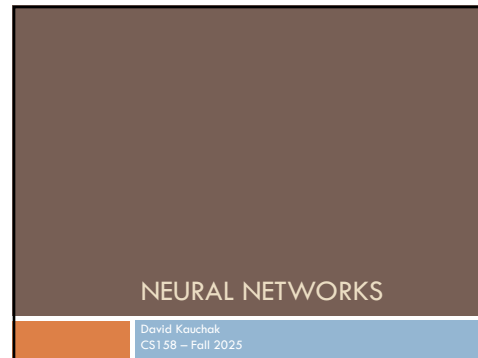
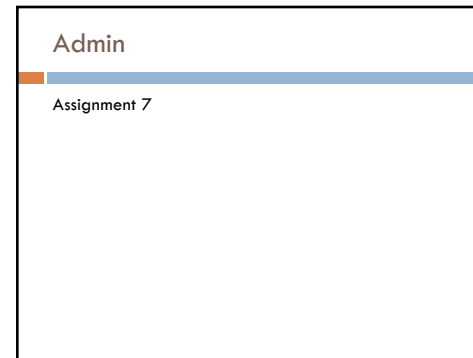
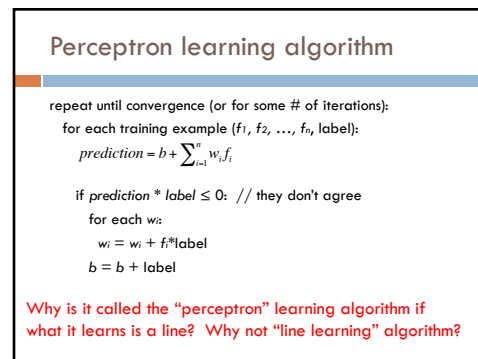




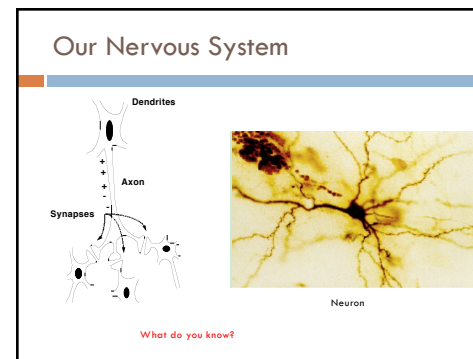
1



2

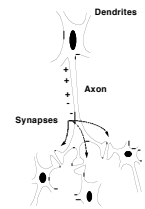


3



4

Our nervous system: the computer science view



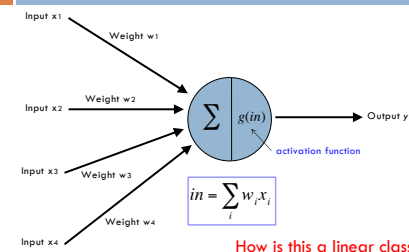
the human brain is a large collection of interconnected neurons

a **NEURON** is a brain cell

- they collect, process, and disseminate electrical signals
- they are connected via synapses
- they **FIRE** depending on the conditions of the neighboring neurons

5

A neuron/perceptron



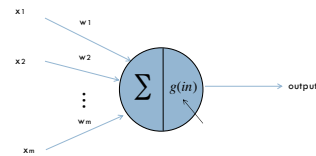
How is this a linear classifier (i.e. perceptron)?

6

Hard threshold = linear classifier

hard threshold:

$$g(in) = \begin{cases} 1 & \text{if } in > -b \\ 0 & \text{otherwise} \end{cases} \quad \text{output} = \begin{cases} 1 & \text{if } \sum_i w_i x_i + b > 0 \\ 0 & \text{otherwise} \end{cases}$$



7

Neural Networks

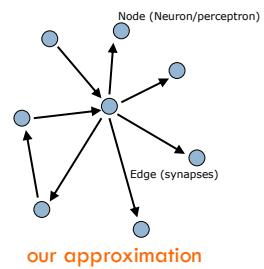
Neural Networks try to mimic the structure and function of our nervous system

People like biologically motivated approaches

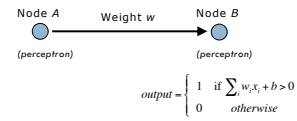


8

Artificial Neural Networks



9



W is the strength of signal sent between A and B.

If A fires and w is **positive**, then A **stimulates** B.

If A fires and w is **negative**, then A **inhibits** B.

10

Other activation functions

hard threshold:

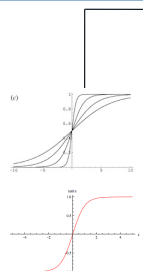
$$g(in) = \begin{cases} 1 & \text{if } in > -b \\ 0 & \text{otherwise} \end{cases}$$

sigmoid

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

tanh x

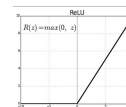
why other threshold functions?



11

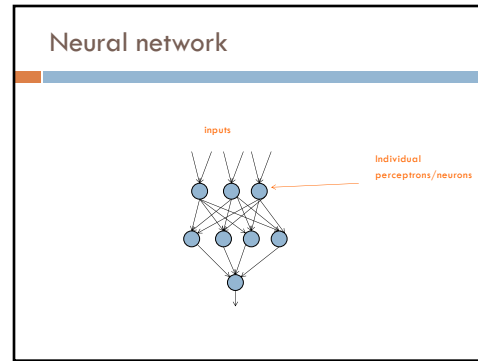
Many other activation functions

Rectified Linear Unit (ReLU)

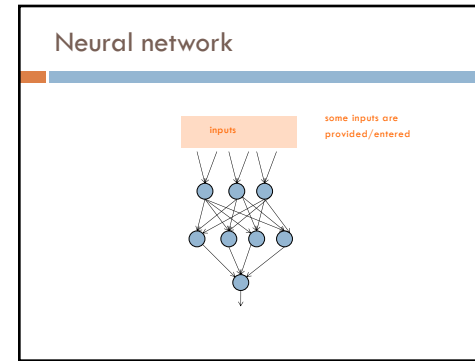


Softmax (for probabilities)

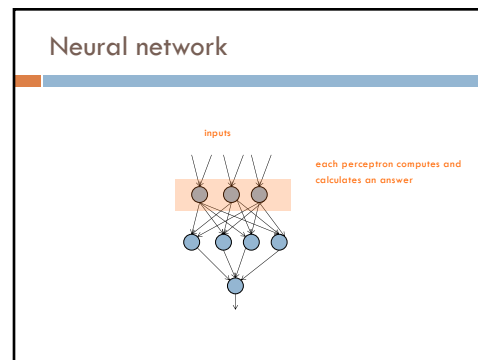
12



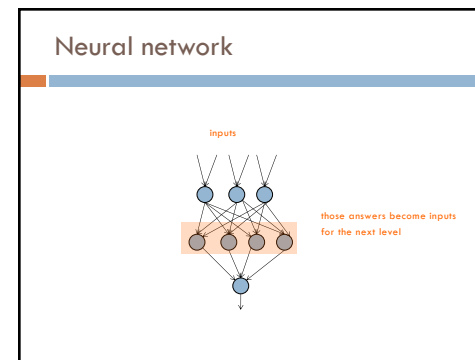
13



14

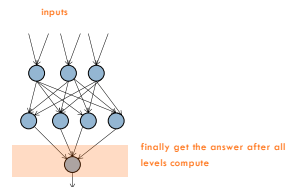


15



16

Neural network



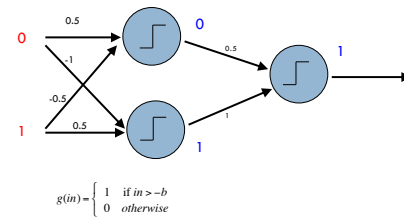
17

Activation spread

<http://www.youtube.com/watch?v=YqZd4ROvZ6I>

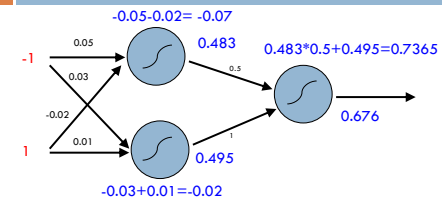
18

Computation (assume 0 bias)

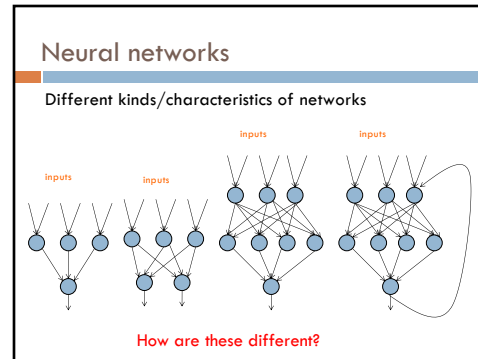


19

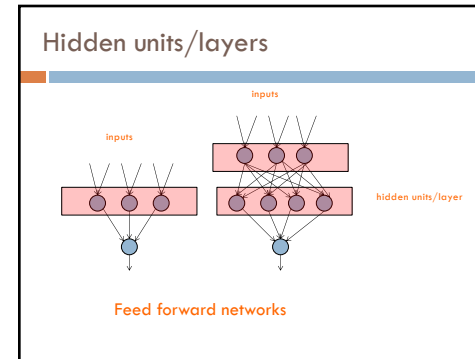
Computation



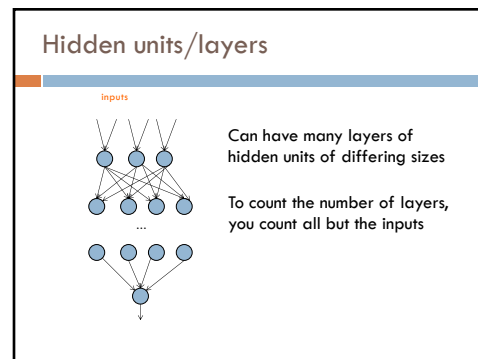
20



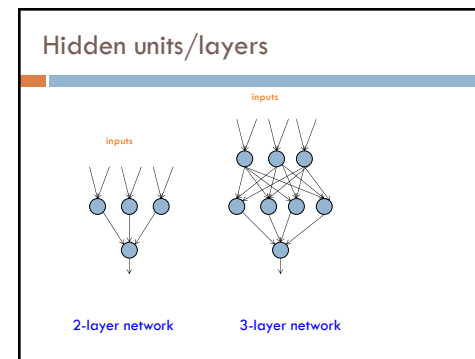
21



22



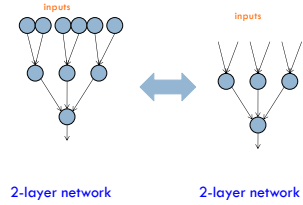
23



24

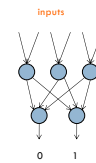
Alternate ways of visualizing

Sometimes the input layer will be drawn with nodes as well



25

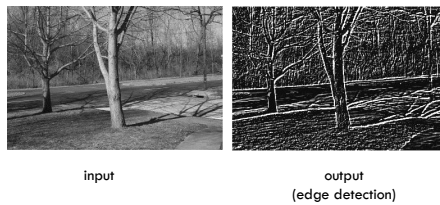
Multiple outputs



Can be used to model multiclass
datasets or more interesting
predictors, e.g. images

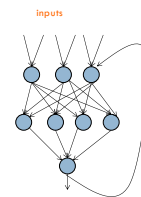
26

Multiple outputs



27

Neural networks



Recurrent network

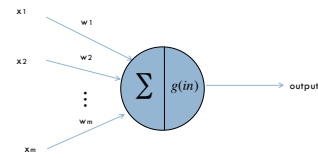
Output is fed back to input

Can support memory!

Good for temporal/sequential
data

28

NN decision boundary

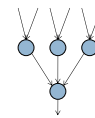


What does the decision boundary of a perceptron look like?

Line (linear set of weights)

29

NN decision boundary



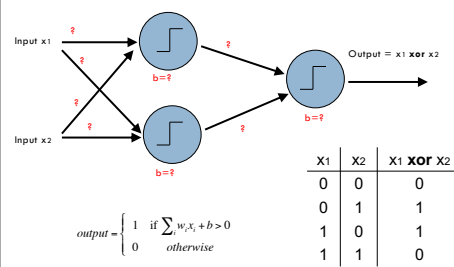
What does the decision boundary of a 2-layer network look like?

Is it linear?

What types of things can and can't it model?

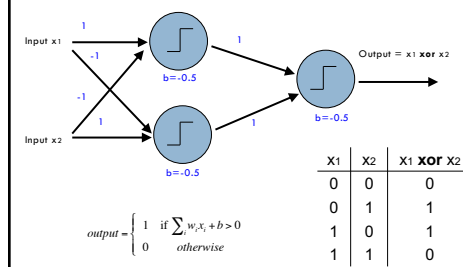
30

XOR



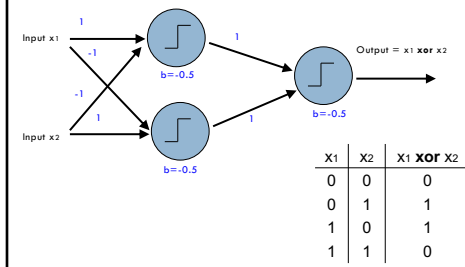
31

XOR



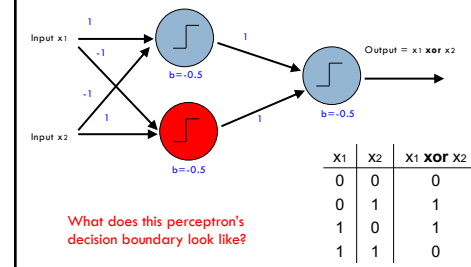
32

What does the decision boundary look like?



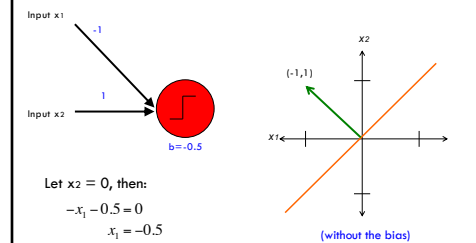
33

What does the decision boundary look like?



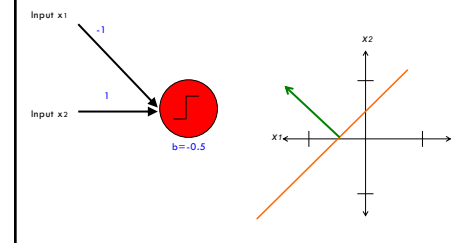
34

NN decision boundary



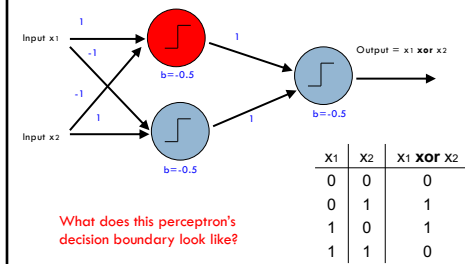
35

NN decision boundary



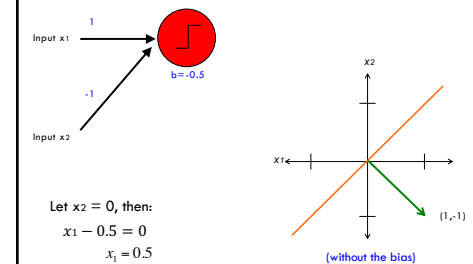
36

What does the decision boundary look like?



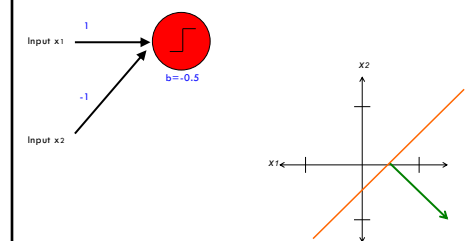
37

NN decision boundary



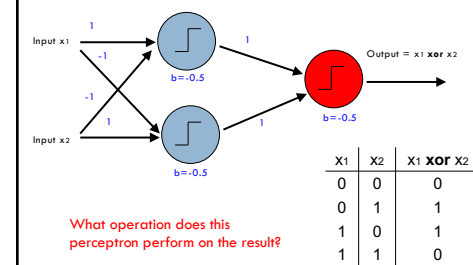
38

NN decision boundary

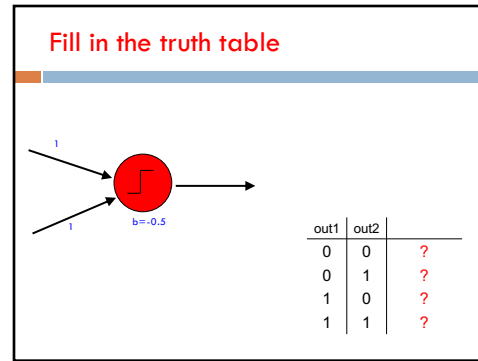


39

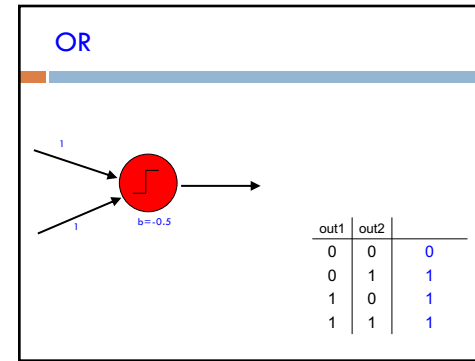
What does the decision boundary look like?



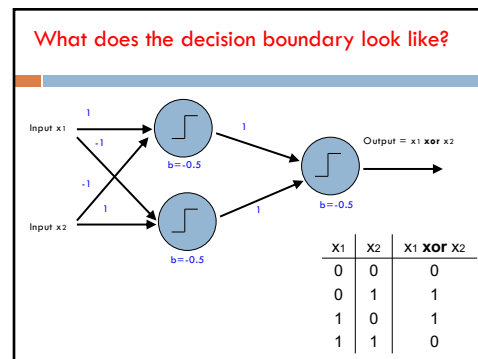
40



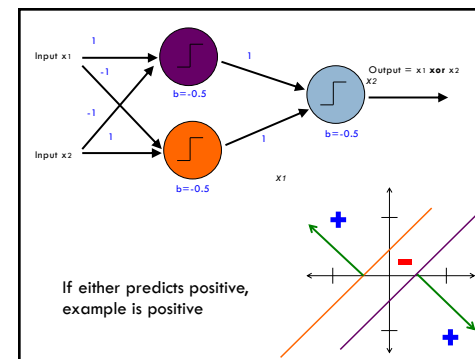
41



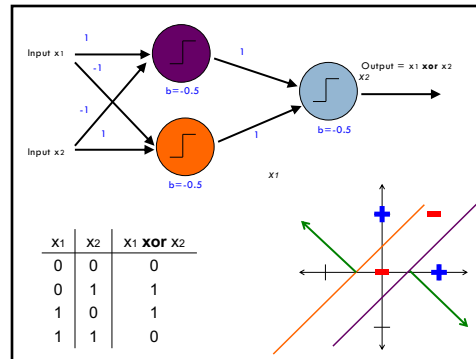
42



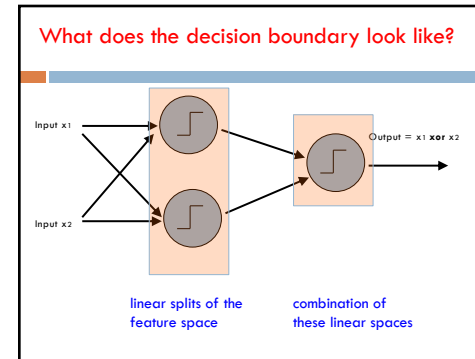
43



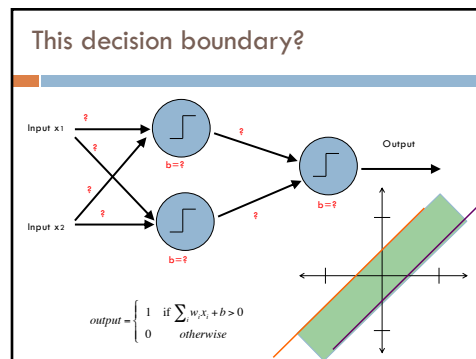
44



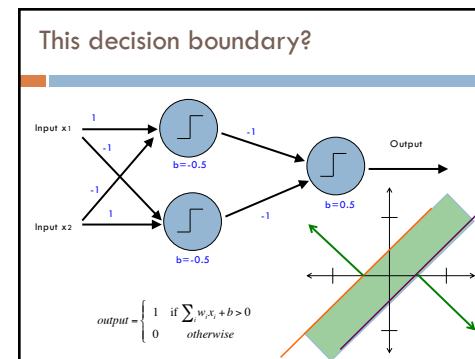
45



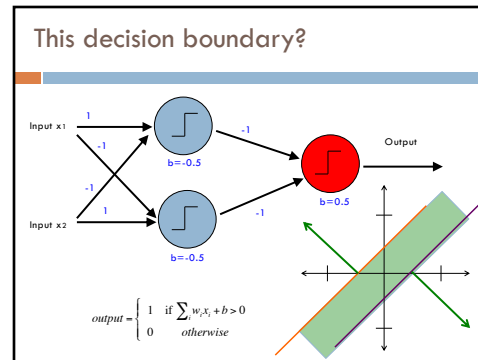
46



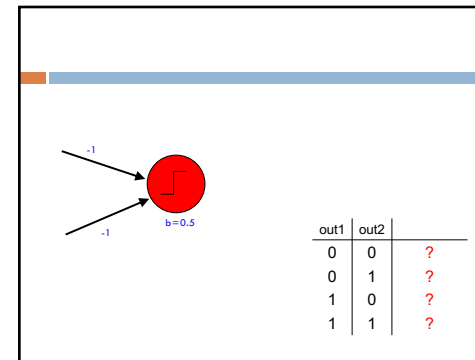
47



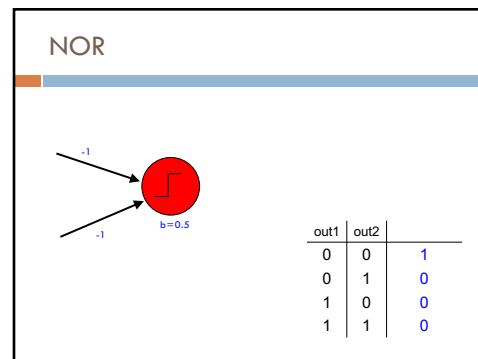
48



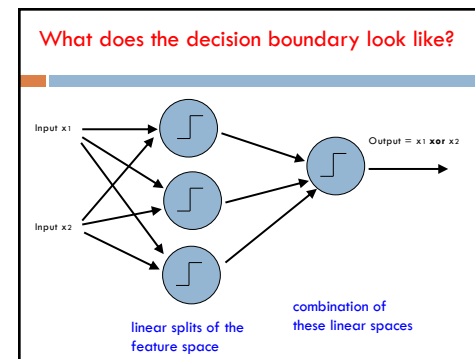
49



50

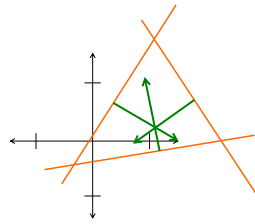


51



52

Three hidden nodes



53

NN decision boundaries

Theorem 9 (Two-Layer Networks are Universal Function Approximators). Let F be a continuous function on a bounded subset of D -dimensional space. Then there exists a two-layer neural network $\hat{F}$ with a finite number of hidden units that approximate F arbitrarily well. Namely, for all x in the domain of F , $|F(x) - \hat{F}(x)| < \epsilon$.

Put simply: **two-layer networks can approximate any function**

54

NN decision boundaries

For DT, as the tree gets larger, the model gets more complex

The same is true for neural networks:
more hidden nodes = more complexity

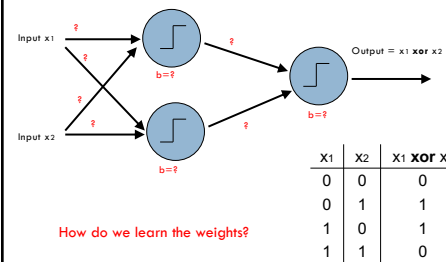
Adding more layers adds even more complexity (and much more quickly)

Good rule of thumb:

$$\text{number of 2-layer hidden nodes} \leq \frac{\text{number of examples}}{\text{number of dimensions}}$$

55

Training



56

Training multilayer networks

perceptron learning: if the perceptron's **output** is different than the expected output, update the weights

gradient descent: compare **output** to label and adjust based on loss function

Any other problem with these for general NNs?



perceptron/
linear model



neural network

57

Learning in multilayer networks

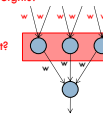
Challenge: for multilayer networks, we don't know what the expected output/error is for the internal nodes!

how do we learn these weights?

expected output?



perceptron/
linear model



neural network

58

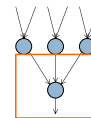
Backpropagation: intuition

Gradient descent method for learning weights by optimizing a loss function

1. calculate output of all nodes
2. calculate the weights for the output layer based on the error
3. "backpropagate" errors through hidden layers

59

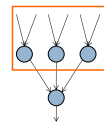
Backpropagation: intuition



We can calculate the actual error here

60

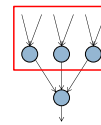
Backpropagation: intuition



Key idea: propagate the error back to this layer

61

Backpropagation: intuition



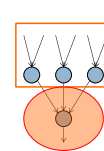
"backpropagate" the error:

Assume all of these nodes were responsible for some of the error

How can we figure out how much they were responsible for?

62

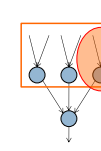
Backpropagation: intuition



error for node is $\sim w_i * \text{error}$

63

Backpropagation: intuition



Calculate as normal using this as the error

64

Backpropagation: the details

Gradient descent method for learning weights by optimizing a **loss function**

1. calculate output of all nodes
2. calculate the updates directly for the output layer
3. "backpropagate" errors through hidden layers

What loss function?

65

Backpropagation: the details

Gradient descent method for learning weights by optimizing a **loss function**

1. calculate output of all nodes
2. calculate the updates directly for the output layer
3. "backpropagate" errors through hidden layers

$$loss = \sum_x \frac{1}{2} (y - \hat{y})^2 \quad \text{squared error}$$

66