

Applied Machine Learning (AML)

Neural Networks

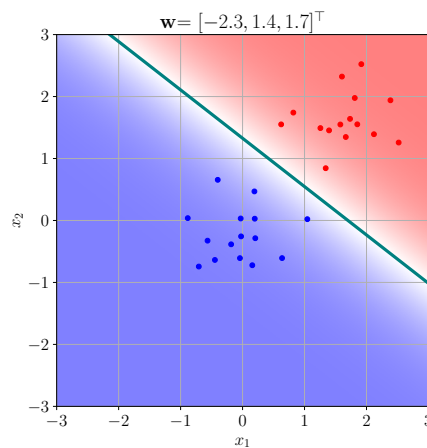
Oisín Mac Aodha • Siddharth N.

Introduction to Neural Networks

Recap - Linear Classifiers

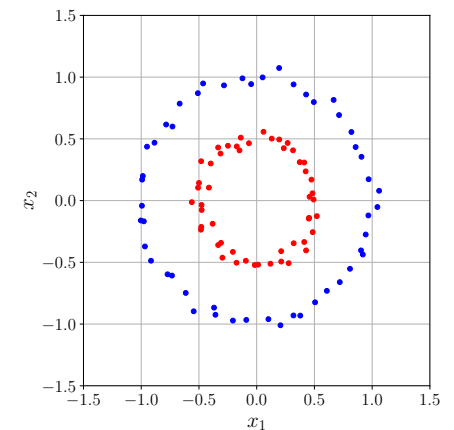
- Logistic regression =
linear weights + logistic function

$$p(y = 1|x) = \sigma(w^T x + b)$$



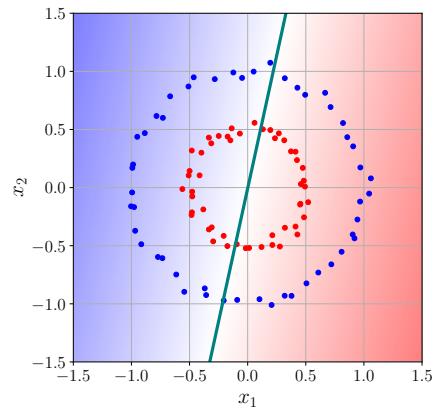
Classifying Non-Linear Data

- There is no linear classifier that will separate the data on the right given these 2D input features.



Classifying Non-Linear Data

- There is no linear classifier that will separate the data on the right given these 2D input features.
- In order to classify it, we can:
 - Use an alternative classifier that can generate *non-linear* decision boundaries.
 - *Transform* our input features so that they are linearly separable.



3

Neural Networks

- The performance of the classification methods we have explored depend on the input features, i.e. having a good ‘representation’ of the problem.
- If we do not have good features, many methods will not be effective, e.g. linear classifiers.
- What if we could *learn* the features from the input data?
- This is what **neural networks** attempt to do.
- We can think of them as a linear method, where we *learn* the features.

4

Neural Network Intuition

- Neural networks allow us to learn more effective features from the raw input data.
- We can think of them as functions that transform our input into something more useful for our task of interest.



5

Biological Neurons

- Neural networks are motivated by a *weak* analogy to the human brain, hence the name **artificial neural networks**.
- **Neurons** (also known as nerve cells) are electrically excitable cells in the nervous system that process and transmit information.
- Neurons are the core components of the brain, spinal cord, and nerves of vertebrates.

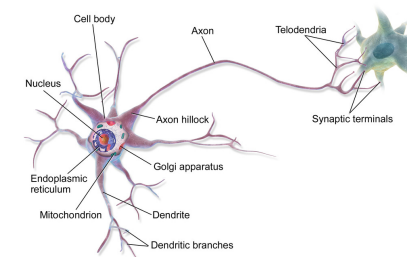
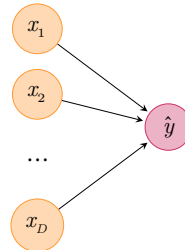


Image by BruceBlaus, CC BY 3.0

6

Artificial Neurons

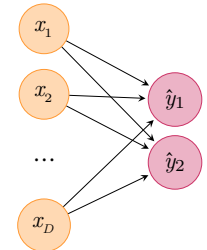
- Each *artificial* neuron is a linear weight vector with a non-linear activation function.
- We compute a neuron's activation as $\hat{y} = g(\mathbf{w}^\top \mathbf{x} + b)$
- Here, $g()$ is a *non-linear activation* function.
- If we used the logistic function this would just be logistic regression.



Note, we are not depicting the bias term b here.

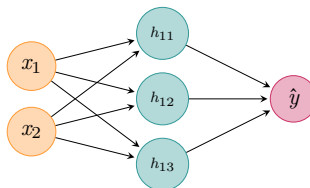
Artificial Neurons

- We can have multiple neurons
 $\hat{y}_1 = g(\mathbf{w}_1^\top \mathbf{x} + b_1)$
 $\hat{y}_2 = g(\mathbf{w}_2^\top \mathbf{x} + b_2)$
- We can present this using a weight *matrix* $\mathbf{W}$ and bias vector $\mathbf{b}$
 $\hat{\mathbf{y}} = g(\mathbf{W}\mathbf{x} + \mathbf{b})$



Single Layer Network

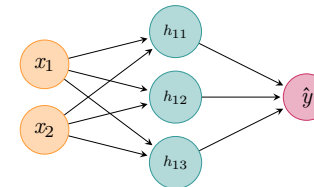
- We can connect multiple neurons (i.e. units) together into a directed acyclic graph.
- This results a **feed-forward neural network**.
- One of the simplest neural networks is a **single layer neural network**.



Single Layer Network

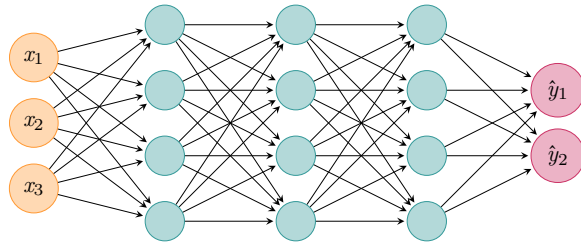
- In a single layer network, we have **input** units, **hidden** units, and **output** units.
- We can represent this function as

$$\hat{y} = g_2(\mathbf{w}_2^\top g_1(\mathbf{W}_1 \mathbf{x} + \mathbf{b}_1) + b_2)$$



Multilayer Neural Network

- Individual units in a network are grouped together into **layers**.
- We can stack multiple layers to form a **multilayer network**, i.e. a multilayer perceptron (MLP).
- Here we see a **fully connected network** with three input features, three hidden layers with four hidden units in each, and two output units.



Non-Linear Activation Functions

- Recall our expression for the single layer neural network
 $\hat{y} = g_2(w_2^T g_1(W_1 x + b_1) + b_2)$
- One might ask why do we need a non-linear activation function $g()$ in $g(Wx + b)$?
- Any sequence of linear layers can be equivalently represented with a single linear layer, i.e.

$$y = W_1 W_2 W_3 x \\ \triangleq W' x$$

Non-Linear Activation Functions - Choices

- **Sigmoid** i.e. Logistic

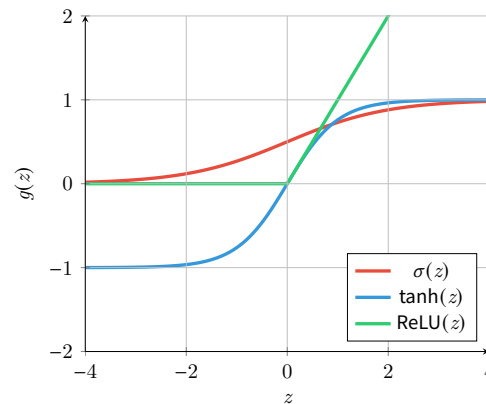
$$\sigma(z) = \frac{1}{1 + \exp(-z)}$$

- **Hyperbolic tangent**

$$\tanh(z) = \frac{\exp(2z) - 1}{\exp(2z) + 1}$$

- **Rectified linear unit**

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



Neural Networks - Expressive Power

- Feed-forward neural nets with non-linear activation functions are **universal function approximators**, i.e. they can approximate any function arbitrarily well.
- In practice, you may need an exponentially large network.
- If you can learn any function, this can just result in overfitting.

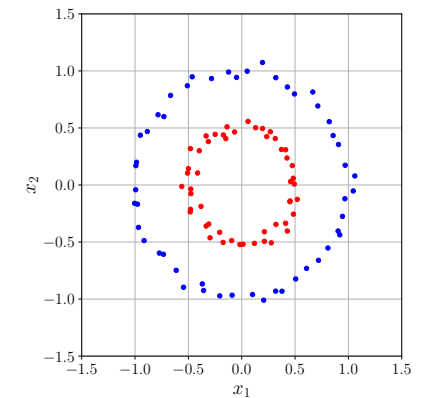
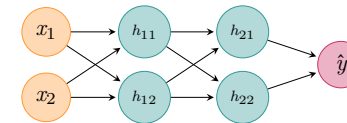
Importance of Network Depth

- A fully connected neural network with *one hidden layer* is a universal function approximator.
- This means it can model any sufficiently smooth function given a suitable number of hidden units.
- However, both experimental and theoretical work have shown that *deeper* neural networks (i.e. ones with more layers) are more effective than *shallow* ones.
- In deeper networks, later layers can leverage the features learned by earlier ones.

Non-Linear Classification with Neural Networks

- Here we revisit the non-linear binary classification problem from earlier.
- We will use the following neural network with **two** hidden layers:

$$\hat{y} = g_3(w_3^T g_2(W_2 g_1(W_1 x + b_1) + b_2) + b_3) \\ = f(x) = f_3(f_2(f_1(x)))$$



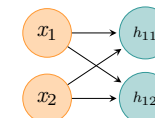
Neural Network Example - Input

- Input x



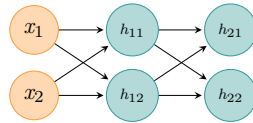
Neural Network Example - First Hidden Layer

- Input x
- First hidden layer
 $h_1 = g_1(W_1 x + b_1)$



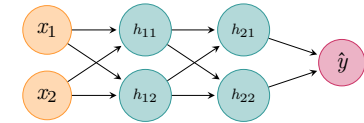
Neural Network Example - Second Hidden Layer

- Input x
- First hidden layer
 $h_1 = g_1(W_1 x + b_1)$
- Second hidden layer
 $h_2 = g_2(W_2 h_1 + b_2)$



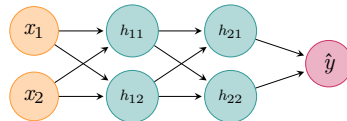
Neural Network Example - Output Layer

- Input x
- First hidden layer
 $h_1 = g_1(W_1 x + b_1)$
- Second hidden layer
 $h_2 = g_2(W_2 h_1 + b_2)$
- Output
 $\hat{y} = g_3(w_3^T h_2 + b_3)$



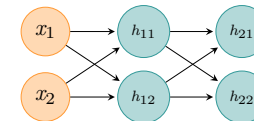
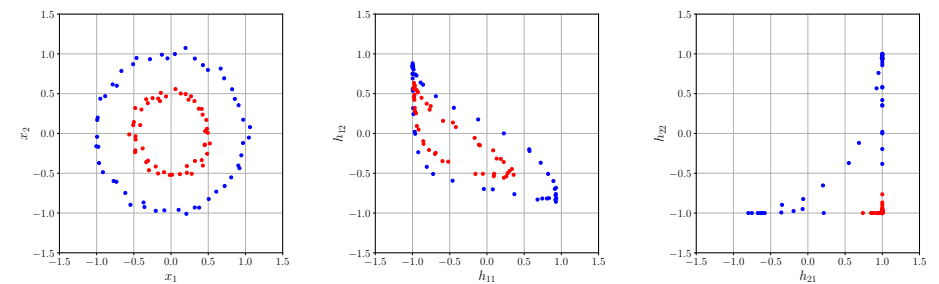
Neural Network Example - Output Layer

- Input x
- First hidden layer
 $h_1 = \tanh(W_1 x + b_1)$
- Second hidden layer
 $h_2 = \tanh(W_2 h_1 + b_2)$
- Output
 $\hat{y} = \sigma(w_3^T h_2 + b_3)$



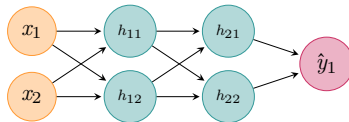
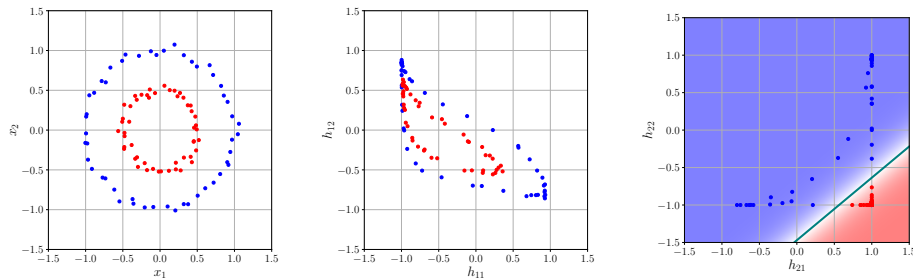
Transforming Features

Here we see the outputs from each layer of the network.



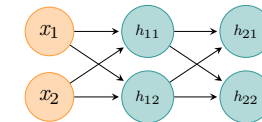
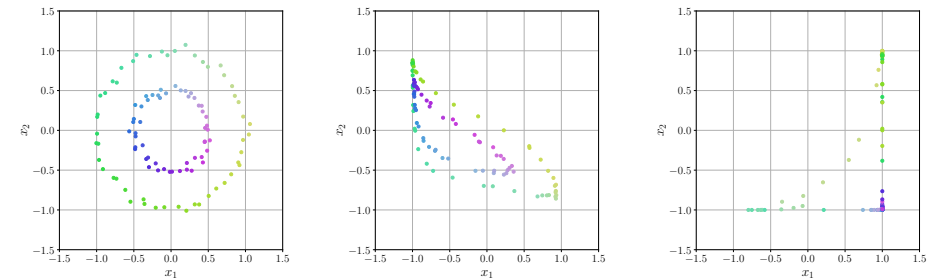
Transforming Features

Now we can use a linear classifier on the final hidden features.



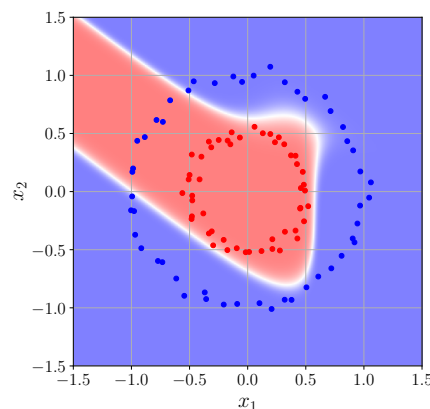
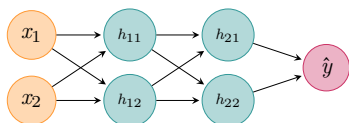
Transforming Features

This is the same as previous, but we have colour coded each individual instance.



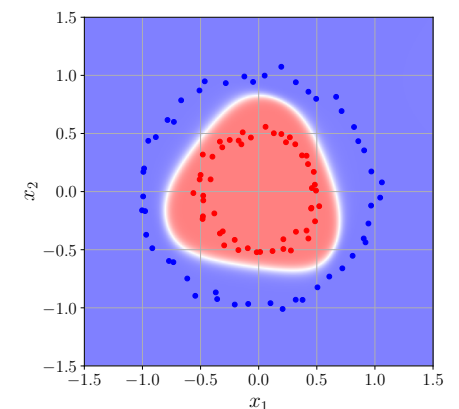
Final Neural Network Predictions

- On the right we see the final network predictions, colour-coded by predicted class.
- Note, in this example we defined a simple, and small, neural network for ease of visualisation.
- Our network did *not* successfully classify *all* the training data.



Final Neural Network Predictions

- With a minor change to the structure of the neural network, we can correctly classify the training data.
- In the example on the right we increased the number of hidden units in the first layer (from two to four).



Training Neural Networks

Training Neural Networks

- The task of training a neural network involves finding the best parameters (i.e. weights) for each unit.
- We will use a loss function $\mathcal{L}(\theta)$ to measure the disagreement between the model prediction $f(x)$ and the ground truth target y .
- During optimisation we will try to find the parameters that minimise the loss.
- We will take the gradient of the loss $\nabla_{\theta} \mathcal{L}$ and use gradient descent to update the parameters.
$$\theta \leftarrow \theta - \eta \cdot \nabla_{\theta} \mathcal{L}$$

Neural Network Parameters

- Recall that a multilayer neural network is a nested set of linear functions with non-linear activations.

$$f(x) = f_L(\dots f_2(f_1(x)))$$

- Each layer f_L has its own weight matrix W_L and bias vector b_L .
- The concatenation of these terms form the model weights that need to be learned.

$$\theta = (W_1, b_1, W_2, b_2, \dots, W_L, b_L)$$

Training Multilayer Neural Networks

- Hidden units make optimising the network weights more complicated as we do not have ground truth targets for them.
- Each hidden activity can affect many output units and can therefore have many separate effects on the error.

Backpropagation

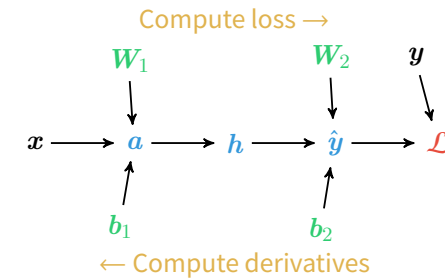
- There is a recursive algorithm for computing the derivatives. It uses the **chain rule** by storing some intermediate terms. This is called **backpropagation**.
- We make use of the layered structure of the network to compute the derivatives, heading backwards from the output layer to the inputs.

Backpropagation Algorithm

- Consists of two main steps:
 - A **forward pass**, in which we compute and store the values at all of the hidden units and the network output.
 - A **backward pass**, in which we calculate the derivatives of each weight, starting at the end of the network, and reusing the previous computation as we move towards the start.

Backpropagation

- We can visualise the computations using a **computation graph**.
- The nodes represent all the inputs and computed quantities, and the edges represent which nodes are computed directly as a function of which other nodes ¹.



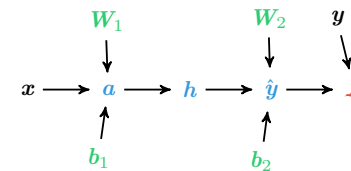
¹Example adapted from Ren and MacKay: CSC 411.

Chain Rule of Calculus

- Suppose we had two functions $u(x)$ and $v(x)$.
- Then $y = u(v(x))$ is a function of a function.
- The **chain rule** of calculus gives us a way to express the derivative of the composition of two differentiable functions $u(x)$ and $v(x)$ in terms of their derivatives.
- For example, if we substitute $s = v(x)$, thus $y = u(s)$, and

$$\frac{dy}{dx} = \frac{dy}{ds} \frac{ds}{dx}$$

Backpropagation - Forward Pass



Forward Pass

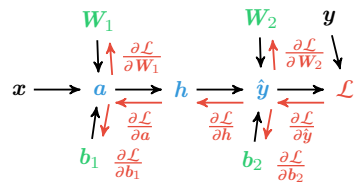
$$a = W_1 x + b_1$$

$$h = g(a)$$

$$\hat{y} = W_2 h + b_2$$

$$\mathcal{L} = \frac{1}{2} \|\hat{y} - y\|^2$$

Backpropagation - Backward Pass



Forward Pass

$$a = W_1 x + b_1$$

$$h = g(a)$$

$$\hat{y} = W_2 h + b_2$$

$$\mathcal{L} = \frac{1}{2} \|\hat{y} - y\|^2$$

Backward Pass

$$\begin{aligned} \frac{\partial \mathcal{L}}{\partial \hat{y}} &= (\hat{y} - y) \\ \frac{\partial \mathcal{L}}{\partial W_2} &= \frac{\partial \mathcal{L}}{\partial \hat{y}} \frac{\partial \hat{y}}{\partial W_2} = \frac{\partial \mathcal{L}}{\partial \hat{y}} h^T \\ \frac{\partial \mathcal{L}}{\partial b_2} &= \frac{\partial \mathcal{L}}{\partial \hat{y}} \frac{\partial \hat{y}}{\partial b_2} = \frac{\partial \mathcal{L}}{\partial \hat{y}} \\ \frac{\partial \mathcal{L}}{\partial h} &= \frac{\partial \mathcal{L}}{\partial \hat{y}} \frac{\partial \hat{y}}{\partial h} \\ \frac{\partial \mathcal{L}}{\partial a} &= \frac{\partial \mathcal{L}}{\partial h} \frac{\partial h}{\partial a} \\ \frac{\partial \mathcal{L}}{\partial W_1} &= \frac{\partial \mathcal{L}}{\partial a} \frac{\partial a}{\partial W_1} = \frac{\partial \mathcal{L}}{\partial a} x^T \\ \frac{\partial \mathcal{L}}{\partial b_1} &= \frac{\partial \mathcal{L}}{\partial a} \frac{\partial a}{\partial b_1} = \frac{\partial \mathcal{L}}{\partial a} \end{aligned}$$

34

Convergence of Neural Networks

- For *logistic regression*, the loss function is conveniently convex. A convex function has just one minimum.
- Multilayer *neural networks* are **non-convex**, and gradient descent may get stuck in local minima during training and never find the global optimum.
- In practice this is not necessarily an issue and we can still apply gradient-based methods and can obtain good solutions for many practical problems of interest.

35

Hyperparameters

Network Structure

- There are several elements of the network that you can change e.g.
 - The number of hidden layers.
 - The number of units in each hidden layer.
 - The type of non-linear activation function e.g. ReLU, sigmoid, ...

Training Schedule

- There also are several aspects of the training procedure that can be changed e.g.
 - The learning rate.
 - The type of optimiser e.g. standard gradient descent, ...
 - How the weights are initialised.
 - When to stop training.

Automatic Differentiation

- The **backpropagation algorithm**, which can be used to compute the gradient of a loss function applied to the output of the network wrt the parameters in each layer.
- This gradient can then be used with any gradient-based optimisation, e.g. gradient descent.
- Manually computing these gradients for anything but small toy problems is too time consuming.
- Instead, we can make use **automatic differentiation** (or **autodiff**). This is a set of automatic techniques to evaluate the derivative of a function.
- Many machine learning frameworks have autodiff functionality built in.

36

37

Automatic Differentiation Example

The following is an example of using autodiff for binary logistic regression.

```
1 import jax.numpy as jnp
2 from jax import grad, nn
3
4 # Define our loss function
5 def nll_loss(X, y, w):
6     pred = nn.sigmoid(X@w)
7     loss_pos = (y==1)*jnp.log(pred)
8     loss_neg = (y==0)*jnp.log(1.0 - pred)
9     loss = -(loss_pos + loss_neg).mean()
10    return loss
11
12
13 # Define our dataset, which has 3 instances
14 # We have already appended a 1.0 to each row of X
15 X = jnp.array([[1.0, 0.5, -0.35],
16               [1.0, -0.1, 0.1],
17               [1.0, -1.2, 1.0]])
18 y = jnp.array([0.0, 0.0, 1.0])
19
20 # This is our initial weight vector w
21 w = jnp.array([0.0, -1.0, 1.0])
22
```

```
23 # (i) Compute the gradient manually
24 # Here we use the derived expression
25 manual_grad = (nn.sigmoid(X@w) - y)@X
26 manual_grad *= (1.0/X.shape[0])
27 print('Manual gradient', jnp.round(manual_grad, 3))
28
29
30 # (ii) Compute the gradient automatically
31 # Evaluate the loss and compute the gradient
32 loss = nll_loss(X, y, w)
33 w_grad = grad(nll_loss, (2))(X, y, w)
34 print('Auto gradient ', jnp.round(w_grad, 3))
35
36
37 # We can take one step of gradient descent
38 learning_rate = 3.0
39 w_update = w - learning_rate*w_grad
40
```

Alternative Network Architectures

Images as Tensors

- We can represent images as **matrices**, where each entry stores the intensity value of a given pixel.



Issues with Fully Connected Neural Networks

- Fully connected networks with high-dimensional inputs have a lot of model weights.
- This results in a very large number of model weights that have to be learned.

- For example, if our input was an image of size 100×100 , this would require 10,000 weights for *each* hidden unit in the first layer.

Shift Invariance

- Fully connected networks are sensitive to the position of the signal of interest in an input image.

0	1	0	0
1	1	1	0
0	1	0	0
0	0	0	0

$$\sum_{i=1}^4 \sum_{j=1}^4 x_{ij} \times w_{ij} = 5$$

0
1
0
0
1
1
1
0
0
1
0
0
0
0
0
0
0
0

$\mathbf{x}_1$

Shift Invariance

- Fully connected networks are sensitive to the position of the signal of interest in an input image.

0	1	0	0
1	1	1	0
0	1	0	0
0	0	0	0

$$\sum_{i=1}^4 \sum_{j=1}^4 x_{ij} \times w_{ij} = 5$$

0
1
0
0
1
1
1
0
0
1
0
0
0
0
0
0
0
0

$\mathbf{x}_1$

0
1
0
0
1
1
1
0
0
1
0
0
0
0
0
0
0
0

$\mathbf{w}$

Shift Invariance

- Fully connected networks are sensitive to the position of the signal of interest in an input image.

0	1	0	0
1	1	1	0
0	1	0	0
0	0	0	0

$$\sum_{i=1}^4 \sum_{j=1}^4 x_{ij} \times w_{ij} = 5$$

0
1
0
0
1
1
1
0
0
1
0
0
0
0
0
0
0
0

$\mathbf{x}_1$

0
1
0
0
1
1
1
0
0
1
0
0
0
0
0
0
0
0

$\mathbf{w}$

0	0	1	0
0	1	1	1
0	0	1	0
0	0	0	0

$$\sum_{i=1}^4 \sum_{j=1}^4 x_{ij} \times w_{ij} = 2$$

0
1
0
0
1
1
1
0
0
1
0
0
0
0
0
0
0
0

$\mathbf{x}_2$

0
1
0
0
1
1
1
0
0
1
0
0
0
0
0
0
0
0

$\mathbf{w}$

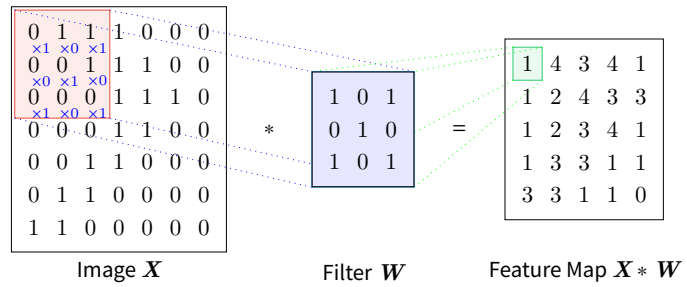
Convolutional Filters

- Constrain each hidden unit to extract features by **sharing** weights across the input.
- For an image X and $K \times K$ weight matrix W (i.e. a filter) we compute the outputs as

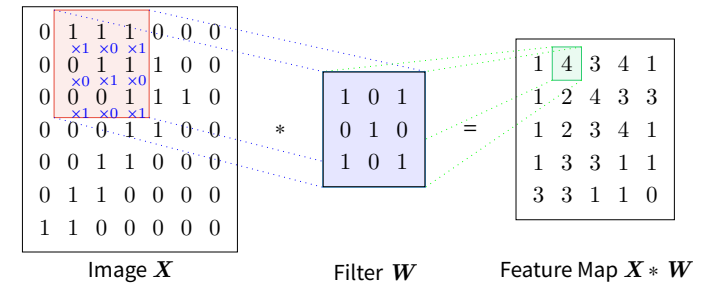
$$h_{ij} = g \left(\sum_{m=1}^K \sum_{n=1}^K w_{m,n} x_{i+m, j+n} + b \right)$$

- The output is a **feature map**, where each entry h_{ij} is the local response of the filter convolved with the image at that location.
- Multiple weight matrices can be used to produce multiple feature maps.

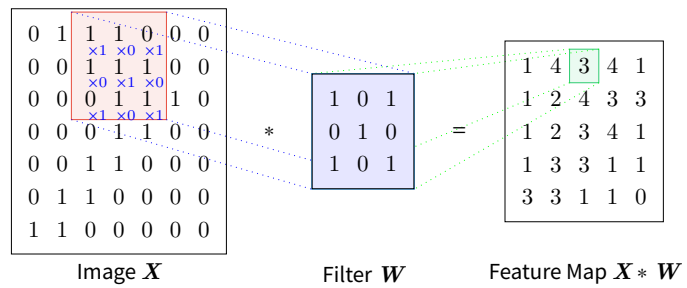
Convolution



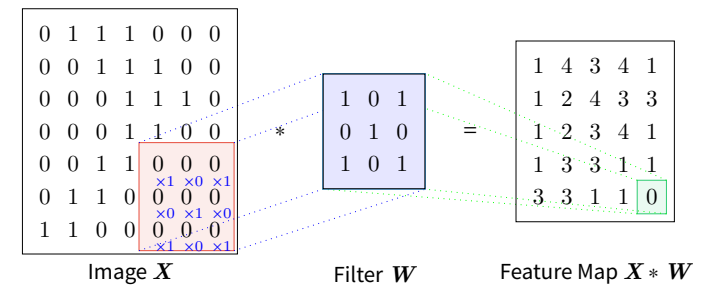
Convolution



Convolution

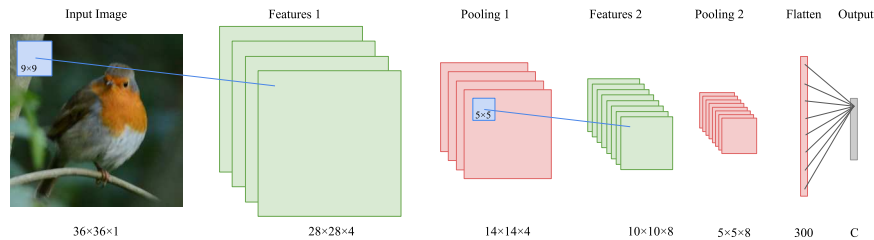


Convolution



Convolutional Neural Network - Example

- A Convolutional Neural Network (CNN) consists of *learnable* convolutional filters and *non-learnable* pooling layers.
- The pooling layers reduce the spatial dimensionality of the feature maps.
- For classification, at the output of the network, we have a fully connected layer which predicts one of C classes.

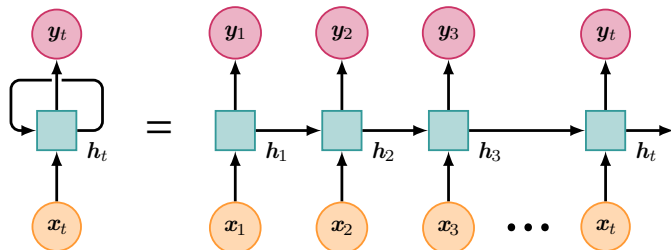


Recurrent Neural Networks

- A model for sequence data (e.g. time series).
- Different network architectures and recurrent units exist, e.g. long short-term-memories (LSTMs).
- In a RNN, each input is processed sequentially, one item at a time.
- Past information is retained through past hidden states.

Recurrent Neural Networks

- In RNNs, the outputs y_t are a function of the current input x_t and the previous hidden state h_{t-1} .



Transformers

- Alternative, and more recent, approach for modelling sequential data.
- Unlike RNNs, Transformers process the entire input all at once.
 - Thus training can be performed in parallel.
 - They are also less susceptible to 'forgetting' information from the past, i.e. better suited to capture *long-range* dependencies.
- Transformers have a special type of unit called a **self-attention** unit. This is used to compute similarity scores between inputs in the input sequence.
- They can also be applied to other data types, e.g. images.

Summary

- Artificial neural networks are a powerful non-linear modelling tool for classification and regression.
- They are *not* biologically plausible models.
- The output of the hidden units are a new representation of the original input data. This can be interpreted as learned features.
- Training makes use of the backpropagation algorithm to compute derivatives.
- Beyond standard fully connected networks, alternative architectures exist for learning from structured input data (e.g. images, audio, text, ...).