

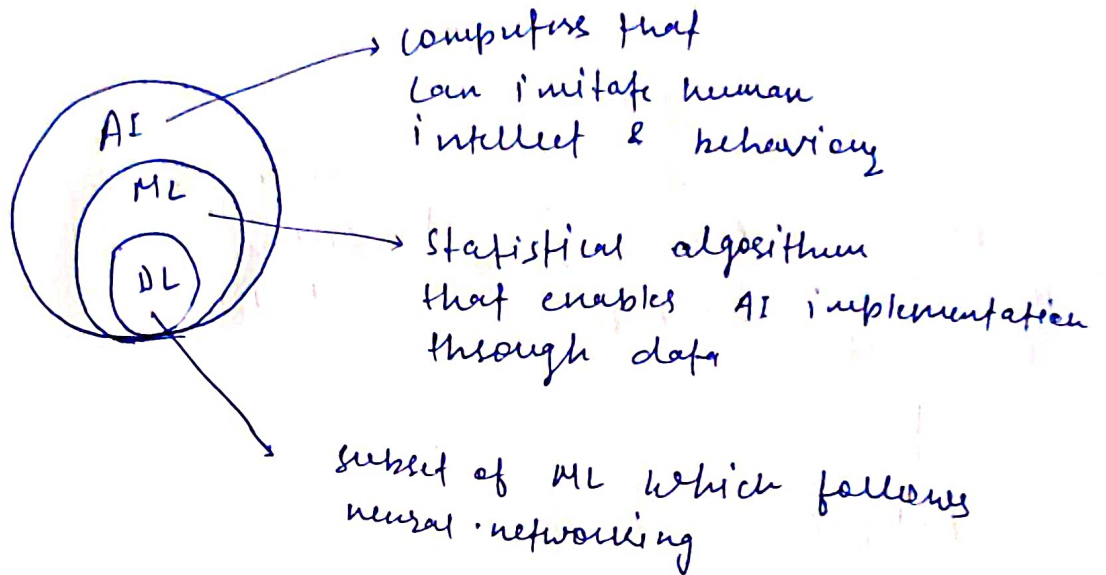
AI

Smart Knowledge

#

Basics

* AI vs ML vs DL



*

ML

Supervised

Classification

- ↳ Logistic
- ↳ Decision tree
- ↳ Random forest
- ↳ KNN
- ↳ SVM
- ↳ Naive Bayes

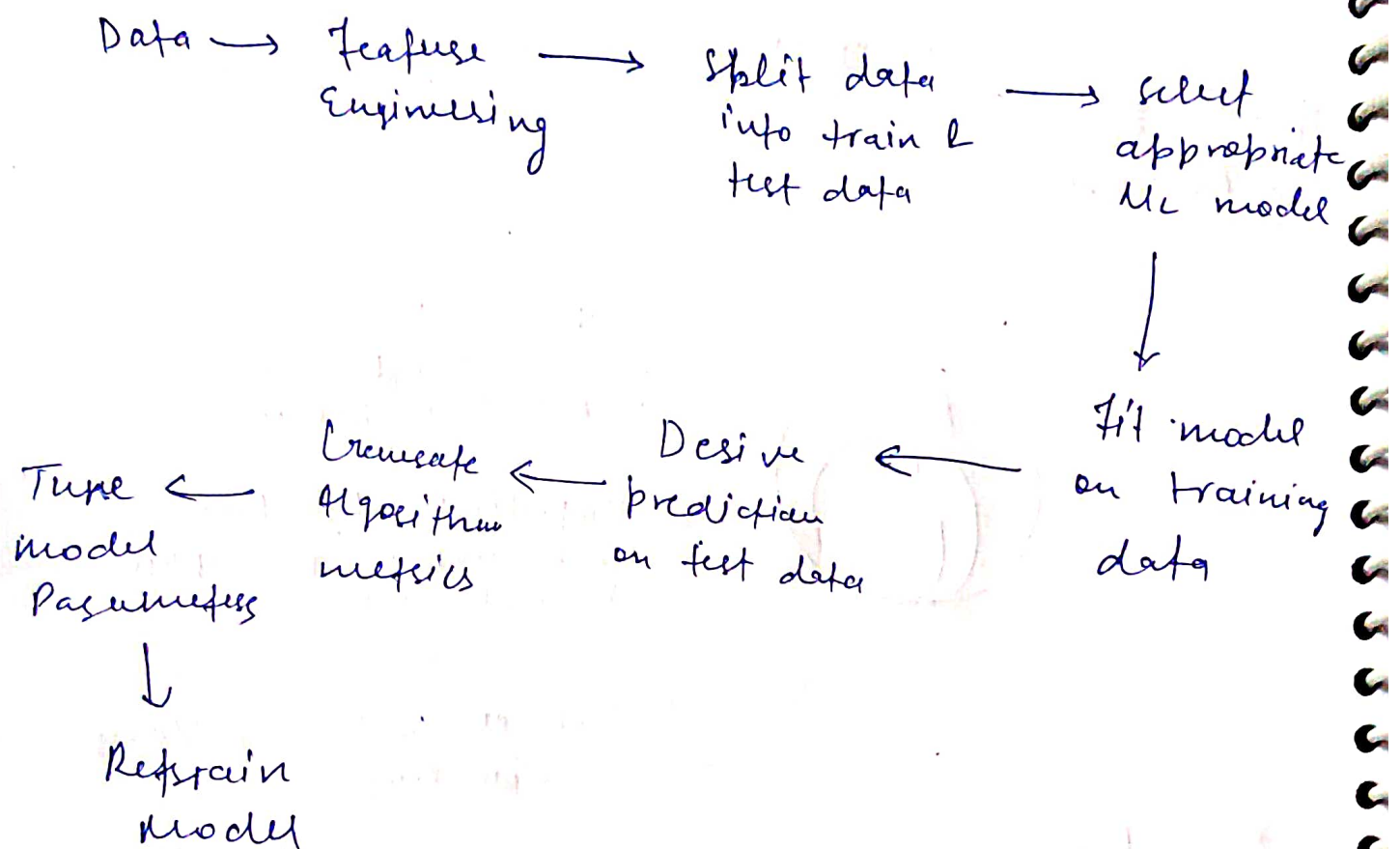
Regression

- ↳ simple linear reg.
- ↳ multiple linear reg.
- ↳ Ridge reg.
- ↳ Lasso reg.

unsupervised

- ↳ K-means clustering
- ↳ Hierarchical clustering

* Prediction pipeline



* Application of AI

- Image classification
- Image De-noising
- Object detection
- Chatbot
- Pose estimation
- Stock price prediction
- Image segmentation
- Text summarization
- Language translation

* Teachable machine search internet

* Working of a neuron

- A neuron accepts a signal and sends a message to the brain to act in a particular way.
- It accepts a signal, generates a particular response in correspondence from its nucleus and then triggers a response.
- The response may pass to another neuron and the cycle proceeds in this manner.

4 Neural Networks

① - Neurons

② - Weights

③ - Bias

④ - Synapse

⑤ - Architecture of a neural network

⑥ - Layers in NN

a) Input layer

Hidden layer

Output layer

⑦ - General equation of a Neural network

⑧ - Activation functions

a) Linear

b) Sigmoid

c) ReLU

d) Leaky ReLU

e) Tanh

f) Softmax

⑨ - Optimizers & Optimization Techniques

a) Gradient descent

b) Stochastic Gradient descent (SGD)

c) Mini batch Gradient descent

- d.) SGD with momentum
- e.) Adagrad
- f.) Adadelta
- g.) RMS Prop
- h.) Adam

⑩ - Loss function -

- a.) MSE
- b.) MAE
- c.) Binary Crossentropy
- d.) Categorical Crossentropy
- e.) Sparse - Categorical Crossentropy

⑪ - Forward Propagation

⑫ - Backward Propagation

⑬ - Epochs

⑭ - Classification Metrics - Accuracy, Precision, Recall, F1-Score

⑮ - Gradient descent from scratch (without tensorflow)

#

Neural network working

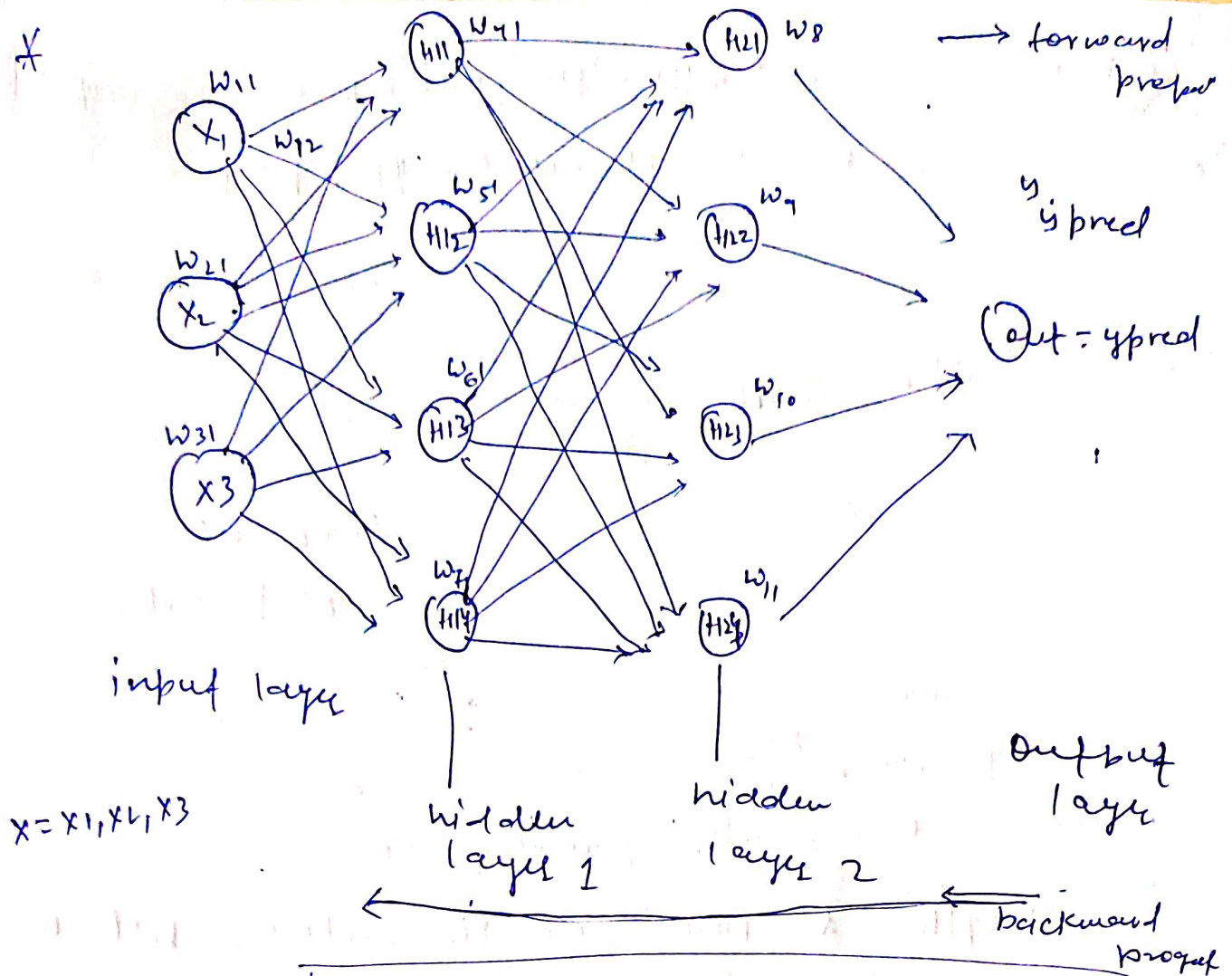
- Neural networks are a set of algorithms, modeled loosely after human brain, that are designed to recognize patterns. They interpret data through a kind of machine perception and process input data to solve a certain problem. The patterns they recognize are numerical, contained in vectors, into which all real world data i.e text, images, sound can fit.
- Neural networks are a class of machine learning algorithms used to model complex patterns in datasets using multiple hidden layers & activation functions. A neural network takes an input, passes it through multiple layers hidden (mini-functions with unique coefficient that must be learned), and outputs a prediction representing the combined input of all the neurons.
- Neural networks are trained iteratively using optimization techniques like gradient descent. After each cycle of training, an error metric is calculated based on the difference between prediction and target. The derivatives of this error metric are calculated & propagated back through the network using a technique called backpropagation. Each neuron's coefficient (weights) are then adjusted relative to how much they contributed to the total error. This process is repeated iteratively until the network error drops below an acceptable threshold.

- In a neural network, each neuron (except those in the input layer) is actually a sum of all its inputs; which are in fact the outputs from the previous layer multiplied by some weights. An additional term called bias is added to this sum. And a non-linear known as activation function is applied to the result.

* Layers in Neural Network

- ① - Input layer - Holds the data your model will train on. Each neuron in the input layer represents a unique attribute in your datasets (e.g. height; hair color, etc).
- ② - Hidden layer - Sits b/w the input & output layers and applies an activation function before passing on the result. There are often multiple hidden layers in a network. In general, hidden layers are typically full-connected layers - each neuron receives input from all the previous layer's neurons & sends its output to every neuron in the next layer.
- ③ - Output layer - The final layer in a ~~network~~ network. It receives inputs from the previous hidden layer, optionally applies an activation function, and returns an output representing your model's prediction.

*



$$x = x_1, x_2, x_3$$

$$Err = \left(\frac{1}{n} \right)^* \sum (y_i - y_{pred_i})^2$$

Total observation

w_{11} = weights

bias (constant term)

$$h_{11-in} = x_1 * w_{11} + x_2 * w_{21} + x_3 * w_{31} + b_{11}$$

$$h_{12-in} = x_1 * w_{12} + x_2 * w_{22} + x_3 * w_{32} + b_{12}$$

Input $\rightarrow$

$$h_{11-out} = \text{Activation}(h_{11-in})$$

$$h_{12-out} = \text{Activation}(h_{12-in})$$

output

$$H_{21-in} = H_{11-out} * w_{41} + H_{12-out} * w_{51} + H_{13-out} * w_{61} + H_{14-out} * w_{71} + b_{21}$$

- The upper table fill now

→ forward propagation

fill now

Iteration - 1

- errors also generating here
- weights & bias are the entities which is going to update in every iteration to minimise errors by backward propagation.
- how weights & bias get updated
Optimization technique.

① - gradient descent

$$w' = w - \alpha * d(Err) / d(w)$$

- w_8, w_9, w_{10}, w_{11} get updated first

- we add bias in case weights are 0, it will have no impact. so bias helps in either shifting of Activation func. to left or right.
- Bias can be also

*

Neuron - A neuron takes a group of weighted inputs, applies an activation function, and returns an output. Inputs to a neuron can either be features from a training set or outputs from previous layers' neurons.

Weight - Weights are applied to the inputs as they travel along synapse to reach the neuron & a bias term is added. The neuron then applies an activation function to the "sum of weighted inputs added to the bias" from each incoming synapse & passes the result on to all the neurons in the next layer.

Bias - Bias helps in controlling the value at which activation function will trigger. Bias value allows you to shift the activation func. to either left or right.

Activation function -

- ① - An activation function is a neural network defines how the weighted sum of the input is transformed into an output from a node or nodes in a layer of the network.
- ② - An activation function is a function that is added into an artificial neural network in order to help the network learn complex patterns in the data.
- ③ - Activation function helps to normalize the output of any input in the range b/w -1 to 1 or 0 to 1
- ④ - decides what is to be fired (activated) to the next neuron.

Synapse - Synapses are like roads in a neural network. They connect inputs to neurons to neurons, & neurons to outputs.

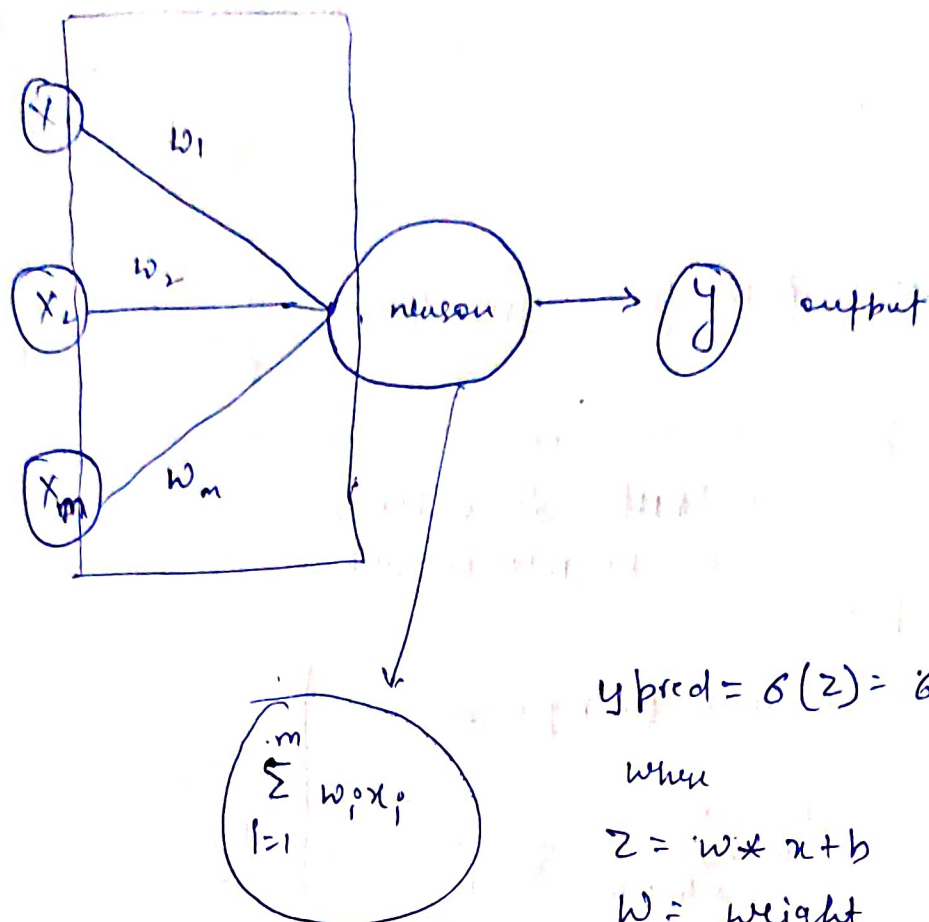
- Optimizers -
- ① - It's very important to ^{change} tweak the weights of the model during the training process, to make our predictions as correct & optimized as possible.
 - ② - This is where optimizers come into picture. They tie together the loss function & model parameters by updating the model in response to the output of the loss function.

Loss function / cost function

- ① - It is a parameter in model's predict function that tells us "how good" the model is at making predictions for a given set of parameters.
- ② - The loss function has its own curve & its own derivatives. The slope of this curve tells us how to change our parameters to make the model more accurate. we use the cost function to update our parameters.

*

General Equation of Neural Network



$$y_{pred} = \sigma(z) = \sigma(\sum_{i=1}^m w_i x_i + b)$$

where

$$z = w * x + b$$

w = weight

x = input

b = bias

σ = Activation function

*

Activation function

- properties

① - Non-Linear

② - Differentiable

- Types

- Linear
- sigmoid
- Tanh
- Relu

- leaky ReLU
- ELU
- softmax

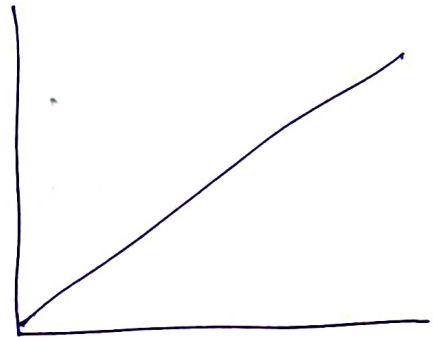
→ Linear Activation function

- Range $-\infty$ to $+\infty$
- It has constant derivative, so gradient has no relation with input

- used in ANNs for regression

↓
Artificial
neural
network

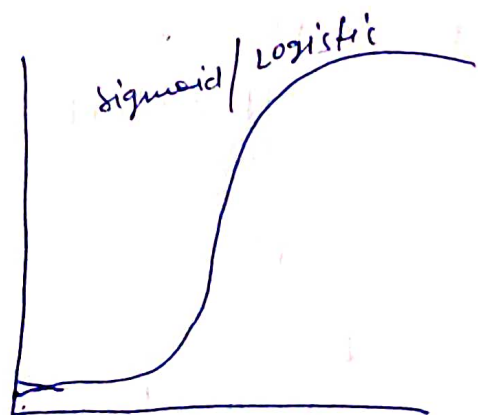
- $f(x) = x$



→ Sigmoid A. f.

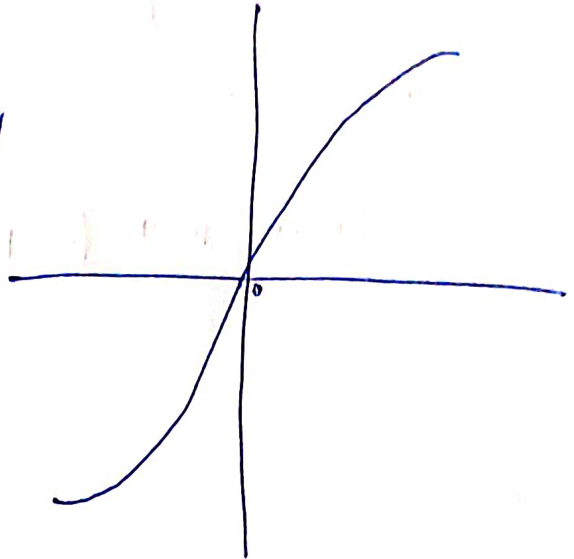
- Range 0 to 1
- Used for binary classification models
- It suffers from vanishing gradient (small) (slope)

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



→ Tanh

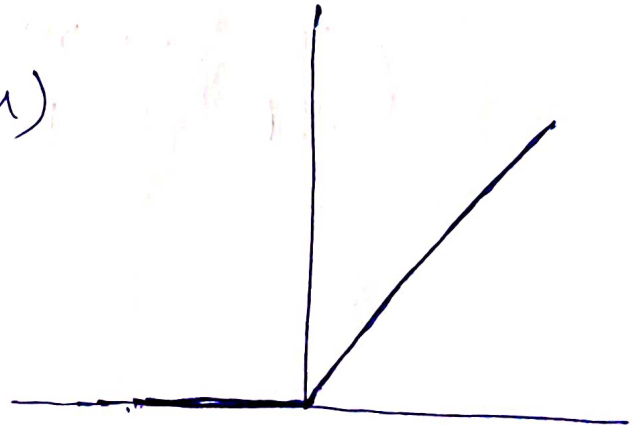
- range -1 to 1
- used in RNNs
- It suffers from vanishing gradient.
- It is zero centered
- $f(x) = \frac{(e^x - e^{-x})}{(e^x + e^{-x})}$



→ RELU

- Relu means rectified linear unit.
- Range 0 to +infinity
- used in hidden layers
- It suffers from dying relu problem

- $f(x) = \max(0, x)$



~~h~~ $h_{ll-out} = \text{Activation}(h_{ll-in})$

$h_{ll-in} = -ve$

$\text{Activation} = \text{relu}$

$h_{ll-out} = 0$

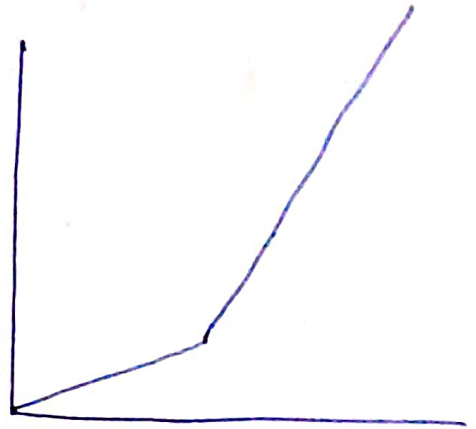
→ so no contribution of this neuron; so it's not activated

∴ suffers dying relu problem

→ Leaky RELU

- Leaky RELU is another variant of RELU
- used in hidden layers
- it overcomes the dying ReLU problem

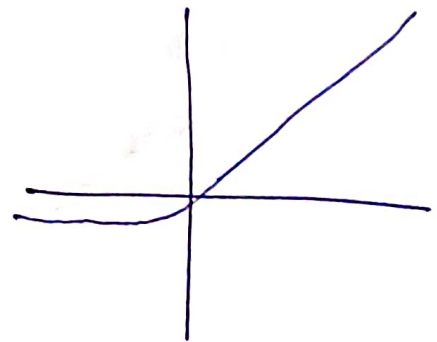
$$f(x) = \max(0.1x, x)$$



→ ELU

- it stands for Exponential Linear Unit
- used in hidden layers
- it overcomes the Dying ReLU problem
- suffers from Exponential Gradient.

$$\begin{cases} x & \text{for } x \geq 0 \\ \alpha(e^x - 1) & \text{for } x < 0 \end{cases}$$



• → Activation function used in output layers is linear

* Vanishing gradient

- like the sigmoid function certain activation function squish an ample input space into a small output space b/w 0 & 1
- Therefore, a large change in the input of the sigmoid function will cause a small change in the output. Hence, the derivative becomes small. For shallow networks with only a few layers that use these activations, this isn't a big problem.
- However, when more layers are used, it can cause the gradient to be too small for training to work effectively.

* Exploding gradient

- Exploding gradient are problems where significant error gradients accumulate & result in very large updates to neural network model weights during training
- An unstable network can result when there are exploding gradients, & the learning cannot be completed.
- The values of the weights can also become so large as to overflow & result in something called NaN values.

X

Optimizers

- point ① & ② already discussed previously

③ - Optimizers shape & mold your model into its most accurate possible form by manoeuvring with the weights. The loss function is the guide to the terrain, telling the optimizer when it's moving in the right or wrong direction.

⑥ Optimizers Functions

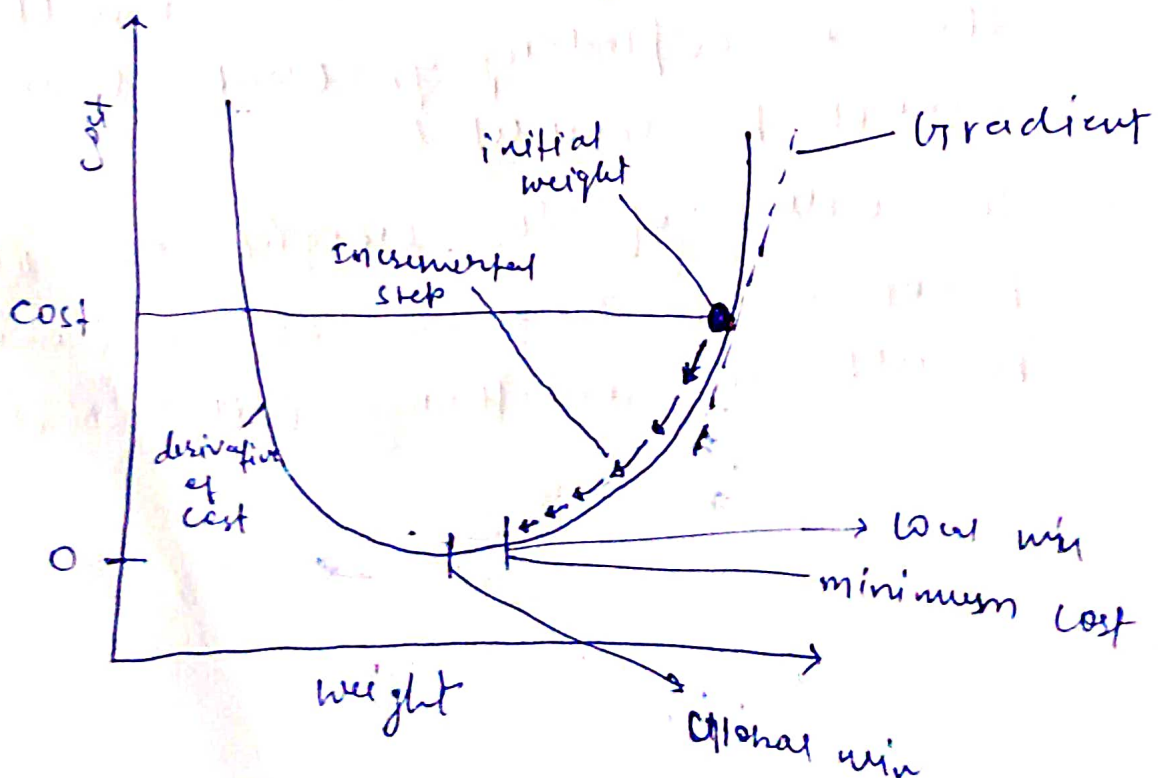
→ Gradient Descent

- It minimizes the given function to its local minimum.
- Gradient Descent iteratively reduces a loss function by moving in the direction opposite to that of steepest ascent.

$$w_{\text{new}} = w_{\text{old}} - \alpha - \frac{\partial(\text{Loss})}{\partial(w_{\text{old}})}$$

↓
Learning Rate

Loss = Cost = Error function

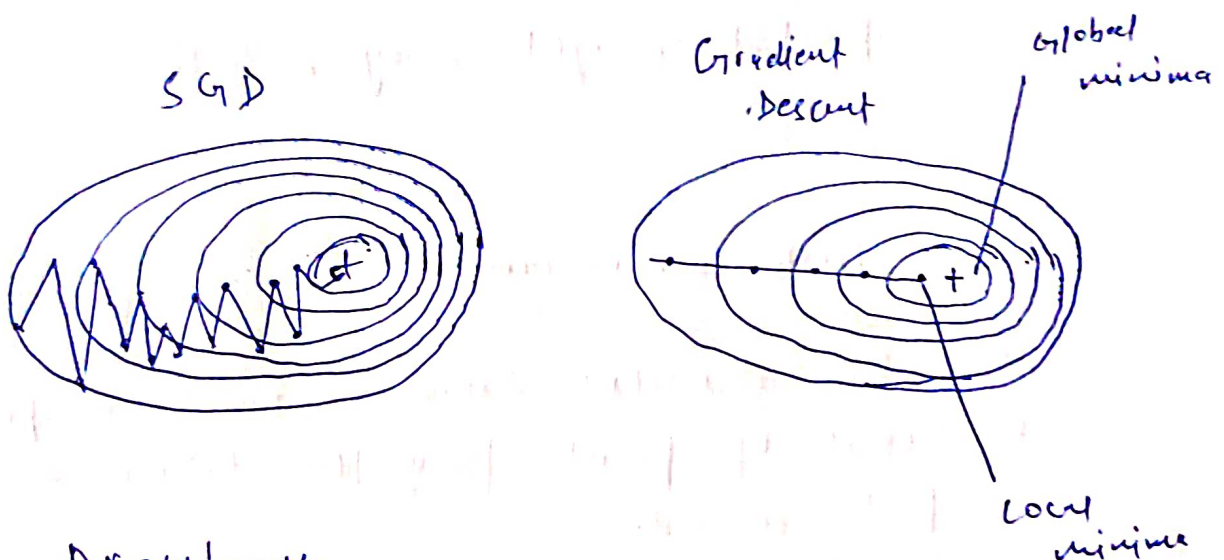


Drawbacks

- Slow in computation
- Requires large memory
- Computationally Expensive

➔ Stochastic Gradient Descent

- In the SGD algorithm derivative is computed taking one point at a time.
- Memory requirement is less than that of GD



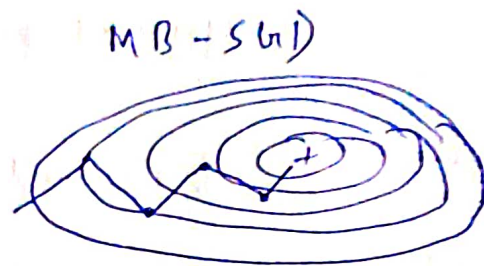
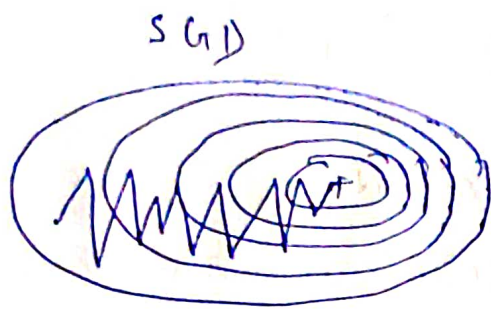
Drawbacks

- Takes a long time to converge
- May stuck at local minima

➔ Mini-batch Gradient Descent

- MB-SGD algorithm takes a batch of points or subset of points from the dataset to compute derivative.
- MB-SGD divides the dataset into various batches & every batch, the parameters are updated.

- less time complexity as compared to SGD & more stab.



Drawbacks

- may stuck at local minima
- updated weights may be too noisy

→ SGD with momentum

- MB-SGD algorithm takes a batch of points or subset of points from the dataset to compute derivative.
- MB-SGD divides the dataset into various batches & after every batch, the parameters are updated.
- The idea is to denoise derivative using exponential weighting average that is to give more weightage to recent updates compared to the previous update.

$$V_{dw} = \beta \cdot V_{dw} + (1 - \beta) \cdot dw$$

$$V_{db} = \beta \cdot V_{dw} + (1 - \beta) \cdot db$$

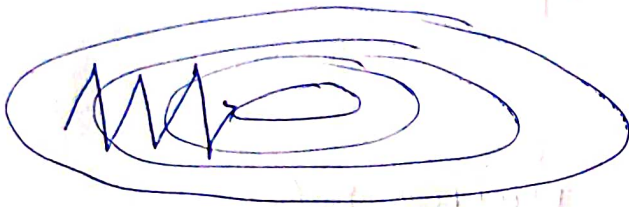
$$w = w - \alpha \cdot V_{dw}$$

$$b = b - \alpha \cdot V_{db}$$

- The problem with SGD is that while it tries to reach minima because of the high oscillation we can't increase the learning rate so it takes time to converge

Drawbacks

- may stuck at local minima
- updated weights may be too noisy



SGD without momentum



SGD with momentum

→ Adam (adaptive Momentum Estimation)

- It is a method that computes adaptive learning rates for each parameter. It stores both the decaying average of the past gradients, similar to momentum & also the decaying average of the past squared gradients, similar to RMS-Prop & Adadelta

$$\begin{cases} V_{dw_t} = \beta V_{dw_{t-1}} + (1-\beta) \frac{dL}{dw_{t-1}} \\ V_{db_t} = \beta V_{db_{t-1}} + (1-\beta) \frac{dL}{db_{t-1}} \end{cases}$$

exponential weighted avg for past gradients

$$s_{dw_t} = \gamma s_{dw_{t-1}} + (1-\gamma) \left(\frac{\partial L}{\partial w_{t-1}} \right)^2$$

exponential
avg for past
squared gradient

$$s_{db_t} = \gamma s_{db_{t-1}} + (1-\gamma) \left(\frac{\partial L}{\partial b_{t-1}} \right)^2$$

$$w_t = w_{t-1} - \frac{\eta}{\sqrt{s_{dw_t} - \epsilon}} * v_{dw_t}$$

$$b_t = b_{t-1} - \frac{\eta}{\sqrt{s_{db_t} - \epsilon}} * v_{db_t}$$

Adam
optimizing

alpha = eta * alpha = learning rate

$$w' = w - \text{alpha} * \left(\frac{\partial L}{\partial w} \right)$$

* Loss functions

loss or cost or error function

- In a neural network; the weights get multiplied with the inputs & then activation function is applied to the element before going to the next layer. Finally we get the predicted value ($\hat{y}$) through the output layer. But prediction is always closer to the actual (y), which we term as error. So we define the loss/cost functions to capture

the errors & try to optimize it through backpropagation.

- The error function that compares actual value with its corresponding predicted value is preferred to as loss function or cost function

• Types

→ MSE (Mean Square Error)

- It is used in Regression
- MSE is the preferred loss function if the distribution of the target variable is Gaussian. Mean square error is defined as the average of the squared differences b/w the predicted & actual values

$$L = \frac{1}{n} \sum_{i=1}^n (y - \hat{y})^2$$

↙ predicted value

n is batch size.

→ MAE (Mean absolute error)

- It is used in Regression
- It is defined as the average of the absolute difference b/w the actual & predicted values.

$$L = \frac{1}{n} \sum_{i=1}^n |y - \hat{y}|$$

→ Binary CrossEntropy

- It is used in Binary classification.
- It measures the difference b/w two probability distributions. If the cross entropy is small, it suggests that two distributions are similar to each other.
- In case of a binary classification the predicted probability is compared to the target/actual (0 or 1). Binary cross entropy calculates a score that provides the negative average difference b/w the actual & predicted probabilities for predicting the class. 1. This score penalizes the probability based on the distance from the expected value.

$$L = -y * \log(\hat{y}) - (1-y) * \log(1-\hat{y})$$

$$= \begin{cases} -y * \log(\hat{y}), & \text{if } y=1 \\ -(1-y) * \log(1-\hat{y}), & \text{if } y=0 \end{cases}$$

→ Categorical cross entropy

- It is used in multiclass classification.
- In case of a multi-class classification

the predicted probability is compared to the target/actual, where each class is assigned a unique integer value $(0, 1, 2, \dots, t)$, assuming data has t unique classes. It calculates a score that provides the negative average difference b/w the actual & the predicted probabilities for all classes

- for multi-class uses entropy; actual targets (y) are one-hot encoded. For a 3-class classification $[[0, 0, 1], [1, 0, 0], [0, 1, 0]]$

$$L = - \sum_{j=1}^t y_{ij} * \log(\hat{y}_{ij})$$

i = no. of rows

j = no. of categories

y_i = is one hot encoded target vector

→ Categorical Cross Entropy

- It is used in Multiclass classification
- In case of a multiclass classification the predicted probability is compared to the target/actual, where class is assigned a unique integer value $(0, 1, 2, \dots, t)$, assuming data has t unique classes. It calculates a score that provides the negative average difference b/w

the actual & predicted probabilities for all classes.

- for multiclass uses entropy, actual targets (y) are one-hot encoded. For a 3-class classification

$[[0, 0, 1], [1, 0, 0], [0, 1, 0]]$

$$L = - \sum_{j=1}^t y_{ij} * \log(\hat{y}_{ij})$$

where i = no. of rows

j = no. of categories

y_i = is one hot encoded target vector.

CNN (Convolutional Neural Network)

- CNN image classifications take an input image, process it & classify it under certain categories (E.g., Dog, Cat, Tiger, Lion).
- Computers see an input image as array of pixels & it depends on the image resolution. Based on the image resolution, it will see $h \times w \times d$ (h = height, w = width, d = dimensions). E.g. An image of $6 \times 6 \times 3$ array of matrix of RGB (3 refers to RGB values) & an image of $4 \times 4 \times 1$ array of matrix of grayscale image.
- When Deep learning CNN model's are applied to train & test, each input image will pass through a series of convolution layers with filters (kernels), Pooling, Flatten layer, fully connected layers (FC). ~~For multiclass~~ and apply softmax function to classify an object with probabilistic values b/w 0 & 1. For multiclass classification, we can use ^{loss} categorical cross entropy, sparse categorical cross entropy.
 - label encoded data
 - used for ? not encoded data

* Convolution

- Let us suppose this is the input matrix of 5×5 & a filter of matrix 3×3 , for those who don't know what a filter is a set of weights applied on an image or a matrix to obtain the required feature.

- In convolutional layer, the image pixels (image matrix) are multiplied with filter matrix to generate a feature matrix.
- Consider the handwritten digit 2 & its pixel representation.



0	0	0	0	0
0	20	0	0	0
0	19.5	0	0	0
0	99.2	0	0	0
0	14.5	279.184	1640	0
0	0	0	0	0

- For sake of simplicity of mathematical operations, consider image matrix of 9 as follows

-1	-1	1	-1	-1
-1	1	1	-1	-1
-1	-1	1	-1	-1
-1	-1	1	-1	-1
-1	-1	1	-1	-1

-1	1	1	1	-1
-1	1	1	1	-1
-1	1	1	1	-1
-1	1	1	1	-1
-1	1	1	1	-1

- Convolution operation is applied using a filter. Let's assume that a 3×3 filter is used for this operation. In convolution layer, the image pixels (image matrix) are multiplied with filter (kernels) matrix to generate feature map.

$$-1 + 1 + 1 - 1 - 1 - 1 + 1 + 1 = -1 = -1/9 = -0.11$$

-1	1	1	1	-1
-1	1	-1	1	-1
-1	1	1	1	-1
-1	-1	-1	1	-1
-1	-1	1	-1	-1
-1	1	1	-1	-1

$\times$

1	1	1
1	-1	1
1	1	1

-0.11	1	-0.11
0.55	0.11	-0.33
-0.33	0.33	-0.33
0.22	-0.11	-0.22
-0.33	-0.33	-0.33

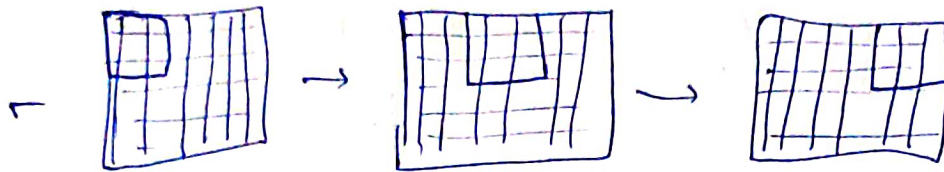
5×3

7×5

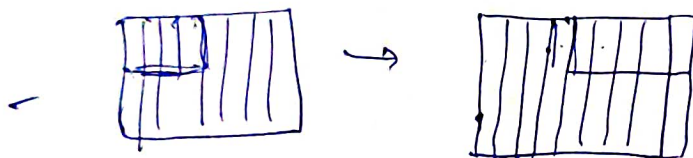
3×3

$$\begin{aligned} \text{img} &= (ih, iw) \\ K &= (Kh, Kw) \\ fm &= (ih - Kh + 1, iw - Kw + 1) \end{aligned}$$

Feature map



Stride = 1 (1 cell movement)



Stride = 2
(2-cell movement)

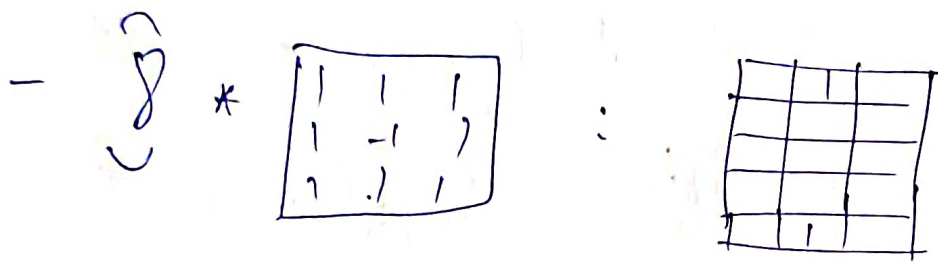
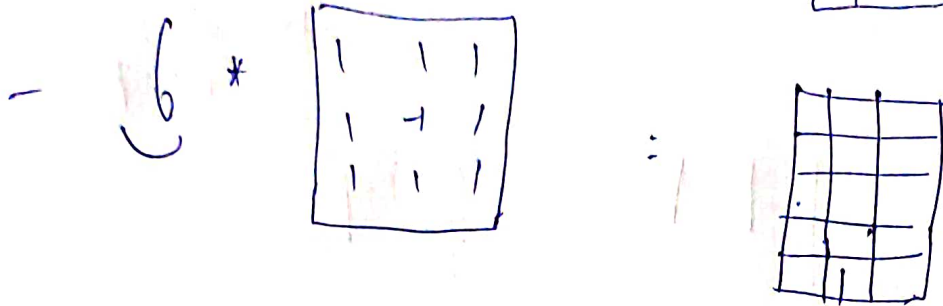
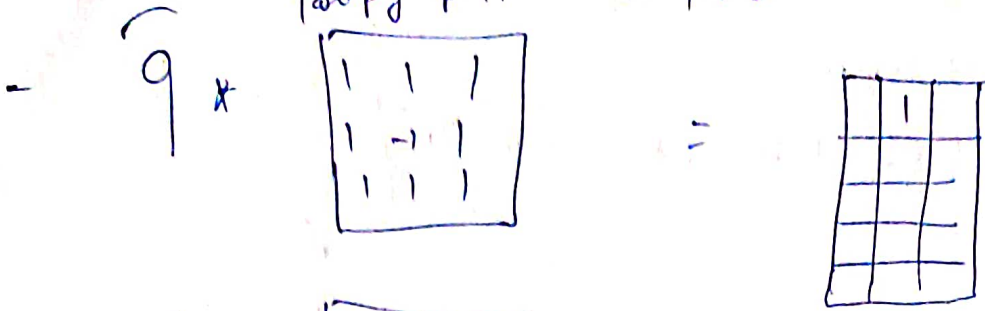
↓
for Stride = 2

$$\begin{aligned} \text{img} &= ih, iw \\ K &= Kh, Kw \\ S &= s \\ fm &= [(ih - Kh)/2 + 1, (iw - Kw)/2 + 1] \end{aligned}$$

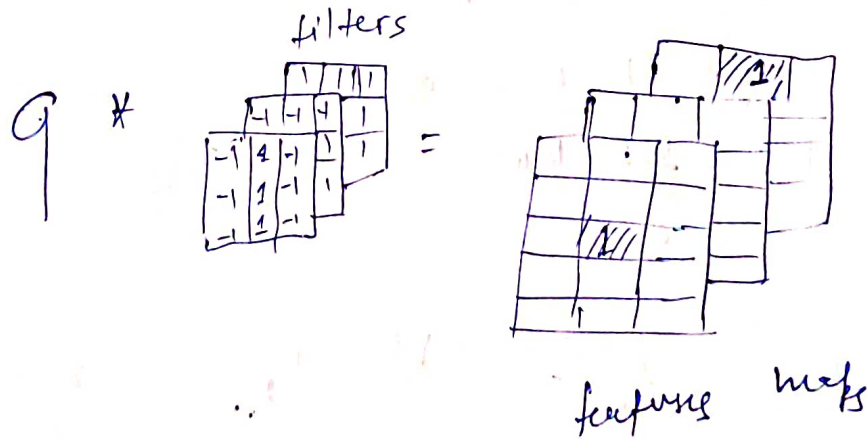
- This is how result of convolution process detects different patterns in the images using filters.

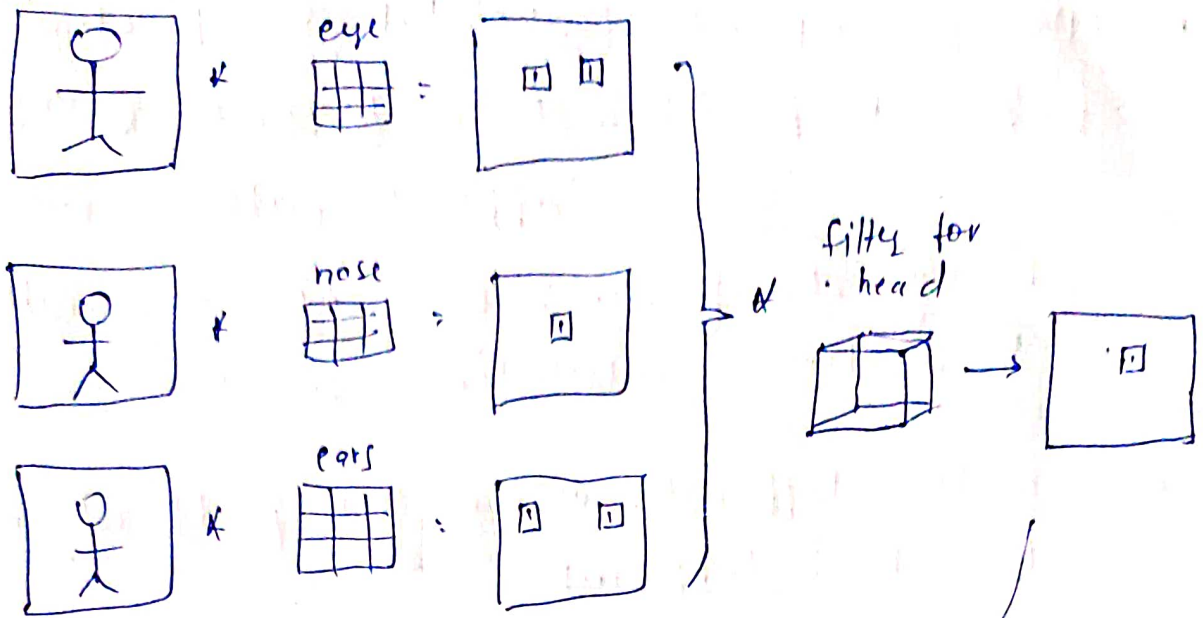
→

loopy pattern detector

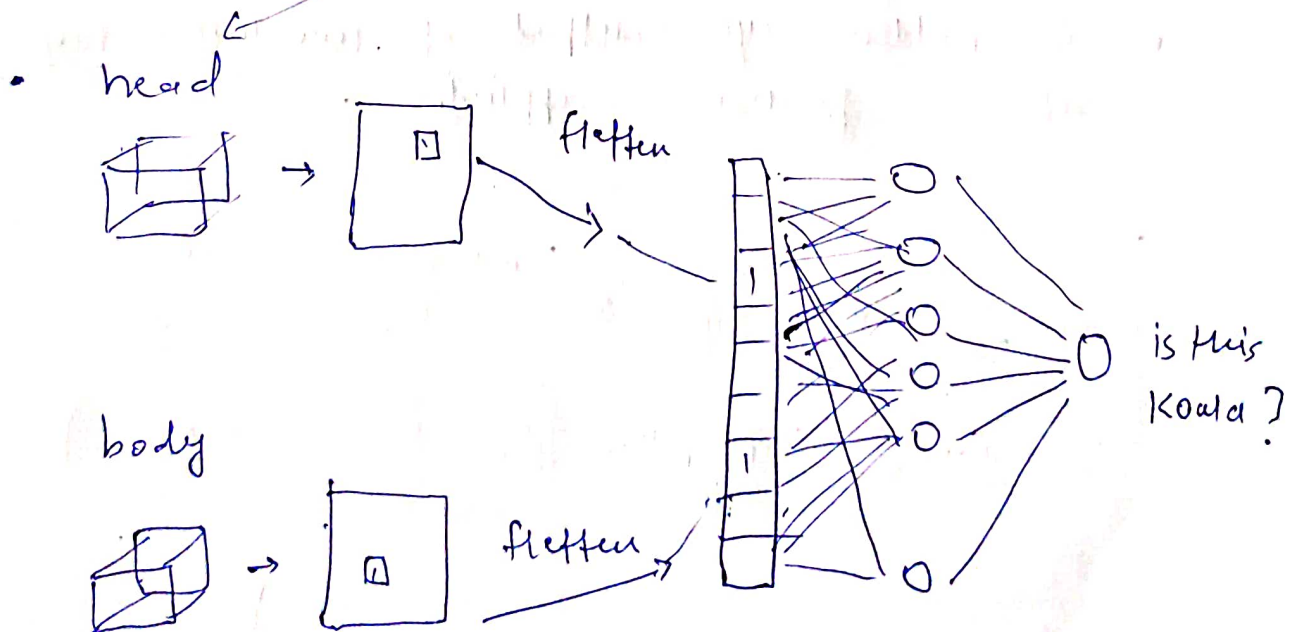


- f filters are used in each convolution layer to detect different features





- In case of above example, filters might end up detecting eyes, nose, ears, Then feature map matrix obtained is again subjected to another convolution possibly to detect head of animal.



- After few repeated convolutional operation, the feature map obtained is flattened & then a Dense layer is applied with sigmoid (binary classification) or softmax (multi classification) activation function.
- with different images the flattened matrix will be different.
- Convolutional layer is responsible for feature extraction
- Dense output layer is responsible for classification

* Activation function is applied to the feature map.

- on hidden layer (output of convolution layer) activation function is applied.

-1	1	1	1	-1
-1	1	1	1	-1
-1	1	1	1	-1
-1	1	1	1	-1
-1	1	1	1	-1
-1	1	1	1	-1

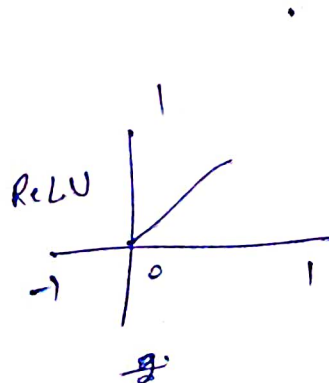
*

1	1	1
-1	1	1
1	1	1

 →
loopy pattern filter

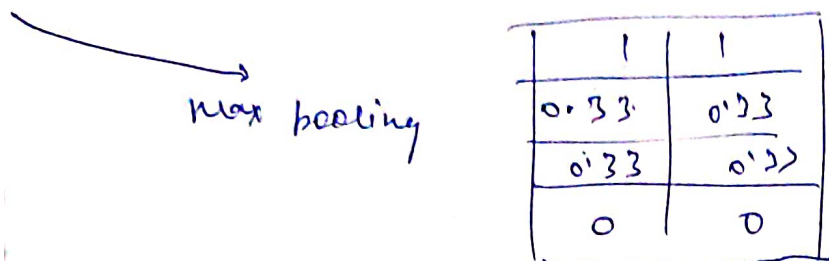
-0.11	1	-0.11
-0.55	0.11	-0.33
-0.33	0.33	-0.33
-0.11	-0.11	-0.22
-0.33	-0.33	-0.22

0	1	0
0	0.11	0
0	0.33	0
0	0	0
0	0	0



* Pooling

- pooling layer is used to reduce the dimension
- on the feature map obtained after convolutional layer, Max Pooling layer is applied to select the maximum value from that window of given stride.

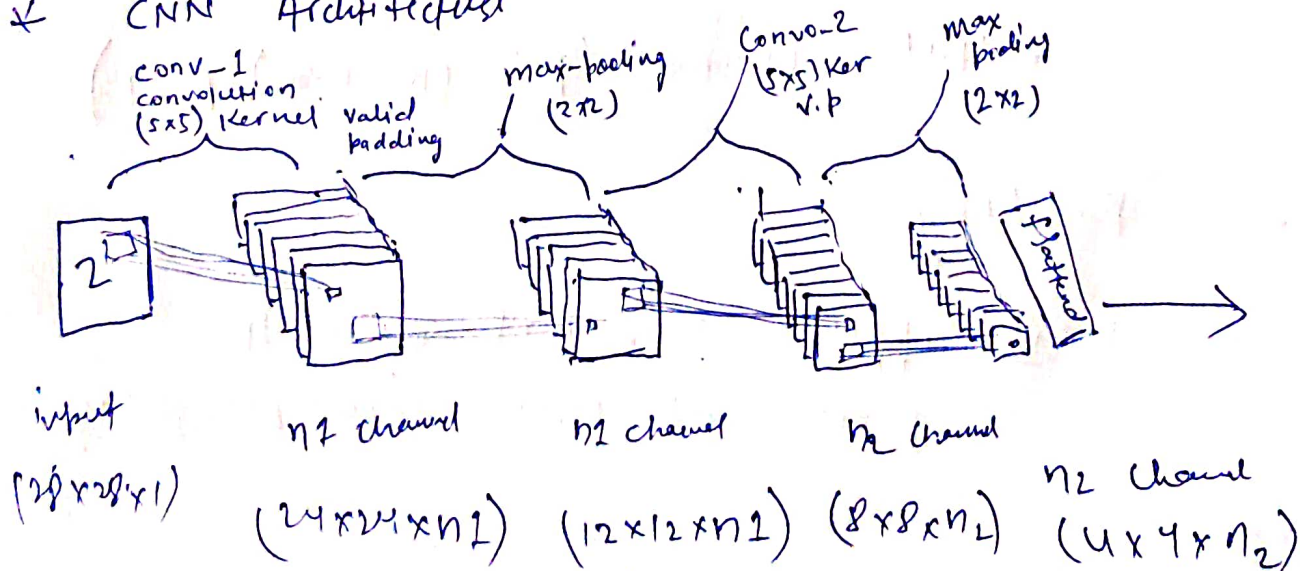


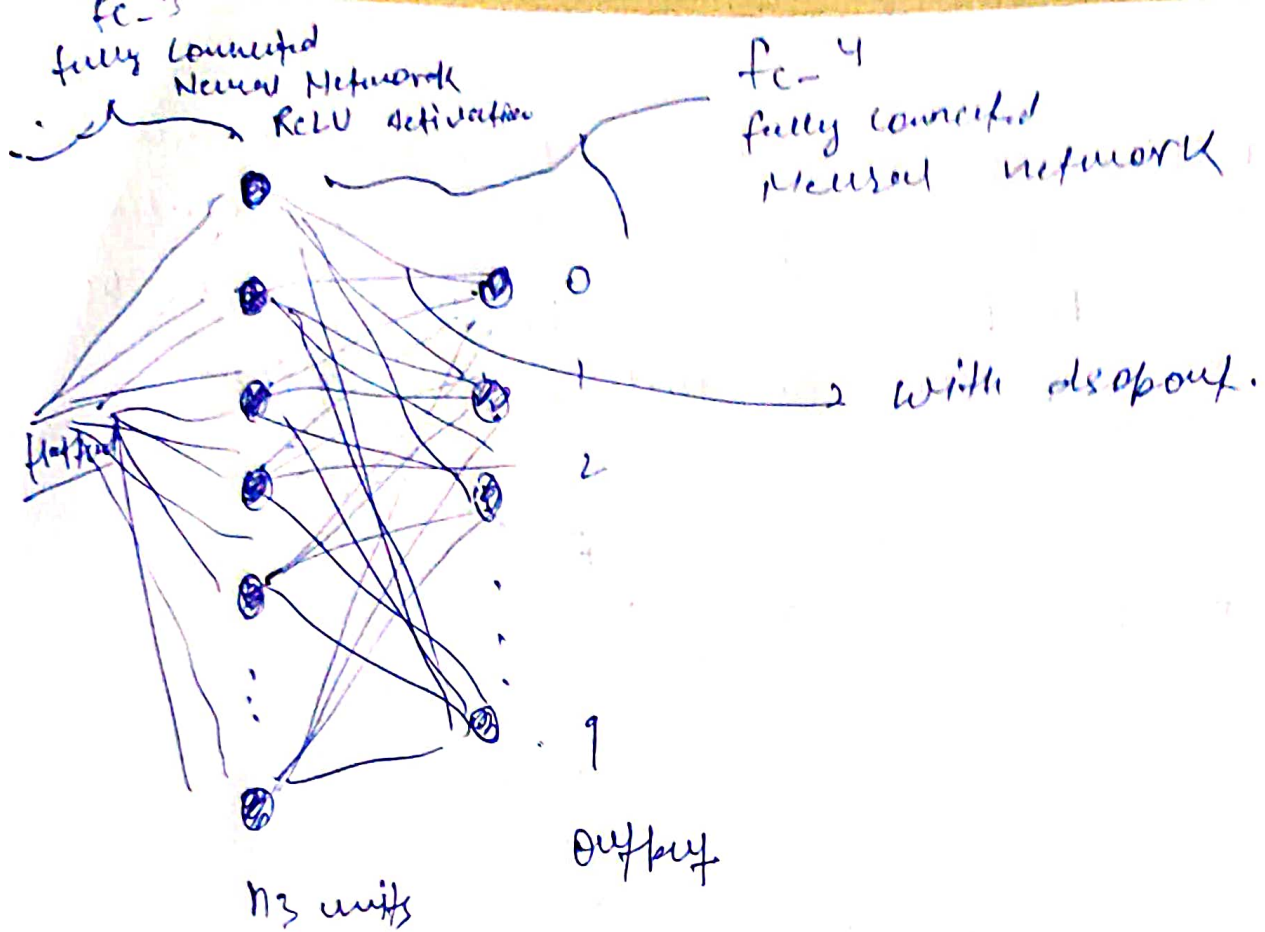
pooling size = $(2, 2)$

stride = 1

- changing the position of q to ~~back~~ back to 1 cell, this type of positional movement is known as translation.
- Max pooling along with convolution helps detect position invariant feature detection.
- pooling also prevents overfitting as there are less parameters.

* CNN Architecture





* Stride

- Stride is number of pixels shifts over the input matrix. when the stride is 1 then we move filter to 1 pixel at a time. when the stride is 2 then we move the filter to 2 pixel at a time & so on.

Input (m, n) , filter (n, n) Stride $= p$

$$\text{output} = (m - n + 1) / p + 1, (m - n + 1) / p + 1$$

* Padding

It is applied so that even the corner pixels get their due weightage.