



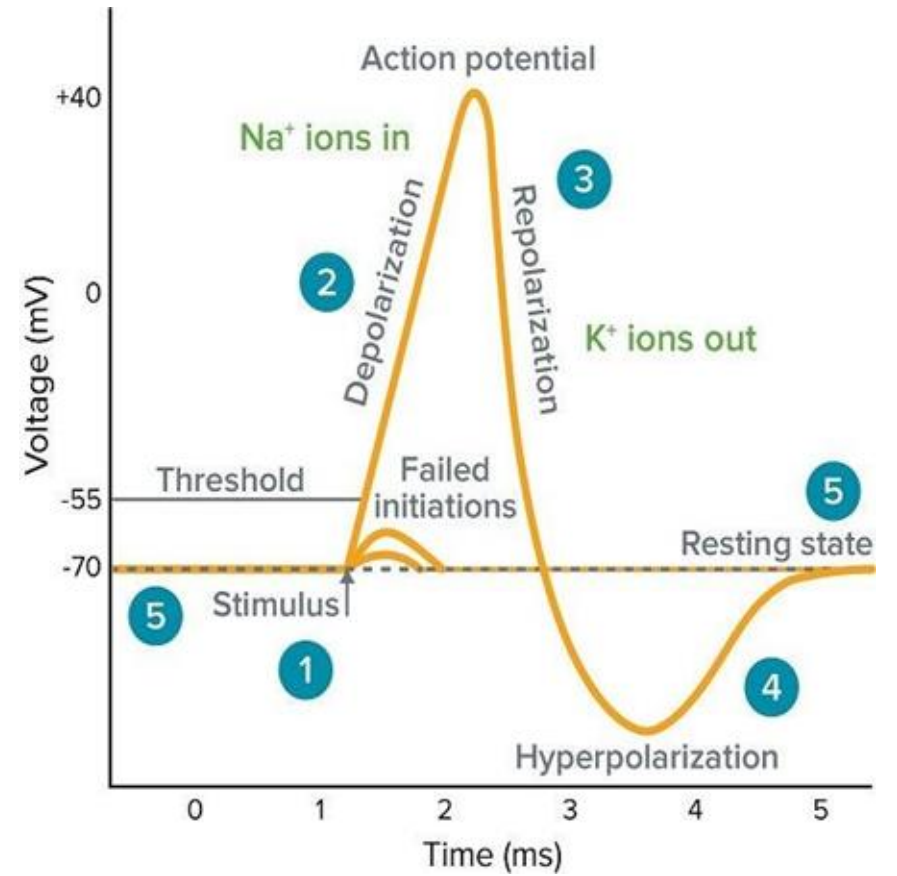
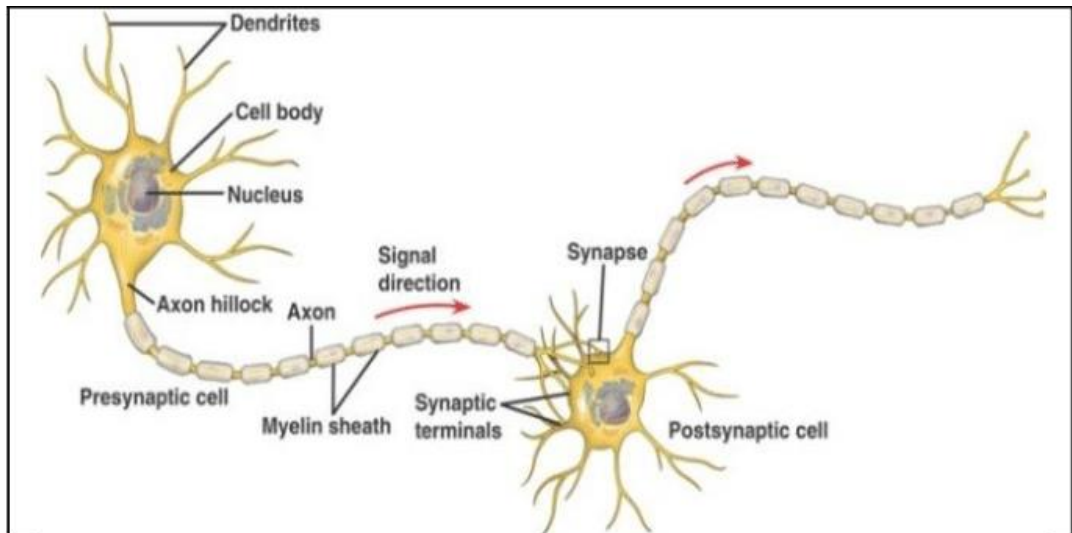
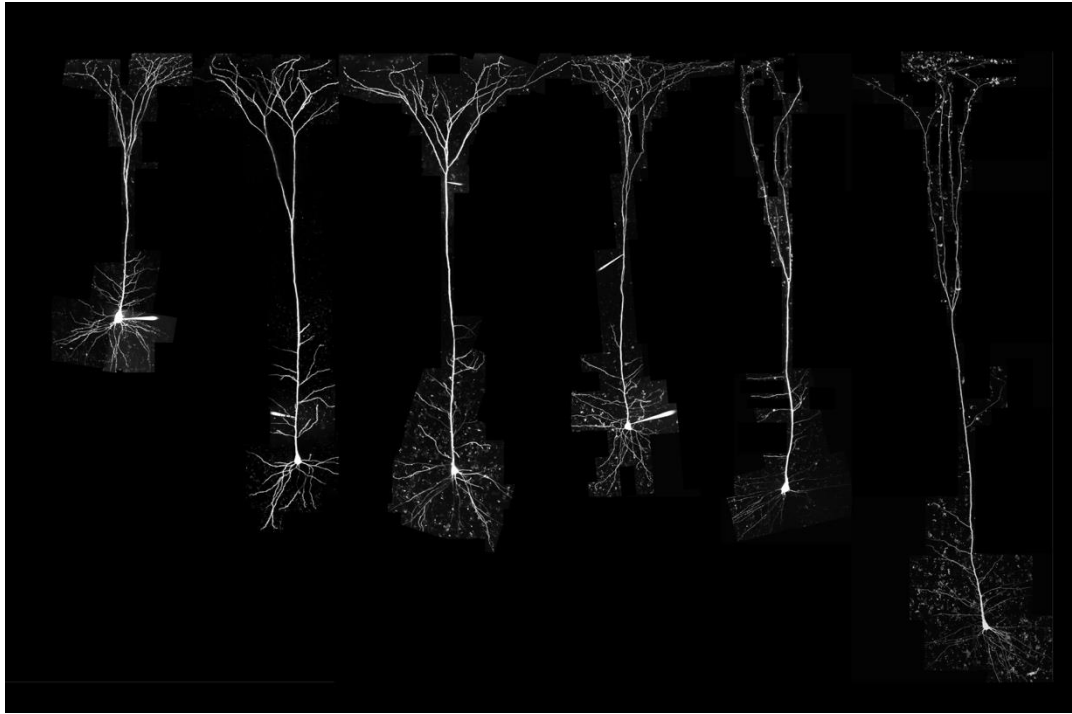
Neural Networks

Guillem Guigó i Corominas

guillem.guigo@uvic.cat

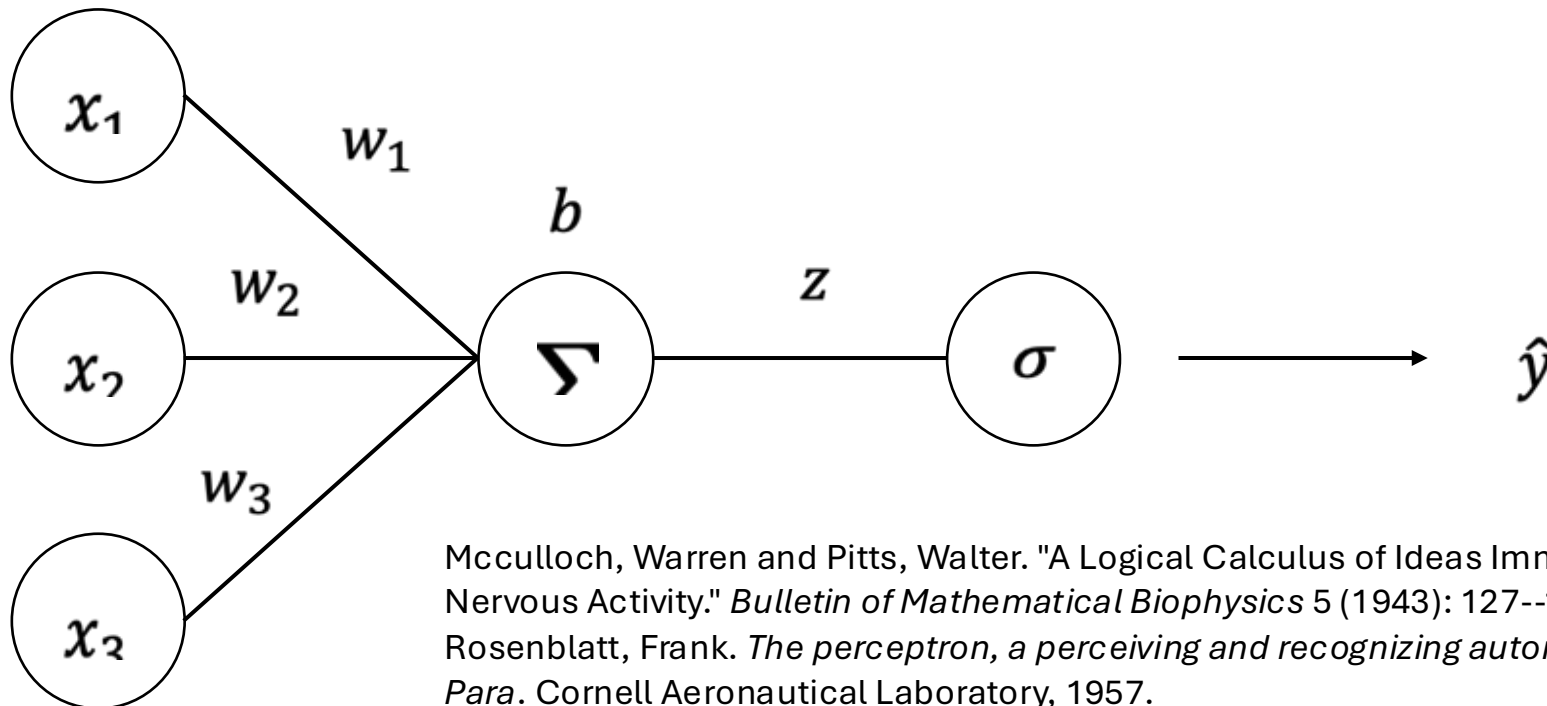


UNIVERSITAT DE VIC
UNIVERSITAT CENTRAL
DE CATALUNYA



Single neuron model: The Perceptron

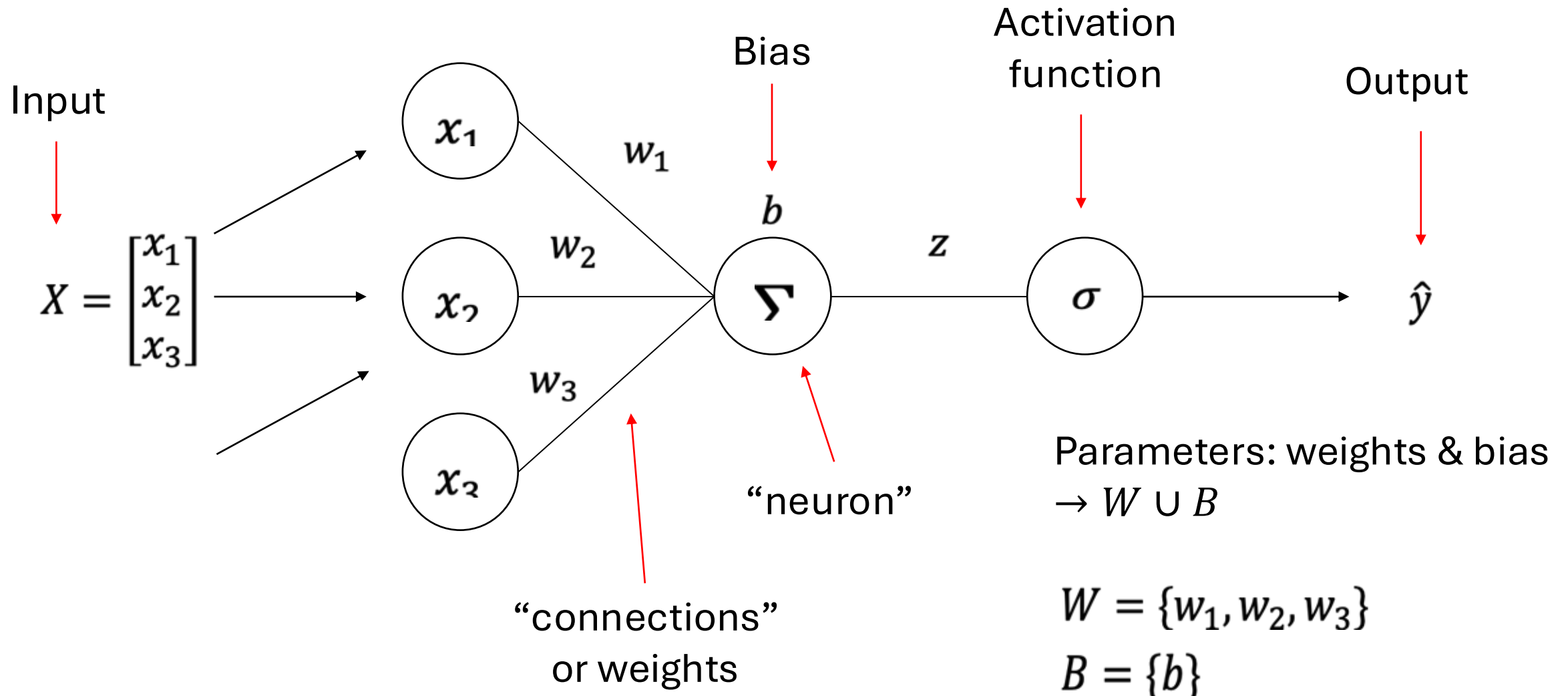
- McCulloch and Pitts (1943) – Model description
- Frank Rosenblatt (1957) – Machine implementation



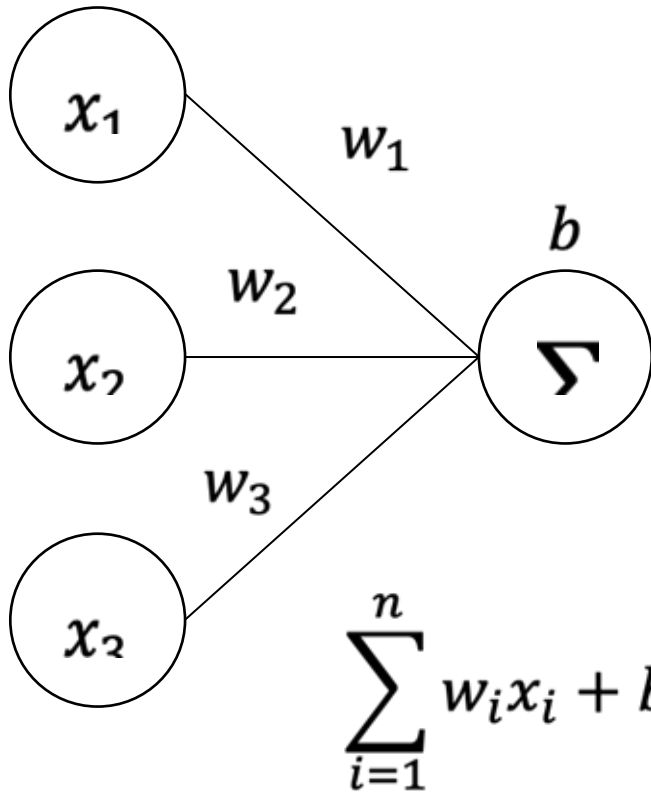
McCulloch, Warren and Pitts, Walter. "A Logical Calculus of Ideas Immanent in Nervous Activity." *Bulletin of Mathematical Biophysics* 5 (1943): 127--147.

Rosenblatt, Frank. *The perceptron, a perceiving and recognizing automaton Project Para.* Cornell Aeronautical Laboratory, 1957.

Single neuron model: The Perceptron



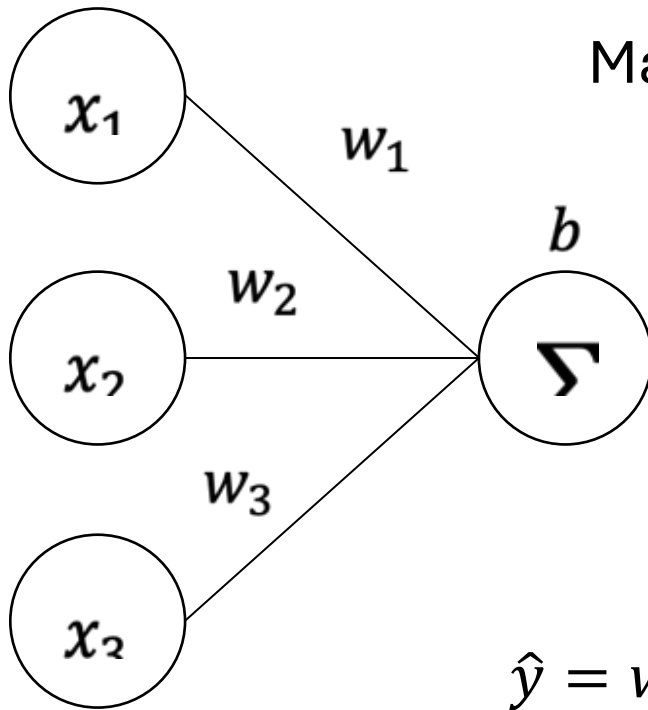
Single neuron model: The Perceptron



Single neuron model: The Perceptron

Linear combination of inputs scaled by the weights

Mathematically equivalent to linear regression



How many trainable parameters does this perceptron have?

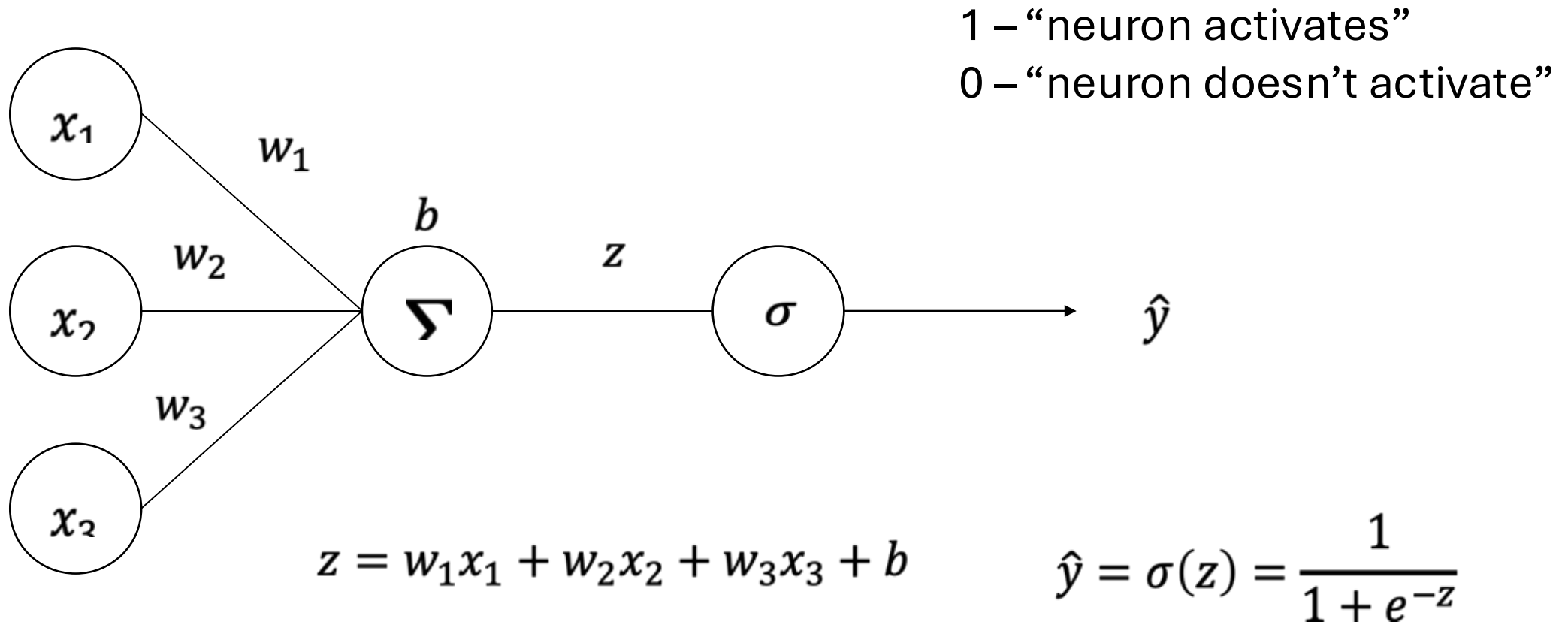
$$\hat{y} = w^T x + b$$

$$x = \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix}$$

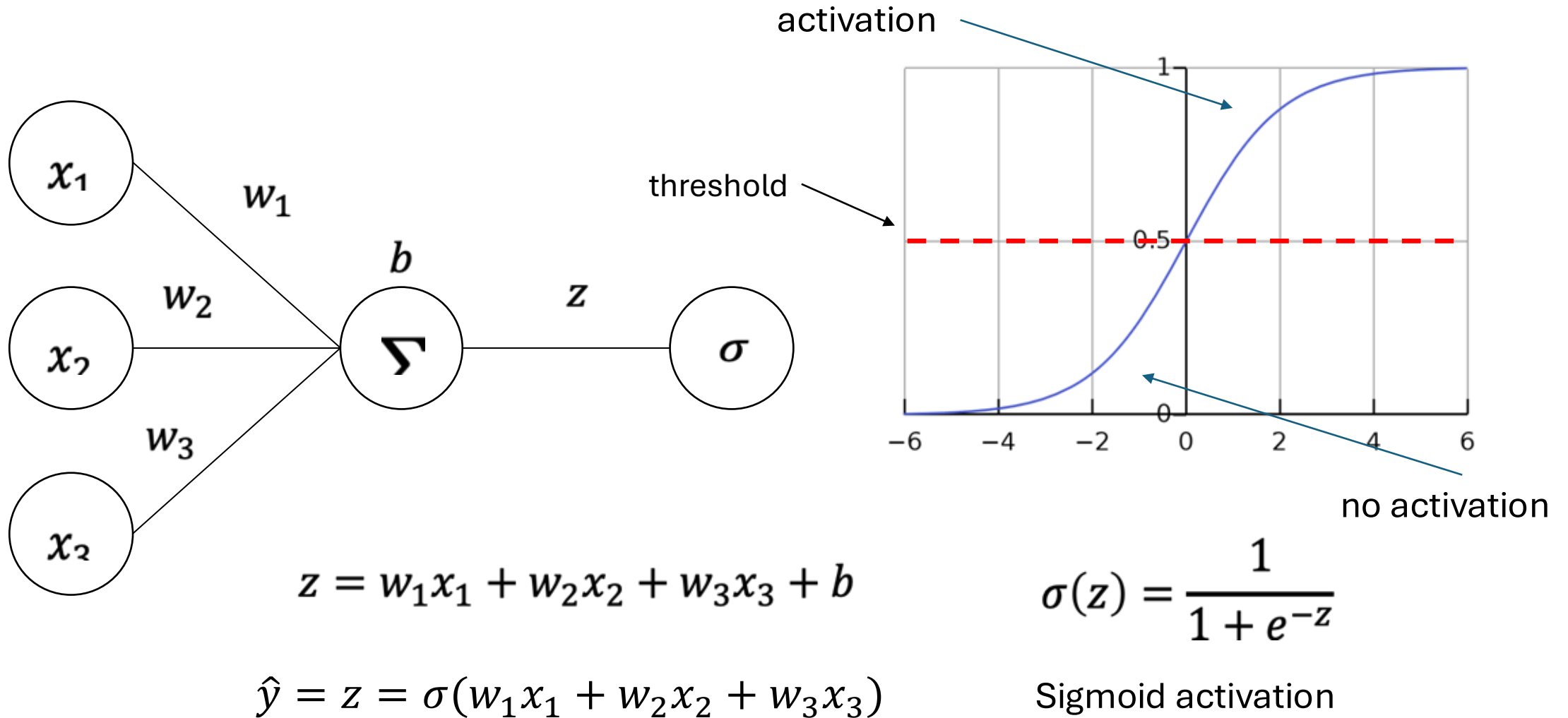
$$w = \begin{bmatrix} w_1 \\ w_2 \\ w_3 \end{bmatrix}$$

$$b = b$$

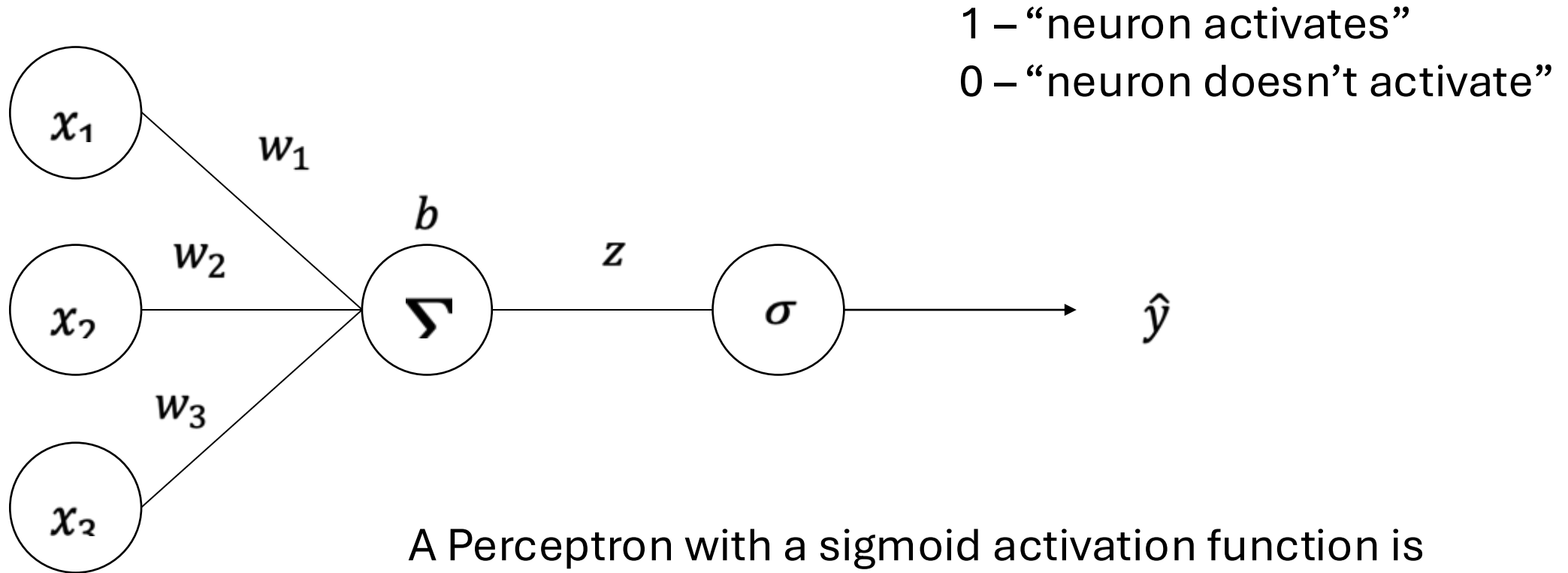
Single neuron model: The Perceptron



Single neuron model: The Perceptron



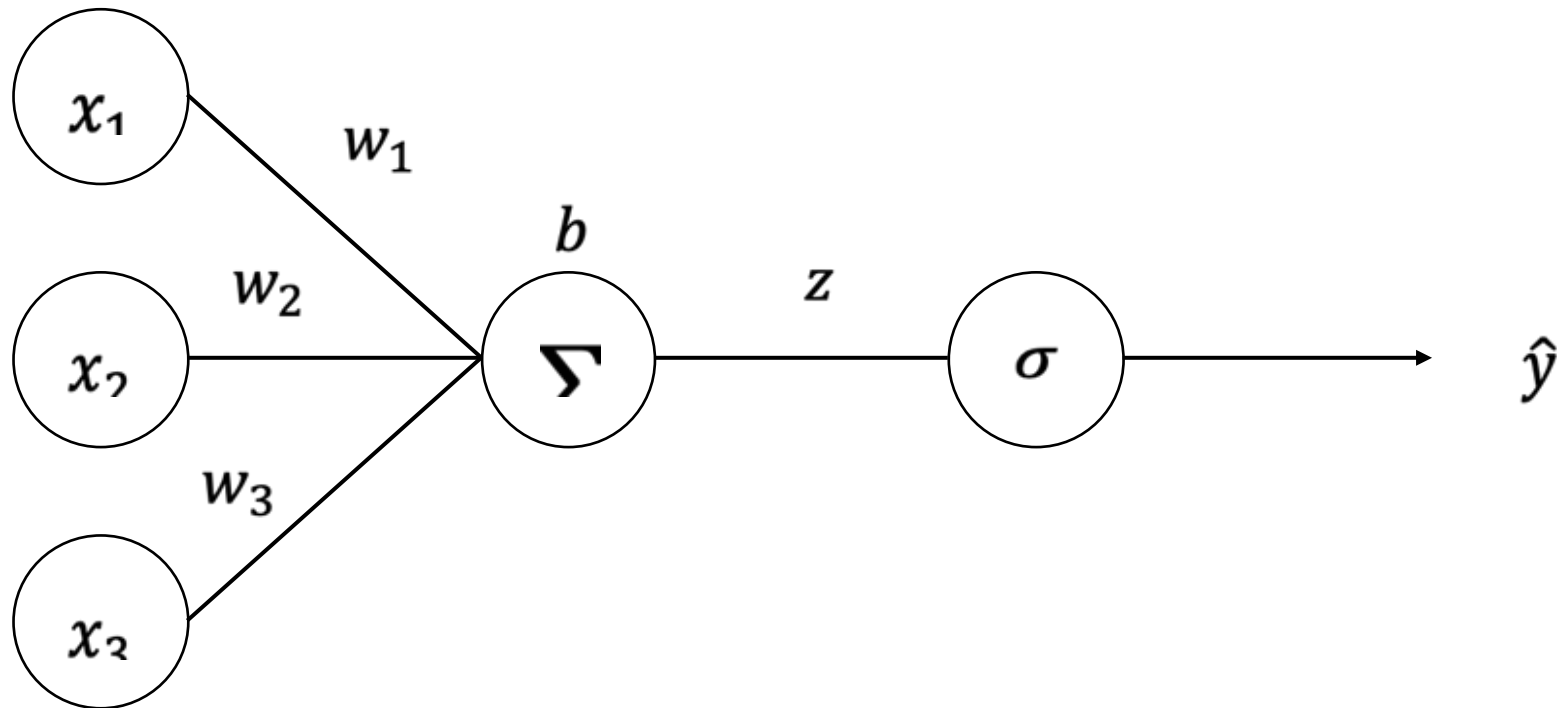
Single neuron model: The Perceptron



A Perceptron with a sigmoid activation function is mathematically equivalent to logistic regression

Single neuron model: The Perceptron

Example: predicting whether someone has a sleep disorder given the quality of the sleep, the daily amount of exercise and the stress level



Single neuron model: The Perceptron

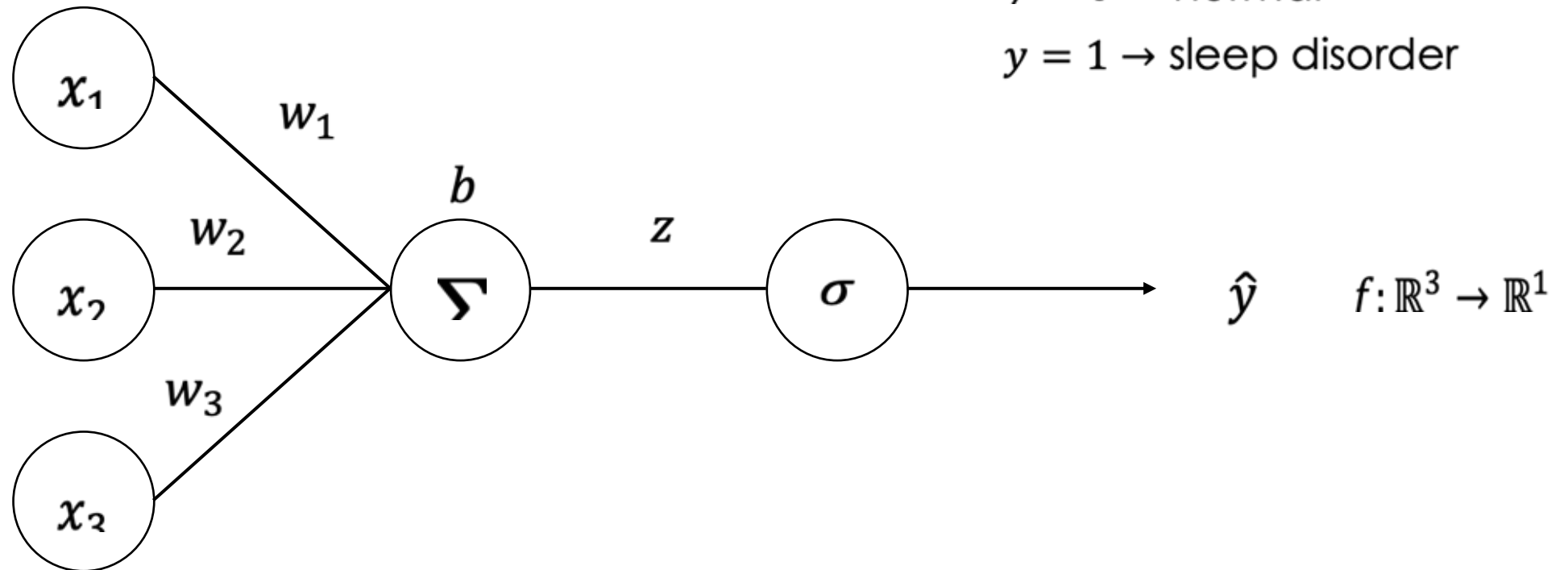
x_1 = sleep quality (1 - 10)

x_2 = exercise (min/daily)

x_3 = stress level (1 - 10)

$y = 0 \rightarrow$ normal

$y = 1 \rightarrow$ sleep disorder



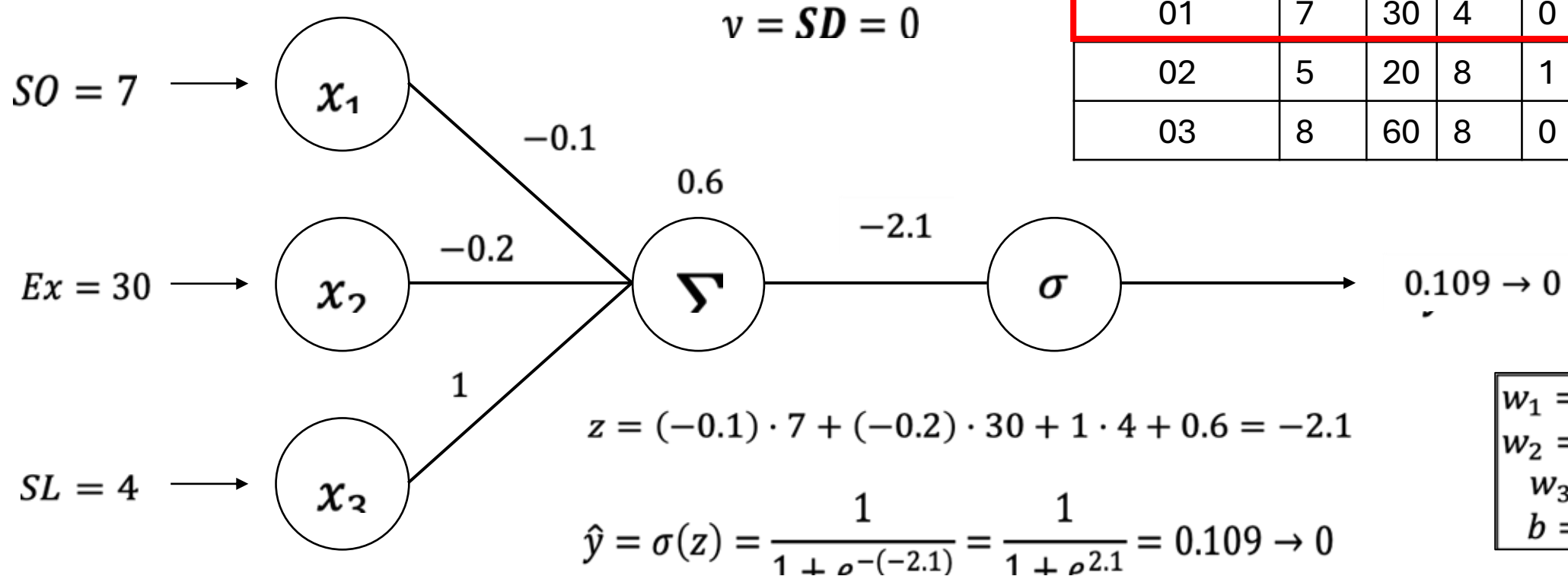
Single neuron model: The Perceptron

x_1 = sleep quality (1 - 10)

x_2 = exercise (min/daily)

x_3 = stress level (1 - 10)

Patient ID	SQ	Ex	SL	SD
01	7	30	4	0
02	5	20	8	1
03	8	60	8	0



$$\begin{aligned} w_1 &= -0.1 \\ w_2 &= -0.2 \\ w_3 &= 1 \\ b &= 0.6 \end{aligned}$$

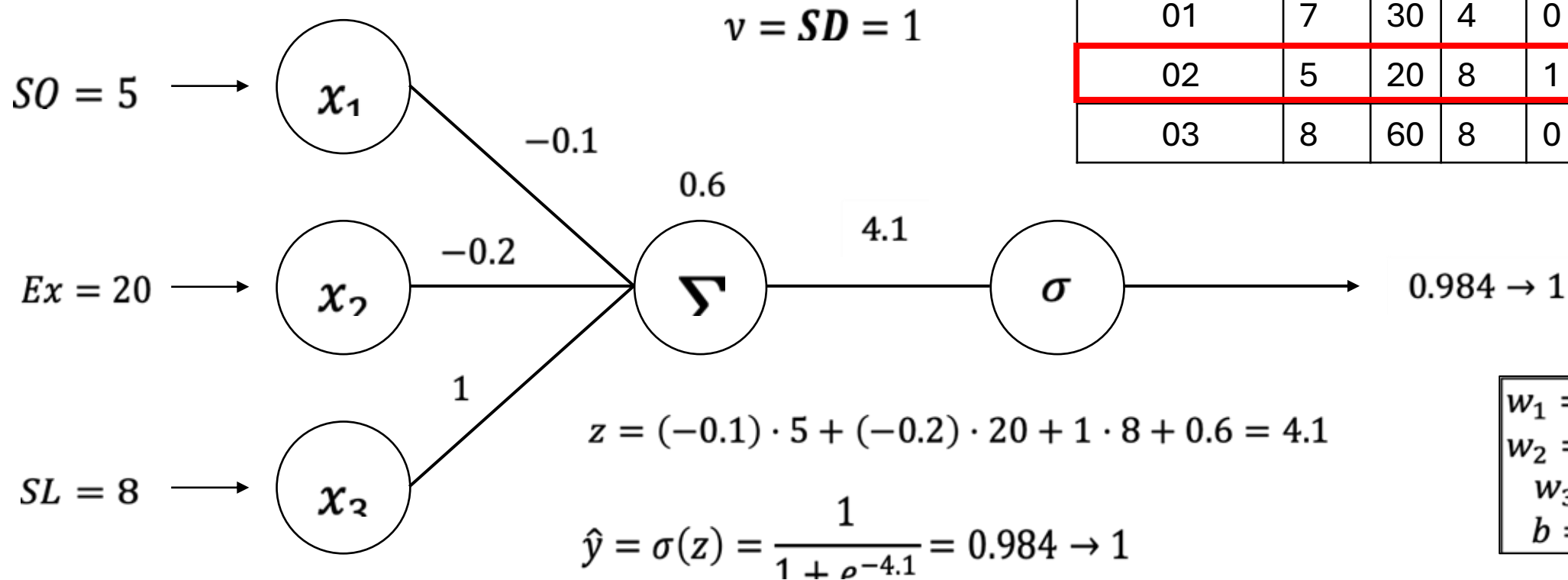
Single neuron model: The Perceptron

x_1 = sleep quality (1 - 10)

x_2 = exercise (min/daily)

x_3 = stress level (1 - 10)

Patient ID	SQ	Ex	SL	SD
01	7	30	4	0
02	5	20	8	1
03	8	60	8	0



$$\begin{aligned} w_1 &= -0.1 \\ w_2 &= -0.2 \\ w_3 &= 1 \\ b &= 0.6 \end{aligned}$$

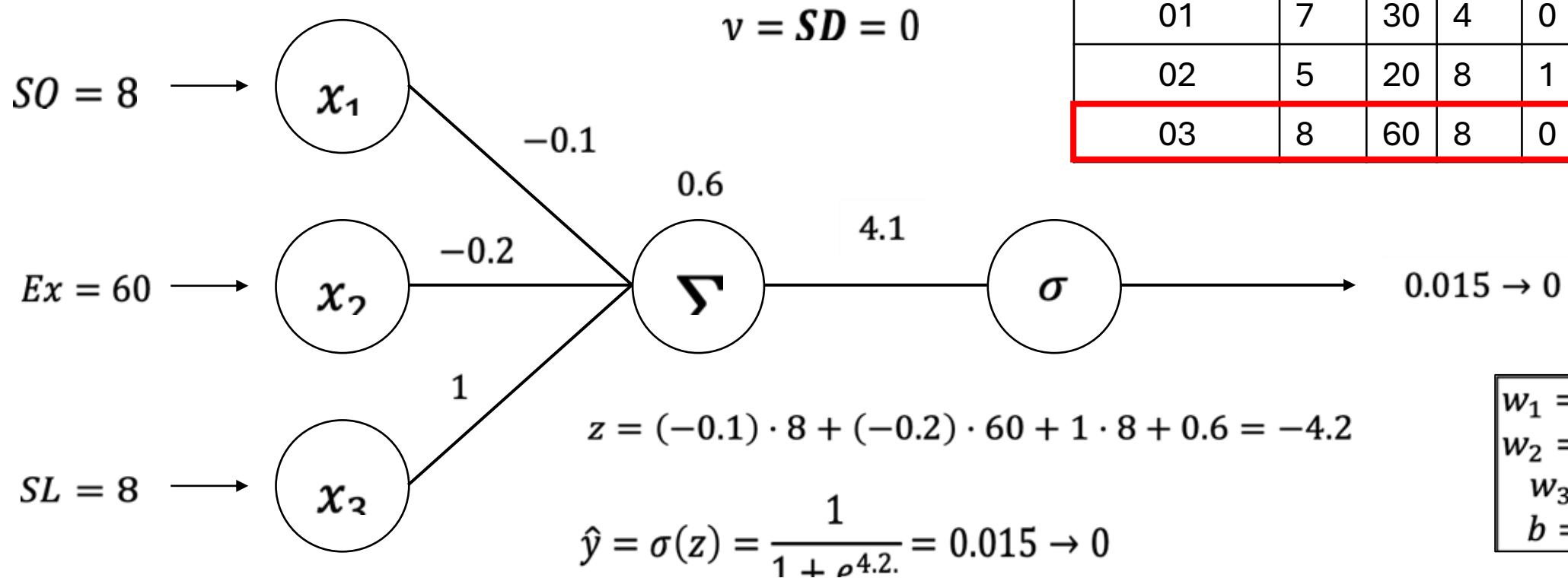
Single neuron model: The Perceptron

x_1 = sleep quality (1 - 10)

x_2 = exercise (min/daily)

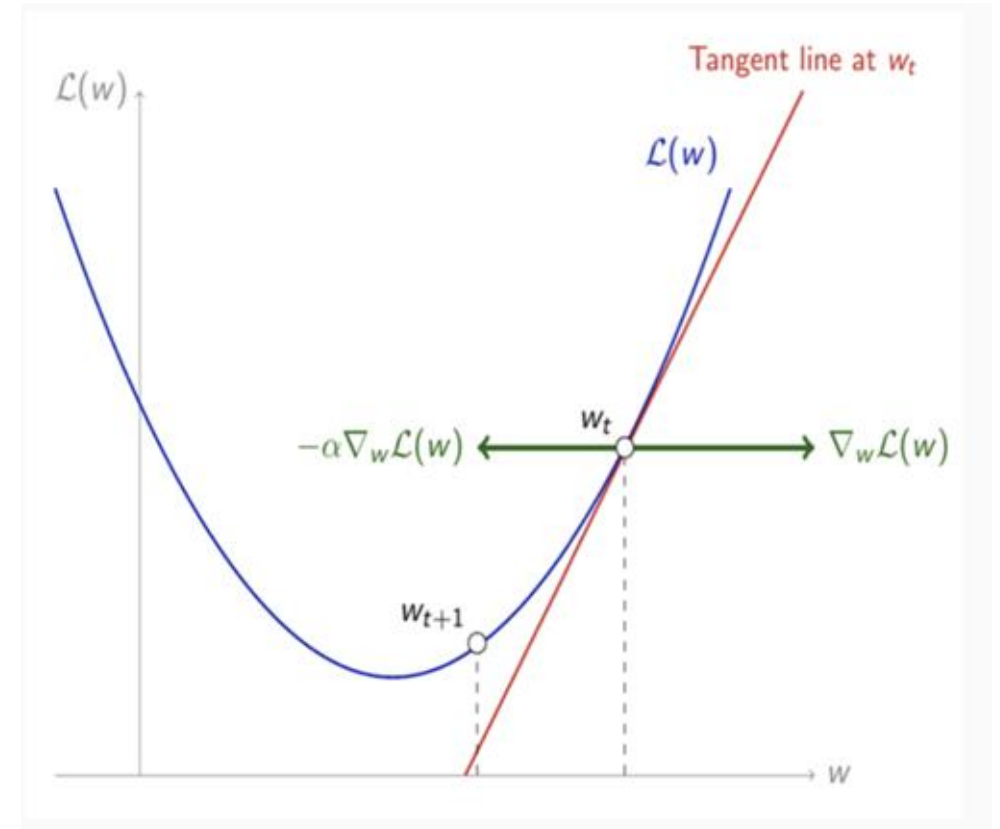
x_3 = stress level (1 - 10)

Patient ID	SQ	Ex	SL	SD
01	7	30	4	0
02	5	20	8	1
03	8	60	8	0

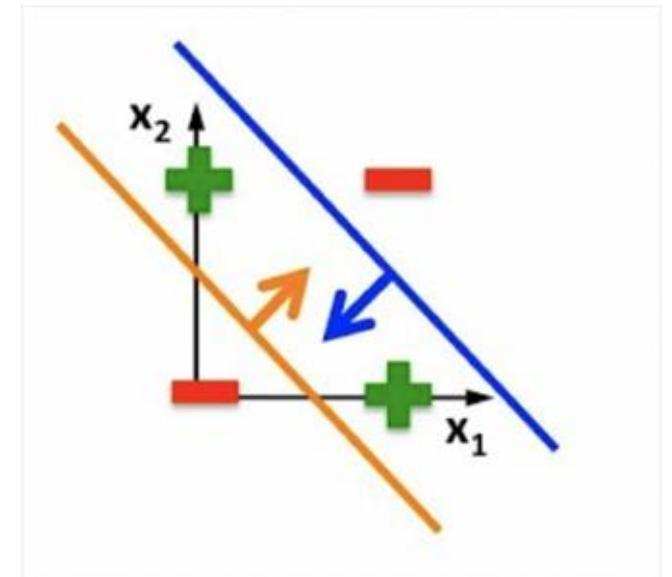
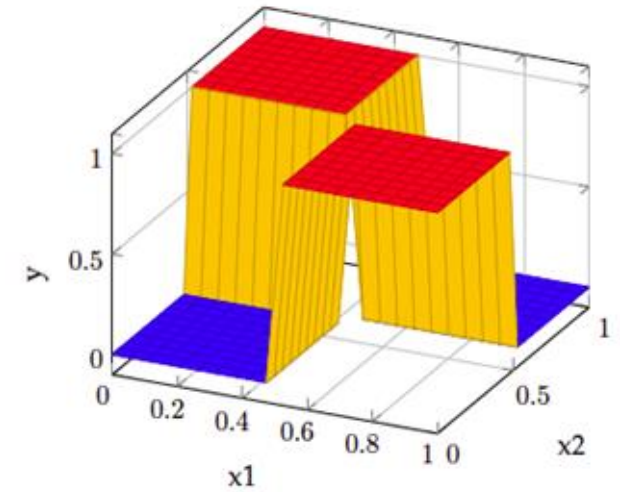
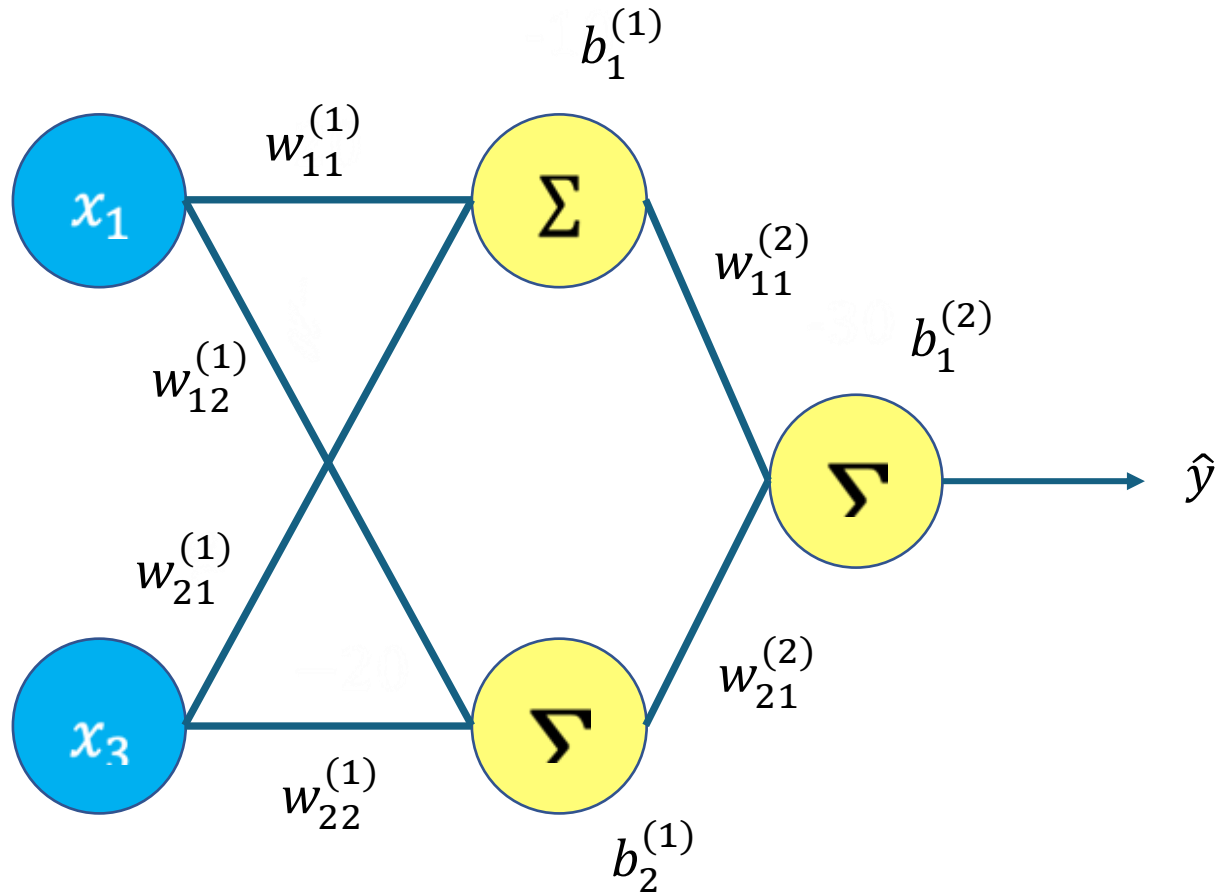


Training a Perceptron

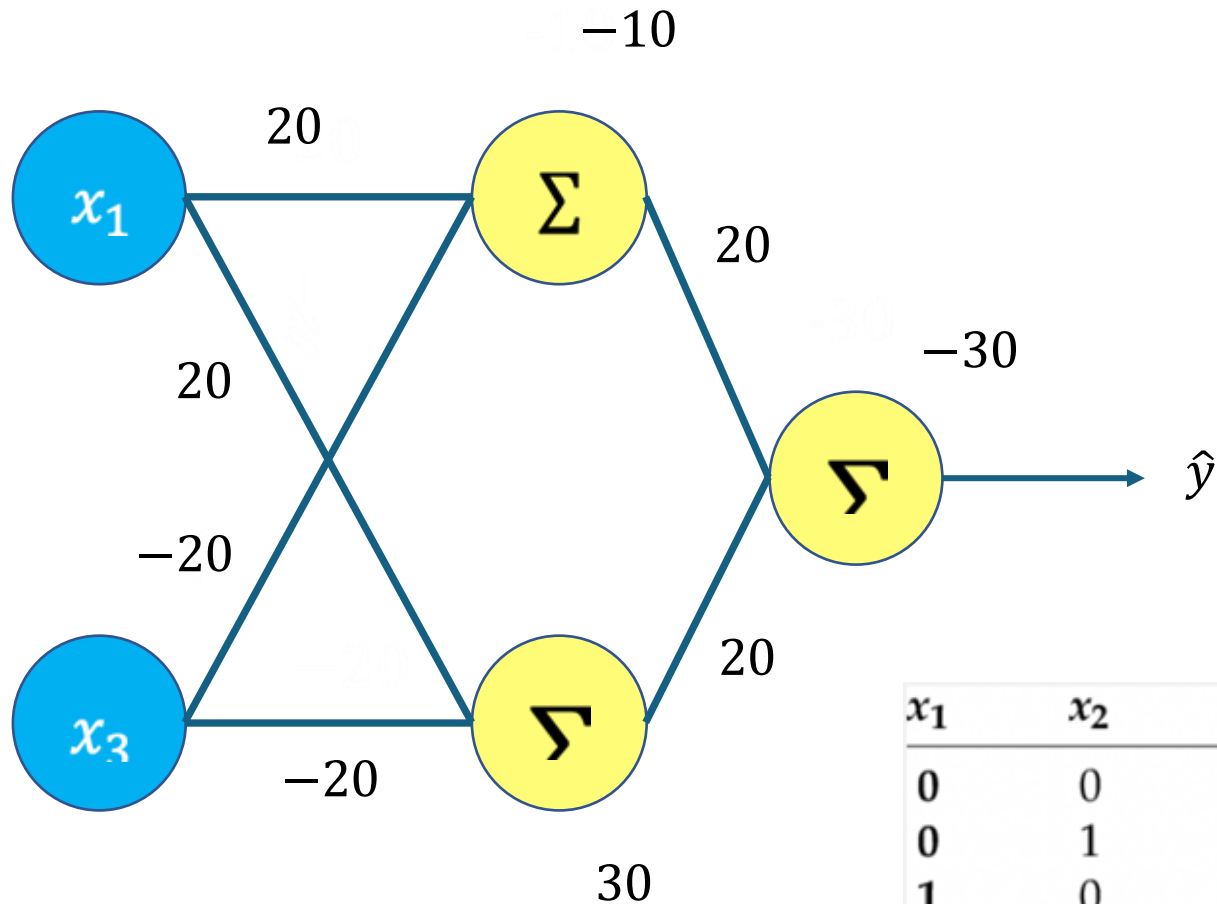
- How do we find w_1, w_2, w_3 and b ?
- Using the training data, we will gradually adjust the values of the parameters such that we minimize the errors in the predictions
- Optimization algorithm: **gradient descent** $\theta_i := \theta_i - \alpha \frac{\partial \mathcal{L}}{\partial \theta_i}$



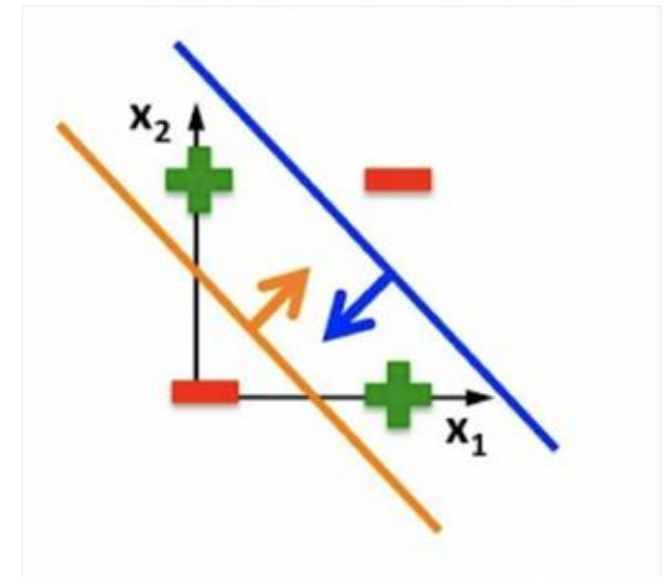
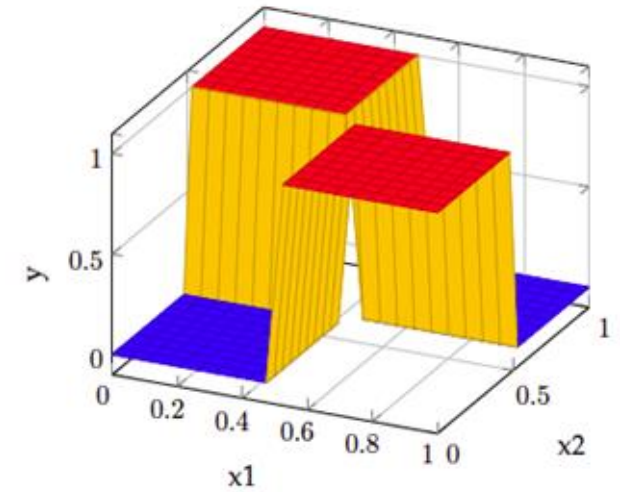
The perceptron and the XOR problem



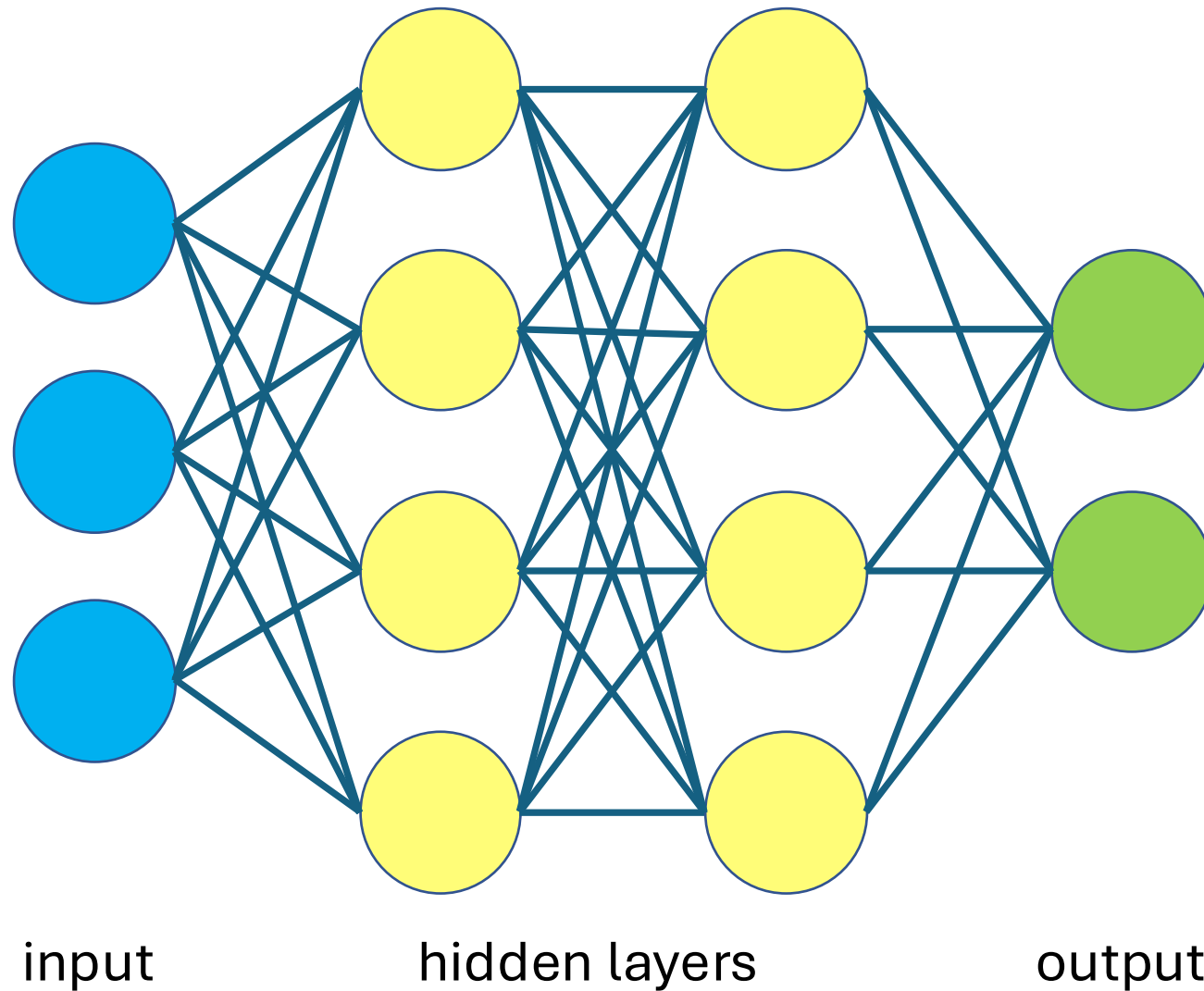
The perceptron and the XOR problem



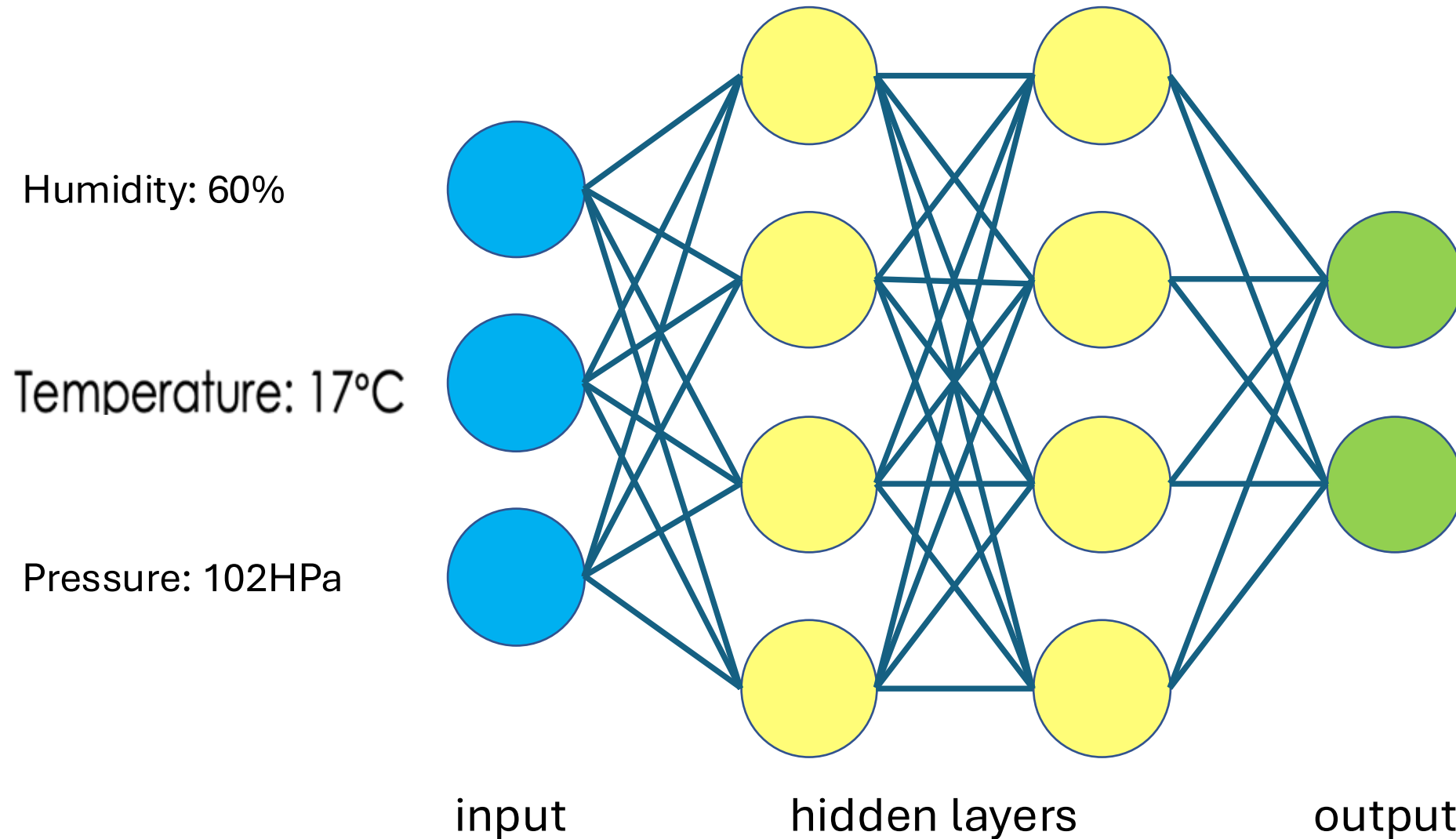
x_1	x_2	$x_1 \text{ XOR } x_2$
0	0	0
0	1	1
1	0	1
1	1	0



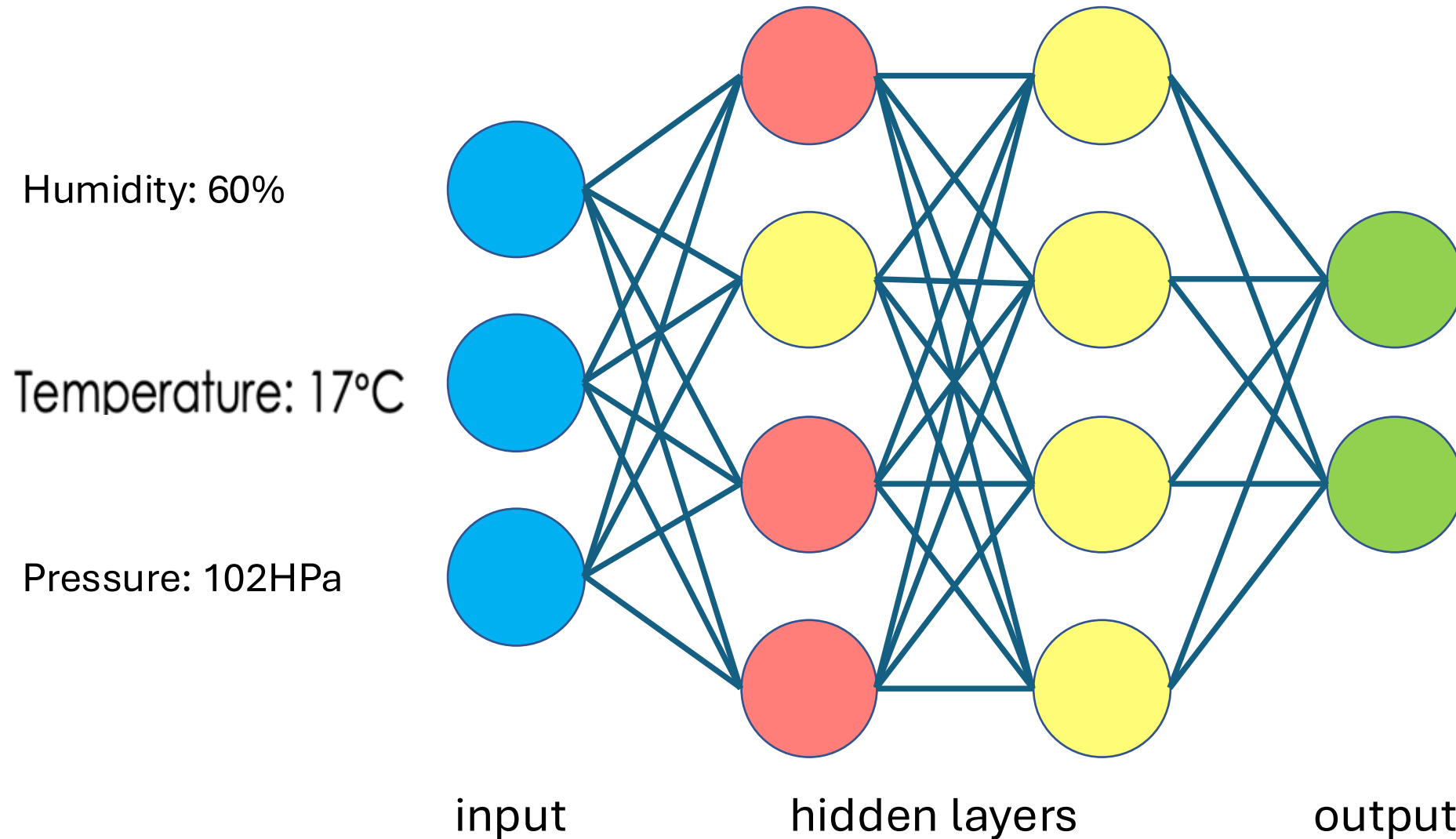
Artificial Neural Network



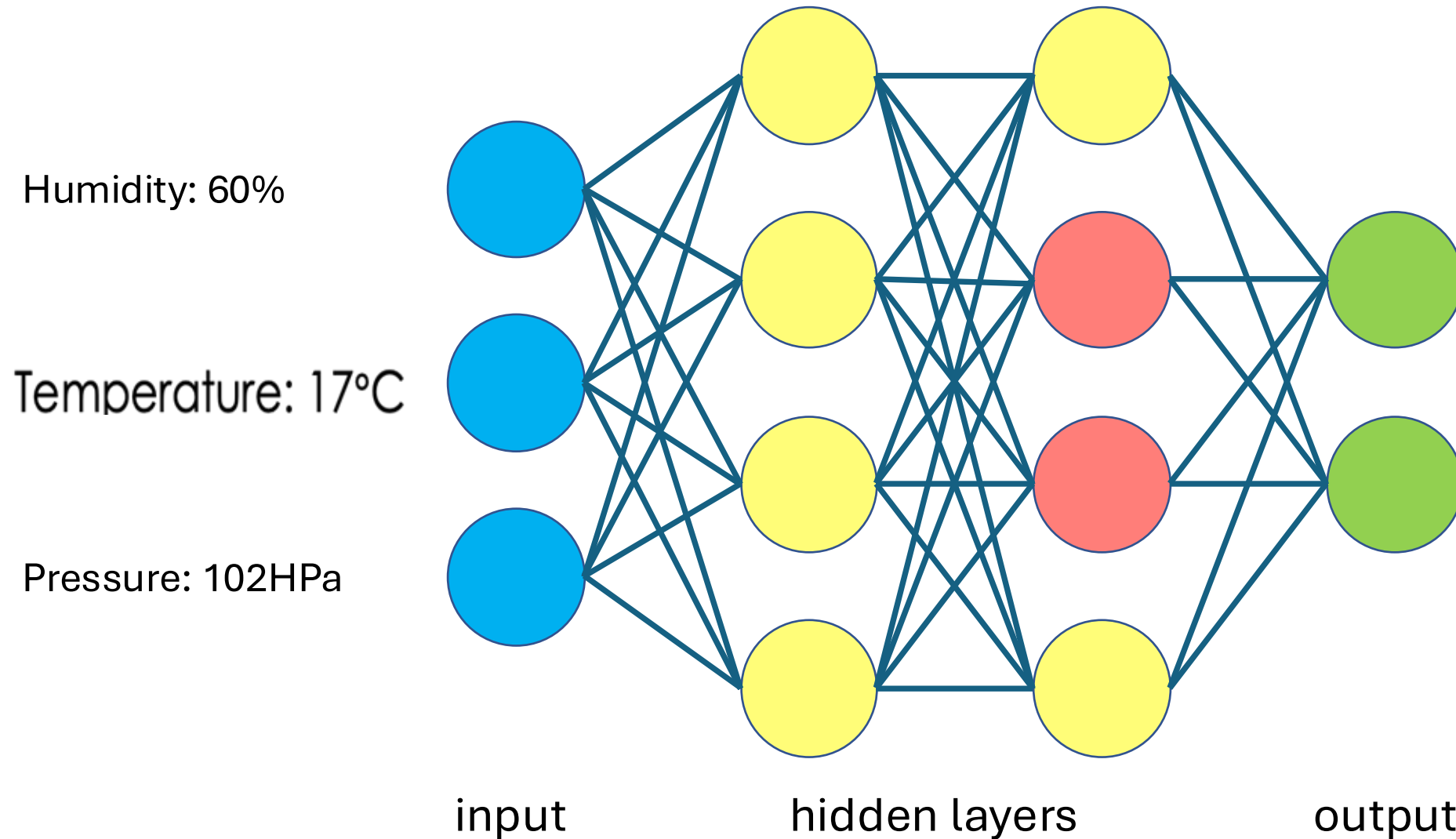
Artificial Neural Network



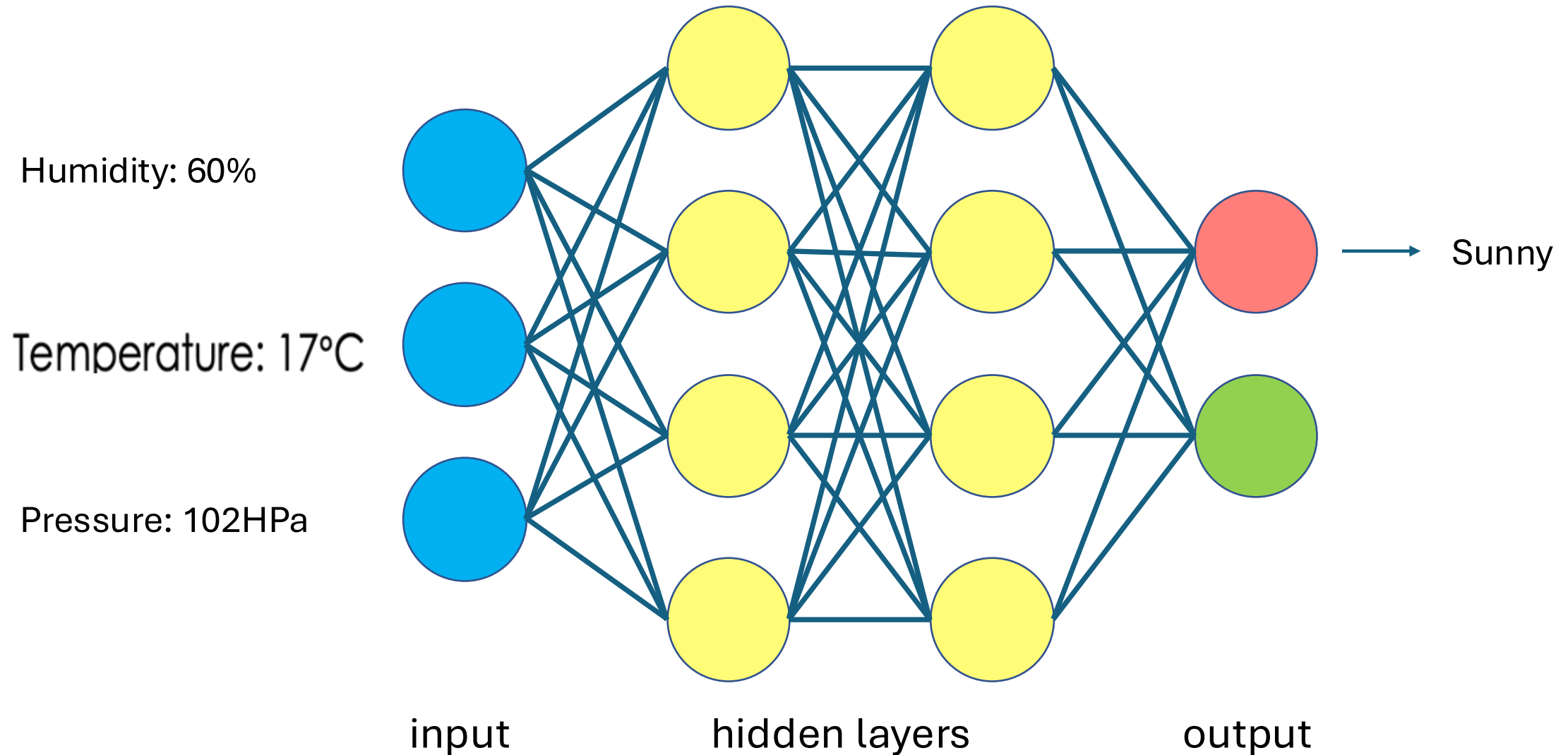
Artificial Neural Network



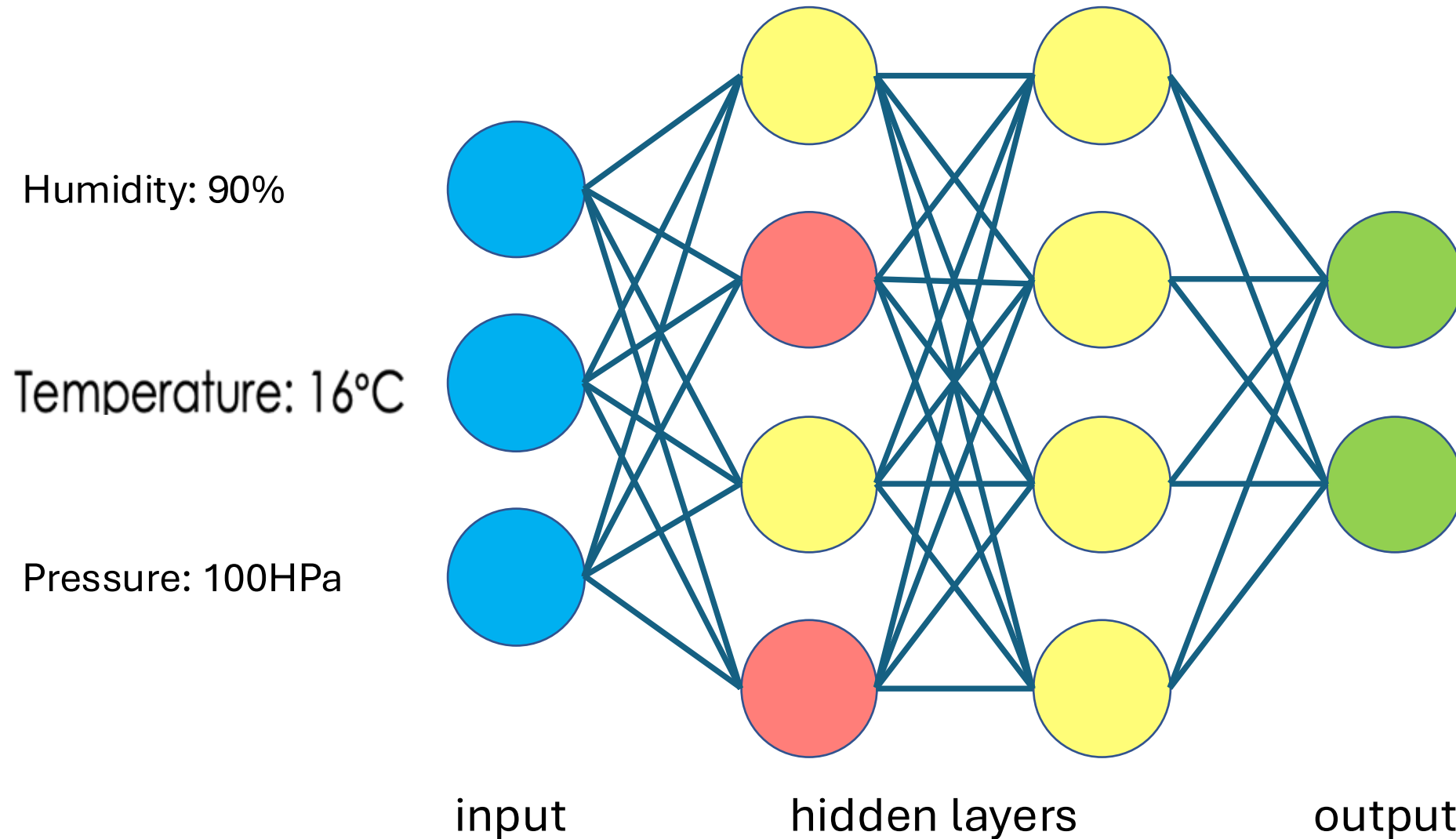
Artificial Neural Network



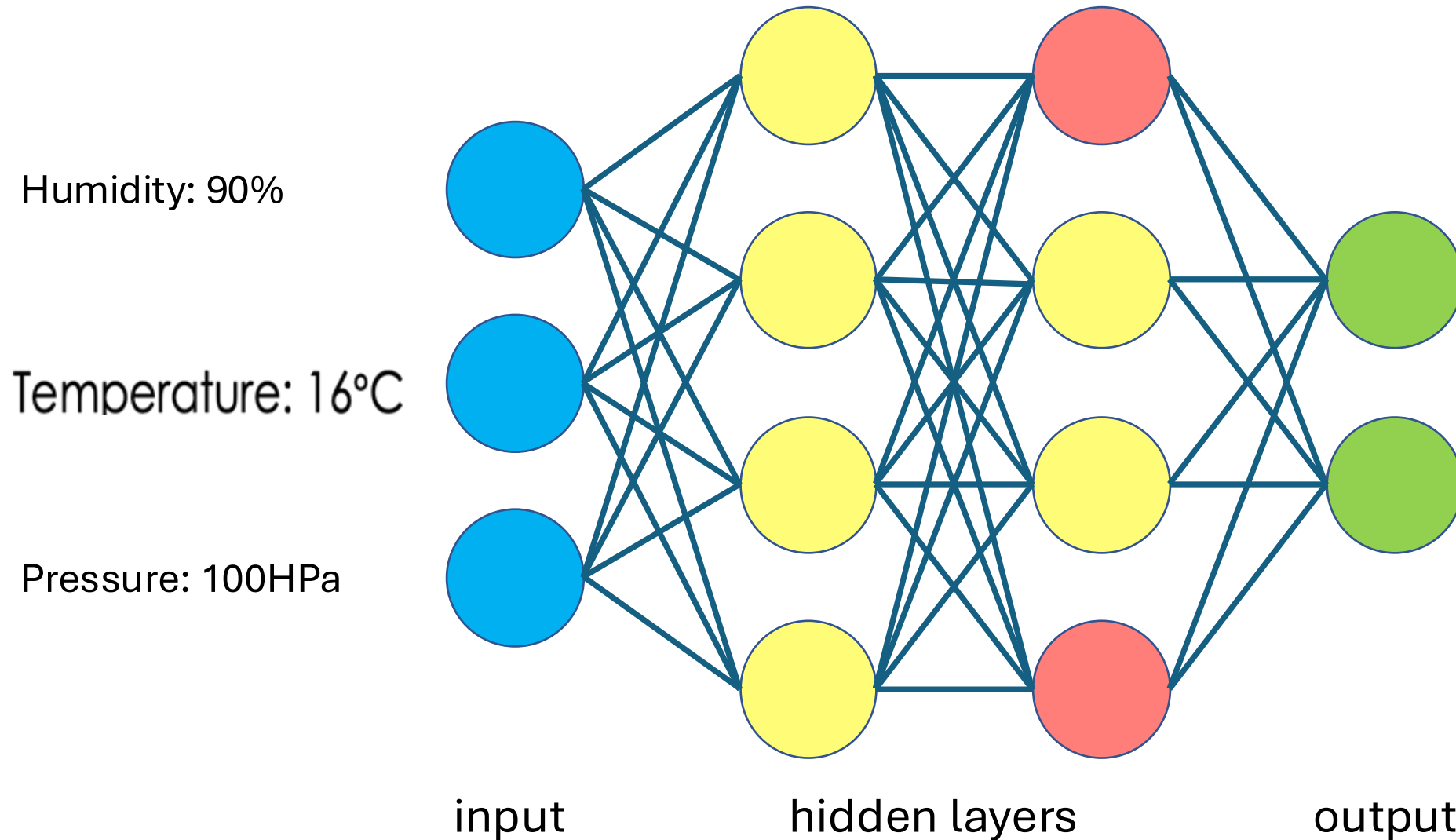
Artificial Neural Network



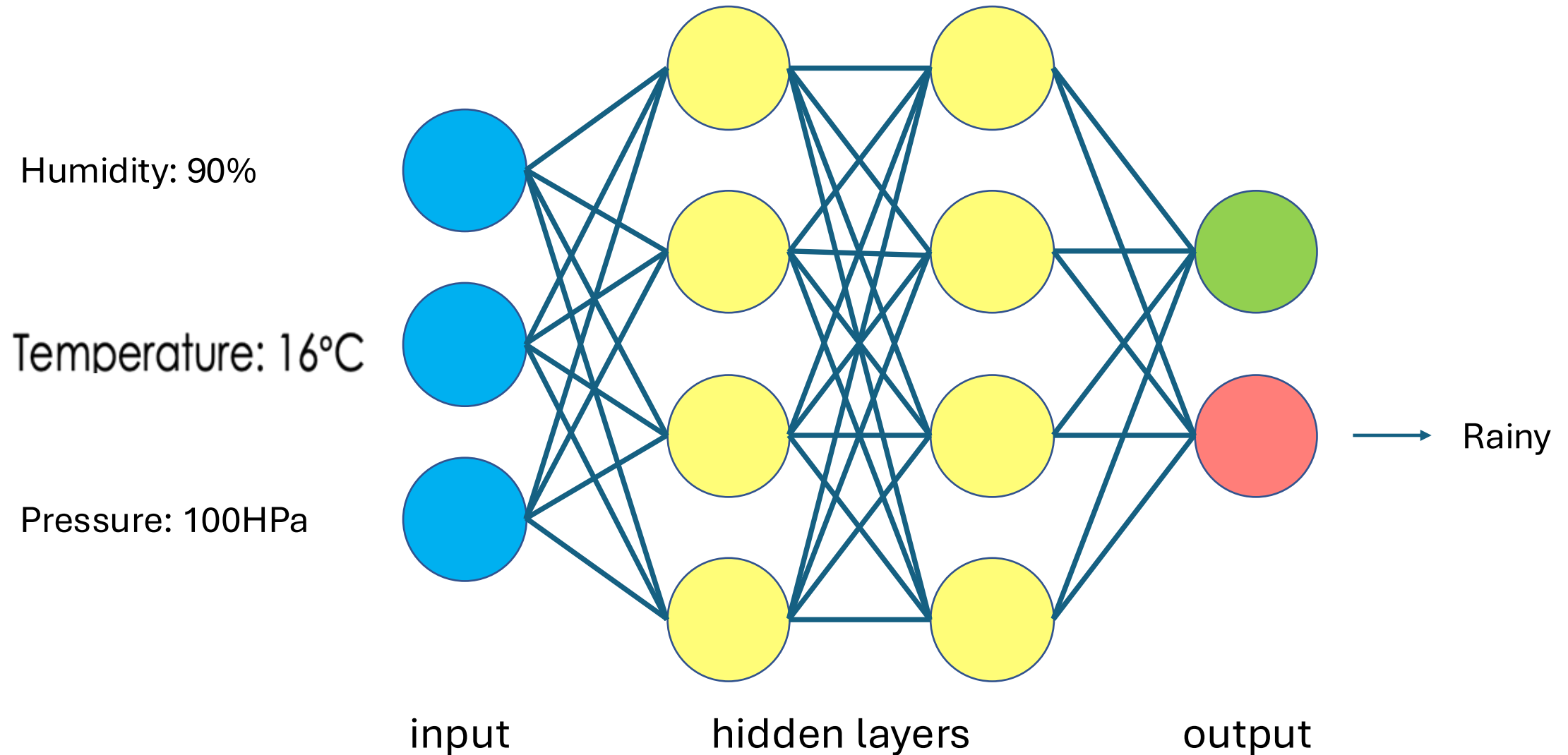
Artificial Neural Network



Artificial Neural Network

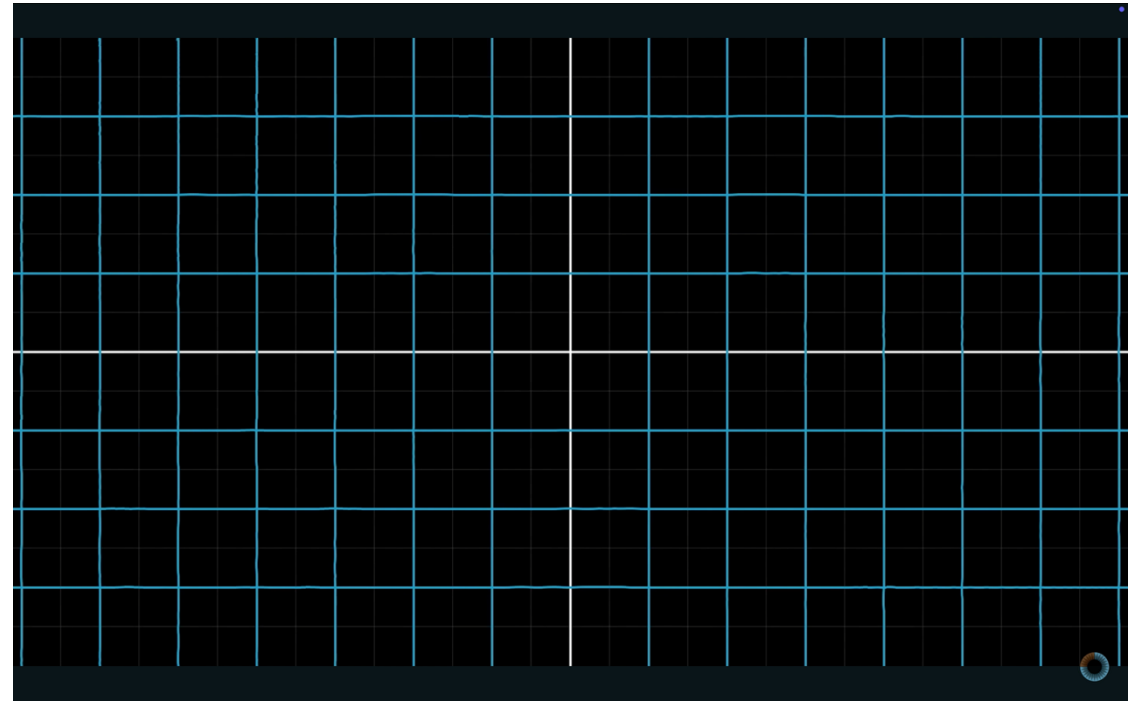
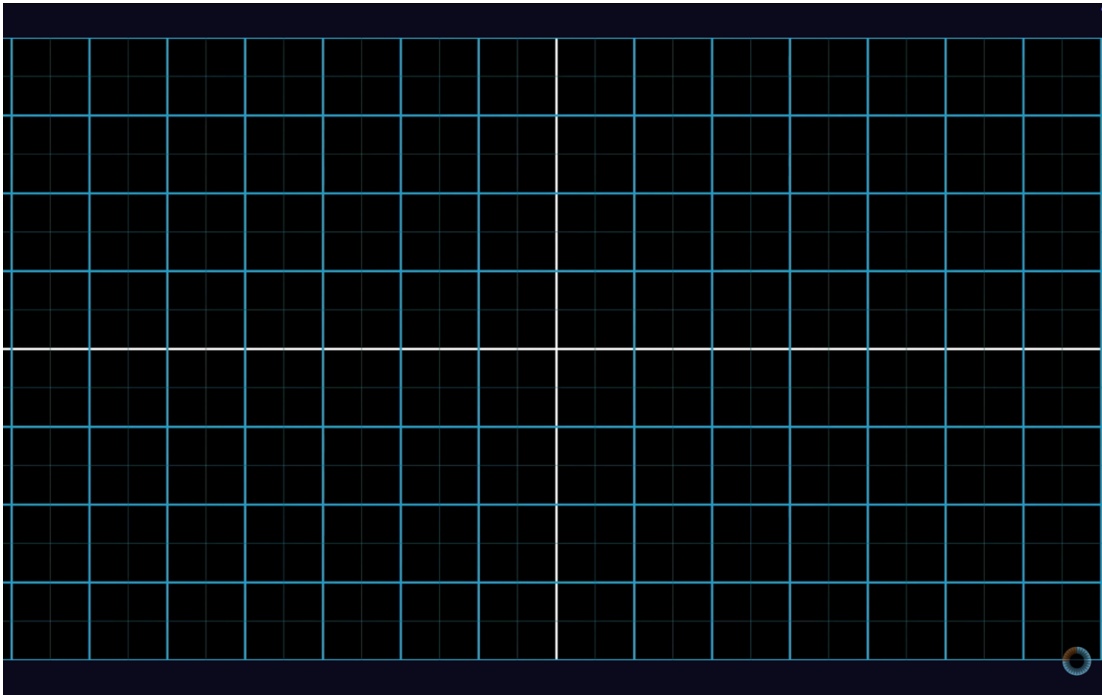


Artificial Neural Network



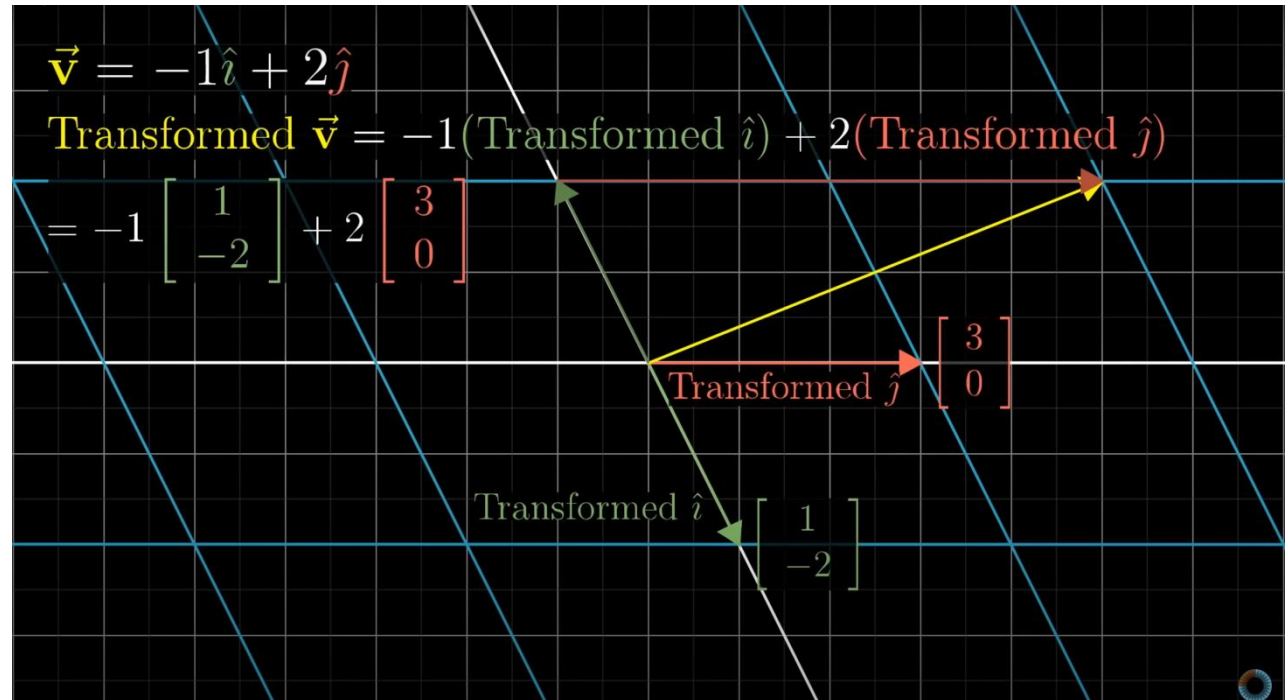
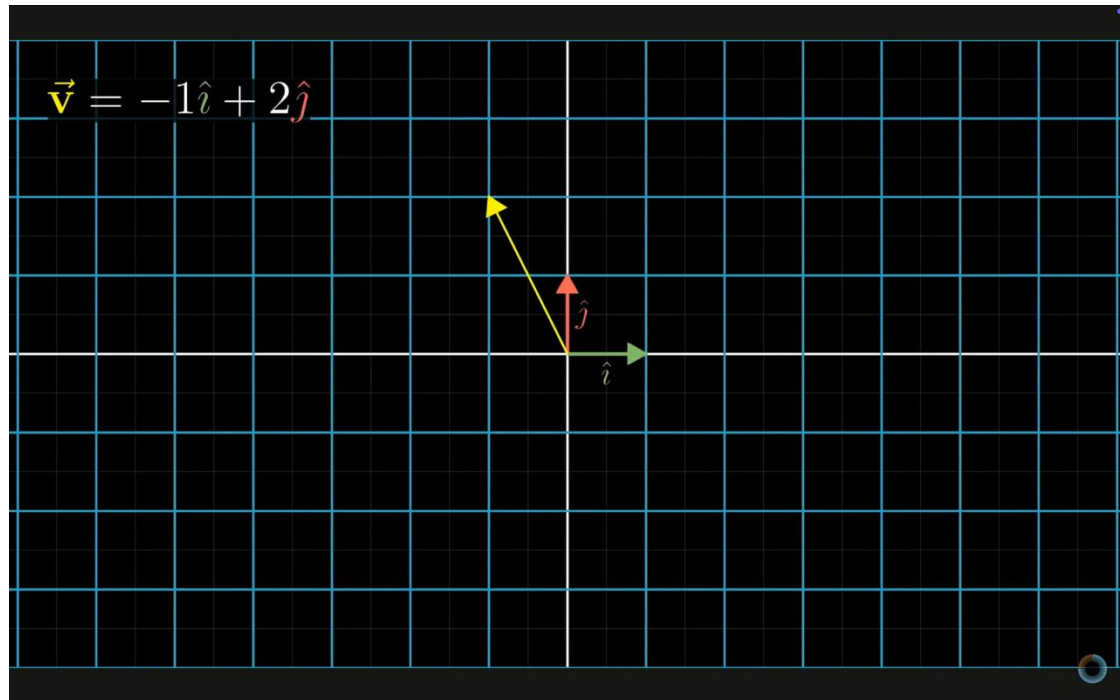
Matrices as linear transformations

- Center remains at origin
- Lines remain parallel

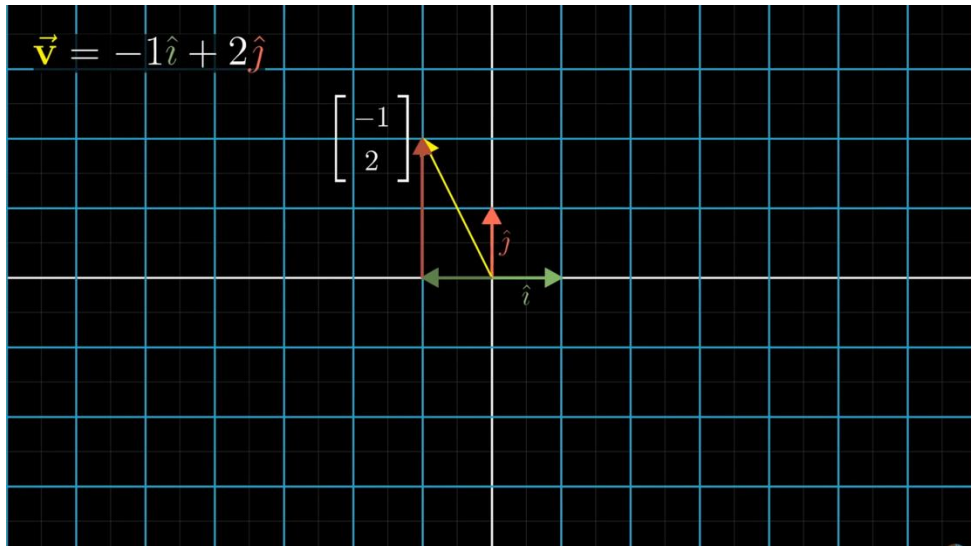


<https://www.youtube.com/@3blue1brown>

Matrices as linear transformations



Matrices as linear transformations



“2x2 Matrix”

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

Where $\hat{i}$ lands Where $\hat{j}$ lands

“2x2 Matrix”

$$\begin{bmatrix} 3 & 2 \\ -2 & 1 \end{bmatrix} \quad \begin{bmatrix} 5 \\ 7 \end{bmatrix}$$

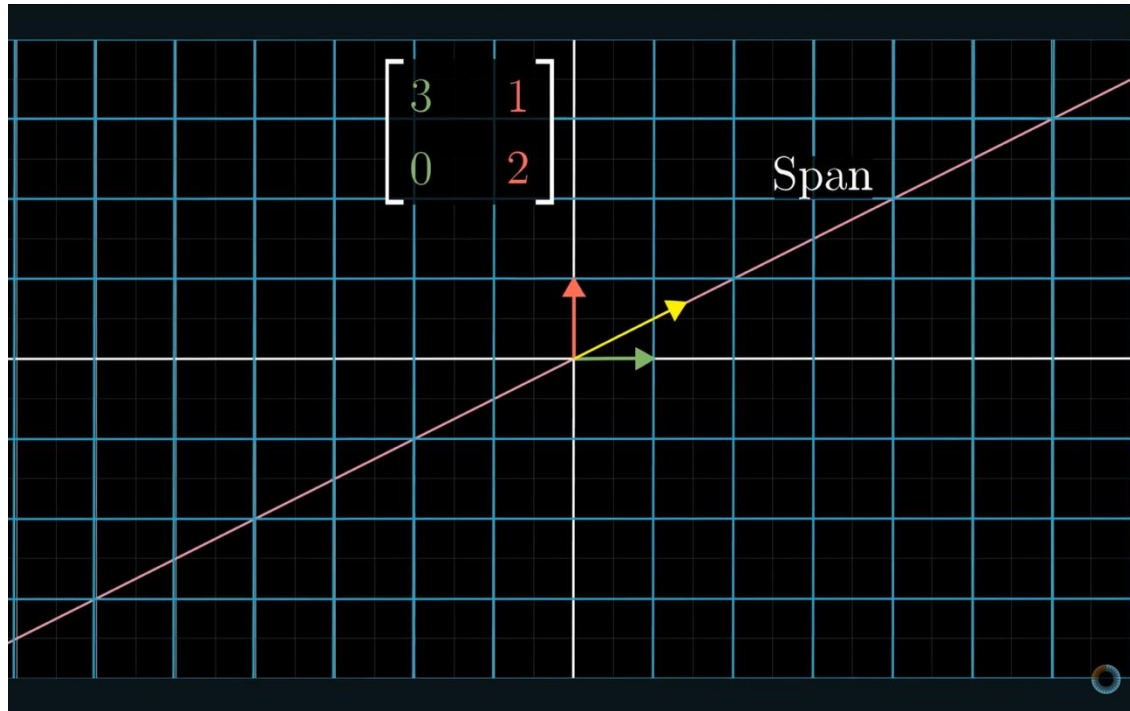
$$5 \begin{bmatrix} 3 \\ -2 \end{bmatrix} + 7 \begin{bmatrix} 2 \\ 1 \end{bmatrix}$$

“2x2 Matrix”

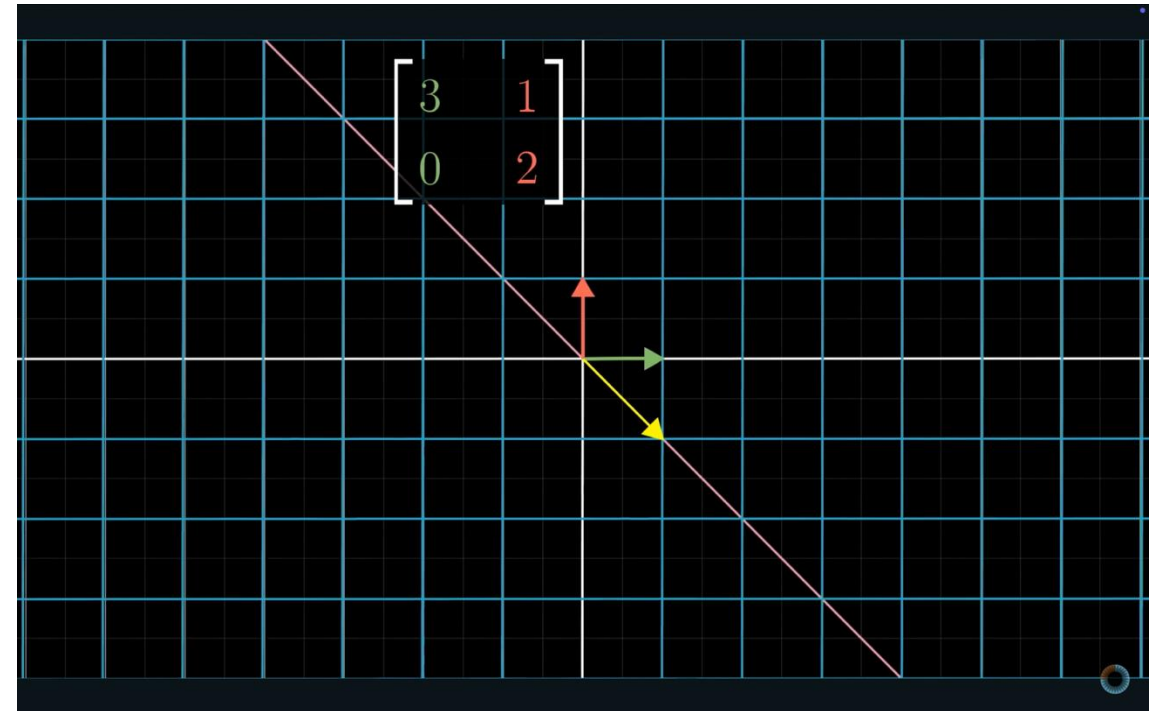
$$\begin{bmatrix} a & b \\ c & d \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix} = x \begin{bmatrix} a \\ c \end{bmatrix} + y \begin{bmatrix} b \\ d \end{bmatrix} = \begin{bmatrix} ax + by \\ cx + dy \end{bmatrix}$$

Matrices as linear transformations

- Eigenvectors are only scaled by the eigenvalues

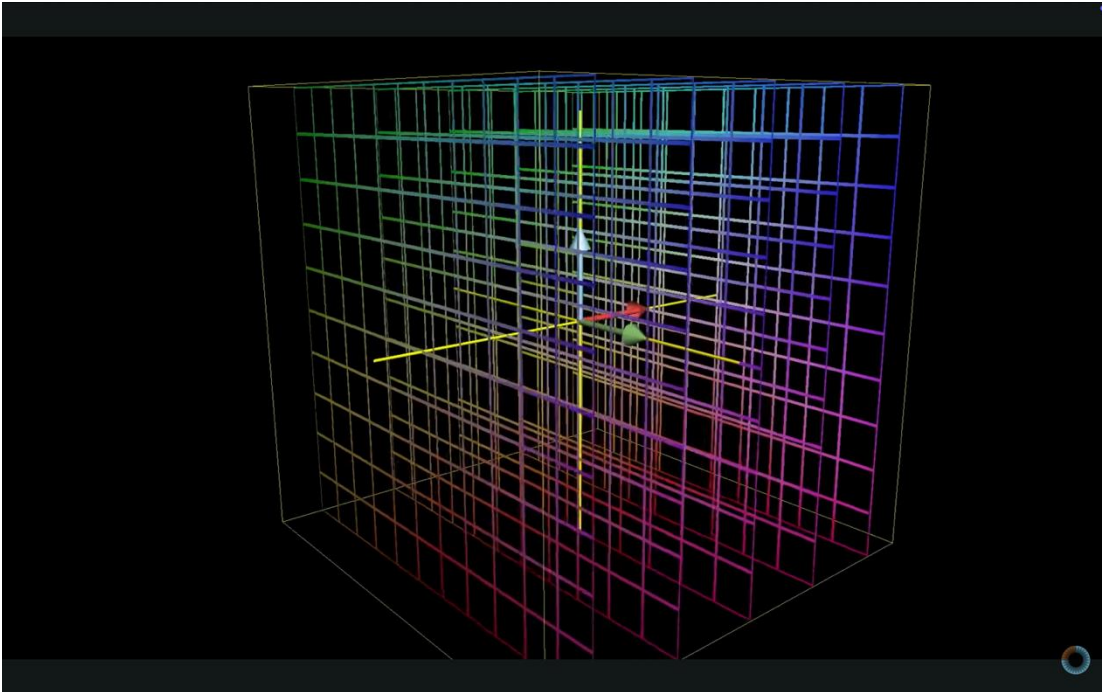


No eigenvector



Eigenvector

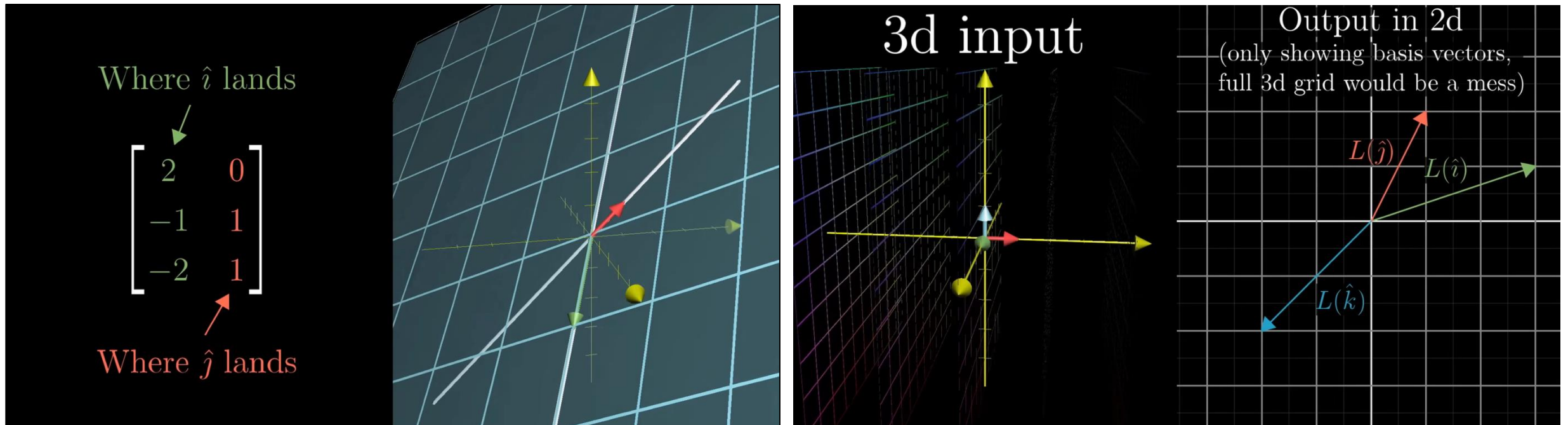
Linear transformations in 3D



$$\underbrace{\begin{bmatrix} 0 & 1 & 2 \\ 3 & 4 & 5 \\ 6 & 7 & 8 \end{bmatrix}}_{\text{Transformation}} \underbrace{\begin{bmatrix} x \\ y \\ z \end{bmatrix}}_{\text{Input vector}} = x \underbrace{\begin{bmatrix} 0 \\ 3 \\ 6 \end{bmatrix}}_{\text{Output vector}} + y \underbrace{\begin{bmatrix} 1 \\ 4 \\ 7 \end{bmatrix}}_{\text{Output vector}} + z \underbrace{\begin{bmatrix} 2 \\ 5 \\ 8 \end{bmatrix}}_{\text{Output vector}}$$

Non-square matrices as transformations between dimensions

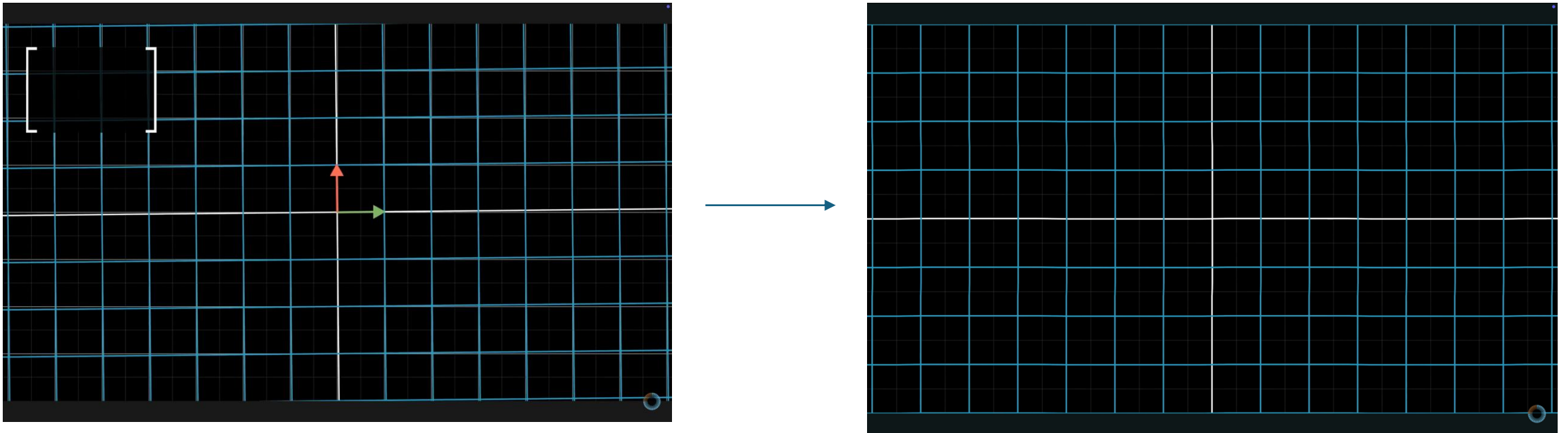
- What transformation does a 3x2 matrix represent? And a 3x2?



Matrix multiplication is linear $\rightarrow$ we can repeatedly apply transformations to a vector between any number of dimensions

Neural Networks

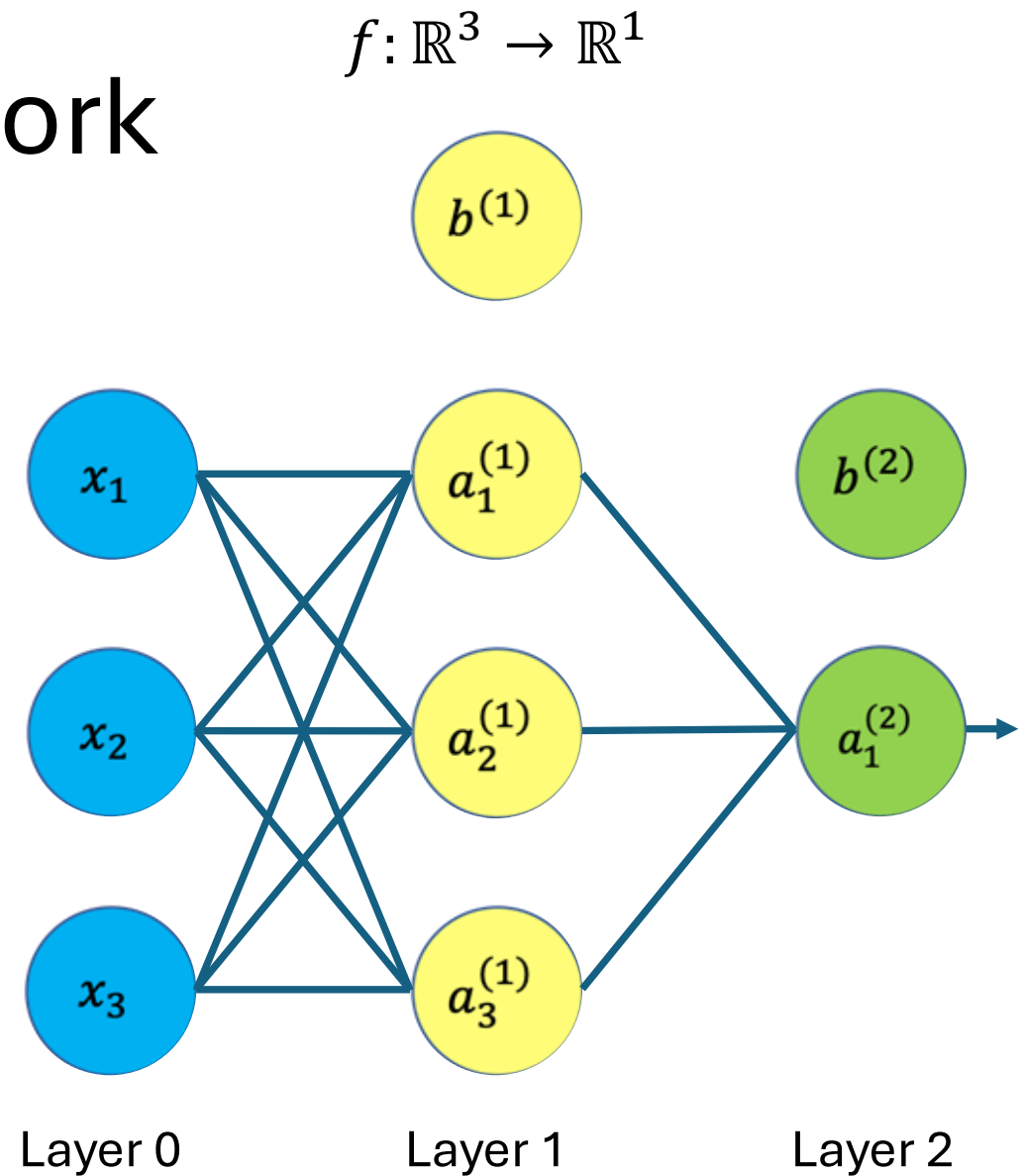
- Each layer of (trainable) weights applies a linear transformation to the input vector, followed by a non-linear activation. The output of a given layer becomes the input to the following layer



Building a neural network

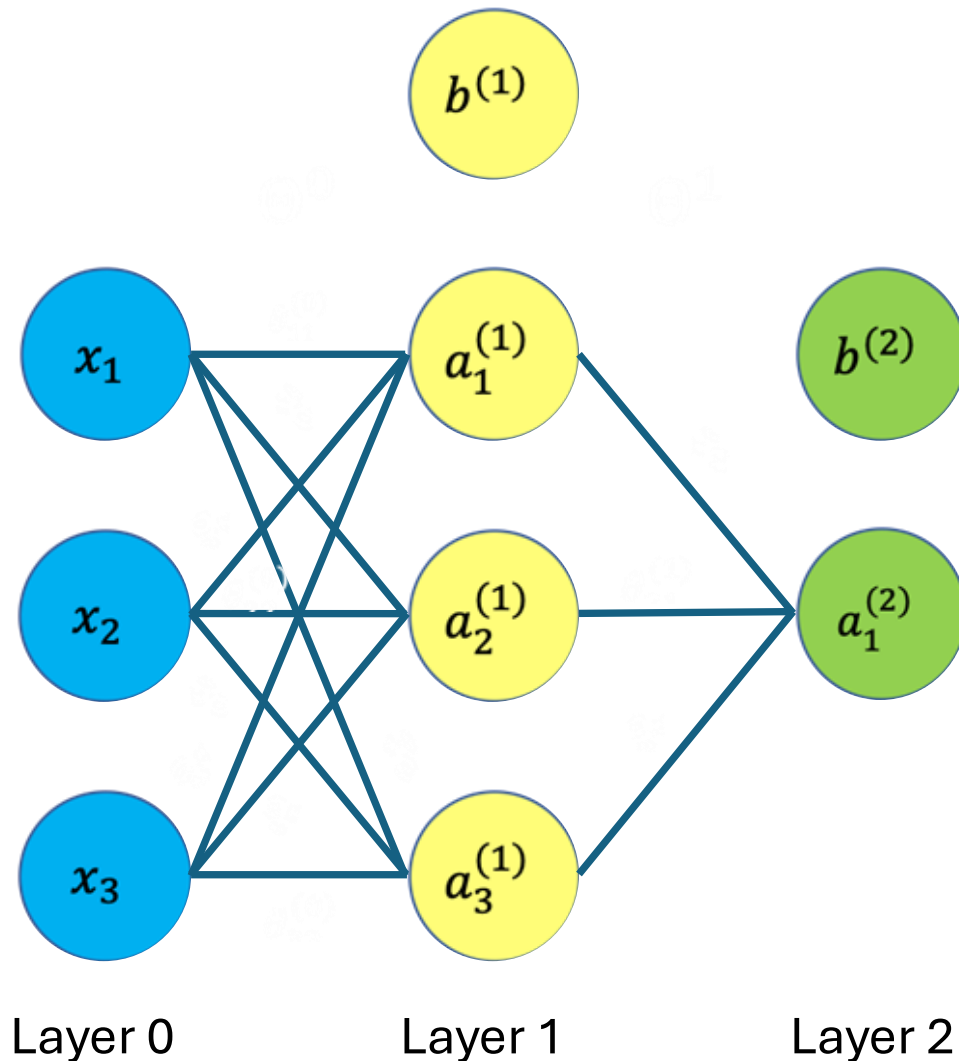
- **Input layer (layer 0):** this layer consists of the input features x_1 , x_2 and x_3
- **Hidden layer (layer 1):** contains values that we don't observe in the training set
- **Output layer (layer 2):** it has a neuron/units (can have more units) that outputs the final value computed by the hypothesis

$$f: \mathbb{R}^M \rightarrow \mathbb{R}^N$$



Building a neural network

$$f: \mathbb{R}^3 \rightarrow \mathbb{R}^1$$



$$a_1^{(1)} = g(\theta_{11}^{(0)} x_1 + \theta_{21}^{(0)} x_2 + \theta_{31}^{(0)} x_3 + b_1^{(1)})$$

$$a_2^{(1)} = g(\theta_{12}^{(0)} x_1 + \theta_{22}^{(0)} x_2 + \theta_{32}^{(0)} x_3 + b_2^{(1)})$$

$$a_3^{(1)} = g(\theta_{13}^{(0)} x_1 + \theta_{23}^{(0)} x_2 + \theta_{33}^{(0)} x_3 + b_3^{(1)})$$

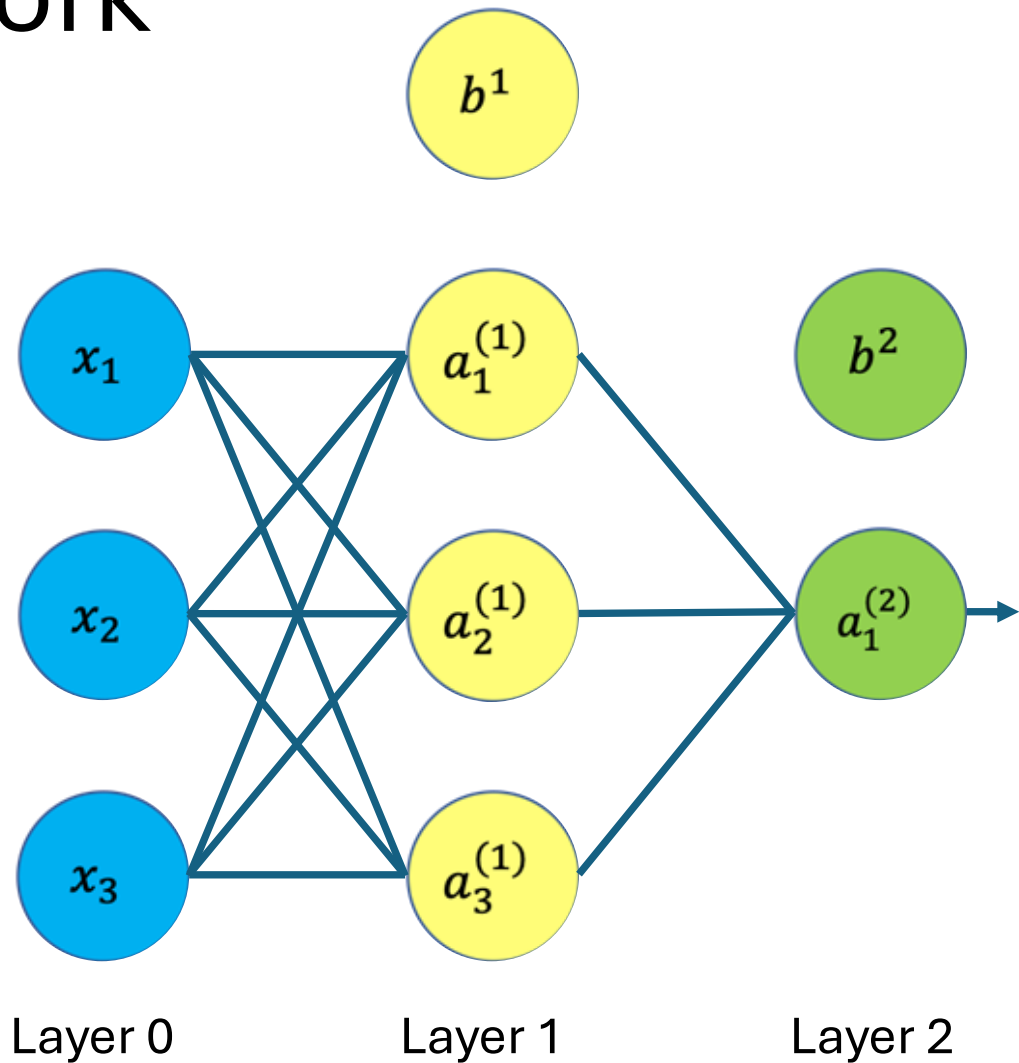
$$h_{\Theta, b}(x) = a_1^{(2)} = \theta_{11}^{(1)} a_1^{(1)} + \theta_{21}^{(1)} a_2^{(1)} + \theta_{31}^{(1)} a_3^{(1)} + b_1^{(2)}$$

$$\Theta^{(0)} = \begin{pmatrix} \theta_{11}^{(0)} & \theta_{12}^{(0)} & \theta_{13}^{(0)} \\ \theta_{21}^{(0)} & \theta_{22}^{(0)} & \theta_{23}^{(0)} \\ \theta_{31}^{(0)} & \theta_{32}^{(0)} & \theta_{33}^{(0)} \end{pmatrix} \quad x = \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} \quad b^{(1)} = \begin{bmatrix} b_1^{(1)} \\ b_2^{(1)} \\ b_3^{(1)} \end{bmatrix}$$

$$a^{(1)} = \begin{bmatrix} a_1^{(1)} \\ a_2^{(1)} \\ a_3^{(1)} \end{bmatrix} = (\Theta^{(0)})^T \cdot x + b^{(1)}$$

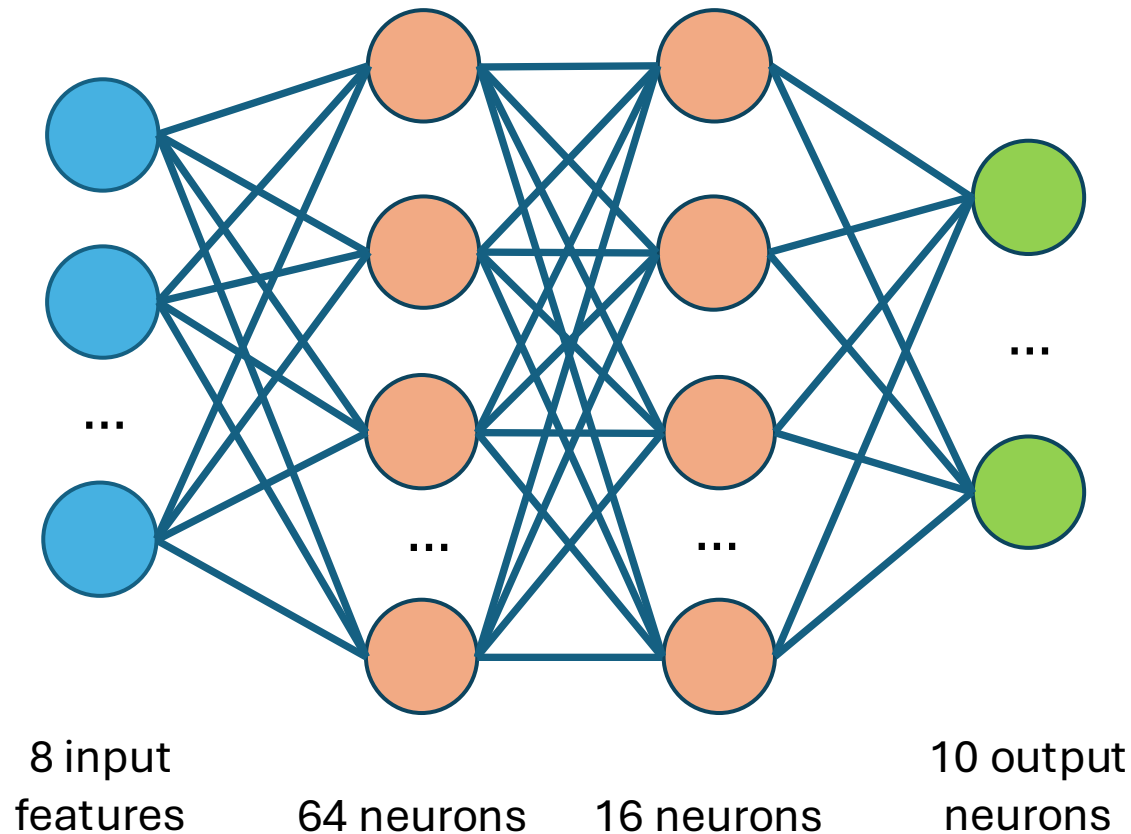
Building a neural network

- In this example the computation of $h_{\Theta,b}(x)$ is known as **forward propagation**
- It starts with the multiplication of the input units with the weights of the first layer.
- It propagates to the hidden layer computing the activation function of the second layer.
- The propagation follows with the activation of the output layer



Number of parameters

- How many parameters does the following neural network have?



- 3 layer neural network with 2 hidden layers
- Sigmoid activation functions
- Softmax on the last layer

Universal approximation theorem

- A feedforward neural network with a single hidden layer containing a finite number of neurons can approximate any continuous function on compact subsets of $\mathbb{R}^n$, given appropriate activation functions
- Arbitrary width
- Arbitrary depth

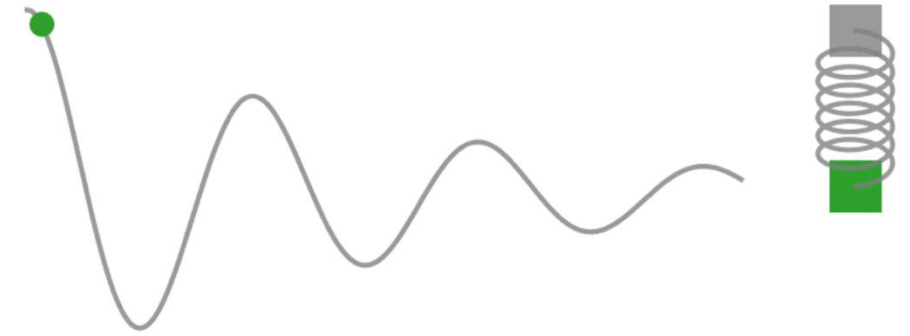
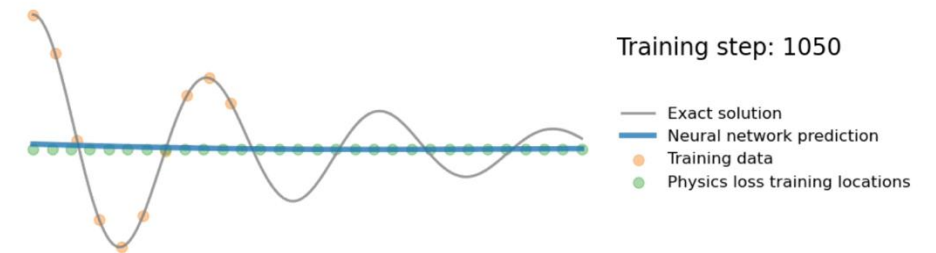


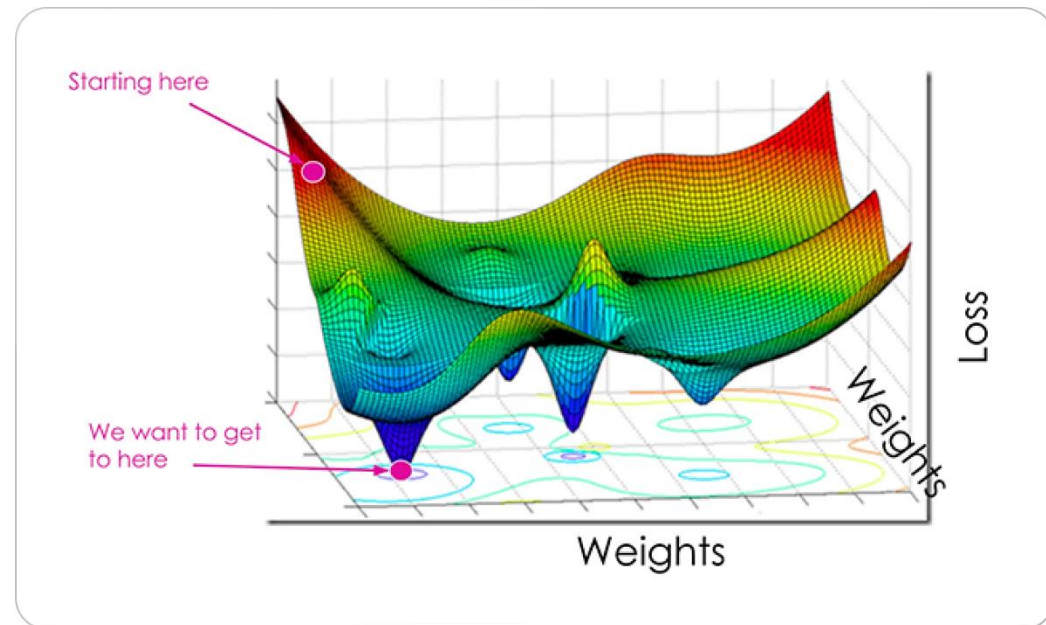
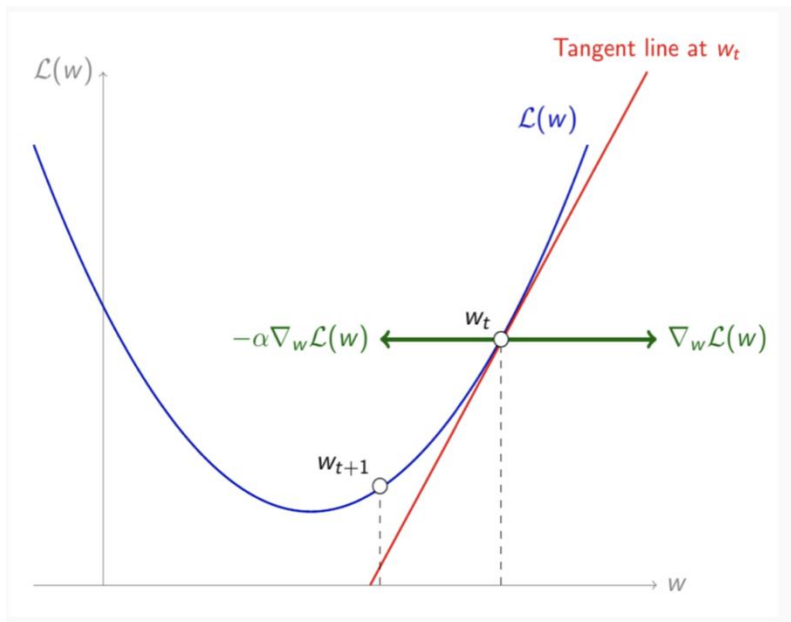
Fig 3: a 1D damped harmonic oscillator



<https://benmoseley.blog/my-research/so-what-is-a-physics-informed-neural-network/>

Learning in neural networks

- The process of finding the combination of weights and biases that minimize the error on the training set – optimization problem
- Generalization: a well-trained model should perform roughly equal on seen vs unseen data



Loss function

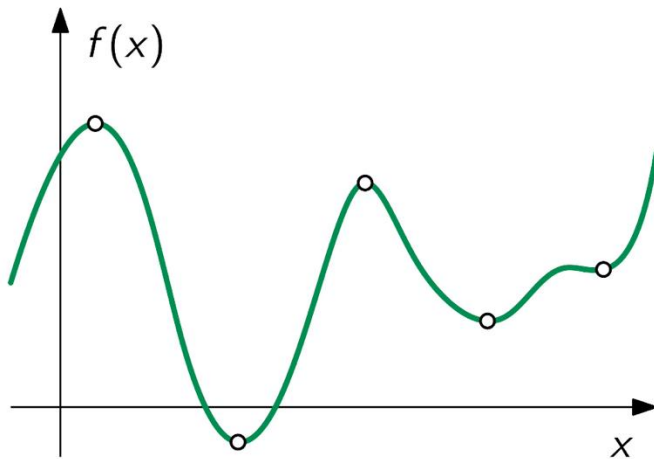
- The loss function gives us a distance (error) between our predictions and the target values

$$\mathcal{L}(w) = \underset{\text{error}}{\text{distance}}(\underset{\substack{\text{input} \\ \text{parameters (weights, biases)}}}{f_{\theta}(x)}, \underset{\text{labels (ground truth)}}{y})$$

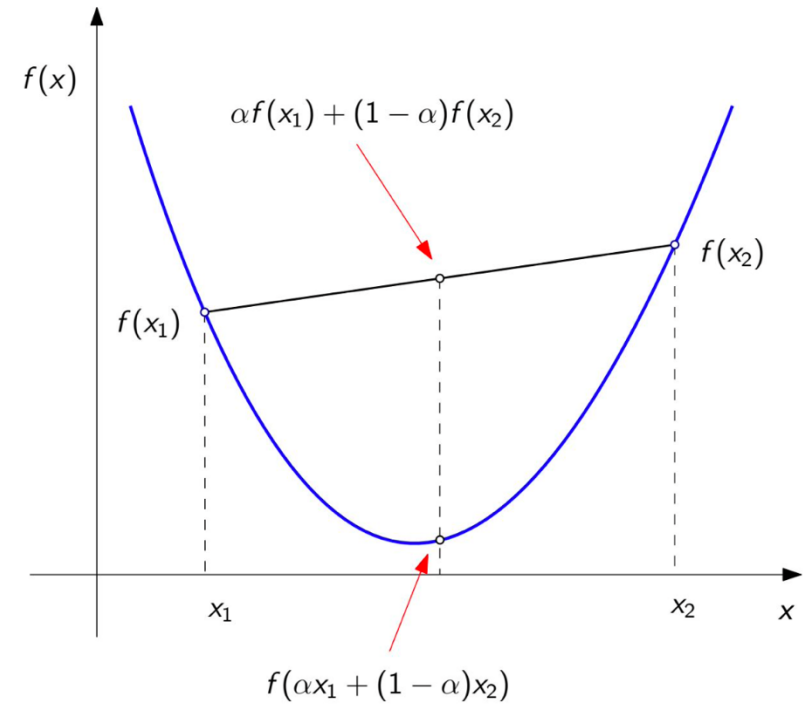
Task	Error type	Loss function	Note
Regression	Mean-squared error	$\frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2$	Easy to learn but sensitive to outliers (MSE, L2 loss)
	Mean absolute error	$\frac{1}{n} \sum_{i=1}^n y_i - \hat{y}_i $	Robust to outliers but not differentiable (MAE, L1 loss)
Classification	Cross entropy = Log loss	$-\frac{1}{n} \sum_{i=1}^n [y_i \log(\hat{y}_i) + (1 - y_i) \log(1 - \hat{y}_i)] =$	Quantify the difference between two probability

Loss function

- Convex functions have a single minimum
- Most loss functions of the problems that we tackle with deep learning are high-dimensional non-convex functions with many local minima



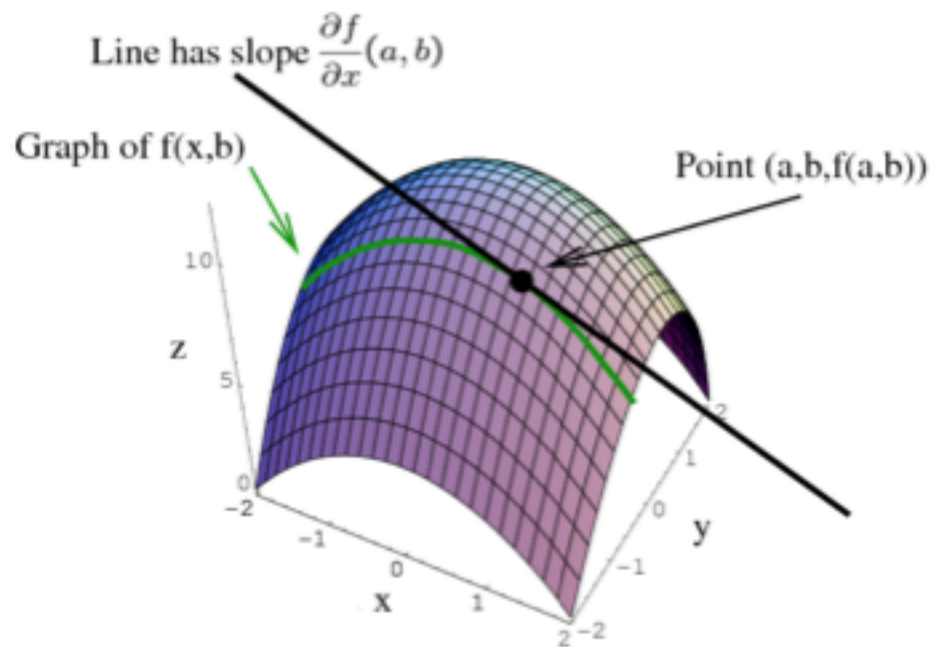
non-convex



convex

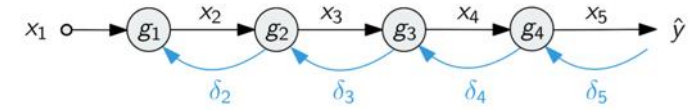
Partial derivatives

- In a multivariate function, the partial derivatives tell us how the output of our function changes with respect to each one of the variables in the input set
- When we derivate with respect to one of the variables, we treat the rest of the variables as constants



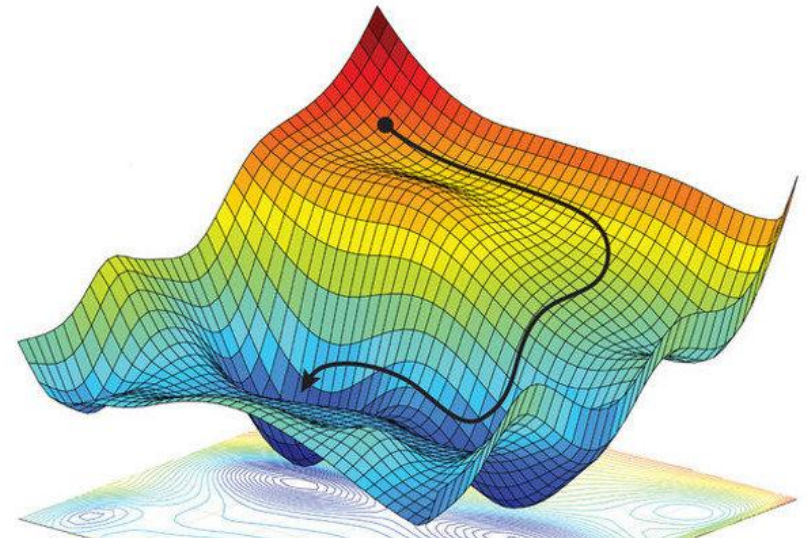
Training a neural network

- **Goal:** Find the set $\{w, b\}$ that minimizes an objective function (loss function) over all the samples in the training set
- Optimization algorithm: **gradient descent** $\rightarrow \theta_i := \theta_i - \alpha \frac{\partial \mathcal{L}}{\partial \theta_i}$
- The gradient of the loss with respect to the parameters of the model is computed with **backpropagation**



Decompose into steps again. Let $\delta_k = \frac{\partial \hat{y}}{\partial x_k}$. **Backpropagation:**

$$\begin{aligned}\delta_5 &= \frac{\partial \hat{y}}{\partial x_5} = 1 \\ \delta_4 &= \frac{\partial \hat{y}}{\partial x_4} = \frac{\partial \hat{y}}{\partial x_5} \frac{\partial x_5}{\partial x_4} = \delta_5 g'_4(x_4) \\ \delta_3 &= \frac{\partial \hat{y}}{\partial x_3} = \frac{\partial \hat{y}}{\partial x_4} \frac{\partial x_4}{\partial x_3} = \delta_4 g'_3(x_3) \\ \delta_2 &= \frac{\partial \hat{y}}{\partial x_2} = \frac{\partial \hat{y}}{\partial x_3} \frac{\partial x_3}{\partial x_2} = \delta_3 g'_2(x_2) \\ \delta_1 &= \frac{\partial \hat{y}}{\partial x_1} = \frac{\partial \hat{y}}{\partial x_2} \frac{\partial x_2}{\partial x_1} = \delta_2 g'_1(x_1)\end{aligned}$$



The chain rule

Easily differentiate compositions of functions.

How does a variation
("difference") on x
affect y ?

$$y = f(z)$$
$$z = g(x)$$

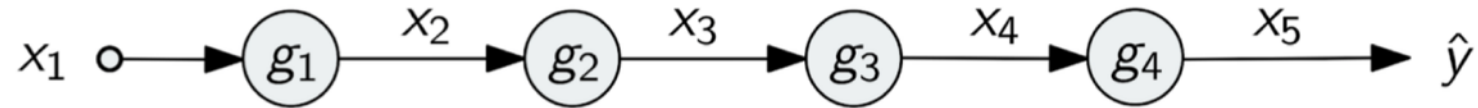
It will depend on the
how y is affected by a
variation in z ...

$$\frac{dy}{dx} = \frac{dy}{dz} \frac{dz}{dx}$$

...scaled (multiplied)
on how z is affected by
a variation in x .

The chain rule

$$\hat{y} = g_4(g_3(g_2(g_1(x_1))))$$



Decompose into steps (**forward propagation**):

$$x_2 = g_1(x_1)$$

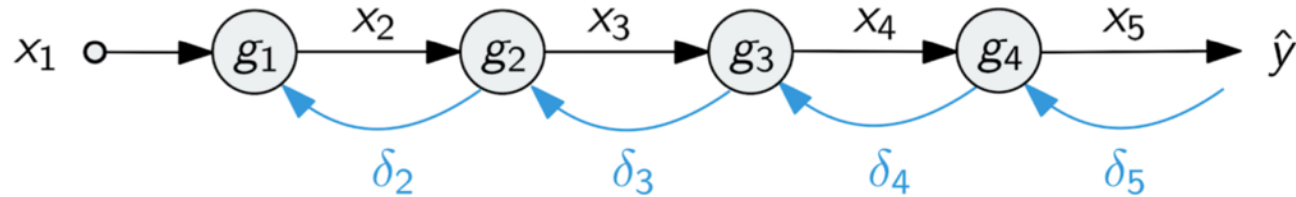
$$x_3 = g_2(x_2)$$

$$x_4 = g_3(x_3)$$

$$\hat{y} = x_5 = g_4(x_4)$$

The chain rule

$$\hat{y} = g_4(g_3(g_2(g_1(x_1))))$$

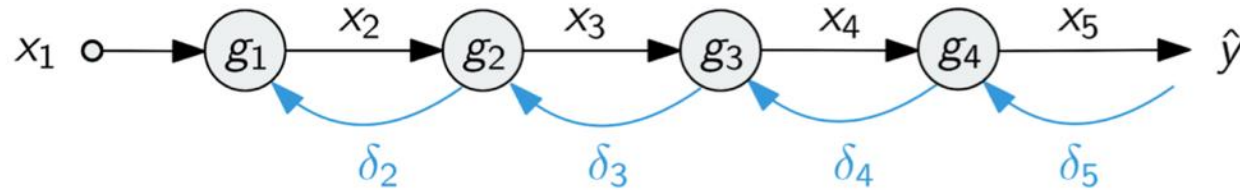


Want to find $\frac{\partial \hat{y}}{\partial x_1}$. Chain rule:

How does a variation
("difference") on the
input affect the
prediction ?

$$\frac{\partial \hat{y}}{\partial x_1} = \frac{\partial \hat{y}}{\partial x_4} \frac{\partial x_4}{\partial x_3} \frac{\partial x_3}{\partial x_2} \frac{\partial x_2}{\partial x_1}$$

The chain rule



Decompose into steps again. Let $\delta_k = \frac{\partial \hat{y}}{\partial x_k}$. **Backpropagation:**

$$\delta_5 = \frac{\partial \hat{y}}{\partial x_5} = 1$$

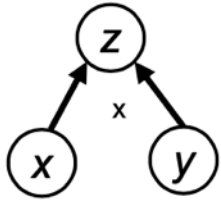
$$\delta_4 = \frac{\partial \hat{y}}{\partial x_4} = \frac{\partial \hat{y}}{\partial x_5} \frac{\partial x_5}{\partial x_4} = \delta_5 g'_4(x_4)$$

$$\delta_3 = \frac{\partial \hat{y}}{\partial x_3} = \frac{\partial \hat{y}}{\partial x_4} \frac{\partial x_4}{\partial x_3} = \delta_4 g'_3(x_3)$$

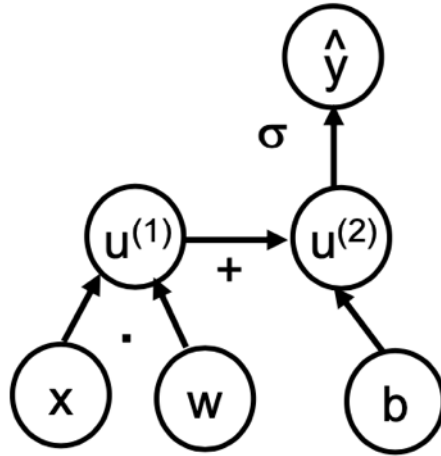
$$\delta_2 = \frac{\partial \hat{y}}{\partial x_2} = \frac{\partial \hat{y}}{\partial x_3} \frac{\partial x_3}{\partial x_2} = \delta_3 g'_2(x_2)$$

$$\delta_1 = \frac{\partial \hat{y}}{\partial x_1} = \frac{\partial \hat{y}}{\partial x_2} \frac{\partial x_2}{\partial x_1} = \delta_2 g'_1(x_1)$$

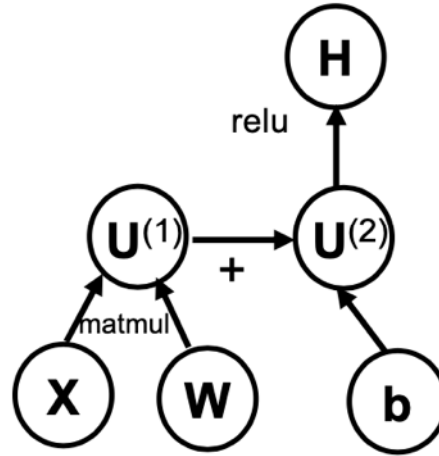
Computational graphs



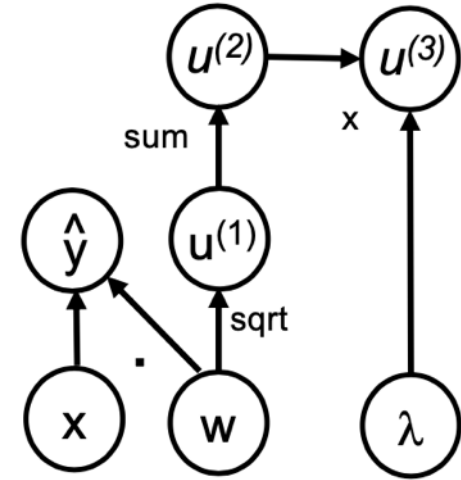
$$z = xy$$



$$\hat{y} = \sigma(x^T w + b)$$



$$H = \max\{0, XW + b\}$$



$$\hat{y} = x^T w$$

$$\lambda \sum_i w_i^2$$

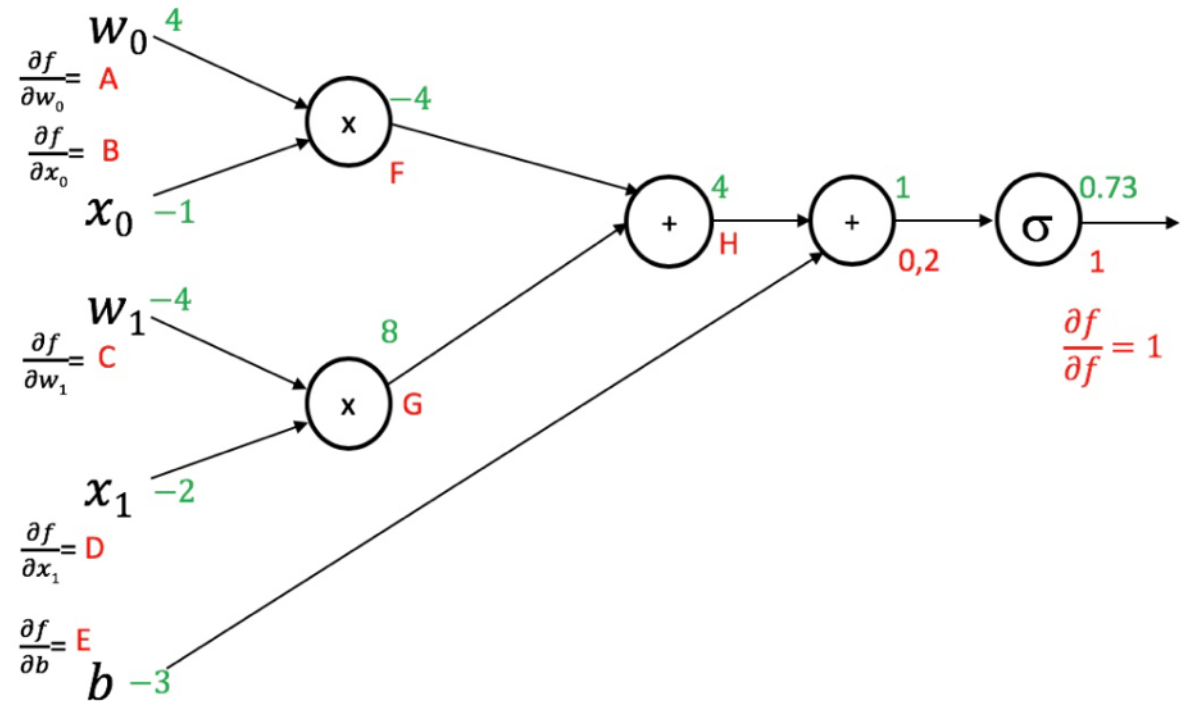
Solving a computational graph

Find the values (from A to H) of the following computational graph:

$$f(x, y, z) = \sigma(w_0x_0 + w_1x_1 + b)$$

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

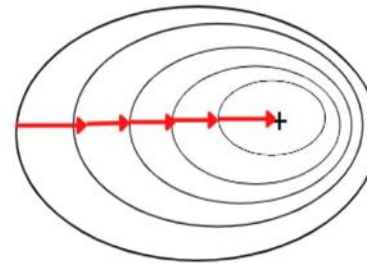
$$\frac{d\sigma(x)}{dx} = (1 - \sigma(x))(\sigma(x))$$



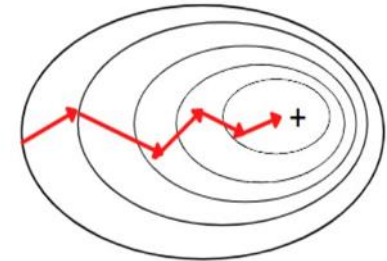
Stochastic Gradient Descent

- Stochastic gradient descent: We compute the gradient with one sample at a time
- Batch gradient descent: we compute the gradient for the entire training set and then take a step towards the minimum
- **Mini-batch gradient descent:** we select a subset of samples (batch) and compute the gradient for that given subset

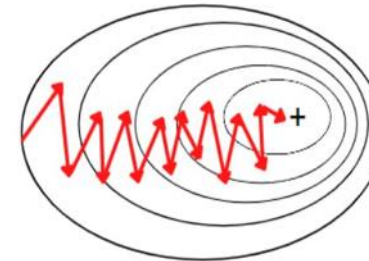
Batch Gradient Descent



Mini-Batch Gradient Descent

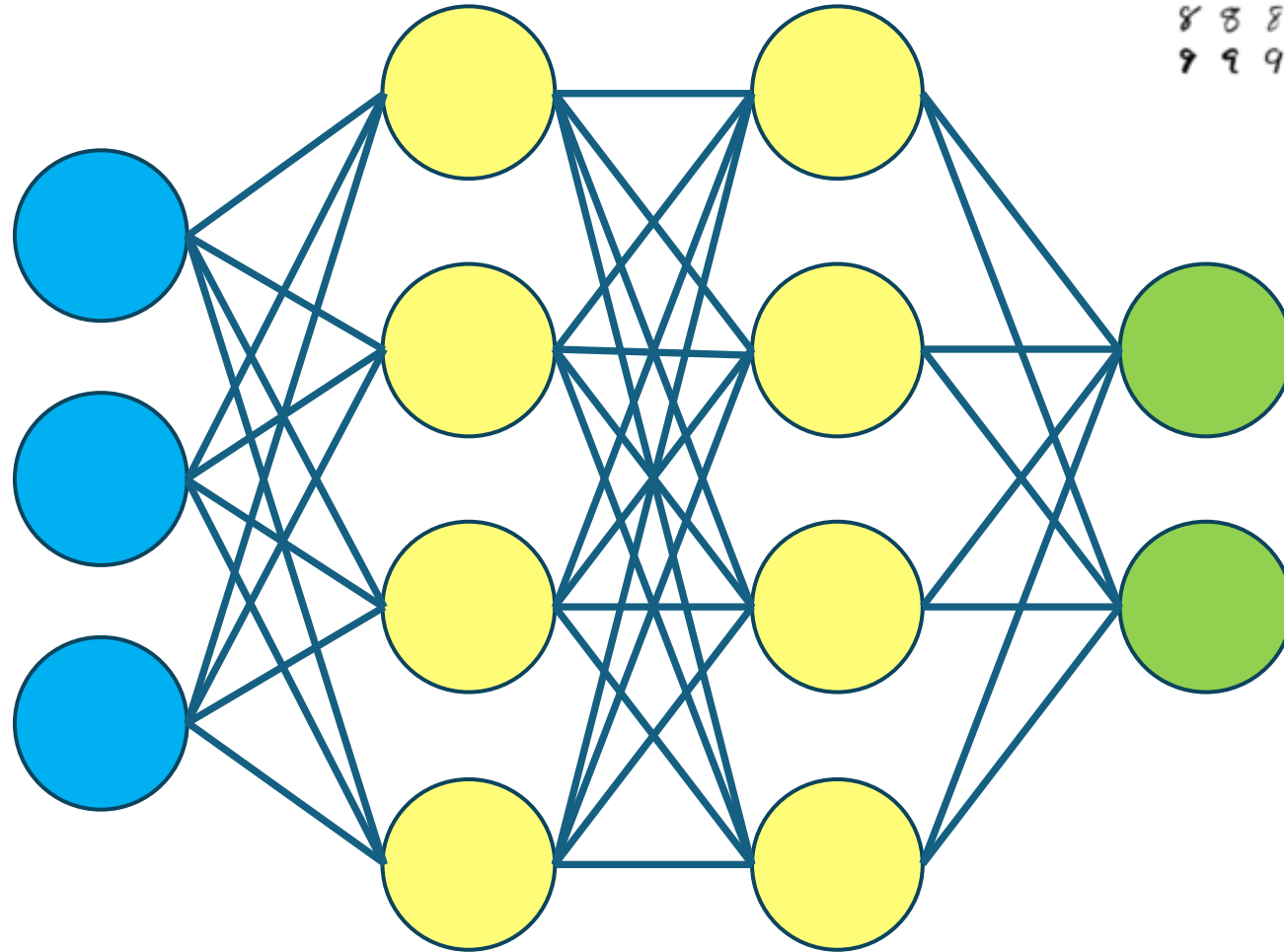
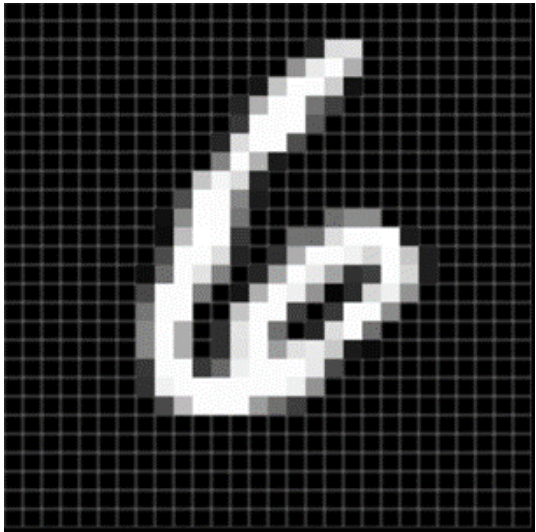


Stochastic Gradient Descent



Training a Neural Network

Input image and perform feed-forward operation



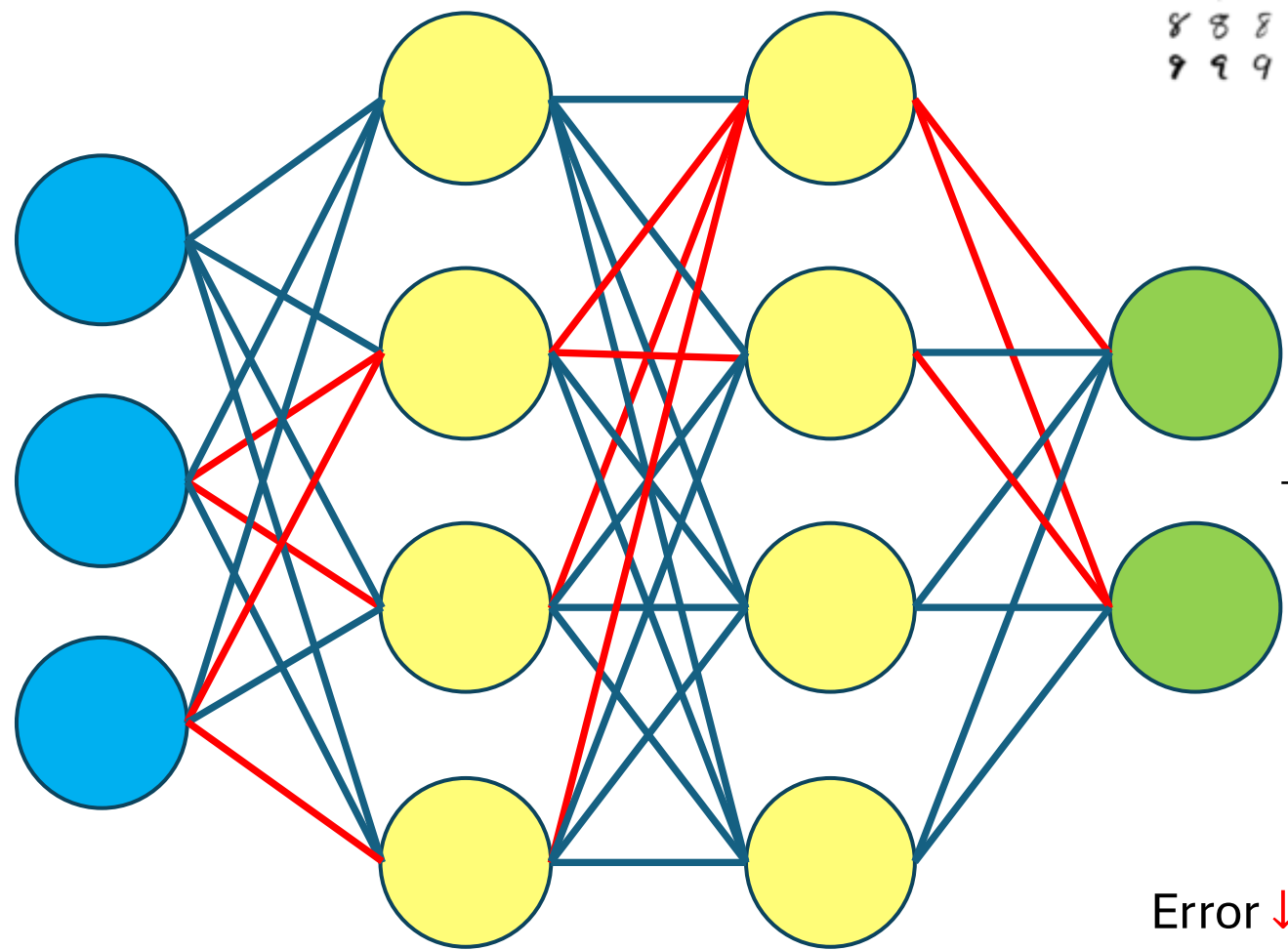
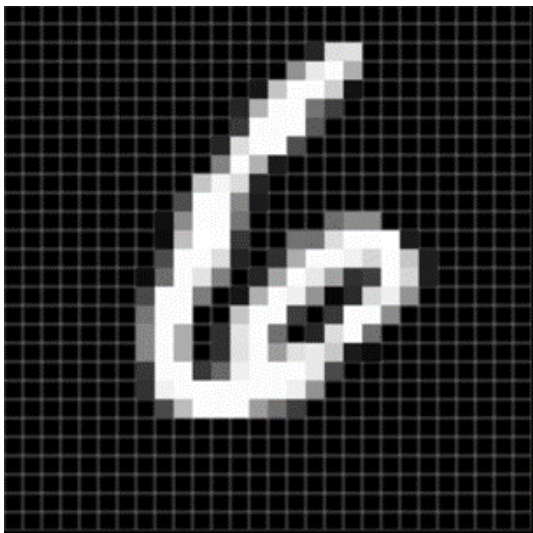
Output

Label

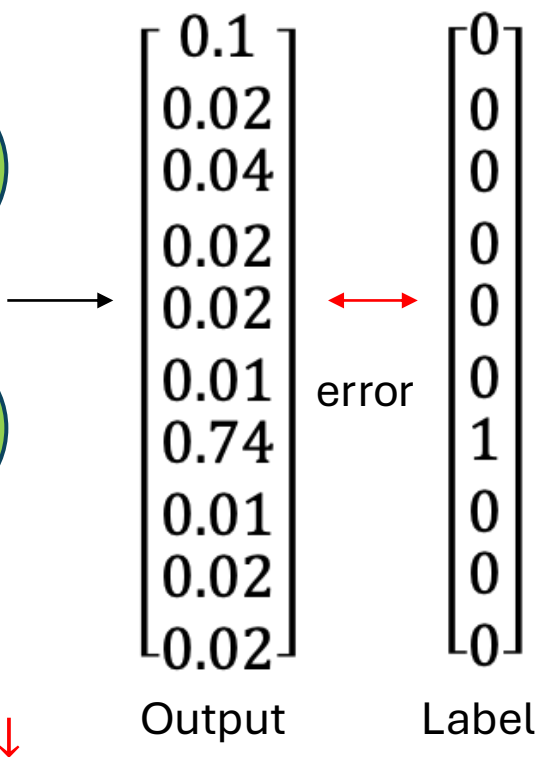
A collection of handwritten digits from 0 to 9, arranged in a 10x10 grid. Each row contains ten variations of a single digit, showing different handwriting styles, slants, and sizes. The digits are written in black ink on a white background.

Training a Neural Network

Update weights



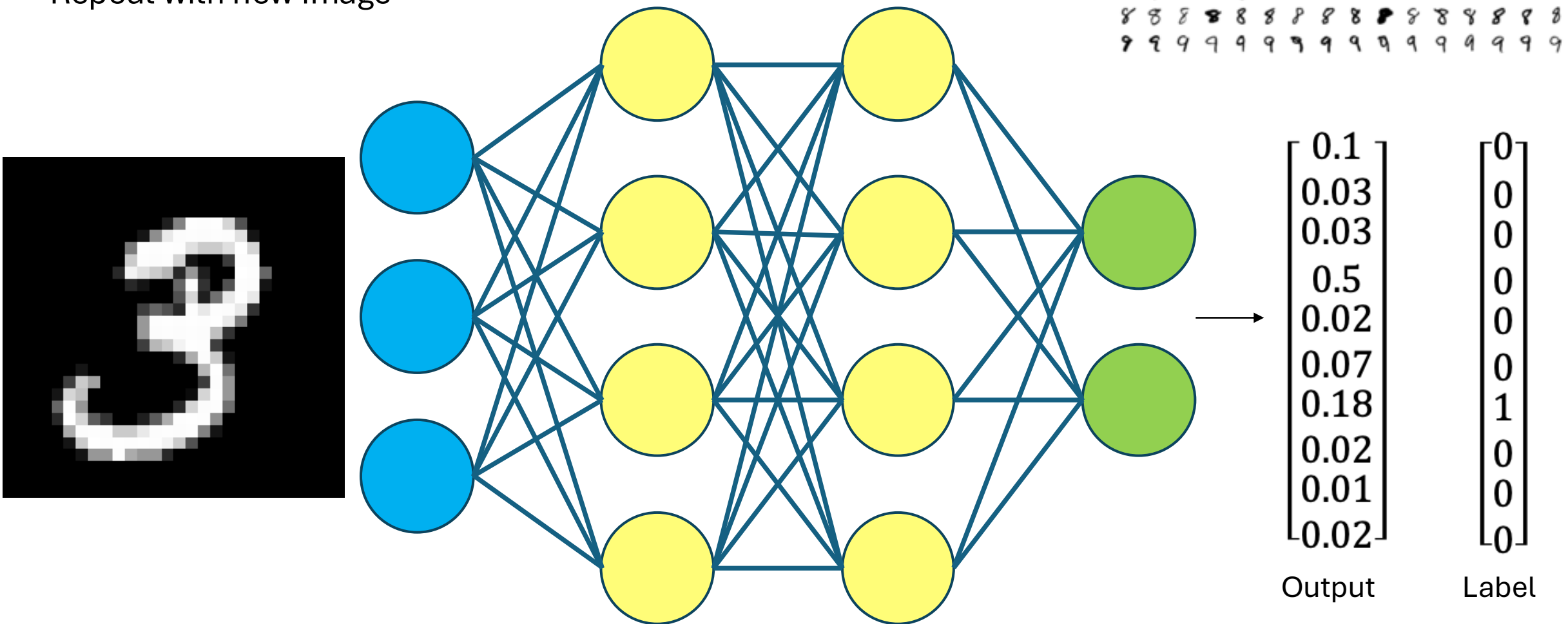
0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0
1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
2	2	2	2	2	2	2	2	2	2	2	2	2	2	2	2
3	3	3	3	3	3	3	3	3	3	3	3	3	3	3	3
4	4	4	4	4	4	4	4	4	4	4	4	4	4	4	4
5	5	5	5	5	5	5	5	5	5	5	5	5	5	5	5
6	6	6	6	6	6	6	6	6	6	6	6	6	6	6	6
7	7	7	7	7	7	7	7	7	7	7	7	7	7	7	7
8	8	8	8	8	8	8	8	8	8	8	8	8	8	8	8
9	9	9	9	9	9	9	9	9	9	9	9	9	9	9	9



Error ↓

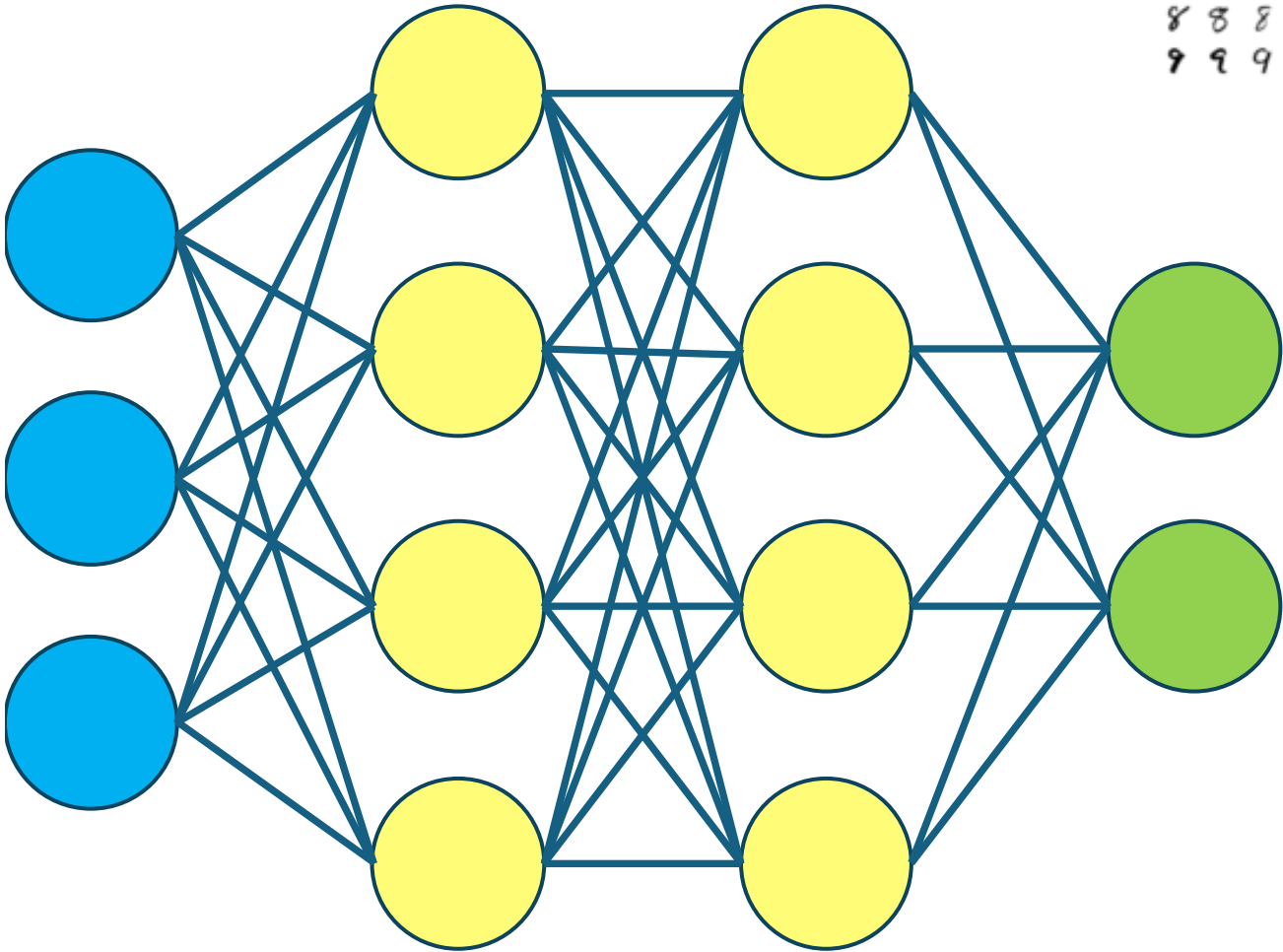
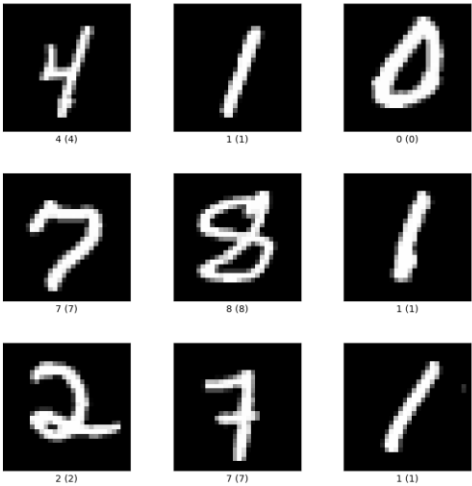
Training a Neural Network

Repeat with new image



Training a Neural Network

Batch



0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0
1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
2	2	2	2	2	2	2	2	2	2	2	2	2	2	2	2
3	3	3	3	3	3	3	3	3	3	3	3	3	3	3	3
4	4	4	4	4	4	4	4	4	4	4	4	4	4	4	4
5	5	5	5	5	5	5	5	5	5	5	5	5	5	5	5
6	6	6	6	6	6	6	6	6	6	6	6	6	6	6	6
7	7	7	7	7	7	7	7	7	7	7	7	7	7	7	7
8	8	8	8	8	8	8	8	8	8	8	8	8	8	8	8
9	9	9	9	9	9	9	9	9	9	9	9	9	9	9	9

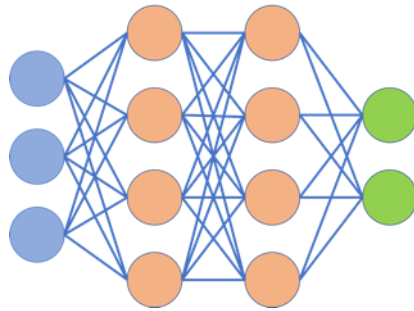
$$\begin{pmatrix} 0.1 & 0.2 & \dots \\ 0.3 & 0.05 & \dots \\ \dots & \dots & \dots \end{pmatrix}$$

Outputs

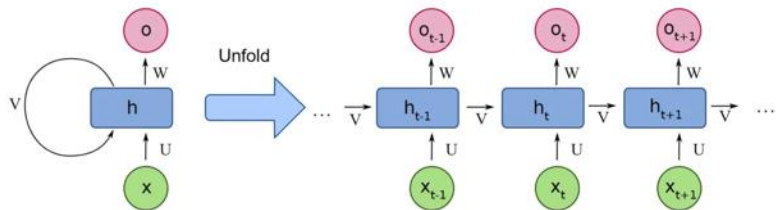
$$\begin{pmatrix} 0 & 0 & \dots \\ 0 & 1 & \dots \\ \dots & \dots & \dots \end{pmatrix}$$

Labels

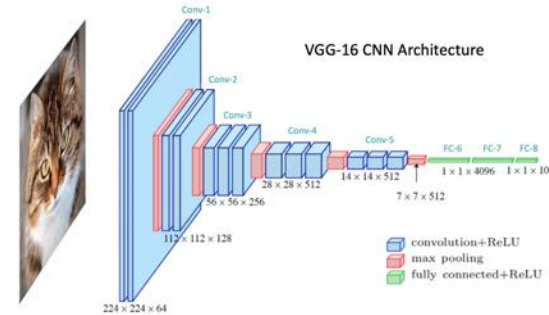
Neural Network architectures



fully connected/linear/dense layers or MLP
→
linear map + non-linear activation functions

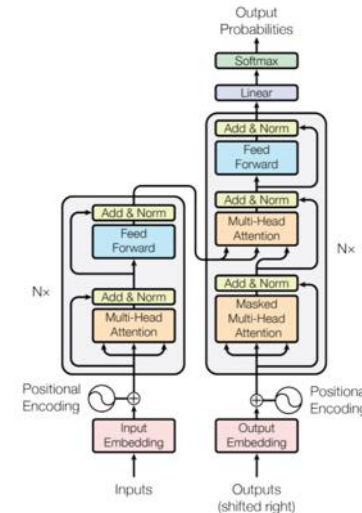


recurrent layers → RNNs, LSTMs, GRUs
time series, natural language processing



convolutional layers → CNNs
computer vision, signal processing,

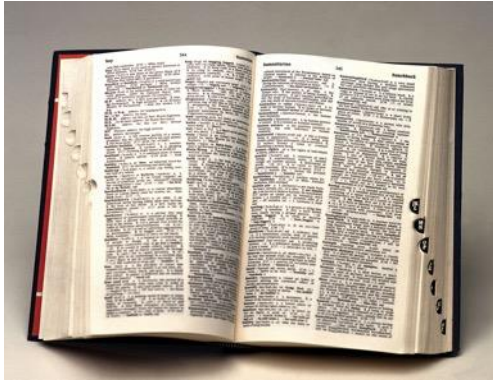
etc.



Transformers
Seq to seq models
Multimodal models
Large language models
Generative models
Etc.

Task: next word prediction

Vocab size



A
Aardvark
Abeam
Abacus
Abandon
...

Embedding dimension

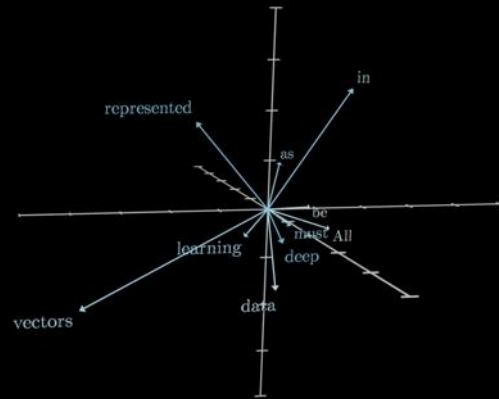
All words, ~ 50k

aah	aardvark	aardwolf	aargh	ab	aback	abacterial	abacus	abalone	abandon	...	zygoid	zygomatic	zygomorphic	zygosis	zygote	zygotic	zyne	zymogen	zymosis	zzz
+1.0	+4.3	+2.0	+0.9	-1.5	+2.9	-1.2	+7.8	+9.2	-2.3	...	+0.6	+1.3	+8.4	-8.5	-8.2	-9.5	+6.6	+5.5	+7.3	+9.5
+5.9	-0.8	+5.6	-7.6	+2.8	-7.1	+8.8	+0.4	-1.7	-4.7	...	-0.9	+1.4	-9.5	+2.3	+2.2	+2.3	+8.8	+3.6	-2.8	-1.2
+3.9	-8.7	+3.3	+3.4	-5.7	-7.3	-3.7	-2.7	+1.4	-1.2	...	-7.9	-5.8	-6.7	+3.0	-4.9	-0.7	-5.1	-6.8	-7.7	+3.1
-7.2	-6.0	-2.6	+6.4	-8.0	+6.7	-8.0	+9.4	-0.6	+9.4	...	+4.7	-9.1	-4.3	-7.5	-4.0	-7.5	-3.6	-1.7	-8.6	+3.8
+1.3	-4.6	+0.5	-8.0	+1.5	+8.5	-3.6	+3.3	-7.3	+4.3	...	-6.3	+1.7	-9.5	+6.5	-9.8	+3.5	-4.6	+4.7	+9.2	-5.0
+1.5	+1.8	+1.4	-5.5	+9.0	-1.0	+6.9	+3.9	-4.0	+6.2	...	+7.5	+1.6	+7.6	+3.8	+4.5	+0.0	+9.0	+2.9	-1.5	+2.1
-9.5	-3.9	+3.2	-4.2	+2.3	-1.4	-7.2	-4.0	+1.4	+1.8	...	+3.0	+3.0	-1.4	+7.9	-2.6	-1.3	+7.8	+6.1	+4.0	-7.9
+8.3	+4.2	+9.9	-6.9	+7.3	-6.7	+2.3	-7.4	+6.9	+6.1	...	-1.8	-8.5	+3.9	-0.9	+4.4	+7.3	+9.4	+7.0	-9.7	-2.8
⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮
-3.7	-2.0	-5.7	-6.2	+8.8	+4.7	-0.2	-5.4	-4.9	-8.8	...	-3.7	+3.9	-2.4	-6.3	-9.4	-8.6	+3.6	-0.9	+0.7	+7.9

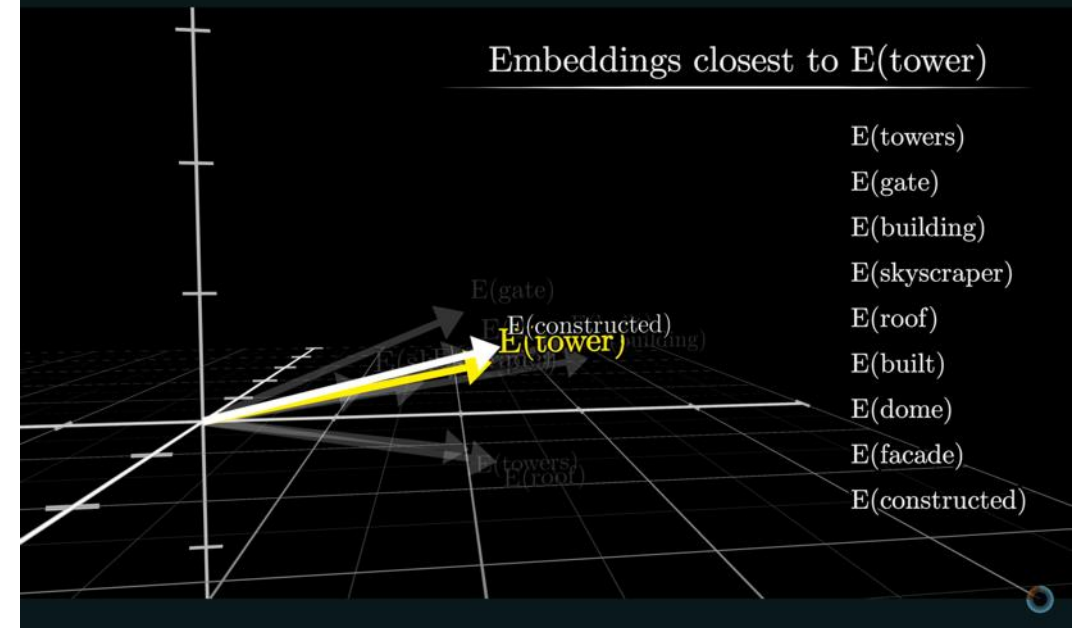
Embedding matrix

Words $\xrightarrow{\text{"Embedding"}}$ Vectors

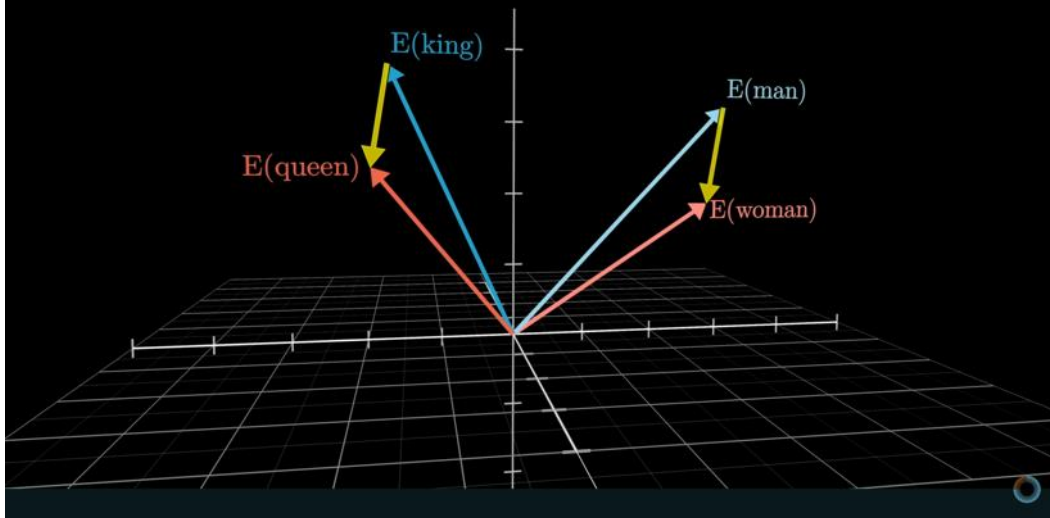
All data in deep learning must be represented as vectors



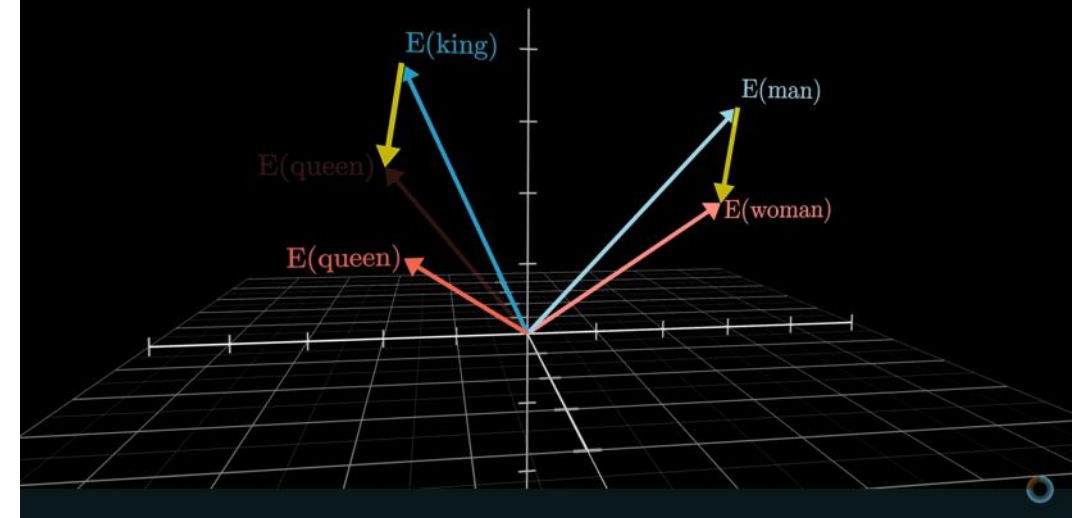
Embeddings closest to E(tower)



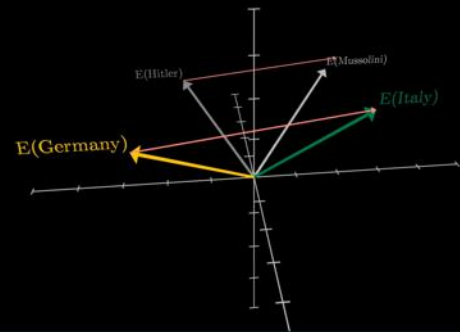
$$E(\text{queen}) \approx E(\text{king}) + E(\text{woman}) - E(\text{man})$$



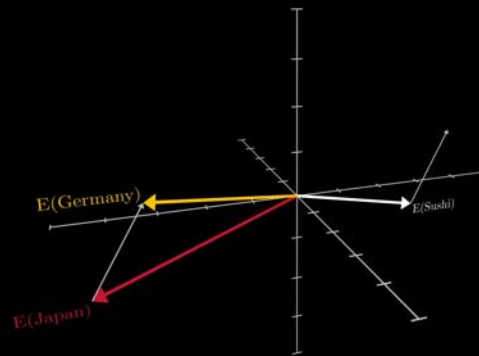
$$E(\text{queen}) \approx E(\text{king}) + E(\text{woman}) - E(\text{man})$$



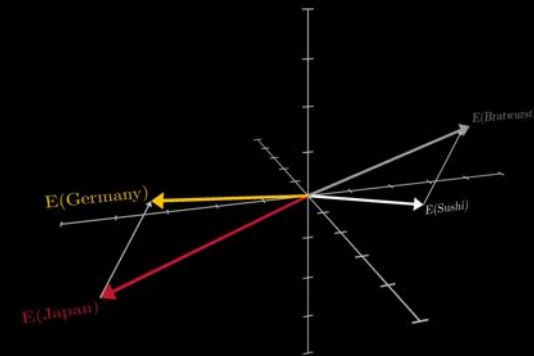
$$E(\text{Hitler}) + E(\text{Italy}) - E(\text{Germany}) \approx E(\text{Mussolini})$$



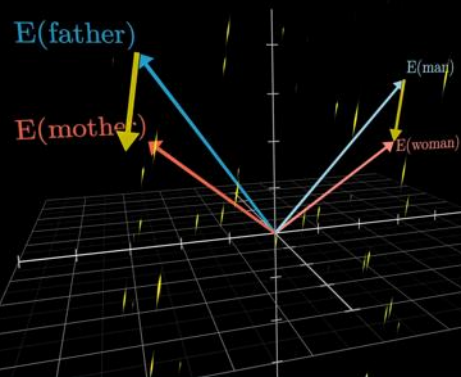
$$E(\text{Sushi}) + E(\text{Germany}) - E(\text{Japan}) \approx E(\text{Bratwurst})$$



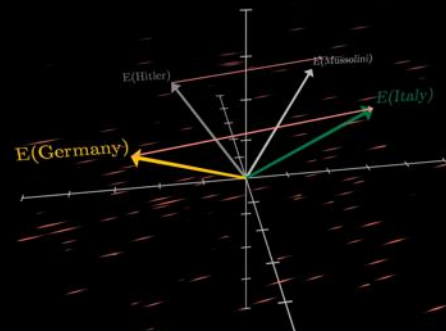
$$E(\text{Sushi}) + E(\text{Germany}) - E(\text{Japan}) \approx E(\text{Bratwurst})$$



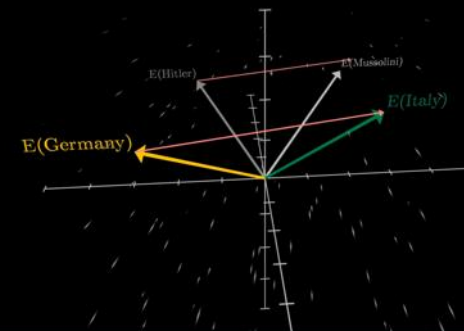
$$E(\text{mother}) - E(\text{father}) \approx E(\text{woman}) - E(\text{man})$$



$$E(\text{Hitler}) + E(\text{Italy}) - E(\text{Germany}) \approx E(\text{Mussolini})$$



$$E(\text{Hitler}) + E(\text{Italy}) - E(\text{Germany}) \approx E(\text{Mussolini})$$



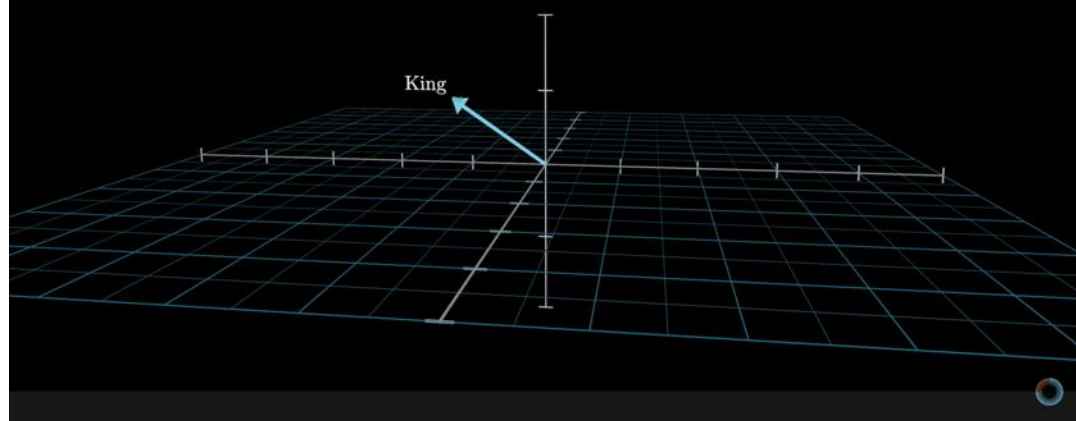
Task: next word prediction

- We can use each word in a sentence to predict the following word; however, we would be missing the context that gives meaning to it
- Context: how the meaning of a word is affected by the meaning of the other words in a given text
 1. Machine learning **model**
 2. Fashion **model**
- Language model: from an initial embedding or snippet we want to obtain a single final embedding, which we can then run through an MLP and predict the probability of each word in the vocabulary appearing next

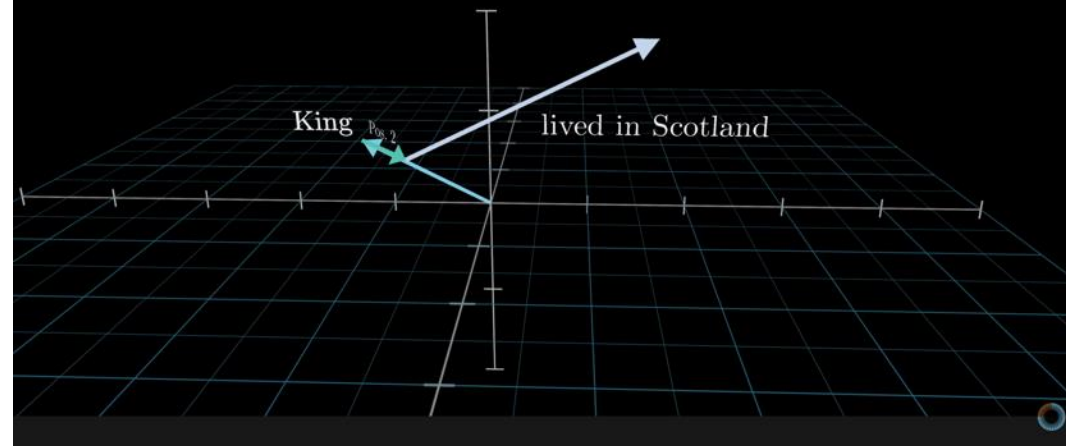
The **King** doth wake tonight and takes his rouse ...

1 2 3 4 5 6 7 8 9 10

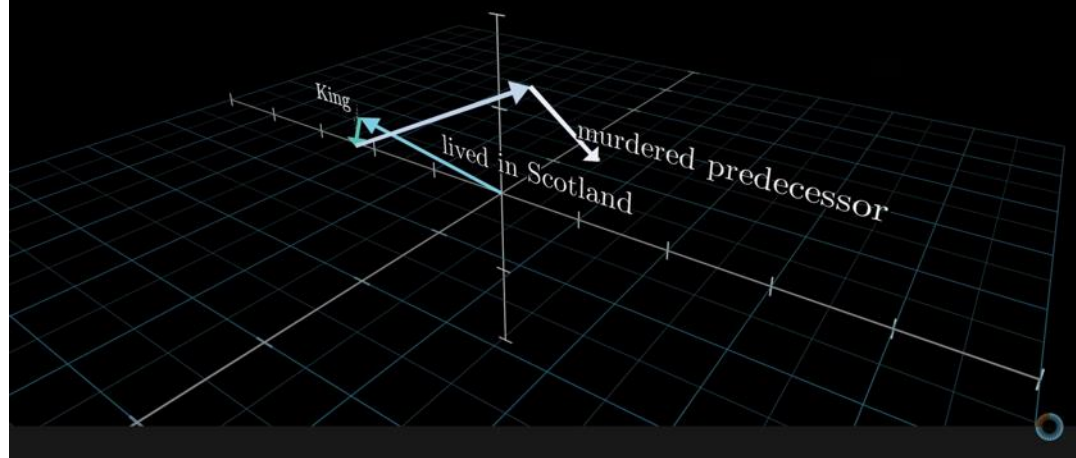
2



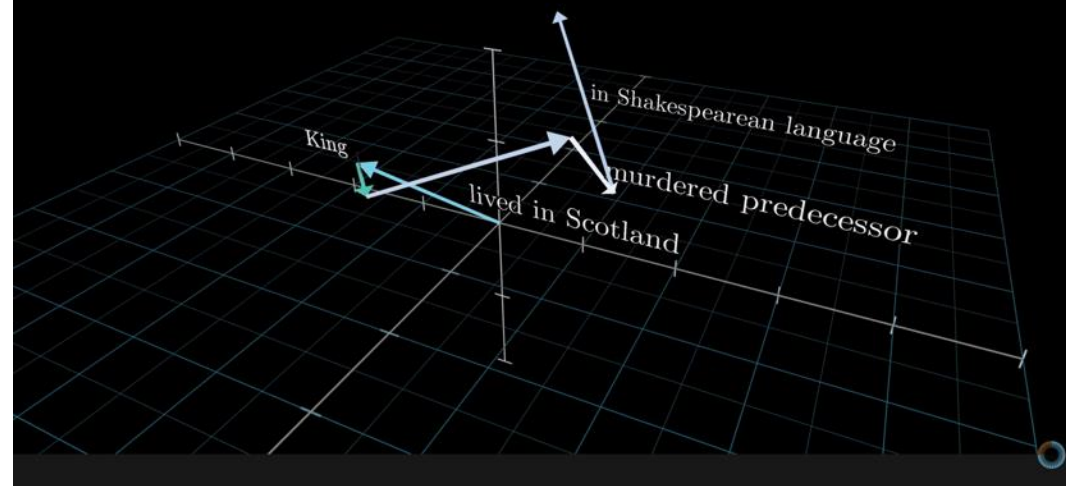
The **King** doth wake tonight and takes his rouse ...



The **King** doth wake tonight and takes his rouse ...



The **King** doth wake tonight and takes his rouse ...



Embedding dimension

Vocab size

Unembedding
matrix

W_U

+3.8	+6.3	-0.2	-7.2	+6.9	+1.5	+4.7	+4.0	...	-4.1
+4.1	-2.7	-2.1	-5.3	-3.1	+8.8	-4.1	-5.0	...	-4.8
-0.5	+6.6	-5.3	-1.4	+2.2	+0.9	+9.4	+3.6	...	+9.2
-1.7	-2.9	-9.0	-6.2	-5.2	-6.3	+5.0	+0.7	...	+6.3
-5.3	-3.4	+4.1	-2.1	-9.3	-1.3	+8.1	-1.8	...	+9.7
+2.8	-2.7	-7.9	+5.7	+4.1	+8.4	-5.6	-7.6	...	-5.9
-6.4	-3.6	+6.3	+0.8	-9.0	-0.7	+3.6	+0.8	...	-5.4
+6.9	+1.2	+4.2	+9.5	-1.4	+7.5	-9.7	-9.2	...	-3.7
⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋱	⋮
+8.4	+3.2	-8.3	+0.8	-2.9	+9.6	-9.6	+2.2	...	-4.2
+9.4	+7.1	+8.2	-9.5	+1.4	-4.0	+6.9	+2.6	...	-7.6
+0.8	+2.6	+9.0	+1.7	+9.3	+9.1	+3.0	+0.1	...	+7.7
-9.3	-7.6	-7.9	+5.1	-3.2	+2.7	+2.0	-2.2	...	+2.9
+8.7	+1.5	+2.3	-8.5	+8.9	+0.5	+6.0	-8.9	...	-4.8
-4.6	+5.8	+2.4	-1.2	-9.7	+9.2	+9.1	-5.6	...	+0.6

softmax

$\begin{bmatrix} -0.8 \\ -5.0 \\ +0.5 \\ +1.5 \\ +3.4 \\ -2.3 \\ +2.5 \end{bmatrix}$



$\begin{bmatrix} e^{x_1/\sum_{n=1}^N e^{x_n}} \\ e^{x_2/\sum_{n=1}^N e^{x_n}} \\ e^{x_3/\sum_{n=1}^N e^{x_n}} \\ e^{x_4/\sum_{n=1}^N e^{x_n}} \\ e^{x_5/\sum_{n=1}^N e^{x_n}} \\ e^{x_6/\sum_{n=1}^N e^{x_n}} \\ e^{x_7/\sum_{n=1}^N e^{x_n}} \end{bmatrix}$

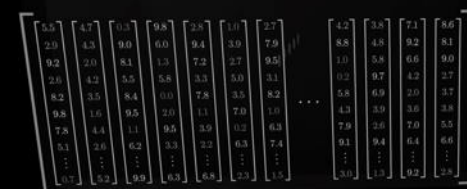
=

$\begin{bmatrix} 0.01 \\ 0.00 \\ 0.03 \\ 0.09 \\ 0.61 \\ 0.00 \\ 0.25 \end{bmatrix}$



Harry Potter was a highly unusual boy ... least favourite teacher, Professor

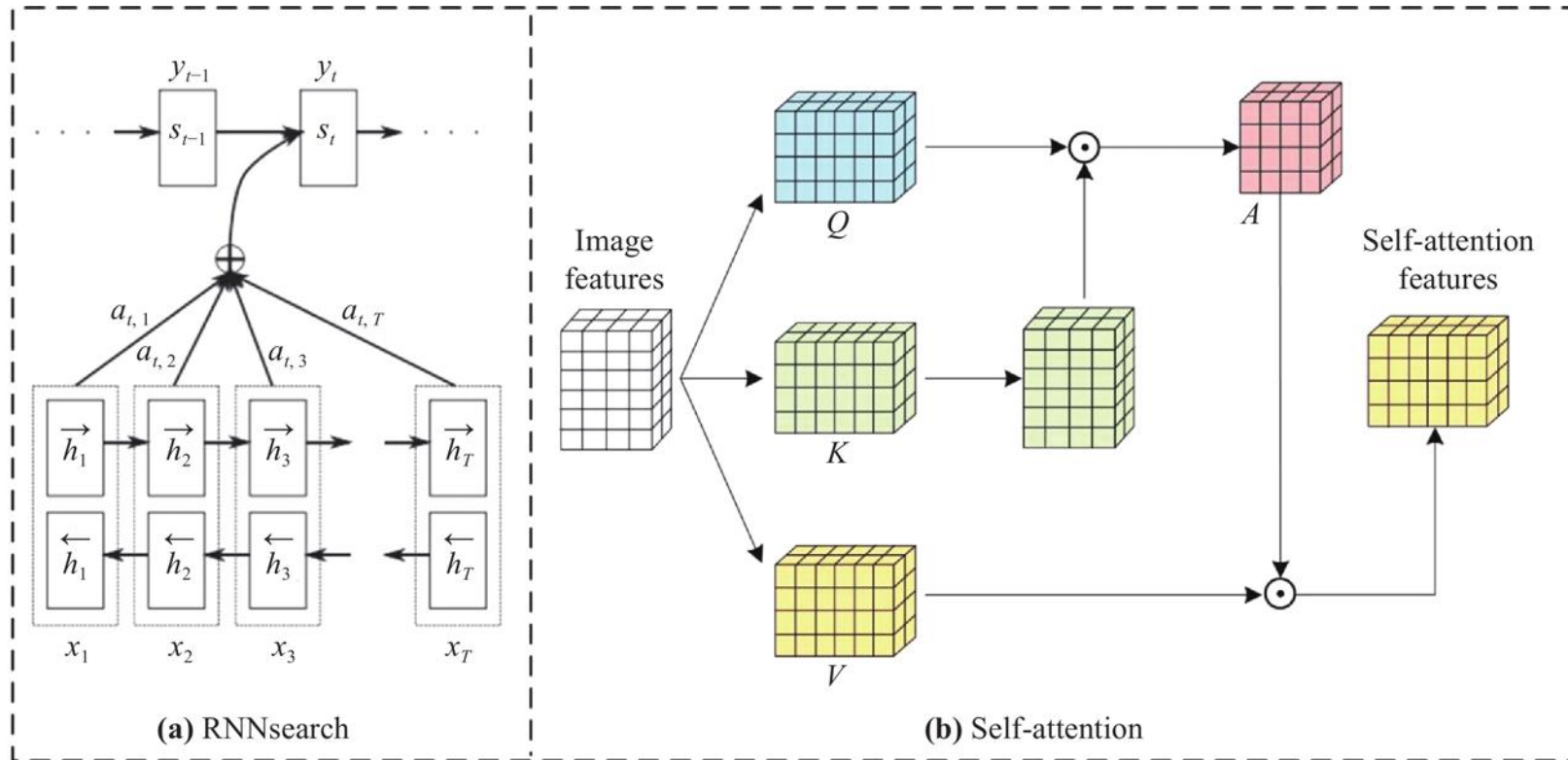
???



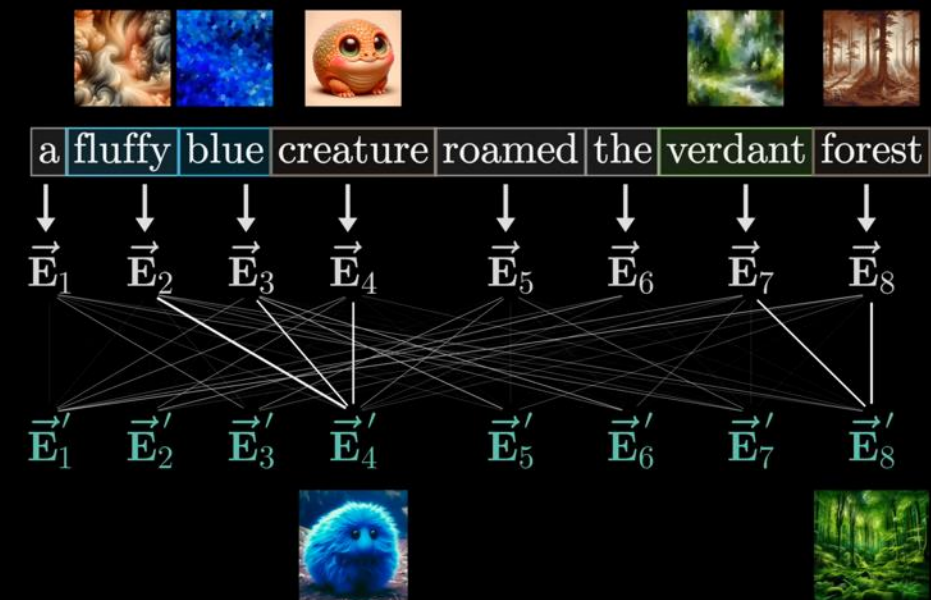
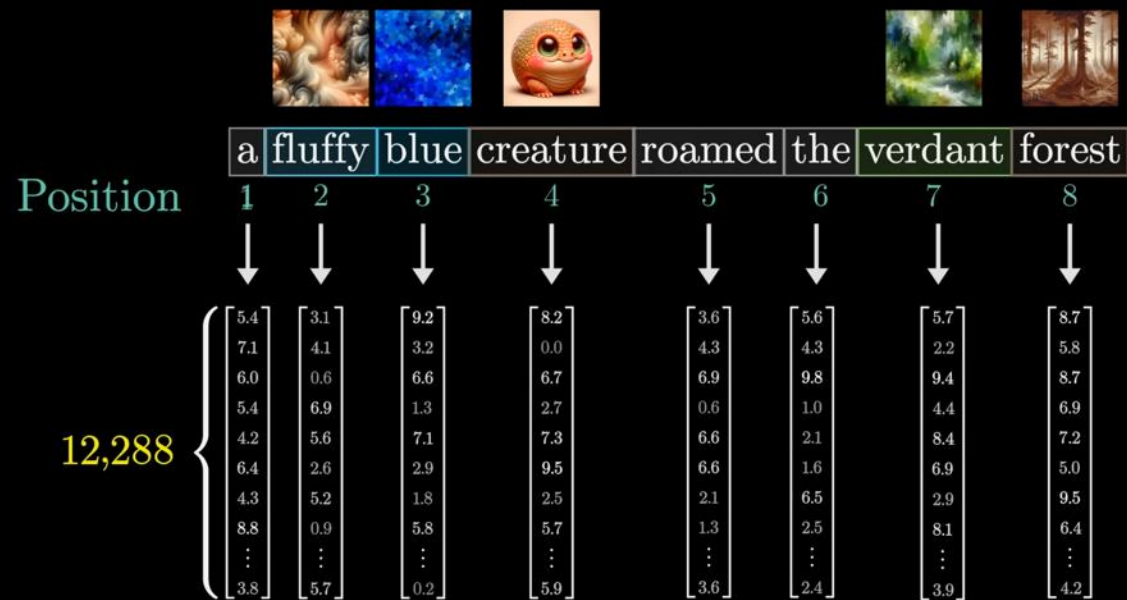
- 0.00 | Snake
- 0.78 | Snake
- 0.00 | Snare
- 0.00 | Treks
- 0.16 | Trelawney
- 0.00 | Trellis
- 0.00 | Quirky
- 0.06 | Quirrell
- 0.00 | Quirt

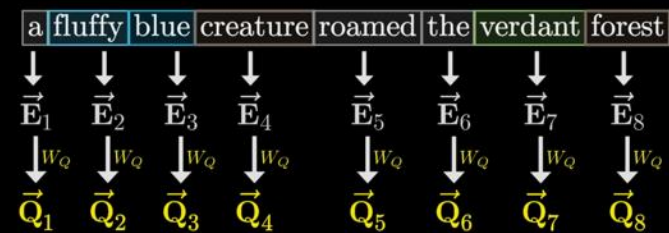
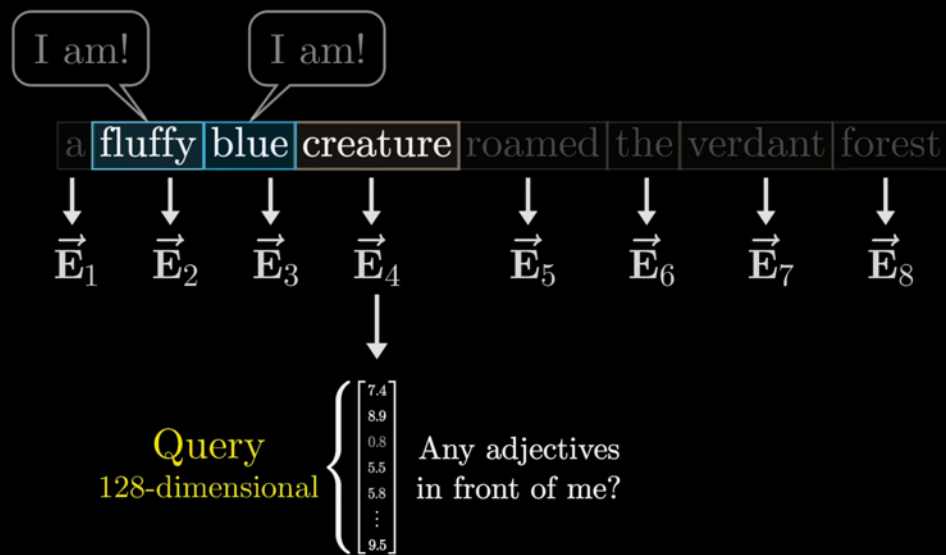
Attention mechanism

The **teacher** asked the **class** a question, and everyone turned to look at **her**



Li, X., Li, M., Yan, P., Li, G., Jiang, Y., Luo, H., & Yin, S. Deep Learning Attention Mechanism in Medical Image Analysis: Basics and Beyonds. *International Journal of Network Dynamics and Intelligence*. 2023, 2(1), 93–116.





Any adjectives
in front of me?

$$\underbrace{\begin{bmatrix} +0.9 & -0.9 & +7.6 & -0.8 & +4.4 & -2.0 & +8.0 & +3.8 & +4.0 & -3.4 & \cdots & +5.1 \\ +2.7 & -5.1 & -6.7 & +5.9 & +9.1 & -0.8 & +1.8 & +7.1 & -0.8 & +8.9 & \cdots & +1.5 \\ +6.4 & +8.1 & +6.2 & -6.7 & +2.6 & -2.0 & -8.7 & -1.5 & -4.8 & +6.9 & \cdots & -9.2 \\ +9.1 & -2.9 & -2.8 & -9.6 & -6.2 & -2.0 & +8.5 & -7.9 & +8.8 & +7.3 & \cdots & -0.9 \\ -3.4 & -5.3 & +2.3 & -9.2 & -9.6 & -1.4 & -8.6 & -4.9 & -5.5 & -4.9 & \cdots & -7.3 \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \ddots & \vdots \\ -9.7 & -7.6 & +2.3 & +9.4 & +9.7 & -1.8 & -6.7 & +2.7 & -0.2 & +9.7 & \cdots & -8.6 \end{bmatrix}}_{W_Q} \begin{bmatrix} 2.9 \\ 2.4 \\ 1.0 \\ 0.2 \\ 9.2 \\ 6.6 \\ 7.8 \\ 2.8 \\ 5.8 \\ 0.6 \\ \vdots \\ 9.7 \end{bmatrix} = \begin{bmatrix} +319.6 \\ -95.2 \\ -21 \\ -152.0 \\ -123.2 \\ \vdots \\ -12.7 \end{bmatrix}$$

$\vec{E}_i$ $\vec{Q}_i$

	a	fluffy	blue	creature	roamed	the	verdant	forest	
	$\downarrow$ $\vec{E}_1$ $\downarrow_{W_Q}$ $\vec{Q}_1$	$\downarrow$ $\vec{E}_2$ $\downarrow_{W_Q}$ $\vec{Q}_2$	$\downarrow$ $\vec{E}_3$ $\downarrow_{W_Q}$ $\vec{Q}_3$	$\downarrow$ $\vec{E}_4$ $\downarrow_{W_Q}$ $\vec{Q}_4$	$\downarrow$ $\vec{E}_5$ $\downarrow_{W_Q}$ $\vec{Q}_5$	$\downarrow$ $\vec{E}_6$ $\downarrow_{W_Q}$ $\vec{Q}_6$	$\downarrow$ $\vec{E}_7$ $\downarrow_{W_Q}$ $\vec{Q}_7$	$\downarrow$ $\vec{E}_8$ $\downarrow_{W_Q}$ $\vec{Q}_8$	
$\boxed{\text{a}} \rightarrow \vec{E}_1 \xrightarrow{W_k} \vec{K}_1$	$\vec{K}_1 \cdot \vec{Q}_1$	$\vec{K}_1 \cdot \vec{Q}_2$	$\vec{K}_1 \cdot \vec{Q}_3$	$\vec{K}_1 \cdot \vec{Q}_4$	$\vec{K}_1 \cdot \vec{Q}_5$	$\vec{K}_1 \cdot \vec{Q}_6$	$\vec{K}_1 \cdot \vec{Q}_7$	$\vec{K}_1 \cdot \vec{Q}_8$	
$\boxed{\text{fluffy}} \rightarrow \vec{E}_2 \xrightarrow{W_k} \vec{K}_2$	$\vec{K}_2 \cdot \vec{Q}_1$	$\vec{K}_2 \cdot \vec{Q}_2$	$\vec{K}_2 \cdot \vec{Q}_3$	$\vec{K}_2 \cdot \vec{Q}_4$	$\vec{K}_2 \cdot \vec{Q}_5$	$\vec{K}_2 \cdot \vec{Q}_6$	$\vec{K}_2 \cdot \vec{Q}_7$	$\vec{K}_2 \cdot \vec{Q}_8$	
$\boxed{\text{blue}} \rightarrow \vec{E}_3 \xrightarrow{W_k} \vec{K}_3$	$\vec{K}_3 \cdot \vec{Q}_1$	$\vec{K}_3 \cdot \vec{Q}_2$	$\vec{K}_3 \cdot \vec{Q}_3$	$\vec{K}_3 \cdot \vec{Q}_4$	$\vec{K}_3 \cdot \vec{Q}_5$	$\vec{K}_3 \cdot \vec{Q}_6$	$\vec{K}_3 \cdot \vec{Q}_7$	$\vec{K}_3 \cdot \vec{Q}_8$	
$\boxed{\text{creature}} \rightarrow \vec{E}_4 \xrightarrow{W_k} \vec{K}_4$	$\vec{K}_4 \cdot \vec{Q}_1$	$\vec{K}_4 \cdot \vec{Q}_2$	$\vec{K}_4 \cdot \vec{Q}_3$	$\vec{K}_4 \cdot \vec{Q}_4$	$\vec{K}_4 \cdot \vec{Q}_5$	$\vec{K}_4 \cdot \vec{Q}_6$	$\vec{K}_4 \cdot \vec{Q}_7$	$\vec{K}_4 \cdot \vec{Q}_8$	
$\boxed{\text{roamed}} \rightarrow \vec{E}_5 \xrightarrow{W_k} \vec{K}_5$	$\vec{K}_5 \cdot \vec{Q}_1$	$\vec{K}_5 \cdot \vec{Q}_2$	$\vec{K}_5 \cdot \vec{Q}_3$	$\vec{K}_5 \cdot \vec{Q}_4$	$\vec{K}_5 \cdot \vec{Q}_5$	$\vec{K}_5 \cdot \vec{Q}_6$	$\vec{K}_5 \cdot \vec{Q}_7$	$\vec{K}_5 \cdot \vec{Q}_8$	
$\boxed{\text{the}} \rightarrow \vec{E}_6 \xrightarrow{W_k} \vec{K}_6$	$\vec{K}_6 \cdot \vec{Q}_1$	$\vec{K}_6 \cdot \vec{Q}_2$	$\vec{K}_6 \cdot \vec{Q}_3$	$\vec{K}_6 \cdot \vec{Q}_4$	$\vec{K}_6 \cdot \vec{Q}_5$	$\vec{K}_6 \cdot \vec{Q}_6$	$\vec{K}_6 \cdot \vec{Q}_7$	$\vec{K}_6 \cdot \vec{Q}_8$	
$\boxed{\text{verdant}} \rightarrow \vec{E}_7 \xrightarrow{W_k} \vec{K}_7$	$\vec{K}_7 \cdot \vec{Q}_1$	$\vec{K}_7 \cdot \vec{Q}_2$	$\vec{K}_7 \cdot \vec{Q}_3$	$\vec{K}_7 \cdot \vec{Q}_4$	$\vec{K}_7 \cdot \vec{Q}_5$	$\vec{K}_7 \cdot \vec{Q}_6$	$\vec{K}_7 \cdot \vec{Q}_7$	$\vec{K}_7 \cdot \vec{Q}_8$	
$\boxed{\text{forest}} \rightarrow \vec{E}_8 \xrightarrow{W_k} \vec{K}_8$	$\vec{K}_8 \cdot \vec{Q}_1$	$\vec{K}_8 \cdot \vec{Q}_2$	$\vec{K}_8 \cdot \vec{Q}_3$	$\vec{K}_8 \cdot \vec{Q}_4$	$\vec{K}_8 \cdot \vec{Q}_5$	$\vec{K}_8 \cdot \vec{Q}_6$	$\vec{K}_8 \cdot \vec{Q}_7$	$\vec{K}_8 \cdot \vec{Q}_8$	

“Attend to”

	a	fluffy	blue	creature	roamed	the
	$\downarrow$ $\vec{E}_1$ $\downarrow_{W_Q}$ $\vec{Q}_1$	$\downarrow$ $\vec{E}_2$ $\downarrow_{W_Q}$ $\vec{Q}_2$	$\downarrow$ $\vec{E}_3$ $\downarrow_{W_Q}$ $\vec{Q}_3$	$\downarrow$ $\vec{E}_4$ $\downarrow_{W_Q}$ $\vec{Q}_4$	$\downarrow$ $\vec{E}_5$ $\downarrow_{W_Q}$ $\vec{Q}_5$	$\downarrow$ $\vec{E}_6$ $\downarrow_{W_Q}$ $\vec{Q}_6$
$\boxed{\text{a}} \rightarrow \vec{E}_1 \xrightarrow{W_k} \vec{K}_1$	.	.	.	.	.	.
$\boxed{\text{fluffy}} \rightarrow \vec{E}_2 \xrightarrow{W_k} \vec{K}_2$	.	.	.	+93.0	.	.
$\boxed{\text{blue}} \rightarrow \vec{E}_3 \xrightarrow{W_k} \vec{K}_3$	.	.	.	+93.4	.	.
$\boxed{\text{creature}} \rightarrow \vec{E}_4 \xrightarrow{W_k} \vec{K}_4$	.	.	.	.	.	.

$$\text{Attention}(Q, K, V) = \text{softmax}\left(\frac{QK^T}{\sqrt{d_k}}\right)V$$

	Q_1	Q_2	Q_3	Q_4	Q_5	$\dots$	Q_n
K_1	$\frac{Q_1 \cdot K_1}{\sqrt{d_k}}$	$\frac{Q_2 \cdot K_1}{\sqrt{d_k}}$	$\frac{Q_3 \cdot K_1}{\sqrt{d_k}}$	$\frac{Q_4 \cdot K_1}{\sqrt{d_k}}$	$\frac{Q_5 \cdot K_1}{\sqrt{d_k}}$	$\dots$	$\frac{Q_n \cdot K_1}{\sqrt{d_k}}$
K_2	$\frac{Q_1 \cdot K_2}{\sqrt{d_k}}$	$\frac{Q_2 \cdot K_2}{\sqrt{d_k}}$	$\frac{Q_3 \cdot K_2}{\sqrt{d_k}}$	$\frac{Q_4 \cdot K_2}{\sqrt{d_k}}$	$\frac{Q_5 \cdot K_2}{\sqrt{d_k}}$	$\dots$	$\frac{Q_n \cdot K_2}{\sqrt{d_k}}$
K_3	$\frac{Q_1 \cdot K_3}{\sqrt{d_k}}$	$\frac{Q_2 \cdot K_3}{\sqrt{d_k}}$	$\frac{Q_3 \cdot K_3}{\sqrt{d_k}}$	$\frac{Q_4 \cdot K_3}{\sqrt{d_k}}$	$\frac{Q_5 \cdot K_3}{\sqrt{d_k}}$	$\dots$	$\frac{Q_n \cdot K_3}{\sqrt{d_k}}$
K_4	$\frac{Q_1 \cdot K_4}{\sqrt{d_k}}$	$\frac{Q_2 \cdot K_4}{\sqrt{d_k}}$	$\frac{Q_3 \cdot K_4}{\sqrt{d_k}}$	$\frac{Q_4 \cdot K_4}{\sqrt{d_k}}$	$\frac{Q_5 \cdot K_4}{\sqrt{d_k}}$	$\dots$	$\frac{Q_n \cdot K_4}{\sqrt{d_k}}$
K_5	$\frac{Q_1 \cdot K_5}{\sqrt{d_k}}$	$\frac{Q_2 \cdot K_5}{\sqrt{d_k}}$	$\frac{Q_3 \cdot K_5}{\sqrt{d_k}}$	$\frac{Q_4 \cdot K_5}{\sqrt{d_k}}$	$\frac{Q_5 \cdot K_5}{\sqrt{d_k}}$	$\dots$	$\frac{Q_n \cdot K_5}{\sqrt{d_k}}$
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\dots$	$\vdots$

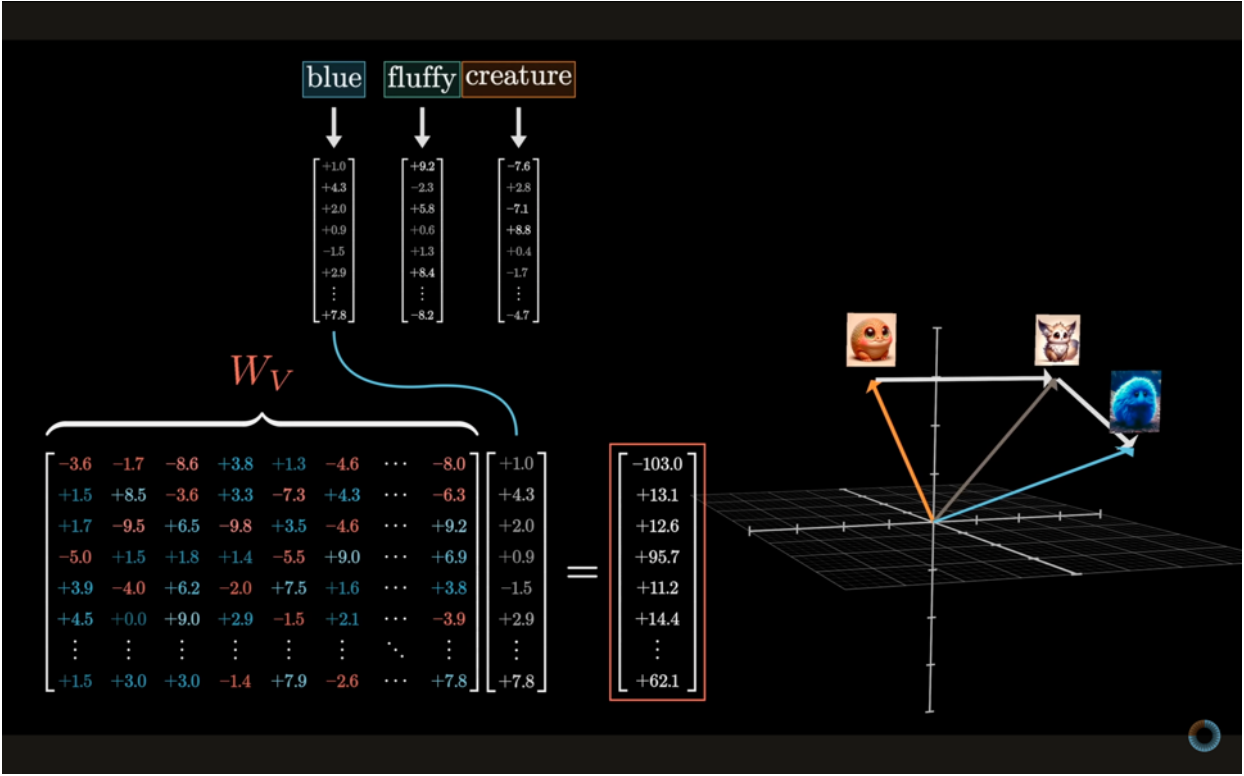
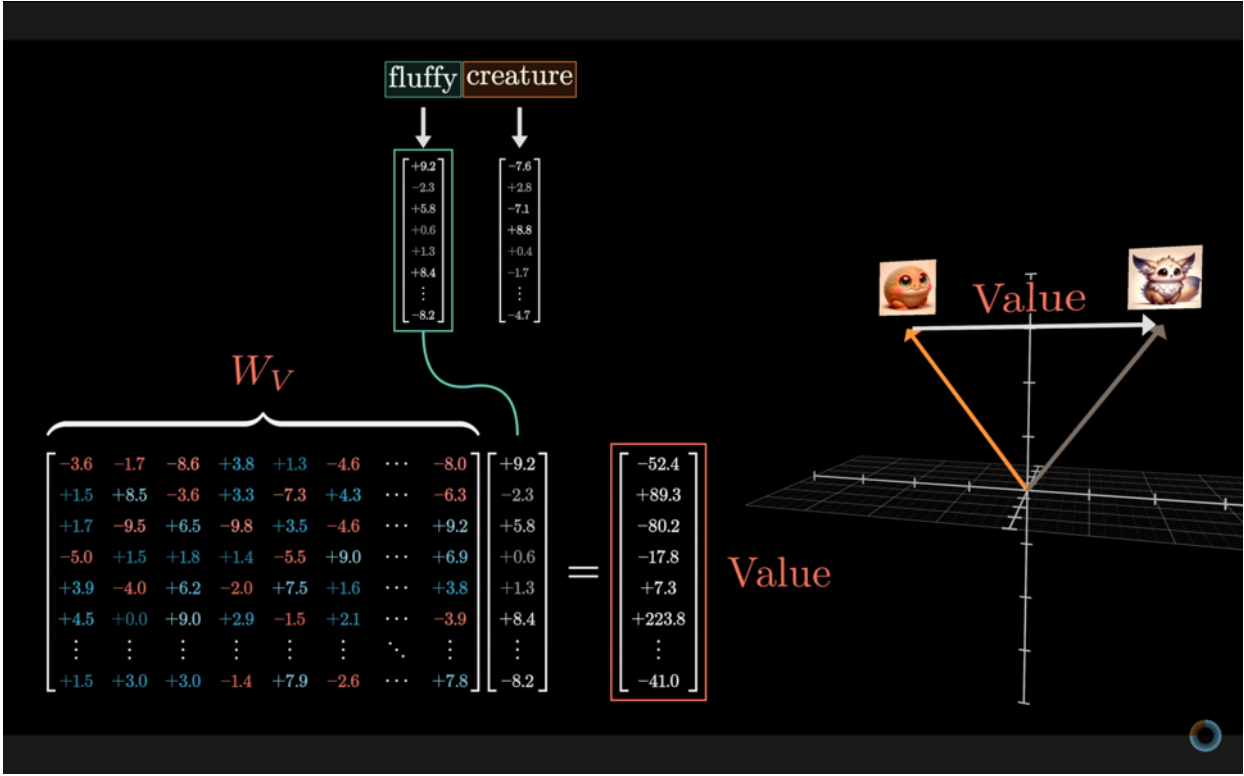
Unnormalized
Attention Pattern

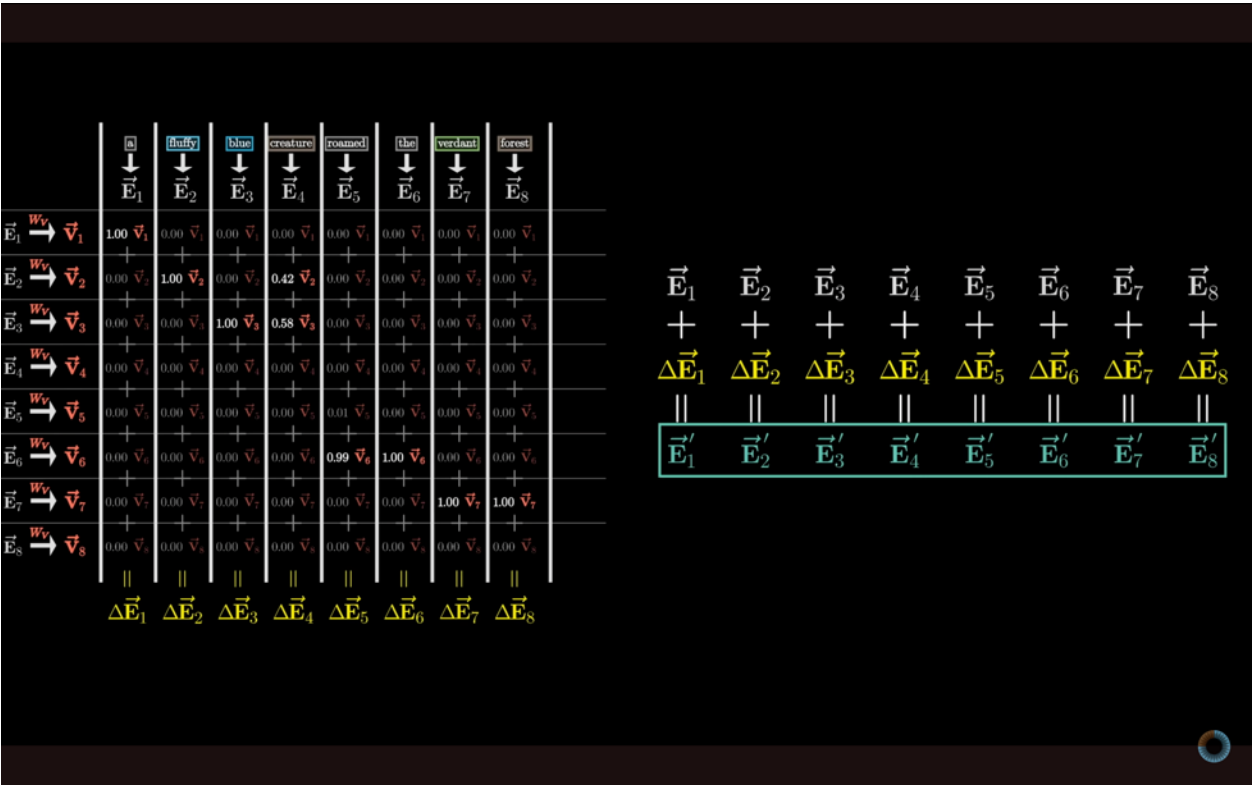
+3.53	+0.80	+1.96	+4.48	+3.74	-1.95
$-\infty$	-0.30	-0.21	+0.82	+0.29	+2.91
$-\infty$	$-\infty$	+0.89	+0.67	+2.99	-0.41
$-\infty$	$-\infty$	$-\infty$	+1.31	+1.73	-1.48
$-\infty$	$-\infty$	$-\infty$	$-\infty$	+3.07	+2.94
$-\infty$	$-\infty$	$-\infty$	$-\infty$	$-\infty$	+0.31

softmax →

Normalized
Attention Pattern

1.00	0.75	0.69	0.92	0.46	0.00
0.00	0.25	0.08	0.02	0.01	0.46
0.00	0.00	0.24	0.02	0.22	0.02
0.00	0.00	0.00	0.04	0.06	0.01
0.00	0.00	0.00	0.00	0.24	0.48
0.00	0.00	0.00	0.00	0.00	0.03





$\vec{E}_1$

$\vec{E}_2$

$\vec{E}_3$

$\vec{E}_4$

$\vec{E}_5$

$\vec{E}_6$

$\vec{E}_7$

$\vec{E}_8$

+

+

+

+

+

+

+

+

$\Delta \vec{E}_1$

$\Delta \vec{E}_2$

$\Delta \vec{E}_3$

$\Delta \vec{E}_4$

$\Delta \vec{E}_5$

$\Delta \vec{E}_6$

$\Delta \vec{E}_7$

$\Delta \vec{E}_8$

$\parallel$

$\parallel$

$\parallel$

$\parallel$

$\parallel$

$\parallel$

$\parallel$

$\parallel$

$\vec{E}'_1$

$\vec{E}'_2$

$\vec{E}'_3$

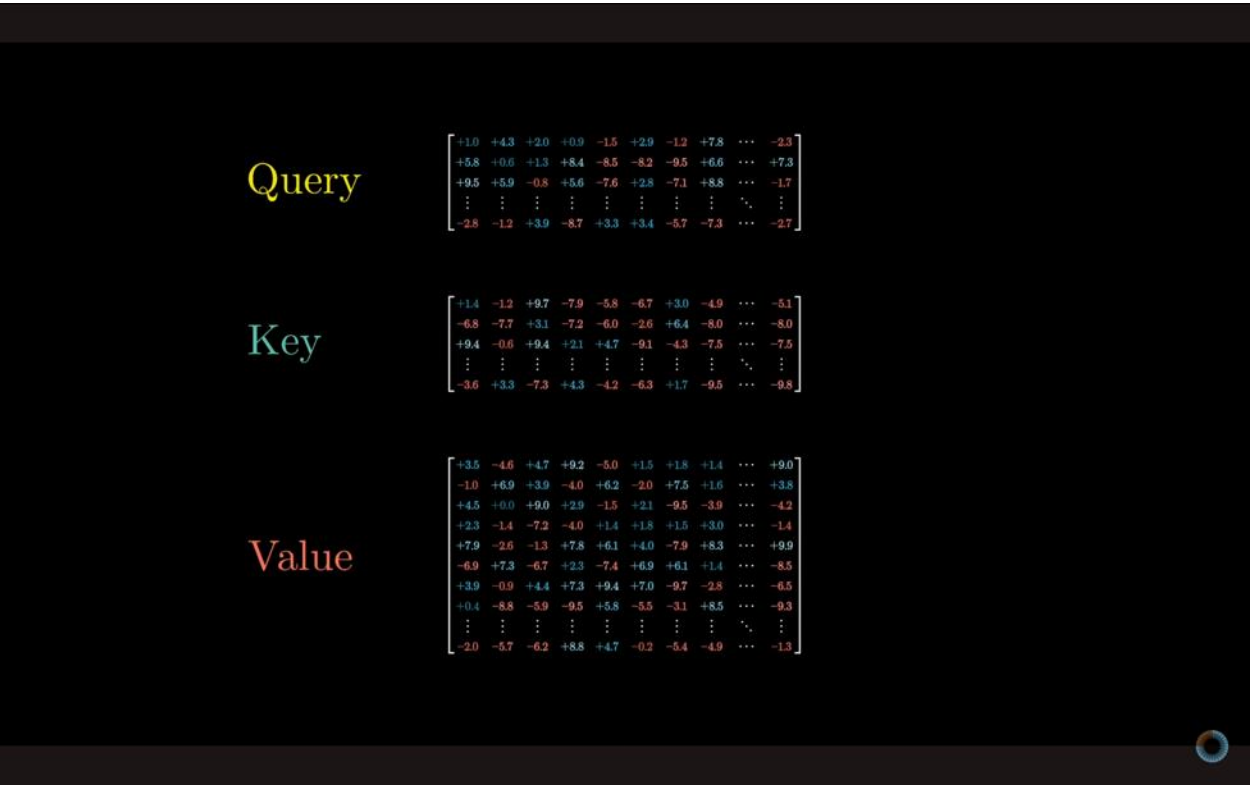
$\vec{E}'_4$

$\vec{E}'_5$

$\vec{E}'_6$

$\vec{E}'_7$

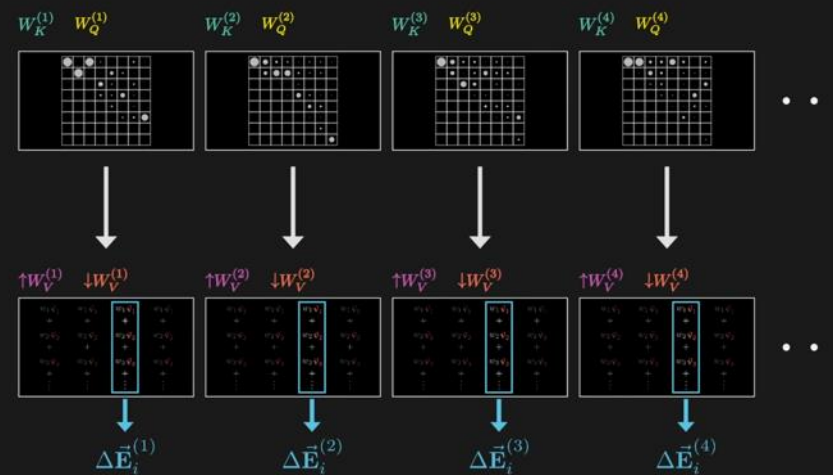
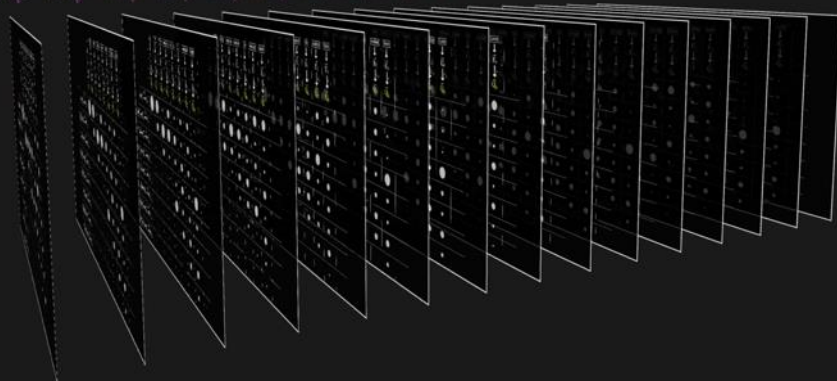
$\vec{E}'_8$



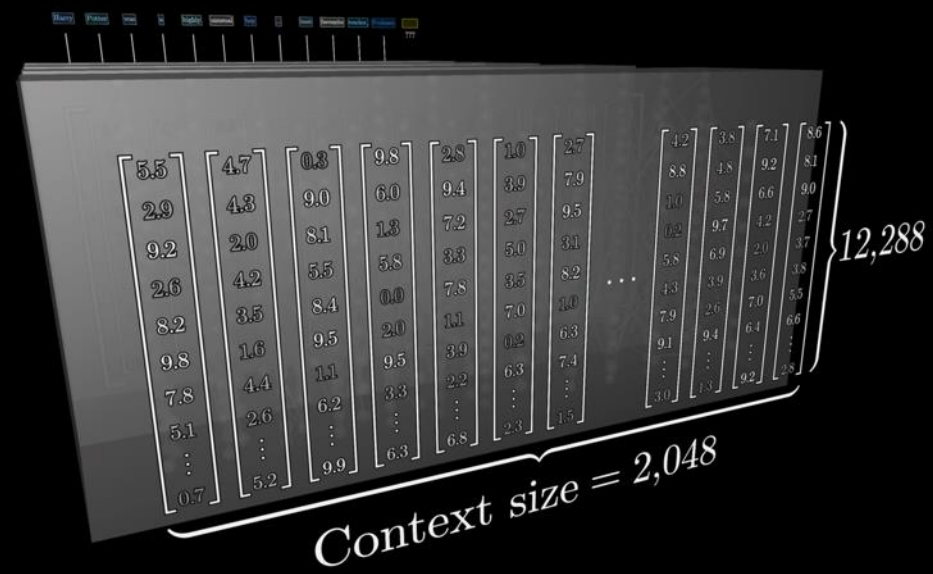
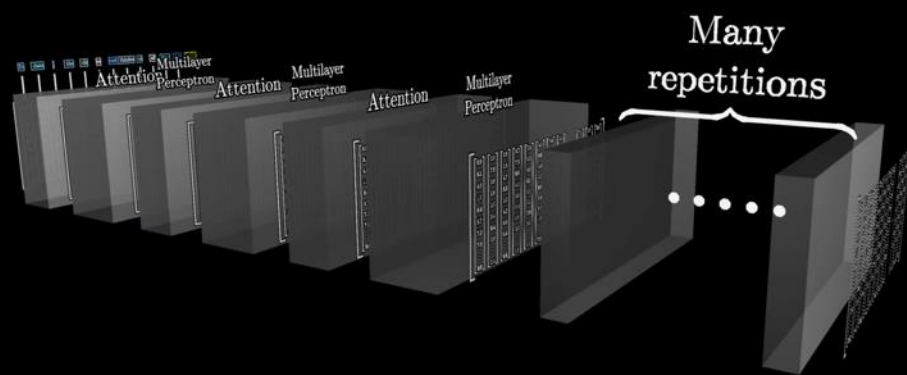
Multi-headed attention

$W_Q^{(1)} W_Q^{(2)} W_Q^{(3)} W_Q^{(4)} W_Q^{(5)} W_Q^{(6)} W_Q^{(7)} W_Q^{(8)} W_Q^{(9)} \dots$
 $W_K^{(1)} W_K^{(2)} W_K^{(3)} W_K^{(4)} W_K^{(5)} W_K^{(6)} W_K^{(7)} W_K^{(8)} W_K^{(9)} \dots$
 $\downarrow W_V^{(1)} \downarrow W_V^{(2)} \downarrow W_V^{(3)} \downarrow W_V^{(4)} \downarrow W_V^{(5)} \downarrow W_V^{(6)} \downarrow W_V^{(7)} \downarrow W_V^{(8)} \downarrow W_V^{(9)} \dots$

96



Original embedding $\vec{E}_i + \Delta \vec{E}_i^{(1)} + \Delta \vec{E}_i^{(2)} + \Delta \vec{E}_i^{(3)} + \Delta \vec{E}_i^{(4)} + \dots$



Transformers

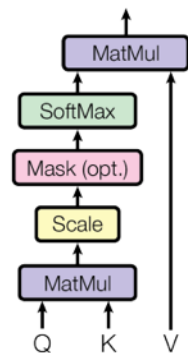
Advantage: we can calculate attention for all the embeddings in a sequence simultaneously. This speeds up computation and allows for massive scalability

$$\text{Attention}(Q, K, V) = \text{softmax}\left(\frac{QK^T}{\sqrt{d_k}}\right)V$$

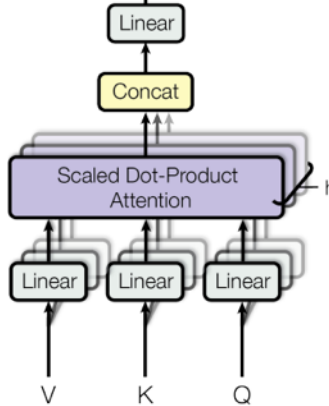
$$\text{MultiHead}(Q, K, V) = \text{Concat}(\text{head}_1, \dots, \text{head}_h)W^O$$

where $\text{head}_i = \text{Attention}(QW_i^Q, KW_i^K, VW_i^V)$

Scaled Dot-Product Attention



Multi-Head Attention



Vaswani, A. "Attention is all you need." *Advances in Neural Information Processing Systems* (2017).

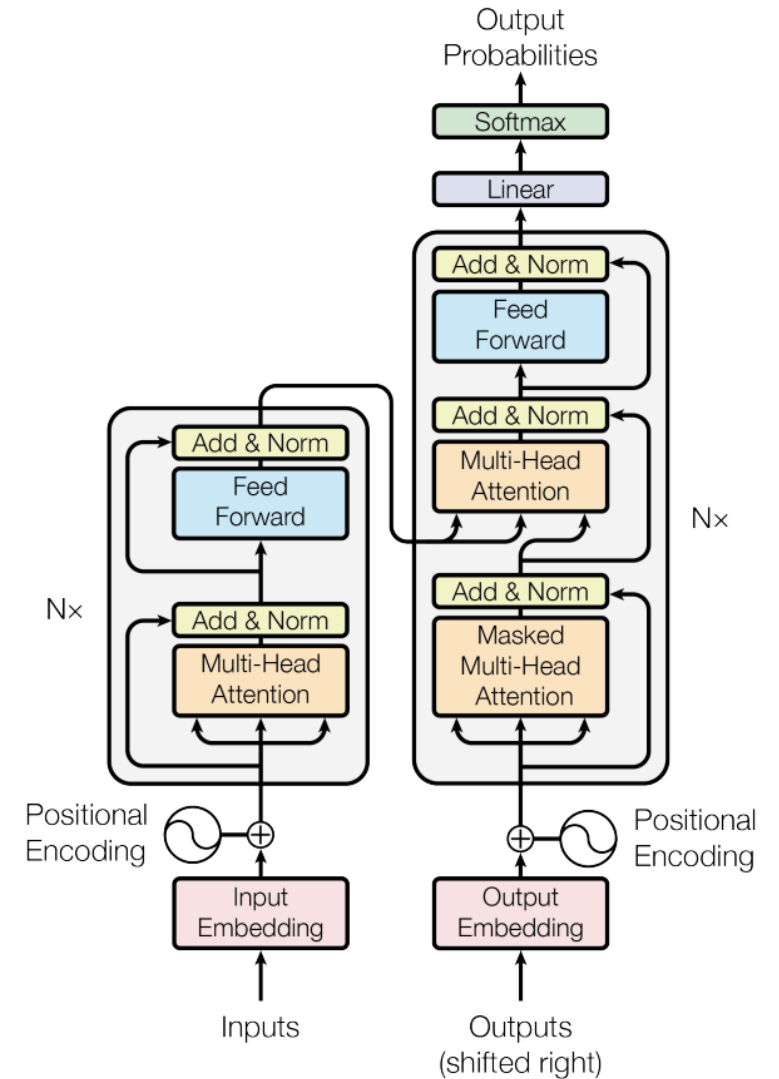


Figure 1: The Transformer - model architecture.