

Complete Backpropagation

A Fully Worked Example — 2-2-2 Network, Sigmoid Units

One forward pass, one backward pass, one weight update

@SSRoy

1. The network and the given values

Two inputs i_1, i_2 ; two hidden neurons h_1, h_2 ; two outputs o_1, o_2 . Every neuron uses the sigmoid $\sigma(x) = \frac{1}{1 + e^{-x}}$. A single bias node feeds each layer.

Inputs	$i_1 = 0.05, i_2 = 0.10$	Targets	$t_1 = 0.01, t_2 = 0.99$
Input→hidden	$w_1=0.15, w_2=0.20, w_3=0.25, w_4=0.30$	Hidden bias	$b_1 = 0.35$
Hidden→output	$w_5=0.40, w_6=0.45, w_7=0.50, w_8=0.55$	Output bias	$b_2 = 0.60$

Key idea

Backpropagation is *one* idea applied repeatedly: the chain rule, run backward.

Forward: $i \rightarrow \text{net} \rightarrow \text{out} \rightarrow E \quad \implies \quad \text{Backward: } E \rightarrow \text{out} \rightarrow \text{net} \rightarrow w$

2. Part A — Forward propagation

Each neuron computes $\text{net} = \sum(\text{weight} \times \text{input}) + \text{bias}$, then $\text{out} = \sigma(\text{net})$.

Hidden layer

$$\text{net}_{h_1} = i_1 w_1 + i_2 w_2 + b_1 = (0.05)(0.15) + (0.10)(0.20) + 0.35 = 0.377500,$$

$$h_1 = \sigma(0.377500) = 0.593270;$$

$$\text{net}_{h_2} = i_1 w_3 + i_2 w_4 + b_1 = (0.05)(0.25) + (0.10)(0.30) + 0.35 = 0.392500,$$

$$h_2 = \sigma(0.392500) = 0.596884.$$

Output layer

$$\text{net}_{o_1} = h_1 w_5 + h_2 w_6 + b_2 = (0.593270)(0.40) + (0.596884)(0.45) + 0.60 = 1.105906,$$

$$o_1 = \sigma(1.105906) = 0.751365;$$

$$\text{net}_{o_2} = h_1 w_7 + h_2 w_8 + b_2 = (0.593270)(0.50) + (0.596884)(0.55) + 0.60 = 1.224921,$$

$$o_2 = \sigma(1.224921) = 0.772928.$$

3. Part B — The error

Squared error, summed over both outputs: $E = \frac{1}{2} \sum (t - o)^2 = E_1 + E_2$.

$$E_1 = \frac{1}{2}(t_1 - o_1)^2 = \frac{1}{2}(0.01 - 0.751365)^2 = \frac{1}{2}(0.741365)^2 = 0.274811,$$

$$E_2 = \frac{1}{2}(t_2 - o_2)^2 = \frac{1}{2}(0.99 - 0.772928)^2 = \frac{1}{2}(0.217072)^2 = 0.023560,$$

$$E = 0.274811 + 0.023560 = \mathbf{0.298371}.$$

4. Part C — The core: gradient of one output weight w_5

The path from w_5 to the error is $w_5 \rightarrow \text{net}_{o_1} \rightarrow o_1 \rightarrow E$. So by the chain rule we need three factors:

$$\frac{\partial E}{\partial w_5} = \underbrace{\frac{\partial E}{\partial o_1}}_{(1)} \cdot \underbrace{\frac{\partial o_1}{\partial \text{net}_{o_1}}}_{(2)} \cdot \underbrace{\frac{\partial \text{net}_{o_1}}{\partial w_5}}_{(3)}$$

(1) Error vs. output. From $E_1 = \frac{1}{2}(o_1 - t_1)^2$:

$$\frac{\partial E}{\partial o_1} = o_1 - t_1 = 0.751365 - 0.01 = 0.741365.$$

(2) Output vs. net — the sigmoid derivative. We *derive* it, not memorise it. With $o_1 = (1 + e^{-\text{net}})^{-1}$,

$$\frac{do_1}{d\text{net}} = -(1 + e^{-\text{net}})^{-2} \cdot (-e^{-\text{net}}) = \frac{e^{-\text{net}}}{(1 + e^{-\text{net}})^2}.$$

Now notice $1 - o_1 = \frac{e^{-\text{net}}}{1 + e^{-\text{net}}}$, so $o_1(1 - o_1) = \frac{1}{1 + e^{-\text{net}}} \cdot \frac{e^{-\text{net}}}{1 + e^{-\text{net}}} = \frac{e^{-\text{net}}}{(1 + e^{-\text{net}})^2}$. Therefore

$$\frac{\partial o_1}{\partial \text{net}_{o_1}} = o_1(1 - o_1) = 0.751365(1 - 0.751365) = 0.751365 \times 0.248635 = 0.186816.$$

(3) Net vs. weight. Since $\text{net}_{o_1} = h_1 w_5 + h_2 w_6 + b_2$, only the w_5 term survives: $\frac{\partial \text{net}_{o_1}}{\partial w_5} = h_1 = 0.593270$.

Combine:

$$\frac{\partial E}{\partial w_5} = (0.741365)(0.186816)(0.593270) = 0.138499 \times 0.593270 = \mathbf{0.082167}.$$

Key idea

Group the first two factors and call it the neuron's **delta**:

$$\delta_o \equiv \frac{\partial E}{\partial \text{net}_o} = (o - t) o(1 - o), \quad \implies \quad \frac{\partial E}{\partial w} = \delta_{(\text{receiving neuron})} \times (\text{input on that weight}).$$

Here $\delta_{o_1} = (0.741365)(0.186816) = 0.138499$ and $\frac{\partial E}{\partial w_5} = \delta_{o_1} h_1$.

5. Part D — All four output-layer gradients

$\delta_{o_1} = 0.138499$ (above). For o_2 : $\frac{\partial E}{\partial o_2} = o_2 - t_2 = -0.217072$, $o_2(1 - o_2) = 0.175510$, so $\delta_{o_2} = (-0.217072)(0.175510) = -0.038098$.

$$\begin{aligned} \frac{\partial E}{\partial w_5} &= \delta_{o_1} h_1 = (0.138499)(0.593270) = 0.082167, & \frac{\partial E}{\partial w_6} &= \delta_{o_1} h_2 = (0.138499)(0.596884) = 0.082668, \\ \frac{\partial E}{\partial w_7} &= \delta_{o_2} h_1 = (-0.038098)(0.593270) = -0.022603, & \frac{\partial E}{\partial w_8} &= \delta_{o_2} h_2 = (-0.038098)(0.596884) = -0.022740. \end{aligned}$$

6. Part E — Backpropagate into the hidden layer

Intuition

A hidden neuron feeds *both* outputs, so error reaches h_1 along two paths ($h_1 \rightarrow o_1 \rightarrow E$ and $h_1 \rightarrow o_2 \rightarrow E$). We must **sum** both contributions before multiplying by the local sigmoid slope.

$$\delta_h = \left(\sum_k \delta_{o_k} w_{h \rightarrow o_k} \right) h(1-h).$$

For h_1 (weights w_5, w_7 carry error back from o_1, o_2):

$$\begin{aligned} \sum &= \delta_{o_1} w_5 + \delta_{o_2} w_7 = (0.138499)(0.40) + (-0.038098)(0.50) = 0.055399 - 0.019049 = 0.036350, \\ h_1(1-h_1) &= 0.593270 \times 0.406730 = \mathbf{0.241301}, \\ \delta_{h_1} &= (0.036350)(0.241301) = \mathbf{0.008771}. \end{aligned}$$

For h_2 (weights w_6, w_8):

$$\begin{aligned} \sum &= \delta_{o_1} w_6 + \delta_{o_2} w_8 = (0.138499)(0.45) + (-0.038098)(0.55) = 0.062324 - 0.020954 = 0.041370, \\ h_2(1-h_2) &= 0.596884 \times 0.403116 = 0.240613, \\ \delta_{h_2} &= (0.041370)(0.240613) = \mathbf{0.009954}. \end{aligned}$$

7. Part F — The four hidden-layer gradients

Same rule, $\frac{\partial E}{\partial w} = \delta_{\text{receiving}} \times \text{input}$, with inputs i_1, i_2 :

$$\begin{aligned} \frac{\partial E}{\partial w_1} &= \delta_{h_1} i_1 = (0.008771)(0.05) = 0.0004386, & \frac{\partial E}{\partial w_2} &= \delta_{h_1} i_2 = (0.008771)(0.10) = 0.0008771, \\ \frac{\partial E}{\partial w_3} &= \delta_{h_2} i_1 = (0.009954)(0.05) = 0.0004977, & \frac{\partial E}{\partial w_4} &= \delta_{h_2} i_2 = (0.009954)(0.10) = 0.0009954. \end{aligned}$$

8. Part G — Update the weights (gradient descent)

$$w_{\text{new}} = w_{\text{old}} - \eta \frac{\partial E}{\partial w}, \quad \eta = 0.5.$$

Intuition

We *subtract* because the gradient points uphill (toward larger error). A positive gradient means “increasing w raises E ,” so we lower w ; a negative gradient (as for w_7, w_8) makes w *rise*. The step size is η .

Weight	Old	Gradient	Change ($-\eta g$)	New
w_1	0.150000	+0.0004386	−0.0002193	0.1497807
w_2	0.200000	+0.0008771	−0.0004386	0.1995614
w_3	0.250000	+0.0004977	−0.0002489	0.2497512
w_4	0.300000	+0.0009954	−0.0004977	0.2995023
w_5	0.400000	+0.0821670	−0.0410835	0.3589165
w_6	0.450000	+0.0826680	−0.0413340	0.4086660
w_7	0.500000	−0.0226025	+0.0113013	0.5113013
w_8	0.550000	−0.0227402	+0.0113701	0.5613701

Biases (single shared bias per layer, as drawn)

The one bias node of a layer feeds *every* neuron in that layer, so its gradient is the *sum* of those neurons' deltas:

$$\begin{aligned}\frac{\partial E}{\partial b_2} &= \delta_{o_1} + \delta_{o_2} = 0.138499 - 0.038098 = 0.100401, & b_2^{\text{new}} &= 0.60 - 0.5(0.100401) = \mathbf{0.549800}, \\ \frac{\partial E}{\partial b_1} &= \delta_{h_1} + \delta_{h_2} = 0.008771 + 0.009954 = 0.018725, & b_1^{\text{new}} &= 0.35 - 0.5(0.018725) = \mathbf{0.340638}.\end{aligned}$$

Intuition

If instead each neuron had its *own* bias (the more common textbook setup), you would *not* sum: $\partial E / \partial b_{o_1} = \delta_{o_1}$, $\partial E / \partial b_{o_2} = \delta_{o_2}$, and so on. Summing is correct *only* because this diagram shares one bias across the layer.

9. Verification — did the error actually drop?

Running one more forward pass with all updated weights and biases:

$$E_{\text{before}} = 0.298371 \rightarrow E_{\text{after one step}} = 0.285751 \quad (\Delta E = -0.012620).$$

The error decreased, confirming the update moved in the descent direction. Repeating forward→backward→update many times (epochs) drives $o_1 \rightarrow 0.01$ and $o_2 \rightarrow 0.99$.

10. The whole algorithm on one page

Key idea

$$\begin{aligned}\frac{\partial E}{\partial w} &= \frac{\partial E}{\partial \text{out}} \cdot \frac{\partial \text{out}}{\partial \text{net}} \cdot \frac{\partial \text{net}}{\partial w}, & \frac{\partial \text{out}}{\partial \text{net}} &= \text{out}(1 - \text{out}), & \frac{\partial \text{net}}{\partial w} &= \text{input}.\end{aligned}$$

Hence $\frac{\partial E}{\partial w} = \delta \times \text{input}$, with $\delta_{\text{output}} = (o - t) o(1 - o)$ and $\delta_{\text{hidden}} = (\sum_k \delta_k w_k) h(1 - h)$.

Vectorised form (what libraries actually run). With activations a , targets t , layer output $\mathbf{o}$, and $\odot$ the elementwise product:

$$\delta^L = (\mathbf{o} - \mathbf{t}) \odot \mathbf{o} \odot (1 - \mathbf{o}), \quad \delta^\ell = (W_{\ell+1}^\top \delta^{\ell+1}) \odot a^\ell \odot (1 - a^\ell), \quad \frac{\partial E}{\partial W_\ell} = \delta^\ell (a^{\ell-1})^\top.$$

Common student mistake

- Using net instead of the sigmoid *output* in $o(1 - o)$ — the slope uses the *activation*, not the pre-activation.
- Forgetting to *sum* both output paths when computing a hidden δ .
- Updating weights *during* the backward pass — compute *all* gradients from the *old* weights first, then update together.
- Sign slip on $(o - t)$ vs. $(t - o)$: with $E = \frac{1}{2}(t - o)^2$ the clean result is $\partial E / \partial o = o - t$.
- Dropping η : the update is not unique without a learning rate. Here $\eta = 0.5$ (chosen for this example; it is *not* given by the diagram).

Forward $\rightarrow$ Error $\rightarrow$ Backpropagation $\rightarrow$ Update. Repeat.