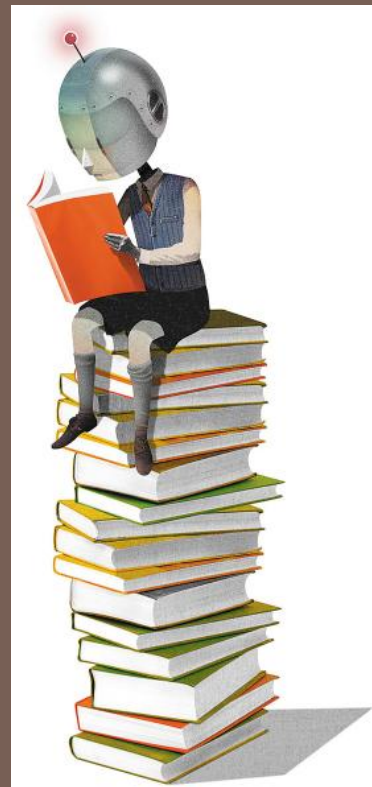


NEURAL MODELS FOR WORD EMBEDDING

Prof. Eduardo Bezerra
(CEFET/RJ)

ebezerra@cefet-rj.br



Summary

2

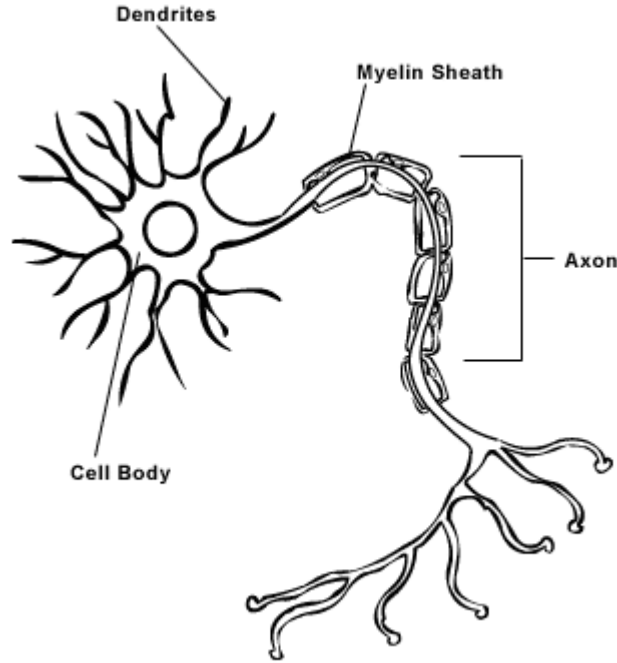
- Neural Nets (quick overview)
- Word2vec
 - ▣ Language models
 - ▣ Skip-gram

3

Neural Nets - Representation

Model of a biological neuron

4



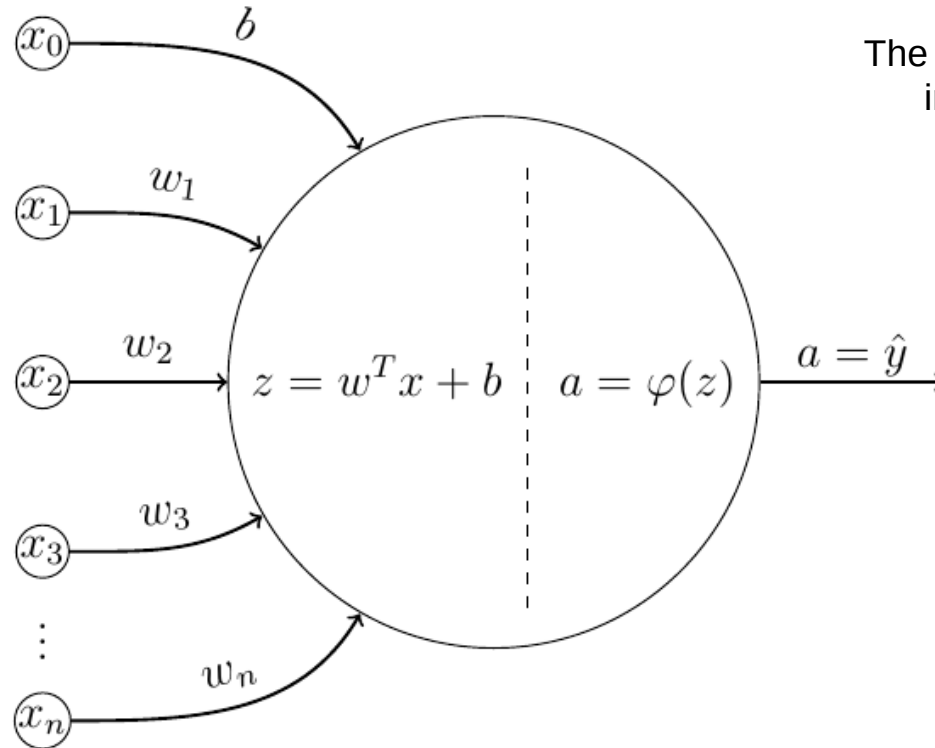
Artificial Neuron

5

- The artificial neuron is a model **inspired** by the real
- In an artificial neuron:
 - ▣ input are features
 - ▣ each feature has an associated **weight**
 - ▣ weighted sum of features is called neuron **pre-activation**
 - ▣ value produced in pre-activation goes through a non-linearity to produce neuron **activation**

Artificial Neuron

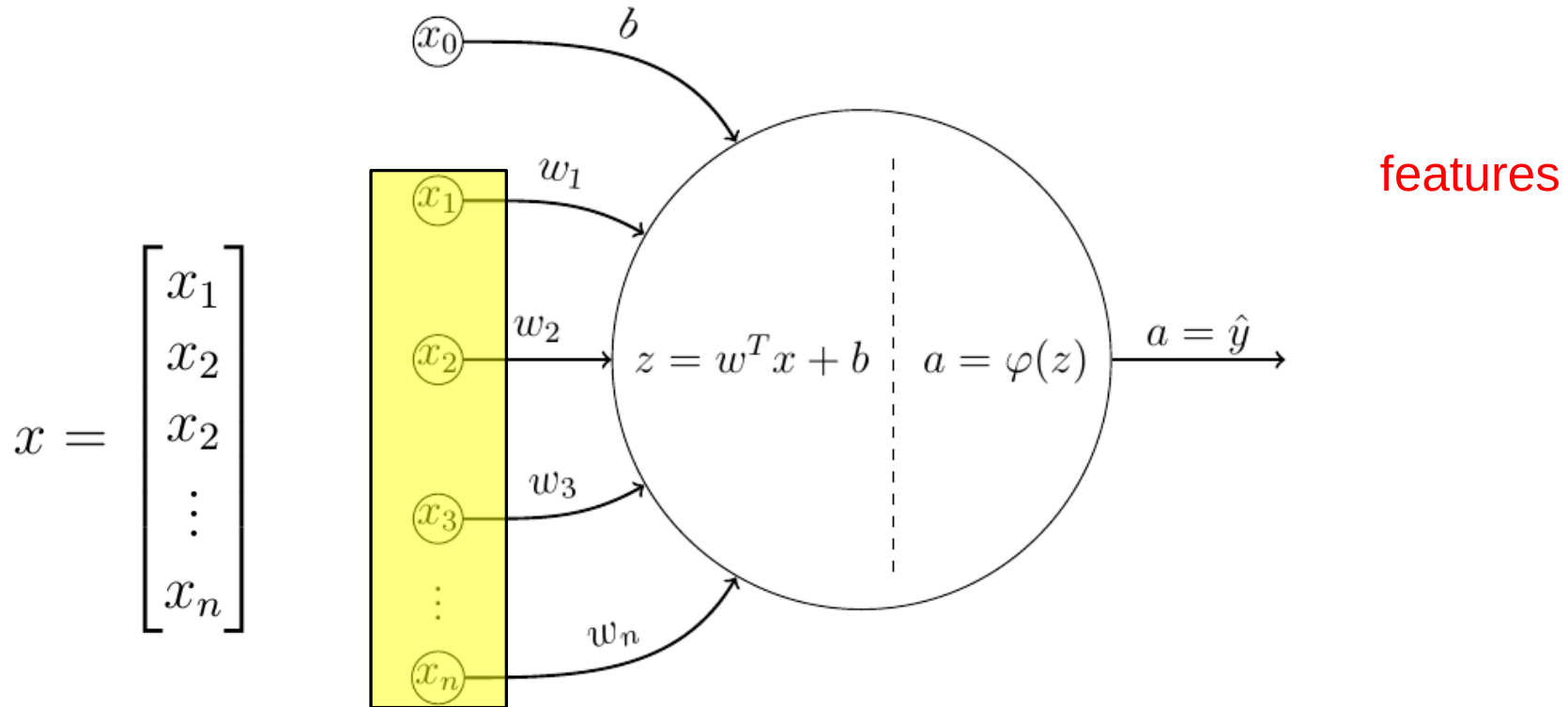
6



The artificial neuron is a model inspired by the real one (biological neuron).

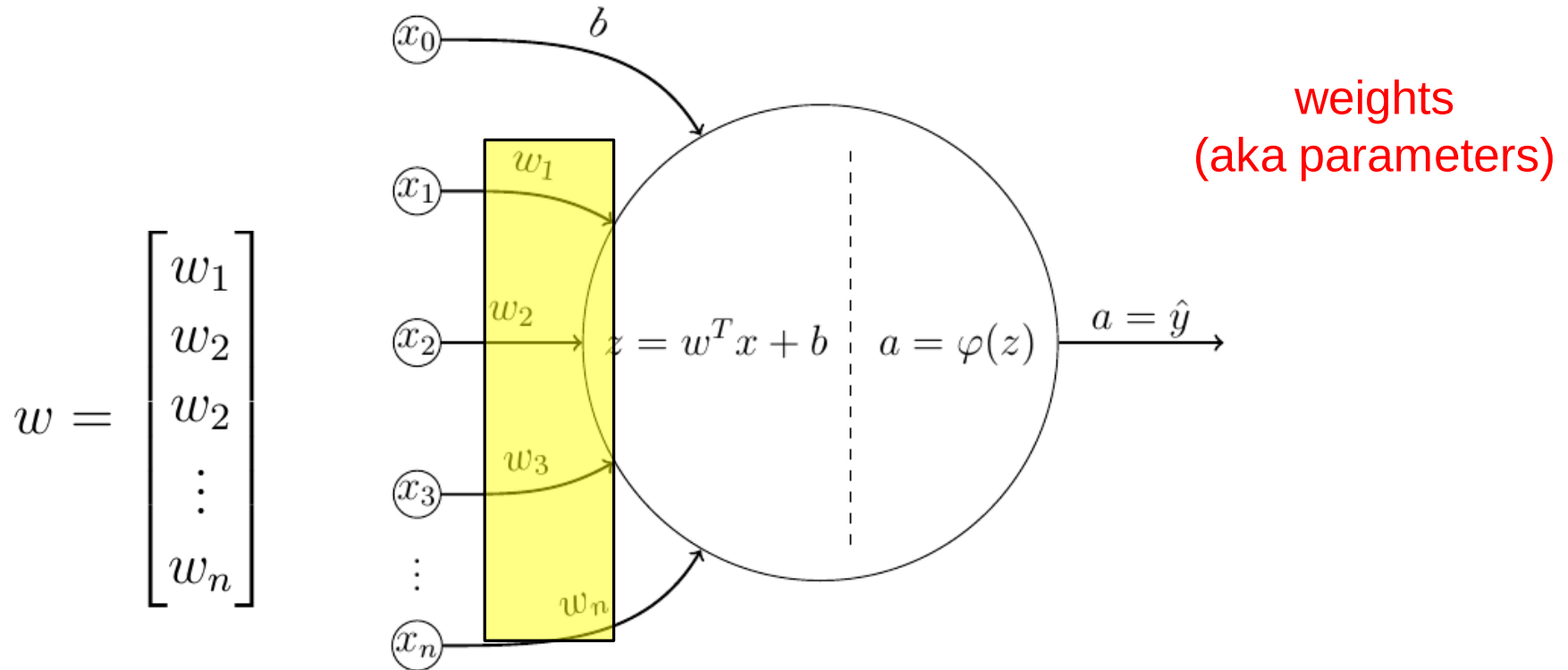
Artificial Neuron - input

7



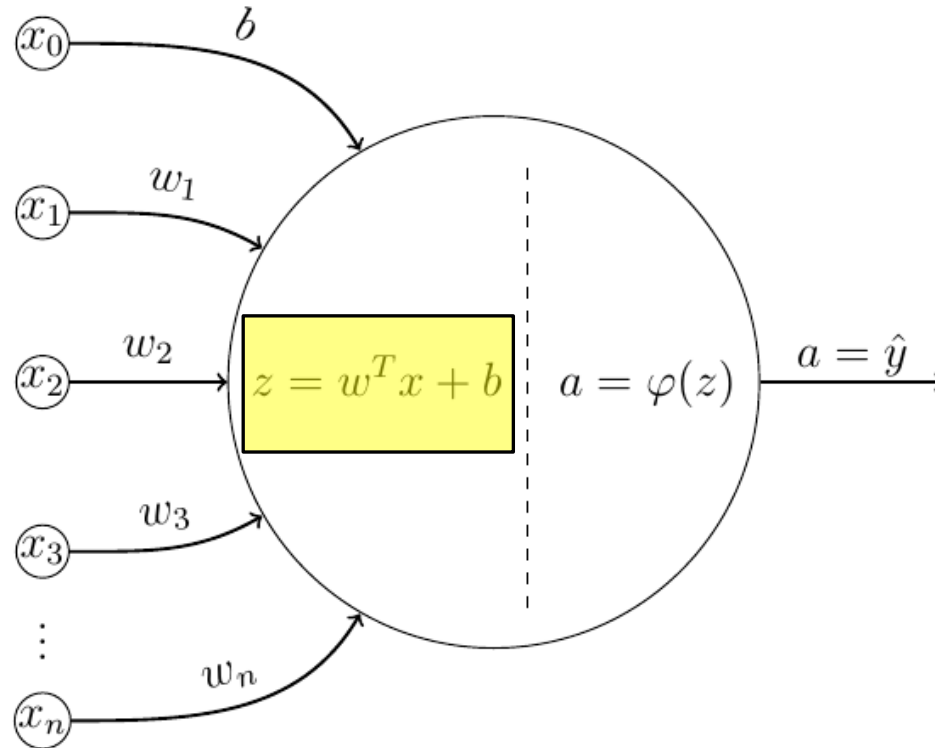
Artificial Neuron – weights

8



Artificial Neuron – pre-activation

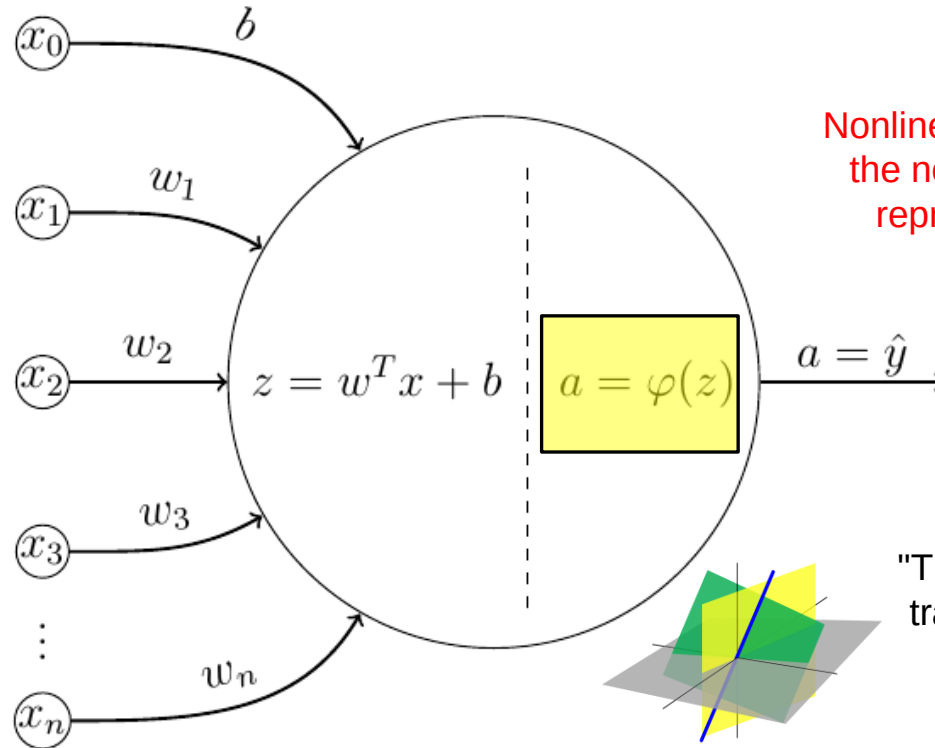
9



pre-activation

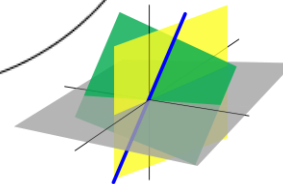
Artificial Neuron – activation function

10



activation

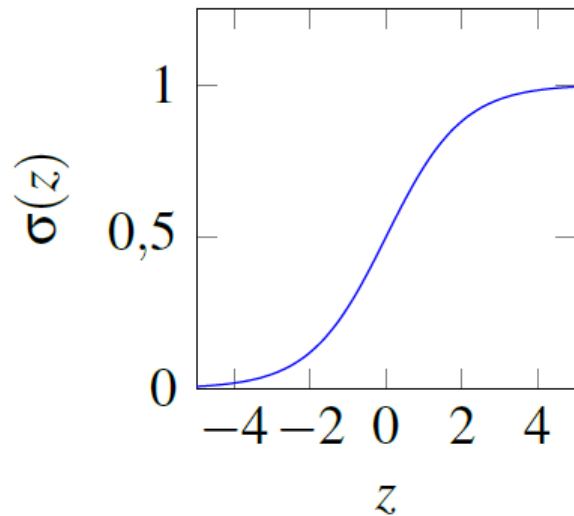
Nonlinearities are required so that the network can learn complex representations of the data.



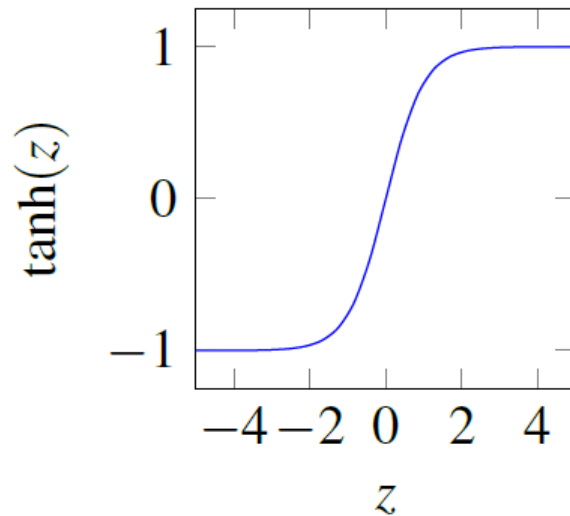
"The composition of linear transformations is also a linear transformation"

Artificial Neuron – activation functions

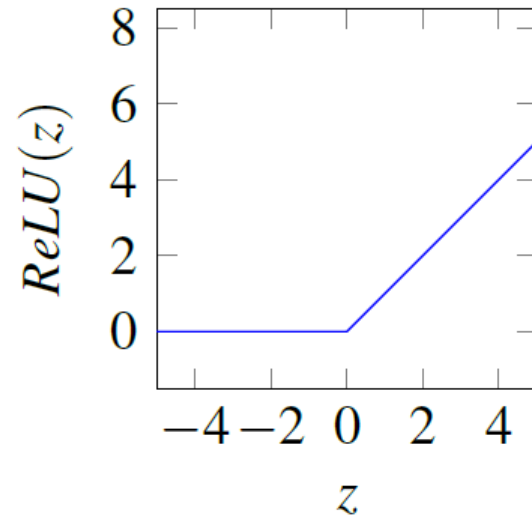
11



(a) Logística.



(b) Hiperbólica tangente.



(c) Retificada linear.

Artificial Neuron – activation functions

12

- Sigmoid Logistic
- Tangent Hyperbolic
- ReLU
- softmax

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

$$\varphi(z) = \tanh(z) = \frac{e^z - e^{-z}}{e^z + e^{-z}}$$

$$\varphi(z) = \text{ReLU}(z) = \max(0, z)$$

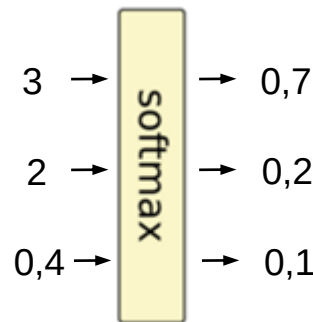
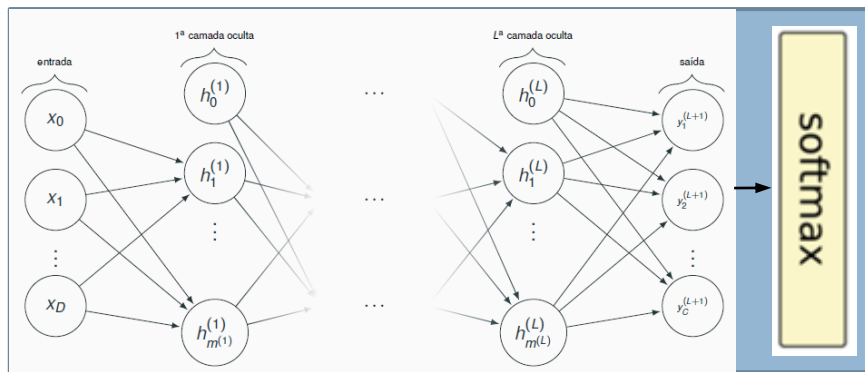
$$y_i = o\left(a_i^{(L+1)}\right) = \frac{e^{a_i^{(L+1)}}}{\sum_{j=1}^C e^{a_j^{(L+1)}}}$$

softmax

$$y_i = o\left(a_i^{(L+1)}\right) = \frac{e^{a_i^{(L+1)}}}{\sum_{j=1}^C e^{a_j^{(L+1)}}}$$

13

- Maps the NN outputs to a **probability distribution**.

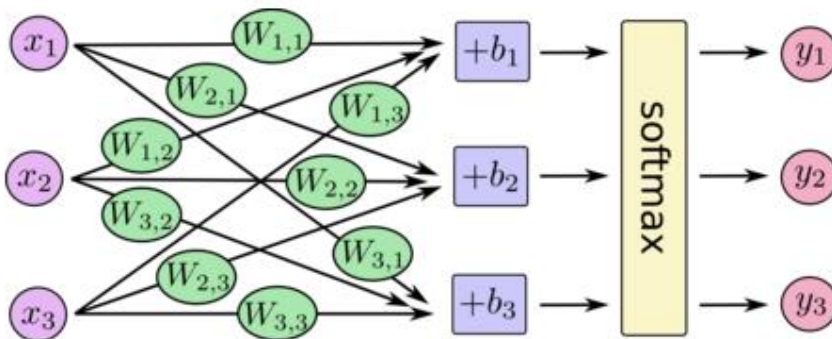


softmax

$$y_i = o \left(a_i^{(L+1)} \right) = \frac{e^{a_i^{(L+1)}}}{\sum_{j=1}^C e^{a_j^{(L+1)}}}$$

14

- Maps the NN outputs to a **probability distribution**.



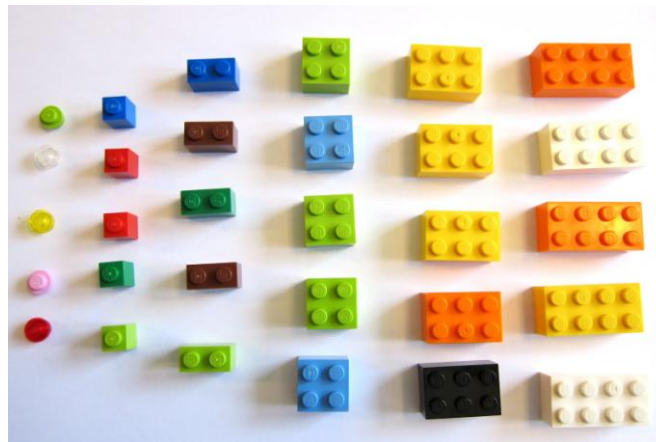
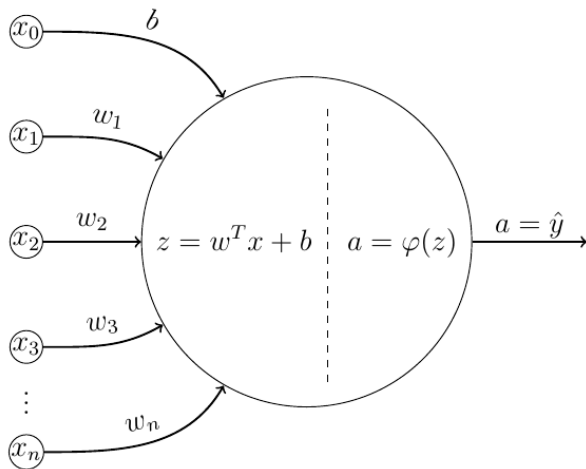
$$\begin{bmatrix} y_1 \\ y_2 \\ y_3 \end{bmatrix} = \text{softmax} \begin{bmatrix} W_{1,1}x_1 + W_{1,2}x_2 + W_{1,3}x_3 + b_1 \\ W_{2,1}x_1 + W_{2,2}x_2 + W_{2,3}x_3 + b_2 \\ W_{3,1}x_1 + W_{3,2}x_2 + W_{3,3}x_3 + b_3 \end{bmatrix}$$

$$\begin{bmatrix} y_1 \\ y_2 \\ y_3 \end{bmatrix} = \text{softmax} \left(\begin{bmatrix} W_{1,1} & W_{1,2} & W_{1,3} \\ W_{2,1} & W_{2,2} & W_{2,3} \\ W_{3,1} & W_{3,2} & W_{3,3} \end{bmatrix} \cdot \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} + \begin{bmatrix} b_1 \\ b_2 \\ b_3 \end{bmatrix} \right)$$

Artificial Neural Net

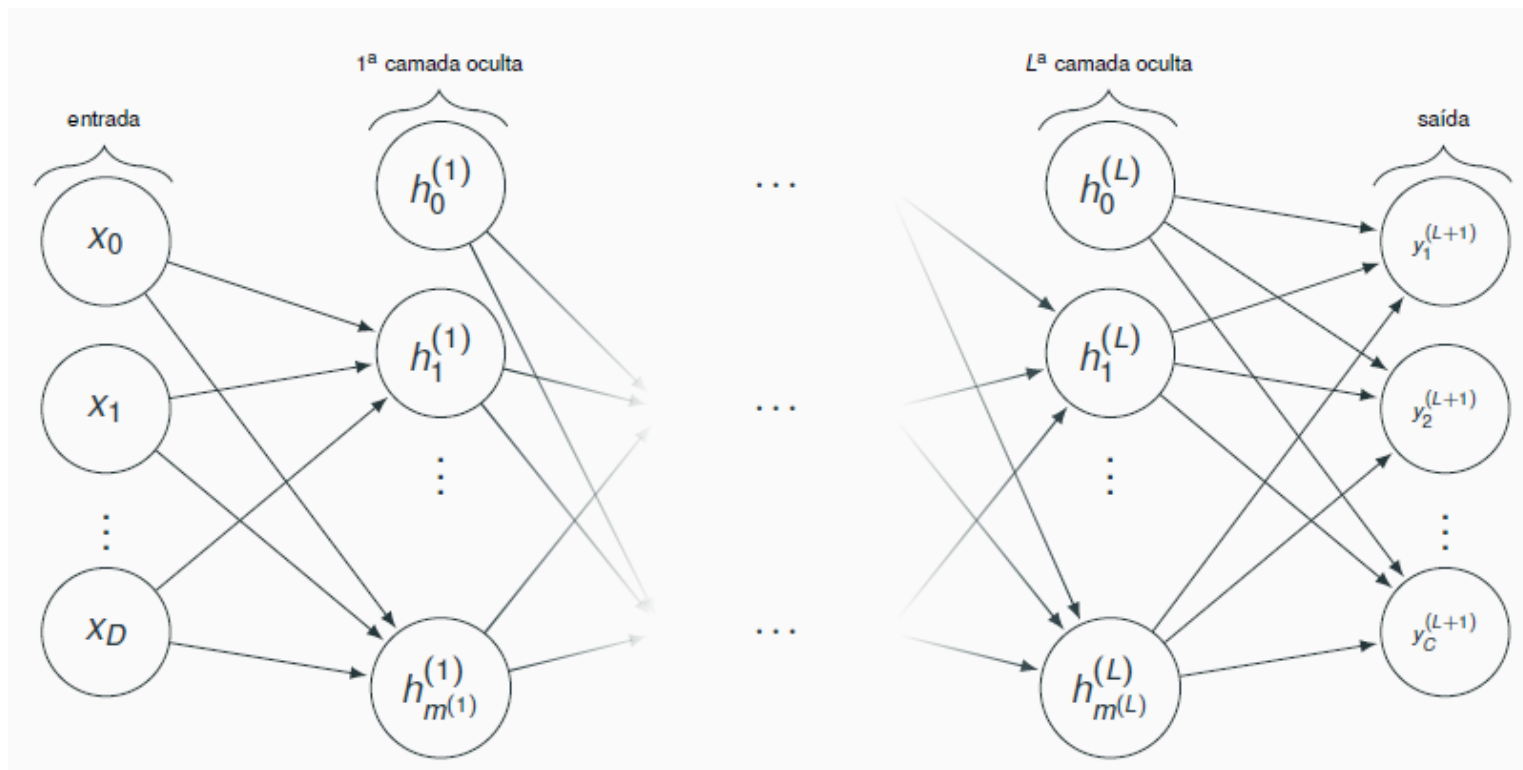
15

- It is possible to architect arbitrarily complex networks using the artificial neuron as the basic component.



Artificial Neural Net

16



Feedforward Neural Network

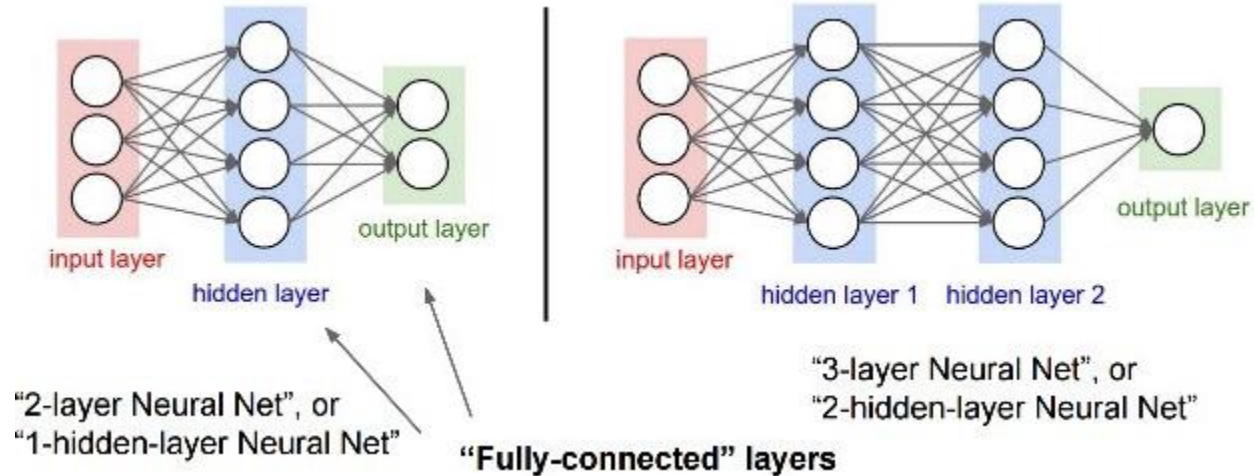
Kinds of layers

17

- In an ANN:
 - ▣ **input layer** are the input patterns (features);
 - ▣ **output layer** is result of the computation performed;
 - ▣ other layers are called **hidden layers**.
- Amount of layers and amount of units (neurons) per layer are part of the ANN architecture.

Kinds of layers - example

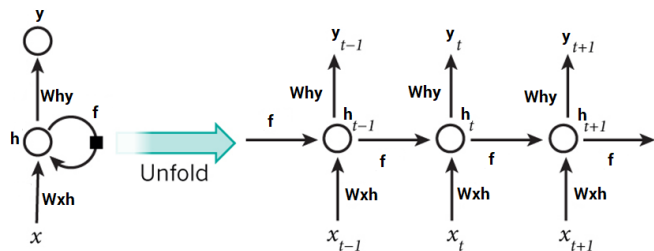
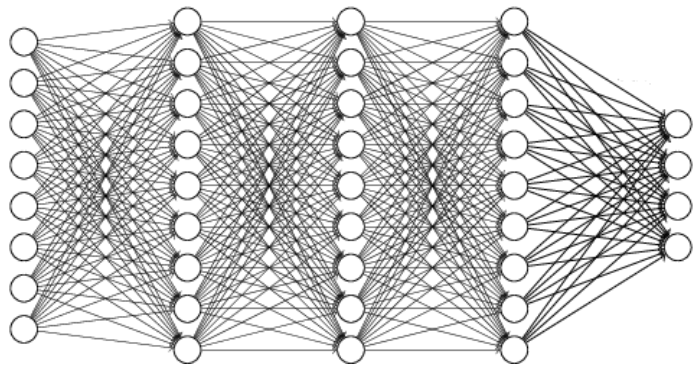
18



Kinds of ANNs

19

- Fully connected x not fully connected
- Recurrent x feed forward



ANNs - notation

20

□ Pre-activation:

$z_i^{[l]}$ $\leftarrow$ identifica a camada
 $\leftarrow$ identifica o neurônio dentro da camada l

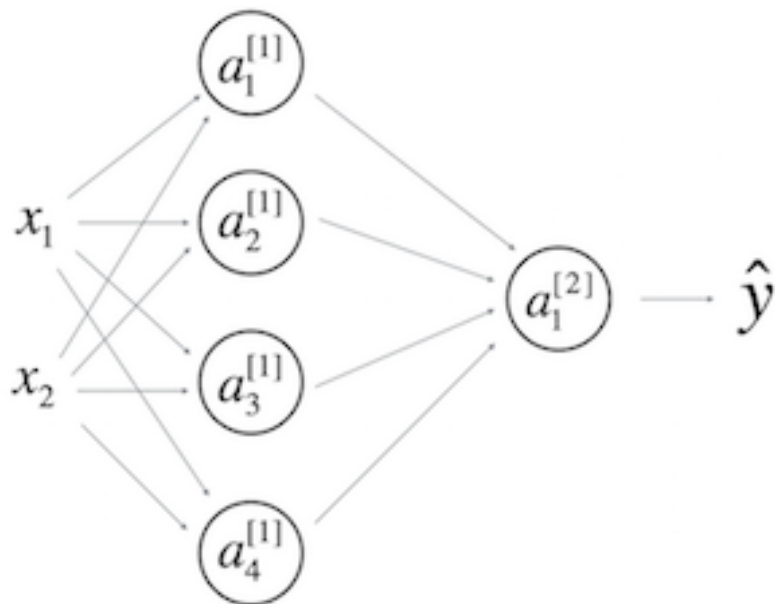
□ Activation:

$a_i^{[l]}$ $\leftarrow$ identifica a camada
 $\leftarrow$ identifica o neurônio dentro da camada l

ANNs – notation - example

21

□ e.g., in a two-layer ANN:



Computations in the first layer:

$$z_1^{[1]} = (w_1^{[1]})^T x + b_1^{[1]}, a_1^{(1)} = \varphi(z_1^{[1]})$$

$$z_2^{[1]} = (w_2^{[1]})^T x + b_2^{[1]}, a_2^{(1)} = \varphi(z_2^{[1]})$$

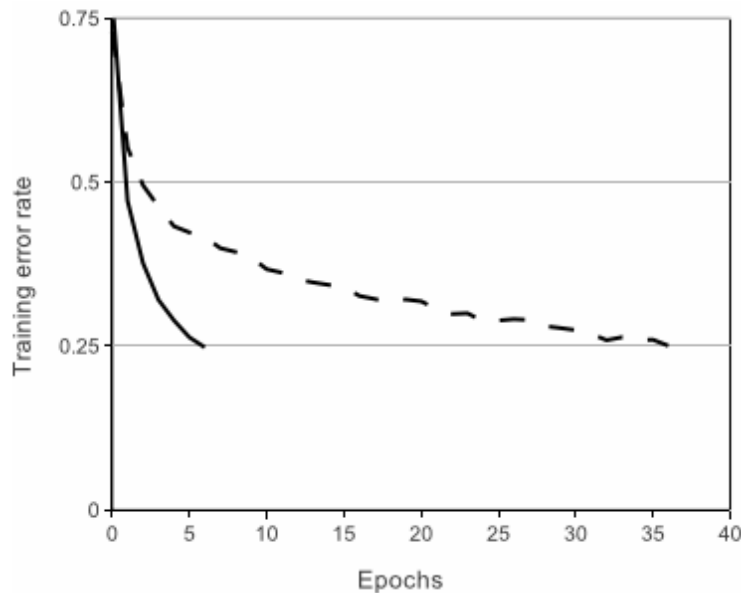
$$z_3^{[1]} = (w_3^{[1]})^T x + b_3^{[1]}, a_3^{(1)} = \varphi(z_3^{[1]})$$

$$z_4^{[1]} = (w_4^{[1]})^T x + b_4^{[1]}, a_4^{(1)} = \varphi(z_4^{[1]})$$

Rectified linear units, ReLU

22

- Convnet of 4 layers with ReLUs.
 - ▣ CIFAR-10 (60K, 32x32 imgs)
 - ▣ It reaches the same error level 6 times faster than tanh.



Neural Nets - Evaluation

Cost function

24

- The **error signal** is the difference between $y^{(i)}$ and output produced for $x^{(i)}$.
 - measures the difference between the network predictions and the desired values.
- A **cost function** must provide a way to measure the error signal for a set of training examples.

Cost function $J(W,b)$

25

- Cost function: error of the network considering all the examples.
 - ▣ function of the set of all weights and bias of the network.

$$J(W,b) = \underbrace{\frac{1}{m} \sum_{i=1}^m J(f(\mathbf{x}^{(i)}; W,b), \mathbf{y}^{(i)})}_{\text{custo dos dados}} + \underbrace{\frac{\lambda}{2m} \sum_k \sum_l (w_k^{[l]})^2}_{\text{custo de regularização}}$$

Some (unregularized) cost functions

26

$$\mathcal{D} = \{(\mathbf{x}^{(i)}, \mathbf{y}^{(i)}) : 1 \leq i \leq m\}$$

$$J_{\text{MSE}} = \frac{1}{m} \sum_{i=1}^m \left[f(\mathbf{x}^{(i)}) - \mathbf{y}^{(i)} \right]^2$$

Mean square error

$$J_{\text{cross entropy}} = \sum_{i=1}^m \left[p(\mathbf{x}^{(i)}) \log q(\mathbf{x}^{(i)}) \right]$$

Cross entropy error

$$J_{\text{KL}} = \sum_{i=1}^m \left[p(\mathbf{x}^{(i)}) \log \frac{p(\mathbf{x}^{(i)})}{q(\mathbf{x}^{(i)})} \right]$$

Kullback-Leibler divergence

Neural Nets - Optimization (learning, training)

Training

28

- Given a training set of the form

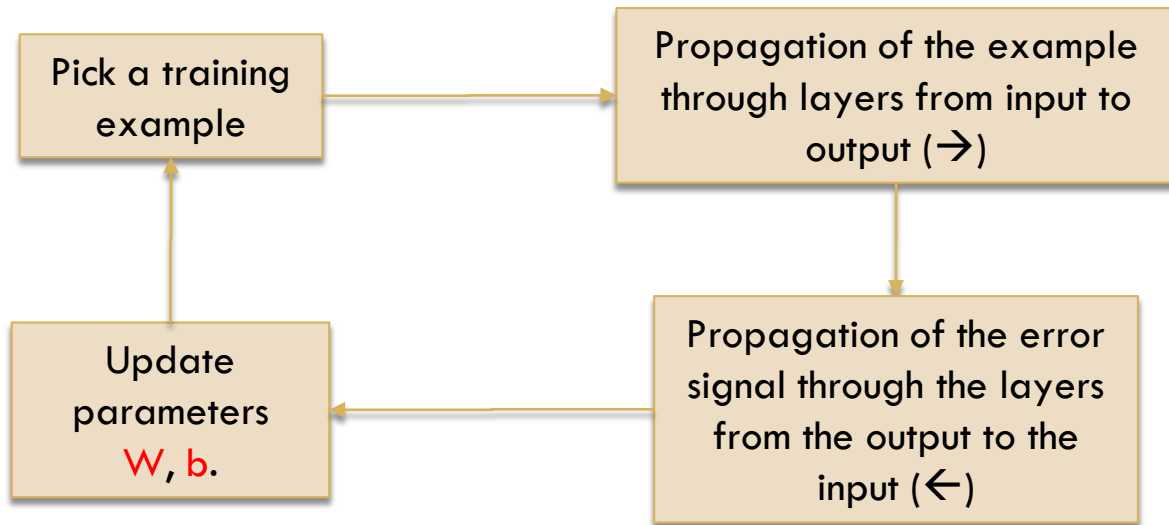
$$\{\mathbf{x}^{(i)}, \mathbf{y}^{(i)} : 1 \leq i \leq m\}$$

- training an ANN corresponds to using this set to adjust the parameters of the network, so that the training error is minimized.
- So, training of an ANN is an optimization problem.

Training

29

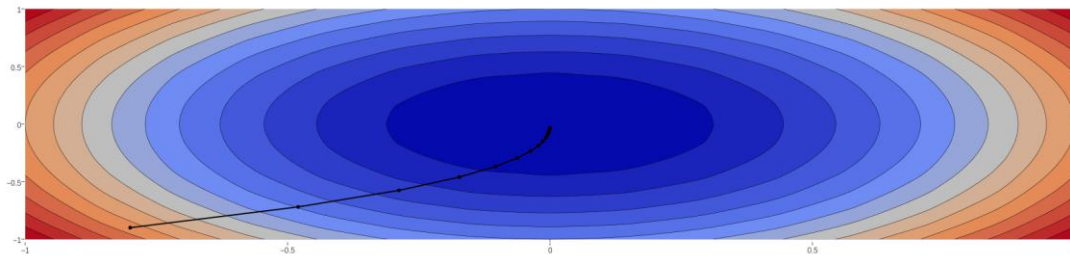
- The **error signal** (computed with a cost function) is used during training to gradually change the weights (parameters), so that the predictions are more accurate.



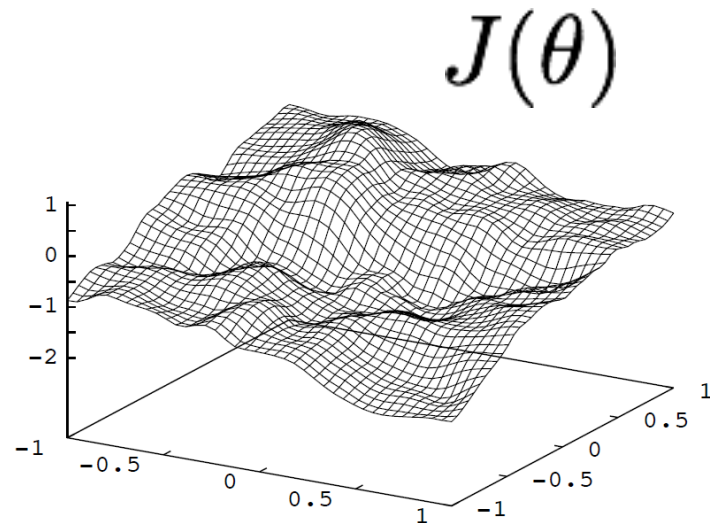
Stochastic Gradient Descent (SGD)

30

- SGD allows to traverse the space of parameters of a neural network in search for a local minimum.



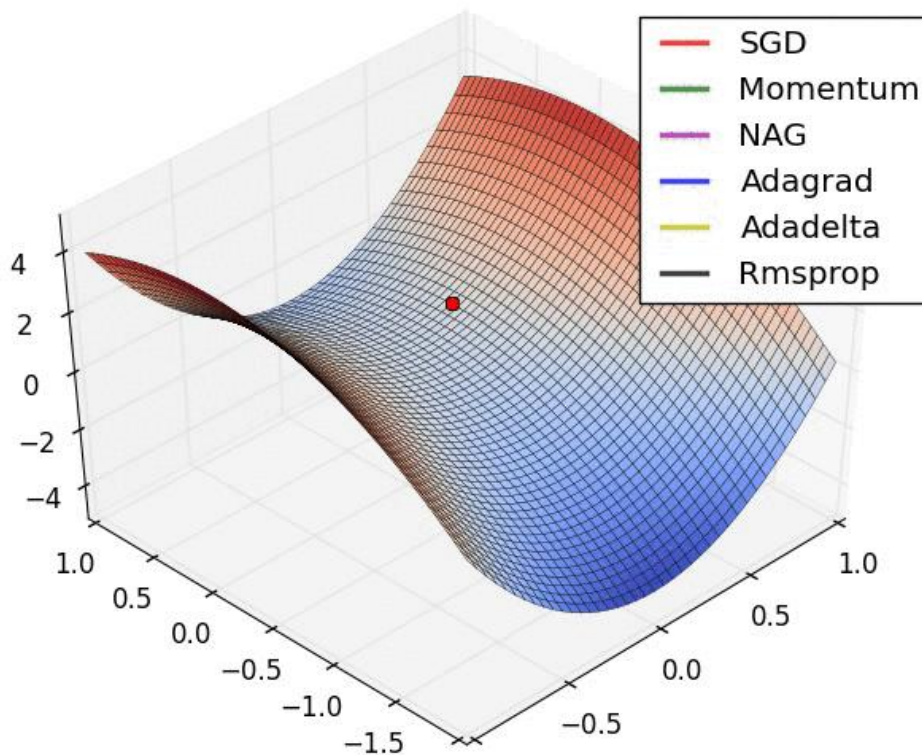
$$\Delta w[t] = -\alpha \frac{\partial}{\partial w[t]} J = -\alpha \nabla J(w[t])$$



Source: David McKay, ITPRNN

Alternatives to SGD

31



Neural Nets - Backpropagation

Error back-propagation

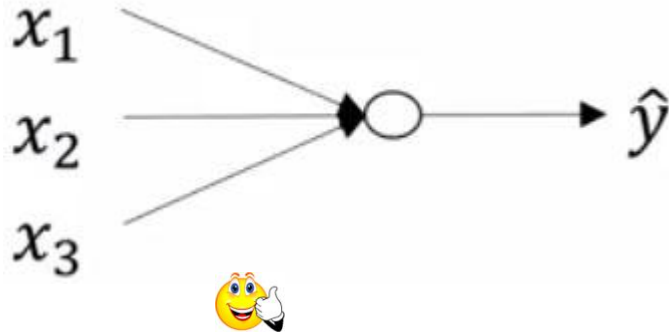
33

- Consider that the output produced for $x^{(i)}$ is different from $y^{(i)}$ (i.e., there is a nonzero **error signal**).
- Credit Assignment Problem
 - ▣ Then it is necessary to determine the responsibility of each parameter (weight) of the network by this error...
 - ▣ ... and change these parameters for the purpose of reducing the error.

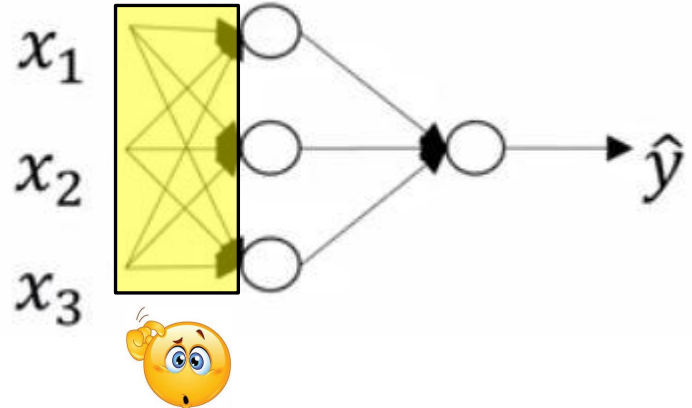
Error back-propagation

34

- In an ANN with at least one hidden layer, there is no direct supervision signal for these layers.



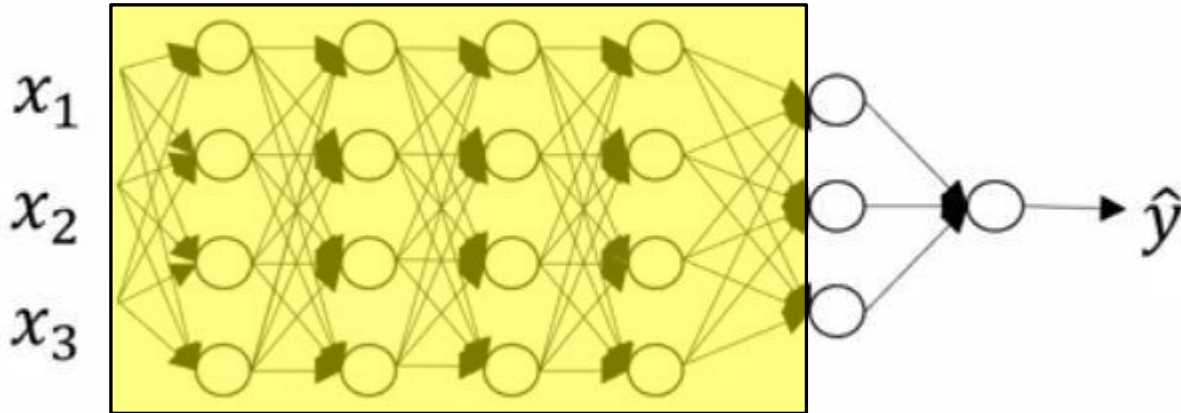
Logistic regression



Error back-propagation

35

- This problem gets worse as the amount of hidden layers increases.



Backpropagation

36

- Algorithm invented many times...
- Recursively propagates the error signal from the output layer through the hidden layers, to the input layer.
- Used in conjunction with some optimization algorithm (e.g., SGD) to gradually minimize network error.

Backpropagation

37

Backpropagation, an abbreviation for “backward propagation of errors”, is a common method of training artificial neural networks used in conjunction with an optimization method such as gradient descent. The method calculates the gradient of a loss function with respect to all the weights in the network. The gradient is fed to the optimization method which in turn uses it to update the weights, in an attempt to minimize the loss function.

--Frederik Kratzert

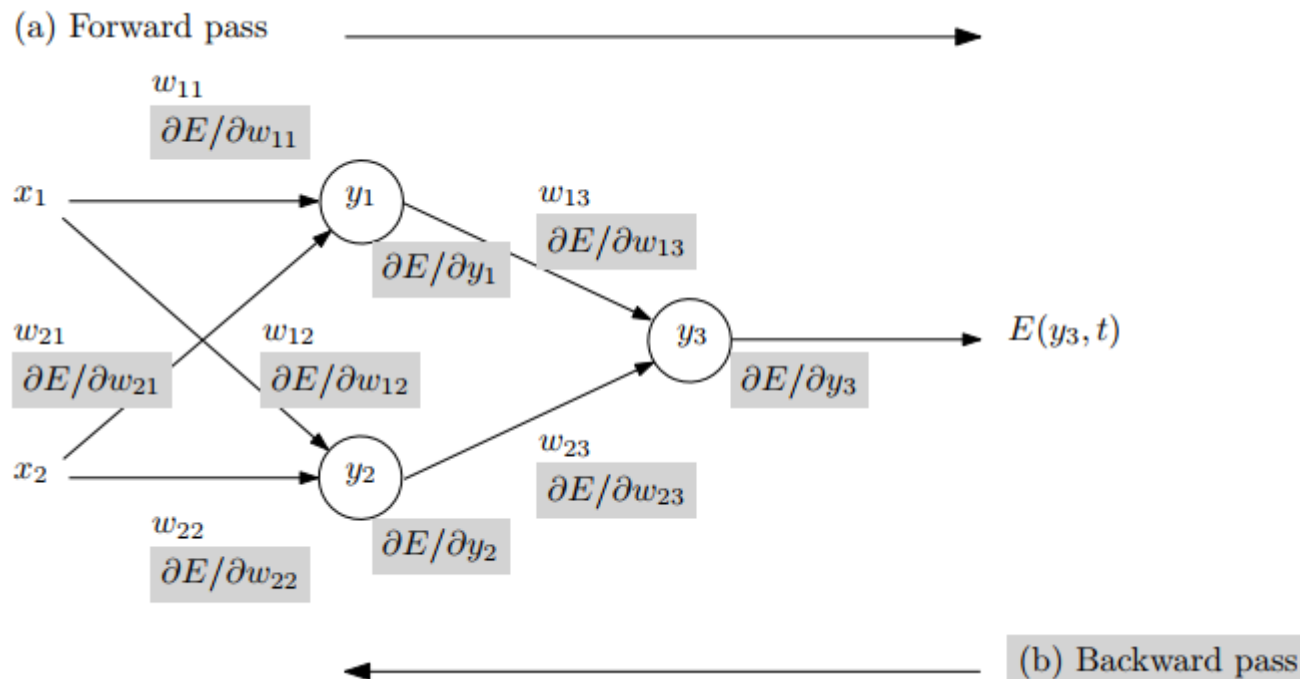
ANN training with Backpropagation

38

- Set initial weights for the network
- Present an example and get the result
- For each layer (from the last):
 - ▣ Calculate the error for each neuron
 - ▣ Update the weights (with SGD or variant)
 - ▣ Propagate this error to the previous layer
- Repeat with other examples until convergence

Backpropagation – example

39



propagation

The diagram illustrates the backpropagation of gradients through an activation function f . It shows three gray circular nodes. The central node is labeled f and contains two red boxes with the partial derivatives $\frac{\partial z}{\partial x}$ and $\frac{\partial z}{\partial y}$. To its left, a green arrow labeled x points from a node to the central node, and a red arrow points back. Below this, another green arrow labeled y points from a node to the central node, with a red arrow pointing back. To the right of the central node, a green arrow labeled z points to another node, with a red arrow pointing back that carries the gradient $\frac{\partial J}{\partial z}$. The chain rule equations are shown next to the red arrows: $\frac{\partial J}{\partial x} = \frac{\partial J}{\partial z} \times \frac{\partial z}{\partial x}$ for the top path and $\frac{\partial J}{\partial y} = \frac{\partial J}{\partial z} \times \frac{\partial z}{\partial y}$ for the bottom path.

$$\frac{\partial J}{\partial x} = \frac{\partial J}{\partial z} \times \frac{\partial z}{\partial x}$$

$$\frac{\partial J}{\partial y} = \frac{\partial J}{\partial z} \times \frac{\partial z}{\partial y}$$

$$\frac{\partial z}{\partial x}$$

$$\frac{\partial z}{\partial y}$$

f

$$z$$

$$\frac{\partial J}{\partial z}$$

41

word2vec

Credit where credit is due

42

- The original papers from Mikolov et al.
- CS 224D: Deep Learning for NLP
 - ▣ https://cs224d.stanford.edu/lecture_notes/notes1.pdf
- McCormick, C. (2016). Word2Vec Tutorial.
 - ▣ <http://www.mccormickml.com>
- Understanding word vectors (Python notebook)
 - ▣ <https://gist.github.com/aparrish/2f562e3737544cf29aaf1af30362f469>

Language Models (Unigrams, Bigrams, etc.)

43

- A model that assigns a probability to a sequence of tokens.
- A good language model gives...
 - ▣ ...(syntactically and semantically) valid sentences a high probability.
 - ▣ ...low probability to nonsense.

Language Models (Unigrams, Bigrams, etc.)

44

- Mathematically, we can apply a LM to any given sequence of n words:

$$P(w_1, w_2, \dots, w_n)$$

Language Models (Unigrams, Bigrams, etc.)

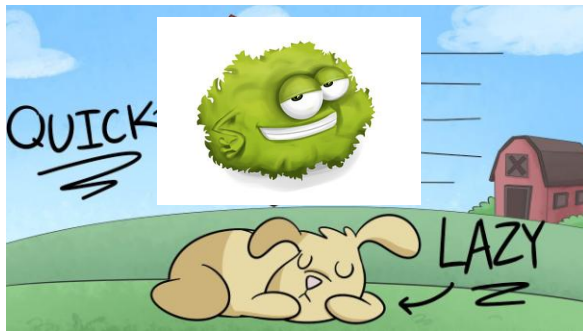
45

- An example:

"The quick brown fox jumps over the lazy dog."

- Another example:

"The quik brown lettuce over jumps the lazy dog."



Language Models (Unigrams, Bigrams, etc.)

46

$$P(w_1, w_2, \dots, w_n) = \prod_{i=1}^n P(w_i) \quad \text{Unigram model}$$

$$P(w_1, w_2, \dots, w_n) = \prod_{i=2}^n P(w_i | w_{i-1}) \quad \text{Bigram model}$$

But, how to **learn** these probabilities?

word2vec

47



<https://research.fb.com/people/mikolov-tomas/>

- Efficient Estimation of Word Representations in Vector Space, September 7th, 2013.
- Distributed Representations of Words and Phrases and their Compositionality, October 16th, 2013.

<https://code.google.com/archive/p/word2vec/>

The word2vec trick

48

- word2vec uses a trick:
 - ▣ A single hidden layer neural network is trained to perform a certain “fake” task.
 - ▣ But this NN is not actually used!
- Instead, the goal is to learn the **weights** of the hidden layer— ***these weights are the “word vectors”***.

The fake tasks

49

- **Skip-gram**: predicting surrounding context words given a center word.
- **CBOW**: predicting a center word from the surrounding context.

50

Skip-gram

Skip-gram

51

- The task: given a specific word w in the middle of a sentence (the input word), look at the words nearby and pick one word at random.
- We then train the network to produce the probability (for every word in the vocabulary) of being nearby w .
 - “nearby” means there is actually a “window size” hyperparameter (typical value: 5)

Skip-gram

52

- Output probabilities are going to relate to how likely it is to find each vocabulary word nearby our input word $\mathbf{w}$.
- For example, say $\mathbf{w} = \text{"Africa"}$
 - ▣ output probabilities should be much higher for related words (e.g. "lion", "zebra") than for unrelated words (e.g. "bear", "kangaroo").

Training examples (word pairs)

53

- The training examples will be **word pairs** taken from the input corpus.
- During training, the network will learn parameters that capture the statistics related to the number of times each pairing shows up.

Training examples (word pairs)

54

- For example, suppose the network gets many more training samples of (“Soviet”, “Union”) than (“Soviet”, “Sasquatch”).
- ▣ After training, if the NN is given the word “Soviet”, it will output a much higher probability for “Union” (or “Russia”) than it will for “Sasquatch”.

Training examples (word pairs)

55

Source Text	Training Samples						
<table><tr><td>The</td><td>quick</td><td>brown</td></tr></table> fox jumps over the lazy dog. ➡	The	quick	brown	(the, quick) (the, brown)			
The	quick	brown					
<table><tr><td>The</td><td>quick</td><td>brown</td><td>fox</td></tr></table> jumps over the lazy dog. ➡	The	quick	brown	fox	(quick, the) (quick, brown) (quick, fox)		
The	quick	brown	fox				
<table><tr><td>The</td><td>quick</td><td>brown</td><td>fox</td><td>jumps</td></tr></table> over the lazy dog. ➡	The	quick	brown	fox	jumps	(brown, the) (brown, quick) (brown, fox) (brown, jumps)	
The	quick	brown	fox	jumps			
<table><tr><td>The</td><td>quick</td><td>brown</td><td>fox</td><td>jumps</td><td>over</td></tr></table> the lazy dog. ➡	The	quick	brown	fox	jumps	over	(fox, quick) (fox, brown) (fox, jumps) (fox, over)
The	quick	brown	fox	jumps	over		

"The quick brown fox jumps over the lazy dog."
window size = 2.

The word highlighted in blue is the input word w .

One hot encoding

56

- But, during training, we cannot provide the NN with the text of the words.
- We need to use some vector representation for each word in the vocabulary → so, we have to use **one hot encoding**.

One hot encoding

57

Rome = [1, 0, 0, 0, 0, 0, ..., 0]

Paris = [0, 1, 0, 0, 0, 0, ..., 0]

Italy = [0, 0, 1, 0, 0, 0, ..., 0]

France = [0, 0, 0, 1, 0, 0, ..., 0]

One hot encoding

58

Rome = $[1, 0, 0, 0, 0, 0, \dots, 0]$

Paris = $[0, 1, 0, 0, 0, 0, \dots, 0]$

Italy = $[0, 0, 1, 0, 0, 0, \dots, 0]$

France = $[0, 0, 0, 1, 0, 0, \dots, 0]$

One hot encoding

59

V = vocabulary size (huge)

Rome = [1, 0, 0, 0, 0, 0, ..., 0]

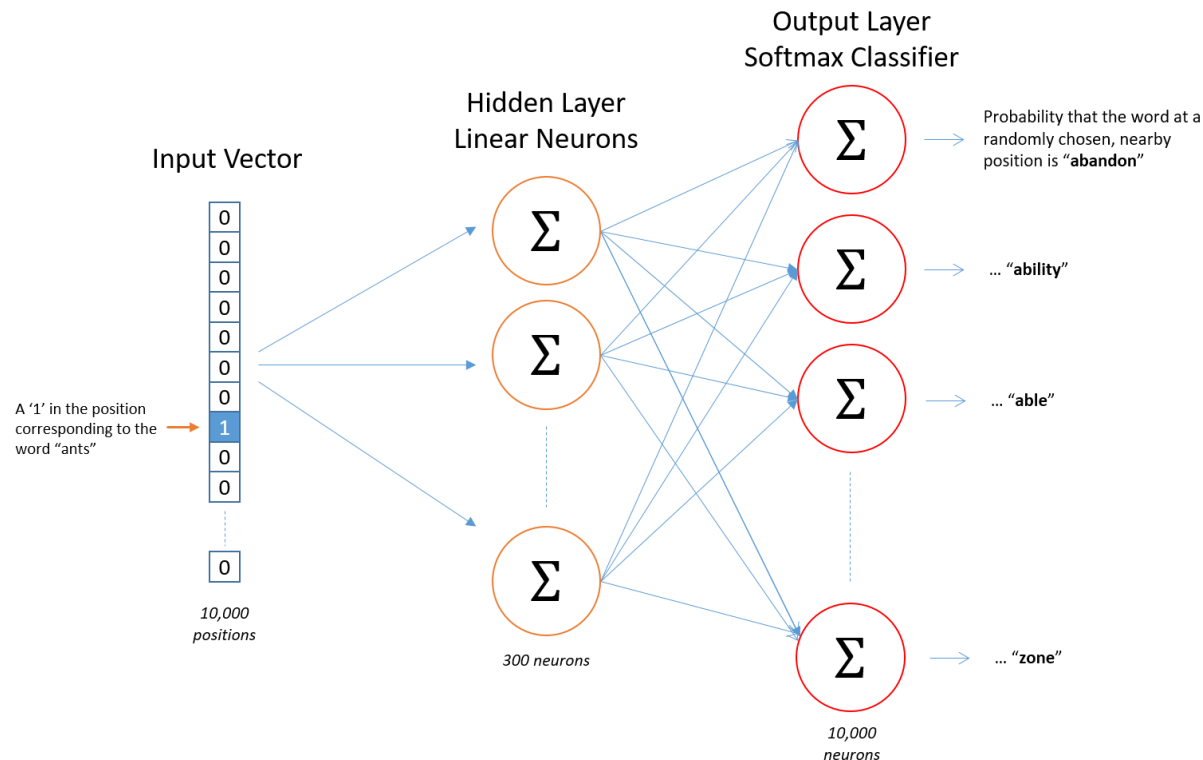
Paris = [0, 1, 0, 0, 0, 0, ..., 0]

Italy = [0, 0, 1, 0, 0, 0, ..., 0]

France = [0, 0, 0, 1, 0, 0, ..., 0]

Skip-gram NN architecture

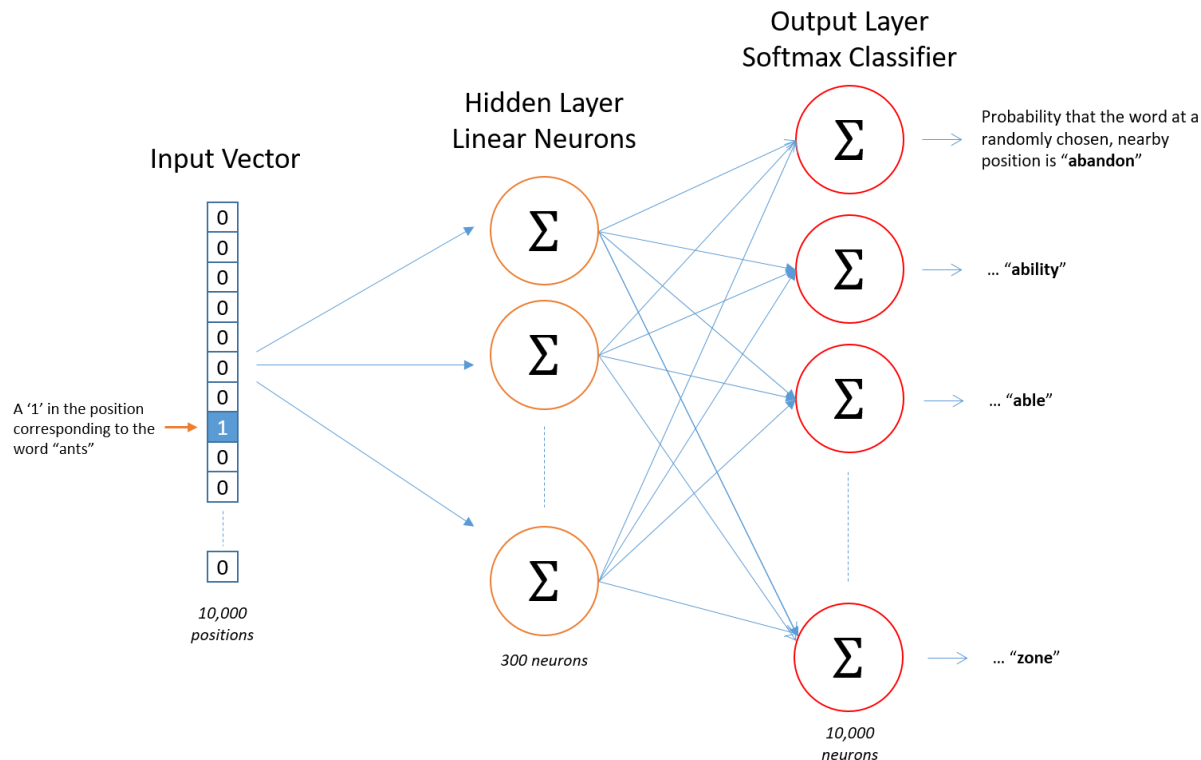
60



"Linear neurons" means there is no activation function on the hidden layer neurons, ...

Skip-gram NN architecture

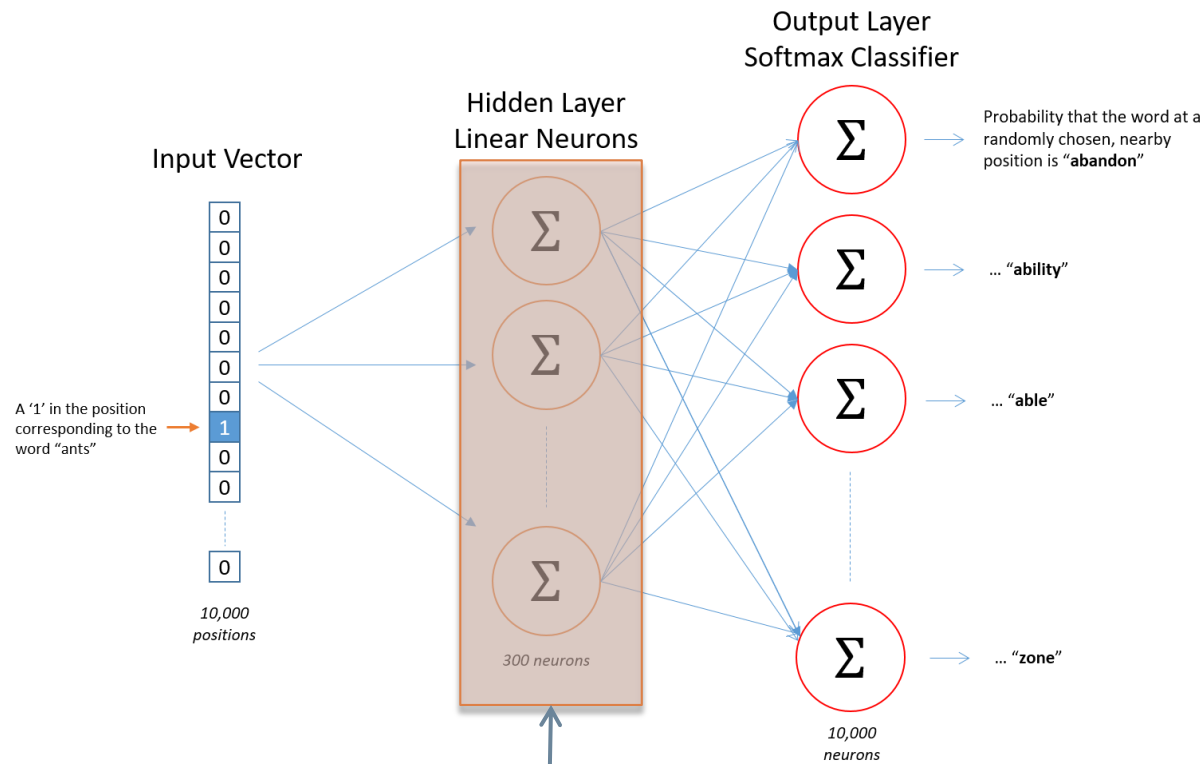
61



..., but the output neurons use softmax.

Skip-gram NN architecture

62



The amount of neurons in the hidden layer (a hyperparameter) determines the size of the embedding.

The hidden layer

63

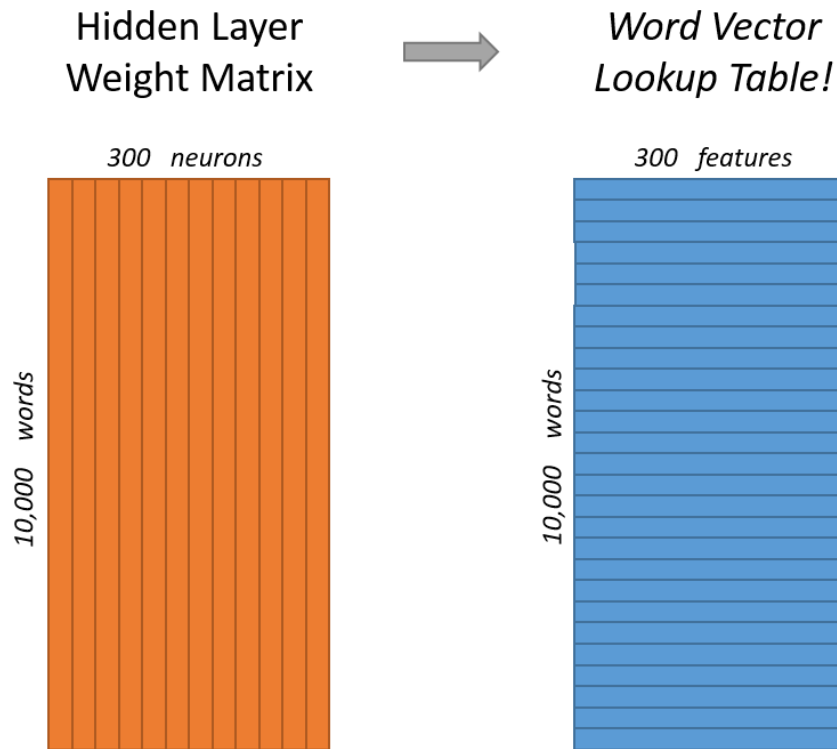
- Multiplying a $1 \times n$ one-hot vector $\mathbf{v}$ by an $n \times d$ matrix $\mathbf{W}$ (the hidden layer) will effectively *select* the row of $\mathbf{W}$ corresponding to the “1” position in $\mathbf{v}$.

$$\begin{bmatrix} 0 & 0 & 0 & 1 & 0 \end{bmatrix} \times \begin{bmatrix} 17 & 24 & 1 \\ 23 & 5 & 7 \\ 4 & 6 & 13 \\ 10 & 12 & 19 \\ 11 & 18 & 25 \end{bmatrix} = \begin{bmatrix} 10 & 12 & 19 \end{bmatrix}$$

So, the output of the hidden layer is just the “word vector” for the input word → the hidden layer plays the role of a **lookup table**.

The hidden layer

64

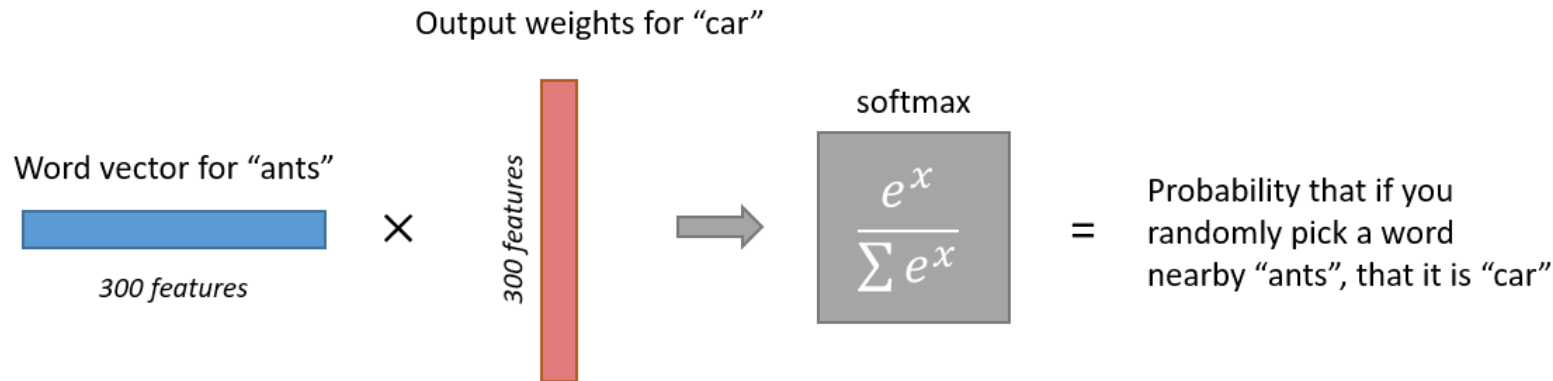


The output of the hidden layer is just the “word vector” for the input word.
W is the matrix we really want!

The output layer

65

- The $1 \times d$ word vector (coming from the hidden layer) then gets fed to the output layer.
- The output layer has $|V|$ neurons and is a **softmax regression classifier**.



Skip-gram captures context similarity

66

- If two words w_i and w_k have similar “contexts” (i.e., similar words are likely to appear around them), then the model needs to output very similar results for them.
 - ▣ One way for the network to do this is if *the word vectors for w_i and w_k are similar*.
- So, if two words have similar contexts, our network is motivated to learn similar word vectors for them.

A drawback of skip-gram

67

- The resulting trained model contains a huge amount of weights.
 - ▣ e.g., $d = 300$ and $|V| = 10,000$, that's 3M weights in the hidden layer and output layer each!
- Training (optimizing) a skip-gram model on a large corpus would be prohibitive.
- In order to circumvent this, some innovations...

Three innovations

68

1. Learn phrases from the corpus.
2. Subsample frequent words to decrease the number of training examples.
3. Modify the optimization objective with a technique called “Negative Sampling”, which causes each training sample to update only a small percentage of the model’s weights.

Learning Phrases

69

- A word pair like “Boston Globe” (a newspaper) has a much different meaning than the individual words “Boston” and “Globe”.
- So it makes sense to treat any occurrences of “Boston Globe” as a single word with its own word vector representation.

Learning Phrases

70

- In an experiment from the 2nd word2vec paper:
 - ▣ a model was trained on 100 billion words from a Google News dataset, using $d=300$.
 - ▣ the addition of phrases to the model reduced the vocabulary size to 3 million "words".
- This (*pre-trained*) model can be freely downloaded.
- Phrase detection is also covered in the paper.

Learning Phrases

71

- Each pass of their phrase detection tool only looks at combinations of 2 words.
- But it can be executed multiple times to get longer phrases.
- So, the first pass will pick up the phrase “New_York”, and then running it again will pick up “New_York_City” as a combination of “New_York” and “City”.

Subsampling frequent words

72

- Two issues with common words like “the”:
 1. When looking at word pairs, (“fox”, “the”) doesn’t tell us much about the meaning of “fox”.
 2. We will have many more samples of (“the”, ...) than we need to learn a good vector for “the”.
- A “subsampling” scheme addresses this...

Subsampling of frequent words

73

- Example: for a window size of 10, if a specific instance of “the” is removed:
 1. As we train on the remaining words, “the” will not appear in any of their context windows.
 2. We'll have 10 fewer training samples where “the” is the input word.
- This has a huge effect on the size of the training set.

Subsampling of frequent words

74

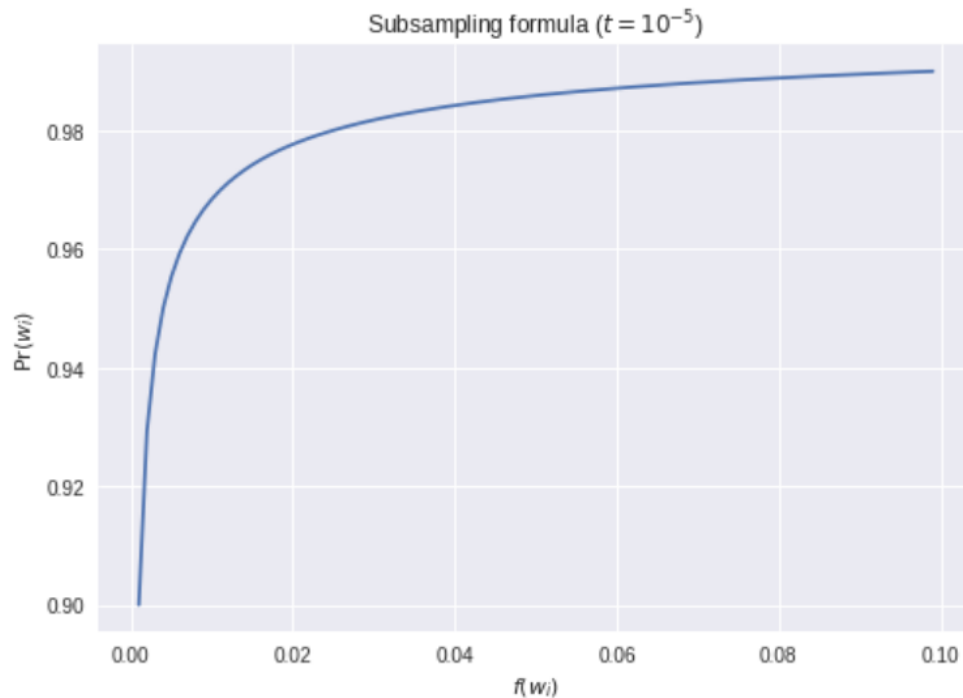
- “subsampling” scheme:
 - ▣ For each word found in the corpus, there is a chance that it will be discarded.
 - ▣ The probability that a word w_i will be discarded is related to its frequency:

$$P(w_i) = 1 - \sqrt{\frac{t}{f(w_i)}}$$

Subsampling of frequent words

75

$$P(w_i) = 1 - \sqrt{\frac{t}{f(w_i)}}$$



Negative sampling

76

- Train a NN means to iterate on: “take each $x^{(i)}$ and adjust all of the weights slightly so that the NN predicts $x^{(i)}$ more accurately”.
- i.e., each training sample will tweak *all* of the weights in the neural network.
- But the word2vec architecture is fully connected!
- Negative sampling addresses this problem.

Negative sampling

77

- Basic idea: modify only a small percentage of the weights in the NN for each training sample.
- When training the network on the word pair (“fox”, “quick”), recall that the “label” (response signal) is a one-hot vector.

Negative sampling

78

- With negative sampling, just a small number of “negative” words (let's say 5) are randomly selected to update the weights.
 - ▣ (a “negative” word is one for which we want the network to output a 0 for).
- We still update the weights for the “positive” word.

Negative sampling - example

79

- Consider Negative Sampling applied to a output layer with a $300 \times 10,000$ weight matrix.
 - ▣ we will just be updating the weights for the positive word, plus the weights for 5 other words that we want to output 0.
 - ▣ That's a total of 6 output neurons, and 1,800 weight values.
 - ▣ That's only 0.06% of the 3M weights in the output layer!

Negative sampling

80

- The “negative samples” are chosen using a “unigram distribution”.
- ▣ The probability for selecting a word as a negative sample is related to its frequency, with more frequent words being more likely to be selected as negative samples.

$$P(w_i) = \frac{f(w_i)^{3/4}}{\sum_{j=0}^n \left(f(w_j)^{3/4} \right)}$$

Negative sampling

81

- The 2nd paper says that
 - ▣ selecting 5-20 words works well for smaller datasets,
 - ▣ only 2-5 words suffice for large datasets.