



Lecture 6: Neural Networks + Backpropagation

Announcements:

- HW #2 is due tonight, uploaded to Gradescope. Please budget time for submission. Please be sure to print out Jupyter Notebooks and .py files for your submission.
- We will upload HW #3 tonight. Since we're going a bit slower than prior iterations of the course, we decided to make it due on Friday, Feb 9, 2024 (instead of Monday, Feb 5, 2024).

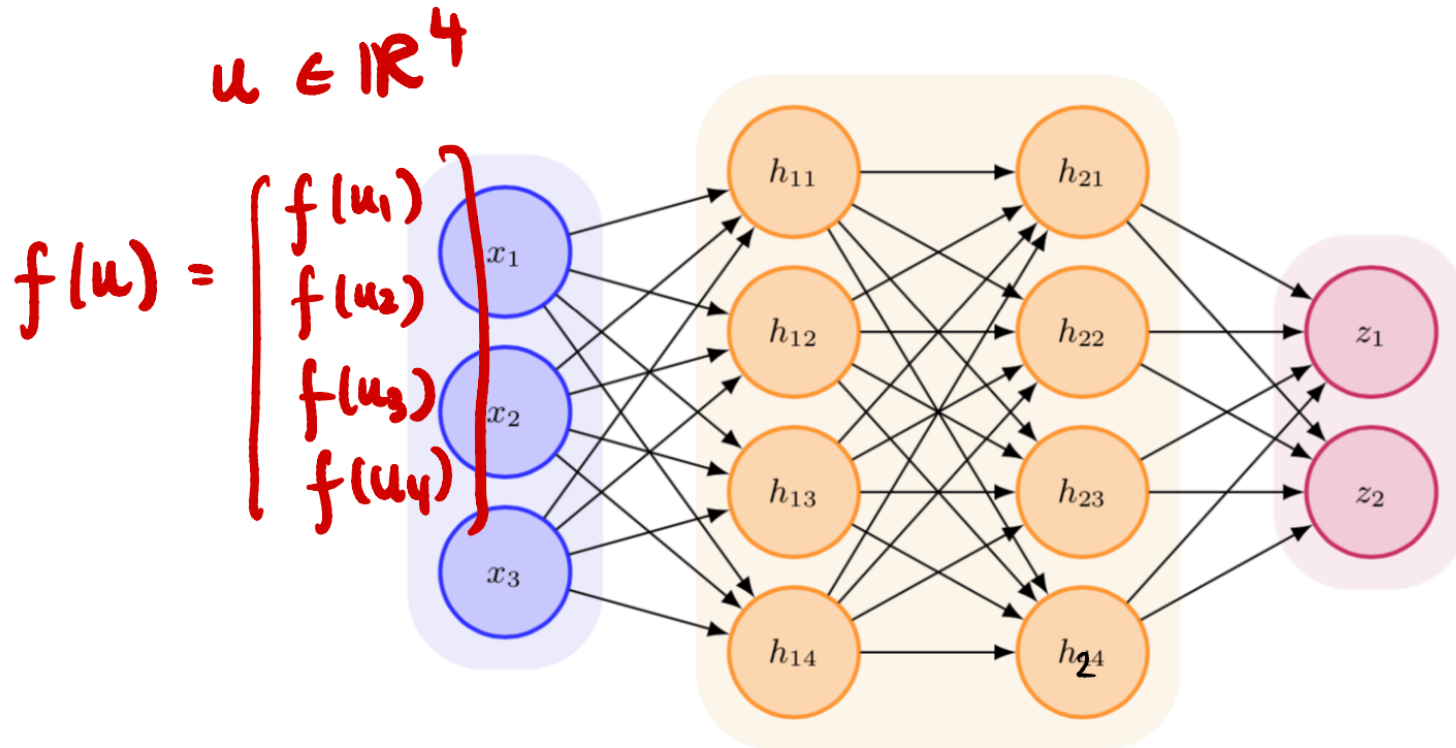
• ★ ★ ★ UPDATE ★ ★ ★

Midterm will be held in class on Feb. 21, 2024.

(seating will be tight.)



Neural networks



First layer: $\mathbf{h}_1 = f(\mathbf{W}_1 \mathbf{x} + \mathbf{b}_1)$

Second layer: $\mathbf{h}_2 = f(\mathbf{W}_2 \mathbf{h}_1 + \mathbf{b}_2)$

Third (output layer): $\mathbf{z} = \mathbf{W}_3 \mathbf{h}_2 + \mathbf{b}_3$

Fully Connected (FC)
Neural Networks

Multilayer Perceptron (MLP)



What nonlinearity (aka activation function or unit) to use?

“One recurring theme throughout neural network design is that the gradient of the cost function must be large and predictable enough to serve as a good guide for the learning algorithm.” (Goodfellow et al., p. 173)

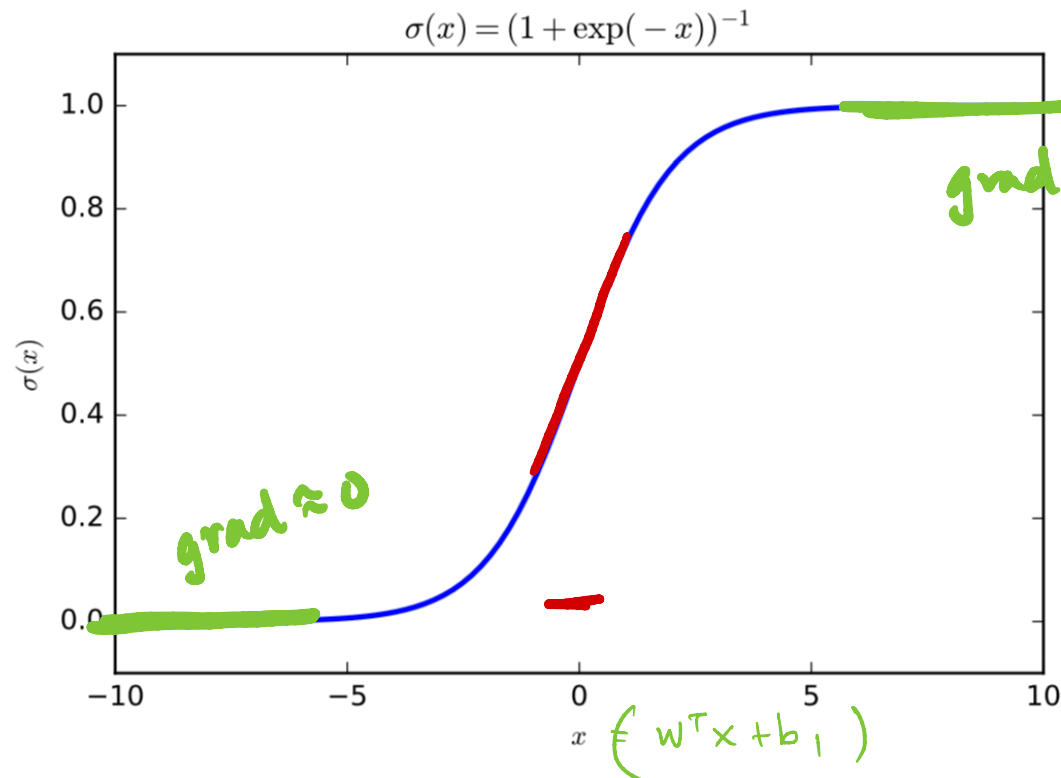


Sigmoid unit

Sigmoid activation, $\sigma(x)$ $= \frac{1}{1+\exp(-x)}$

$$f(w_1 x + b_1) = \sigma(w_1 x + b_1)$$

```
f = lambda x: 1.0 / (1.0 + np.exp(-x))
```



Its derivative is:

$$\frac{d\sigma(x)}{dx} = \sigma(x)(1 - \sigma(x))$$



Sigmoid unit

“One recurring theme throughout neural network design is that the gradient of the cost function must be large and predictable enough to serve as a good guide for the learning algorithm.” (Goodfellow et al., p. 173)

$$\mathcal{L}(w)$$

$$\sigma(w) = \sigma(w^T x + b)$$

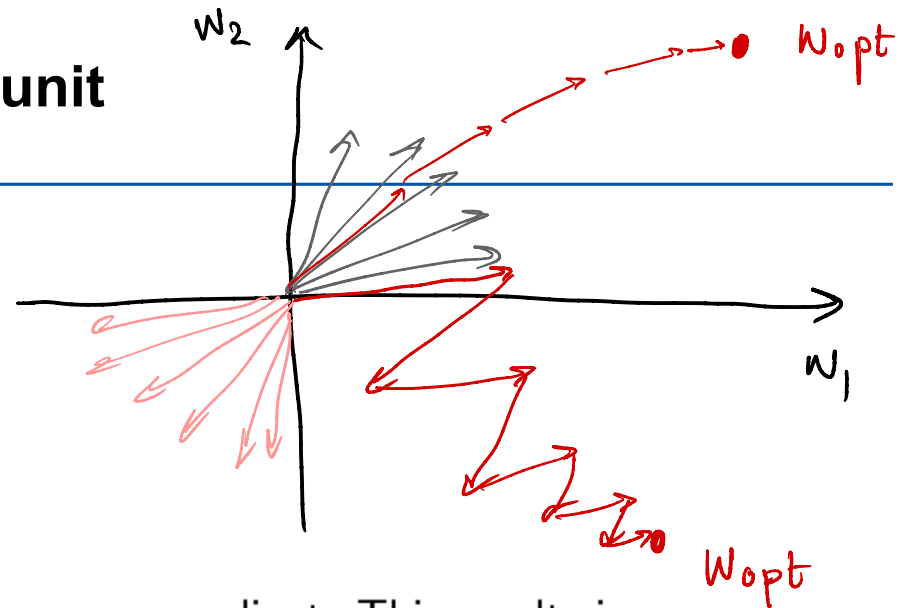
$$w \leftarrow w - \epsilon \cdot \underbrace{\frac{\partial \mathcal{L}}{\partial w}}$$

$$\frac{\partial \mathcal{L}}{\partial w} = \underbrace{\frac{\partial \sigma(w)}{\partial w}} \frac{\partial \mathcal{L}}{\partial \sigma(w)}$$

a very small #



Sigmoid unit



Sigmoid activation, $\sigma(x) = \frac{1}{1+\exp(-x)}$

Pros:

- Around $x = 0$, the unit behaves linearly.
- It is differentiable everywhere.

Cons:

- At extremes, the unit saturates and thus has zero gradient. This results in no learning with gradient descent.
- The sigmoid unit is non-negative. SGD with the sigmoid unit zig-zags.

$$x \xrightarrow{(\tilde{w}, \tilde{b}), \tau} h_1 \xrightarrow{(\tilde{w}, \tilde{b}), \tau} h_2 \rightarrow$$

$$h_1 = \sigma(\tilde{W}x + \tilde{b})$$

Consider $f(\mathbf{w}) = \sigma(\mathbf{w}^T \mathbf{h}_1 + b)$. Defining $z = \mathbf{w}^T \mathbf{h}_1 + b$, the derivative with respect to $\mathbf{w}$, the parameters, is: $\frac{\partial f(\mathbf{w})}{\partial \mathbf{w}} = \sigma(z)(1 - \sigma(z)) \mathbf{h}_1$

$$\frac{\partial f(\mathbf{w})}{\partial \mathbf{w}} = \sigma(z)(1 - \sigma(z)) \mathbf{h}_1 = \frac{\partial z}{\partial \mathbf{w}} \frac{\partial f(\mathbf{w})}{\partial z}$$

If $x \geq 0$ (e.g., if the input units all had a sigmoidal output), then the gradient has all positive entries. Let's say we had some gradient, $\frac{\partial \mathcal{L}}{\partial f}$, which can be positive or negative. Then $\frac{\partial \mathcal{L}}{\partial \mathbf{w}}$ will have all positive or negative entries.

$$f = \sigma(z)$$

$$\frac{\partial f}{\partial z} = \sigma(z)(1 - \sigma(z))$$

$$z = \mathbf{w}^T \mathbf{h}_1 + b$$

$$\frac{\partial z}{\partial \mathbf{w}} = \mathbf{h}_1$$

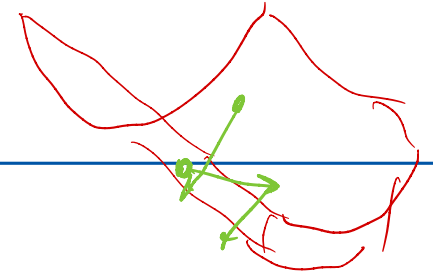
This can result in zig-zagging during gradient descent.

$\frac{\partial \mathcal{L}}{\partial f(\mathbf{w})}$ is a scalar
 ≥ 0
 ≤ 0

$\frac{\partial \mathcal{L}}{\partial \mathbf{w}} = \frac{\partial f(\mathbf{w})}{\partial \mathbf{w}} \frac{\partial \mathcal{L}}{\partial f(\mathbf{w})}$
 $\downarrow$
 all entries in this gradient are either ≤ 0 , or ≥ 0



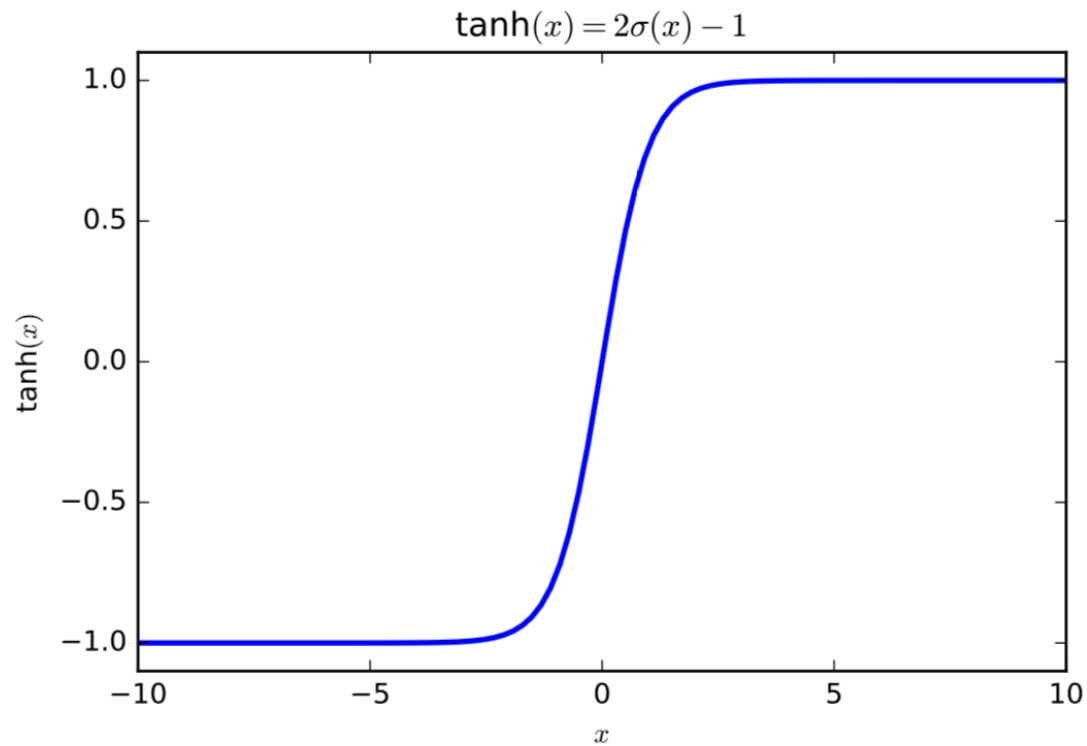
Hyperbolic tangent



Hyperbolic tangent, $\tanh(x) = 2\sigma(x) - 1$

The hyperbolic tangent is a zero-centered sigmoid-looking activation.

```
f = lambda x: np.tanh(x)
```



Its derivative is:

$$\frac{d \tanh(x)}{dx} = 1 - \tanh^2(x)$$



Hyperbolic tangent

Hyperbolic tangent, $\tanh(x) = 2\sigma(x) - 1$

Pros:

- Around $x = 0$, the unit behaves linearly.
- It is differentiable everywhere.
- It is zero-centered. *It's both pos. & neg. so it isn't constrained to zig-zag.*

Cons:

- Like the sigmoid unit, when a unit saturates, i.e., its values grow larger or smaller, the unit saturates and no additional learning occurs.

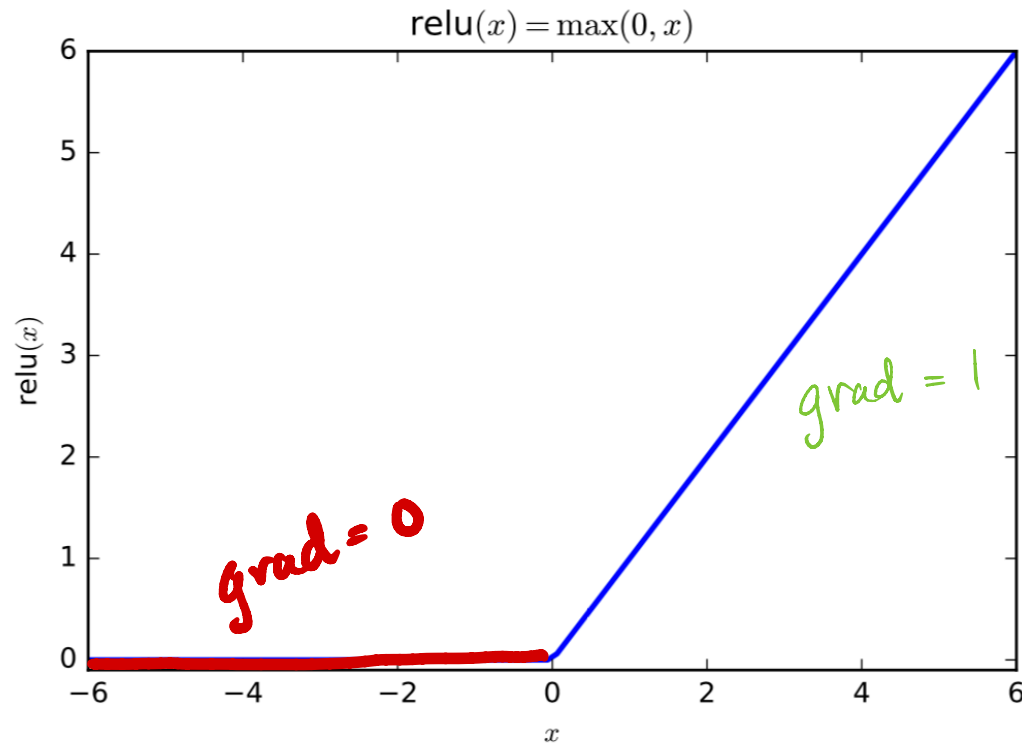


ReLU unit

$$f(x) = x$$

Rectified linear unit, $\text{ReLU}(x) = \max(0, x)$

```
f = lambda x: x * (x > 0)
```



Its derivative is:

$$\frac{d\text{ReLU}(x)}{dx} = \begin{cases} 1 & x > 0 \\ 0 & x < 0 \end{cases}$$

This function is not differentiable at $x = 0$. However, we can define its subgradient by setting the derivative to be between $[0, 1]$ at $x = 0$.



ReLU unit

$$z = Wx + b$$

Rectified linear unit, $\text{ReLU}(x) = \max(0, x)$

$$\text{relu}(z) = \begin{bmatrix} 1 \\ 0 \\ 4 \\ 0 \end{bmatrix} \quad z = \begin{bmatrix} 1 \\ -0.03 \\ 4 \\ -0.5 \end{bmatrix}$$

Pros:

- In practice, learning with the ReLU unit converges faster than sigmoid and tanh. *AlexNet, relu was 6x faster than tanh.*
- When a unit is active, it behaves as a linear unit.
- The derivative at all points, except $x = 0$, is 0 or 1. When $x > 0$, the gradients are large, and not scaled by second order effects.
- There is no saturation if $x > 0$.

Cons:

- $\text{ReLU}(x)$, like sigmoid, is non-negative. SGD with $\text{ReLU}()$ therefore zig-zags.
- $\text{ReLU}(x)$ is not differentiable at $x = 0$. However, in practice, this is not a large issue. A heuristic when evaluating $\left. \frac{d\text{ReLU}(x)}{dx} \right|_{x=0}$ is to return the left derivative (0) or the right derivative (1); this is reasonable given digital computation is subject to numerical error.
- Learning does not happen for examples that have zero activation. This can be fixed by e.g., using a leaky ReLU or maxout unit.

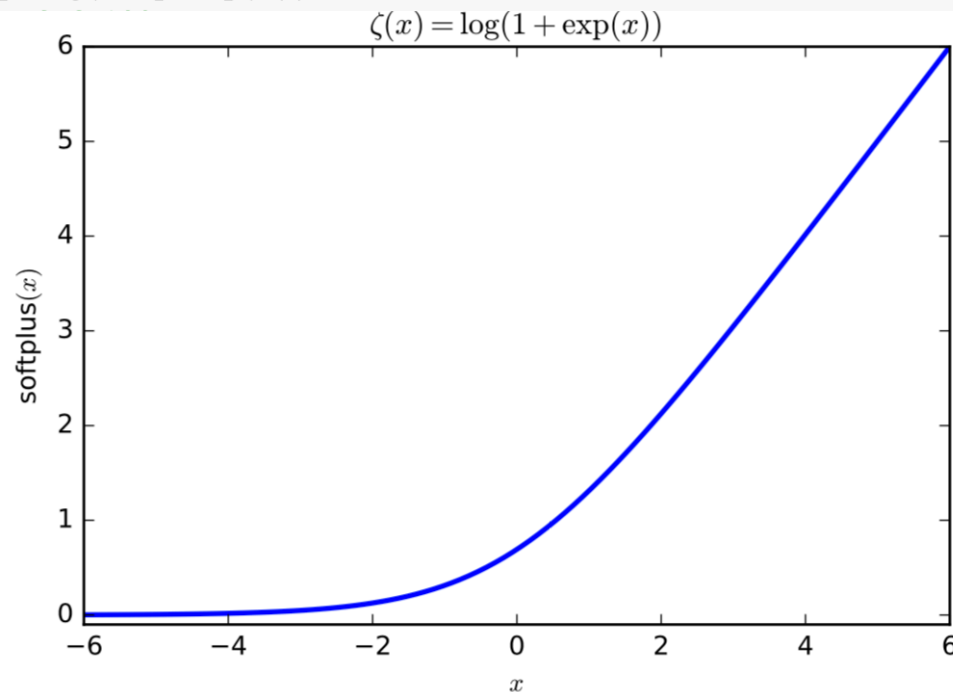


Softplus unit

Softplus unit, $\zeta(x) = \log(1 + \exp(x))$

One may consider using the softplus function, $\zeta(x) = \log(1 + e^x)$, in place of $\text{ReLU}(x)$. Intuitively, this ought to work well as it resembles $\text{ReLU}(x)$ and is differentiable everywhere. However, empirically, it performs worse than $\text{ReLU}(x)$.

```
f = lambda x: np.log(1+np.exp(x))
```



Its derivative is:

$$\frac{d\zeta(x)}{dx} = \sigma(x)$$



Softplus unit

“One might expect it to have an advantage over the [ReLU] due to being differentiable everywhere or due to saturating less completely, but empirically it does not.” (Goodfellow et al., p. 191)

Neuron	MNIST	CIFAR10	NISTP	NORB
<i>With unsupervised pre-training</i>				
Rectifier	1.20%	49.96%	32.86%	16.46%
Tanh	1.16%	50.79%	35.89%	17.66%
Softplus	1.17%	49.52%	33.27%	19.19%
<i>Without unsupervised pre-training</i>				
Rectifier	1.43%	50.86%	32.64%	16.40%
Tanh	1.57%	52.62%	36.46%	19.29%
Softplus	1.77%	53.20%	35.48%	17.68%

Glorot et al., 2011a

relu

softplus

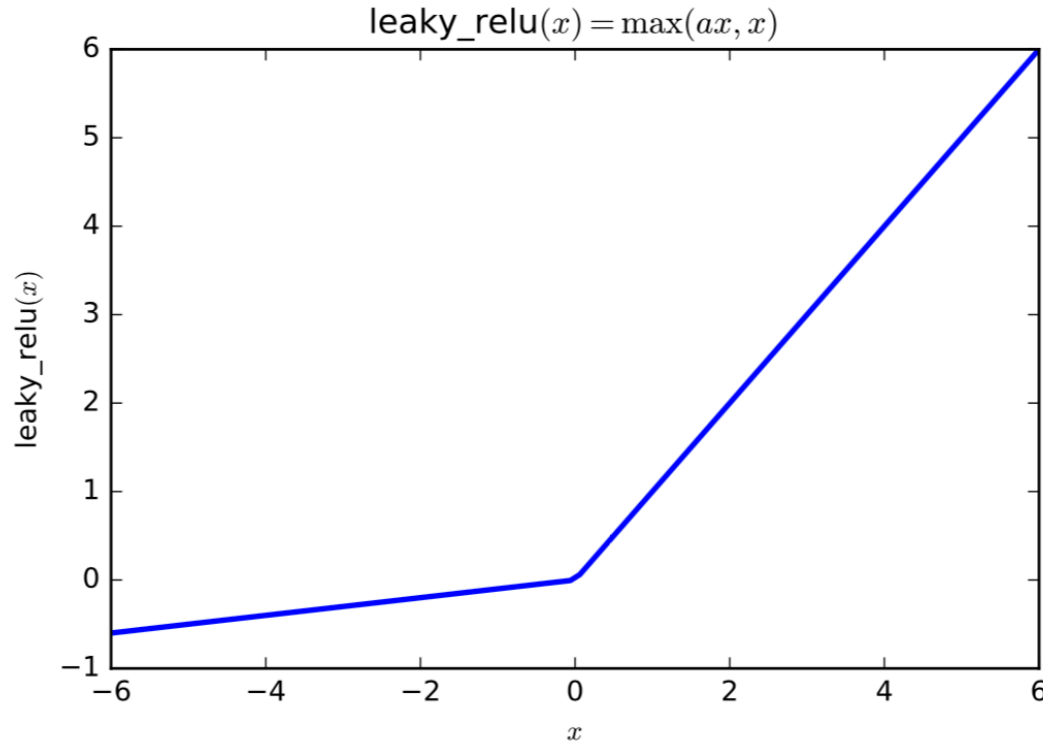


Leaky ReLU / PReLU unit

Leaky rectified linear unit, $f(x) = \max(\alpha x, x)$

```
f = lambda x: x * (x>0) + 0.1*x * (x<0)
```

$\alpha = 0.01$



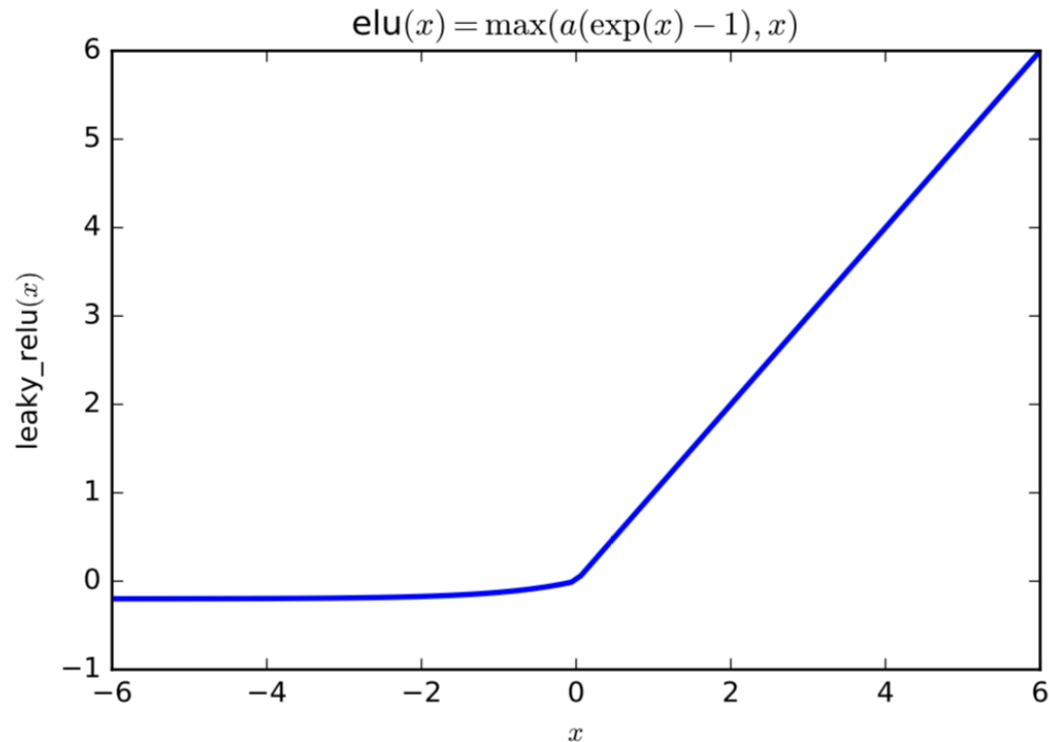
The leaky ReLU avoids the stopping of learning when $x < 0$. α may be treated as a selected hyperparameter, or it may be a parameter to be optimized in learning, in which case it is typically called the “PReLU” for parametrized rectified linear unit.



ELU unit

Exponential linear unit, $f(x) = \max(\alpha(\exp(x) - 1), x)$

```
f = lambda x: x * (x>0) + 0.2*(np.exp(x) - 1) * (x<0)
```



The exponential linear unit avoids the stopping of learning when $x < 0$. A con of this activation function is that it requires computation of the exponential, which is more expensive.

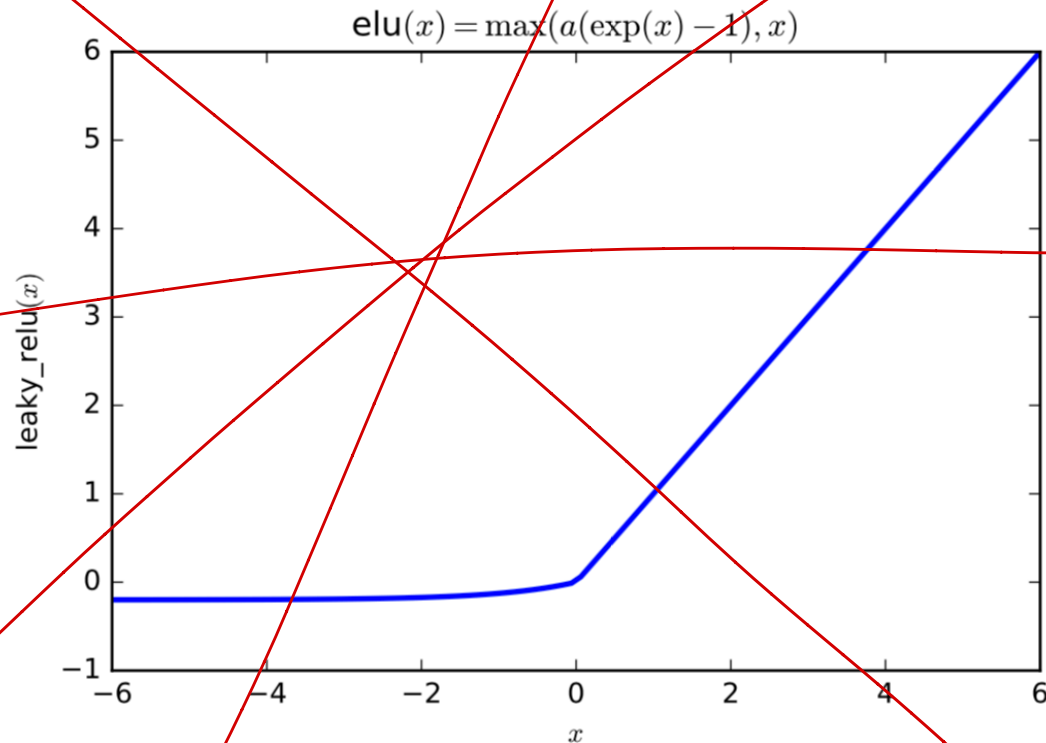


IGNORE

Swish activation function

Exponential linear unit, $f(x) = \max(\alpha(\exp(x) - 1), x)$

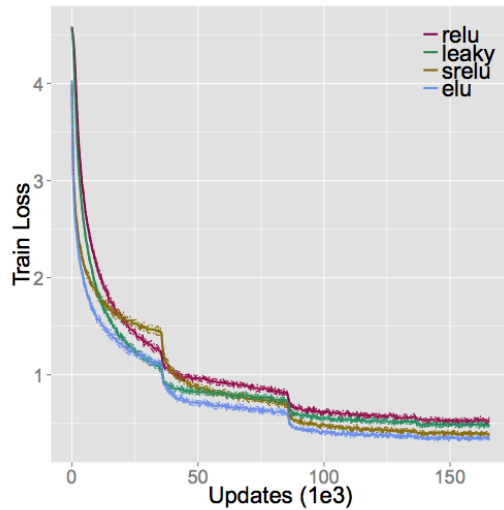
```
f = lambda x: x * (x > 0) + 0.2 * (np.exp(x) - 1) * (x < 0)
```



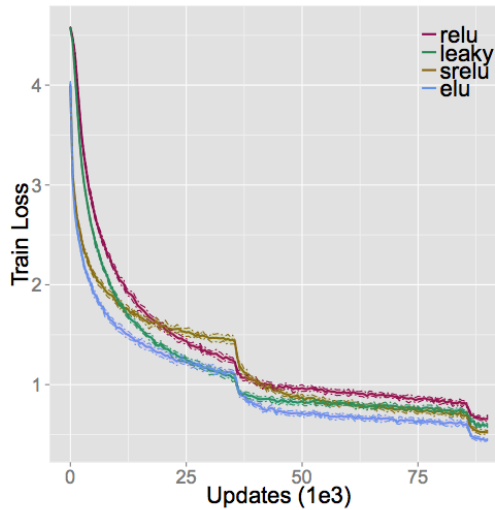
The exponential linear unit avoids the stopping of learning when $x < 0$. A con of this activation function is that it requires computation of the exponential, which is more expensive.



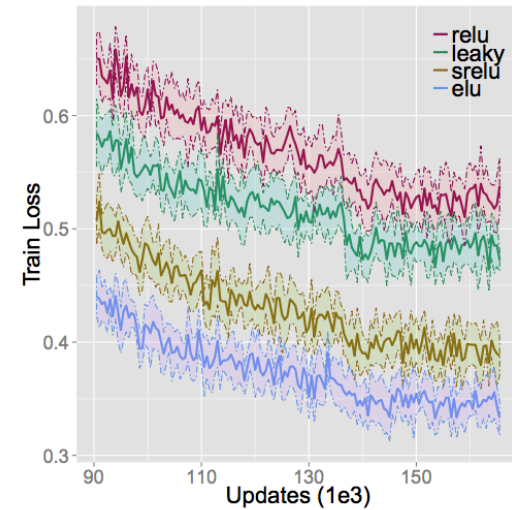
Which LU do I use?



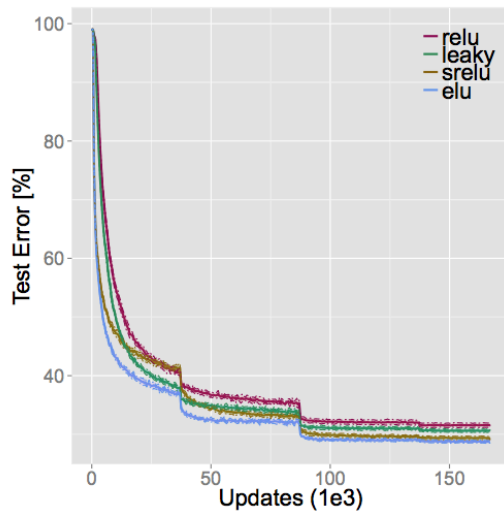
(a) Training loss



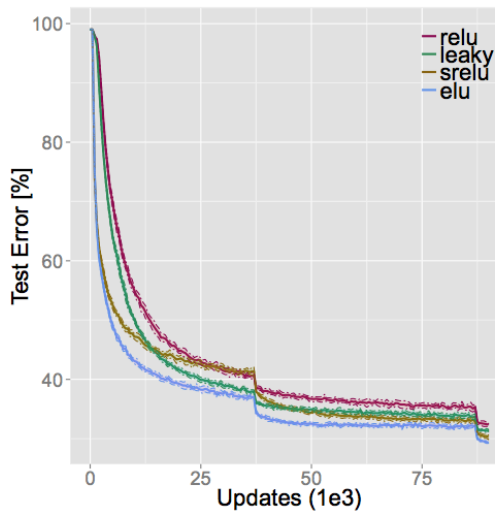
(b) Training loss (start)



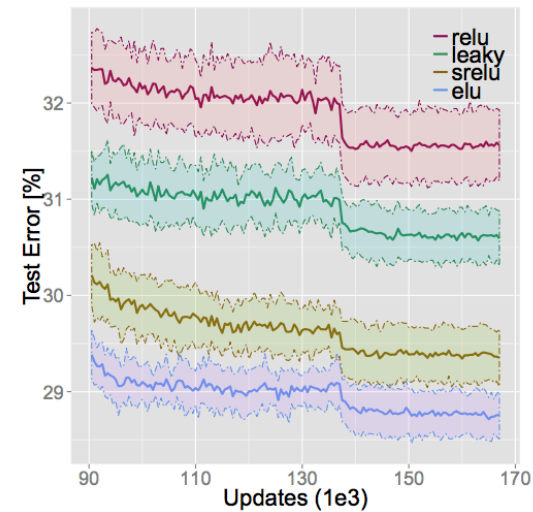
(c) Training loss (end)



(d) Test error



(e) Test error (start)



(f) Test error (end)

Figure 4: Comparison of ReLUs, LReLUs, and SReLUs on CIFAR-100. Panels (a-c) show the
Clevert et al., 2015



Maxout unit

Not tested. See textbook for more details.

Maxout unit

A generalization of the ReLU and PReLU units is the maxout unit, where:

$$\text{maxout}(\mathbf{x}) = \max(\mathbf{w}_1^T \mathbf{x} + b_1, \mathbf{w}_2^T \mathbf{x} + b_2)$$

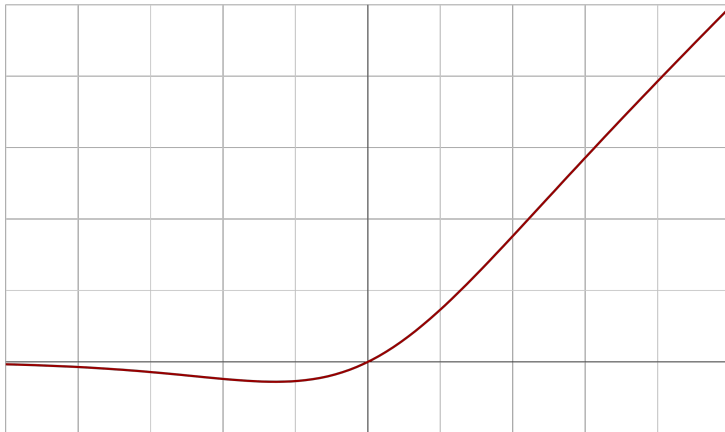
This can be generalized to more than two components. If $\mathbf{w}_1 = \mathbf{0}$ and $b_1 = 0$, this is the rectified linear unit.



More recently...

$$\text{swish}(x) = x \text{ sigmoid}(\beta x) = \frac{x}{1 + e^{-\beta x}}.$$

$= x \cdot \sigma(\beta x)$



$$\text{GELU}(x) = xP(X \leq x) = x\Phi(x) = x \cdot \frac{1}{2} \left[1 + \text{erf}(x/\sqrt{2}) \right]$$

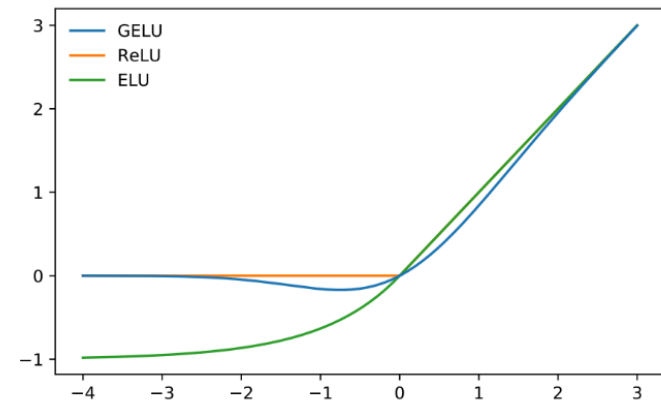


Figure 1: The GELU ($\mu = 0, \sigma = 1$), ReLU, and ELU ($\alpha = 1$).



What activation function do I use?

In practice...

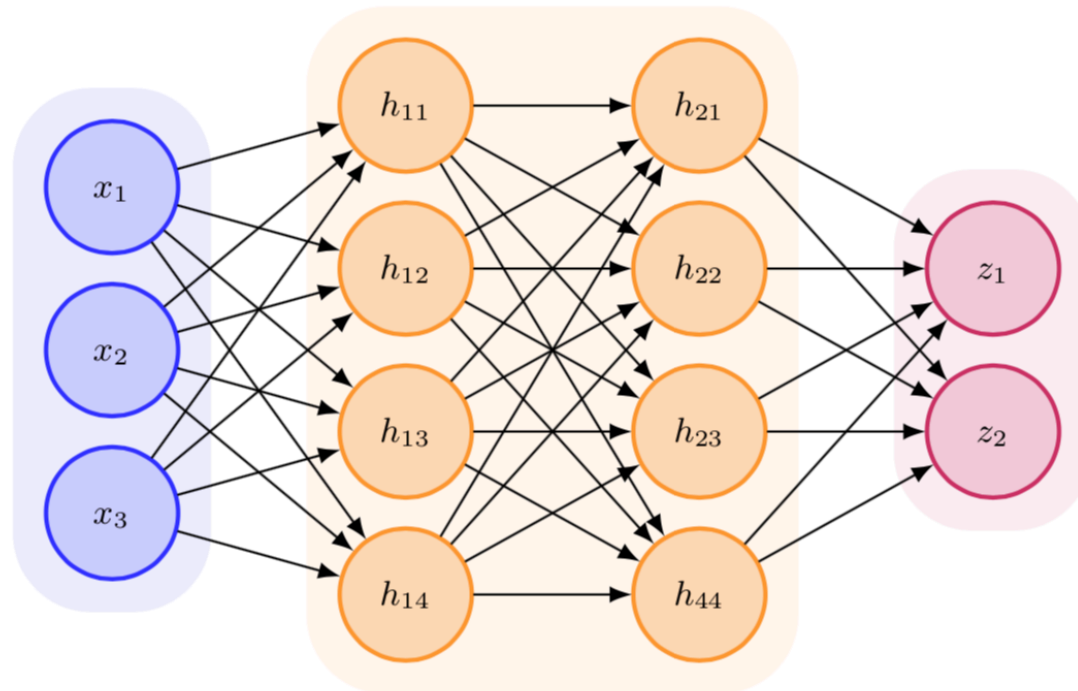
In practice...

- The ReLU unit is very popular.
- The sigmoid unit is almost never used; tanh is preferred.
- It may be worth trying out leaky ReLU / PReLU / ELU / maxout for your application.
- This is an active area of research.

swish, GELU



Back to NN architecture



- Layer 1: $\mathbf{h}_1 = f(\mathbf{W}_1 \mathbf{x} + \mathbf{b}_1)$
- Layer 2: $\mathbf{h}_2 = f(\mathbf{W}_2 \mathbf{h}_1 + \mathbf{b}_2)$
- $\vdots$
- Layer N: $\mathbf{z} = \mathbf{W}_N \mathbf{h}_{N-1} + \mathbf{b}_N$

What output activation do we use?



The output activation interacts with the cost function

large ^{positive} $z \rightarrow$ class 1 ; $y^{(i)} = 1$

What outputs and cost functions?

large negative $z \rightarrow$ class 0 ; $y^{(i)} = 0$

There are several options to process the output scores, z , to ultimately arrive at a cost function. The choice of output units interacts with what cost function to use.

$$\sigma(z) = \Pr\{x^{(i)} \text{ belongs to class 1}\}$$

Example: Consider a neural network that produces a single score, z , for binary classification. As the output unit, we choose the sigmoid nonlinearity, so that $\hat{y} = \sigma(z)$. On a given example, $y^{(i)}$ is either 0 or 1, and $\hat{y}^{(i)} = \sigma(z^{(i)})$ can be interpreted as the algorithm's probability the output is in class 1. Is it better to use mean-square error or cross-entropy (i.e., corresponding to maximum-likelihood estimation) as the cost function? For n examples:

$$\text{MSE} = \frac{1}{2} \sum_{i=1}^n \left(y^{(i)} - \sigma(z^{(i)}) \right)^2$$

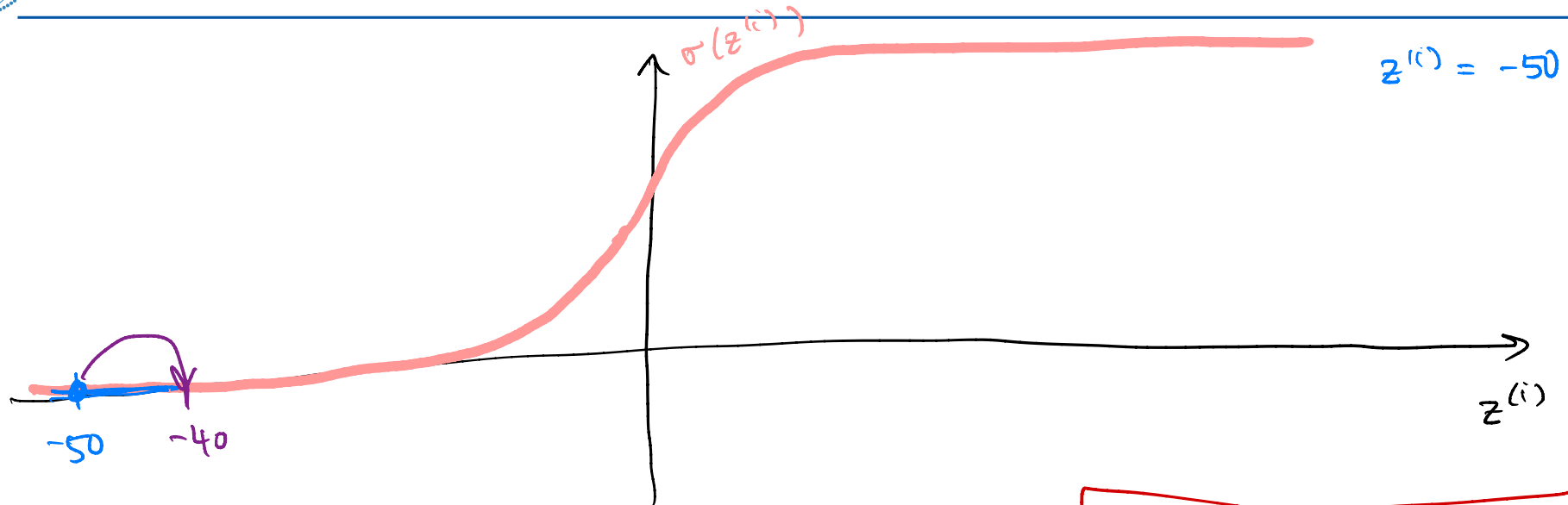
$y^{(i)} = 1$

$$\text{CE} = - \sum_{i=1}^n \left[y^{(i)} \log \sigma(z^{(i)}) + (1 - y^{(i)}) \log(1 - \sigma(z^{(i)})) \right]$$



A picture for intuition

$$y^{(i)} = 1$$



$$\text{MSE for ex. } i: \text{MSE}^{(i)} = (y^{(i)} - \sigma(z^{(i)}))^2$$

$$\frac{\partial \text{MSE}^{(i)}}{\partial z^{(i)}} = \underbrace{\frac{\partial \sigma(z^{(i)})}{\partial z^{(i)}}}_{\text{grad}} \frac{\partial \text{MSE}^{(i)}}{\partial \sigma(z^{(i)})}$$

$$\text{when } z^{(i)} = -50$$

$$\text{grad} \approx 0$$

$$(y^{(i)} - \sigma(-50))^2$$

$$(1 - 0)^2 \approx 1$$

$$(y^{(i)} - \sigma(-40))^2$$

$$\approx 1$$



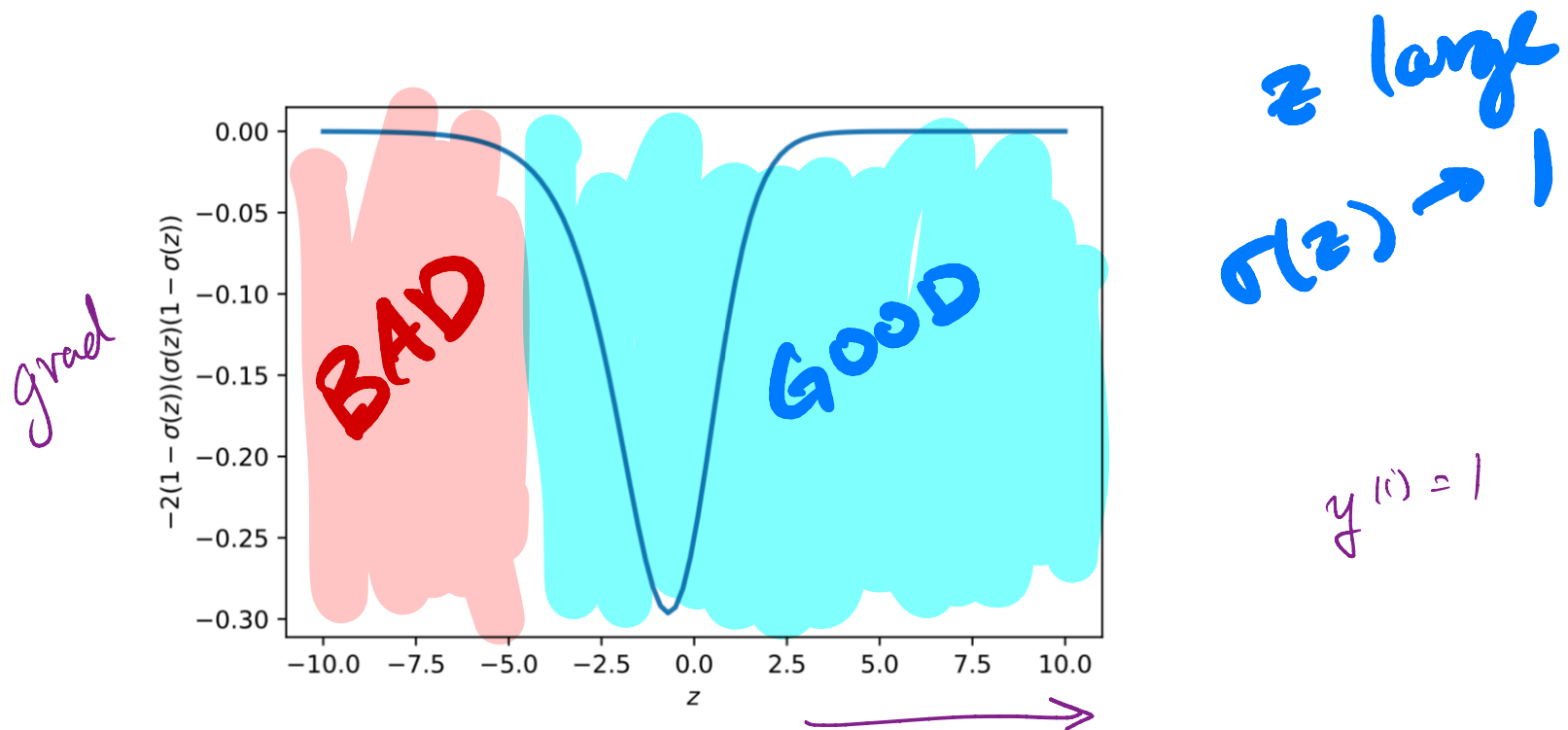
The output activation interacts with the cost function

What outputs and cost functions? (cont.)

Example (cont): Consider just one example, where $y^{(i)} = 1$. For this example,

$$\frac{\partial \text{MSE}}{\partial z^{(i)}} = -2(y^{(i)} - \sigma(z^{(i)}))(\sigma(z^{(i)})(1 - \sigma(z^{(i)})))$$

This derivative looks like the following:



When z is very negative, indicating a large MSE, the gradient saturates to zero, and no learning occurs.



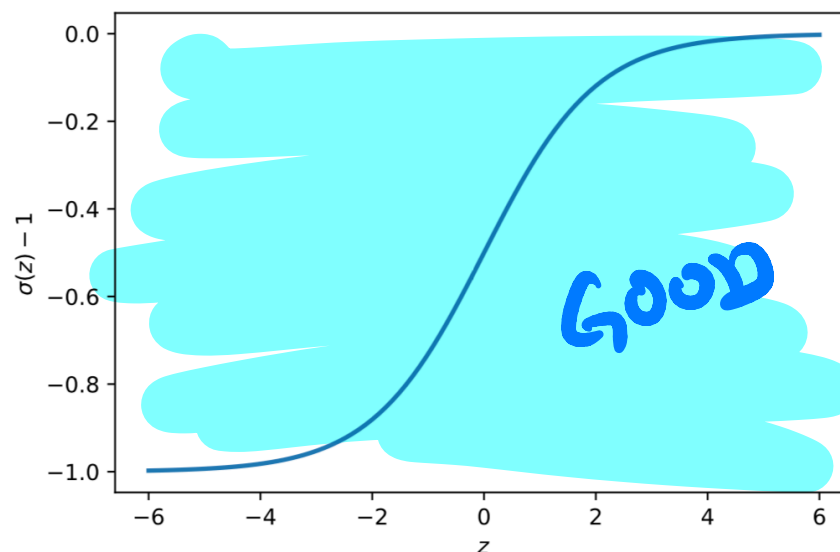
The output activation interacts with the cost function

What outputs and cost functions? (cont.)

Example (cont): Now let's consider the cross-entropy. For one example, where $y^{(i)} = 1$,

$$\frac{\partial \text{CE}}{\partial z^{(i)}} = \sigma(z^{(i)}) - 1$$

This derivative looks like the following:



Notice that when z is very negative, learning will occur, and it will only "stall" when z gets close to the right answer.



Output activations

NOT TESTED

- Linear output units: $\hat{\mathbf{y}} = \mathbf{z}$.

These output units typically specify the conditional mean of a Gaussian distribution, i.e.,

$$p(\mathbf{y}|\mathbf{z}) = \mathcal{N}(\mathbf{z}, \mathbf{I})$$

and in this case, MLE estimation is equivalent to minimizing squared error.



Output activations

- Sigmoid outputs: $\hat{y} = \sigma(\mathbf{z})$.

These outputs are typically used in binary classification to approximate a Bernoulli distribution.

Question: Why not use the following output?

$$\Pr(y = 1|\mathbf{x}) = \max(0, \min(1, \mathbf{z}))$$



Output activations

- Softmax output: $\hat{\mathbf{y}}_i = \text{softmax}_i(\mathbf{z})$.

The softmax is the generalization of the sigmoid output to multiple classes.

A softmax output activation is fairly common.

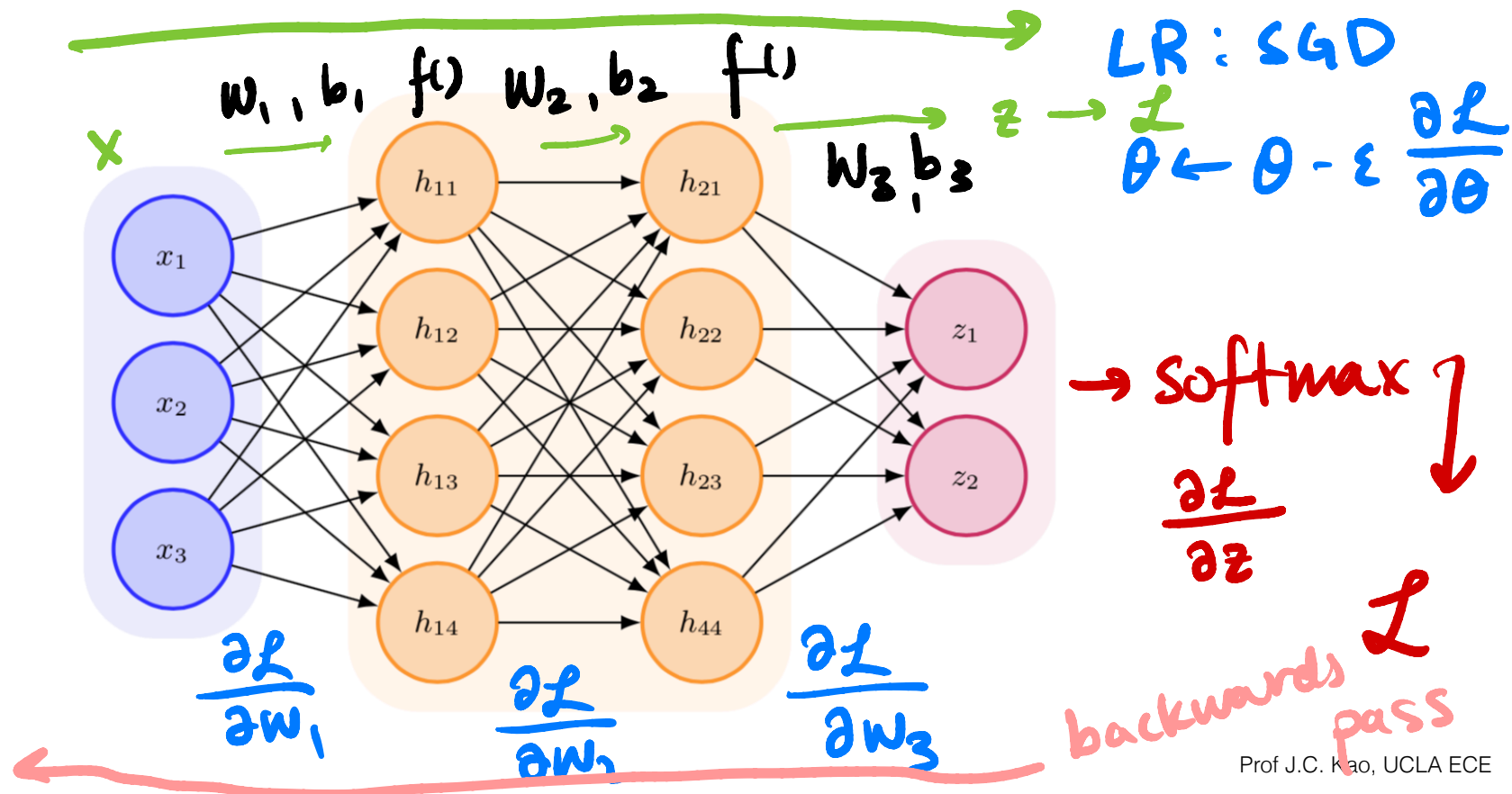


What next?

Now that we've defined all the elements of the neural network, the question now becomes: how do we learn its parameters?

The short answer is that we use versions of gradient descent.

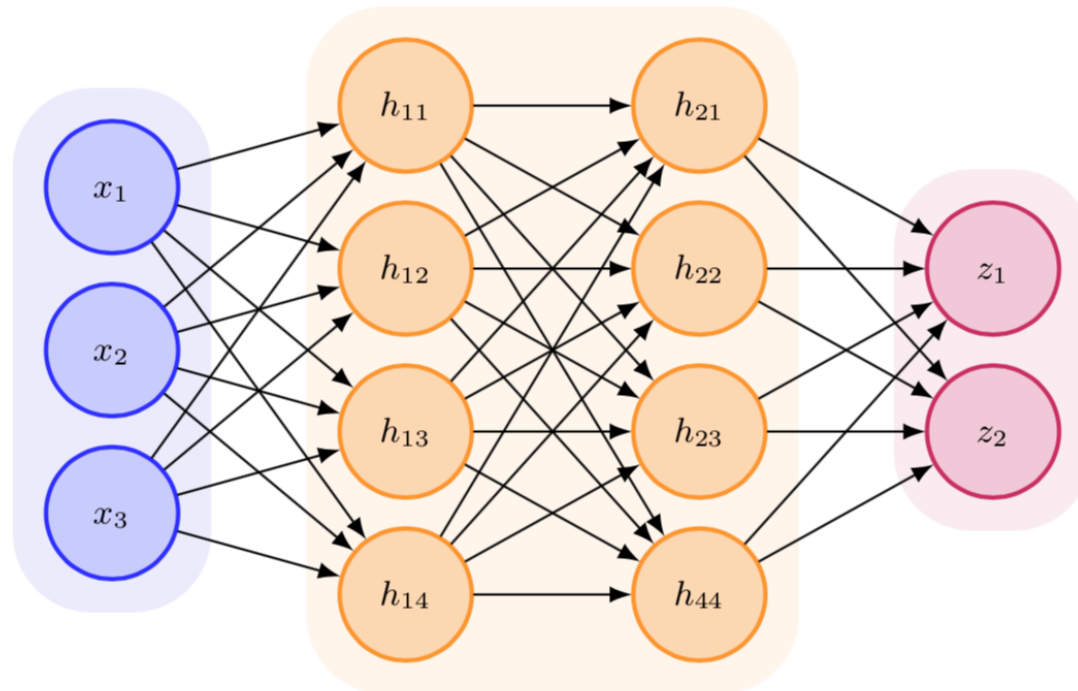
However, neural networks architectures have units that are several layers from the output. In these scenarios, how do we arrive at the gradient?





Backpropagation

In this lecture, we'll introduce backpropagation as a technique to calculate the gradient of the loss function with respect to parameters in a neural network.

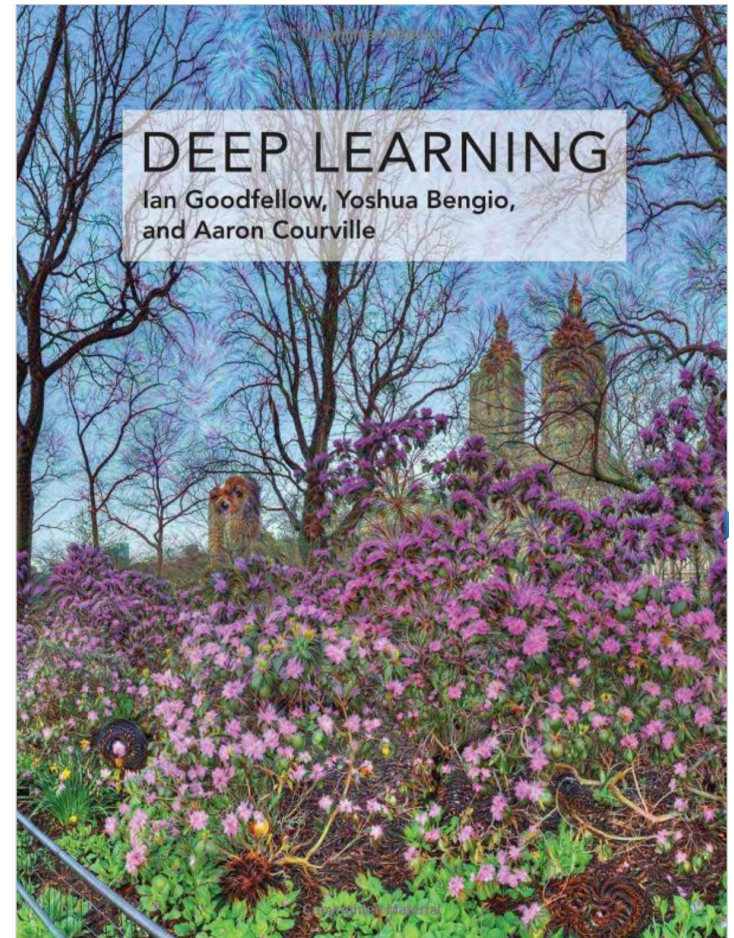




Reading

Reading:

Deep Learning, 6.5-6.6



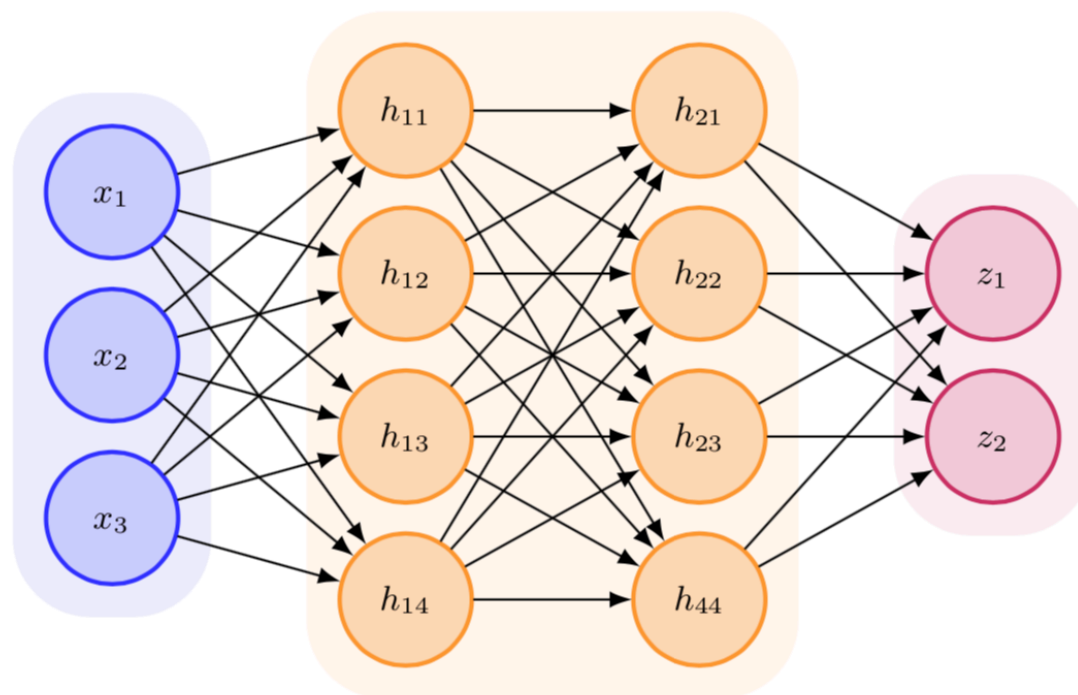


What next?

Now that we've defined all the elements of the neural network, the question now becomes: how do we learn its parameters?

The short answer is that we use versions of gradient descent.

However, neural networks architectures have units that are several layers from the output. In these scenarios, how do we arrive at the gradient?





Backpropagation

Intuitively, backpropagation is the application of the chain rule for derivatives.

Motivation for backpropagation

To do gradient descent, we need to calculate the gradient of the objective with respect to the parameters, θ . However, in a neural network, some parameters are not directly connected to the output. How do we calculate then the gradient of the objective with respect to these parameters? Backpropagation answers this question, and is a general application of the chain rule for derivatives.



Backpropagation

Nomenclature

Forward propagation:

- Forward propagation is the act of calculating the values of the hidden and output units of the neural network given an input.
- It involves taking input x , propagating it through each hidden unit sequentially, until you arrive at the output y . From forward propagation, we can also calculate the cost function $J(\theta)$.
- In this manner, the forward propagated signals are the activations.
- With the input as the “start” and the output as the “end,” information propagates in a forward manner.

Backpropagation (colloquially called backprop):

- As its name suggests, now information is passed backward through the network, from the cost function and outputs to the inputs.
- The signal that is backpropagated are the gradients.
- It enables the calculation of gradients at every stage going back to the input layer.



Backpropagation

Why do we need backpropagation?

- Backpropagation is computationally efficient because we will find that most of the terms used in backpropagation can be cached from the forward pass.
- Informally (at least for me) sometimes taking analytical multivariate gradients can be challenging. Backpropagation breaks down these multivariate gradients into easier steps. *Operationalizes gradients.*

A few further notes on backpropagation.

- Backpropagation is not the learning algorithm. It's the method of computing gradients.
- Backpropagation is not specific to multilayer neural networks, but a general way to compute derivatives of functions.

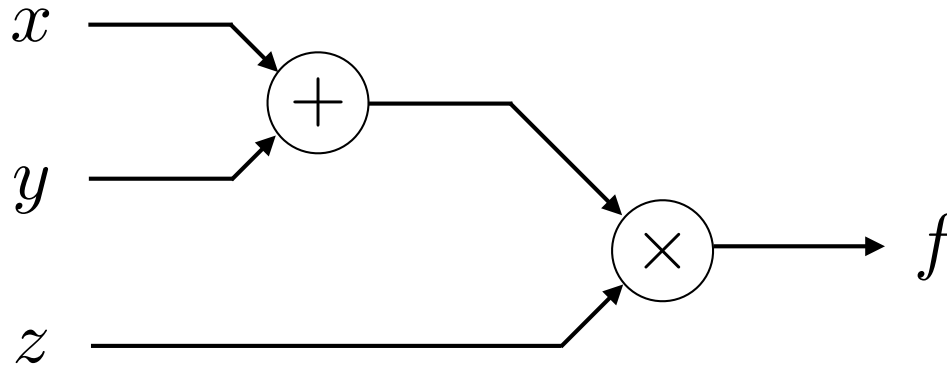


A simple example

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

$$\frac{\partial f}{\partial z} = x + y$$

$$\frac{\partial f}{\partial x} = \frac{\partial f}{\partial y} = z$$

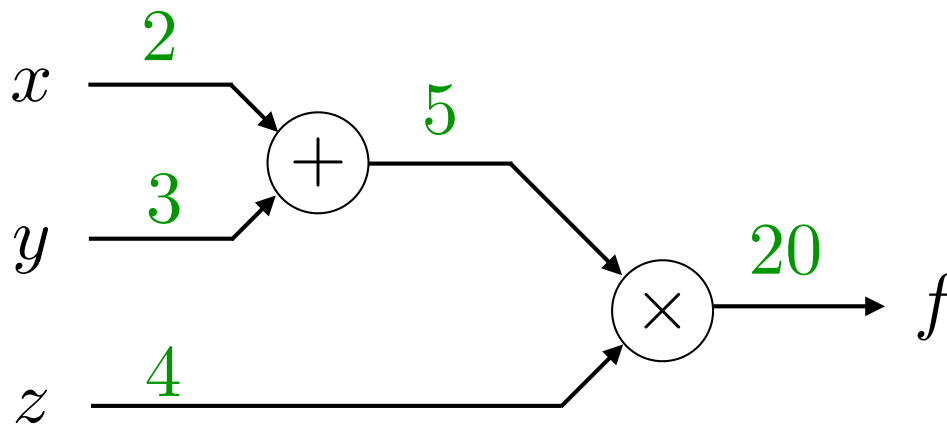




A simple example

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

Forward propagation:





A simple example

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

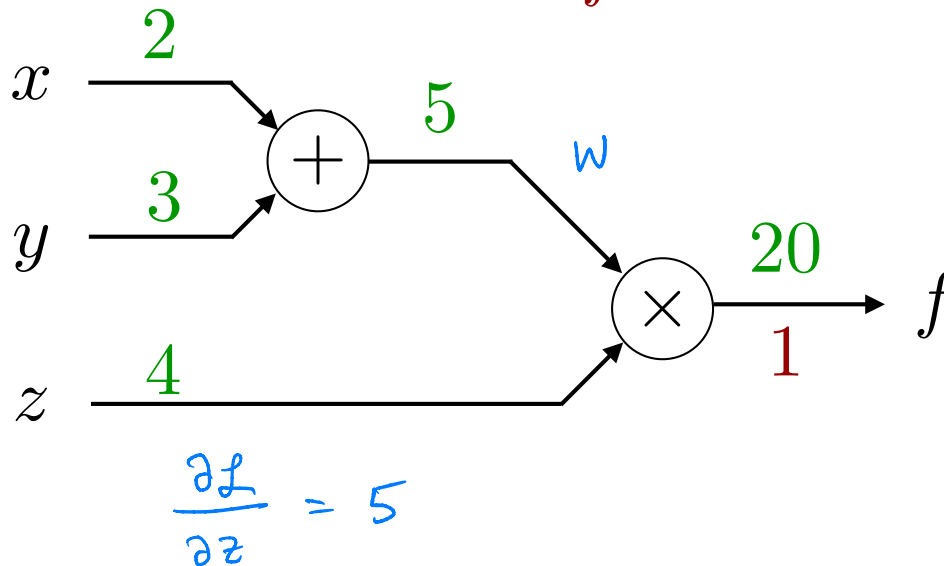
$$\frac{\partial \mathcal{L}}{\partial z} = \frac{\partial f}{\partial z} \frac{\partial \mathcal{L}}{\partial f}$$

(Handwritten annotations: a blue bracket above $\frac{\partial f}{\partial z}$ is labeled "=5", and a red bracket above $\frac{\partial \mathcal{L}}{\partial f}$ is labeled "=1")

$$f = w \cdot z \Rightarrow \frac{\partial f}{\partial z} = w$$

Backpropagation:

$$\frac{\partial \mathcal{L}}{\partial f} = 1$$



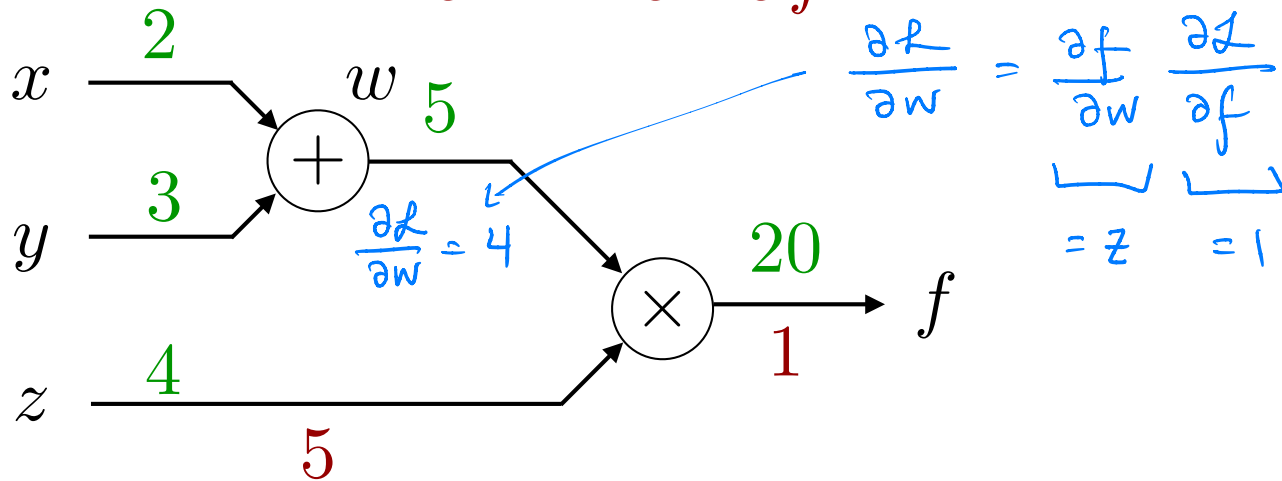


A simple example

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

Backpropagation:

$$\frac{\partial \mathcal{L}}{\partial z} = \frac{\partial f}{\partial z} \frac{\partial \mathcal{L}}{\partial f}$$





A simple example

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

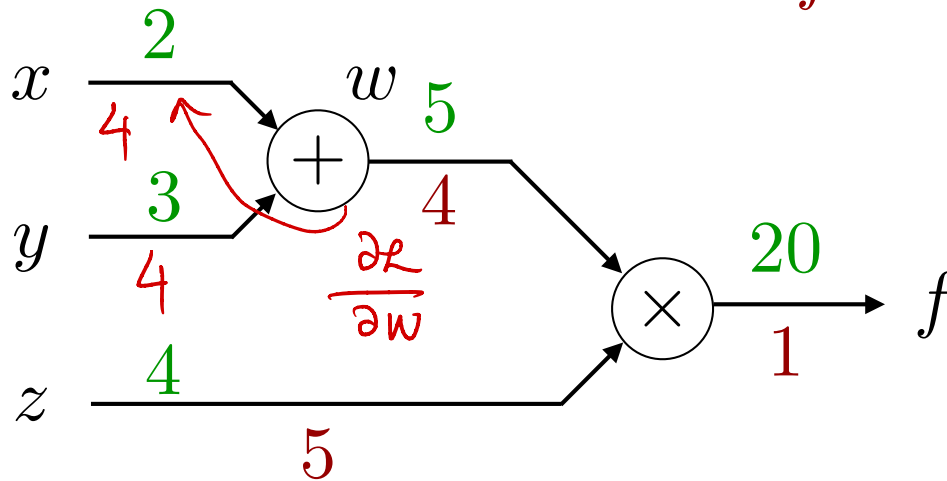
Backpropagation:

$$\frac{\partial \mathcal{L}}{\partial w} = \frac{\partial f}{\partial w} \frac{\partial \mathcal{L}}{\partial f}$$

$$w = x + y$$

$$\frac{\partial w}{\partial x} = 1$$

$$\frac{\partial w}{\partial y} = 1$$



$$\frac{\partial \mathcal{L}}{\partial x} = \frac{\partial w}{\partial x} \frac{\partial \mathcal{L}}{\partial w}$$

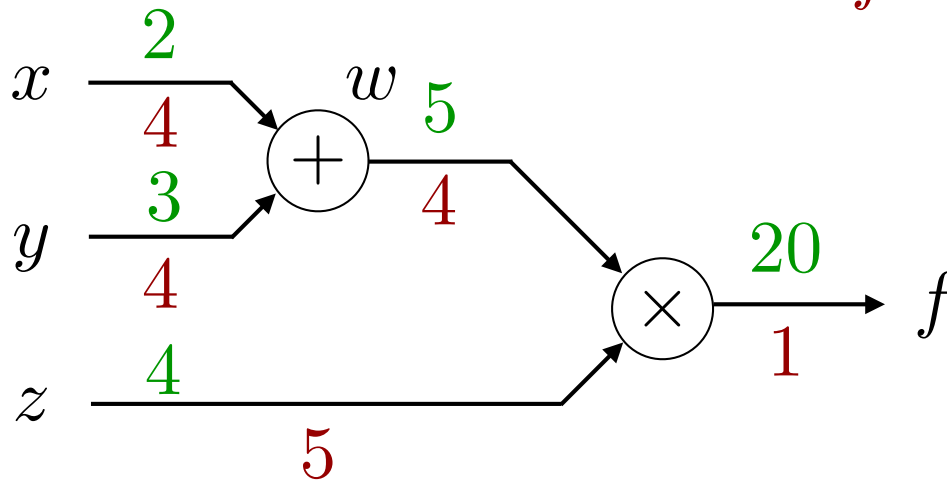


A simple example

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

Backpropagation:

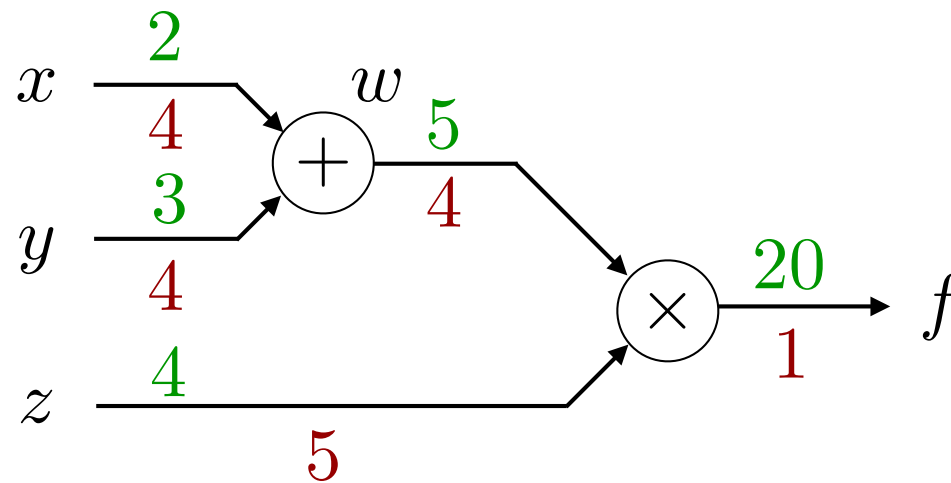
$$\frac{\partial \mathcal{L}}{\partial x} = \frac{\partial w}{\partial x} \frac{\partial f}{\partial w} \frac{\partial \mathcal{L}}{\partial f}$$





Idea: computational graphs apply to gradients

In the **forward pass**, we apply a function to the node inputs to calculate an output. In the **backward pass**, we take the **upstream derivative** and apply a **local gradient** to calculate the backpropagated derivative.

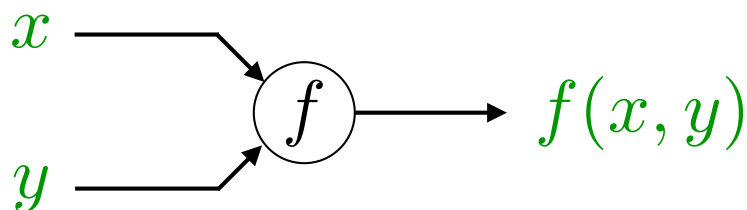




Idea: computational graphs apply to gradients

In the **forward pass**, we apply a function to the node inputs to calculate an output.
In the **backward pass**, we take the **upstream derivative** and apply a **local gradient** to calculate the backpropagated derivative.

Forward pass:

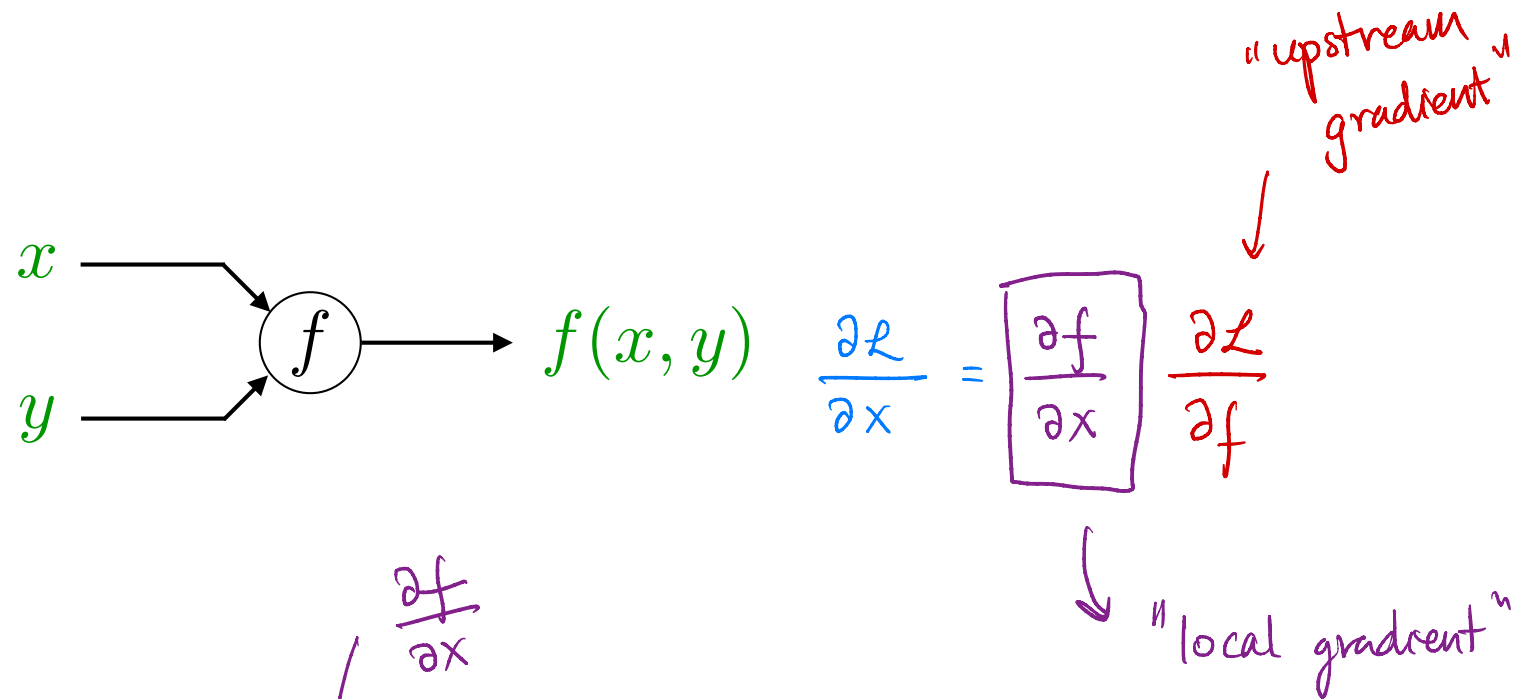




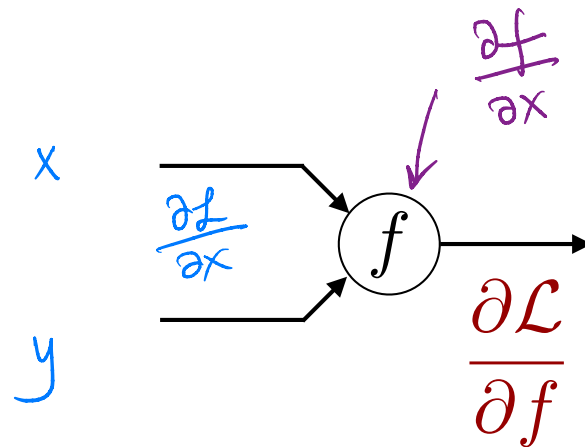
Idea: computational graphs apply to gradients

In the **forward pass**, we apply a function to the node inputs to calculate an output.
In the **backward pass**, we take the **upstream derivative** and apply a **local gradient** to calculate the backpropagated derivative.

Forward pass:



Backward pass:

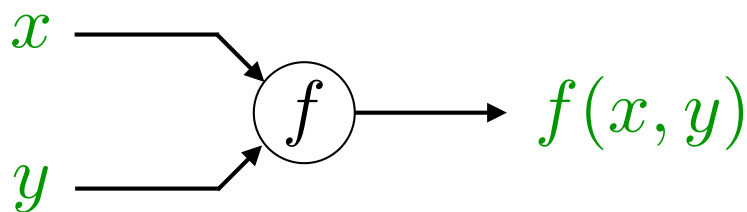




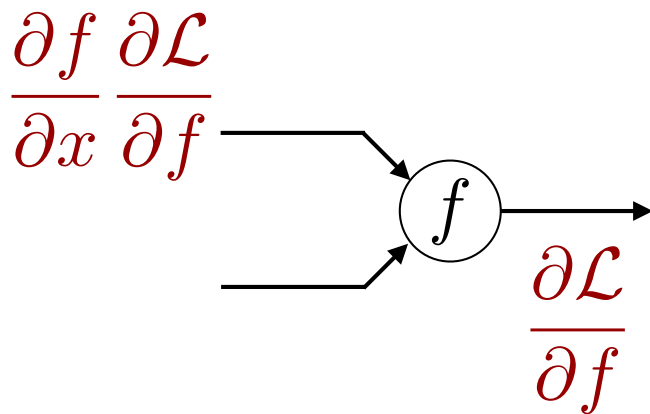
Idea: computational graphs apply to gradients

In the **forward pass**, we apply a function to the node inputs to calculate an output.
In the **backward pass**, we take the **upstream derivative** and apply a **local gradient** to calculate the backpropagated derivative.

Forward pass:



Backward pass:

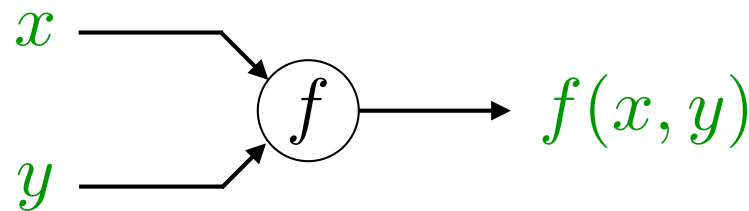




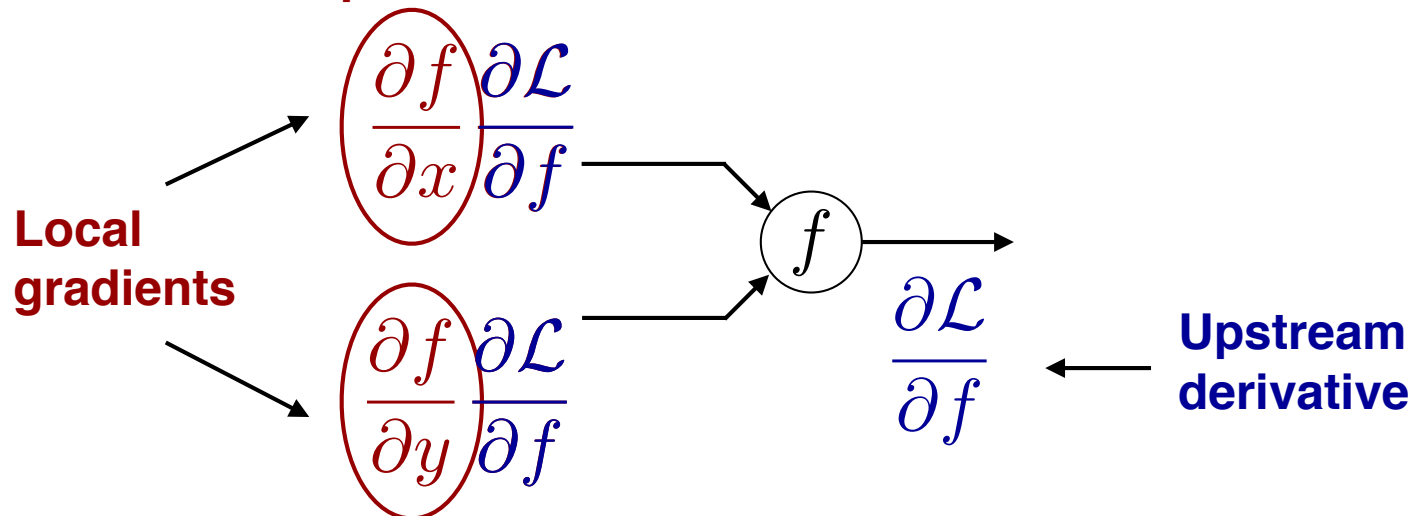
Idea: computational graphs apply to gradients

In the **forward pass**, we apply a function to the node inputs to calculate an output.
In the **backward pass**, we take the **upstream derivative** and apply a **local gradient** to calculate the backpropagated derivative.

Forward pass:

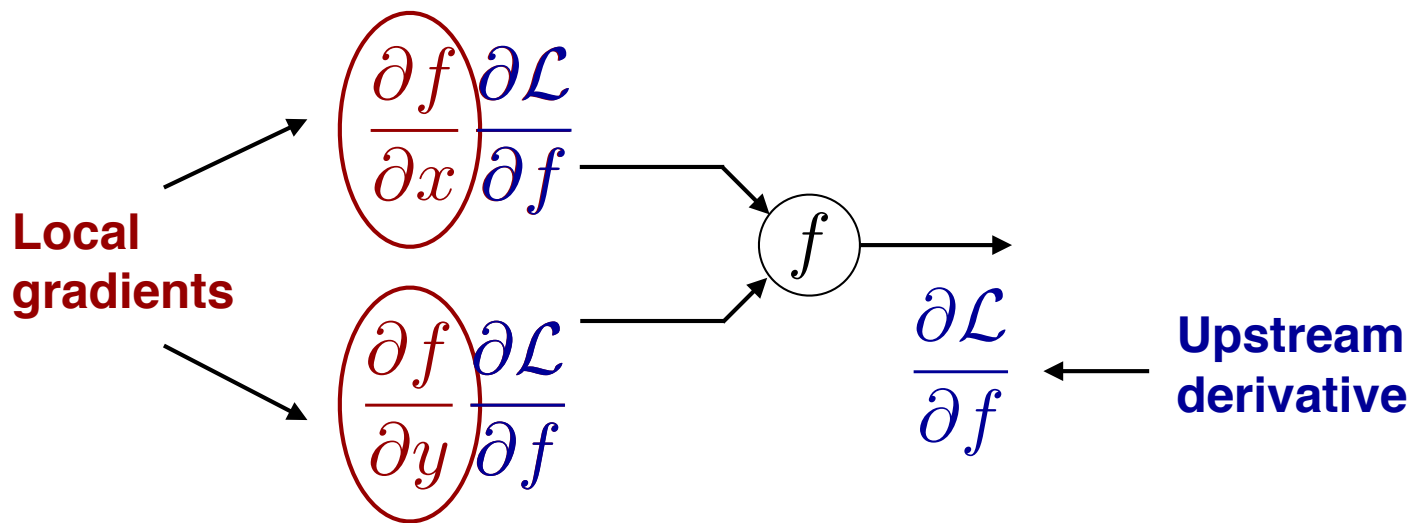


Backward pass:





Idea: computational graphs apply to gradients



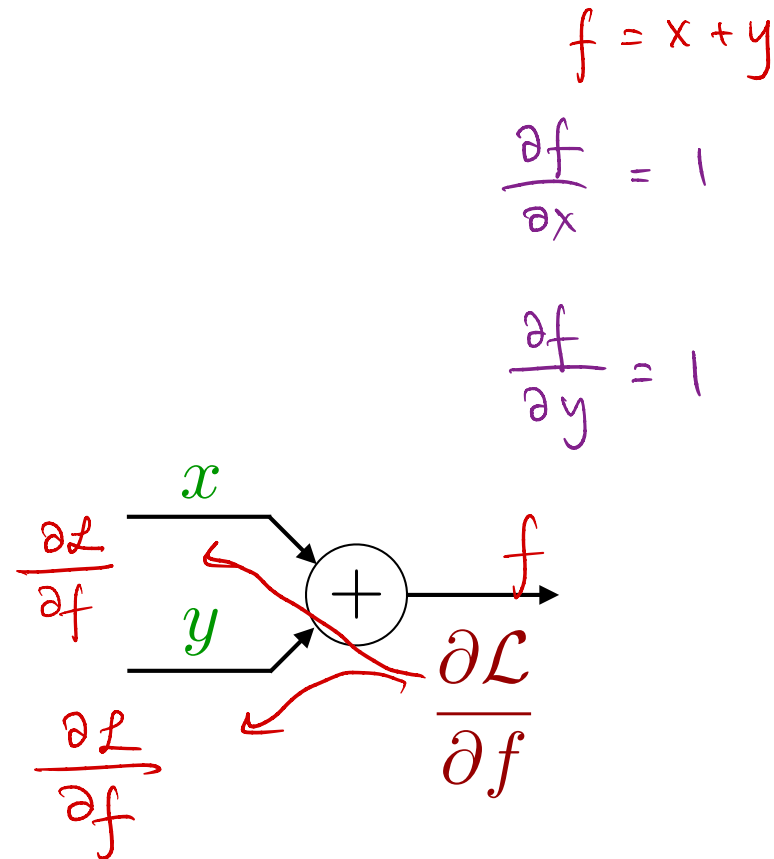
The basic intuition of backpropagation is that we break up the calculation of the gradient into **small and simple steps**. Each of these nodes in the graph is a straightforward gradient calculation, where we multiply an **input** (the “upstream derivative”) with a **local gradient** (an application of the chain rule).

Composing all of these gradients together returns the overall gradient.



A gate view of gradients

Interpreting backpropagation as gradient “gates”:

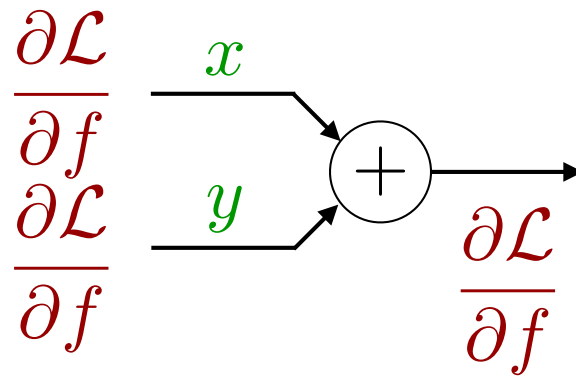




A gate view of gradients

Interpreting backpropagation as gradient “gates”:

Add gate: distributes the gradient

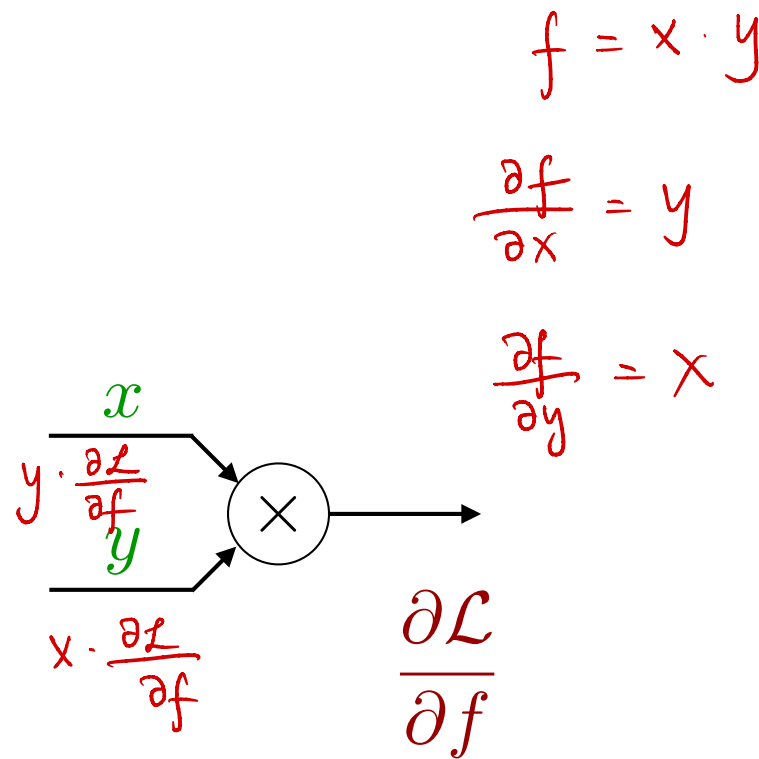




A gate view of gradients

Interpreting backpropagation as gradient “gates”:

Add gate: distributes the gradient



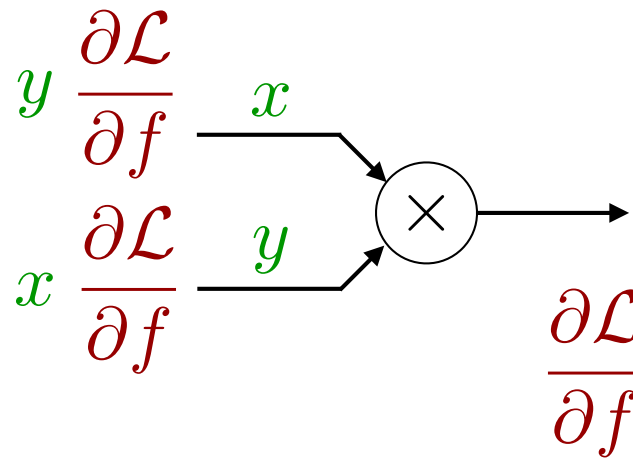


A gate view of gradients

Interpreting backpropagation as gradient “gates”:

Add gate: distributes the gradient

Mult gate: switches the gradient





A gate view of gradients

Interpreting backpropagation as gradient “gates”:

Add gate: distributes the gradient

Mult gate: switches the gradient

$$\text{relu}(x) = \max(x, 0)$$

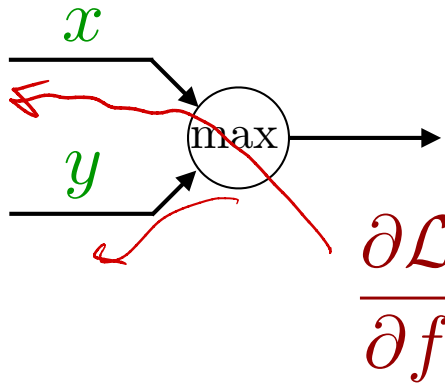
$$f = \max(x, y)$$

$$\frac{\partial f}{\partial x} = \begin{cases} 1, & \text{if } x > y \\ 0, & \text{else} \end{cases}$$

$$= \mathbb{I}\{x > y\}$$

$$\mathbb{I}(x > y) \frac{\partial \mathcal{L}}{\partial f}$$

$$\mathbb{I}(y > x) \frac{\partial \mathcal{L}}{\partial f}$$





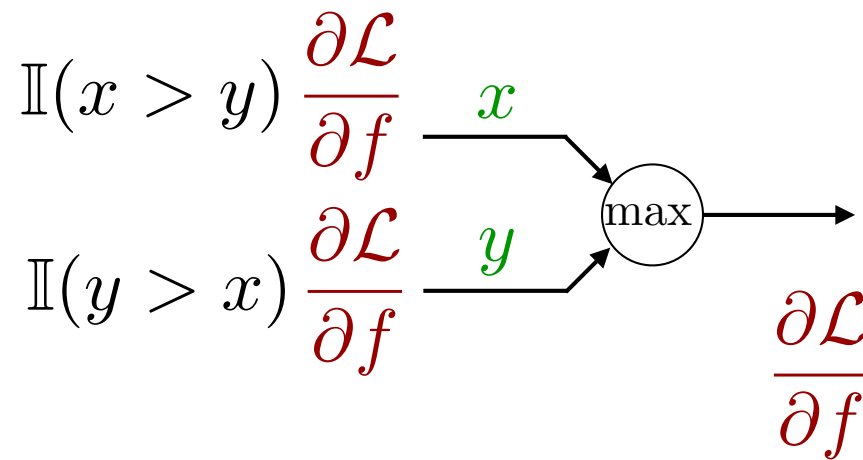
A gate view of gradients

Interpreting backpropagation as gradient “gates”:

Add gate: distributes the gradient

Mult gate: switches the gradient

Max gate: routes the gradient





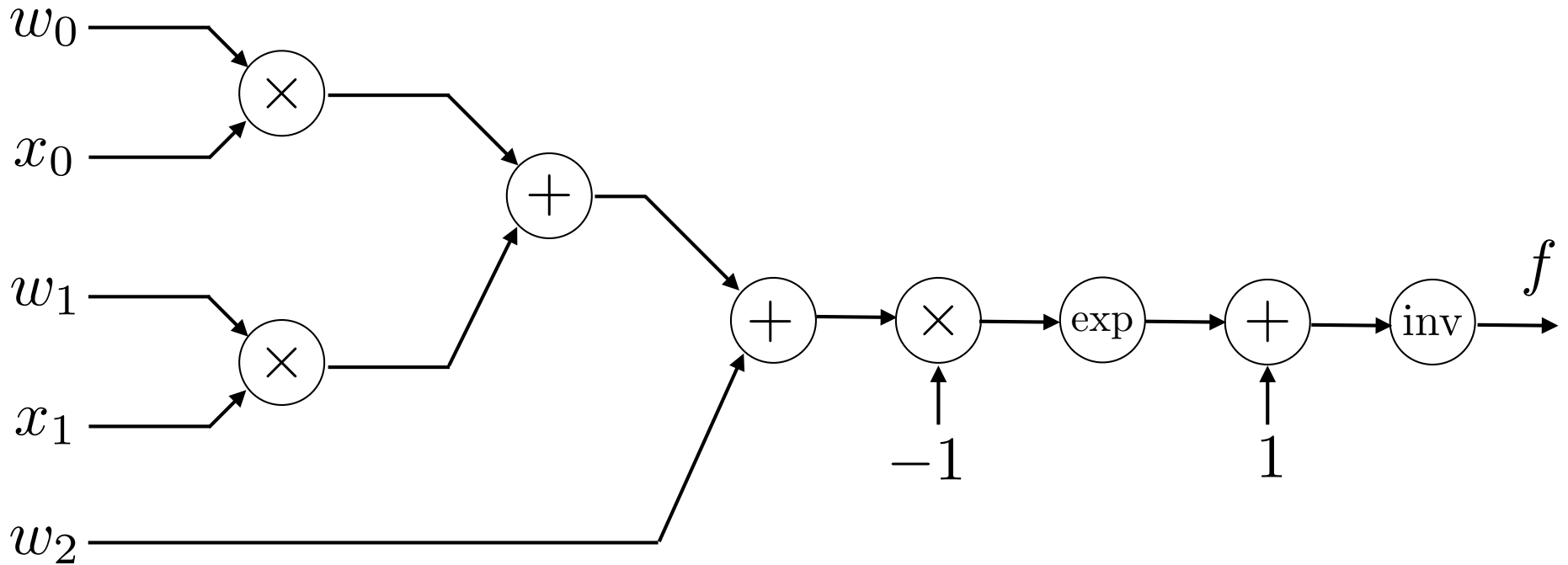
A more involved scalar example

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



A more involved scalar example

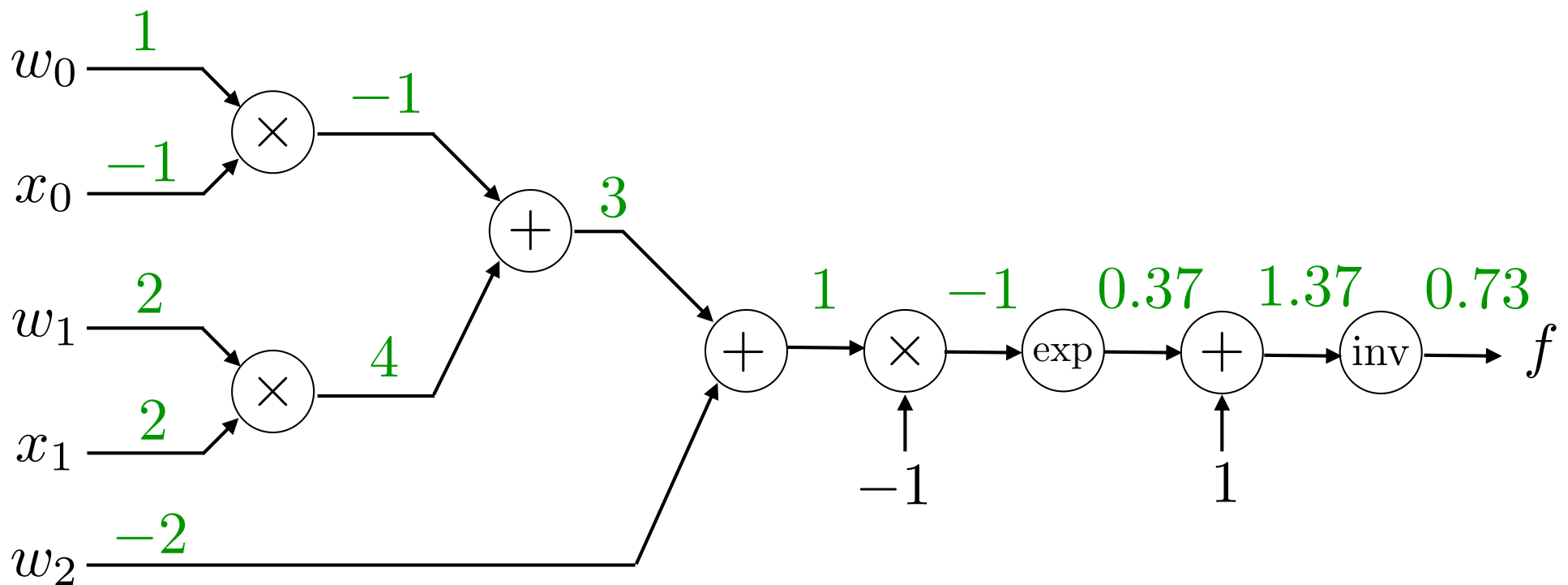
$$f(\mathbf{w}, \mathbf{x}) = \frac{1}{1 + \exp(-(w_0x_0 + w_1x_1 + w_2))}$$





A more involved scalar example

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





A more involved scalar example

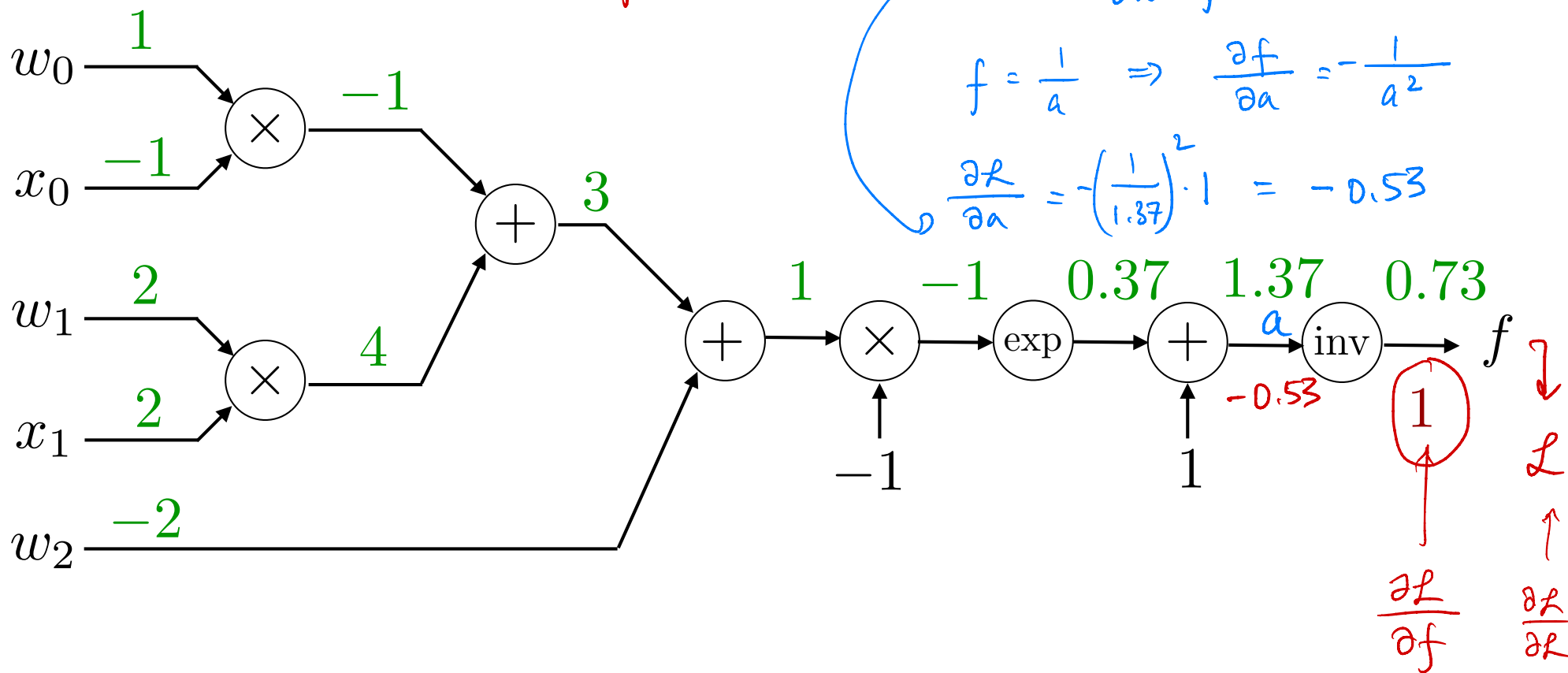
$$f(\mathbf{w}, \mathbf{x}) = \frac{1}{1 + \exp(-(w_0x_0 + w_1x_1 + w_2))}$$

$$\frac{\partial \mathcal{L}}{\partial f} = 1$$

$$\frac{\partial \mathcal{L}}{\partial a} = \frac{\partial f}{\partial a} \frac{\partial \mathcal{L}}{\partial f}$$

$$f = \frac{1}{a} \Rightarrow \frac{\partial f}{\partial a} = -\frac{1}{a^2}$$

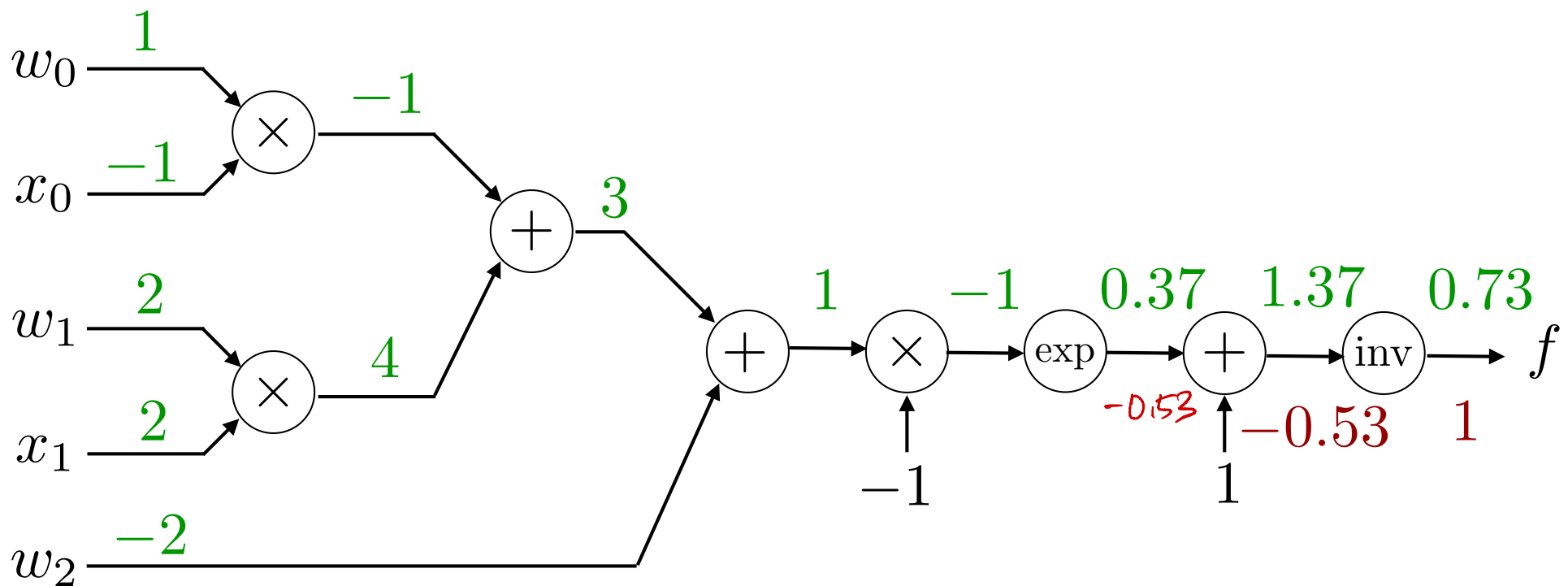
$$\frac{\partial \mathcal{L}}{\partial a} = -\left(\frac{1}{1.37}\right)^2 \cdot 1 = -0.53$$





A more involved scalar example

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



$$\frac{\partial \mathcal{L}}{\partial z} = \frac{\partial f}{\partial z} \frac{\partial \mathcal{L}}{\partial f}$$

$$\frac{\partial f}{\partial z} = -\frac{1}{z^2}$$



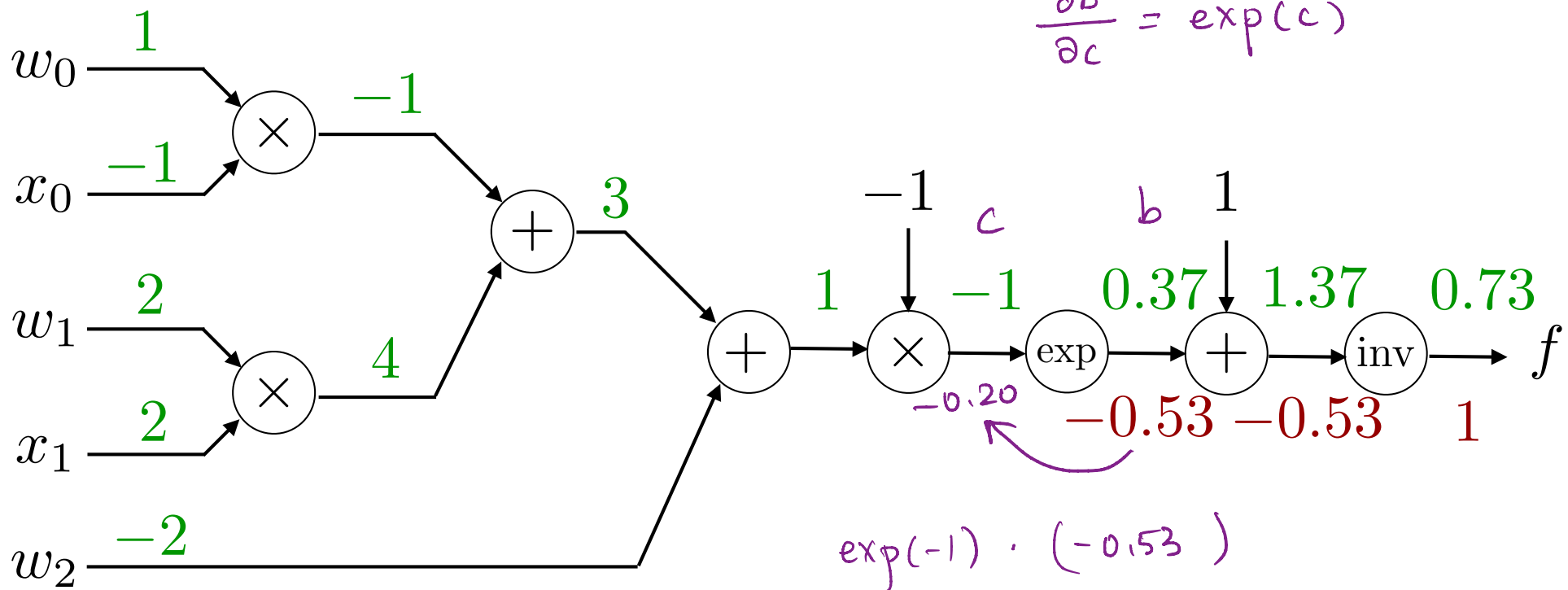
A more involved scalar example

$$f(x) = e^x$$
$$\frac{\partial f}{\partial x} = e^x$$

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

$$b = \exp(c) = e^c$$

$$\frac{\partial b}{\partial c} = \exp(c)$$

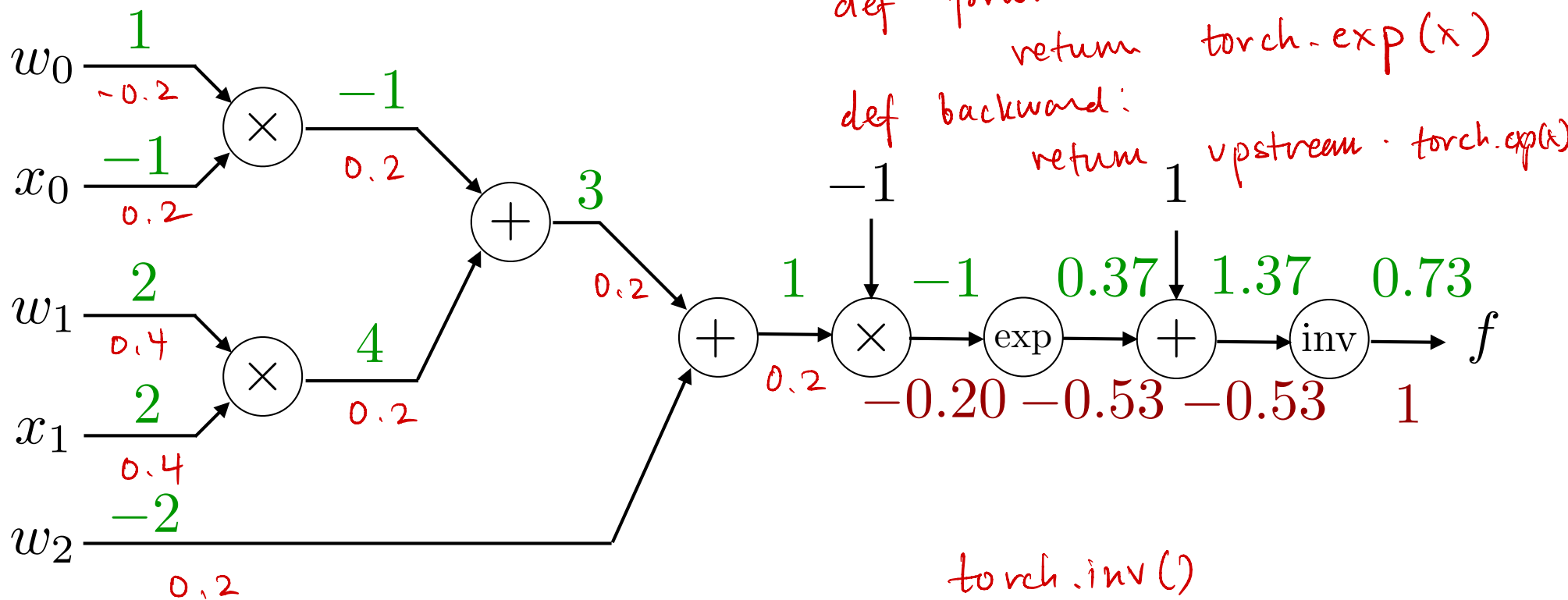




A more involved scalar example

f.backward()

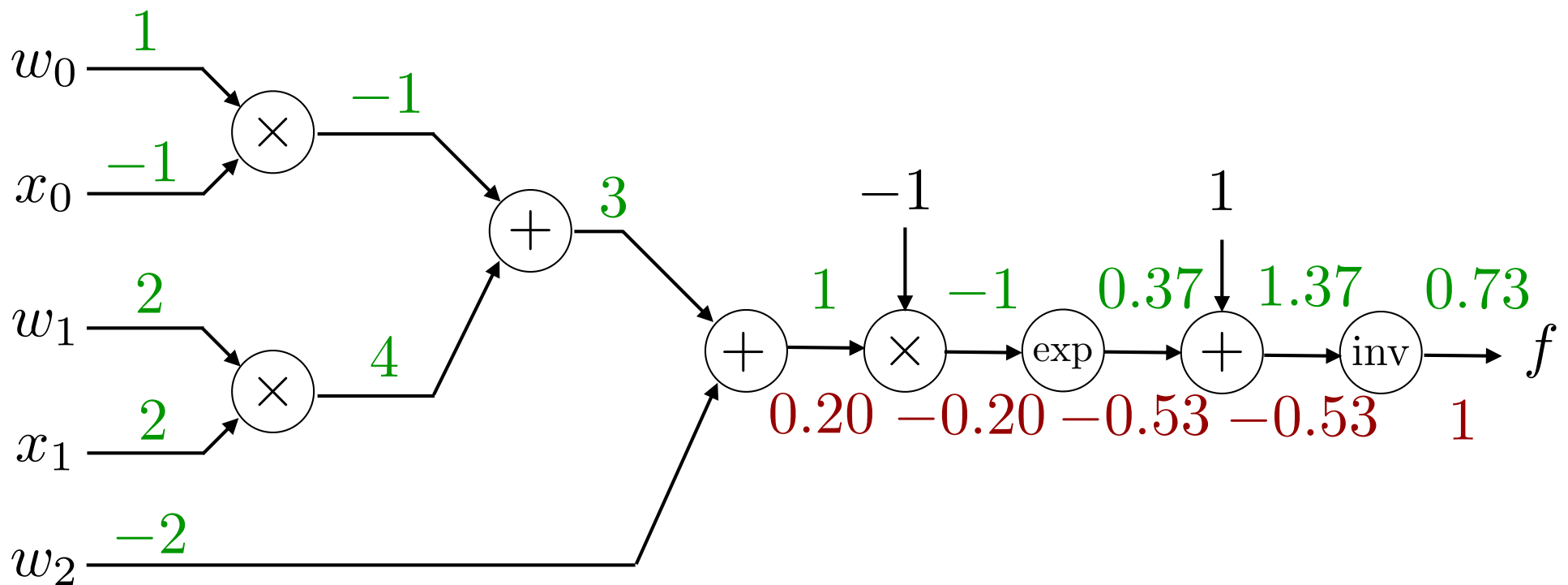
$$f(\mathbf{w}, \mathbf{x}) = \frac{1}{1 + \exp(-(w_0x_0 + w_1x_1 + w_2))}$$





A more involved scalar example

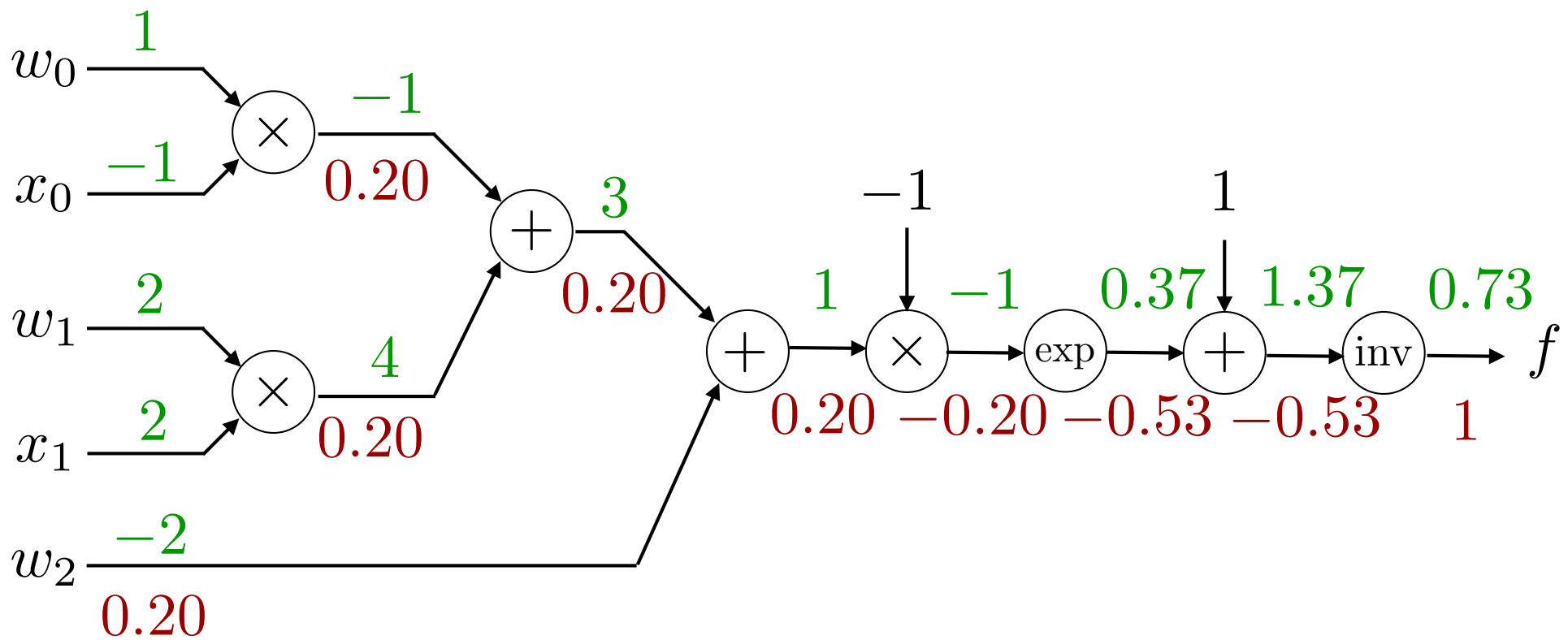
$$f(\mathbf{w}, \mathbf{x}) = \frac{1}{1 + \exp(-(w_0x_0 + w_1x_1 + w_2))}$$





A more involved scalar example

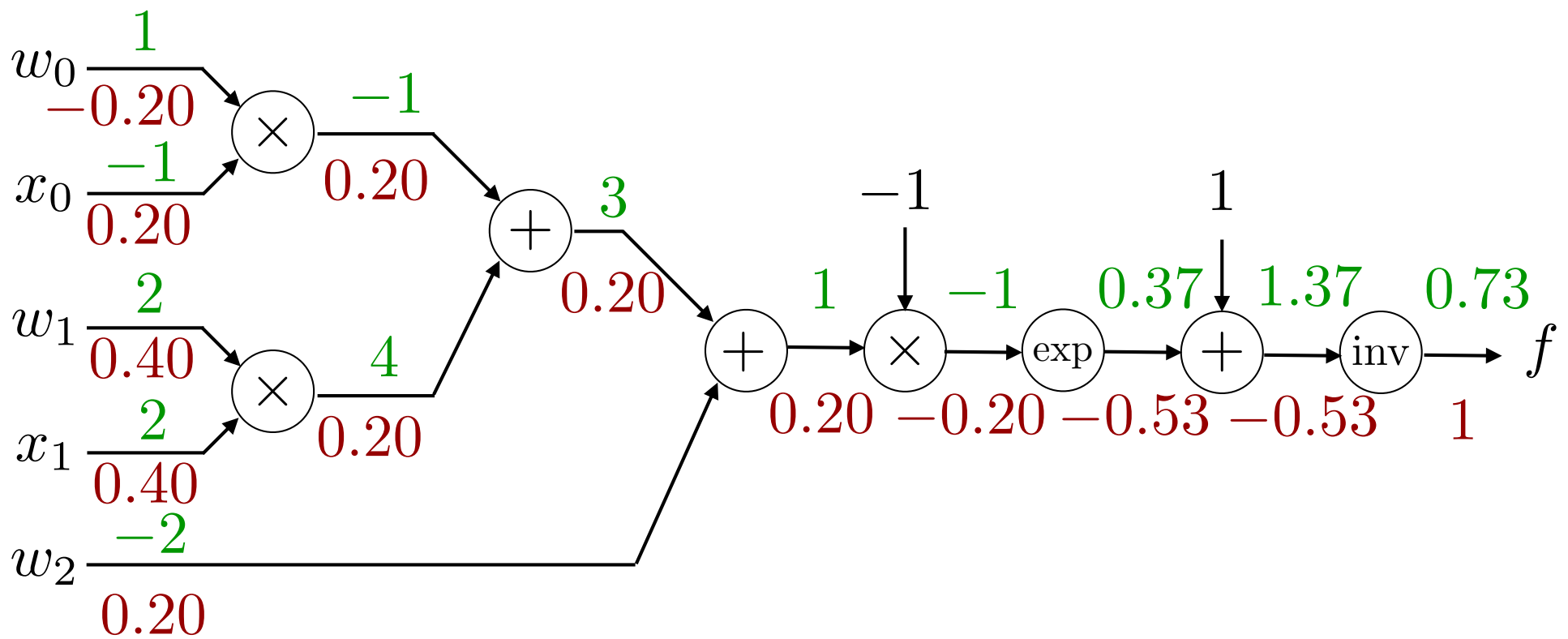
$$f(\mathbf{w}, \mathbf{x}) = \frac{1}{1 + \exp(-(w_0x_0 + w_1x_1 + w_2))}$$





A more involved scalar example

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





You can take any gradient this way

With backpropagation, as long as you can **break the computation into components where you know the local gradients**, you can take the gradient of anything.



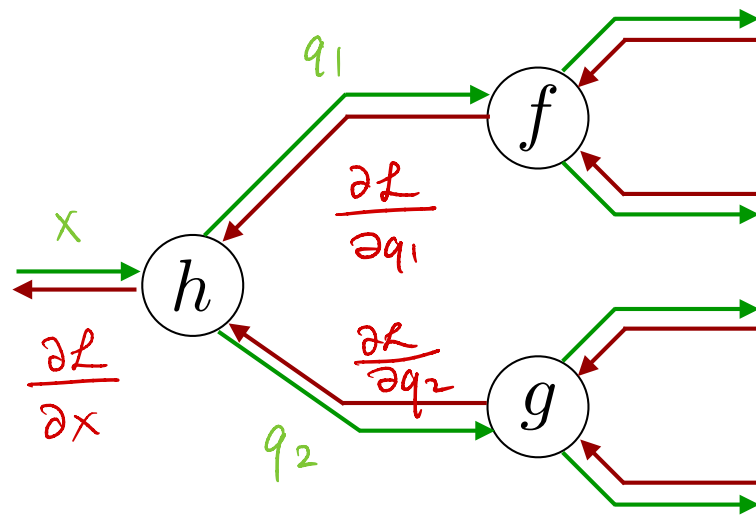
What happens when two gradient paths converge?

$$q_1 = h(x)$$

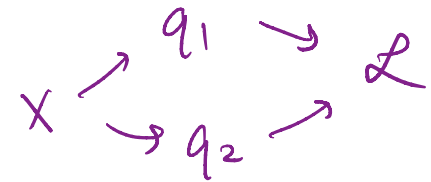
$$q_2 = h(x)$$

$$\frac{\partial \mathcal{L}}{\partial x} = \sum_{i=1}^n \frac{\partial q_i}{\partial x} \frac{\partial \mathcal{L}}{\partial q_i}$$

Law of
total derivatives



$$h(x) = x$$



$$\frac{\partial \mathcal{L}}{\partial x} = \frac{\partial \mathcal{L}}{\partial q_1} + \frac{\partial \mathcal{L}}{\partial q_2}$$