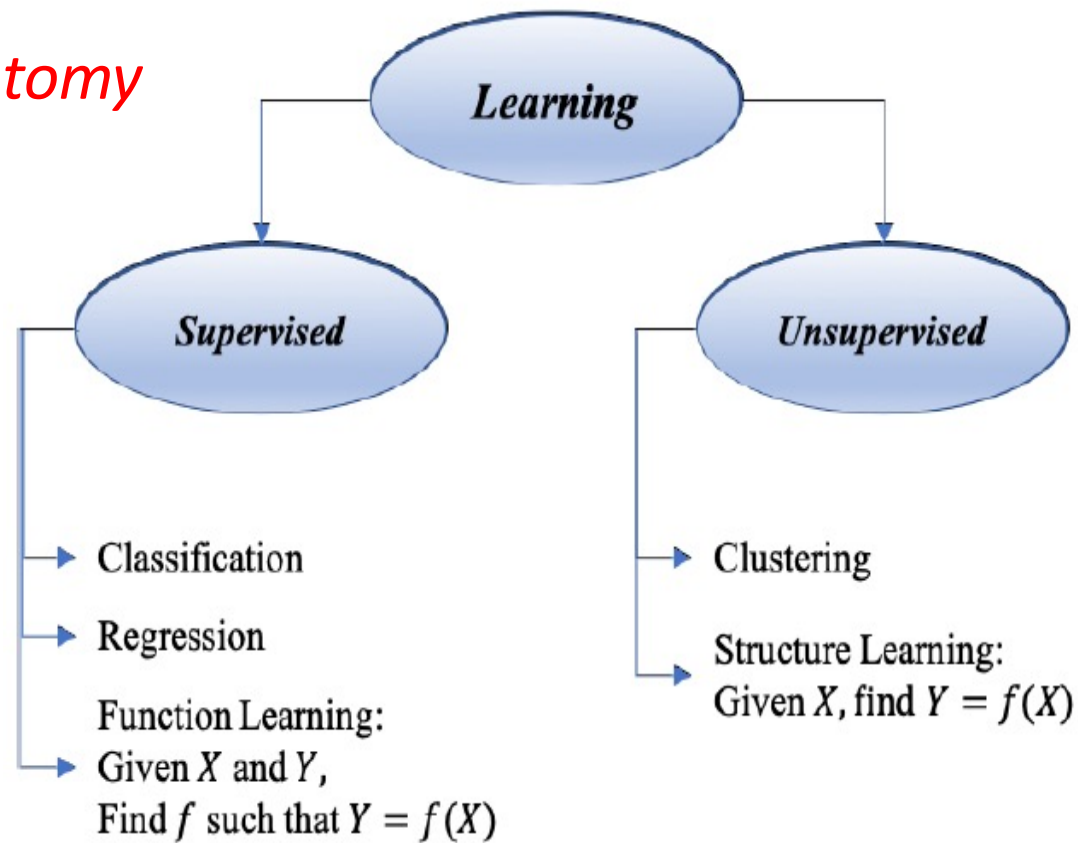


Learning Dichotomy

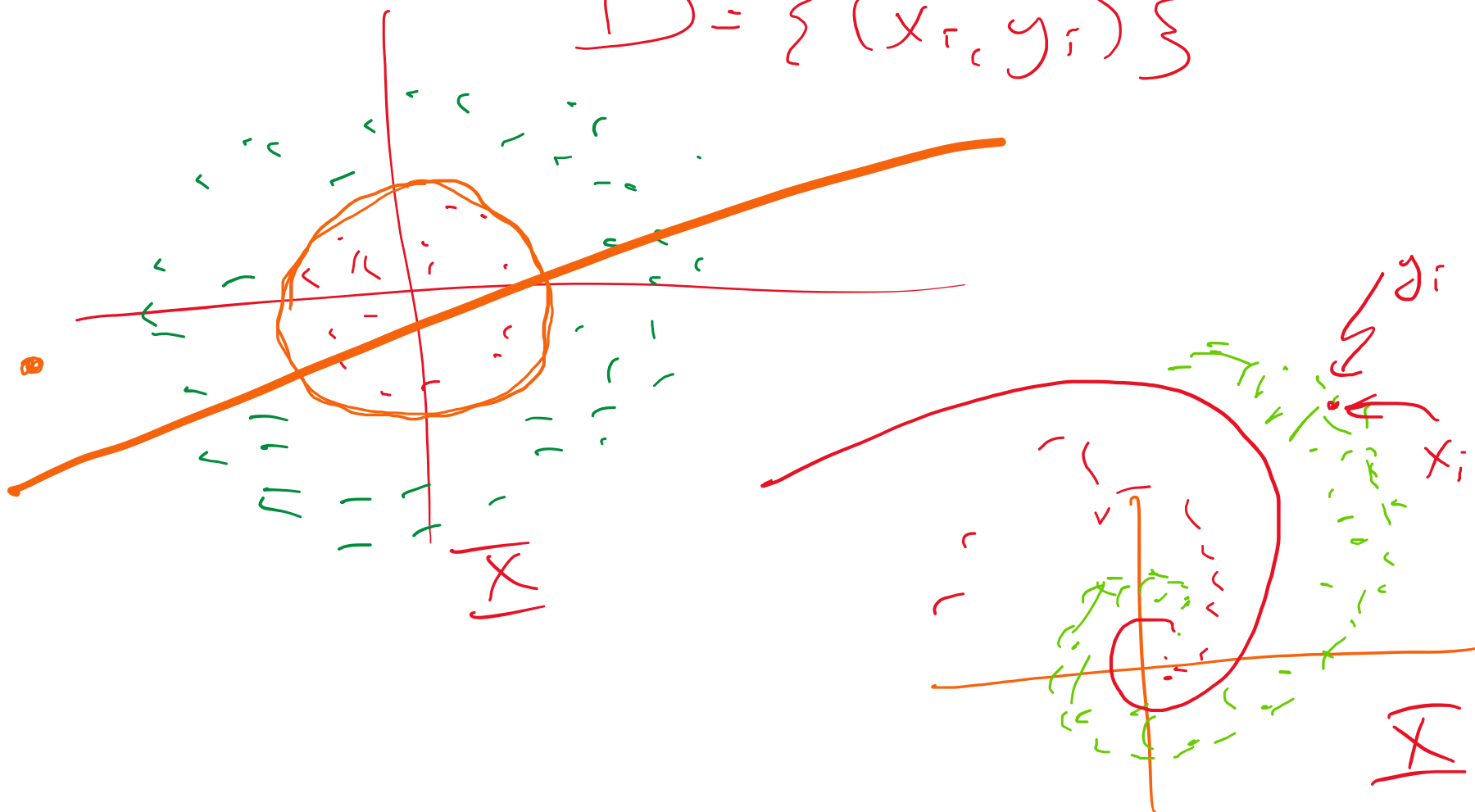


Labelled data
Learn a function

Pattern
Structure

Well...at least one more
- re-enforcement learning

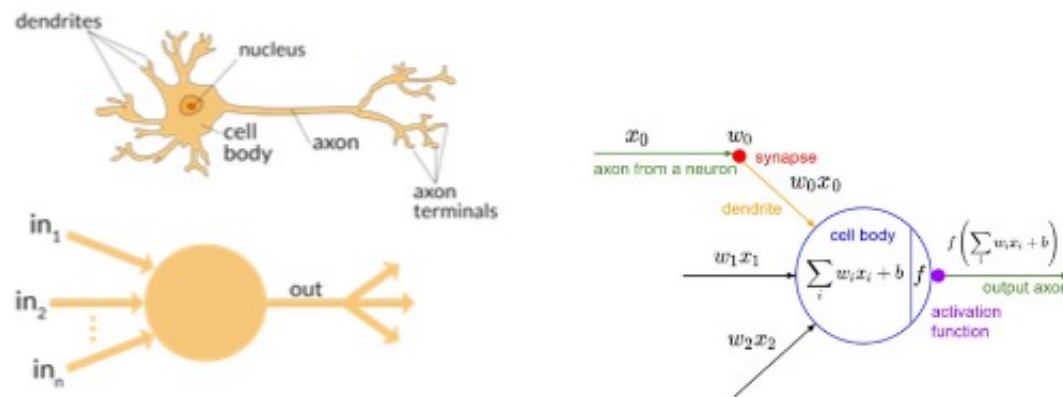
$$D = \{ (x_i, y_i) \}$$



And now for something completely different

Artificial Neural Networks

Neural Networks (NN)



Suddenly ANN is best – why here why now? - I thought it was a crazy idea when I first heard of it.

Advancements of technology

- Stochastic Gradient Descent.
- The computational GPU (Graphic Processing Unites), and the huge evolution of computations speed.

The problem – an example data set USPS Figs –

A supervised learning problem

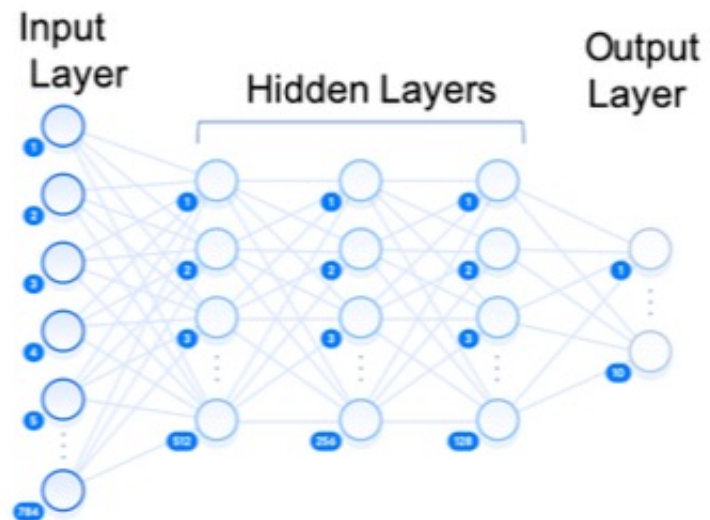


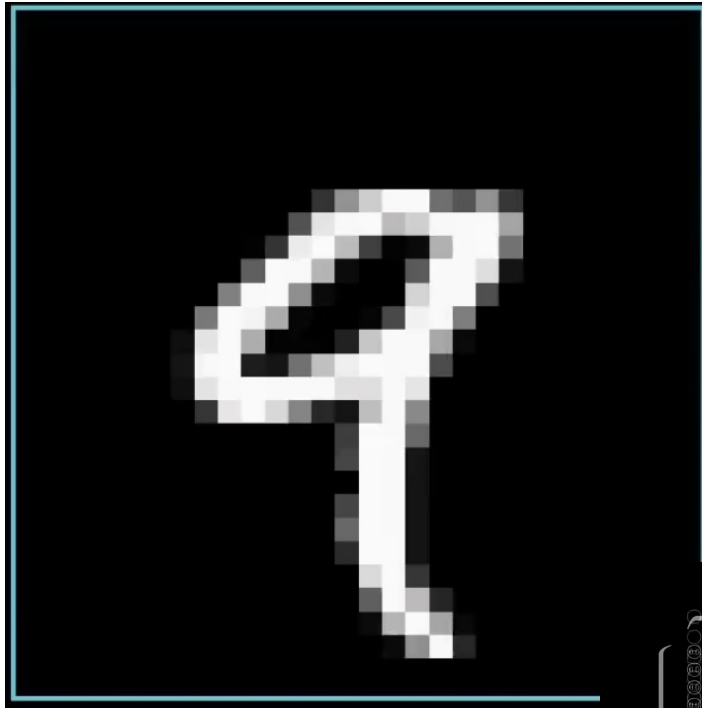
- input layer
- hidden inner layers
- output layer

- hidden inner layers

- hidden inner layers

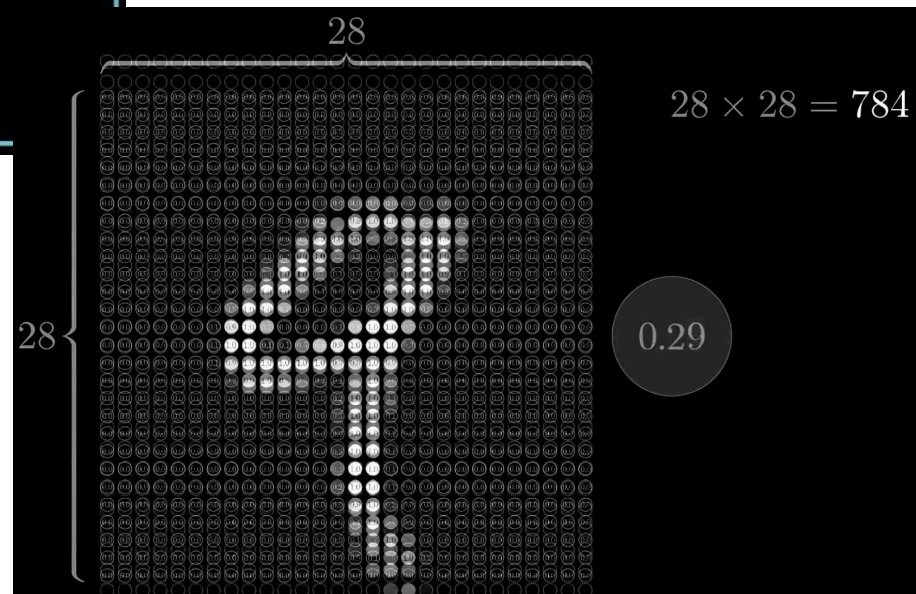
-output layer





This datum is a 784 numbers between 0 and 1.

Data $\{(x_i, y_i)\}_{i=1..784}$

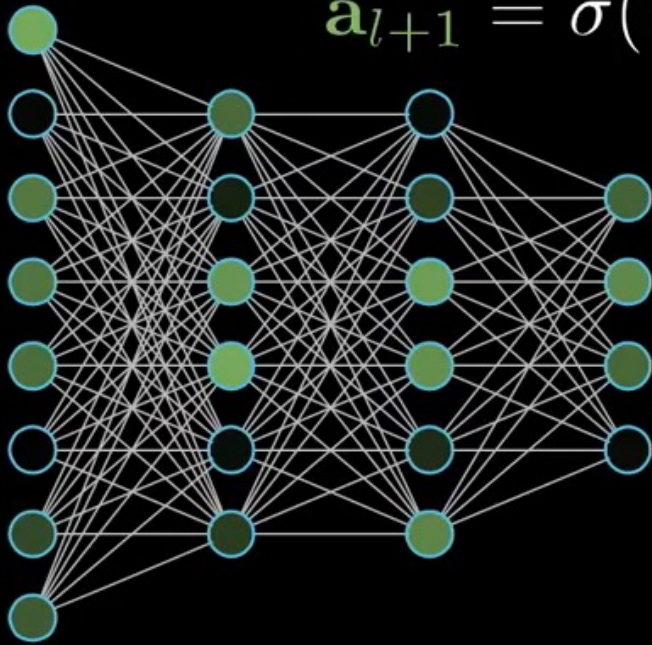


The data is a collection on ordered pairs, vector valued figures together with a symbol

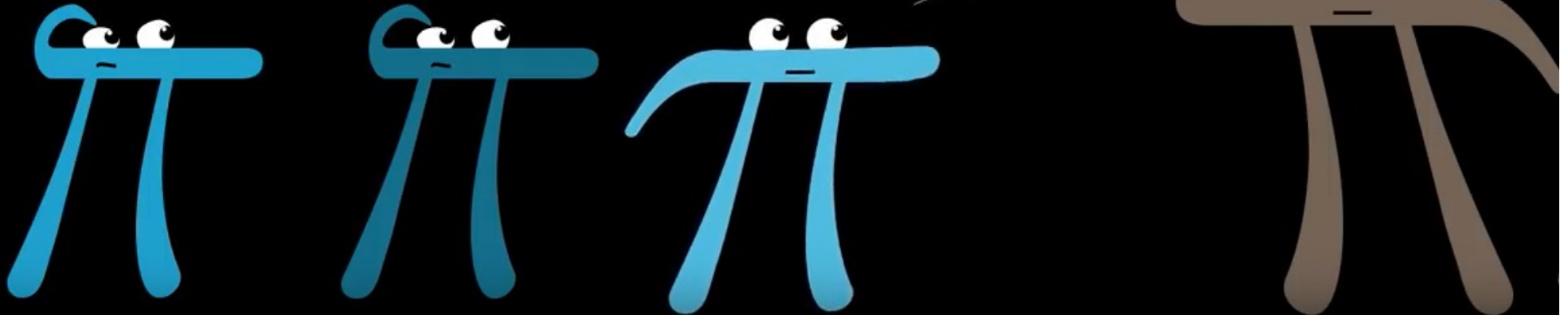
(0, 0)	(6, 6)	(3, 3)	(6, 6)	(7, 7)	(8, 8)	(0, 0)	(9, 9)
(5, 5)	(4, 4)	(3, 3)	(6, 6)	(5, 5)	(8, 8)	(9, 9)	(5, 5)
(4, 4)	(4, 4)	(7, 7)	(2, 2)	(0, 0)	(3, 3)	(2, 2)	(8, 8)
(9, 9)	(1, 1)	(9, 9)	(2, 2)	(2, 2)	(7, 7)	(9, 9)	(4, 4)
(8, 8)	(7, 7)	(4, 4)	(1, 1)	(3, 3)	(1, 1)	(5, 5)	(3, 3)
(2, 2)	(3, 3)	(9, 9)	(0, 0)	(9, 9)	(9, 9)	(1, 1)	(5, 5)
(8, 8)	(4, 4)	(7, 7)	(7, 7)	(4, 4)	(4, 4)	(4, 4)	(2, 2)
(0, 0)	(7, 7)	(2, 2)	(4, 4)	(8, 8)	(2, 2)	(6, 6)	(9, 9)
(9, 9)	(2, 2)	(8, 8)	(7, 7)	(6, 6)	(1, 1)	(1, 1)	(2, 2)
(3, 3)	(9, 9)	(1, 1)	(6, 6)	(5, 5)	(1, 1)	(1, 1)	(0, 0)

$$\mathbf{a}_{l+1} = \sigma(W_l \mathbf{a}_l + b_l)$$

Machine learning
Neural network

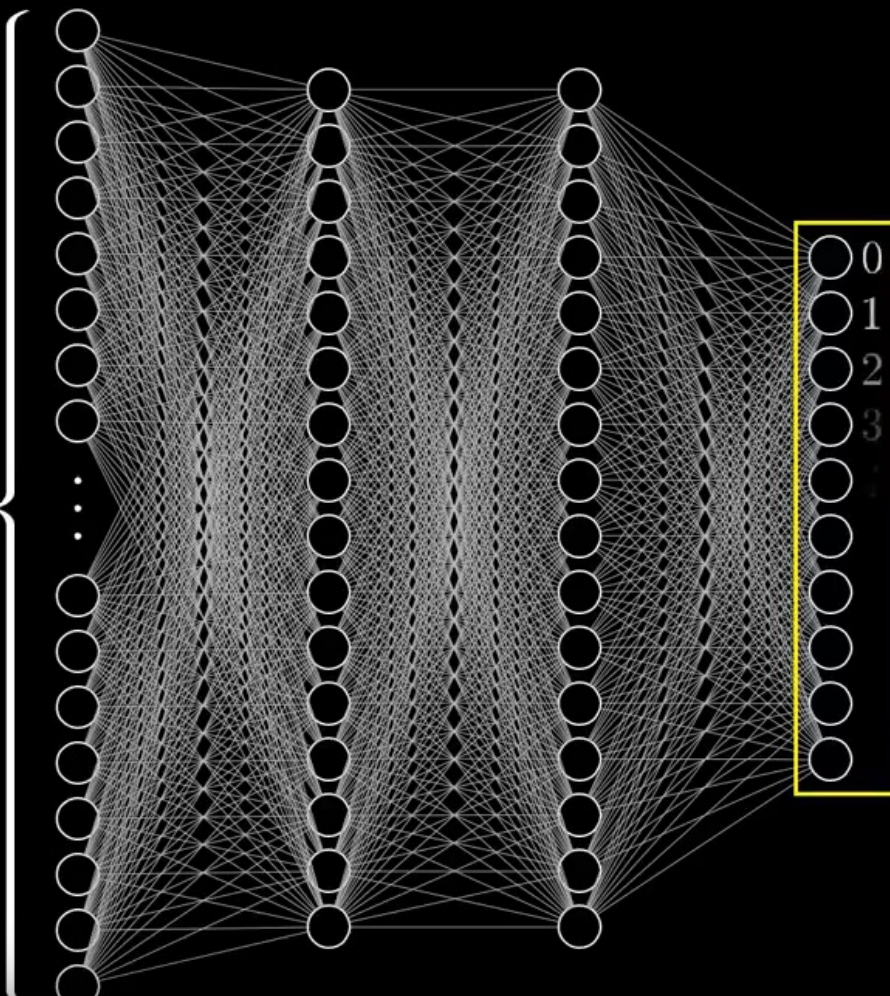


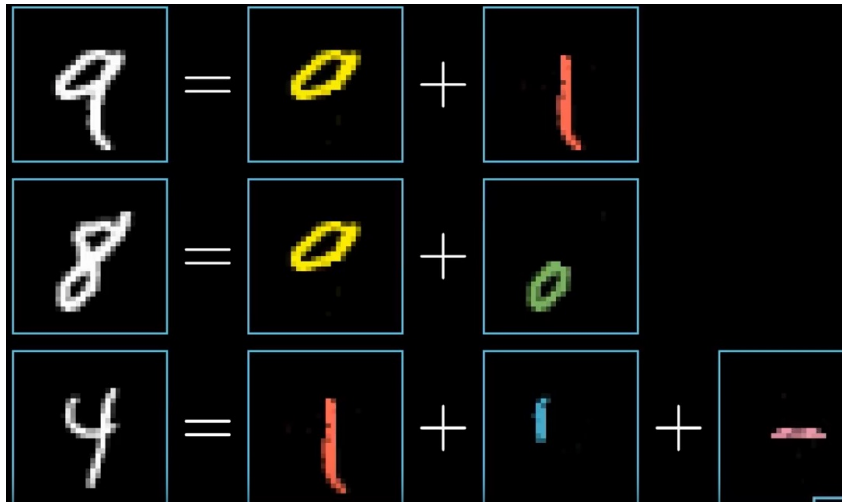
Why the layers?



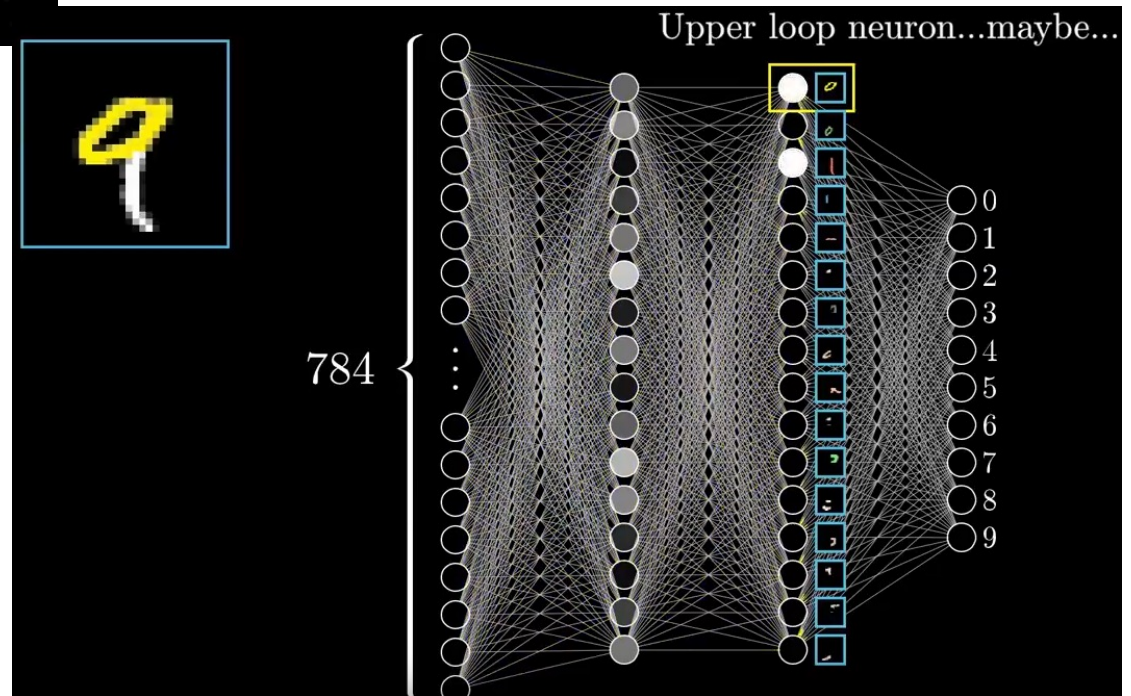


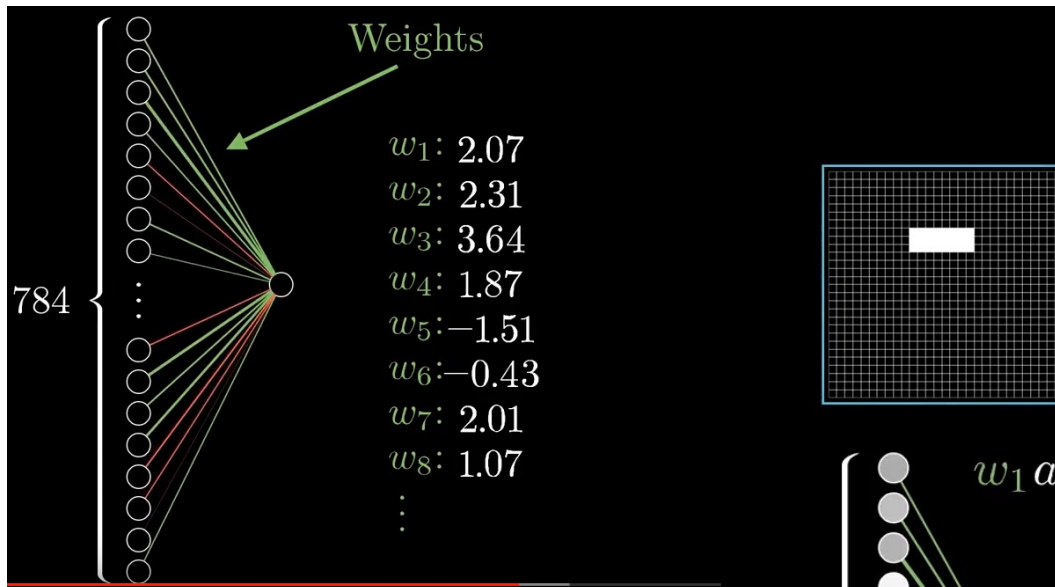
784





What do you wish in the “neurons”?

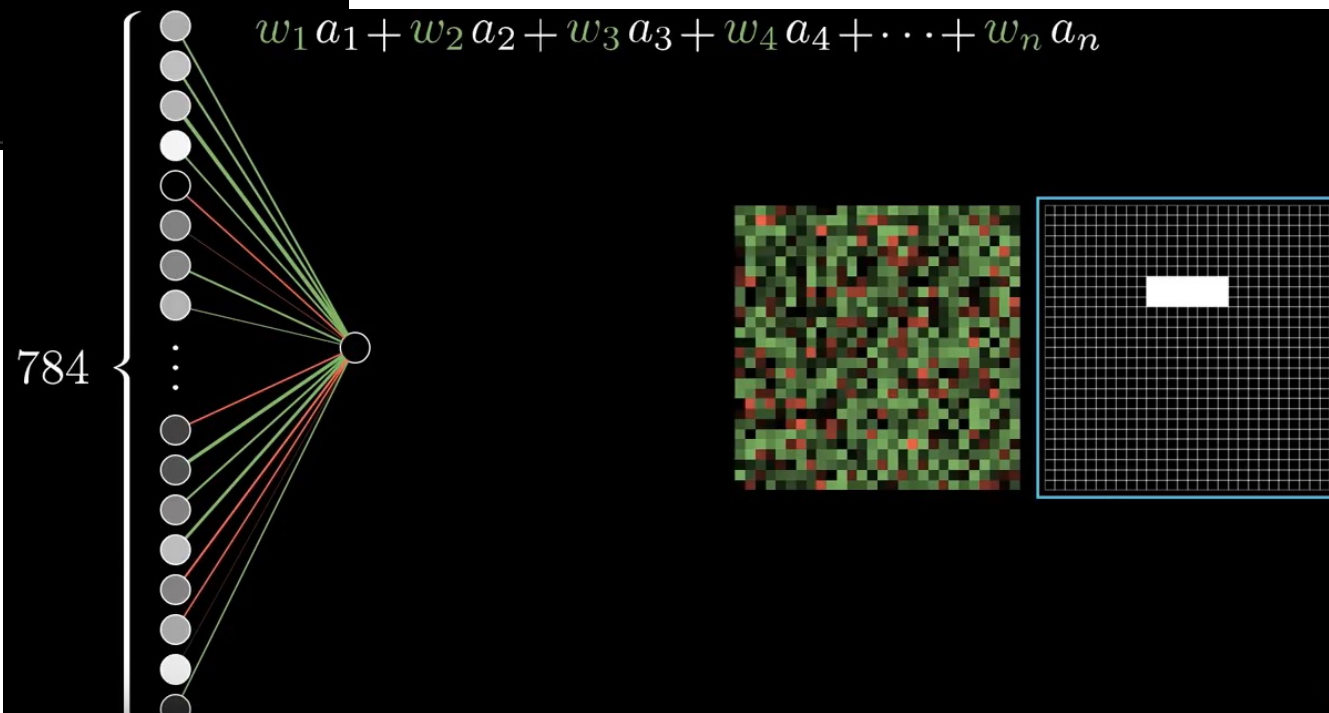


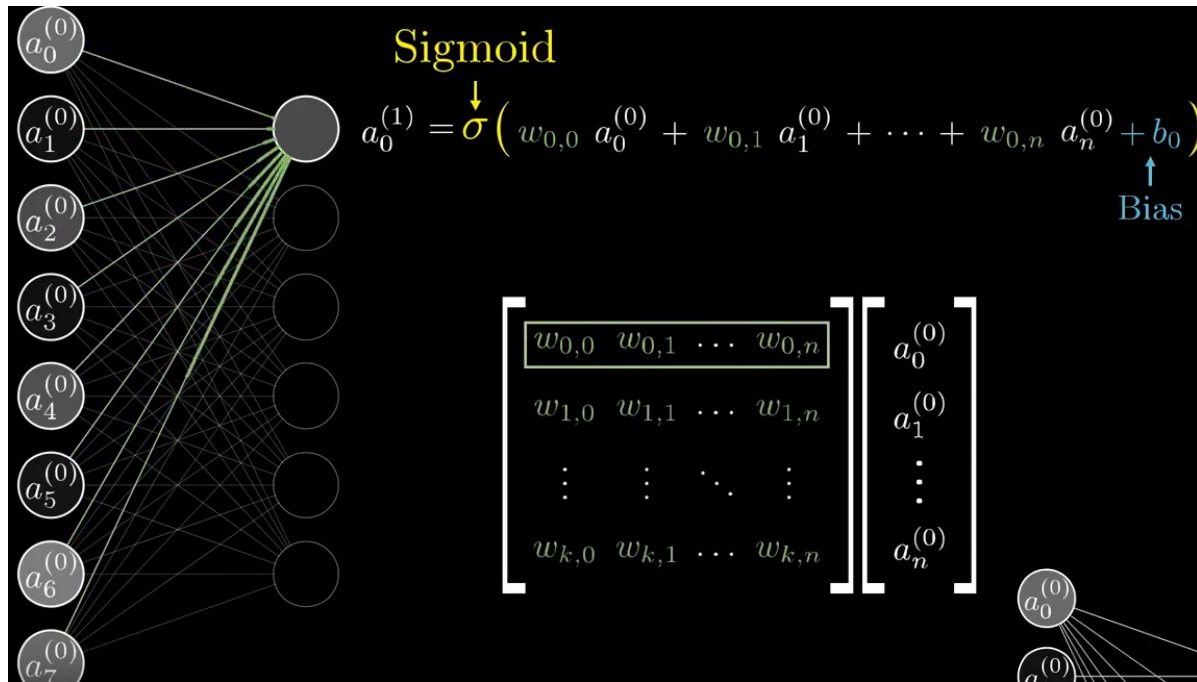


What do you wish in the “neurons”?

Vs

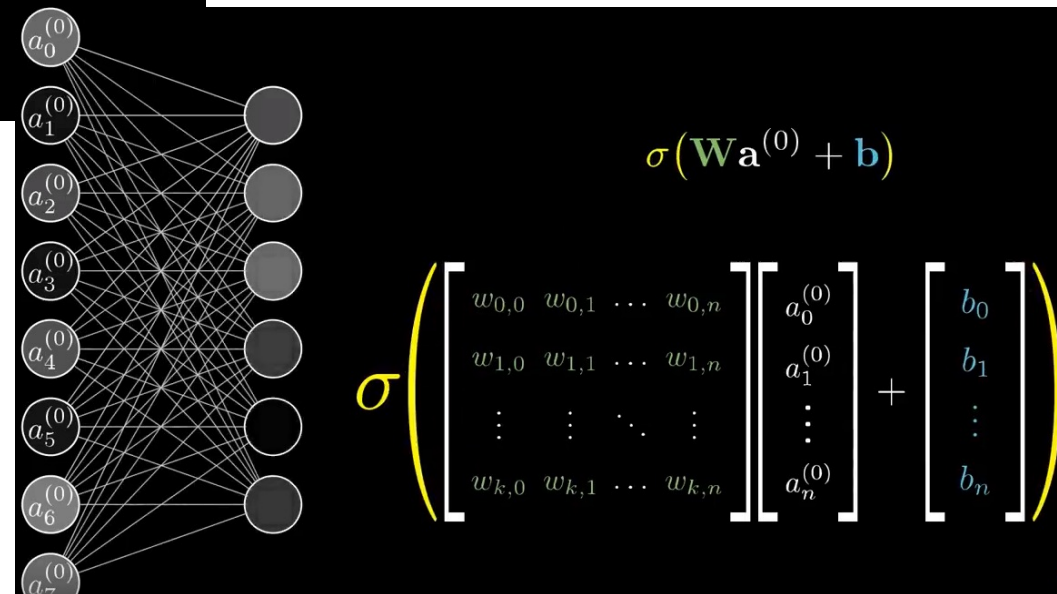
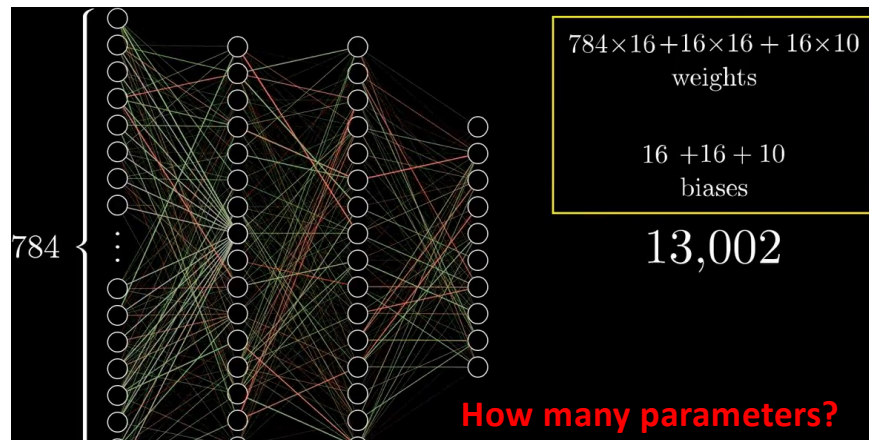
What do we get in the neurons.





The ANN – Mathematical Process

The matrix form of a layer, and composition



The basic nodal statement.

$$y = \sigma \left(\sum_{i=1}^n w_i x_i + b \right)$$

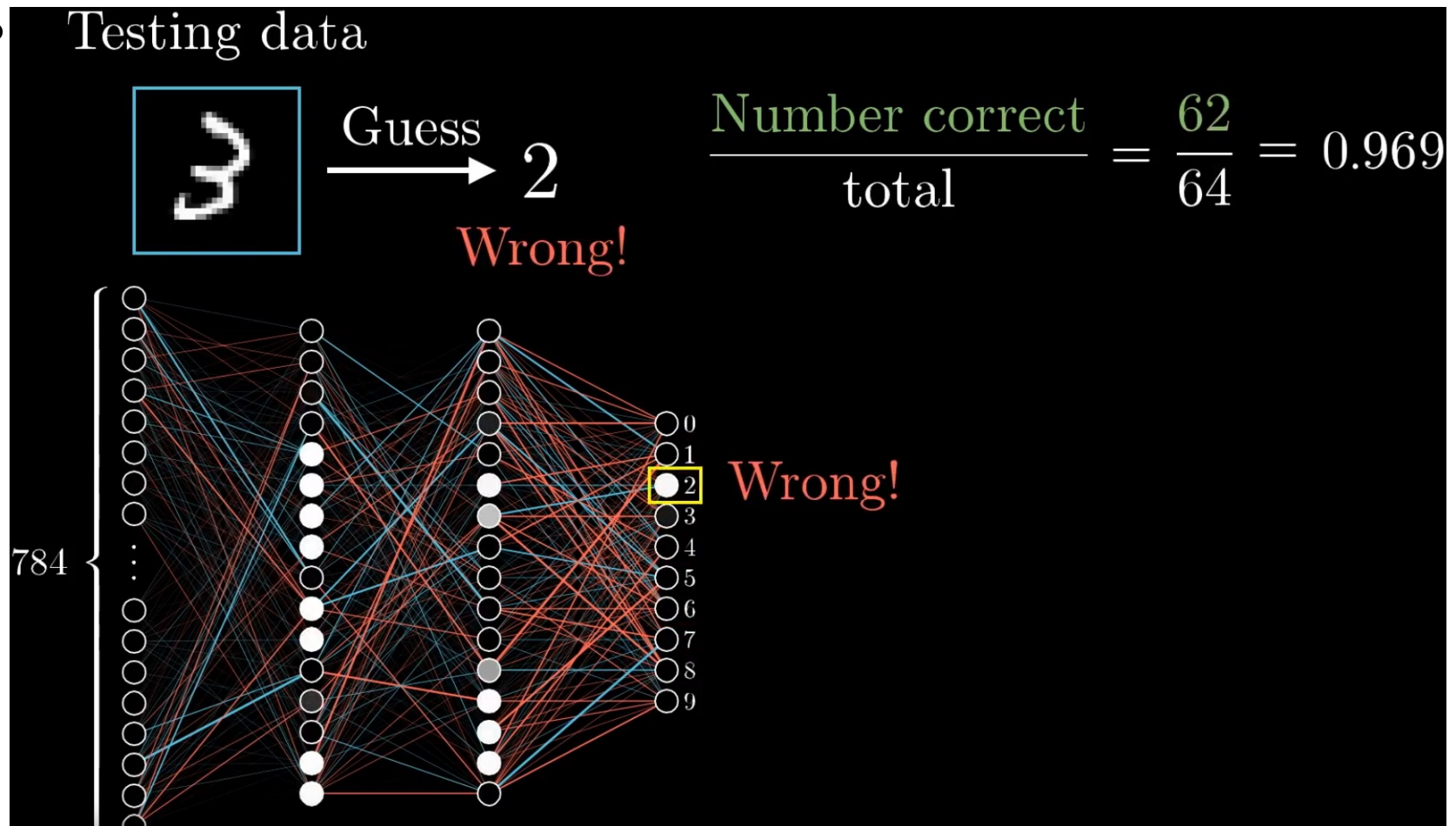
Favorite activation functions

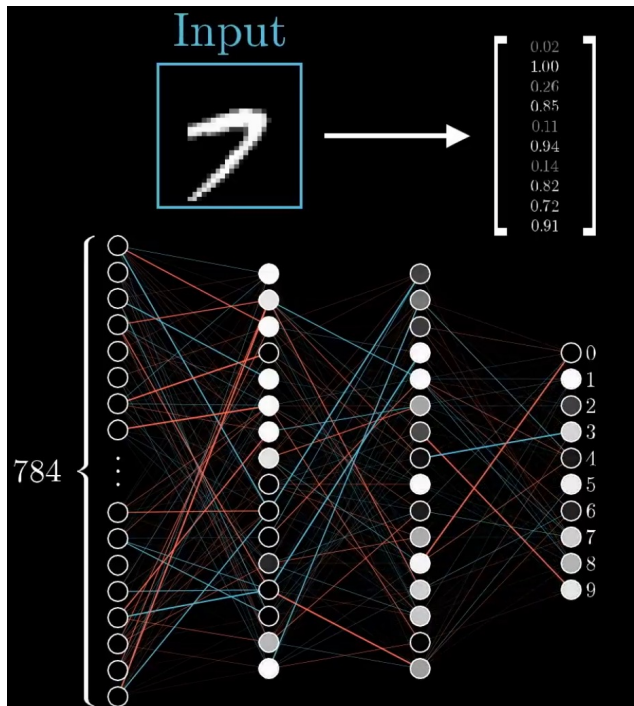
- Linear function: $\sigma(x) = ax + b$, with a and b are constants.
- Sigmoid (Logistic) function: $\sigma(x) = \frac{1}{1+e^{-x}}$.
- Hyperbolic Tangent: $\sigma(x) = \tanh(x)$.
- Rectified Linear Unit (ReLU): $\sigma(x) = \max\{0, x\}$.

What's the cost?

Vs what?

Vs what we want?





Neural network function

Input: 784 numbers (pixels)

Output: 10 numbers

Parameters:

Average cost of
all training data...

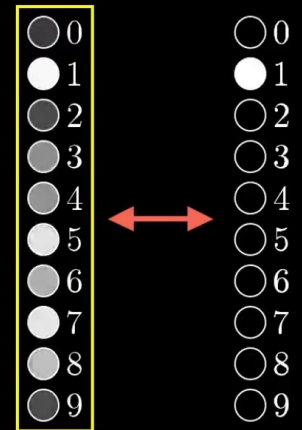
Cost of



$$\left\{ \begin{array}{l} (0.23 - 0.00)^2 + \\ (0.98 - 1.00)^2 + \\ (0.31 - 0.00)^2 + \\ (0.56 - 0.00)^2 + \\ (0.58 - 0.00)^2 + \\ (0.90 - 0.00)^2 + \\ (0.69 - 0.00)^2 + \\ (0.91 - 0.00)^2 + \\ (0.76 - 0.00)^2 + \\ (0.31 - 0.00)^2 \end{array} \right.$$

Let's write out the cost function

What's the "cost"
of this difference?



Utter trash

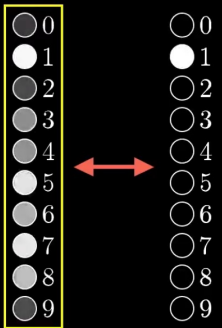
Average cost of all training data...

Cost of  {

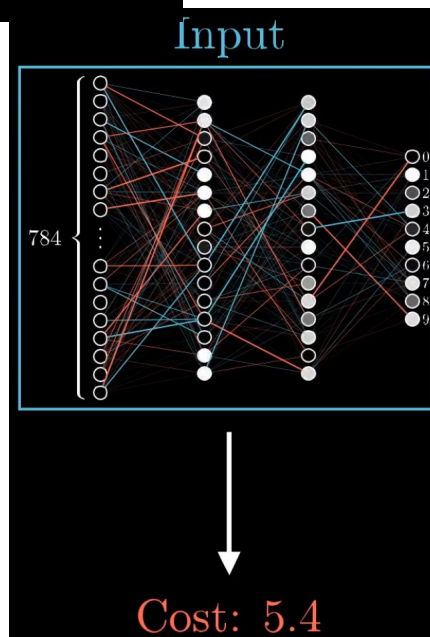
$$\begin{aligned}
 & (0.23 - 0.00)^2 + \\
 & (0.98 - 1.00)^2 + \\
 & (0.31 - 0.00)^2 + \\
 & (0.56 - 0.00)^2 + \\
 & (0.58 - 0.00)^2 + \\
 & (0.90 - 0.00)^2 + \\
 & (0.69 - 0.00)^2 + \\
 & (0.91 - 0.00)^2 + \\
 & (0.76 - 0.00)^2 + \\
 & (0.31 - 0.00)^2
 \end{aligned}$$

What's the "cost" of this difference?

Utter trash



Let's write out the cost function



Cost function

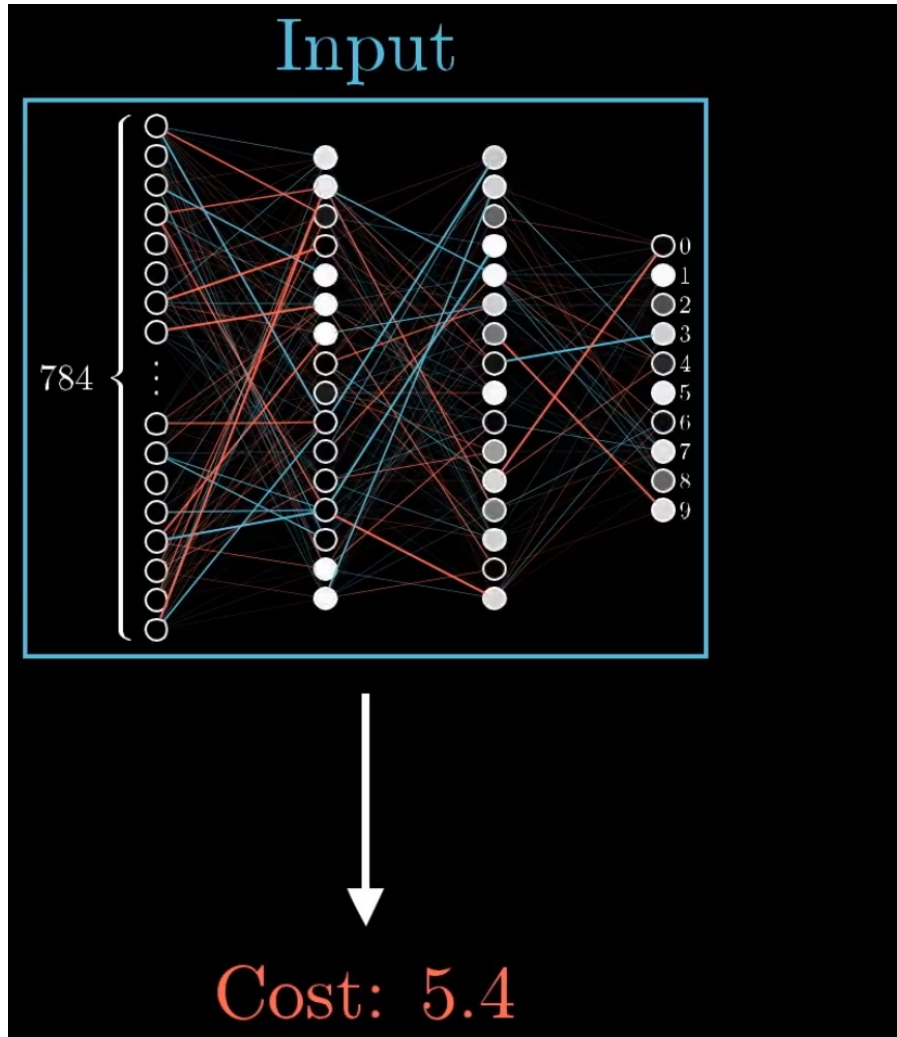
Input: 13,002 weights/biases

Output: 1 number (the cost)

Parameters: Many, many, many training examples

(, 2)

Let's write out the cost function



Cost function

Input: 13,002 weights/biases

Output: 1 number (the cost)

Parameters: Many, many, many training examples

$(\boxed{2}, 2)$

Cost function

$$C(w_1, w_2, \dots, w_{13,002})$$

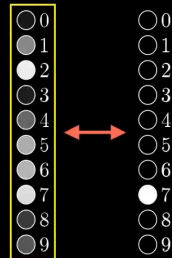
Weights and biases

Average cost of all training data...

Cost of 7

$$\begin{aligned} & (0.10 - 0.00)^2 + \\ & (0.54 - 0.00)^2 + \\ & (0.94 - 0.00)^2 + \\ & (0.12 - 0.00)^2 + \\ & (0.40 - 0.00)^2 + \\ & (0.69 - 0.00)^2 + \\ & (0.72 - 0.00)^2 + \\ & (0.87 - 1.00)^2 + \\ & (0.24 - 0.00)^2 + \\ & (0.25 - 0.00)^2 \end{aligned}$$

What's the "cost" of this difference?

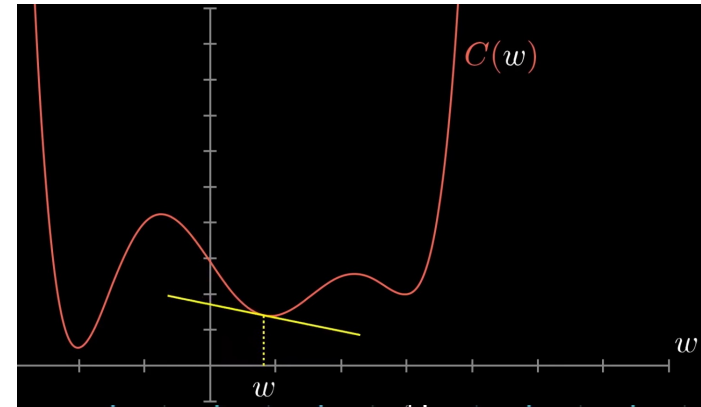
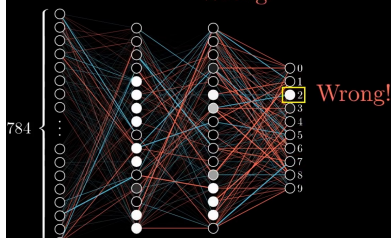


Utter trash

Testing data

Guess 2 Wrong!

$$\frac{\text{Number correct}}{\text{total}} = \frac{62}{64} = 0.969$$

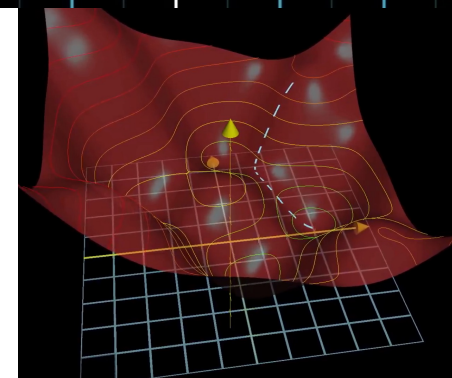


Input space

"Gradient", the direction of steepest increase

$$\nabla C(x, y)$$

Which direction decreases $C(x, y)$ most quickly?



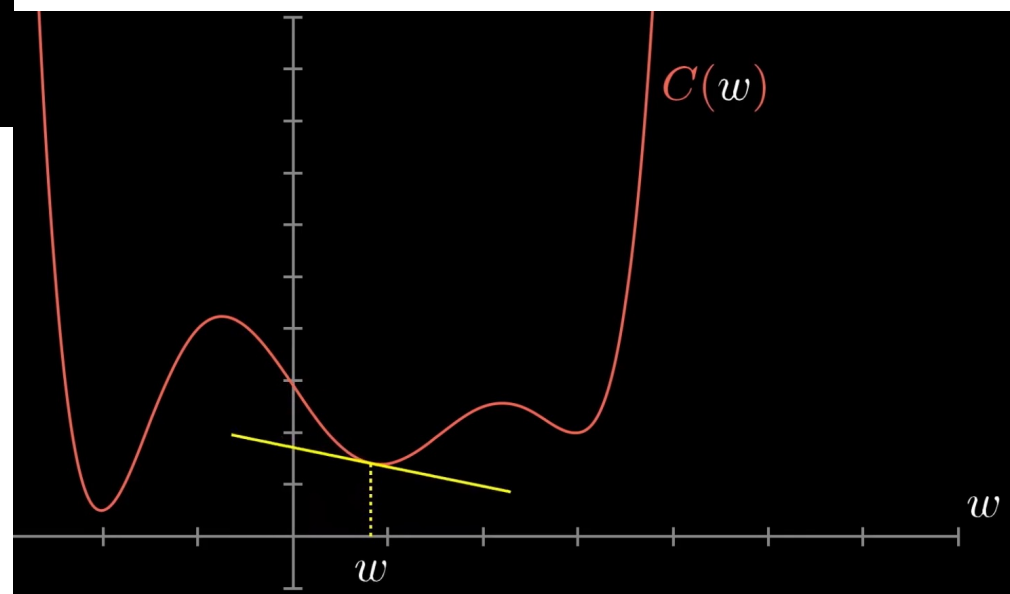
Optimization methods

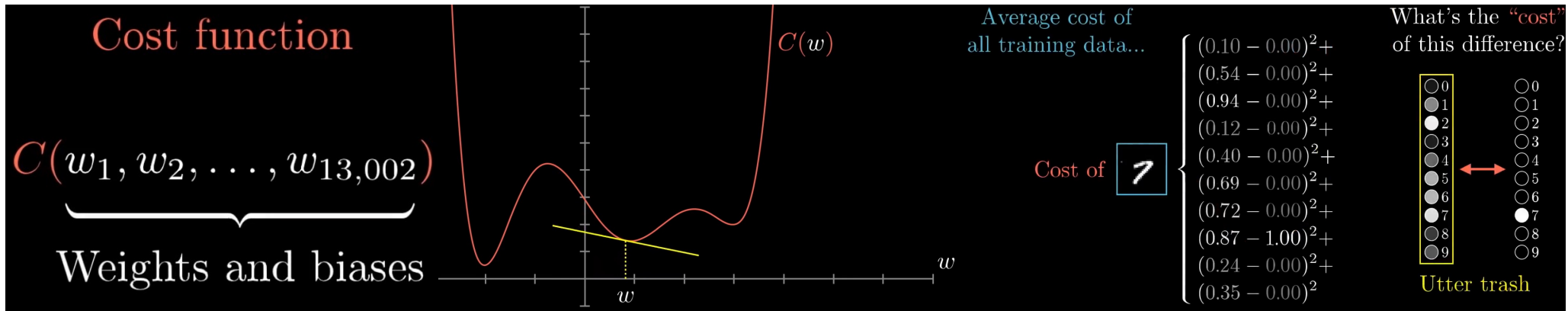
Cost function

$$C(\underbrace{w_1, w_2, \dots, w_{13,002}}_{\text{Weights and biases}})$$

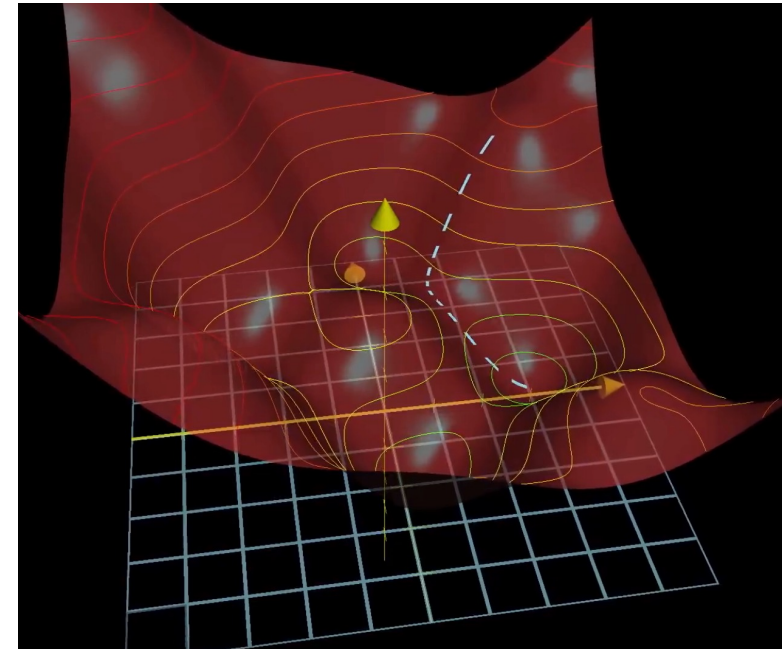
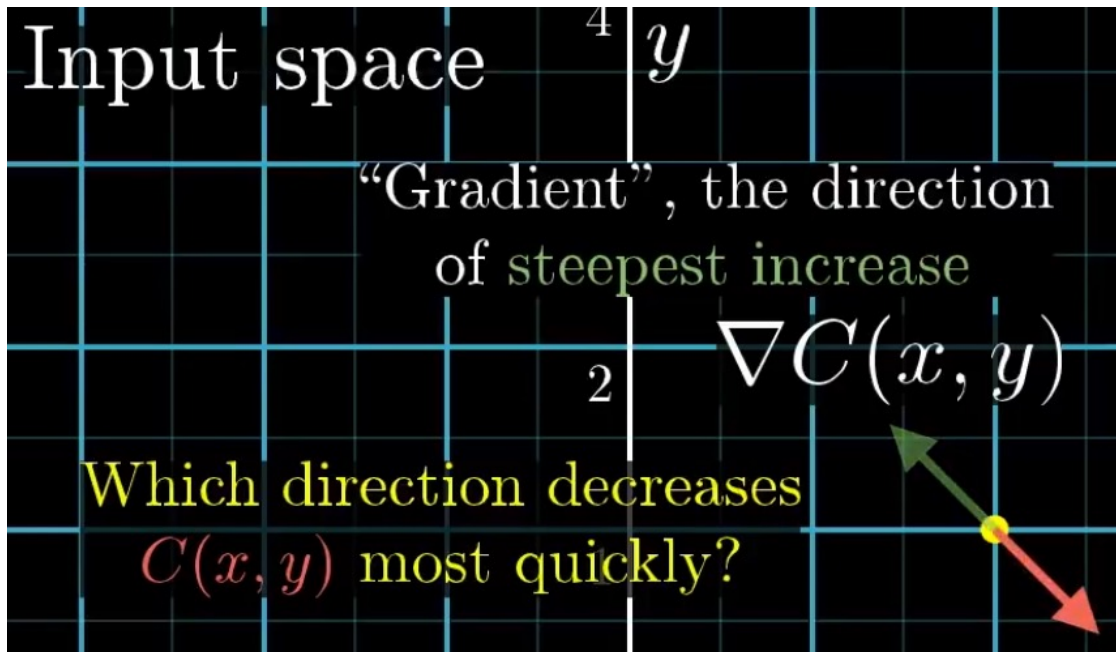
Weights and biases

An optimization problem

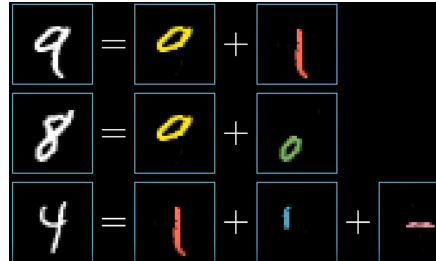




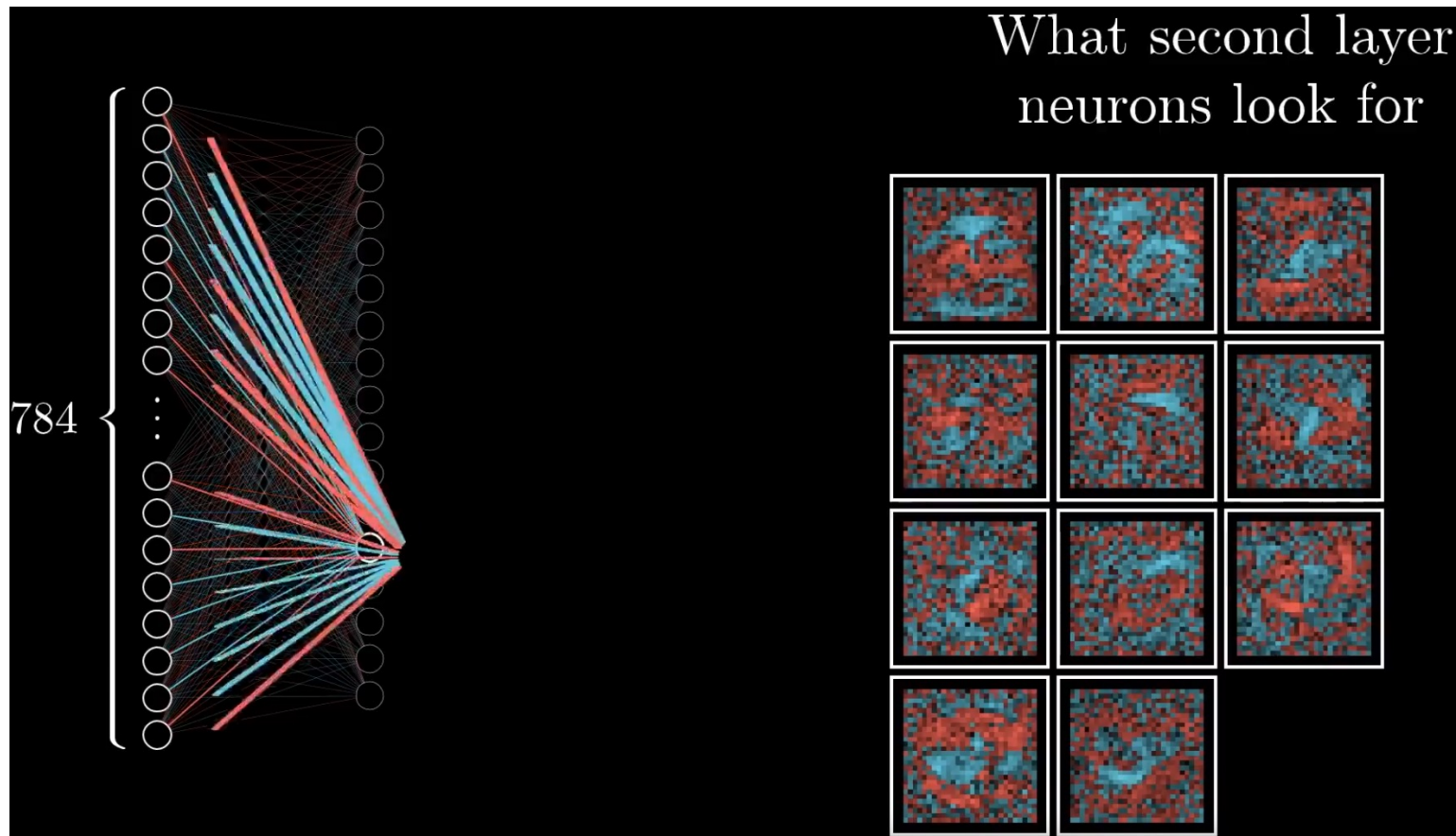
Optimization methods



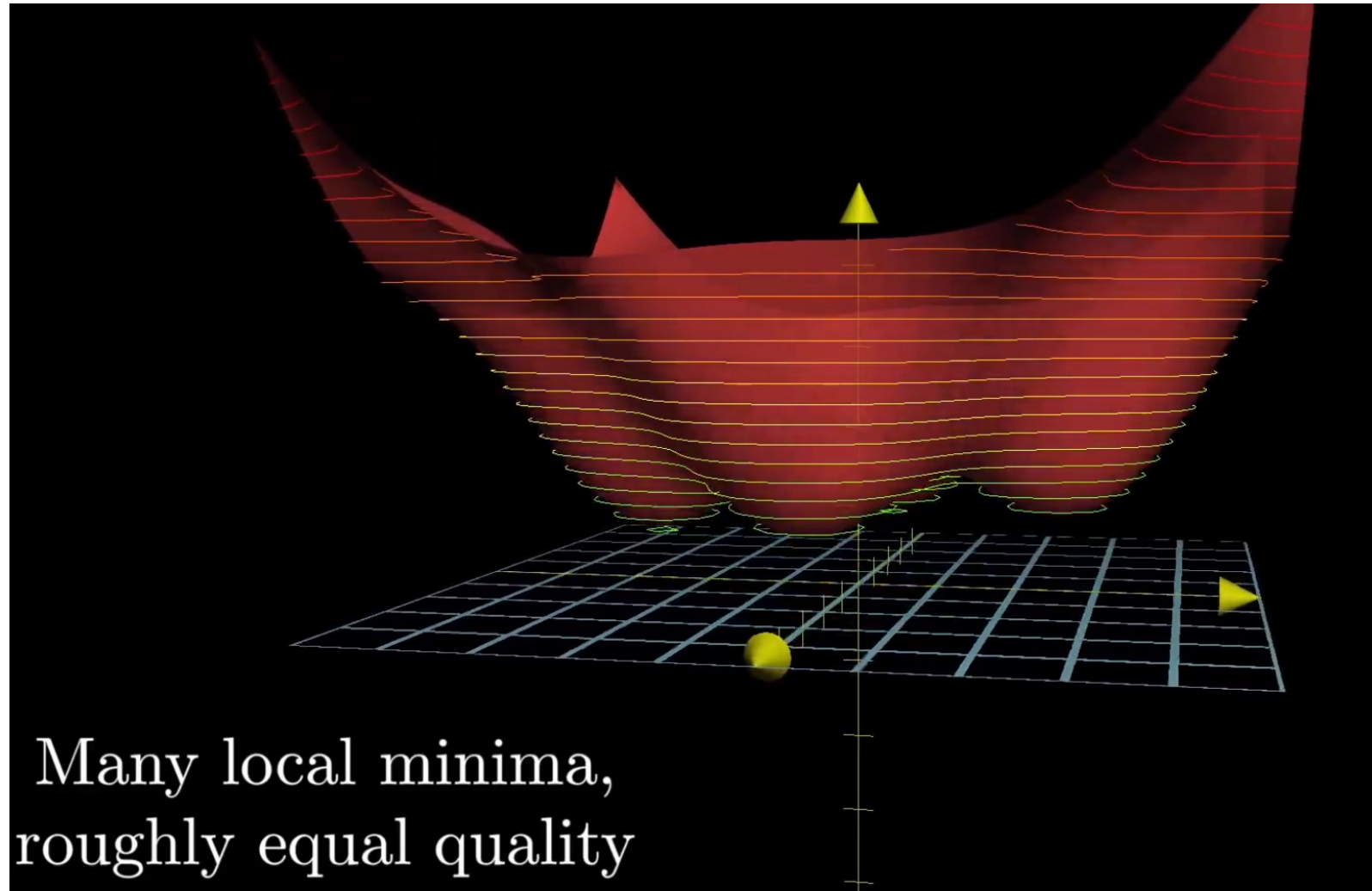
What do you wish in the “neurons”?



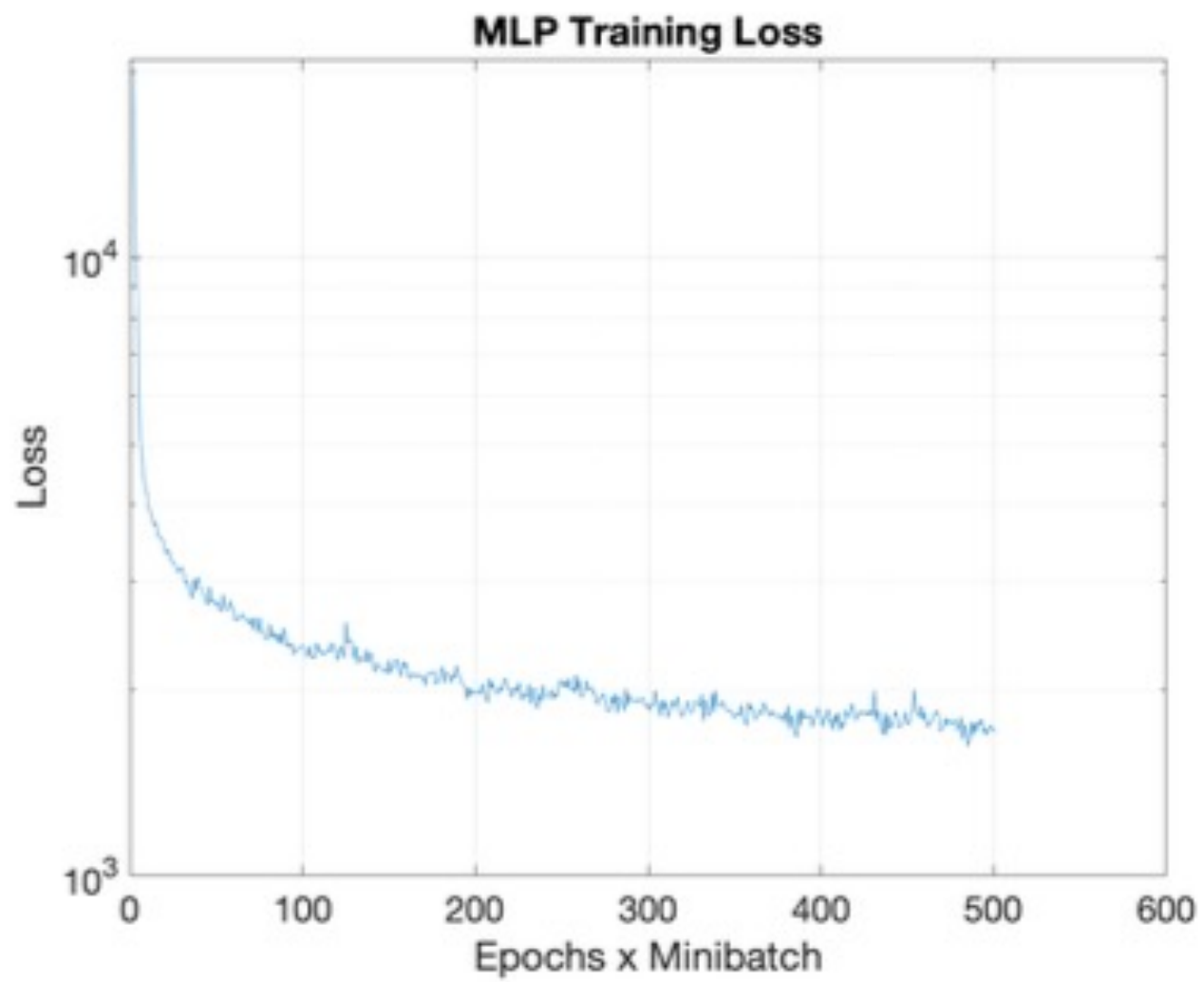
What do we actually get?



MANY local minima of the VERY high dimensional so how important is it to find the best One?

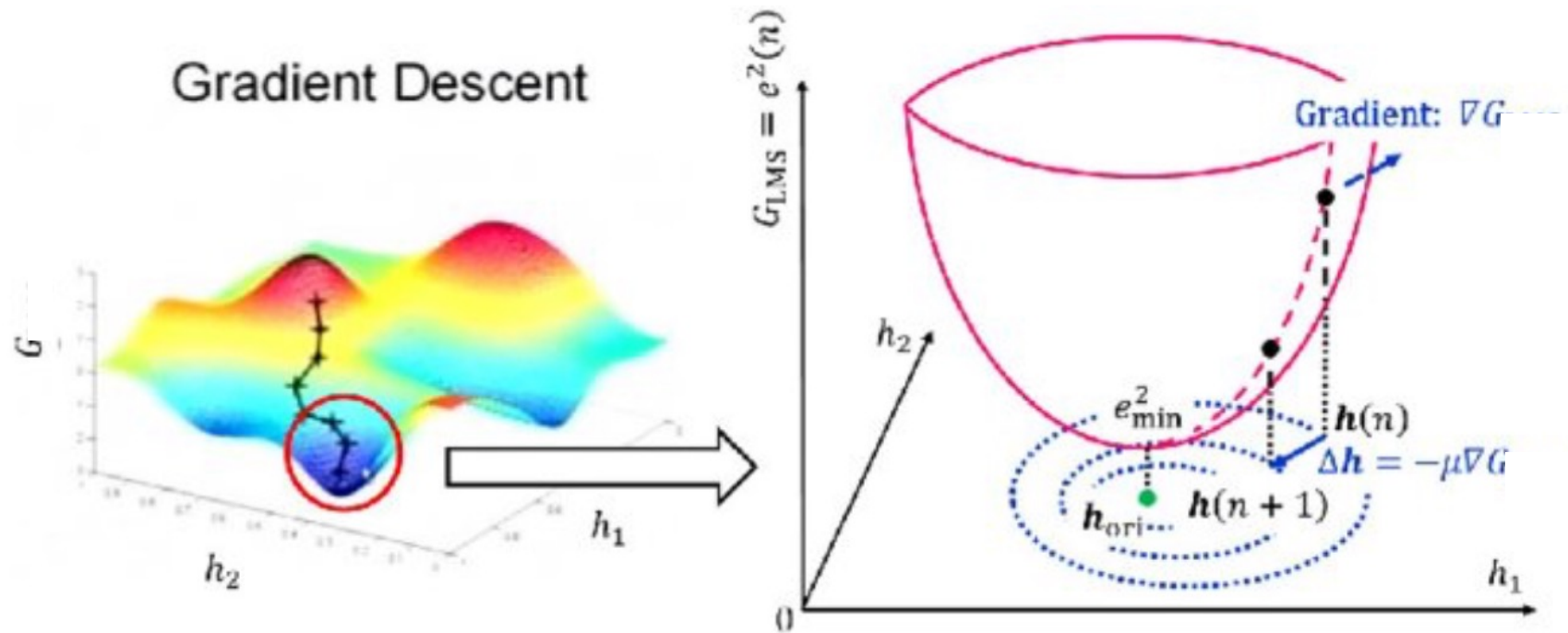


On training

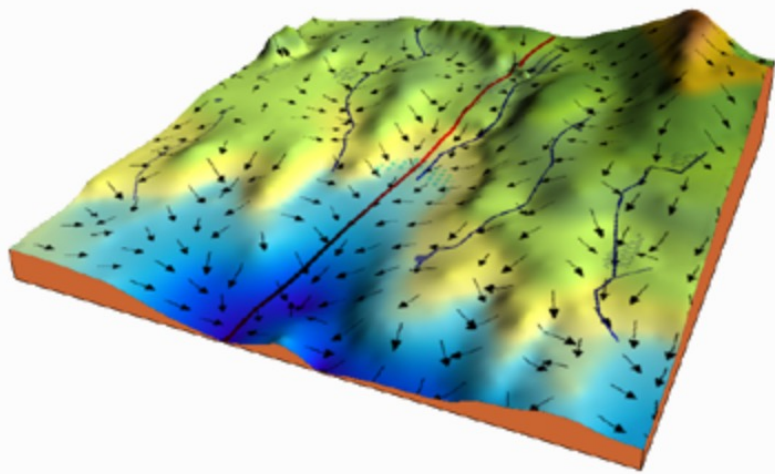


General smooth optimization, $G(h): \mathbb{R}^d \rightarrow \mathbb{R}$

By general gradient descent



More on gradient descent



Given the cost function:

$$f(m, b) = \frac{1}{N} \sum_{i=1}^n (y_i - (mx_i + b))^2$$

The gradient can be calculated as:

$$f'(m, b) = \begin{bmatrix} \frac{df}{dm} \\ \frac{df}{db} \end{bmatrix} = \begin{bmatrix} \frac{1}{N} \sum -2x_i(y_i - (mx_i + b)) \\ \frac{1}{N} \sum -2(y_i - (mx_i + b)) \end{bmatrix}$$

$$\mathbf{w}_{j+1}(\delta) = \mathbf{w}_j - \delta \nabla f_k(\mathbf{w}_j)$$

$$\theta^{(t+1)} = \theta^{(t)} - \alpha \cdot \nabla_{\theta} L(\theta^{(t)}, \mathbf{y})$$

In this equation:

- $\theta^{(t)}$ is our current estimate of θ^* at the t th iteration
- α is the learning rate
- $\nabla_{\theta} L(\theta^{(t)}, \mathbf{y})$ is the gradient of the loss function
- We compute the next estimate $\theta^{(t+1)}$ by subtracting the product of α and $\nabla_{\theta} L(\theta, \mathbf{y})$ computed at $\theta^{(t)}$

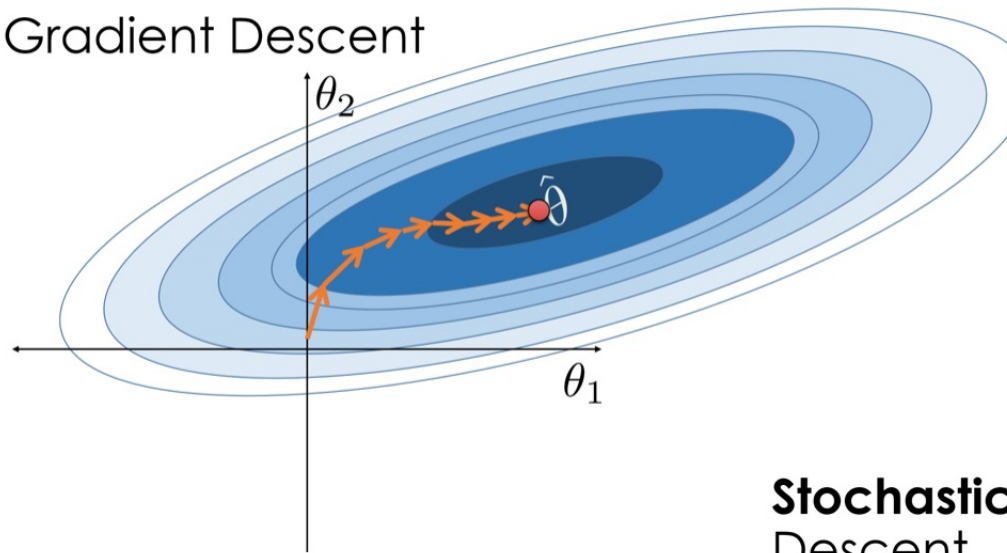
Stochastic Gradient Descent

$$\mathbf{w}_{j+1}(\delta) = \mathbf{w}_j - \delta \nabla f_K(\mathbf{w}_j)$$

Batch (several data points)

$$K \in [k_1, k_2, \dots, k_p]$$

Gradient Descent



Stochastic Gradient Descent

