

NEURAL NETWORKS FROM SCRATCH

LECTURE 02 - NEURAL NETWORKS

Hendrik Borras, Robin Janssen
{hendrik.borras, robin.janssen}@ziti.uni-heidelberg.de,
HAWAI Lab, Institute of Computer Engineering
Heidelberg University

RECAP

Example: Polynomial Fit

Parameters and training

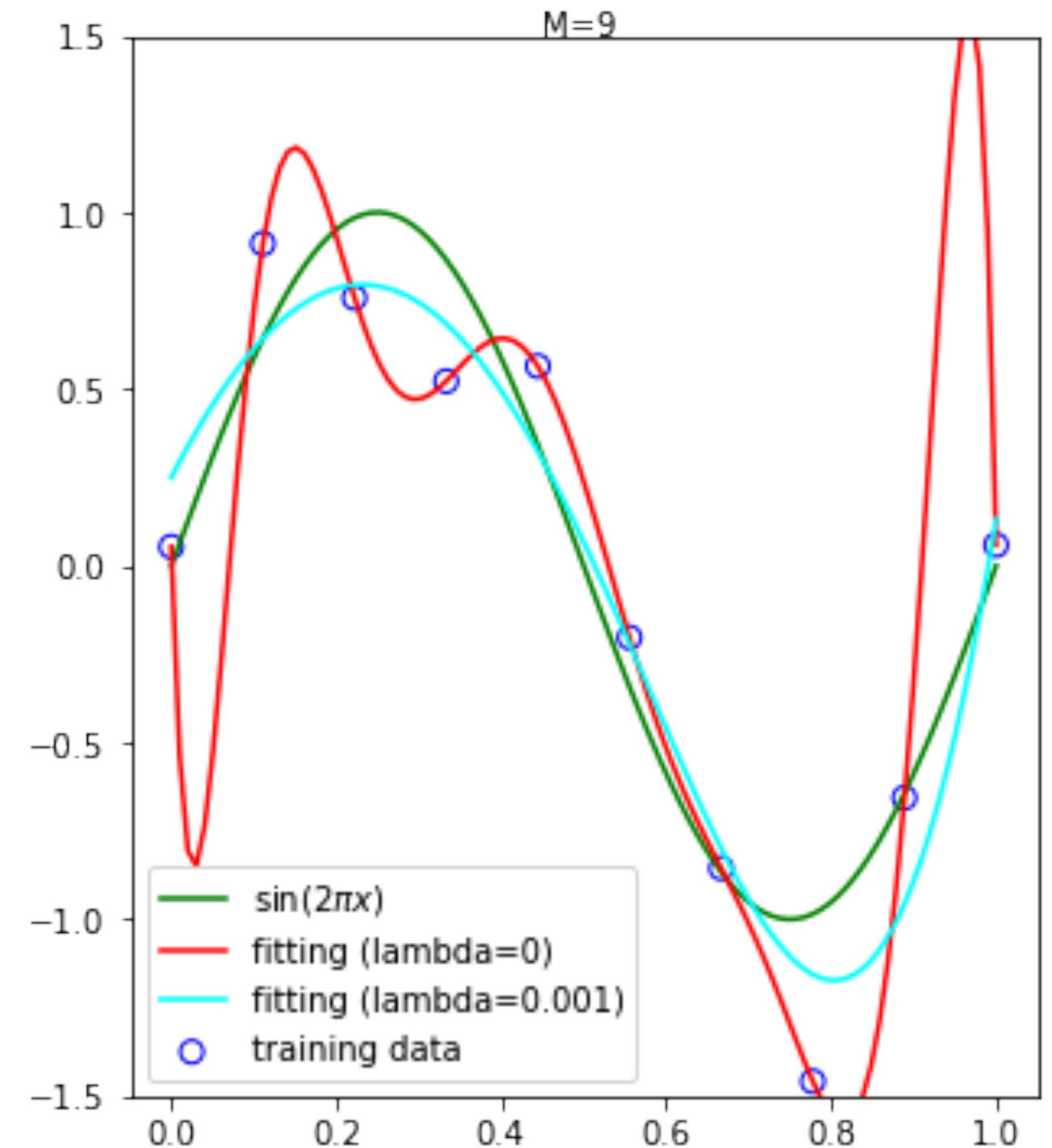
Model selection

Overfitting and regularization

Polynomials are not universal approximators

But Artificial Neural Networks (ANNs) are

Subject of today



GLOSSARY

(Mini-)Batch

Number of samples processed before each update

SGD

Stochastic Gradient Descent: Iterative gradient-based optimization using random batches

Hyperparameters

Settings you choose before training (e.g. learning rate, batch size)

Loss

Measure of how wrong the model's predictions are ("Error"/"Objective")

Overfitting

Model learns specifics of the training set (noise) instead of general patterns

Regularisation

Methods to keep models simpler and prevent overfitting

Train/Val/Test split

Separate data for fitting, tuning, and evaluating

ARTIFICIAL NEURAL NETWORKS

ARTIFICIAL NEURAL NETWORKS (ANNs)

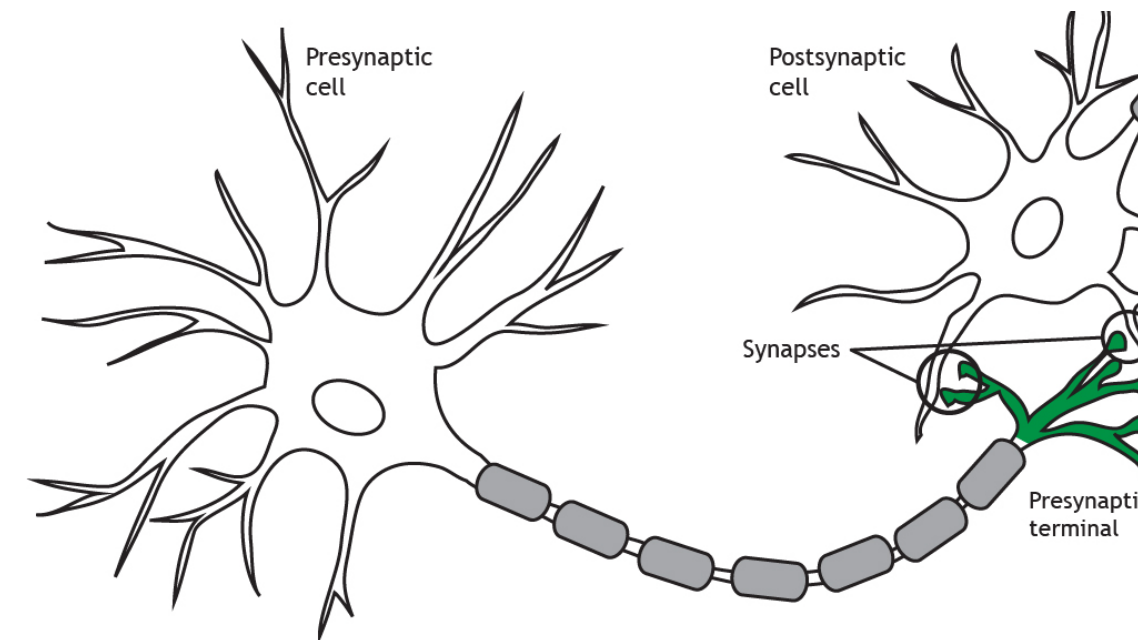
ANNs ?= Machine Learning

Historically no, today kinda yes

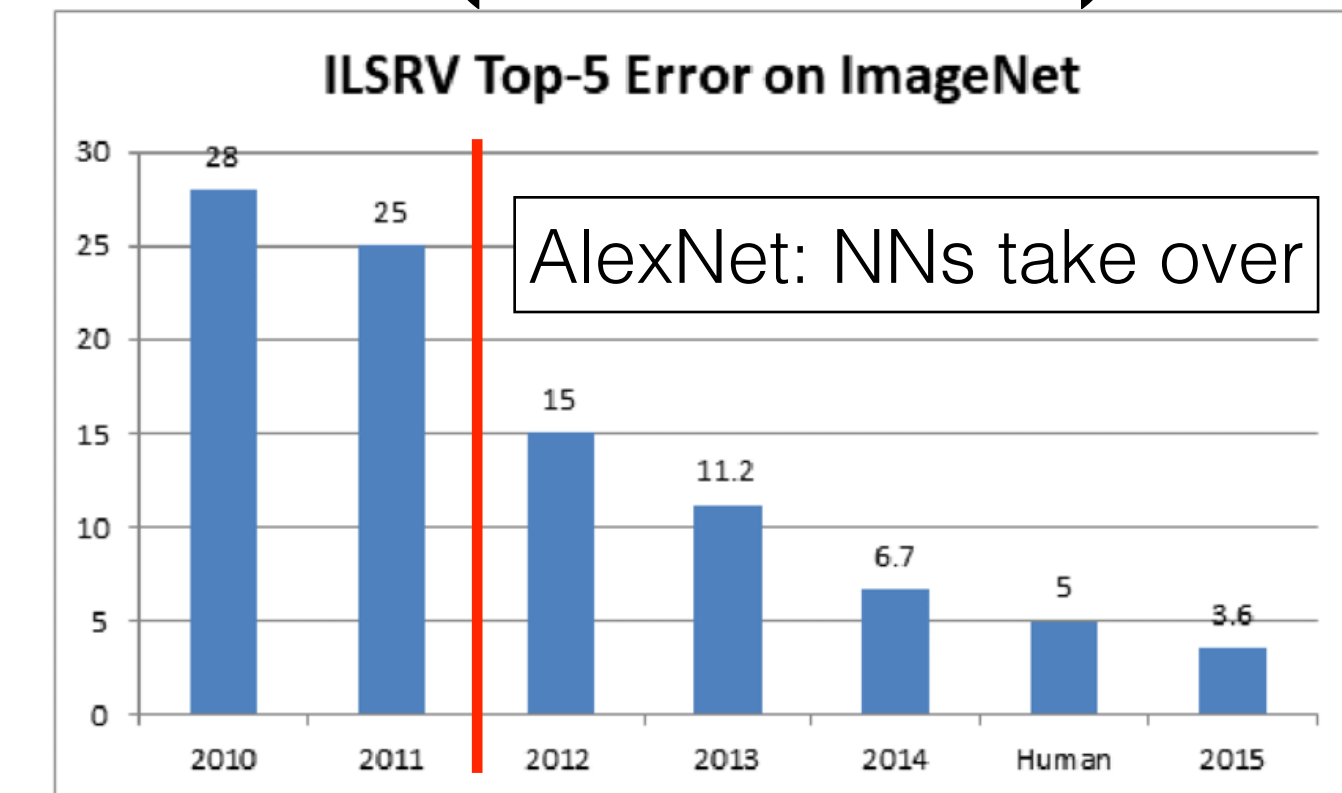
Kind of inspired by biology

!= spiking neural networks

c.f. “non-differentiable”



Biological Synapse Structure



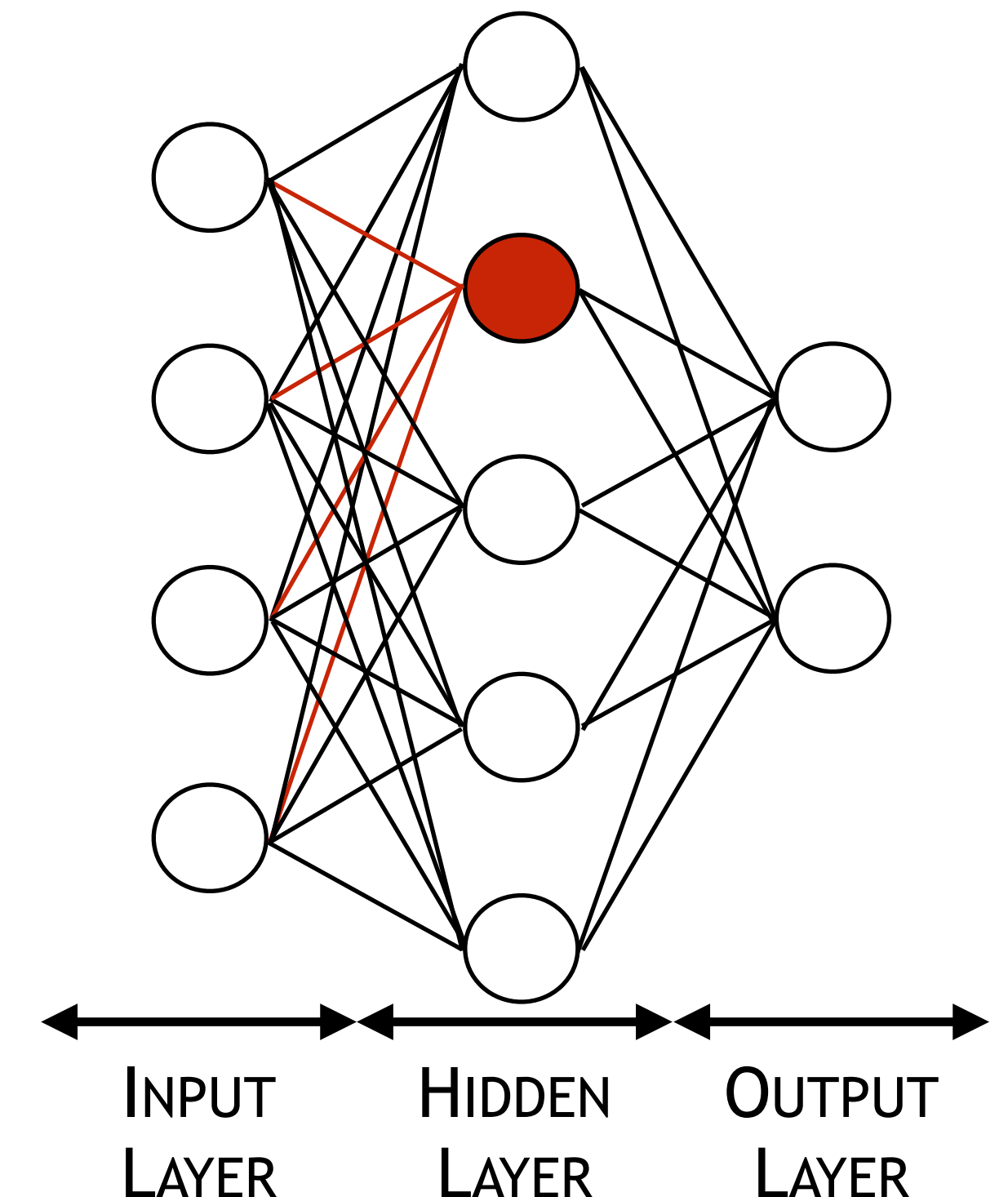
More complex problems require more complex models

Informal term of “model capacity” (or expressivity)

ANNs are universal function approximators

Universal approximation theorem(s):

Neural networks can represent a wide variety of interesting functions with appropriate weights



MULTI-LAYER PERCEPTRON (MLP)

For neuron k of a given layer

$$y_k = f\left(\sum_j (w_{k,j} \cdot x_j) + b_k\right)$$

f : non-linear function
(sigmoid, reLU, ...)

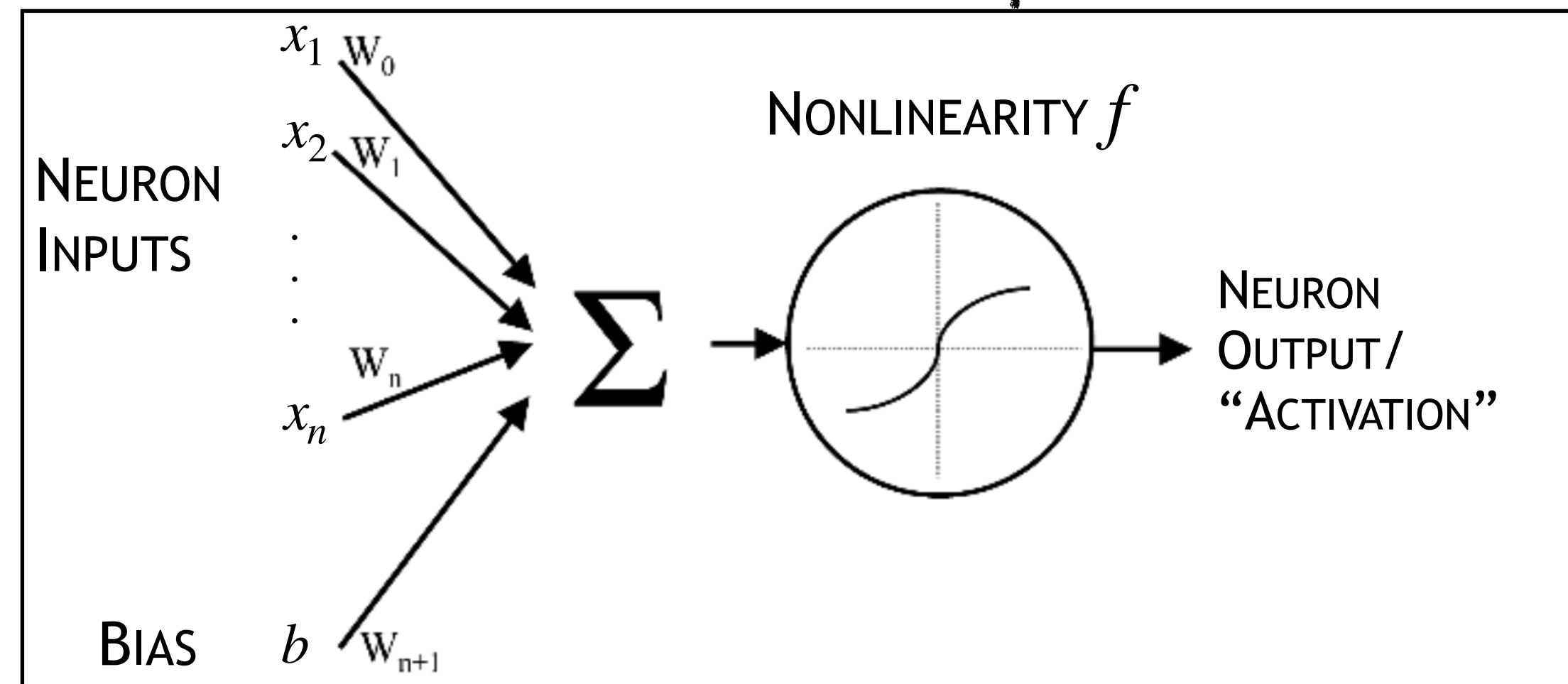
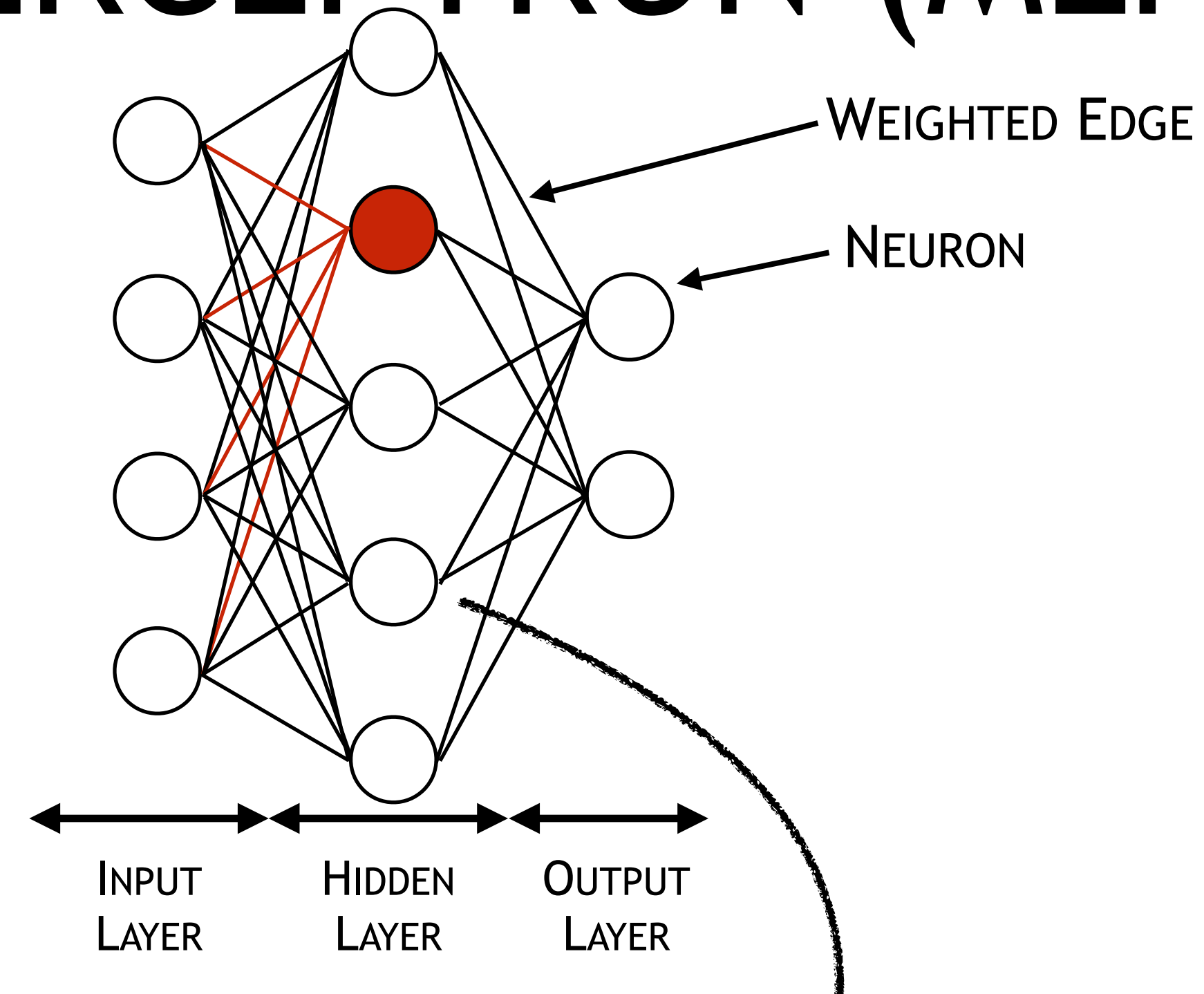
$\mathbf{W}$: weight matrix

$\mathbf{x}$: activation vector

$\mathbf{b}$: bias vector (hidden)

Vector notation for layer l

$$\mathbf{x}_l = f(\mathbf{W}_l \cdot \mathbf{x}_{l-1} + \mathbf{b}_l)$$



INTERLUDE: INTUITIVE VIEW

Learning logical operations

$$\text{AND: } x_1 + x_2 > 1.5 \rightarrow x_1 + x_2 - 1.5 = 0$$

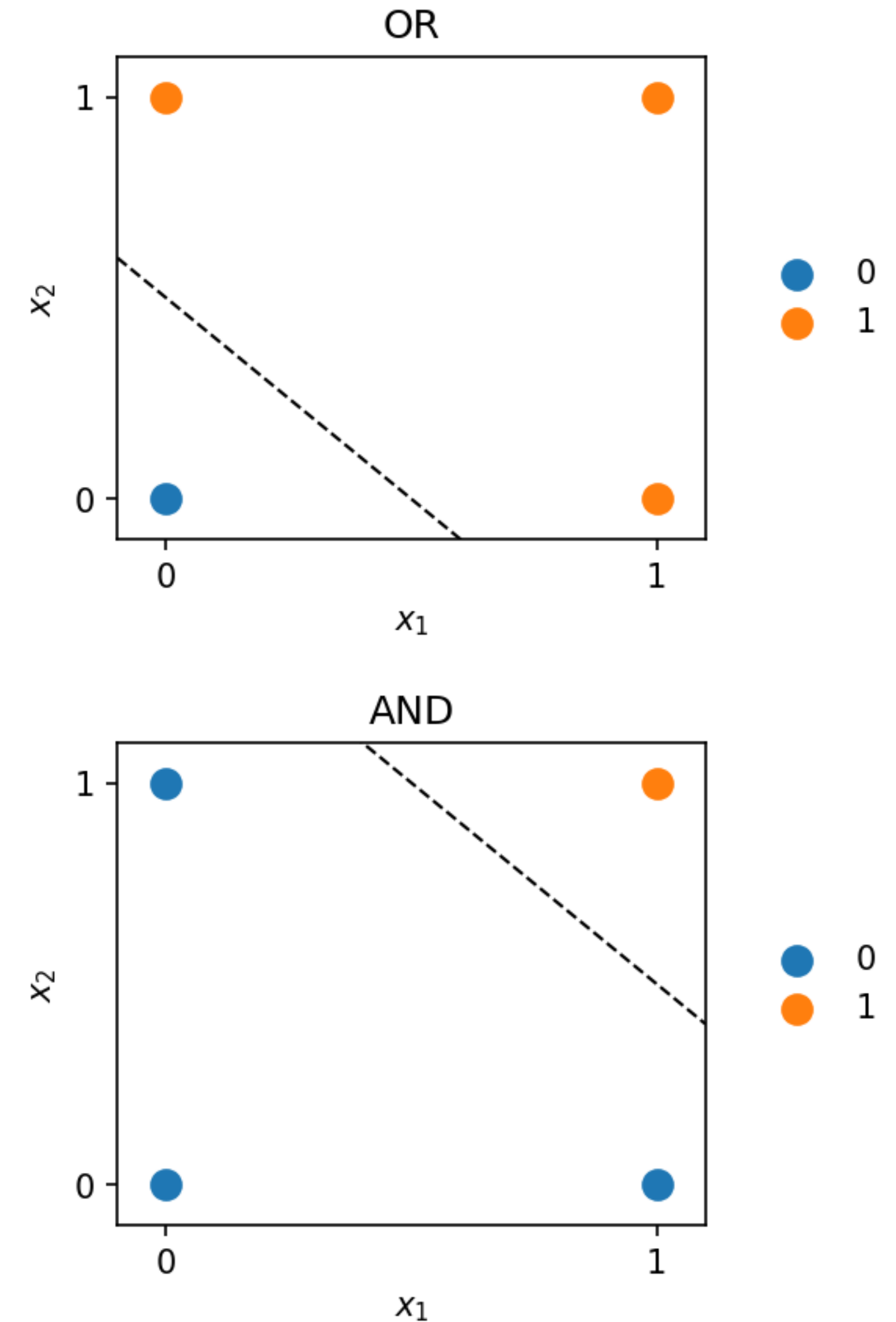
$$\text{OR: } x_1 + x_2 > 0.5 \rightarrow x_1 + x_2 - 0.5 = 0$$

A single neuron can do the same

$$y = f(w_1x_1 + w_2x_2 + b)$$

$$\text{AND: } w_1 = w_2 = 1, b = -1.5$$

$$\text{OR: } w_1 = w_2 = 1, b = -0.5$$



INTERLUDE: INTUITIVE VIEW

Learning XOR

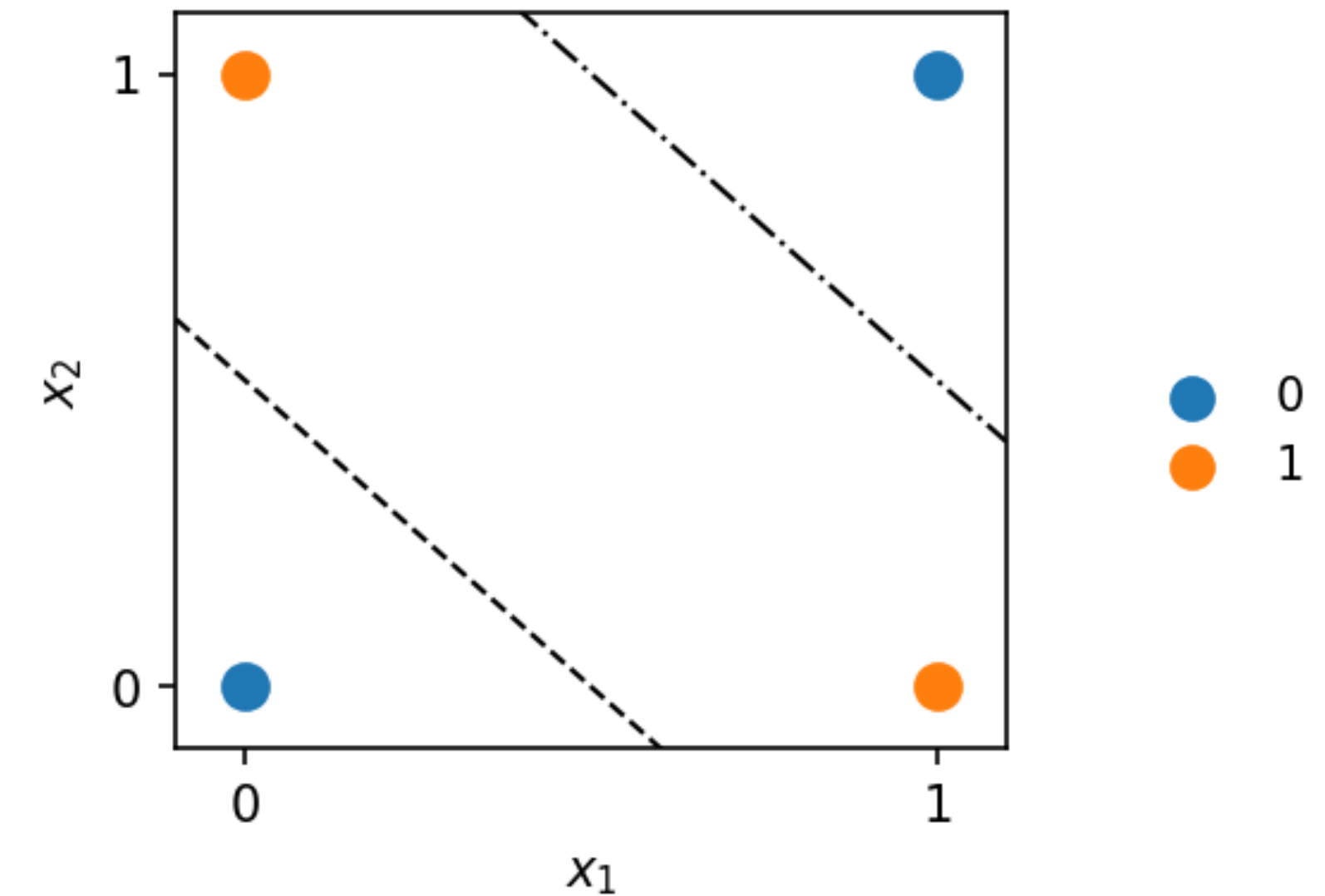
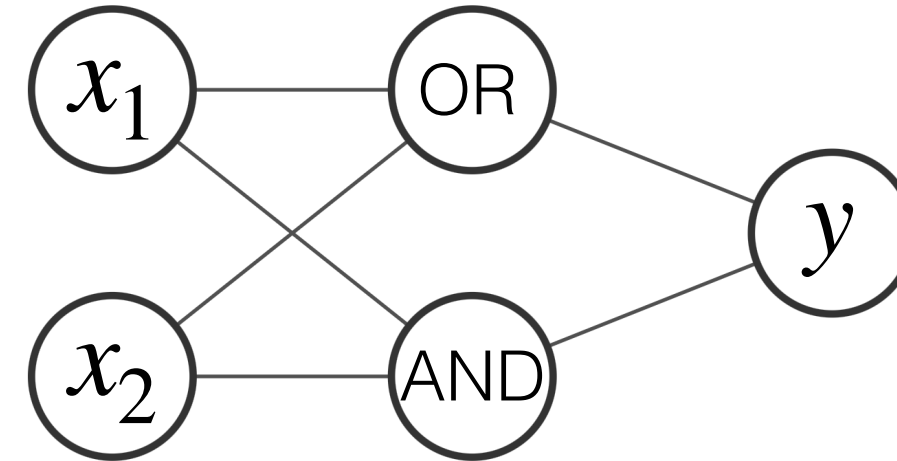
XOR: Not Linearly Separable

But two (three) neurons can still solve the problem:

$$y_{\text{OR}} = f(w_{11}x_1 + w_{12}x_2 + b_1)$$

$$y_{\text{AND}} = f(w_{21}x_1 + w_{22}x_2 + b_2)$$

$$y = f(w_{\text{OR}}y_{\text{OR}} + w_{\text{AND}}y_{\text{AND}} + b)$$



Nonlinearities are essential

Without them, this is just a sum of linear functions (still linear)

EXAMPLE NONLINEARITIES

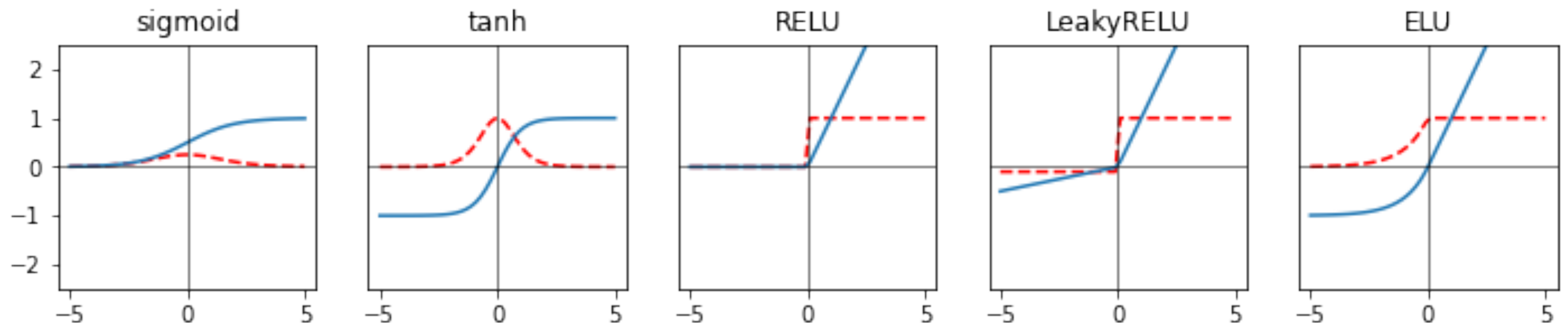
sigmoid: $f(x) = \frac{1}{1 + e^{-x}}$ $\Rightarrow$ output in range $[0,1]$

tanh: $f(x) = \frac{e^x - e^{-x}}{e^x + e^{-x}}$ $\Rightarrow$ output in range $[-1,1]$

ReLU: $f(x) = \max(x,0)$ $\Rightarrow$ no negative output

LeakyReLU: $f(x) = \begin{cases} x; x \geq 0 \\ \alpha x; x < 0 \end{cases}$ $\Rightarrow$ no clamping to zero for negative inputs

ELU: $f(x) = \begin{cases} x; x \geq 0 \\ e^x - 1; x < 0 \end{cases}$ $\Rightarrow$ smoother gradient



Basically any non-linear function can be used

INTERLUDE: INTUITIVE VIEW

XOR: Why not use a polynomial again?

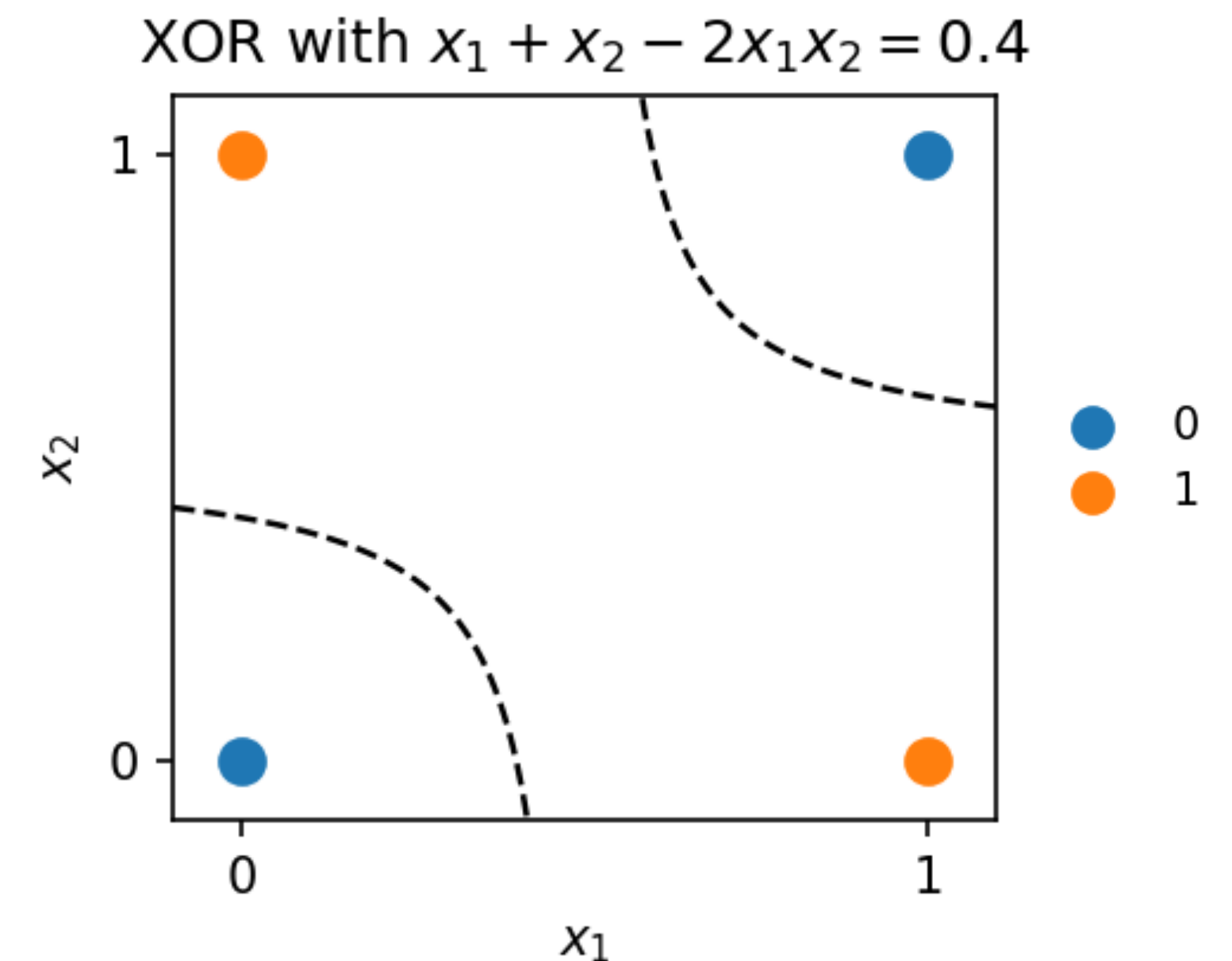
E.g.: $x_1 + x_2 - 2x_1x_2 = 0.4$ works!

But: Combinatorial scaling, $N_{\text{terms}} \sim \binom{n_{\text{data}} + d_{\text{poly}}}{d_{\text{poly}}}$

Curse of dimensionality

One pixel = one dimension

MNIST: $28 \times 28 \text{px} \rightarrow n_{\text{data}} = 784$



Dims	Terms
n=3, d=3	20
n=5, d=3	56
n=10, d=3	286
n=100, d=3	~160k

INTERLUDE: INTUITIVE VIEW

High dimensions are weird: In high d ...

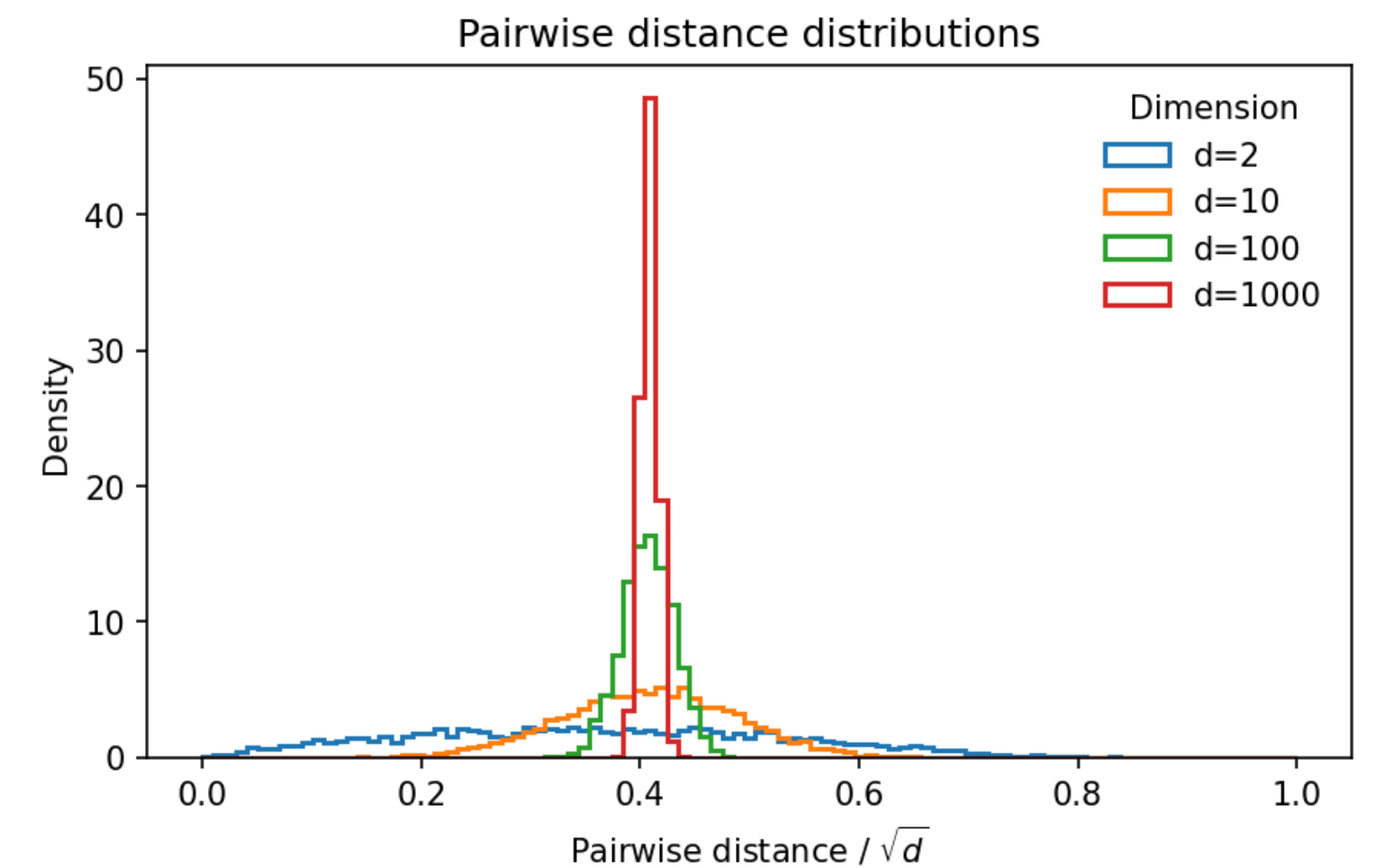
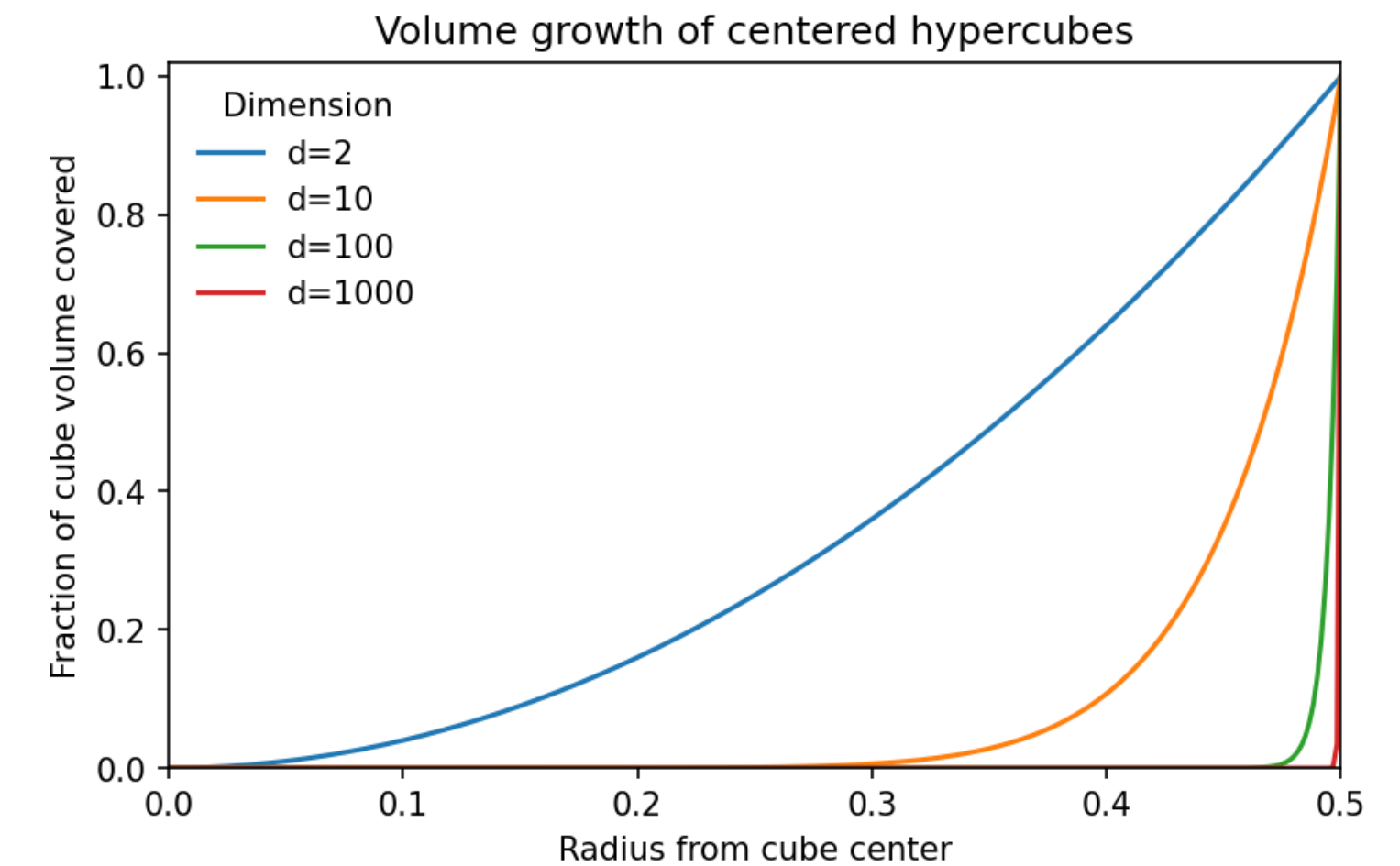
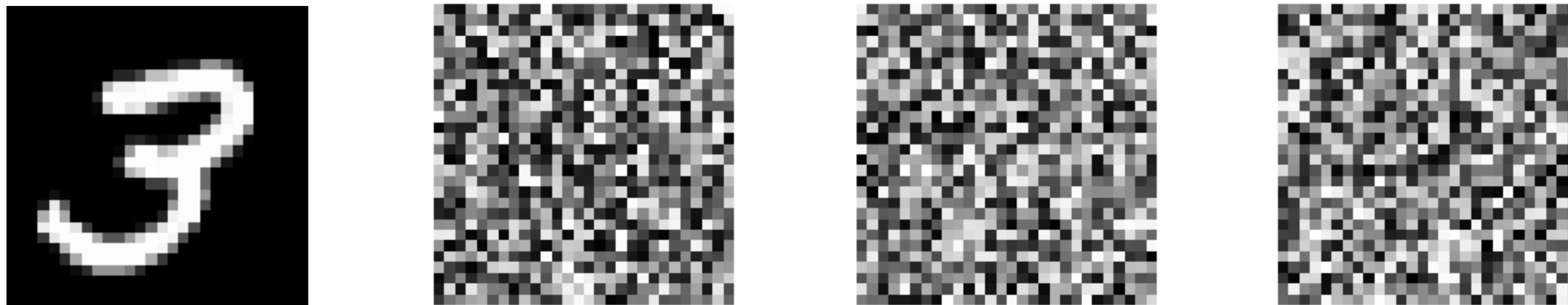
almost all of the volume of a hypersphere/cube is close to its surface

All points are equally “far away”

Our only hope: Lower-dimensional manifolds

TL;DR: Neural networks are “smart mappers”

Random samples in 784-D (uniform[0,1])



VECTOR AND MATRIX NOTATION

Matrix **W**, composed of elements $w_{k,j}$

Matrix = bold uppercase

Matrix element $w_{k,j}$ has row k , column j

Vector **x**, composed of elements x_i

Vector = bold lowercase

Vectors are vertical, use $\mathbf{x}^T$ for horizontal vectors

Matrix-vector multiplication

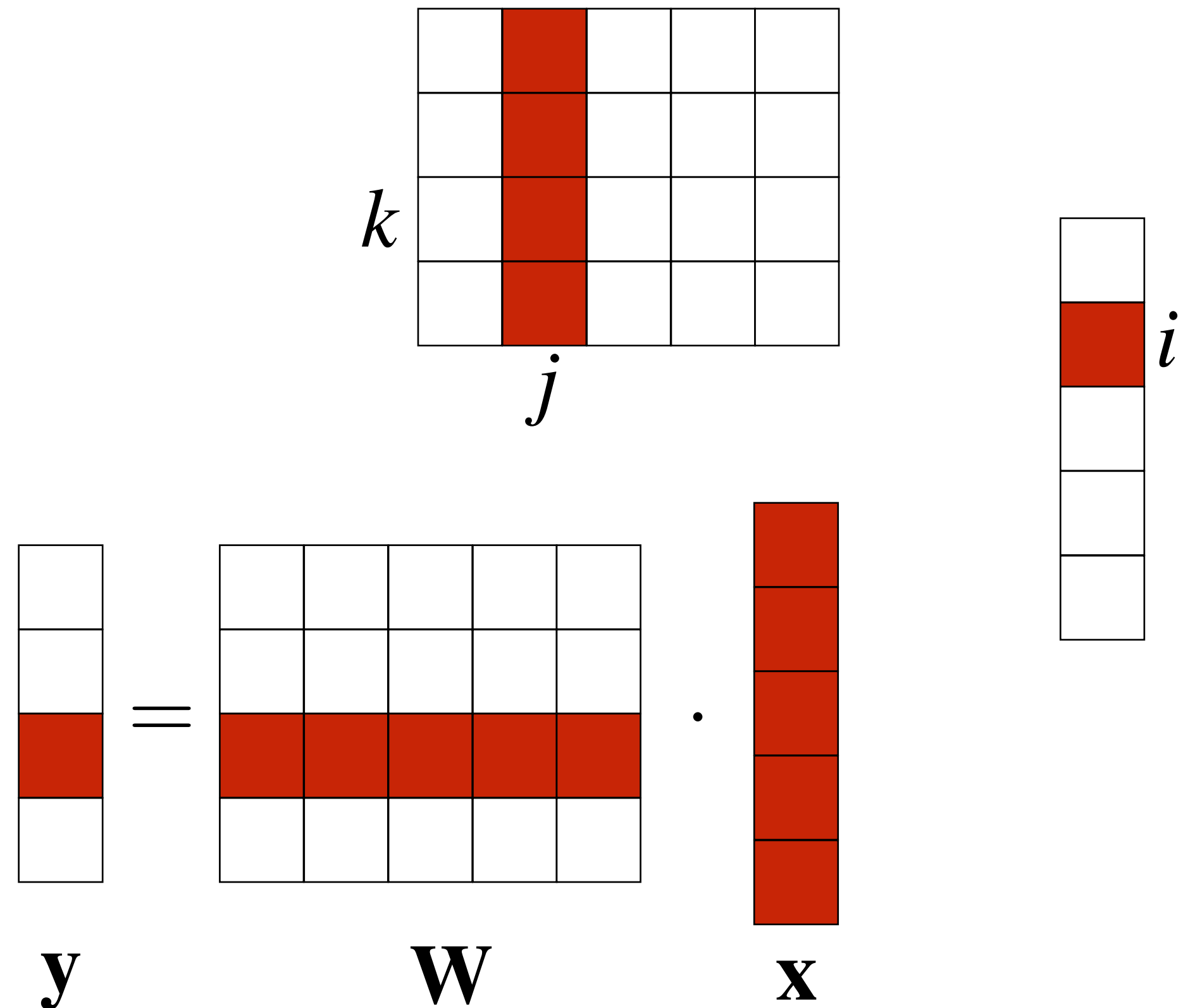
Length of the vector equals the number of columns of the matrix

$$y_k = \sum_j (w_{k,j} \cdot x_j), \text{ resp. } \mathbf{y} = \mathbf{W} \cdot \mathbf{x}$$

Vector-vector multiplication

$$\text{Dot product: } a = \sum_j (b_j \cdot c_j), \text{ resp. } a = \mathbf{b} \cdot \mathbf{c} = \mathbf{c} \cdot \mathbf{b} = \mathbf{b}^T \mathbf{c}$$

$$\text{Outer product: } a = \mathbf{c} \otimes \mathbf{b} = \mathbf{bc}^T$$



MULTI-LAYER PERCEPTRON (MLP)

For neuron k of a given layer

$$y_k = f\left(\sum_j (w_{k,j} \cdot x_j) + b_k\right)$$

f : non-linear function
(sigmoid, reLU, ...)

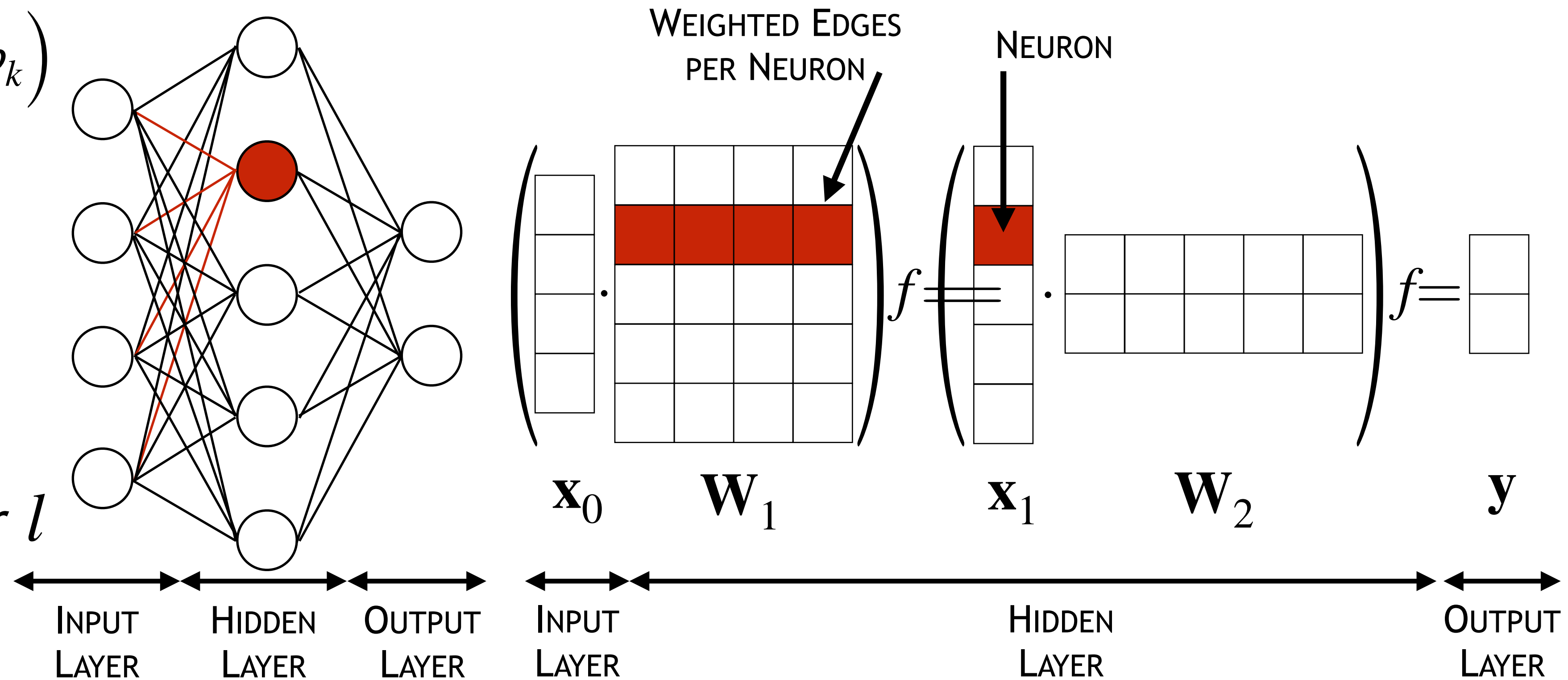
$\mathbf{W}$: weight matrix

$\mathbf{x}$: activation vector

$\mathbf{b}$: bias vector (hidden)

Vector notation for layer l

$$\mathbf{x}_l = f(\mathbf{W}_l \cdot \mathbf{x}_{l-1} + \mathbf{b}_l)$$



FORWARD PROP ON ONE SLIDE

(Deep) Neural Networks: L stacked processing units, where each unit computes an activation function

$x_l = f(\mathbf{W}_l \oplus \mathbf{x}_{l-1} + \mathbf{b}_l)$, for nonlinear activation function $f(\cdot)$, linear operation $\oplus$, weight matrix $\mathbf{W}$, input activations $\mathbf{x}$, and bias $\mathbf{b}$ of layer l

Bias vector $\mathbf{b}$ is usually encoded in the weight matrix $\mathbf{W}$ by introducing another constant activation element (e.g., $x_0 = 1$)

Then a complete MLP with L layers is

$$\mathbf{y}(\mathbf{W}, \mathbf{x}_0) = \mathbf{x}_L = \mathbf{W}_L \oplus f(\mathbf{W}_{L-1} \oplus f(\dots \oplus f(\mathbf{W}_1 \oplus \mathbf{x}_0)) \dots)$$

LAST LAYER FOR CLASSIFICATION TASKS

For classification tasks: output of last layer uncalibrated and difficult to interpret

Softmax provides us with a multinomial categorical output

See logistic function for binomial categories

Softmax: $s(x) = \frac{e^{x_i}}{\sum_{j=1}^K e^{x_j}}$, for K classes

Usually used in classification tasks to normalize the set of output activations to 1

Do not confuse softmax with a mathematically sound probability

Sometimes numerically unstable (divisions by large numbers)

$$\frac{e^{x_i}}{\sum_{j=1}^K e^{x_j}} = \frac{C e^{x_i}}{C \sum_{j=1}^K e^{x_j}} = \frac{e^{x+\log C}}{\sum_{j=1}^K e^{x_j + \log C}}, \text{ for } \log C = -\max_j x_j$$

See also log-softmax

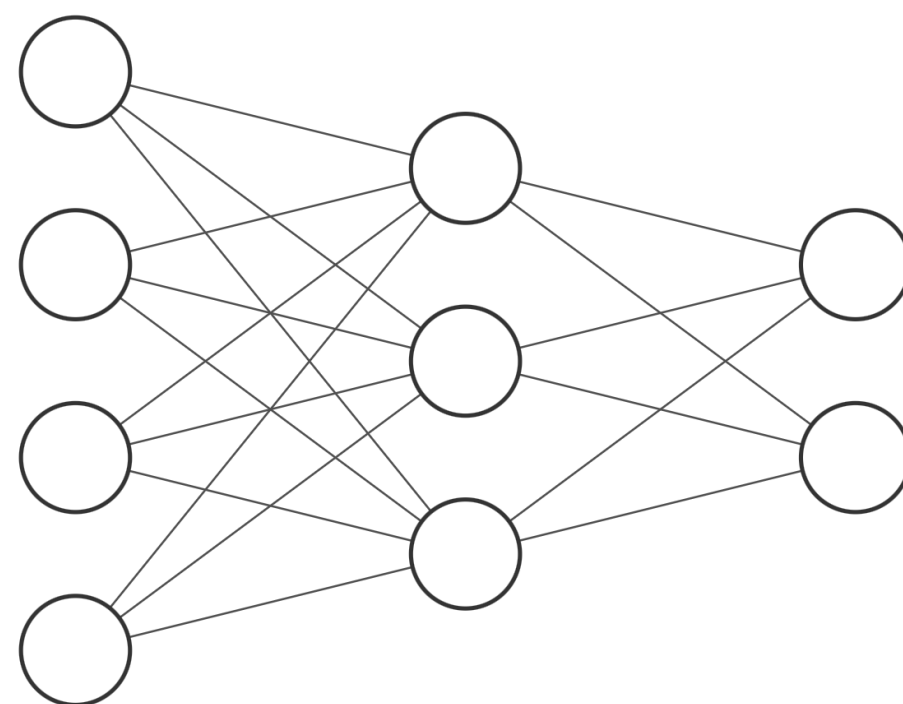
x	argmax(x)	softmax(x)	norm(x)
1	0,0 %	7,0 %	13,1 %
0,1	0,0 %	2,9 %	1,3 %
3	100,0 %	51,9 %	39,2 %
0,1	0,0 %	2,9 %	1,3 %
0,4	0,0 %	3,9 %	5,2 %
2	0,0 %	19,1 %	26,1 %
1	0,0 %	7,0 %	13,1 %
0,05	0,0 %	2,7 %	0,7 %
0	0,0 %	2,6 %	0,0 %

CALCULATING A NEURAL NETWORK

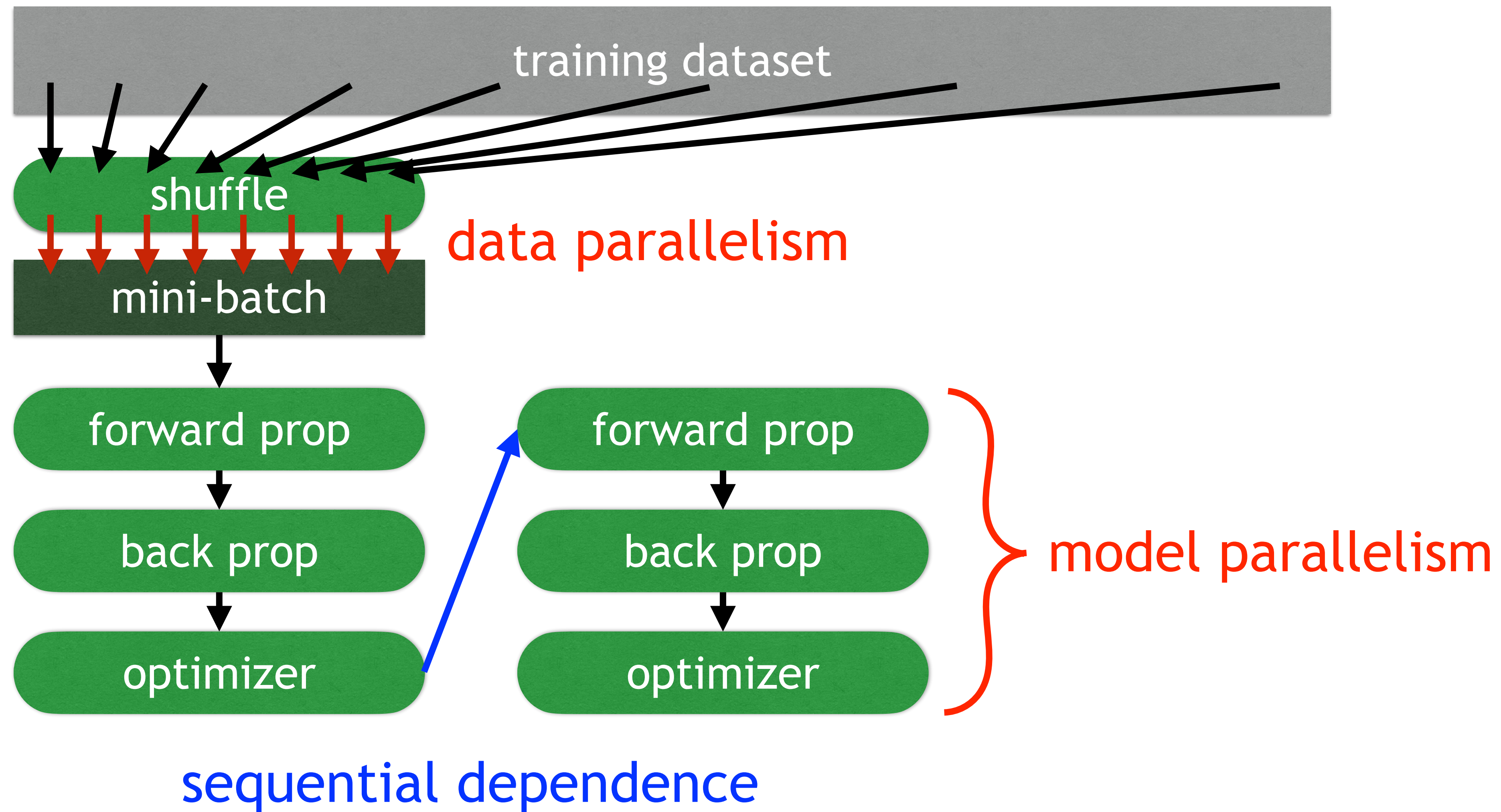
Forward pass

$$\begin{array}{c}
 \mathbf{x}^{0T} \\
 \begin{array}{|c|c|c|c|}
 \hline
 x_0 & x_1 & x_2 & x_3 \\
 \hline
 \end{array}
 \end{array}
 \bullet
 \begin{array}{c}
 \mathbf{W}^0 \\
 \begin{array}{|c|c|c|}
 \hline
 w_{0,0} & w_{0,1} & w_{0,2} \\
 \hline
 w_{1,0} & w_{1,1} & w_{1,2} \\
 \hline
 w_{2,0} & w_{2,1} & w_{2,2} \\
 \hline
 w_{3,0} & w_{3,1} & w_{3,2} \\
 \hline
 \end{array}
 \end{array}
 =
 \begin{array}{|c|c|c|}
 \hline
 x_0 & x_1 & x_2 \\
 \hline
 \end{array}
 +
 \begin{array}{|c|c|c|}
 \hline
 b_0 & b_1 & b_2 \\
 \hline
 \end{array}
 \begin{array}{c}
 \mathbf{b}^{0T} \\
 \uparrow \\
 \begin{array}{|c|c|c|}
 \hline
 x_0 & x_1 & x_2 \\
 \hline
 \end{array}
 \end{array}
 \begin{array}{c}
 \sigma(\mathbf{x}) = \frac{1}{1 + e^{-x}} \\
 \uparrow
 \end{array}
 =
 \begin{array}{c}
 \begin{array}{|c|c|c|}
 \hline
 x_0 & x_1 & x_2 \\
 \hline
 \end{array} \\
 \mathbf{x}^{1T}
 \end{array}
 \bullet
 \begin{array}{c}
 \mathbf{W}^1 \\
 \begin{array}{|c|c|}
 \hline
 w_{0,0} & w_{0,1} \\
 \hline
 w_{1,0} & w_{1,1} \\
 \hline
 w_{2,0} & w_{2,1} \\
 \hline
 \end{array}
 \end{array}
 =
 \begin{array}{|c|c|}
 \hline
 b_0 & b_1 \\
 \hline
 \end{array}
 +
 \begin{array}{|c|c|}
 \hline
 x_0 & x_1 \\
 \hline
 \end{array}
 \begin{array}{c}
 \mathbf{b}^{1T} \\
 \downarrow \\
 \begin{array}{|c|c|}
 \hline
 y_0 & y_1 \\
 \hline
 \end{array}
 \end{array}
 =
 \begin{array}{|c|c|}
 \hline
 y_0 & y_1 \\
 \hline
 \end{array}
 \begin{array}{c}
 \mathcal{L}(\hat{\mathbf{y}}, \mathbf{y})
 \end{array}$$

$$\mathbf{y} = s(\mathbf{W}^1(\sigma(\mathbf{W}^0\mathbf{x}^0 + \mathbf{b}^0)) + \mathbf{b}^1)$$



TRAINING OF DEEP NEURAL NETWORKS



BACK PROP ON ONE SLIDE

Data set containing N input-target pairs: $\mathcal{D} = \{(\mathbf{x}_1, t_1), \dots, (\mathbf{x}_N, t_N)\}$

Training ANNs

adjust randomly initialized weights $\mathbf{W}$ to solve a given task by minimizing a loss function $\mathcal{L}$ using gradient-based optimization

$$\mathcal{L}(\mathbf{W}; \mathcal{D}) = \sum_{n=1}^N l(y(\mathbf{W}, \mathbf{x}_n), t_n) + \lambda r(\mathbf{W})$$

data term l that penalizes wrong prediction (error function)

regularizer $r(\mathbf{W})$ such as ℓ^1 -norm or ℓ^2 -norm, trade-off hyperparameter λ

Backpropagation: compute gradient for input-target pair and minimize the loss function by iteratively calculating

$$\mathbf{W} := \mathbf{W} - \eta \nabla_{\mathbf{W}} \mathcal{L}(\mathbf{W}; \mathcal{D}), \text{ for } \nabla_{\mathbf{x}} = \left(\frac{\partial}{\partial x_1}, \dots, \frac{\partial}{\partial x_n} \right) \text{ and learning rate } \eta$$

PARTIAL DERIVATIVES SOUGHT

Consider exemplary function $f(x, y, z, w) = xy + \max(z, w)$

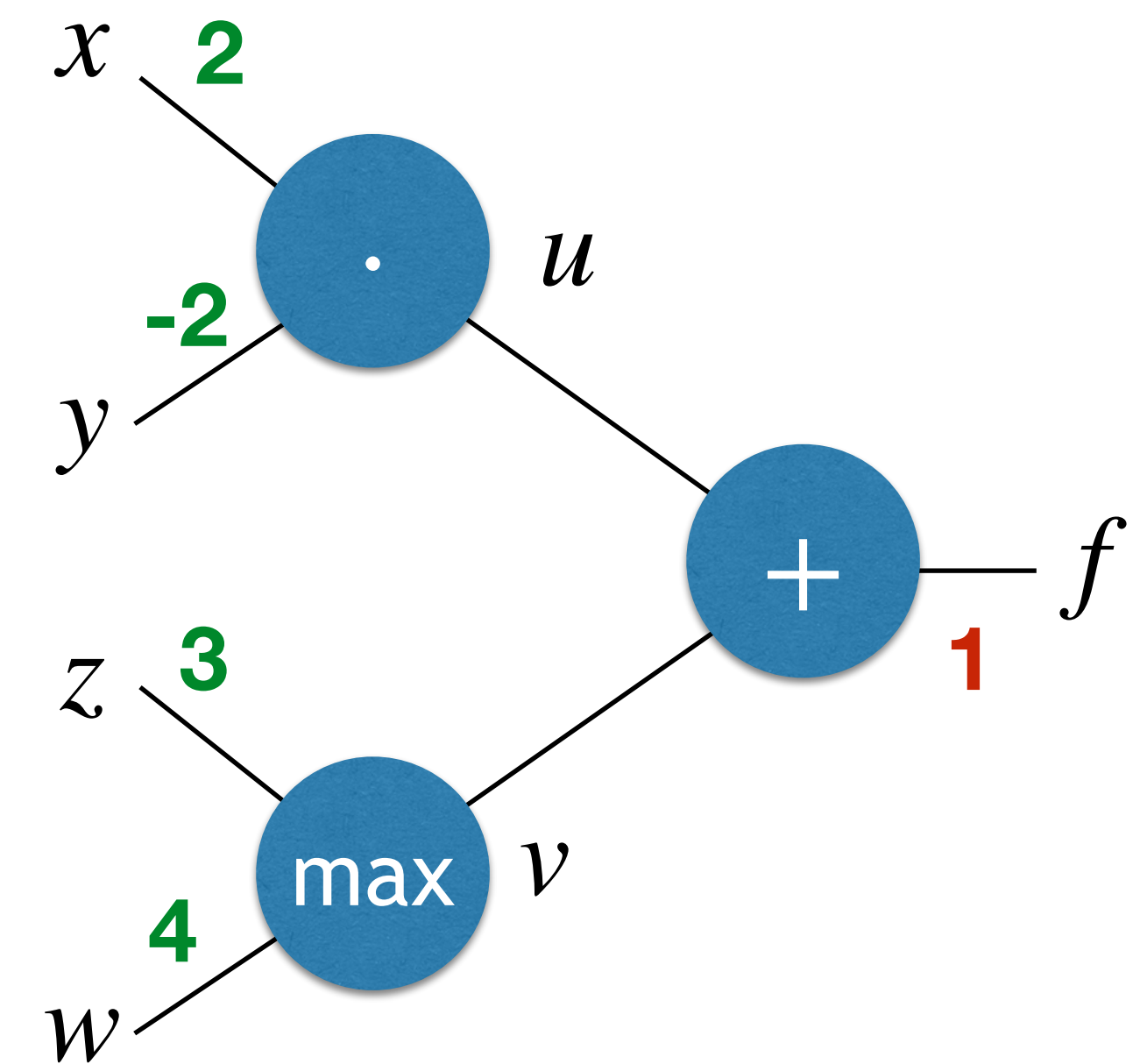
$$\frac{\partial}{\partial x} f, \frac{\partial}{\partial y} f, \frac{\partial}{\partial z} f, \frac{\partial}{\partial w} f \quad ?$$

Consider $f(x, y, z, w) = \text{add}(\text{mul}(x, y), \max(z, w))$

$$\text{add}(u, v) = u + v; \frac{\partial \text{add}}{\partial u} = \frac{\partial \text{add}}{\partial v} = 1$$

$$\text{mul}(x, y) = xy; \frac{\partial \text{mul}}{\partial x} = y; \frac{\partial \text{mul}}{\partial y} = x$$

$$\max(z, w) = \begin{cases} z; & \text{if } z \geq w \\ w; & \text{else} \end{cases}; \frac{\partial \max}{\partial z} = \begin{cases} 1; & \text{if } z \geq w \\ 0; & \text{else} \end{cases}; \frac{\partial \max}{\partial w} = \begin{cases} 0; & \text{if } z \geq w \\ 1; & \text{else} \end{cases}$$



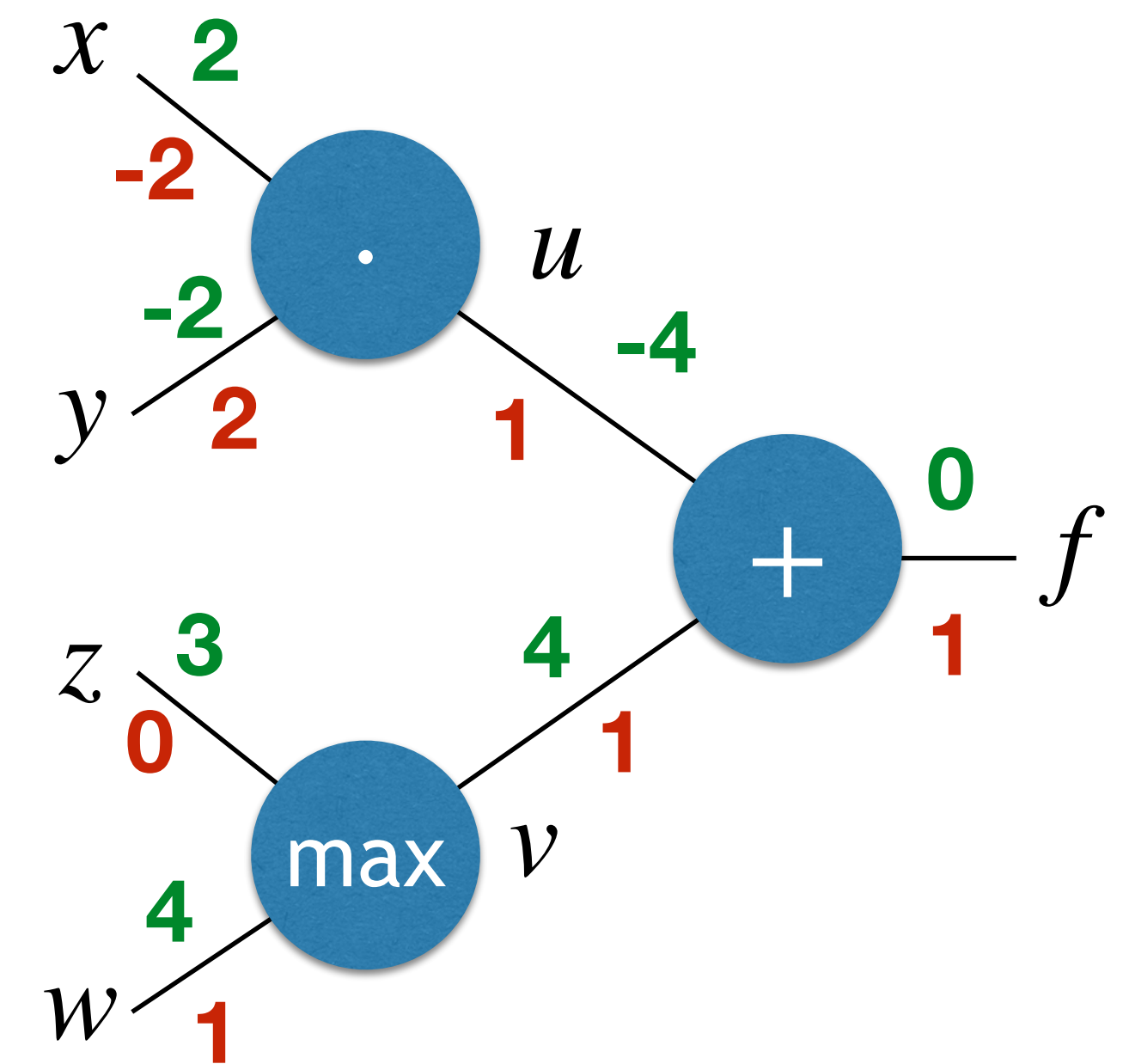
CHAIN RULE TO THE RESCUE

Chain rule of calculus: $\frac{\partial f}{\partial x} = \frac{\partial f}{\partial u} \frac{\partial u}{\partial x}$

$$\frac{\partial f}{\partial u} = \frac{\partial \text{add}(u, v)}{\partial u} = 1; \frac{\partial u}{\partial x} = \frac{\partial \text{mul}(x, y)}{\partial x} = y$$

$$\Rightarrow \frac{\partial f}{\partial x} = 1 \cdot y$$

Others in an analogous manner



$$\theta := \theta - \eta \cdot \frac{\partial \mathcal{L}}{\partial \theta}$$

GRADIENT BEHAVIOR

add gate: gradient distributor

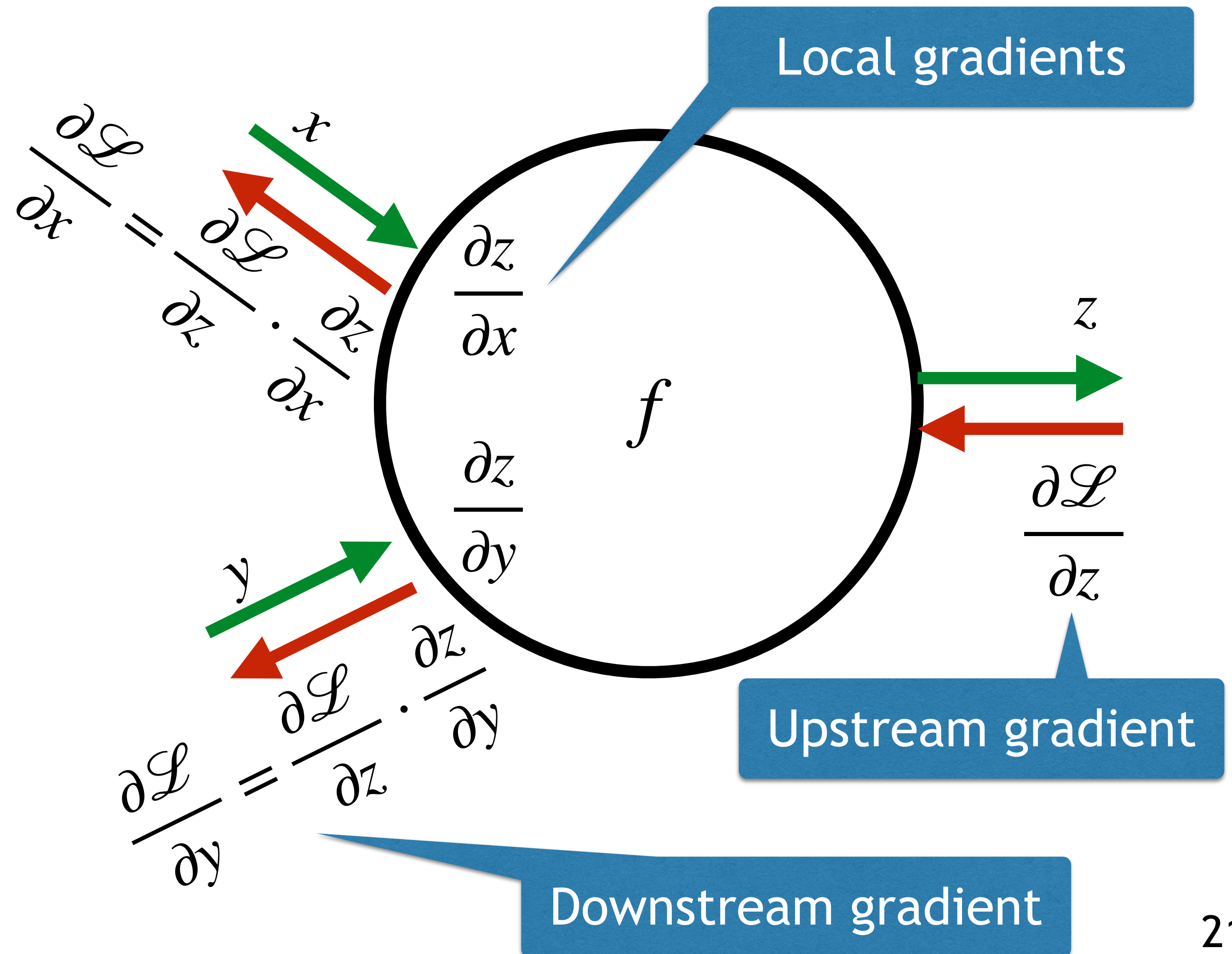
Both downstream gradients equal upstream gradient

mul gate: gradient switcher

Downstream gradients multiplied with other input

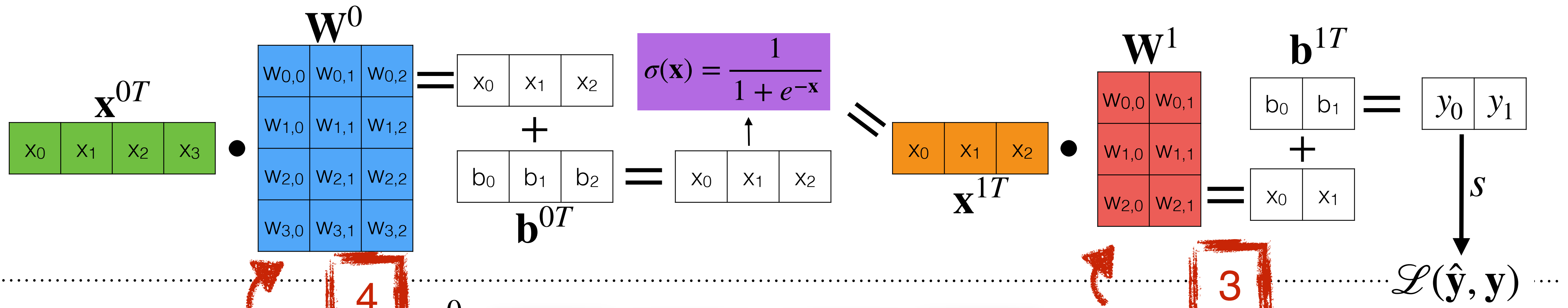
max gate: gradient router

Downstream gradients depend on comparison of inputs

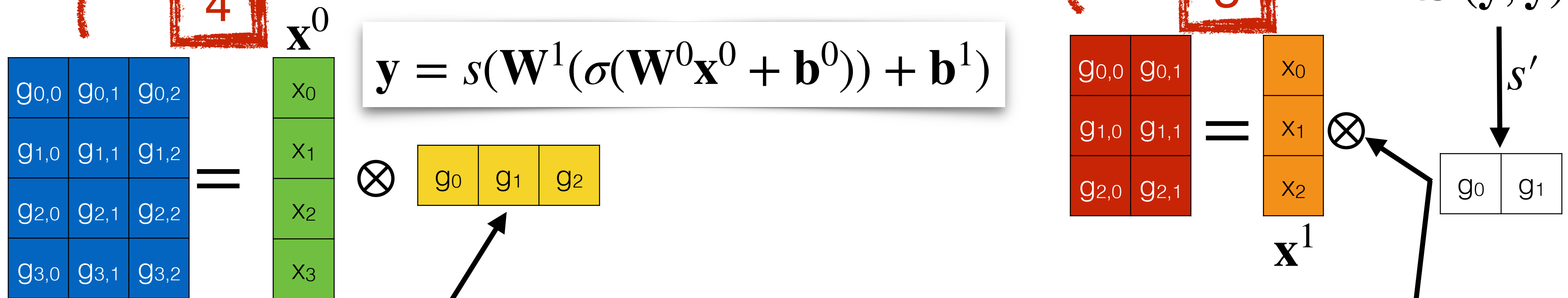


CALCULATING A NEURAL NETWORK

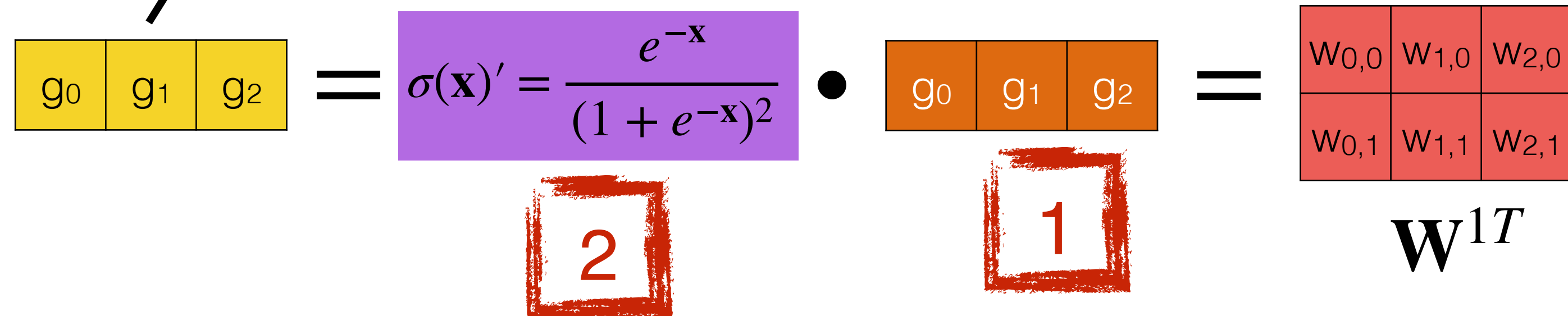
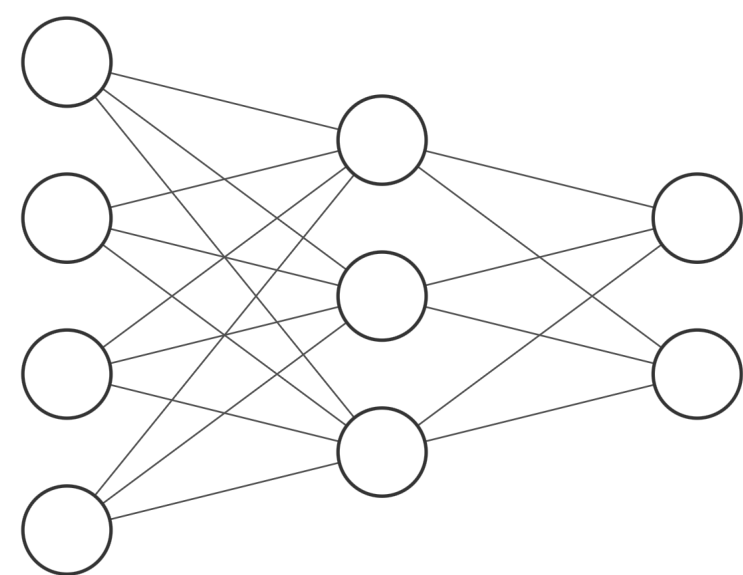
Forward pass



Backward
weight update

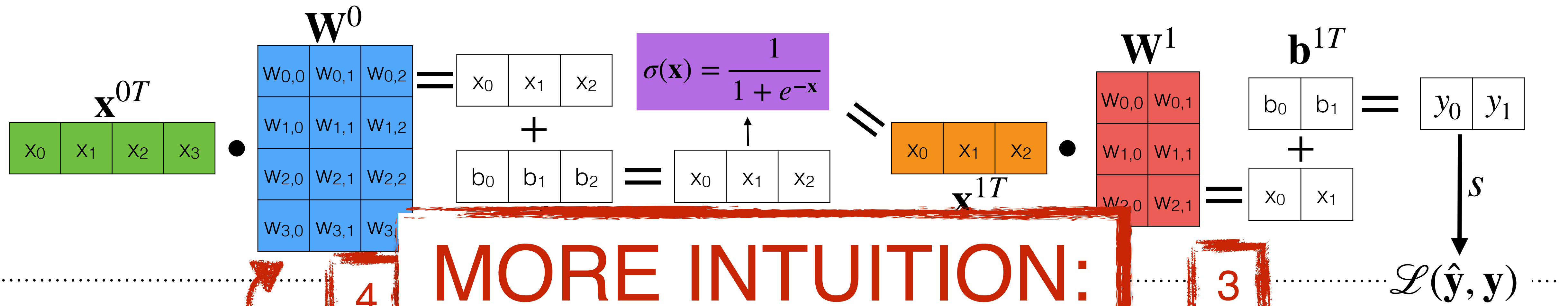


Backward
act. flow

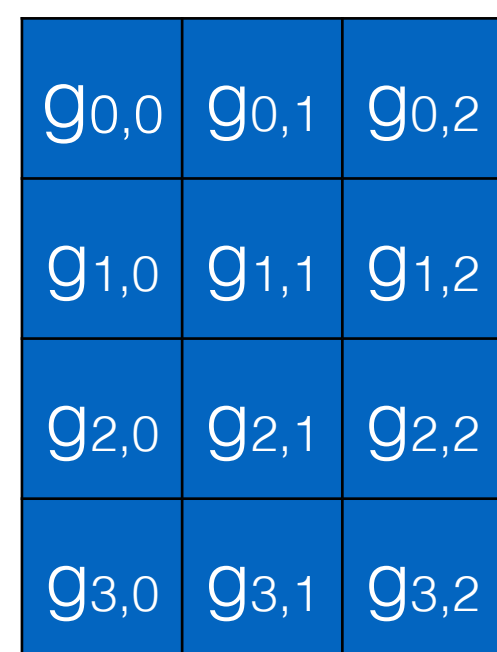


CALCULATING A NEURAL NETWORK

Forward pass

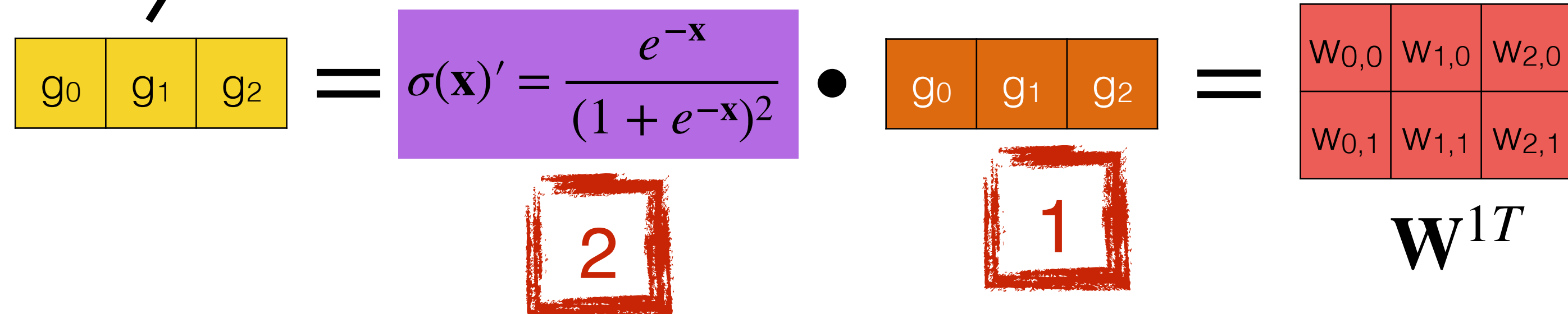
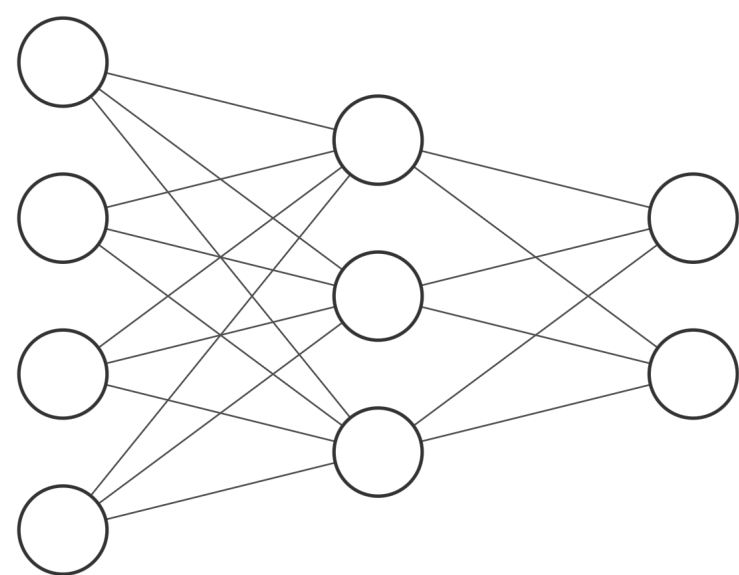


Backward
weight update



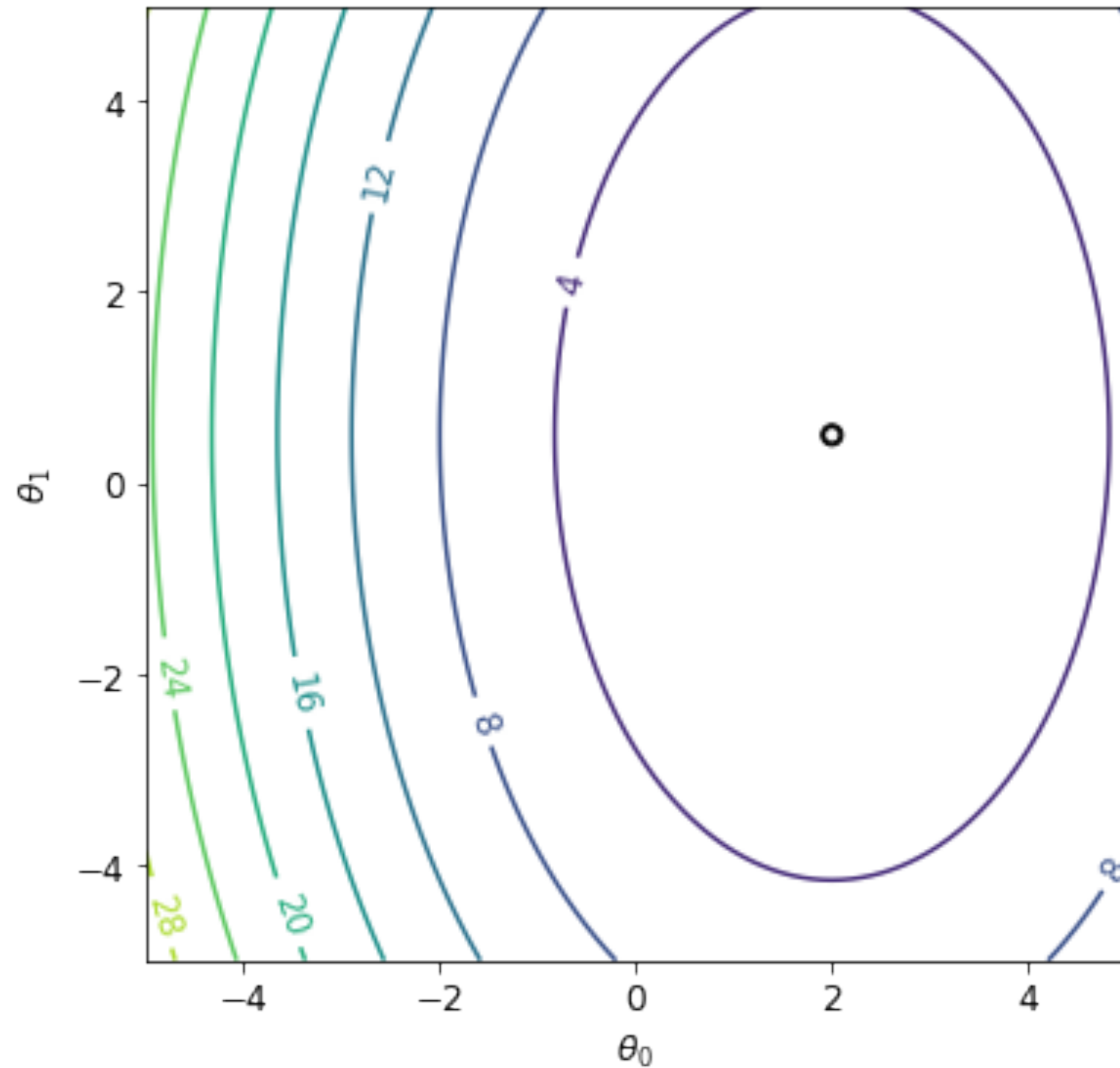
**MORE INTUITION:
3B1B NEURAL
NETWORK SERIES**

Backward
act. flow

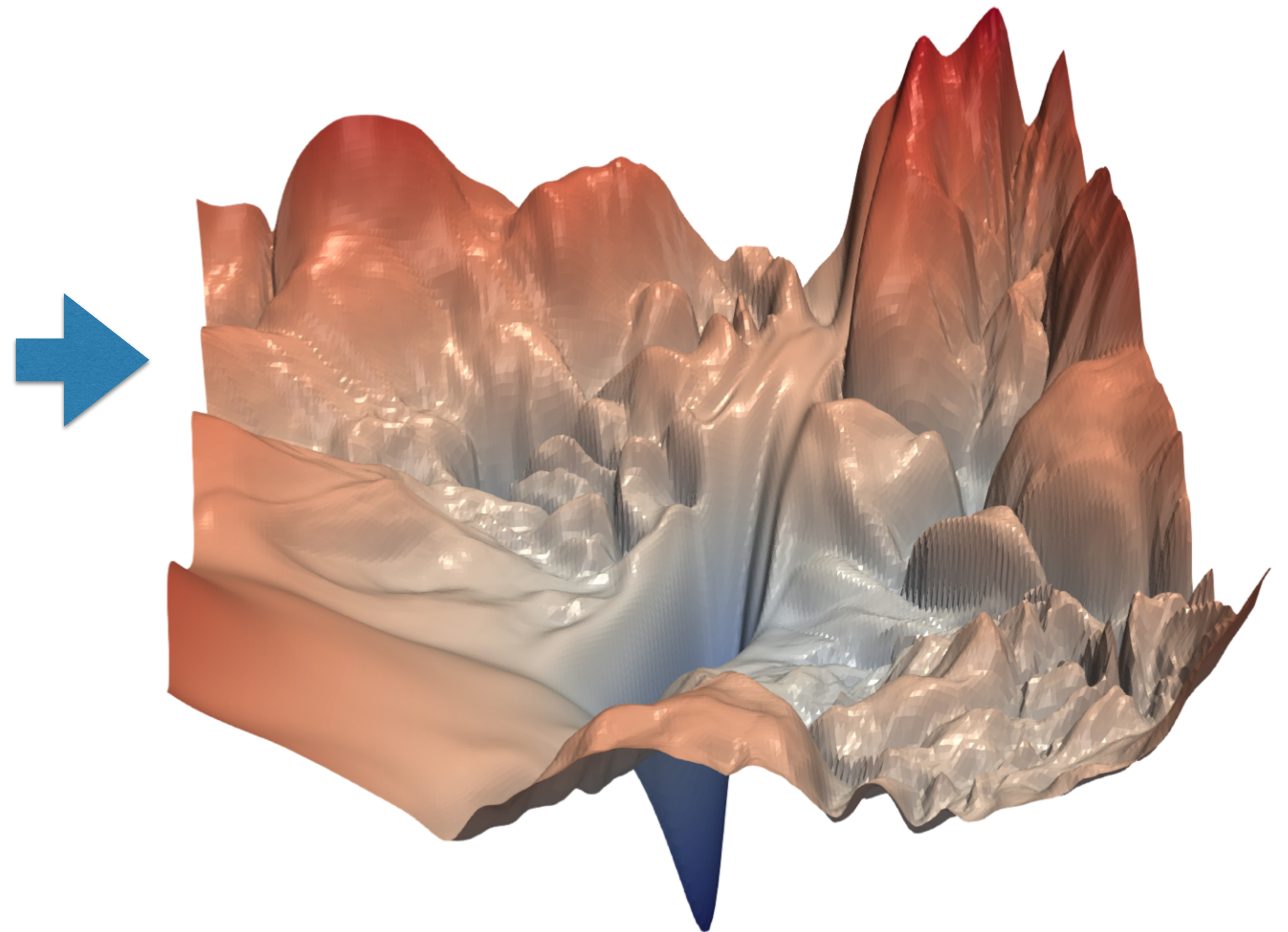


EXAMPLE LOSS LANDSCAPES IN MODERN ANNS

Convex problem



Highly non-convex problem



Hao Li, Zheng Xu, Gavin Taylor, Christoph Studer, and Tom Goldstein. 2018. Visualizing the loss landscape of neural nets. In 32nd International Conference on Neural Information Processing Systems (NIPS'18)

WRAPPING UP

SUMMARY

MLP: The prototypical neural network

ANNs are smart mappers

- Universal approximation capability

- Nonlinearities are crucial

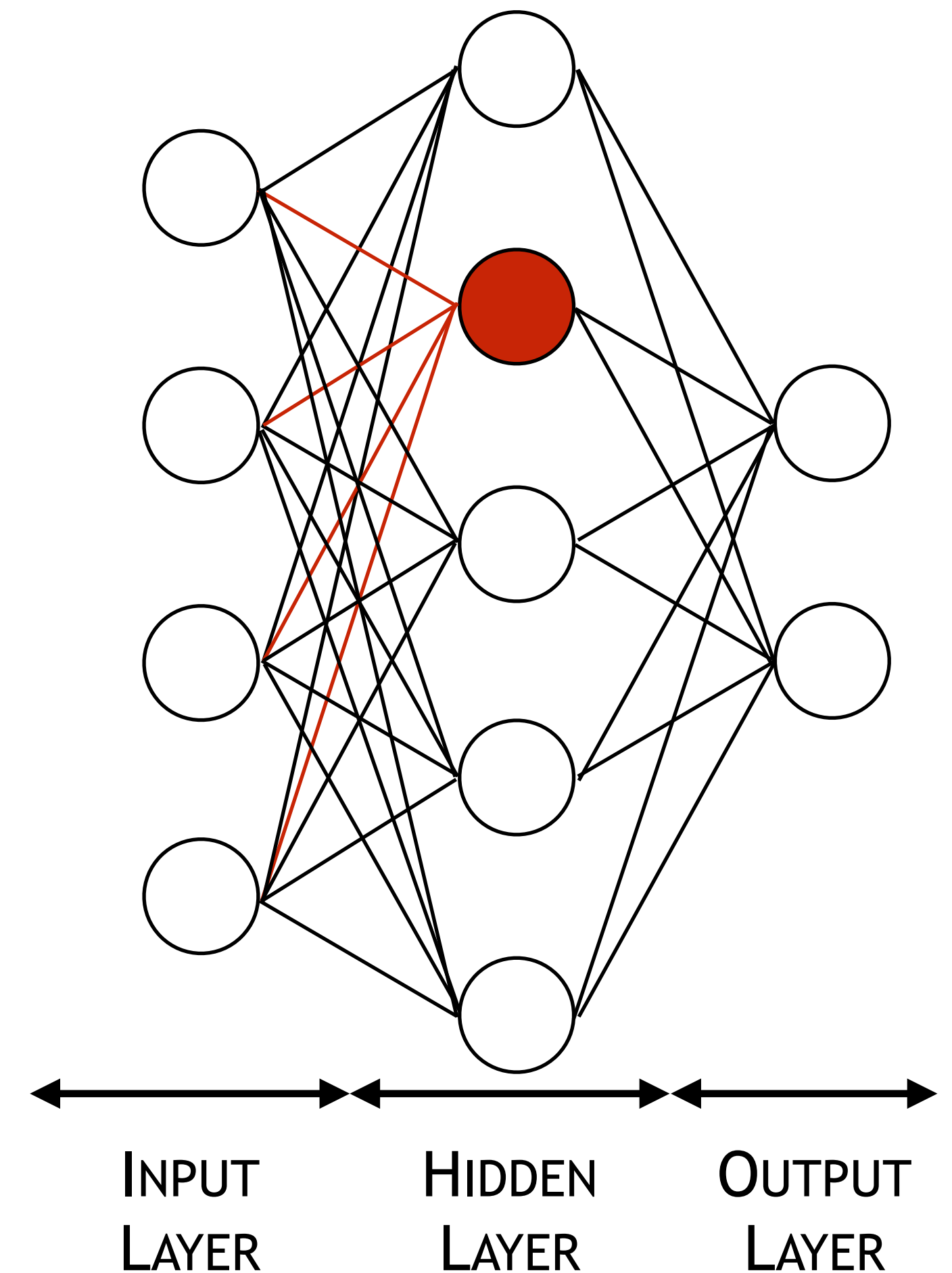
- Manifold learning to beat the curse of dimensionality

NN training

- Forward pass: Matmuls + nonlinearities

- Backward pass: Partial derivatives of weight w.r.t. loss

- Iterative weight updates



EXERCISE 02

Online at the materials page

Neural network from scratch: Forward pass, backward pass, weight update

Training and investigation of the LR

To be submitted via e-mail by:

Wednesday 9:00

Discussion after the next lecture



https://hawaii.ziti.uni-heidelberg.de/teaching/ap_nn_from_scratch_materials_wise2025/

DISCUSSION EXERCISE 01