

Neural networks and the backpropagation algorithm

S. Marchand-Maillet

Computer Science – University of Geneva

1 introduction

Given a function $f : \mathbb{R}^D \rightarrow \mathbb{R}^d$, we wish to approximate (regress) it from training data $\mathcal{X} \times \mathcal{Y} = \{(\mathbf{x}_i, \mathbf{y}_i)\}_{i=1}^N$ (where $\mathbf{y}_i = f(\mathbf{x}_i)$) using a mapping

$$\begin{aligned} \Phi_{\theta} : \mathbb{R}^D &\rightarrow \mathbb{R}^d \\ \mathbf{x}_i &\mapsto \Phi_{\theta}(\mathbf{x}_i) \end{aligned}$$

parameterized with parameters θ (using notation from statistics).

In machine learning terms, f is the mapping that underlies the phenomenon we wish to learn and is unknown. We can access annotated samples $\mathcal{X} \times \mathcal{Y} = \{(\mathbf{x}_i, \mathbf{y}_i)\}_{i=1}^N$ and we wish to train a learner Φ_{θ} by adjusting its parameters θ to minimize a loss function $\mathcal{L}(\mathcal{X} \times \mathcal{Y}; \theta)$.

Here, we present $\Phi_{\theta}(\cdot)$ as a feedforward neural network taking D input values and producing d output values. We will further consider the practical case of the [squared loss](#):

$$\mathcal{L}(\mathcal{X} \times \mathcal{Y}; \theta) = \frac{1}{2} \frac{1}{N} \sum_{i=1}^N \|\Phi_{\theta}(\mathbf{x}_i) - \mathbf{y}_i\|_2^2 = \frac{1}{2N} \sum_{i=1}^N \sum_{k=1}^d (\Phi_{\theta}(\mathbf{x}_i)_{[k]} - \mathbf{y}_i_{[k]})^2 \quad (1)$$

Note. When the label $\mathbf{y}_i$ is of dimension $d = 1$, the squared loss reduces to the classical

$$\mathcal{L}(\mathcal{X} \times \mathcal{Y}; \theta) = \frac{1}{2N} \sum_{i=1}^N (\hat{y}_i - y_i)^2 = \frac{1}{2N} \sum_{i=1}^N (\Phi_{\theta}(\mathbf{x}_i) - y_i)^2$$

A Neural Network model is a Directed Acyclic Graph (DAG) of units (referred to as “neurons”) organized in layers (see figure 2 for an example). In the particular case of a [feedforward](#) network, conventionally, connections between neurons only exist from one layer to the next.

1.1 “Neuron” model and notation

A generic neuron is located at layer $l \in \llbracket L \rrbracket$ in a network of L layers and takes its inputs from the outputs of other neurons. To ease the presentation and because this is the dominant “feedforward” model, we assume that the neuron takes its inputs from the output of the neurons in the previous layer $l - 1$ (see figure 1)^(a).

A neuron (l, k) located at index $k \in \llbracket N_l \rrbracket$ of layer $l \in \llbracket L \rrbracket$ comprising N_l neurons:

1. receives its inputs from weighted connections with weights w_{jk}^l
2. aggregates these inputs into a linear sum into its [activation](#) $a_k^l = \sum_{j=0}^{N_{l-1}} w_{jk}^l \phi_k^{l-1}$
3. outputs the result of its [activation function](#) $\phi_k^l(\cdot)$ applied to its activation: $\phi_k^l = \phi_k^l(a_k^l)$

Note that in this model, the weights are attached to the neuron itself. They form the adjustable parameters of the model

$$\theta = \{w_{jk}^l \mid j \in \llbracket N_{l-1} \rrbracket, k \in \llbracket N_l \rrbracket, l \in \llbracket L \rrbracket\}$$

The [hyperparameters](#) of the model concern its structure and choice of activation functions, i.e $L > 0$, $\{N_l\}_{l=1}^L$ and $\{\phi_k^l\}_{k,l=1}^{N_l,L}$.

Also following the feedforward structure, we fix the following boundary conditions for the model:

$$\begin{cases} \phi_k^0 = \mathbf{x}_{i[k]} & N_0 = D & (\text{input}) \\ \phi_{\theta}(\mathbf{x}_i) = \Phi^L(\mathbf{x}_i) & N_L = d & (\text{output}) \end{cases}$$

^(a)A more generic formulation using the notation for graph adjacency is left as exercise.

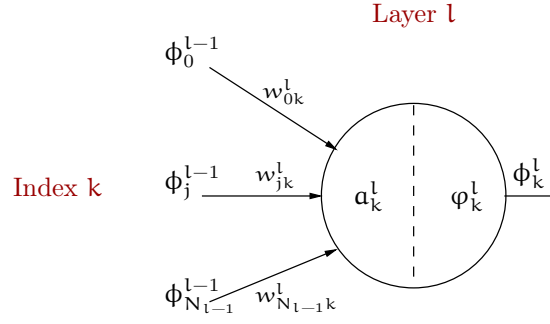


Figure 1: A generic neuron in a feedforward network, taking input from the previous layer

where the notation $\phi_k^l(\mathbf{x}_i)$ indicates the output of neuron (l, k) when the network takes $\mathbf{x}_i$ as input. $a_k^l(\mathbf{x}_i)$ is defined similarly.

Using variable homogenization, it is also customary to consider that $\phi_0^l = 1$ where a [bias](#) (hence w_0^l) should be introduced in order to escape from the pure linear activation.

Note. The [logistic neuron](#) defines its activation functions as the [logistic function](#):

$$\phi_k^l(z) = \frac{1}{1 + e^{-z}} \quad \forall k \in \llbracket N_l \rrbracket \quad \forall l \in \llbracket L \rrbracket$$

The present study is more general. It is left as exercise to model the [logistic regression](#) as a neural network.

As a result we have the following system of equations:

$$\left\{ \begin{array}{ll} \phi_k^l = \phi_k^l(a_k^l) & \forall k \in \llbracket N_l \rrbracket \quad \forall l \in \llbracket L \rrbracket \\ a_k^l = \sum_{j=0}^{N_{l-1}} w_{jk}^l \phi_j^{l-1} & \forall k \in \llbracket N_l \rrbracket \quad \forall l \in \llbracket L \rrbracket \\ \phi_k^0 = \mathbf{x}_i[k] & \forall k \in \llbracket D \rrbracket \\ \phi_0^l = 1 & \forall l \in \llbracket L-1 \rrbracket \\ \phi_{\theta}(\mathbf{x}_i)[k] = \phi_k^l(\mathbf{x}_i) & \forall k \in \llbracket d \rrbracket \quad \forall i \in \llbracket N \rrbracket \\ \mathcal{L}(\mathcal{X} \times \mathcal{Y}; \theta) = \frac{1}{2N} \sum_{i=1}^N \sum_{k=1}^d (\phi_{\theta}(\mathbf{x}_i)[k] - \mathbf{y}_i[k])^2 & \end{array} \right.$$

to adjust the parameters of the model $\theta = \{w_{jk}^l \mid j \in \llbracket N_{l-1} \rrbracket, k \in \llbracket N_l \rrbracket, l \in \llbracket L \rrbracket\}$.

Note. In this document, the assumption is that of a DAG. The connectivity may be more general than “feedforward” (e.g with “skip-connections” connecting layers that are not successive). Only the notation would change.

1.2 Training the network using the backpropagation algorithm

The aim of training is the minimization of the approximation (regression, prediction) error (the “loss”):

$$\theta^* = \underset{\theta \in \Theta}{\operatorname{argmin}} \mathcal{L}(\mathcal{X} \times \mathcal{Y}; \theta)$$

As indicated, the parameters of the models are the weights $\theta = \{w_{jk}^l\}_{j,k,l}$.

The algorithm used for training is the [gradient descent](#) over the error function to adjust the parameters. In other words, the simple structure of the algorithm is as shown in algorithm 1.

The key element of the algorithm is therefore the computation of all values of

$$\frac{\partial \mathcal{L}}{\partial w_{jk}^l} \quad \forall j \in \llbracket N_{l-1} \rrbracket, k \in \llbracket N_l \rrbracket, l \in \llbracket L \rrbracket$$

as the variation of the loss against the modification of weights w_{jk}^l .

Algorithm 1 The base structure of the backpropagation algorithm

```

1: procedure BACKPROPAGATION( $\mathcal{X}, \mathcal{Y}$ )
2:    $w_{jk}^{l,(0)} \leftarrow \text{random}()$   $\forall j \in \llbracket N_{l-1} \rrbracket, k \in \llbracket N_l \rrbracket, l \in \llbracket L \rrbracket$  ▷ Initialize random weights
3:   for iteration  $t = 1$  until  $t = T$  do ▷ or until convergence
4:     for  $i \in \llbracket N \rrbracket$  do ▷ Gradient iteration for data  $\mathbf{x}_i$ 
5:       for  $j \in \llbracket N_{l-1} \rrbracket, k \in \llbracket N_l \rrbracket, l \in \llbracket L \rrbracket$  do
6:          $w_{jk}^{l,(t)} \leftarrow w_{jk}^{l,(t-1)} - \frac{1}{N} \eta \frac{\partial \mathcal{L}}{\partial w_{jk}^l}(\mathbf{x}_i)$  ▷  $\eta > 0$  is the gradient step (learning rate)

```

The name “backpropagation” for this algorithm comes from the realization that when the network has more than one layer ($L > 1$), there is no direct access to $\frac{\partial \mathcal{L}}{\partial w_{jk}^l}$. The chain rule should be used and the gradient propagated back (“backpropagated”) from the end of the network towards its input.

In the specific case of the squared loss, given $p \in \llbracket N_{l-1} \rrbracket, r \in \llbracket N_l \rrbracket, m \in \llbracket L \rrbracket$ we write

$$\frac{\partial \mathcal{L}}{\partial w_{pr}^m} = \frac{1}{N} \sum_{i=1}^N \sum_{k=1}^d \frac{\partial \phi_{\theta[k]}(\mathbf{x}_i)}{\partial w_{pr}^m} (\phi_{\theta}(\mathbf{x}_i)_{[k]} - \mathbf{y}_i_{[k]})$$

where $\phi_{\theta[k]}$ is the k^{th} component function of $\phi_{\theta}(\cdot)$. By definition, $\phi_{\theta[k]}(\mathbf{x}_i) = \phi_k^L(\mathbf{x}_i)$ so we can write

$$\frac{\partial \phi_{\theta[k]}(\mathbf{x}_i)}{\partial w_{pr}^m} = \frac{\partial \phi_k^L(\mathbf{x}_i)}{\partial w_{pr}^m} \stackrel{[1]}{=} \left[\frac{\partial \phi_k^L}{\partial a_k^L} \cdot \frac{\partial a_k^L}{\partial w_{pr}^m} \right] (\mathbf{x}_i)$$

[1]: using the chain rule for derivation.

Now,

$$\frac{\partial \phi_k^L}{\partial w_{pr}^m} = \frac{\partial \phi_k^L}{\partial a_k^L} \cdot \frac{\partial a_k^L}{\partial w_{pr}^m} \quad (2)$$

$$\frac{\partial a_k^L}{\partial w_{pr}^m} = \sum_{j=0}^{N_{l-1}} w_{jk}^L \frac{\partial \phi_j^{L-1}}{\partial w_{pr}^m} \quad (3)$$

$$\Rightarrow \frac{\partial \phi_k^L}{\partial w_{pr}^m} = \frac{\partial \phi_k^L}{\partial a_k^L} \cdot \sum_{j=0}^{N_{l-1}} w_{jk}^L \frac{\partial \phi_j^{L-1}}{\partial w_{pr}^m}$$

With the following particular cases:

$$\frac{\partial a_r^m}{\partial w_{pr}^m} = \frac{\partial}{\partial w_{pr}^m} \left[\sum_{j=0}^{N_{m-1}} w_{jr}^m \phi_j^{m-1} \right] = \phi_p^{m-1}$$

and for $p = 0$,

$$\frac{\partial a_r^m}{\partial w_{0r}^m} = \phi_0^{m-1} = 1 \quad \forall m \in \llbracket L \rrbracket$$

or, if $k \neq r$ or $1 < l < m$, w_{pr}^m is not involved in the computation of a_k^L and

$$\frac{\partial a_k^L}{\partial w_{pr}^m} = \frac{\partial}{\partial w_{pr}^m} \left[\sum_{j=0}^{N_{l-1}} w_{jk}^L \phi_j^{L-1} \right] = 0 \quad (4)$$

finally, if $m = 1$ and $p \in \llbracket D \rrbracket$,

$$\frac{\partial a_r^1}{\partial w_{pr}^1}(\mathbf{x}_i) = \frac{\partial}{\partial w_{pr}^1} \left[\sum_{j=0}^{N_0} w_{jr}^1 \phi_j^0 \right] (\mathbf{x}_i) = \phi_p^0(\mathbf{x}_i) = \mathbf{x}_i[p]$$

This set of equations (in particular (3) and (2)) creates a recursive chain of computation for derivatives, starting from layer L backpropagating values towards the first (input) layer. Note that equation (4) confirms that the backtracking of derivation of $\frac{\partial a_k^L}{\partial w_{pr}^m}$ concern only paths in the network containing w_{pr}^m and stopping at layer m (equation (4)). We can therefore define the **recursive** algorithm 2 where the **highlighted** functions ensure the recursive backpropagation.

Clearly, the weights $\{w_{jk}^l\}_{jkl}$ may be stored in 3D arrays. Similarly, the values for $\{a_k^l(\mathbf{x}_i)\}_{ikl}$ and $\{\phi_k^l(\mathbf{x}_i)\}_{ikl}$ may be stored using matrices (for every $\mathbf{x}_i$) or 3D arrays. Algebraic operations (addition, multiplication and more) can be defined over these so-called “tensors” and accelerated using GPUs (or TPUs).

Algorithm 2 The **recursive** backpropagation algorithm

```

1: function  $\frac{\partial \phi_k^l}{\partial w_{pr}^m}(\mathbf{x}_i)$  ▷ Use chain rule since  $\phi_k^l = \varphi_k^l(a_k^l)$ 
2:   return  $\frac{\partial \phi_k^l}{\partial a_k^l} \cdot \frac{\partial a_k^l}{\partial w_{pr}^m}(\mathbf{x}_i)$ 
3:
4: function  $\frac{\partial \phi_k^l}{\partial a_k^l}(\mathbf{x}_i)$  ▷ Activation function
5:   Depends on  $l, k$  and the definition of  $\varphi_k^l(z)$ 
6:   e.g if  $\varphi_k^l(z) = z$  return 1
7:   e.g if  $\varphi_k^l(z) = z^2$  return  $2a_k^l(\mathbf{x}_i)$ 
8:   e.g if  $\varphi_k^l(z) = \frac{1}{1+e^{-z}}$  return  $\varphi(a_k^l(\mathbf{x}_i))(1 - \varphi(a_k^l(\mathbf{x}_i)))$ 
9:
10: function  $\frac{\partial a_k^l}{\partial w_{pr}^m}(\mathbf{x}_i)$  ▷ Activation
11:   if  $l > m$  then return  $\sum_{j=0}^{N_{l-1}} w_{jk}^l \frac{\partial \phi_j^{l-1}}{\partial w_{pr}^m}(\mathbf{x}_i)$ 
12:   if  $l = m$  then
13:     if  $k = r$  then return  $\phi_p^{m-1}(\mathbf{x}_i)$ 
14:     else return 0
15:   if  $l = 1$  then return  $\mathbf{x}_{i[p]}$ 
16:
17: function  $\frac{\partial \mathcal{L}}{\partial w_{jk}^l}(\mathbf{x}_i)$ 
18:   return  $\sum_{r=1}^d \frac{\partial \phi_r^l}{\partial w_{jk}^l}(\mathbf{x}_i) (\phi_r^l(\mathbf{x}_i) - \mathbf{y}_{i[k]})$ 
19:
20: procedure BACKPROPAGATION( $\mathcal{X}, \mathcal{Y}$ )
21:    $w_{jk}^{l,(0)} \leftarrow \text{random}() \ \forall j \in \llbracket N_{l-1} \rrbracket, k \in \llbracket N_l \rrbracket, l \in \llbracket L \rrbracket$  ▷ Initialize random weights
22:   for iteration  $t = 1$  until  $t = T$  do ▷ or until convergence
23:     for  $i \in \llbracket N \rrbracket$  do
24:       Forward pass computes and stores values for  $a_k^l(\mathbf{x}_i)$  and  $\phi_k^l(\mathbf{x}_i) \ \forall k, l$  ▷ Forward pass for  $\mathbf{x}_i$ 
25:       for  $j \in \llbracket N_{l-1} \rrbracket, k \in \llbracket N_l \rrbracket, l \in \llbracket L \rrbracket$  do
26:          $w_{jk}^{l,(t)} \leftarrow w_{jk}^{l,(t-1)} - \frac{1}{N} \eta \frac{\partial \mathcal{L}}{\partial w_{jk}^l}(\mathbf{x}_i)$  ▷  $\eta > 0$  is the gradient step (learning rate)

```

1.3 Example

1.3.1 Quadratic regression

Given $\mathcal{X} \times \mathcal{Y}$, a quadratic regression considers the squared loss and a learner of the quadratic family:

$$\phi_{\theta}(\mathbf{x}_i)_{[k]} = (\boldsymbol{\omega}_{2k}^T \mathbf{x}_i)^2 + \boldsymbol{\omega}_{1k}^T \mathbf{x}_i + \boldsymbol{\omega}_{0k}^T \mathbf{1}_D \quad (5)$$

ϕ_{θ} may be represented with a 3-layer feedforward network with $N_1 = 2$, $\varphi_1^1(z) = z$, $\varphi_2^1(z) = z^2$, $\varphi_2^2(z) = z$, as represented in figure 2.

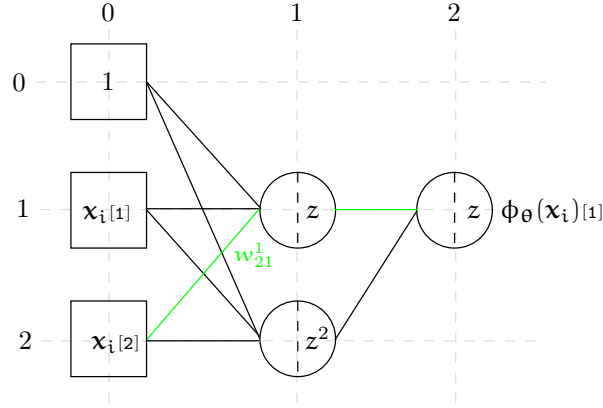


Figure 2: A neural network performing quadratic regression ($D = 2$, $d = 1$). Square nodes simply output their values. The green path shows the chain rule used for computing $\frac{\partial \mathcal{L}}{\partial w_{21}^1}$.

Following our notation, we obtain for the case $D = 2$, $d = 1$:

$$\phi_{\theta}(\mathbf{x}_i) = w_{11}^2 \overbrace{(w_{01}^1 + w_{11}^1 \mathbf{x}_i[1] + w_{21}^1 \mathbf{x}_i[2])}^{a_1^1} + w_{21}^2 \overbrace{(w_{02}^1 + w_{12}^1 \mathbf{x}_i[1] + w_{22}^1 \mathbf{x}_i[2])}^{a_2^1} \quad (6)$$

$$= w_{11}^2 \varphi_1^1(a_1^1) + w_{21}^2 \varphi_2^1(a_2^1) \quad (7)$$

Considering a squared loss, one obtains

$$\frac{\partial \mathcal{L}}{\partial w_{jk}} = \frac{1}{N} \sum_{i=1}^N \frac{\partial \phi_{\theta}}{\partial w_{jk}^1}(\mathbf{x}_i) (\phi_{\theta}(\mathbf{x}_i) - y_i)$$

and

$$\begin{aligned} \frac{\partial \phi_{\theta}}{\partial w_{01}^1}(\mathbf{x}_i) &= w_{11}^2 & \frac{\partial \phi_{\theta}}{\partial w_{11}^1}(\mathbf{x}_i) &= w_{11}^2 \mathbf{x}_i[1] & \frac{\partial \phi_{\theta}}{\partial w_{21}^1}(\mathbf{x}_i) &= w_{11}^2 \mathbf{x}_i[2] \\ \frac{\partial \phi_{\theta}}{\partial w_{02}^1}(\mathbf{x}_i) &= 2w_{21}^2 (w_{02}^1 + w_{12}^1 \mathbf{x}_i[1] + w_{22}^1 \mathbf{x}_i[2]) & \frac{\partial \phi_{\theta}}{\partial w_{12}^1}(\mathbf{x}_i) &= 2w_{21}^2 \mathbf{x}_i[1] (w_{02}^1 + w_{12}^1 \mathbf{x}_i[1] + w_{22}^1 \mathbf{x}_i[2]) \\ \frac{\partial \phi_{\theta}}{\partial w_{22}^1}(\mathbf{x}_i) &= 2w_{21}^2 \mathbf{x}_i[2] (w_{02}^1 + w_{12}^1 \mathbf{x}_i[1] + w_{22}^1 \mathbf{x}_i[2]) & \frac{\partial \phi_{\theta}}{\partial w_{11}^2}(\mathbf{x}_i) &= w_{01}^1 + w_{11}^1 \mathbf{x}_i[1] + w_{21}^1 \mathbf{x}_i[2] \\ \frac{\partial \phi_{\theta}}{\partial w_{21}^2}(\mathbf{x}_i) &= (w_{02}^1 + w_{12}^1 \mathbf{x}_i[1] + w_{22}^1 \mathbf{x}_i[2])^2 \end{aligned}$$

where one can see that the development of the chain rule follows paths according to equations derived in section 1.2. For example (see figure 2):

$$\frac{\partial \phi_{\theta}}{\partial w_{21}^1}(\mathbf{x}_i) = \frac{\partial \phi_1^2}{\partial a_1^2} \cdot \frac{\partial a_1^2}{\partial \phi_1^1} \cdot \frac{\partial \phi_1^1}{\partial a_1^1} \cdot \frac{\partial a_1^1}{\partial w_{21}^1}(\mathbf{x}_i) = 1 \cdot w_{11}^2 \cdot 1 \cdot \mathbf{x}_i[2]$$

These are the only non-zero terms from the chain

$$\frac{\partial \phi_{\theta}}{\partial w_{21}^1}(\mathbf{x}_i) = \frac{\partial \phi_1^2}{\partial a_1^2} \cdot \sum_{j=0}^{N_1} \left[\frac{\partial a_1^2}{\partial \phi_j^1} \cdot \frac{\partial \phi_j^1}{\partial a_j^1} \cdot \frac{\partial a_j^1}{\partial w_{21}^1} \right] (\mathbf{x}_i)$$

since $\frac{\partial a_2^1}{\partial w_{21}^1} = 0$ by equation (4).

Note. Developing equation (5) for this case, one can see that the neural model is over-specified: many combinations of weights lead to the same solution. To fix a unique solution, one can e.g fix and not update (clamp) w_{02}^1 , w_{11}^2 and w_{21}^2 .