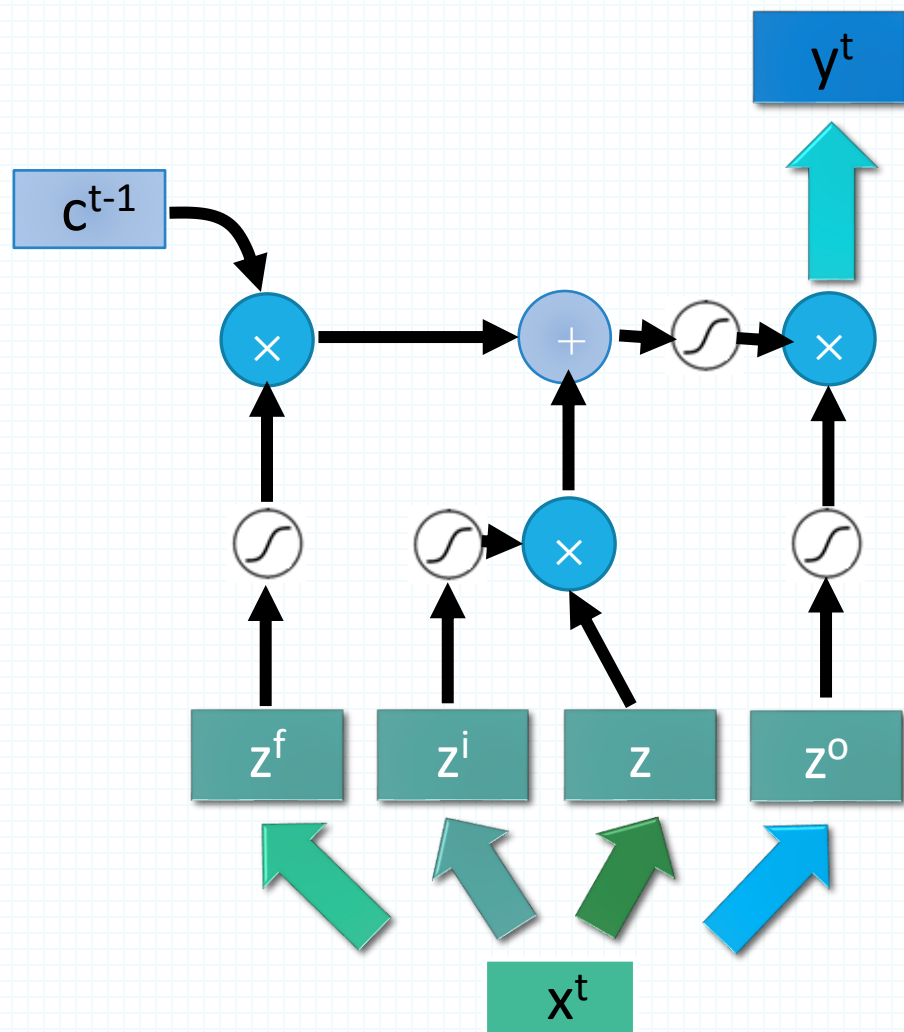
Mimic open and close gate

$$c' = g(z)f(z_i) + cf(z_f)$$

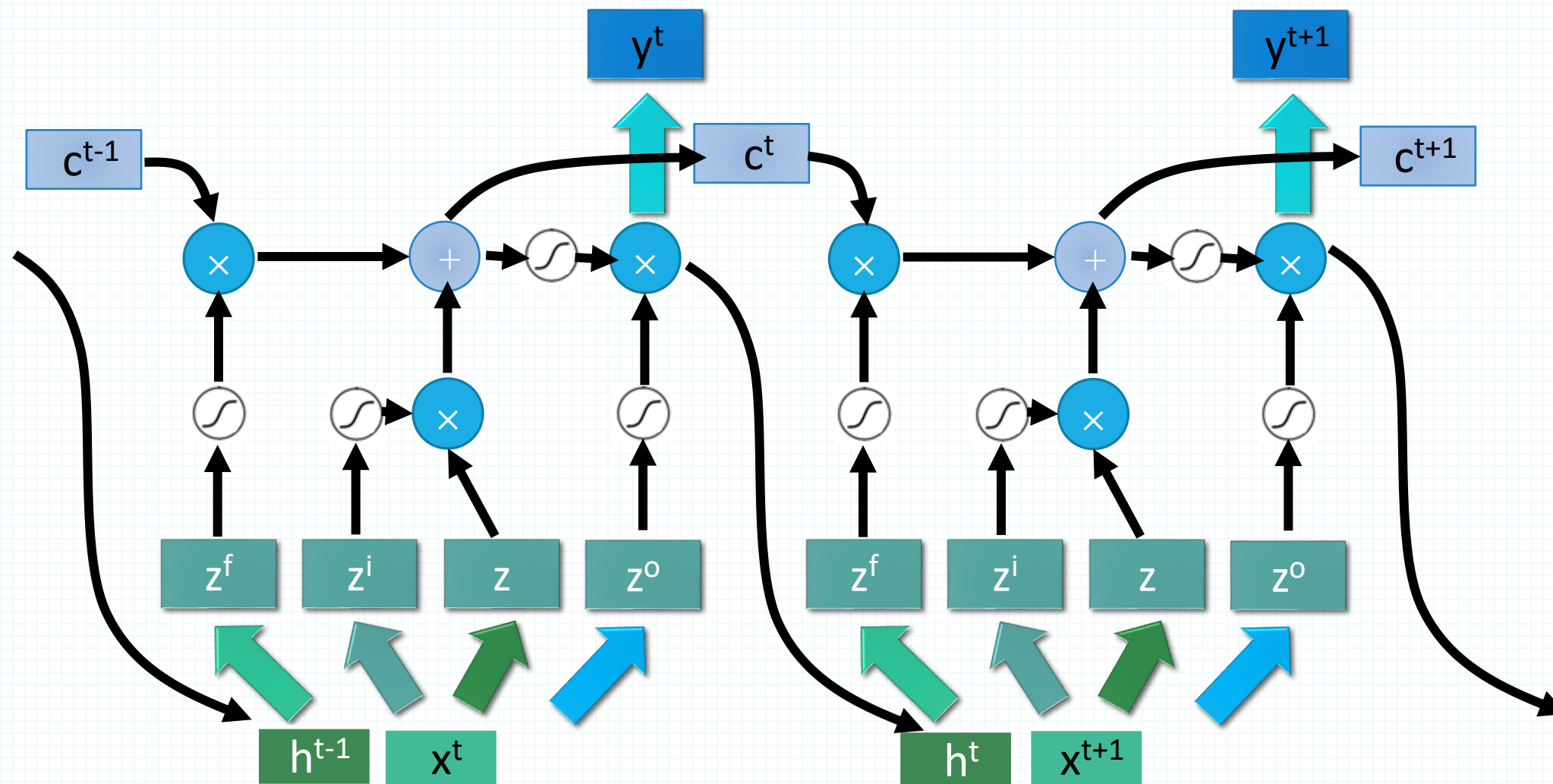
# LSTM



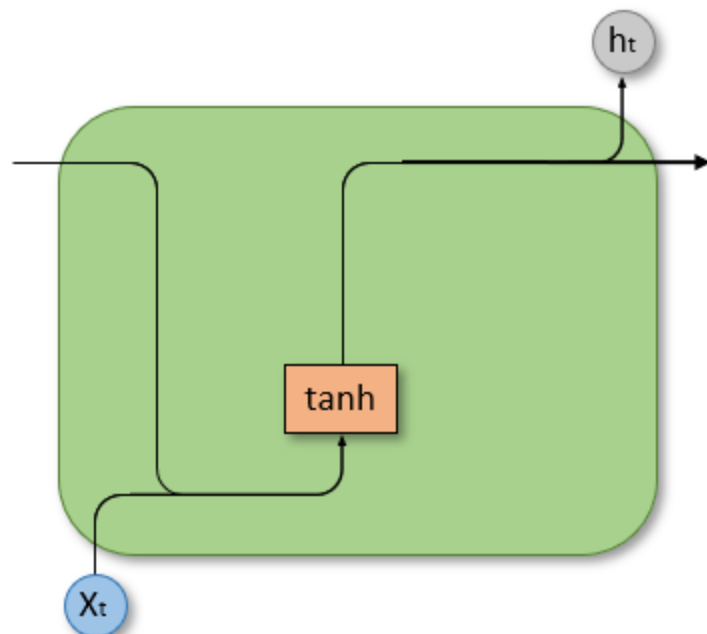
# LSTM



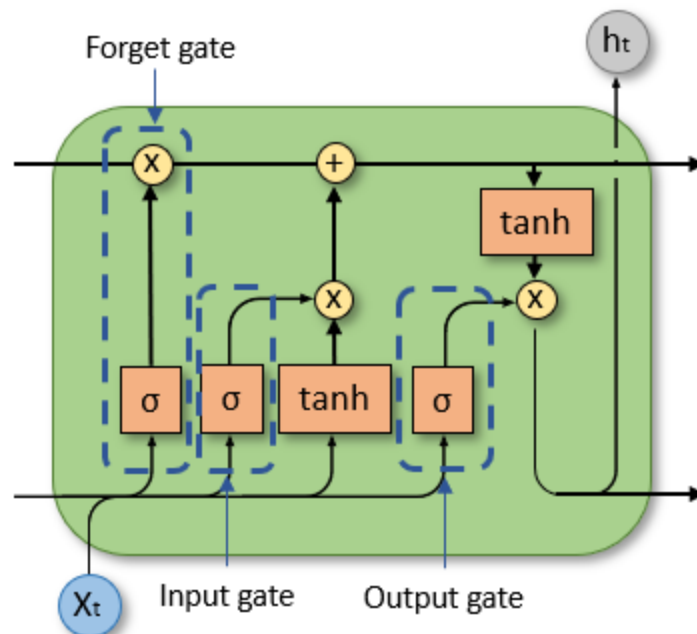
# LSTM



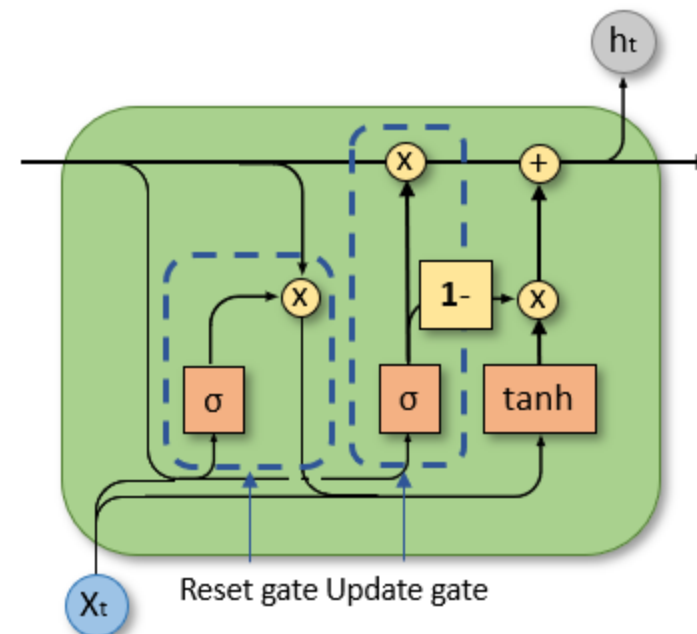
# LSTM vs GRU



RNN



LSTM



GRU

# Limitations of RNN/LSTM

---

- Sequential computation
- Slow training
- Difficult long-range dependencies
- Hard to parallelize

These limitations motivated the development of Attention and Transformers.

# Encoder-Decoder Model (Seq2Seq Model)

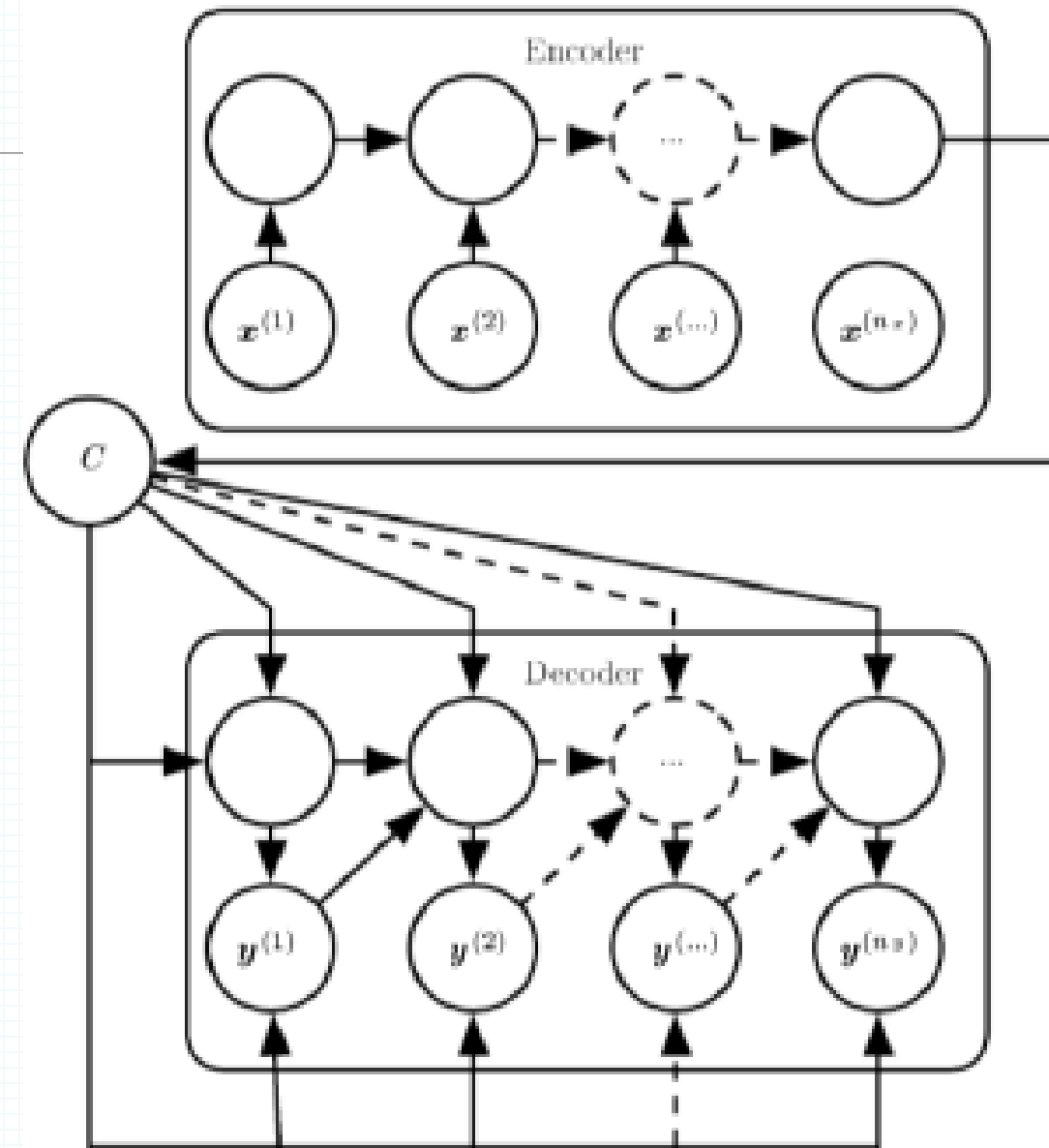
---

# Encoder-Decoder Model

- $X^{(i)}$ :  $i^{\text{th}}$  input
- $Y^{(i)}$ :  $i^{\text{th}}$  output
- $C$ : context (embedding)

## Usage:

- Machine translation
- Question answering
- Dialog



# Image Caption Generation

- Input an image, but output a sequence of words

[Kelvin Xu, arXiv'15][Li Yao, ICCV'15]

