



CS 639: Foundation Models **Deep Learning I**

Fred Sala

University of Wisconsin-Madison

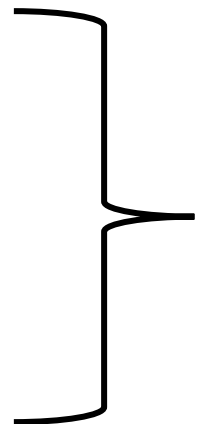
Jan. 27, 2026



Announcements

- **Midterm: Weds. March 11th**
 - More details as we get closer
- **Resources**
 - <https://mlstory.org/> : fun book by Hardt and Recht
- **Class roadmap:**

Tuesday Jan. 27	Deep Learning I
Thursday Jan. 20	Deep Learning II
Tuesday Feb. 3	Self-Supervised Learning
Thursday Feb. 5	Guest Lecture



Review + Start FMs

Outline

- **Neural Networks**

- Perceptrons, setup, training, limitations, multilayer perceptrons, loss functions (mostly from last time)

- **Convolutional Neural Networks**

- Motivation, convolutional layers, CNN architectures

- **Sequence Models**

- Recurrent neural networks, architecture, LSTMs, alternatives, training tricks

Outline

- **Neural Networks**

- Perceptrons, setup, training, limitations, multilayer perceptrons, loss functions (mostly from last time)

- **Convolutional Neural Networks**

- Motivation, convolutional layers, CNN architectures

- **Sequence Models**

- Recurrent neural networks, architecture, LSTMs, alternatives, training tricks

From Last Time: Perceptron

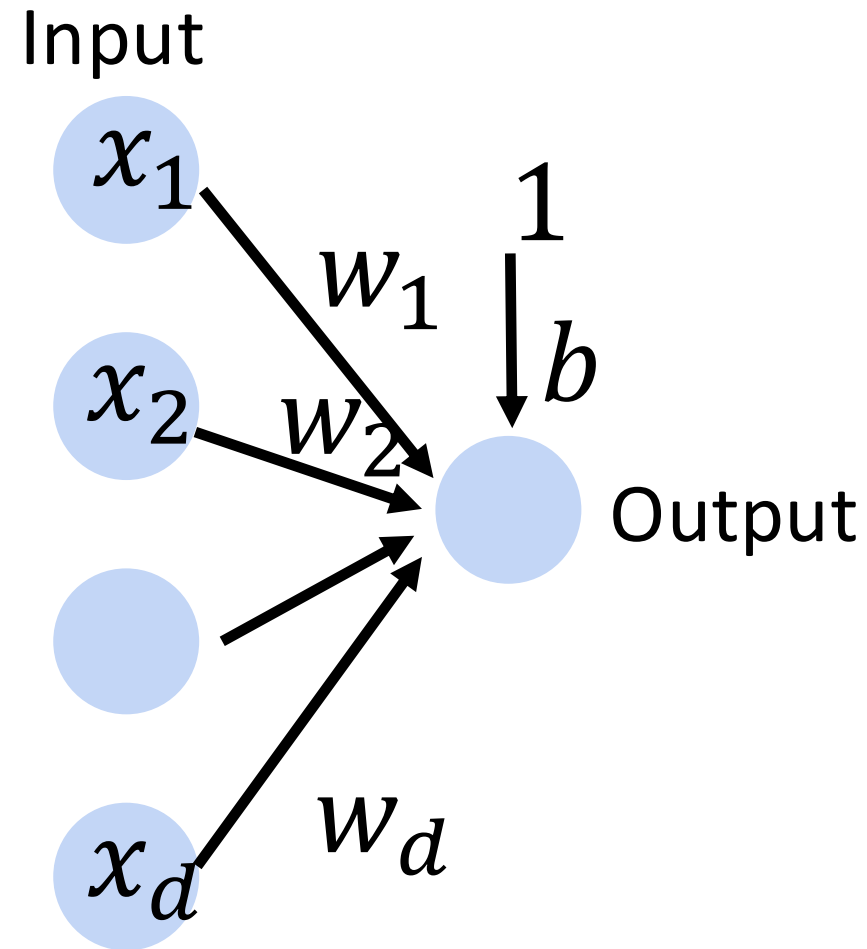
- Given input $\mathbf{x}$, weight $\mathbf{w}$ and bias b , perceptron outputs:

$$o = \sigma(\langle \mathbf{w}, \mathbf{x} \rangle + b)$$

$$\sigma(x) = \begin{cases} 1 & \text{if } x > 0 \\ 0 & \text{otherwise} \end{cases}$$

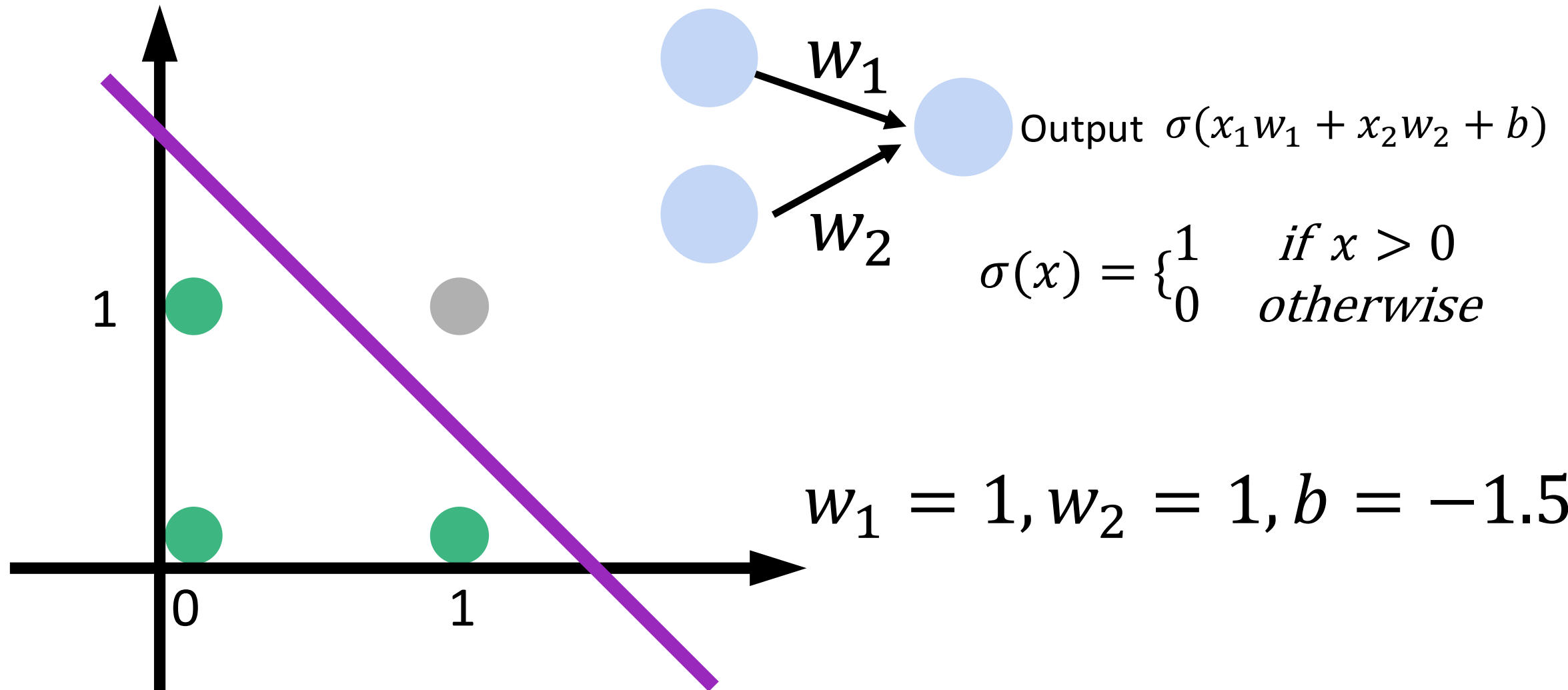
Activation function

Cats vs. dogs?



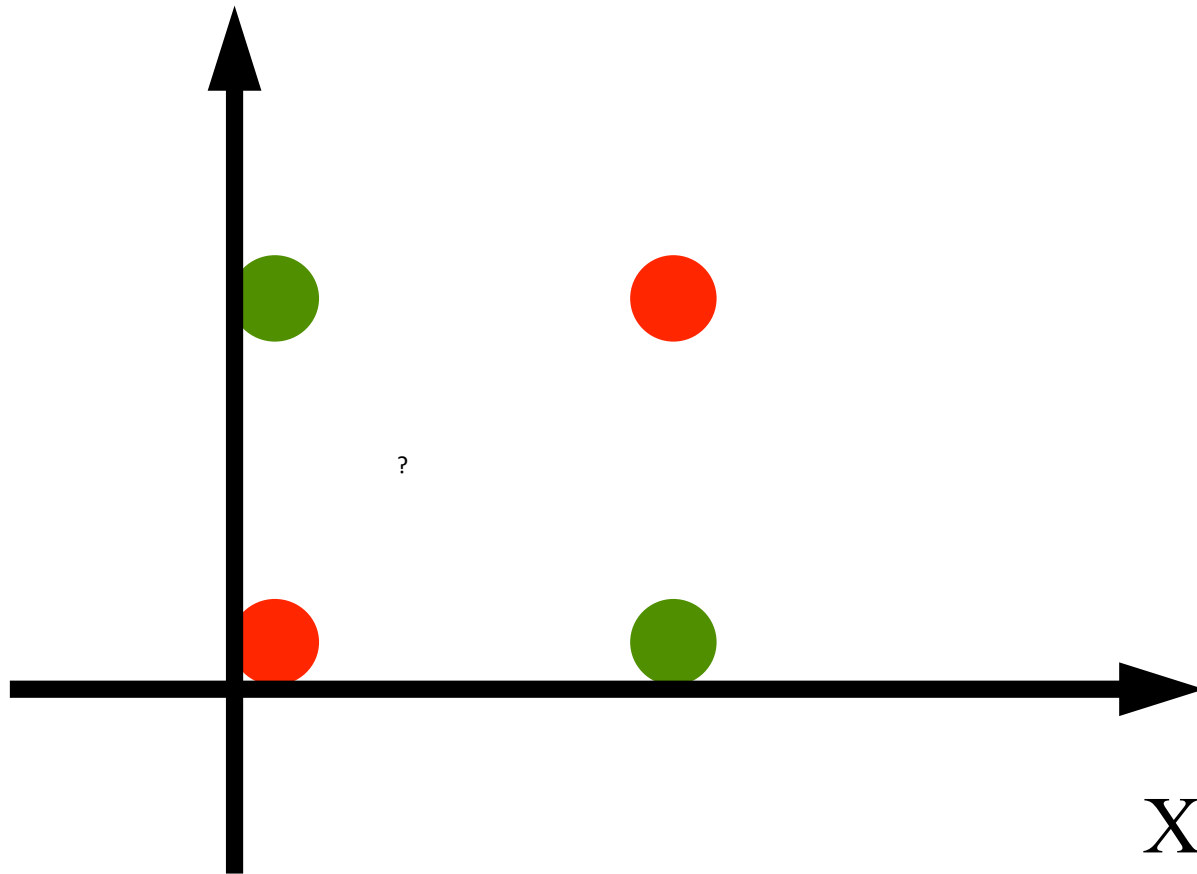
Learning **AND** function using perceptron

The perceptron can learn an AND function



The limited power of a single neuron

The perceptron cannot learn an **XOR** function
(neurons can only generate linear separators)



$$x_1 = 1, x_2 = 1, y = 0$$

$$x_1 = 1, x_2 = 0, y = 1$$

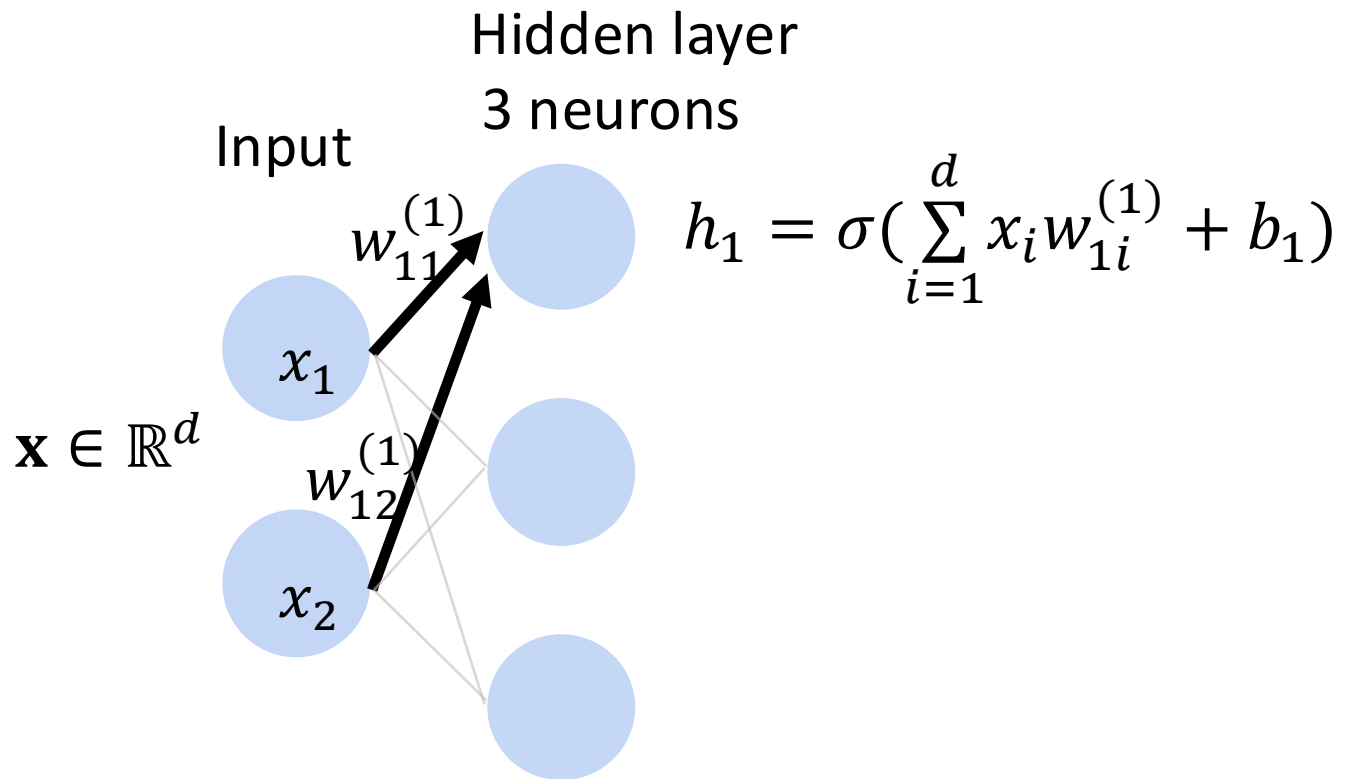
$$x_1 = 0, x_2 = 1, y = 1$$

$$x_1 = 0, x_2 = 0, y = 0$$

$$\text{XOR}(x_1, x_2) = (x_1 \wedge \neg x_2) \vee (\neg x_1 \wedge x_2)$$

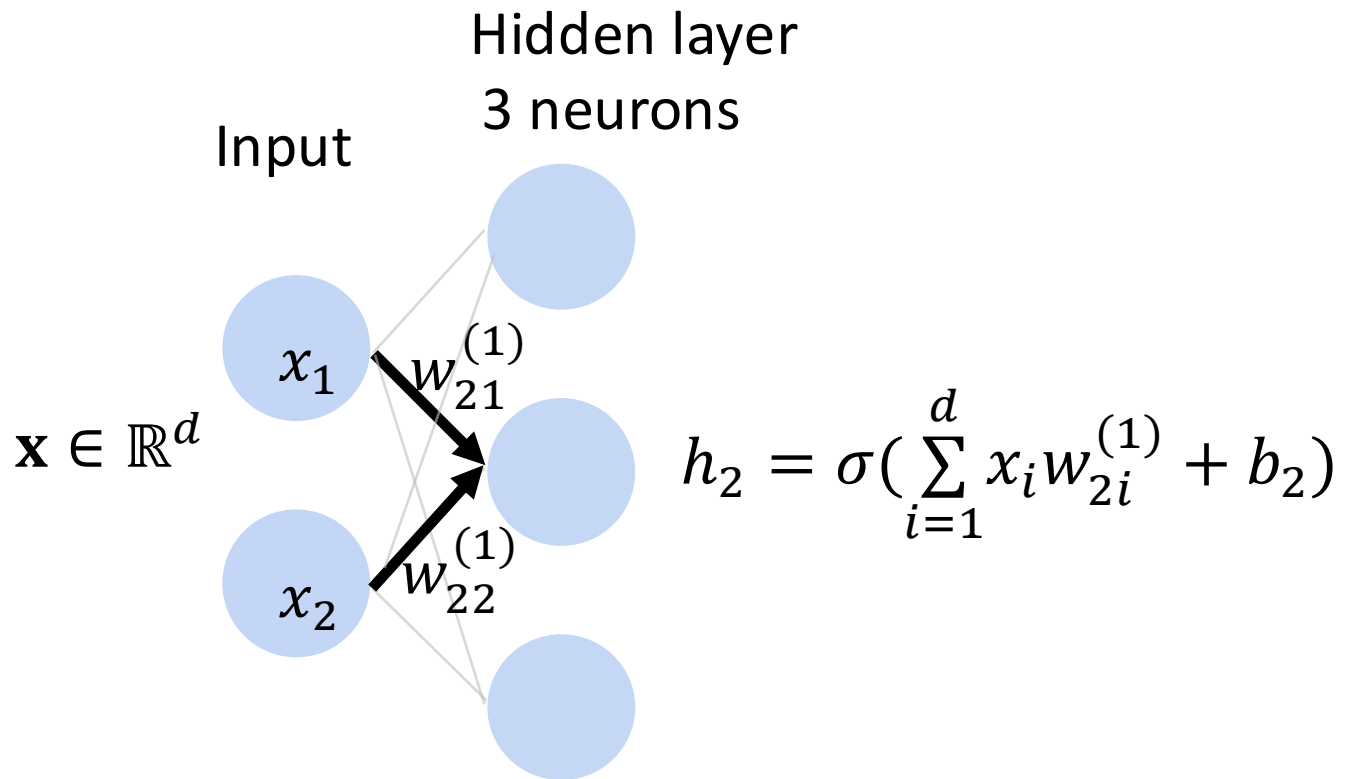
Multi-layer perceptron: Example

- Standard way to connect Perceptrons
- Example: 1 hidden layer, 1 output layer, depth = 2



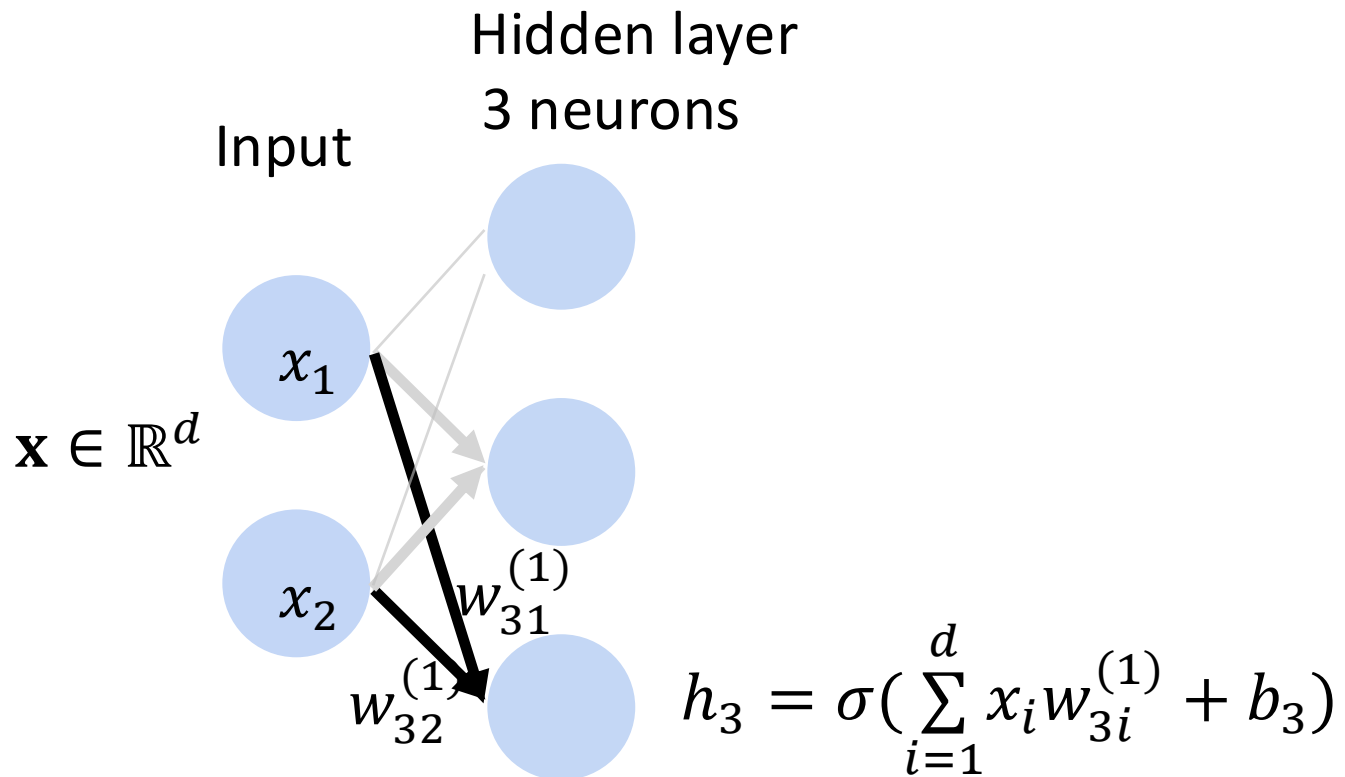
Multi-layer perceptron: Example

- Standard way to connect Perceptrons
- Example: 1 hidden layer, 1 output layer, depth = 2



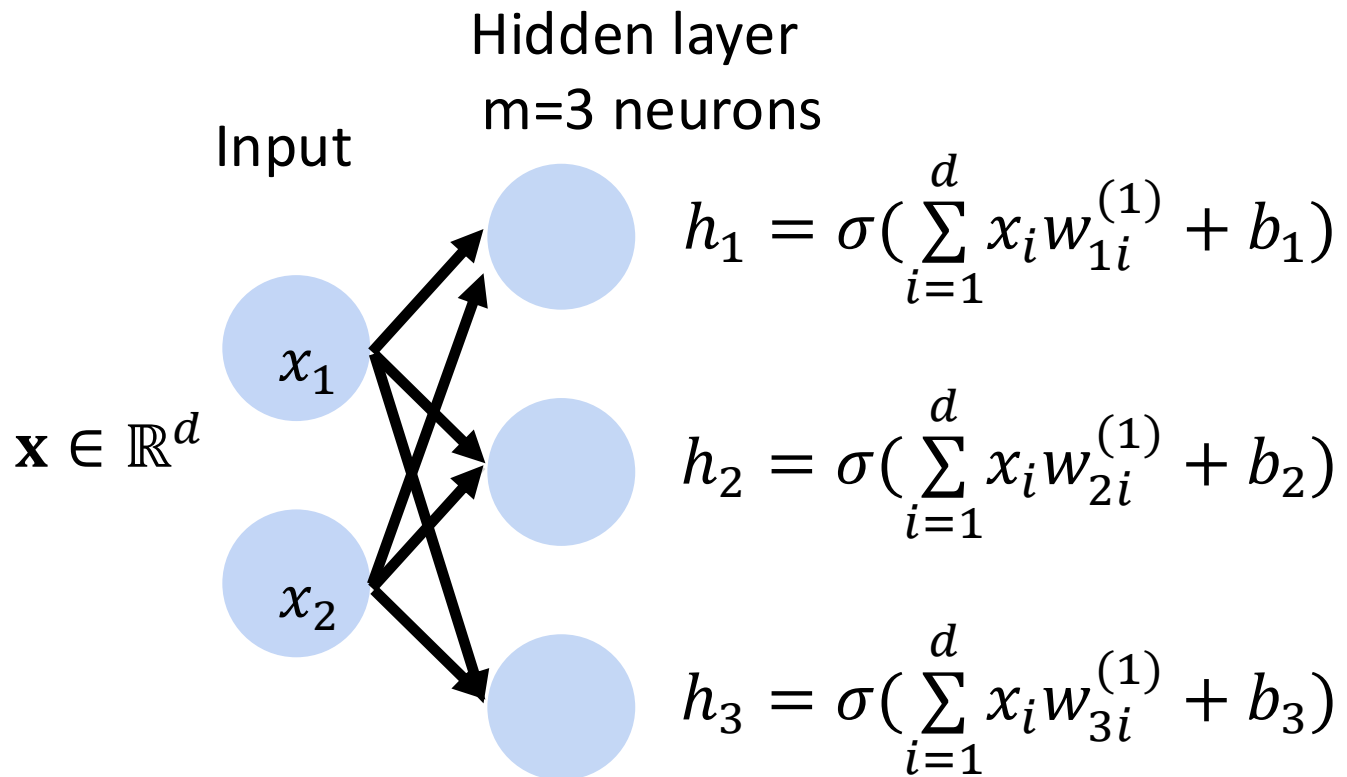
Multi-layer perceptron: Example

- Standard way to connect Perceptrons
- Example: 1 hidden layer, 1 output layer, depth = 2



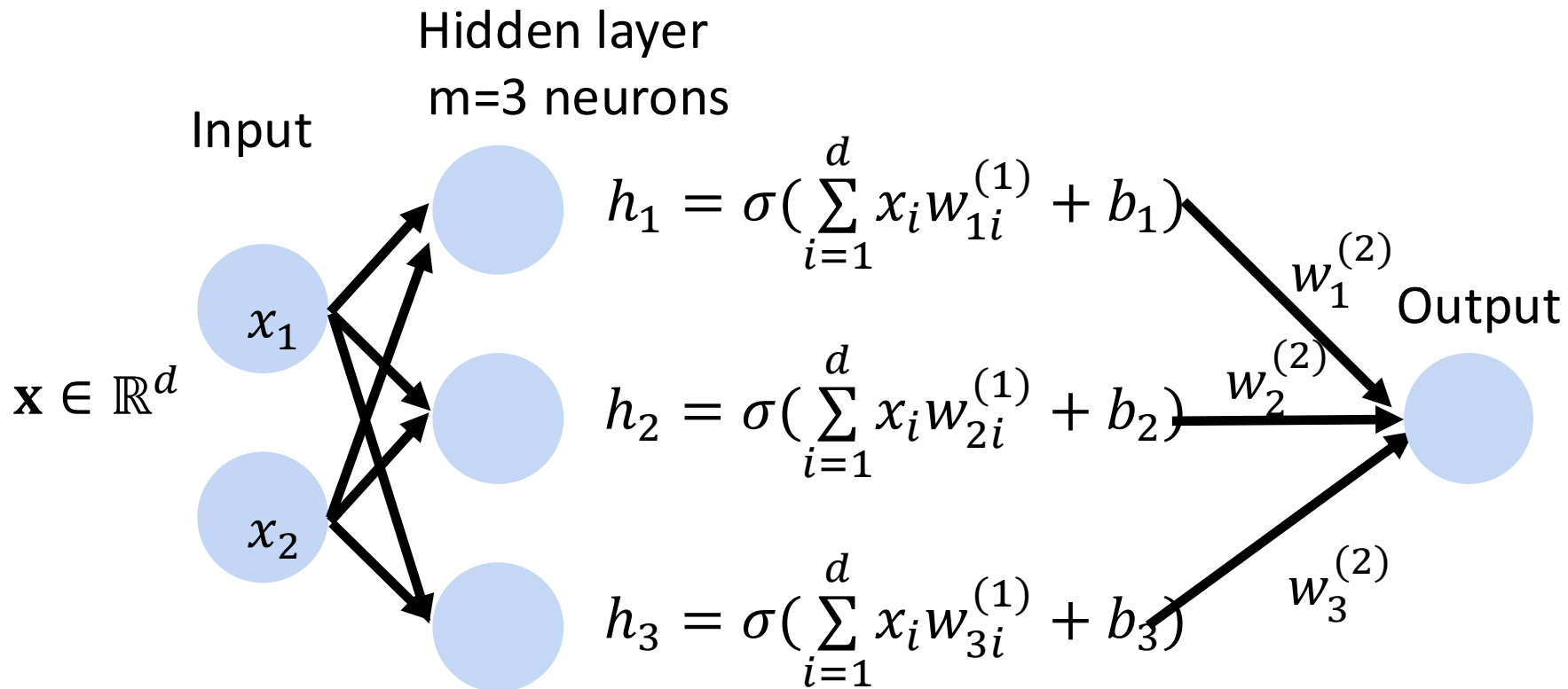
Multi-layer perceptron: Example

- Standard way to connect Perceptrons
- Example: 1 hidden layer, 1 output layer, depth = 2



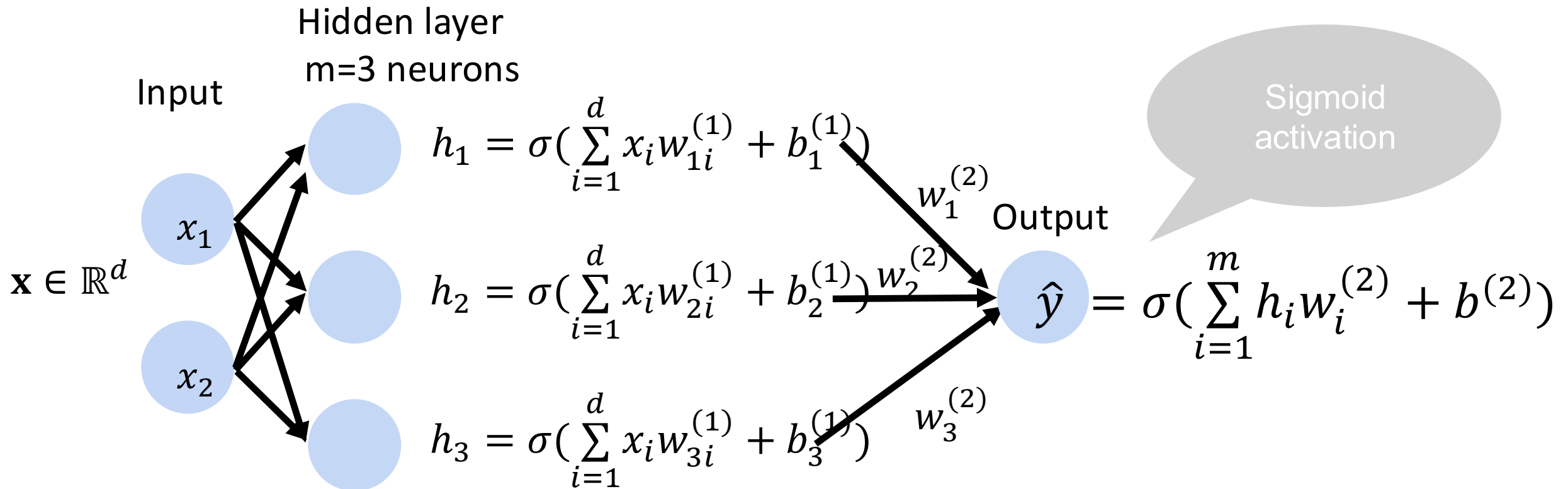
Multi-layer perceptron: Example

- Standard way to connect Perceptrons
- Example: 1 hidden layer, 1 output layer, depth = 2



Multi-layer perceptron: Example

- Standard way to connect Perceptrons
- Example: 1 hidden layer, 1 output layer, depth = 2

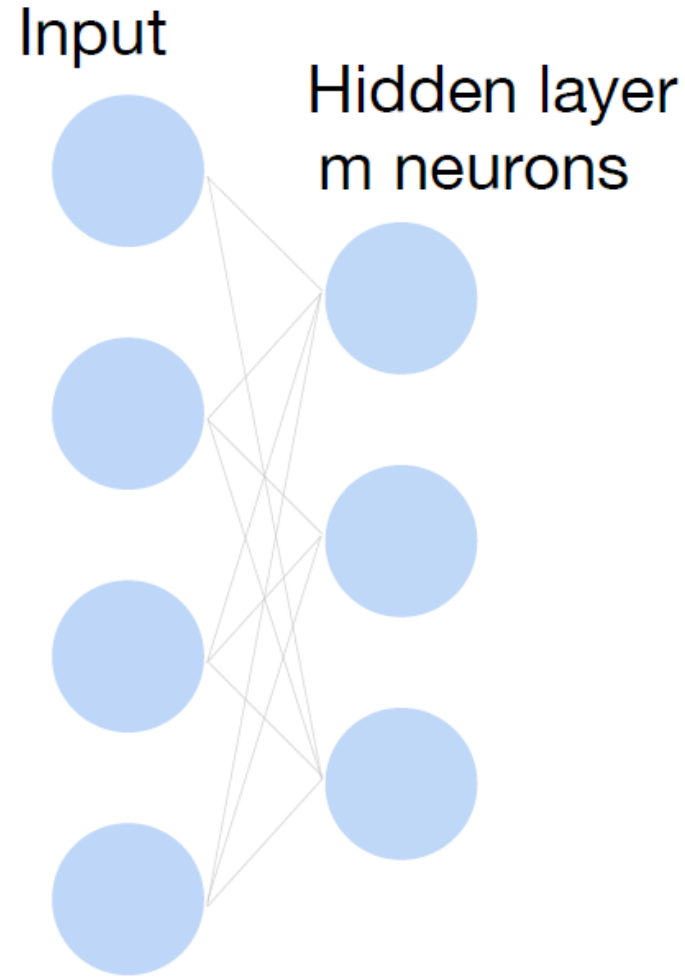


Multi-layer perceptron: Matrix Notation

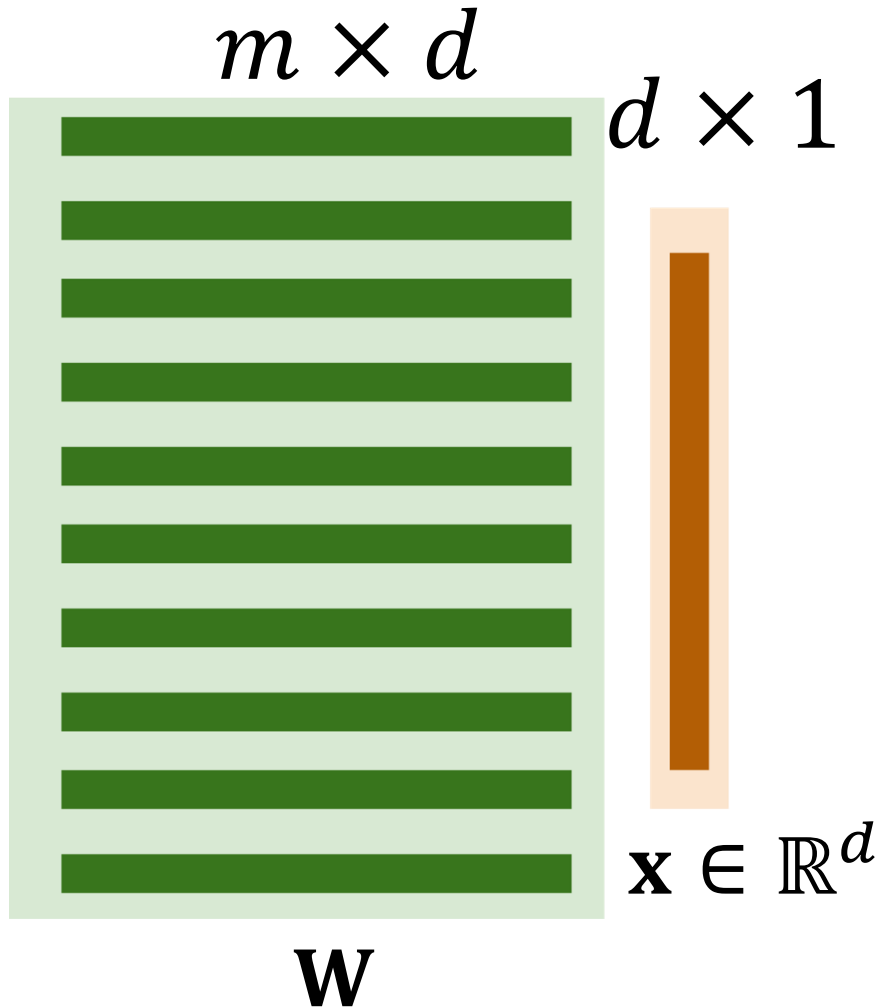
- Input $\mathbf{x} \in \mathbb{R}^d$
- Hidden $\mathbf{W}^{(1)} \in \mathbb{R}^{m \times d}$, $\mathbf{b}^{(1)} \in \mathbb{R}^m$
- Intermediate output

$$\mathbf{h} = \sigma(\mathbf{W}^{(1)}\mathbf{x} + \mathbf{b}^{(1)})$$

$$\mathbf{h} \in \mathbb{R}^m$$

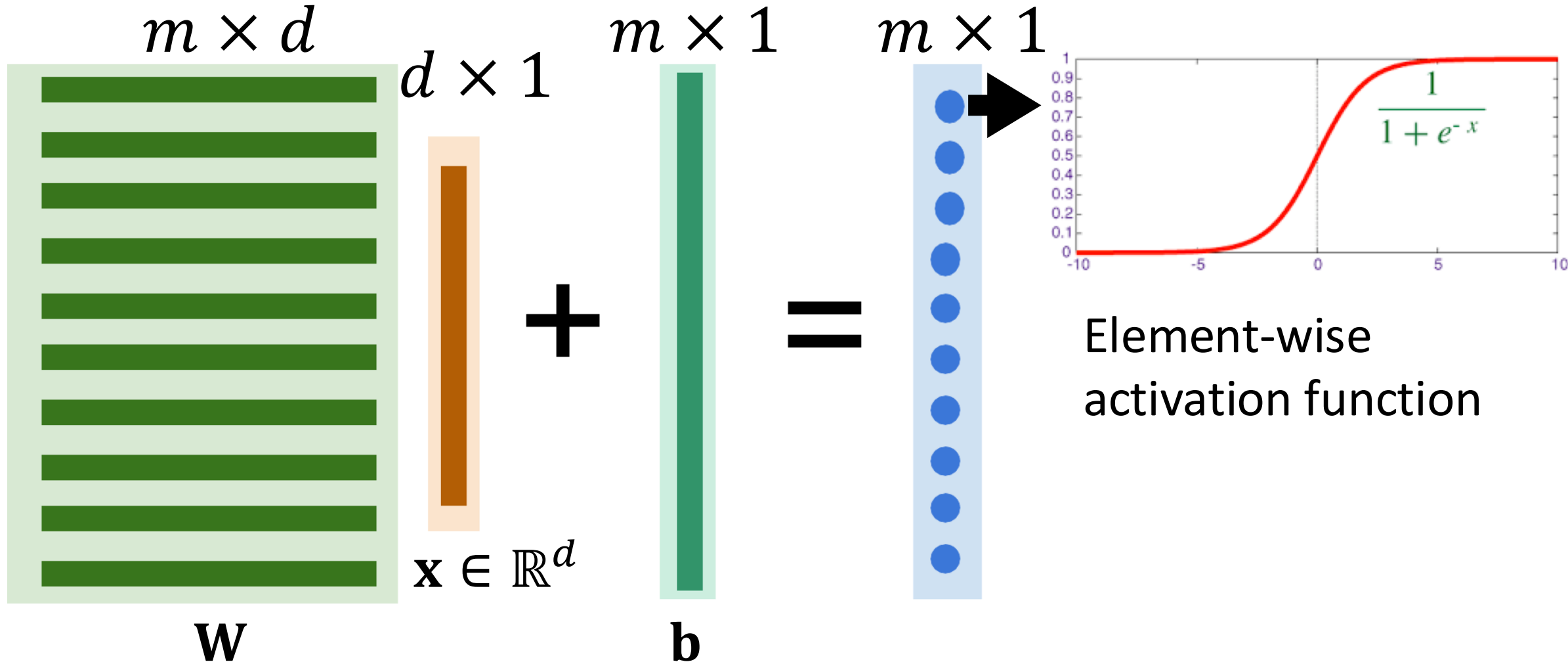


Multi-layer perceptron: **Matrix Notation**



Multi-layer perceptron: Matrix Notation

Key elements: linear operations + Nonlinear activations

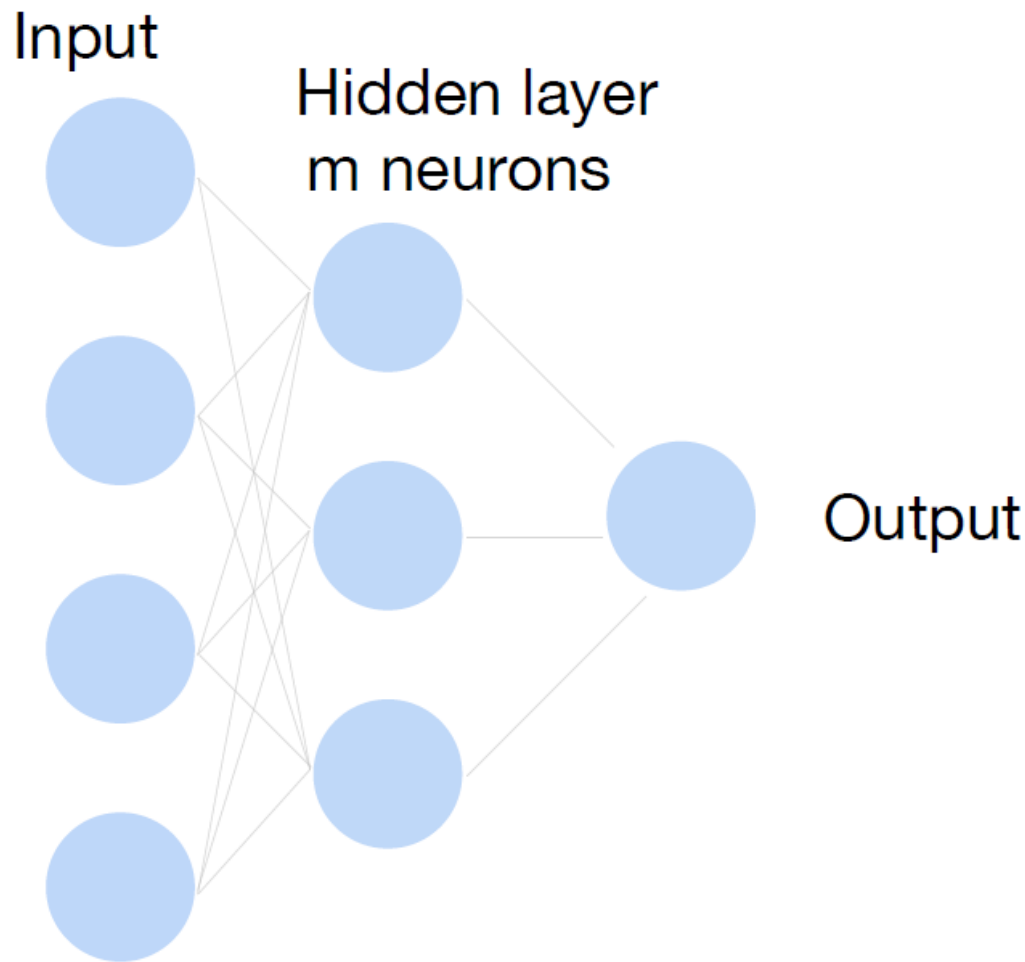


Multi-layer perceptron: Matrix Notation

- Input $\mathbf{x} \in \mathbb{R}^d$
- Hidden $\mathbf{W}^{(1)} \in \mathbb{R}^{m \times d}$, $\mathbf{b}^{(1)} \in \mathbb{R}^m$
- Intermediate output

$$\mathbf{h} = \sigma(\mathbf{W}^{(1)}\mathbf{x} + \mathbf{b}^{(1)})$$

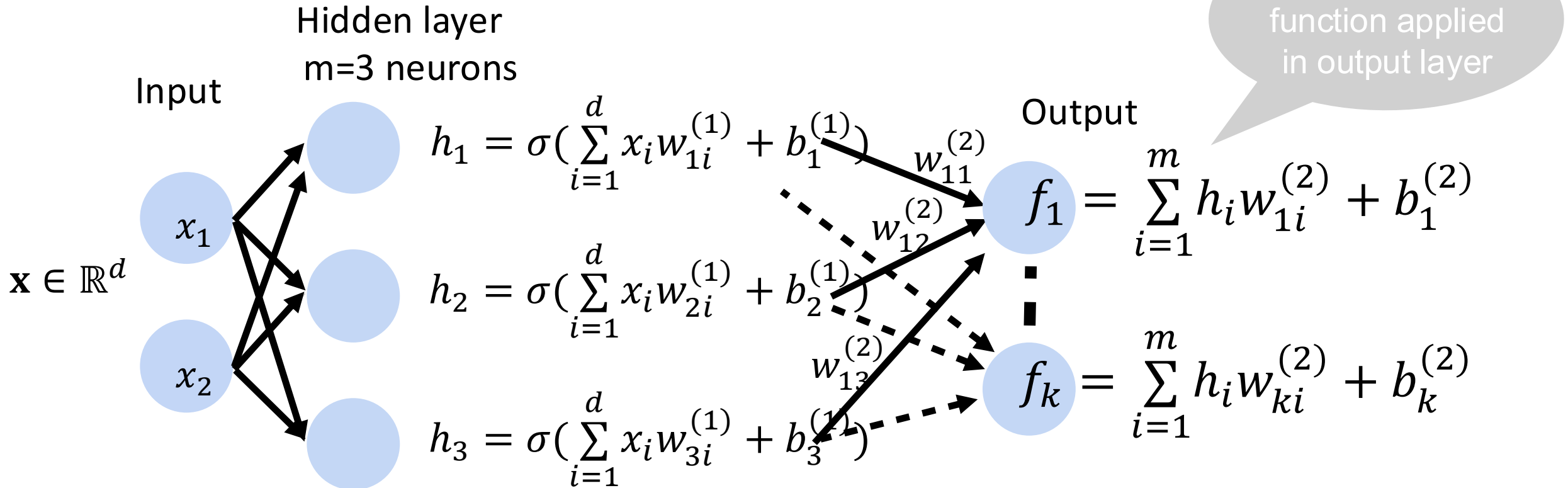
$$\hat{y} = \sigma(\mathbf{w}^{(2)}\mathbf{h} + \mathbf{b}^{(2)})$$



Neural network for k-way classification

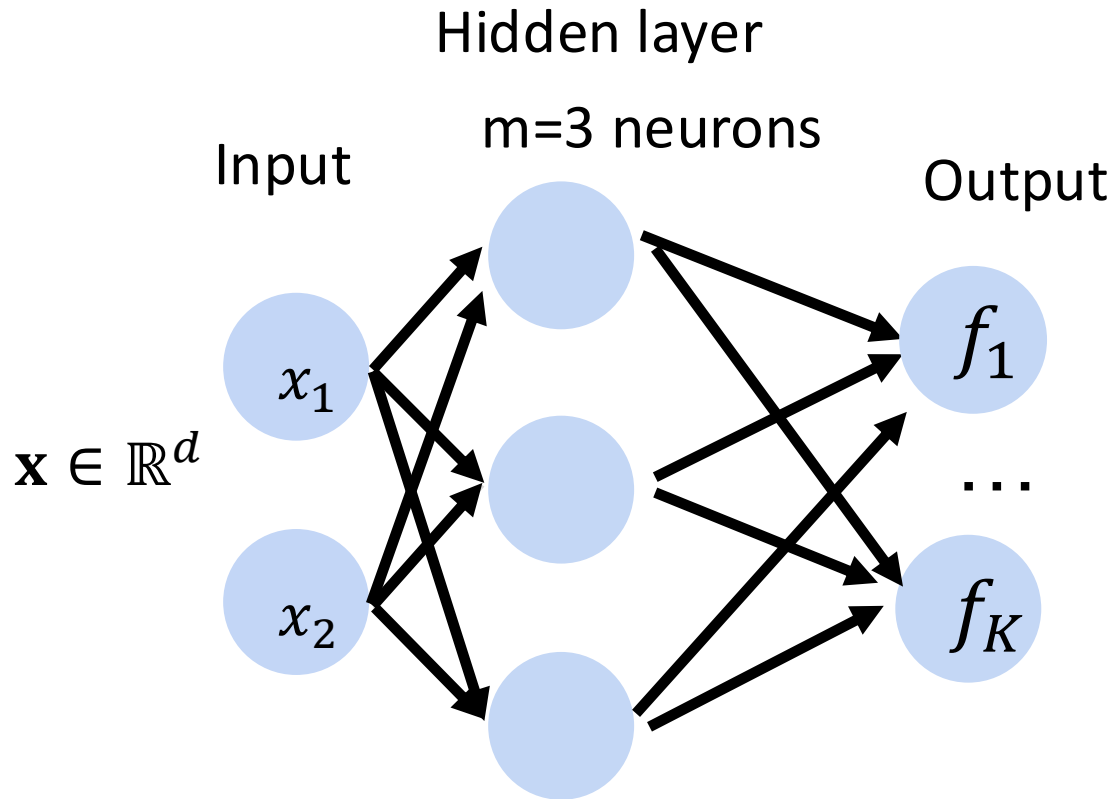
- K outputs in the final layer

Multi-class classification (e.g., ImageNet with K=1000)



Softmax

Turns outputs f into probabilities (sum up to 1 across K classes)



$$\begin{aligned} p(y|\mathbf{x}) &= \textit{softmax}(f) \\ &= \frac{\exp(f_y(x))}{\sum_{k=1}^K \exp(f_k(x))} \end{aligned}$$

Softmax

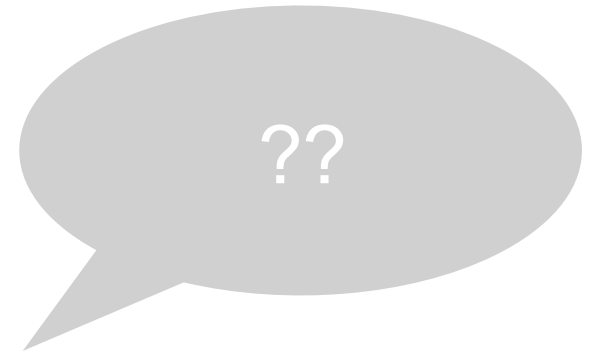
Turns outputs f into probabilities (sum up to 1 across K classes)

Output
layer

$$\begin{bmatrix} 1.3 \\ 5.1 \\ 2.2 \\ 0.7 \\ 1.1 \end{bmatrix}$$

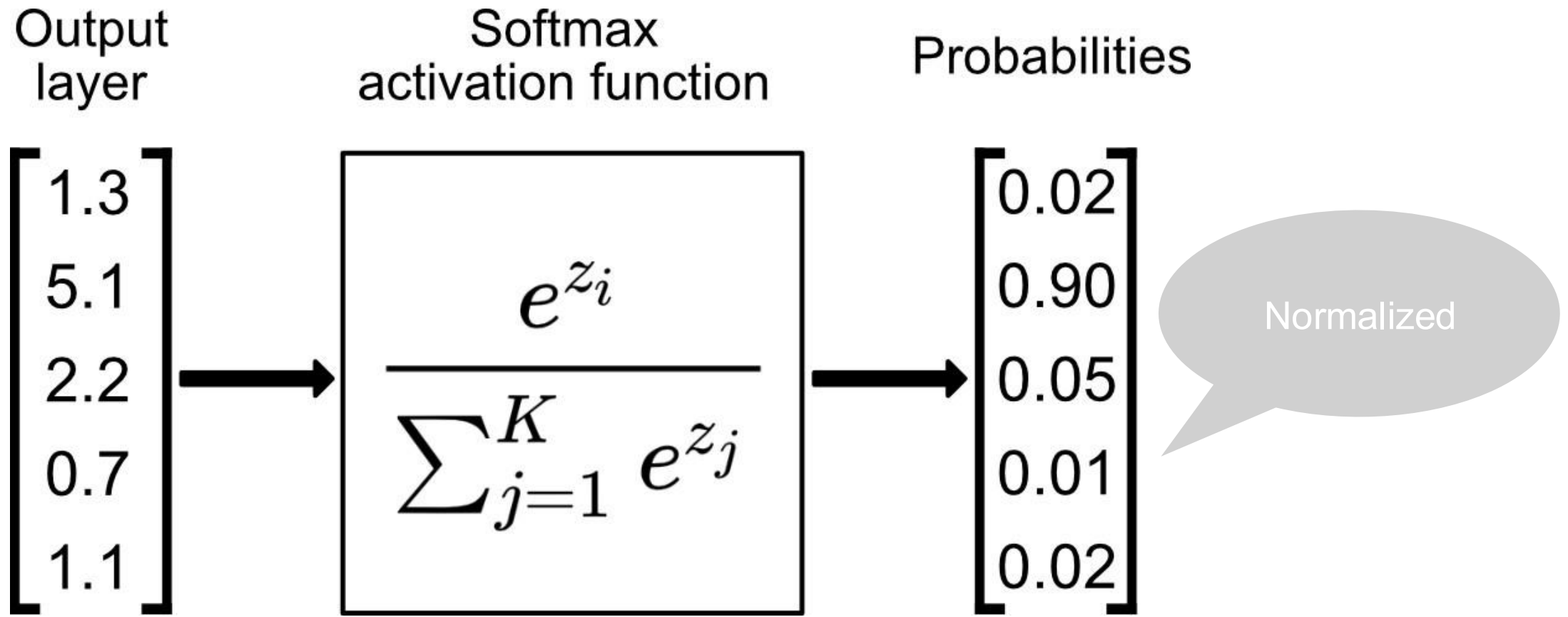
Softmax
activation function

$$\frac{e^{z_i}}{\sum_{j=1}^K e^{z_j}}$$



Softmax

Turns outputs f into probabilities (sum up to 1 across K classes)



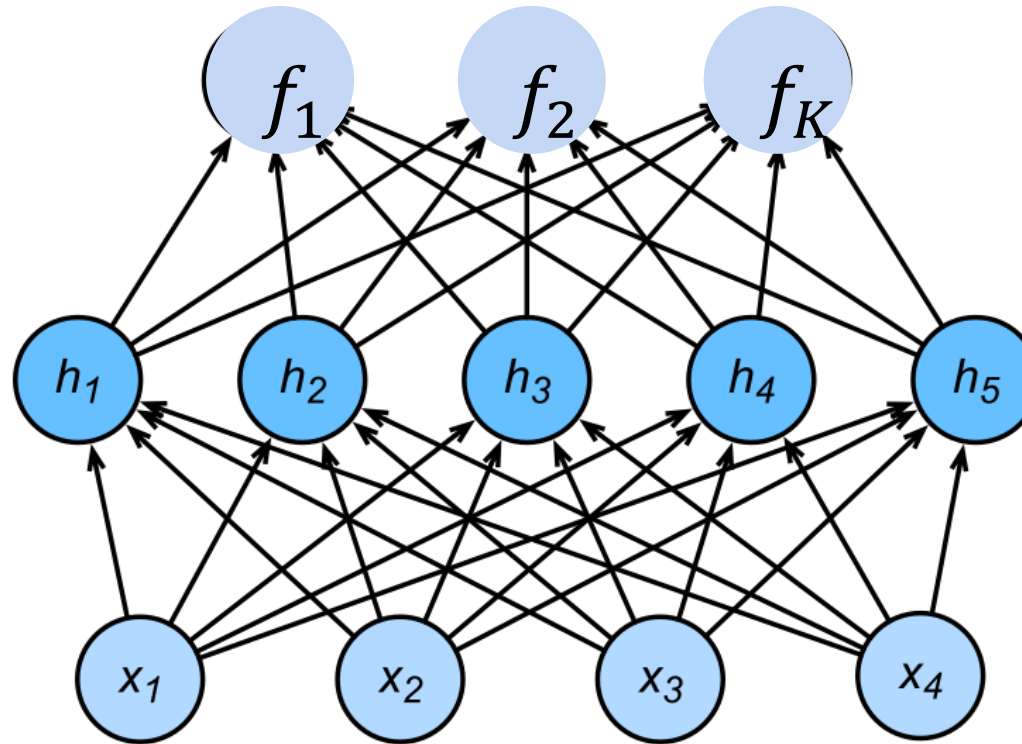
More complicated neural networks

$$p_1, p_2, \dots, p_K = \text{softmax}(f_1, f_2, \dots, f_K)$$

Output layer

Hidden layer

Input layer



More complicated neural networks

- Input $\mathbf{x} \in \mathbb{R}^d$

- Hidden $\mathbf{W}^{(1)} \in \mathbb{R}^{m \times d}, \mathbf{b}^{(1)} \in \mathbb{R}^m$

$$p_1, p_2, \dots, p_K = \text{softmax}(f_1, f_2, \dots, f_K)$$

Output layer

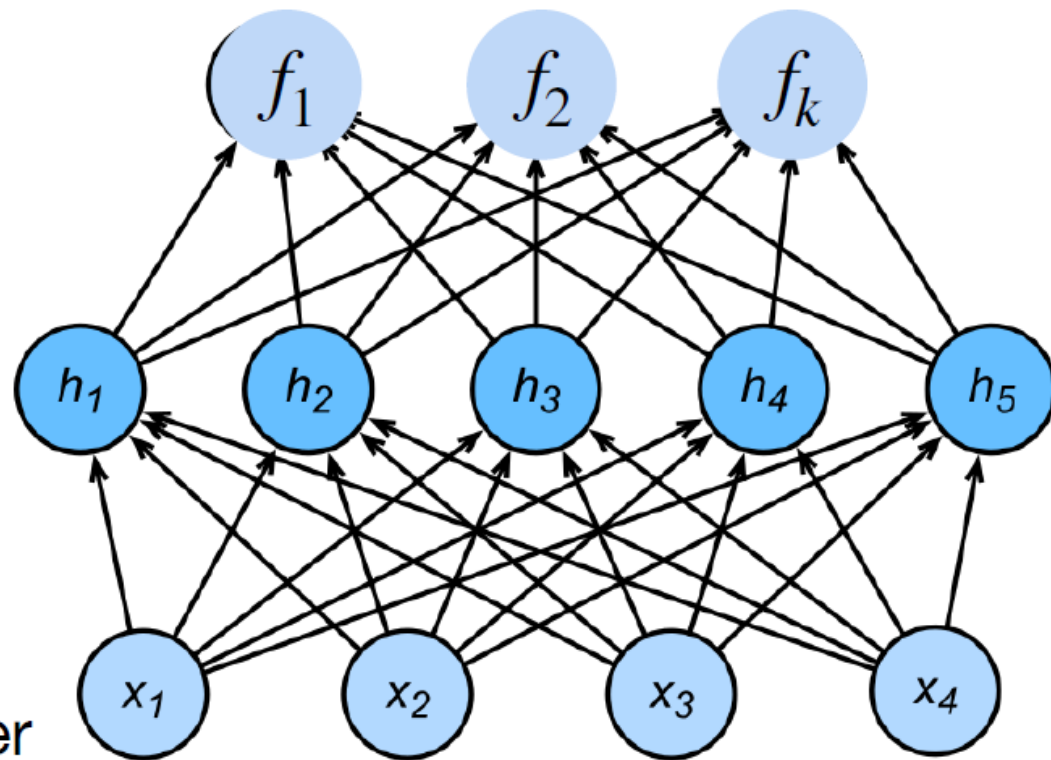
$$\mathbf{h} = \sigma(\mathbf{W}^{(1)}\mathbf{x} + \mathbf{b}^{(1)})$$

$$\mathbf{f} = \mathbf{W}^{(2)}\mathbf{h} + \mathbf{b}^{(2)}$$

$$\mathbf{p} = \text{softmax}(\mathbf{f})$$

Hidden layer

Input layer



More complicated neural networks: multiple hidden layers

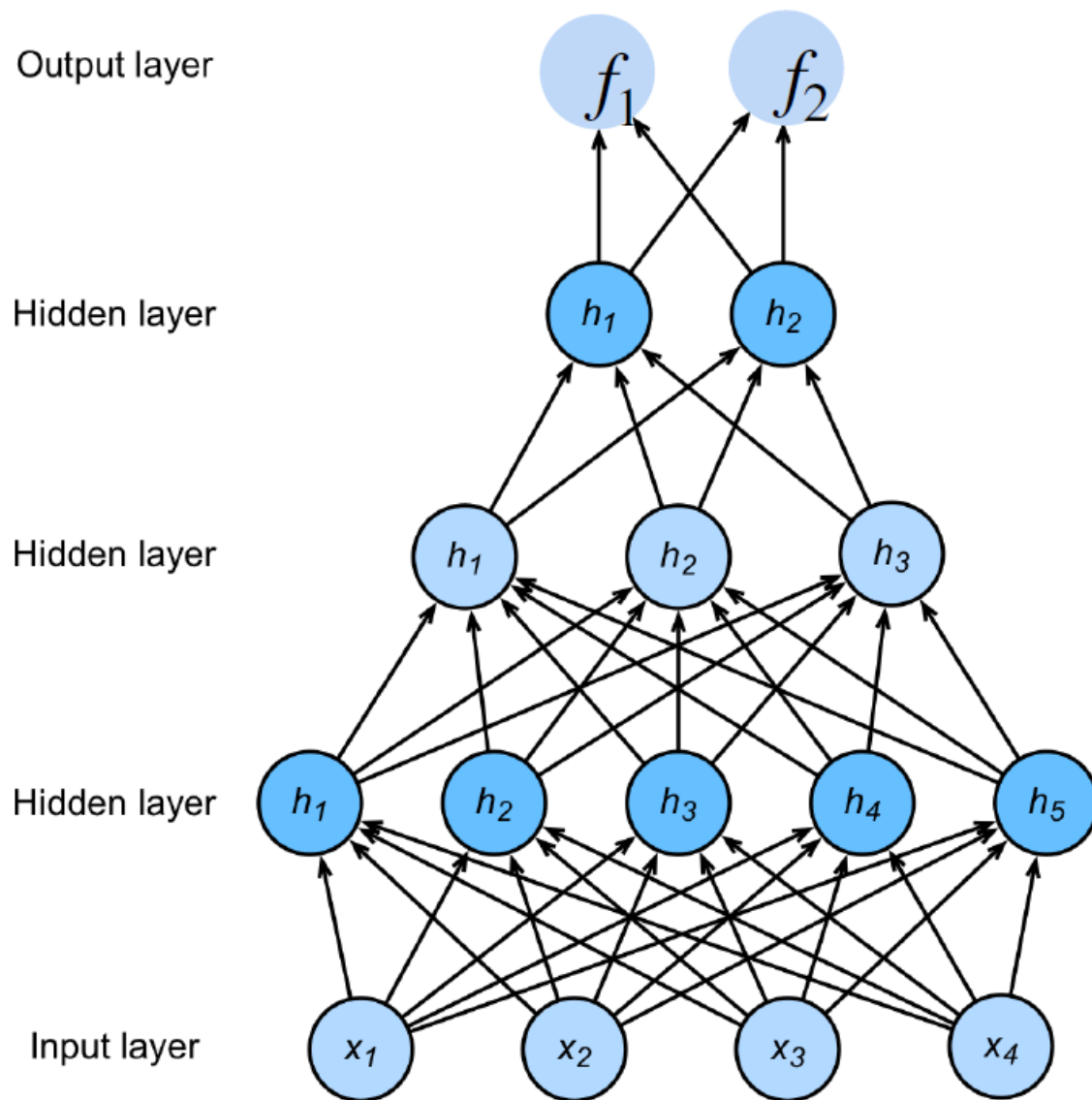
$$\mathbf{h}_1 = \sigma(\mathbf{W}^{(1)}\mathbf{x} + \mathbf{b}^{(1)})$$

$$\mathbf{h}_2 = \sigma(\mathbf{W}^{(2)}\mathbf{h}_1 + \mathbf{b}^{(2)})$$

$$\mathbf{h}_3 = \sigma(\mathbf{W}^{(3)}\mathbf{h}_2 + \mathbf{b}^{(3)})$$

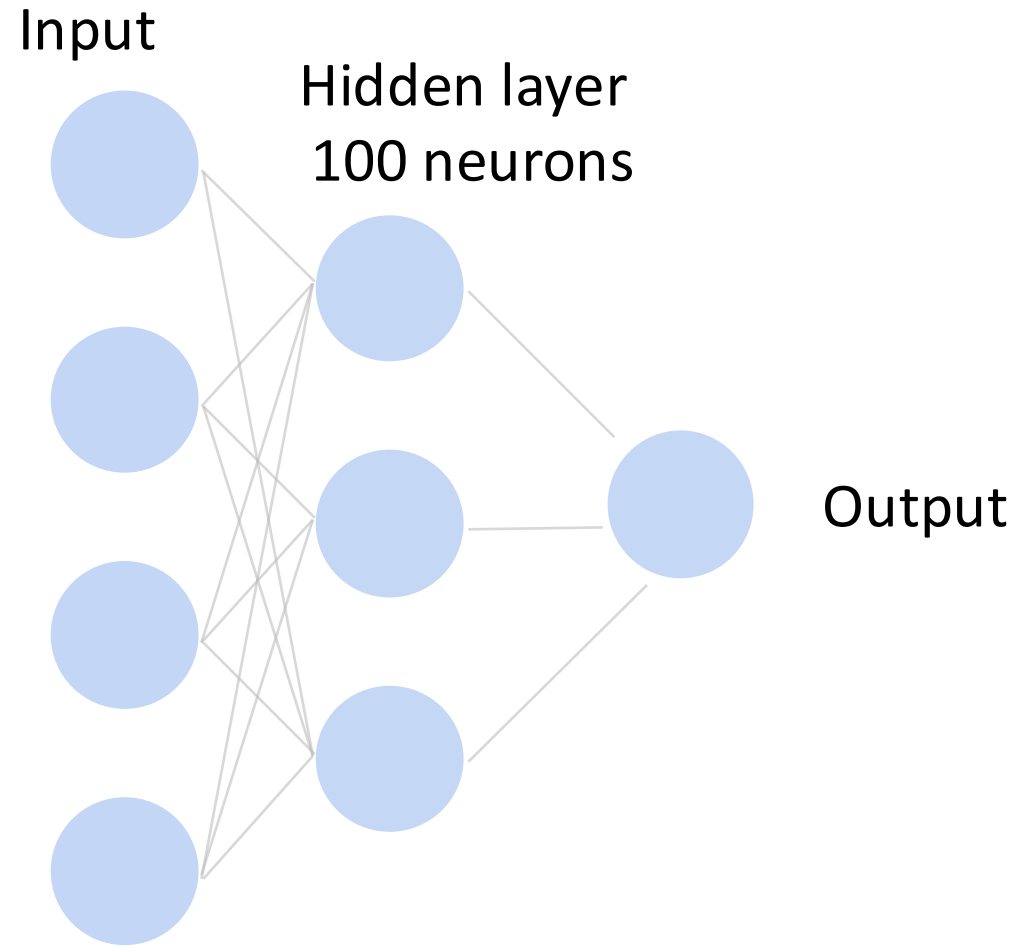
$$\mathbf{f} = \mathbf{W}^{(4)}\mathbf{h}_3 + \mathbf{b}^{(4)}$$

$$\mathbf{p} = \text{softmax}(\mathbf{f})$$



How to train a neural network?

Classify cats vs. dogs



How to train a neural network? Binary classification

$\mathbf{x} \in \mathbb{R}^d$ One training data point in the training set D

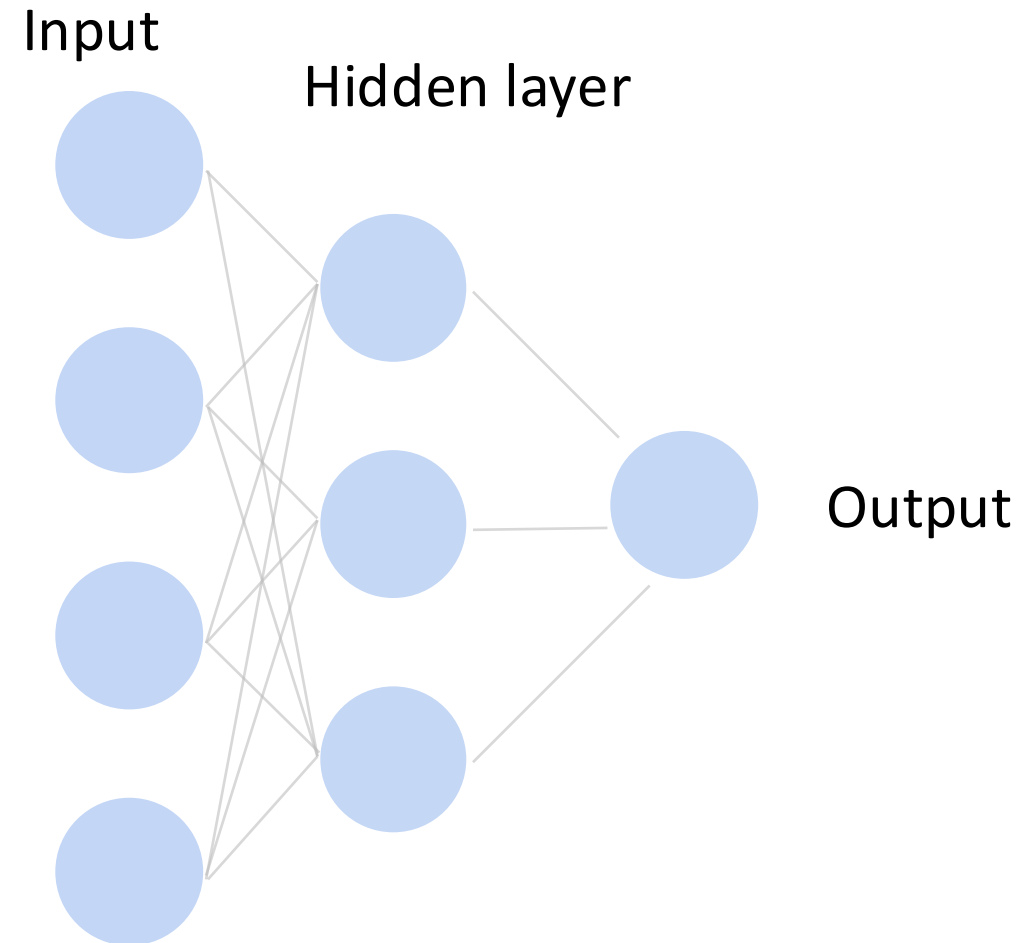
$\hat{y} \in [0,1]$ Model output for example $\mathbf{x}$

(This is a function of all weights W : $\hat{y} = g(W)$)

y Ground truth label for example $\mathbf{x}$

Learning by matching the output to the label

**We want $\hat{y} \rightarrow 1$ when $y = 1$,
and $\hat{y} \rightarrow 0$ when $y = 0$**



How to train a neural network? Binary classification

Loss function: $\frac{1}{|D|} \sum_{(\mathbf{x}, y) \in D} \ell(\mathbf{x}, y)$

Per-sample loss:

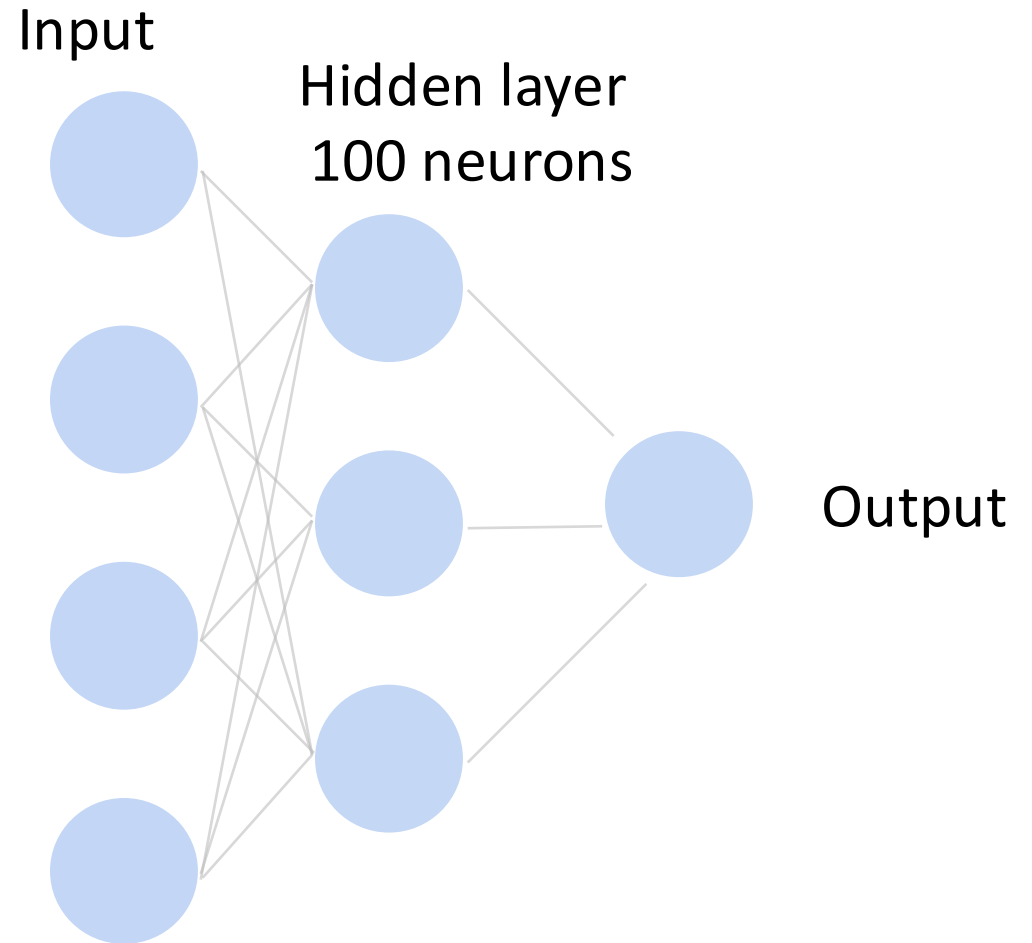
$$\ell(\mathbf{x}, y) = -y \log(\hat{y}) - (1 - y) \log(1 - \hat{y})$$



Negative log likelihood

Also known as **binary cross-entropy loss**

- We'll use all the time for LLMs



How to train a neural network? Multiclass

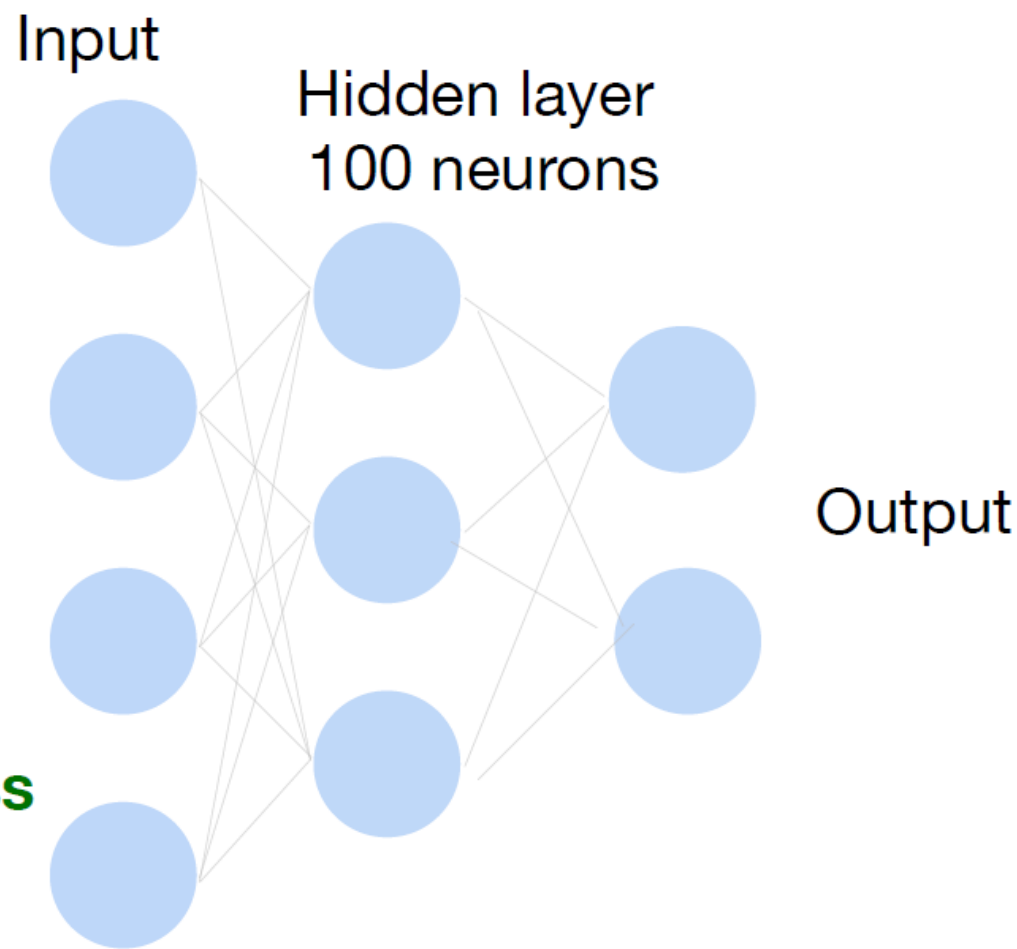
Loss function: $\frac{1}{|D|} \sum_{(\mathbf{x}, y) \in D} \ell(\mathbf{x}, y)$

Per-sample loss:

$$\ell(\mathbf{x}, y) = \sum_{k=1}^K -Y_k \log p_k = -\log p_y$$

where Y is one-hot encoding of y

Also known as **cross-entropy loss**
or **softmax loss**



Training Neural Networks

- Algorithm:

- Get

$$D = \{(x^{(1)}, y^{(1)}), \dots, (x^{(n)}, y^{(n)})\}$$

- Initialize weights

- Until stopping criteria met,

- For each training point $(x^{(i)}, y^{(i)})$

- Compute: $f_{\text{network}}(x^{(d)})$

← Forward Pass

- Compute gradient: $\nabla L^{(i)}(w) = \left[\frac{\partial L^{(d)}}{\partial w_0}, \frac{\partial L^{(d)}}{\partial w_1}, \dots, \frac{\partial L^{(d)}}{\partial w_m} \right]^T$

← Backward Pass

- Update weights:

$$w \leftarrow w - \alpha \nabla L^{(i)}(w)$$



Break & Questions

Outline

- **Neural Networks**

- Perceptrons, setup, training, limitations, multilayer perceptrons, loss functions (mostly from last time)

- **Convolutional Neural Networks**

- Motivation, convolutional layers, CNN architectures

- **Sequence Models**

- Recurrent neural networks, architecture, LSTMs, alternatives, training tricks

How to classify Cats vs. dogs?

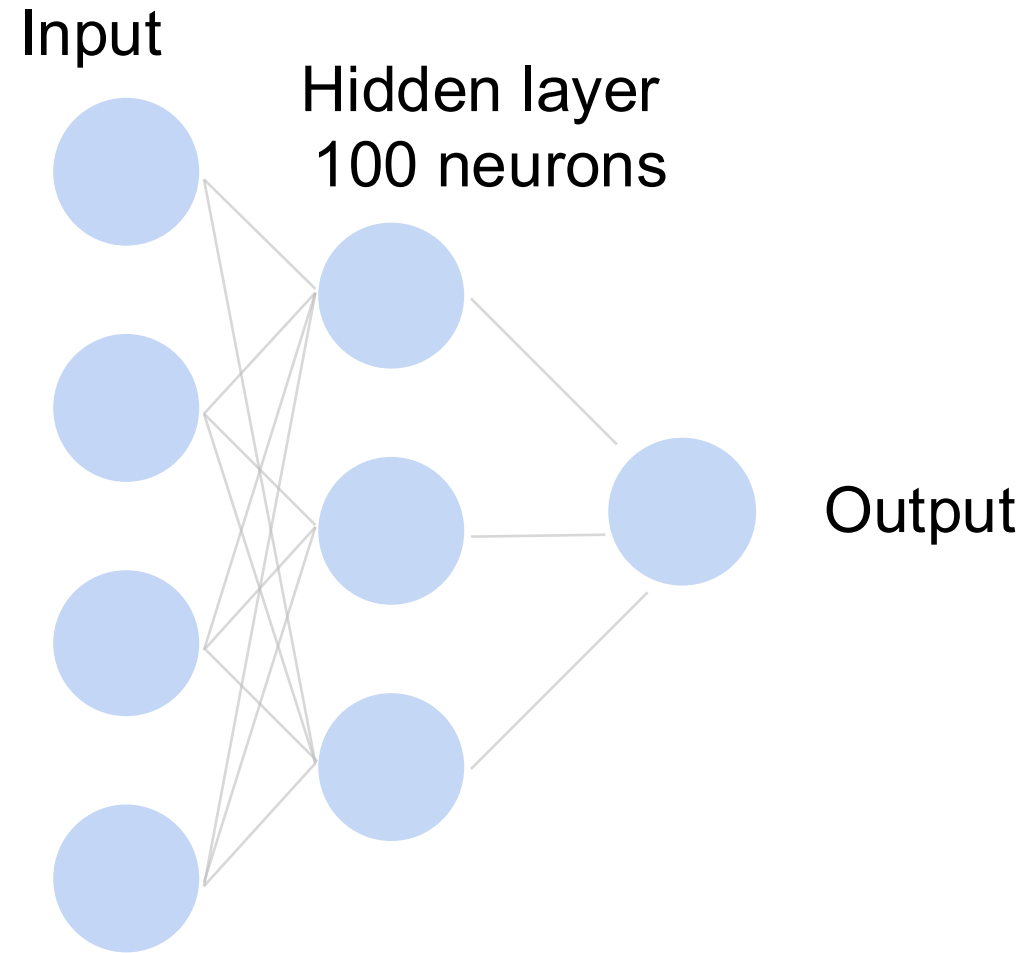


Dual
12MP
wide-angle and
telephoto cameras

**36M floats in a RGB
image!**

Fully Connected Networks

Cats vs. dogs?



$\sim 36\text{M elements} \times 100 = \sim \mathbf{3.6B}$ parameters!

Neural Networks: Convolution Layers

- Notation:
 - $X: n_h \times n_w$ input matrix
 - $W: k_h \times k_w$ kernel matrix
 - b : bias (a scalar)
 - $Y: () \times ()$ output matrix
- As usual W, b are learnable parameters

0	1	2
3	4	5
6	7	8

 *

0	1
2	3

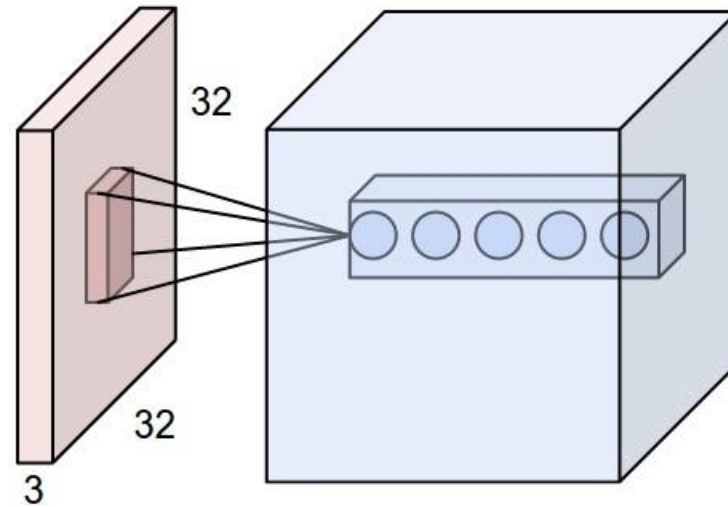
 =

19	25
37	43

Neural Networks: Convolution NNs

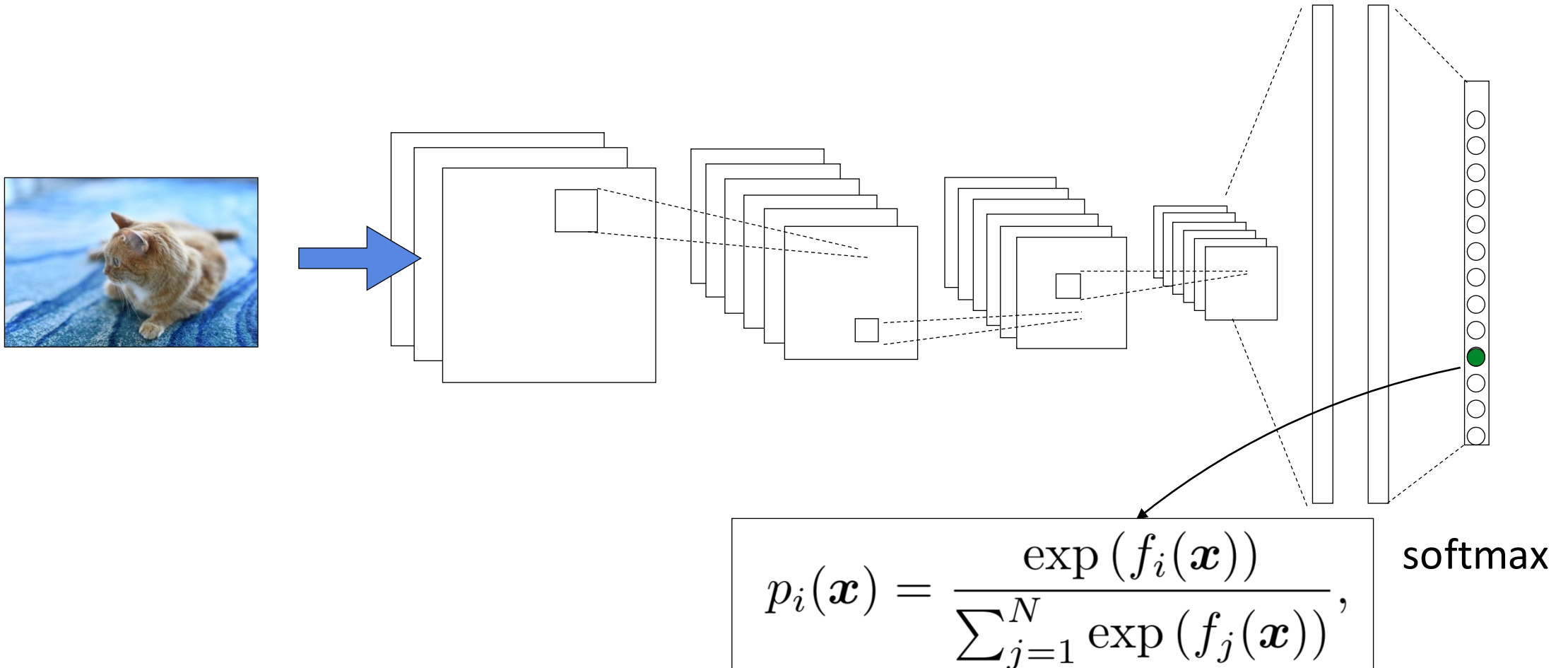
- Properties

- Input: volume $c_i \times n_h \times n_w$ (channels x height x width)
- Hyperparameters: # of kernels/filters c_o , size $k_h \times k_w$, stride $s_h \times s_w$, zero padding $p_h \times p_w$
- Output: volume $c_o \times m_h \times m_w$ (channels x height x width)
- Parameters: $k_h \times k_w \times c_i$ per filter, total $(k_h \times k_w \times c_i) \times c_o$



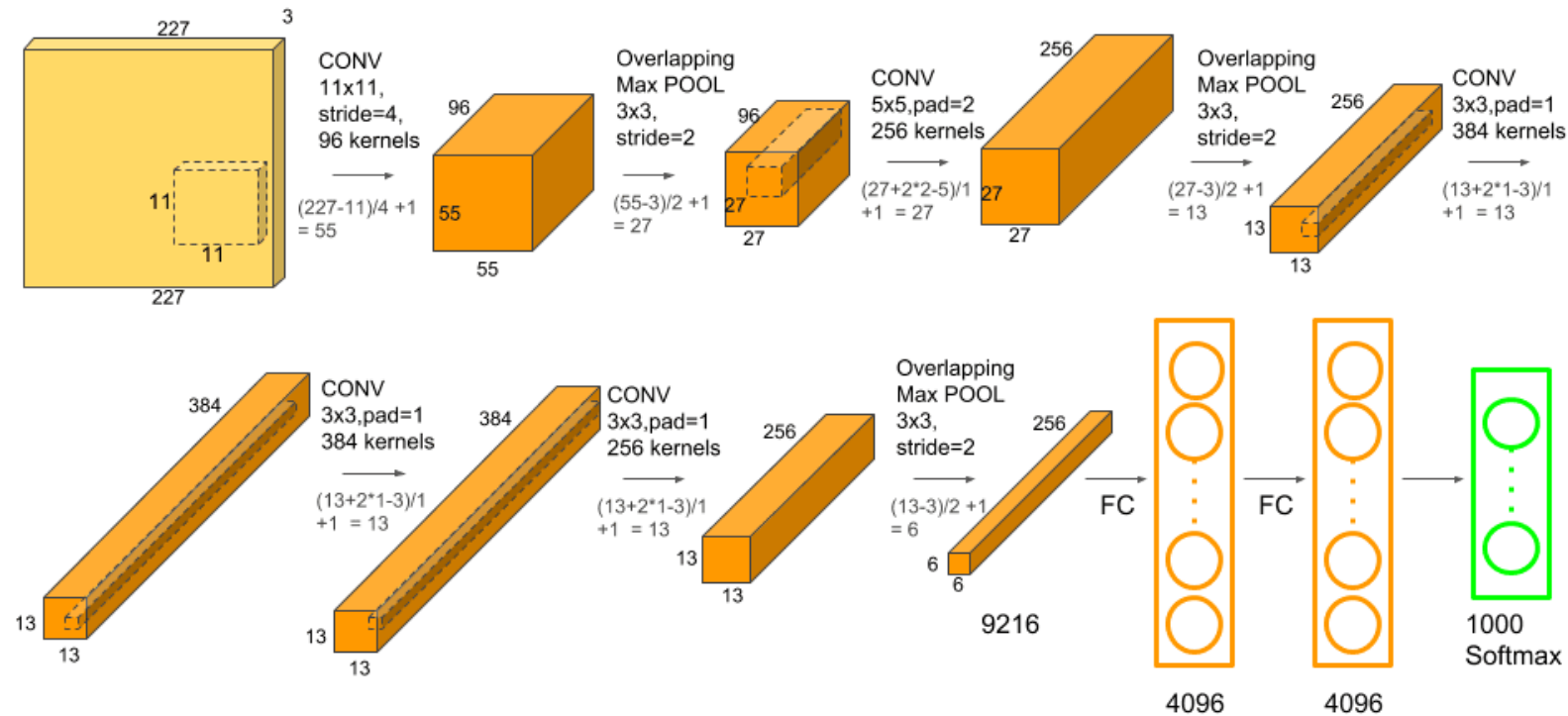
Training a CNN

- Q: so we have a bunch of layers. How do we train?
- A: same as before. Apply softmax at the end, use backprop.



CNN Architectures: AlexNet

- First of the major advancements: AlexNet
- Wins 2012 ImageNet competition
- Major trends: deeper, bigger LeNet





Break & Questions

Outline

- **Neural Networks**

- Perceptrons, setup, training, limitations, multilayer perceptrons, loss functions (mostly from last time)

- **Convolutional Neural Networks**

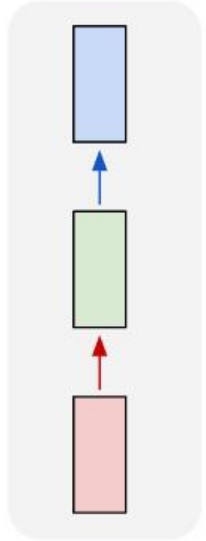
- Motivation, convolutional layers, CNN architectures

- **Sequence Models**

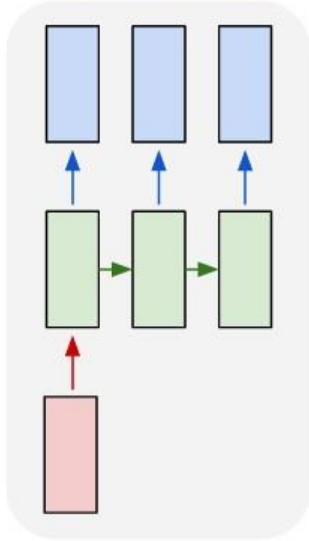
- Recurrent neural networks, architecture, LSTMs, alternatives, training tricks

Tasks We Can Handle with NNs?

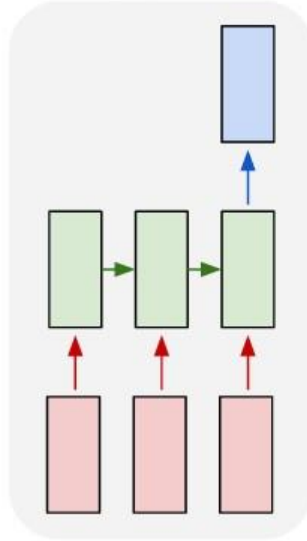
one to one



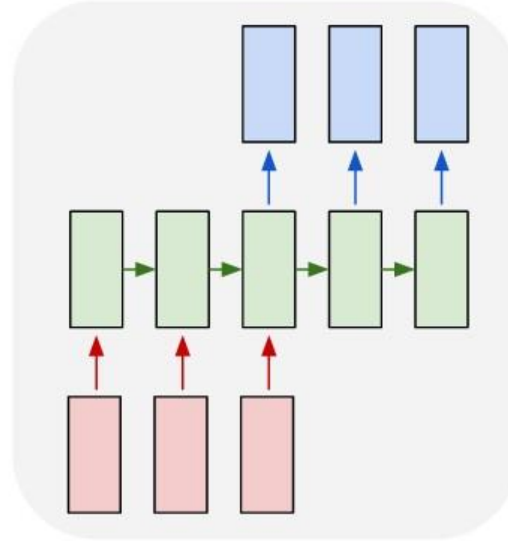
one to many



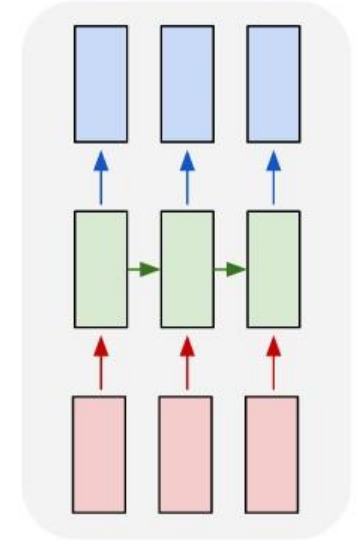
many to one



many to many



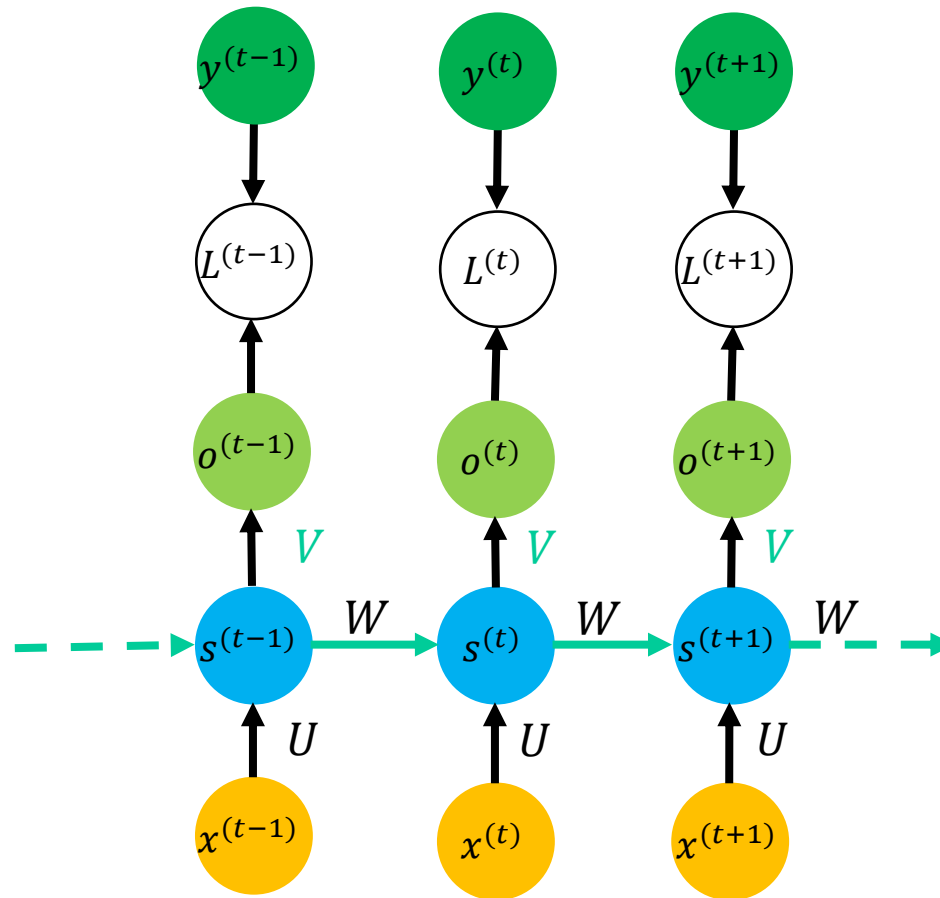
many to many



- Mostly talked about (1) so far
 - Others: need a new kind of model

Neural Networks: Simple RNNs

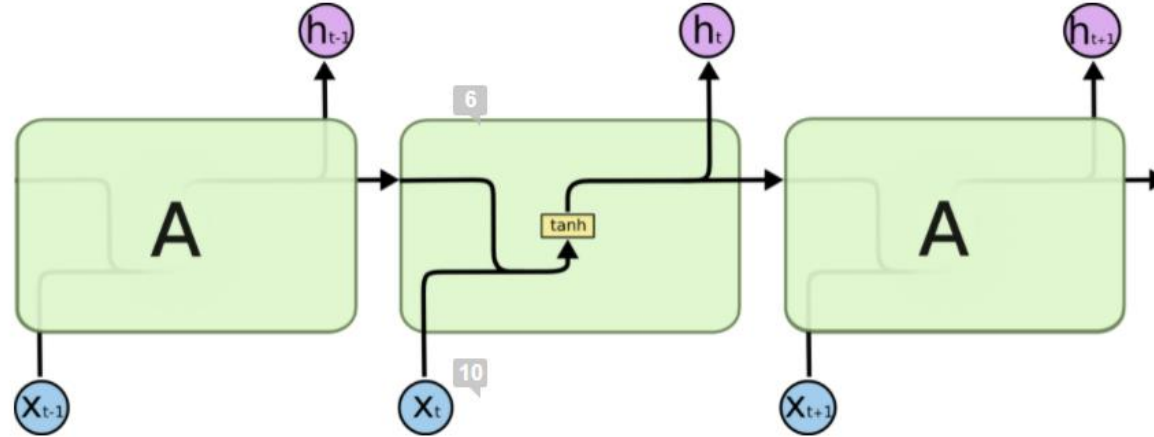
- Classical RNN variant:



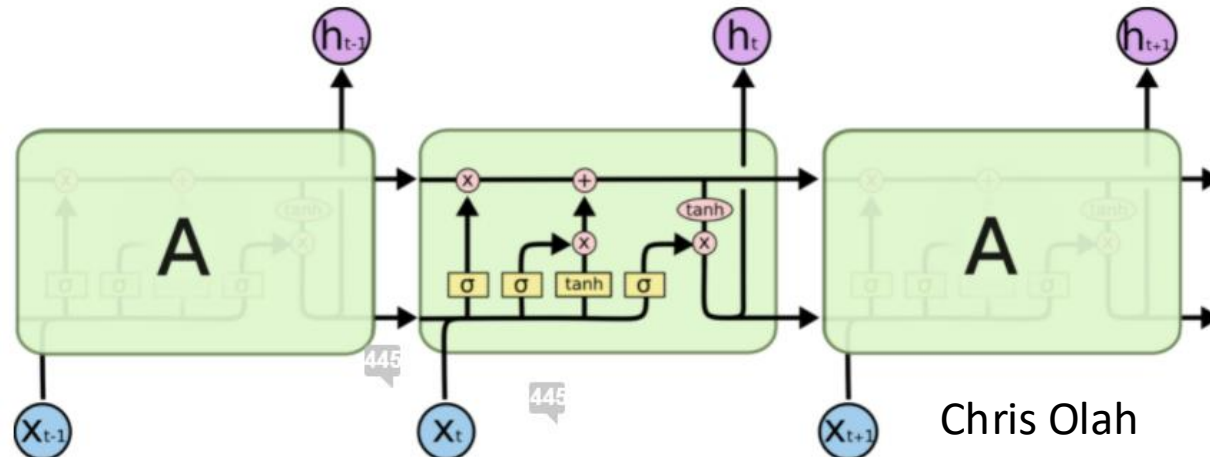
$$\begin{aligned} a^{(t)} &= b + Ws^{(t-1)} + Ux^{(t)} \\ s^{(t)} &= \tanh(a^{(t)}) \\ o^{(t)} &= c + Vs^{(t)} \\ \hat{y}^{(t)} &= \text{softmax}(o^{(t)}) \\ L^{(t)} &= \text{CrossEntropy}(y^{(t)}, \hat{y}^{(t)}) \end{aligned}$$

Neural Networks: LSTMs

- RNN: can write structure as:

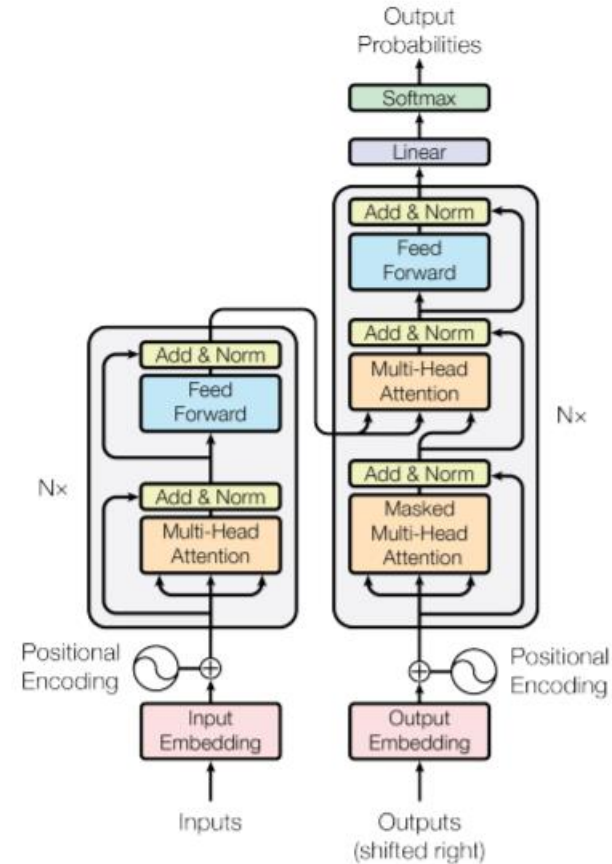


- Long Short-Term Memory: deals with problem. Cell:



Neural Networks: Transformers

- Initial goal for an architecture: **encoder-decoder**
 - Get **rid of recurrence**
 - Replace with **self-attention**
- Architecture
 - The famous picture you've seen
 - Centered on self-attention blocks



Data Augmentation

Augmentation: transform + add new samples to dataset

- Transformations: based on domain
- Idea: build **invariances** into the model
 - **Ex:** if all images have same alignment, model learns to use it
- Keep the label the same!



Data Augmentation: Examples

Examples of transformations for images

- **Crop** (and zoom)
- **Color** (change contrast/brightness)
- **Rotations+** (translate, stretch, shear, etc)

Many more possibilities. Combine as well!

Q: how to deal with this at **test time**?

- A: transform, test, average

