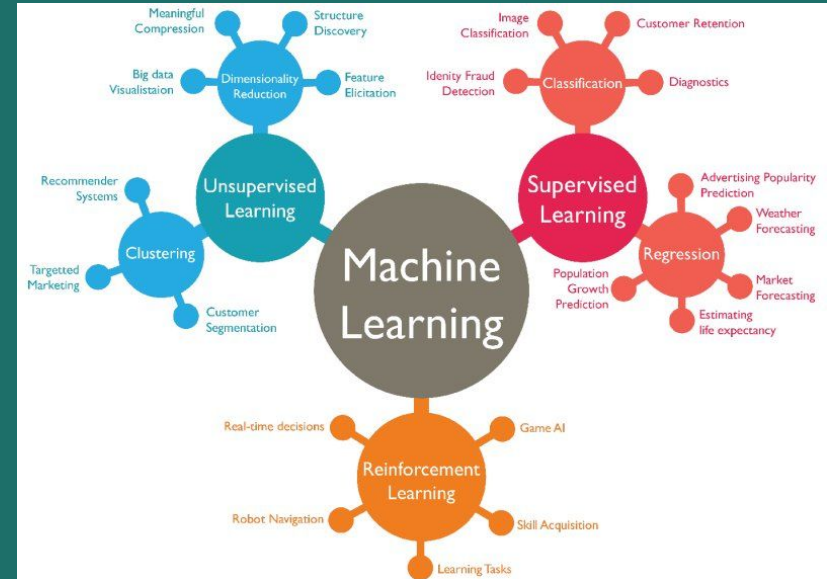


An Introduction to Artificial Neural Networks

Birds Eye View



What is Machine Learning?

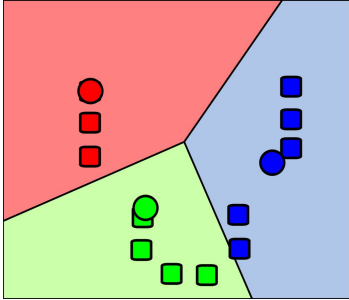
Field of study that gives computers the ability to learn **without being explicitly programmed**.

Defined by Arthur Samuel (1959)

Practically speaking:

- Techniques for fitting data.
- It is not restricted by a single or a set of simple mathematical expressions.

Two Types of machine learning

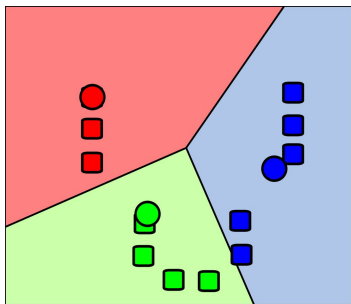


Unsupervised learning

looks for previously undetected patterns in a data set with no pre-existing labels and with a minimum of human supervision, e.g.

- principal component analysis
- clustering methods (k-means, self-organizing maps)
- autoencoders

Two Types of machine learning



Unsupervised learning

looks for previously undetected patterns in a data set with no pre-existing labels and with a minimum of human supervision, e.g.

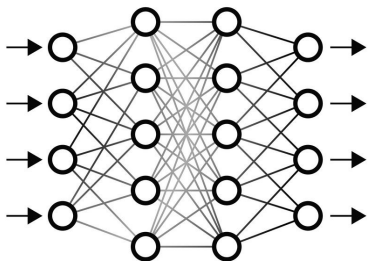
- principal component analysis
- clustering methods (k-means, self-organizing maps)
- autoencoders

Supervised learning



maps an input to an output based on example input-output pairs, e.g.

- regression
- decision trees and random forests
- support vector machines
- artificial neural networks



Reasons to use ML for Science

- Do it **better**
 - e.g. Convective parameterizations in models are not perfect, use ML to make them more accurate
- Do it **faster**
 - e.g. Radiation code in models is very slow (but we know the right answer) - use ML methods to speed things up
- Do **something new**
 - e.g. Go looking for non-linear relationships you didn't know were there

Example uses in Climate Science

- Making climate models better
 - Parameterizations
 - Speed-up computations
- Better predictions
 - Post-processing dynamical (physics-based) forecasts
 - Purely data-driven forecasts
- Sources of predictability
 - process-based understanding
 - identifying dynamical model errors/biases
- Feature identification
 - Counting clouds, labeling cloud types in images
 - Counting/identifying extreme events
- Summarizing a lot of data
 - Dimension reduction
 - Cluster/group behaviour

The Basic Idea

Find the **best function** $F(X)$ for fitting the input to the output

VARIABLE(S)

$F(X)$

learned non-linear
relationship

VARIABLE(S)

Linear Regression

Imagine you have a paired data set **x** [input] and **y** [output].

If **x** and **y** are linearly related, their *modeled* linear relationship would look something like this...

$$\begin{array}{c} \text{prediction} \\ \mathbf{y}_{\square} \end{array} = \mathbf{a}_0 + \mathbf{a}_1 \begin{array}{c} \mathbf{x}_{\square} \\ \text{y-intercept} \quad \text{slope} \end{array}$$

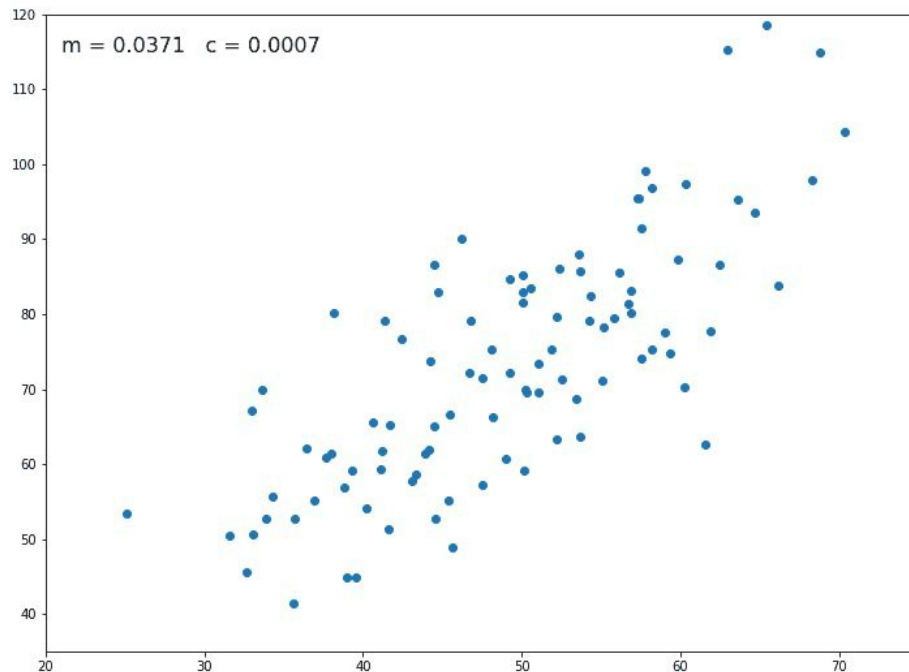
Linear Regression

$$\text{prediction } \hat{y}_i = a_0 + a_1 x_i$$

y-intercept slope

Find a_0 & a_1 to minimize MSE

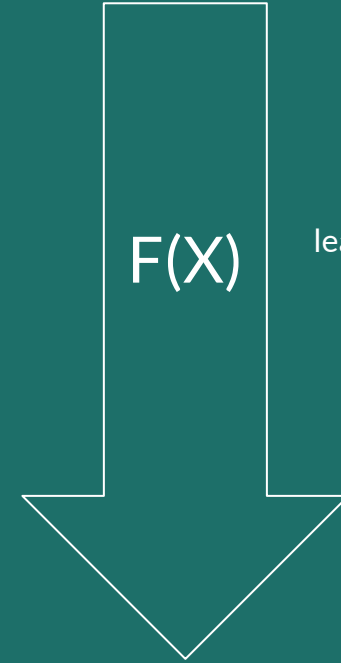
$$\text{MSE} = \frac{1}{n} \sum (\hat{y}_i - y_i)^2$$



We have an equation for a_0 and a_1 ... but what about when we don't?

Artificial Neural Network (ANN)

VARIABLE(S)



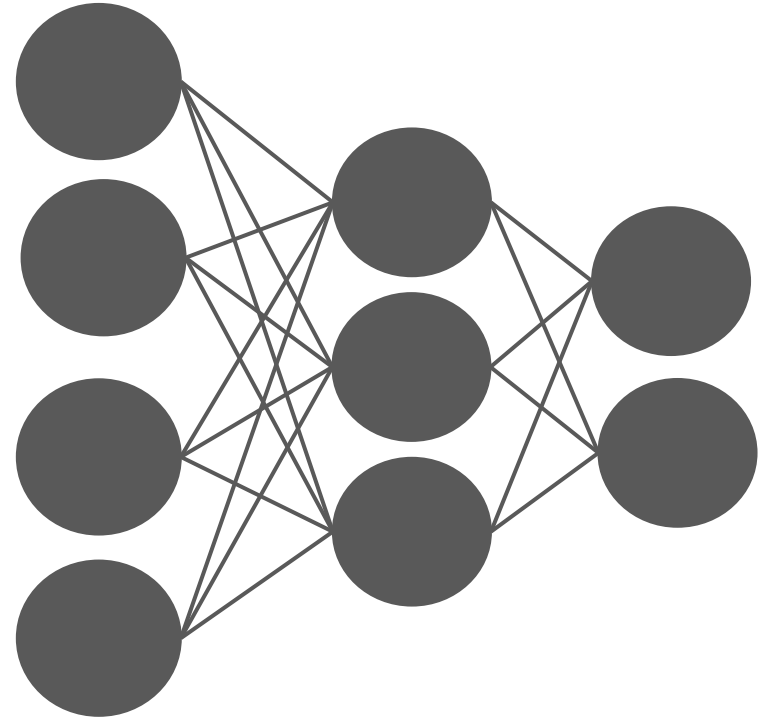
$F(X)$

learned non-linear
relationship

VARIABLE(S)

Linear regression as an ANN

$$\hat{y}_x = a_0 + a_1 x$$



Linear regression as an ANN

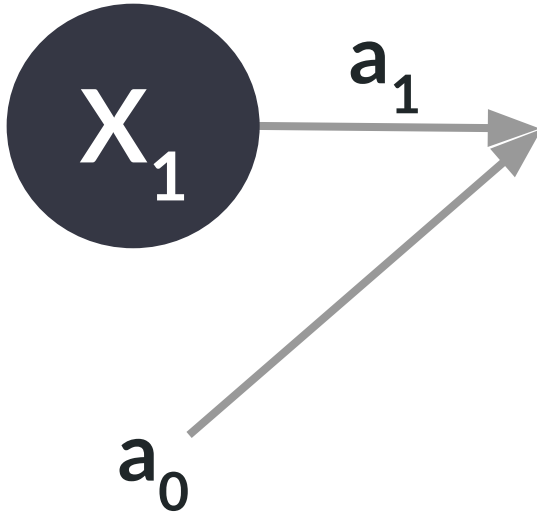
$$\hat{y} = a_0 + a_1 x$$



$$= X_1$$

Linear regression as an ANN

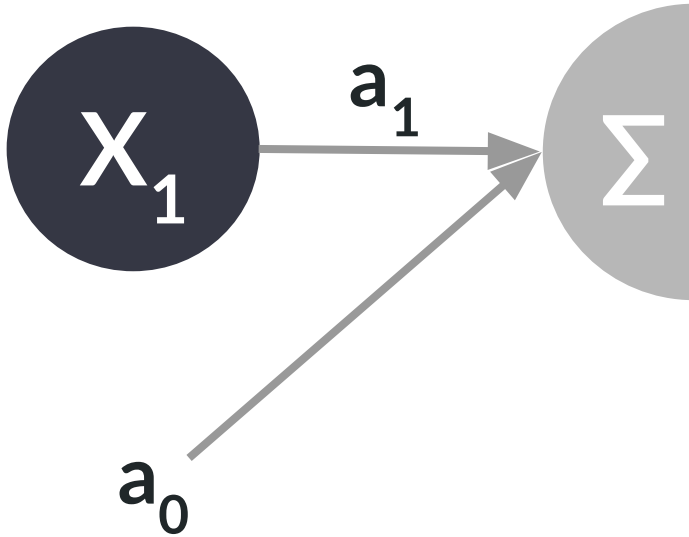
$$\hat{y} = a_0 + a_1 x$$



$$= x_1 a_1 + a_0$$

Linear regression as an ANN

$$\hat{y} = a_0 + a_1 x$$

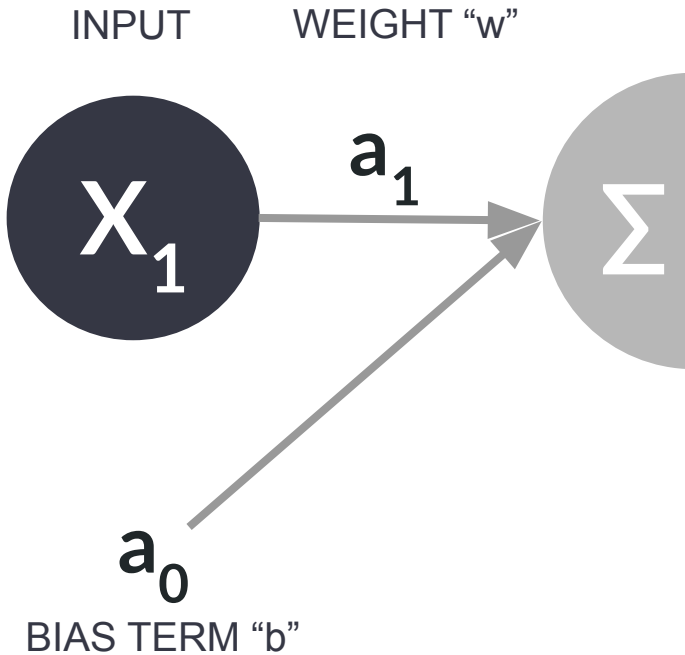


$$= x_1 a_1 + a_0$$

Linear regression as an ANN

$$\hat{y} = a_0 + a_1 x$$

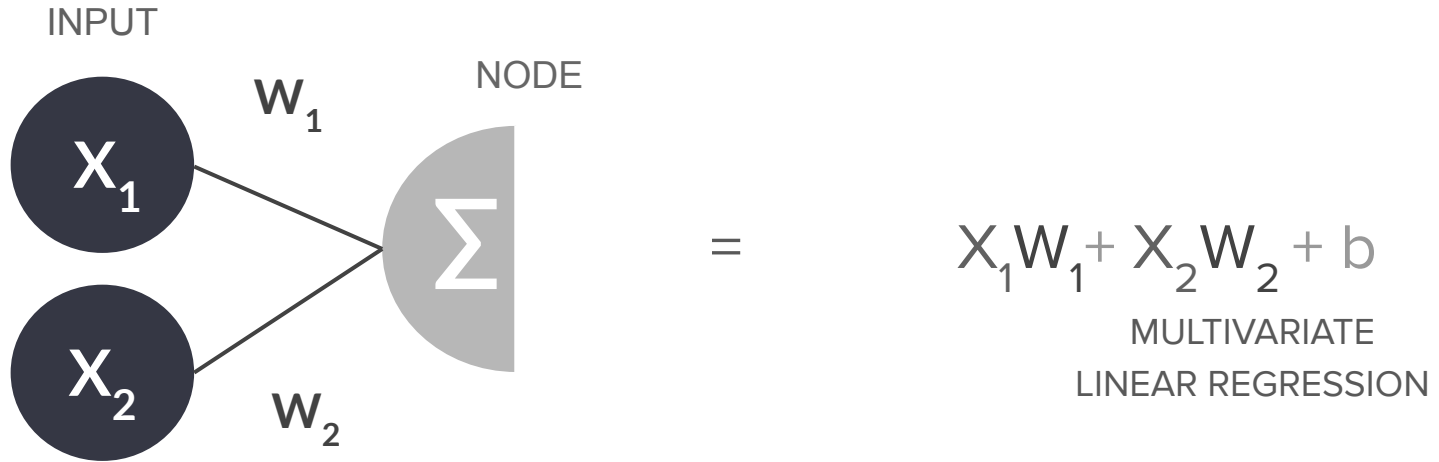
ARCHITECTURE



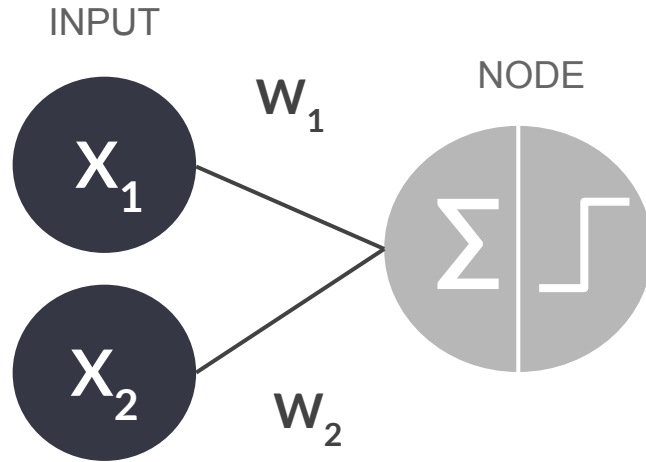
OUTPUT/PREDICTION

$$= X_1 a_1 + a_0$$

Artificial neural networks



Artificial neural networks



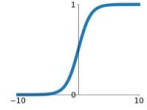
$$= f_{\text{activation}}(x_1 w_1 + x_2 w_2 + b)$$

linear regression with non-linear mapping by an “activation function”

Activation Functions

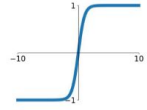
Sigmoid

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



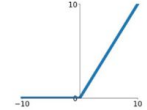
tanh

$$\tanh(x)$$

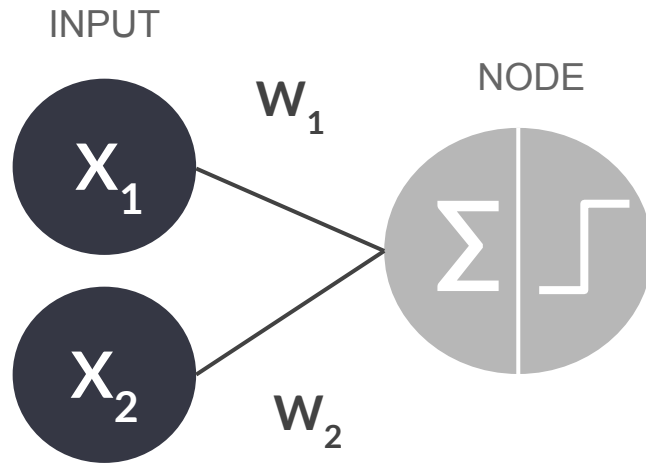


ReLU

$$\max(0, x)$$



Artificial neural networks



$$= f_{\text{activation}}(X_1 W_1 + X_2 W_2 + b)$$

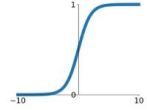
Activation function determines, if information is moving forward from that specific node.

This is the step that allows for nonlinearity in these algorithms, without activation all we would be doing is linear algebra!

Activation Functions

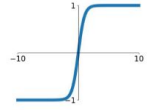
Sigmoid

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



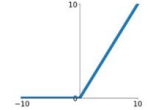
tanh

$$\tanh(x)$$



ReLU

$$\max(0, x)$$

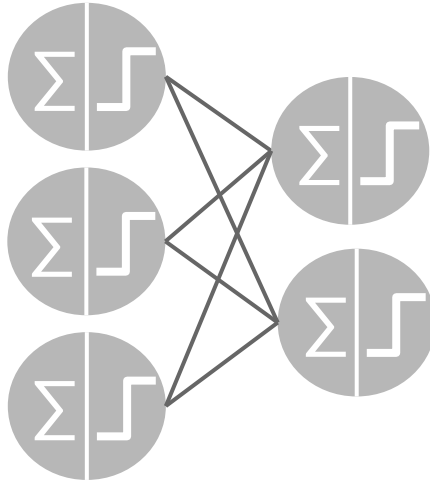


INPUT

X_1

X_2

X_3



HIDDEN LAYERS

PREDICTIONS



Training an ANN is determining the weights “w” and bias “b”

Once trained, you have an array of weights and biases which can be used for prediction on *new* data



Terminology

units/nodes/neurons: computational unit connected by weights to all other nodes

Weights/Bias: strength of the connection between nodes/shifting term

Activation Function: function that maps the output of a node; determines whether a node is “activated”

(Hidden) Layer: layer of nodes (not in input or output layer)

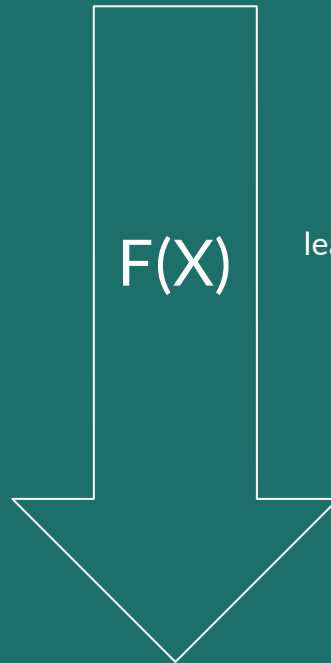
Deep: training a NN with **3+** layers

The Basic Idea

Find the **best function** $F(X)$ for fitting the input to the output

... but what do we mean by **best**?

VARIABLE(S)



$F(X)$

learned non-linear
relationship

VARIABLE(S)

Splitting Your Data

Training Data

- to build the ML model (i.e. determine the parameters)

Validation Data

- to evaluate the ML model on unseen data
- to optimize parameters that are choices by the user

Testing Data

- to evaluate the final ML model on data not used for training or evaluation

Training Data

Validation Data

Testing Data

A
L
L

T
H
E

D
A
T
A

How to Split Data

TRAINING data
VALIDATION data
TEST data



Randomly Split



Chronologically Split

How to Split Data

Climate data is often autocorrelated, so we split data **chronologically**

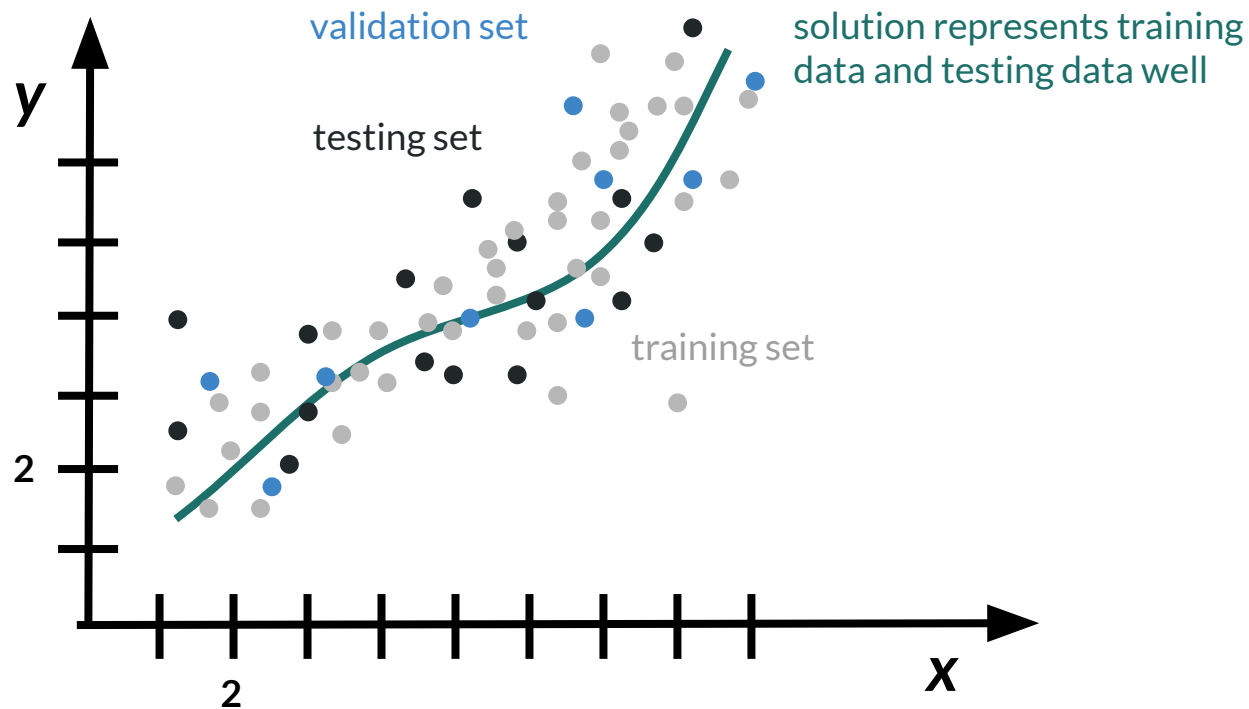
Example: Full dataset from 1900-2000

- Training: 1900-1980
- Validation: 1981-1990
- Testing: 1991- 2000



Chronologically Split

Machine Learning Approach



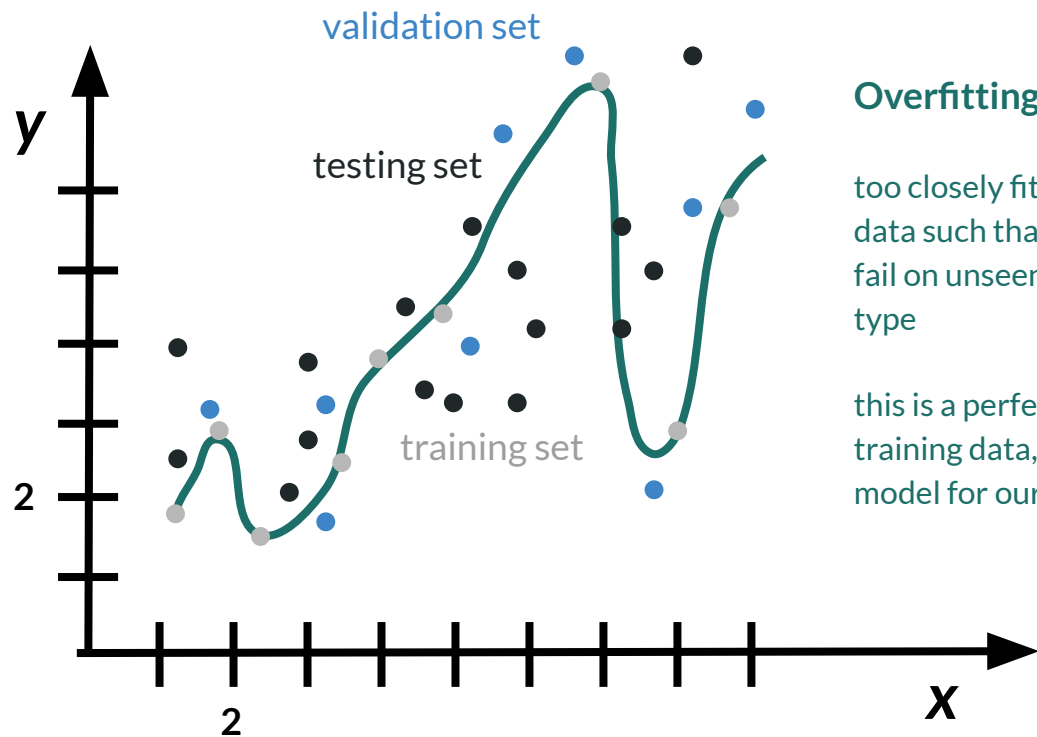
Training Data

Validation Data

Testing Data

ALL THE DATA

Machine Learning Approach



Overfitting

too closely fitting the training data such that the model will fail on unseen data of the same type

this is a perfect model for the training data, but a very poor model for our testing data

Training Data

Validation Data

Testing Data

ALL THE DATA

How do we quantify a good prediction?

cost/loss functions

$$\text{MSE} = \frac{1}{n} \sum (\hat{y}_i - y_i)^2$$
$$\text{loss/cost} = \frac{1}{n} \sum (\text{prediction} - \text{actual})^2$$

Find a quantity you want to minimize (the “loss”) to help determine the weights

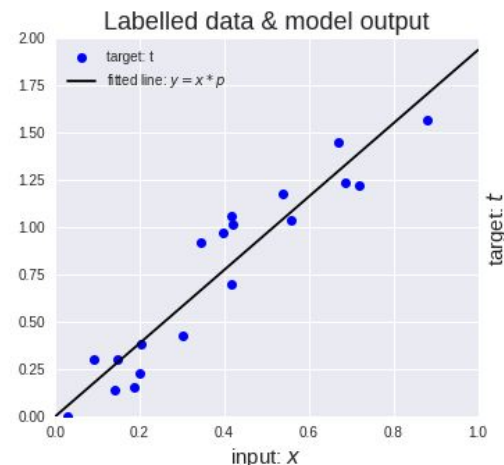
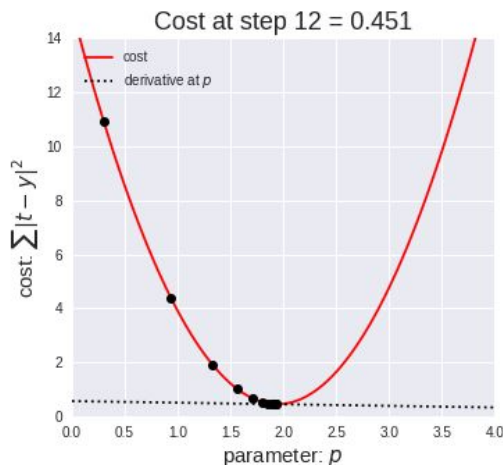
How do we find the minima in real life?

Let's start with an easy linear example

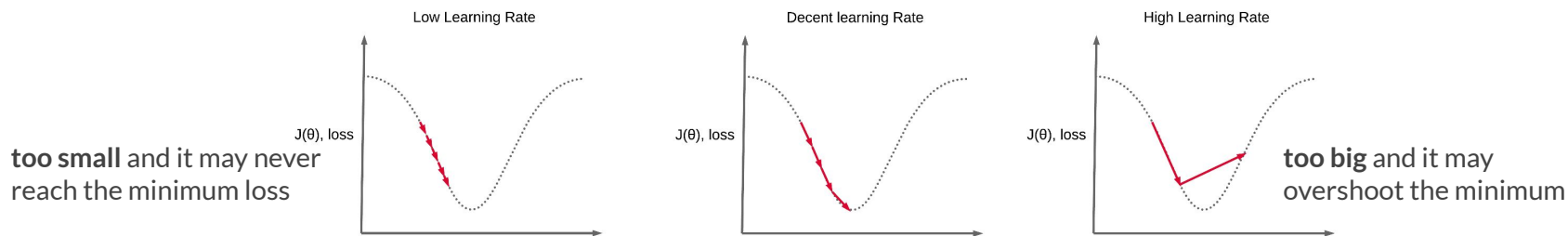
$$\tilde{y}_t = a_0 + a_1 x_t$$

An example loss function (MSE)

$$J = \frac{1}{N} \sum_{i=1}^N (\tilde{y}_i - y_i)^2$$



Gradient Descent



Gradient: gives you the direction of greatest/fastest increase

Descent: let's follow the greatest *decrease*...

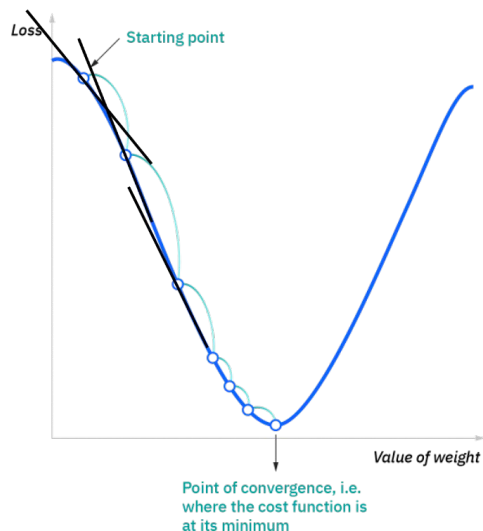
Gradient Descent: a method used to minimize the loss/cost function; algorithm to update the weights to follow greatest loss decrease

Learning Rate: how “quickly” the neural network tries to minimize the loss function

Gradient descent mathematically

$$J = \frac{1}{N} \sum_{i=1}^N (\tilde{y}_i - y_i)^2$$

$$\tilde{y}_t = a_0 + a_1 x_t$$

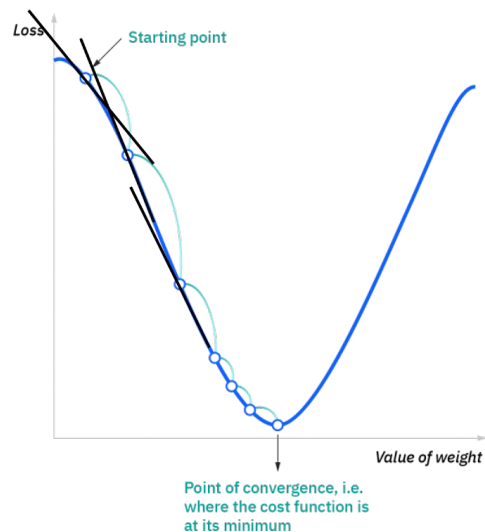


- MSE (our loss function) depends on the linear function and its weights: a_0 and a_1
- The **gradient** is calculated as partial derivatives of the loss function with relation to the weights

$$J = J(a_0, a_1)$$

$$\nabla J_t = \begin{bmatrix} \frac{\partial J}{\partial a_0} \\ \frac{\partial J}{\partial a_1} \end{bmatrix}$$

Partial derivatives for MSE (linear regression)



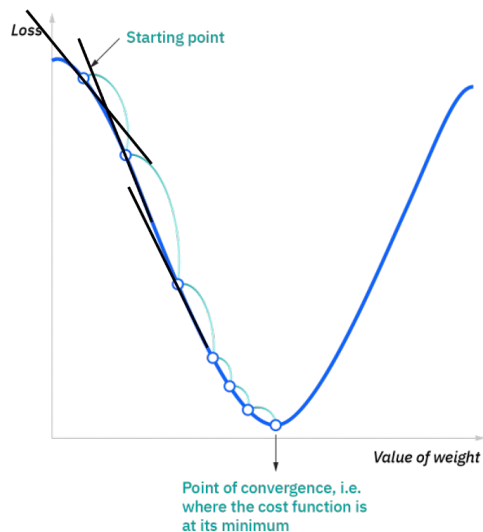
$$J = \frac{1}{N} \sum_{i=1}^N (\tilde{y}_i - y_i)^2 = \frac{1}{N} \sum_{i=1}^N ((a_0 + a_1 x_i) - y_i)^2$$

From this, we can calculate

$$\begin{aligned} \frac{\partial J}{\partial a_0} &= \frac{\partial}{\partial a_0} \left(\frac{1}{N} \sum_{i=1}^N ((a_0 + a_1 x_i) - y_i)^2 \right) = \frac{2}{N} \sum_{i=1}^N ((a_0 + a_1 x_i) - y_i) \\ &= \frac{2}{N} \sum_{i=1}^N (\tilde{y}_i - y_i) \end{aligned}$$

$$\begin{aligned} \frac{\partial J}{\partial a_1} &= \frac{\partial}{\partial a_1} \left(\frac{1}{N} \sum_{i=1}^N ((a_0 + a_1 x_i) - y_i)^2 \right) = \frac{2}{N} \sum_{i=1}^N ((a_0 + a_1 x_i) - y_i) x_i \\ &= \frac{2}{N} \sum_{i=1}^N (\tilde{y}_i - y_i) x_i \end{aligned}$$

Gradient descent mathematically



After taking the derivatives, the rest is easy!

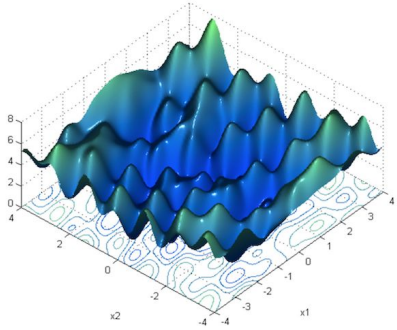
$$\begin{aligned}a_{0,new} &= a_0 - \lambda \frac{\partial \text{Loss}}{\partial a_0} \\ a_{1,new} &= a_1 - \lambda \frac{\partial \text{Loss}}{\partial a_1}\end{aligned}$$

The only other thing one must pay attention to is the learning rate λ (how big of a step to take). Too small and finding the right weights takes forever, too big and you might miss the minimum.



So far so good - this all looks super easy! What's the catch?

We calculated the gradients by hand because we knew the functional form we wanted to fit...



But the whole point of an ANN is that you do not need to assume the optimal functional form of the equation ahead of time - it's probably really complex!

So what do we do now...?

Backpropagation:

how gradients are calculated for gradient descent

- Short for “back propagation of errors”
- A method for determining the gradient of the loss function when you don't know the functional form
- **The method calculates the gradient of the loss function with respect to the neural network weights.**
- Why only certain activation functions are allowed - must be differentiable!

$$J = \frac{1}{N} \sum_{i=1}^N (\tilde{y}_i - y_i)^2$$

$$w_{k,new} = w_k - \lambda \frac{\partial \text{Loss}}{\partial w_k}$$

The partial derivative of the PREDICTED y with respect to the weights w .

$$\frac{\partial \text{Loss}}{\partial w_k} = \frac{\partial}{\partial w_k} \left(\frac{1}{n} \sum_{i=1}^n (\tilde{y}_i - y_i)^2 \right) = \frac{2}{n} \sum_{i=1}^n (\tilde{y}_i - y_i) \boxed{\frac{\partial \tilde{y}_i}{\partial w_k}}$$

What we need to know!

Backpropagation:

how gradients are calculated for gradient descent

3Blue1Brown Neural Network Series (Chapter 4):

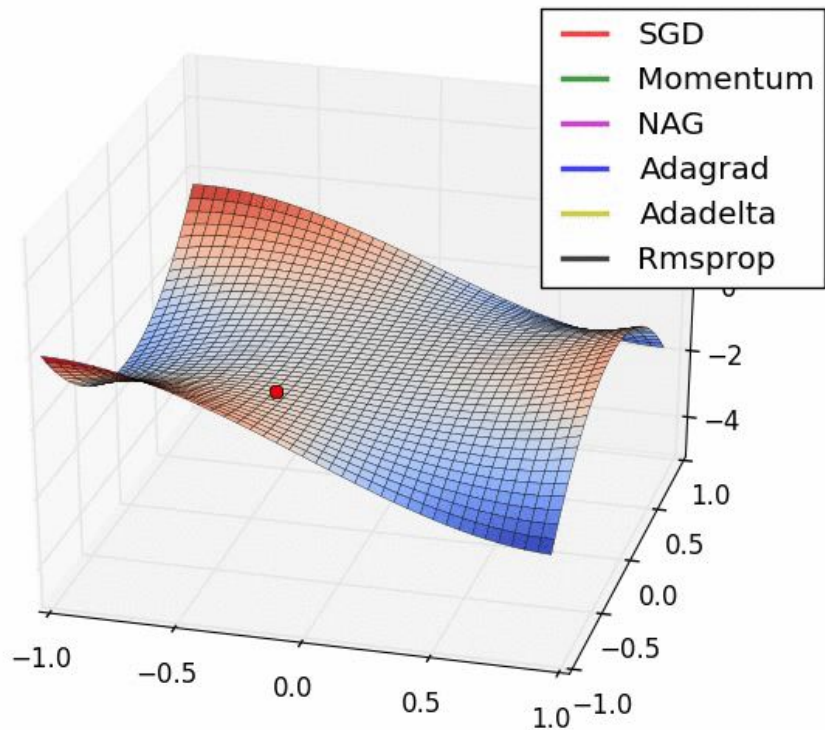
https://www.youtube.com/watch?v=tLeHLnjs5U8&list=PLZHQObOWTQDNU6R1_67000Dx_ZCJB-3pi&index=4

Neural Networks & Deep Learning *Free Book* (Chapter 2):

<http://neuralnetworksanddeeplearning.com/chap2.html>

The Chain Rule!!

$$\frac{\partial \text{Loss}}{\partial w_k} = \frac{\partial}{\partial w_k} \left(\frac{1}{n} \sum_{i=1}^n (\tilde{y}_i - y_i)^2 \right) = \frac{2}{n} \sum_{i=1}^n (\tilde{y}_i - y_i) \boxed{\frac{\partial \tilde{y}_i}{\partial w_k}}$$

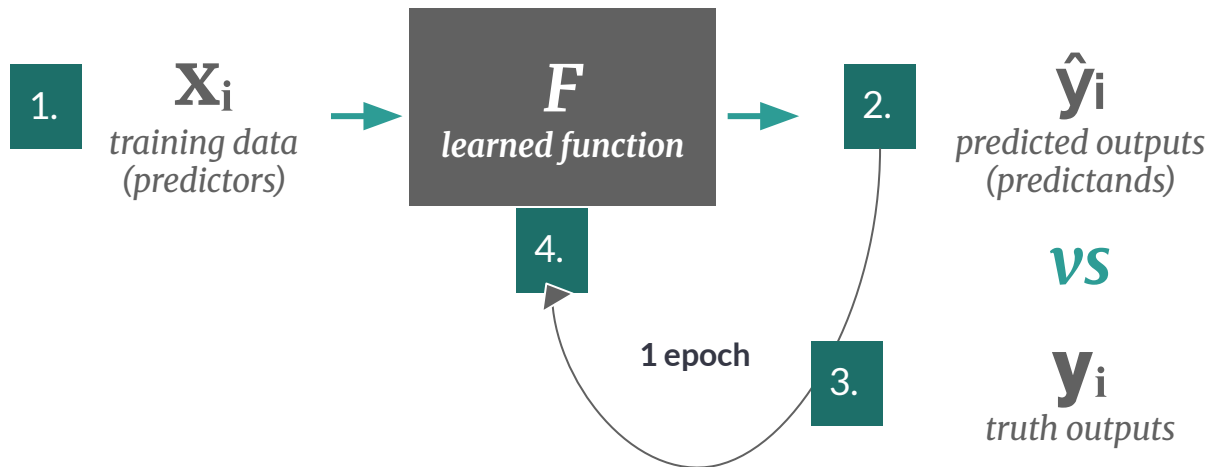
Gradient Descent Algorithms

Figure by Alec Radford

Machine Learning Approach

- 1) Training data is put into the machine learning model's learned function.
- 2) The model outputs its predictions.
- 3) These predictions are compared with their truth values.
- 4) The weights are updated to decrease the error between predicted and truth outputs.

This cycle continues.





Terminology

Batch size: number of samples in a “chunk” of training data that is iterated before updating the weights and biases

Epoch: One complete pass through the entire training dataset.

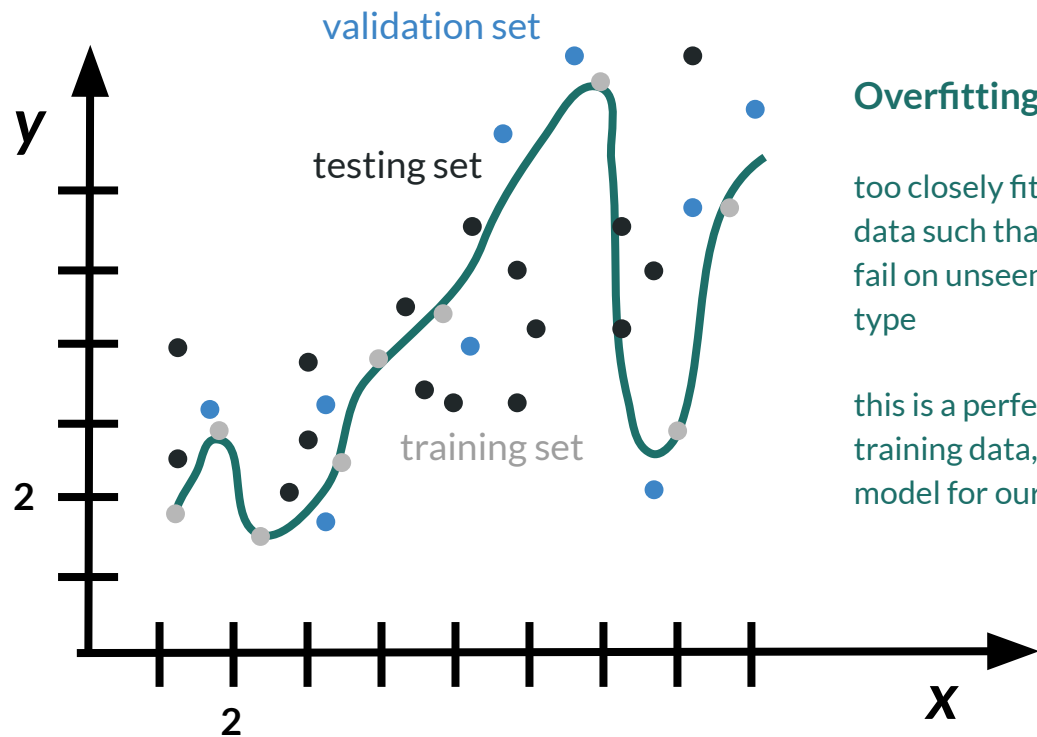
Example:

Samples: $N = 100$

Batch Size: $N = 20$

1 Epoch = 5 batches of 20 samples

How do we know if we're overfitting?



Overfitting

too closely fitting the training data such that the model will fail on unseen data of the same type

this is a perfect model for the training data, but a very poor model for our testing data

Training Data

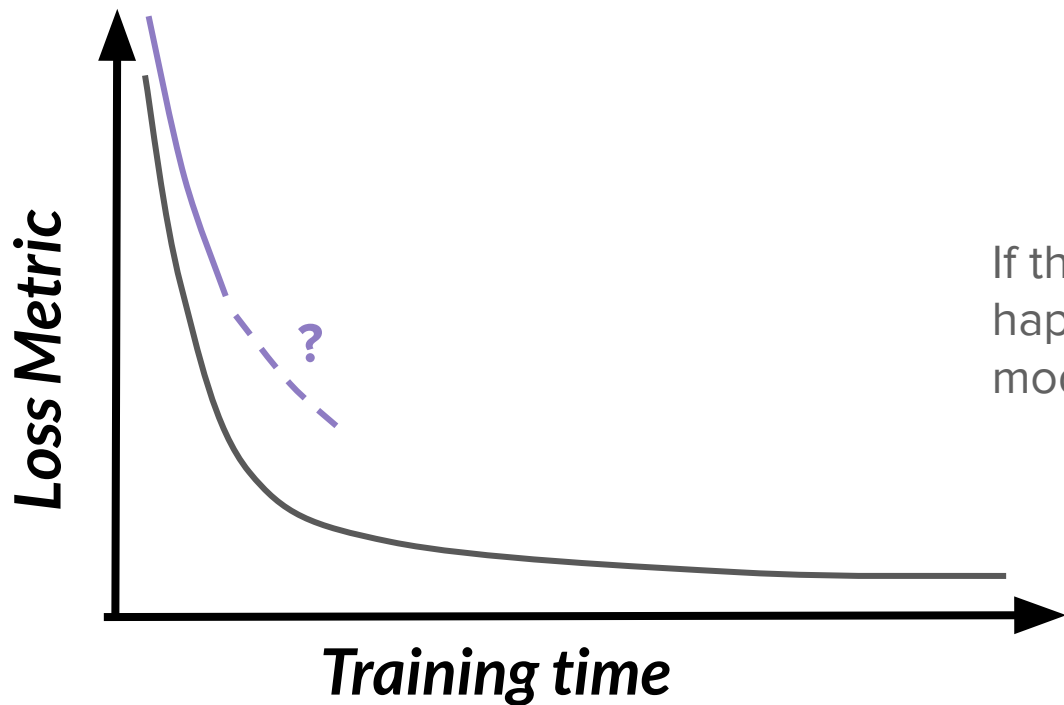
Validation Data

Testing Data

ALL THE DATA

What does **overfitting** look like?

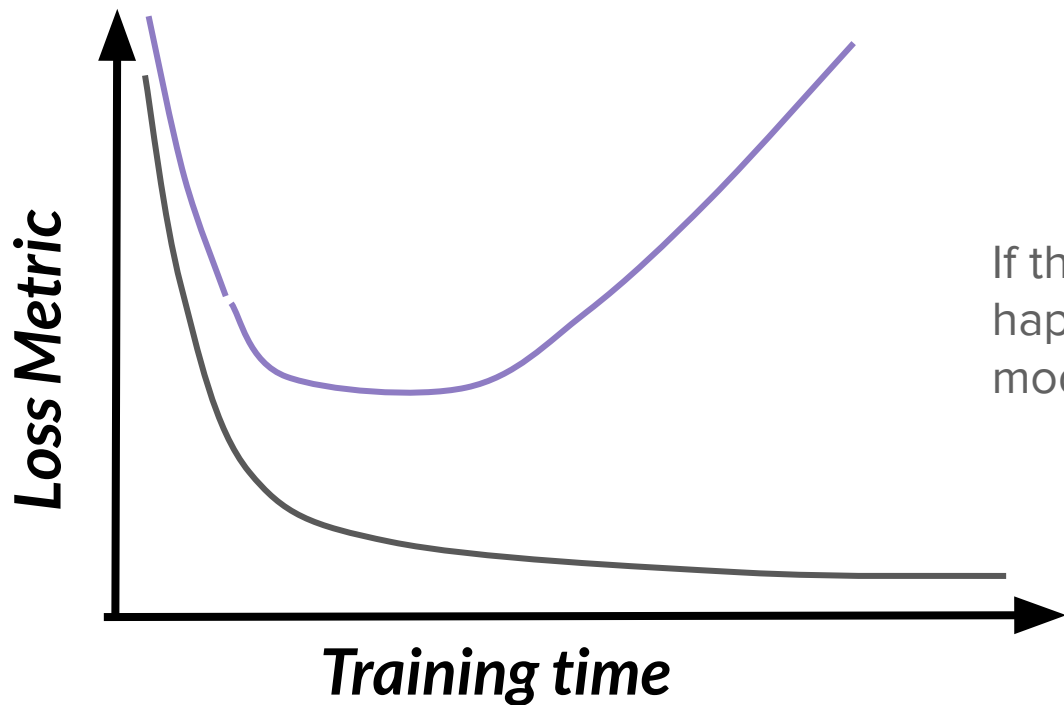
TRAINING data
VALIDATION data



If the model is overfitting, what will happen to the performance of the model on **validation**?

What does **overfitting** look like?

TRAINING data
VALIDATION data



If the model is overfitting, what will happen to the performance of the model on **validation**?

Common ways to address overfitting:

Ridge Regression/L2

- Adds a term to the loss/error function that forces coefficients to share the weight

$$\text{Loss} = \text{RMSE} + \lambda \sum_{i=1}^p w_i^2$$

Ridge Parameter

How much to
penalize large
weights

Shrinkage Term

Forces the individual
coefficients to be small

>> sharing of weight across
coefficients

Common ways to address overfitting:

LASSO Regression/L1

- Adds a term to the loss/error function that forces some weights to zero

$$\text{Loss} = \text{RMSE} + \lambda \sum_{i=1}^p |w_i|$$

LASSO Parameter

How much to
penalize large
weights

Shrinkage Term

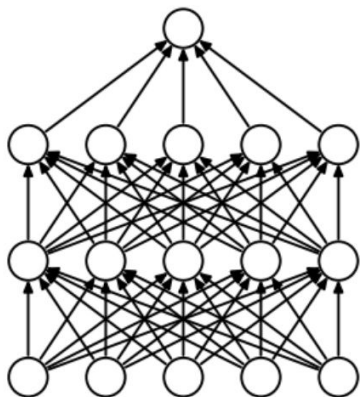
Punishes high values but
actually sets them to exactly
zero if not important

>> helpful for determining
which features are **most**
important

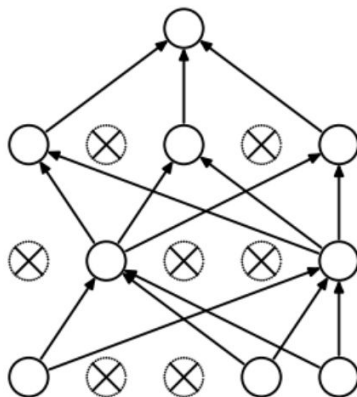
Common ways to address overfitting:

Dropout (only during training)

- Randomly remove or “drop out” nodes, requiring that network to still make accurate predictions in spite of losing these nodes.
- Forces nodes within a layer to take on more or less responsibility for the inputs



(a) Standard Neural Net



(b) After applying dropout.

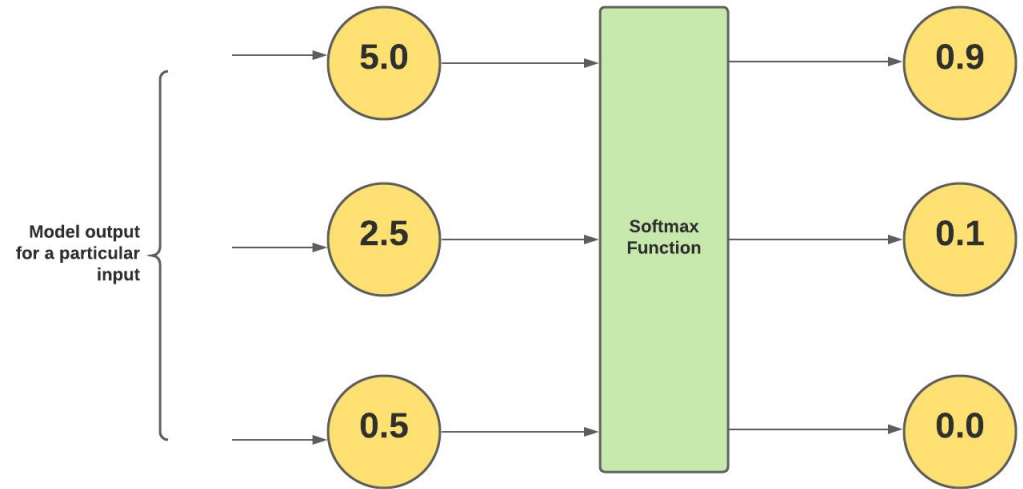
requires the user chose the percent to drop-out

- *(often ~50-80%)*

requires more epochs during training

Softmax: Let's make the output a probability!

- Converts “raw” output to probabilities for classification tasks
- After the softmax, output values are between 0 and 1
- The class with the highest probability gives your answer



Oh, is that all?



ANN considerations

Train the data the right number of times (epochs)....

If we iterate too few times, our model will not have time to find the optimal fit. If we iterate too many times, we will overfit the training data.

Choosing the right number of layers and nodes (architecture)....

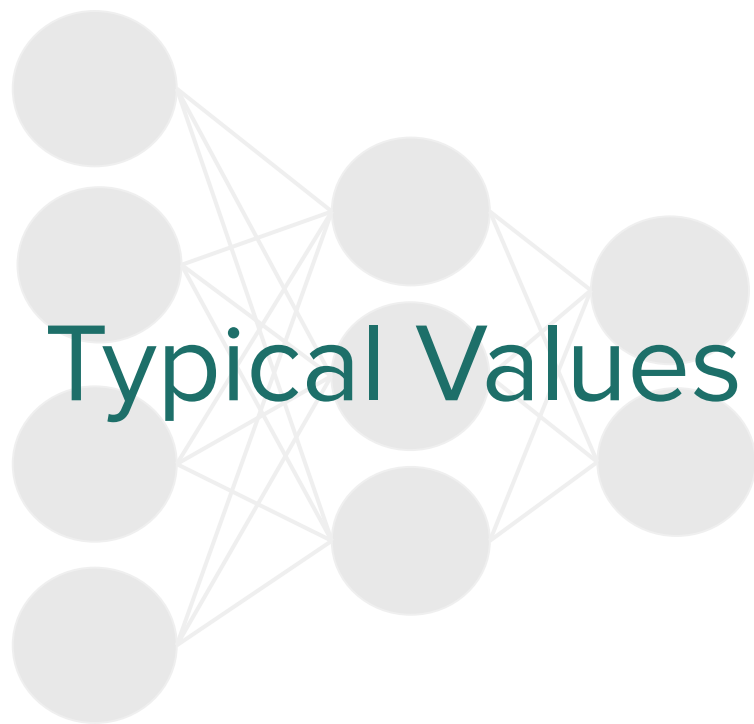
Fewer layers/nodes allow for easier interpretability. More layers/nodes allow for more nonlinearity.

Choose an activation function most appropriate for your solution....

Identify the right learning rate....

A learning rate too large may cause the solution to skip between local minima. A learning rate too small may get stuck in one local minimum, and also may take a longer to learn.

Others include... batch size, gradient descent algorithm, regularization, data preparation



Learning Rate: 0.01, 0.001., 0.0001

Batch Size: 32, 64, 128

Nodes: 4, 8, ..., 256

ENSO PREDICTION

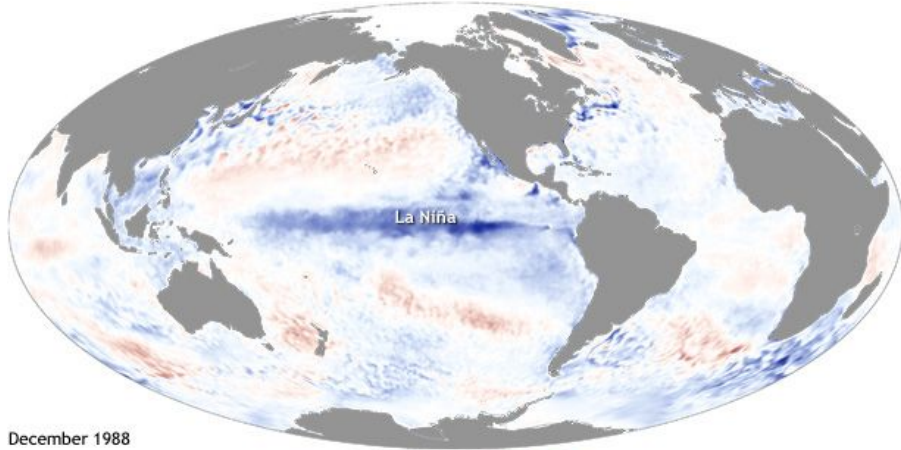
Find the **best function** that maps SSTs to the Nino3.4 index

Tropical SSTs

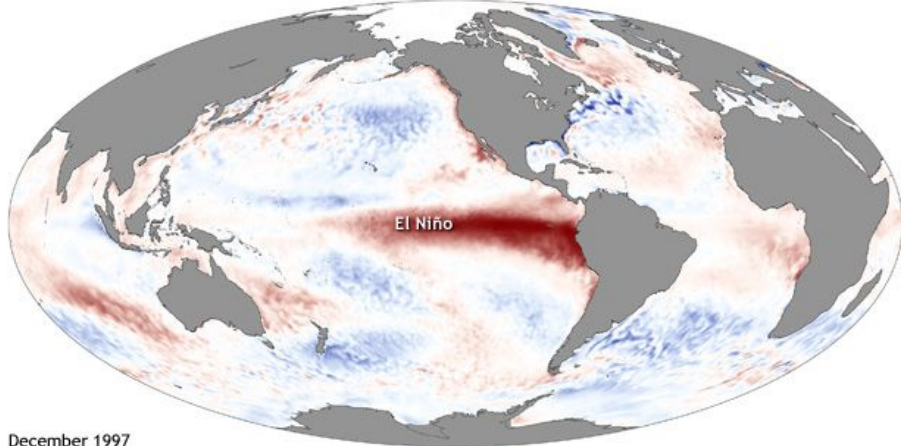
$F(X)$

learned non-linear
relationship

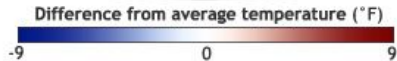
La Nina or El Nino



December 1988



December 1997

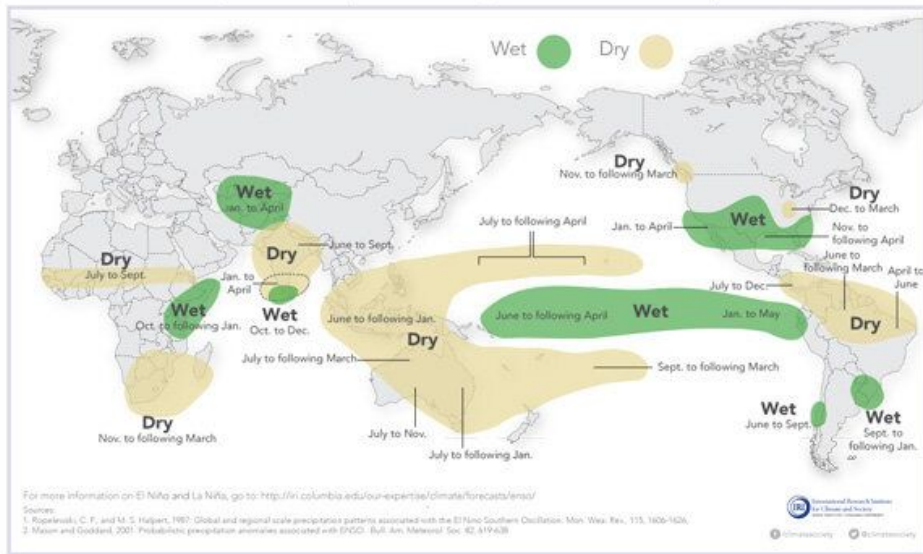


What is the El Nino Southern Oscillation (ENSO)?

SST anomalies during a strong La Niña and El Niño. Maps by NOAA Climate.gov, based on data provided by NOAA View.

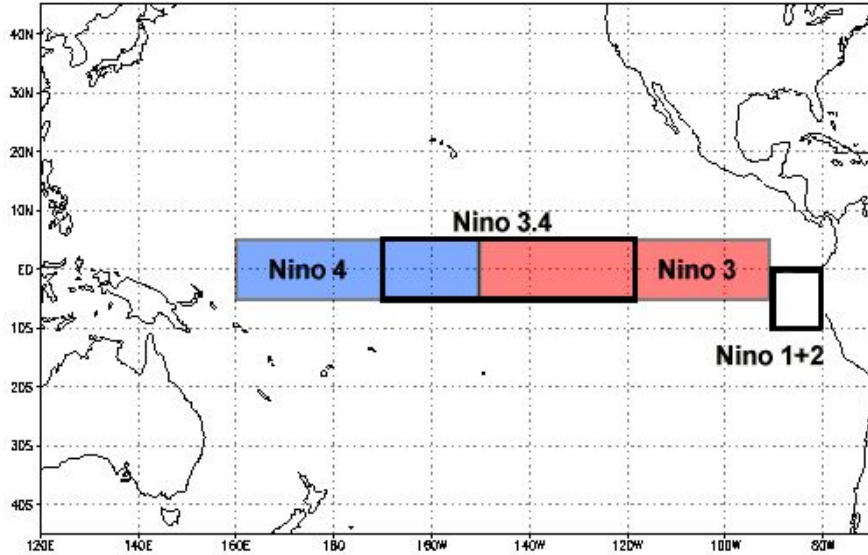
El Niño and Rainfall

El Niño conditions in the tropical Pacific are known to shift rainfall patterns in many different parts of the world. Although they vary somewhat from one El Niño to the next, the strongest shifts remain fairly consistent in the regions and seasons shown on the map below.



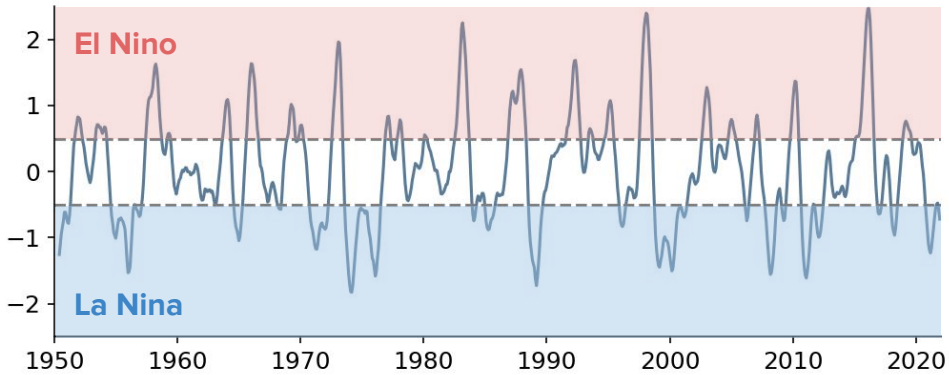
Why do we care?

Typical rainfall patterns during El Niño events. Such teleconnections are *likely* during El Niño events, but not *certain*. Map by [IRI](#).



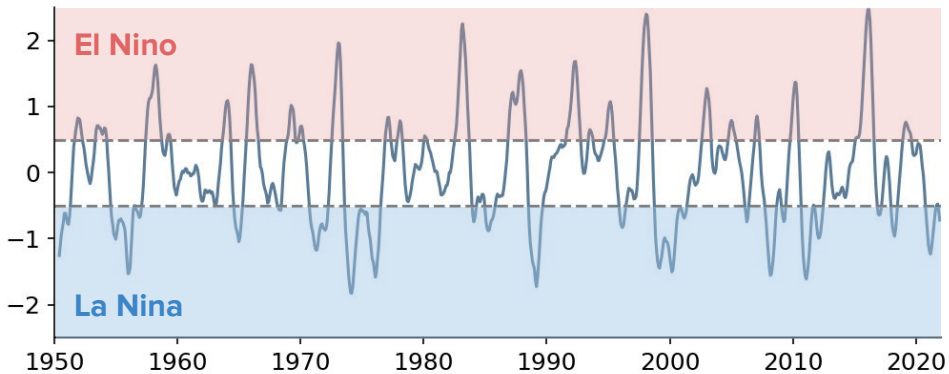
How do we define ENSO?

Nino3.4 Index

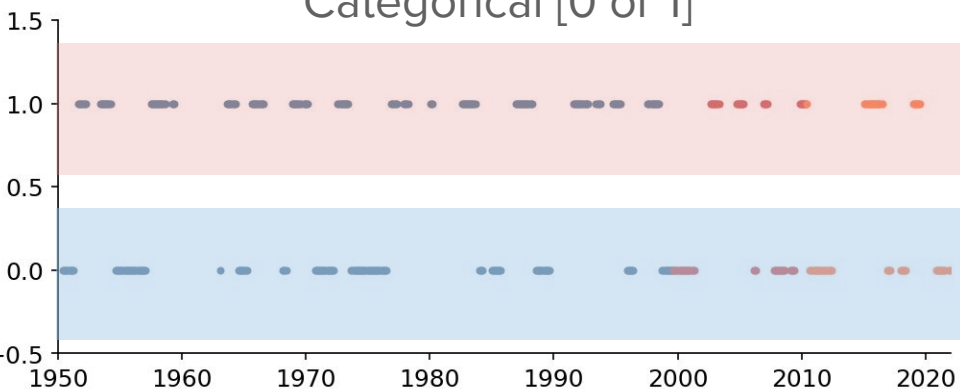


How do we define ENSO?

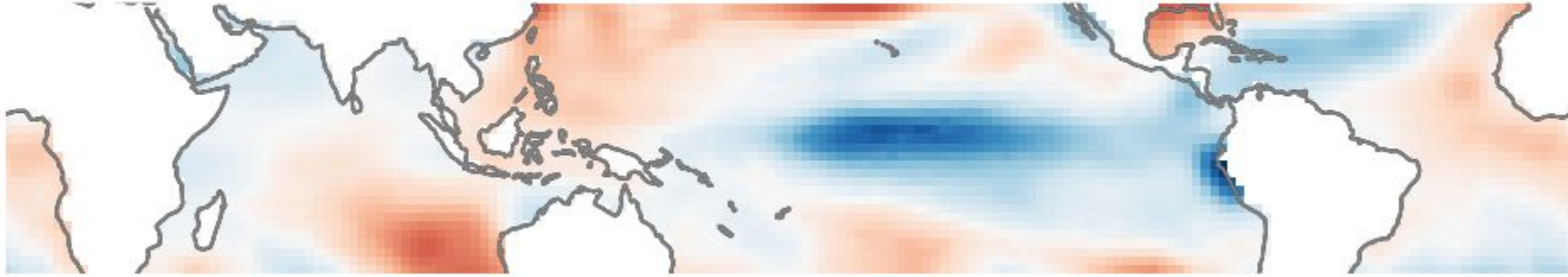
Nino3.4 Index



Categorical [0 or 1]



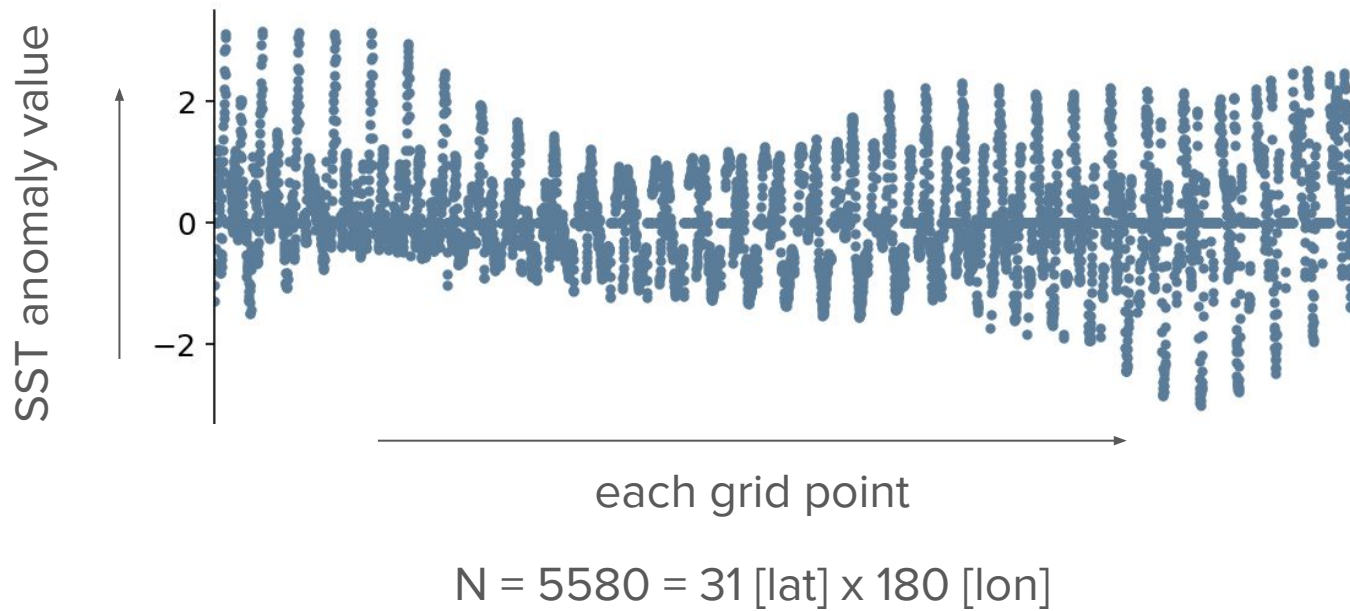
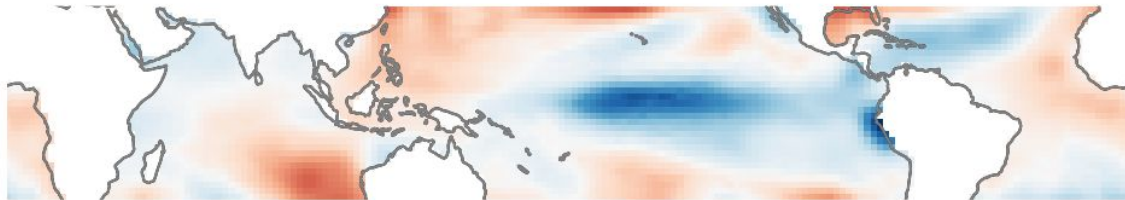
What will we predict?

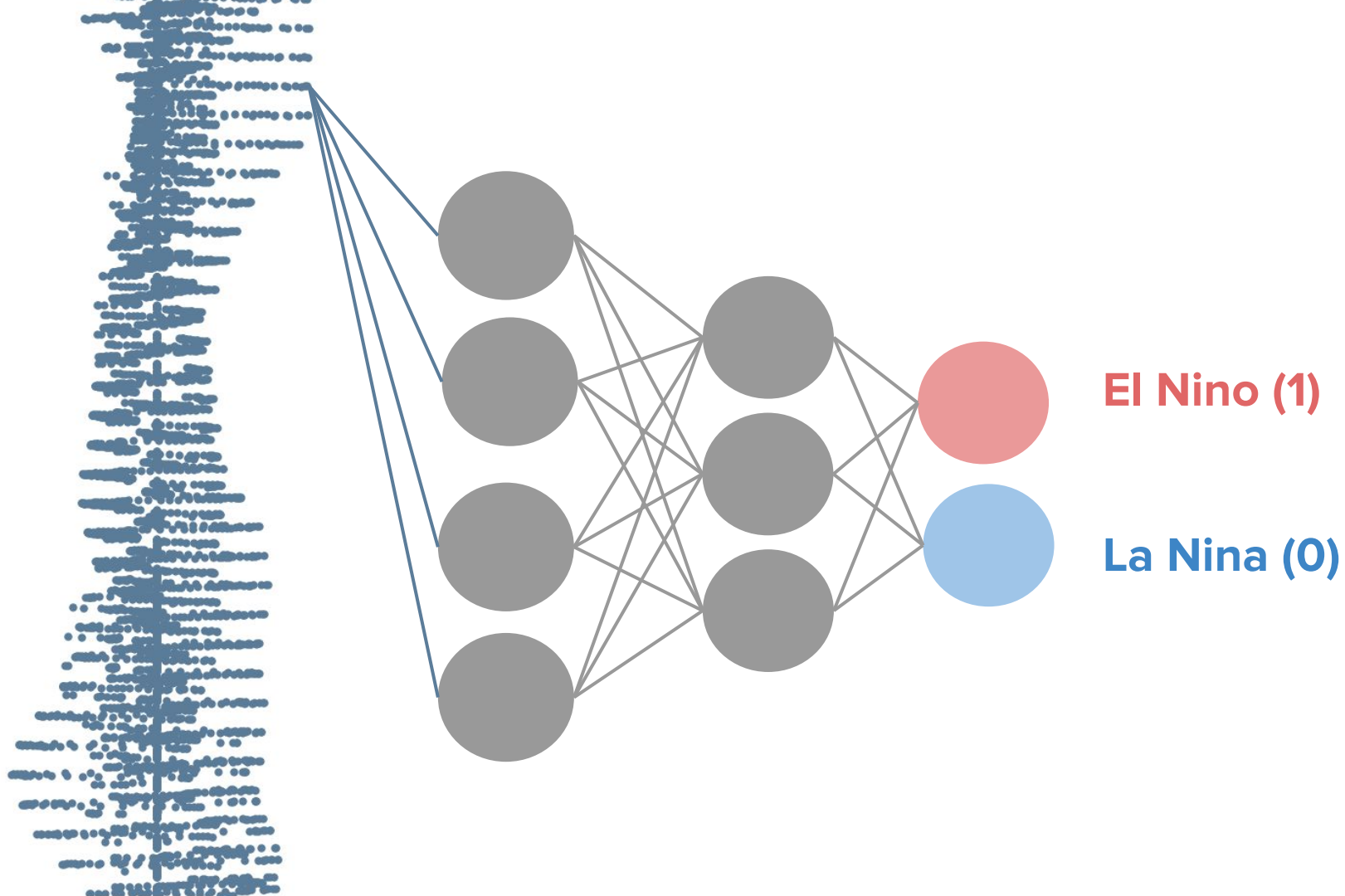


?

La Nina (0)

El Nino (1)





model.summary()

Model: "model_1"

Layer (type)	Output Shape	Param #
input_2 (InputLayer)	[(None, 5580)]	0
dense_3 (Dense)	(None, 12)	66972
dense_4 (Dense)	(None, 12)	156
dense_5 (Dense)	(None, 2)	26
Total params: 67,154		
Trainable params: 67,154		
Non-trainable params: 0		

How do we get the parameter #?

$$5580 \times 12 = 66960$$

BUT we also need to add the bias terms for each node

$$66960 + 12 = \mathbf{66972}$$

model.summary()

Model: "model_1"

Layer (type)	Output Shape	Param #
input_2 (InputLayer)	[(None, 5580)]	0
dense_3 (Dense)	(None, 12)	66972
dense_4 (Dense)	(None, 12)	156
dense_5 (Dense)	(None, 2)	26
Total params: 67,154		
Trainable params: 67,154		
Non-trainable params: 0		

How do we get the parameter #?

$$5580 \times 12 = 66960$$

BUT we also need to add the bias terms for each node

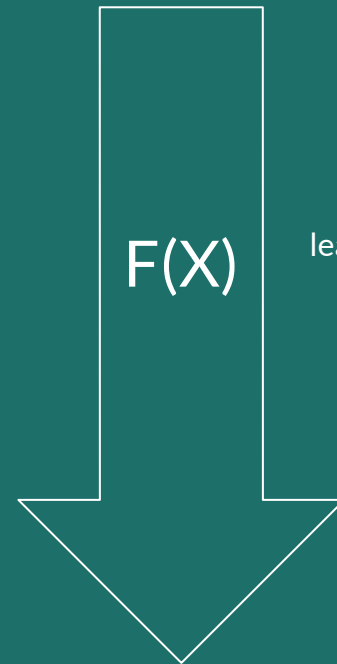
$$66960 + 12 = \mathbf{66972}$$

$$12 \times 12 = 144 + 12 = \mathbf{156}$$

$$12 \times 2 = 24 + 2 = \mathbf{26}$$

ENSO EXAMPLE

Tropical SSTs



$F(X)$

learned non-linear
relationship

La Nina or El Nino

During this exercise, try to:

- Train a “good” model
 - I’ll leave the definition up to you
- Train a “bad” model
 - E.g. an overfit model = bad model
- Try some non-standard values for various hyperparameters
 - Does the ANN respond as you expected?

Some resources I recommend to get started!

[3Blue1Brown Videos](#)

[Hands-On Machine Learning with Scikit-Learn, Keras & Tensorflow](#) (book)

Blogs! (IBM, AWS, Towards Data Science)

[Coursera ML Course by Andrew Ng](#)

Tensorflow/Keras and Pytorch are the go-to python machine learning libraries

- their websites have example code
- Tensorflow has a smaller learning curve