

Neural Networks with Fully Connected Layers in Practice

Fundamentals of Artificial Intelligence

Fabien Cromieres

Kyoto University

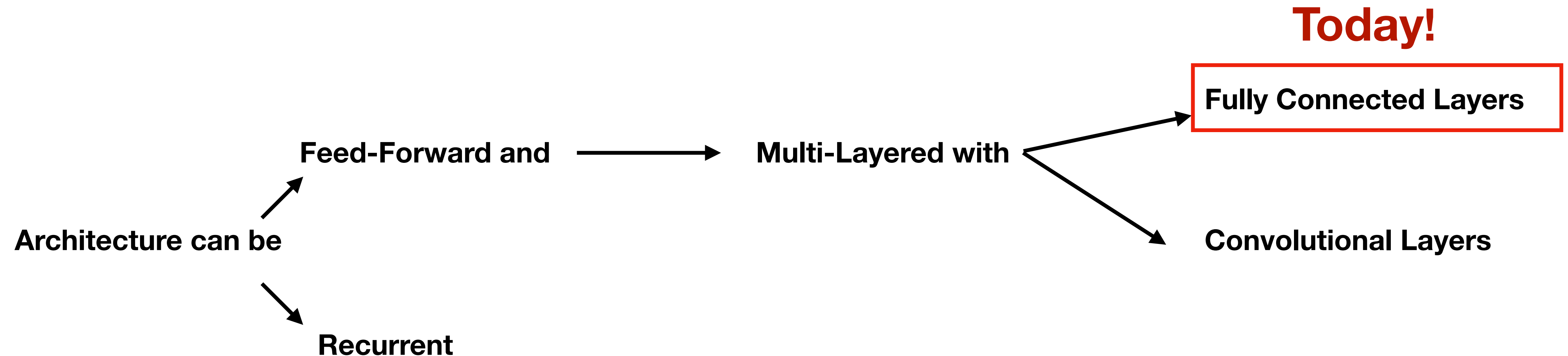
<http://lotus.kuee.kyoto-u.ac.jp/~fabien/lectures/IA/>

Program for today:

- 1- Discuss the mathematical representation of Fully-connected layers in Neural Networks
- 2- Define and train a real Neural Network in a Jupyter Notebook

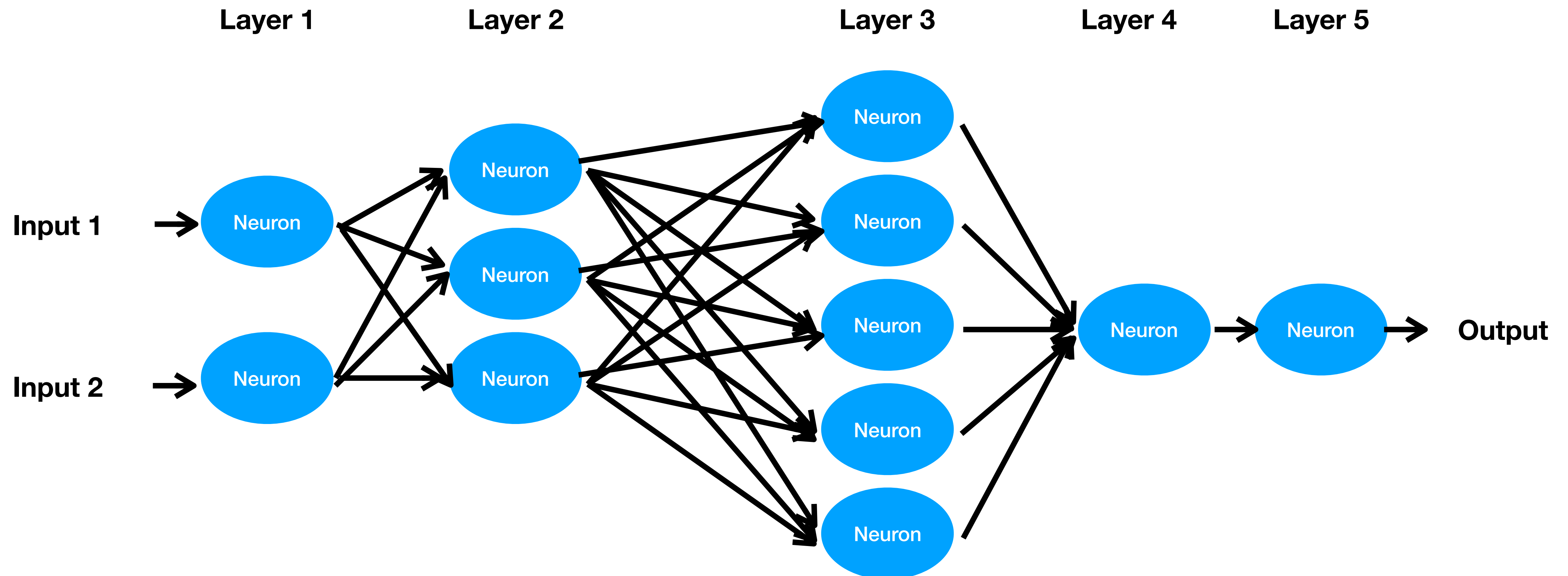
Neural Network Architectures

- We are still here:



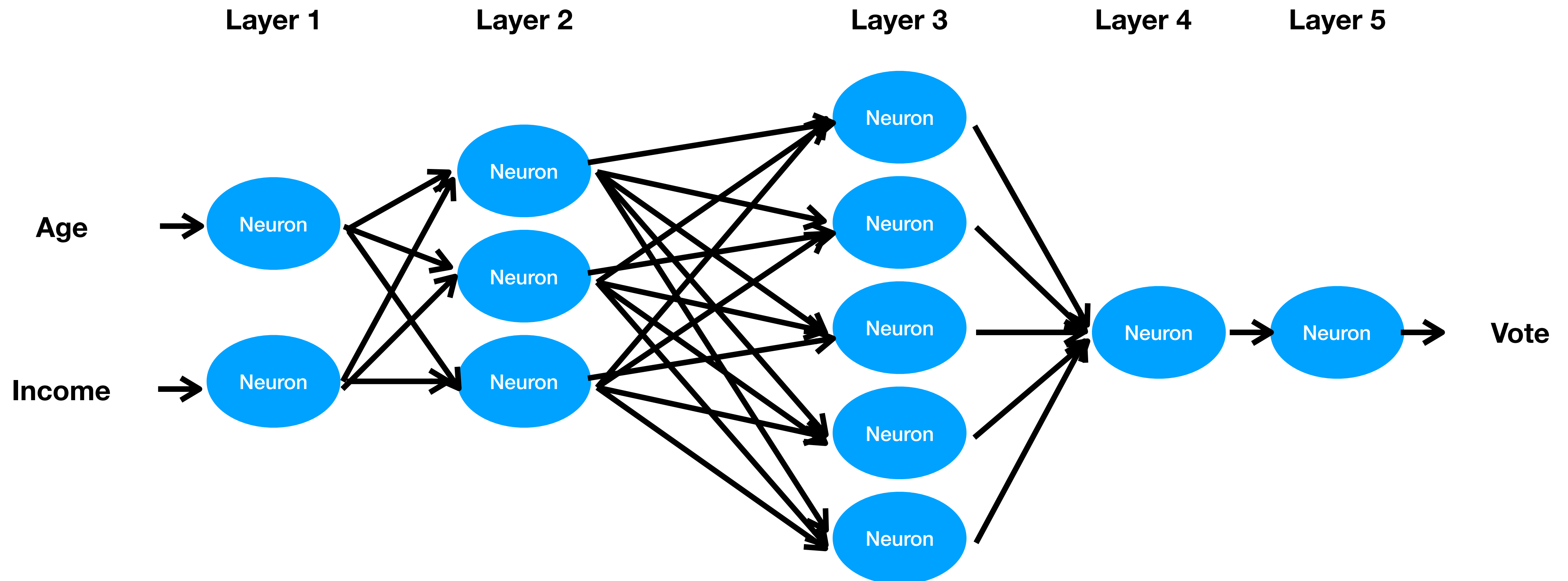
Feed-Forward networks with fully connected layers

- Therefore, we are going to consider this type of Neural Network:



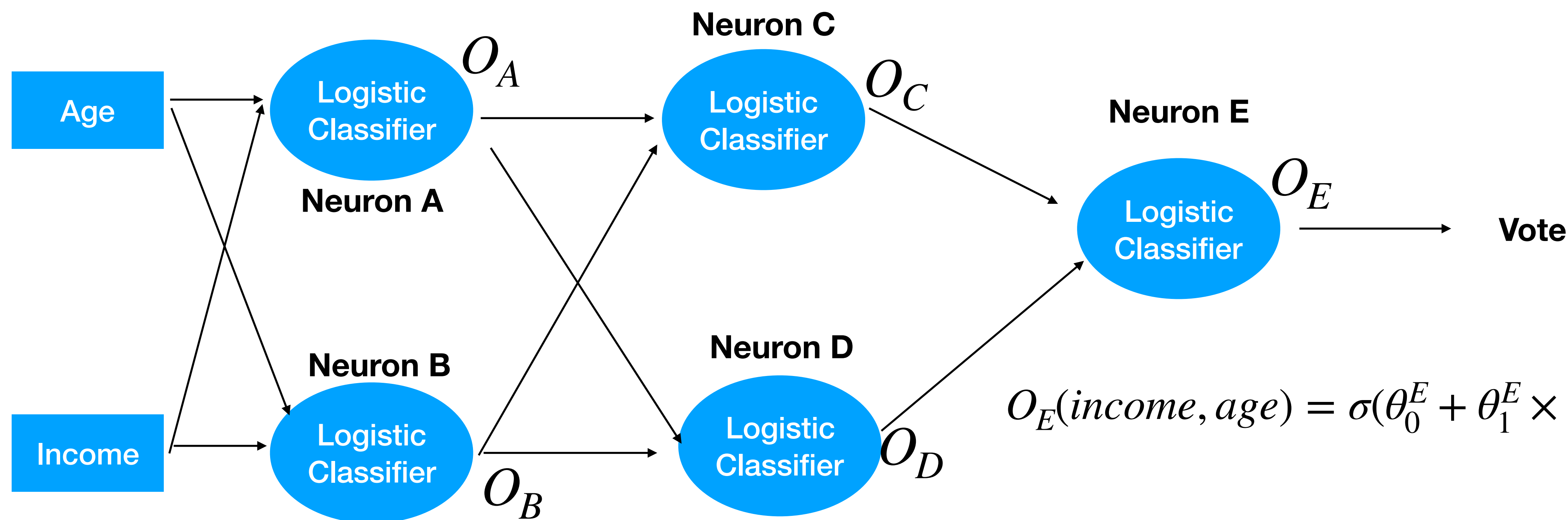
Feed-Forward networks with fully connected layers

- Therefore, we are going to consider this type of Neural Network:



Keeping in mind what this type of graph mean

$$O_A(\text{income}, \text{age}) = \sigma(\theta_0^A + \theta_1^A \times \text{income} + \theta_2^A \times \text{age}) \quad O_C(\text{income}, \text{age}) = \sigma(\theta_0^C + \theta_1^C \times O_A + \theta_2^C \times O_B)$$



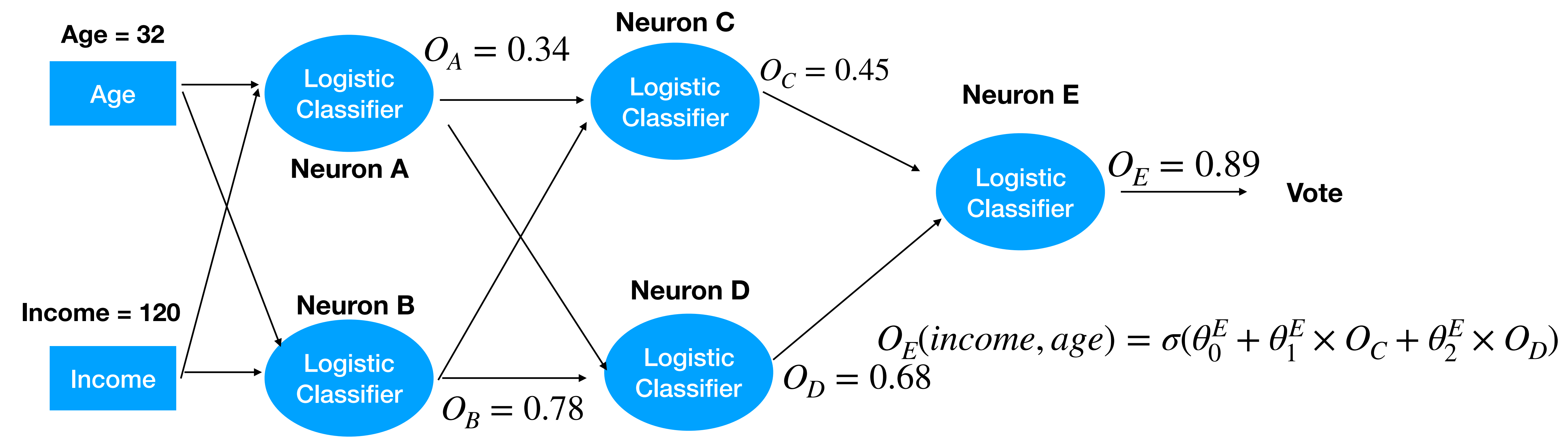
$$O_E(\text{income}, \text{age}) = \sigma(\theta_0^E + \theta_1^E \times O_C + \theta_2^E \times O_D)$$

$$O_B(\text{income}, \text{age}) = \sigma(\theta_0^B + \theta_1^B \times \text{income} + \theta_2^B \times \text{age})$$

$$O_D(\text{income}, \text{age}) = \sigma(\theta_0^D + \theta_1^D \times O_A + \theta_2^D \times O_B)$$

Keeping in mind what this type of graph mean

$$O_A(\text{income}, \text{age}) = \sigma(\theta_0^A + \theta_1^A \times \text{income} + \theta_2^A \times \text{age}) \qquad O_C(\text{income}, \text{age}) = \sigma(\theta_0^C + \theta_1^C \times O_A + \theta_2^C \times O_B)$$



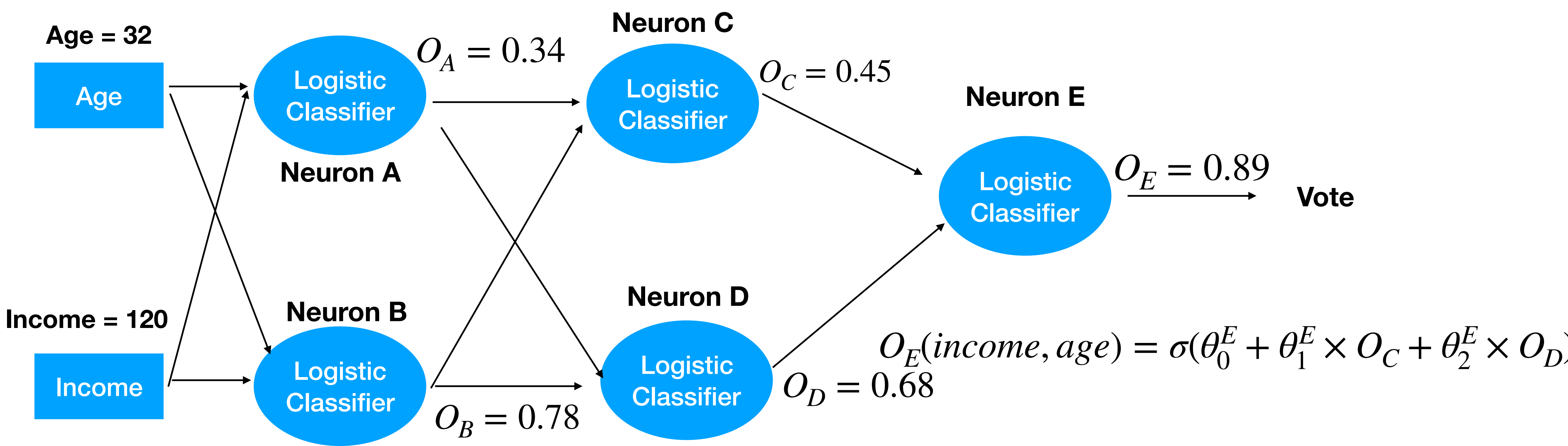
$$O_E(\text{income}, \text{age}) = \sigma(\theta_0^E + \theta_1^E \times O_C + \theta_2^E \times O_D)$$

$$O_B(\text{income}, \text{age}) = \sigma(\theta_0^B + \theta_1^B \times \text{income} + \theta_2^B \times \text{age}) \qquad O_D(\text{income}, \text{age}) = \sigma(\theta_0^D + \theta_1^D \times O_A + \theta_2^D \times O_B)$$

Keeping in mind what this type of graph mean

—> Each Neural Network architecture defines a function of the input with parameters θ

$$O_A(\text{income}, \text{age}) = \sigma(\theta_0^A + \theta_1^A \times \text{income} + \theta_2^A \times \text{age}) \qquad O_C(\text{income}, \text{age}) = \sigma(\theta_0^C + \theta_1^C \times O_A + \theta_2^C \times O_B)$$



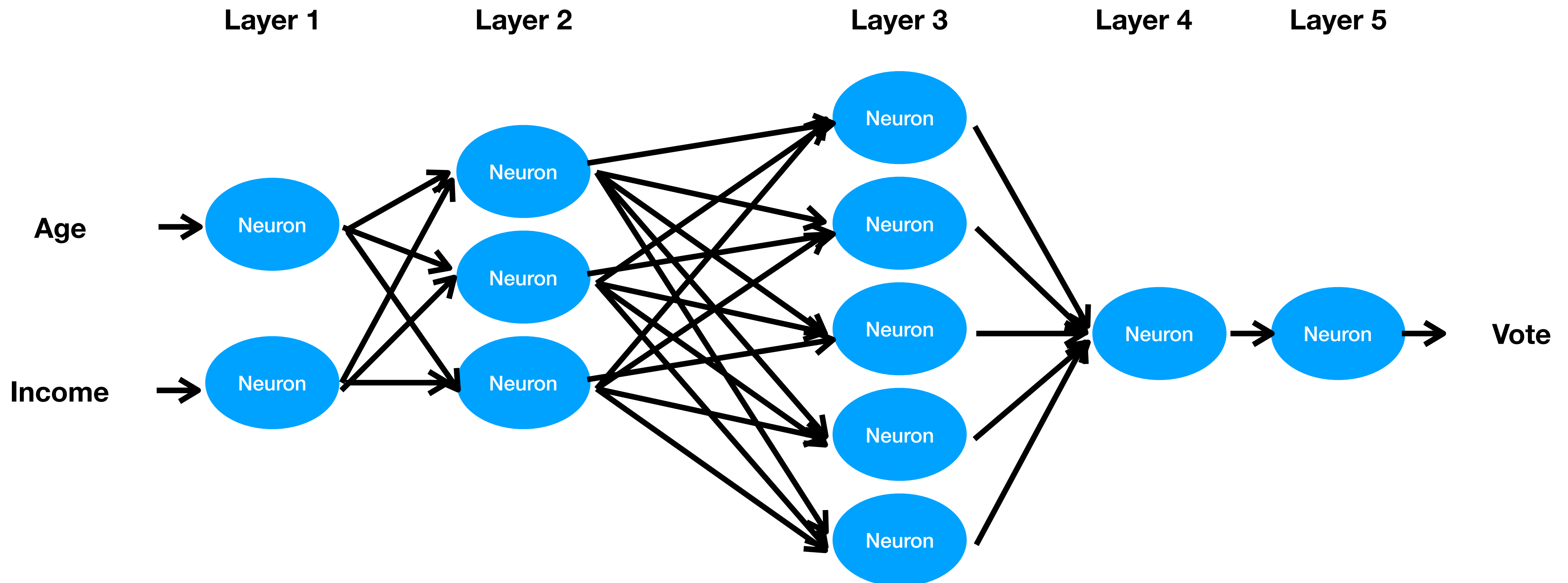
$$O_E(\text{income}, \text{age}) = \sigma(\theta_0^E + \theta_1^E \times O_C + \theta_2^E \times O_D)$$

$$O_B(\text{income}, \text{age}) = \sigma(\theta_0^B + \theta_1^B \times \text{income} + \theta_2^B \times \text{age}) \qquad O_D(\text{income}, \text{age}) = \sigma(\theta_0^D + \theta_1^D \times O_A + \theta_2^D \times O_B)$$

Feed-Forward networks with fully connected layers

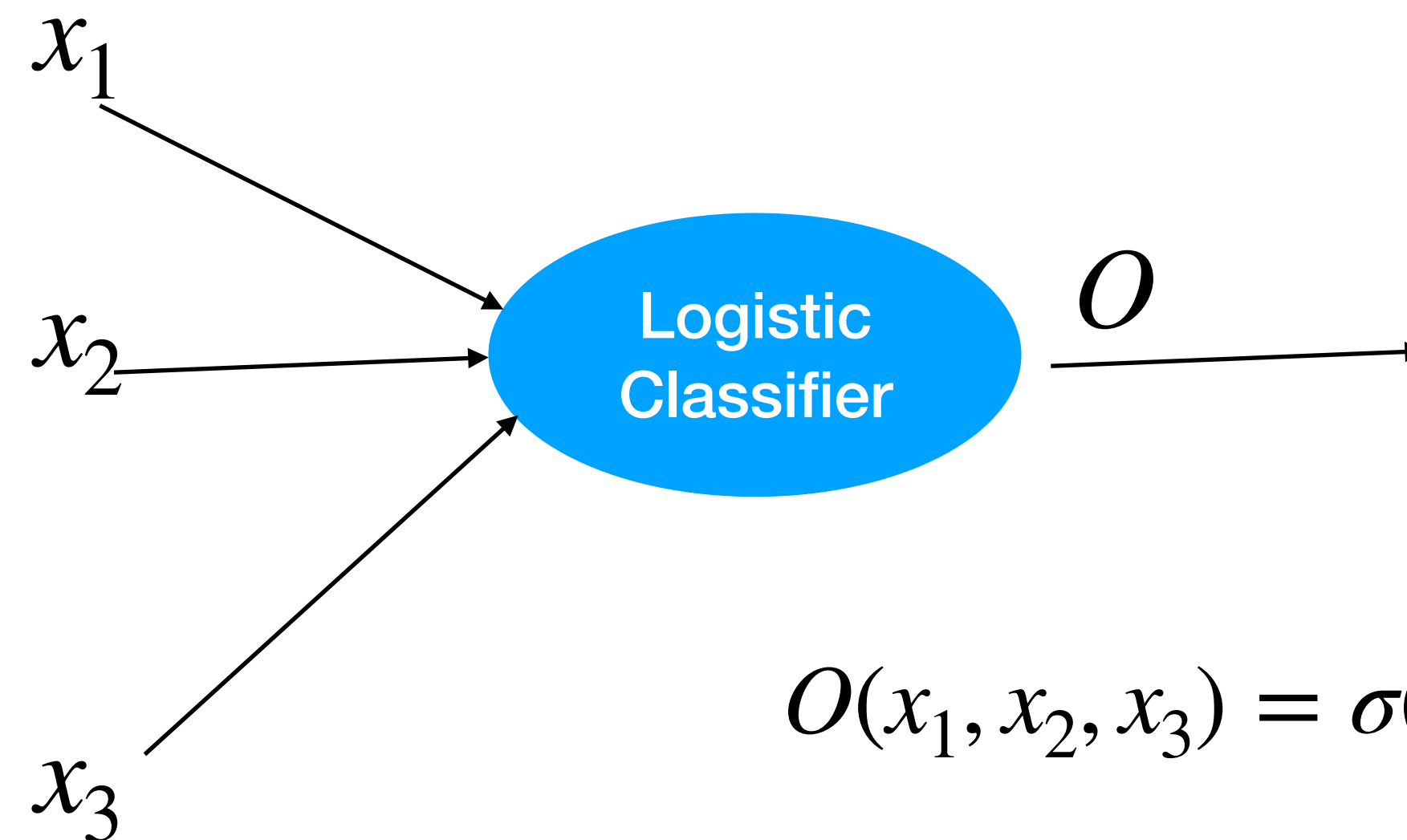
—> Each Neural Network architecture defines a function of the input with parameters θ

- Therefore, this is just a visual way of defining a complicated parameterized function of ***Vote*** given ***Age*** and ***Income***:



Parameters

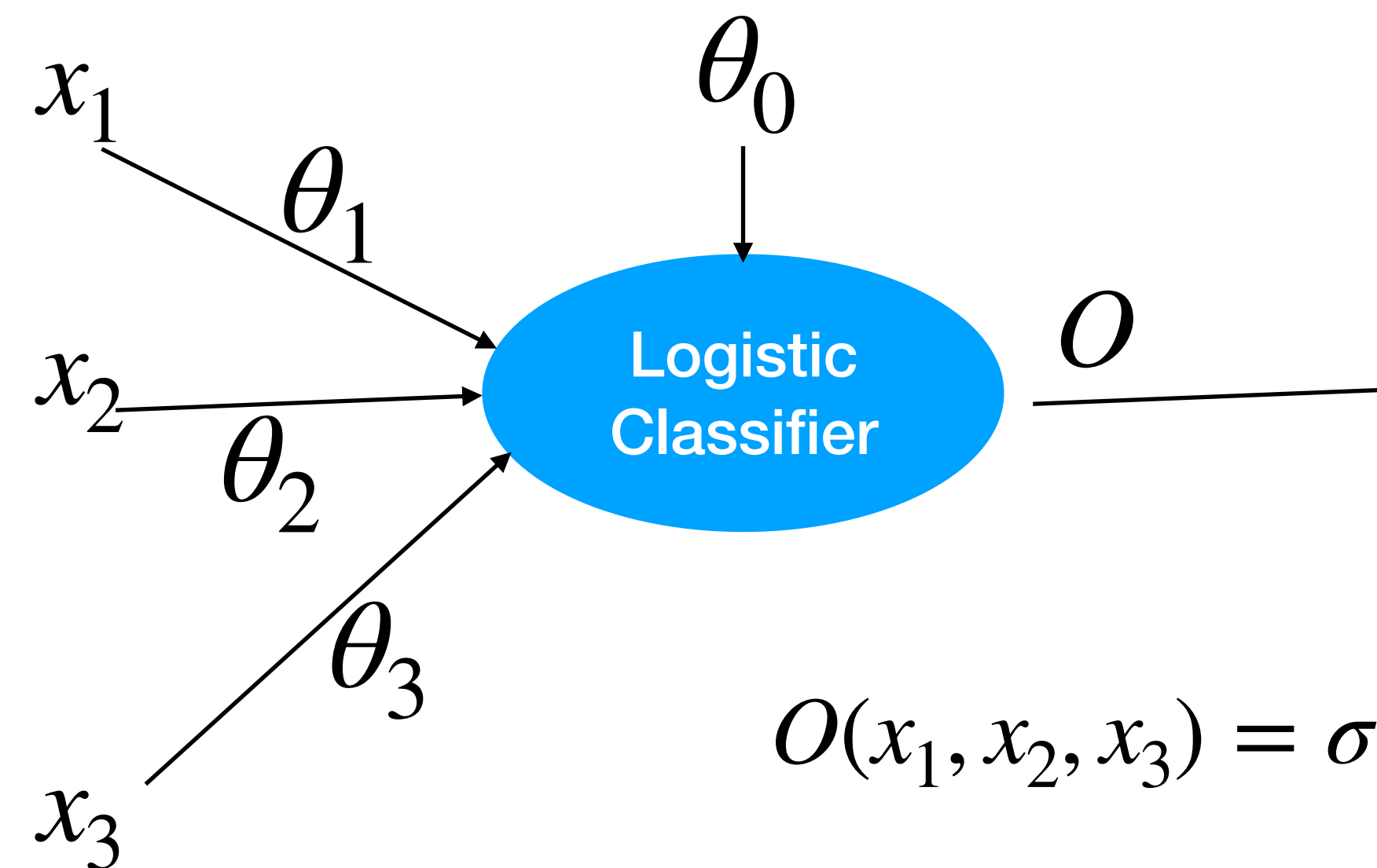
- If a neuron has N inputs, it has N+1 parameters



$$O(x_1, x_2, x_3) = \sigma(\theta_0 + \theta_1 \times x_1 + \theta_2 \times x_2 + \theta_3 \times x_3)$$

Parameters

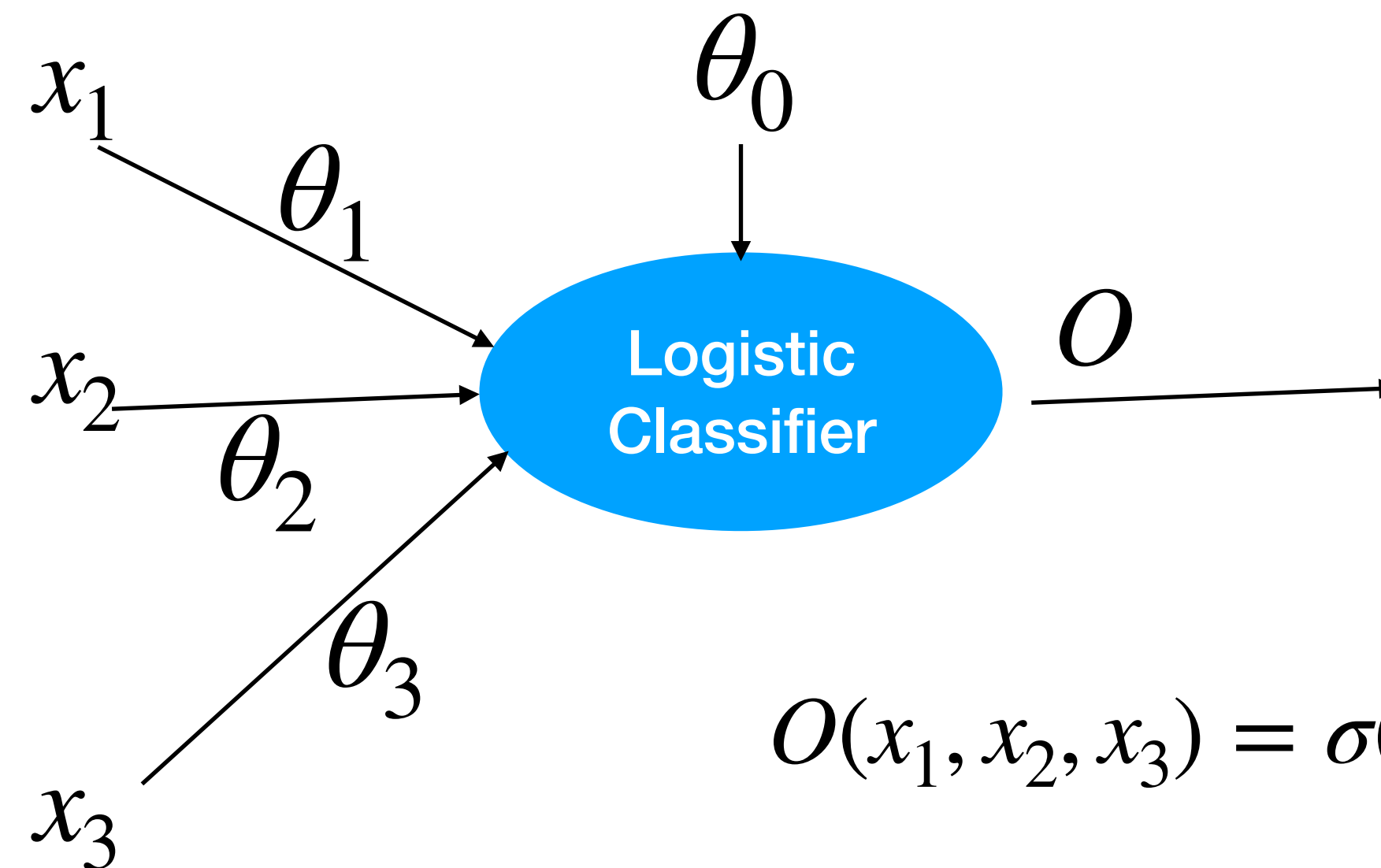
- If a neuron has N inputs, it has N+1 parameters
- Visually, we can associate a parameter to each input, and show θ_0 separately



$$O(x_1, x_2, x_3) = \sigma(\theta_0 + \theta_1 \times x_1 + \theta_2 \times x_2 + \theta_3 \times x_3)$$

Parameters: Terminology

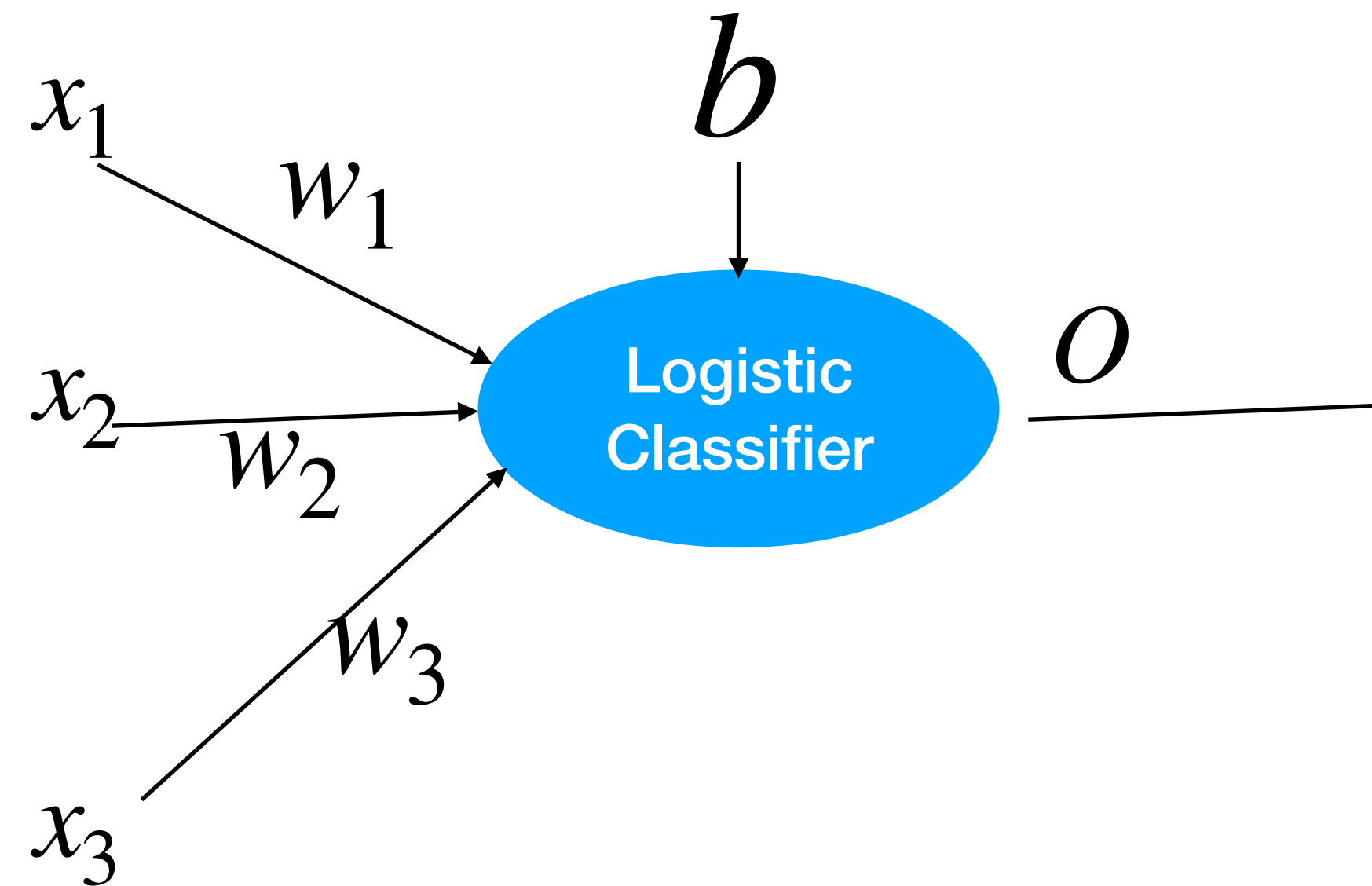
- $\theta_1, \theta_2, \theta_3$ are often called the ***weights*** of the neuron
- θ_0 is often called the ***bias*** of the neuron



$$O(x_1, x_2, x_3) = \sigma(\theta_0 + \theta_1 \times x_1 + \theta_2 \times x_2 + \theta_3 \times x_3)$$

Parameters: Terminology

- $\theta_1, \theta_2, \theta_3$ are often called the **weights** of the neuron
 - They are therefore often also noted w_1, w_2, w_3
- θ_0 is often called the **bias** of the neuron
 - It is often noted **b**



$$O(x_1, x_2, x_3) = \sigma(b + w_1 \times x_1 + w_2 \times x_2 + w_3 \times x_3)$$

Linear Algebra

- In order to discuss the way Neural Networks are actually implemented, we need to discuss some mathematical concepts
- We discuss:
 - Vector scalar products (a.k.a inner product)
 - Matrices
 - Matrix-vector multiplication
 - Matrix-Matrix multiplication
- Warning: it is a bit ambitious to explain that in 30 minutes. Hopefully you understand most of it.

Row Vectors and Column Vectors

- We have seen that **vectors** are just a “list” of numbers
- We are actually going to distinguish 2 “types” of vectors: **row vectors** and **column vectors**



Column vectors

- Column vectors of the same dimension can be added

(5-dimensional)
Column vector

0.2
3.6
2.1
5.3
-2.2

(5-dimensional)
Column vector

1.1
-2.0
1.1
-5.3
-1.0

+

=

(5-dimensional)
Column vector

1.3
1.6
3.2
0.0
-3.2

Column vectors

- Column vectors can be multiplied by a number

$$\begin{array}{ccc} & \begin{array}{c} \text{(5-dimensional)} \\ \text{Column vector} \end{array} & \begin{array}{c} \text{(5-dimensional)} \\ \text{Column vector} \end{array} \\ 2.0 & \times & \begin{array}{c} 1.1 \\ -2.0 \\ 1.1 \\ -5.3 \\ -1.0 \end{array} = \begin{array}{c} 2.2 \\ -4.0 \\ 2.2 \\ -10.6 \\ -2.0 \end{array} \end{array}$$

Row vectors

- Row vectors have the same operations as Column vectors
- They can be added:

$$\begin{array}{c} \text{(5-dimensional)} \\ \text{Row vector} \end{array} \quad \begin{array}{|c|c|c|c|c|} \hline 0.2 & 3.6 & 2.1 & 5.3 & -2.2 \\ \hline \end{array} \quad + \quad \begin{array}{c} \text{(5-dimensional)} \\ \text{Row vector} \end{array} \quad \begin{array}{|c|c|c|c|c|} \hline 1.1 & -2.0 & 1.1 & -5.3 & -1.0 \\ \hline \end{array} \quad = \quad \begin{array}{c} \text{(5-dimensional)} \\ \text{Row vector} \end{array} \quad \begin{array}{|c|c|c|c|c|} \hline 1.3 & 1.6 & 3.2 & 0.0 & -3.2 \\ \hline \end{array}$$

- They can be multiplied by a number:

$$2.0 \quad \times \quad \begin{array}{|c|c|c|c|c|} \hline 0.2 & 3.6 & 2.1 & 5.3 & -2.2 \\ \hline \end{array} \quad = \quad \begin{array}{|c|c|c|c|c|} \hline 0.4 & 7.2 & 4.2 & 10.6 & -4.4 \\ \hline \end{array}$$

Row vectors and Column Vectors

- Row vectors and Column vectors cannot be added together

(5-dimensional)
Row vector

0.2 3.6 2.1 5.3 -2.2

+

(5-dimensional)
Column vector

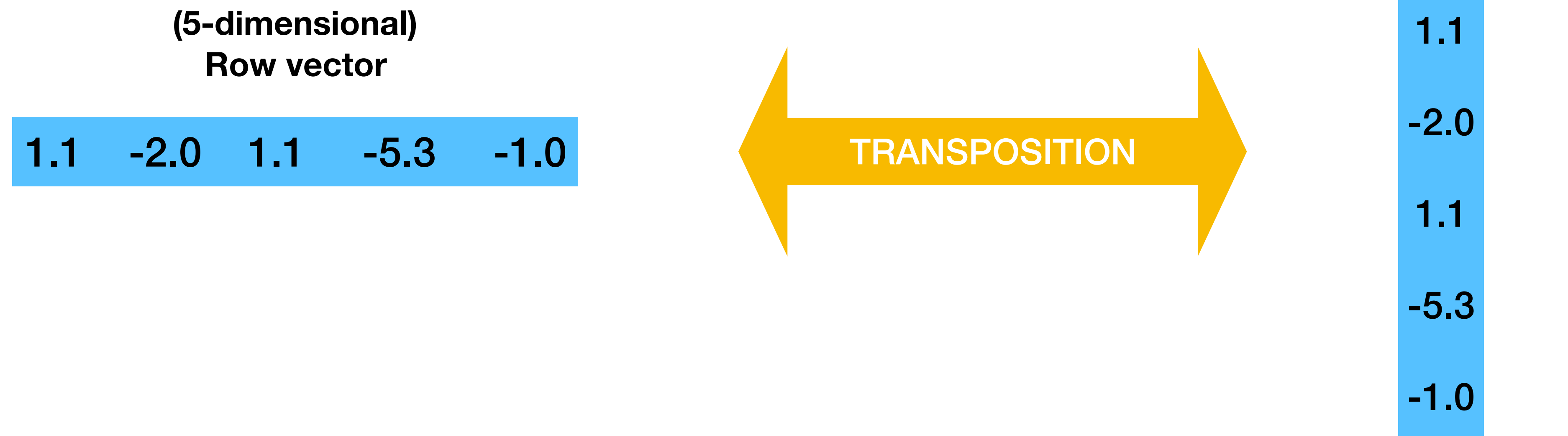
2.2
-4.0
2.2
-10.6
-2.0

=

X

Transposition

- However, *row vectors* can be transformed in *column vectors* by an operation called ***transposition*** (and vice-versa)



Transposition

- Transposition is usually noted with a ***T*** in exponent:

The diagram illustrates the process of transposition for a 1D array. It shows three stages of the operation:

- A horizontal blue bar representing a 1D array with the values: 1.1, -2.0, 1.1, -5.3, -1.0. A large black **T** is positioned to its right.
- An equals sign (=) follows.
- A vertical blue bar representing a 2D array (column vector) with the values: 1.1, -2.0, 1.1, -5.3, -1.0. A thick green vertical line is positioned to its right.
- Another equals sign (=) follows.
- A second horizontal blue bar, identical to the first, representing the original 1D array: 1.1, -2.0, 1.1, -5.3, -1.0. A large black **T** is positioned to its left.

Inner product (a.k.a Scalar product)

- We can compute the *inner product* of a *row vector* and a *column vector* to obtain a single number
- It consists in taking the product of the number dimension by dimension and then taking the sum

$$\begin{array}{c} \text{(3-dimensional)} \\ \text{Row vector} \end{array} \begin{array}{|c|c|c|} \hline 2 & 0.0 & -1.0 \\ \hline \end{array} \bullet \begin{array}{c} \text{(3-dimensional)} \\ \text{Column vector} \end{array} \begin{array}{|c|} \hline 1.5 \\ \hline 1.0 \\ \hline 2.0 \\ \hline \end{array} = 2 \times 1.5 + 0 \times 1 + -1 \times 2 = 1.0$$

Inner product (a.k.a Scalar product)

- The vectors should have the **same dimension**
- You should have the row vector on the left, and the column vector on the right
 - (if it is the opposite, the operation is called *outer product* and gives a different result)

$$\begin{array}{c} \text{(3-dimensional)} \\ \text{Row vector} \\ \hline 2 \quad 0.0 \quad -1.0 \end{array} \quad \bullet \quad \begin{array}{c} \text{(3-dimensional)} \\ \text{Column vector} \\ \hline 1.5 \\ 1.0 \\ 2.0 \end{array} = 1.0$$

Inner product (a.k.a Scalar product)

$$\begin{bmatrix} 2 & 0.0 & -1.0 \end{bmatrix} \bullet \begin{bmatrix} 1.5 \\ 1.0 \\ 2.0 \end{bmatrix} = 1.0$$

You can try to visualize this as the column vector lying down on the row vector to produce the number

Inner product (a.k.a Scalar product)

$$\begin{bmatrix} 1.5 & 1.0 & 2.0 \end{bmatrix} \cdot \begin{bmatrix} 2 \\ 0.0 \\ -1.0 \end{bmatrix} = 1.0$$

$3.0 + 0.0 + -2.0 \rightarrow 1.0$

You can try to visualize this as the column vector lying down on the row vector to produce the number

Why inner product is important?

- It allows us to **express linear functions efficiently**
- And remember that linear functions are one of the **fundamental components** of Machine Learning

Linear regression

$$f(x, y) = \theta_0 + \theta_1 \times x + \theta_2 \times y$$

Logistic Classifier

$$score(x, y) = \theta_0 + \theta_1 \times x + \theta_2 \times y$$

$$prediction = \sigma(score)$$

Neuron (same formula as Logistic Classifier)

$$output(x, y) = \sigma(\theta_0 + \theta_1 \times x + \theta_2 \times y)$$

Why inner product is important?

$$\text{vote}(\text{age}, \text{income}) = \sigma(\theta_0 + \theta_1 \times \text{age} + \theta_2 \times \text{income})$$

$$\begin{bmatrix} 1 & \text{income} & \text{age} \end{bmatrix} \bullet \begin{bmatrix} \Theta_0 \\ \Theta_1 \\ \Theta_2 \end{bmatrix} = \theta_0 + \theta_1 \times \text{income} + \theta_2 \times \text{age}$$

$$\text{vote}(\text{age}, \text{income}) = \sigma\left(\begin{bmatrix} 1 & \text{income} & \text{age} \end{bmatrix} \bullet \begin{bmatrix} \Theta_0 \\ \Theta_1 \\ \Theta_2 \end{bmatrix} \right)$$

Representing many inner product at once

$$\text{vote}(\text{age}, \text{income}) = \sigma(\theta_0 + \theta_1 \times \text{age} + \theta_2 \times \text{income})$$

1	income ₀	age ₀
1	income ₁	age ₁
1	income ₂	age ₂
1	income ₃	age ₃

 $\bullet$

Θ_0
Θ_1
Θ_2

 $=$

Representing many inner product at once

$$\text{vote}(\text{age}, \text{income}) = \sigma(\theta_0 + \theta_1 \times \text{age} + \theta_2 \times \text{income})$$

1	income ₀	age ₀
1	income ₁	age ₁
1	income ₂	age ₂
1	income ₃	age ₃

 $\bullet$

Θ_0
Θ_1
Θ_2

 $=$

Representing many inner product at once

$$\text{vote}(\text{age}, \text{income}) = \sigma(\theta_0 + \theta_1 \times \text{age} + \theta_2 \times \text{income})$$

The diagram illustrates the vectorization of the inner product formula. On the left, a matrix of features is shown with four rows and three columns. The columns are labeled θ_0 , θ_1 , and θ_2 at the top. The rows represent four data points, each with a bias term of 1, an income value, and an age value. Red arrows indicate the dot product of each row with the parameter vector θ . The matrix is followed by a dot product symbol $\cdot$ and an equals sign $=$. On the right, the resulting linear combinations are listed for each row.

	θ_0	θ_1	θ_2
1	income ₀	age ₀	
1	income ₁	age ₁	
1	income ₂	age ₂	
1	income ₃	age ₃	

$\cdot$

$=$

$\theta_0 + \theta_1 \times \text{income}_0 + \theta_2 \times \text{age}_0$

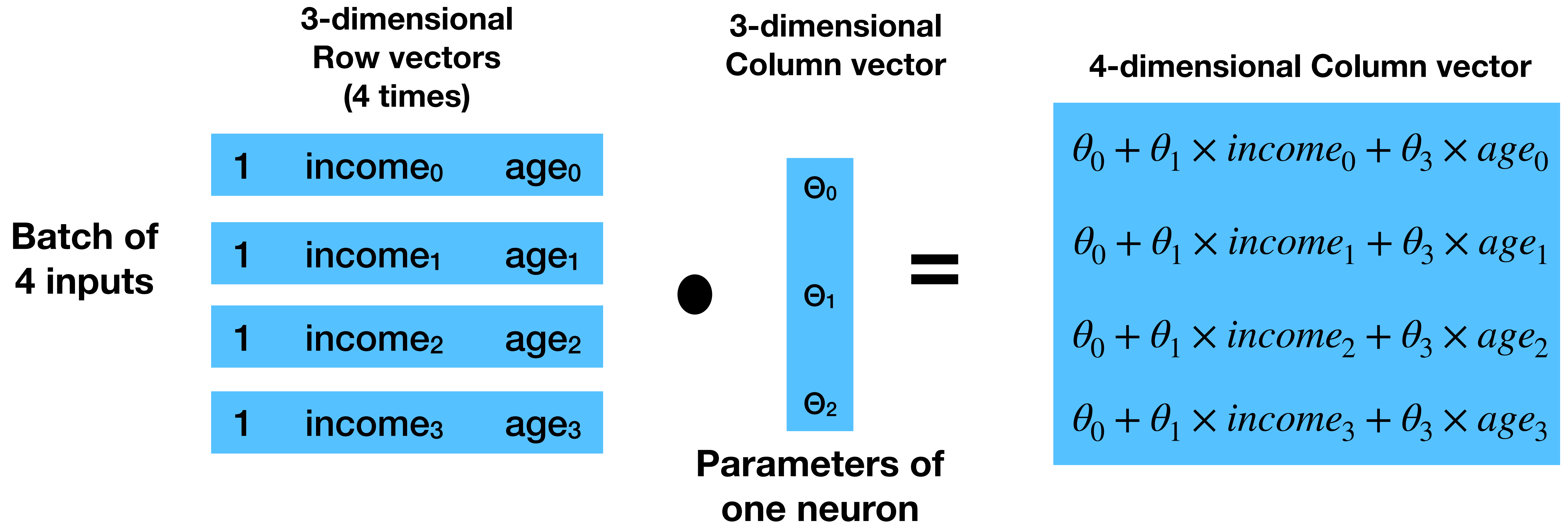
$\theta_0 + \theta_1 \times \text{income}_1 + \theta_2 \times \text{age}_1$

$\theta_0 + \theta_1 \times \text{income}_2 + \theta_2 \times \text{age}_2$

$\theta_0 + \theta_1 \times \text{income}_3 + \theta_2 \times \text{age}_3$

Representing many inner product at once

$$vote(age, income) = \sigma(\theta_0 + \theta_1 \times age + \theta_2 \times income)$$



This represents the computation of the score output of a single neuron for 4 inputs at the same time!

Representing many inner product at once

$$vote(age, income) = \sigma(\theta_0 + \theta_1 \times age + \theta_2 \times income)$$

3-dimensional
Row vectors
(4 times)

Batch of
4 inputs

1	income ₀	age ₀
1	income ₁	age ₁
1	income ₂	age ₂
1	income ₃	age ₃

This is called a 4x3 matrix

3-dimensional
Column vector

$$\begin{bmatrix} \Theta_0 \\ \Theta_1 \\ \Theta_2 \end{bmatrix}$$

4-dimensional Column vector

$$\begin{bmatrix} \theta_0 + \theta_1 \times income_0 + \theta_3 \times age_0 \\ \theta_0 + \theta_1 \times income_1 + \theta_3 \times age_1 \\ \theta_0 + \theta_1 \times income_2 + \theta_3 \times age_2 \\ \theta_0 + \theta_1 \times income_3 + \theta_3 \times age_3 \end{bmatrix}$$

Representing many inner product at once

$$vote(age, income) = \sigma(\theta_0 + \theta_1 \times age + \theta_2 \times income)$$

3-dimensional
Row vectors
(4 times)

1	income ₀	age ₀
1	income ₁	age ₁
1	income ₂	age ₂
1	income ₃	age ₃

3-dimensional
Column vector

$$\begin{bmatrix} \Theta_0 \\ \Theta_1 \\ \Theta_2 \end{bmatrix}$$

=

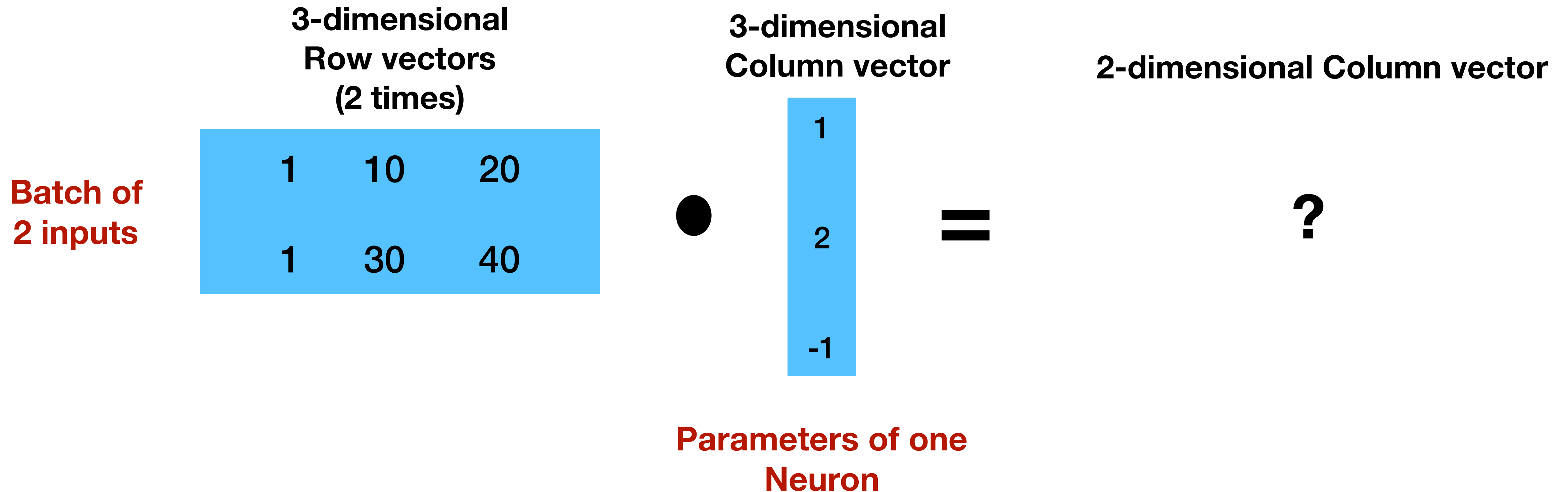
4-dimensional Column vector

$$\begin{bmatrix} \theta_0 + \theta_1 \times income_0 + \theta_3 \times age_0 \\ \theta_0 + \theta_1 \times income_1 + \theta_3 \times age_1 \\ \theta_0 + \theta_1 \times income_2 + \theta_3 \times age_2 \\ \theta_0 + \theta_1 \times income_3 + \theta_3 \times age_3 \end{bmatrix}$$

This is called a 4x3 matrix

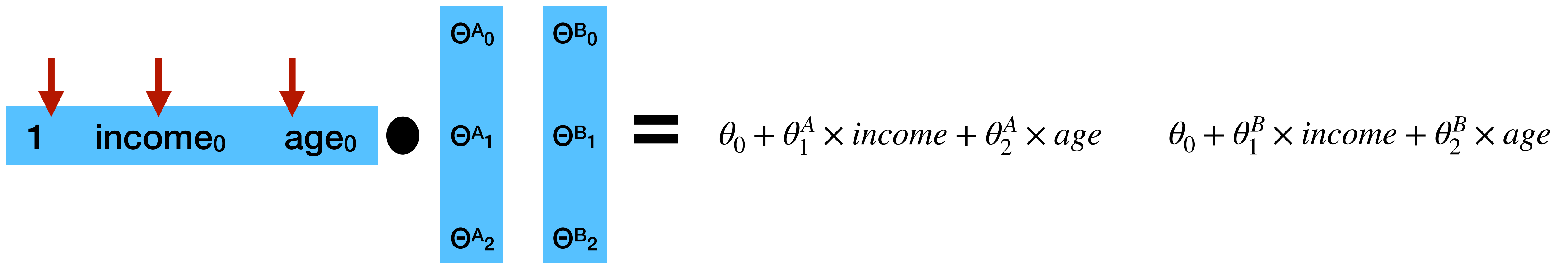
Representing many inner product at once

$$vote(age, income) = \sigma(\theta_0 + \theta_1 \times age + \theta_2 \times income)$$



Representing many inner product at once

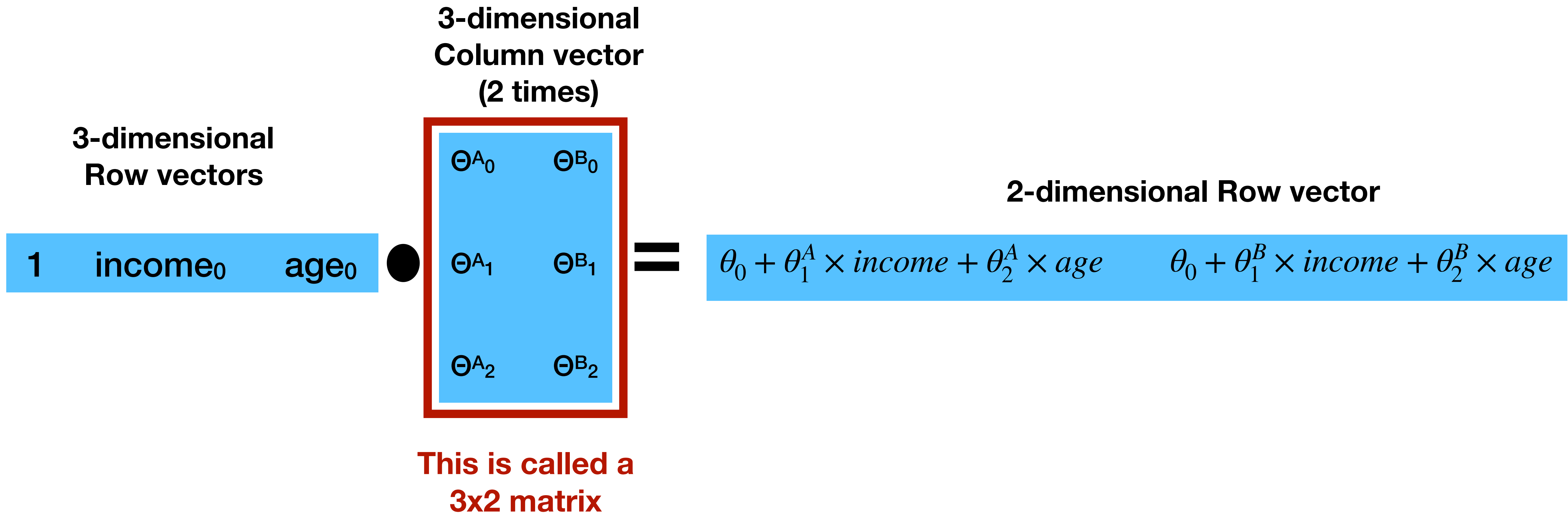
$$\text{vote}(\text{age}, \text{income}) = \sigma(\theta_0 + \theta_1 \times \text{income} + \theta_2 \times \text{age})$$



Parameters for
two different
neurons A and B

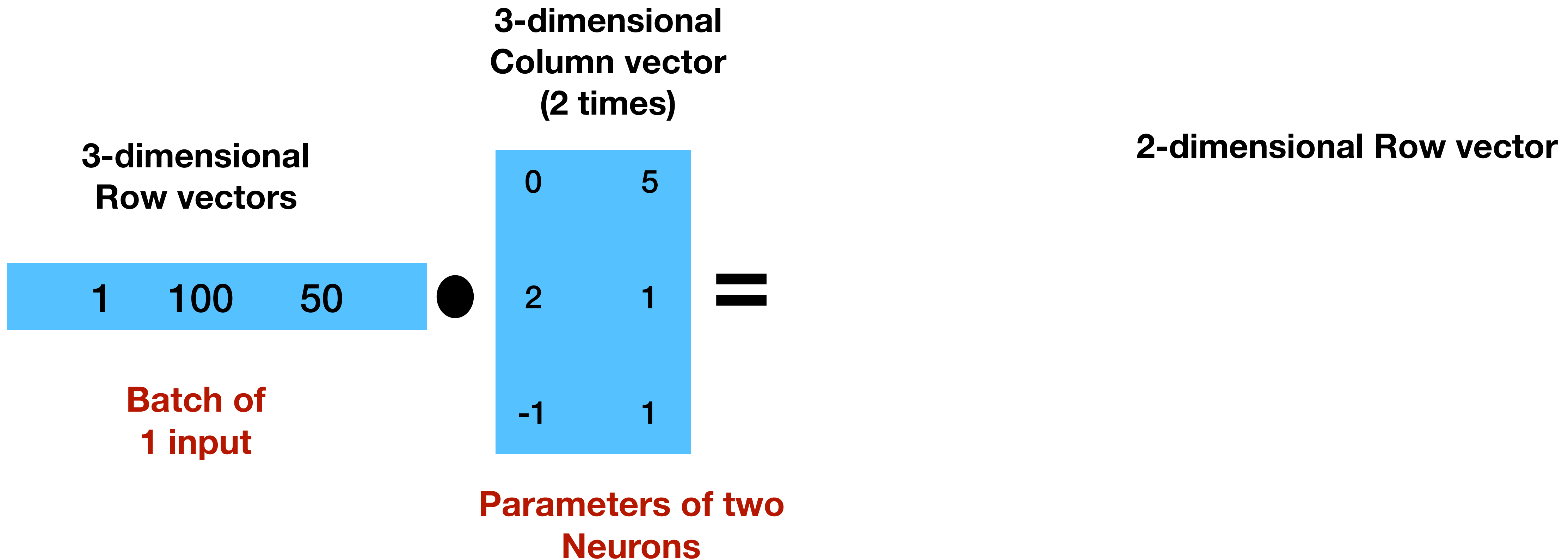
Representing many inner product at once

$$vote(age, income) = \sigma(\theta_0 + \theta_1 \times age + \theta_2 \times income)$$



Representing many inner product at once

$$\text{vote}(\text{age}, \text{income}) = \sigma(\theta_0 + \theta_1 \times \text{age} + \theta_2 \times \text{income})$$



Representing many inner product at once

- If we combine the 2 aspects of having several inputs and several neurons, we have what is called a ***matrix multiplication***

4x3 matrix

1	income ₀	age ₀
1	income ₁	age ₁
1	income ₂	age ₂
1	income ₃	age ₃

Batch of 4 inputs

3x2 matrix

Θ^A_0	Θ^B_0
Θ^A_1	Θ^B_1
Θ^A_2	Θ^B_2

Parameters for
2 neurons

4x2 matrix

$\theta_0 + \theta^A_1 \times \text{income}_0 + \theta^A_2 \times \text{age}_0$	$\theta_0 + \theta^B_1 \times \text{income}_0 + \theta^B_2 \times \text{age}_0$
$\theta_0 + \theta^A_1 \times \text{income}_1 + \theta^A_2 \times \text{age}_1$	$\theta_0 + \theta^B_1 \times \text{income}_1 + \theta^B_2 \times \text{age}_1$
$\theta_0 + \theta^A_1 \times \text{income}_2 + \theta^A_2 \times \text{age}_2$	$\theta_0 + \theta^B_1 \times \text{income}_2 + \theta^B_2 \times \text{age}_2$
$\theta_0 + \theta^A_1 \times \text{income}_3 + \theta^A_2 \times \text{age}_3$	$\theta_0 + \theta^B_1 \times \text{income}_3 + \theta^B_2 \times \text{age}_3$

Outputs of the
2 neurons for each of the 4 inputs

Matrix Multiplication

4x3 matrix

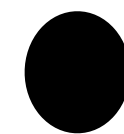
1	income ₀	age ₀
1	income ₁	age ₁
1	income ₂	age ₂
1	income ₃	age ₃

Batch of 4 inputs

3x2 matrix

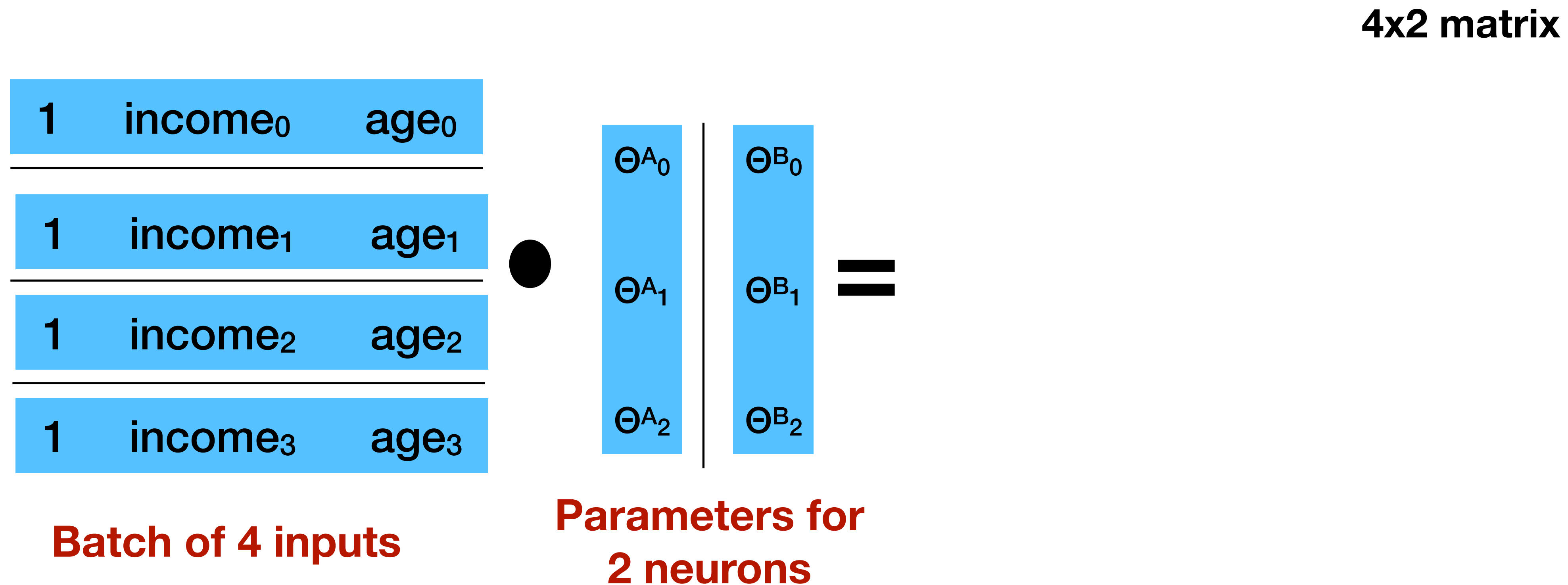
Θ^A_0	Θ^B_0
Θ^A_1	Θ^B_1
Θ^A_2	Θ^B_2

Parameters for
2 neurons

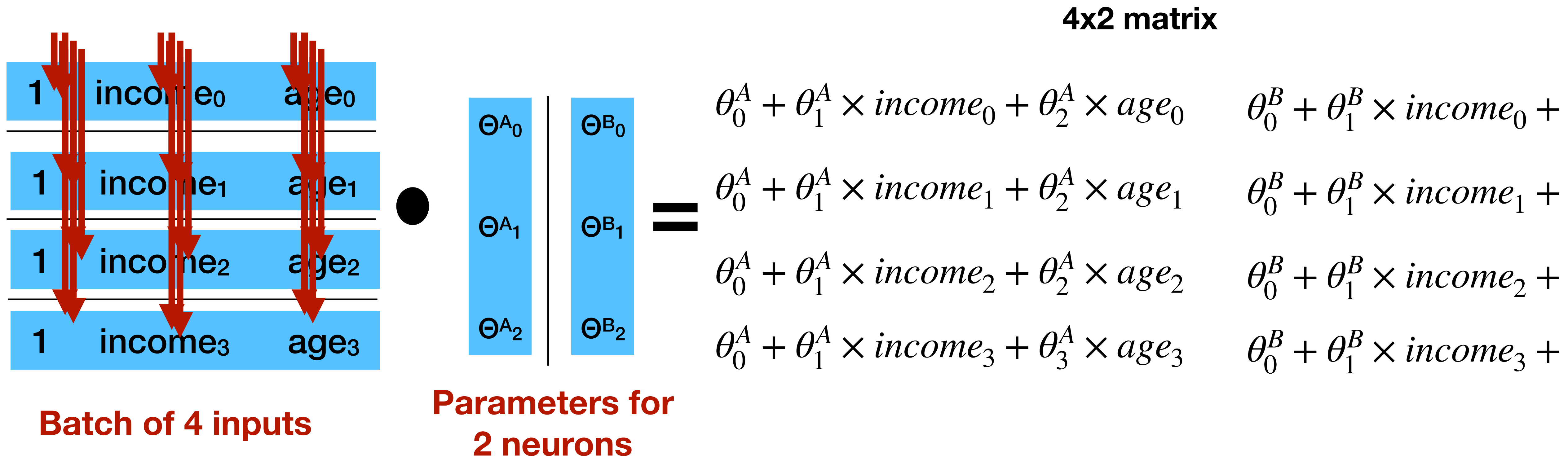


4x2 matrix

Matrix Multiplication



Matrix Multiplication



Matrix Multiplication

- We represent many inputs as a matrix
- We represent many neurons as a matrix
- Matrix multiplication compute the output score of all the neurons for all the inputs

4x3 matrix

1	income ₀	age ₀
1	income ₁	age ₁
1	income ₂	age ₂
1	income ₃	age ₃

Batch of 4 inputs

3x2 matrix

Θ^A_0	Θ^B_0
Θ^A_1	Θ^B_1
Θ^A_2	Θ^B_2

**Parameters for
2 neurons**

4x3 matrix

$\theta_0 + \theta^A_1 \times income_0 + \theta^A_2 \times age_0$	$\theta_0 + \theta^B_1 \times income_0 + \theta^B_2 \times age_0$
$\theta_0 + \theta^A_1 \times income_1 + \theta^A_2 \times age_1$	$\theta_0 + \theta^B_1 \times income_1 + \theta^B_2 \times age_1$
$\theta_0 + \theta^A_1 \times income_2 + \theta^A_2 \times age_2$	$\theta_0 + \theta^B_1 \times income_2 + \theta^B_2 \times age_2$
$\theta_0 + \theta^A_1 \times income_3 + \theta^A_2 \times age_3$	$\theta_0 + \theta^B_1 \times income_3 + \theta^B_2 \times age_3$

**Outputs of the
2 neurons for each of the 4 inputs**

Matrix Multiplication

2x3 matrix

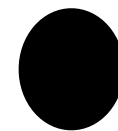
1	100	50
1	10	20

Batch of 2 inputs

3x2 matrix

1	0
2	1
1	-1

**Parameters for
2 neurons**

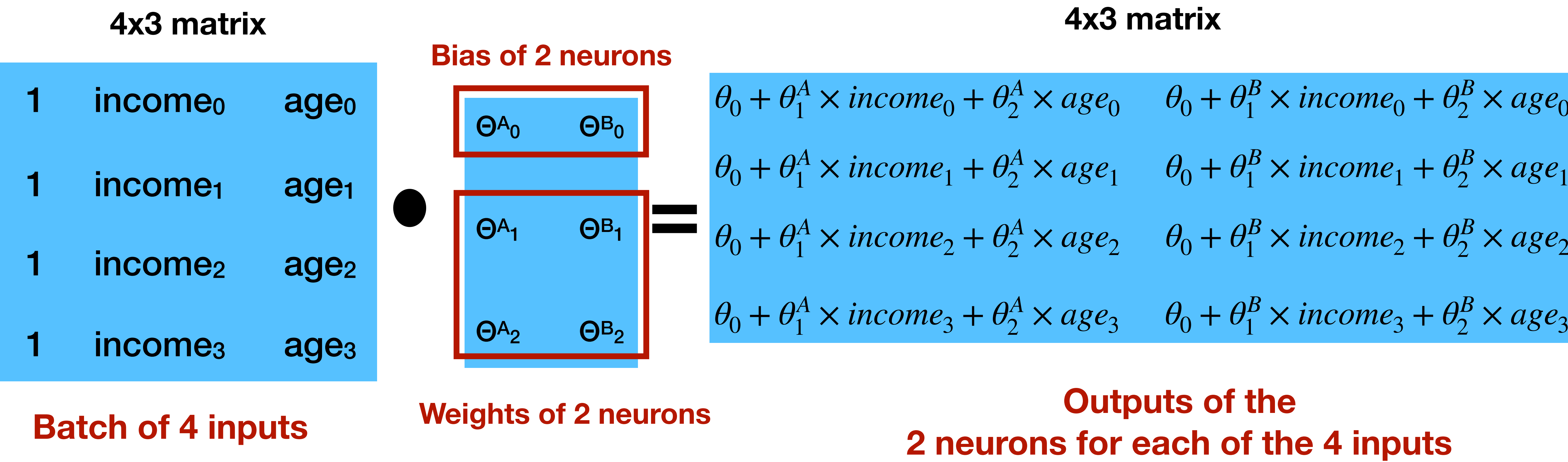


2x2 matrix

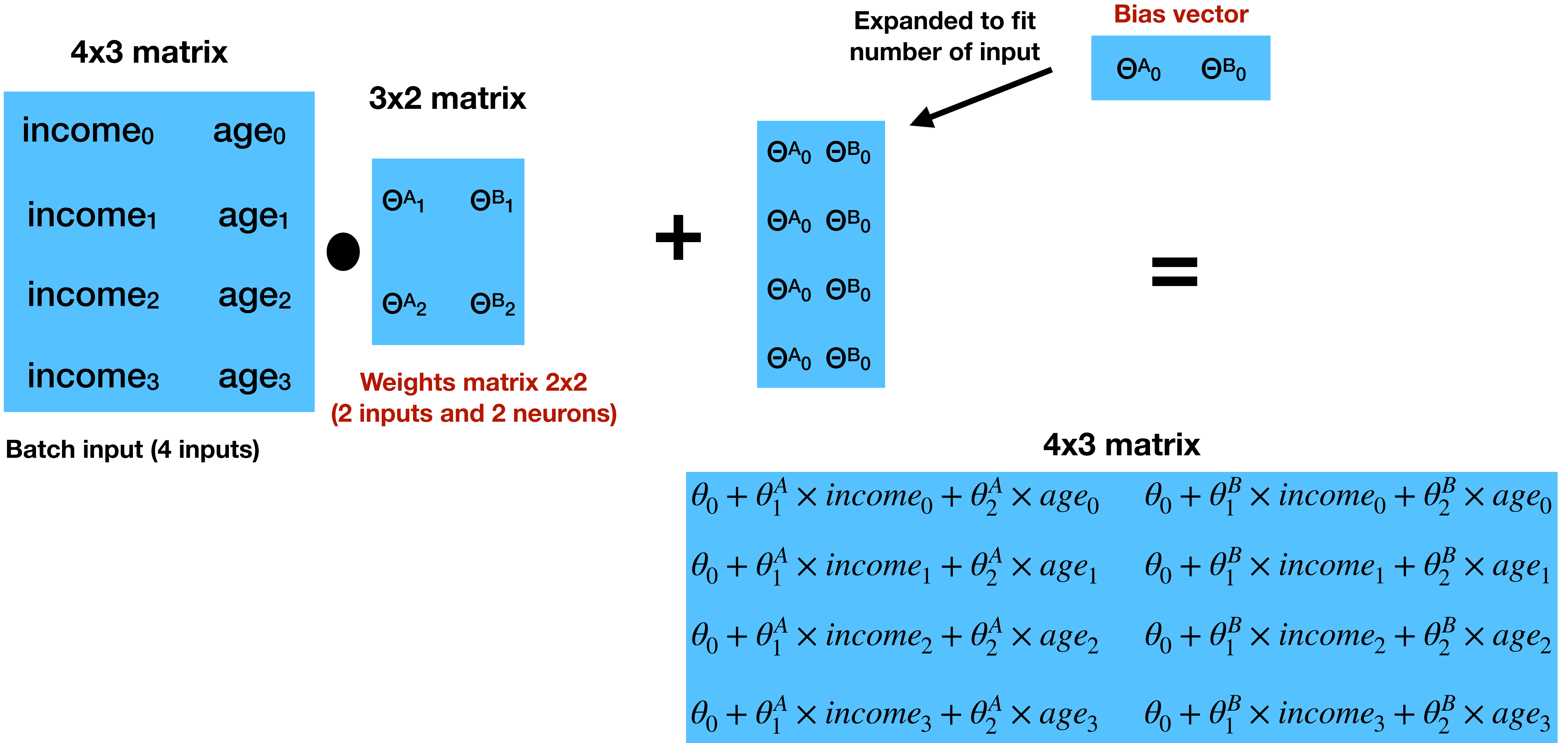
**Outputs of the
2 neurons for each of the 2 inputs**

Computing the output of many neurons for many inputs

- In practice, we separate the **weights** from the **bias**



Computing the output of many neurons for many inputs



Matrix Multiplication

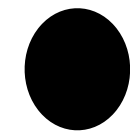
**Bias for
2 neurons**

1	0
---	---

2x2 matrix

2x2 matrix

100	50
10	20



2x2 matrix

2	1
1	-1

=

Batch of 2 inputs

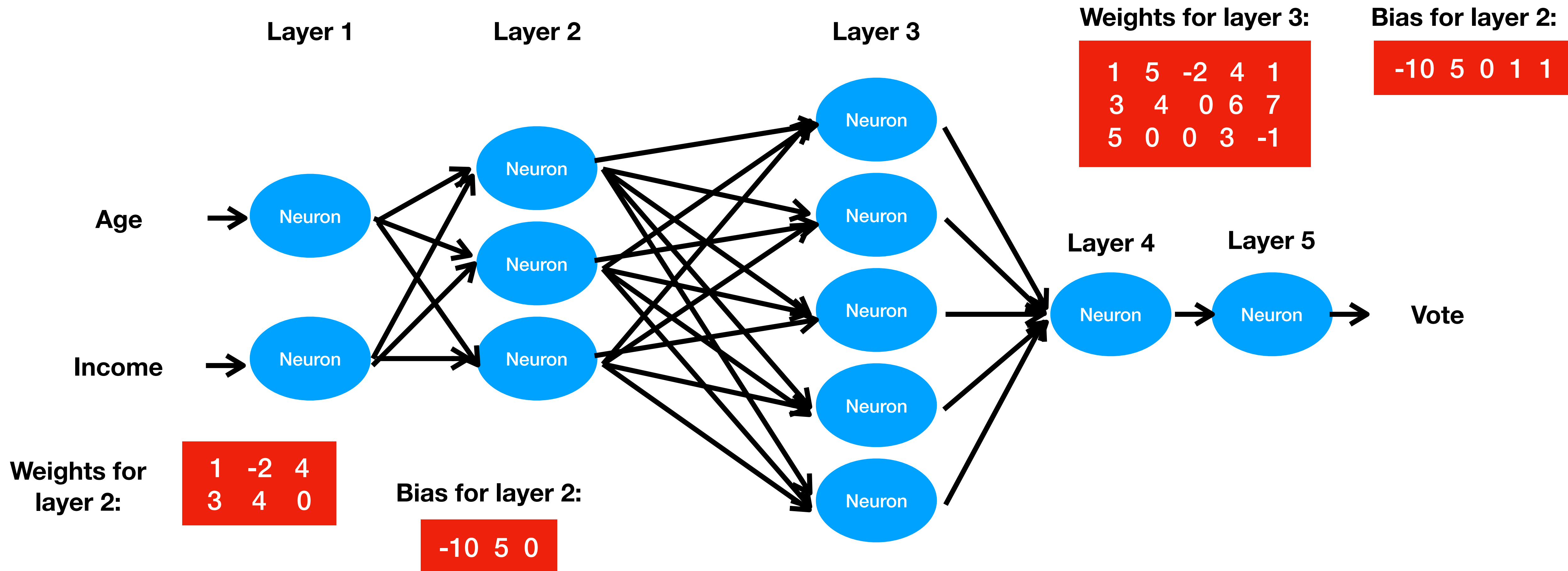
**Weights for
2 neurons**

**Score outputs of the
2 neurons for each of the 2 inputs**

- Important: using these matrix representations only work if all the neurons have the same input and are fully connected
- Then, in practice, a **Fully Connected layer of N neurons with K input** can be represented by a **matrix of weights W** of shape $K \times N$, and a **bias vector b** of dimension N .

Feed-Forward networks with fully connected layers

- Then, in practice, a **Fully Connected layer of N neurons with K input** can be represented by a ***matrix of weights* W** of shape $K \times N$, and a ***bias vector* b** of dimension N .



Random restarts

- Because the weights of a neural network are initialized randomly, the result of a training will change every time we reinitialize the network
- Therefore, to obtain best performances, it is common to train a networks several times with different random initializations, and keep the best result

Regularization methods for Neural Network

- When we train a network with many neurons, the danger of overfitting is large
- There are a few technics that are very efficient at preventing this:
 - Early Stopping
 - Dropout
 - Weight Decay
 - Stochastic Gradient Descent

Early stopping

- We keep a validation set separate from the training data
- We fix a ***patience*** number (typically patience = 10 or 20)
- During training, if we see no improvement on validation set after ***patience*** measures, we stop the training

Dropout

- During training, we add random noise to disturb the network
- In practice, we randomly “cut” a certain proportion of connections (typically 10% to 50%)

Weight Decay

- At every iteration in the training, the weights are scaled down by a factor $d < 1$
- Equivalent to a L2-Regularization

Stochastic Gradient Descent

- Instead of computing the gradient of the loss for all the training data, we compute it for a subsampled part of the training data
- This is actually done anyway to get faster training
- But it is also beneficial to prevent overfitting even if you could afford to compute the gradient for all the examples at once.