

Indian Institute of Information Technology, Allahabad



Neural Networks

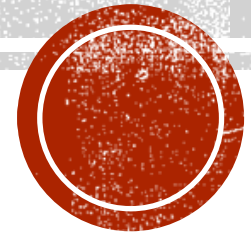
By

Dr. Shiv Ram Dubey

Computer Vision and Biometrics Lab

Department of Information Technology

Indian Institute of Information Technology, Allahabad



DISCLAIMER

The content (text, image, and graphics) used in this slide are adopted from many sources for Academic purposes. Broadly, the sources have been given due credit appropriately. However, there is a chance of missing out some original primary sources. The authors of this material do not claim any copyright of such material.



TODAY'S CLASS

Optimization

Gradient Descent & Back propagation

Update rule

Neural networks

Activation Function

Data Pre-processing



OPTIMIZATION

Optimization is the process of finding the set of parameters W that minimize the loss function.

Strategy #1: First very bad idea solution: Random search:

Simply try out many different random weights and keep track of what works best.

Strategy #2: Random local search:

Start out with a random W , generate random changes δW to it and if the loss at the changed $W + \delta W$ is lower, we will perform an update.

Strategy #3: Following the gradients:

There is no need to randomly search for a good direction: this direction is related to the gradient of the loss function.



GRADIENT DESCENT

The procedure of repeatedly evaluating the **gradient of loss function** and then performing a **parameter update**.

Vanilla (Original) Gradient Descent:

```
while True:
    weights_grad = evaluate_gradient(loss_fun, data, weights)
    weights += - step_size * weights_grad # perform parameter update
```

Mini-batch Gradient Descent (MGD):

```
while True:
    data_batch = sample_training_data(data, 256) # sample 256 examples
    weights_grad = evaluate_gradient(loss_fun, data_batch, weights)
    weights += - step_size * weights_grad # perform parameter update
```

Stochastic Gradient Descent (SGD):

Special case of MGD when mini-batch contains only a single example



INTERPRETATION OF THE GRADIENT

Interpretation. Derivatives indicate the rate of change of a function with respect to that variable surrounding an infinitesimally small region near a particular point:

$$\frac{df(x)}{dx} = \lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h}$$



INTERPRETATION OF THE GRADIENT

Interpretation. Derivatives indicate the rate of change of a function with respect to that variable surrounding an infinitesimally small region near a particular point:

$$\frac{df(x)}{dx} = \lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h}$$

$$f(x, y) = x + y \quad \rightarrow \quad \frac{\partial f}{\partial x} = \quad \frac{\partial f}{\partial y} =$$



INTERPRETATION OF THE GRADIENT

Interpretation. Derivatives indicate the rate of change of a function with respect to that variable surrounding an infinitesimally small region near a particular point:

$$\frac{df(x)}{dx} = \lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h}$$

$$f(x, y) = x + y \quad \rightarrow \quad \frac{\partial f}{\partial x} = 1 \quad \frac{\partial f}{\partial y} = 1$$



INTERPRETATION OF THE GRADIENT

Interpretation. Derivatives indicate the rate of change of a function with respect to that variable surrounding an infinitesimally small region near a particular point:

$$\frac{df(x)}{dx} = \lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h}$$

$$f(x, y) = x + y \quad \rightarrow \quad \frac{\partial f}{\partial x} = 1 \quad \frac{\partial f}{\partial y} = 1$$

$$f(x, y) = xy \quad \rightarrow \quad \frac{\partial f}{\partial x} = \quad \frac{\partial f}{\partial y} =$$



INTERPRETATION OF THE GRADIENT

Interpretation. Derivatives indicate the rate of change of a function with respect to that variable surrounding an infinitesimally small region near a particular point:

$$\frac{df(x)}{dx} = \lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h}$$

$$f(x, y) = x + y \quad \rightarrow \quad \frac{\partial f}{\partial x} = 1 \quad \frac{\partial f}{\partial y} = 1$$

$$f(x, y) = xy \quad \rightarrow \quad \frac{\partial f}{\partial x} = y \quad \frac{\partial f}{\partial y} = x$$



INTERPRETATION OF THE GRADIENT

Interpretation. Derivatives indicate the rate of change of a function with respect to that variable surrounding an infinitesimally small region near a particular point:

$$\frac{df(x)}{dx} = \lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h}$$

$$f(x, y) = x + y \quad \rightarrow \quad \frac{\partial f}{\partial x} = 1 \quad \frac{\partial f}{\partial y} = 1$$

$$f(x, y) = xy \quad \rightarrow \quad \frac{\partial f}{\partial x} = y \quad \frac{\partial f}{\partial y} = x$$

$$f(x, y) = \max(x, y) \quad \rightarrow \quad \frac{\partial f}{\partial x} = \quad \frac{\partial f}{\partial y} =$$



INTERPRETATION OF THE GRADIENT

Interpretation. Derivatives indicate the rate of change of a function with respect to that variable surrounding an infinitesimally small region near a particular point:

$$\frac{df(x)}{dx} = \lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h}$$

$$f(x, y) = x + y \quad \rightarrow \quad \frac{\partial f}{\partial x} = 1 \quad \frac{\partial f}{\partial y} = 1$$

$$f(x, y) = xy \quad \rightarrow \quad \frac{\partial f}{\partial x} = y \quad \frac{\partial f}{\partial y} = x$$

$$f(x, y) = \max(x, y) \quad \rightarrow \quad \frac{\partial f}{\partial x} = 1(x \geq y) \quad \frac{\partial f}{\partial y} = 1(y \geq x)$$



COMPOUND EXPRESSIONS WITH CHAIN RULE

$$f(x, y, z) = (x + y)z$$



COMPOUND EXPRESSIONS WITH CHAIN RULE

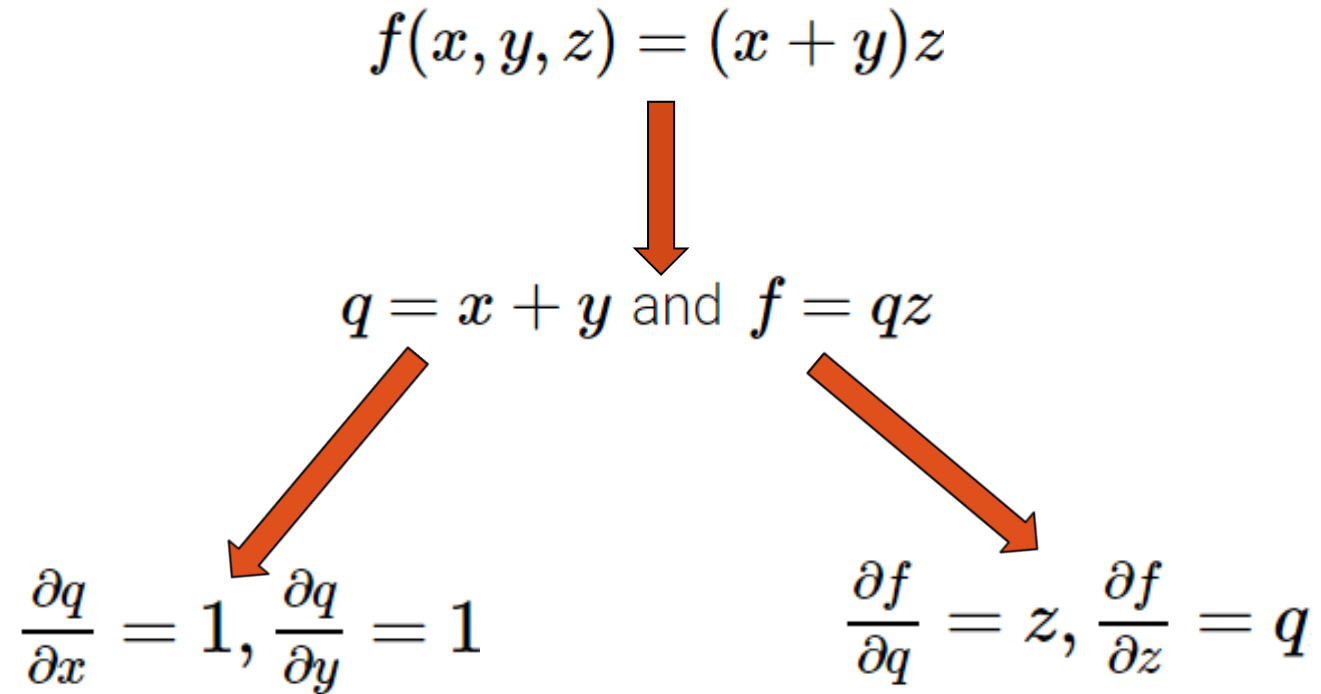
$$f(x, y, z) = (x + y)z$$



$$q = x + y \text{ and } f = qz$$



COMPOUND EXPRESSIONS WITH CHAIN RULE



COMPOUND EXPRESSIONS WITH CHAIN RULE

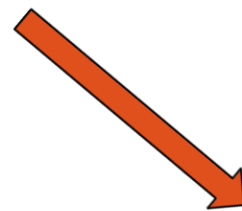
$$f(x, y, z) = (x + y)z$$



$$q = x + y \text{ and } f = qz$$



$$\frac{\partial q}{\partial x} = 1, \frac{\partial q}{\partial y} = 1$$



$$\frac{\partial f}{\partial q} = z, \frac{\partial f}{\partial z} = q$$

Chain rule:

$$\frac{\partial f}{\partial x} = \frac{\partial f}{\partial q} \frac{\partial q}{\partial x}$$



COMPOUND EXPRESSIONS WITH CHAIN RULE

$$f(x, y, z) = (x + y)z$$

↓

$$q = x + y \text{ and } f = qz$$

$$\frac{\partial f}{\partial q} = z, \frac{\partial f}{\partial z} = q$$

$$\frac{\partial q}{\partial x} = 1, \frac{\partial q}{\partial y} = 1$$

$$\frac{\partial f}{\partial x} = \frac{\partial f}{\partial q} \frac{\partial q}{\partial x}$$



COMPOUND EXPRESSIONS WITH CHAIN RULE

$$f(x, y, z) = (x + y)z$$

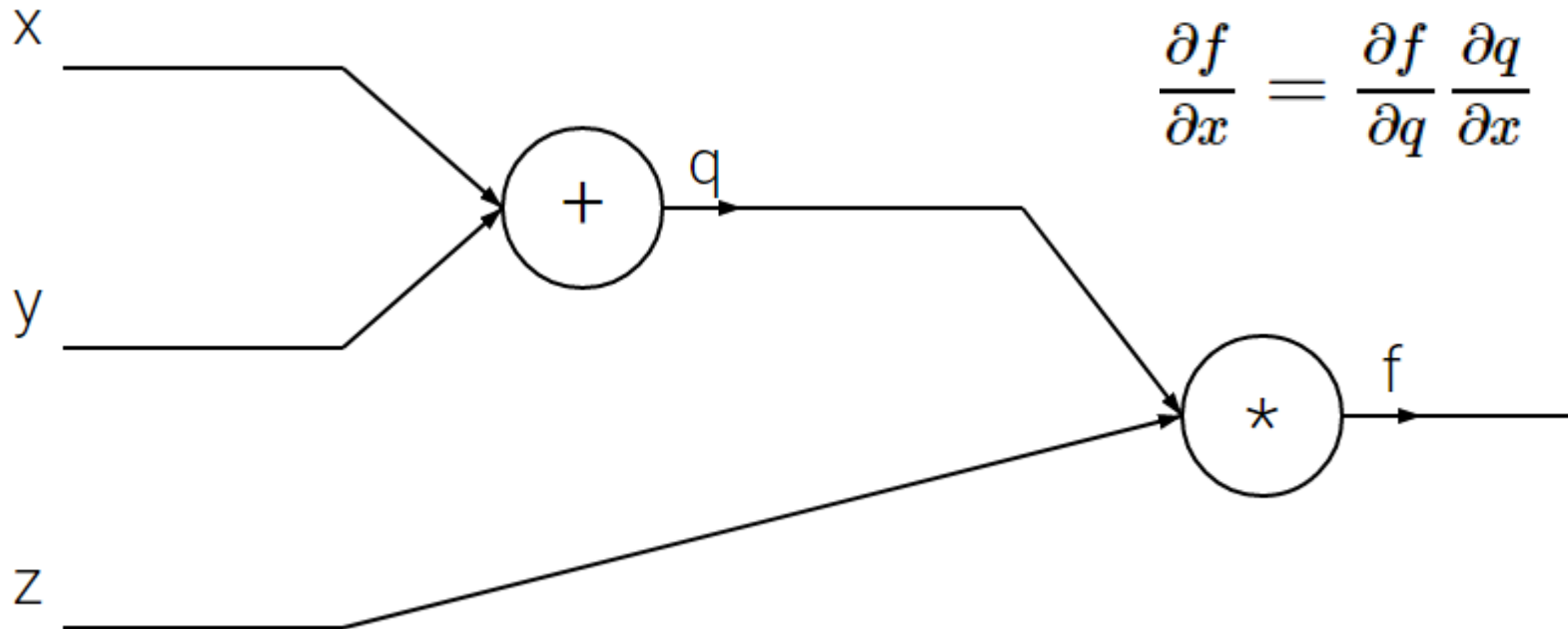


$$q = x + y \text{ and } f = qz$$

$$\frac{\partial f}{\partial q} = z, \frac{\partial f}{\partial z} = q$$

$$\frac{\partial q}{\partial x} = 1, \frac{\partial q}{\partial y} = 1$$

$$\frac{\partial f}{\partial x} = \frac{\partial f}{\partial q} \frac{\partial q}{\partial x}$$



COMPOUND EXPRESSIONS WITH CHAIN RULE

$$f(x, y, z) = (x + y)z$$

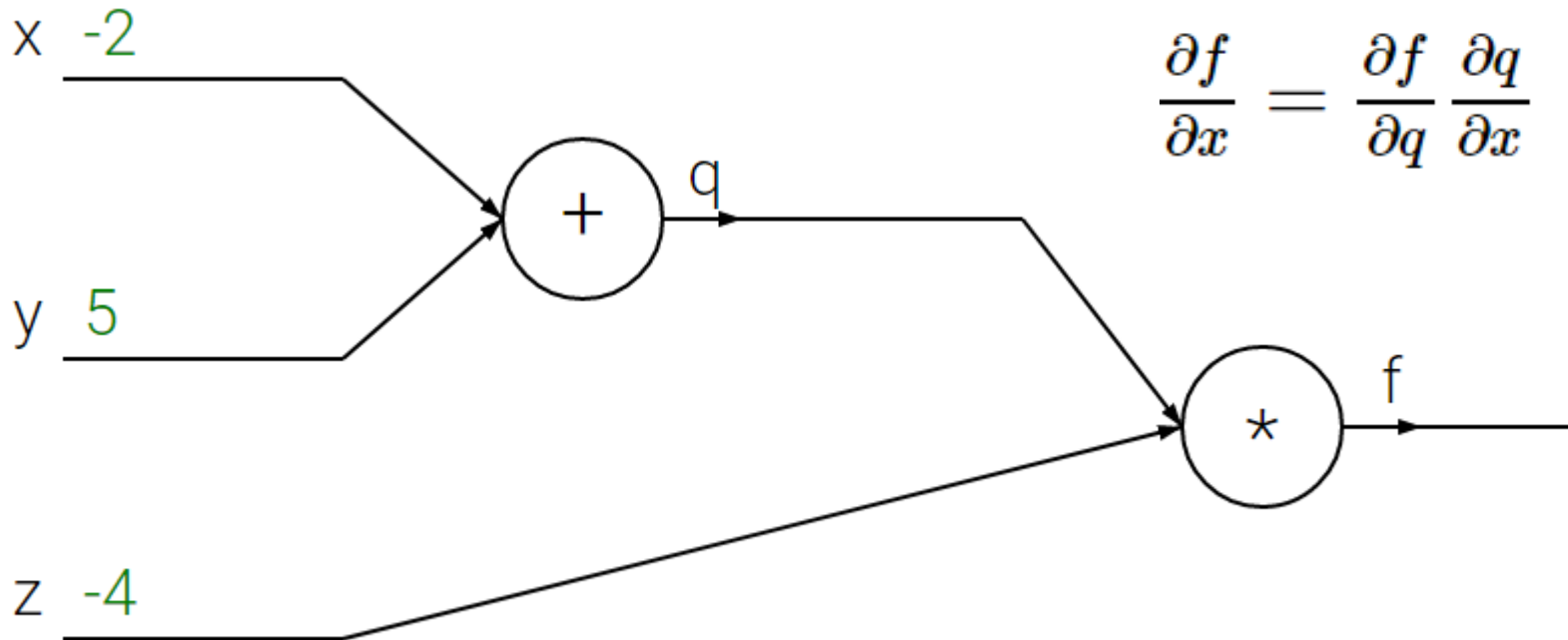


$$q = x + y \text{ and } f = qz$$

$$\frac{\partial f}{\partial q} = z, \frac{\partial f}{\partial z} = q$$

$$\frac{\partial q}{\partial x} = 1, \frac{\partial q}{\partial y} = 1$$

$$\frac{\partial f}{\partial x} = \frac{\partial f}{\partial q} \frac{\partial q}{\partial x}$$



COMPOUND EXPRESSIONS WITH CHAIN RULE

$$f(x, y, z) = (x + y)z$$

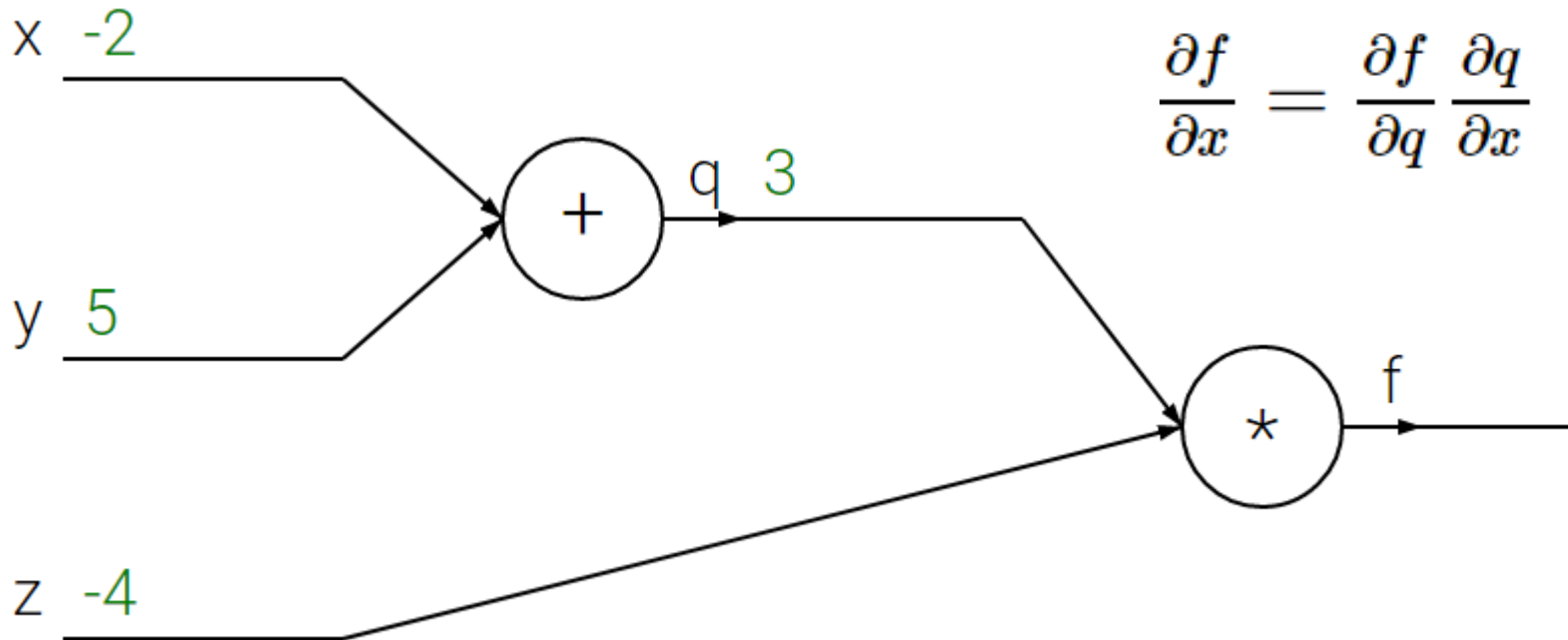


$$q = x + y \text{ and } f = qz$$

$$\frac{\partial f}{\partial q} = z, \frac{\partial f}{\partial z} = q$$

$$\frac{\partial q}{\partial x} = 1, \frac{\partial q}{\partial y} = 1$$

$$\frac{\partial f}{\partial x} = \frac{\partial f}{\partial q} \frac{\partial q}{\partial x}$$



COMPOUND EXPRESSIONS WITH CHAIN RULE

$$f(x, y, z) = (x + y)z$$

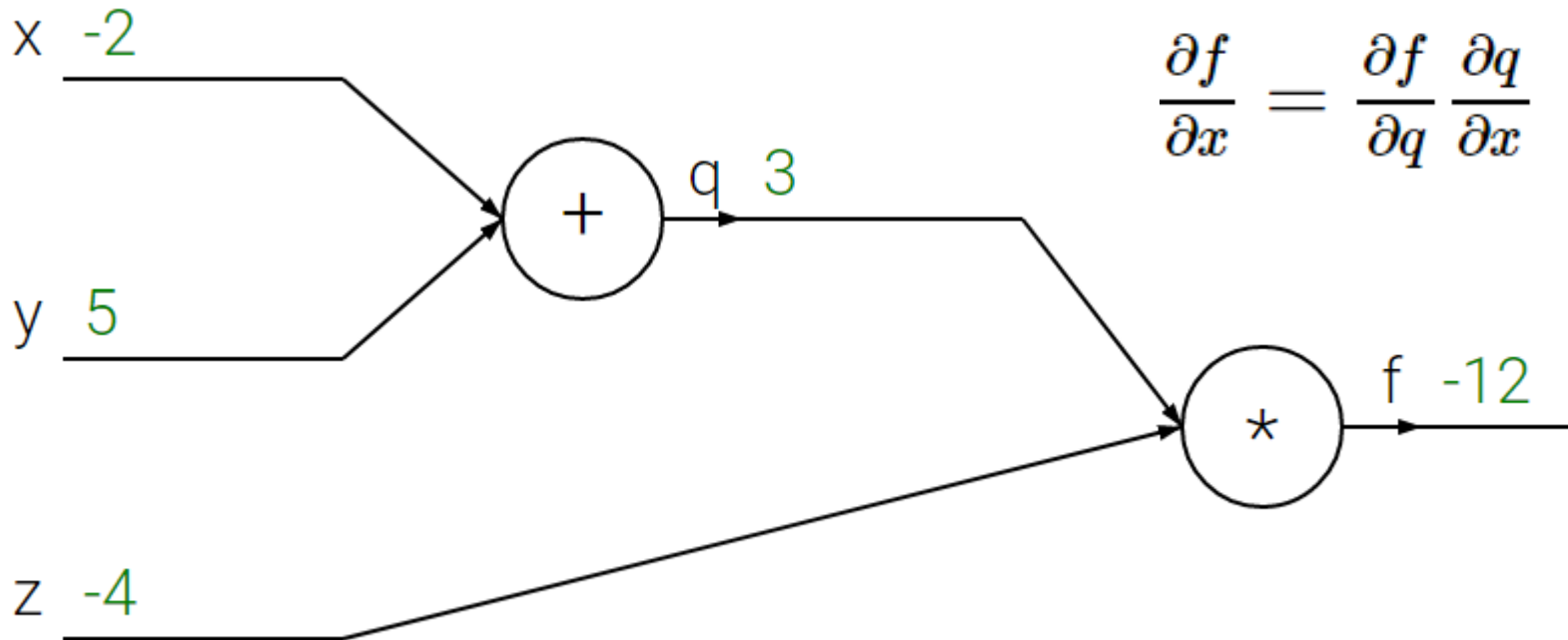


$$q = x + y \text{ and } f = qz$$

$$\frac{\partial f}{\partial q} = z, \frac{\partial f}{\partial z} = q$$

$$\frac{\partial q}{\partial x} = 1, \frac{\partial q}{\partial y} = 1$$

$$\frac{\partial f}{\partial x} = \frac{\partial f}{\partial q} \frac{\partial q}{\partial x}$$



COMPOUND EXPRESSIONS WITH CHAIN RULE

$$f(x, y, z) = (x + y)z$$

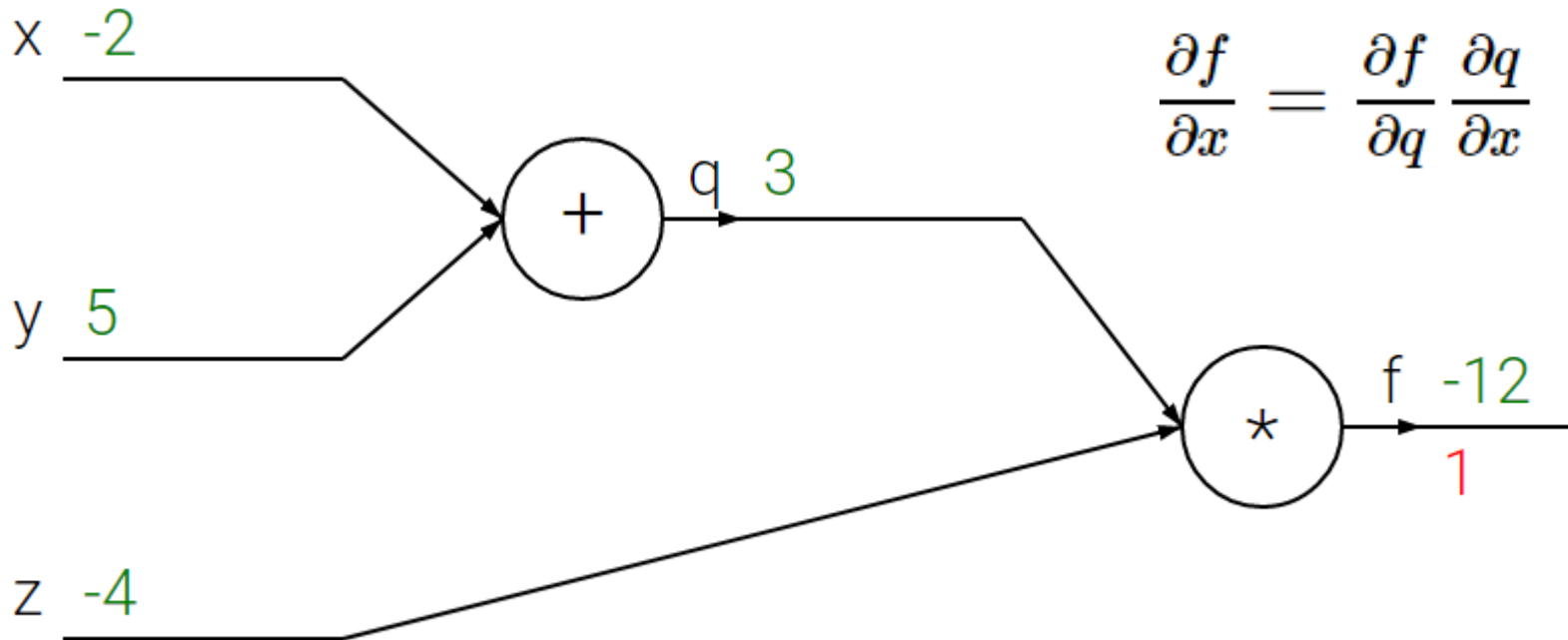


$$q = x + y \text{ and } f = qz$$

$$\frac{\partial f}{\partial q} = z, \frac{\partial f}{\partial z} = q$$

$$\frac{\partial q}{\partial x} = 1, \frac{\partial q}{\partial y} = 1$$

$$\frac{\partial f}{\partial x} = \frac{\partial f}{\partial q} \frac{\partial q}{\partial x}$$



COMPOUND EXPRESSIONS WITH CHAIN RULE

$$f(x, y, z) = (x + y)z$$

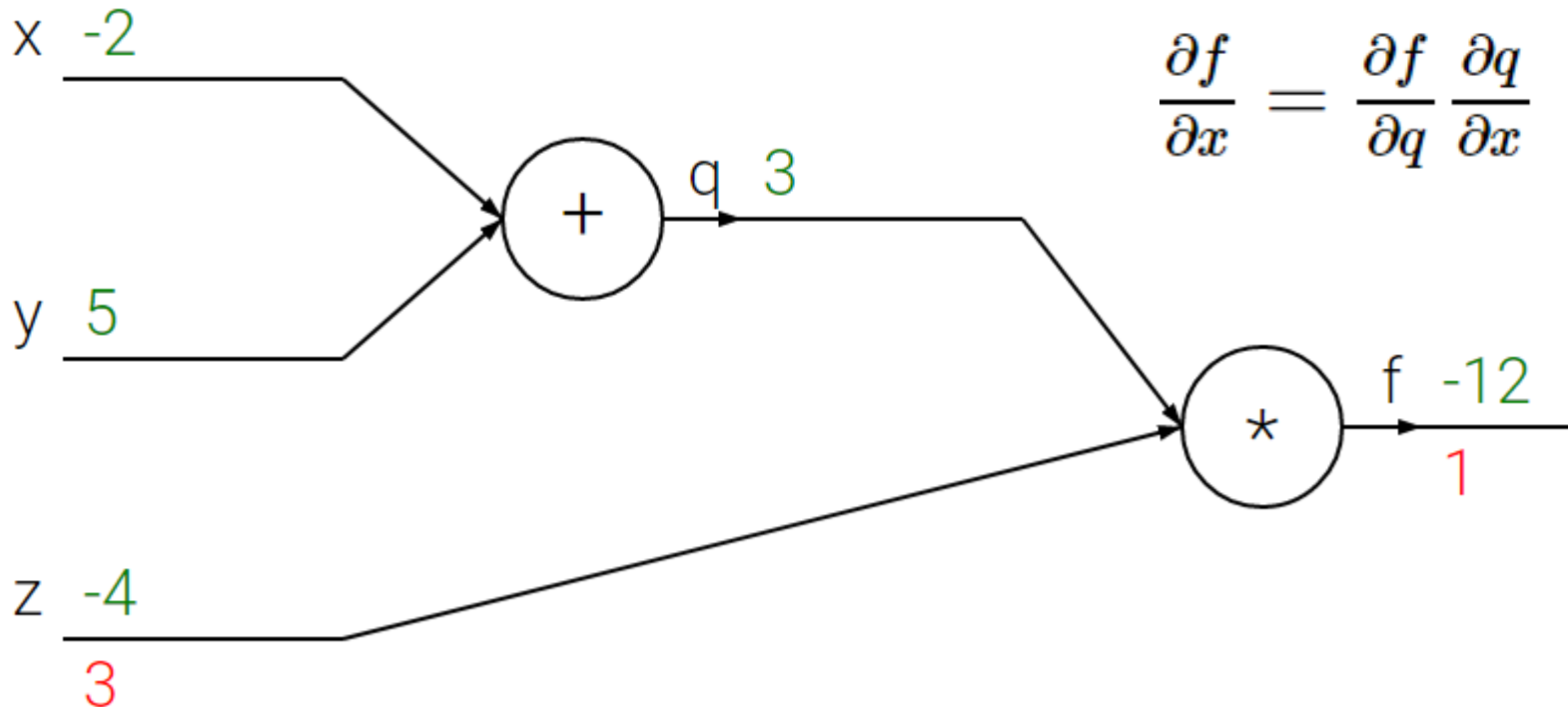


$$q = x + y \text{ and } f = qz$$

$$\frac{\partial f}{\partial q} = z, \frac{\partial f}{\partial z} = q$$

$$\frac{\partial q}{\partial x} = 1, \frac{\partial q}{\partial y} = 1$$

$$\frac{\partial f}{\partial x} = \frac{\partial f}{\partial q} \frac{\partial q}{\partial x}$$



COMPOUND EXPRESSIONS WITH CHAIN RULE

$$f(x, y, z) = (x + y)z$$

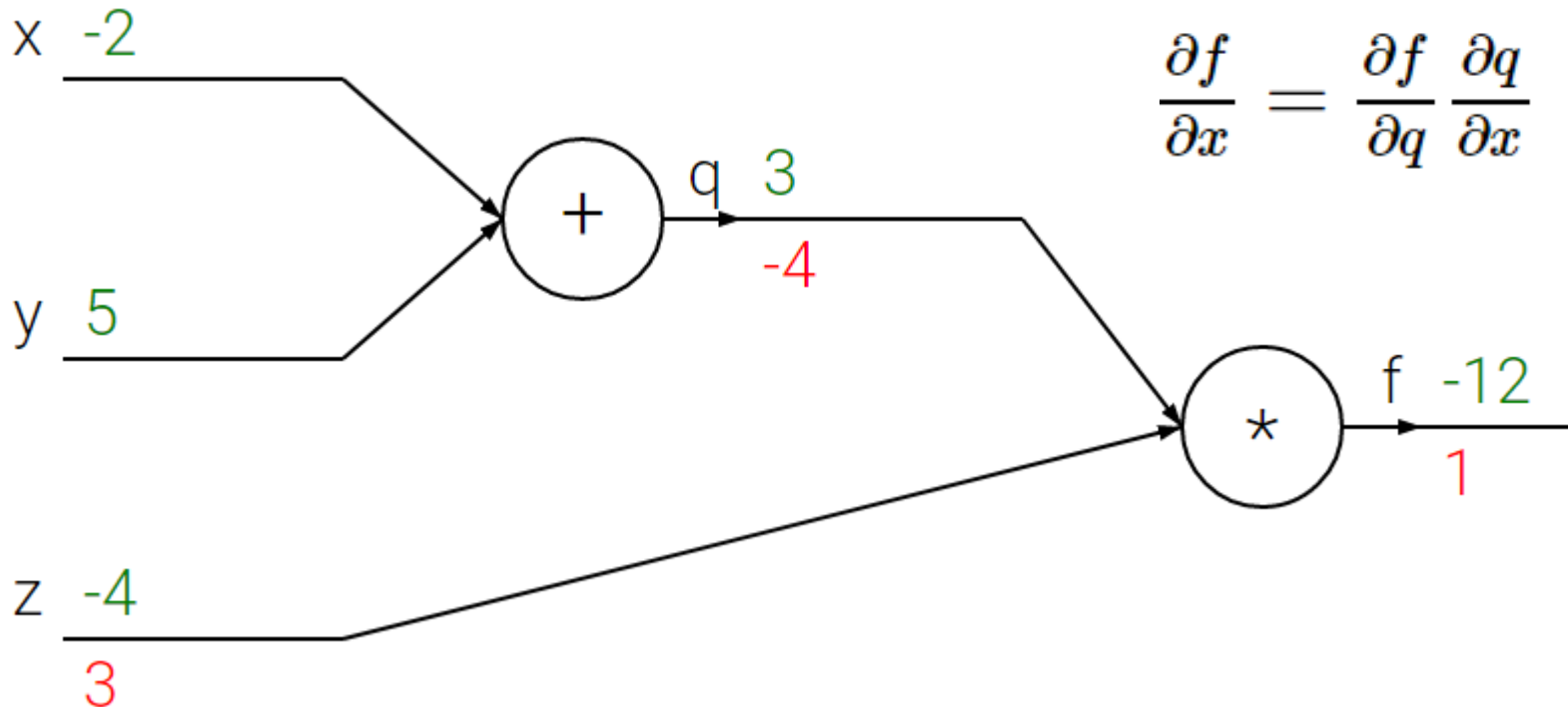


$$q = x + y \text{ and } f = qz$$

$$\frac{\partial f}{\partial q} = z, \frac{\partial f}{\partial z} = q$$

$$\frac{\partial q}{\partial x} = 1, \frac{\partial q}{\partial y} = 1$$

$$\frac{\partial f}{\partial x} = \frac{\partial f}{\partial q} \frac{\partial q}{\partial x}$$



COMPOUND EXPRESSIONS WITH CHAIN RULE

$$f(x, y, z) = (x + y)z$$

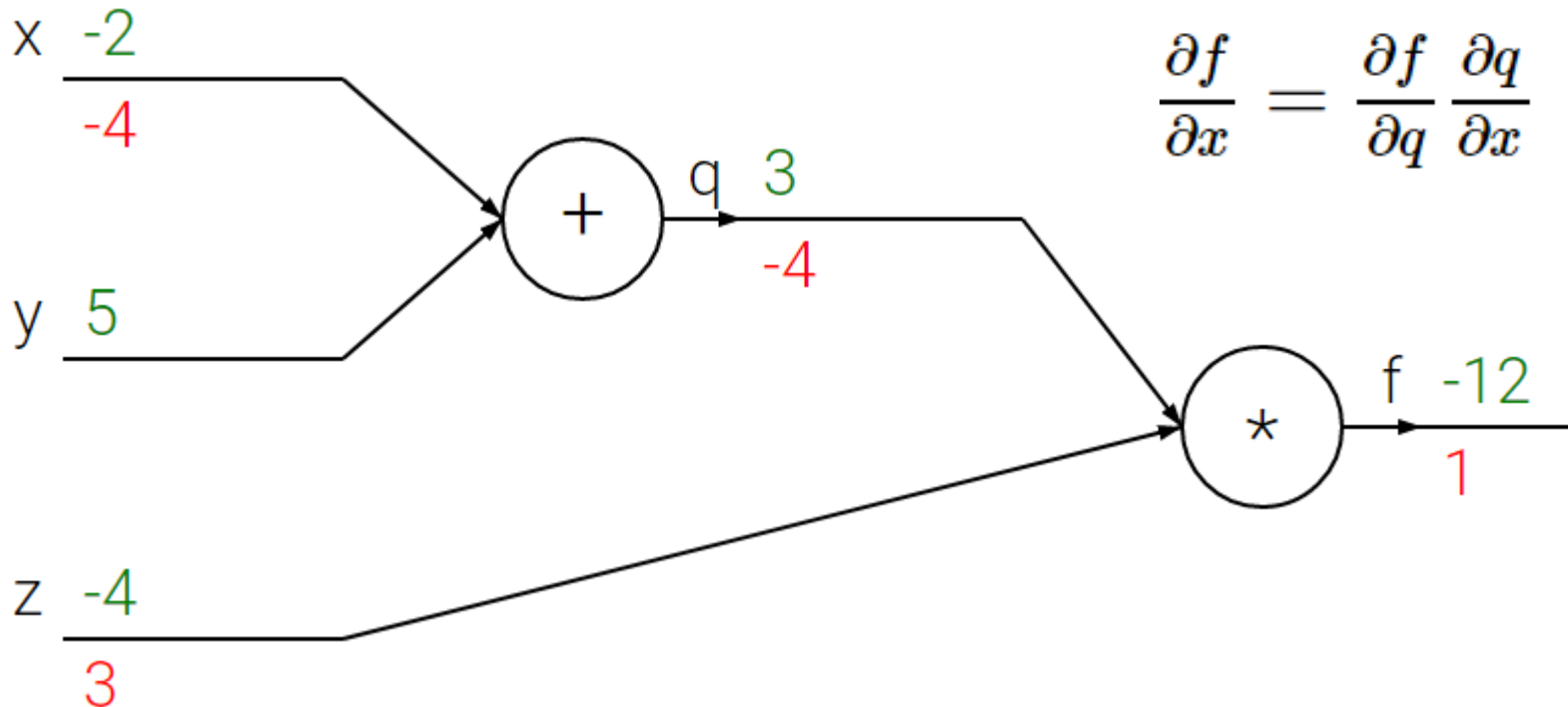


$$q = x + y \text{ and } f = qz$$

$$\frac{\partial f}{\partial q} = z, \frac{\partial f}{\partial z} = q$$

$$\frac{\partial q}{\partial x} = 1, \frac{\partial q}{\partial y} = 1$$

$$\frac{\partial f}{\partial x} = \frac{\partial f}{\partial q} \frac{\partial q}{\partial x}$$



COMPOUND EXPRESSIONS WITH CHAIN RULE

$$f(x, y, z) = (x + y)z$$

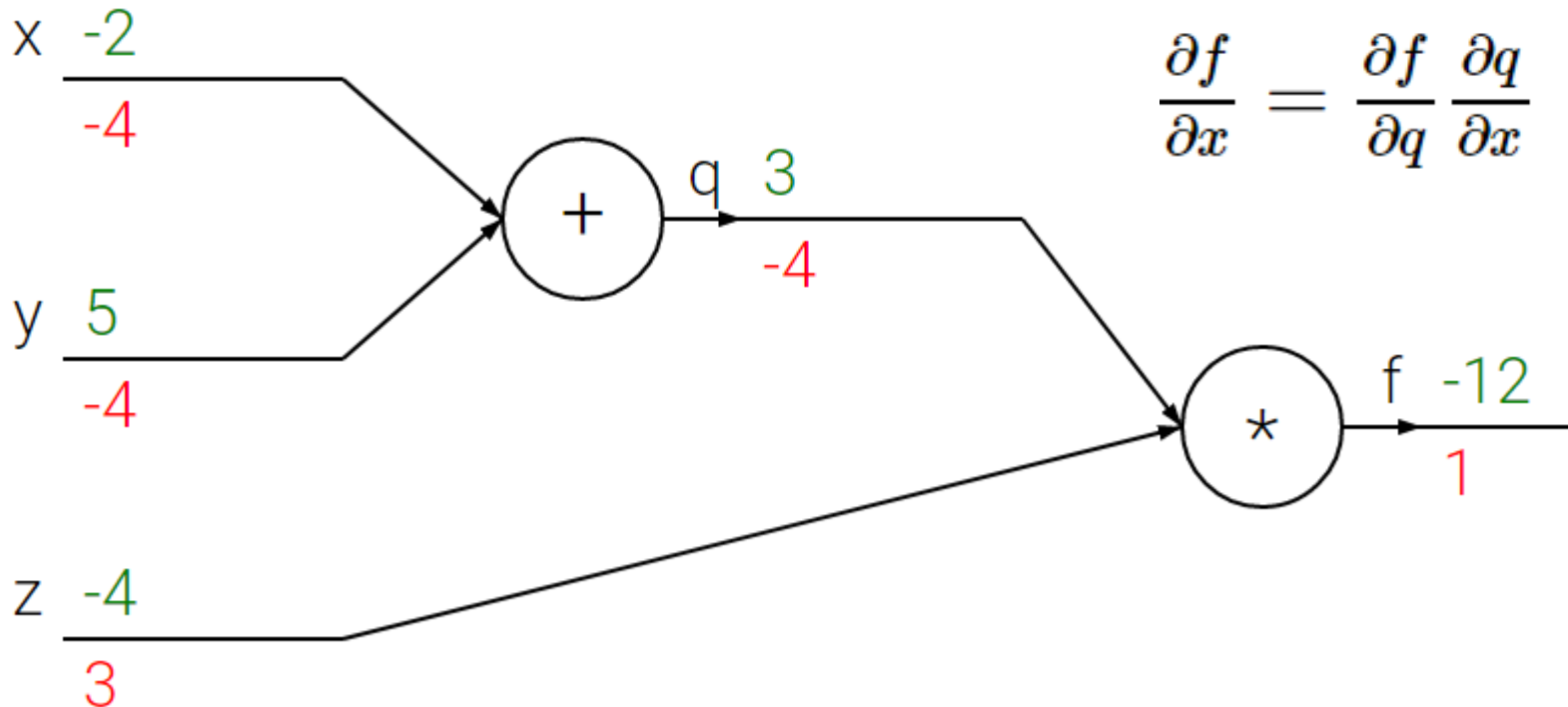


$$q = x + y \text{ and } f = qz$$

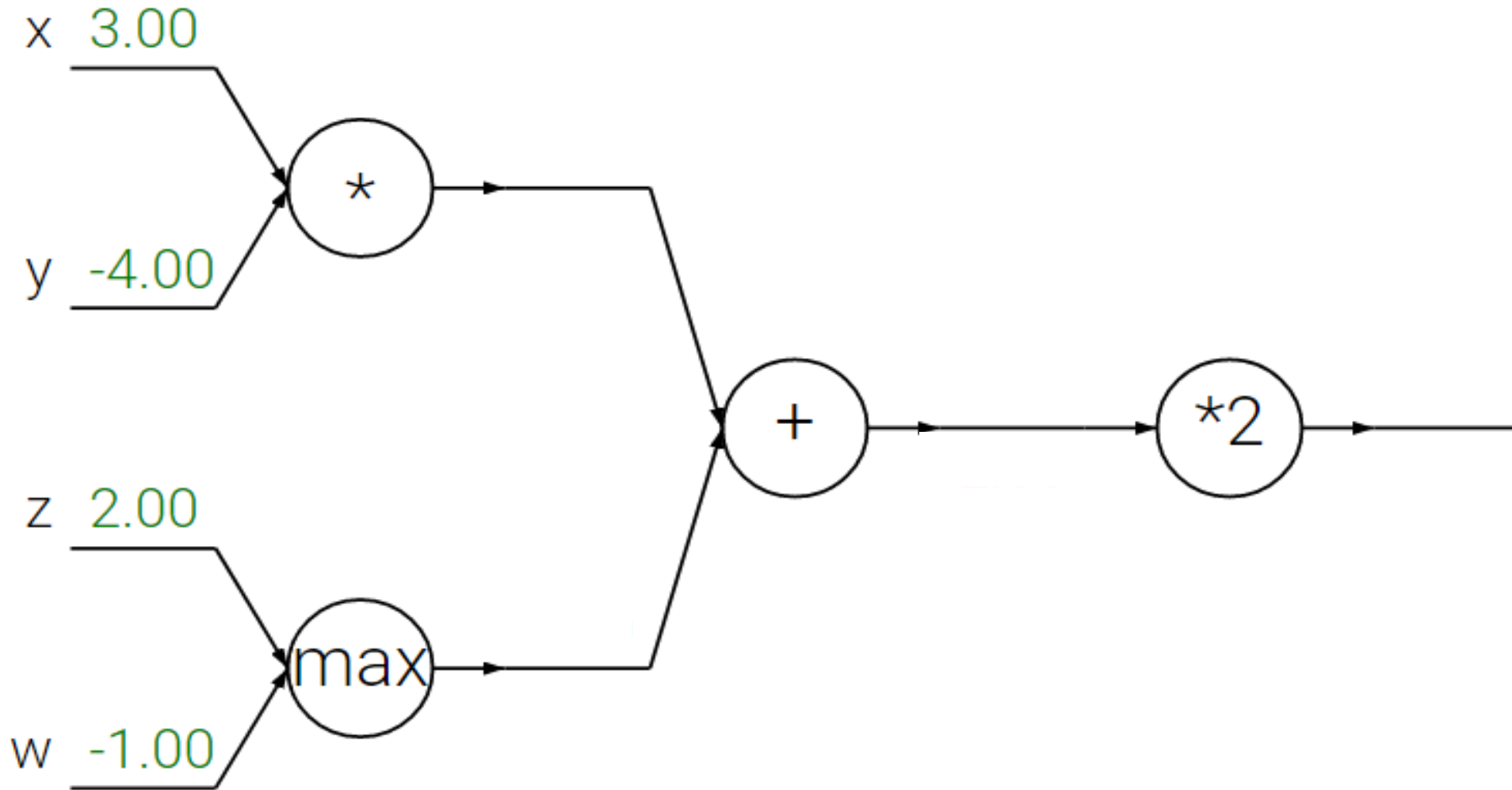
$$\frac{\partial f}{\partial q} = z, \frac{\partial f}{\partial z} = q$$

$$\frac{\partial q}{\partial x} = 1, \frac{\partial q}{\partial y} = 1$$

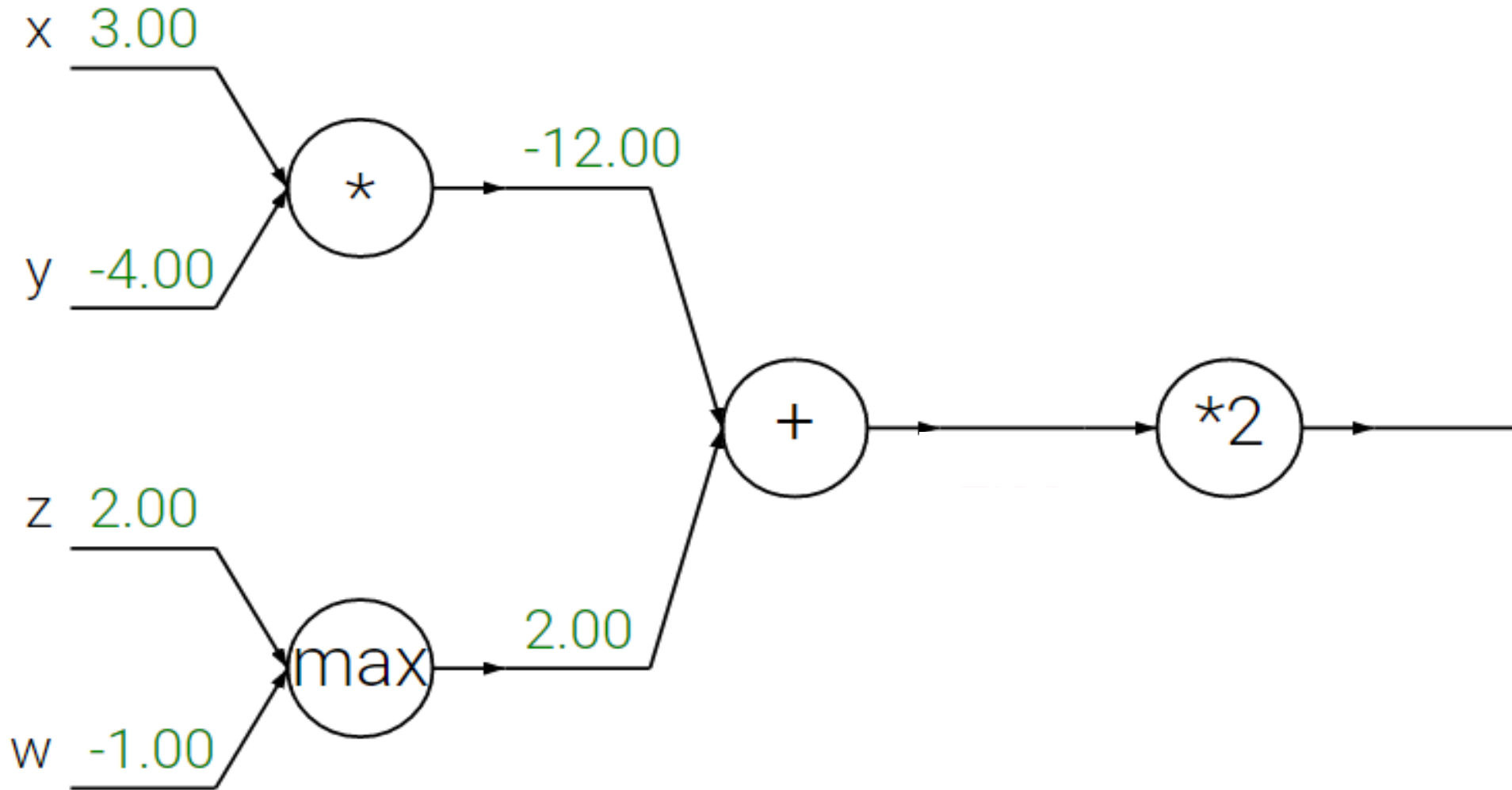
$$\frac{\partial f}{\partial x} = \frac{\partial f}{\partial q} \frac{\partial q}{\partial x}$$



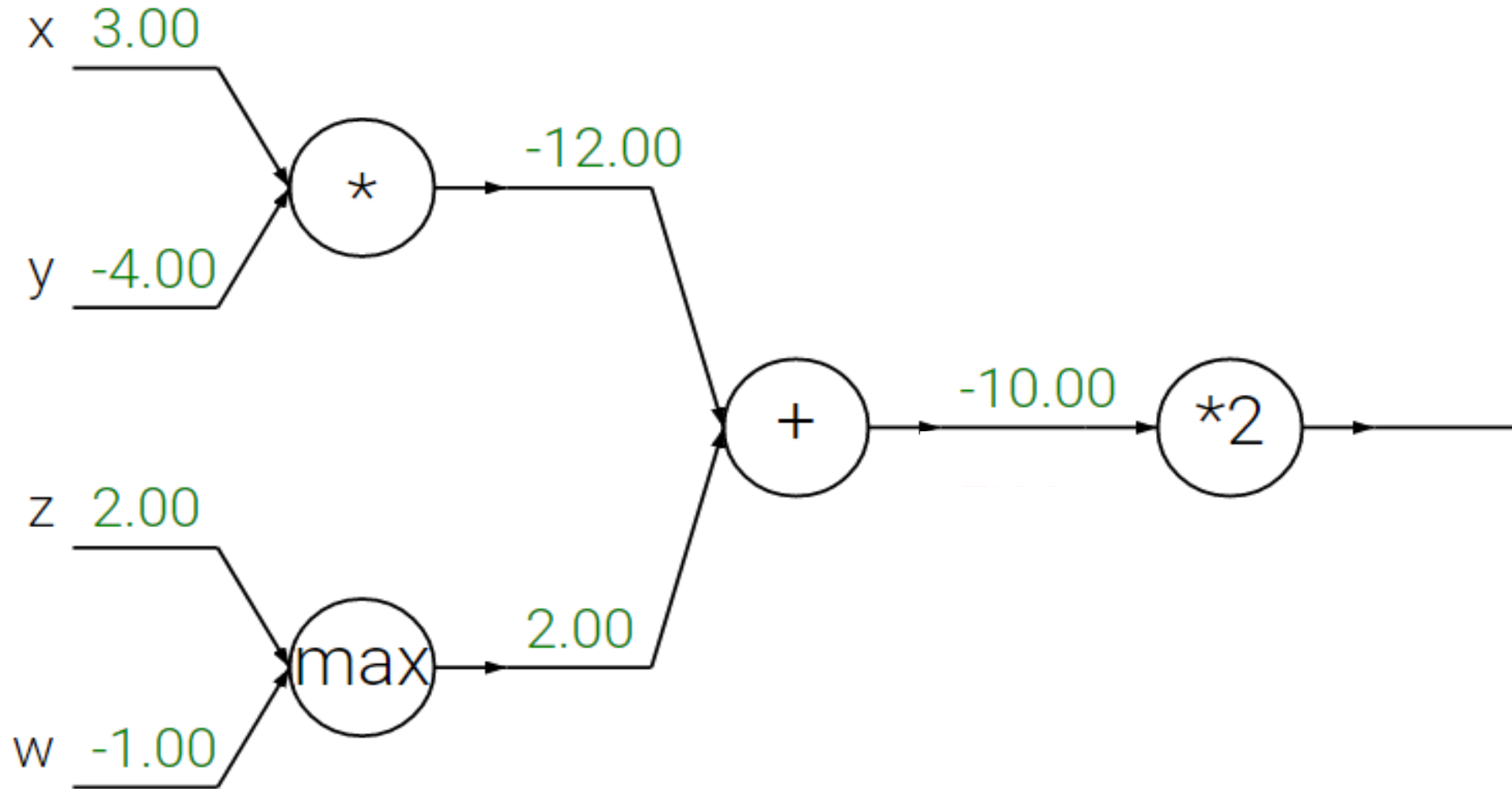
FORWARD AND BACKWARD PASS



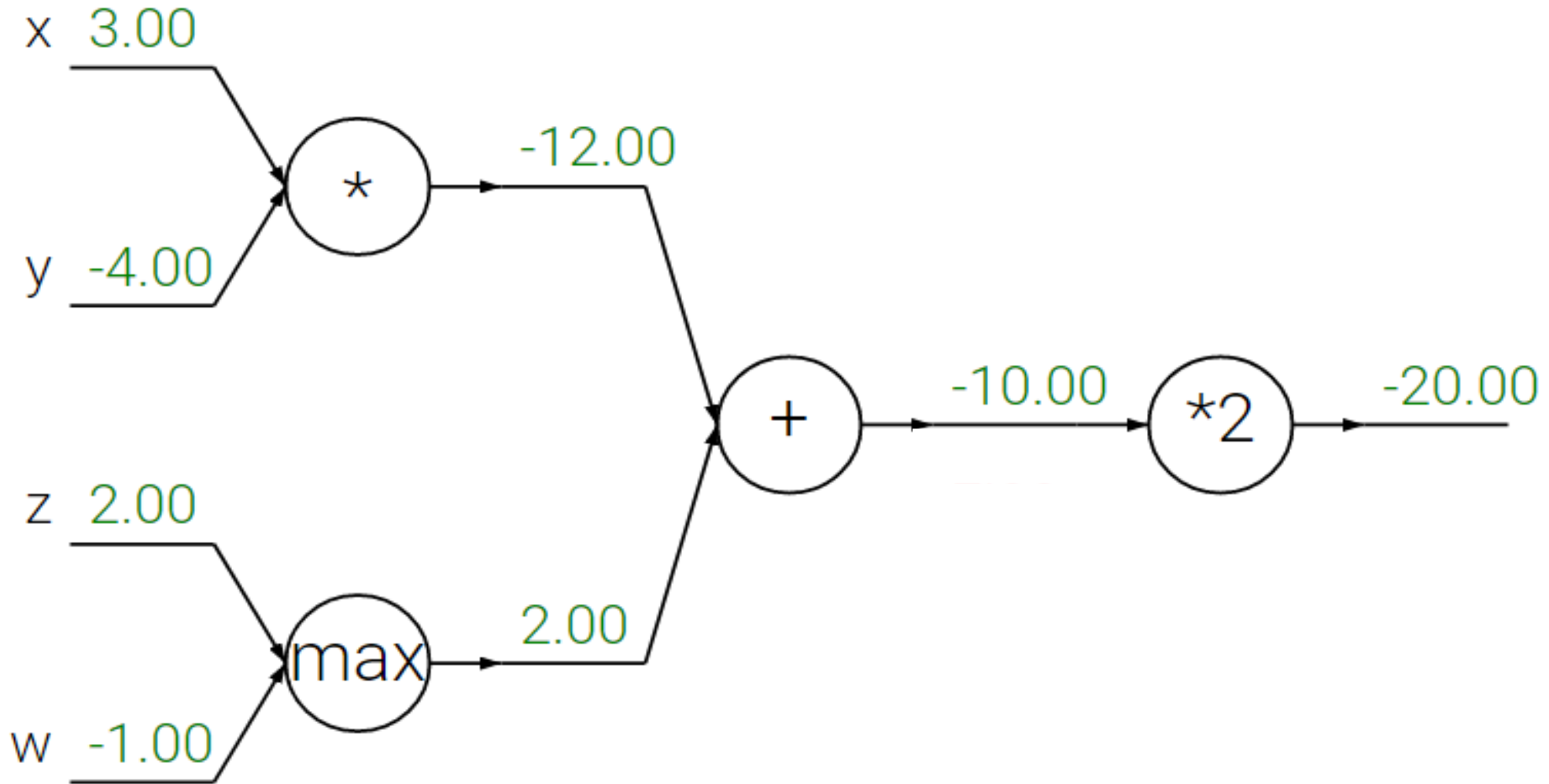
FORWARD AND BACKWARD PASS



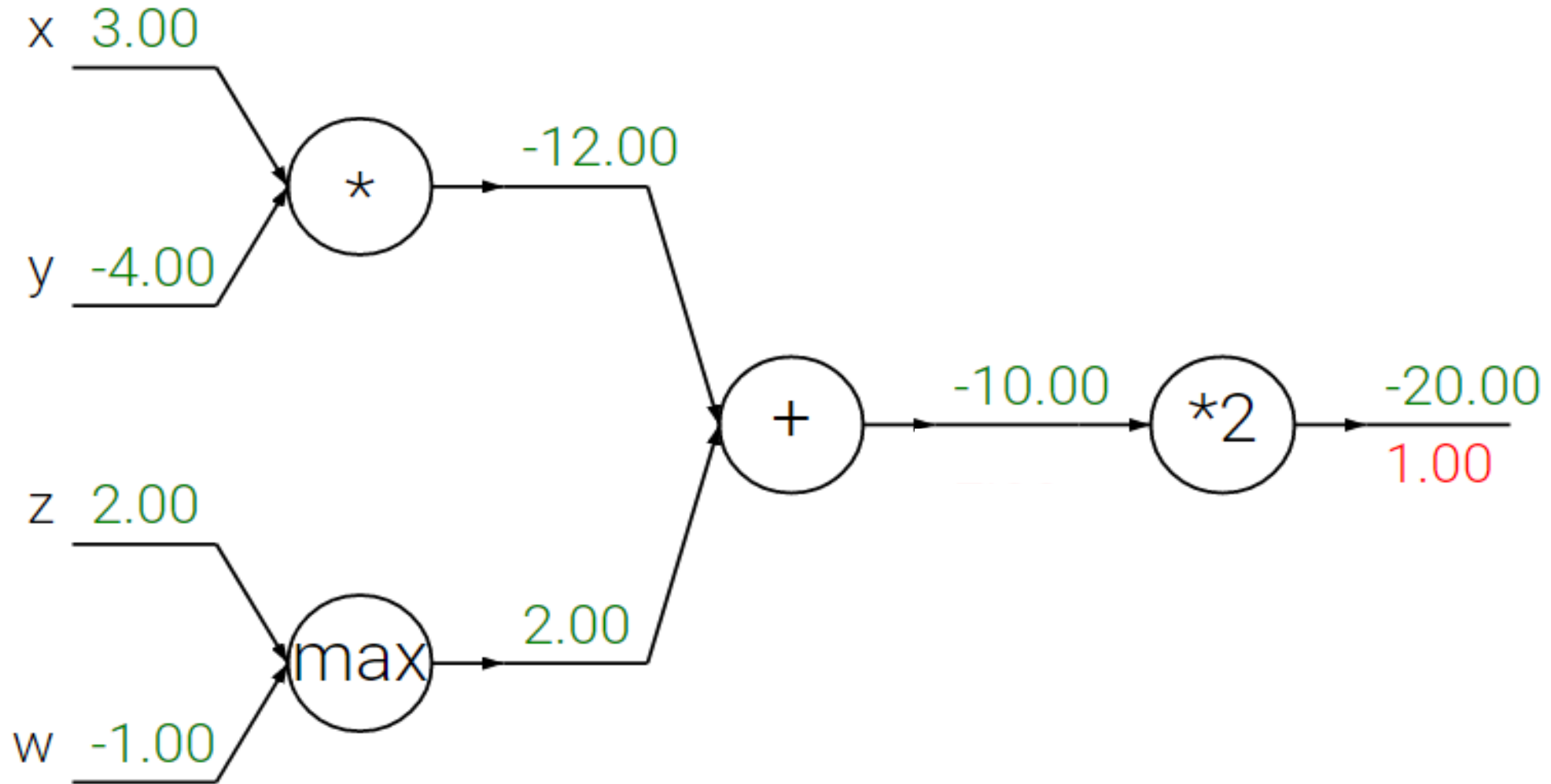
FORWARD AND BACKWARD PASS



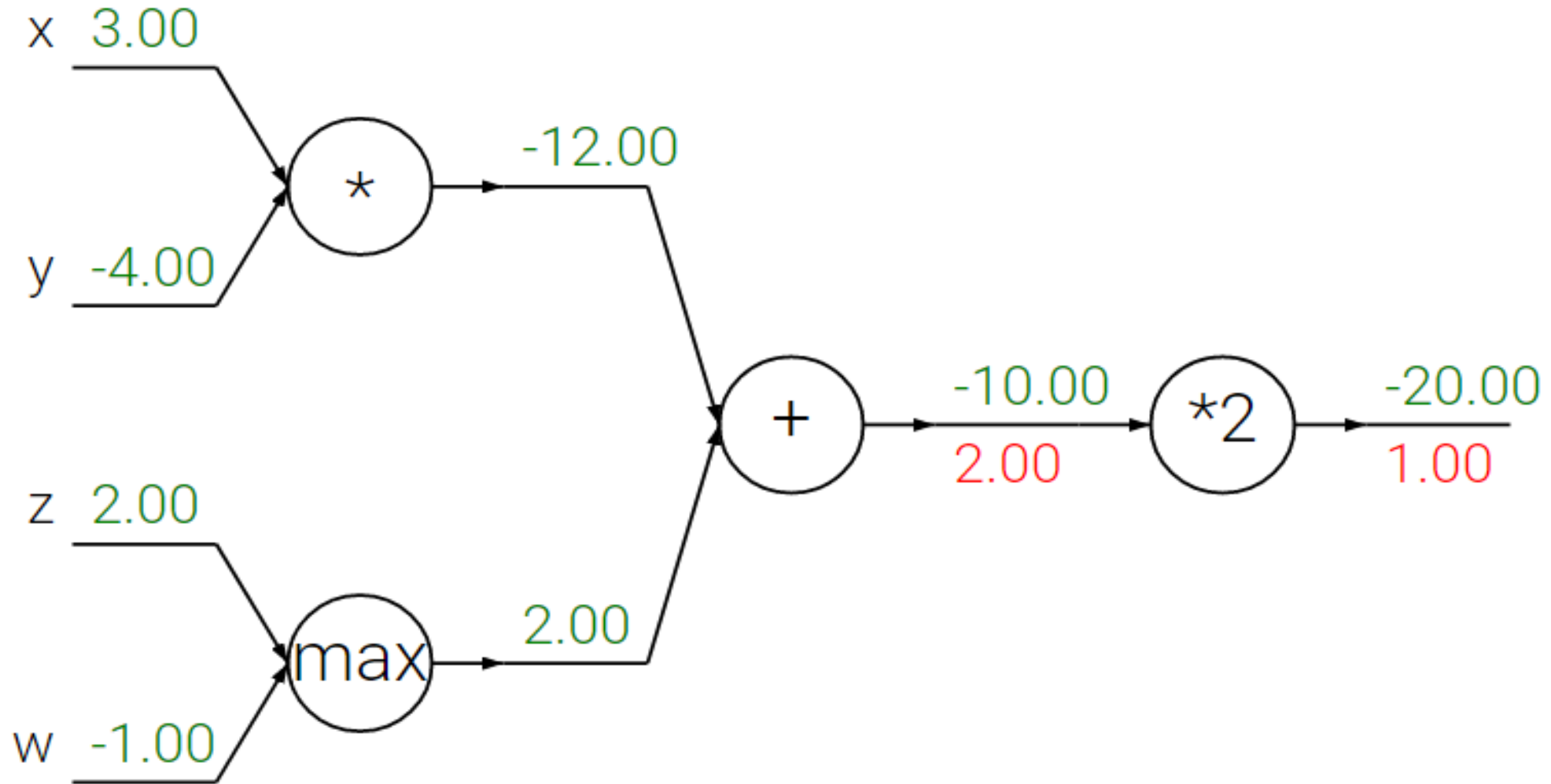
FORWARD AND BACKWARD PASS



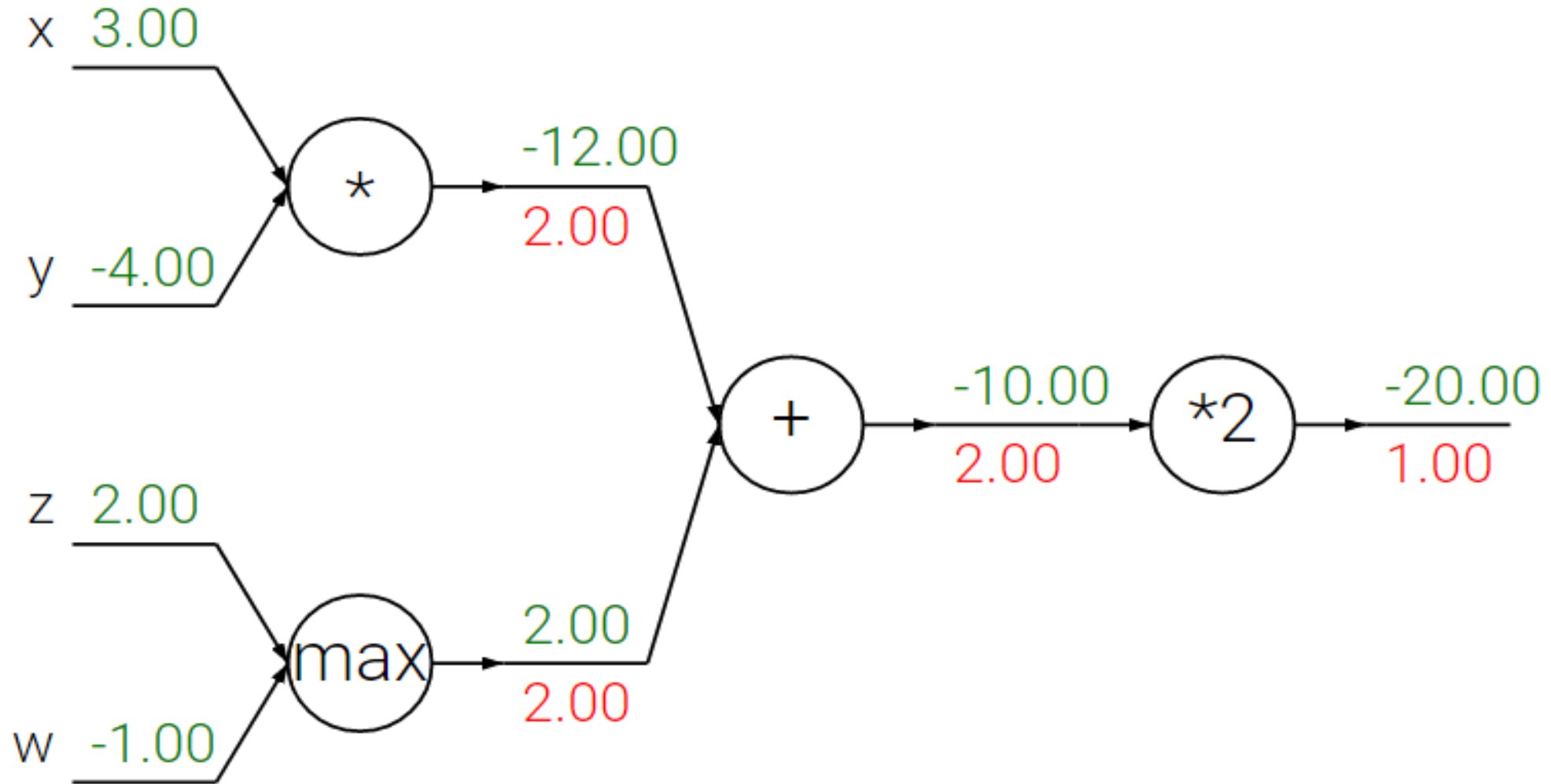
FORWARD AND BACKWARD PASS



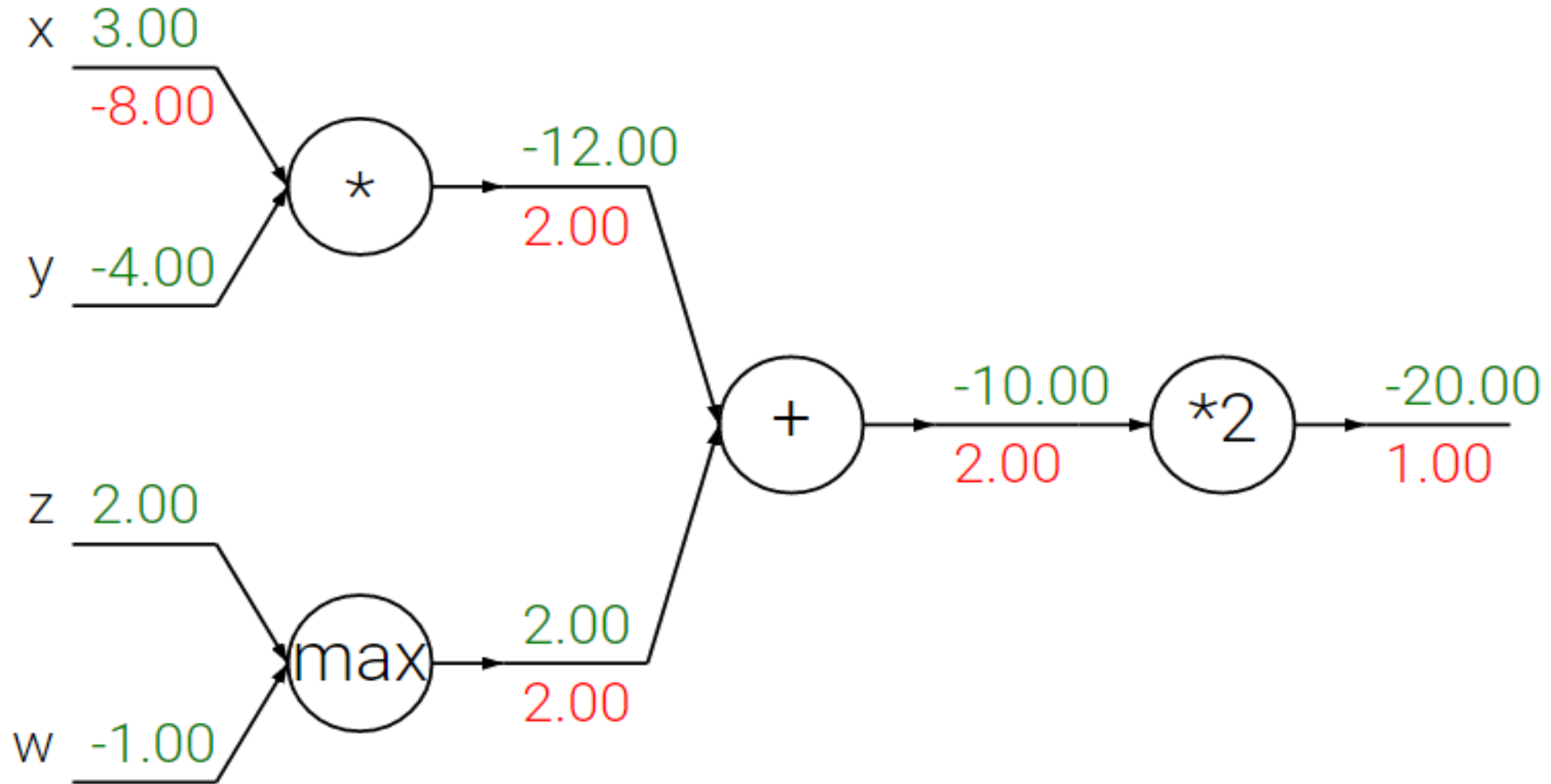
FORWARD AND BACKWARD PASS



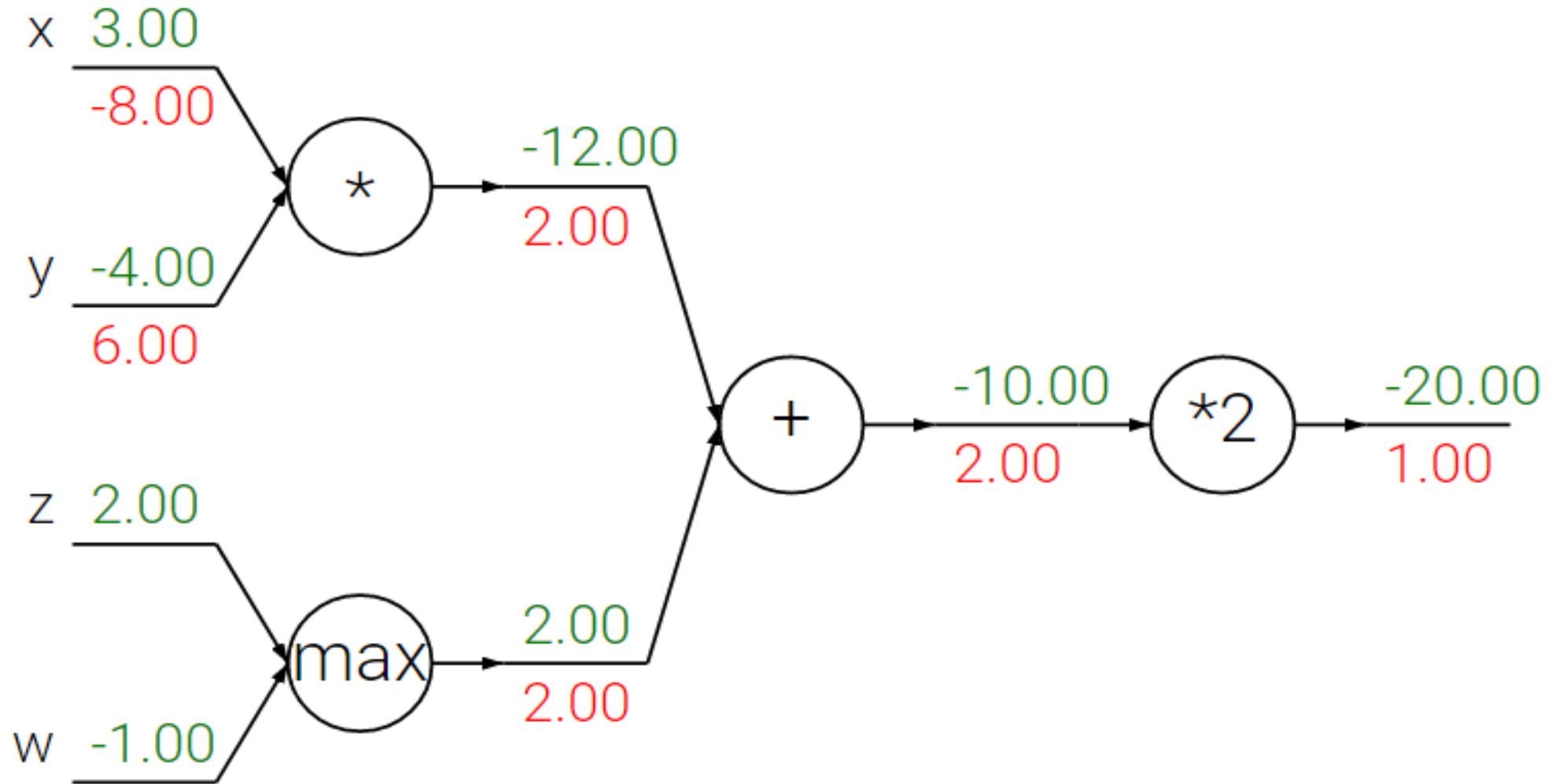
FORWARD AND BACKWARD PASS



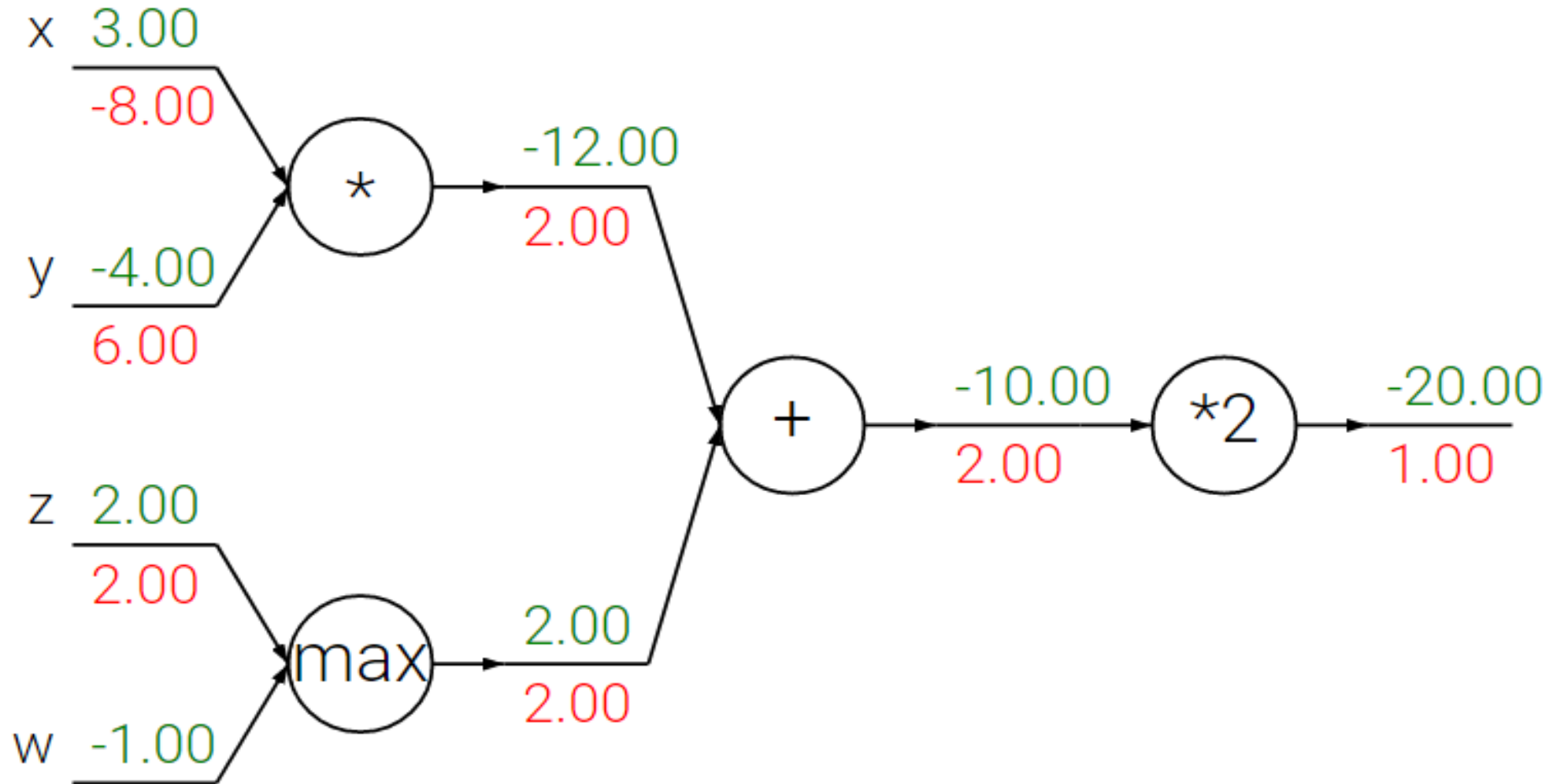
FORWARD AND BACKWARD PASS



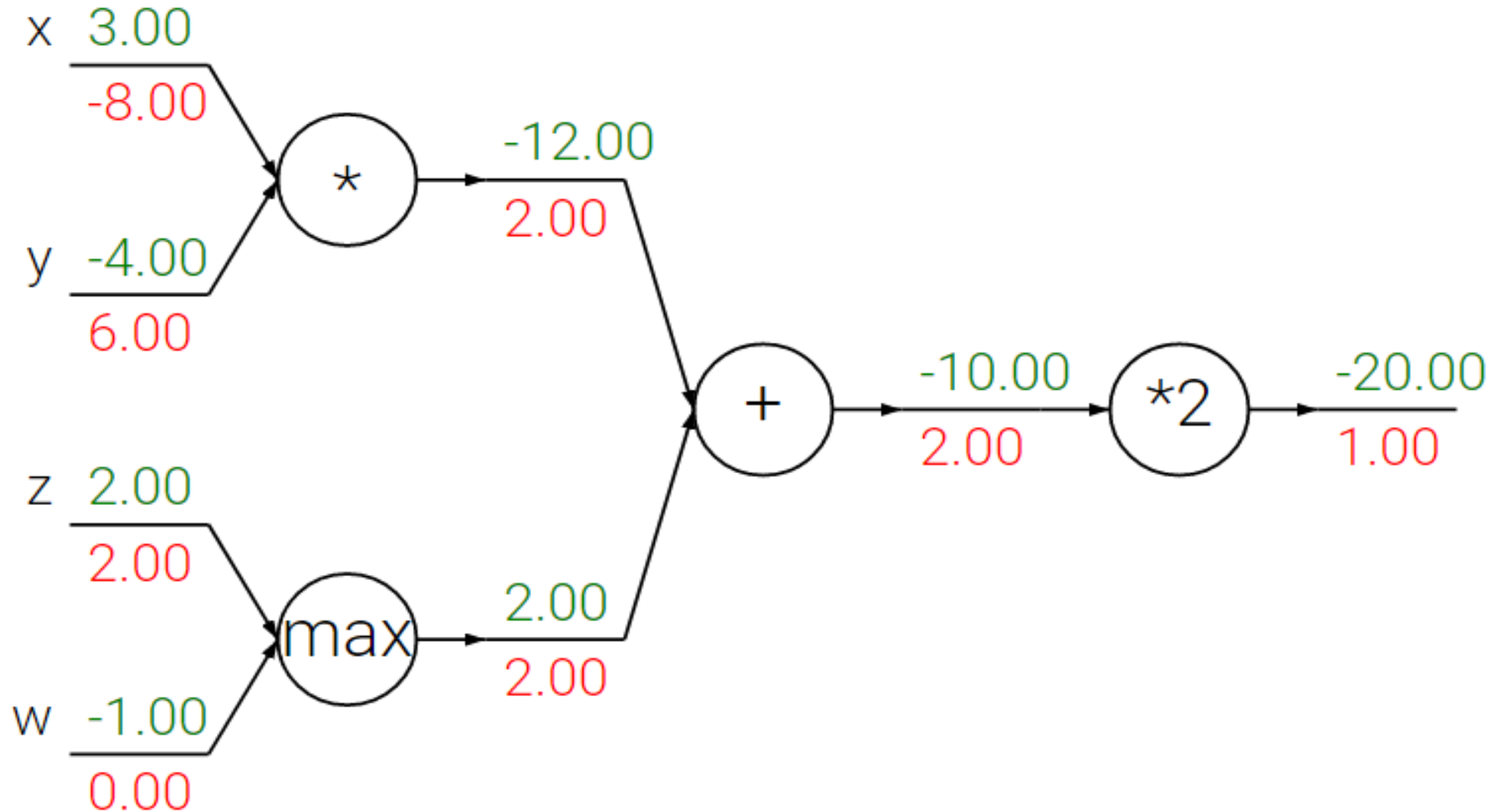
FORWARD AND BACKWARD PASS



FORWARD AND BACKWARD PASS

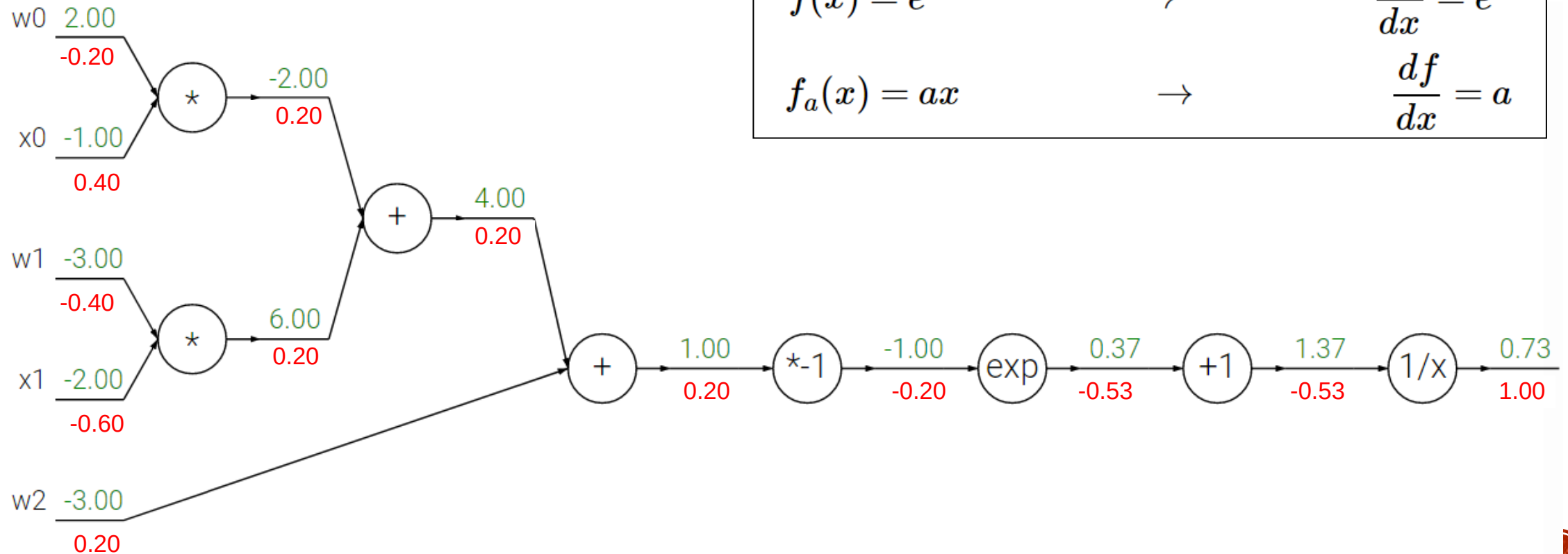


FORWARD AND BACKWARD PASS



SIGMOID EXAMPLE

$$f(w, x) = \frac{1}{1 + e^{-(w_0x_0 + w_1x_1 + w_2)}}$$



$f(x) = \frac{1}{x}$	$\rightarrow$	$\frac{df}{dx} = -1/x^2$
$f_c(x) = c + x$	$\rightarrow$	$\frac{df}{dx} = 1$
$f(x) = e^x$	$\rightarrow$	$\frac{df}{dx} = e^x$
$f_a(x) = ax$	$\rightarrow$	$\frac{df}{dx} = a$



SVM LOSS: GRADIENT

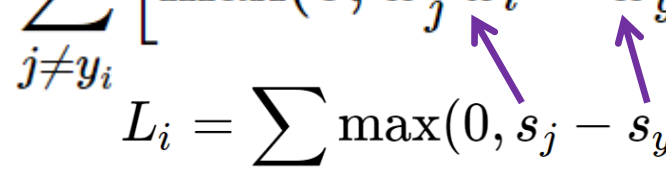
SVM loss function for a single datapoint (without regularization):

$$L_i = \sum_{j \neq y_i} \left[\max(0, w_j^T x_i - w_{y_i}^T x_i + \Delta) \right]$$



SVM LOSS: GRADIENT

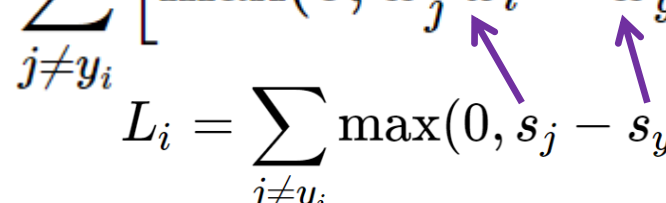
SVM loss function for a single datapoint (without regularization):

$$L_i = \sum_{j \neq y_i} \left[\max(0, w_j^T x_i - w_{y_i}^T x_i + \Delta) \right]$$
$$L_i = \sum_{j \neq y_i} \max(0, s_j - s_{y_i} + \Delta)$$




SVM LOSS: GRADIENT

SVM loss function for a single datapoint (without regularization):

$$L_i = \sum_{j \neq y_i} \left[\max(0, w_j^T x_i - w_{y_i}^T x_i + \Delta) \right]$$
$$L_i = \sum_{j \neq y_i} \max(0, s_j - s_{y_i} + \Delta)$$


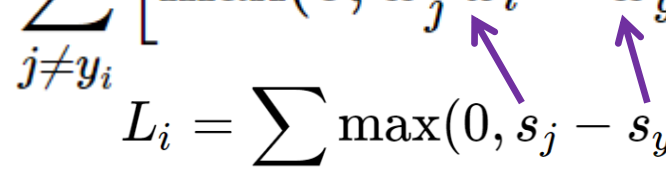
Gradient w.r.t. w_{y_i} :

$$\nabla_{w_{y_i}} L_i = - \left(\sum_{j \neq y_i} 1(w_j^T x_i - w_{y_i}^T x_i + \Delta > 0) \right) x_i$$



SVM LOSS: GRADIENT

SVM loss function for a single datapoint (without regularization):

$$L_i = \sum_{j \neq y_i} \left[\max(0, w_j^T x_i - w_{y_i}^T x_i + \Delta) \right]$$
$$L_i = \sum_{j \neq y_i} \max(0, s_j - s_{y_i} + \Delta)$$


Gradient w.r.t. w_{y_i} :

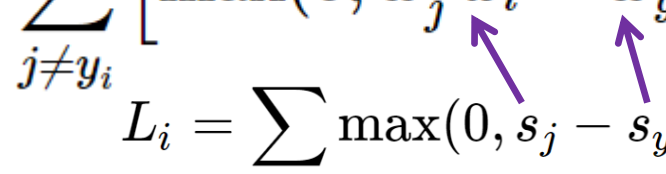
$$\nabla_{w_{y_i}} L_i = - \left(\sum_{j \neq y_i} 1(w_j^T x_i - w_{y_i}^T x_i + \Delta > 0) \right) x_i$$

Count of the number of classes that
didn't meet the desired margin



SVM LOSS: GRADIENT

SVM loss function for a single datapoint (without regularization):

$$L_i = \sum_{j \neq y_i} \left[\max(0, w_j^T x_i - w_{y_i}^T x_i + \Delta) \right]$$
$$L_i = \sum_{j \neq y_i} \max(0, s_j - s_{y_i} + \Delta)$$


Gradient w.r.t. w_{y_i} :

$$\nabla_{w_{y_i}} L_i = - \underbrace{\left(\sum_{j \neq y_i} 1(w_j^T x_i - w_{y_i}^T x_i + \Delta > 0) \right)}_{\text{Count of the number of classes that didn't meet the desired margin}} x_i$$

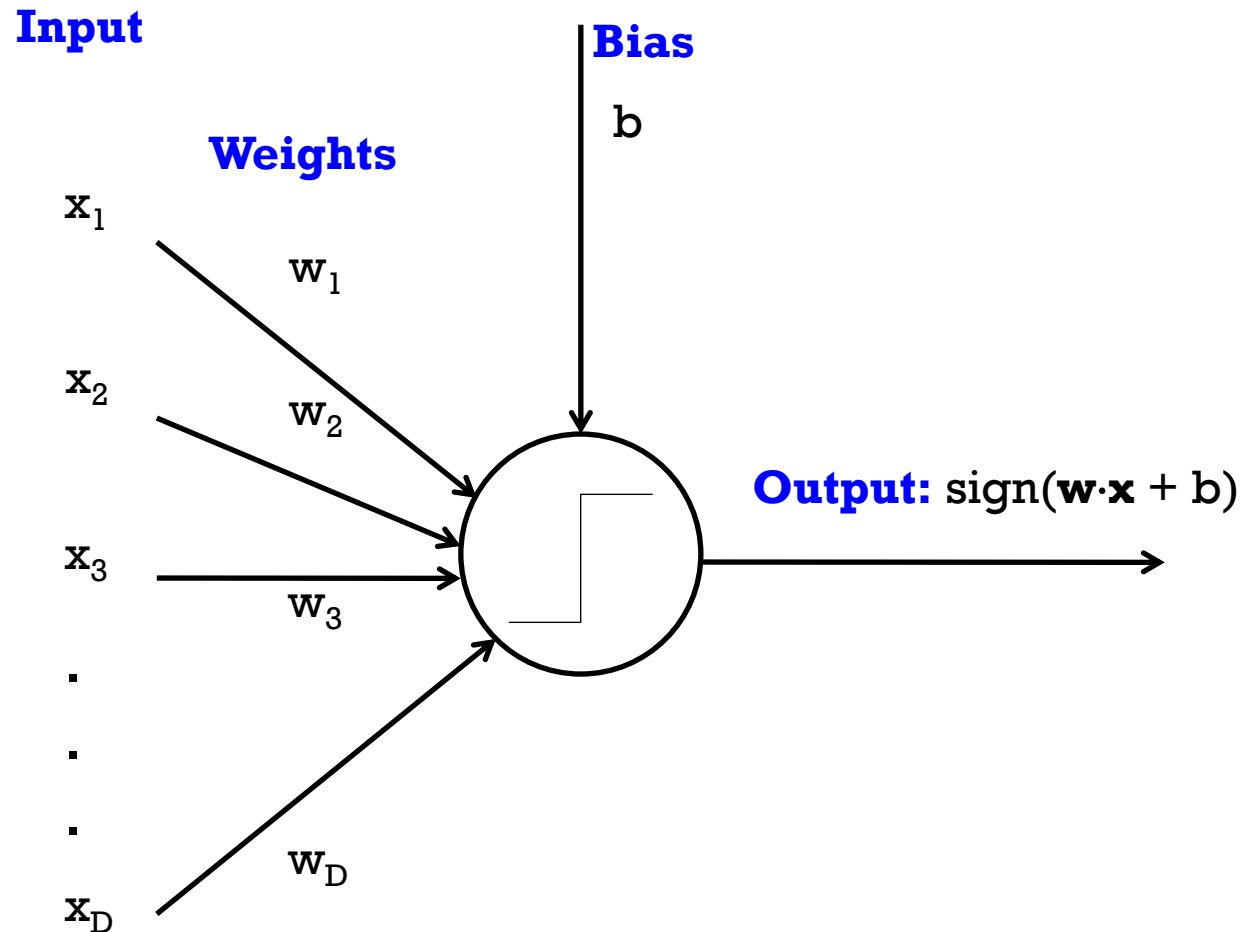
Gradient for the other rows where $j \neq y_i$:

$$\nabla_{w_j} L_i = 1(w_j^T x_i - w_{y_i}^T x_i + \Delta > 0) x_i$$



PERCEPTRON

- Supervised learning of binary classifier



SINGLE NEURON AS A LINEAR CLASSIFIER

Binary Softmax classifier (*Logistic Regression*) $\sigma(\sum_i w_i x_i + b)$



SINGLE NEURON AS A LINEAR CLASSIFIER

Binary Softmax classifier (*Logistic Regression*)

$$\sigma(\sum_i w_i x_i + b)$$



Probability of one of the classes:

$$P(y_i = 1 \mid x_i; w)$$



SINGLE NEURON AS A LINEAR CLASSIFIER

Binary Softmax classifier (*Logistic Regression*)

$$\sigma(\sum_i w_i x_i + b)$$



Probability of one of the classes:

$$P(y_i = 1 \mid x_i; w)$$

Probability of the other class would be:

$$P(y_i = 0 \mid x_i; w) = 1 - P(y_i = 1 \mid x_i; w)$$



SINGLE NEURON AS A LINEAR CLASSIFIER

Binary Softmax classifier (*Logistic Regression*)

$$\sigma(\sum_i w_i x_i + b)$$



Probability of one of the classes:

$$P(y_i = 1 \mid x_i; w)$$

Probability of the other class would be:

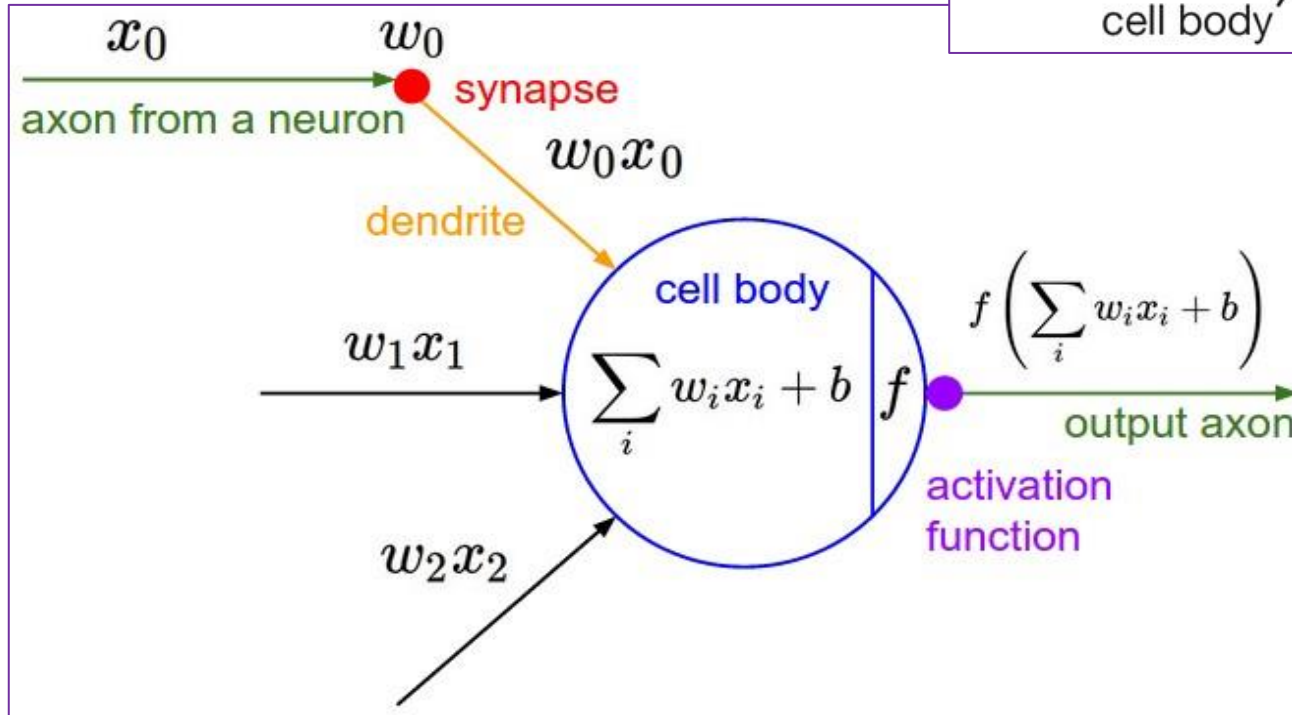
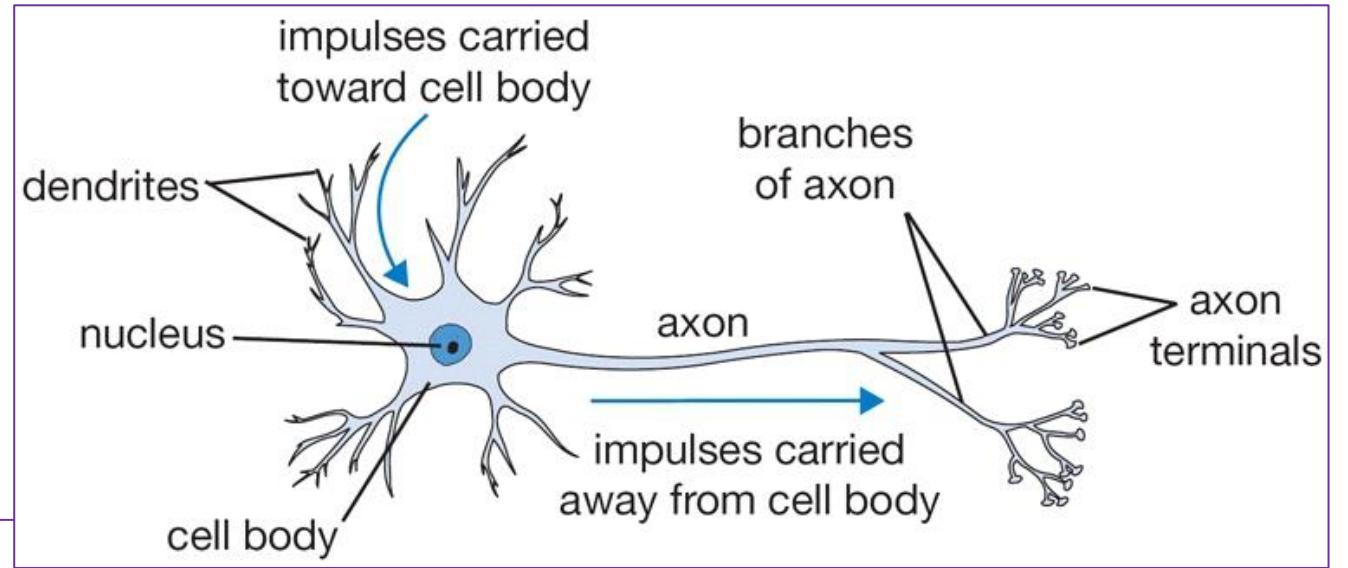
$$P(y_i = 0 \mid x_i; w) = 1 - P(y_i = 1 \mid x_i; w)$$

Binary SVM classifier:

Alternatively, we could attach a max-margin hinge loss to the output of the neuron and train it to become a binary Support Vector Machine.

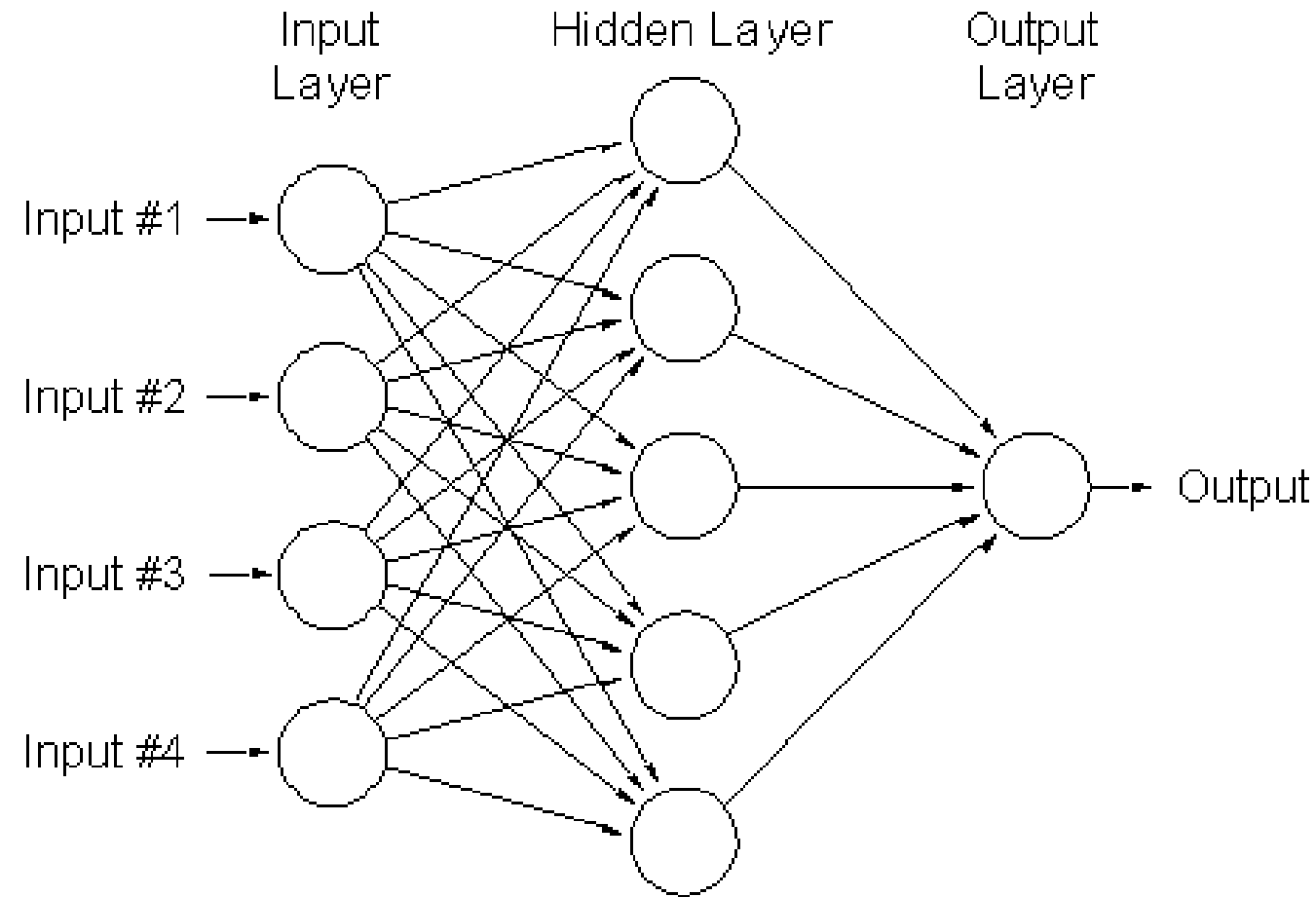


LOOSE INSPIRATION: HUMAN NEURONS



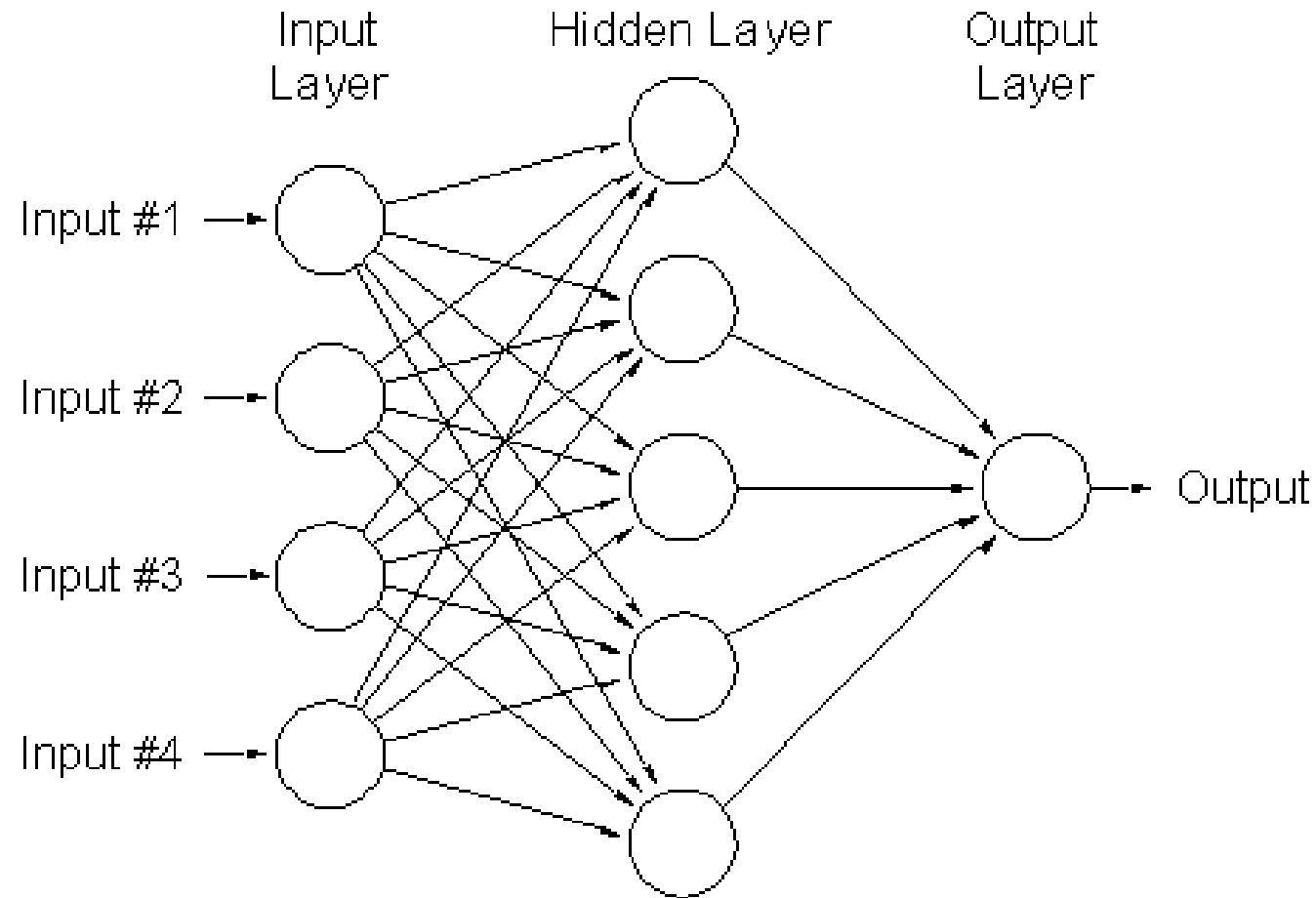
MULTI-LAYER NEURAL NETWORKS

- Network with a hidden layer:



MULTI-LAYER NEURAL NETWORKS

- Network with a hidden layer:

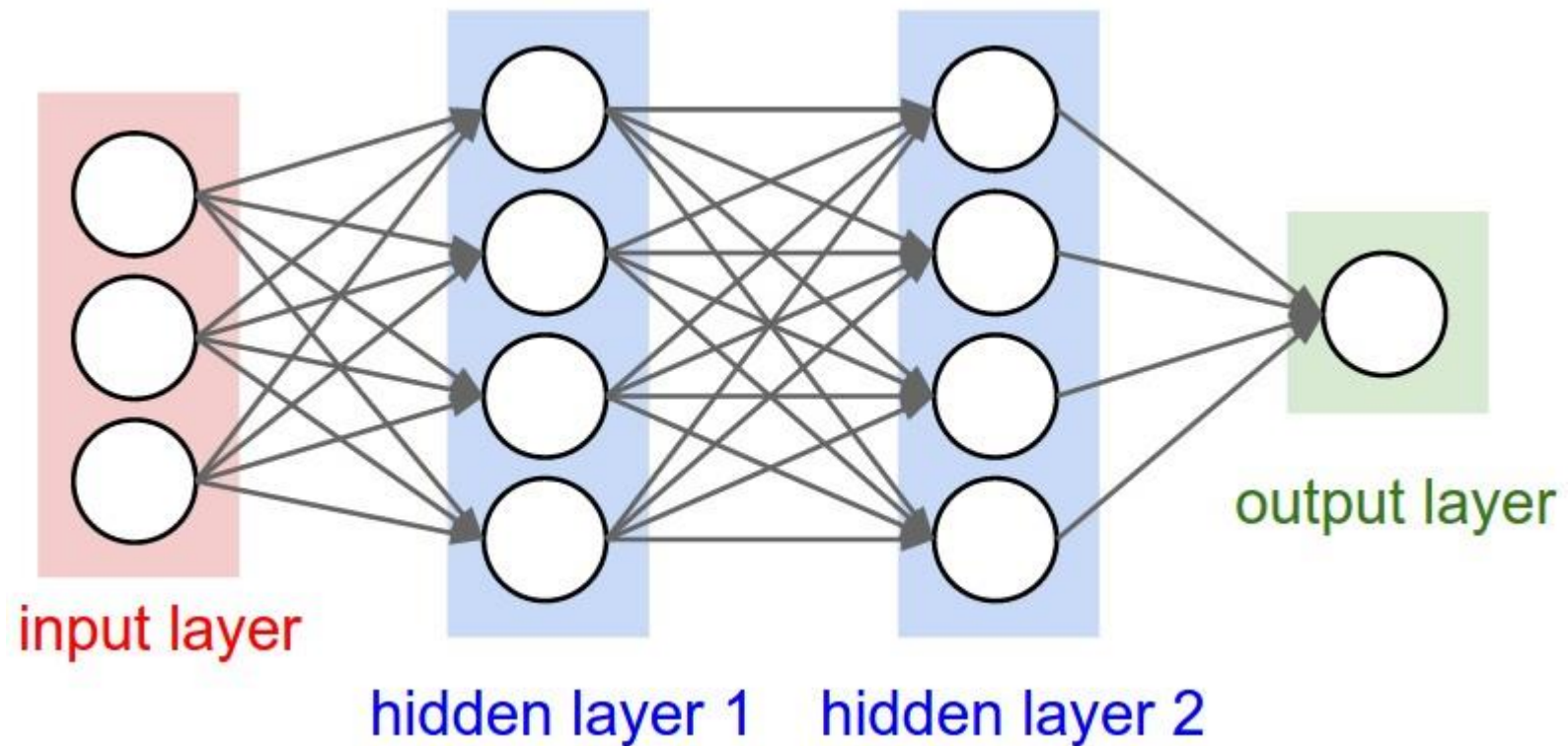


- Can represent nonlinear functions (provided each perceptron has a nonlinearity)

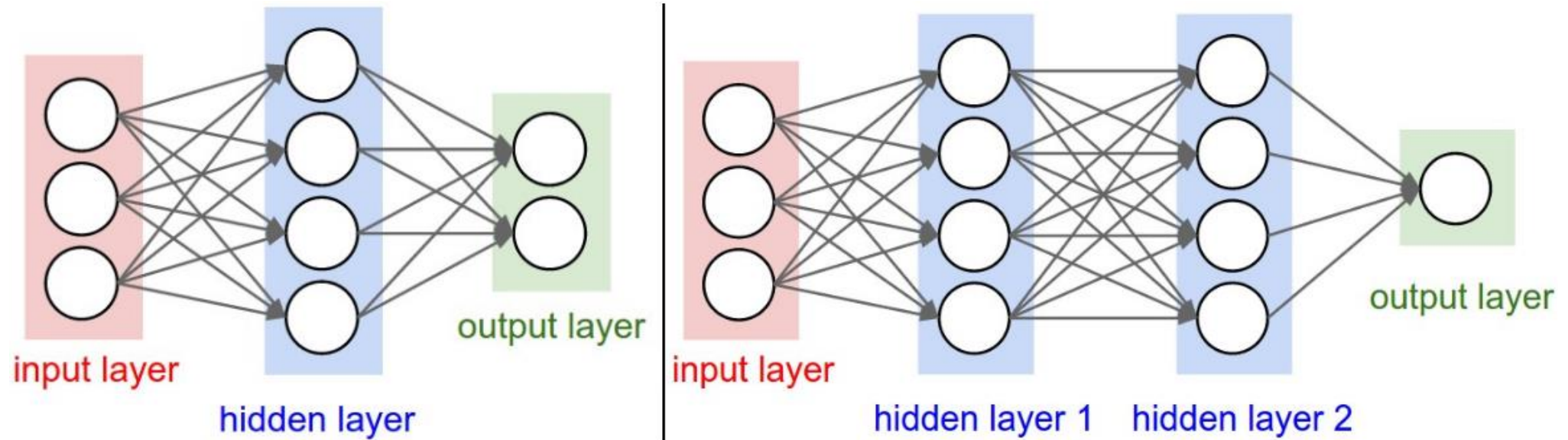


MULTI-LAYER NEURAL NETWORKS

- Beyond a single hidden layer:



SIZING NEURAL NETWORKS



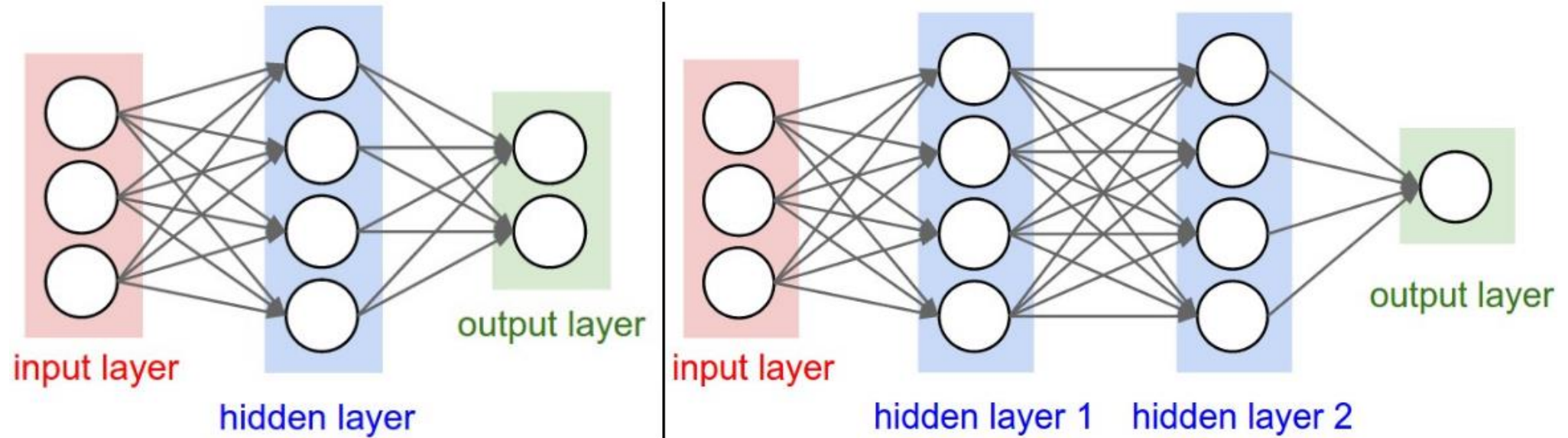
First network (left):

No. of neurons (not counting the inputs):

No. of learnable parameters:



SIZING NEURAL NETWORKS



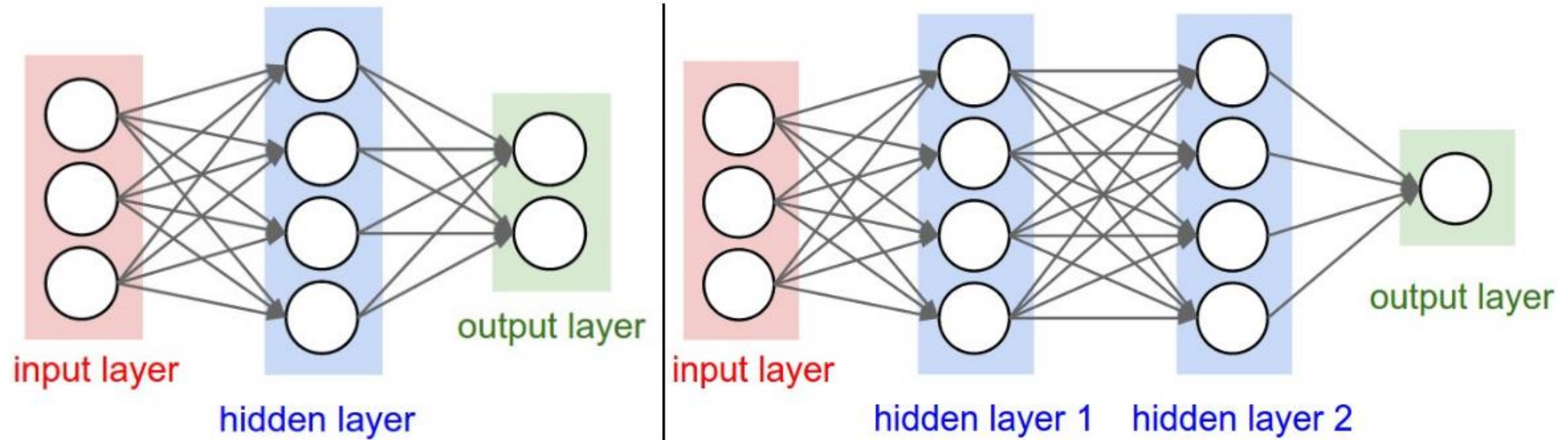
First network (left):

No. of neurons (not counting the inputs): $4 + 2 = 6$

No. of learnable parameters:



SIZING NEURAL NETWORKS



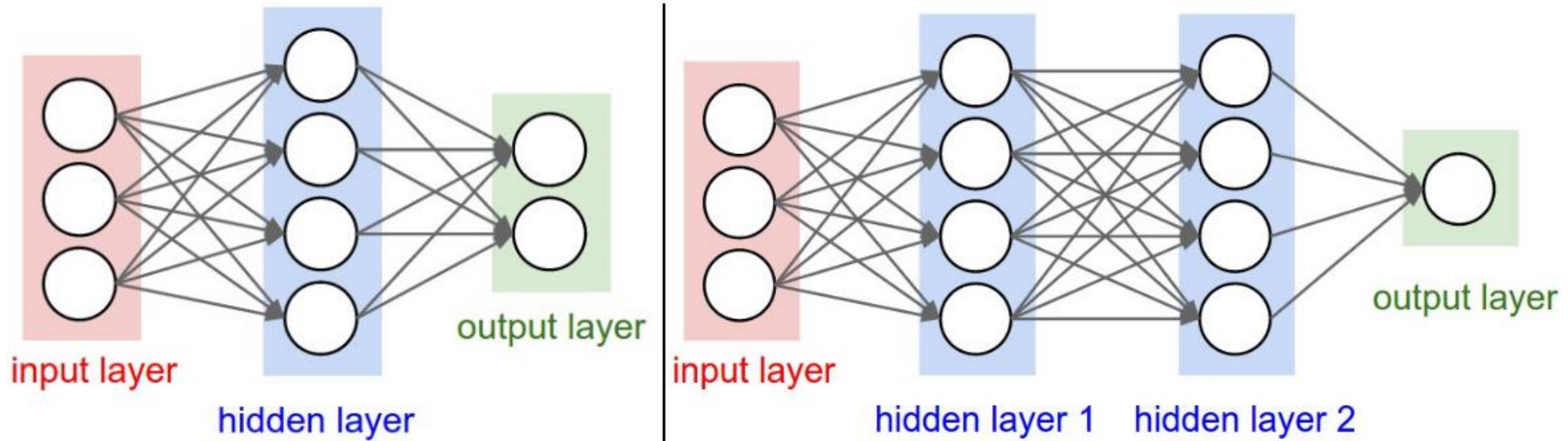
First network (left):

No. of neurons (not counting the inputs): $4 + 2 = 6$

No. of learnable parameters: $[3 \times 4] + [4 \times 2] = 20$ weights +
 $4 + 2 = 6$ biases = 26.



SIZING NEURAL NETWORKS



First network (left):

No. of neurons (not counting the inputs): $4 + 2 = 6$

No. of learnable parameters: $[3 \times 4] + [4 \times 2] = 20$ weights +
 $4 + 2 = 6$ biases = 26.

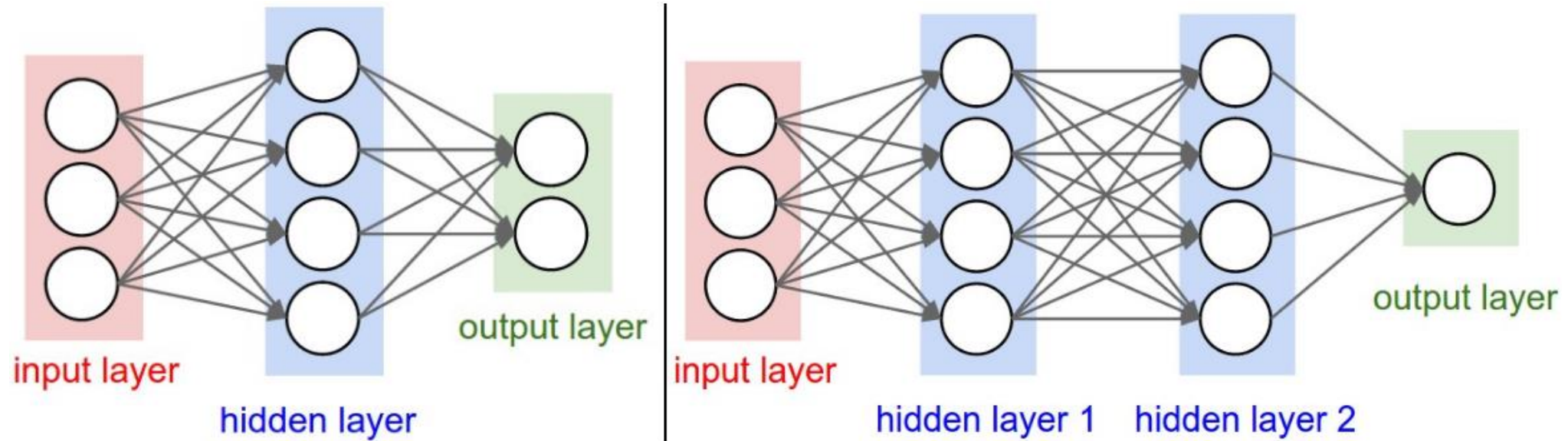
Second network (right):

No. of neurons (not counting the inputs):

No. of learnable parameters:



SIZING NEURAL NETWORKS



First network (left):

No. of neurons (not counting the inputs): $4 + 2 = 6$

No. of learnable parameters: $[3 \times 4] + [4 \times 2] = 20$ weights +
 $4 + 2 = 6$ biases = 26.

Second network (right):

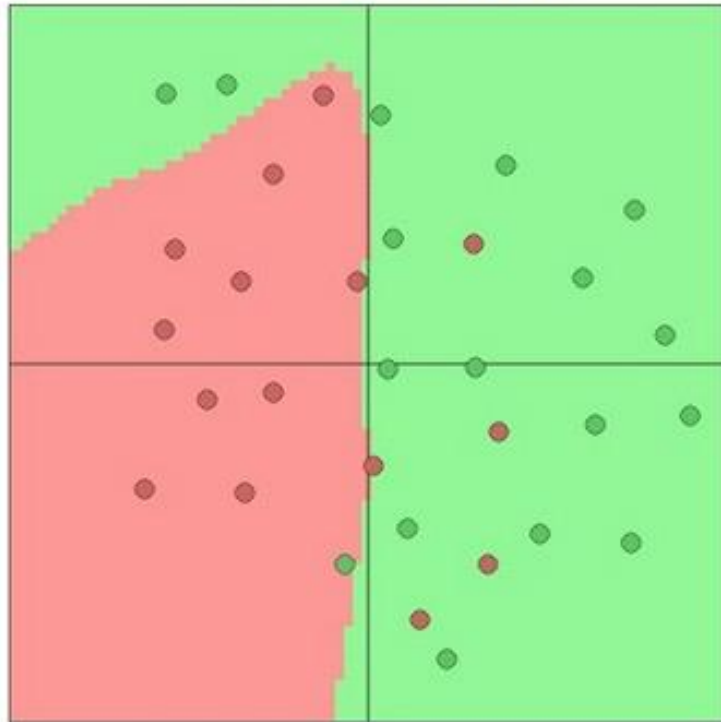
No. of neurons (not counting the inputs): $4 + 4 + 1 = 9$

No. of learnable parameters: $[3 \times 4] + [4 \times 4] + [4 \times 1] = 32$ weights +
 $4 + 4 + 1 = 9$ biases = 41.

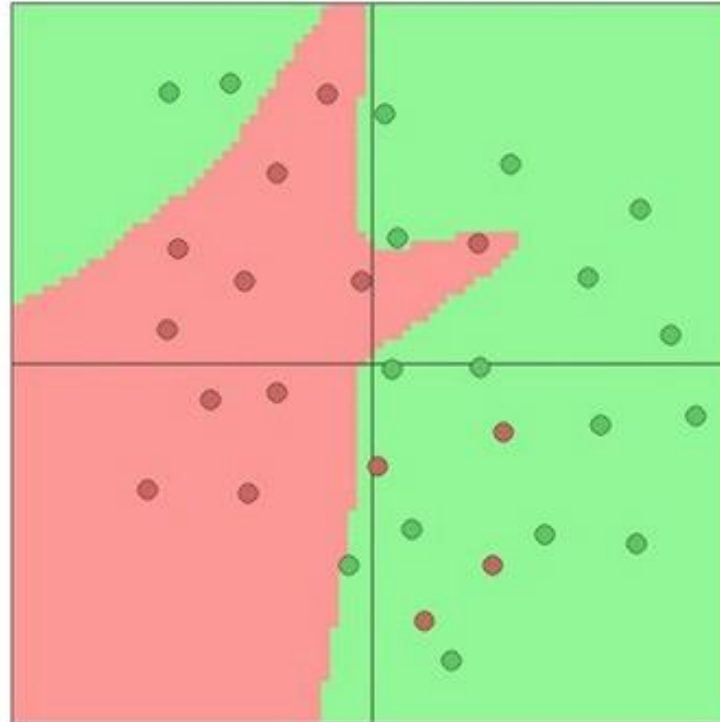


MULTI-LAYER NEURAL NETWORKS

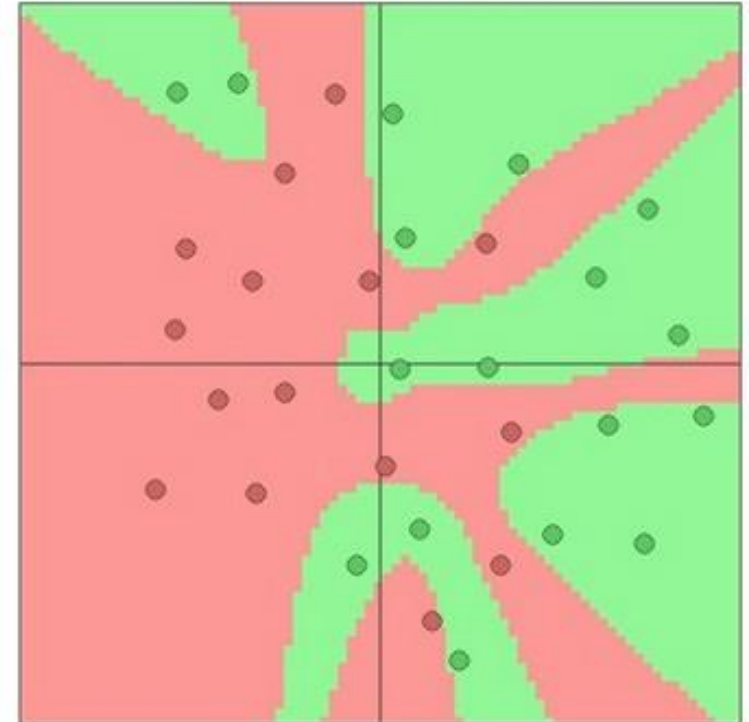
3 hidden neurons



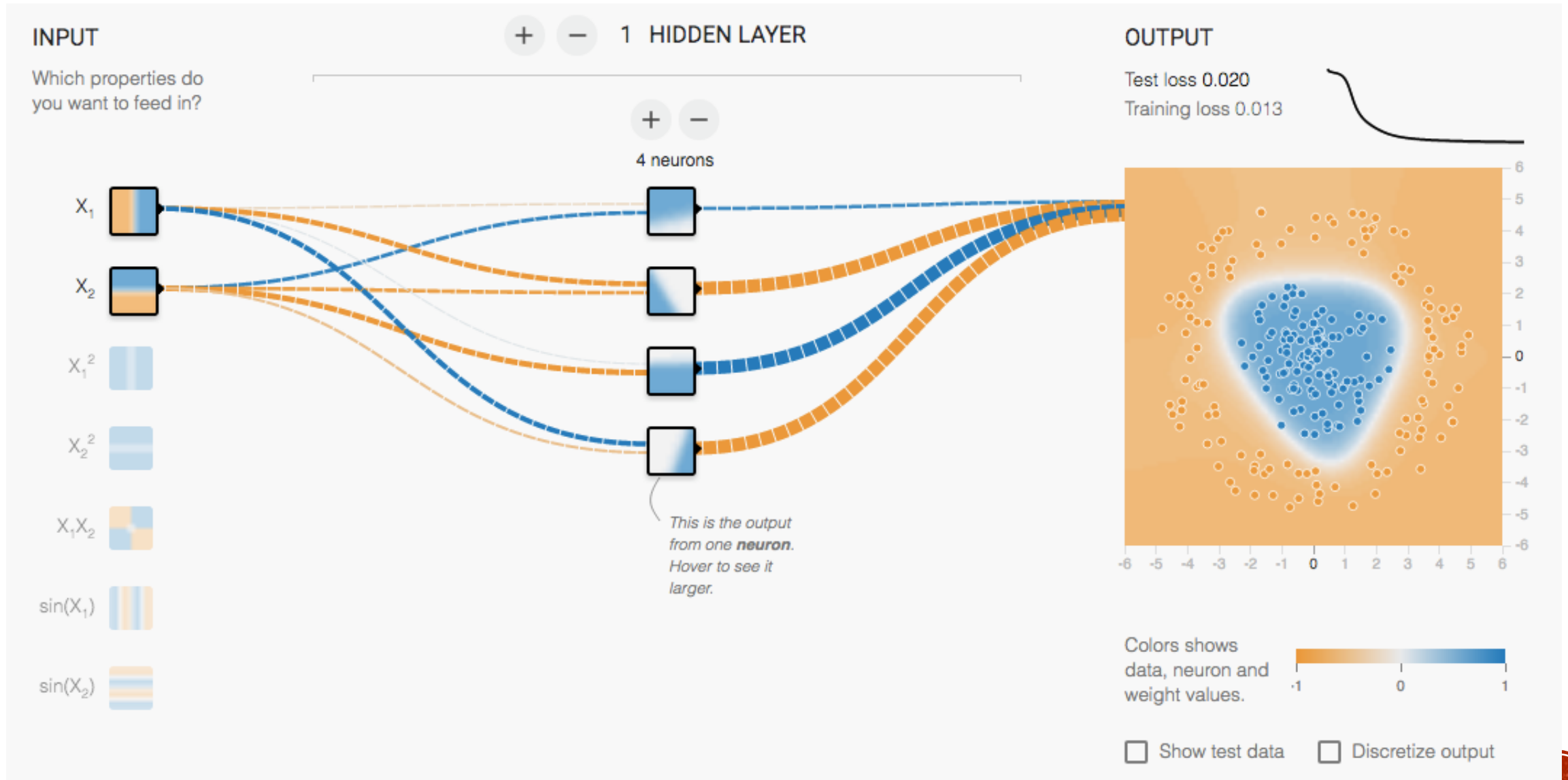
6 hidden neurons



20 hidden neurons



MULTI-LAYER NETWORK DEMO



<http://playground.tensorflow.org/>



TRAINING OF MULTI-LAYER NETWORKS

- Find network weights to minimize the error between true and estimated outputs of training examples:

$$E(\mathbf{w}) = \sum_{j=1}^N (y_j - f_{\mathbf{w}}(\mathbf{x}_j))^2$$



TRAINING OF MULTI-LAYER NETWORKS

- Find network weights to minimize the error between true and estimated outputs of training examples:

$$E(\mathbf{w}) = \sum_{j=1}^N (y_j - f_{\mathbf{w}}(\mathbf{x}_j))^2$$

- Update weights by **gradient descent**: $\mathbf{w} \leftarrow \mathbf{w} - \alpha \frac{\partial E}{\partial \mathbf{w}}$



TRAINING OF MULTI-LAYER NETWORKS

- Find network weights to minimize the error between true and estimated outputs of training examples:

$$E(\mathbf{w}) = \sum_{j=1}^N (y_j - f_{\mathbf{w}}(\mathbf{x}_j))^2$$

- Update weights by **gradient descent**: $\mathbf{w} \leftarrow \mathbf{w} - \alpha \frac{\partial E}{\partial \mathbf{w}}$
- **Back-propagation**: gradients are computed in the direction from output to input layers and combined using chain rule



NEURAL NETWORKS: PROS AND CONS

Pros

- Flexible and general function approximation framework
- Can build extremely powerful models by adding more layers



NEURAL NETWORKS: PROS AND CONS

Pros

- Flexible and general function approximation framework
- Can build extremely powerful models by adding more layers

Cons

- Hard to analyze theoretically (e.g., training is prone to local optima)
- Huge amount of training data, computing power may be required to get good performance
- The space of implementation choices are huge (network architectures, parameters)



Activation Functions

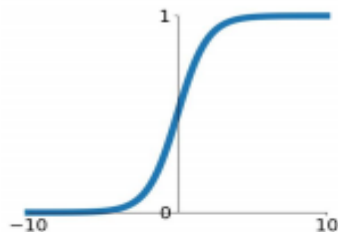


NON-LINEARITY LAYER

Activation Functions

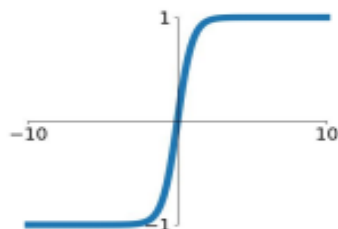
Sigmoid

$$\sigma(x) = \frac{1}{1+e^{-x}}$$



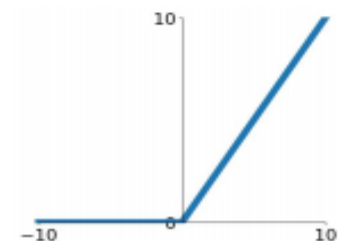
tanh

$$\tanh(x)$$



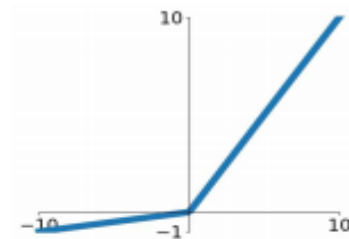
ReLU

$$\max(0, x)$$



Leaky ReLU

$$\max(0.1x, x)$$

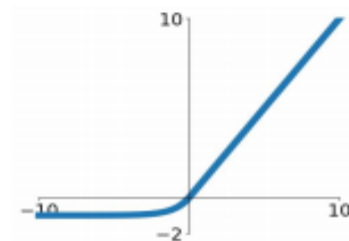


Maxout

$$\max(w_1^T x + b_1, w_2^T x + b_2)$$

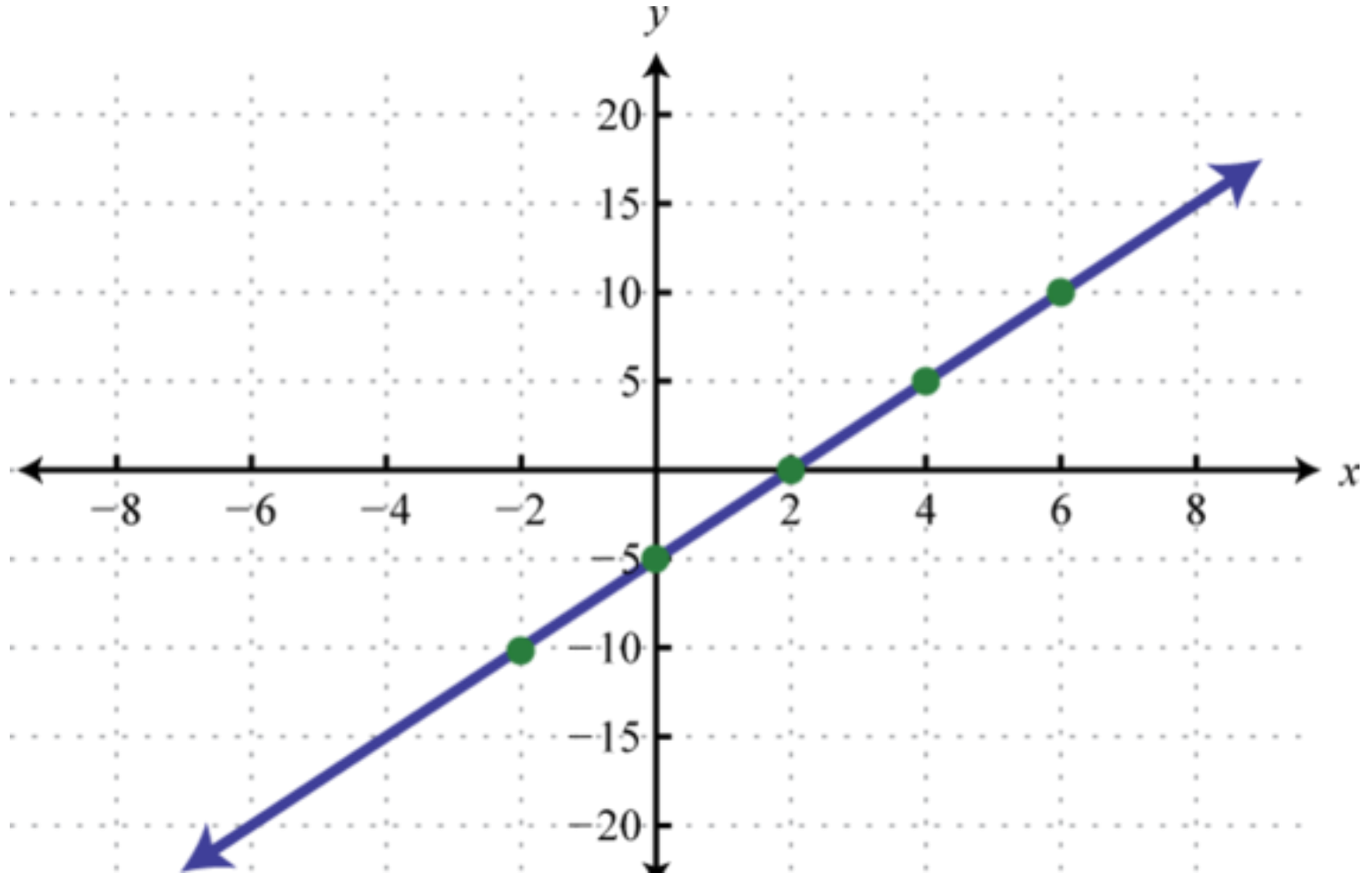
ELU

$$\begin{cases} x & x \geq 0 \\ \alpha(e^x - 1) & x < 0 \end{cases}$$



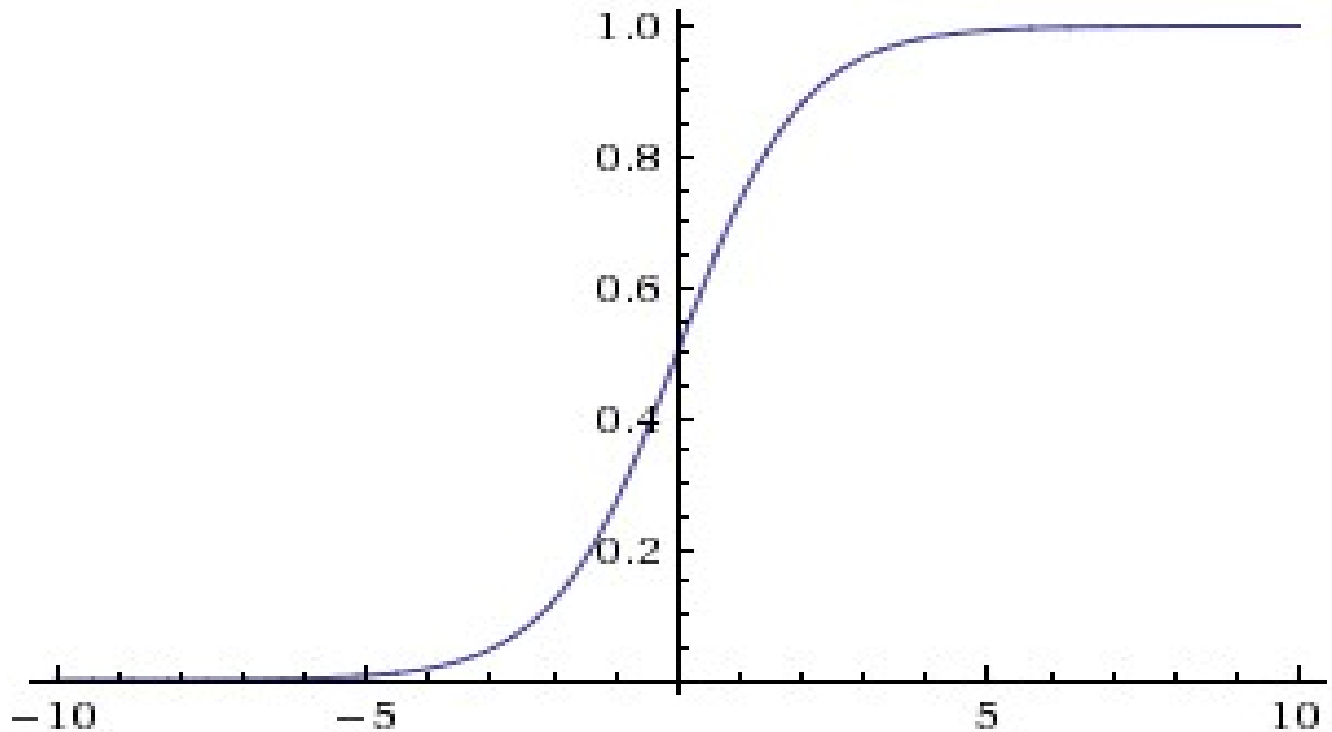
ACTIVATION FUNCTIONS: LINEAR

- Simplest activation function
- Does not include any non-linearity.



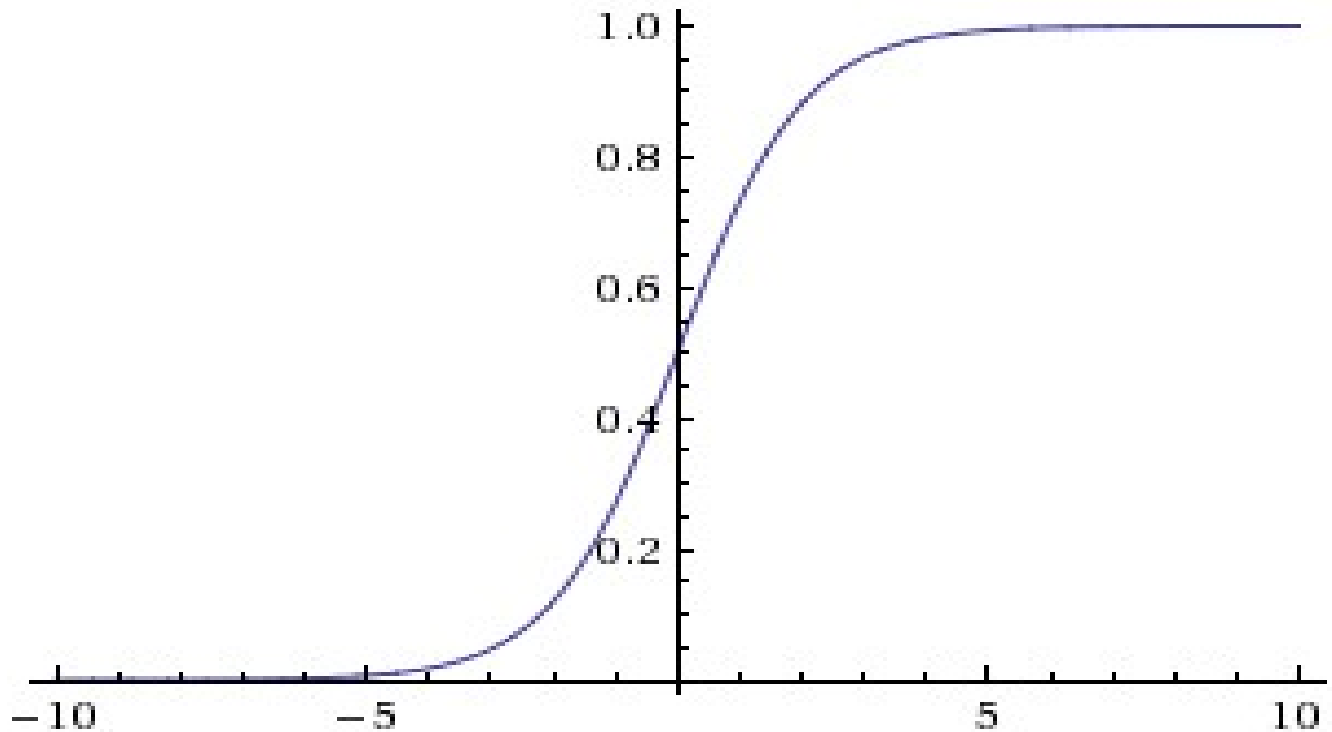
ACTIVATION FUNCTIONS: SIGMOID

$$\sigma(x) = 1/(1 + e^{-x})$$



ACTIVATION FUNCTIONS: SIGMOID

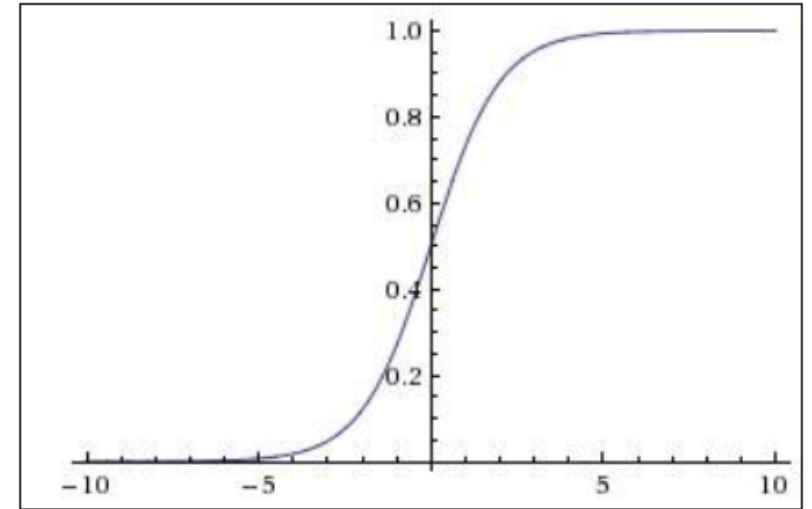
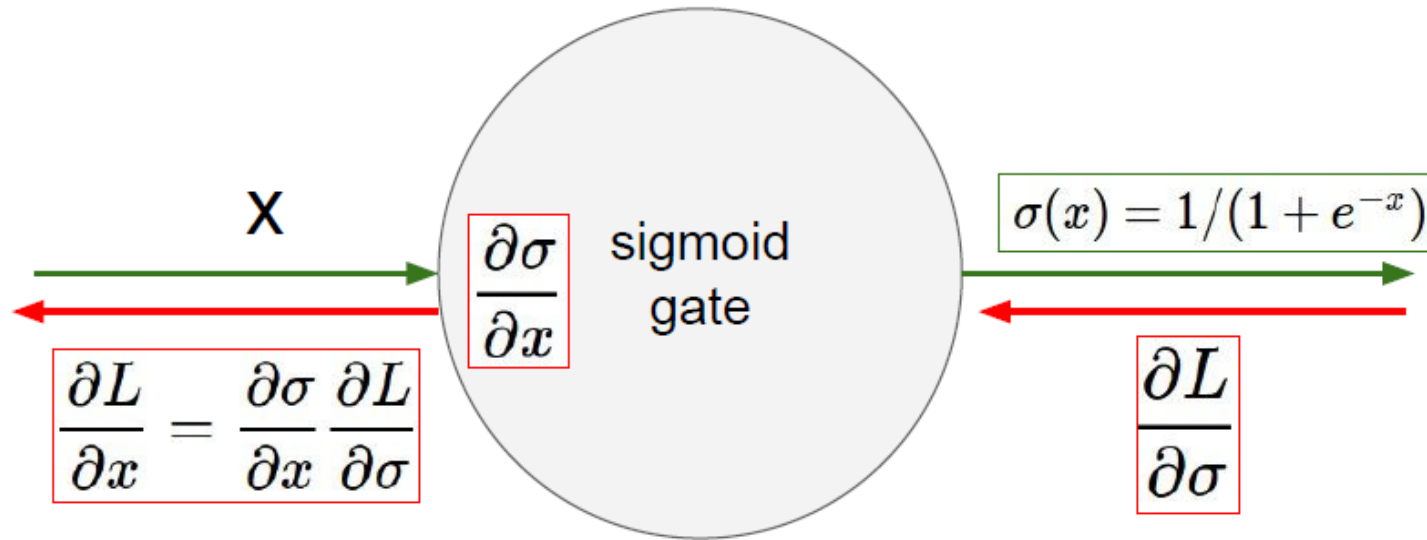
$$\sigma(x) = 1/(1 + e^{-x})$$



- Sigmoids saturate and kill gradients.



ACTIVATION FUNCTIONS: SIGMOID



What happens when $x = -10$?

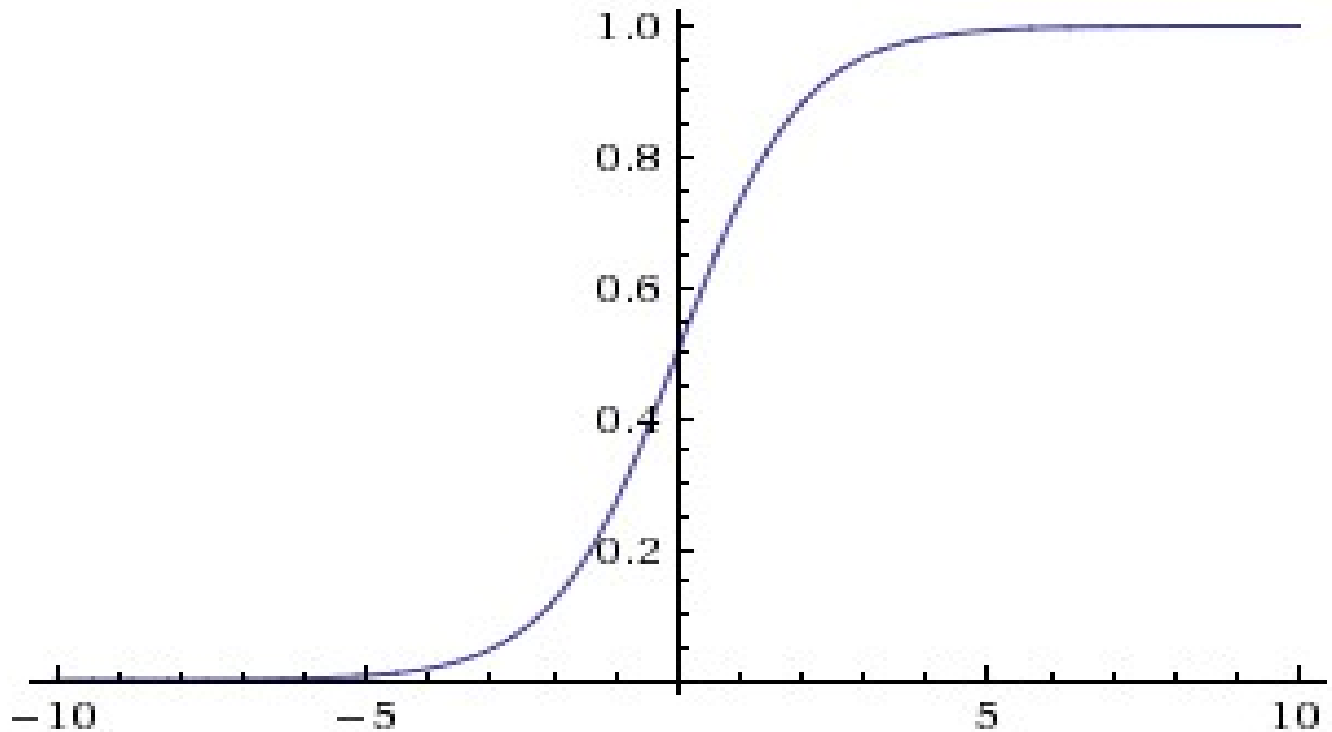
What happens when $x = 0$?

What happens when $x = 10$?



ACTIVATION FUNCTIONS: SIGMOID

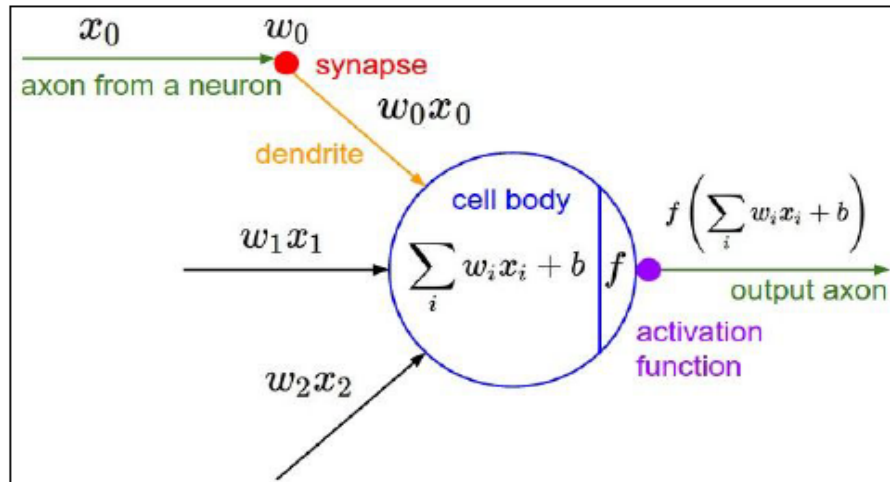
$$\sigma(x) = 1 / (1 + e^{-x})$$



- Sigmoids saturate and kill gradients.
- Sigmoid outputs are not zero-centered.



Consider what happens when the input to a neuron (x) is always positive:



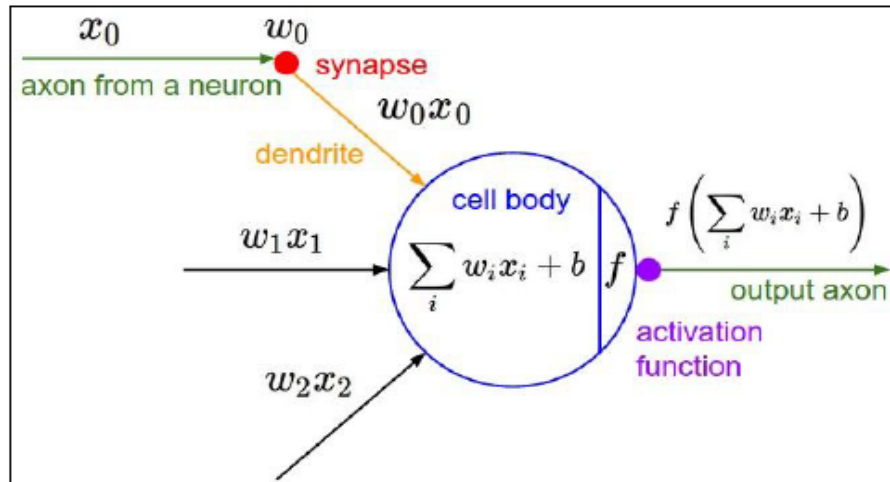
$$f\left(\sum_i w_i x_i + b\right)$$

What can we say about the gradients on $\mathbf{w}$?

Always all positive or all negative
(this is also why you want zero-mean data!)



Consider what happens when the input to a neuron (x) is always positive:



$$f\left(\sum_i w_i x_i + b\right)$$

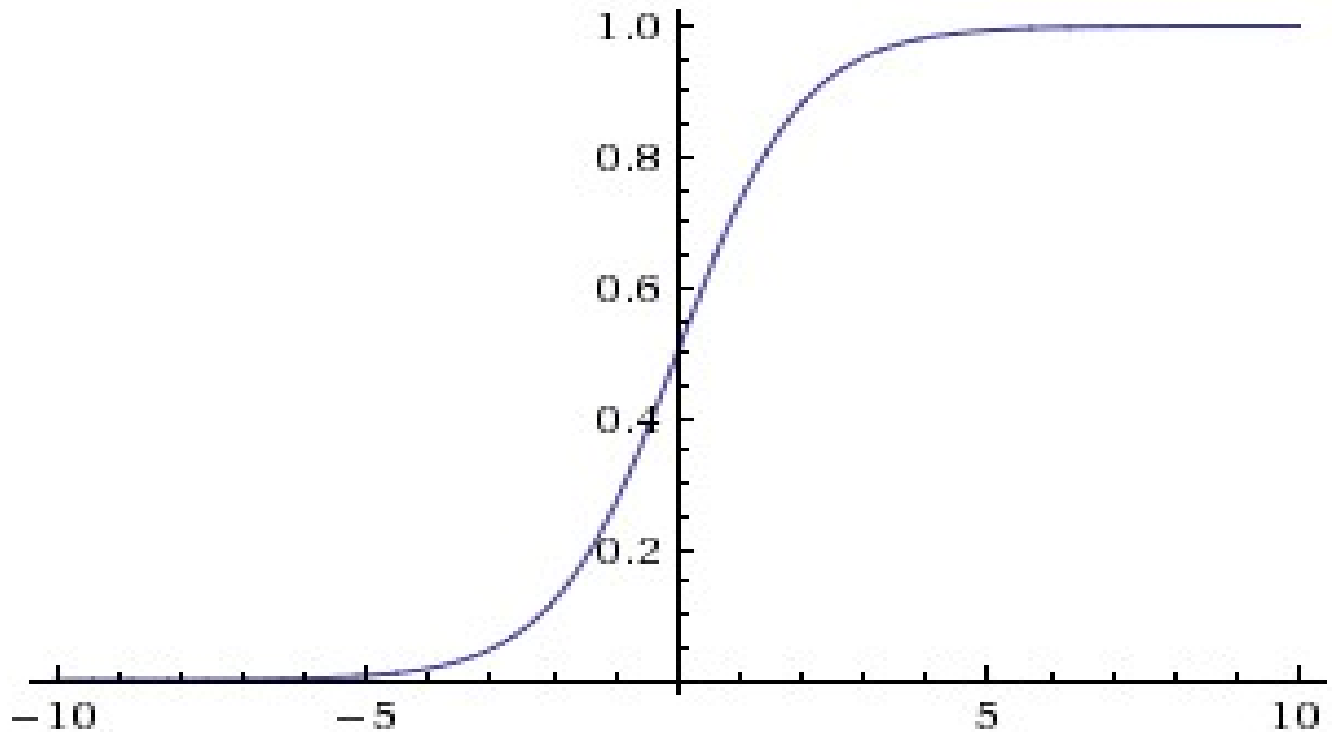
What can we say about the gradients on $\mathbf{w}$?

Always all positive or all negative
(this is also why you want zero-mean data!)



ACTIVATION FUNCTIONS: SIGMOID

$$\sigma(x) = 1/(1 + e^{-x})$$

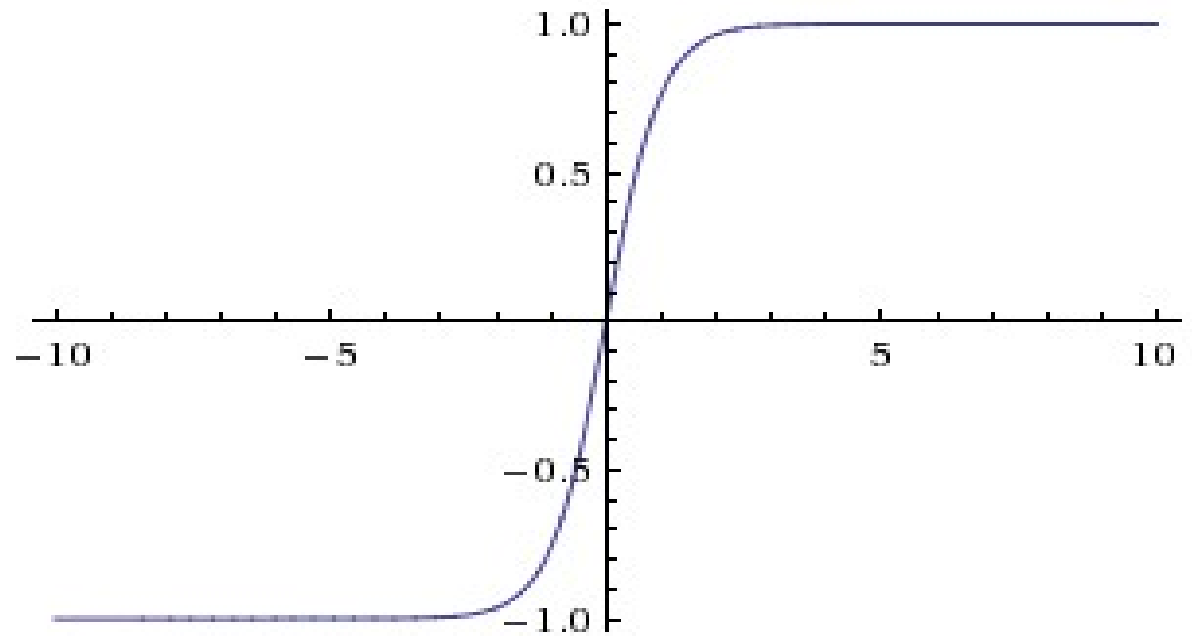


- Sigmoids saturate and kill gradients.
- Sigmoid outputs are not zero-centered.
- $\text{Exp}()$ is a bit compute expensive.



ACTIVATION FUNCTIONS: TANH

$$\tanh(x) = \frac{e^x - e^{-x}}{e^x + e^{-x}}$$



[LeCun et al., 1991]

Source: <http://cs231n.github.io>



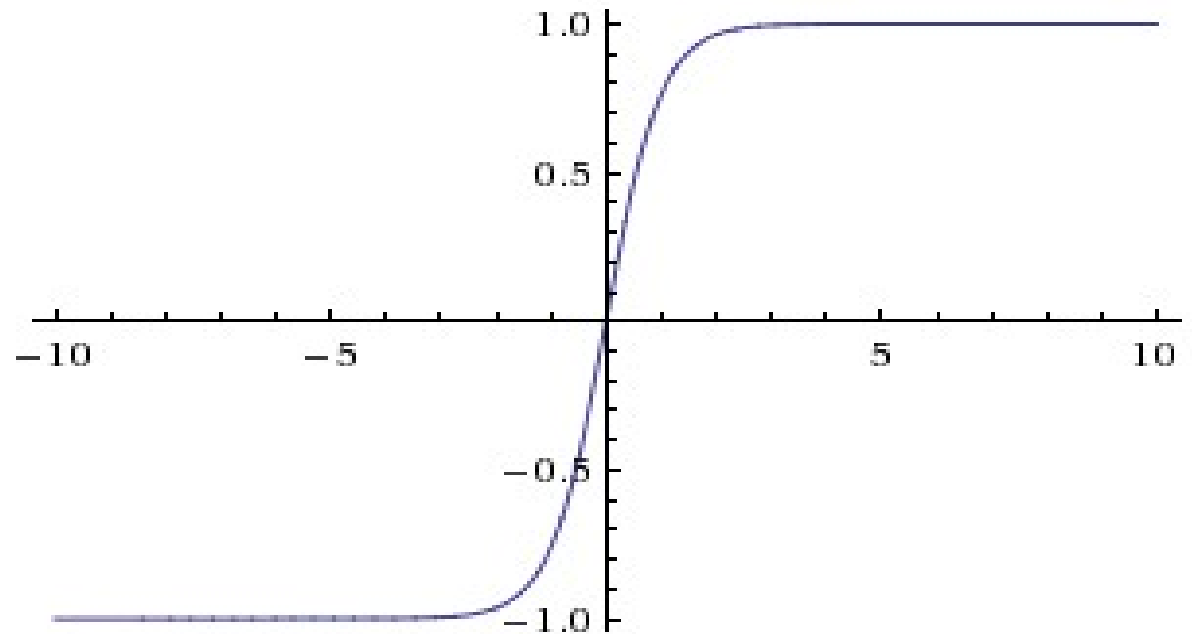
ACTIVATION FUNCTIONS: TANH

$$\tanh(x) = \frac{e^x - e^{-x}}{e^x + e^{-x}}$$

tanh neuron is simply a scaled sigmoid neuron

$$\tanh(x) = 2\sigma(2x) - 1.$$

↑
Sigmoid



[LeCun et al., 1991]

Source: <http://cs231n.github.io>



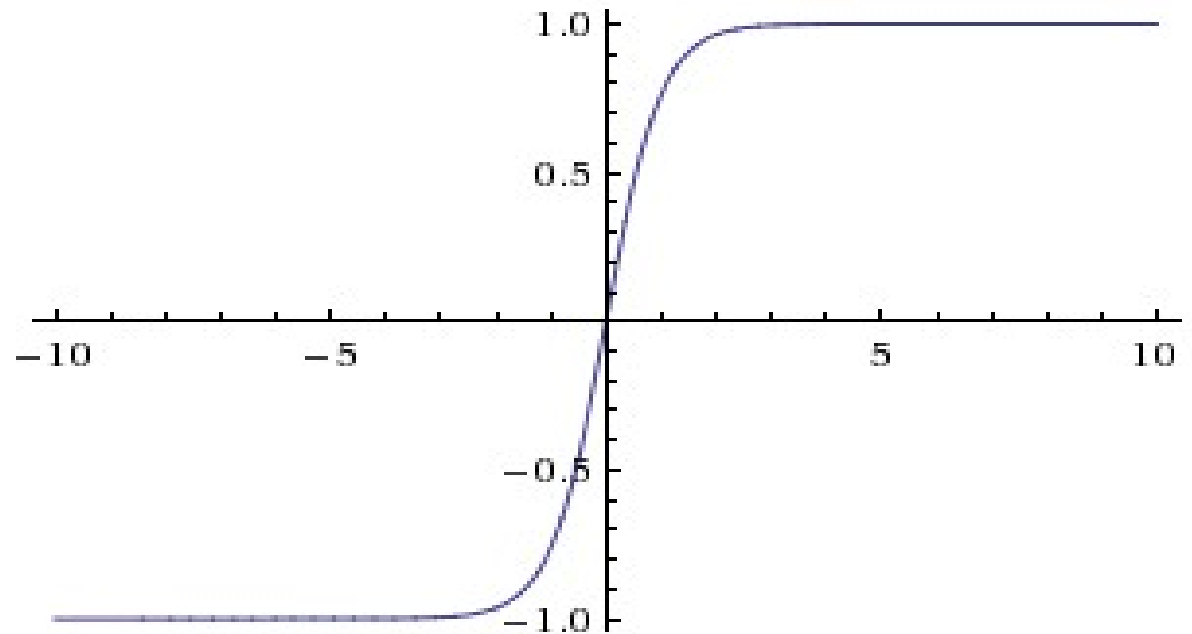
ACTIVATION FUNCTIONS: TANH

$$\tanh(x) = \frac{e^x - e^{-x}}{e^x + e^{-x}}$$

tanh neuron is simply a scaled sigmoid neuron

$$\tanh(x) = 2\sigma(2x) - 1.$$

↑
Sigmoid



Like the sigmoid neuron, its activations saturate.

Unlike the sigmoid neuron its output is zero-centered.

In practice the *tanh non-linearity is always preferred to the sigmoid nonlinearity.*

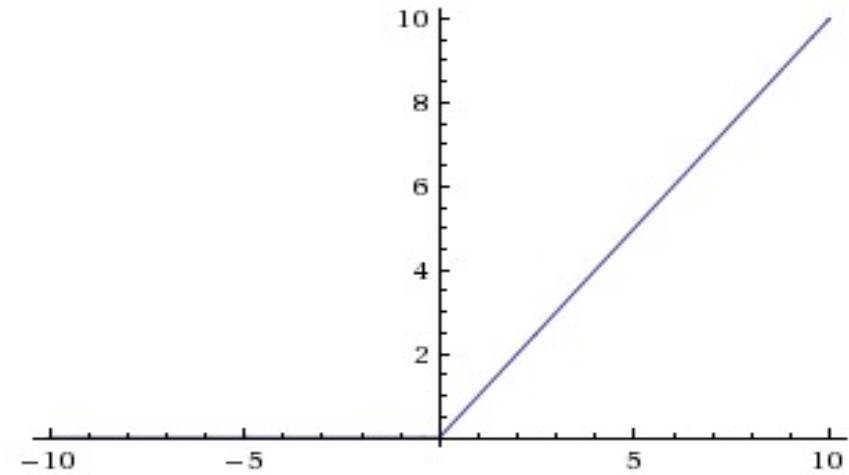
[LeCun et al., 1991]

Source: <http://cs231n.github.io>



ACTIVATION FUNCTIONS: RELU

$$f(x) = \max(0, x)$$



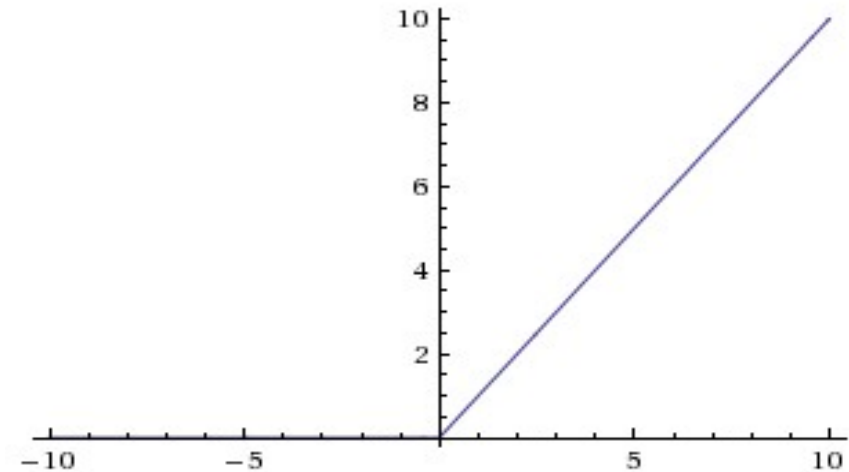
[Krizhevsky et al., 2012]

Source: <http://cs231n.github.io>



ACTIVATION FUNCTIONS: RELU

$$f(x) = \max(0, x)$$



ReLU is 6 times faster in the convergence of stochastic gradient descent compared to the sigmoid/tanh (Krizhevsky et al.).

ReLU is simple as compared to tanh/sigmoid that involve expensive operations (exponentials, etc.)

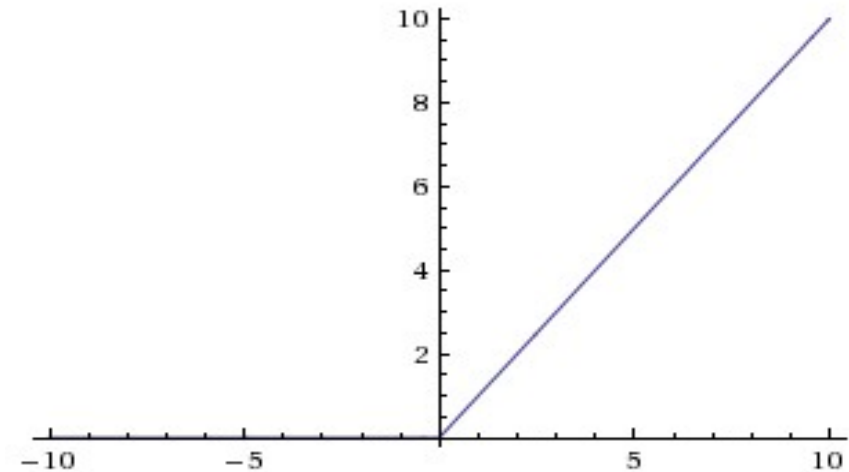
[Krizhevsky et al., 2012]

Source: <http://cs231n.github.io>



ACTIVATION FUNCTIONS: RELU

$$f(x) = \max(0, x)$$



ReLU is 6 times faster in the convergence of stochastic gradient descent compared to the sigmoid/tanh (Krizhevsky et al.).

ReLU is simple as compared to tanh/sigmoid that involve expensive operations (exponentials, etc.)

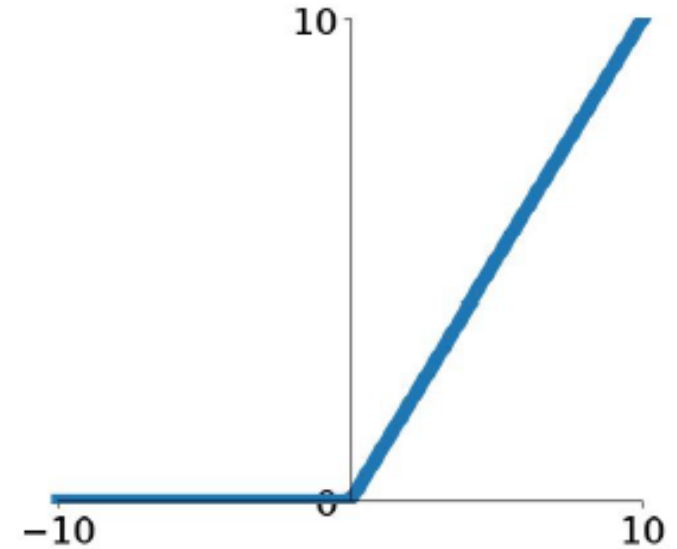
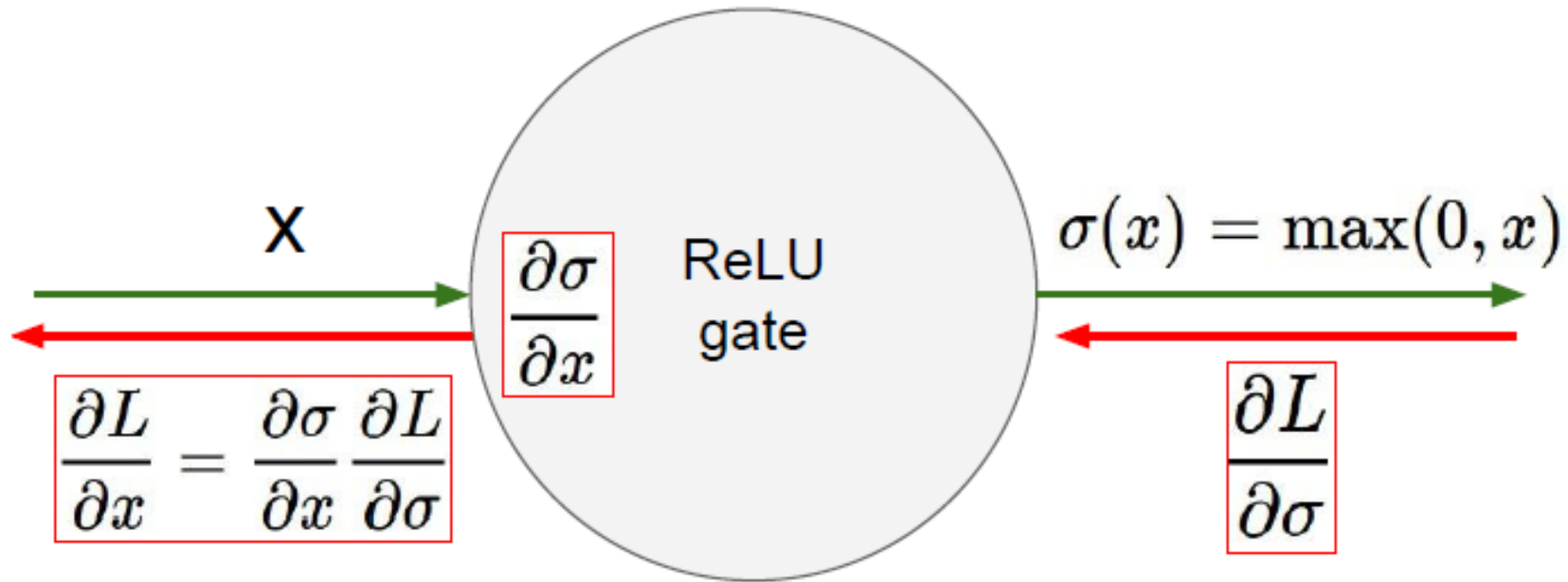
Dying ReLU problem: a large gradient flowing through a ReLU neuron could cause the weights to update in such a way that the neuron will never activate on any datapoint again.

[Krizhevsky et al., 2012]

Source: <http://cs231n.github.io>



ACTIVATION FUNCTIONS: RELU



What happens when $x = -10$?

What happens when $x = 0$?

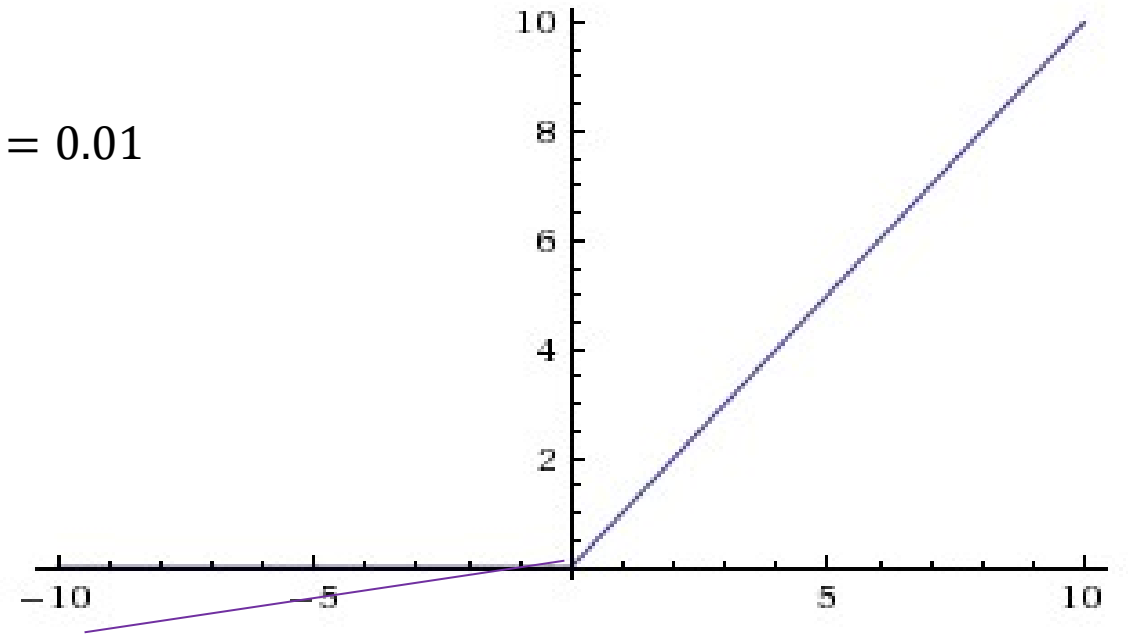
What happens when $x = 10$?



ACTIVATION FUNCTIONS: LEAKY RELU

$$f(x) = \begin{cases} \alpha x, & x < 0 \\ x, & x \geq 0 \end{cases}$$

$$\alpha = 0.01$$



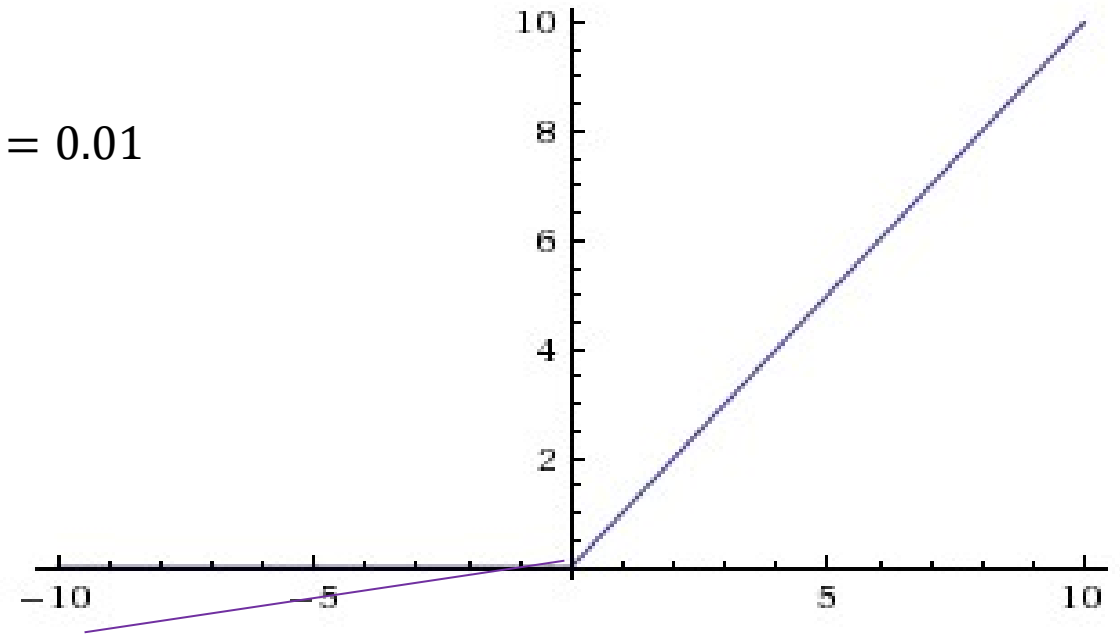
[Mass et al., 2013]

Source: <http://cs231n.github.io>



ACTIVATION FUNCTIONS: LEAKY RELU

$$f(x) = \begin{cases} \alpha x, & x < 0 \\ x, & x \geq 0 \end{cases} \quad \alpha = 0.01$$

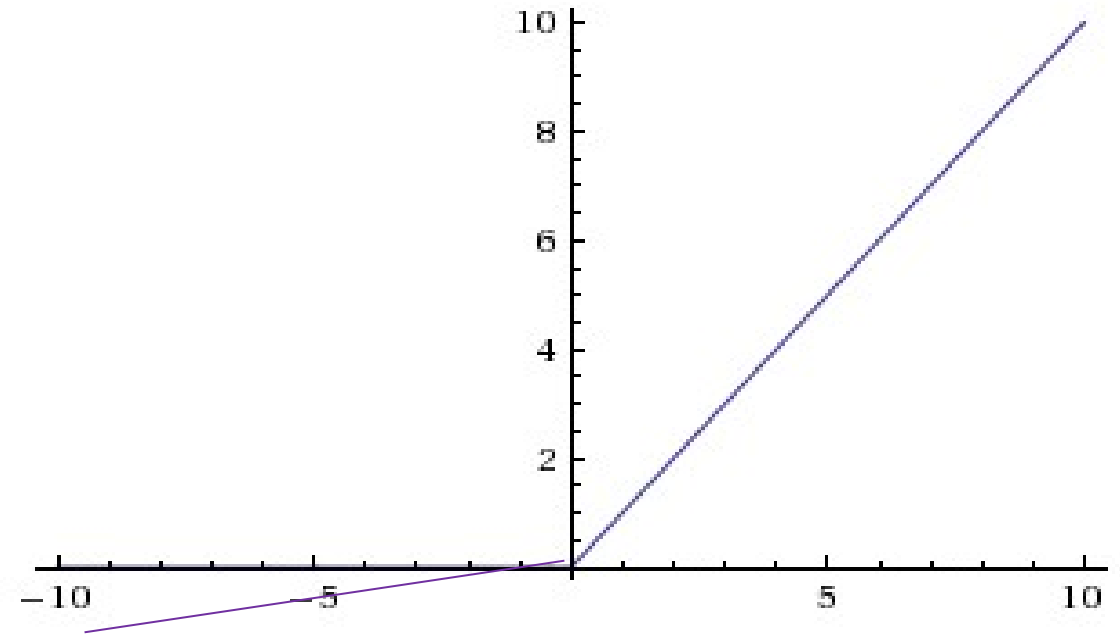


Succeeded in some cases, but the results are not always consistent.



ACTIVATION FUNCTIONS: PARAMETRIC RELU

$$f(x) = \begin{cases} \alpha x, & x < 0 \\ x, & x \geq 0 \end{cases}$$



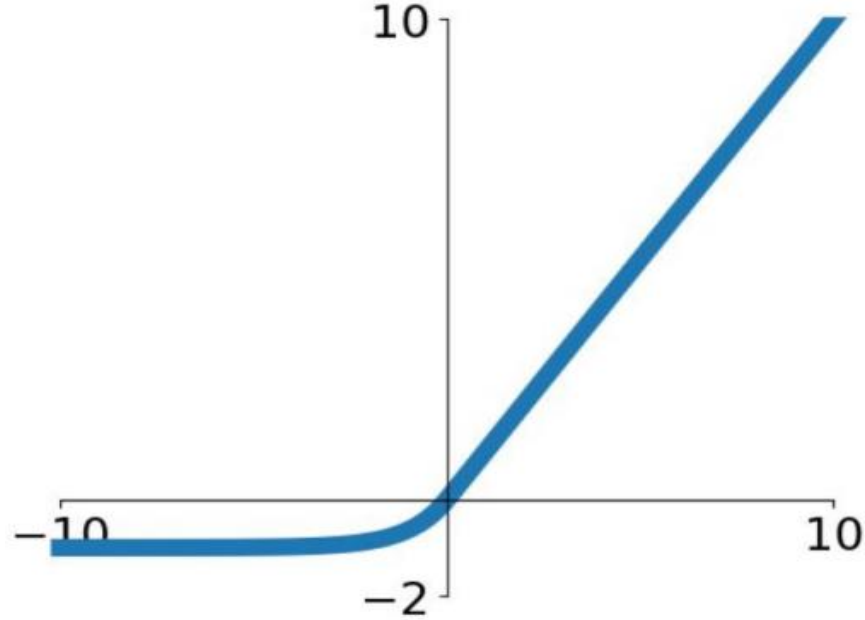
In PReLU, the slope in the negative region is considered as a parameter of each neuron and learnt from data.

He, K., Zhang, X., Ren, S., & Sun, J. (2015). Delving deep into rectifiers: Surpassing human-level performance on imagenet classification. *IEEE international conference on computer vision (CVPR)*.

Source: <http://cs231n.github.io>



ACTIVATION FUNCTIONS: ELU

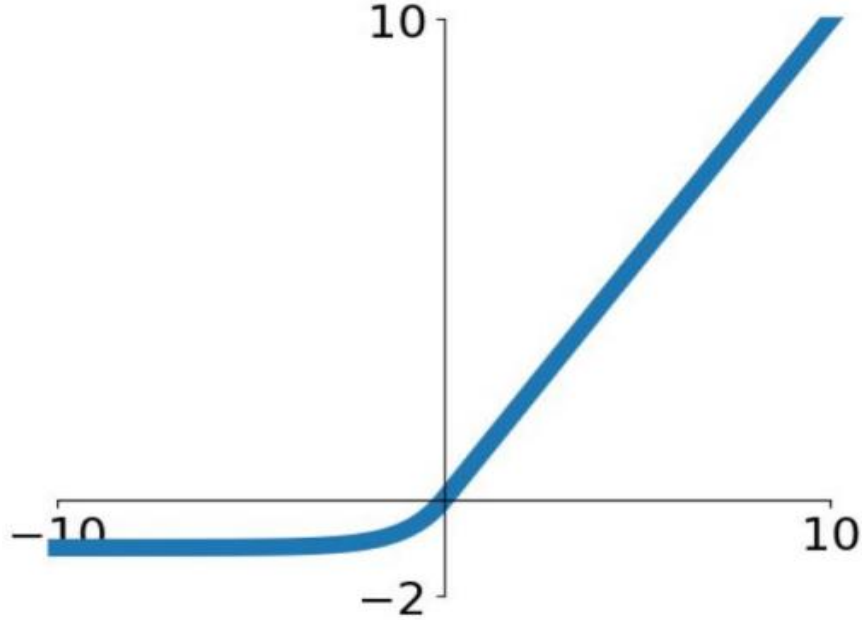


$$f(x) = \begin{cases} x & \text{if } x > 0 \\ \alpha (\exp(x) - 1) & \text{if } x \leq 0 \end{cases}$$

- Exponential Linear Unit



ACTIVATION FUNCTIONS: ELU

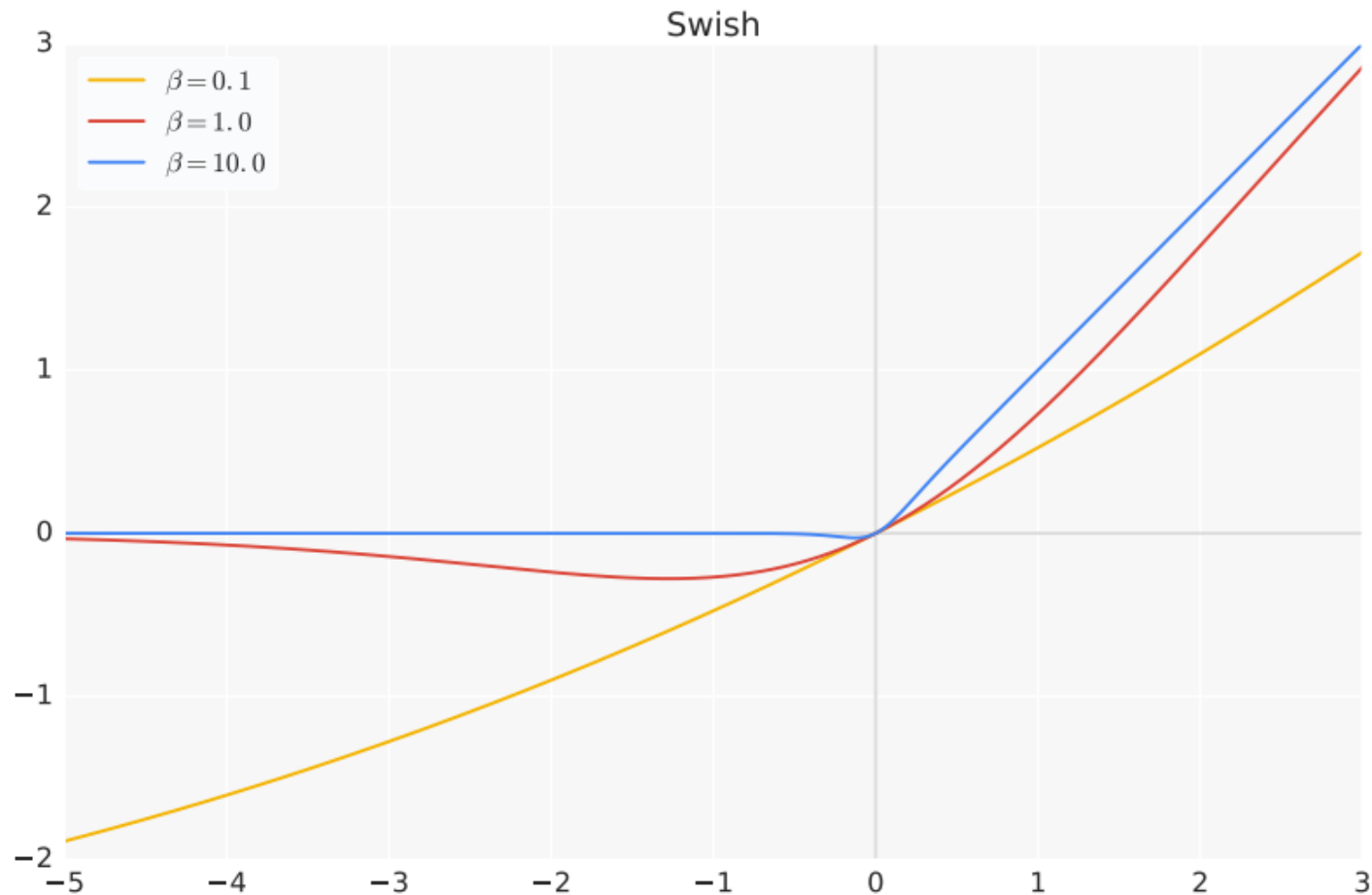


$$f(x) = \begin{cases} x & \text{if } x > 0 \\ \alpha (\exp(x) - 1) & \text{if } x \leq 0 \end{cases}$$

- Exponential Linear Unit
- All benefits of ReLU
- Negative saturation regime compared with Leaky ReLU adds some robustness to noise
- Computation requires `exp()`



ACTIVATION FUNCTIONS: SWISH

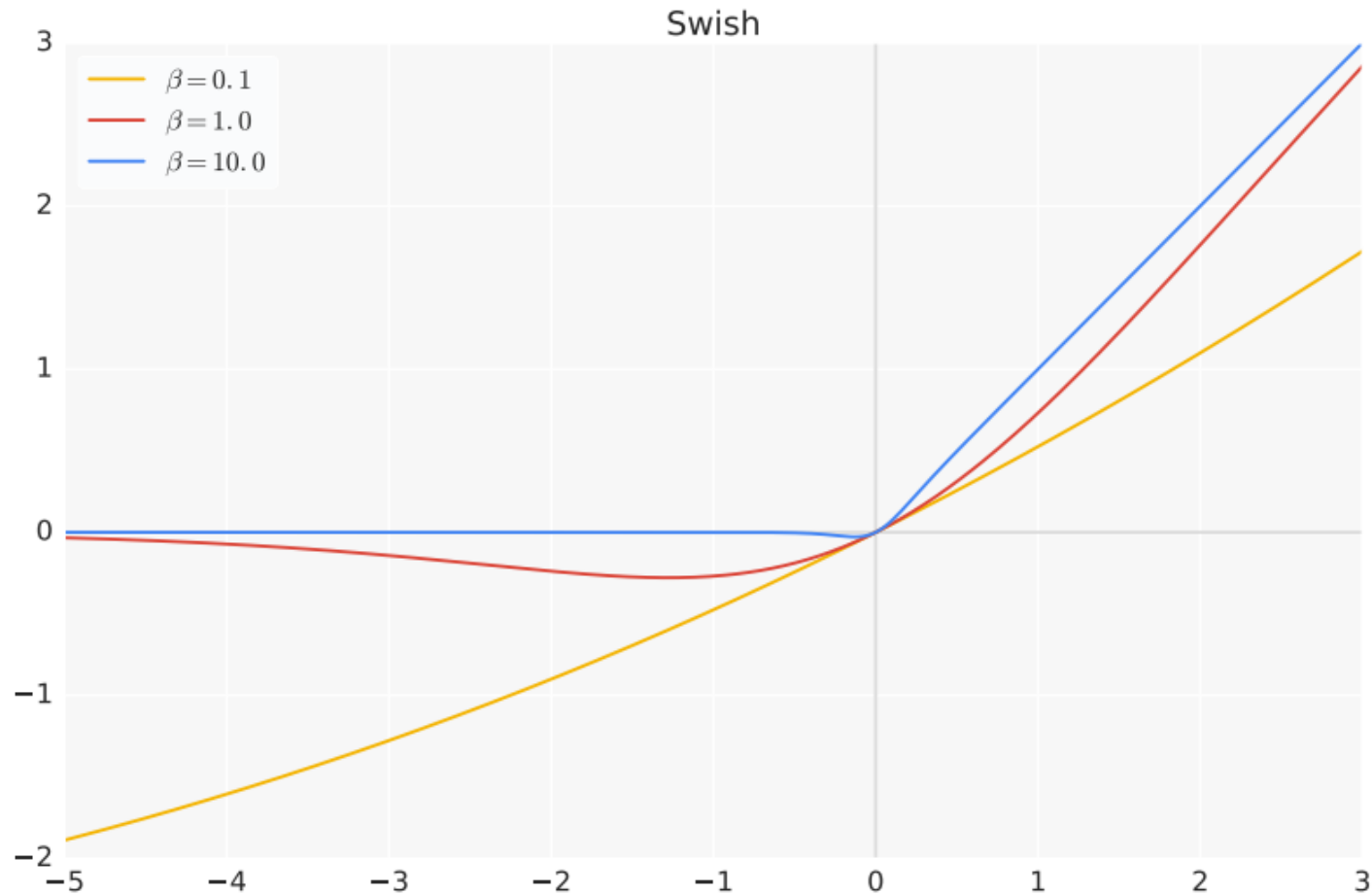


$$f(x) = x \cdot \text{sigmoid}(\beta x)$$

- ReLU is special case of Swish



ACTIVATION FUNCTIONS: SWISH



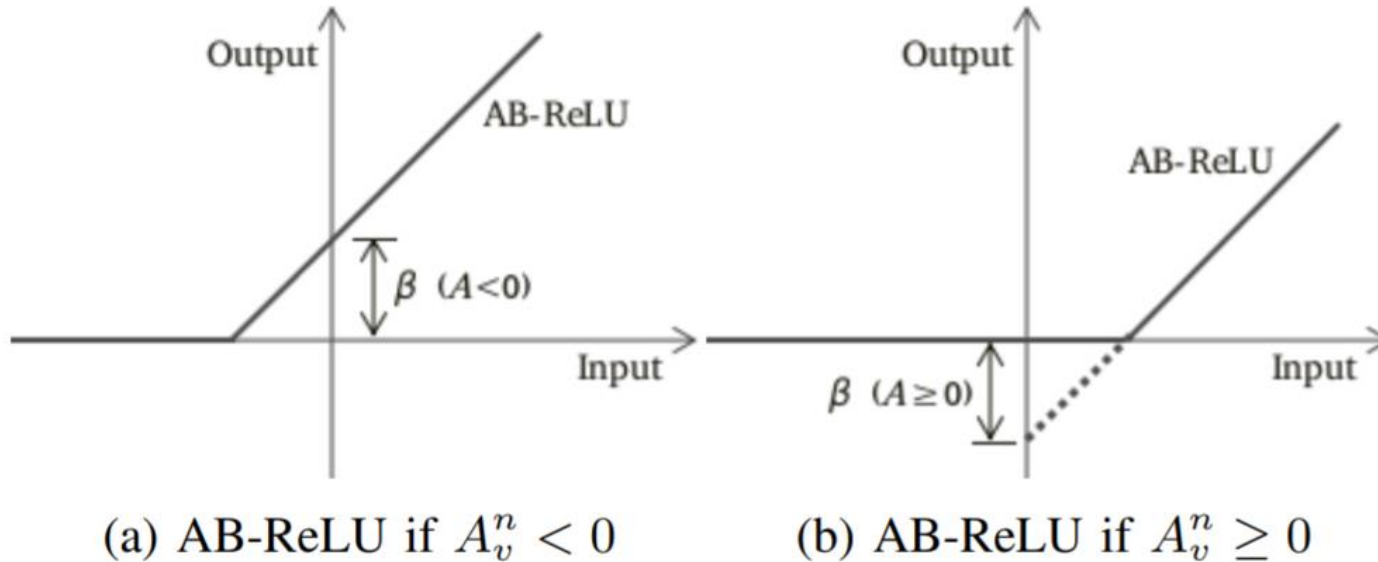
$$f(x) = x \cdot \text{sigmoid}(\beta x)$$

- ReLU is special case of Swish

CIFAR-10 accuracy

Model	ResNet	WRN	DenseNet
LReLU	94.2	95.6	94.7
PReLU	94.1	95.1	94.5
Softplus	94.6	94.9	94.7
ELU	94.1	94.1	94.4
SELU	93.0	93.2	93.9
GELU	94.3	95.5	94.8
ReLU	93.8	95.3	94.8
Swish-1	94.7	95.5	94.8
Swish	94.5	95.5	94.8

ACTIVATION FUNCTIONS: ABRELU



$$I_v^{n+1}(\rho) = \begin{cases} I_v^n(\rho) - \beta, & \text{if } I_v^n(\rho) - \beta > 0 \\ 0, & \text{otherwise} \end{cases}$$

$$\beta = \alpha \times A_v^n$$

average of input volume

Average Biased ReLU (ABReLU)



ACTIVATION FUNCTIONS: IN PRACTICE

- Use ReLU. Be careful with your learning rates
- Try out AReLU/Swish/
- Try out Leaky ReLU but performance might not be stable
- Try out tanh but don't expect much
- Don't use sigmoid



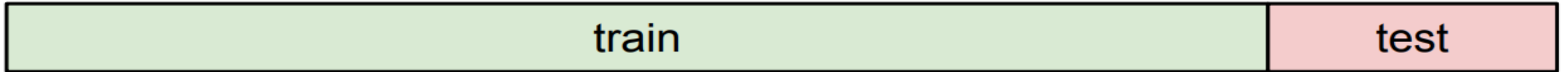
DATASET PREPARATION

TRAIN/VAL/TEST SETS



IN GENERAL PEOPLE DO: TRAIN/TEST

- Split data into train and test,
- Choose hyperparameters that work best on test data



IN GENERAL PEOPLE DO: TRAIN/TEST

- Split data into train and test,
- Choose hyperparameters that work best on test data



BAD: No idea how algorithm will perform on new data



K-FOLD VALIDATION

- Split data into folds,
- Try each fold as validation and average the results

fold 1	fold 2	fold 3	fold 4	fold 5	test
fold 1	fold 2	fold 3	fold 4	fold 5	test
fold 1	fold 2	fold 3	fold 4	fold 5	test



K-FOLD VALIDATION

- Split data into folds,
- Try each fold as validation and average the results

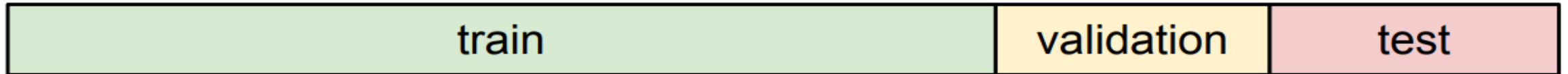
fold 1	fold 2	fold 3	fold 4	fold 5	test
fold 1	fold 2	fold 3	fold 4	fold 5	test
fold 1	fold 2	fold 3	fold 4	fold 5	test

Useful for small datasets, but not used too frequently in deep learning



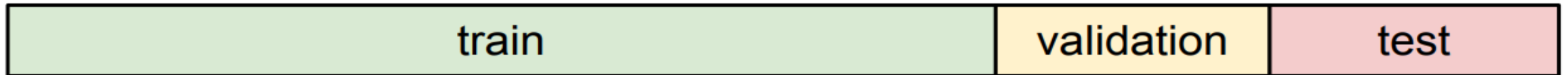
BETTER APPROACH: TRAIN/VAL/TEST SETS

- Split data into **train**, **val**, and **test**;
- Choose hyperparameters on val and evaluate on test



BETTER APPROACH: TRAIN/VAL/TEST SETS

- Split data into **train**, **val**, and **test**;
- Choose hyperparameters on val and evaluate on test



Division can be done based on the size of dataset:

- Roughly 10k or 10% whichever is less for val and test sets.
- Rest in train set.

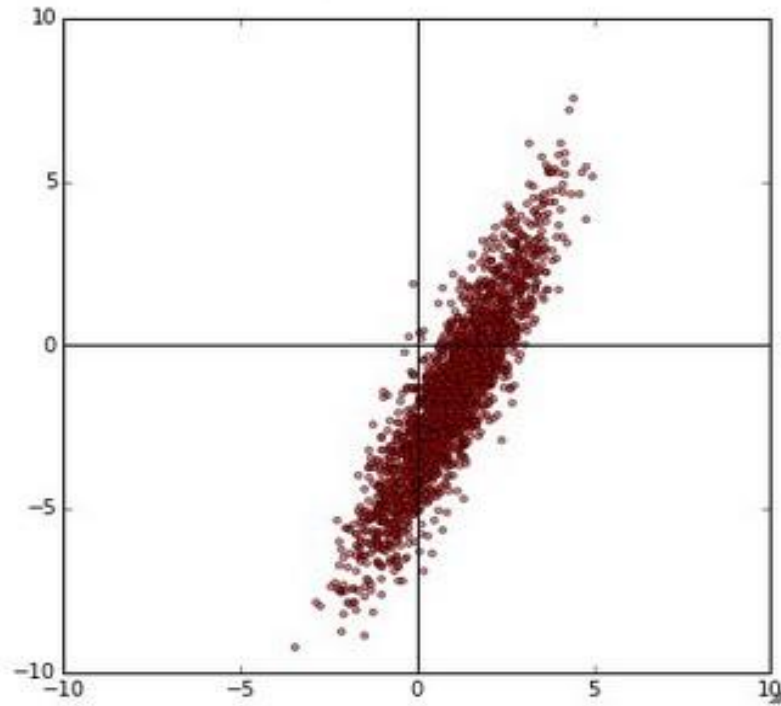


Data Preprocessing

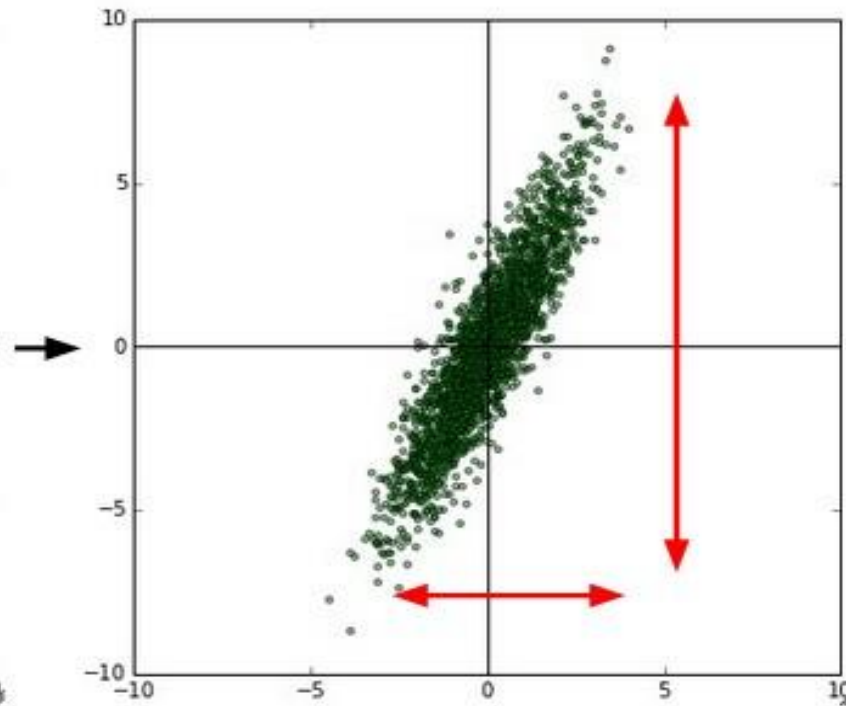


DATA PREPROCESSING

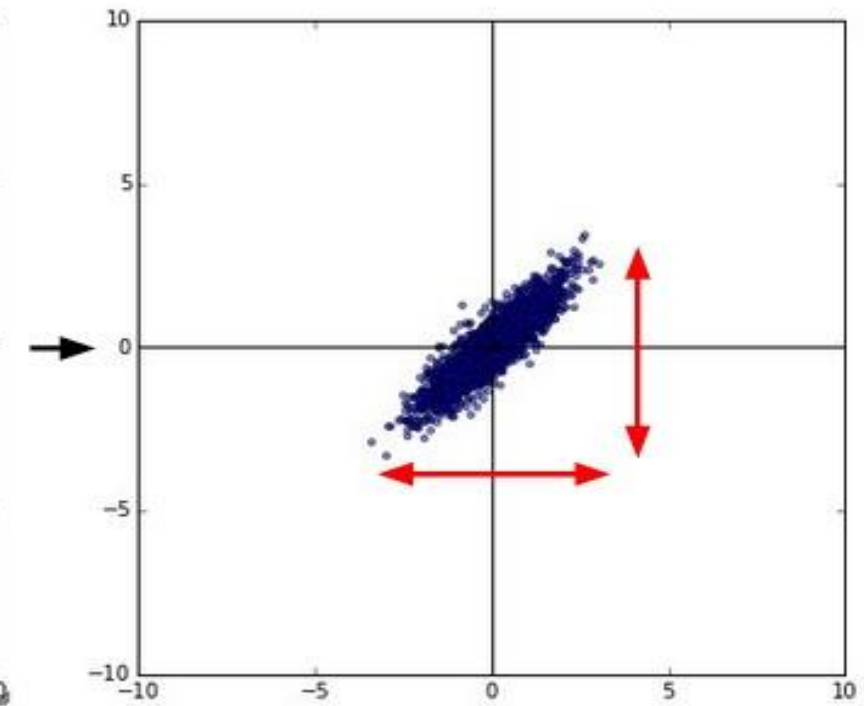
original data



zero-centered data

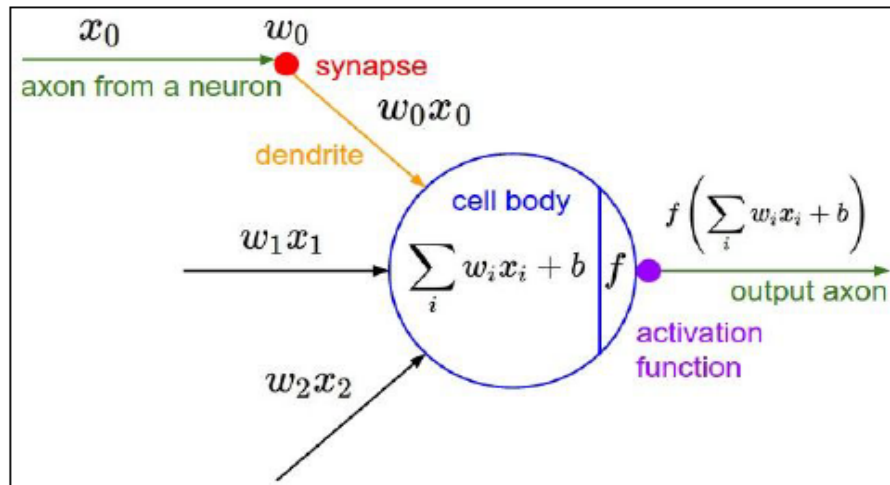


normalized data



DATA PREPROCESSING

Consider what happens when the input to a neuron (x) is always positive:



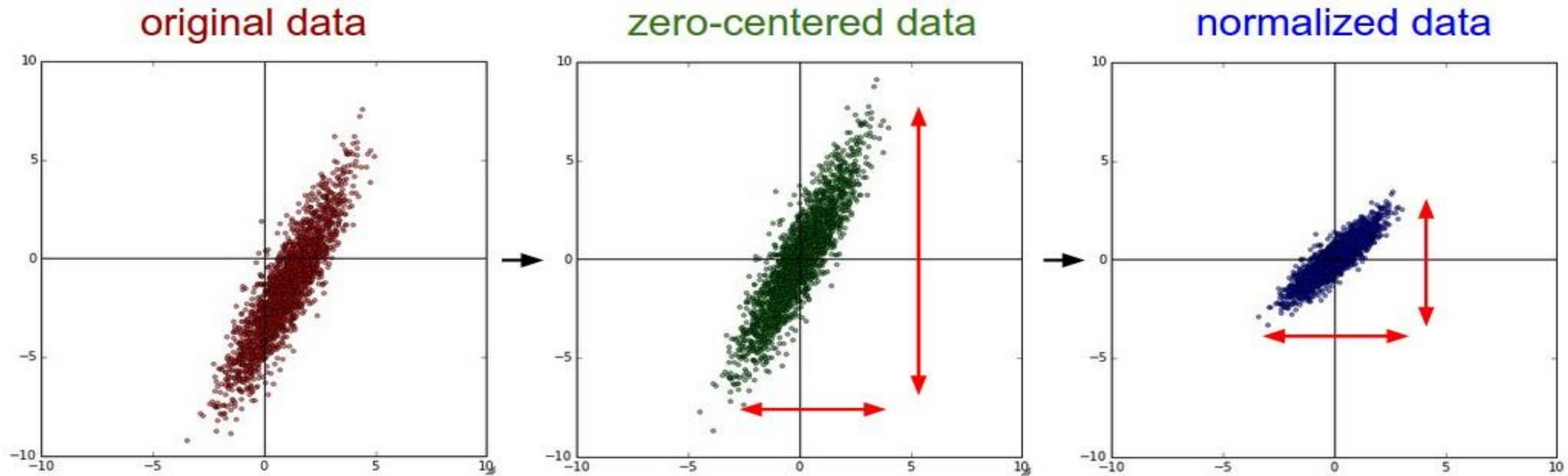
$$f\left(\sum_i w_i x_i + b\right)$$

What can we say about the gradients on $\mathbf{w}$?

Always all positive or all negative
(this is also why you want zero-mean data!)



DATA PREPROCESSING



In practice for Images: only centering is preferred

e.g. consider CIFAR-10 example with [32,32,3] images

- Subtract the mean image (e.g. AlexNet)
(mean image = [32,32,3] array)
- Subtract per-channel mean (e.g. VGGNet, ResNet, etc.)
(mean along each channel = 3 numbers)

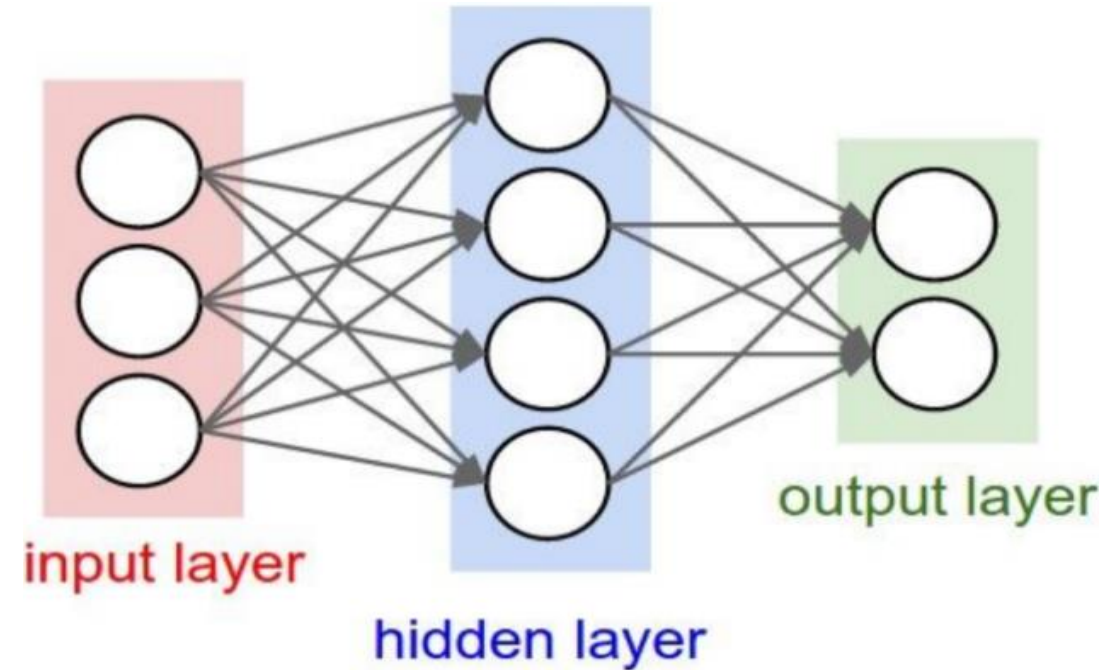


Weight Initialization



WEIGHT INITIALIZATION: CONSTANT

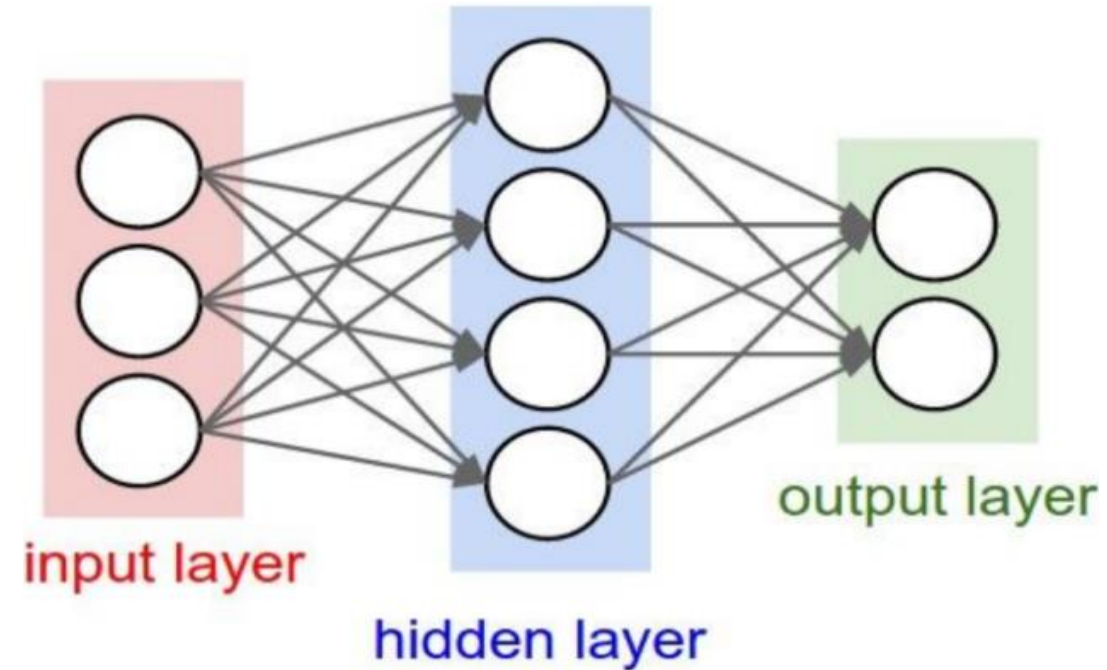
Q: what happens when $W = \text{Constant init}$ is used?



WEIGHT INITIALIZATION: CONSTANT

Q: what happens when $W = \text{Constant init}$ is used?

- Every neuron will compute the same output and undergo the exact same parameter updates.
- There is no source of asymmetry between neurons if their weights are initialized to be the same.



WEIGHT INITIALIZATION: GAUSSIAN

First idea: **Small random numbers**

(Gaussian with zero mean and $1e-2$ standard deviation)

Symmetry breaking: Weights are different for different neurons



WEIGHT INITIALIZATION: GAUSSIAN

First idea: **Small random numbers**

(Gaussian with zero mean and $1e-2$ standard deviation)

Symmetry breaking: Weights are different for different neurons

Works ~okay for small networks, but problems with deeper networks,

i.e. Almost all neurons will become zero

-> gradient diminishing problem.



WEIGHT INITIALIZATION: GAUSSIAN

First idea: **Small random numbers**

(Gaussian with zero mean and $1e-2$ standard deviation)

Symmetry breaking: Weights are different for different neurons

Works ~okay for small networks, but problems with deeper networks,

i.e. Almost all neurons will become zero

-> gradient diminishing problem.

Increase the standard deviation to 1



WEIGHT INITIALIZATION: GAUSSIAN

First idea: **Small random numbers**

(Gaussian with zero mean and $1e-2$ standard deviation)

Symmetry breaking: Weights are different for different neurons

Works ~okay for small networks, but problems with deeper networks,

i.e. Almost all neurons will become zero

-> gradient diminishing problem.

Increase the standard deviation to 1

Almost all neurons completely saturated, either -1 or 1. Gradients will be all zero.

-> gradient diminishing problem.



WEIGHT INITIALIZATION: GAUSSIAN

Lets look at
some
activation
statistics

E.g. 10-layer net with
500 neurons on each
layer, using tanh
non-linearities, and
initializing as
described in last slide.

```
# assume some unit gaussian 10-D input data
D = np.random.randn(1000, 500)
hidden_layer_sizes = [500]*10
nonlinearities = ['tanh']*len(hidden_layer_sizes)
```

```
act = {'relu':lambda x:np.maximum(0,x), 'tanh':lambda x:np.tanh(x)}
Hs = {}
for i in xrange(len(hidden_layer_sizes)):
    X = D if i == 0 else Hs[i-1] # input at this layer
    fan_in = X.shape[1]
    fan_out = hidden_layer_sizes[i]
    W = np.random.randn(fan_in, fan_out) * 0.01 # layer initialization

    H = np.dot(X, W) # matrix multiply
    H = act[nonlinearities[i]](H) # nonlinearity
    Hs[i] = H # cache result on this layer
```

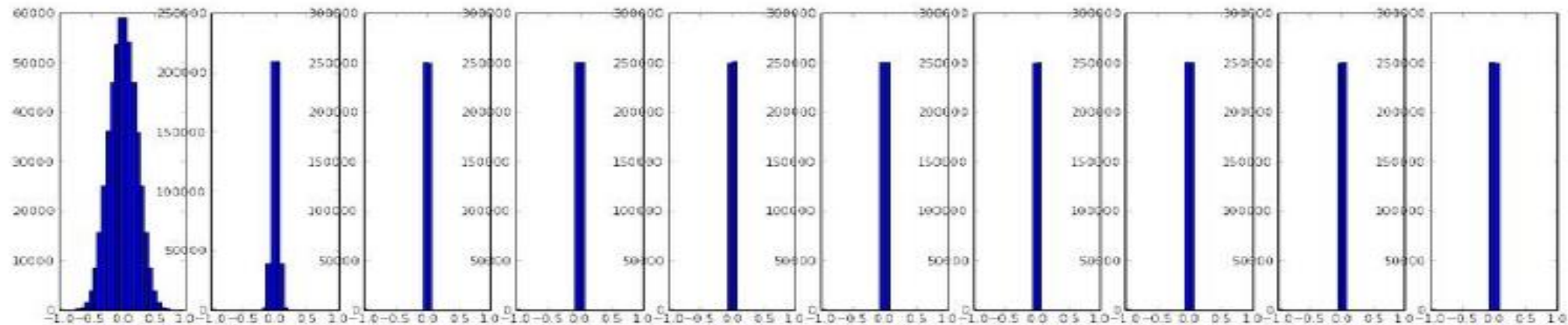
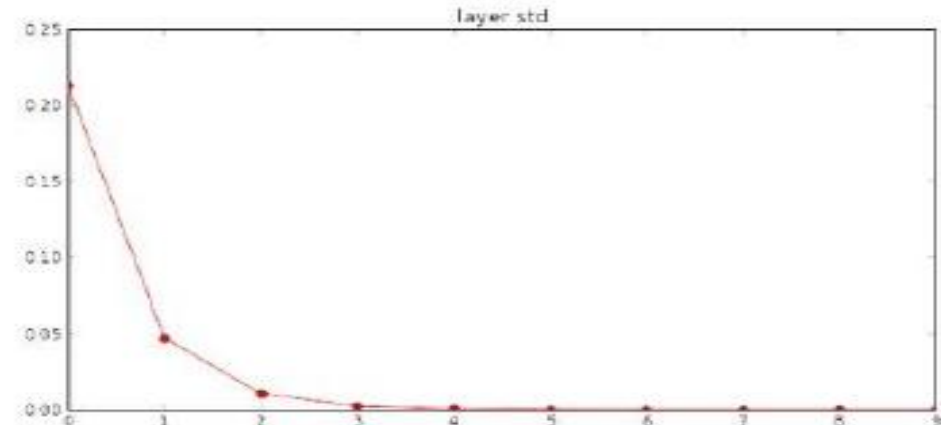
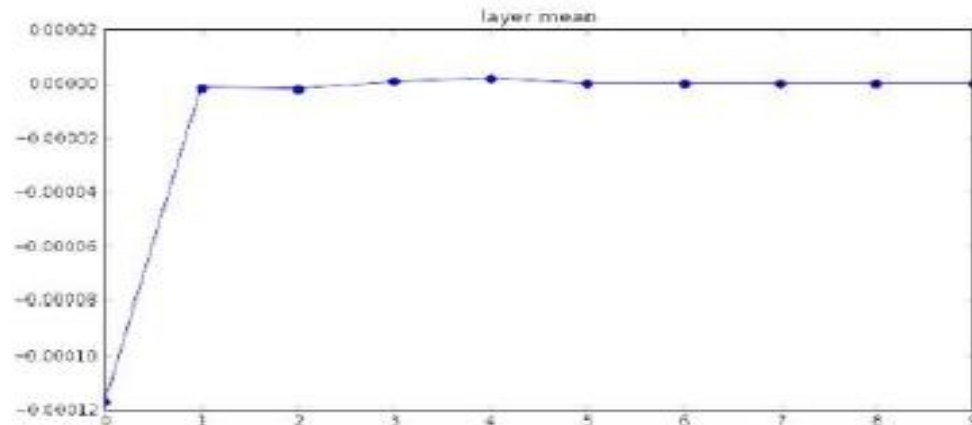
```
# look at distributions at each layer
print 'input layer had mean %f and std %f' % (np.mean(D), np.std(D))
layer_means = [np.mean(H) for i,H in Hs.iteritems()]
layer_stds = [np.std(H) for i,H in Hs.iteritems()]
for i,H in Hs.iteritems():
    print 'hidden layer %d had mean %f and std %f' % (i+1, layer_means[i], layer_stds[i])

# plot the means and standard deviations
plt.figure()
plt.subplot(121)
plt.plot(Hs.keys(), layer_means, 'ob-')
plt.title('layer mean')
plt.subplot(122)
plt.plot(Hs.keys(), layer_stds, 'or-')
plt.title('layer std')

# plot the raw distributions
plt.figure()
for i,H in Hs.iteritems():
    plt.subplot(1,len(Hs),i+1)
    plt.hist(H.ravel(), 30, range=(-1,1))
```


WEIGHT INITIALIZATION: GAUSSIAN

```
input layer had mean 0.000927 and std 0.998388
hidden layer 1 had mean -0.000117 and std 0.213081
hidden layer 2 had mean -0.000001 and std 0.047551
hidden layer 3 had mean -0.000002 and std 0.010630
hidden layer 4 had mean 0.000001 and std 0.002378
hidden layer 5 had mean 0.000002 and std 0.000532
hidden layer 6 had mean -0.000000 and std 0.000119
hidden layer 7 had mean 0.000000 and std 0.000026
hidden layer 8 had mean -0.000000 and std 0.000006
hidden layer 9 had mean 0.000000 and std 0.000001
hidden layer 10 had mean -0.000000 and std 0.000000
```

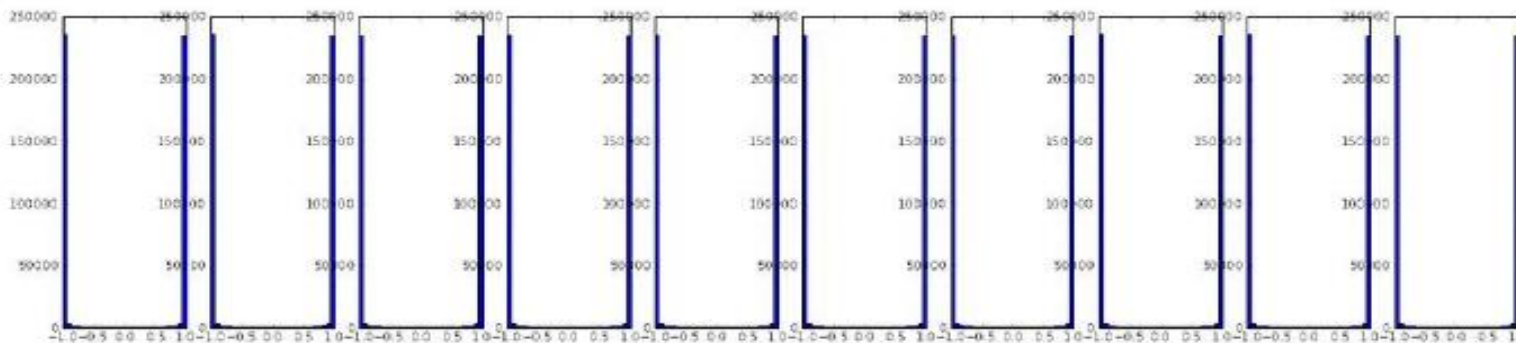
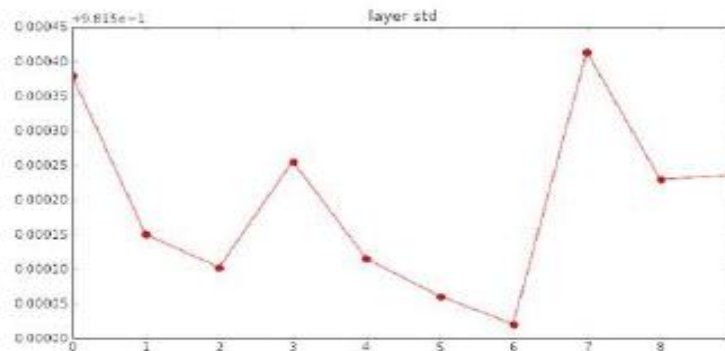
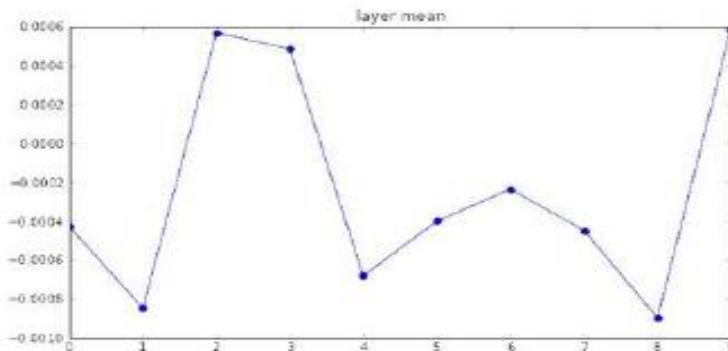


WEIGHT INITIALIZATION: GAUSSIAN

```
W = np.random.randn(fan_in, fan_out) * 1.0 # layer initialization
```

input layer had mean 0.001800 and std 1.001311
hidden layer 1 had mean -0.000430 and std 0.981879
hidden layer 2 had mean -0.000849 and std 0.981649
hidden layer 3 had mean 0.000566 and std 0.981601
hidden layer 4 had mean 0.000483 and std 0.981755
hidden layer 5 had mean -0.000682 and std 0.981614
hidden layer 6 had mean -0.000401 and std 0.981560
hidden layer 7 had mean -0.000237 and std 0.981520
hidden layer 8 had mean -0.000448 and std 0.981913
hidden layer 9 had mean -0.000899 and std 0.981728
hidden layer 10 had mean 0.000584 and std 0.981736

*1.0 instead of *0.01



Almost all neurons completely saturated, either -1 and 1. Gradients will be all zero.



WEIGHT INITIALIZATION: XAVIER

Calibrating the variances with $1/\text{sqrt}(\text{fan_in})$

```
W = np.random.randn(fan_in, fan_out)/np.sqrt(fan_in)
```

Reasonable initialization.

(Mathematical derivation assumes linear activations)

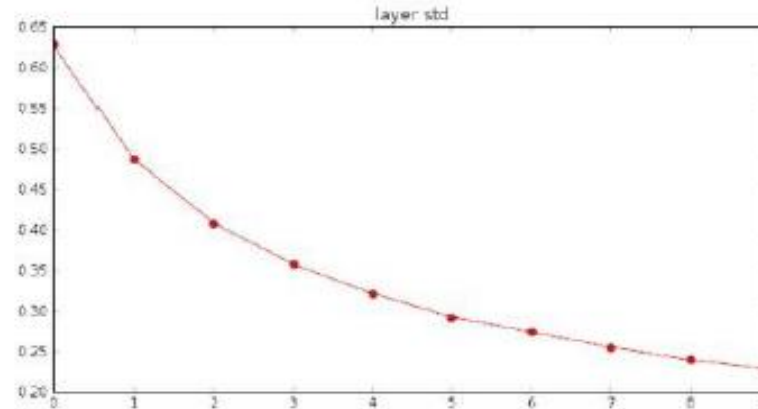
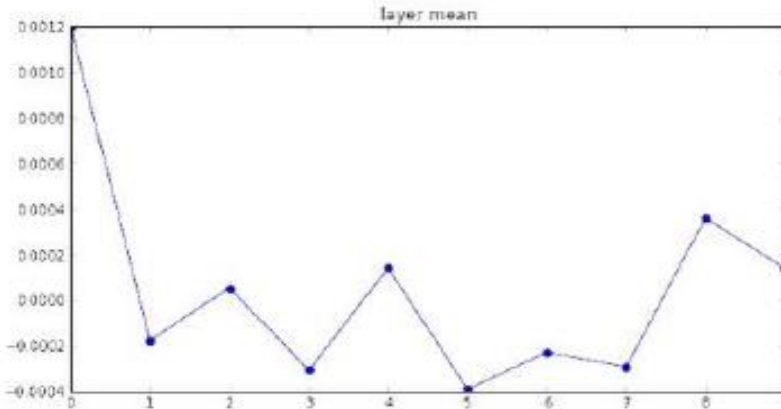


WEIGHT INITIALIZATION: XAVIER

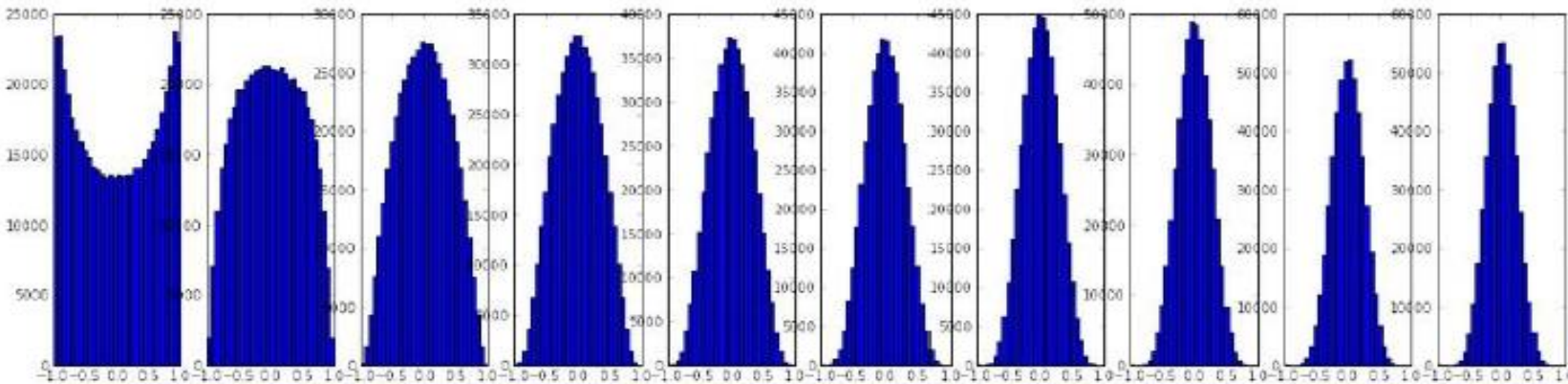
input layer had mean 0.001800 and std 1.001311
hidden layer 1 had mean 0.001198 and std 0.627953
hidden layer 2 had mean -0.000175 and std 0.486051
hidden layer 3 had mean 0.000055 and std 0.407723
hidden layer 4 had mean -0.000306 and std 0.357108
hidden layer 5 had mean 0.000142 and std 0.320917
hidden layer 6 had mean -0.000389 and std 0.292116
hidden layer 7 had mean -0.000228 and std 0.273387
hidden layer 8 had mean -0.000291 and std 0.254935
hidden layer 9 had mean 0.000361 and std 0.239266
hidden layer 10 had mean 0.000139 and std 0.228008

```
W = np.random.randn(fan_in, fan_out) / np.sqrt(fan_in) # layer initialization
```

“Xavier initialization”
[Glorot et al., 2010]



Reasonable initialization.
(Mathematical derivation
assumes linear activations)

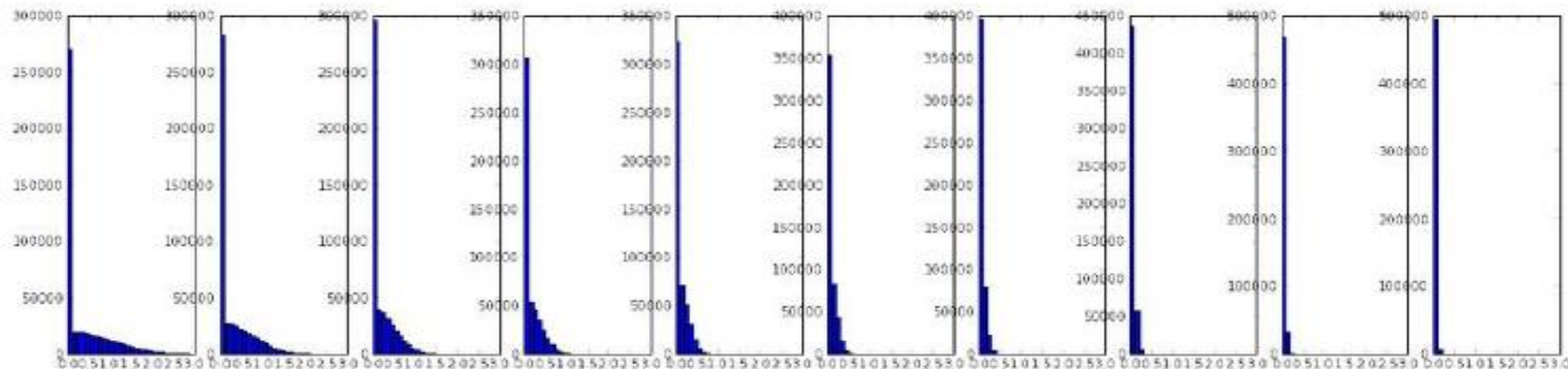
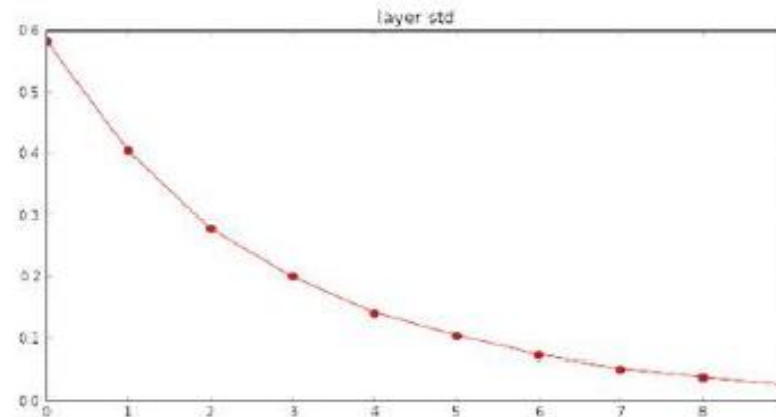
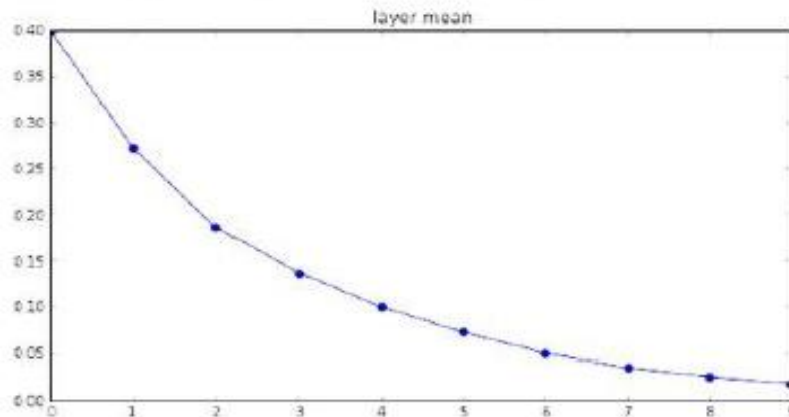


WEIGHT INITIALIZATION: XAVIER

input layer had mean 0.000501 and std 0.999444
hidden layer 1 had mean 0.398623 and std 0.582273
hidden layer 2 had mean 0.272352 and std 0.403795
hidden layer 3 had mean 0.186076 and std 0.276912
hidden layer 4 had mean 0.136442 and std 0.198685
hidden layer 5 had mean 0.099568 and std 0.140299
hidden layer 6 had mean 0.072234 and std 0.103280
hidden layer 7 had mean 0.049775 and std 0.072748
hidden layer 8 had mean 0.035138 and std 0.051572
hidden layer 9 had mean 0.025404 and std 0.038583
hidden layer 10 had mean 0.018408 and std 0.026076

```
W = np.random.randn(fan_in, fan_out) / np.sqrt(fan_in) # layer initialization
```

but when using the ReLU nonlinearity it breaks.

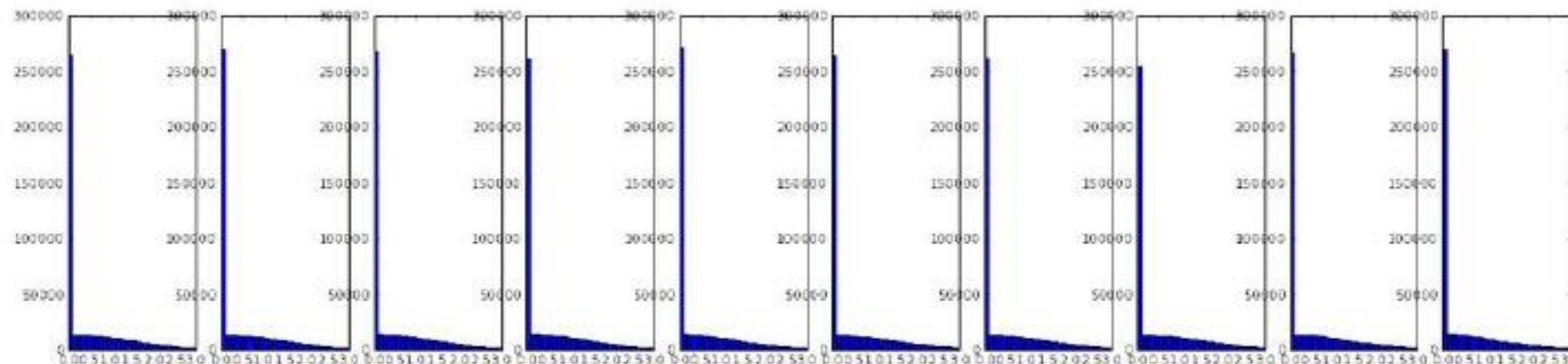
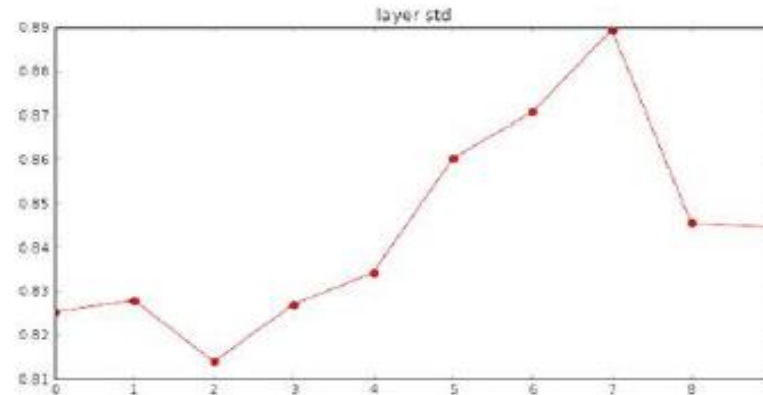
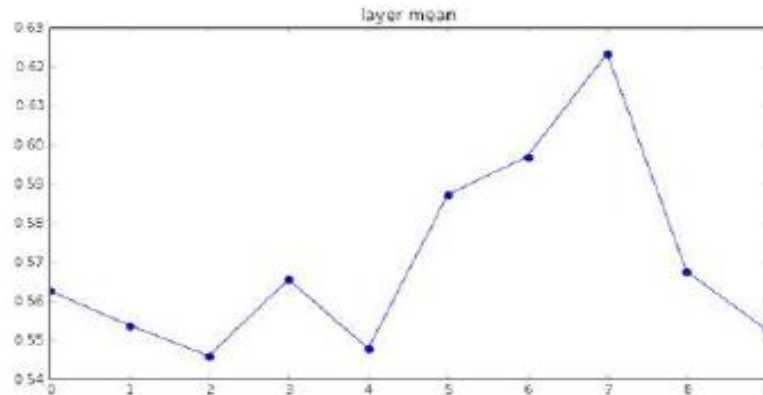


WEIGHT INITIALIZATION: HE

input layer had mean 0.000501 and std 0.999444
hidden layer 1 had mean 0.562488 and std 0.825232
hidden layer 2 had mean 0.553614 and std 0.827835
hidden layer 3 had mean 0.545867 and std 0.813855
hidden layer 4 had mean 0.565396 and std 0.826902
hidden layer 5 had mean 0.547678 and std 0.834092
hidden layer 6 had mean 0.587103 and std 0.860035
hidden layer 7 had mean 0.596867 and std 0.870610
hidden layer 8 had mean 0.623214 and std 0.889348
hidden layer 9 had mean 0.567498 and std 0.845357
hidden layer 10 had mean 0.552531 and std 0.844523

```
W = np.random.randn(fan_in, fan_out) / np.sqrt(fan_in/2) # layer initialization
```

He et al., 2015
(note additional /2)



ACKNOWLEDGEMENT

Thanks to the following courses and corresponding researchers for making their teaching/research material online

- Convolutional Neural Networks for Visual Recognition, Stanford University
- Deep Learning, Stanford University
- Introduction to Deep Learning, University of Illinois at Urbana-Champaign
- Introduction to Deep Learning, Carnegie Mellon University
- Natural Language Processing with Deep Learning, Stanford University
- And Many More Publicly Available Resources



Questions?

