

Convolutional Neural Networks



Multi-layer Neural Network

$$a^{(1)} = x$$

$$z^{(2)} = \Theta^{(1)} a^{(1)}$$

$$a^{(2)} = g(z^{(2)})$$

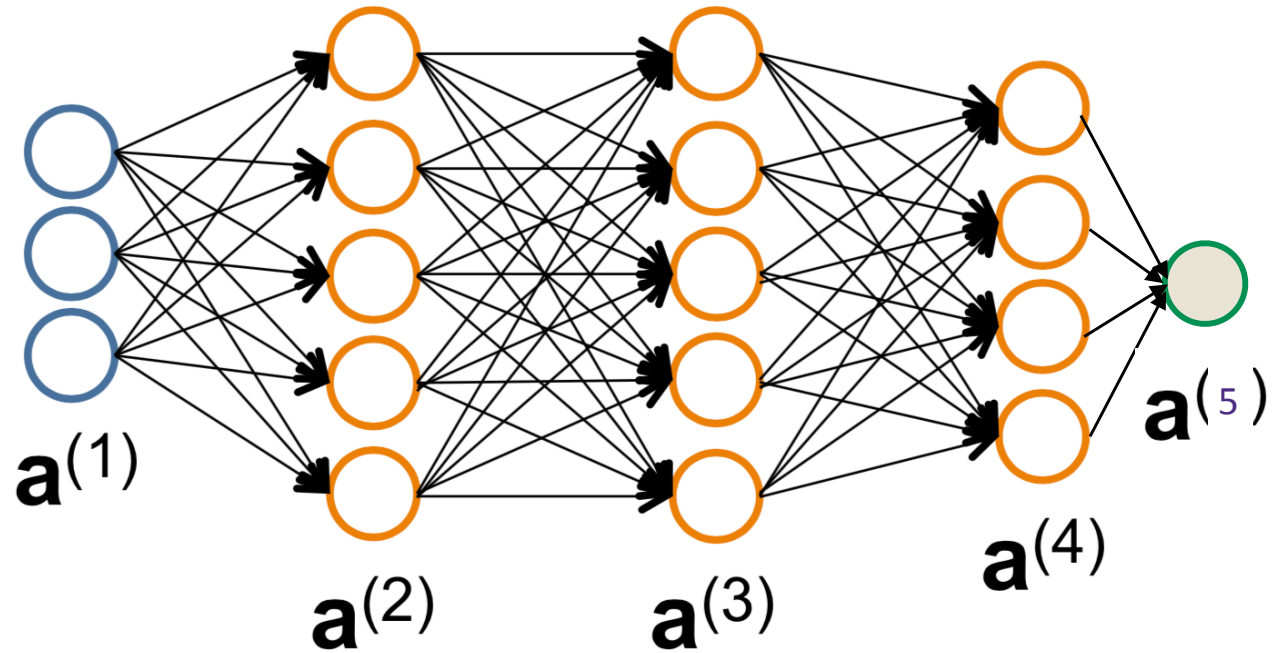
$$\vdots$$

$$z^{(l+1)} = \Theta^{(l)} a^{(l)}$$

$$a^{(l+1)} = g(z^{(l+1)})$$

$$\vdots$$

$$\hat{y} = a^{(L+1)}$$



$$L(y, \hat{y}) = y \log(\hat{y}) + (1 - y) \log(1 - \hat{y})$$

$$g(z) = \frac{1}{1 + e^{-z}}$$

Binary
Logistic
Regression

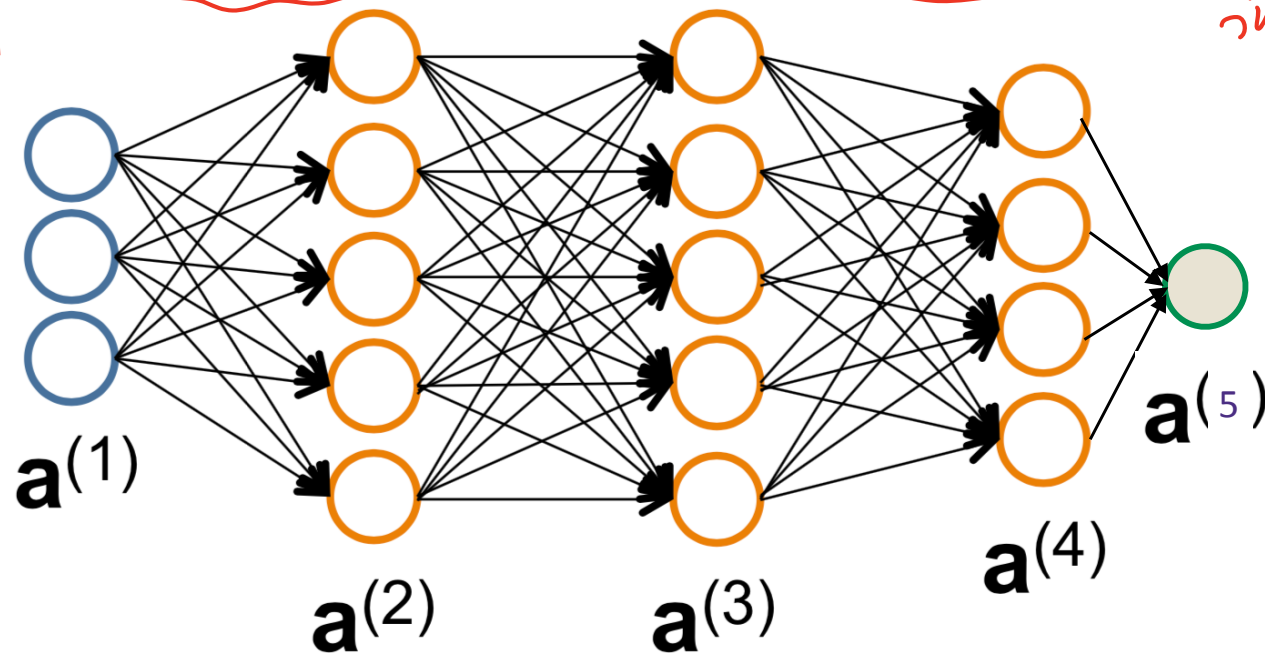
Neural Network Architecture

The neural network architecture is defined by the number of layers, and the number of nodes in each layer, but also by allowable edges.

depth

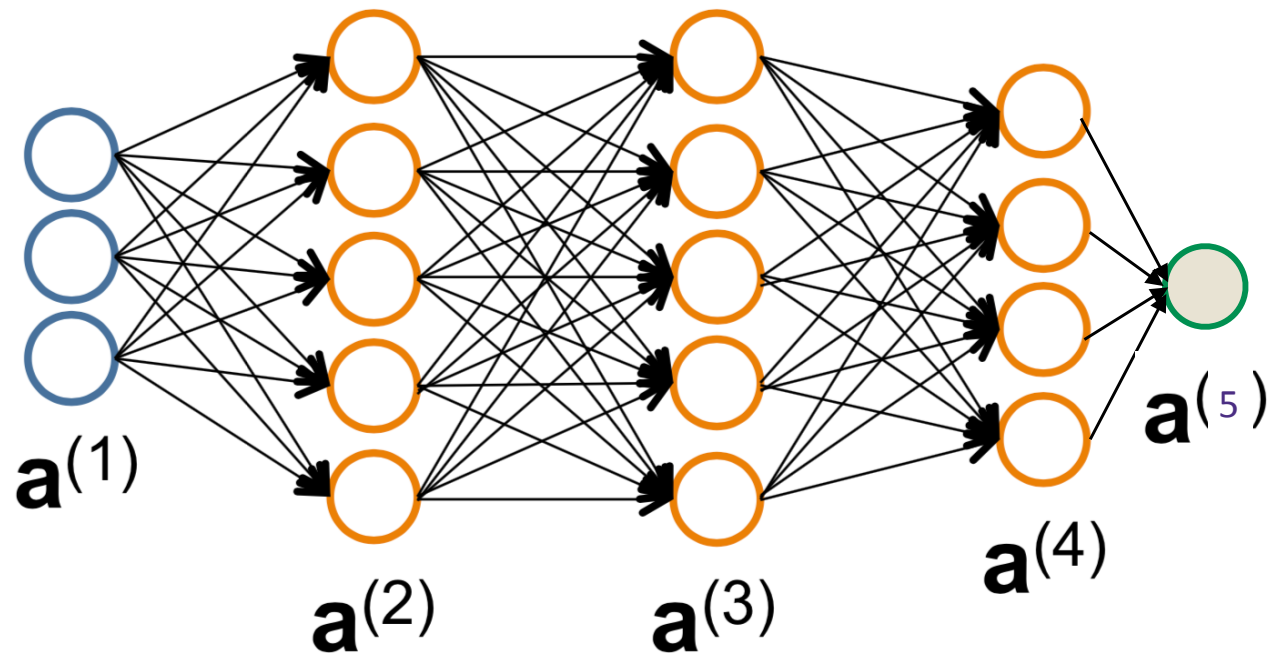
subject prior

width



Neural Network Architecture

The neural network architecture is defined by the number of layers, and the number of nodes in each layer, but also by **allowable edges**.



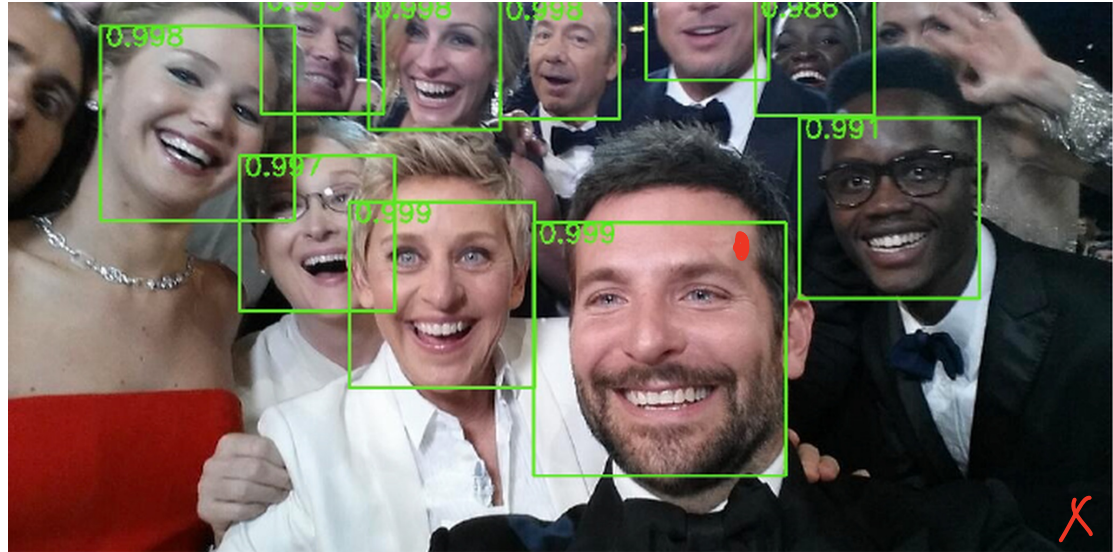
We say a layer is **Fully Connected (FC)** if all linear mappings from the current layer to the next layer are permissible.

$$\mathbf{a}^{(k+1)} = g(\Theta \mathbf{a}^{(k)}) \quad \text{for any } \Theta \in \mathbb{R}^{n_{k+1} \times n_k}$$

A lot of parameters!! $n_1 n_2 + n_2 n_3 + \cdots + n_L n_{L+1}$

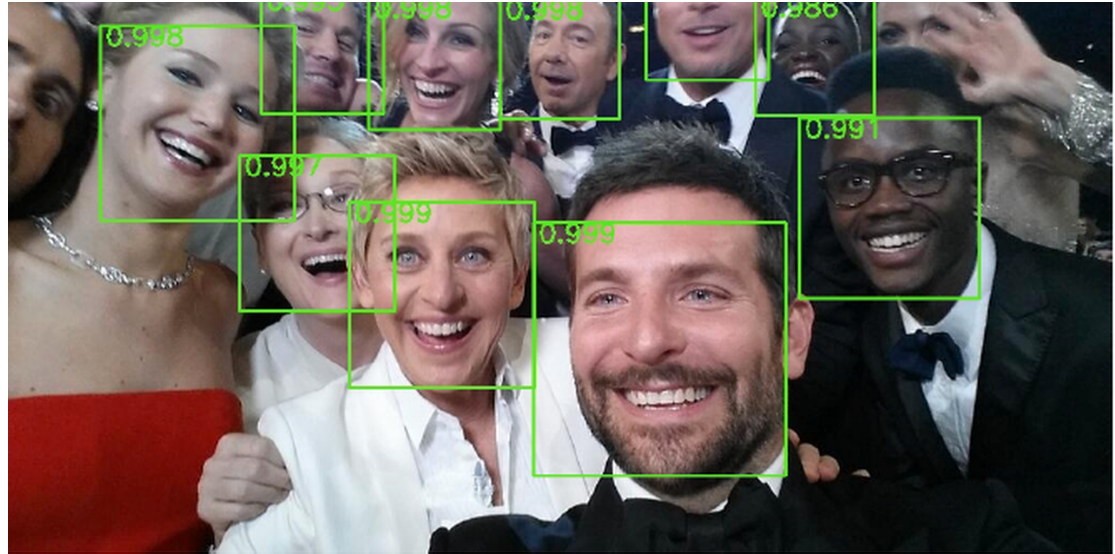
Neural Network Architecture

Objects are often **localized in space** so to find the faces in an image, not every pixel is important for classification—makes sense to drag a window across an image.

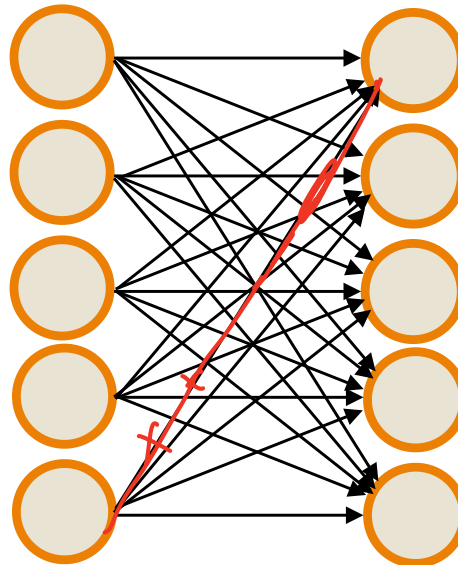


Neural Network Architecture

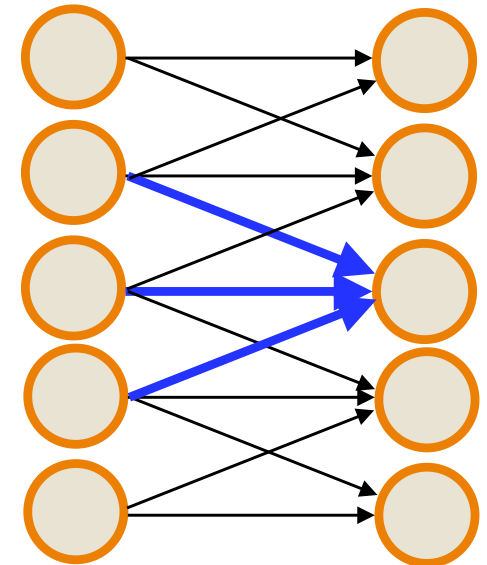
Objects are often **localized in space** so to find the faces in an image, not every pixel is important for classification—makes sense to drag a window across an image.



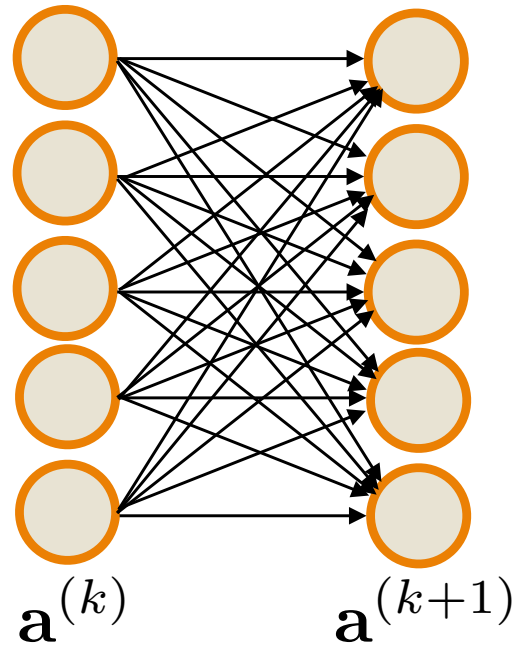
Similarly, to identify edges or other local structure, it makes sense to only look at **local information**



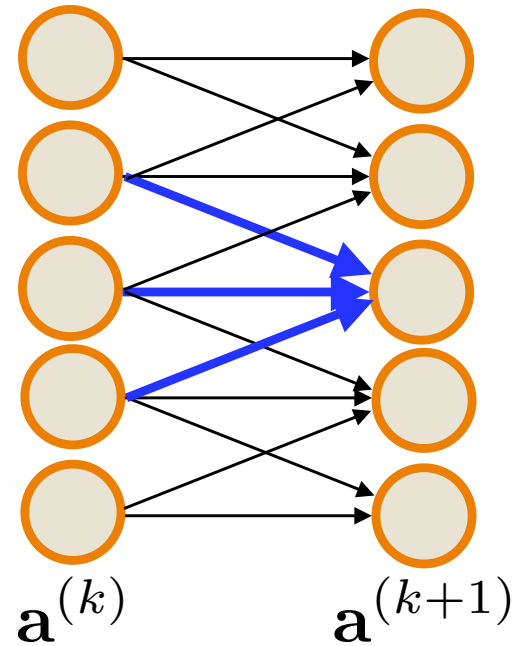
vs.



Neural Network Architecture



vs.



$$\begin{bmatrix} \Theta_{0,0} & \Theta_{0,1} & \Theta_{0,2} & \Theta_{0,3} & \Theta_{0,4} \\ \Theta_{1,0} & \Theta_{1,1} & \Theta_{1,2} & \Theta_{1,3} & \Theta_{1,4} \\ \Theta_{2,0} & \Theta_{2,1} & \Theta_{2,2} & \Theta_{2,3} & \Theta_{2,4} \\ \Theta_{3,0} & \Theta_{3,1} & \Theta_{3,2} & \Theta_{3,3} & \Theta_{3,4} \\ \Theta_{4,0} & \Theta_{4,1} & \Theta_{4,2} & \Theta_{4,3} & \Theta_{4,4} \end{bmatrix}$$

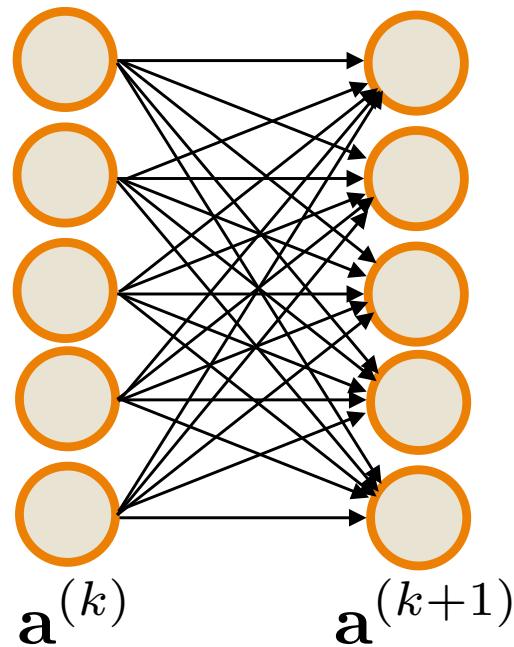
Parameters: n^2

$$\begin{bmatrix} \Theta_{0,0} & \Theta_{0,1} & 0 & 0 & 0 \\ \Theta_{1,0} & \Theta_{1,1} & \Theta_{1,2} & 0 & 0 \\ 0 & \Theta_{2,1} & \Theta_{2,2} & \Theta_{2,3} & 0 \\ 0 & 0 & \Theta_{3,2} & \Theta_{3,3} & \Theta_{3,4} \\ 0 & 0 & 0 & \Theta_{4,3} & \Theta_{4,4} \end{bmatrix}$$

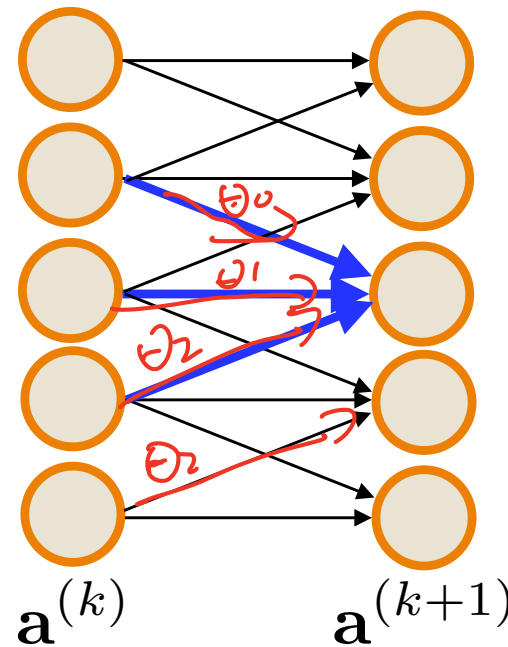
$3n - 2$

$$\mathbf{a}_i^{(k+1)} = g \left(\sum_{j=0}^{n-1} \Theta_{i,j} \mathbf{a}_j^{(k)} \right)$$

Neural Network Architecture



vs.



Mirror/share local weights everywhere
(e.g., structure equally likely to be anywhere in image)

$$\begin{bmatrix} \Theta_{0,0} & \Theta_{0,1} & \Theta_{0,2} & \Theta_{0,3} & \Theta_{0,4} \\ \Theta_{1,0} & \Theta_{1,1} & \Theta_{1,2} & \Theta_{1,3} & \Theta_{1,4} \\ \Theta_{2,0} & \Theta_{2,1} & \Theta_{2,2} & \Theta_{2,3} & \Theta_{2,4} \\ \Theta_{3,0} & \Theta_{3,1} & \Theta_{3,2} & \Theta_{3,3} & \Theta_{3,4} \\ \Theta_{4,0} & \Theta_{4,1} & \Theta_{4,2} & \Theta_{4,3} & \Theta_{4,4} \end{bmatrix}$$

Parameters: n^2

$$\begin{bmatrix} \Theta_{0,0} & \Theta_{0,1} & 0 & 0 & 0 \\ \Theta_{1,0} & \Theta_{1,1} & \Theta_{1,2} & 0 & 0 \\ 0 & \Theta_{2,1} & \Theta_{2,2} & \Theta_{2,3} & 0 \\ 0 & 0 & \Theta_{3,2} & \Theta_{3,3} & \Theta_{3,4} \\ 0 & 0 & 0 & \Theta_{4,3} & \Theta_{4,4} \end{bmatrix}$$

$3n - 2$

$$\begin{bmatrix} \theta_1 & \theta_2 & 0 & 0 & 0 \\ \theta_0 & \theta_1 & \theta_2 & 0 & 0 \\ 0 & \theta_0 & \theta_1 & \theta_2 & 0 \\ 0 & 0 & \theta_0 & \theta_1 & \theta_2 \\ 0 & 0 & 0 & \theta_0 & \theta_1 \end{bmatrix}$$

3

$$\mathbf{a}_i^{(k+1)} = g \left(\sum_{j=0}^{n-1} \Theta_{i,j} \mathbf{a}_j^{(k)} \right)$$

$$\mathbf{a}_i^{(k+1)} = g \left(\sum_{j=0}^{m-1} \theta_j \mathbf{a}_{i+j}^{(k)} \right)$$

Neural Network Architecture

Fully Connected (FC) Layer

$$\begin{bmatrix} \Theta_{0,0} & \Theta_{0,1} & \Theta_{0,2} & \Theta_{0,3} & \Theta_{0,4} \\ \Theta_{1,0} & \Theta_{1,1} & \Theta_{1,2} & \Theta_{1,3} & \Theta_{1,4} \\ \Theta_{2,0} & \Theta_{2,1} & \Theta_{2,2} & \Theta_{2,3} & \Theta_{2,4} \\ \Theta_{3,0} & \Theta_{3,1} & \Theta_{3,2} & \Theta_{3,3} & \Theta_{3,4} \\ \Theta_{4,0} & \Theta_{4,1} & \Theta_{4,2} & \Theta_{4,3} & \Theta_{4,4} \end{bmatrix}$$

Convolutional (CONV) Layer (1 filter)

$$\begin{bmatrix} \theta_1 & \theta_2 & 0 & 0 & 0 \\ \theta_0 & \theta_1 & \theta_2 & 0 & 0 \\ 0 & \theta_0 & \theta_1 & \theta_2 & 0 \\ 0 & 0 & \theta_0 & \theta_1 & \theta_2 \\ 0 & 0 & 0 & \theta_0 & \theta_1 \end{bmatrix} \quad m=3$$

$$\mathbf{a}_i^{(k+1)} = g \left(\sum_{j=0}^{n-1} \Theta_{i,j} \mathbf{a}_j^{(k)} \right)$$

$$\mathbf{a}_i^{(k+1)} = g \left(\sum_{j=0}^{m-1} \theta_j \mathbf{a}_{i+j}^{(k)} \right) = g([\theta * \mathbf{a}^{(k)}]_i)$$

Convolution*

$\mathbf{a}^{(k)}$: k^{th} input
 $\mathbf{a}_i^{(k+1)}$: $k+1^{\text{th}}$ layer input at location i

$\theta = (\theta_0, \dots, \theta_{m-1}) \in \mathbb{R}^m$ is referred to as a “filter”

Example (1d convolution)

$$(\theta * x)_i = \sum_{j=0}^{m-1} \theta_j x_{i+j}$$

1	1	1	0	0
---	---	---	---	---

Input $x \in \mathbb{R}^n$

1	0	1
---	---	---

Filter $\theta \in \mathbb{R}^m$

--	--	--

Output $\theta * x$

Example (1d convolution)

1	1	1	0	0
---	---	---	---	---

Input $x \in \mathbb{R}^n$

$$(\theta * x)_i = \sum_{j=0}^{m-1} \theta_j x_{i+j}$$

1	0	1
---	---	---

Filter $\theta \in \mathbb{R}^m$

1 _{x1}	1 _{x0}	1 _{x1}	0	0
-----------------	-----------------	-----------------	---	---

2		
---	--	--

Output $\theta * x$

Example (1d convolution)

1	1	1	0	0
---	---	---	---	---

Input $x \in \mathbb{R}^n$

$$(\theta * x)_i = \sum_{j=0}^{m-1} \theta_j x_{i+j}$$

1	0	1
---	---	---

Filter $\theta \in \mathbb{R}^m$

1	1 _{x1}	1 _{x0}	0 _{x1}	0
---	-----------------	-----------------	-----------------	---

2	1	
---	---	--

Output $\theta * x$

Example (1d convolution)

$$(\theta * x)_i = \sum_{j=0}^{m-1} \theta_j x_{i+j}$$

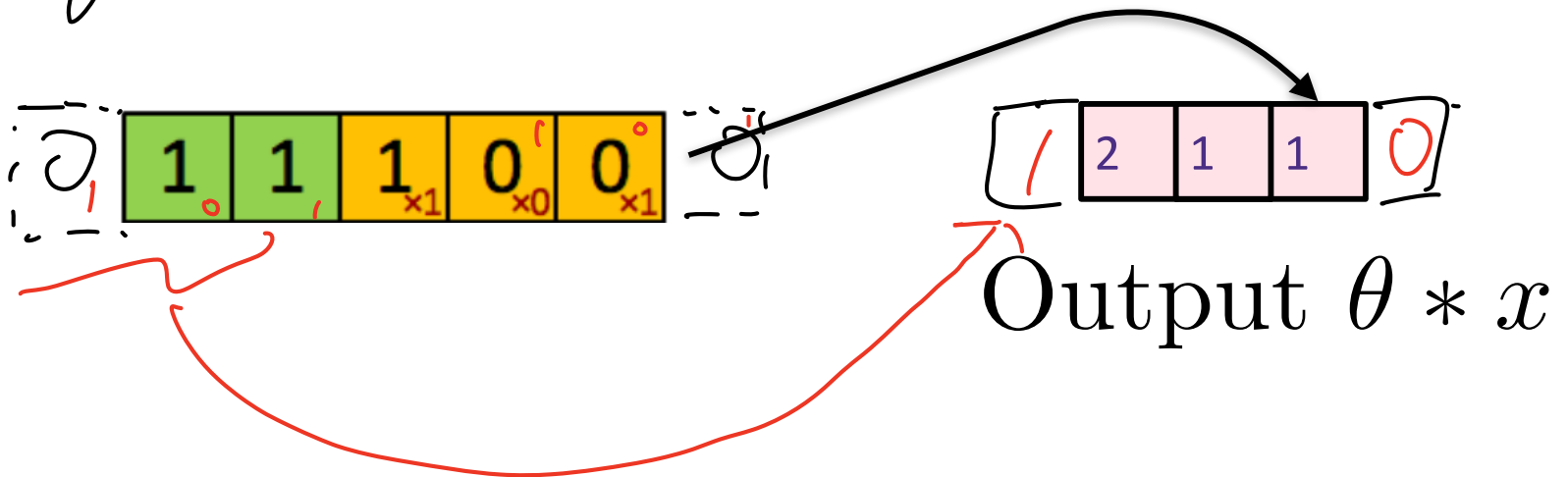
padding

1	1	1	0	0
---	---	---	---	---

Input $x \in \mathbb{R}^n$
 $\theta_0 \quad \theta_1 \quad \theta_2$

1	0	1
---	---	---

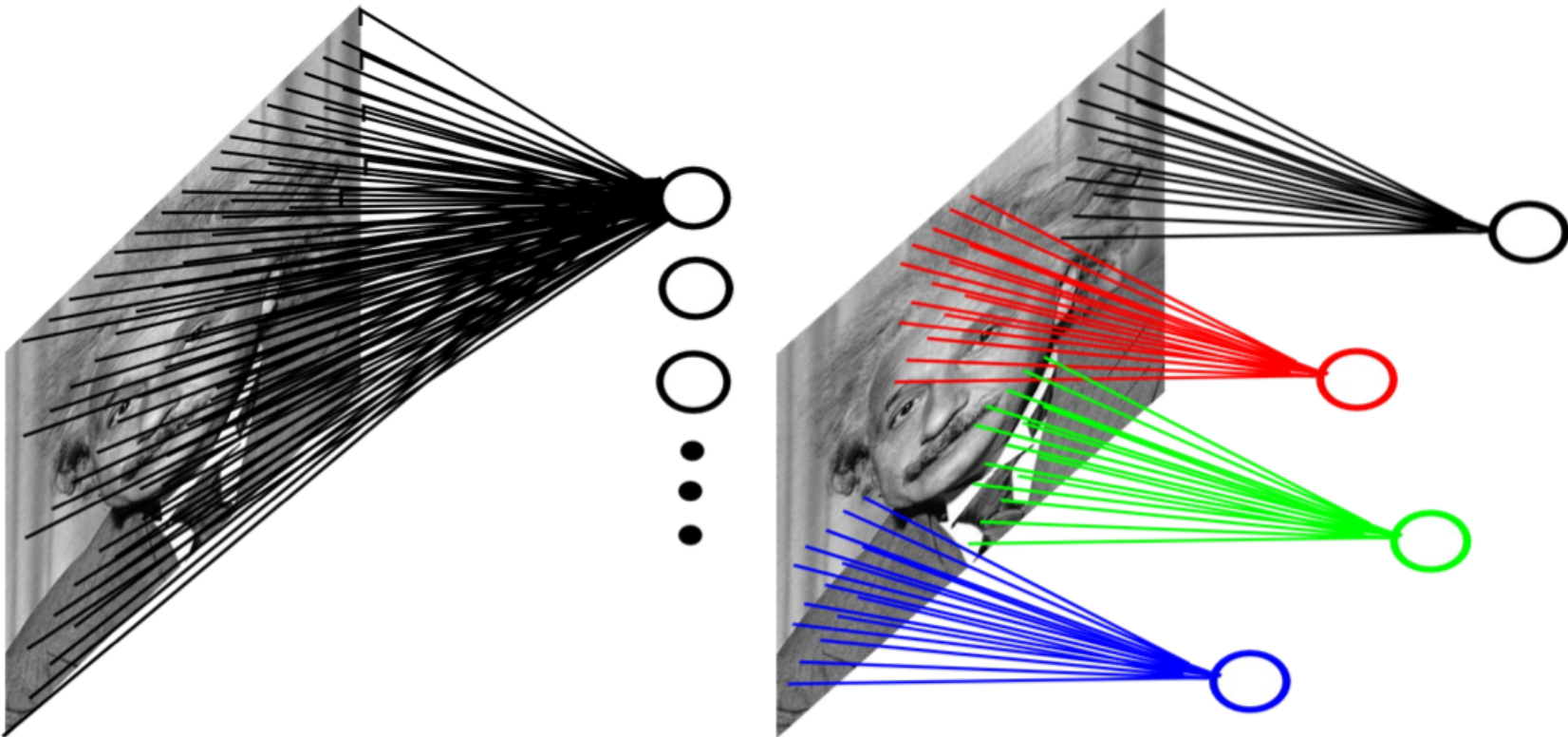
Filter $\theta \in \mathbb{R}^m$



2d Convolution Layer

■ Example: 200x200 image

- ▶ Fully-connected, 400,000 hidden units = 16 billion parameters
- ▶ Locally-connected, 400,000 hidden units 10x10 fields = 40 million params
- ▶ Local connections capture local dependencies



Convolution of images (2d convolution)

$1 \times 1 \rightarrow 7 \times 7$

$$(I * K)(i, j) = \sum_m \sum_n I(i + m, j + n) K(m, n)$$

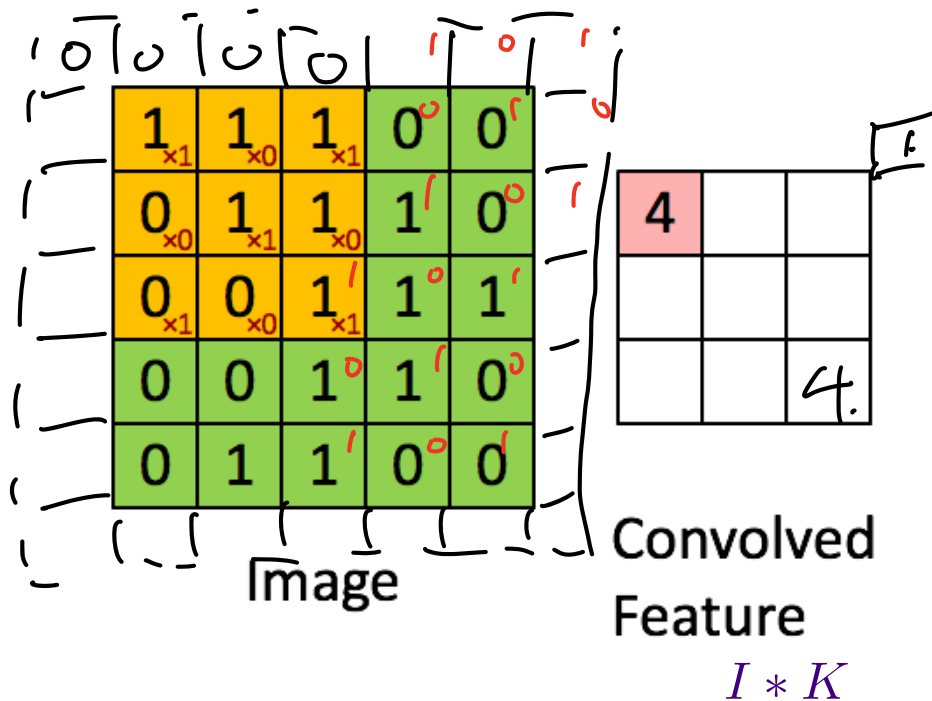
2D padding

1	1	1	0	0
0	1	1	1	0
0	0	1	1	1
0	0	1	1	0
0	1	1	0	0

Image I

1	0	1
0	1	0
1	0	1

Filter K



Convolution of images

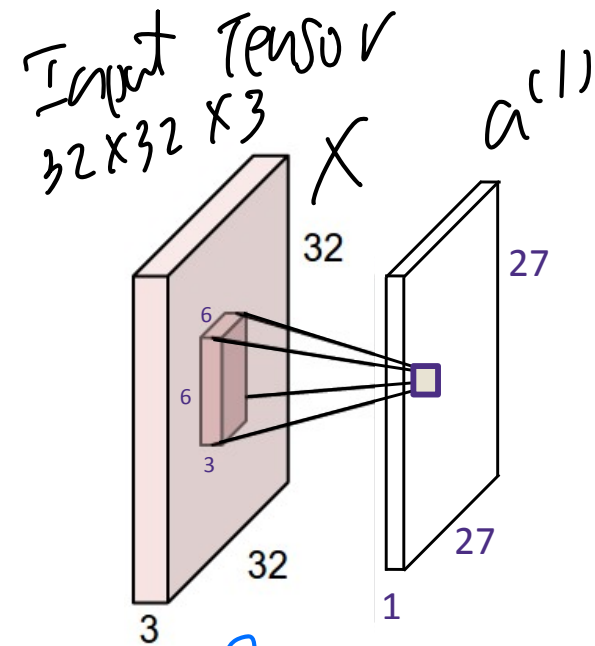
$$(I * K)(i, j) = \sum_m \sum_n I(i + m, j + n) K(m, n)$$

Image I



Operation	Filter K	Convolved Image $I * K$
Edge detection	$\begin{bmatrix} 1 & 0 & -1 \\ 0 & 0 & 0 \\ -1 & 0 & 1 \end{bmatrix}$	
	$\begin{bmatrix} 0 & 1 & 0 \\ 1 & -4 & 1 \\ 0 & 1 & 0 \end{bmatrix}$	
	$\begin{bmatrix} -1 & -1 & -1 \\ -1 & 8 & -1 \\ -1 & -1 & -1 \end{bmatrix}$	
Sharpen	$\begin{bmatrix} 0 & -1 & 0 \\ -1 & 5 & -1 \\ 0 & -1 & 0 \end{bmatrix}$	
Box blur (normalized)	$\frac{1}{9} \begin{bmatrix} 1 & 1 & 1 \\ 1 & 1 & 1 \\ 1 & 1 & 1 \end{bmatrix}$	
Gaussian blur (approximation)	$\frac{1}{16} \begin{bmatrix} 1 & 2 & 1 \\ 2 & 4 & 2 \\ 1 & 2 & 1 \end{bmatrix}$	

Stacking convolved images

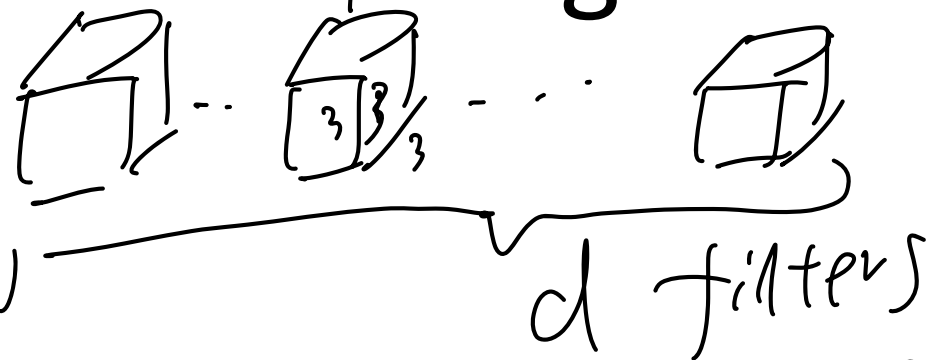
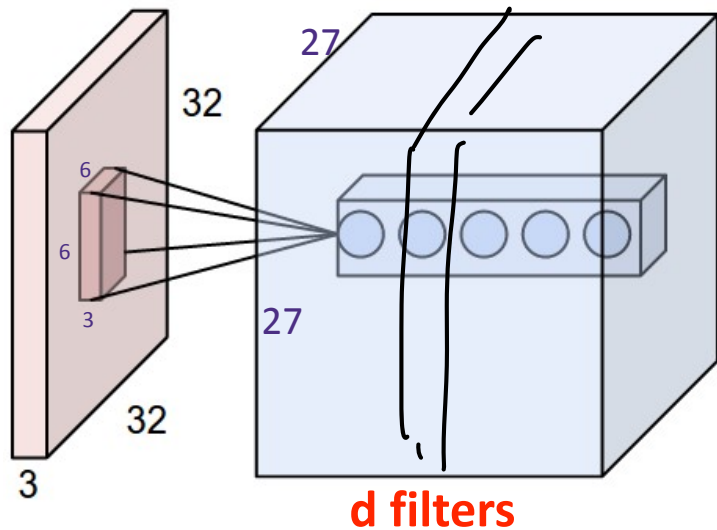


R G B
 channels

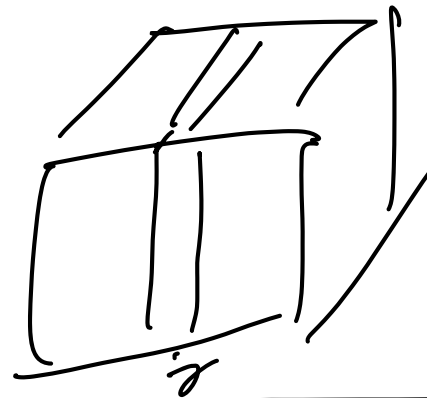
$$a^{(1)} = g \left(\sum_{d=1}^3 x[i, j, d] * K[i, j, d] \right)$$

$$x \in \mathbb{R}^{n \times n \times r}$$

Stacking convolved images



Repeat with d filters!



Second layer

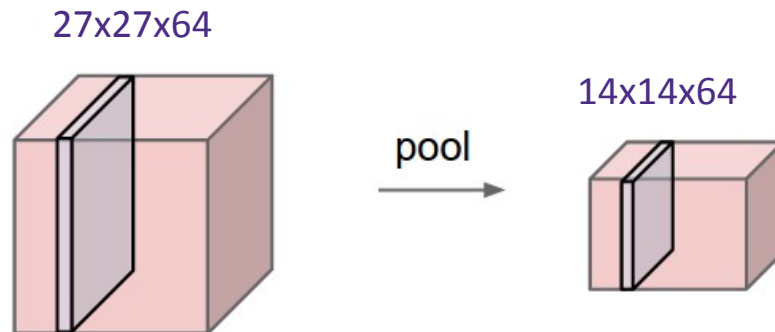
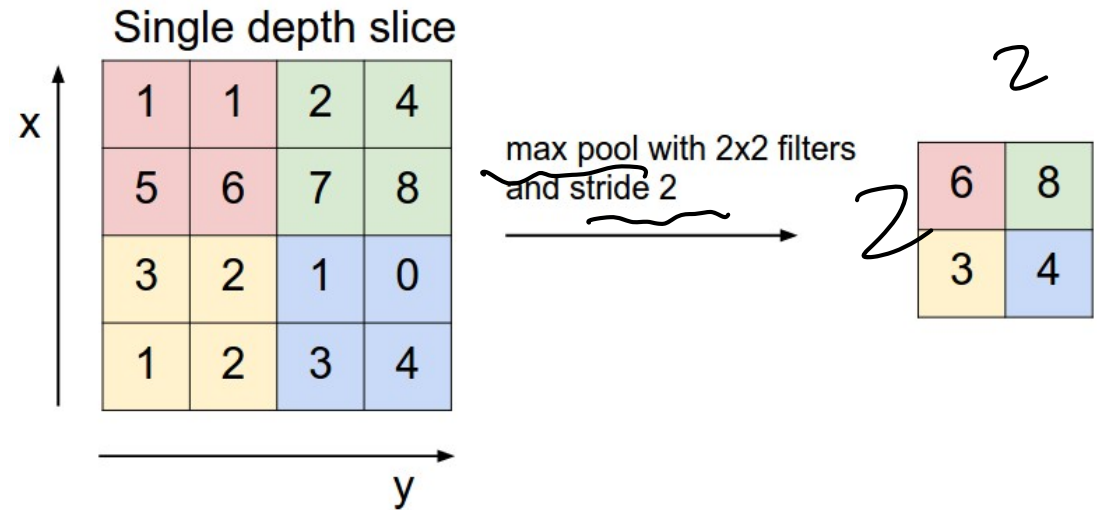


$$a_j^{(2)} = g\left(\sum_{i=1}^d a_{[:, :, i]}^{(1)} * K_i^{(2)}[:, :, i]\right)$$

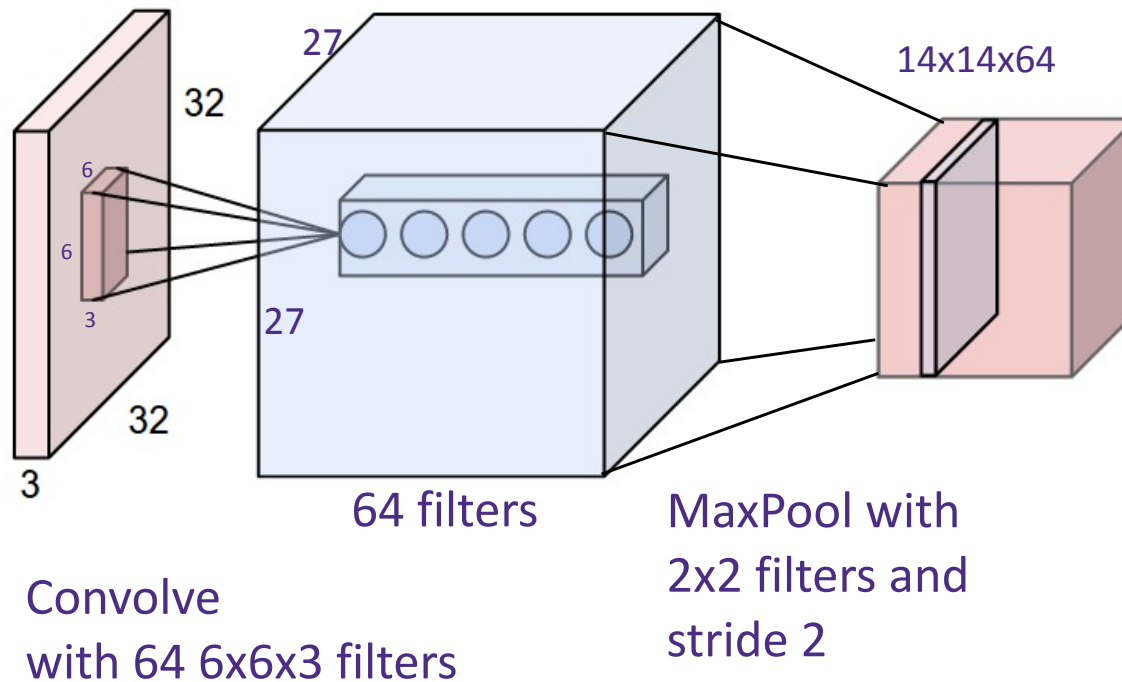
Pooling

average
pooling

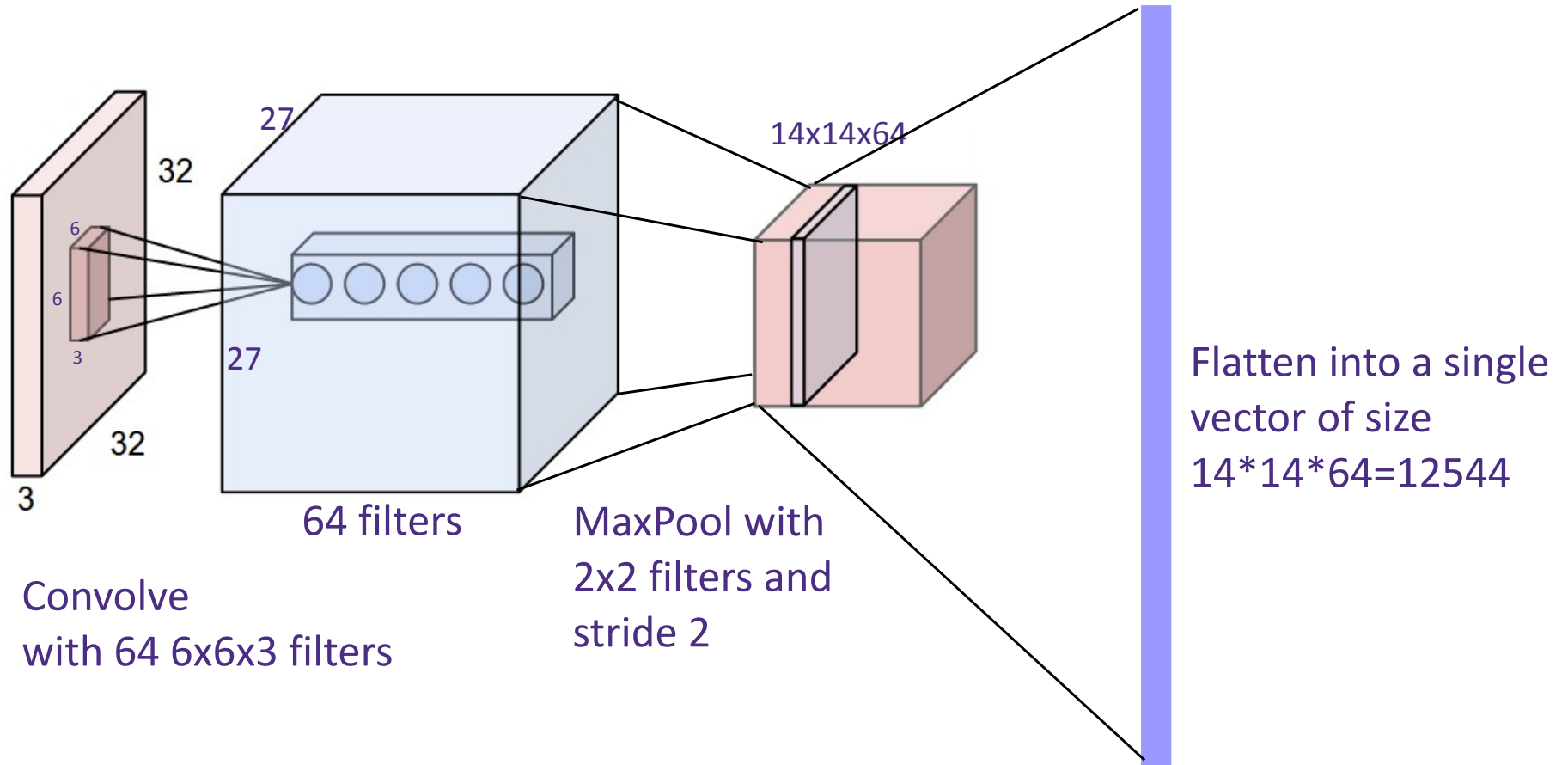
Pooling reduces the dimension and can be interpreted as “This filter had a high response in this general region”



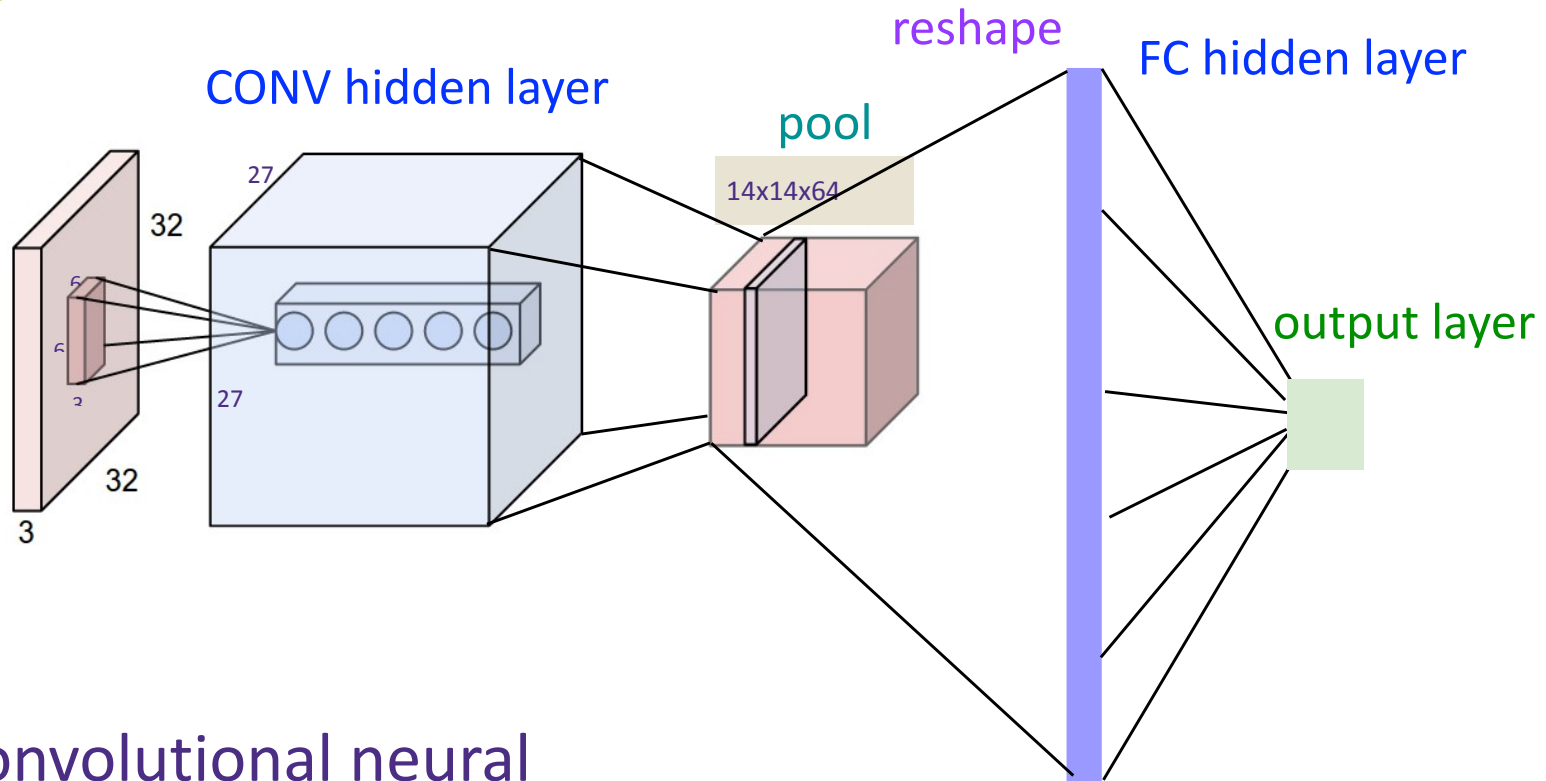
Pooling Convolution layer



Flattening

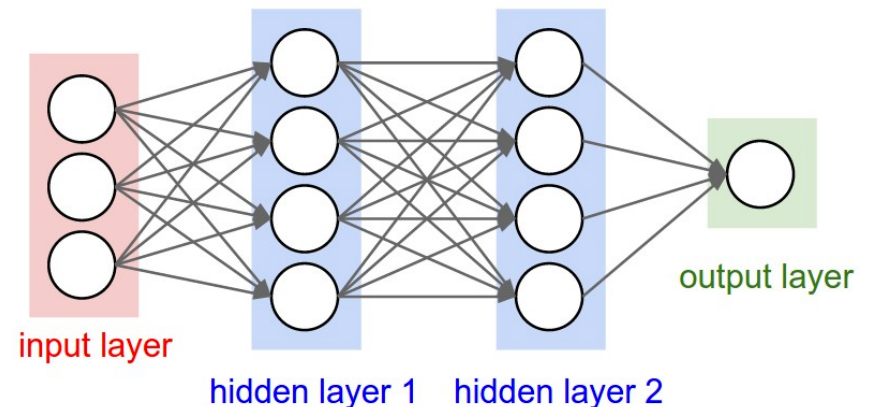


Training Convolutional Networks

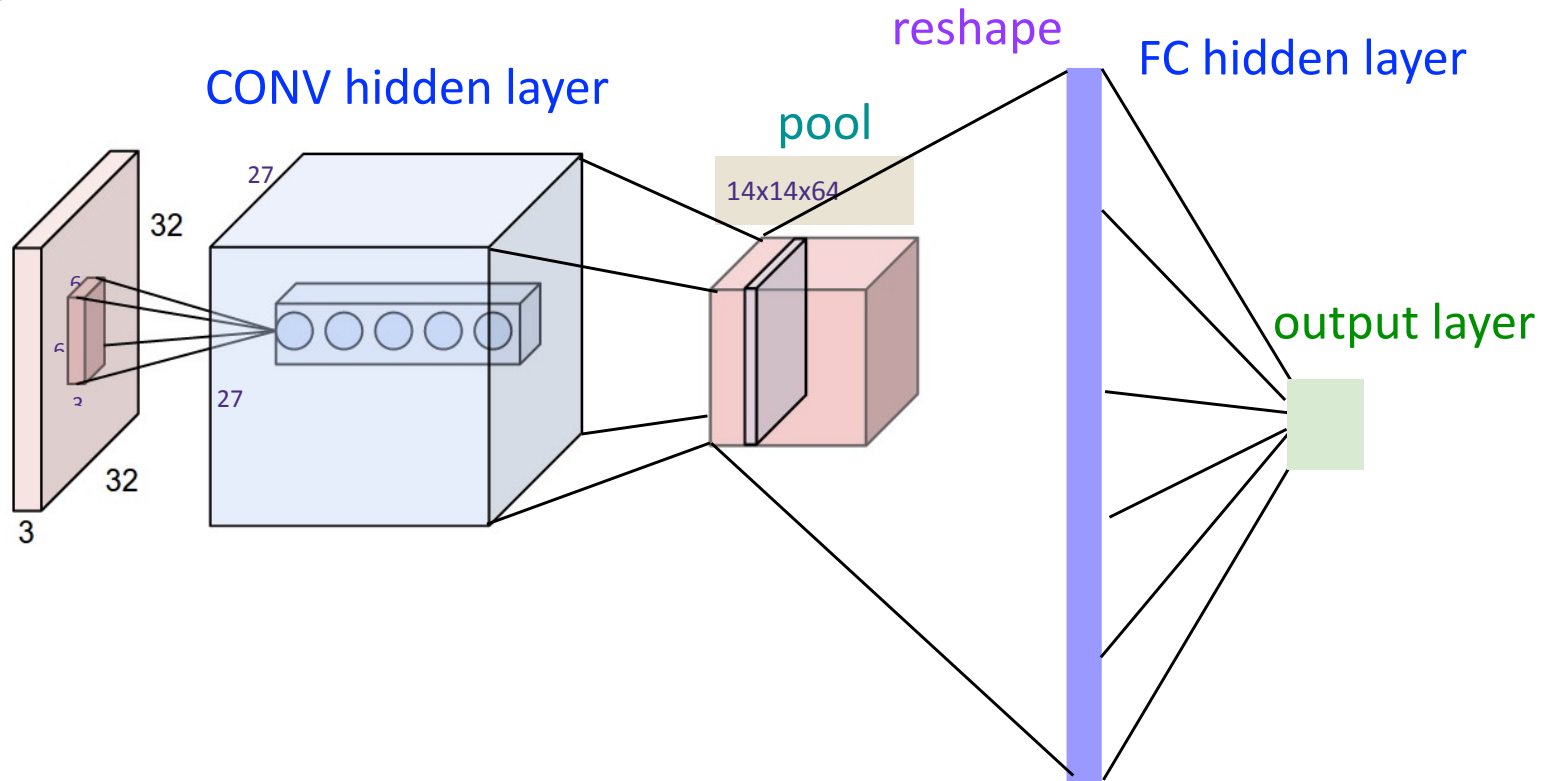


Recall: Convolutional neural networks (CNN) are just regular fully connected (FC) neural networks with some connections removed.

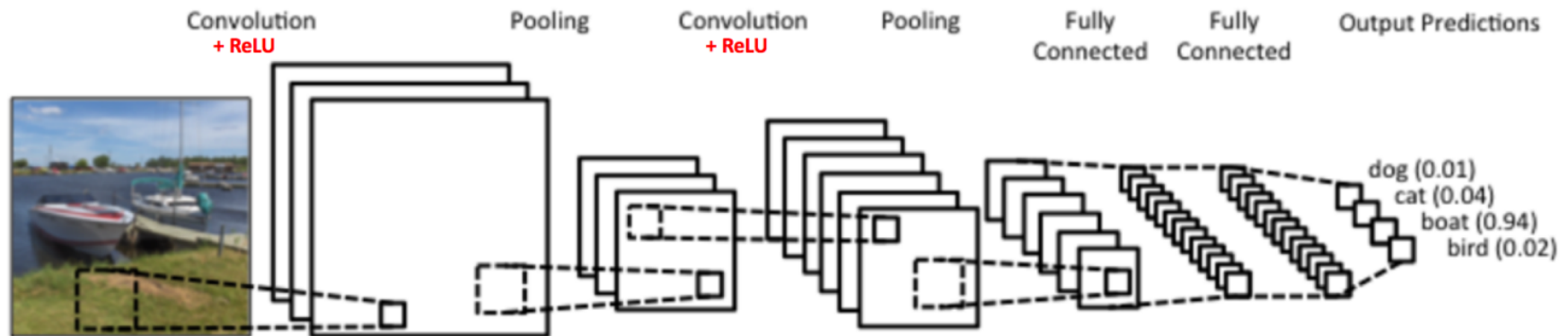
Train with SGD!

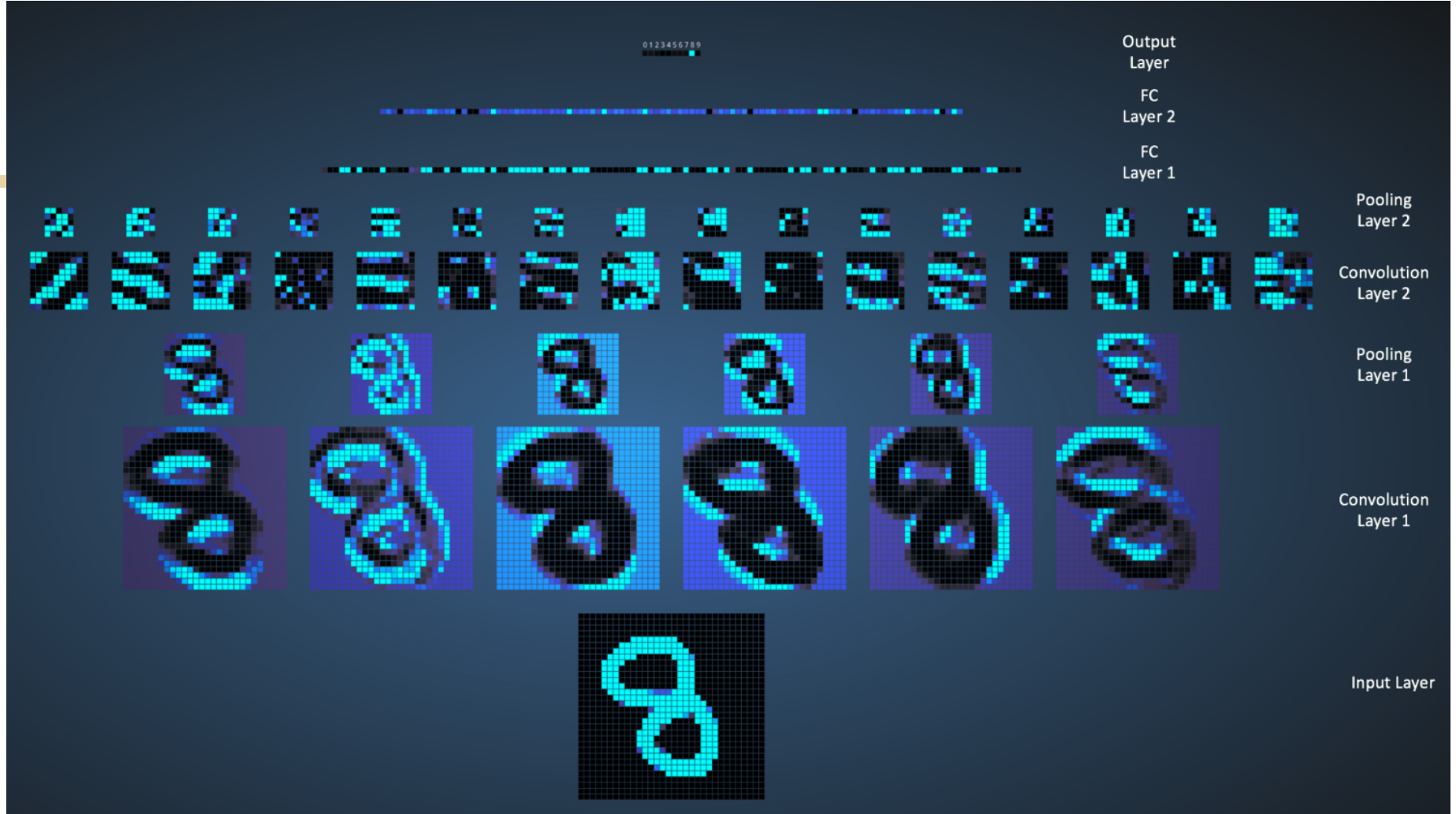


Training Convolutional Networks

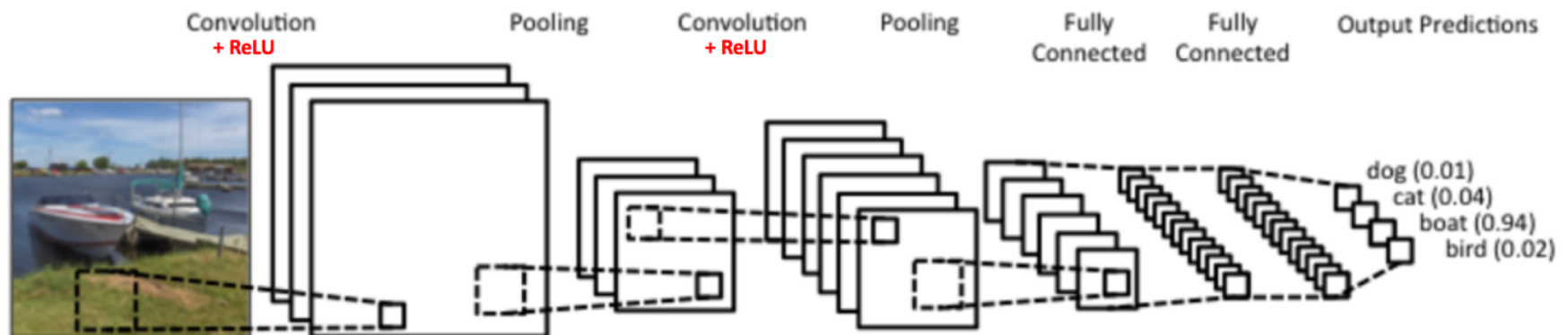


Real example network: LeNet





Real example network: LeNet



Famous CNNs

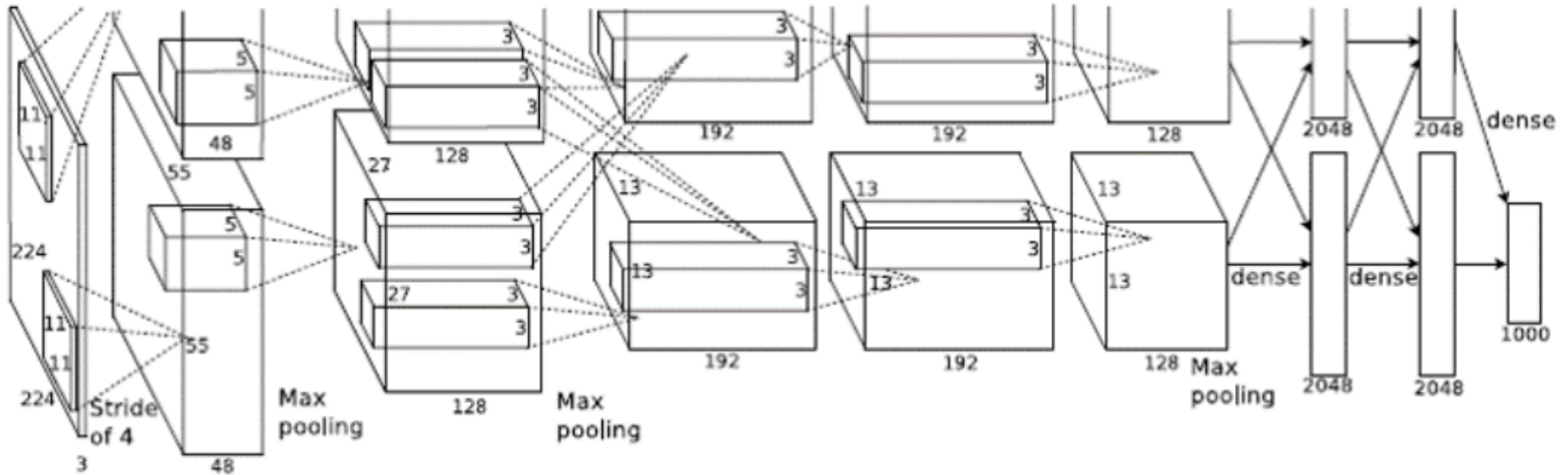
W

ImageNet Dataset

~14 million images, 20k classes



Deng et al. "Imagenet: a large scale hierarchical image database" '09



Krizhevsky, Sutskever, Hinton “ImageNet Claasification with Deep Convolutional Neural Networks”, NIPS 2012.

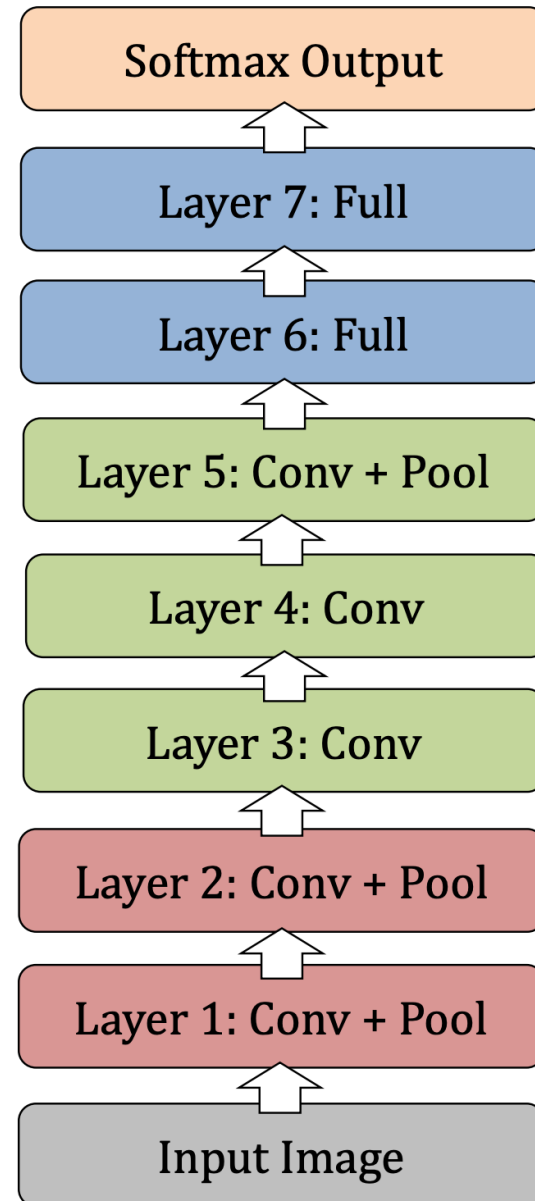
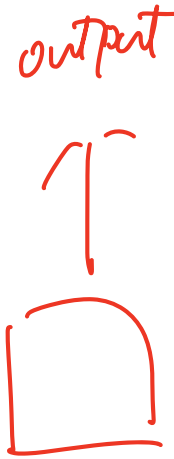
AlexNet

8 layers, ~60M parameters

Top5 error: 18.2%

Techniques used:

ReLU activation, overlapping pooling,
dropout, ensemble (create 10
patches by cropping and average the
predictions), data-augmentation
(intensity of RGB channels)



[From Rob Fergus' CIFAR 2016 tutorial]

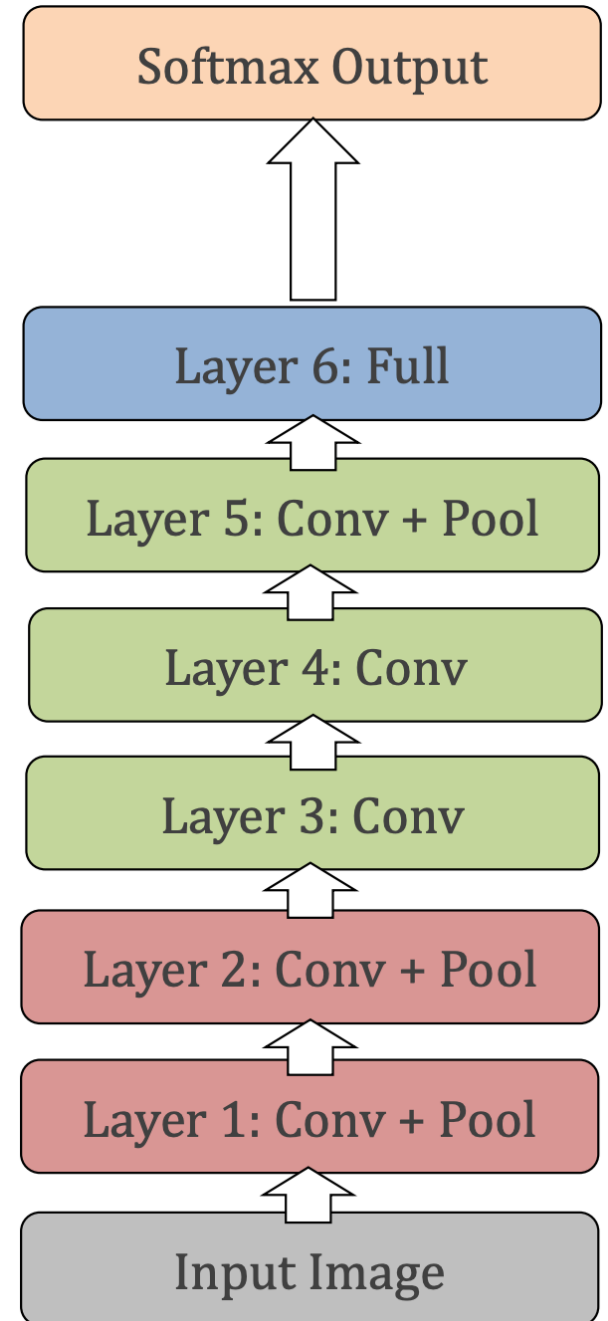
AlexNet

Remove top fully-connected layer 7

Drop ~16 million parameters

1.1% drop in performance

[From Rob Fergus' CIFAR 2016 tutorial]

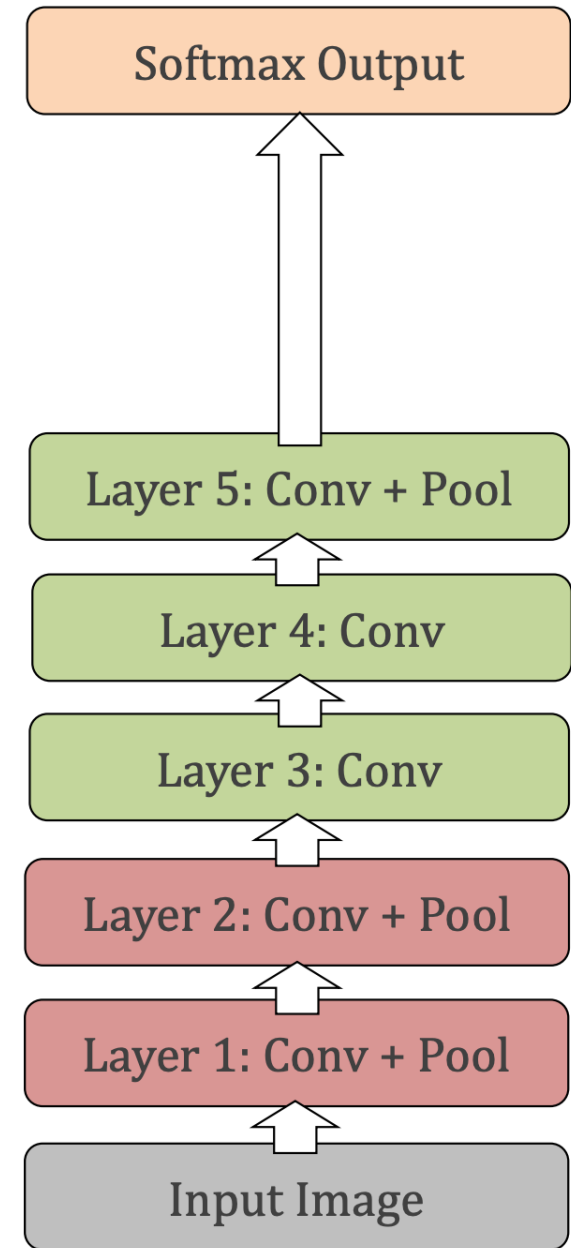


AlexNet

Remove both fully connected
layers 6 and 7

Drop ~50 million parameters

5.7% drop in performance



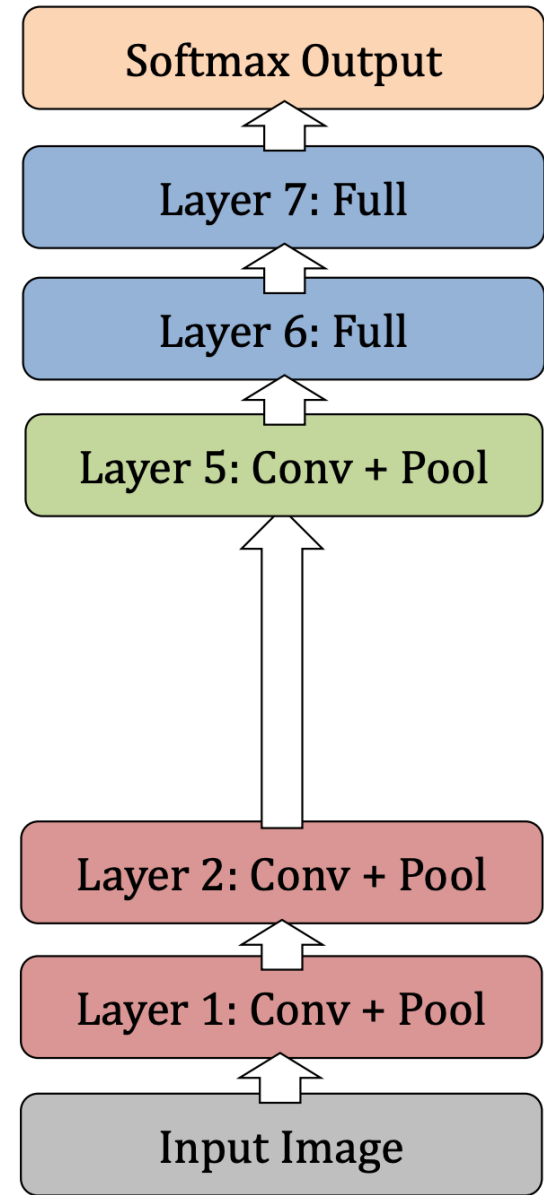
[From Rob Fergus' CIFAR 2016 tutorial]

AlexNet

Remove upper convolutio / feature extractor layers (layer 3 and 4)

Drop ~1 million parameters

3% drop in performance



[From Rob Fergus' CIFAR 2016 tutorial]

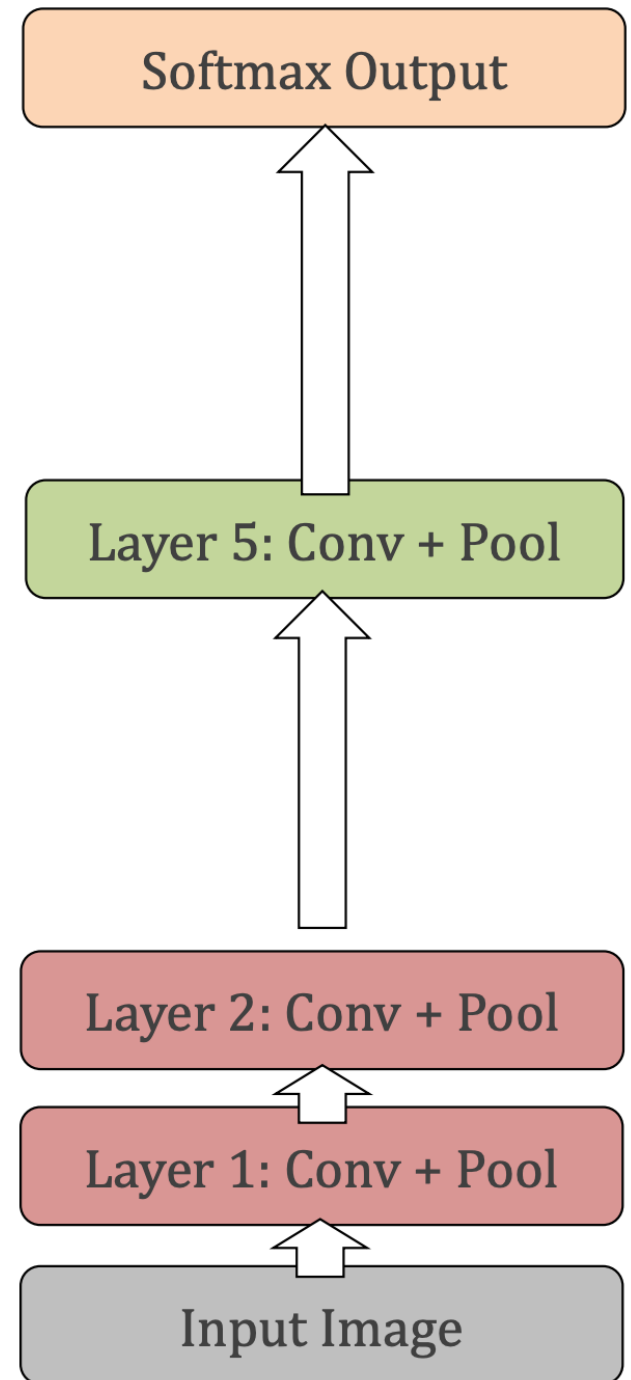
AlexNet

Remove top fully connected layer
6,7 and upper convolution layers
3,4.

33.5% drop in performance.

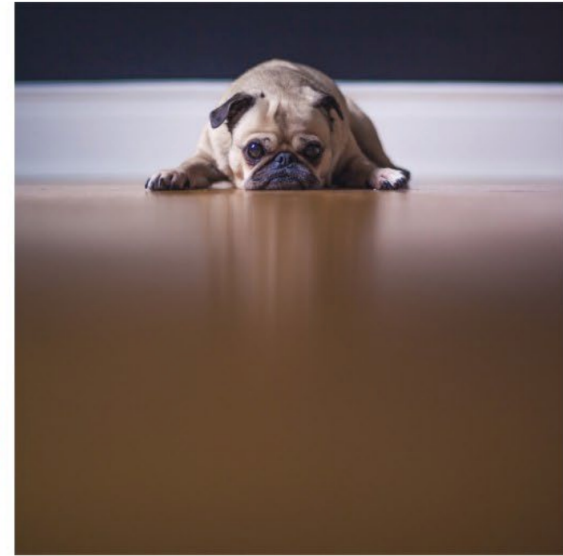
Depth of the network is the key.

[From Rob Fergus' CIFAR 2016 tutorial]



GoogLeNet

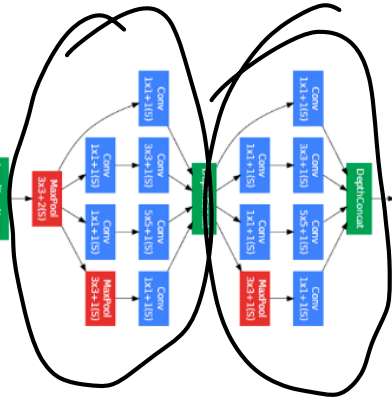
Motivation: multiscale nature of images



Large kernel for global features, and **smaller kernel** for local features.

Idea: have multiple different-size kernels at any layer.

[Going Deep with Convolutions, Szegedy et al. '14]

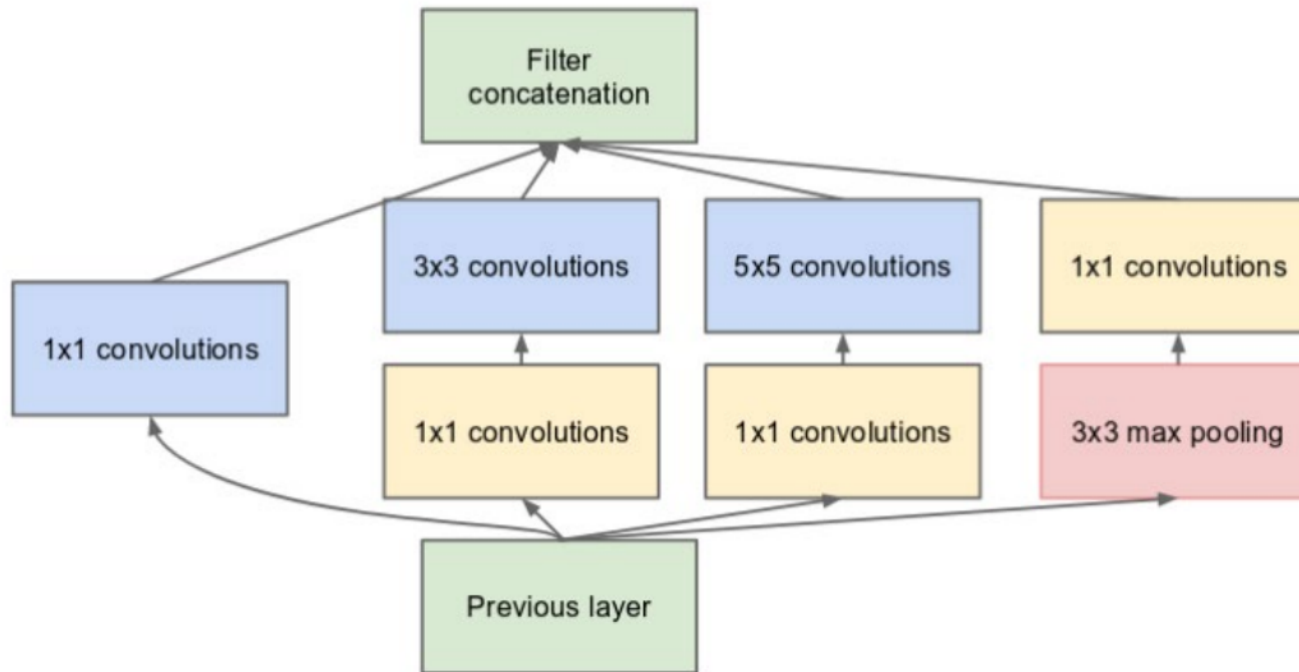


Large kernel for global features, and **smaller kernel** for local features.

Idea: have multiple different-size kernels at any layer.

[Going Deep with Convolutions, Szegedy et al. '14]

Inception Module



Multiple filter scales at each layer

Dimensionality reduction to keep computational requirements down

[Going Deep with Convolutions, Szegedy et al. '14]

Residual Networks

Motivation: extremely deep nets are hard to train (gradient explosion/
vanishing)

decreased learning rate

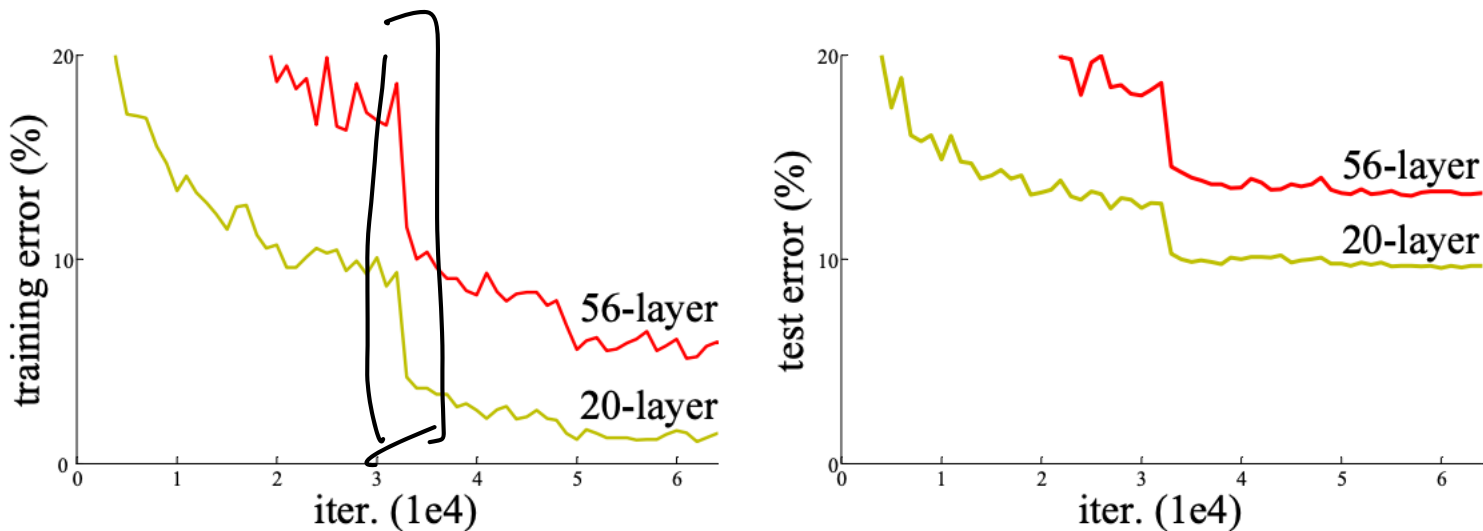


Figure 1. Training error (left) and test error (right) on CIFAR-10 with 20-layer and 56-layer “plain” networks. The deeper network has higher training error, and thus test error. Similar phenomena on ImageNet is presented in Fig. 4.

Residual Networks

$$W_H \sigma(W_{H-1} \dots \underbrace{\sigma(W_1 x) \dots}_{\text{shallow}})$$

deep

Idea: identity shortcut, skip one or more layers.

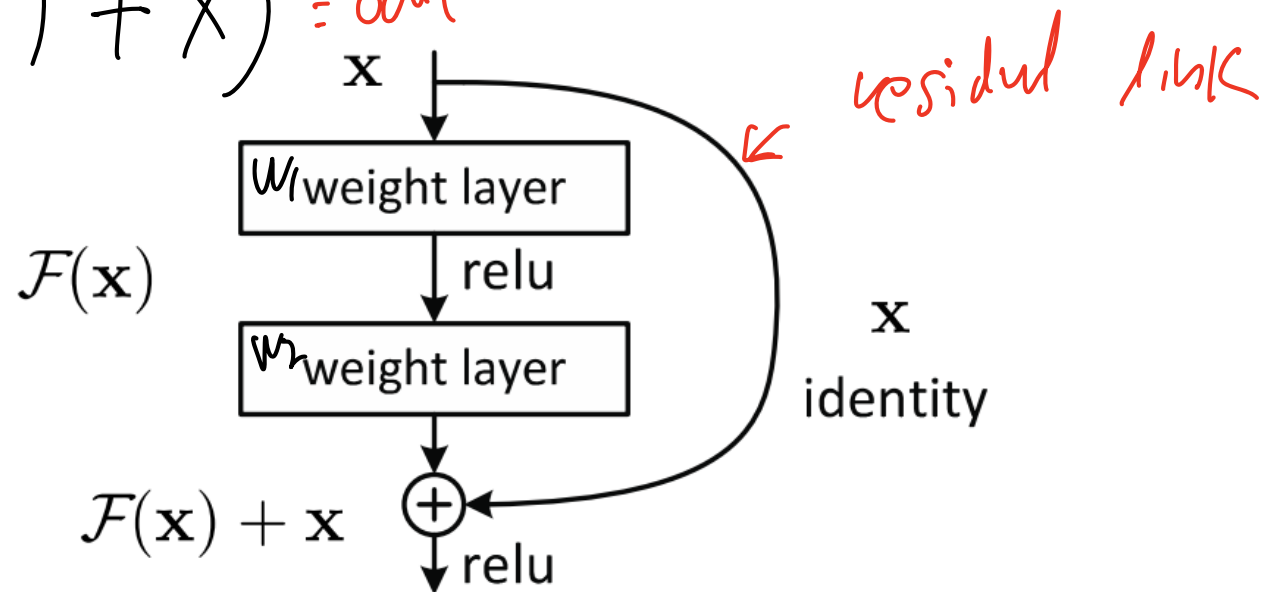
Justification: network can easily simulate shallow network ($F \approx 0$), so performance should not degrade by going deeper.

$$\sigma(W_2 \sigma(W_1 x) + x) = \text{output}$$

$$\sigma(x) = x$$

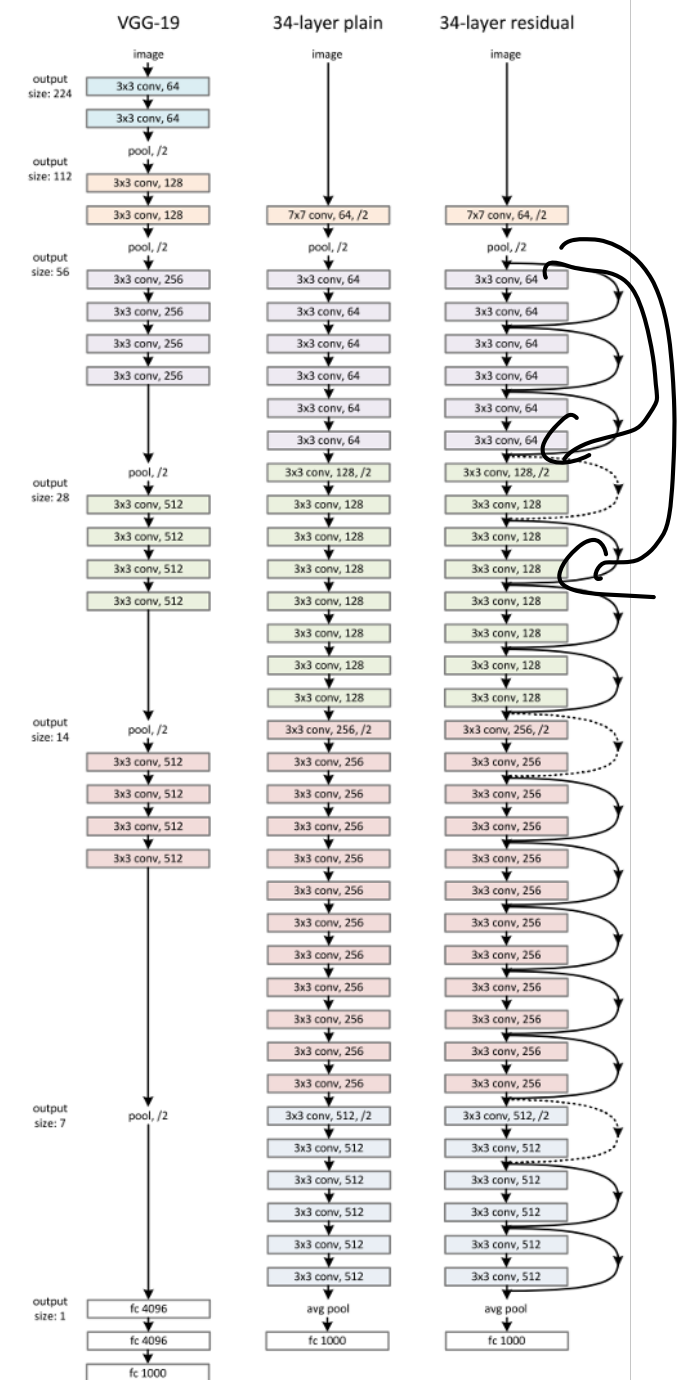
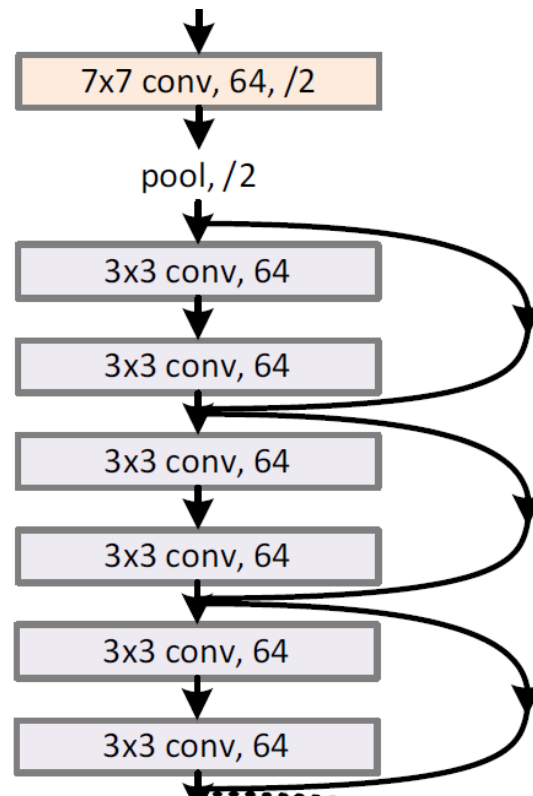
$$W_1, W_2 \approx 0$$

$$\text{output} \approx x$$



Residual Networks

- 3.57% top-5 error on ImageNet
- First deep network with > 100 layers.
- Widely used in many domains (AlphaGo)



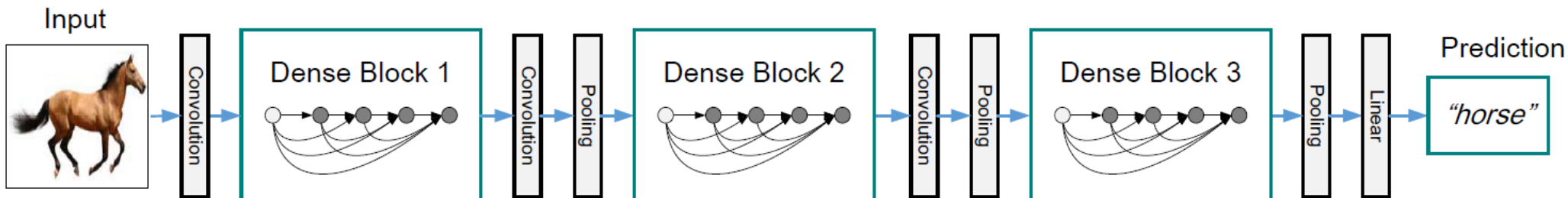
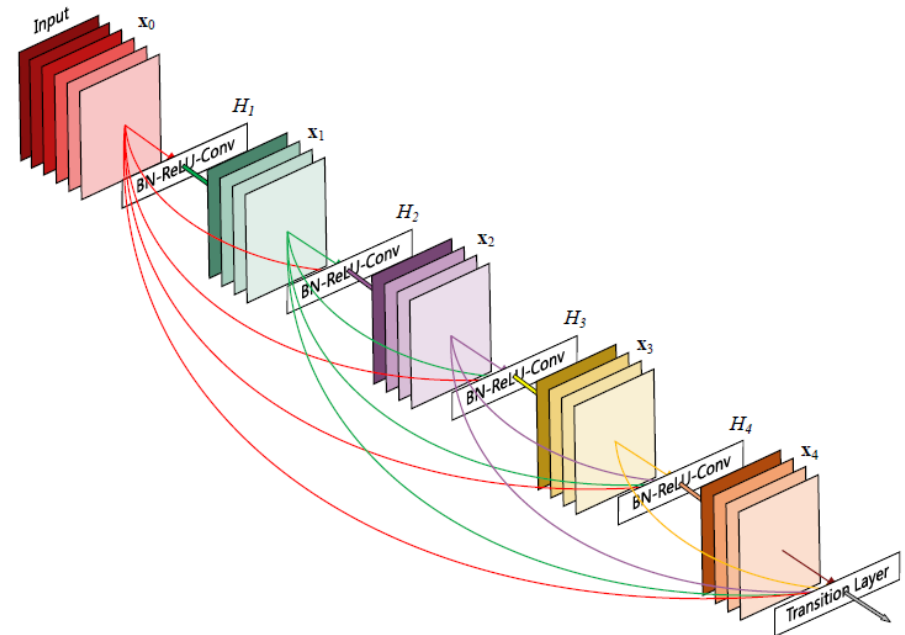
[He, Zhang, Ren, Sun, '16]

Densely Connected Network

Idea: explicit forward output of layer to all future layers (by concatenation)

Intuition: helps vanishing gradients, encourage reuse features (reduce parameter count)

Issues: network maybe too wide, need to be careful about memory consumption



Neural Architecture / Hyper-Parameter Search

Many design choices:

- Number of layers, width, kernel size, pooling, connections, etc.
- Normalization, learning rate, batch size, etc.

Strategies:

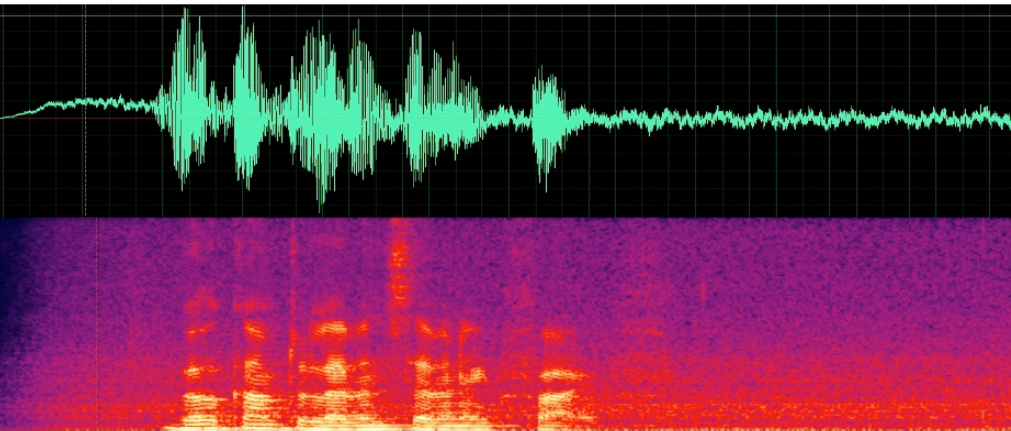
- Grid search
- Random search [Bergstra & Bengio '12]
- Bandit-based [Li et al. '16]
- Gradient-based (DARTS) [Liu et al. '19]
- Neural tangent kernel [Xu et al. '21]
- ...

Scaling law M^p

Recurrent Neural Networks



Sequence Data



检测语言 英语 中文 德语

中文 (简体) 英语 日语

Deep learning is a popular area in AI.



38 / 5000



深度学习是AI的热门领域。



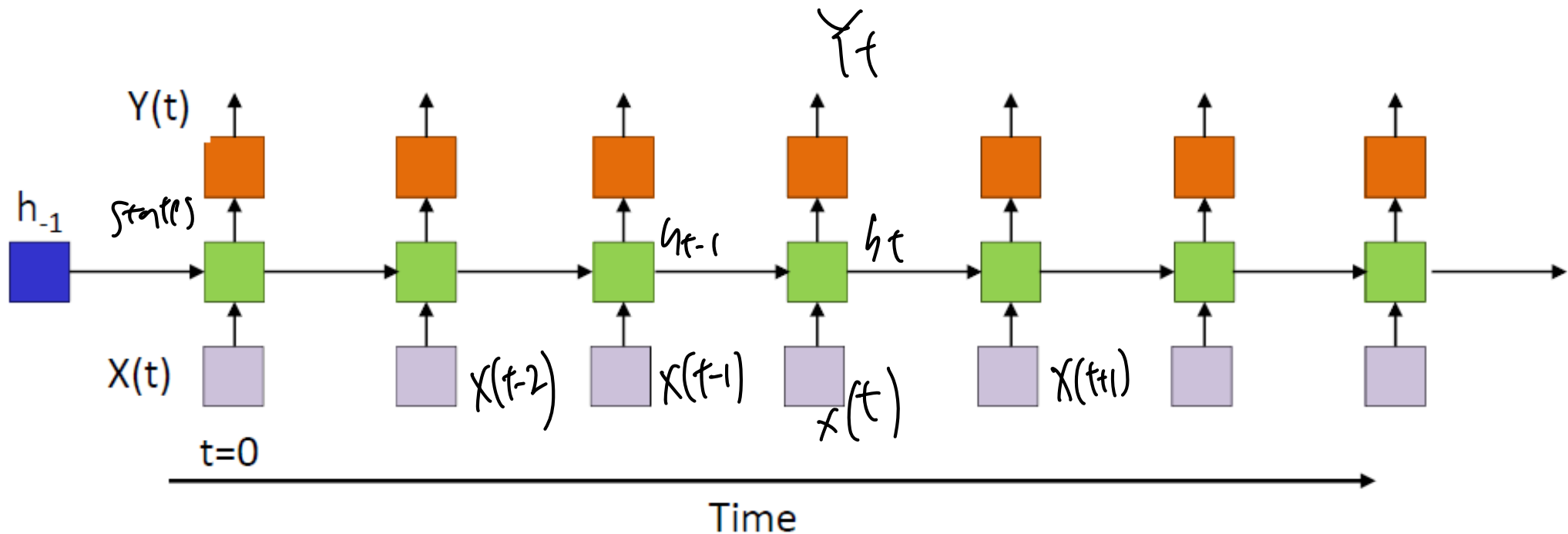
Shēndù xuéxí shì AI de rènmén lǐngyù.



State-Space Model

h_t : information till time t

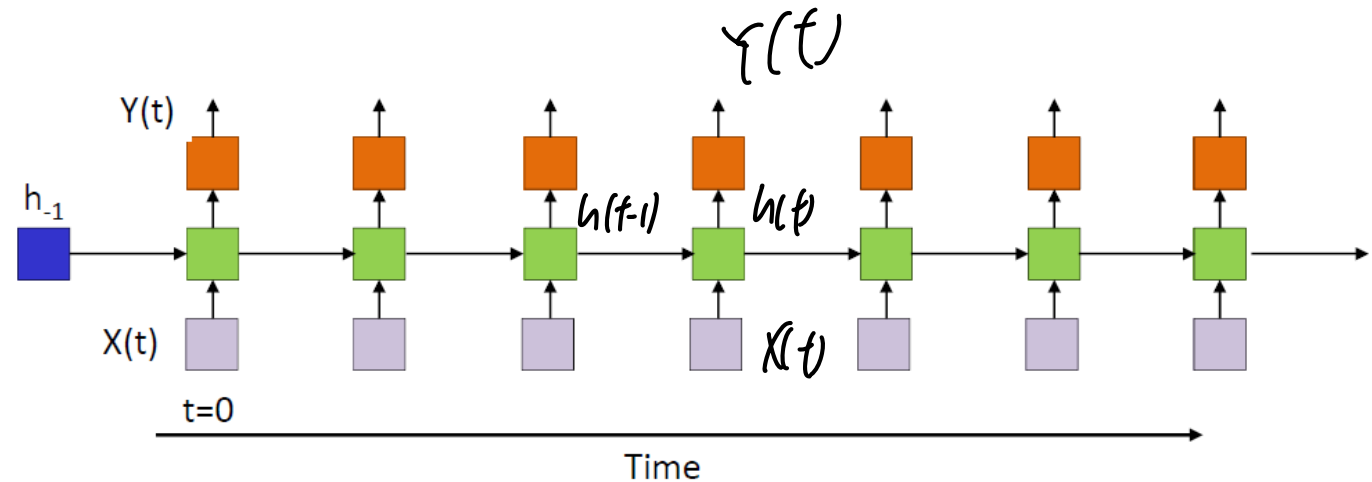
- h_t : hidden state
- X_t : input
- Y_t : output
- $Y_t, h_t = f(h_{t-1}, X_t; \theta)$
- h_{-1} : initial state $\subset \mathcal{O}$



Recurrent Neural Network

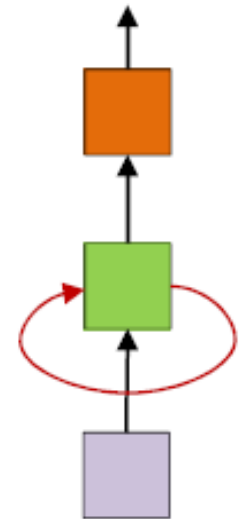
$X \rightarrow Y$

- h_t : hidden state
- X_t : input
- Y_t : output
- $Y_t, h_t = f(h_{t-1}, X_t; \theta)$
- h_{-1} : initial state



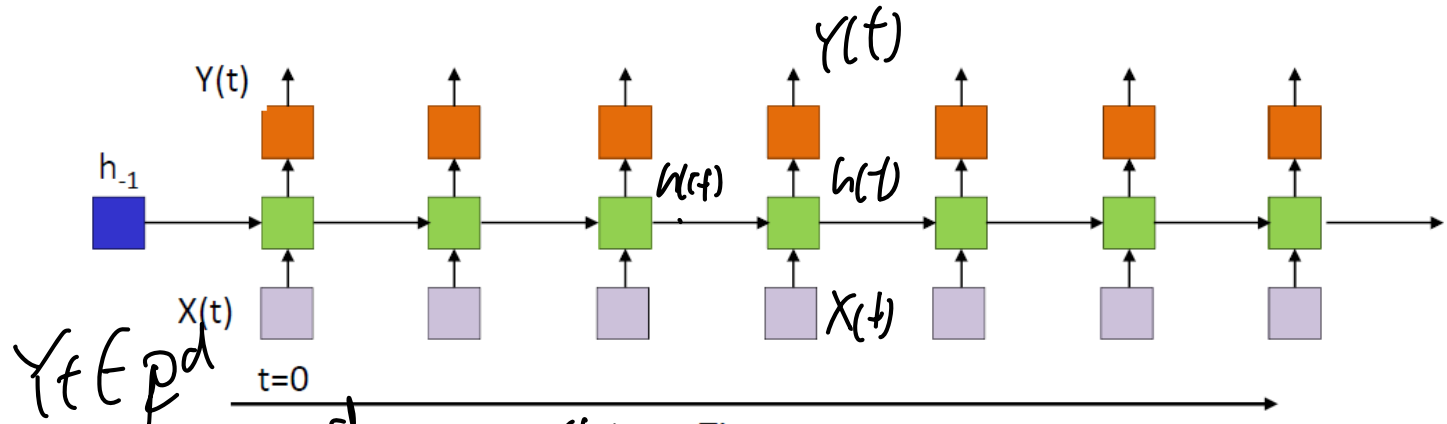
Fully-connect NN vs. RNN

- h_t : a vector summarizes all past inputs (a.k.a. “memory”)
- h_{-1} affects the entire dynamics (typically set to zero)
- X_t affects all the outputs and states after t
- Y_t depends on $X_0, \dots, X_t$



Recurrent Neural Network

- h_t : hidden state
- X_t : input
- Y_t : output
- $Y_t, h_t = f(h_{t-1}, X_t; \theta)$
- h_{-1} : initial state

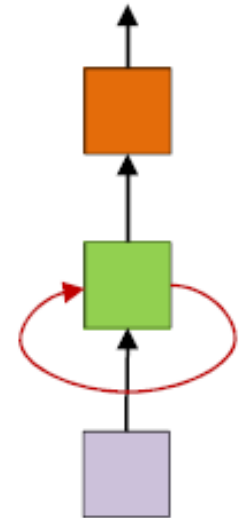


$Y_t \in \mathbb{R}^d$
 $X_t \in \mathbb{R}^d$
 $h_t \in \mathbb{R}^r$
 $b^{(1)} \in \mathbb{R}^r$
 $W^{(1)}: r \times d$
 $W^{(11)}: r \times r$

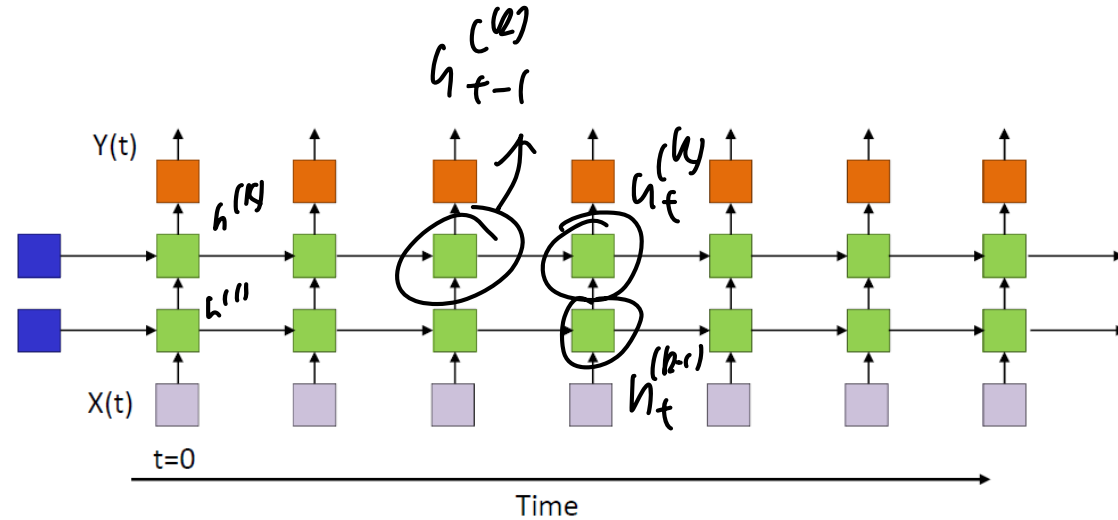
Fully-connect NN vs. RNN

- RNN can be viewed as repeated applying fully-connected NNs
- $h_t = \sigma_1(W^{(1)}X_t + W^{(11)}h_{t-1} + b^{(1)})$
- $Y_t = \sigma_2(W^{(2)}h_t + b^{(2)})$
- σ_1, σ_2 are activation functions (sigmoid, ReLU, tanh, etc)

$W^{(2)}: d \times r$
 $b^{(2)}: d$

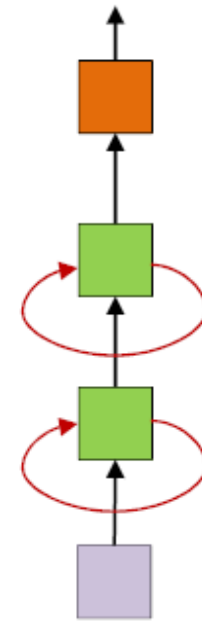


Recurrent Neural Network



Stack K layers of fully-connected NN

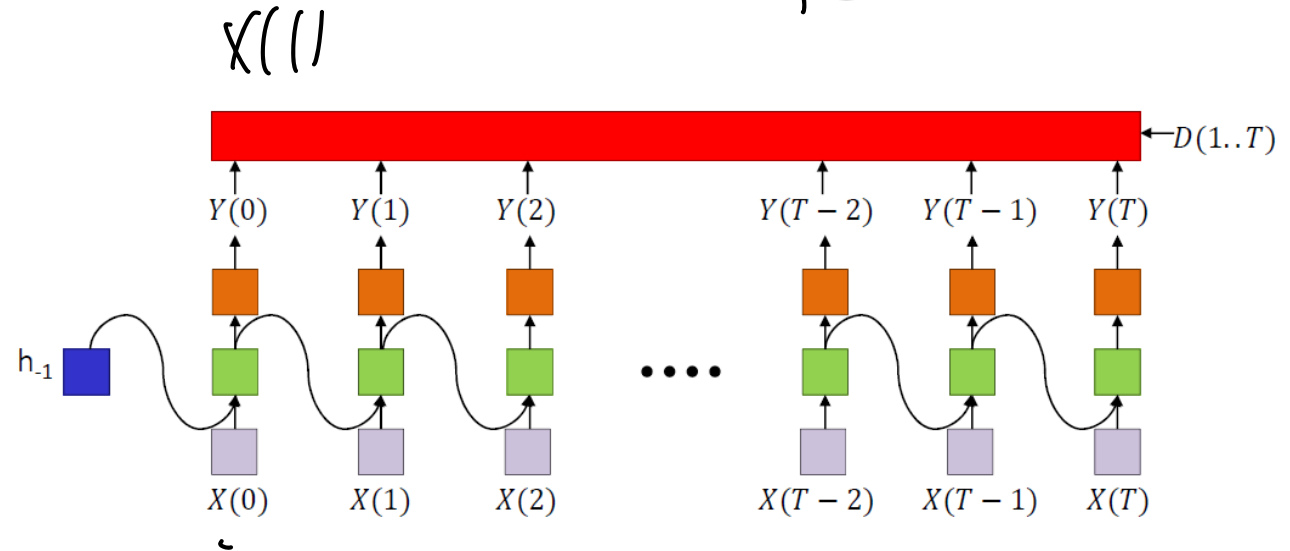
- $h_t^{(k)}$: hidden state
- X_t : input
- Y_t : output
- $h_t^{(1)} = f_1^{(1)}(h_{t-1}^{(1)}, X_t; \theta)$
- $h_t^{(k)} = f_1^{(k)}(h_{t-1}^{(k)}, h_t^{(k-1)}; \theta)$
- $Y_t = f_2(h_t^{(K)}; \theta)$
- $h_{-1}^{(k)}$: initial states



Training Recurrent Neural Network

$$\sum_{t=1}^T (Y(t) - X(t))^2$$

- h_t : hidden state
- X_t : input
- Y_t : output
- $Y_t, h_t = f(h_{t-1}, X_t; \theta)$
- h_{-1} : initial state

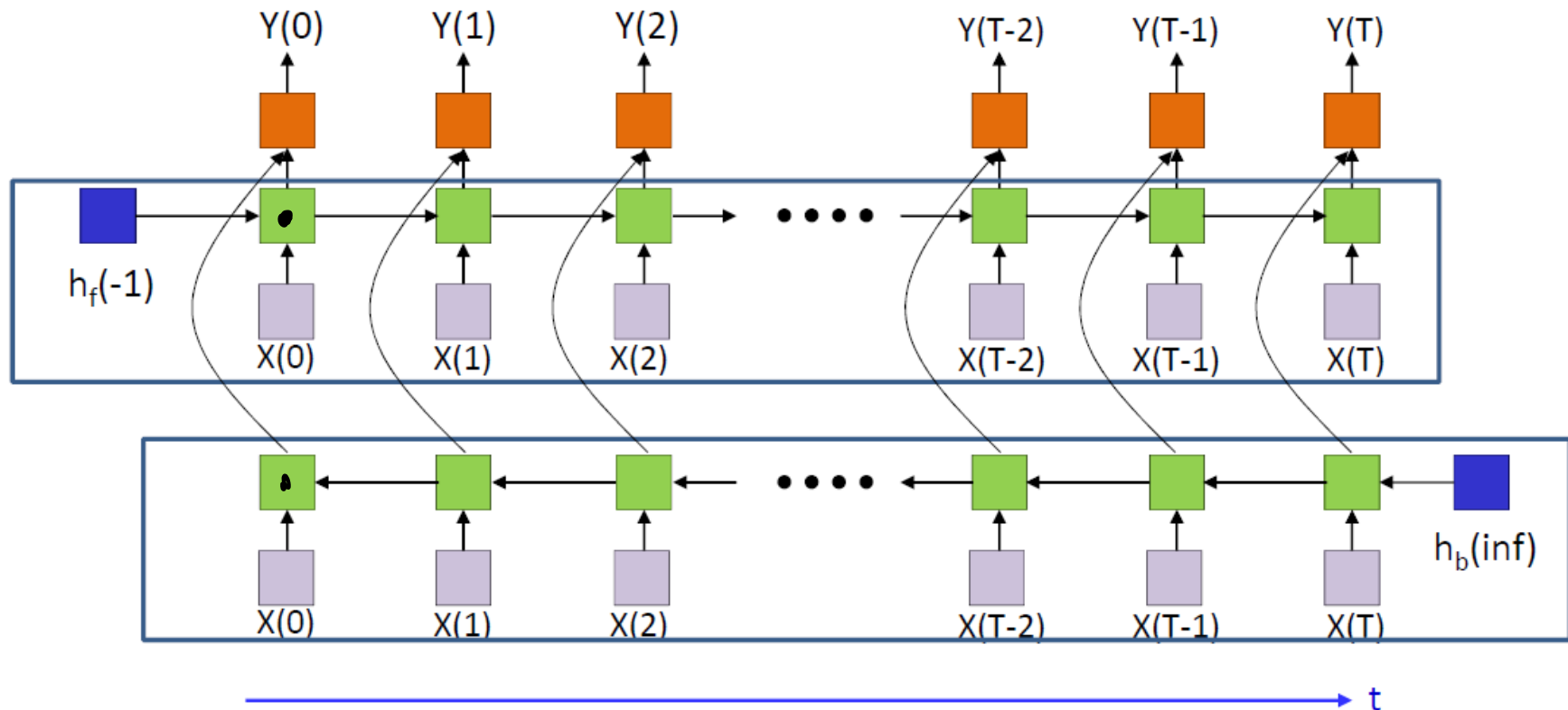


- Data: $\{(X_t, D_t)\}_{t=1}^T$ (RNN can handle more general data format)
- Loss $L(\theta) = \sum_{t=1}^T \ell(Y_t, D_t)$
- Goal: learn θ by gradient-based method
 - Back propagation

Extensions

What if Y_t depends on the entire inputs?

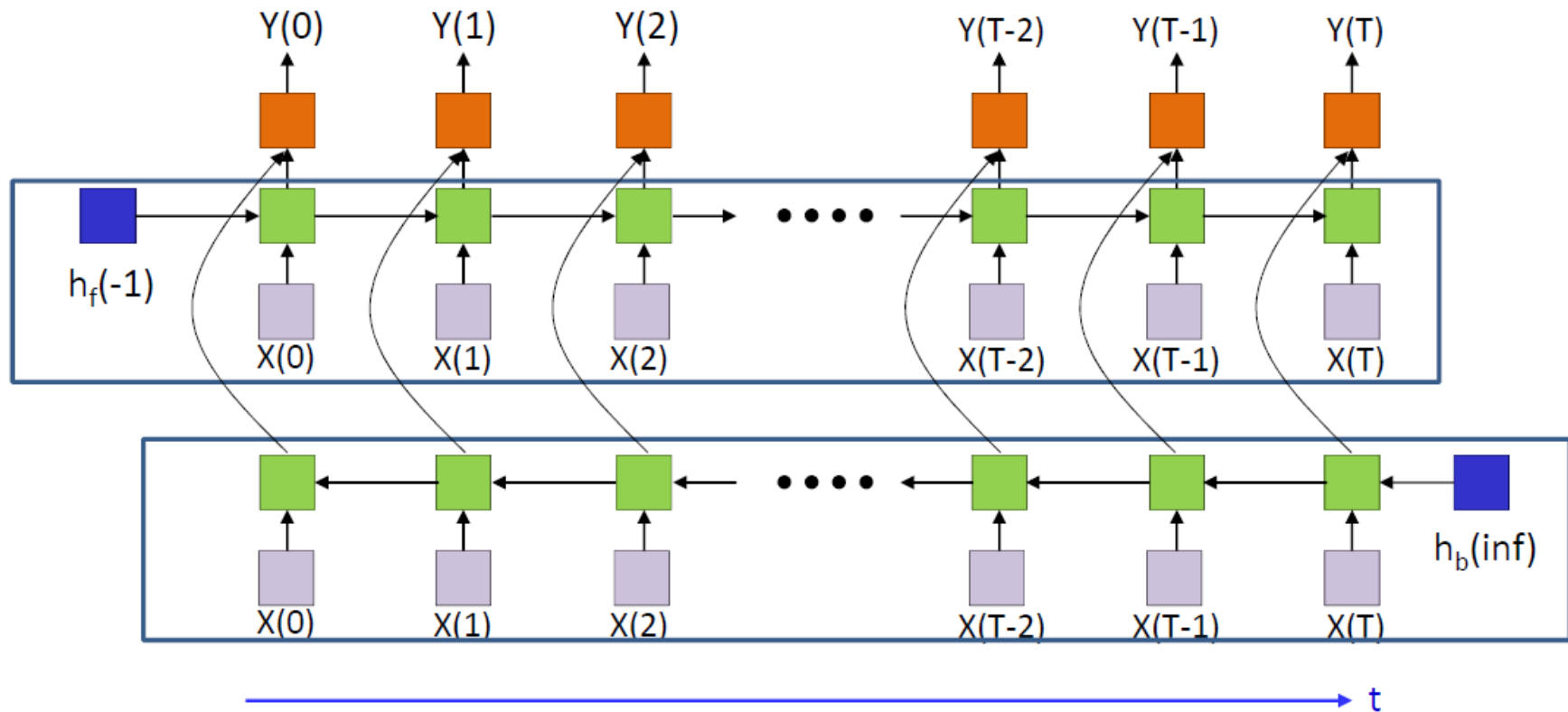
- Biredirectional RNN:
 - AN RNN for forward dependencies: $t = 0, \dots, T$
 - An RNN for backward dependencies: $t = T, \dots, 0$
 - $Y_t = f_2(h_t^f, h_t^b; \theta)$



Extensions

RNN for sequence classification (sentiment analysis)

- $Y = \max_t Y_t$
- Cross-entropy loss



Practical issues of RNN $\delta(z) = z$

Linear RNN derivation

$$\begin{aligned} h_t &= W^{(1)} h_{t-1} + W^{(1)} x_t \\ h_k &= W^{(1)} x_k + W^{(1)} h_{k-1} \\ &= W^{(1)} x_k + W^{(1)} (W^{(1)} x_{k-1} + W^{(1)} h_{k-2}) \\ &= \underbrace{(W^{(1)})^{k+1}} h_{-1} + \sum_{i=0}^k \underbrace{(W^{(1)})^{k-i}} W^{(1)} x_i \end{aligned}$$

if

$\lambda_{\max}(W^{(1)}) > 1 \rightarrow h_k \text{ exp large}$

$\lambda_{\max}(W^{(1)}) < 1 \rightarrow h_k \text{ exp small}$
weak dependency

forgetting

on h_{-1}

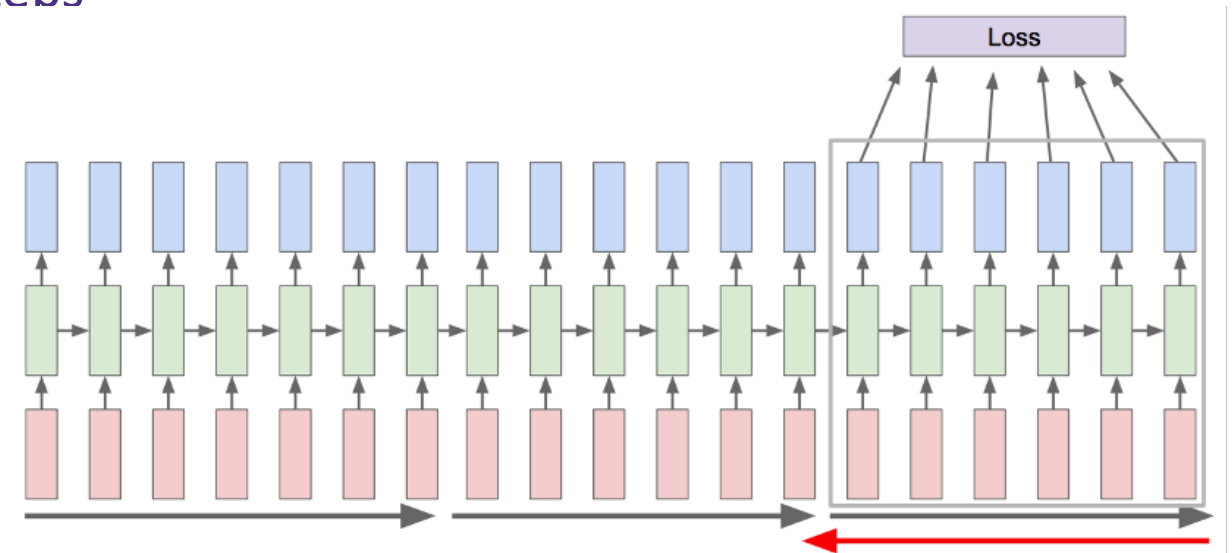
Practical issues of RNN: training

Gradient explosion and gradient vanishing

$$\text{gradient} \propto (w^{(1)})^L \prod_{t=0}^L \delta'(h_t)$$

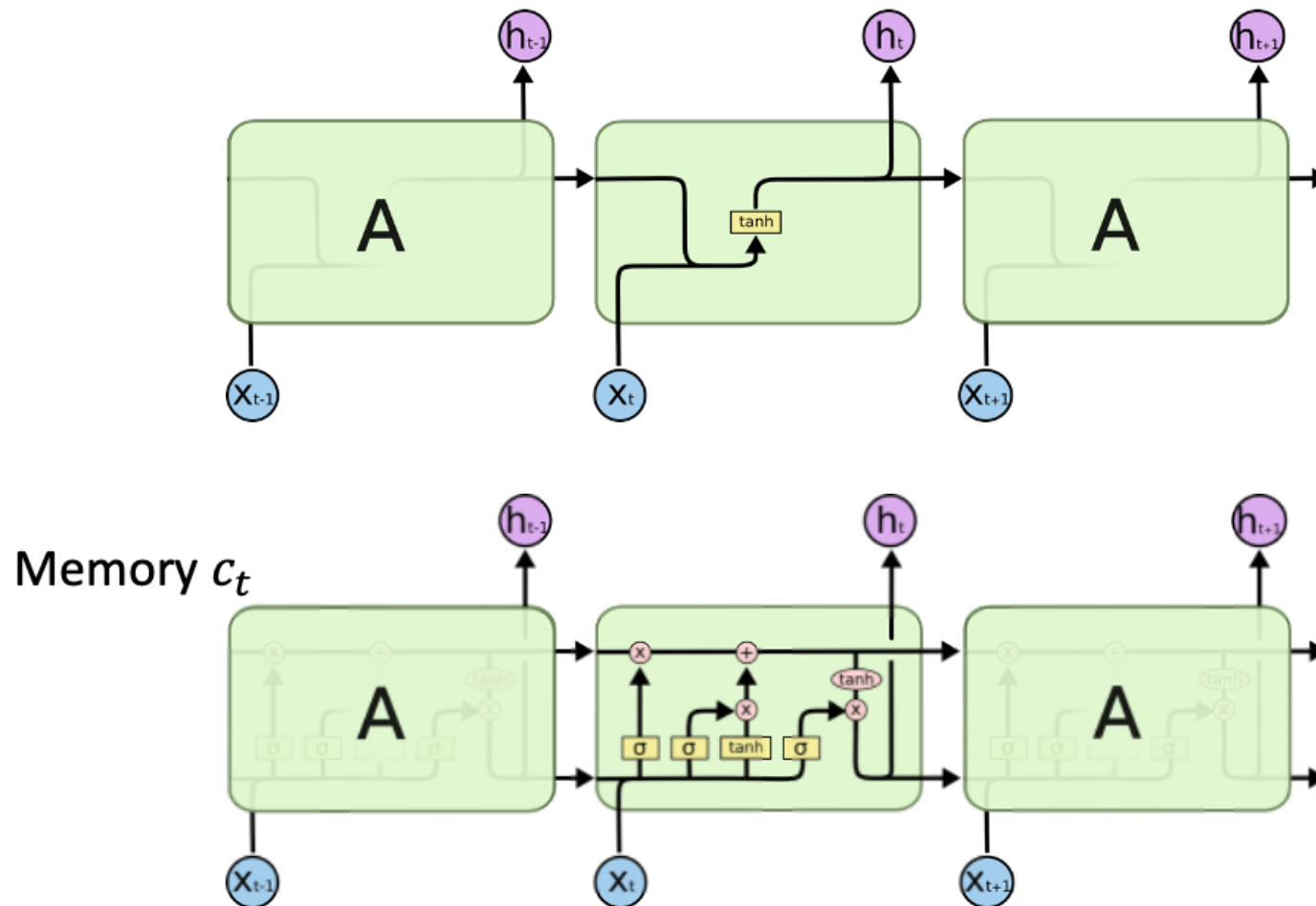
Techniques for avoiding gradient explosion

- Gradient clipping
- Identity initialization
- Truncated backprop through time
 - Only backprop for a few steps



Preserve Long-Term Memory

- Difficult for RNN to preserve long-term memory
 - The hidden state h_t is constantly being written (short-term memory)
 - Use a separate cell to maintain long-term memory



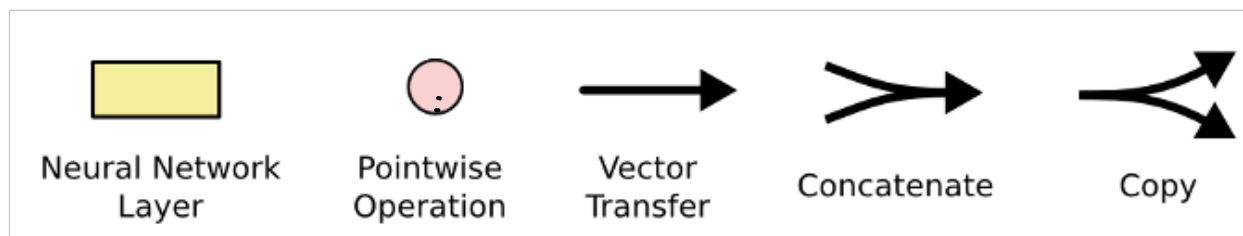
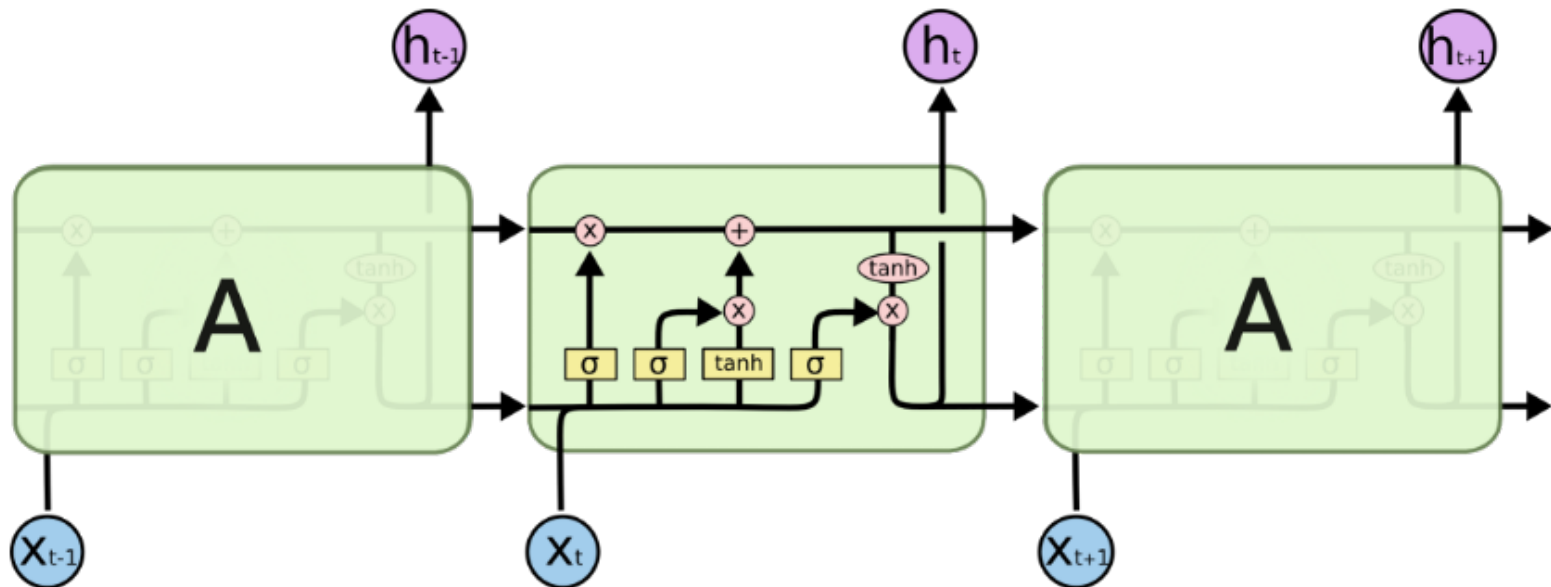
Long Short-Term Memory Network

LSTM (Hochreiter & Schmidhuber, '97)

- RNN architecture for learning long-term dependencies
- σ : layer with sigmoid activation

gating:
pointwise product
 $a, b \text{ vectors} \in \mathbb{R}^n$

$$(a) \cdot (b) = \begin{pmatrix} a_1 b_1 \\ \vdots \\ a_n b_n \end{pmatrix}$$



Long Short-Term Memory Network

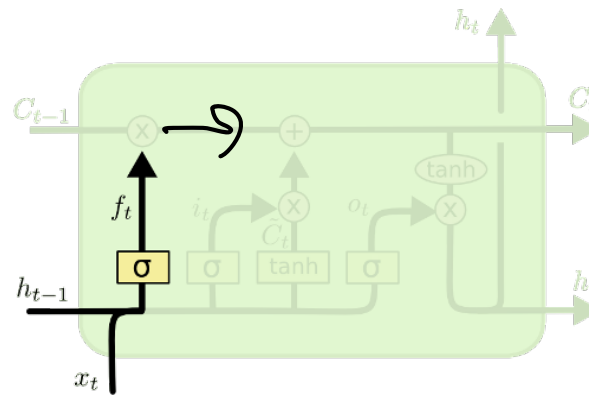
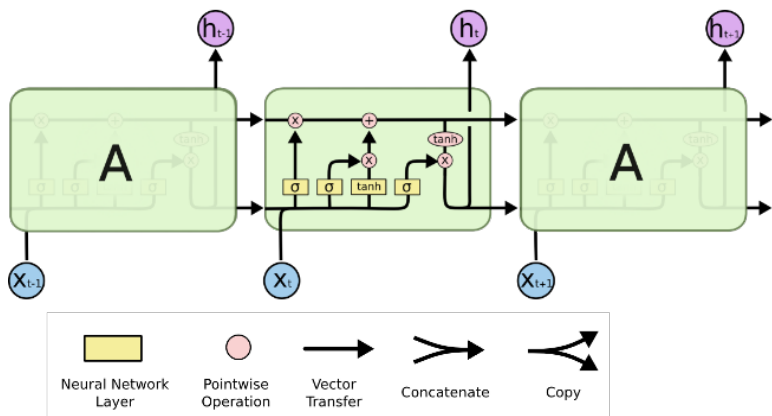
$$c_t \in \mathbb{R}^d$$

$$f_t(i) \in [0, 1]$$

Forget gate $f_t \in \mathbb{R}^d$

- f_t outputs whether we want to “forget” things in c_t
 - Compute $c_{t-1} \odot f_t$ (element-wise)
 - $f_t(i) \rightarrow 0$: want to forget $c_t(i)$
 - $f_t(i) \rightarrow 1$: we want to keep the information in $c_t(i)$

$$\begin{pmatrix} c_{t-1}(1) \cdot f_t(1) \\ \vdots \\ c_{t-1}(i) \cdot f_t(i) \\ \vdots \end{pmatrix}$$



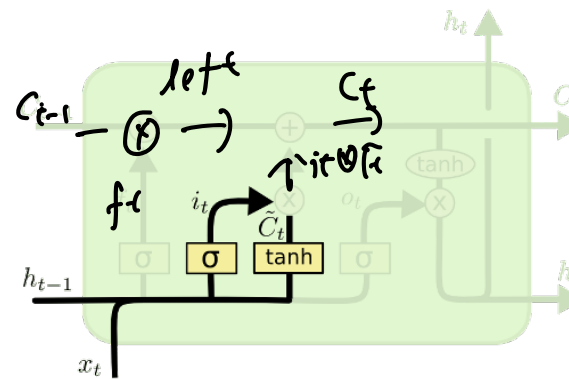
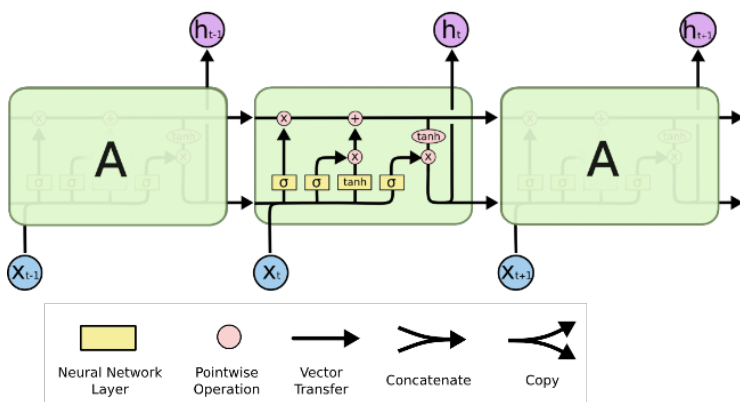
forget-invariant

$$f_t = \sigma(W_f \cdot [h_{t-1}, x_t] + b_f)$$

Long Short-Term Memory Network

Input gate i_t

- i_t extracts useful information from X_t to update memory
 - $\tilde{C}_t$: information from X_t to update memory
 - i_t : which dimension in the memory should be updated by X_t
 - $i_t(j) \rightarrow 1$: we want to use the information in $\tilde{C}_t(j)$ to update memory
 - $i_t(t) \rightarrow 0$: $\tilde{C}_t(j)$ should not contribute to memory



$$i_t = \sigma(W_i \cdot [h_{t-1}, x_t] + b_i)$$

$$\tilde{C}_t = \tanh(W_C \cdot [h_{t-1}, x_t] + b_C)$$

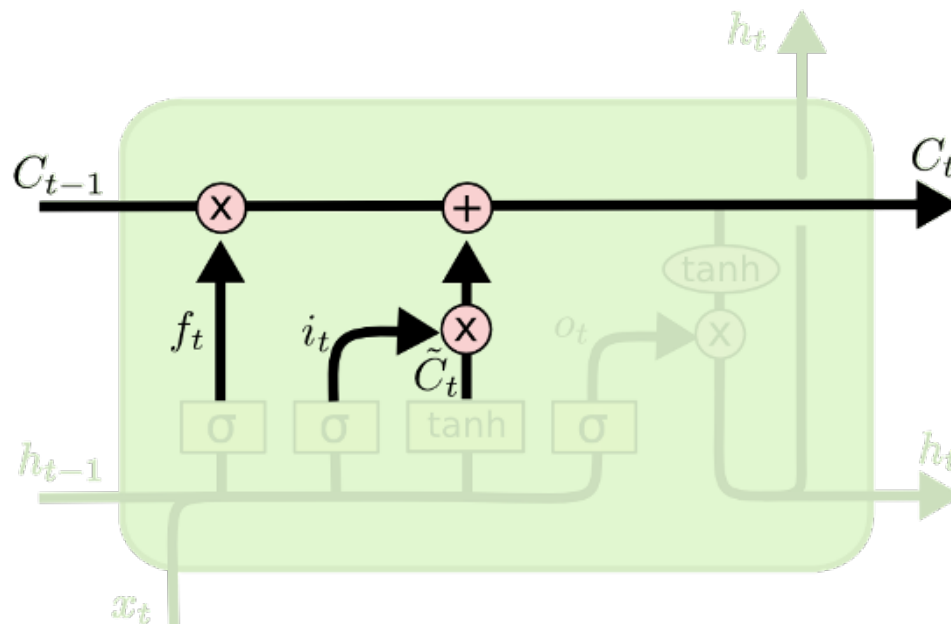
$$o_t = \sigma(W_o \cdot [h_{t-1}, x_t] + b_o)$$

$$h_t = C_t \odot \tanh(W_o \cdot [h_{t-1}, x_t] + b_o)$$

Long Short-Term Memory Network

Memory update

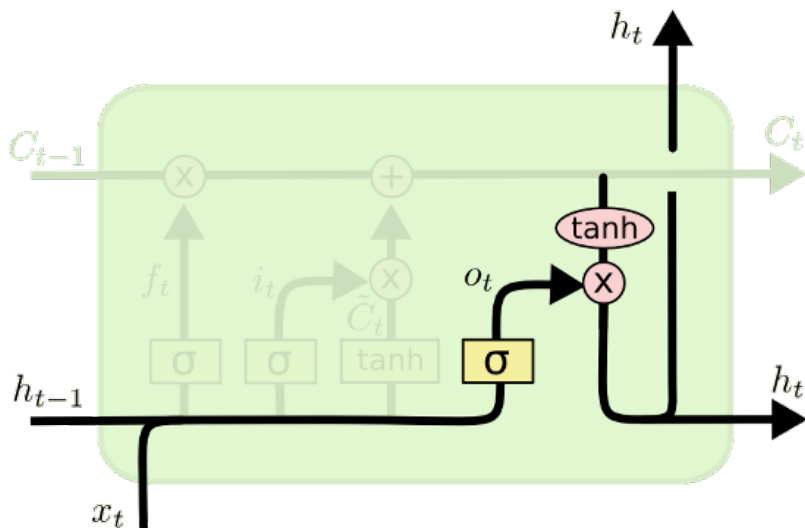
- $c_t = f_t \odot c_{t-1} + i_t \odot \tilde{c}_t$
- f_t forget gate; i_t input gate
- $f_t \odot c_{t-1}$: drop useless information in old memory
- $i_t \odot \tilde{c}_t$: add selected new information from current input



Long Short-Term Memory Network

Output gate $o_t \in \mathbb{R}^d$ $h_t \in \mathbb{R}^d$

- Next hidden state $h_t = o_t \odot \tanh(c_t)$
 - $\tanh(c_t)$: non-linear transformation over all past information
 - o_t : choose important dimensions for the next state
 - $o_t(j) \rightarrow 1$: $\tanh(c_t(j))$ is important for the next state
 - $o_t(j) \rightarrow 0$: $\tanh(c_t(j))$ is not important



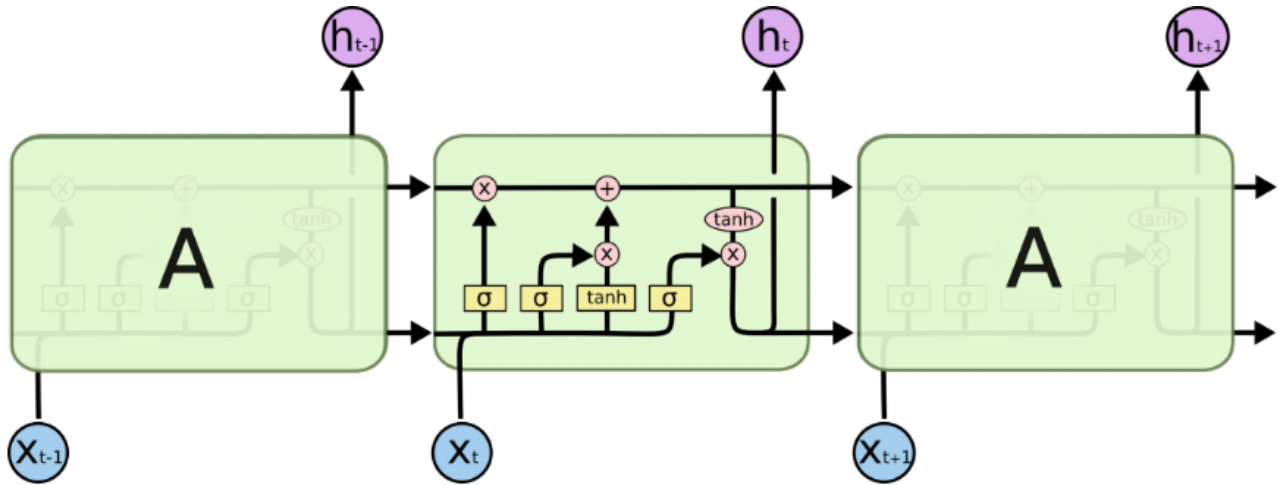
$$o_t = \sigma(W_o [h_{t-1}, x_t] + b_o)$$

$$h_t = o_t \odot \tanh(C_t)$$

Long Short-Term Memory Network

- $h_t = o_t \odot \tanh(c_t)$
- $c_t = f_t \odot c_{t-1} + i_t \odot \tilde{c}_t$
- $Y_t = g(h_t)$

gates: time-dependent



Remarks:

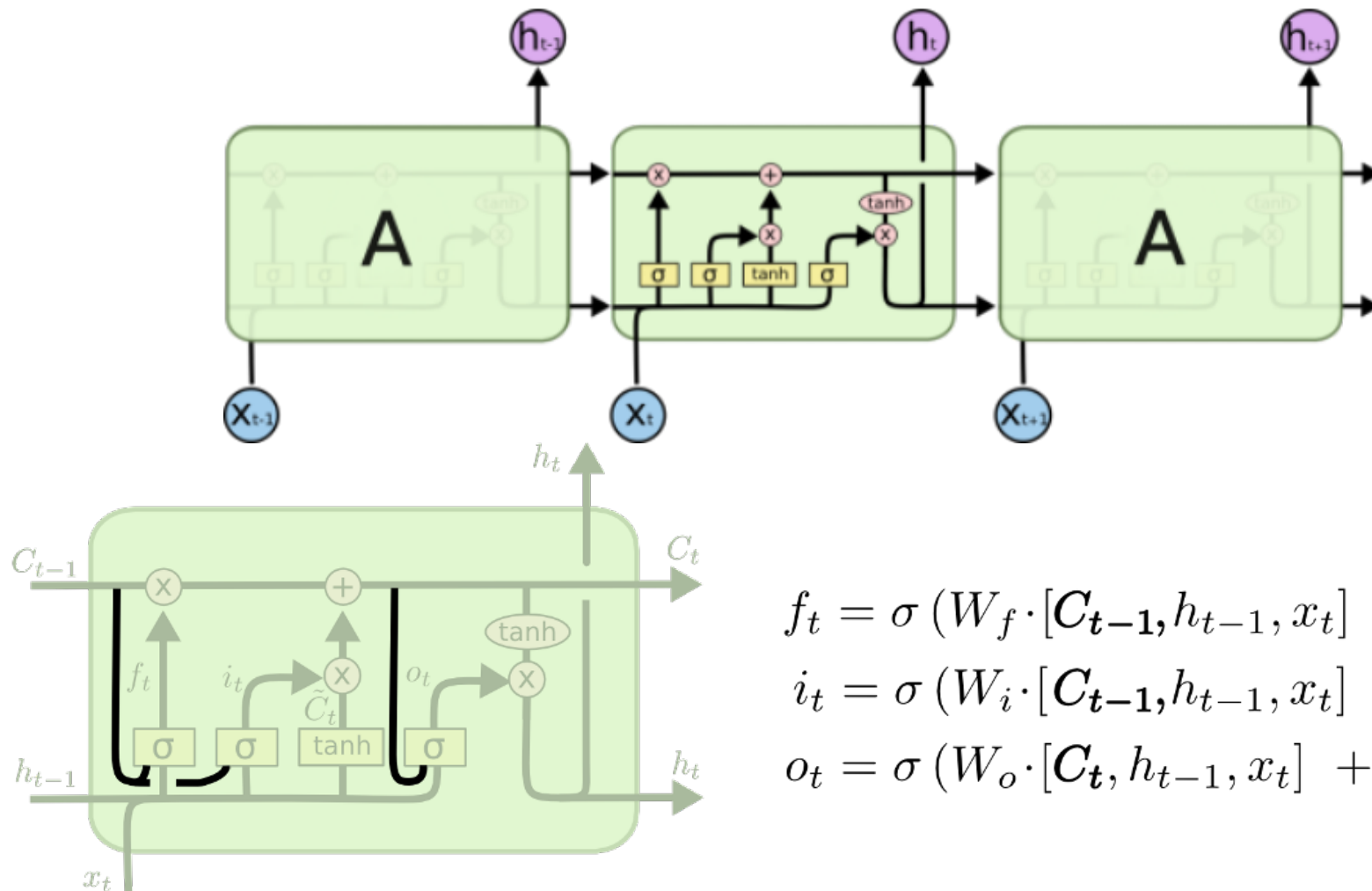
1. No more matrix multiplications for c_t
2. LSTM does not have guarantees for gradient explosion/vanishing
3. LSTM is the dominant architecture for sequence modeling from '13 - '16.
4. Why tanh

instead $W^{(11)} h_t$

LSTM Variant

Peephole Connections (Gers & Schmidhuber '00)

- Allow gates to take in c_t information



$$f_t = \sigma(W_f \cdot [C_{t-1}, h_{t-1}, x_t] + b_f)$$

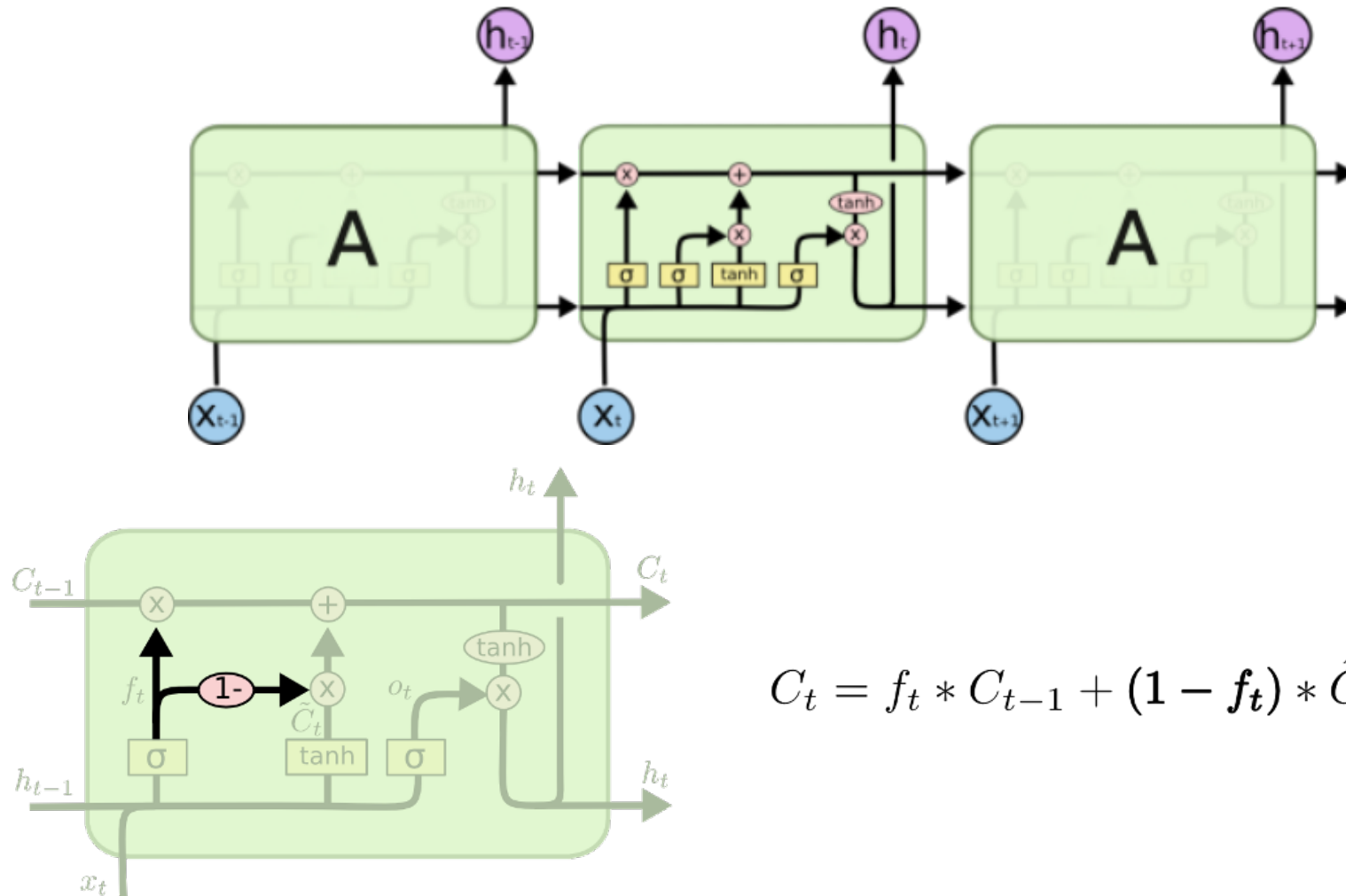
$$i_t = \sigma(W_i \cdot [C_{t-1}, h_{t-1}, x_t] + b_i)$$

$$o_t = \sigma(W_o \cdot [C_t, h_{t-1}, x_t] + b_o)$$

LSTM Variant

Simplified LSTM

- Assume $i_t = 1 - f_t$
- Only two gates are needed: fewer parameters

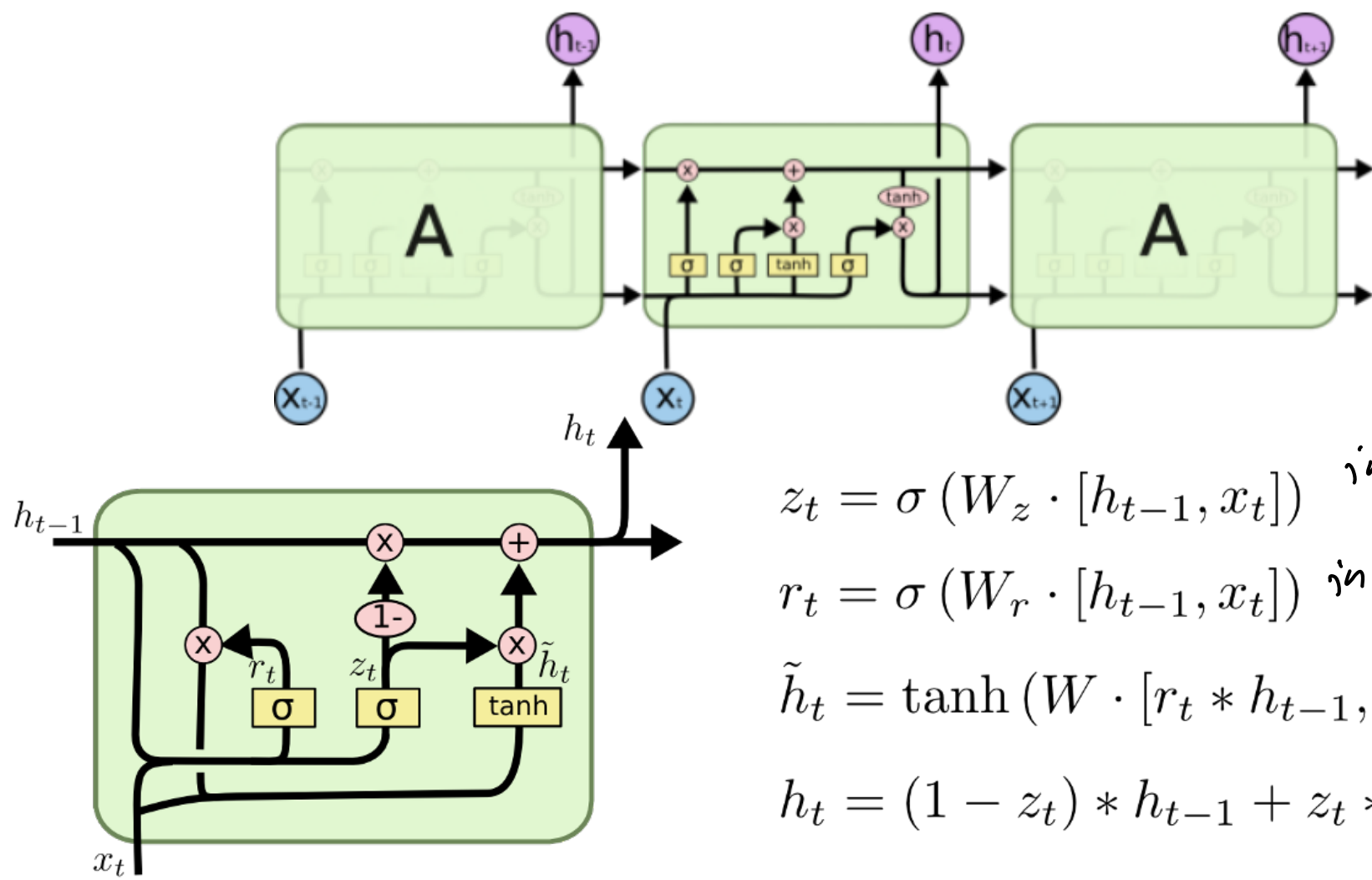


$$C_t = f_t * C_{t-1} + (1 - f_t) * \tilde{C}_t$$

LSTM Variant

Gated Recurrent Unit (GRU, Cho et al. '14)

- Merge h_t and c_t : much fewer parameters



$$z_t = \sigma(W_z \cdot [h_{t-1}, x_t]) \quad \text{input 1}$$

$$r_t = \sigma(W_r \cdot [h_{t-1}, x_t]) \quad \text{input 2}$$

$$\tilde{h}_t = \tanh(W \cdot [r_t * h_{t-1}, x_t])$$

$$h_t = (1 - z_t) * h_{t-1} + z_t * \tilde{h}_t$$

$$\underline{(x_1, \dots, x_n)}$$

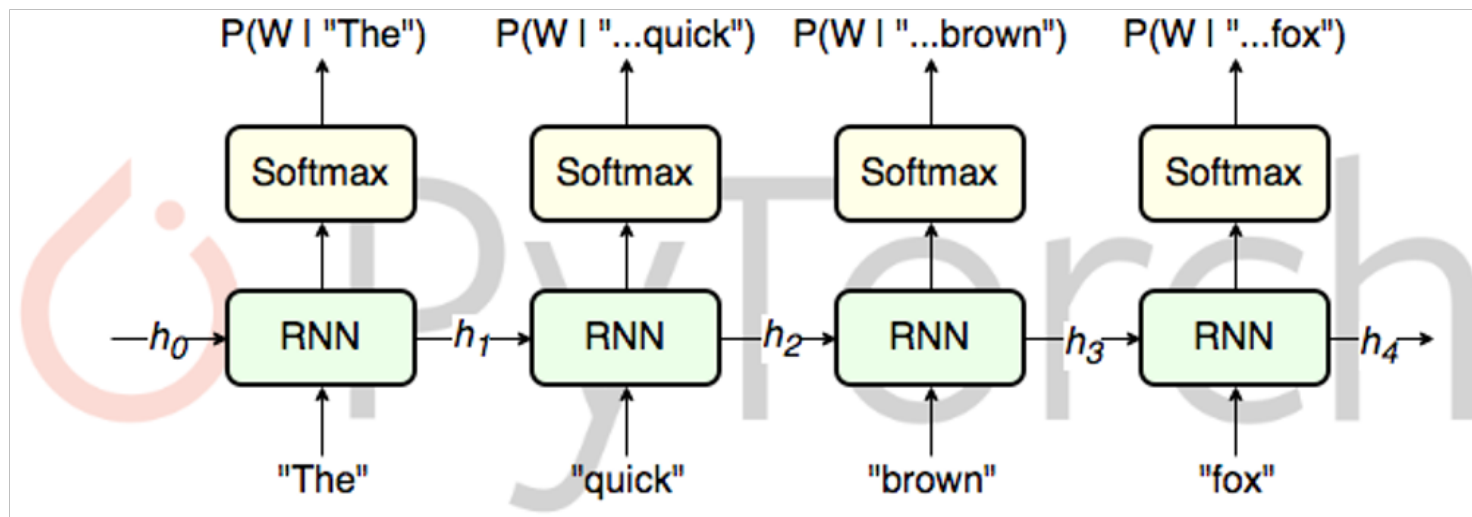
- X : a sentence
- Sequential generation

- X_t : word at position t .
- Y_t : softmax over all words

- Data: a collection of texts:
 - Wiki

$$P(X_1 | X_0; \theta) \cdot P(X_2 | X_0, X_1; \theta)$$

(Query) $[]$
 t arr t+1
 [[1 ...]]
 # of freq word



LSTM application: text classification

Bi-directional LSTM and then run softmax on the final hidden state.

