

Lecture 6:

Activation Functions & Normalization Layers

Administrative: Assignment 1

Due today! 11:59pm

- K-Nearest Neighbor
- Linear classifiers: SVM, Softmax

Administrative: Assignment 2

Out soon

- Multi-layer Neural Networks,
- Image Features,
- Optimizers

Administrative: Fridays

This Friday

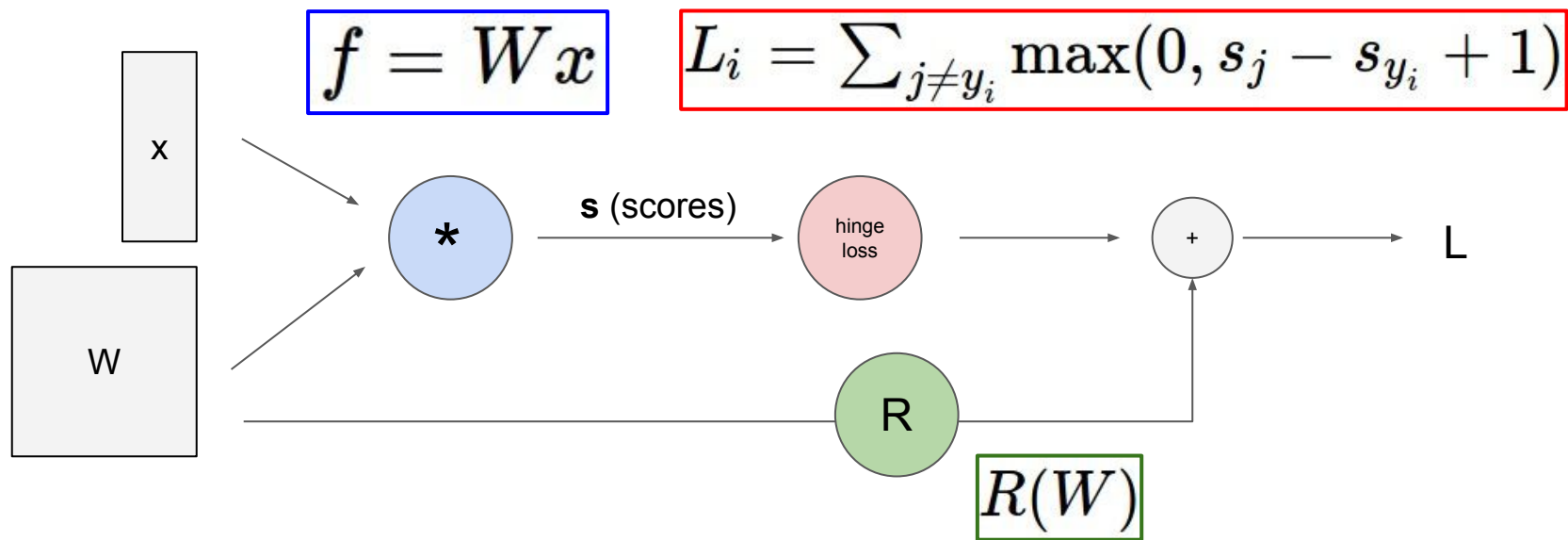
Convolutions & Vectorization

Project

Fill out partner form by eod if you want to!

Where we are now...

Computational graphs



Where we are now...

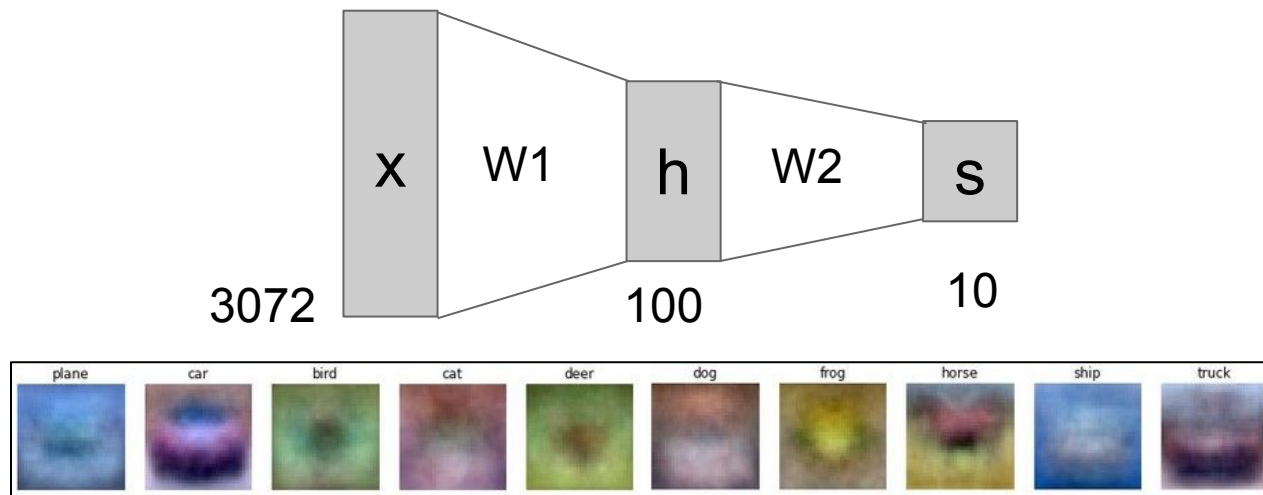
Neural Networks

Linear score function:

$$f = Wx$$

2-layer Neural Network

$$f = W_2 \max(0, W_1 x)$$



Where we are now...

Convolutional Neural Networks

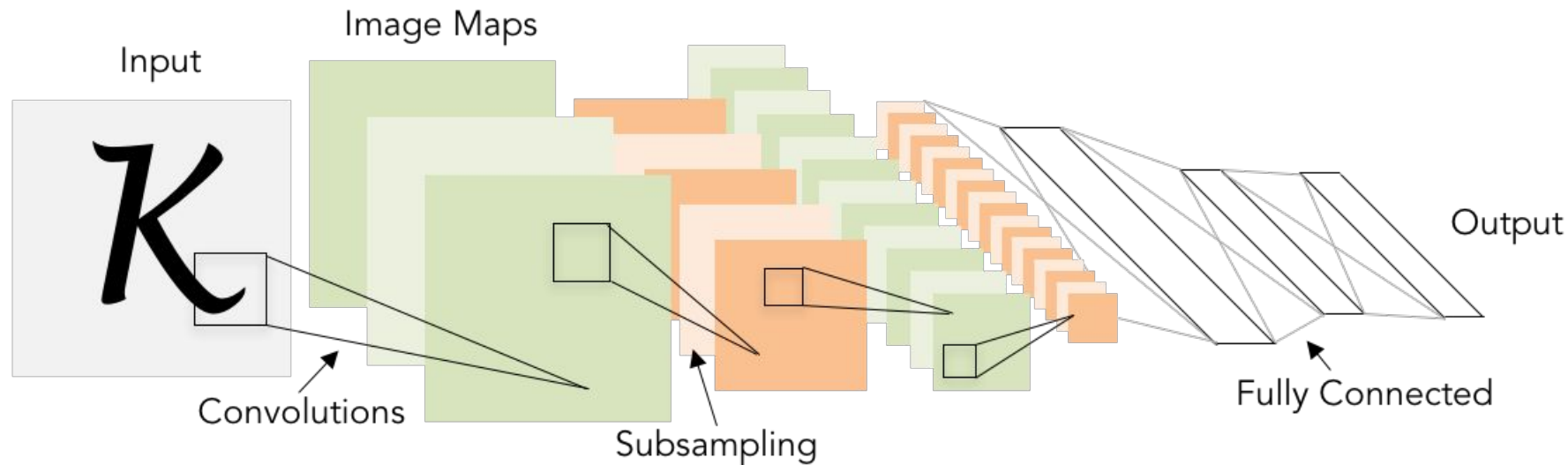
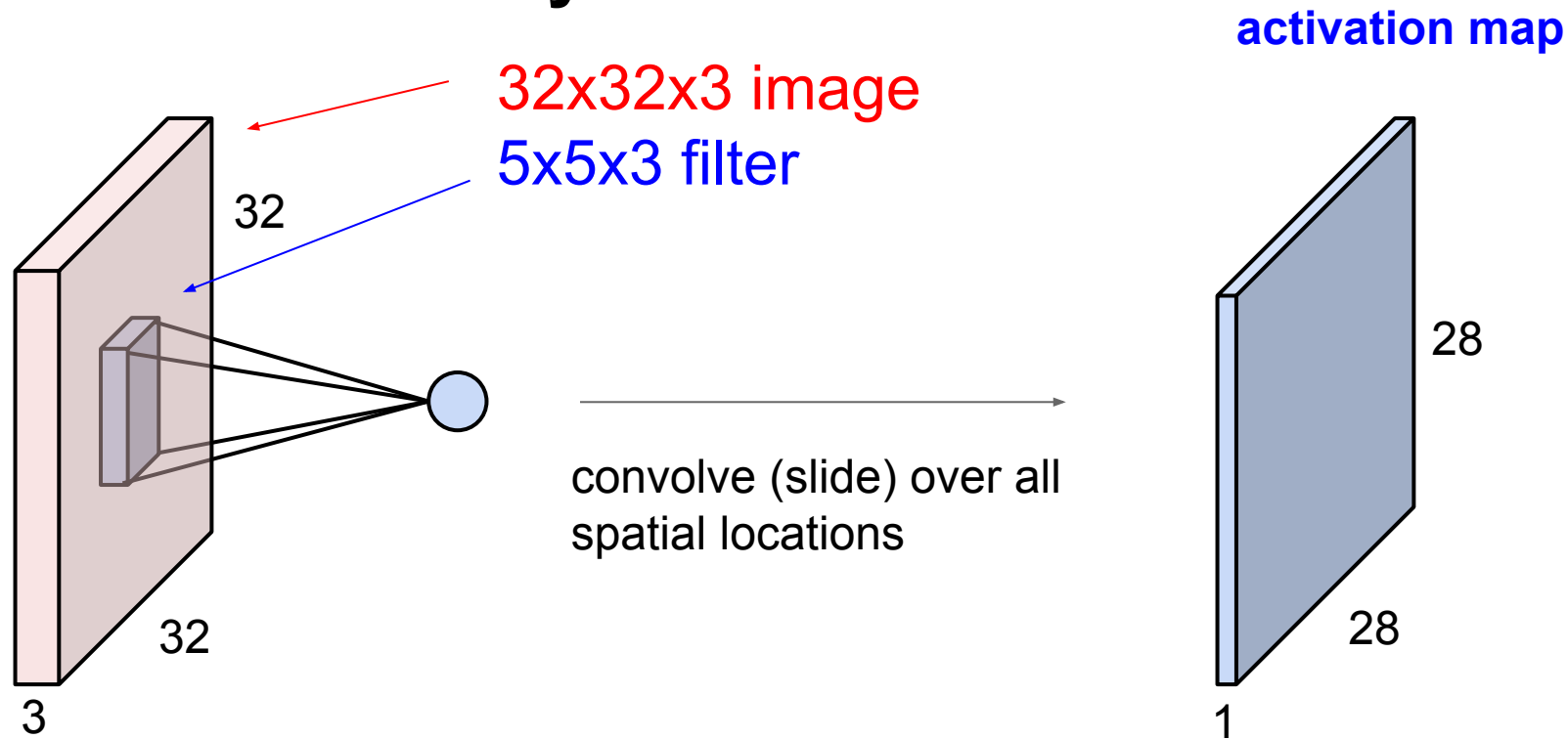
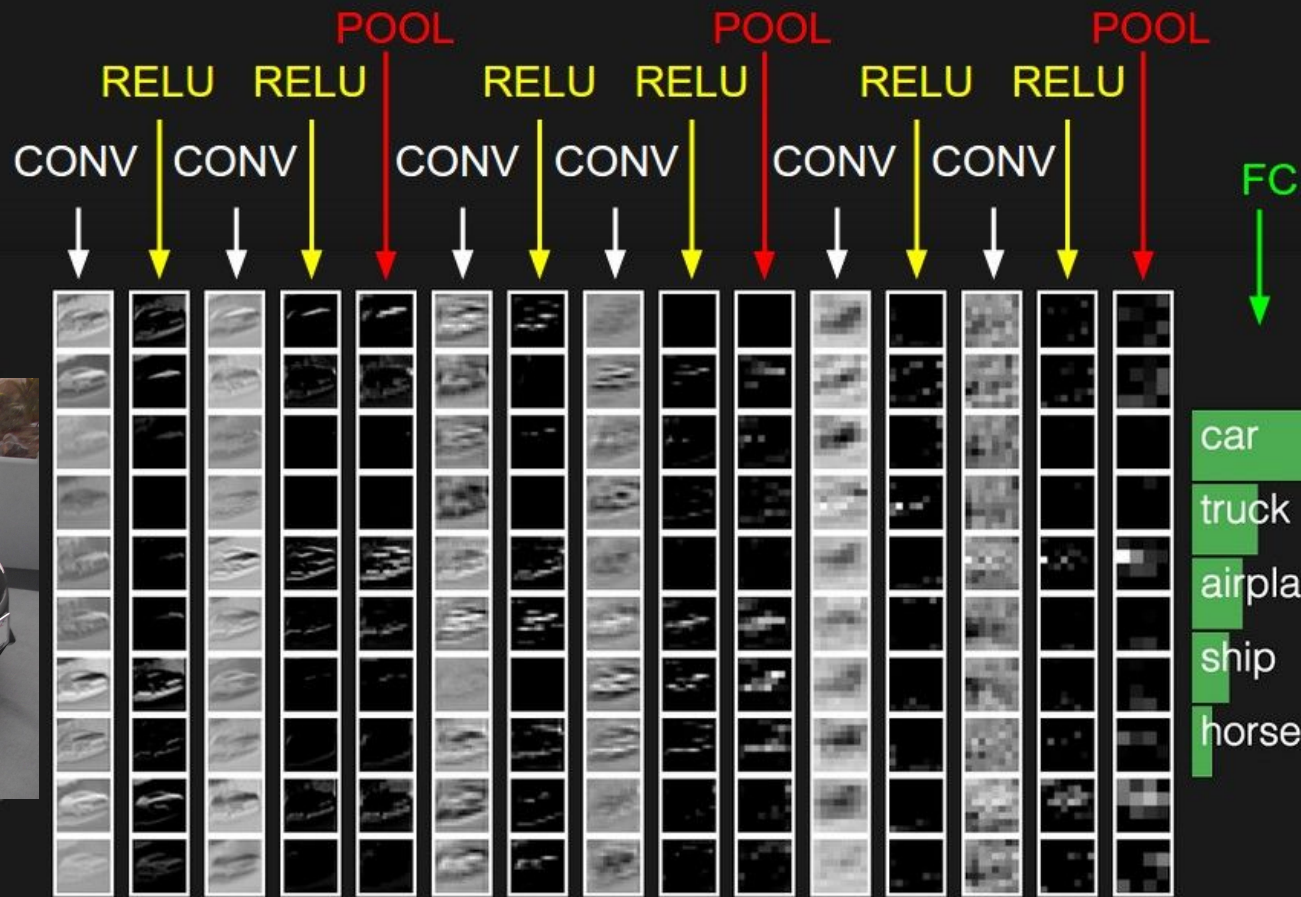


Illustration of LeCun et al. 1998 from CS231n 2017 Lecture 1

Where we are now...

Convolutional Layer





- car
- truck
- airplane
- ship
- horse

Today:

- Activation Function
- Normalization layers

There are a few more steps before we start training

1. **One time setup**

activation functions, preprocessing, weight initialization, regularization, gradient checking

2. **Training dynamics**

babysitting the learning process, parameter updates, hyperparameter optimization

3. **Evaluation**

model ensembles, test-time augmentation, transfer learning

Today's agenda

- Briefly revisit backprop
- Finish CNNs
- Activation Functions
- Data Preprocessing
- Weight Initialization
- Normalization Layers

Backprop with Vectors

Jacobian is **sparse**:
off-diagonal entries
always zero! Never
explicitly form
Jacobian -- instead
use **implicit**
multiplication

4D input x:

$\begin{bmatrix} 1 \\ -2 \\ 3 \\ -1 \end{bmatrix}$

$$f(x) = \max(0, x) \\ (\textit{elementwise})$$

4D output z:

$\begin{bmatrix} 1 \\ 0 \\ 3 \\ 0 \end{bmatrix}$

4D dL/dx:

$\begin{bmatrix} 4 \\ 0 \\ 5 \\ 0 \end{bmatrix}$

$\begin{bmatrix} dz/dx & dL/dz \end{bmatrix}$

$\begin{bmatrix} 1 & 0 & 0 & 0 \end{bmatrix} \begin{bmatrix} 4 \\ -1 \\ 5 \\ 9 \end{bmatrix}$

4D dL/dz:

$\begin{bmatrix} 4 \\ -1 \\ 5 \\ 9 \end{bmatrix}$

Upstream
gradient

Backprop with Vectors

4D input x:

$\begin{bmatrix} 1 \\ -2 \\ 3 \\ -1 \end{bmatrix}$

$$f(x) = \max(0, x) \quad (\text{elementwise})$$

4D output z:

$\begin{bmatrix} 1 \\ 0 \\ 3 \\ 0 \end{bmatrix}$

Jacobian is **sparse**:
off-diagonal entries
always zero! Never
explicitly form
Jacobian -- instead
use **implicit**
multiplication

4D dL/dx :

$\begin{bmatrix} 4 \\ 0 \\ 5 \\ 0 \end{bmatrix}$

$[dz/dx] [dL/dz]$

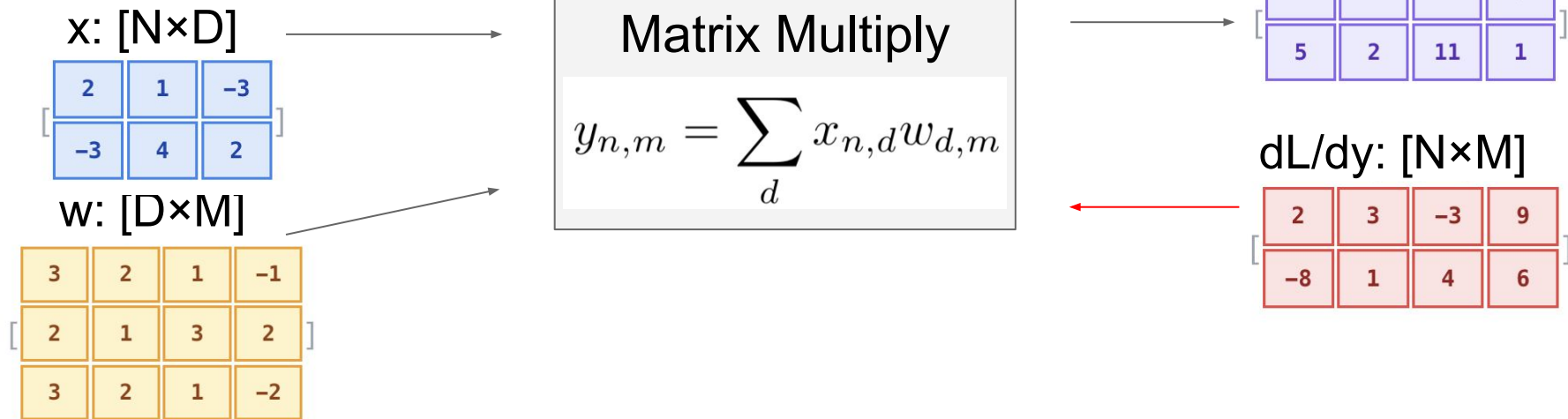
$$\left(\frac{\partial L}{\partial x} \right)_i = \begin{cases} \left(\frac{\partial L}{\partial z} \right)_i & \text{if } x_i > 0 \\ 0 & \text{otherwise} \end{cases}$$

4D dL/dz :

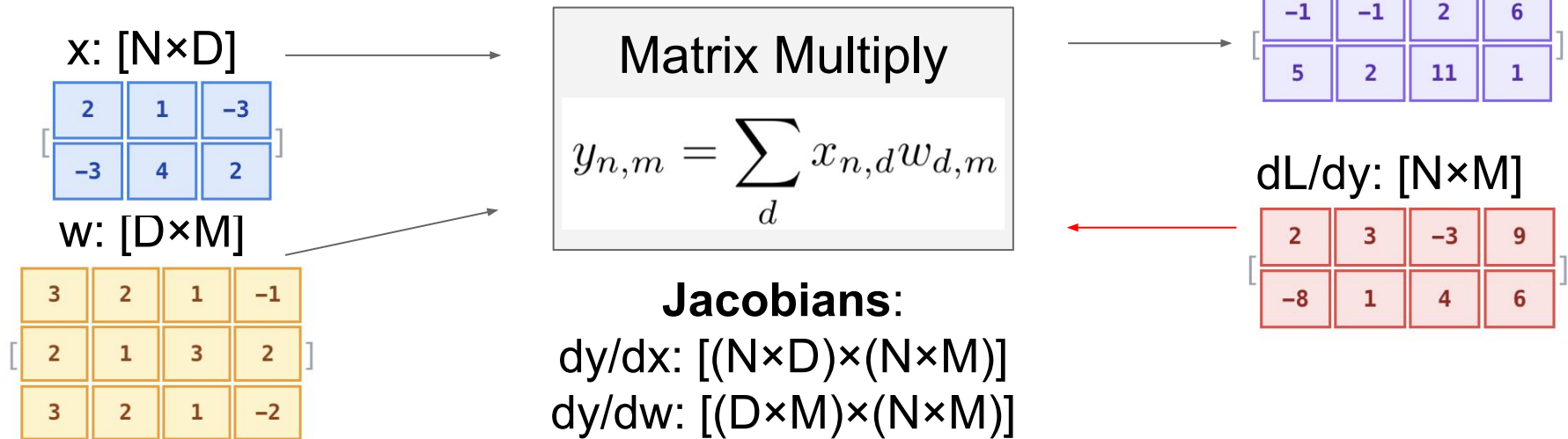
$\begin{bmatrix} 4 \\ -1 \\ 5 \\ 9 \end{bmatrix}$

Upstream
gradient

Backprop with Matrices

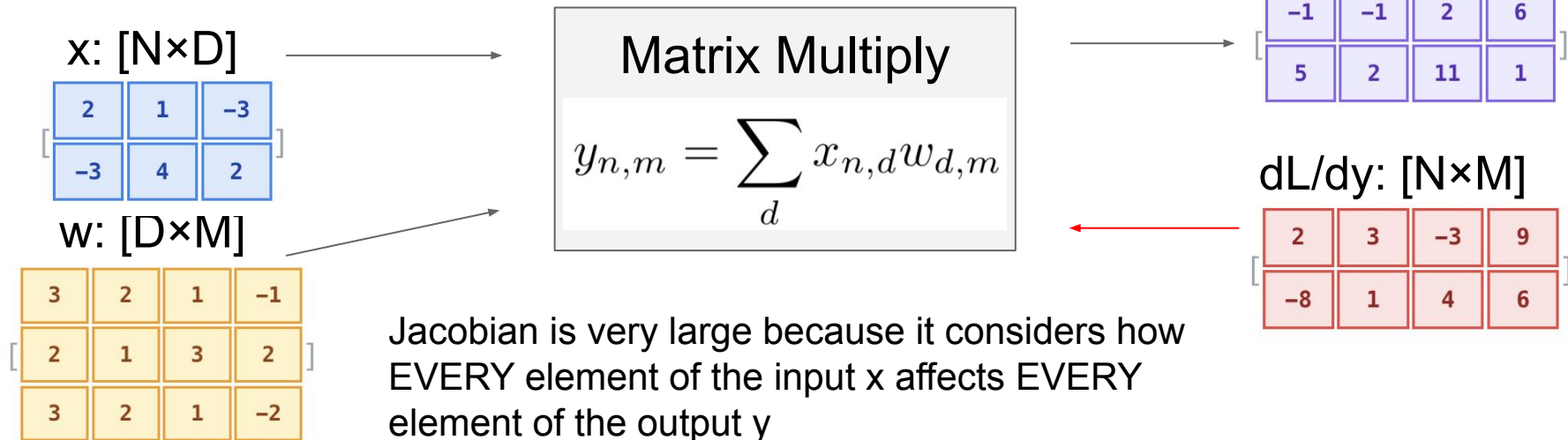


Backprop with Matrices



For a neural net we may have
 $N=64, D=M=4096$
Each Jacobian takes ~256 GB of
memory! Must work with them implicitly!

Backprop with Matrices



Jacobian is very large because it considers how EVERY element of the input x affects EVERY element of the output y

Can we get dL/dx more efficiently, by just looking at what parts of the input ACTUALLY affect which parts of the output?

$$\begin{array}{c} \mathbf{x} [2 \times 3] \\ \begin{bmatrix} 2 & 1 & -3 \\ -1 & 4 & 2 \end{bmatrix} \end{array} \times \begin{array}{c} \mathbf{w} [3 \times 4] \\ \begin{bmatrix} 3 & 2 & 1 & -1 \\ 2 & 1 & 3 & 2 \\ 3 & 2 & 1 & -2 \end{bmatrix} \end{array} = \begin{array}{c} \mathbf{y} [2 \times 4] \\ \begin{bmatrix} -1 & -1 & 2 & 6 \\ 5 & 2 & 11 & 1 \end{bmatrix} \end{array}$$

Diagram illustrating matrix multiplication. A cyan arrow points to the element 1 in the first row, second column of matrix $\mathbf{x}$. Another cyan arrow points to the element 11 in the second row, third column of matrix $\mathbf{y}$.

$$\begin{array}{c} \frac{\partial L}{\partial \mathbf{y}} [2 \times 4] \\ \begin{bmatrix} 2 & 3 & -3 & 9 \\ -8 & 1 & 4 & 6 \end{bmatrix} \end{array}$$

Let's look at how a single element of the input affects a single element of the output!

$$\begin{array}{c} \mathbf{x} [2 \times 3] \\ \begin{bmatrix} 2 & 1 & -3 \\ -1 & 4 & 2 \end{bmatrix} \end{array} \times \begin{array}{c} \mathbf{w} [3 \times 4] \\ \begin{bmatrix} 3 & 2 & 1 & -1 \\ 2 & 1 & 3 & 2 \\ 3 & 2 & 1 & -2 \end{bmatrix} \end{array} = \begin{array}{c} \mathbf{y} [2 \times 4] \\ \begin{bmatrix} -1 & -1 & 2 & 6 \\ 5 & 2 & 11 & 1 \end{bmatrix} \end{array}$$

Diagram illustrating matrix multiplication. A cyan arrow points to the element 1 in the first row, second column of matrix $\mathbf{x}$. Another cyan arrow points to the element 11 in the second row, third column of matrix $\mathbf{y}$.

What is $\frac{\partial y_{0,2}}{\partial x_{0,1}}$?

$$\begin{array}{c} \frac{\partial L}{\partial \mathbf{y}} [2 \times 4] \\ \begin{bmatrix} 2 & 3 & -3 & 9 \\ -8 & 1 & 4 & 6 \end{bmatrix} \end{array}$$

$$\begin{array}{c} \mathbf{x} [2 \times 3] \\ \begin{bmatrix} 2 & 1 & -3 \\ -3 & 4 & 2 \end{bmatrix} \end{array} \times \begin{array}{c} \mathbf{w} [3 \times 4] \\ \begin{bmatrix} 3 & 2 & 1 & -1 \\ 2 & 1 & 3 & 2 \\ 3 & 2 & 1 & -2 \end{bmatrix} \end{array} = \begin{array}{c} \mathbf{y} [2 \times 4] \\ \begin{bmatrix} -1 & -1 & 2 & 6 \\ 5 & 2 & 11 & 1 \end{bmatrix} \end{array}$$

$$\begin{array}{c} \frac{\partial L}{\partial \mathbf{y}} [2 \times 4] \\ \begin{bmatrix} 2 & 3 & -3 & 9 \\ -8 & 1 & 4 & 6 \end{bmatrix} \end{array}$$

What is $\frac{\partial y_{0,2}}{\partial x_{0,1}}$?

Looking at the formula for $y_{0,2}$, what's the partial derivative with respect to $x_{0,1}$?

$$y_{0,2} = x_{0,0} \cdot w_{0,2} + x_{0,1} \cdot w_{1,2} + x_{0,2} \cdot w_{2,2}$$

$$\frac{\partial y_{0,2}}{\partial x_{0,1}} = w_{1,2} = 3$$

$$\begin{array}{c} \mathbf{x} [2 \times 3] \\ \begin{bmatrix} 2 & 1 & -3 \\ -3 & 4 & 2 \end{bmatrix} \end{array} \times \begin{array}{c} \mathbf{w} [3 \times 4] \\ \begin{bmatrix} 3 & 2 & 1 & -1 \\ 2 & 1 & 3 & 2 \\ 3 & 2 & 1 & -2 \end{bmatrix} \end{array} = \begin{array}{c} \mathbf{y} [2 \times 4] \\ \begin{bmatrix} -1 & -1 & 2 & 6 \\ 5 & 2 & 11 & 1 \end{bmatrix} \end{array}$$

$$\begin{array}{c} \frac{\partial L}{\partial \mathbf{y}} [2 \times 4] \\ \begin{bmatrix} 2 & 3 & -3 & 9 \\ -8 & 1 & 4 & 6 \end{bmatrix} \end{array}$$

What is $\frac{\partial y_{0,2}}{\partial x_{0,1}}$?

What is $\frac{\partial y}{\partial x_{0,1}}$?

Looking at the formula for $y_{0,2}$, what's the partial derivative with respect to $x_{0,1}$?

$$y_{0,2} = x_{0,0} \cdot w_{0,2} + x_{0,1} \cdot w_{1,2} + x_{0,2} \cdot w_{2,2}$$

$$\frac{\partial y_{0,2}}{\partial x_{0,1}} = w_{1,2} = 3$$

$$\begin{array}{c} \mathbf{x} [2 \times 3] \\ \begin{bmatrix} 2 & 1 & -3 \\ -3 & 4 & 2 \end{bmatrix} \end{array} \times \begin{array}{c} \mathbf{w} [3 \times 4] \\ \begin{bmatrix} 3 & 2 & 1 & -1 \\ 2 & 1 & 3 & 2 \\ 3 & 2 & 1 & -2 \end{bmatrix} \end{array} = \begin{array}{c} \mathbf{y} \\ \begin{bmatrix} x_{00}w_{00} + x_{01}w_{10} + x_{02}w_{20} & x_{00}w_{01} + x_{01}w_{11} + x_{02}w_{21} & x_{00}w_{02} + x_{01}w_{12} + x_{02}w_{22} & x_{00}w_{03} + x_{01}w_{13} + x_{02}w_{23} \\ x_{10}w_{00} + x_{11}w_{10} + x_{12}w_{20} & x_{10}w_{01} + x_{11}w_{11} + x_{12}w_{21} & x_{10}w_{02} + x_{11}w_{12} + x_{12}w_{22} & x_{10}w_{03} + x_{11}w_{13} + x_{12}w_{23} \end{bmatrix} \end{array}$$

$$\begin{array}{c} \frac{\partial L}{\partial y} [2 \times 4] \\ \begin{bmatrix} 2 & 3 & -3 & 9 \\ -8 & 1 & 4 & 6 \end{bmatrix} \end{array}$$

What is $\frac{\partial y_{0,2}}{\partial x_{0,1}}$?

What is $\frac{\partial y}{\partial x_{0,1}}$?

Looking at the formula for $y_{0,2}$, what's the partial derivative with respect to $x_{0,1}$?

$$y_{0,2} = x_{0,0} \cdot w_{0,2} + x_{0,1} \cdot w_{1,2} + x_{0,2} \cdot w_{2,2}$$

$$\frac{\partial y_{0,2}}{\partial x_{0,1}} = w_{1,2} = 3$$

$$\begin{array}{c} \mathbf{x} [2 \times 3] \\ \begin{bmatrix} 2 & 1 & -3 \\ -3 & 4 & 2 \end{bmatrix} \end{array} \times \begin{array}{c} \mathbf{w} [3 \times 4] \\ \begin{bmatrix} 3 & 2 & 1 & -1 \\ 2 & 1 & 3 & 2 \\ 3 & 2 & 1 & -2 \end{bmatrix} \end{array} = \begin{array}{c} \mathbf{y} \\ \begin{bmatrix} x_{00}w_{00} + x_{01}w_{10} + x_{02}w_{20} & x_{00}w_{01} + x_{01}w_{11} + x_{02}w_{21} & x_{00}w_{02} + x_{01}w_{12} + x_{02}w_{22} & x_{00}w_{03} + x_{01}w_{13} + x_{02}w_{23} \\ x_{10}w_{00} + x_{11}w_{10} + x_{12}w_{20} & x_{10}w_{01} + x_{11}w_{11} + x_{12}w_{21} & x_{10}w_{02} + x_{11}w_{12} + x_{12}w_{22} & x_{10}w_{03} + x_{11}w_{13} + x_{12}w_{23} \end{bmatrix} \end{array}
 \end{array}$$

$$\begin{array}{c} \frac{\partial L}{\partial \mathbf{y}} [2 \times 4] \\ \begin{bmatrix} 2 & 3 & -3 & 9 \\ -8 & 1 & 4 & 6 \end{bmatrix} \end{array}$$

What is $\frac{\partial y_{0,2}}{\partial x_{0,1}}$?

What is $\frac{\partial y}{\partial x_{0,1}}$?

Looking at the formula for $y_{0,2}$, what's the partial derivative with respect to $x_{0,1}$?

$$y_{0,2} = x_{0,0} \cdot w_{0,2} + x_{0,1} \cdot w_{1,2} + x_{0,2} \cdot w_{2,2}$$

$$\frac{\partial y_{0,2}}{\partial x_{0,1}} = w_{1,2} = 3$$

There are multiple local gradients since $x_{0,1}$ affects multiple outputs:

$$\frac{\partial y_{0,0}}{\partial x_{0,1}} = w_{1,0} = 2$$

$$\frac{\partial y_{0,1}}{\partial x_{0,1}} = w_{1,1} = 1$$

$$\frac{\partial y_{0,2}}{\partial x_{0,1}} = w_{1,2} = 3$$

$$\frac{\partial y_{0,3}}{\partial x_{0,1}} = w_{1,3} = 2$$

$$\frac{\partial y_{1,m}}{\partial x_{0,1}} = 0 \text{ (} x_{0,1} \text{ doesn't affect row 1)}$$

$$\begin{array}{c} \mathbf{x} [2 \times 3] \\ \begin{bmatrix} 2 & 1 & -3 \\ -3 & 4 & 2 \end{bmatrix} \end{array} \times \begin{array}{c} \mathbf{w} [3 \times 4] \\ \begin{bmatrix} 3 & 2 & 1 & -1 \\ 2 & 1 & 3 & 2 \\ 3 & 2 & 1 & -2 \end{bmatrix} \end{array} = \begin{array}{c} \mathbf{y} \\ \begin{bmatrix} x_{00}w_{00} + x_{01}w_{10} + x_{02}w_{20} & x_{00}w_{01} + x_{01}w_{11} + x_{02}w_{21} & x_{00}w_{02} + x_{01}w_{12} + x_{02}w_{22} & x_{00}w_{03} + x_{01}w_{13} + x_{02}w_{23} \\ x_{10}w_{00} + x_{11}w_{10} + x_{12}w_{20} & x_{10}w_{01} + x_{11}w_{11} + x_{12}w_{21} & x_{10}w_{02} + x_{11}w_{12} + x_{12}w_{22} & x_{10}w_{03} + x_{11}w_{13} + x_{12}w_{23} \end{bmatrix} \end{array}
 \end{array}$$

$$\begin{array}{c} \frac{\partial L}{\partial \mathbf{y}} [2 \times 4] \\ \begin{bmatrix} 2 & 3 & -3 & 9 \\ -8 & 1 & 4 & 6 \end{bmatrix} \end{array}$$

What is $\frac{\partial y_{0,2}}{\partial x_{0,1}}$?

What is $\frac{\partial y}{\partial x_{0,1}}$?

What is $\frac{\partial L}{\partial x_{0,1}}$?

Looking at the formula for $y_{0,2}$, what's the partial derivative with respect to $x_{0,1}$?

$$y_{0,2} = x_{0,0} \cdot w_{0,2} + x_{0,1} \cdot w_{1,2} + x_{0,2} \cdot w_{2,2}$$

$$\frac{\partial y_{0,2}}{\partial x_{0,1}} = w_{1,2} = 3$$

There are multiple local gradients since $x_{0,1}$ affects multiple outputs:

$$\frac{\partial y_{0,0}}{\partial x_{0,1}} = w_{1,0} = 2$$

$$\frac{\partial y_{0,1}}{\partial x_{0,1}} = w_{1,1} = 1$$

$$\frac{\partial y_{0,2}}{\partial x_{0,1}} = w_{1,2} = 3$$

$$\frac{\partial y_{0,3}}{\partial x_{0,1}} = w_{1,3} = 2$$

$$\frac{\partial y_{1,m}}{\partial x_{0,1}} = 0 \text{ (} x_{0,1} \text{ doesn't affect row 1)}$$

$$\begin{array}{c} \mathbf{x} [2 \times 3] \\ \begin{bmatrix} 2 & 1 & -3 \\ -3 & 4 & 2 \end{bmatrix} \end{array} \times \begin{array}{c} \mathbf{w} [3 \times 4] \\ \begin{bmatrix} 3 & 2 & 1 & -1 \\ 2 & 1 & 3 & 2 \\ 3 & 2 & 1 & -2 \end{bmatrix} \end{array} = \begin{array}{c} \mathbf{y} \\ \begin{bmatrix} x_{00}w_{00} + x_{01}w_{10} + x_{02}w_{20} & x_{00}w_{01} + x_{01}w_{11} + x_{02}w_{21} & x_{00}w_{02} + x_{01}w_{12} + x_{02}w_{22} & x_{00}w_{03} + x_{01}w_{13} + x_{02}w_{23} \\ x_{10}w_{00} + x_{11}w_{10} + x_{12}w_{20} & x_{10}w_{01} + x_{11}w_{11} + x_{12}w_{21} & x_{10}w_{02} + x_{11}w_{12} + x_{12}w_{22} & x_{10}w_{03} + x_{11}w_{13} + x_{12}w_{23} \end{bmatrix} \end{array} \begin{array}{c} \frac{\partial L}{\partial \mathbf{y}} [2 \times 4] \\ \begin{bmatrix} 2 & 3 & -3 & 9 \\ -8 & 1 & 4 & 6 \end{bmatrix} \end{array}$$

What is $\frac{\partial y_{0,2}}{\partial x_{0,1}}$?

What is $\frac{\partial y}{\partial x_{0,1}}$?

What is $\frac{\partial L}{\partial x_{0,1}}$?

Looking at the formula for $y_{0,2}$, what's the partial derivative with respect to $x_{0,1}$?

$$y_{0,2} = x_{0,0} \cdot w_{0,2} + x_{0,1} \cdot w_{1,2} + x_{0,2} \cdot w_{2,2}$$

$$\frac{\partial y_{0,2}}{\partial x_{0,1}} = w_{1,2} = 3$$

There are multiple local gradients since $x_{0,1}$ affects multiple outputs:

$$\frac{\partial y_{0,0}}{\partial x_{0,1}} = w_{1,0} = 2$$

$$\frac{\partial y_{0,1}}{\partial x_{0,1}} = w_{1,1} = 1$$

$$\frac{\partial y_{0,2}}{\partial x_{0,1}} = w_{1,2} = 3$$

$$\frac{\partial y_{0,3}}{\partial x_{0,1}} = w_{1,3} = 2$$

$$\frac{\partial y_{1,m}}{\partial x_{0,1}} = 0 \text{ (} x_{0,1} \text{ doesn't affect row 1)}$$

Sum of: (how much $x_{0,1}$ affects each y) $\times$ (how much each y affects L)

$$\begin{array}{c} \mathbf{x} [2 \times 3] \\ \begin{bmatrix} 2 & 1 & -3 \\ -3 & 4 & 2 \end{bmatrix} \end{array} \times \begin{array}{c} \mathbf{w} [3 \times 4] \\ \begin{bmatrix} 3 & 2 & 1 & -1 \\ 2 & 1 & 3 & 2 \\ 3 & 2 & 1 & -2 \end{bmatrix} \end{array} = \begin{array}{c} \mathbf{y} \\ \begin{bmatrix} x_{00}w_{00} + x_{01}w_{10} + x_{02}w_{20} & x_{00}w_{01} + x_{01}w_{11} + x_{02}w_{21} & x_{00}w_{02} + x_{01}w_{12} + x_{02}w_{22} & x_{00}w_{03} + x_{01}w_{13} + x_{02}w_{23} \\ x_{10}w_{00} + x_{11}w_{10} + x_{12}w_{20} & x_{10}w_{01} + x_{11}w_{11} + x_{12}w_{21} & x_{10}w_{02} + x_{11}w_{12} + x_{12}w_{22} & x_{10}w_{03} + x_{11}w_{13} + x_{12}w_{23} \end{bmatrix} \end{array} \begin{array}{c} \frac{\partial L}{\partial \mathbf{y}} [2 \times 4] \\ \begin{bmatrix} 2 & 3 & -3 & 9 \\ -8 & 1 & 4 & 6 \end{bmatrix} \end{array}$$

What is $\frac{\partial y_{0,2}}{\partial x_{0,1}}$?

What is $\frac{\partial y}{\partial x_{0,1}}$?

What is $\frac{\partial L}{\partial x_{0,1}}$?

Looking at the formula for $y_{0,2}$, what's the partial derivative with respect to $x_{0,1}$?

$$y_{0,2} = x_{0,0} \cdot w_{0,2} + x_{0,1} \cdot w_{1,2} + x_{0,2} \cdot w_{2,2}$$

$$\frac{\partial y_{0,2}}{\partial x_{0,1}} = w_{1,2} = 3$$

There are multiple local gradients since $x_{0,1}$ affects multiple outputs:

$$\frac{\partial y_{0,0}}{\partial x_{0,1}} = w_{1,0} = 2$$

$$\frac{\partial y_{0,1}}{\partial x_{0,1}} = w_{1,1} = 1$$

$$\frac{\partial y_{0,2}}{\partial x_{0,1}} = w_{1,2} = 3$$

$$\frac{\partial y_{0,3}}{\partial x_{0,1}} = w_{1,3} = 2$$

$$\frac{\partial y_{1,m}}{\partial x_{0,1}} = 0 \text{ (} x_{0,1} \text{ doesn't affect row 1)}$$

Sum of: (how much $x_{0,1}$ affects each y) $\times$ (how much each y affects L)

$$\begin{aligned} &= \frac{\partial y_{0,0}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,0}} + \frac{\partial y_{0,1}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,1}} + \frac{\partial y_{0,2}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,2}} + \frac{\partial y_{0,3}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,3}} \\ &+ 0 \cdot \frac{\partial L}{\partial y_{1,0}} + 0 \cdot \frac{\partial L}{\partial y_{1,1}} + 0 \cdot \frac{\partial L}{\partial y_{1,2}} + 0 \cdot \frac{\partial L}{\partial y_{1,3}} \end{aligned}$$

$$\begin{array}{c} \mathbf{x} [2 \times 3] \\ \begin{bmatrix} 2 & 1 & -3 \\ -3 & 4 & 2 \end{bmatrix} \end{array} \times \begin{array}{c} \mathbf{w} [3 \times 4] \\ \begin{bmatrix} 3 & 2 & 1 & -1 \\ 2 & 1 & 3 & 2 \\ 3 & 2 & 1 & -2 \end{bmatrix} \end{array} = \begin{array}{c} \mathbf{y} \\ \begin{bmatrix} x_{00}w_{00} + x_{01}w_{10} + x_{02}w_{20} & x_{00}w_{01} + x_{01}w_{11} + x_{02}w_{21} & x_{00}w_{02} + x_{01}w_{12} + x_{02}w_{22} & x_{00}w_{03} + x_{01}w_{13} + x_{02}w_{23} \\ x_{10}w_{00} + x_{11}w_{10} + x_{12}w_{20} & x_{10}w_{01} + x_{11}w_{11} + x_{12}w_{21} & x_{10}w_{02} + x_{11}w_{12} + x_{12}w_{22} & x_{10}w_{03} + x_{11}w_{13} + x_{12}w_{23} \end{bmatrix} \end{array} \begin{array}{c} \frac{\partial L}{\partial \mathbf{y}} [2 \times 4] \\ \begin{bmatrix} 2 & 3 & -3 & 9 \\ -8 & 1 & 4 & 6 \end{bmatrix} \end{array}$$

What is $\frac{\partial y_{0,2}}{\partial x_{0,1}}$?

What is $\frac{\partial y}{\partial x_{0,1}}$?

What is $\frac{\partial L}{\partial x_{0,1}}$?

Looking at the formula for $y_{0,2}$, what's the partial derivative with respect to $x_{0,1}$?

$$y_{0,2} = x_{0,0} \cdot w_{0,2} + x_{0,1} \cdot w_{1,2} + x_{0,2} \cdot w_{2,2}$$

$$\frac{\partial y_{0,2}}{\partial x_{0,1}} = w_{1,2} = 3$$

There are multiple local gradients since $x_{0,1}$ affects multiple outputs:

$$\frac{\partial y_{0,0}}{\partial x_{0,1}} = w_{1,0} = 2$$

$$\frac{\partial y_{0,1}}{\partial x_{0,1}} = w_{1,1} = 1$$

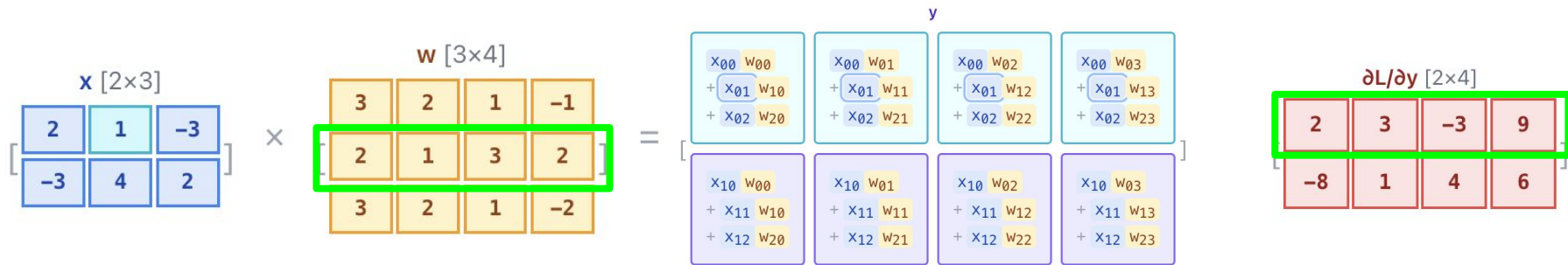
$$\frac{\partial y_{0,2}}{\partial x_{0,1}} = w_{1,2} = 3$$

$$\frac{\partial y_{0,3}}{\partial x_{0,1}} = w_{1,3} = 2$$

$$\frac{\partial y_{1,m}}{\partial x_{0,1}} = 0 \text{ (} x_{0,1} \text{ doesn't affect row 1)}$$

Sum of: (how much $x_{0,1}$ affects each y) $\times$ (how much each y affects L)

$$\begin{aligned} &= \frac{\partial y_{0,0}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,0}} + \frac{\partial y_{0,1}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,1}} + \frac{\partial y_{0,2}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,2}} + \frac{\partial y_{0,3}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,3}} \\ &+ 0 \cdot \frac{\partial L}{\partial y_{1,0}} + 0 \cdot \frac{\partial L}{\partial y_{1,1}} + 0 \cdot \frac{\partial L}{\partial y_{1,2}} + 0 \cdot \frac{\partial L}{\partial y_{1,3}} \\ &= w_{1,0} \cdot \frac{\partial L}{\partial y_{0,0}} + w_{1,1} \cdot \frac{\partial L}{\partial y_{0,1}} + w_{1,2} \cdot \frac{\partial L}{\partial y_{0,2}} + w_{1,3} \cdot \frac{\partial L}{\partial y_{0,3}} \end{aligned}$$



What is $\frac{\partial y_{0,2}}{\partial x_{0,1}}$?

Looking at the formula for $y_{0,2}$, what's the partial derivative with respect to $x_{0,1}$?

$$y_{0,2} = x_{0,0} \cdot w_{0,2} + x_{0,1} \cdot w_{1,2} + x_{0,2} \cdot w_{2,2}$$

$$\frac{\partial y_{0,2}}{\partial x_{0,1}} = w_{1,2} = 3$$

What is $\frac{\partial y}{\partial x_{0,1}}$?

There are multiple local gradients since $x_{0,1}$ affects multiple outputs:

$$\frac{\partial y_{0,0}}{\partial x_{0,1}} = w_{1,0} = 2$$

$$\frac{\partial y_{0,1}}{\partial x_{0,1}} = w_{1,1} = 1$$

$$\frac{\partial y_{0,2}}{\partial x_{0,1}} = w_{1,2} = 3$$

$$\frac{\partial y_{0,3}}{\partial x_{0,1}} = w_{1,3} = 2$$

$$\frac{\partial y_{1,m}}{\partial x_{0,1}} = 0 \text{ (} x_{0,1} \text{ doesn't affect row 1)}$$

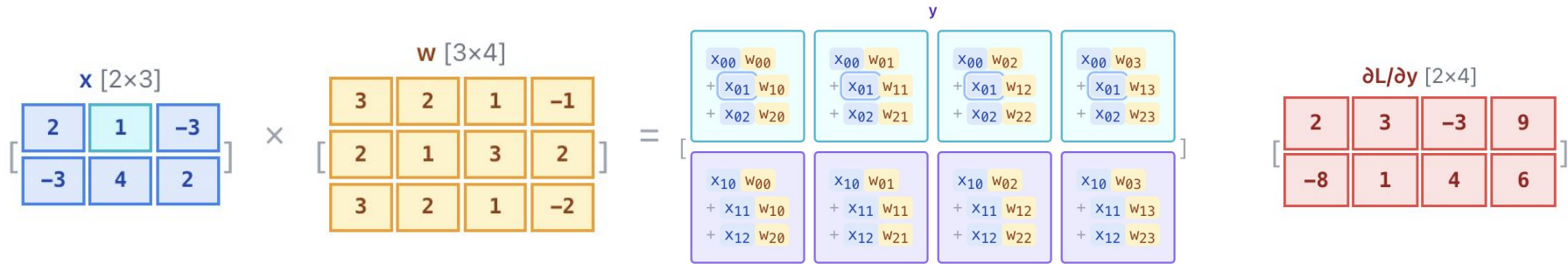
What is $\frac{\partial L}{\partial x_{0,1}}$?

Sum of: (how much $x_{0,1}$ affects each y) $\times$ (how much each y affects L)

$$= \frac{\partial y_{0,0}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,0}} + \frac{\partial y_{0,1}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,1}} + \frac{\partial y_{0,2}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,2}} + \frac{\partial y_{0,3}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,3}} + 0 \cdot \frac{\partial L}{\partial y_{1,0}} + 0 \cdot \frac{\partial L}{\partial y_{1,1}} + 0 \cdot \frac{\partial L}{\partial y_{1,2}} + 0 \cdot \frac{\partial L}{\partial y_{1,3}}$$

$$= w_{1,0} \cdot \frac{\partial L}{\partial y_{0,0}} + w_{1,1} \cdot \frac{\partial L}{\partial y_{0,1}} + w_{1,2} \cdot \frac{\partial L}{\partial y_{0,2}} + w_{1,3} \cdot \frac{\partial L}{\partial y_{0,3}}$$

$$= w[\text{row } 1] \cdot \frac{\partial L}{\partial y}[\text{row } 0]$$



What is $\frac{\partial L}{\partial x_{0,1}}$?

$$\frac{\partial L}{\partial x_{0,1}} = \frac{\partial L}{\partial y} [\text{row } 0] \cdot w[\text{row } 1]$$

Sum of: (how much $x_{0,1}$ affects each y) $\times$ (how much each y affects L)

$$\begin{aligned} &= \frac{\partial y_{0,0}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,0}} + \frac{\partial y_{0,1}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,1}} + \frac{\partial y_{0,2}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,2}} + \frac{\partial y_{0,3}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,3}} \\ &+ 0 \cdot \frac{\partial L}{\partial y_{1,0}} + 0 \cdot \frac{\partial L}{\partial y_{1,1}} + 0 \cdot \frac{\partial L}{\partial y_{1,2}} + 0 \cdot \frac{\partial L}{\partial y_{1,3}} \\ &= w_{1,0} \cdot \frac{\partial L}{\partial y_{0,0}} + w_{1,1} \cdot \frac{\partial L}{\partial y_{0,1}} + w_{1,2} \cdot \frac{\partial L}{\partial y_{0,2}} + w_{1,3} \cdot \frac{\partial L}{\partial y_{0,3}} \\ &= w[\text{row } 1] \cdot \frac{\partial L}{\partial y} [\text{row } 0] \end{aligned}$$

$$\begin{array}{c} \mathbf{x} [2 \times 3] \\ \begin{bmatrix} 2 & 1 & -3 \\ -3 & 4 & 2 \end{bmatrix} \end{array} \times \begin{array}{c} \mathbf{w} [3 \times 4] \\ \begin{bmatrix} 3 & 2 & 1 & -1 \\ 2 & 1 & 3 & 2 \\ 3 & 2 & 1 & -2 \end{bmatrix} \end{array} = \begin{array}{c} \mathbf{y} \\ \begin{bmatrix} x_{00}w_{00} + x_{01}w_{10} + x_{02}w_{20} & x_{00}w_{01} + x_{01}w_{11} + x_{02}w_{21} & x_{00}w_{02} + x_{01}w_{12} + x_{02}w_{22} & x_{00}w_{03} + x_{01}w_{13} + x_{02}w_{23} \\ x_{10}w_{00} + x_{11}w_{10} + x_{12}w_{20} & x_{10}w_{01} + x_{11}w_{11} + x_{12}w_{21} & x_{10}w_{02} + x_{11}w_{12} + x_{12}w_{22} & x_{10}w_{03} + x_{11}w_{13} + x_{12}w_{23} \end{bmatrix} \end{array}
 \end{array}$$

$$\begin{array}{c} \frac{\partial L}{\partial \mathbf{y}} [2 \times 4] \\ \begin{bmatrix} 2 & 3 & -3 & 9 \\ -8 & 1 & 4 & 6 \end{bmatrix} \end{array}$$

What is $\frac{\partial L}{\partial x_{0,1}}$?

$$\frac{\partial L}{\partial x_{0,1}} = \frac{\partial L}{\partial y} [\text{row } 0] \cdot w[\text{row } 1]$$

$$\frac{\partial L}{\partial x_{n,d}} = \frac{\partial L}{\partial y} [\text{row } n] \cdot w[\text{row } d]$$

Sum of: (how much $x_{0,1}$ affects each y) $\times$ (how much each y affects L)

$$\begin{aligned}
 &= \frac{\partial y_{0,0}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,0}} + \frac{\partial y_{0,1}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,1}} + \frac{\partial y_{0,2}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,2}} + \frac{\partial y_{0,3}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,3}} \\
 &+ 0 \cdot \frac{\partial L}{\partial y_{1,0}} + 0 \cdot \frac{\partial L}{\partial y_{1,1}} + 0 \cdot \frac{\partial L}{\partial y_{1,2}} + 0 \cdot \frac{\partial L}{\partial y_{1,3}} \\
 &= w_{1,0} \cdot \frac{\partial L}{\partial y_{0,0}} + w_{1,1} \cdot \frac{\partial L}{\partial y_{0,1}} + w_{1,2} \cdot \frac{\partial L}{\partial y_{0,2}} + w_{1,3} \cdot \frac{\partial L}{\partial y_{0,3}} \\
 &= w[\text{row } 1] \cdot \frac{\partial L}{\partial y} [\text{row } 0]
 \end{aligned}$$

$$\begin{array}{c} \mathbf{x} [2 \times 3] \\ \begin{bmatrix} 2 & 1 & -3 \\ -3 & 4 & 2 \end{bmatrix} \end{array} \times \begin{array}{c} \mathbf{w} [3 \times 4] \\ \begin{bmatrix} 3 & 2 & 1 & -1 \\ 2 & 1 & 3 & 2 \\ 3 & 2 & 1 & -2 \end{bmatrix} \end{array} = \begin{array}{c} \mathbf{y} \\ \begin{bmatrix} x_{00}w_{00} + x_{01}w_{10} + x_{02}w_{20} & x_{00}w_{01} + x_{01}w_{11} + x_{02}w_{21} & x_{00}w_{02} + x_{01}w_{12} + x_{02}w_{22} & x_{00}w_{03} + x_{01}w_{13} + x_{02}w_{23} \\ x_{10}w_{00} + x_{11}w_{10} + x_{12}w_{20} & x_{10}w_{01} + x_{11}w_{11} + x_{12}w_{21} & x_{10}w_{02} + x_{11}w_{12} + x_{12}w_{22} & x_{10}w_{03} + x_{11}w_{13} + x_{12}w_{23} \end{bmatrix} \end{array}
 \end{array}$$

$$\frac{\partial L}{\partial \mathbf{y}} [2 \times 4] = \begin{bmatrix} 2 & 3 & -3 & 9 \\ -8 & 1 & 4 & 6 \end{bmatrix}$$

What is $\frac{\partial L}{\partial x_{0,1}}$?

$$\frac{\partial L}{\partial x_{0,1}} = \frac{\partial L}{\partial y} [\text{row } 0] \cdot w[\text{row } 1]$$

$$\frac{\partial L}{\partial x_{n,d}} = \frac{\partial L}{\partial y} [\text{row } n] \cdot w[\text{row } d]$$

$$\begin{array}{c} \frac{\partial L}{\partial \mathbf{x}} [2 \times 3] \\ \begin{bmatrix} \frac{\partial L}{\partial y_0} \cdot w[0] & \frac{\partial L}{\partial y_0} \cdot w[1] & \frac{\partial L}{\partial y_0} \cdot w[2] \\ \frac{\partial L}{\partial y_1} \cdot w[0] & \frac{\partial L}{\partial y_1} \cdot w[1] & \frac{\partial L}{\partial y_1} \cdot w[2] \end{bmatrix} \end{array}$$

Sum of: (how much $x_{0,1}$ affects each y) $\times$ (how much each y affects L)

$$\begin{aligned} &= \frac{\partial y_{0,0}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,0}} + \frac{\partial y_{0,1}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,1}} + \frac{\partial y_{0,2}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,2}} + \frac{\partial y_{0,3}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,3}} \\ &+ 0 \cdot \frac{\partial L}{\partial y_{1,0}} + 0 \cdot \frac{\partial L}{\partial y_{1,1}} + 0 \cdot \frac{\partial L}{\partial y_{1,2}} + 0 \cdot \frac{\partial L}{\partial y_{1,3}} \\ &= w_{1,0} \cdot \frac{\partial L}{\partial y_{0,0}} + w_{1,1} \cdot \frac{\partial L}{\partial y_{0,1}} + w_{1,2} \cdot \frac{\partial L}{\partial y_{0,2}} + w_{1,3} \cdot \frac{\partial L}{\partial y_{0,3}} \\ &= w[\text{row } 1] \cdot \frac{\partial L}{\partial y} [\text{row } 0] \end{aligned}$$

$$\begin{array}{c} \mathbf{x} [2 \times 3] \\ \begin{bmatrix} 2 & 1 & -3 \\ -3 & 4 & 2 \end{bmatrix} \end{array} \times \begin{array}{c} \mathbf{w} [3 \times 4] \\ \begin{bmatrix} 3 & 2 & 1 & -1 \\ 2 & 1 & 3 & 2 \\ 3 & 2 & 1 & -2 \end{bmatrix} \end{array} = \begin{array}{c} \mathbf{y} \\ \begin{bmatrix} x_{00}w_{00} + x_{01}w_{10} + x_{02}w_{20} & x_{00}w_{01} + x_{01}w_{11} + x_{02}w_{21} & x_{00}w_{02} + x_{01}w_{12} + x_{02}w_{22} & x_{00}w_{03} + x_{01}w_{13} + x_{02}w_{23} \\ x_{10}w_{00} + x_{11}w_{10} + x_{12}w_{20} & x_{10}w_{01} + x_{11}w_{11} + x_{12}w_{21} & x_{10}w_{02} + x_{11}w_{12} + x_{12}w_{22} & x_{10}w_{03} + x_{11}w_{13} + x_{12}w_{23} \end{bmatrix} \end{array}
 \end{array}$$

$$\begin{array}{c} \frac{\partial L}{\partial \mathbf{y}} [2 \times 4] \\ \begin{bmatrix} 2 & 3 & -3 & 9 \\ -8 & 1 & 4 & 6 \end{bmatrix} \end{array}$$

What is $\frac{\partial L}{\partial x_{0,1}}$?

$$\begin{array}{c} \frac{\partial L}{\partial \mathbf{y}} [2 \times 4] \\ \begin{bmatrix} \frac{\partial L}{\partial y_{0,0}} & \frac{\partial L}{\partial y_{0,1}} & \frac{\partial L}{\partial y_{0,2}} & \frac{\partial L}{\partial y_{0,3}} \\ \frac{\partial L}{\partial y_{1,0}} & \frac{\partial L}{\partial y_{1,1}} & \frac{\partial L}{\partial y_{1,2}} & \frac{\partial L}{\partial y_{1,3}} \end{bmatrix} \end{array} \times \begin{array}{c} \mathbf{w}^T [4 \times 3] \\ \begin{bmatrix} w_{0,0} & w_{1,0} & w_{2,0} \\ w_{0,1} & w_{1,1} & w_{2,1} \\ w_{0,2} & w_{1,2} & w_{2,2} \\ w_{0,3} & w_{1,3} & w_{2,3} \end{bmatrix} \end{array} = \begin{array}{c} \frac{\partial L}{\partial \mathbf{x}} [2 \times 3] \\ \begin{bmatrix} \frac{\partial L}{\partial y_0} \cdot w_{0,0} & \frac{\partial L}{\partial y_0} \cdot w_{1,0} & \frac{\partial L}{\partial y_0} \cdot w_{2,0} \\ \frac{\partial L}{\partial y_1} \cdot w_{0,0} & \frac{\partial L}{\partial y_1} \cdot w_{1,0} & \frac{\partial L}{\partial y_1} \cdot w_{2,0} \end{bmatrix} \end{array}$$

Sum of: (how much $x_{0,1}$ affects each y) $\times$ (how much each y affects L)

$$\begin{aligned} &= \frac{\partial y_{0,0}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,0}} + \frac{\partial y_{0,1}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,1}} + \frac{\partial y_{0,2}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,2}} + \frac{\partial y_{0,3}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,3}} \\ &+ 0 \cdot \frac{\partial L}{\partial y_{1,0}} + 0 \cdot \frac{\partial L}{\partial y_{1,1}} + 0 \cdot \frac{\partial L}{\partial y_{1,2}} + 0 \cdot \frac{\partial L}{\partial y_{1,3}} \\ &= w_{1,0} \cdot \frac{\partial L}{\partial y_{0,0}} + w_{1,1} \cdot \frac{\partial L}{\partial y_{0,1}} + w_{1,2} \cdot \frac{\partial L}{\partial y_{0,2}} + w_{1,3} \cdot \frac{\partial L}{\partial y_{0,3}} \\ &= \mathbf{w}[\text{row } 1] \cdot \frac{\partial L}{\partial \mathbf{y}}[\text{row } 0] \end{aligned}$$

$$\begin{array}{c} \mathbf{x} [2 \times 3] \\ \begin{bmatrix} 2 & 1 & -3 \\ -3 & 4 & 2 \end{bmatrix} \end{array} \times \begin{array}{c} \mathbf{w} [3 \times 4] \\ \begin{bmatrix} 3 & 2 & 1 & -1 \\ 2 & 1 & 3 & 2 \\ 3 & 2 & 1 & -2 \end{bmatrix} \end{array} = \begin{array}{c} \mathbf{y} \\ \begin{bmatrix} x_{00}w_{00} + x_{01}w_{10} + x_{02}w_{20} & x_{00}w_{01} + x_{01}w_{11} + x_{02}w_{21} & x_{00}w_{02} + x_{01}w_{12} + x_{02}w_{22} & x_{00}w_{03} + x_{01}w_{13} + x_{02}w_{23} \\ x_{10}w_{00} + x_{11}w_{10} + x_{12}w_{20} & x_{10}w_{01} + x_{11}w_{11} + x_{12}w_{21} & x_{10}w_{02} + x_{11}w_{12} + x_{12}w_{22} & x_{10}w_{03} + x_{11}w_{13} + x_{12}w_{23} \end{bmatrix} \end{array}
 \end{array}$$

$$\begin{array}{c} \frac{\partial L}{\partial \mathbf{y}} [2 \times 4] \\ \begin{bmatrix} 2 & 3 & -3 & 9 \\ -8 & 1 & 4 & 6 \end{bmatrix} \end{array}$$

$$\frac{\partial L}{\partial \mathbf{x}} = \frac{\partial L}{\partial \mathbf{y}} \cdot \mathbf{w}^T$$

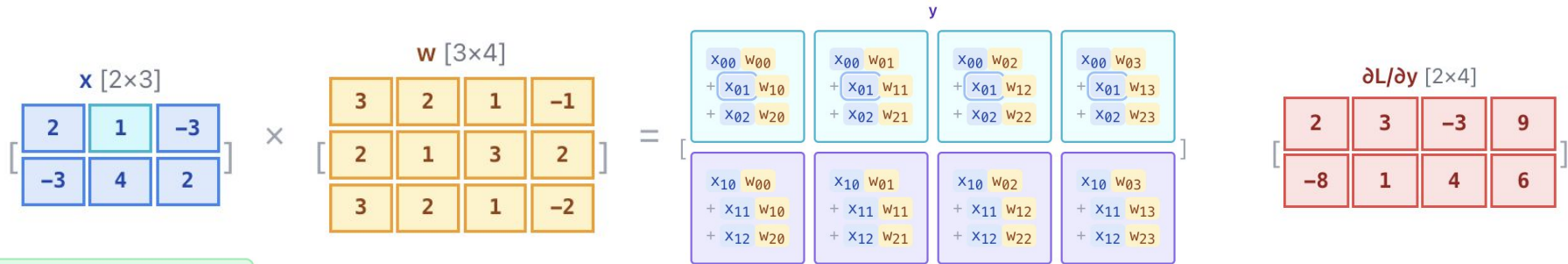
What is $\frac{\partial L}{\partial x_{0,1}}$?

$$\begin{array}{c} \frac{\partial L}{\partial \mathbf{y}} [2 \times 4] \\ \begin{bmatrix} \frac{\partial L}{\partial y_{0,0}} & \frac{\partial L}{\partial y_{0,1}} & \frac{\partial L}{\partial y_{0,2}} & \frac{\partial L}{\partial y_{0,3}} \\ \frac{\partial L}{\partial y_{1,0}} & \frac{\partial L}{\partial y_{1,1}} & \frac{\partial L}{\partial y_{1,2}} & \frac{\partial L}{\partial y_{1,3}} \end{bmatrix} \end{array} \times \begin{array}{c} \mathbf{w}^T [4 \times 3] \\ \begin{bmatrix} w_{0,0} & w_{1,0} & w_{2,0} \\ w_{0,1} & w_{1,1} & w_{2,1} \\ w_{0,2} & w_{1,2} & w_{2,2} \\ w_{0,3} & w_{1,3} & w_{2,3} \end{bmatrix} \end{array} =$$

$$\begin{array}{c} \frac{\partial L}{\partial \mathbf{x}} [2 \times 3] \\ \begin{bmatrix} \frac{\partial L}{\partial y_0} \cdot w_{0,0} & \frac{\partial L}{\partial y_0} \cdot w_{1,0} & \frac{\partial L}{\partial y_0} \cdot w_{2,0} \\ \frac{\partial L}{\partial y_1} \cdot w_{0,0} & \frac{\partial L}{\partial y_1} \cdot w_{1,0} & \frac{\partial L}{\partial y_1} \cdot w_{2,0} \end{bmatrix} \end{array}$$

Sum of: (how much $x_{0,1}$ affects each y) $\times$ (how much each y affects L)

$$\begin{aligned} &= \frac{\partial y_{0,0}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,0}} + \frac{\partial y_{0,1}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,1}} + \frac{\partial y_{0,2}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,2}} + \frac{\partial y_{0,3}}{\partial x_{0,1}} \cdot \frac{\partial L}{\partial y_{0,3}} \\ &+ 0 \cdot \frac{\partial L}{\partial y_{1,0}} + 0 \cdot \frac{\partial L}{\partial y_{1,1}} + 0 \cdot \frac{\partial L}{\partial y_{1,2}} + 0 \cdot \frac{\partial L}{\partial y_{1,3}} \\ &= w_{1,0} \cdot \frac{\partial L}{\partial y_{0,0}} + w_{1,1} \cdot \frac{\partial L}{\partial y_{0,1}} + w_{1,2} \cdot \frac{\partial L}{\partial y_{0,2}} + w_{1,3} \cdot \frac{\partial L}{\partial y_{0,3}} \\ &= \mathbf{w}[\text{row } 1] \cdot \frac{\partial L}{\partial \mathbf{y}}[\text{row } 0] \end{aligned}$$



$$\frac{\partial L}{\partial x} = \frac{\partial L}{\partial y} \cdot w^T$$

Summary: Two Ways to Compute the Same Thing

For $y = x \cdot w$ where x is [NxD], w is [DxM], y is [NxM]:

Explicit Jacobian Approach

$$\frac{\partial L}{\partial x} = \text{Jacobian} \frac{\partial y}{\partial x} \cdot \frac{\partial L}{\partial y}$$

Jacobian size: $[(N \cdot D) \times (N \cdot M)]$ – mostly zeros!

Dimensions: $(N \cdot D) \times (N \cdot M) \times (N \cdot M) \rightarrow (N \cdot D)$

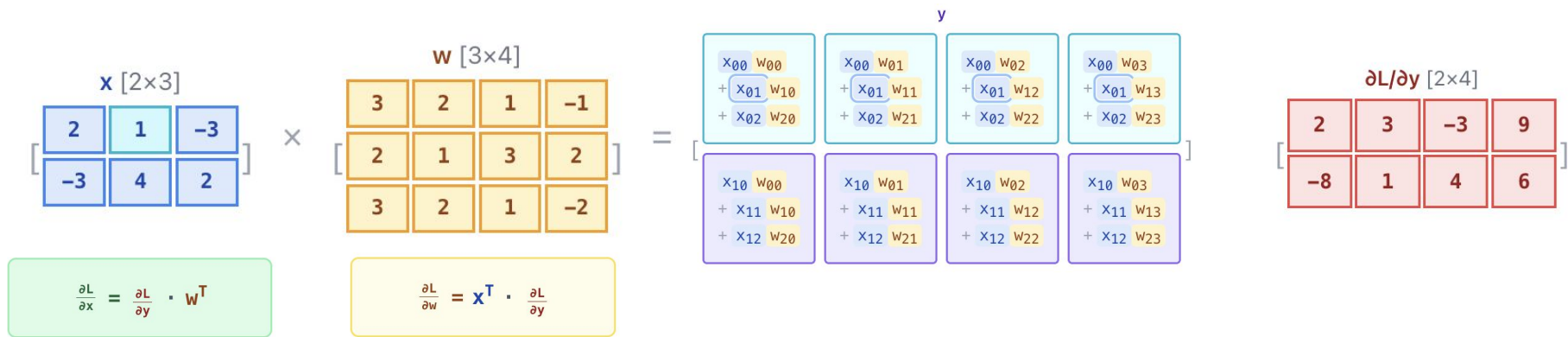
Large sparse matrix — wasteful

Implicit Matrix Multiply

$$\frac{\partial L}{\partial x} = \frac{\partial L}{\partial y} \cdot w^T$$

Dimensions: $(N \times M) \times (M \times D) \rightarrow (N \times D)$

Small dense matrix — efficient



Summary: Two Ways to Compute the Same Thing

For $y = x \cdot w$ where x is $[N \times D]$, w is $[D \times M]$, y is $[N \times M]$:

Explicit Jacobian Approach

$$\frac{\partial L}{\partial x} = \text{Jacobian} \frac{\partial y}{\partial x} \cdot \frac{\partial L}{\partial y}$$

Jacobian size: $[(N \cdot D) \times (N \cdot M)]$ – mostly zeros!

Dimensions: $(N \cdot D) \times (N \cdot M) \times (N \cdot M) \rightarrow (N \cdot D)$

Large sparse matrix – wasteful

Implicit Matrix Multiply

$$\frac{\partial L}{\partial x} = \frac{\partial L}{\partial y} \cdot w^T$$

Dimensions: $(N \times M) \times (M \times D) \rightarrow (N \times D)$

Small dense matrix – efficient

$x [2 \times 3]$

2	1	-3
-3	4	2

×

$w [3 \times 4]$

3	2	1	-1
2	1	3	2
3	2	1	-2

=

y

$x_{00} w_{00} + x_{01} w_{10} + x_{02} w_{20}$	$x_{00} w_{01} + x_{01} w_{11} + x_{02} w_{21}$	$x_{00} w_{02} + x_{01} w_{12} + x_{02} w_{22}$	$x_{00} w_{03} + x_{01} w_{13} + x_{02} w_{23}$
$x_{10} w_{00} + x_{11} w_{10} + x_{12} w_{20}$	$x_{10} w_{01} + x_{11} w_{11} + x_{12} w_{21}$	$x_{10} w_{02} + x_{11} w_{12} + x_{12} w_{22}$	$x_{10} w_{03} + x_{11} w_{13} + x_{12} w_{23}$

$\frac{\partial L}{\partial y} [2 \times 4]$

2	3	-3	9
-8	1	4	6

$\frac{\partial L}{\partial x} = \frac{\partial L}{\partial y} \cdot w^T$

$\frac{\partial L}{\partial w} = x^T \cdot \frac{\partial L}{\partial y}$

These formulas are easy to remember: they are the only way to make shapes match up!

Summary: Two Ways to Compute the Same Thing

For $y = x \cdot w$ where x is $[N \times D]$, w is $[D \times M]$, y is $[N \times M]$:

Explicit Jacobian Approach

$$\frac{\partial L}{\partial x} = \text{Jacobian } \frac{\partial y}{\partial x} \cdot \frac{\partial L}{\partial y}$$

Jacobian size: $[(N \cdot D) \times (N \cdot M)]$ — mostly zeros!

Dimensions: $(N \cdot D) \times (N \cdot M) \times (N \cdot M) \rightarrow (N \cdot D)$

Large sparse matrix — wasteful

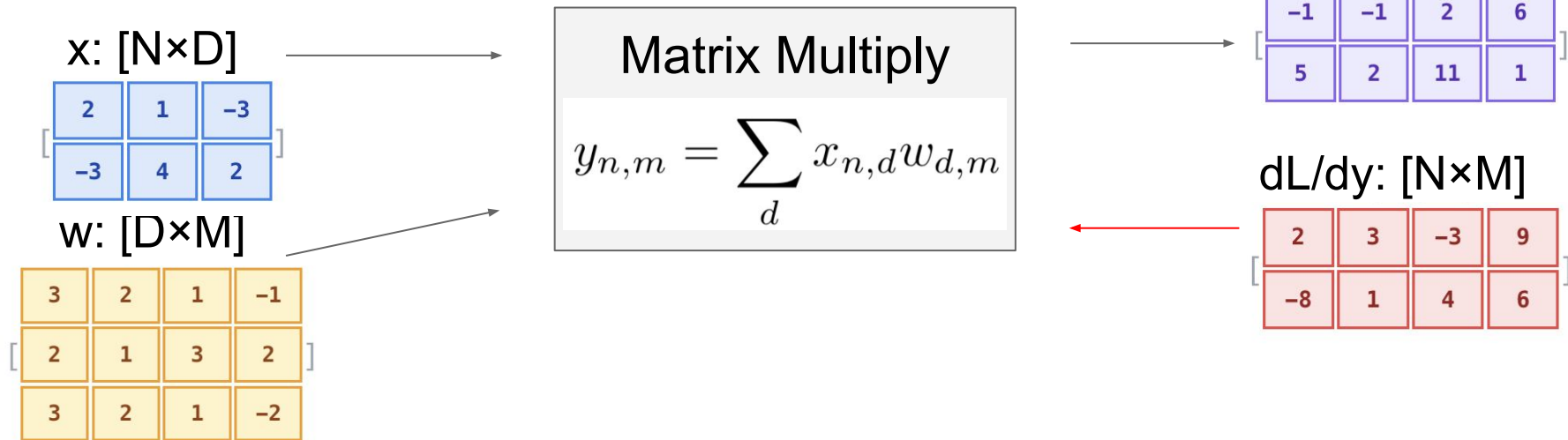
Implicit Matrix Multiply

$$\frac{\partial L}{\partial x} = \frac{\partial L}{\partial y} \cdot w^T$$

Dimensions: $(N \times M) \times (M \times D) \rightarrow (N \times D)$

Small dense matrix — efficient

Backprop with Matrices



$[N \times D] \quad [N \times M] \quad [M \times D]$

$$\frac{\partial L}{\partial x} = \left(\frac{\partial L}{\partial y} \right) w^T$$

$[D \times M] \quad [D \times N] \quad [N \times M]$

$$\frac{\partial L}{\partial w} = x^T \left(\frac{\partial L}{\partial y} \right)$$

Today's agenda

- Briefly revisit backprop
- Finish CNNs
- Activation Functions
- Data Preprocessing
- Weight Initialization
- Normalization Layers

Convolution Layer

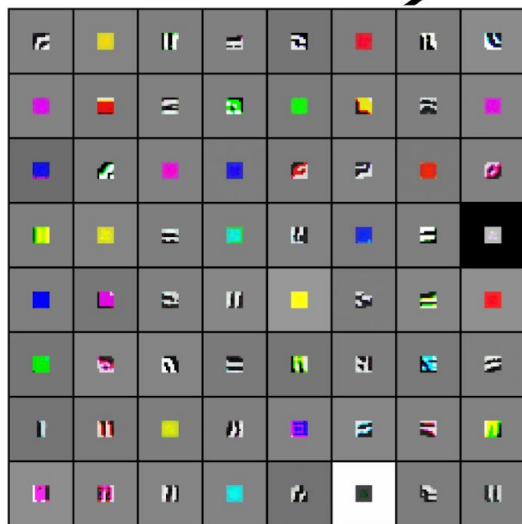
consider a second, **green** filter



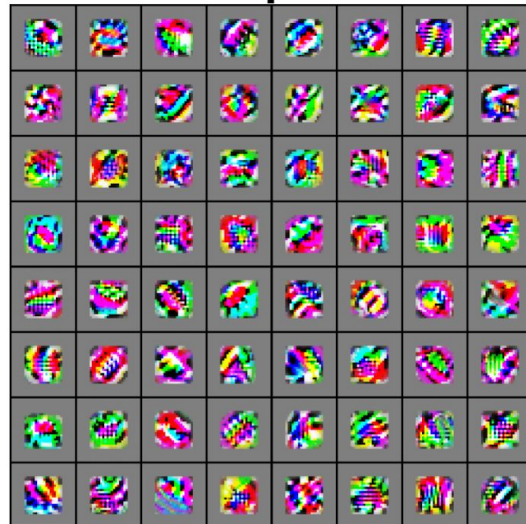
Preview

[Zeiler and Fergus 2013]

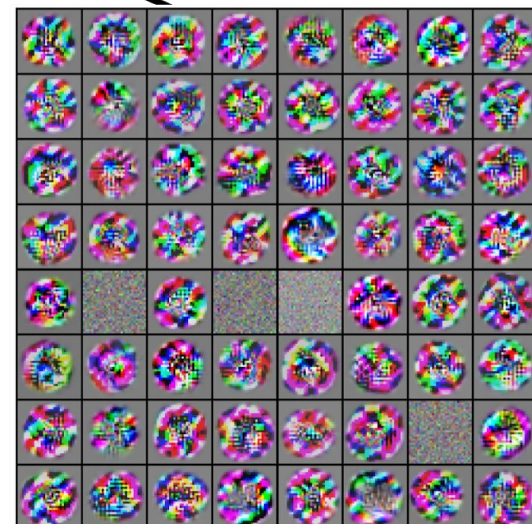
Visualization of VGG-16 by Lane McIntosh. VGG-16 architecture from [Simonyan and Zisserman 2014].



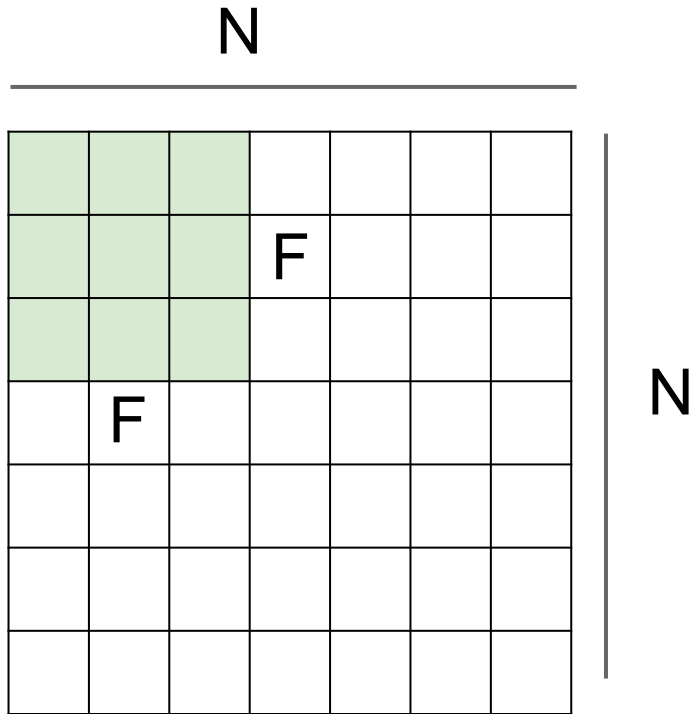
VGG-16 Conv1_1



VGG-16 Conv3_2



VGG-16 Conv5_3



Output size:
 $(N - F) / \text{stride} + 1$

e.g. $N = 7, F = 3$:

stride 1 $\Rightarrow (7 - 3) / 1 + 1 = 5$

stride 2 $\Rightarrow (7 - 3) / 2 + 1 = 3$

stride 3 $\Rightarrow (7 - 3) / 3 + 1 = 2.33 \therefore \backslash$

In practice: Common to zero pad the border

0	0	0	0	0	0			
0								
0								
0								
0								

e.g. input 7x7

3x3 filter, applied with **stride 1**

pad with 1 pixel border => what is the output?

(recall:)

$$(N - F) / \text{stride} + 1$$

In practice: Common to zero pad the border

0	0	0	0	0	0			
0								
0								
0								
0								

e.g. input 7x7

3x3 filter, applied with **stride 1**

pad with 1 pixel border => what is the output?

7x7 output!

(recall:)

$$(N + 2P - F) / \text{stride} + 1$$

In practice: Common to zero pad the border

0	0	0	0	0	0			
0								
0								
0								
0								

e.g. input 7x7

3x3 filter, applied with **stride 1**

pad with 1 pixel border => what is the output?

7x7 output!

in general, common to see CONV layers with stride 1, filters of size $F \times F$, and zero-padding with $(F-1)/2$. (will preserve size spatially)

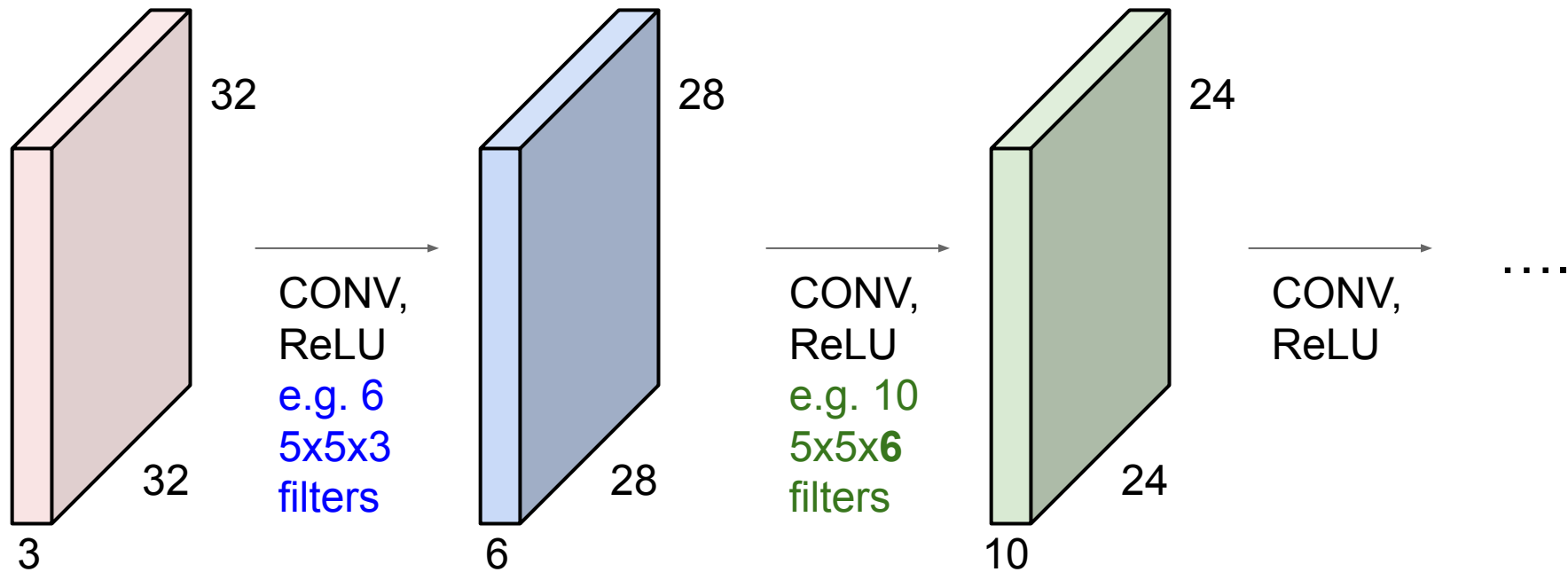
e.g. $F = 3 \Rightarrow$ zero pad with 1

$F = 5 \Rightarrow$ zero pad with 2

$F = 7 \Rightarrow$ zero pad with 3

Remember back to...

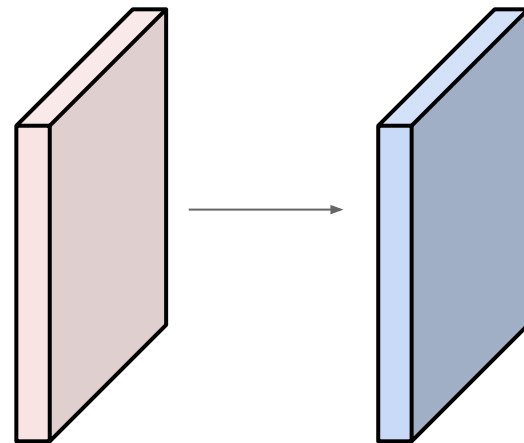
E.g. 32x32 input convolved repeatedly with 5x5 filters shrinks volumes spatially!
(32 $\rightarrow$ 28 $\rightarrow$ 24 ...).



Examples time:

Input volume: **32x32x3**

10 5x5 filters with stride 1, pad 2



Let's assume output size is $H \times W \times D$.
What is D ?

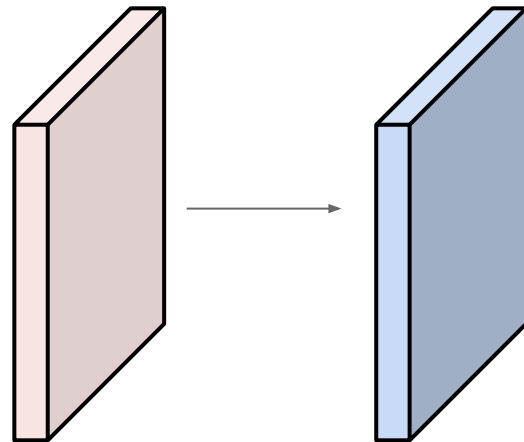
Examples time:

Input volume: **32x32x3**

10 5x5 filters with stride 1, pad 2

Let's assume output size is $H \times W \times D$.

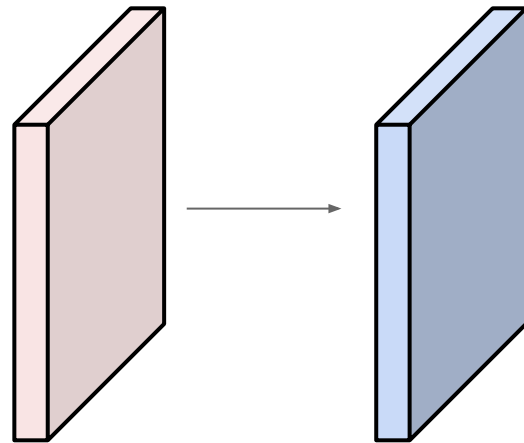
What is D ? **10**



Examples time:

Input volume: **32x32x3**

10 5x5 filters with stride 1, pad 2



Let's assume output size is $H \times W \times D$.

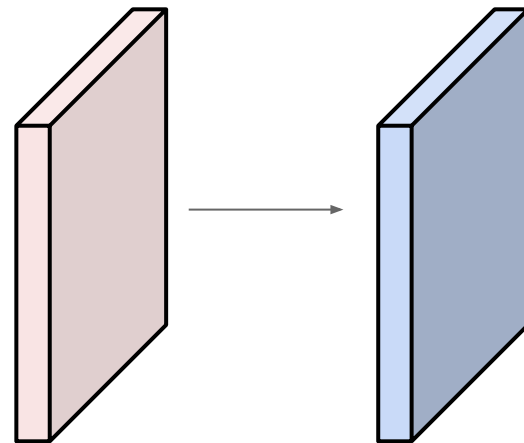
What is D ? **10**

What is H or W ?

Examples time:

Input volume: **32x32x3**

10 **5x5** filters with stride **1**, pad **2**



Let's assume output size is $H \times W \times D$.

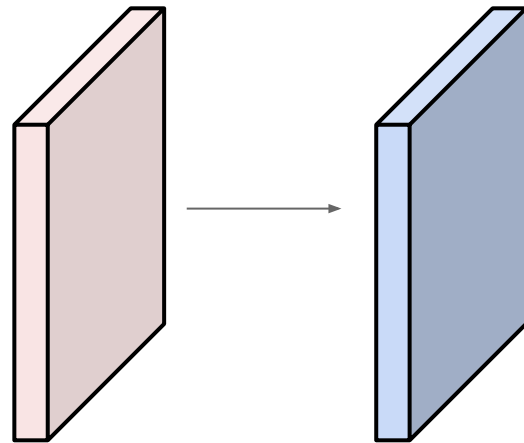
What is D ? 10

What is H or W ? $(32 + 2 * 2 - 5) / 1 + 1 = 32$

Examples time:

Input volume: **32x32x3**

10 **5x5** filters with stride **1**, pad **2**



Let's assume output size is $H \times W \times D$.

What is D ? **10**

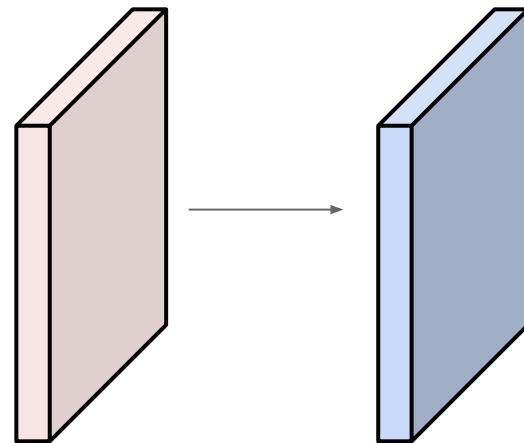
What is H or W ? $(32 + 2 * 2 - 5) / 1 + 1 = 32$

So the total output size is: **32x32x10**

Examples time:

Input volume: **32x32x3**

10 5x5 filters with stride 1, pad 2

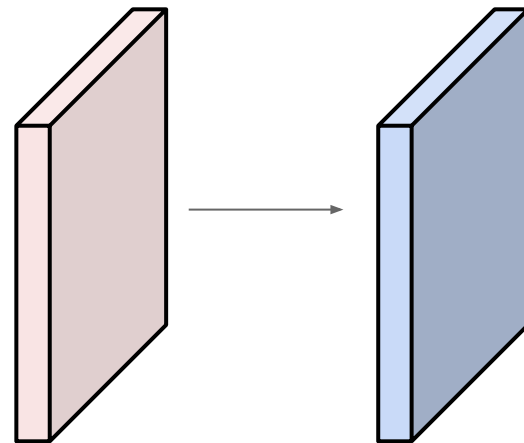


Number of parameters in this layer?

Examples time:

Input volume: **32x32x3**

10 **5x5** filters with stride 1, pad 2



Number of parameters in this layer?

each filter has $5*5*3 + 1 = 76$ params (+1 for bias)

=> $76*10 = 760$

Convolution layer: summary

Let's assume input is $W_1 \times H_1 \times C$

Conv layer needs 4 hyperparameters:

- Number of filters **K**
- The filter size **F**
- The stride **S**
- The zero padding **P**

This will produce an output of $W_2 \times H_2 \times K$
where:

- $W_2 = (W_1 - F + 2P)/S + 1$
- $H_2 = (H_1 - F + 2P)/S + 1$

Number of parameters: F^2KC and K biases

Convolution layer: summary

Common settings:

Let's assume input is $W_1 \times H_1 \times C$

Conv layer needs 4 hyperparameters:

- Number of filters **K**
- The filter size **F**
- The stride **S**
- The zero padding **P**

K = (powers of 2, e.g. 32, 64, 128, 512)

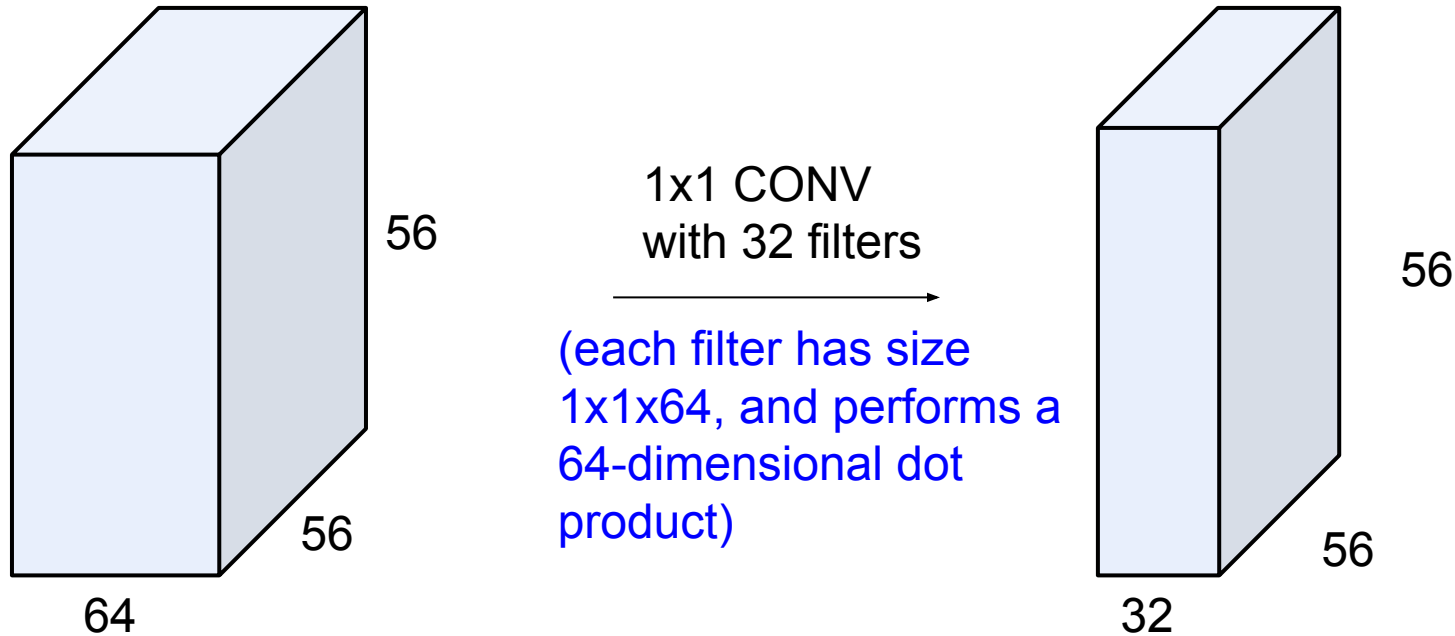
- **F** = 3, **S** = 1, **P** = 1
- **F** = 5, **S** = 1, **P** = 2
- **F** = 5, **S** = 2, **P** = ? (whatever fits)
- **F** = 1, **S** = 1, **P** = 0

This will produce an output of $W_2 \times H_2 \times K$
where:

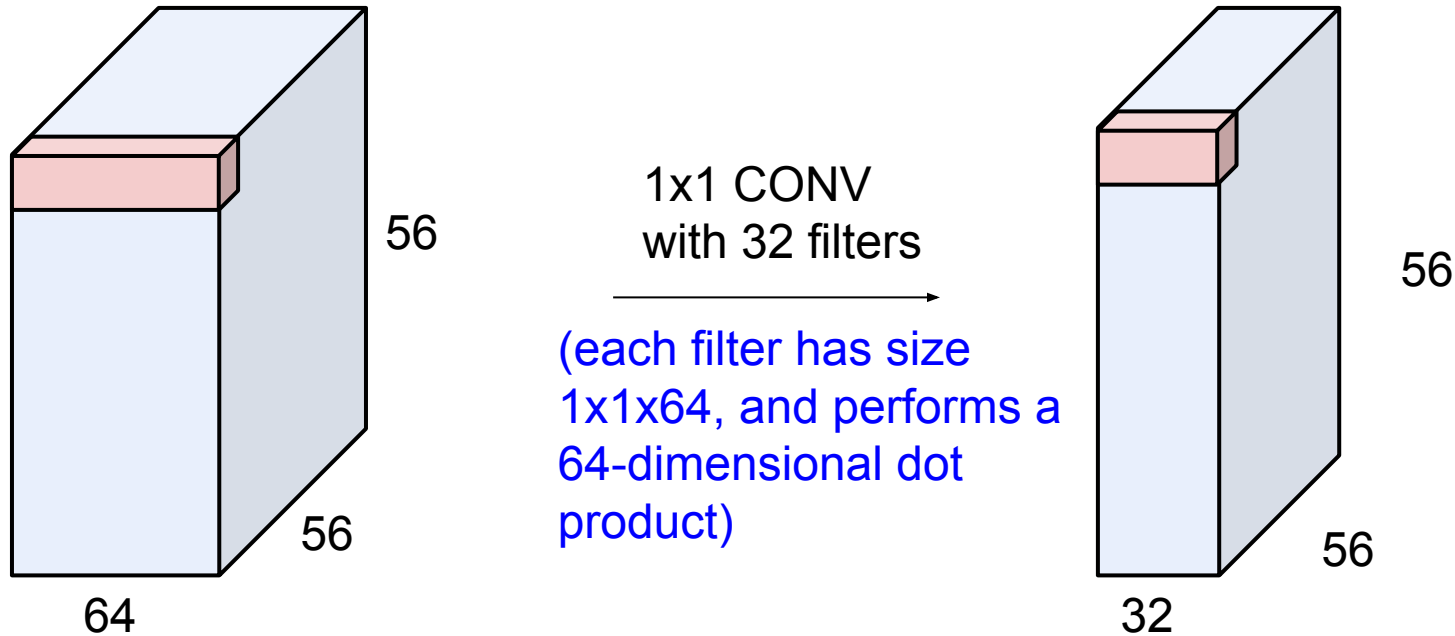
- $W_2 = (W_1 - F + 2P)/S + 1$
- $H_2 = (H_1 - F + 2P)/S + 1$

Number of parameters: F^2CK and K biases

(btw, 1x1 convolution layers are very useful)



(btw, 1x1 convolution layers are a very useful)



Example: CONV layer in PyTorch

Conv2d

```
CLASS torch.nn.Conv2d(in_channels, out_channels, kernel_size, stride=1, padding=0, dilation=1, groups=1, bias=True)
```

[SOURCE]

Applies a 2D convolution over an input signal composed of several input planes.

In the simplest case, the output value of the layer with input size (N, C_{in}, H, W) and output $(N, C_{out}, H_{out}, W_{out})$ can be precisely described as:

$$\text{out}(N_i, C_{out,j}) = \text{bias}(C_{out,j}) + \sum_{k=0}^{C_{in}-1} \text{weight}(C_{out,j}, k) \star \text{input}(N_i, k)$$

where $\star$ is the valid 2D **cross-correlation** operator, N is a batch size, C denotes a number of channels, H is a height of input planes in pixels, and W is width in pixels.

- `stride` controls the stride for the cross-correlation, a single number or a tuple.
- `padding` controls the amount of implicit zero-paddings on both sides for `padding` number of points for each dimension.
- `dilation` controls the spacing between the kernel points; also known as the à trous algorithm. It is harder to describe, but this [link](#) has a nice visualization of what `dilation` does.
- `groups` controls the connections between inputs and outputs. `in_channels` and `out_channels` must both be divisible by `groups`. For example,
 - At `groups=1`, all inputs are convolved to all outputs.
 - At `groups=2`, the operation becomes equivalent to having two conv layers side by side, each seeing half the input channels, and producing half the output channels, and both subsequently concatenated.
 - At `groups= in_channels`, each input channel is convolved with its own set of filters, of size: $\begin{bmatrix} C_{out} \\ C_{in} \end{bmatrix}$.

The parameters `kernel_size`, `stride`, `padding`, `dilation` can either be:

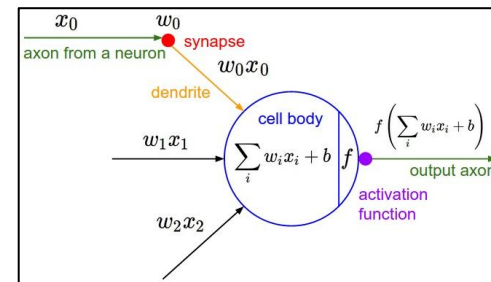
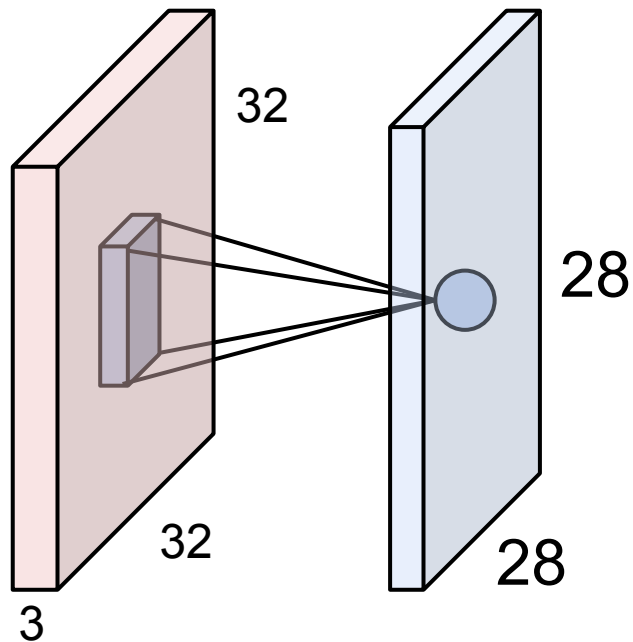
- a single `int` – in which case the same value is used for the height and width dimension
- a `tuple` of two ints – in which case, the first `int` is used for the height dimension, and the second `int` for the width dimension

PyTorch is licensed under [BSD 3-clause](#).

Conv layer needs 4 hyperparameters:

- Number of filters **K**
- The filter size **F**
- The stride **S**
- The zero padding **P**

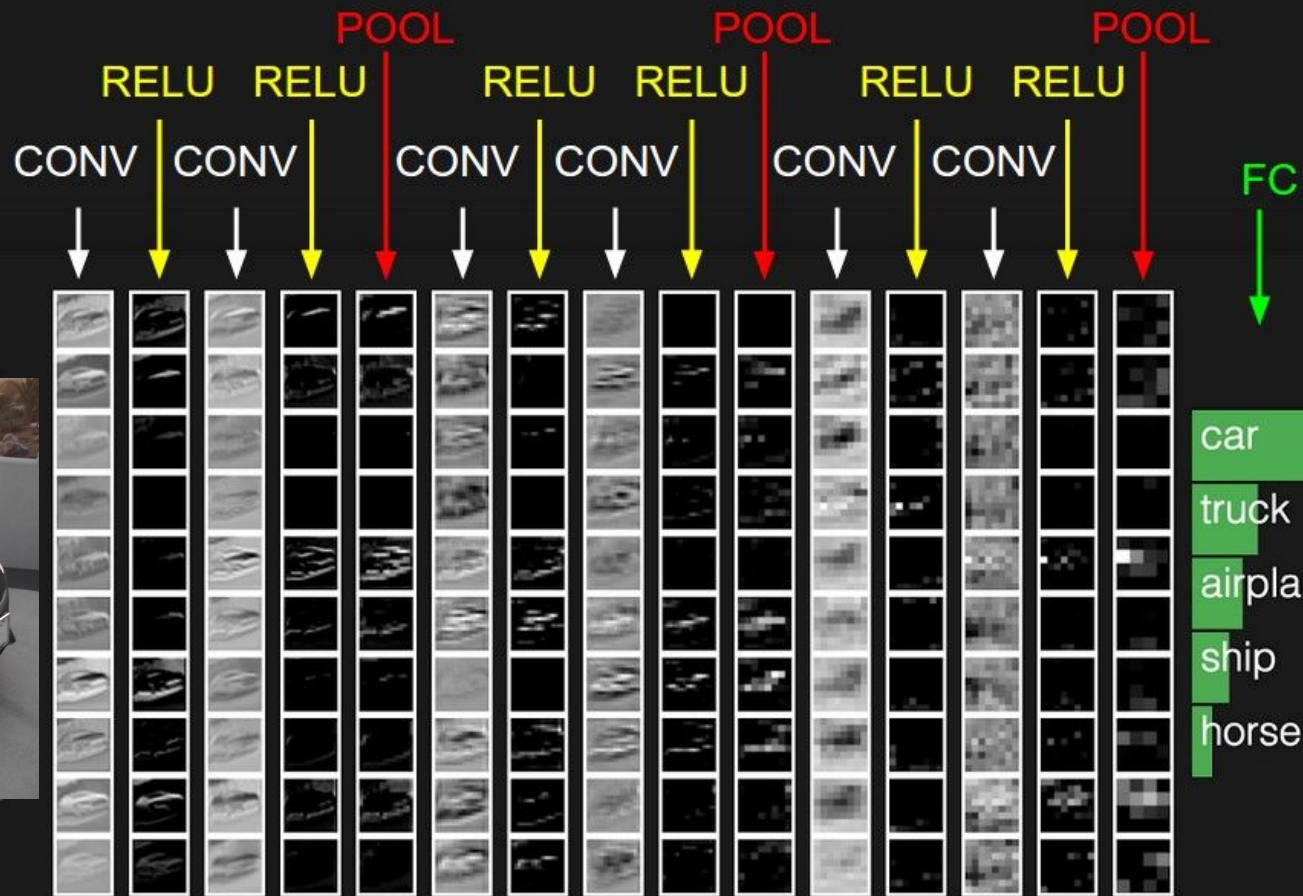
Receptive field



An activation map is a 28x28 sheet of neuron outputs:

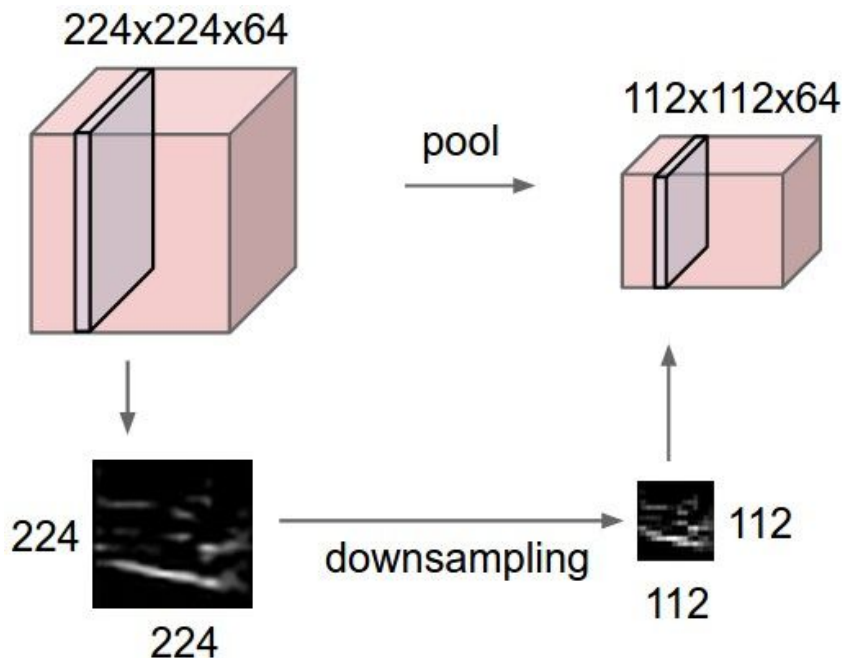
1. Each is connected to a small region in the input
2. All of them share parameters

“5x5 filter” -> “5x5 receptive field for each neuron”

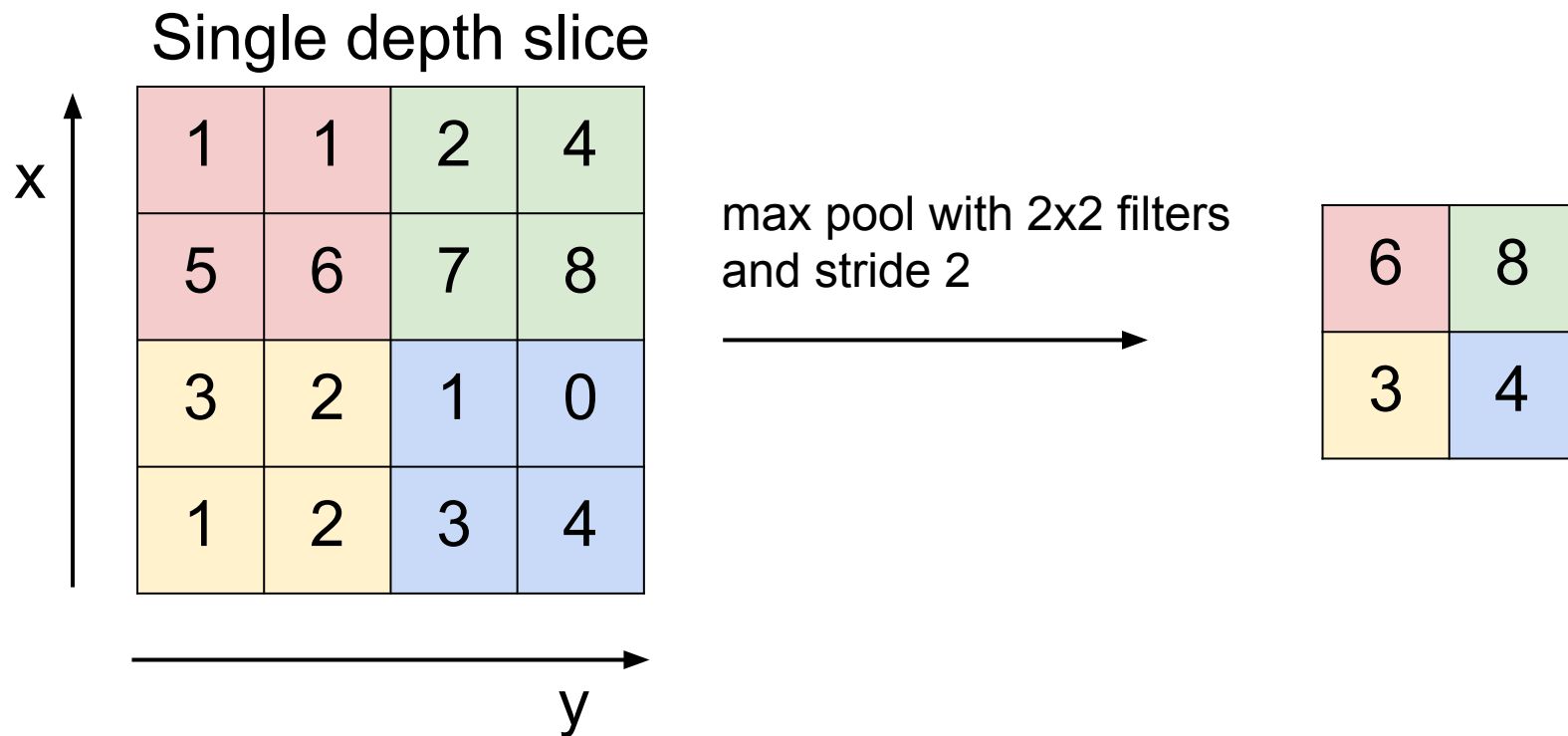


Pooling layer

- makes the representations smaller and more manageable
- operates over each activation map independently:



MAX POOLING



Pooling layer: summary

Let's assume input is $W_1 \times H_1 \times C$

Conv layer needs 2 hyperparameters:

- The spatial extent **F**
- The stride **S**

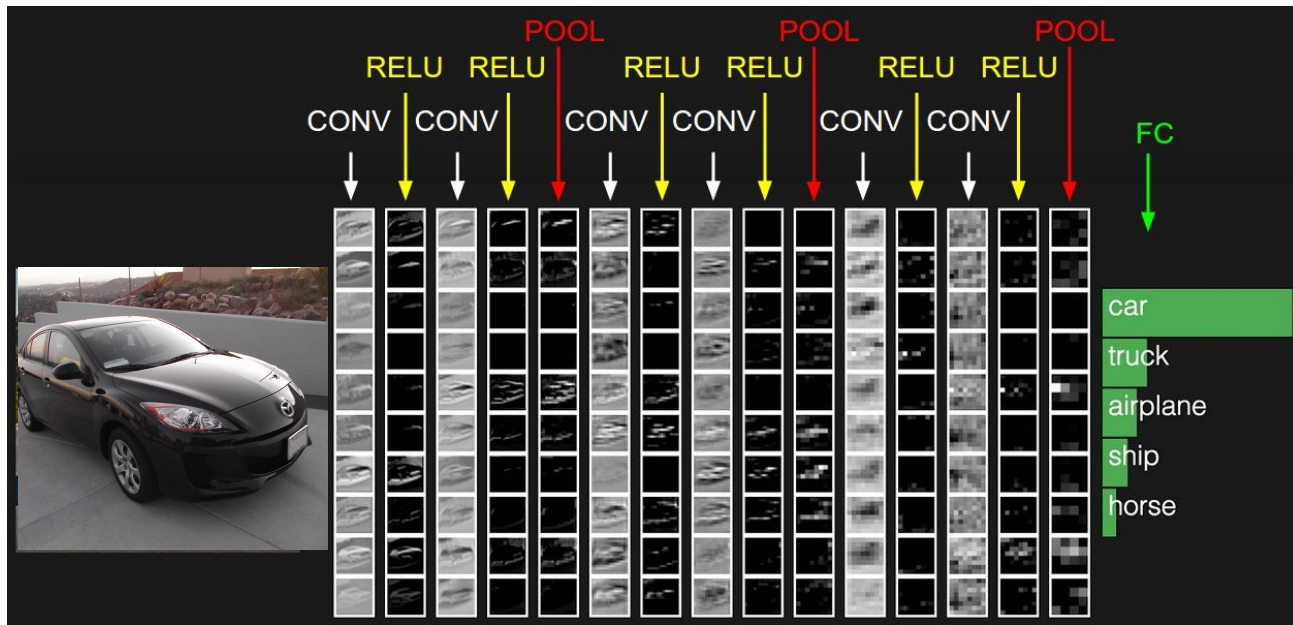
This will produce an output of $W_2 \times H_2 \times C$ where:

- $W_2 = (W_1 - F) / S + 1$
- $H_2 = (H_1 - F) / S + 1$

Number of parameters: 0

Fully Connected Layer (FC layer)

- Contains neurons that connect to the entire input volume, as in ordinary Neural Networks



Summary

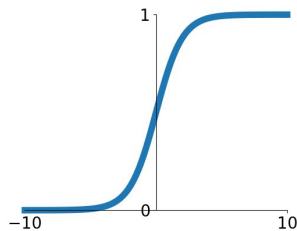
- ConvNets stack CONV, POOL, FC layers
- Trend towards smaller filters and deeper architectures
- Trend towards getting rid of POOL/FC layers (just CONV)
- Between 2012-2016 architectures looked like
 $[(\text{CONV-RELU})^N \text{- POOL?}]^M \text{- (FC-RELU)}^K, \text{SOFTMAX}$
where N is usually up to ~5, M is large, $0 \leq K \leq 2$.
 - but recent advances such as ResNet/GoogLeNet have challenged this paradigm

Activation Functions

Activation Functions

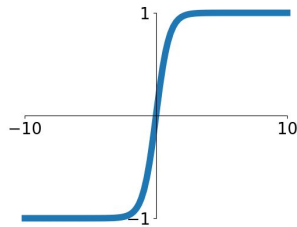
Sigmoid

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



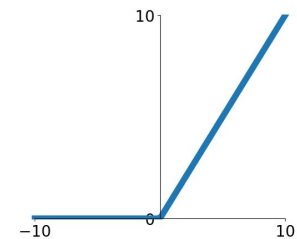
tanh

$$\tanh(x)$$



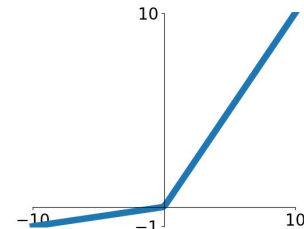
ReLU

$$\max(0, x)$$



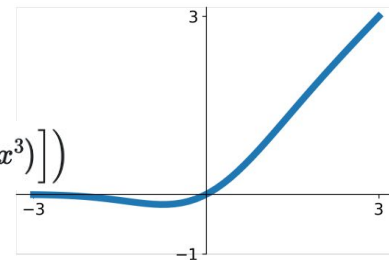
Leaky ReLU

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



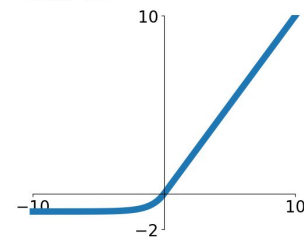
GeLU

$$0.5x \left(1 + \tanh \left[\sqrt{2/\pi} (x + 0.044715x^3) \right] \right)$$

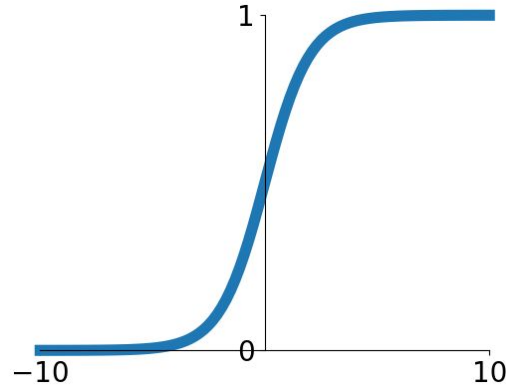


ELU

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



Sigmoid

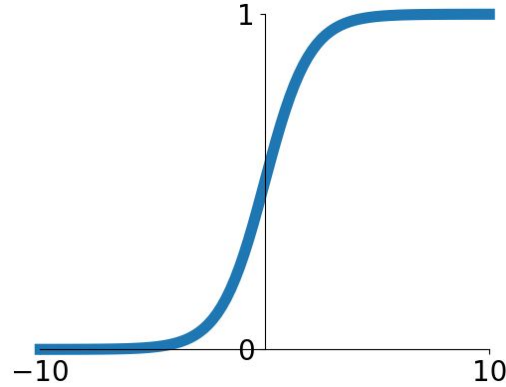


Sigmoid

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

- Squashes numbers to range [0,1]
- Historically popular since they have nice interpretation as a saturating “firing rate” of a neuron

Sigmoid



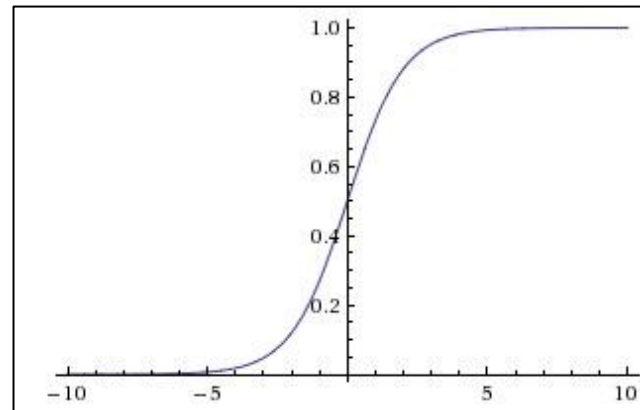
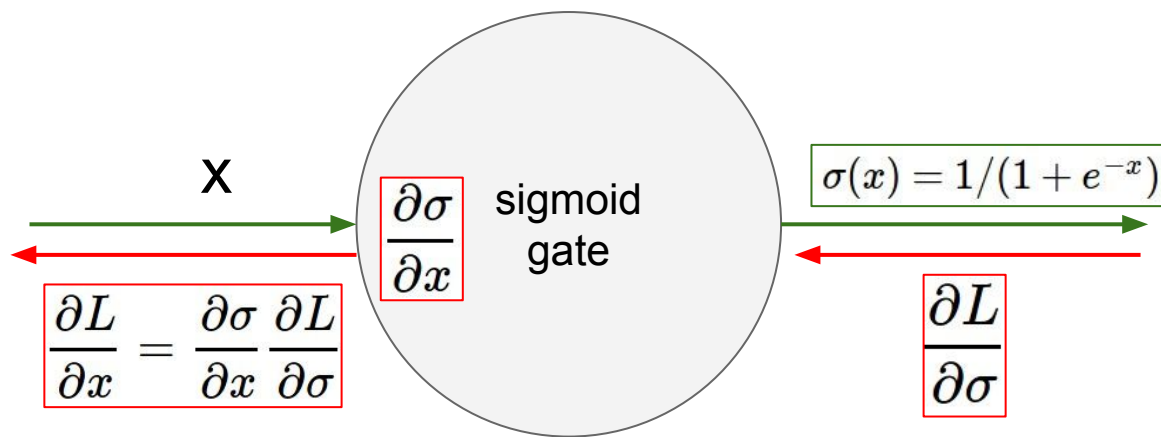
Sigmoid

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

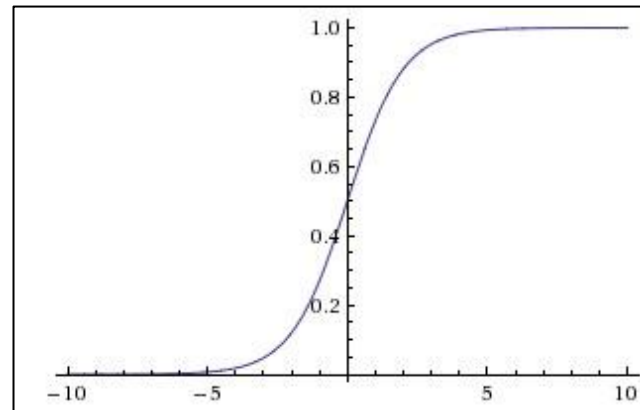
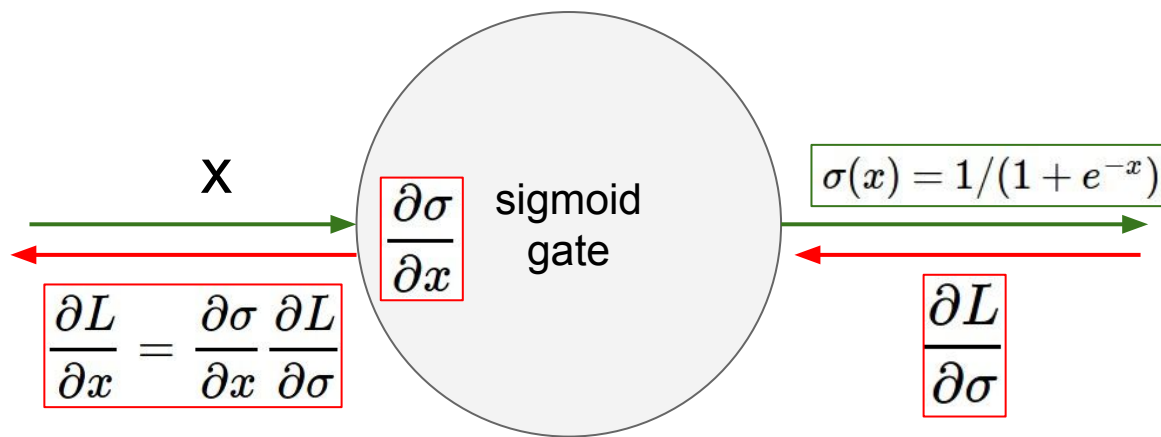
- Squashes numbers to range [0,1]
- Historically popular since they have nice interpretation as a saturating “firing rate” of a neuron

3 problems:

1. Saturated neurons “kill” the gradients

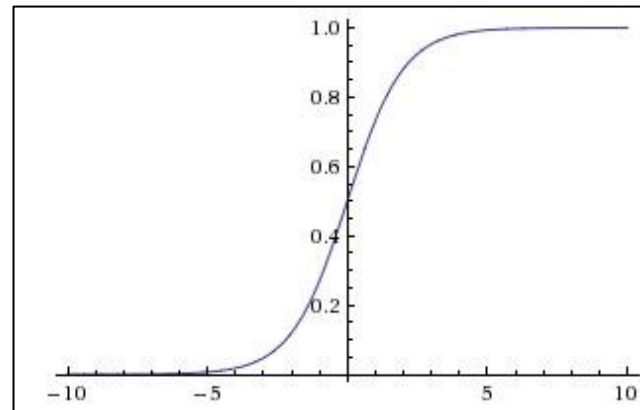
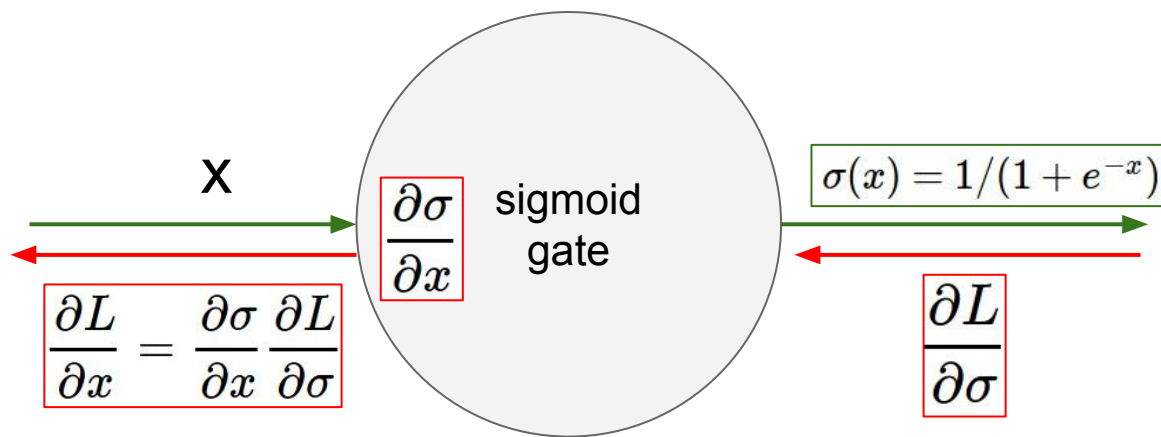


$$\frac{\partial \sigma(x)}{\partial x} = \sigma(x) (1 - \sigma(x))$$



What happens when $x = -10$?

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

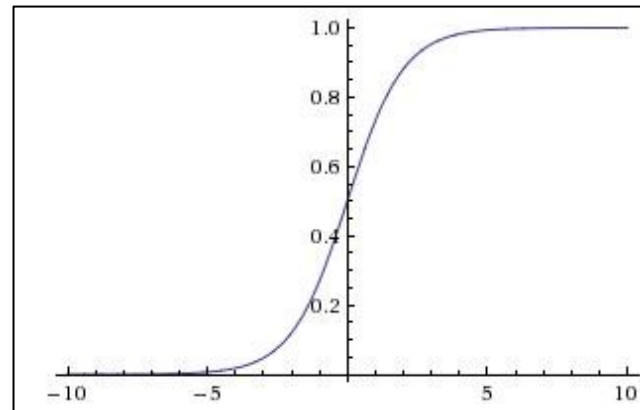
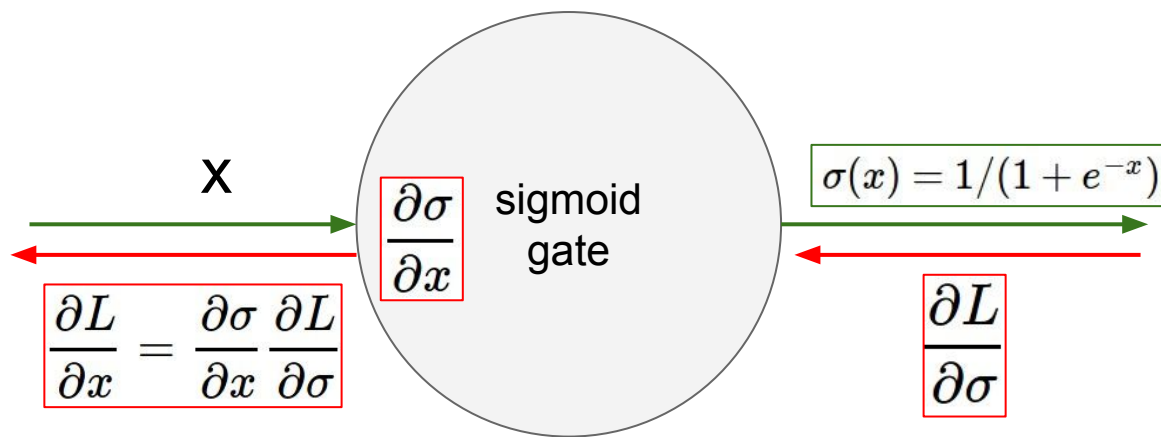


What happens when $x = -10$?

$$\sigma(x) \approx 0$$

$$\frac{\partial \sigma(x)}{\partial x} = \sigma(x) (1 - \sigma(x)) = 0(1 - 0) = 0$$

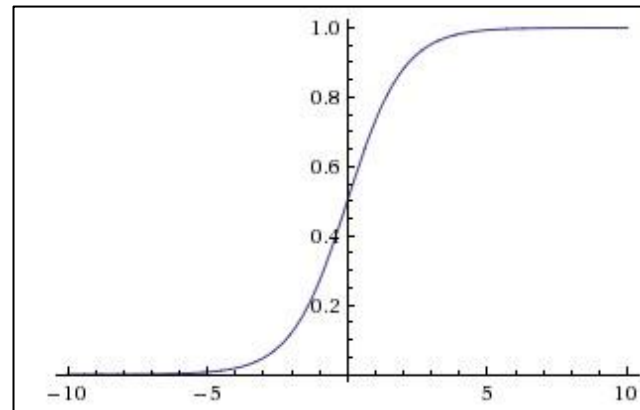
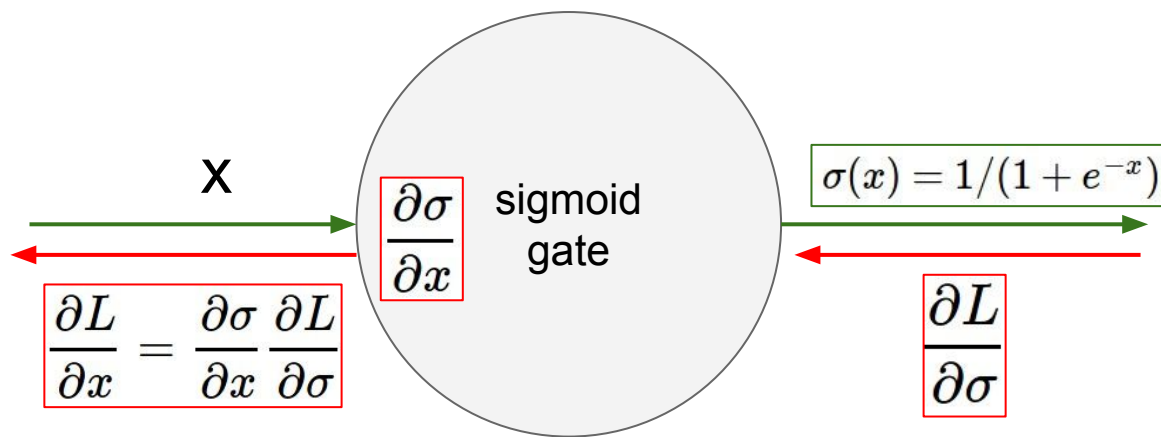
$$\frac{\partial \sigma(x)}{\partial x} = \sigma(x) (1 - \sigma(x))$$



What happens when $x = -10$?

What happens when $x = 0$?

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

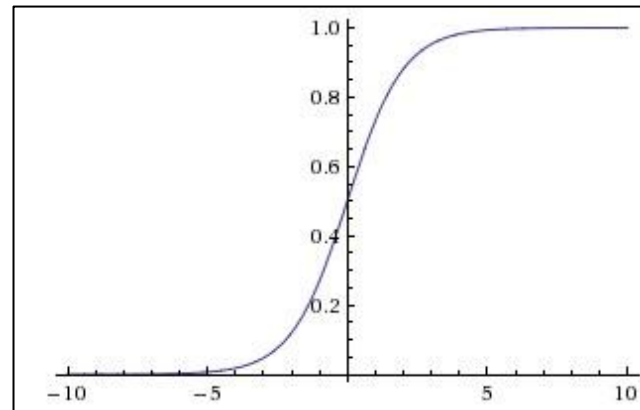
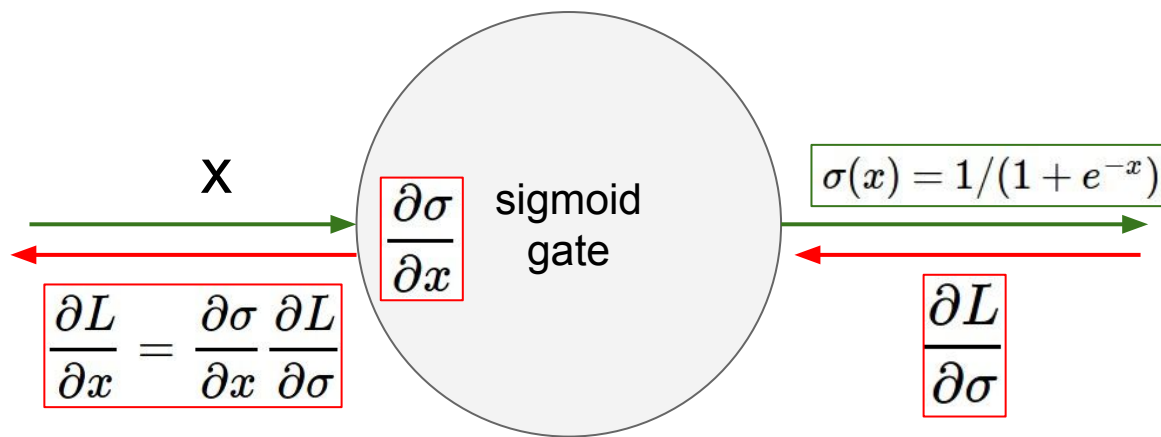


What happens when $x = -10$?

What happens when $x = 0$?

What happens when $x = 10$?

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



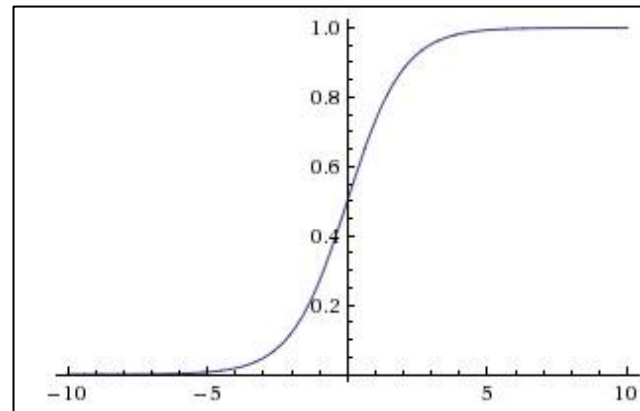
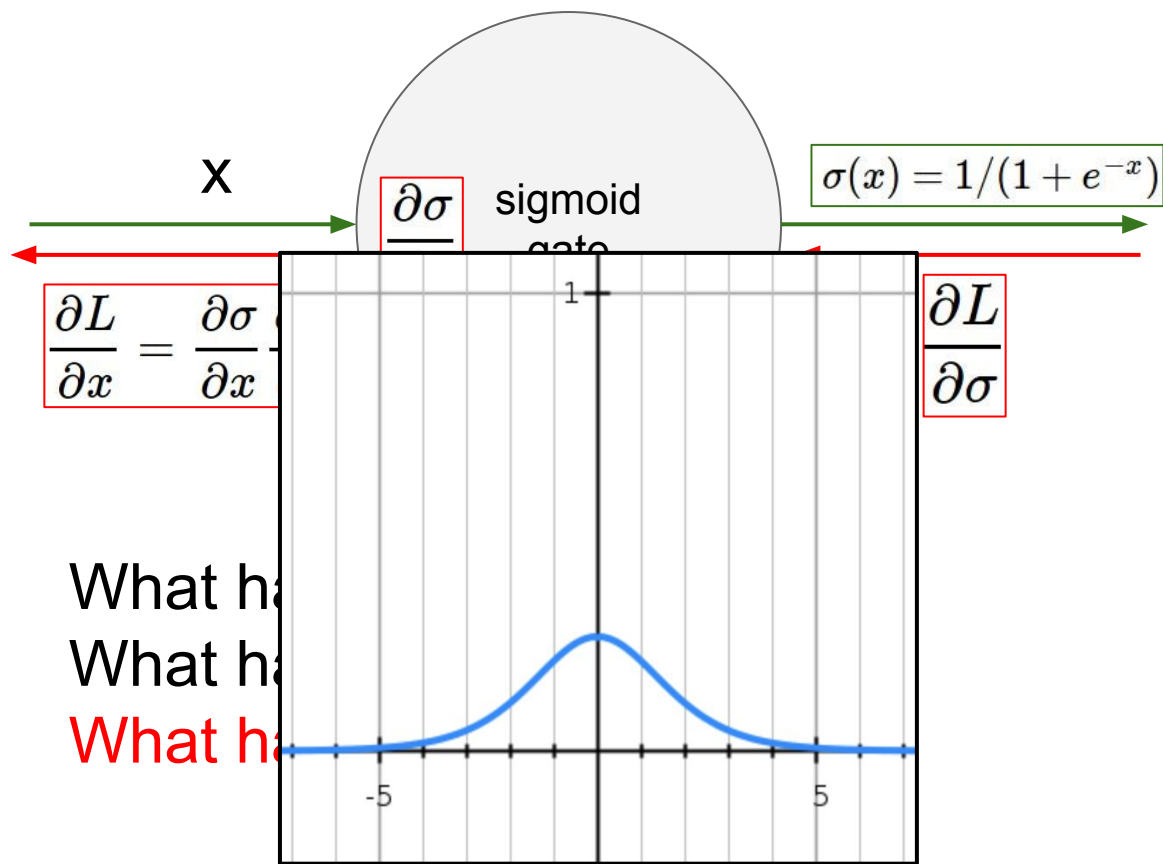
What happens when $x = -10$?

What happens when $x = 0$?

What happens when $x = 10$?

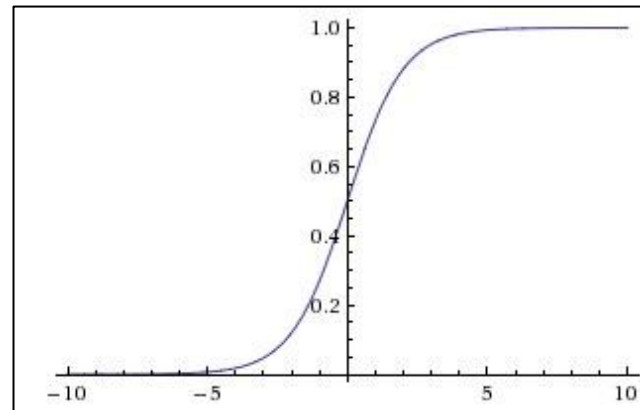
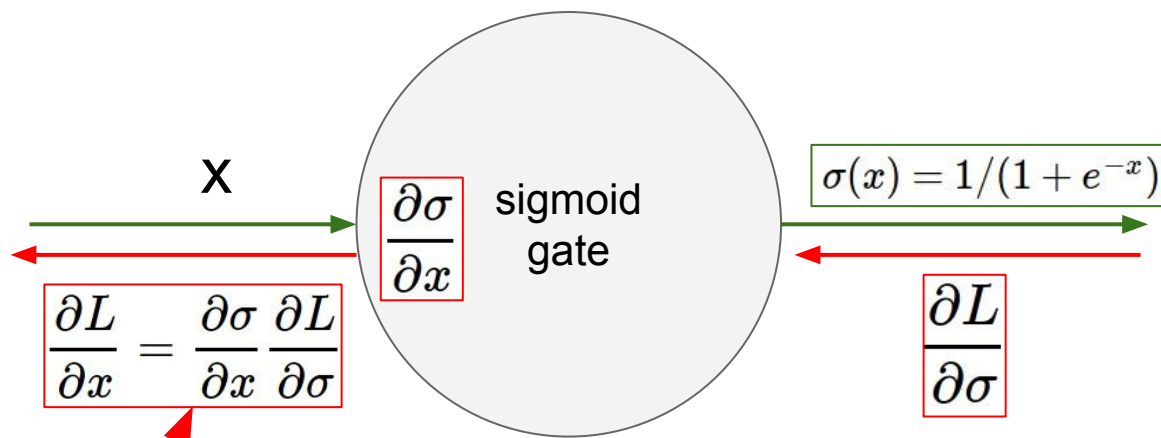
$$\frac{\partial \sigma(x)}{\partial x} = \sigma(x) (1 - \sigma(x))$$

$$\sigma(x) \approx 1 \quad \frac{\partial \sigma(x)}{\partial x} = \sigma(x) (1 - \sigma(x)) = 1(1 - 1) = 0$$



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

What has
What has
What has

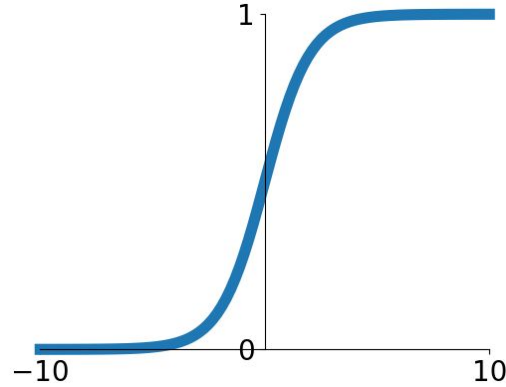


$$\frac{\partial \sigma(x)}{\partial x} = \sigma(x) (1 - \sigma(x))$$

Why is this a problem?

If all the gradients flowing back will be zero and weights will never change

Sigmoid



Sigmoid

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

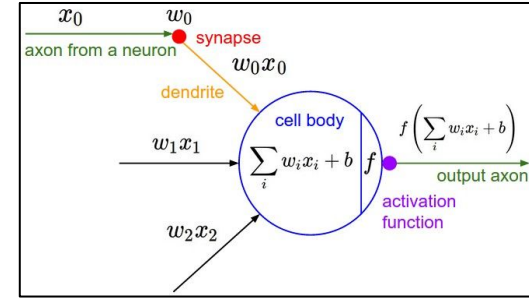
- Squashes numbers to range [0,1]
- Historically popular since they have nice interpretation as a saturating “firing rate” of a neuron

3 problems:

1. Saturated neurons “kill” the gradients
2. Sigmoid outputs are not zero-centered

Consider what happens when the input to a neuron is always positive...

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



What can we say about the gradients for that neuron (i.e. a row of $\mathbf{w}$)?

Neuron/row = (vector size n)

Input = $\mathbf{x}$ (vector size n , all positive)

Output = z (what shape?)

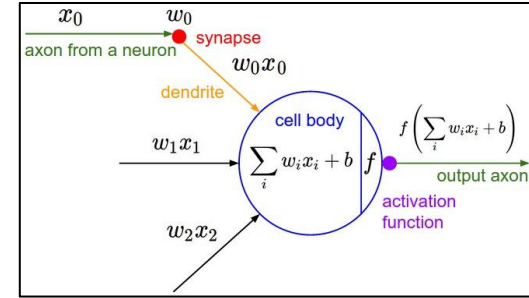
dL/dz = (what shape?)

$dL/d\mathbf{w}$ = (what shape?)

$$\frac{\partial L}{\partial \mathbf{w}} = \frac{\partial z}{\partial \mathbf{w}} \cdot \frac{\partial L}{\partial z}$$

Consider what happens when the input to a neuron is always positive...

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



What can we say about the gradients for that neuron (i.e. a row of $\mathbf{w}$)?

Neuron/row = (vector size n)

Input = $\mathbf{x}$ (vector size n , all positive)

Output = z (what shape?)

dL/dz = (what shape?)

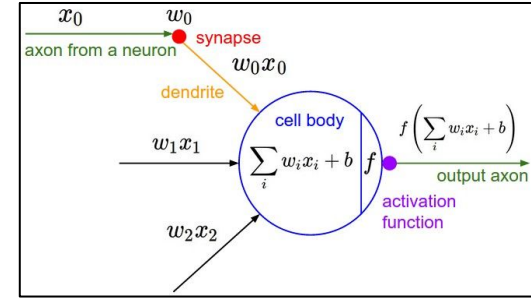
$dL/d\mathbf{w}$ = (what shape?)

$$\frac{\partial L}{\partial \mathbf{w}} = \frac{\partial z}{\partial \mathbf{w}} \cdot \boxed{\frac{\partial L}{\partial z}}$$

Scalar (+ or -)

Consider what happens when the input to a neuron is always positive...

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



What can we say about the gradients for that neuron (i.e. a row of $\mathbf{w}$)?

Neuron/row = (vector size n)

Input = $\mathbf{x}$ (vector size n , all positive)

Output = z (what shape?)

dL/dz = (what shape?)

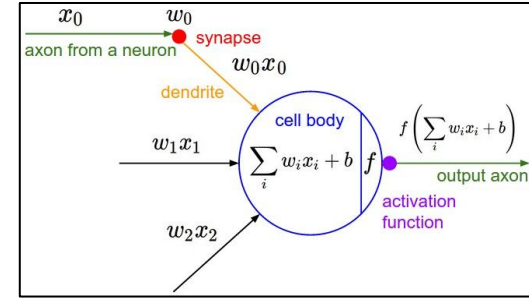
dL/dw = (what shape?)

$$\frac{\partial L}{\partial \mathbf{w}} = \boxed{\frac{\partial z}{\partial \mathbf{w}}} \cdot \boxed{\frac{\partial L}{\partial z}}$$

X (all +) Scalar (+ or -)

Consider what happens when the input to a neuron is always positive...

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



What can we say about the gradients for that neuron (i.e. a row of $\mathbf{w}$)?

Neuron/row = (vector size n)

Input = $\mathbf{x}$ (vector size n , all positive)

Output = z (what shape?)

dL/dz = (what shape?)

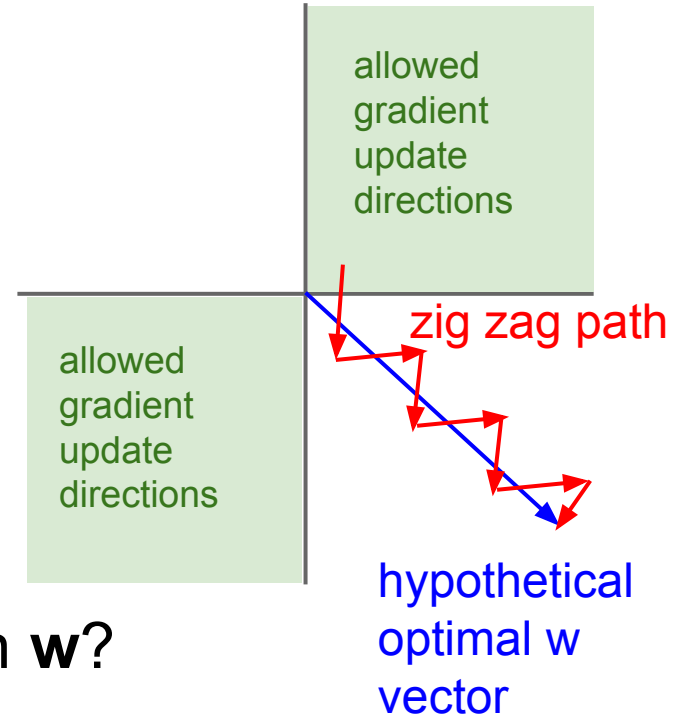
dL/dw = (what shape?)

$$\boxed{\frac{\partial L}{\partial w}} = \boxed{\frac{\partial z}{\partial w}} \cdot \boxed{\frac{\partial L}{\partial z}}$$

(all + or -) X (all +) Scalar (+ or -)

Consider what happens when the input to a neuron is always positive...

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

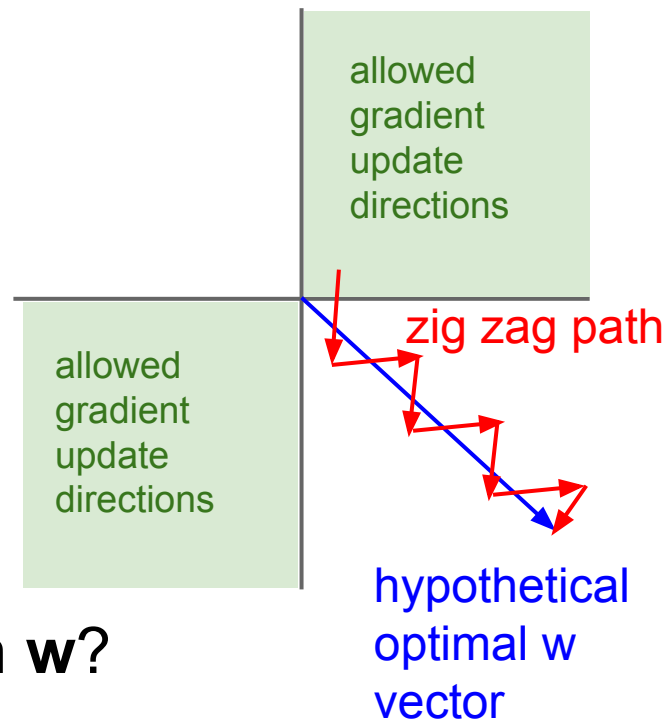


What can we say about the gradients on $\mathbf{w}$?

Always all positive or all negative :(

Consider what happens when the input to a neuron is always positive...

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

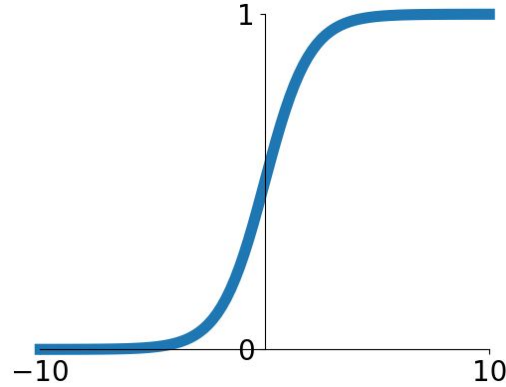


What can we say about the gradients on $\mathbf{w}$?

Always all positive or all negative :(

(For a single element! Minibatches help)

Sigmoid



Sigmoid

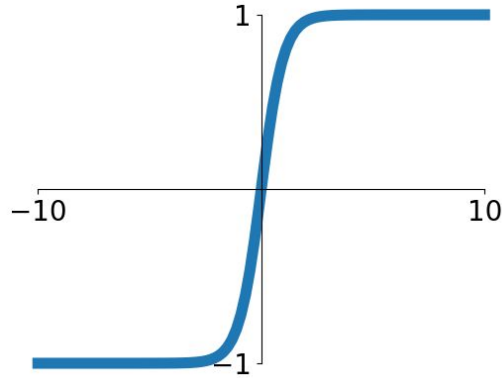
$$\sigma(x) = 1/(1 + e^{-x})$$

- Squashes numbers to range [0,1]
- Historically popular since they have nice interpretation as a saturating “firing rate” of a neuron

3 problems:

1. Saturated neurons “kill” the gradients
2. Sigmoid outputs are not zero-centered
3. $\exp()$ is a bit compute expensive

Tanh

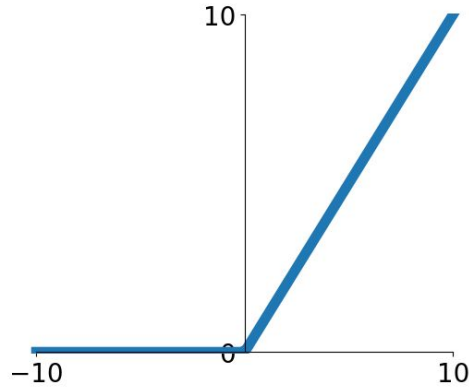


$\tanh(x)$

- Squashes numbers to range $[-1,1]$
- zero centered (nice)
- still kills gradients when saturated :(

[LeCun et al., 1991]

ReLU

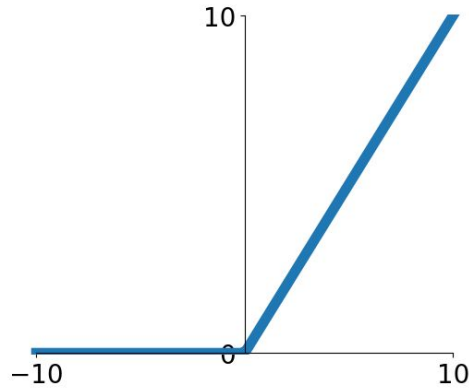


- Computes $f(x) = \max(0, x)$
- Does not saturate (in +region)
- Very computationally efficient
- Converges much faster than sigmoid/tanh in practice (e.g. 6x)

ReLU
(Rectified Linear Unit)

[Krizhevsky et al., 2012]

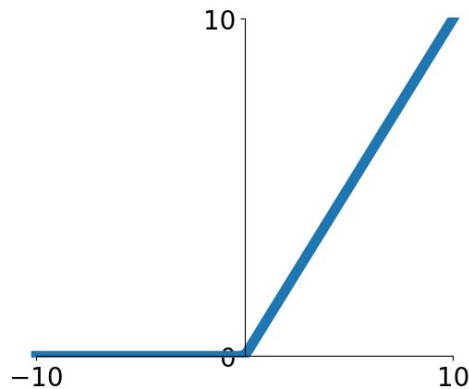
ReLU



ReLU
(Rectified Linear Unit)

- Computes $f(x) = \max(0, x)$
- Does not saturate (in +region)
- Very computationally efficient
- Converges much faster than sigmoid/tanh in practice (e.g. 6x)
- Not zero-centered output

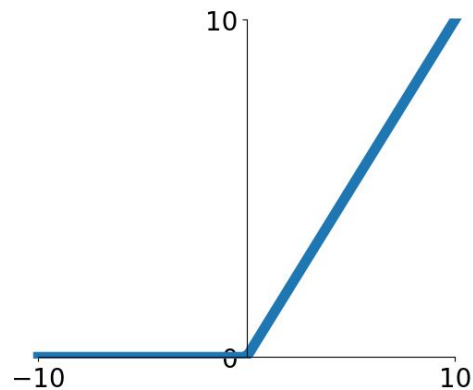
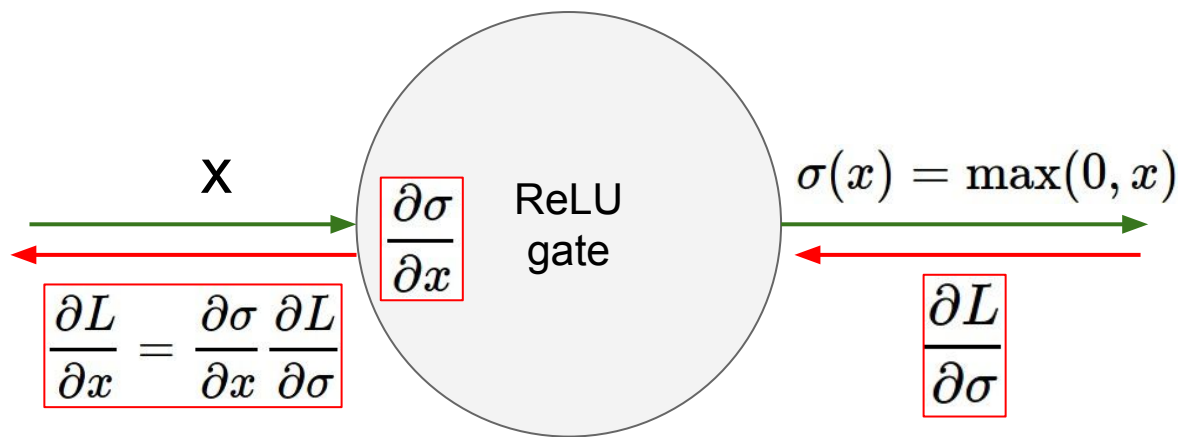
ReLU



ReLU
(Rectified Linear Unit)

- Computes $f(x) = \max(0, x)$
- Does not saturate (in +region)
- Very computationally efficient
- Converges much faster than sigmoid/tanh in practice (e.g. 6x)
- Not zero-centered output
- An annoyance:

hint: what is the gradient when $x < 0$?

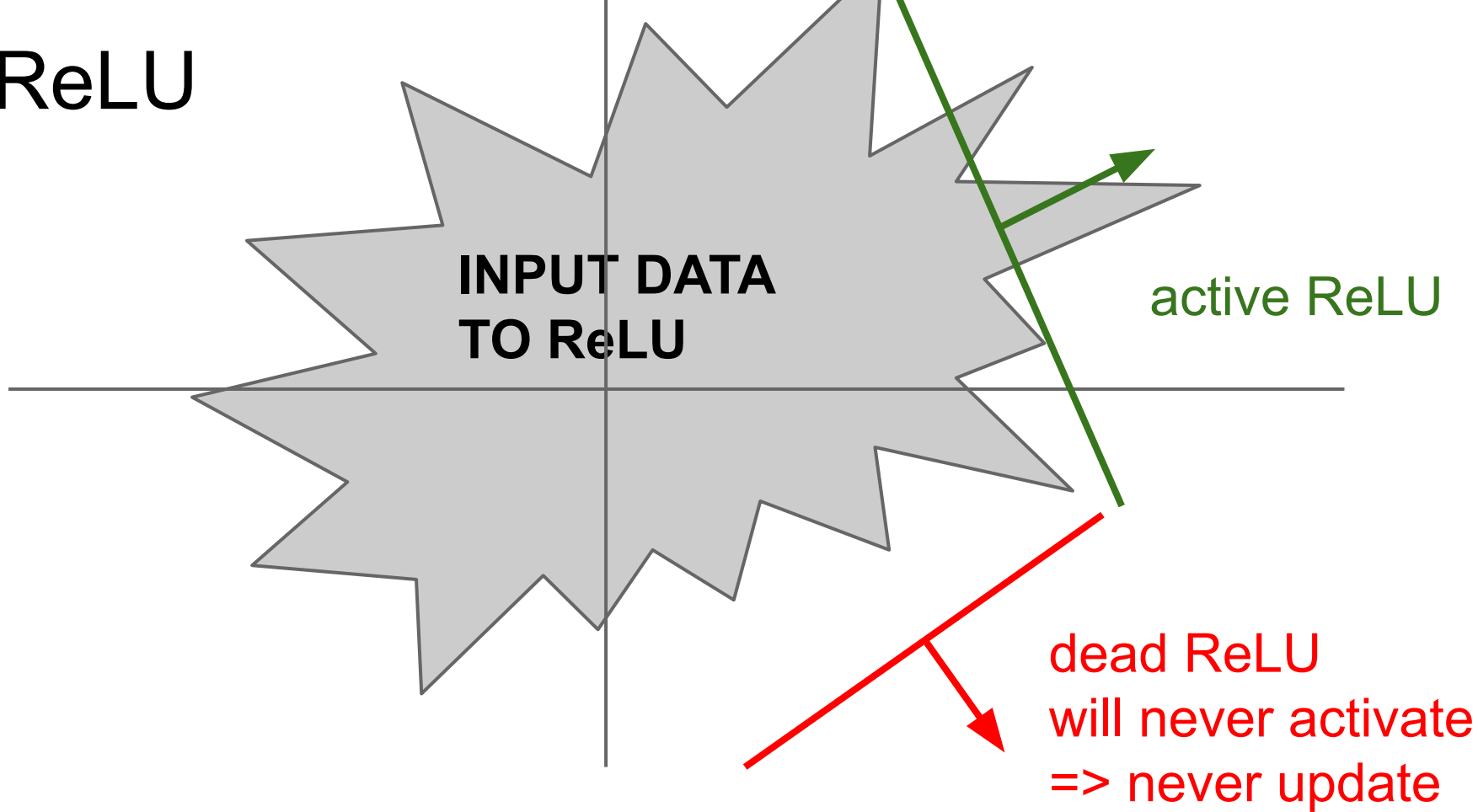


What happens when $x = -10$?

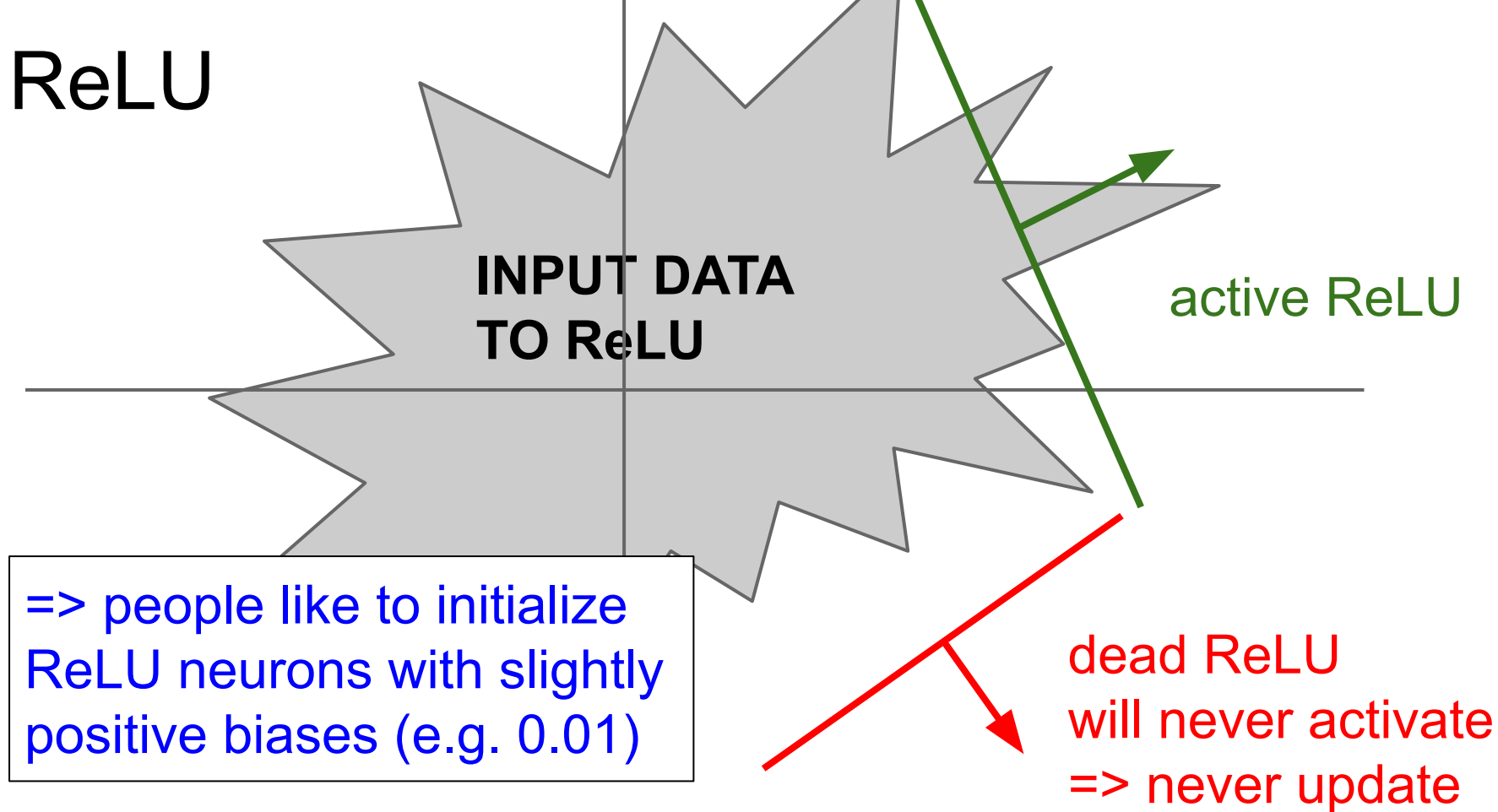
What happens when $x = 0$?

What happens when $x = 10$?

ReLU



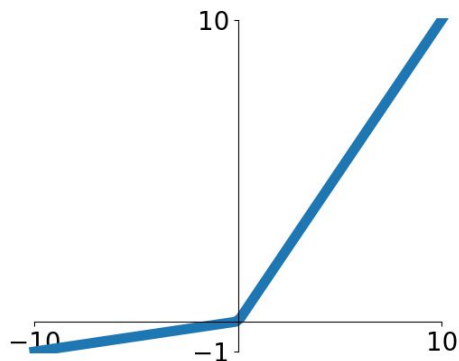
ReLU



Leaky ReLU

[Mass et al., 2013]

[He et al., 2015]



- Does not saturate
- Computationally efficient
- Converges much faster than sigmoid/tanh in practice! (e.g. 6x)
- **will not “die”.**

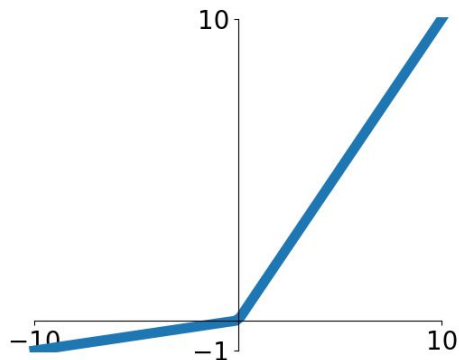
Leaky ReLU

$$f(x) = \max(0.01x, x)$$

Leaky ReLU

[Mass et al., 2013]

[He et al., 2015]



Leaky ReLU

$$f(x) = \max(0.01x, x)$$

- Does not saturate
- Computationally efficient
- Converges much faster than sigmoid/tanh in practice! (e.g. 6x)
- **will not “die”.**

Parametric Rectifier (PReLU)

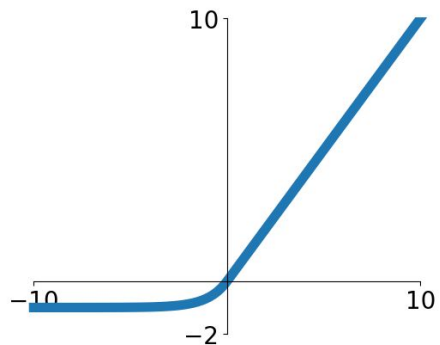
$$f(x) = \max(\alpha x, x)$$

backprop into α
(parameter)

ELU

[Clevert et al., 2015]

Exponential Linear Units (ELU)

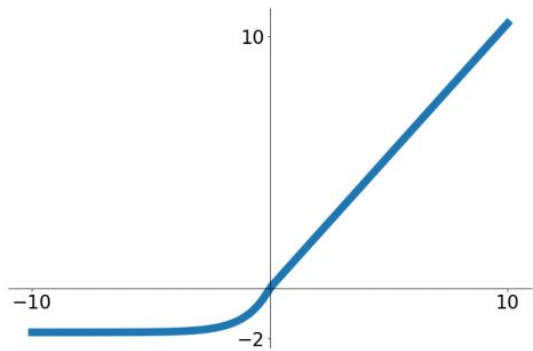


- All benefits of ReLU
- Closer to zero mean outputs
- Computation requires $\exp()$
- Negative saturation can kill gradients for large negative

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

(Alpha default = 1)

Scaled Exponential Linear Units (SELU)



- Scaled version of ELU that works better for deep networks
- “Self-normalizing” property;
- Can train deep SELU networks without BatchNorm
 - (will discuss more later)

$$f(x) = \begin{cases} \lambda x & \text{if } x > 0 \\ \lambda \alpha (e^x - 1) & \text{otherwise} \end{cases}$$

$\alpha = 1.6733, \lambda = 1.0507$

Maxout “Neuron”

[Goodfellow et al., 2013]

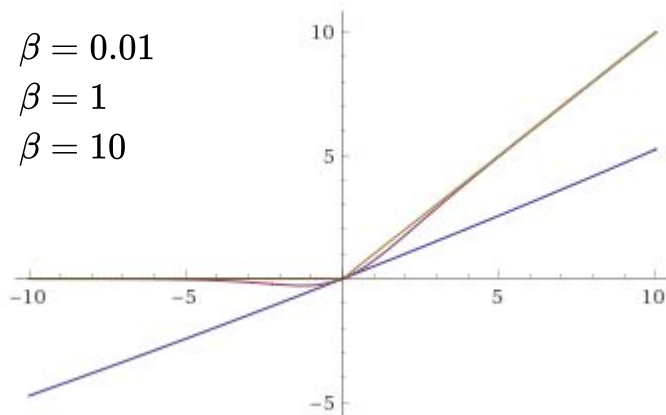
- Does not have the basic form of dot product -> nonlinearity
- Generalizes ReLU and Leaky ReLU
- Linear Regime! Does not saturate! Does not die!

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

Problem: doubles the number of parameters/weights :(

Swish

[Ramachandran et al. 2018]



- They trained a neural network to generate and test out different non-linearities.
- Swish outperformed all other options for CIFAR-10 accuracy

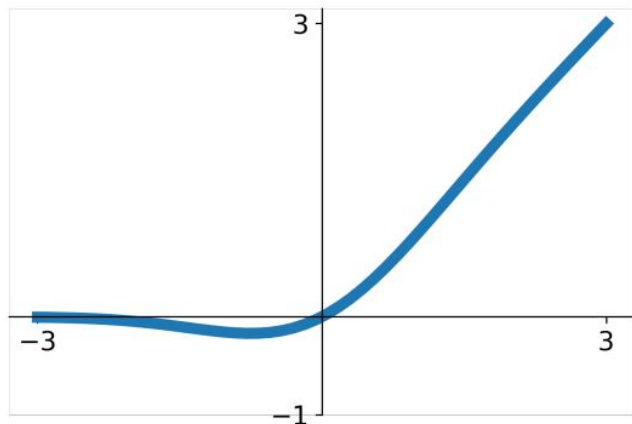
Swish

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

GeLU

[Hendrycks and Gimpel, Gaussian Error Linear Units (GELUs), 2016]

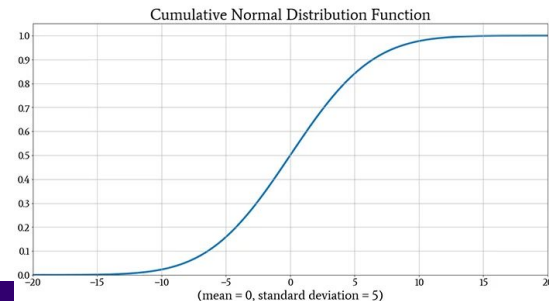
Gaussian Error Linear Units



$X \sim N(0, 1)$

$$\text{gelu}(x) = xP(X \leq x) = \frac{x}{2}(1 + \text{erf}(x/\sqrt{2}))$$
$$\approx x\sigma(1.702x)$$

- Idea: combine the best parts of sigmoid and relu.
- GELU transitions more smoothly
- It weights each input according to the Gaussian (Normal) CDF



GeLU

[Hendrycks and Gimpel, Gaussian Error Linear Units (GELUs), 2016]

Mathematically, $\text{GELU}(x) = x \times \Phi(x)$

where $\Phi(x)$ is the cumulative distribution function (CDF) of a standard normal distribution $N(0,1)$

$$\Phi(x) = \frac{1}{2} \left[1 + \text{erf} \left(\frac{x}{\sqrt{2}} \right) \right]$$

$$\text{GELU}(x) \approx 0.5 x \left[1 + \tanh \left(\sqrt{\frac{2}{\pi}} (x + 0.044715 x^3) \right) \right]$$

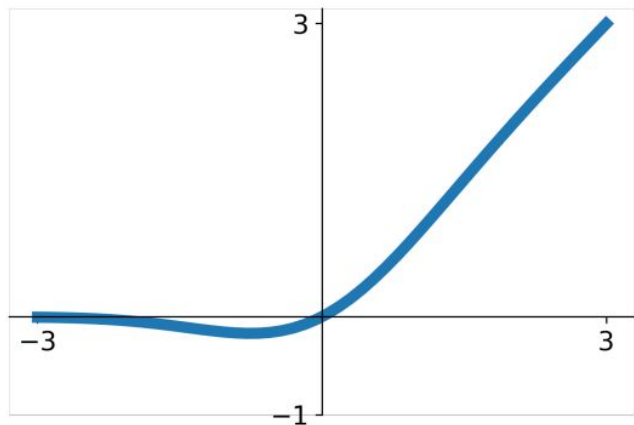
It is approximated as



Activation Functions

[Hendrycks and Gimpel, Gaussian Error Linear Units (GELUs), 2016]

GeLU



$X \sim N(0, 1)$

$$\begin{aligned} \text{gelu}(x) &= xP(X \leq x) = \frac{x}{2}(1 + \text{erf}(x/\sqrt{2})) \\ &\approx x\sigma(1.702x) \end{aligned}$$

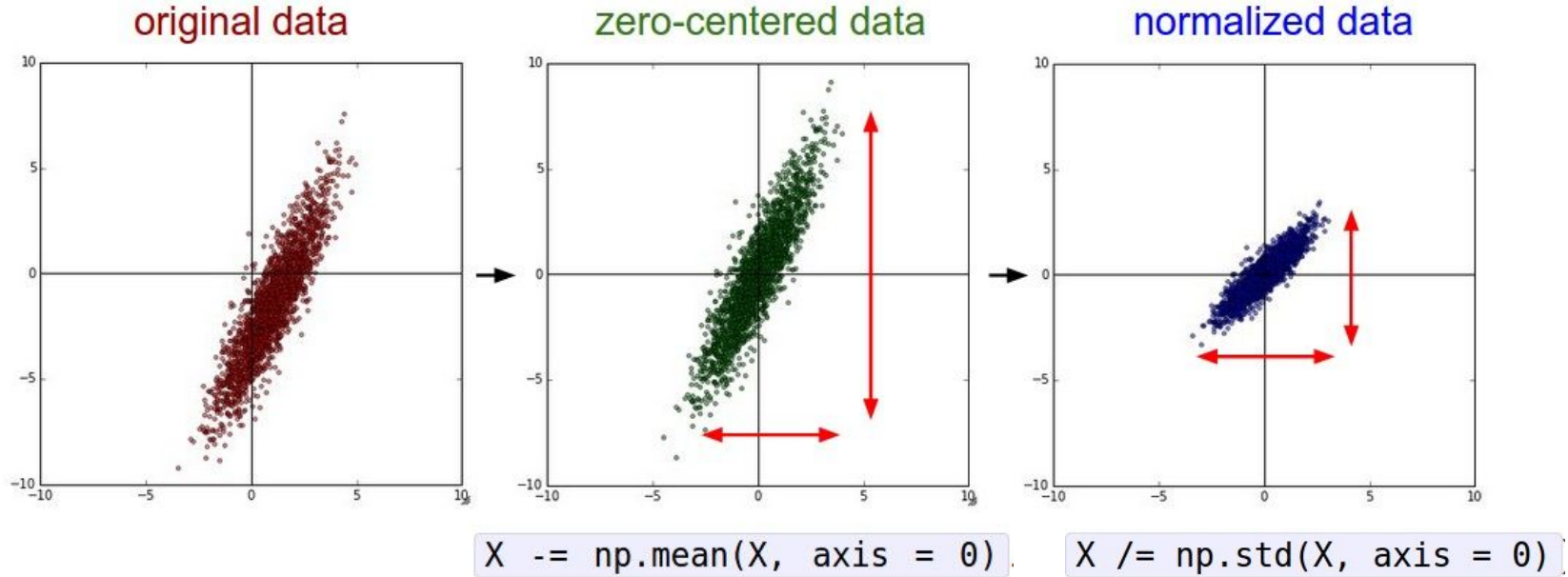
- More continuous version of ReLU. It scales the input x by a probability factor dependent on x
- Avoid dead neurons
- Empirically stable training
- Very common in Transformers (BERT, GPT, ViT)

TLDR: In practice:

- Use **ReLU**. Be careful with your learning rates
- Use **GeLU** when using transformers
- Try out **Leaky ReLU / Maxout / ELU / SELU**
 - To squeeze out some marginal gains
- Don't use **sigmoid** or **tanh**

Data Preprocessing

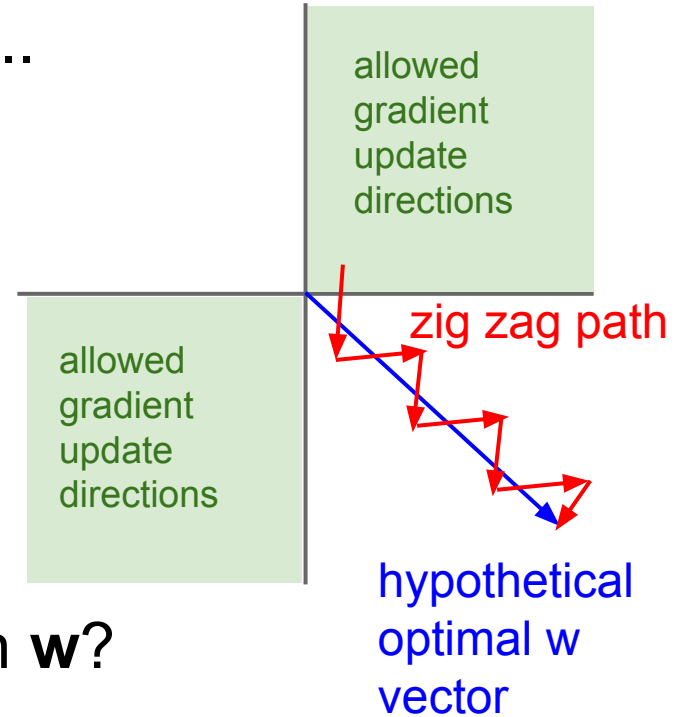
Data Preprocessing



(Assume X [NxD] is data matrix,
each example in a row)

Remember: Consider what happens when the input to a neuron is always positive...

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

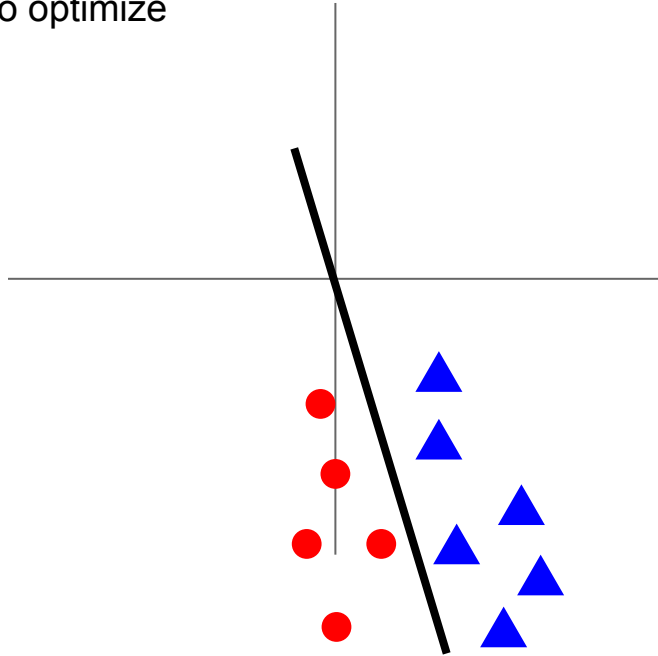


What can we say about the gradients on $\mathbf{w}$?

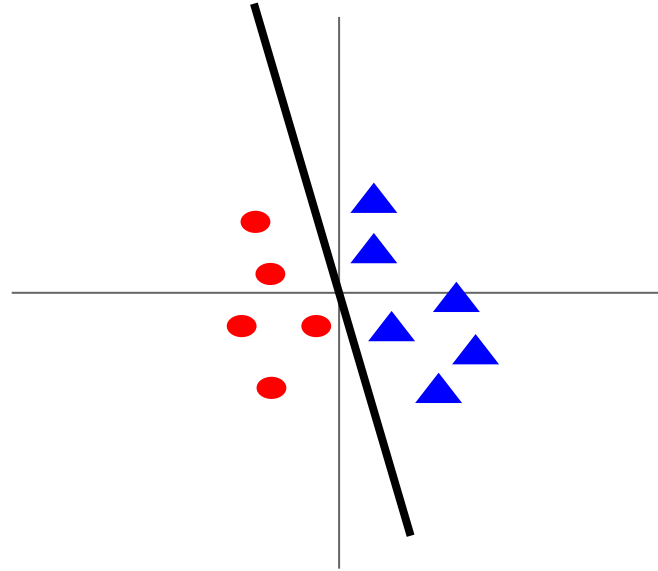
Always all positive or all negative :(
(this is also why you want zero-mean data!)

Data Preprocessing

Before normalization: classification loss very sensitive to changes in weight matrix; hard to optimize

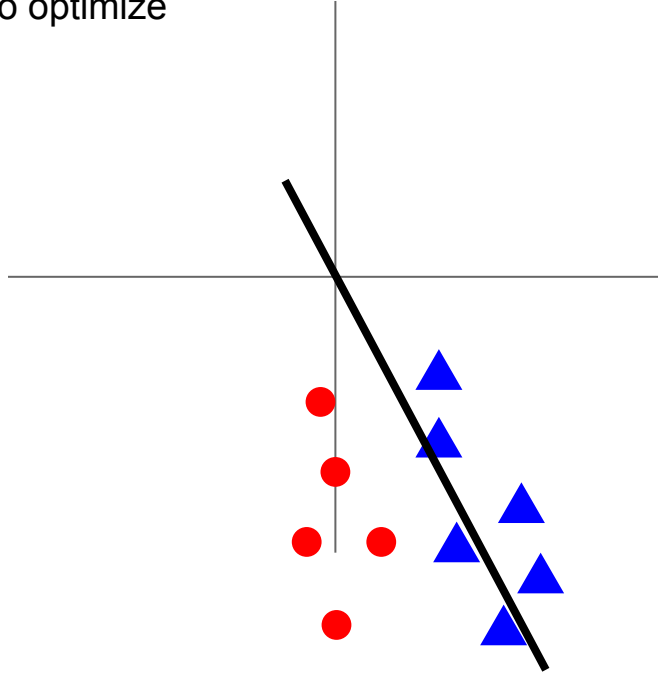


After normalization: less sensitive to small changes in weights; easier to optimize

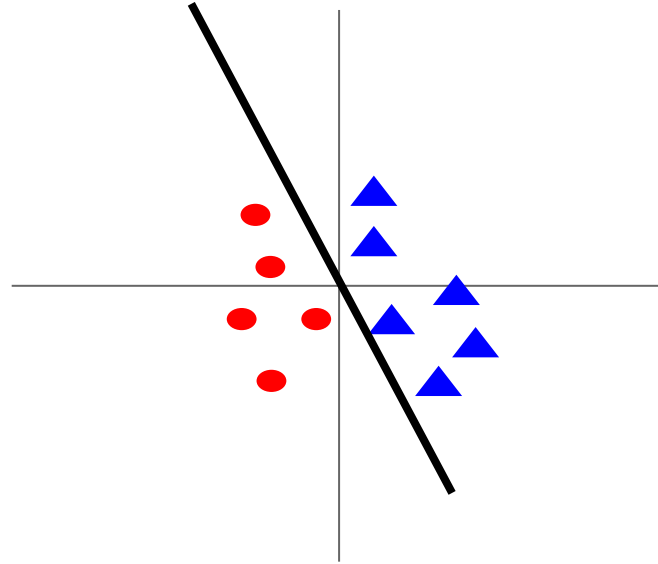


Data Preprocessing

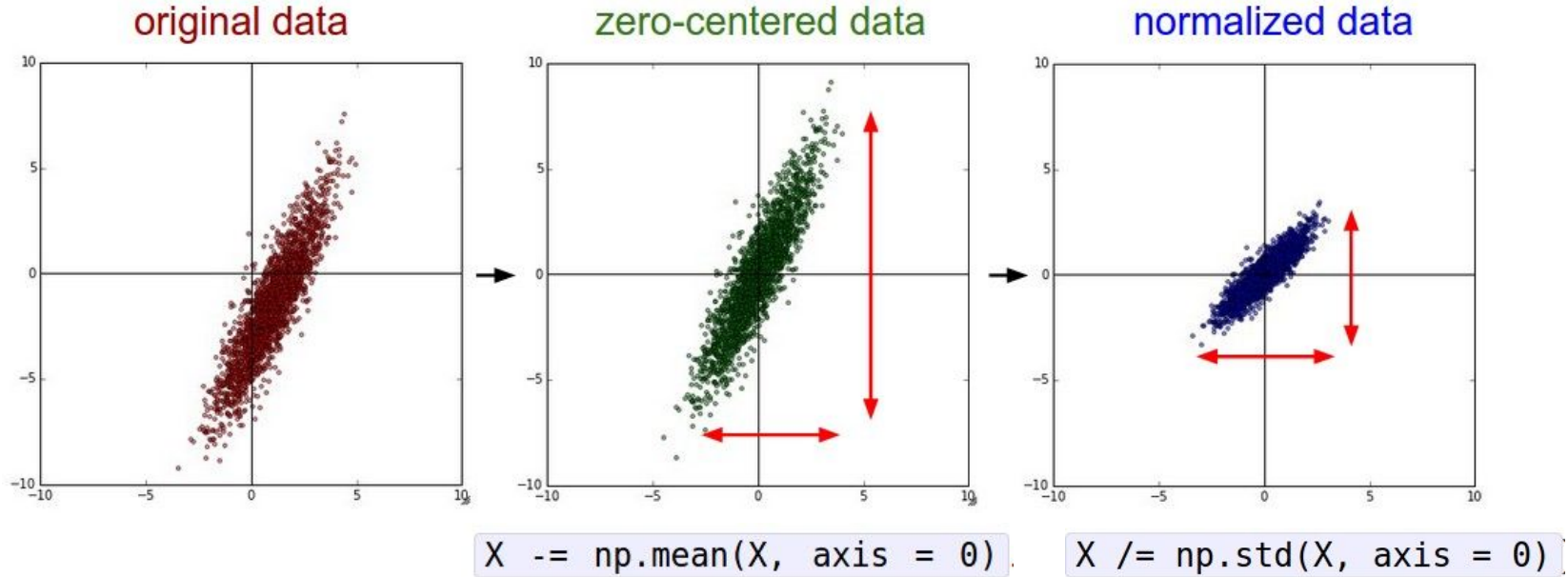
Before normalization: classification loss very sensitive to changes in weight matrix; hard to optimize



After normalization: less sensitive to small changes in weights; easier to optimize



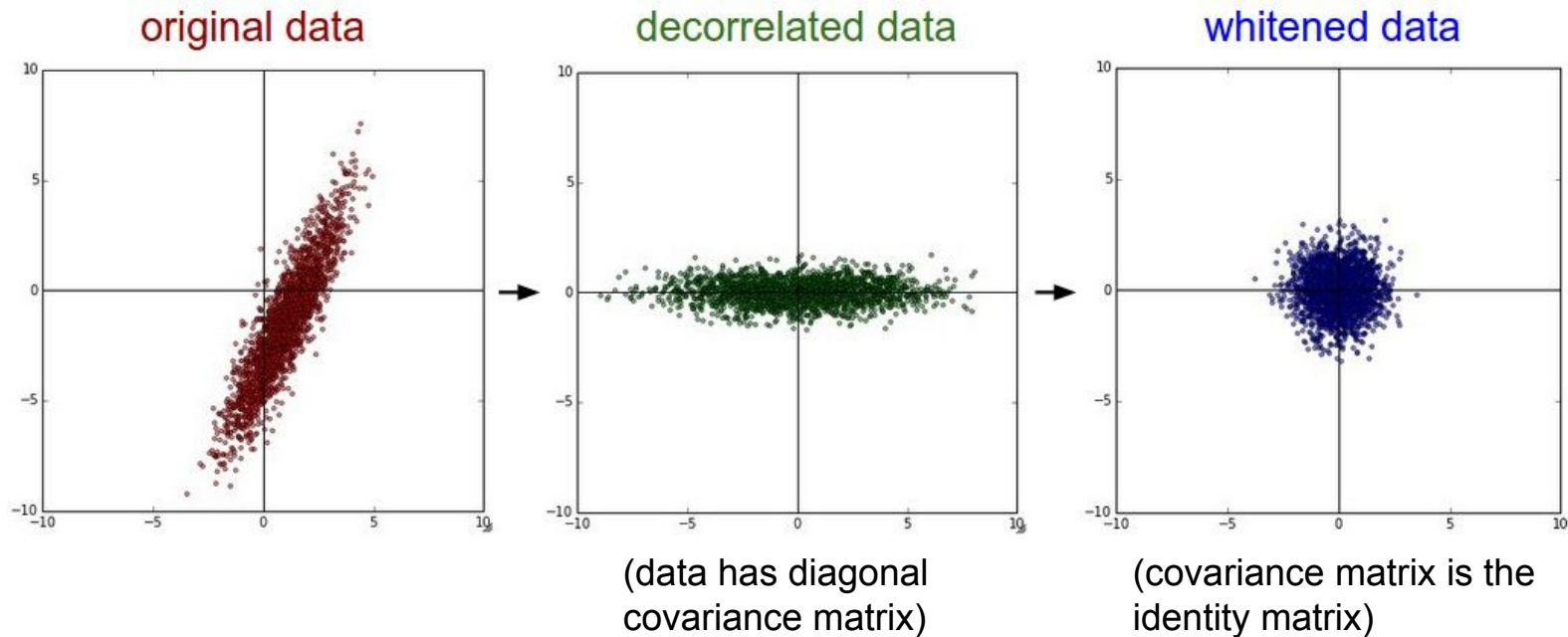
Data Preprocessing



(Assume X [NxD] is data matrix, each example in a row)

Data Preprocessing

In practice, you may also see **PCA** and **Whitening** of the data



TLDR: In practice for Images: center only

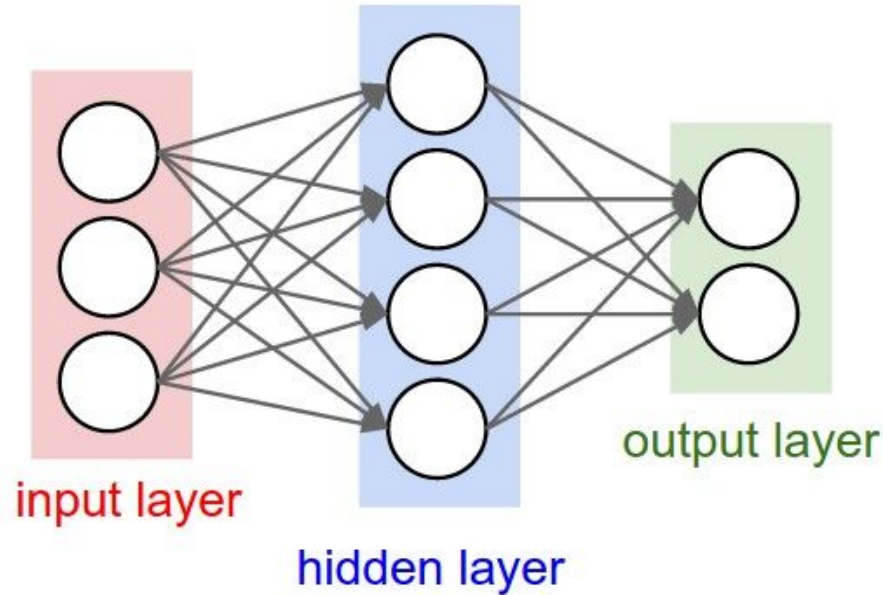
e.g. consider CIFAR-10 example with [32,32,3] images

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

Not common
to do PCA or
whitening

Weight Initialization

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



- First idea: **Small random numbers**
(gaussian with zero mean and $1e-2$ standard deviation)

```
W = 0.01 * np.random.randn(Din, Dout)
```

- First idea: **Small random numbers**
(gaussian with zero mean and $1e-2$ standard deviation)

```
W = 0.01 * np.random.randn(Din, Dout)
```

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

Weight Initialization: Activation statistics

```
dims = [4096] * 7      Forward pass for a 6-layer
hs = []                net with hidden size 4096
x = np.random.randn(16, dims[0])
for Din, Dout in zip(dims[:-1], dims[1:]):
    W = 0.01 * np.random.randn(Din, Dout)
    x = np.tanh(x.dot(W))
    hs.append(x)
```

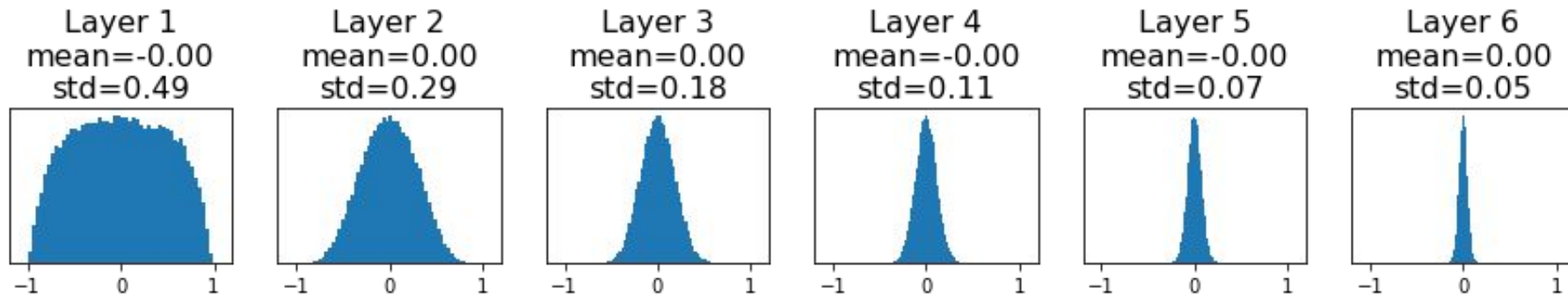
What will happen to the activations for the last layer?

Weight Initialization: Activation statistics

```
dims = [4096] * 7      Forward pass for a 6-layer
hs = []                net with hidden size 4096
x = np.random.randn(16, dims[0])
for Din, Dout in zip(dims[:-1], dims[1:]):
    W = 0.01 * np.random.randn(Din, Dout)
    x = np.tanh(x.dot(W))
    hs.append(x)
```

All activations tend to zero for deeper network layers

Q: What do the gradients dL/dW look like?



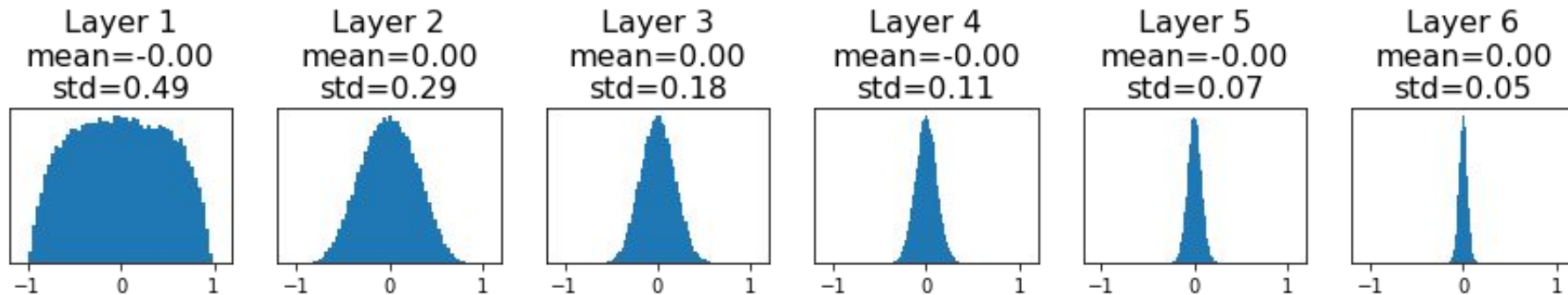
Weight Initialization: Activation statistics

```
dims = [4096] * 7      Forward pass for a 6-layer
hs = []                net with hidden size 4096
x = np.random.randn(16, dims[0])
for Din, Dout in zip(dims[:-1], dims[1:]):
    W = 0.01 * np.random.randn(Din, Dout)
    x = np.tanh(x.dot(W))
    hs.append(x)
```

All activations tend to zero for deeper network layers

Q: What do the gradients dL/dW look like?

A: All zero, no learning =(



Weight Initialization: Activation statistics

```
dims = [4096] * 7      Increase std of initial
hs = []                weights from 0.01 to 0.05
x = np.random.randn(16, dims[0])
for Din, Dout in zip(dims[:-1], dims[1:]):
    W = 0.05 * np.random.randn(Din, Dout)
    x = np.tanh(x.dot(W))
    hs.append(x)
```

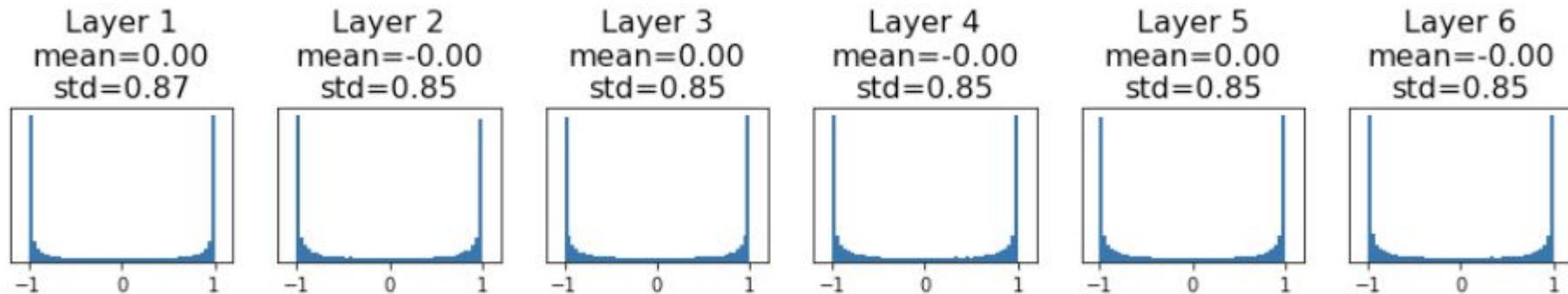
What will happen to the activations when the weights are initialized with larger values?

Weight Initialization: Activation statistics

```
dims = [4096] * 7      Increase std of initial
hs = []                weights from 0.01 to 0.05
x = np.random.randn(16, dims[0])
for Din, Dout in zip(dims[:-1], dims[1:]):
    W = 0.05 * np.random.randn(Din, Dout)
    x = np.tanh(x.dot(W))
    hs.append(x)
```

All activations saturate

Q: What do the gradients look like?



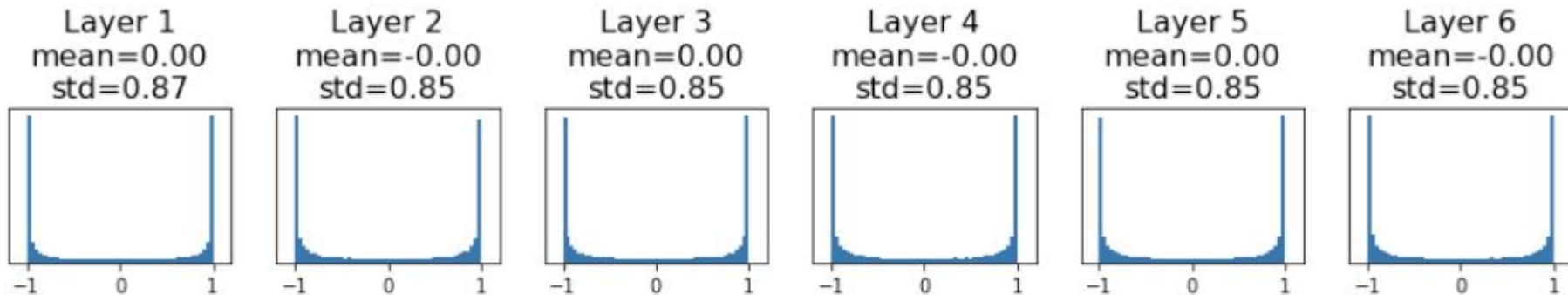
Weight Initialization: Activation statistics

```
dims = [4096] * 7      Increase std of initial
hs = []                weights from 0.01 to 0.05
x = np.random.randn(16, dims[0])
for Din, Dout in zip(dims[:-1], dims[1:]):
    W = 0.05 * np.random.randn(Din, Dout)
    x = np.tanh(x.dot(W))
    hs.append(x)
```

All activations saturate

Q: What do the gradients look like?

A: Local gradients all zero, no learning =(



Weight Initialization: “Xavier” Initialization

```
dims = [4096] * 7          "Xavier" initialization:
hs = []                    std = 1/sqrt(Din)
x = np.random.randn(16, dims[0])
for Din, Dout in zip(dims[:-1], dims[1:]):
    W = np.random.randn(Din, Dout) / np.sqrt(Din)
    x = np.tanh(x.dot(W))
    hs.append(x)
```

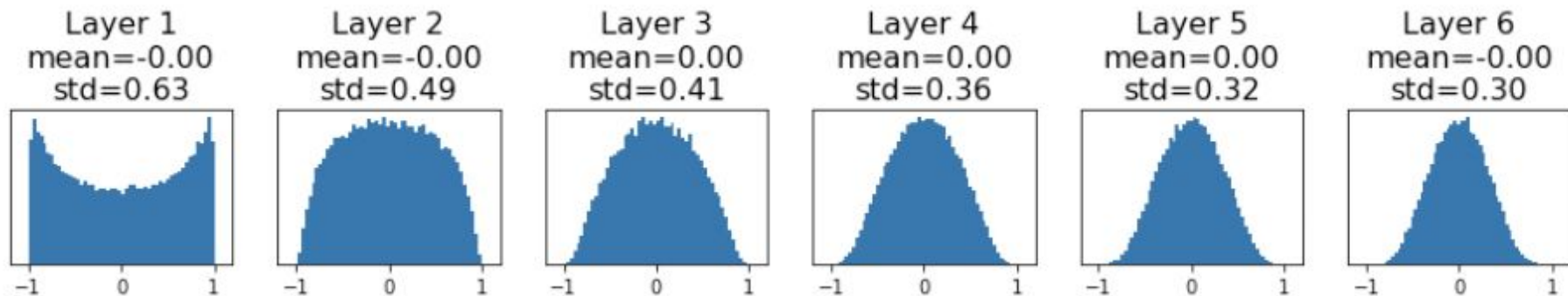
Glorot and Bengio, “Understanding the difficulty of training deep feedforward neural networks”, AISTAT 2010

Weight Initialization: “Xavier” Initialization

```
dims = [4096] * 7
hs = []
x = np.random.randn(16, dims[0])
for Din, Dout in zip(dims[:-1], dims[1:]):
    W = np.random.randn(Din, Dout) / np.sqrt(Din)
    x = np.tanh(x.dot(W))
    hs.append(x)
```

“Xavier” initialization:
 $\text{std} = 1/\sqrt{\text{Din}}$

“Just right”: Activations are nicely scaled for all layers!



Glorot and Bengio, “Understanding the difficulty of training deep feedforward neural networks”, AISTAT 2010

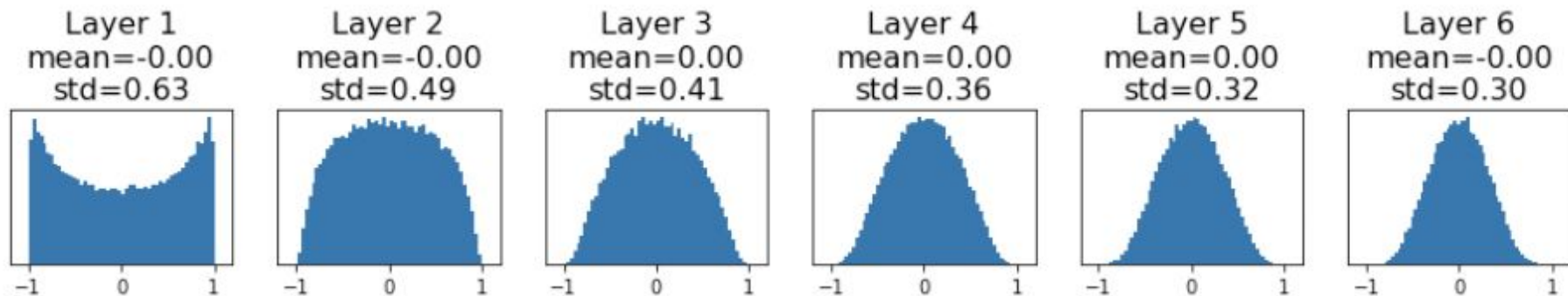
Weight Initialization: “Xavier” Initialization

```
dims = [4096] * 7
hs = []
x = np.random.randn(16, dims[0])
for Din, Dout in zip(dims[:-1], dims[1:]):
    W = np.random.randn(Din, Dout) / np.sqrt(Din)
    x = np.tanh(x.dot(W))
    hs.append(x)
```

“Xavier” initialization:
 $\text{std} = 1/\sqrt{\text{Din}}$

“Just right”: Activations are nicely scaled for all layers!

For conv layers, Din is $\text{filter_size}^2 * \text{input_channels}$



Glorot and Bengio, “Understanding the difficulty of training deep feedforward neural networks”, AISTAT 2010

Weight Initialization: “Xavier” Initialization

```
dims = [4096] * 7
hs = []
x = np.random.randn(16, dims[0])
for Din, Dout in zip(dims[:-1], dims[1:]):
    W = np.random.randn(Din, Dout) / np.sqrt(Din)
    x = np.tanh(x.dot(W))
    hs.append(x)
```

“Xavier” initialization:
std = 1/sqrt(Din)

“Just right”: Activations are nicely scaled for all layers!

For conv layers, Din is $\text{filter_size}^2 * \text{input_channels}$

Let: $y = x_1 w_1 + x_2 w_2 + \dots + x_{D_{in}} w_{D_{in}}$

Glorot and Bengio, “Understanding the difficulty of training deep feedforward neural networks”, AISTAT 2010

Weight Initialization: “Xavier” Initialization

```
dims = [4096] * 7
hs = []
x = np.random.randn(16, dims[0])
for Din, Dout in zip(dims[:-1], dims[1:]):
    W = np.random.randn(Din, Dout) / np.sqrt(Din)
    x = np.tanh(x.dot(W))
    hs.append(x)
```

“Xavier” initialization:
std = $1/\sqrt{D_{in}}$

“Just right”: Activations are nicely scaled for all layers!

For conv layers, D_{in} is $\text{filter_size}^2 * \text{input_channels}$

Let: $y = x_1 w_1 + x_2 w_2 + \dots + x_{D_{in}} w_{D_{in}}$

Assume: $\text{Var}(x_1) = \text{Var}(x_2) = \dots = \text{Var}(x_{D_{in}})$

Glorot and Bengio, “Understanding the difficulty of training deep feedforward neural networks”, AISTAT 2010

Weight Initialization: “Xavier” Initialization

```
dims = [4096] * 7
hs = []
x = np.random.randn(16, dims[0])
for Din, Dout in zip(dims[:-1], dims[1:]):
    W = np.random.randn(Din, Dout) / np.sqrt(Din)
    x = np.tanh(x.dot(W))
    hs.append(x)
```

“Xavier” initialization:
std = $1/\sqrt{D_{in}}$

“Just right”: Activations are nicely scaled for all layers!

For conv layers, D_{in} is $\text{filter_size}^2 * \text{input_channels}$

Let: $y = x_1 w_1 + x_2 w_2 + \dots + x_{D_{in}} w_{D_{in}}$

Assume: $\text{Var}(x_1) = \text{Var}(x_2) = \dots = \text{Var}(x_{D_{in}})$

We want: $\text{Var}(y) = \text{Var}(x_i)$

Glorot and Bengio, “Understanding the difficulty of training deep feedforward neural networks”, AISTAT 2010

Weight Initialization: “Xavier” Initialization

```
dims = [4096] * 7
hs = []
x = np.random.randn(16, dims[0])
for Din, Dout in zip(dims[:-1], dims[1:]):
    W = np.random.randn(Din, Dout) / np.sqrt(Din)
    x = np.tanh(x.dot(W))
    hs.append(x)
```

“Xavier” initialization:
std = $1/\sqrt{D_{in}}$

“Just right”: Activations are nicely scaled for all layers!

For conv layers, D_{in} is $\text{filter_size}^2 * \text{input_channels}$

Let: $y = x_1 w_1 + x_2 w_2 + \dots + x_{D_{in}} w_{D_{in}}$

Assume: $\text{Var}(x_1) = \text{Var}(x_2) = \dots = \text{Var}(x_{D_{in}})$

We want: $\text{Var}(y) = \text{Var}(x_i)$

$\text{Var}(y) = \text{Var}(x_1 w_1 + x_2 w_2 + \dots + x_{D_{in}} w_{D_{in}})$
[substituting value of y]

Weight Initialization: “Xavier” Initialization

```
dims = [4096] * 7
hs = []
x = np.random.randn(16, dims[0])
for Din, Dout in zip(dims[:-1], dims[1:]):
    W = np.random.randn(Din, Dout) / np.sqrt(Din)
    x = np.tanh(x.dot(W))
    hs.append(x)
```

“Xavier” initialization:
std = $1/\sqrt{\text{Din}}$

“Just right”: Activations are nicely scaled for all layers!

For conv layers, Din is $\text{filter_size}^2 * \text{input_channels}$

Let: $y = x_1 w_1 + x_2 w_2 + \dots + x_{\text{Din}} w_{\text{Din}}$

Assume: $\text{Var}(x_1) = \text{Var}(x_2) = \dots = \text{Var}(x_{\text{Din}})$

We want: $\text{Var}(y) = \text{Var}(x_i)$

$$\begin{aligned}\text{Var}(y) &= \text{Var}(x_1 w_1 + x_2 w_2 + \dots + x_{\text{Din}} w_{\text{Din}}) \\ &= \text{Din } \text{Var}(x_i w_i) \\ &[\text{Assume all } x_i, w_i \text{ are iid}]\end{aligned}$$

Weight Initialization: “Xavier” Initialization

```
dims = [4096] * 7
hs = []
x = np.random.randn(16, dims[0])
for Din, Dout in zip(dims[:-1], dims[1:]):
    W = np.random.randn(Din, Dout) / np.sqrt(Din)
    x = np.tanh(x.dot(W))
    hs.append(x)
```

“Xavier” initialization:
std = $1/\sqrt{D_{in}}$

“Just right”: Activations are nicely scaled for all layers!

For conv layers, D_{in} is $\text{filter_size}^2 * \text{input_channels}$

Let: $y = x_1 w_1 + x_2 w_2 + \dots + x_{D_{in}} w_{D_{in}}$

Assume: $\text{Var}(x_1) = \text{Var}(x_2) = \dots = \text{Var}(x_{D_{in}})$

We want: $\text{Var}(y) = \text{Var}(x_i)$

$$\begin{aligned}\text{Var}(y) &= \text{Var}(x_1 w_1 + x_2 w_2 + \dots + x_{D_{in}} w_{D_{in}}) \\ &= D_{in} \text{Var}(x_i w_i) \\ &= D_{in} \text{Var}(x_i) \text{Var}(w_i)\end{aligned}$$

[Assume all x_i , w_i are zero mean]

Weight Initialization: “Xavier” Initialization

```
dims = [4096] * 7
hs = []
x = np.random.randn(16, dims[0])
for Din, Dout in zip(dims[:-1], dims[1:]):
    W = np.random.randn(Din, Dout) / np.sqrt(Din)
    x = np.tanh(x.dot(W))
    hs.append(x)
```

“Xavier” initialization:
std = $1/\sqrt{\text{Din}}$

“Just right”: Activations are nicely scaled for all layers!

For conv layers, Din is $\text{filter_size}^2 * \text{input_channels}$

Let: $y = x_1 w_1 + x_2 w_2 + \dots + x_{\text{Din}} w_{\text{Din}}$

Assume: $\text{Var}(x_1) = \text{Var}(x_2) = \dots = \text{Var}(x_{\text{Din}})$

We want: $\text{Var}(y) = \text{Var}(x_i)$

$$\begin{aligned}\text{Var}(y) &= \text{Var}(x_1 w_1 + x_2 w_2 + \dots + x_{\text{Din}} w_{\text{Din}}) \\ &= \text{Din } \text{Var}(x_i w_i) \\ &= \text{Din } \text{Var}(x_i) \text{Var}(w_i) \\ &\text{[Assume all } x_i, w_i \text{ are iid]}\end{aligned}$$

So, $\text{Var}(y) = \text{Var}(x_i)$ only when $\text{Var}(w_i) = 1/\text{Din}$

Glorot and Bengio, “Understanding the difficulty of training deep feedforward neural networks”, AISTAT 2010

Weight Initialization: What about ReLU?

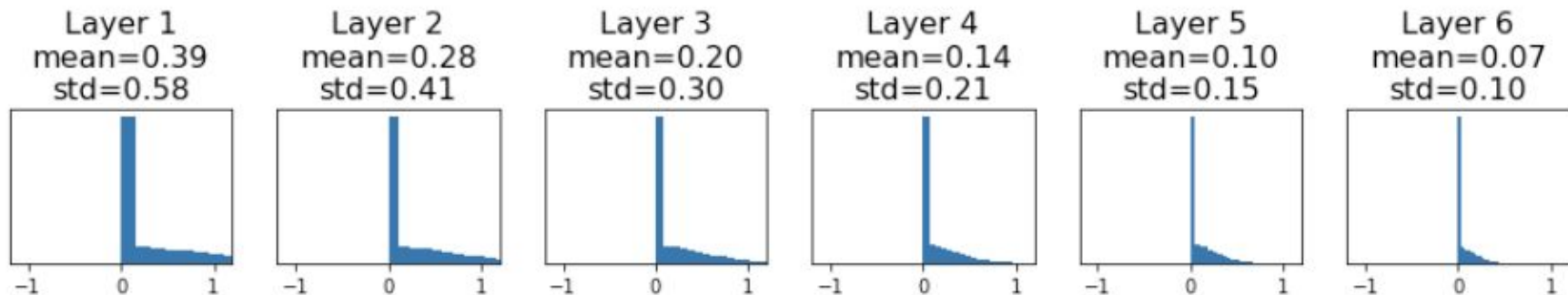
```
dims = [4096] * 7      Change from tanh to ReLU
hs = []
x = np.random.randn(16, dims[0])
for Din, Dout in zip(dims[:-1], dims[1:]):
    W = np.random.randn(Din, Dout) / np.sqrt(Din)
    x = np.maximum(0, x.dot(W))
    hs.append(x)
```

Weight Initialization: What about ReLU?

```
dims = [4096] * 7      Change from tanh to ReLU
hs = []
x = np.random.randn(16, dims[0])
for Din, Dout in zip(dims[:-1], dims[1:]):
    W = np.random.randn(Din, Dout) / np.sqrt(Din)
    x = np.maximum(0, x.dot(W))
    hs.append(x)
```

Xavier assumes zero centered activation function

Activations collapse to zero again, no learning =(

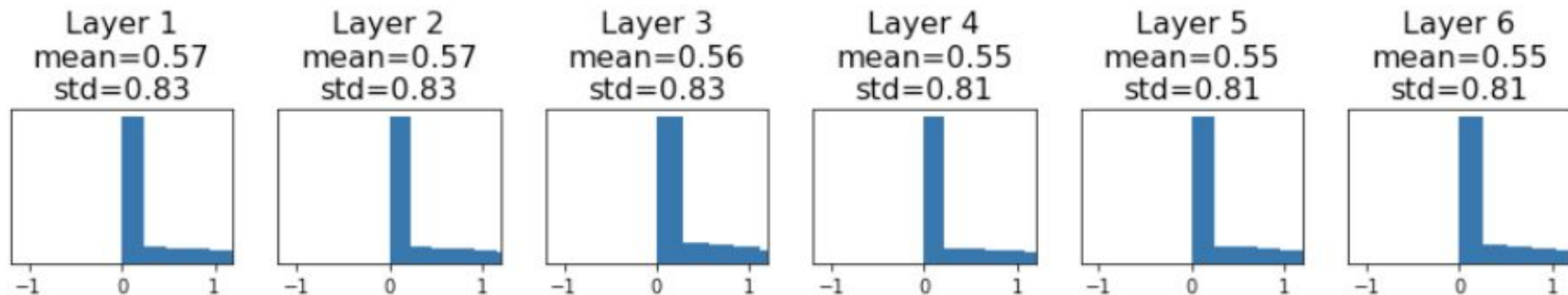


Weight Initialization: Kaiming / MSRA Initialization

```
dims = [4096] * 7
hs = []
x = np.random.randn(16, dims[0])
for Din, Dout in zip(dims[:-1], dims[1:]):
    W = np.random.randn(Din, Dout) * np.sqrt(2/Din)
    x = np.maximum(0, x.dot(W))
    hs.append(x)
```

ReLU correction: $\text{std} = \sqrt{2 / \text{Din}}$

“Just right”: Activations are nicely scaled for all layers!



He et al, “Delving Deep into Rectifiers: Surpassing Human-Level Performance on ImageNet Classification”, ICCV 2015

Proper initialization is an active area of research...

Understanding the difficulty of training deep feedforward neural networks

by Glorot and Bengio, 2010

Exact solutions to the nonlinear dynamics of learning in deep linear neural networks by Saxe et al, 2013

Random walk initialization for training very deep feedforward networks by Sussillo and Abbott, 2014

Delving deep into rectifiers: Surpassing human-level performance on ImageNet classification by He et al., 2015

Data-dependent Initializations of Convolutional Neural Networks by Krähenbühl et al., 2015

All you need is a good init, Mishkin and Matas, 2015

Fixup Initialization: Residual Learning Without Normalization, Zhang et al, 2019

The Lottery Ticket Hypothesis: Finding Sparse, Trainable Neural Networks, Frankle and Carbin, 2019

Batch Normalization

Batch Normalization

[Ioffe and Szegedy, 2015]

“you want zero-mean unit-variance activations? just make them so.”

consider a batch of activations at some layer. To make each dimension zero-mean unit-variance, apply:

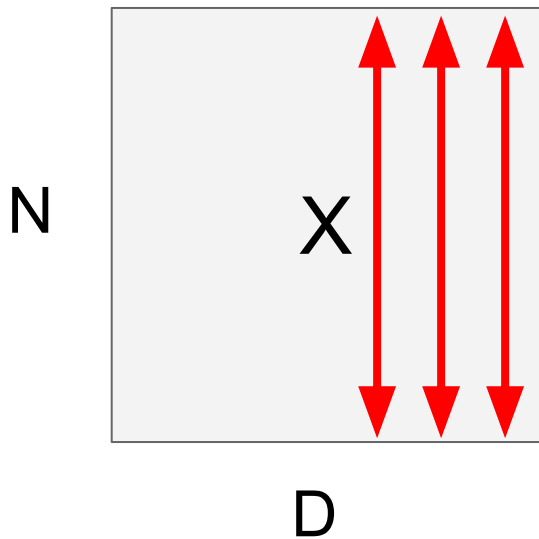
$$\hat{x}^{(k)} = \frac{x^{(k)} - \mathbb{E}[x^{(k)}]}{\sqrt{\text{Var}[x^{(k)}]}}$$

this is a vanilla
differentiable function...

Batch Normalization

[Ioffe and Szegedy, 2015]

Input: $x : N \times D$



$$\mu_j = \frac{1}{N} \sum_{i=1}^N x_{i,j}$$

Per-channel mean,
shape is D

$$\sigma_j^2 = \frac{1}{N} \sum_{i=1}^N (x_{i,j} - \mu_j)^2$$

Per-channel var,
shape is D

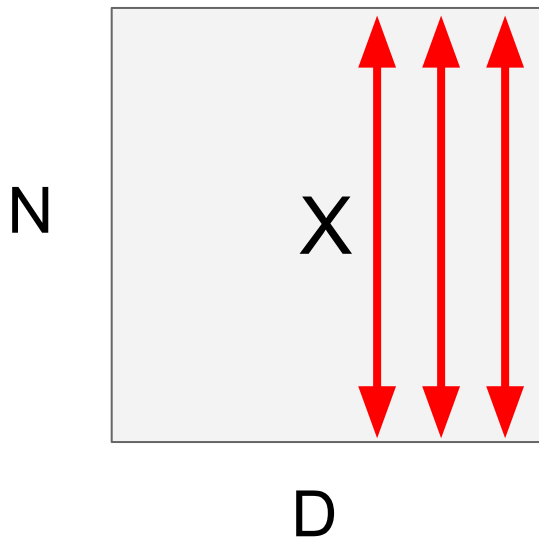
$$\hat{x}_{i,j} = \frac{x_{i,j} - \mu_j}{\sqrt{\sigma_j^2 + \epsilon}}$$

Normalized x,
Shape is N x D

Batch Normalization

[Ioffe and Szegedy, 2015]

Input: $x : N \times D$



$$\mu_j = \frac{1}{N} \sum_{i=1}^N x_{i,j}$$

Per-channel mean,
shape is D

$$\sigma_j^2 = \frac{1}{N} \sum_{i=1}^N (x_{i,j} - \mu_j)^2$$

Per-channel var,
shape is D

$$\hat{x}_{i,j} = \frac{x_{i,j} - \mu_j}{\sqrt{\sigma_j^2 + \epsilon}}$$

Normalized x,
Shape is N x D

Problem: What if zero-mean, unit
variance is too hard a constraint?

Batch Normalization

[Ioffe and Szegedy, 2015]

Input: $x : N \times D$

Learnable scale and shift parameters:

$$\gamma, \beta : D$$

Learning $\gamma = \sigma$,
 $\beta = \mu$ will recover the
identity function!

$$\mu_j = \frac{1}{N} \sum_{i=1}^N x_{i,j}$$

Per-channel mean,
shape is D

$$\sigma_j^2 = \frac{1}{N} \sum_{i=1}^N (x_{i,j} - \mu_j)^2$$

Per-channel var,
shape is D

$$\hat{x}_{i,j} = \frac{x_{i,j} - \mu_j}{\sqrt{\sigma_j^2 + \epsilon}}$$

Normalized x,
Shape is N x D

$$y_{i,j} = \gamma_j \hat{x}_{i,j} + \beta_j$$

Output,
Shape is N x D

Batch Normalization: Test-Time

Estimates depend on minibatch;
can't do this at test-time!

Input: $x : N \times D$

Learnable scale and shift parameters:

$$\gamma, \beta : D$$

Learning $\gamma = \sigma$,
 $\beta = \mu$ will recover the
identity function!

$$\mu_j = \frac{1}{N} \sum_{i=1}^N x_{i,j} \quad \text{Per-channel mean, shape is D}$$
$$\sigma_j^2 = \frac{1}{N} \sum_{i=1}^N (x_{i,j} - \mu_j)^2 \quad \text{Per-channel var, shape is D}$$

$$\hat{x}_{i,j} = \frac{x_{i,j} - \mu_j}{\sqrt{\sigma_j^2 + \epsilon}} \quad \text{Normalized x, Shape is N x D}$$

$$y_{i,j} = \gamma_j \hat{x}_{i,j} + \beta_j \quad \text{Output, Shape is N x D}$$

Batch Normalization: Test-Time

Input: $x : N \times D$

$$\mu_j = \text{(Running) average of values seen during training}$$

Per-channel mean,
shape is D

Learnable scale and shift parameters:

$$\gamma, \beta : D$$

$$\sigma_j^2 = \text{(Running) average of values seen during training}$$

Per-channel var,
shape is D

During testing batchnorm becomes a linear operator!
Can be fused with the previous fully-connected or conv layer

$$\hat{x}_{i,j} = \frac{x_{i,j} - \mu_j}{\sqrt{\sigma_j^2 + \epsilon}}$$

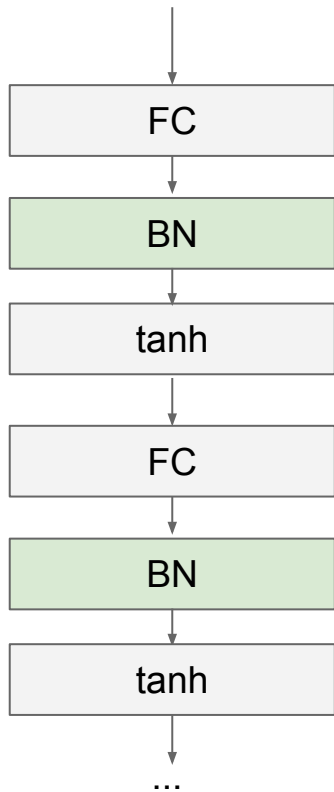
Normalized x,
Shape is N x D

$$y_{i,j} = \gamma_j \hat{x}_{i,j} + \beta_j$$

Output,
Shape is N x D

Batch Normalization

[Ioffe and Szegedy, 2015]

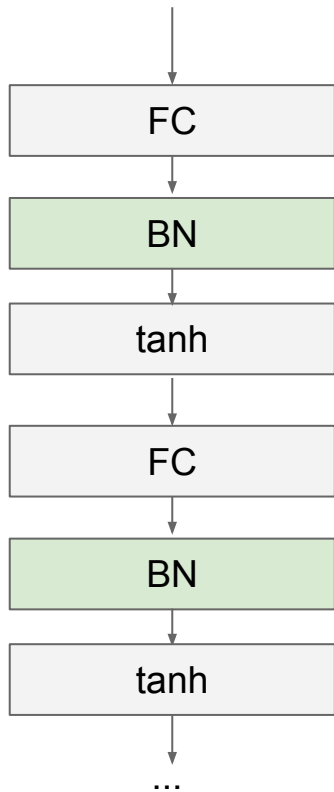


Usually inserted after Fully Connected or Convolutional layers, and before nonlinearity.

$$\hat{x}^{(k)} = \frac{x^{(k)} - \mathbb{E}[x^{(k)}]}{\sqrt{\text{Var}[x^{(k)}]}}$$

Batch Normalization

[Ioffe and Szegedy, 2015]



- Makes deep networks **much** easier to train!
- Improves gradient flow
- Allows higher learning rates, faster convergence
- Networks become more robust to initialization
- Acts as regularization during training
- Zero overhead at test-time: can be fused with conv!
- Behaves differently during training and testing: this is a very common source of debugging!

Batch Normalization for convolutions

Batch Normalization for
fully-connected networks

$$\mathbf{x}: \mathbf{N} \times \mathbf{D}$$

Normalize 

$$\boldsymbol{\mu}, \boldsymbol{\sigma}: 1 \times \mathbf{D}$$

$$\gamma, \beta: 1 \times \mathbf{D}$$

$$\mathbf{y} = \gamma (\mathbf{x} - \boldsymbol{\mu}) / \boldsymbol{\sigma} + \beta$$

Batch Normalization for
convolutional networks
(Spatial Batchnorm, BatchNorm2D)

$$\mathbf{x}: \mathbf{N} \times \mathbf{C} \times \mathbf{H} \times \mathbf{W}$$

Normalize   

$$\boldsymbol{\mu}, \boldsymbol{\sigma}: 1 \times \mathbf{C} \times 1 \times 1$$

$$\gamma, \beta: 1 \times \mathbf{C} \times 1 \times 1$$

$$\mathbf{y} = \gamma (\mathbf{x} - \boldsymbol{\mu}) / \boldsymbol{\sigma} + \beta$$

Layer Normalization

Layer Normalization for MLPs

Batch Normalization for
fully-connected networks

$$\mathbf{x}: \mathbf{N} \times \mathbf{D}$$

Normalize



$$\boldsymbol{\mu}, \boldsymbol{\sigma}: 1 \times \mathbf{D}$$

$$\boldsymbol{\gamma}, \boldsymbol{\beta}: 1 \times \mathbf{D}$$

$$\mathbf{y} = \boldsymbol{\gamma} (\mathbf{x} - \boldsymbol{\mu}) / \boldsymbol{\sigma} + \boldsymbol{\beta}$$

Layer Normalization for
fully-connected networks
Same behavior at train and test!
Often used in transformers

$$\mathbf{x}: \mathbf{N} \times \mathbf{D}$$

Normalize



$$\boldsymbol{\mu}, \boldsymbol{\sigma}: \mathbf{N} \times 1$$

$$\boldsymbol{\gamma}, \boldsymbol{\beta}: 1 \times \mathbf{D}$$

$$\mathbf{y} = \boldsymbol{\gamma} (\mathbf{x} - \boldsymbol{\mu}) / \boldsymbol{\sigma} + \boldsymbol{\beta}$$

Ba, Kiros, and Hinton, "Layer Normalization", arXiv 2016

Instance Normalization

Instance Normalization for Convolutions

Batch Normalization for
convolutional networks

$\mathbf{x}$: $N \times C \times H \times W$

Normalize



$\boldsymbol{\mu}, \boldsymbol{\sigma}$: $1 \times C \times 1 \times 1$

$\boldsymbol{\gamma}, \boldsymbol{\beta}$: $1 \times C \times 1 \times 1$

$$\mathbf{y} = \boldsymbol{\gamma} (\mathbf{x} - \boldsymbol{\mu}) / \boldsymbol{\sigma} + \boldsymbol{\beta}$$

Instance Normalization for
convolutional networks
Same behavior at train / test!

$\mathbf{x}$: $N \times C \times H \times W$

Normalize



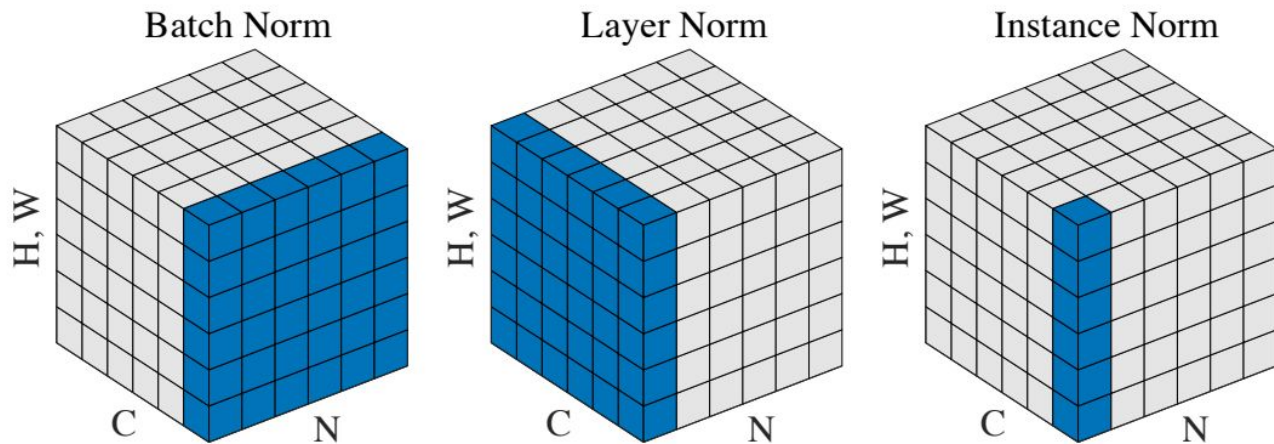
$\boldsymbol{\mu}, \boldsymbol{\sigma}$: $N \times C \times 1 \times 1$

$\boldsymbol{\gamma}, \boldsymbol{\beta}$: $1 \times C \times 1 \times 1$

$$\mathbf{y} = \boldsymbol{\gamma} (\mathbf{x} - \boldsymbol{\mu}) / \boldsymbol{\sigma} + \boldsymbol{\beta}$$

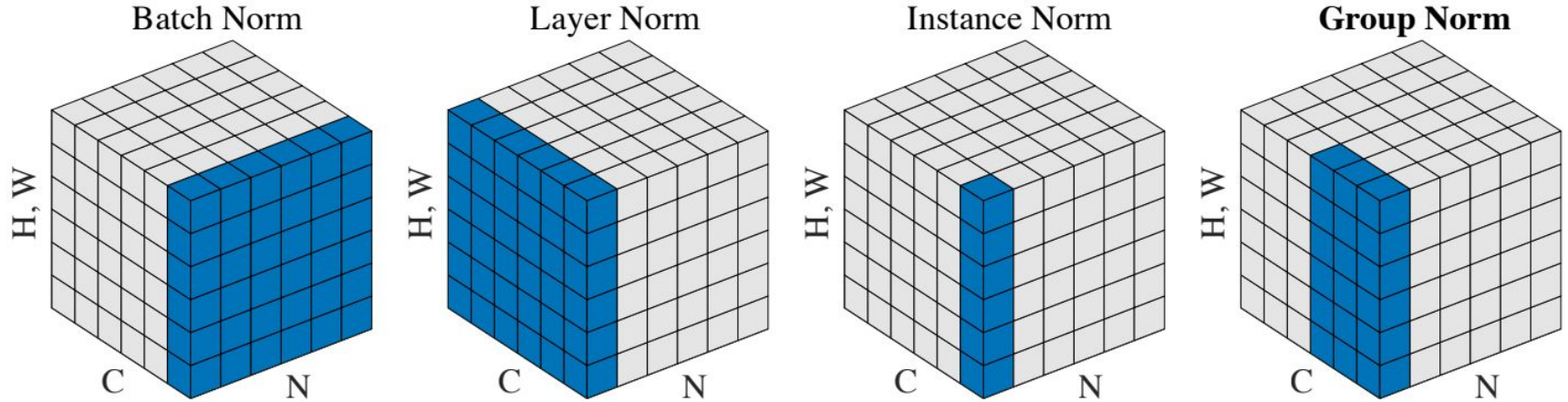
Ulyanov et al, Improved Texture Networks: Maximizing Quality and Diversity in Feed-forward Stylization and Texture Synthesis, CVPR 2017

Comparison of Normalization Layers



Wu and He, "Group Normalization", ECCV 2018

Group Normalization



Wu and He, "Group Normalization", ECCV 2018

Summary

TLDRs

We looked in detail at:

- Activation Functions (use ReLU or GeLU)
- Data Preprocessing (images: subtract mean)
- Weight Initialization (use Xavier/He init)
- Batch Normalization (use this!)
- Layer Normalization (used in transformers!)

Next time: Optimizers

- Parameter update schemes
- Learning rate schedules
- Gradient checking
- Regularization (Dropout etc.)
- Babysitting learning
- Evaluation (Ensembles etc.)
- Hyperparameter Optimization