

* Deep Learning

• Perceptron :-

~~X~~ Single Layer Neural Network

→ This is a type of neural network that performs binary classification that maps input features to output decision.

→ Types :

1) Single Layer Perceptron :-

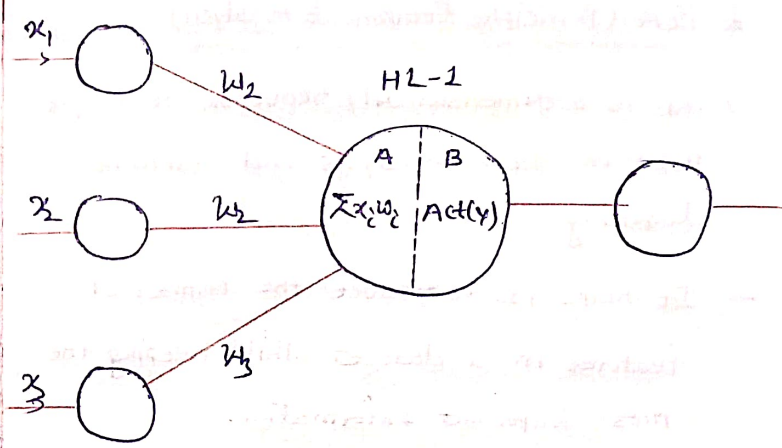
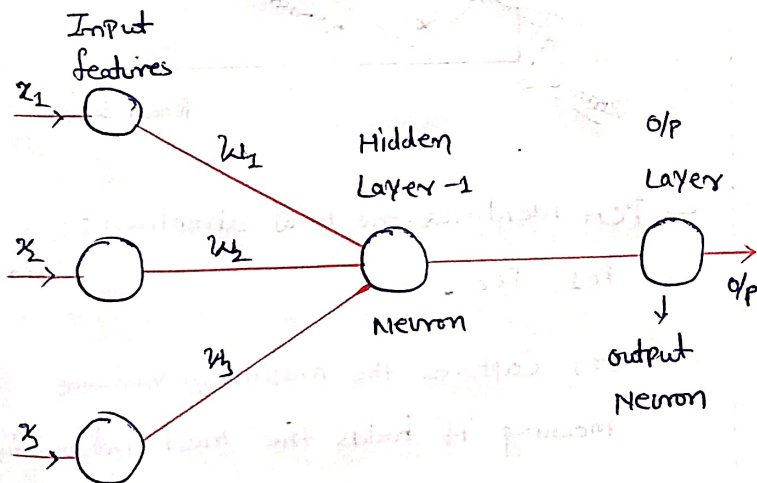
→ This is limited to learning linearly separable patterns.

→ It consists single layer of i/p nodes that are fully connected to a layer of output nodes.

2) Multi-Layer Perceptron :-

→ Consists two or more layers.

→ Adept at handling more complex patterns.



Weighted sum -

$$Z = \sum x_i w_i = w_1 x_1 + x_2 w_2 + \dots + x_n w_n$$

$$= w^T x$$

→ Now Weighted sum passes through a Activation fn that classify the o/p.

$$h(z) = \begin{cases} 0, & z < \text{Threshold} \\ 1, & z > \text{Threshold} \end{cases}$$

↓
heaviside step fn

→ To Enhance flexibility we use bias that helps perceptron to do adjustment independent of input.

$$Z = \sum x_i w_i + b$$

↓
bias

∴ for binary classification we can use sigmoid fn -

$$\text{sigmoid} = \frac{1}{1 + e^{-z}}$$

$$= \frac{1}{1 + e^{-(\sum x_i w_i + b)}}$$

→ This Complete process is called forward propagation.

→ during training, perceptron's weights are adjusted to minimize the difference b/w the predicted output and Actual output.

So To update the weights

Back Propagation happens.

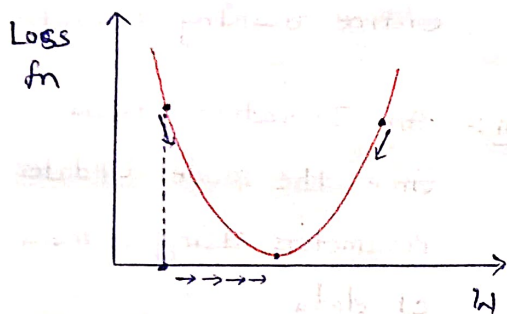
Here we use optimizers.

* Back Propagation -

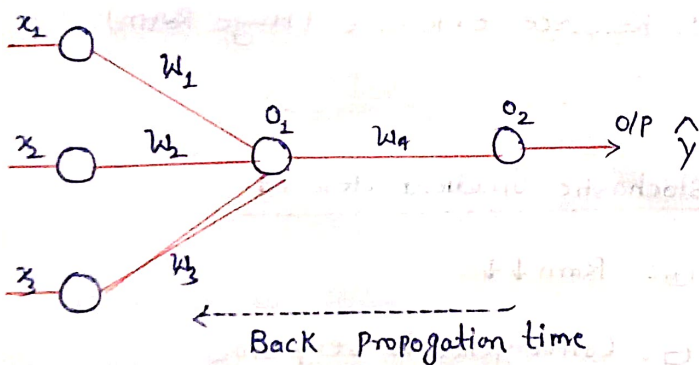
1) Weight updation formula -

$$w_{\text{new}} = w_{\text{old}} - \eta \frac{dL}{dw_{\text{old}}}$$

Learning Rate



2) chain Rule of differentiation:-



$$w_{4\text{new}} = w_{4\text{old}} - \eta \frac{dL}{dw_{4\text{old}}}$$

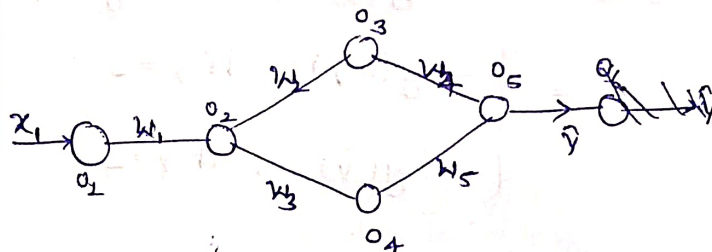
$$\frac{\partial L}{\partial w_{4\text{old}}} = \frac{\partial L}{\partial o_2} * \frac{\partial o_2}{\partial w_{4\text{old}}}$$

$$\therefore z_{o_2} = o_1 w_4 + b_2$$

$$b_{2\text{new}} = b_{2\text{old}} - \eta \frac{\partial L}{\partial b_{2\text{old}}}$$

$$\Rightarrow w_{1\text{new}} = w_{1\text{old}} - \eta \frac{\partial L}{\partial w_{1\text{old}}}$$

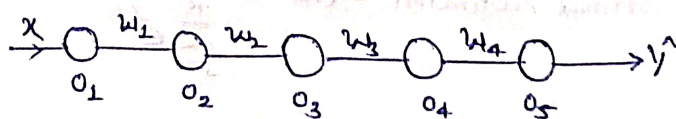
$$\frac{\partial L}{\partial w_{10}} = \frac{\partial L}{\partial o_2} * \frac{\partial o_2}{\partial o_1} * \frac{\partial o_1}{\partial w_{10}}$$



$$\frac{\partial L}{\partial w_{10}} = \left(\frac{\partial L}{\partial o_5} * \frac{\partial o_5}{\partial o_3} * \frac{\partial o_3}{\partial o_2} * \frac{\partial o_2}{\partial w_{10}} \right) + \left(\frac{\partial L}{\partial o_5} * \frac{\partial o_5}{\partial o_4} * \frac{\partial o_4}{\partial o_2} * \frac{\partial o_2}{\partial w_{10}} \right)$$

3) Vanishing Gradient problem :

→ This occurs when gradients become extremely small during back propagation, causing early layers to learn very slowly or stop learning.



$$o_5 = \sigma [(o_4 * w_4) + b]$$

↓
sigmoid fn

$$\therefore 0 < \sigma(y) < 1$$

$$0 < \sigma'(y) < 0.25$$

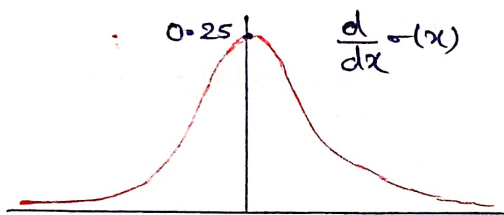
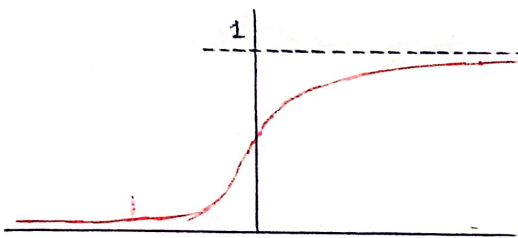
$$\frac{\partial L}{\partial W_{1n}} \rightarrow 0$$

$$W_{1n} \approx W_{10}$$

Commonly used Activation fn:-

1) Sigmoid function

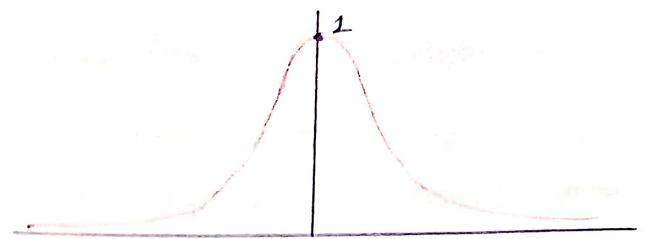
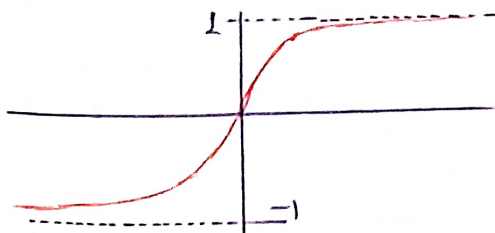
$$\sigma(x) = \frac{1}{1+e^{-x}}$$



- ⊖:
- 1) Prone to gradient ~~dis~~ vanishing
 - 2) Function o/p is not Zero centered
 - 3) Power operation are relatively time Consuming

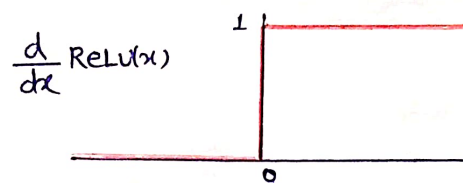
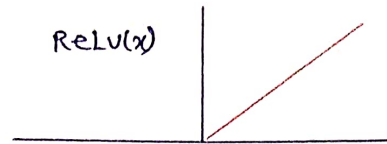
2) tanh function :-

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



3) ReLU function :- (Rectified Linear unit)

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

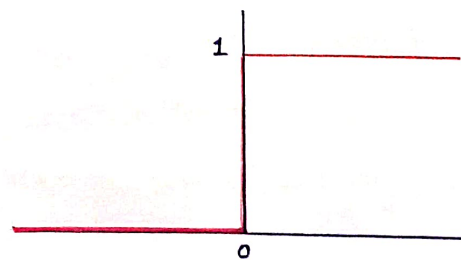
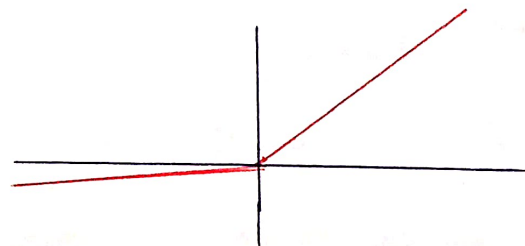


Dead Neuron
Problem

- ⊖:
- 1) Not a Zero centric
 - 2) When the input is 0, ReLU is Completely Inactive.

4) Leaky ReLU function -

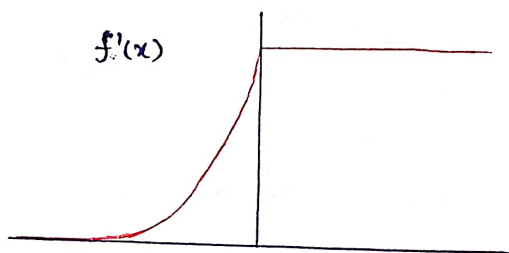
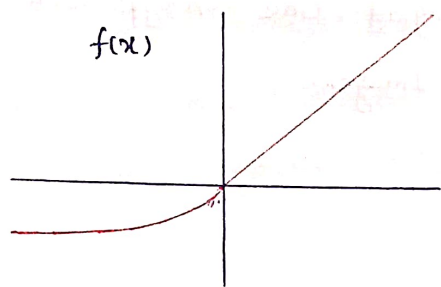
$$f(x) = \max(0.01x, x)$$



→ Solve the dead ReLU Problem.

5) ELU (Exponential Linear units)

$$f(x) = \begin{cases} x & , \text{ if } x > 0 \\ \alpha(e^x - 1) & , \text{ otherwise} \end{cases}$$

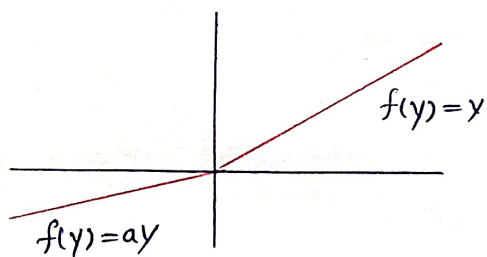


→ No dead ReLU issue

→ Zero centered

6) PReLU (Parametric ReLU)

$$f(y) = \begin{cases} y & , y > 0 \\ ay & , y \leq 0 \end{cases}$$



if $a = 0 \Rightarrow \text{ReLU}$

if $a > 0 \Rightarrow \text{Leaky ReLU}$

if a is a learnable parameter,
 f becomes PReLU

In Hidden Layer we use -

- 1) ReLU
- 2) PReLU
- 3) ELU

In output Layer we can use -

- 1) Sigmoid } Binary classification
- 2) softmax }
- 3) Linear Activation fn } Regression

* Loss function :-

A) Regression :-

1) Mean Squared Err (MSE)

$$\text{Cost fn} = \frac{1}{2n} \sum_{i=1}^n (y_i - \hat{y}_i)^2$$

⊕ : 1) differentiable fn

2) It has only one local/global minima

3) Converge faster

⊖ : 1) Not Robust to outliers.

2) Mean Absolute Err (MAE)

$$\text{Cost fn} = \frac{1}{n} \sum_{i=1}^n |y_i - \hat{y}_i|$$

⊕ : 1) Robust to outliers.

⊖ : 1) Not differentiable

3) Huber Loss

$$\text{Loss} = \begin{cases} \frac{1}{2} (y - \hat{y})^2, & \text{if } |y - \hat{y}| \leq \delta \\ \delta |y - \hat{y}| - \frac{1}{2} \delta^2, & \text{otherwise} \end{cases}$$

B) Classification:-

1) Cross Entropy:-

• Binary Cross Entropy -

↳ for Binary Classification

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

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

$$\hat{y} = \frac{1}{1+e^{-x}}$$

• Categorical Cross Entropy -

↳ for multiclass Classification

$$y_i = [y_{i1}, y_{i2}, y_{i3}, \dots, y_{ic}]$$

$$\text{Loss} = - \sum_{j=1}^c y_{ij} * \ln(\hat{y}_{ij})$$

Where — $\hat{y}_{ij}$ = Softmax Activation
↳ o/p Layer

$$\text{Softmax Activation } \sigma(z) = \frac{e^{z_i}}{\sum_{j=1}^K e^{z_j}}$$

* Optimizers:-

1) Gradient descent

$$W_{\text{new}} = W_{\text{old}} - \eta \frac{\partial L}{\partial W_{\text{old}}}$$

• Batch:- A batch refers to a subset of the entire training dataset that is processed at once during a single forward and backward pass through the neural network.

• Epoch:- An Epoch signifies the Completion of one full cycle through the entire training dataset.

• Iteration:- An Iteration occurs Each time the model updates its Parameters using a mini-batch of data.

→ disadvantage of Gradient descent

1) Resource Extensive (Huge RAM)

2) Stochastic Gradient descent

⊕: Ram ↓ ↓

⊖: Convergence is very slow

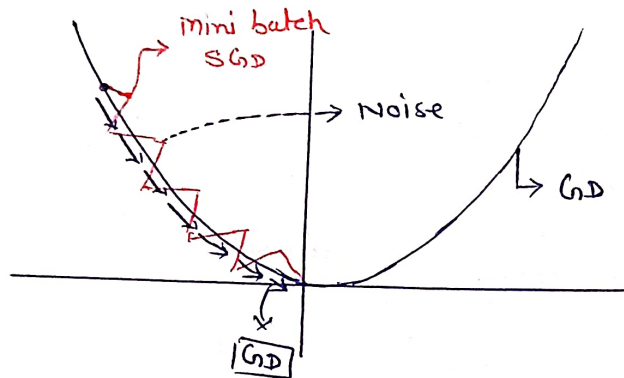
→ In this We train model ~~by~~ ~~the~~ With the 1 ~~batch~~ ~~of~~ records

— Time Complexity ↑ ↑

3) Mini Batch SGD (Stochastic Gradient descent)

→ In this we select batch size and train model by batches.

- ⊕:
- 1) Resource Intensive
 - 2) Convergence will be better
 - 3) Time Complexity will Improve.



4) SGD With momentum :-

$$w_n = w_0 - \eta \frac{\partial L}{\partial w_0}$$

$$\Rightarrow w_t = w_{t-1} - \eta \frac{\partial L}{\partial w_{t-1}}$$

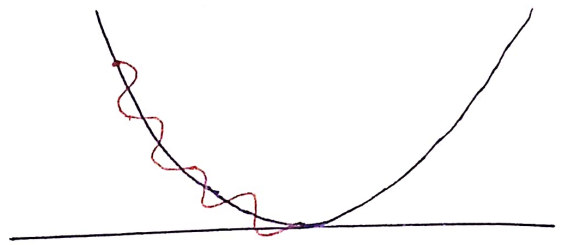
Exponential Weighted Average -

t_1	t_2	t_3		t_n
a_1	a_2	a_3	----	a_n

$$v_{t_1} = a$$

$$v_{t_2} = \beta * v_{t_1} + (1-\beta) * a_2$$

β is a Hyperparameter that tells the importance of current value.



$$w_t = w_{t-1} - v_{dw} * \eta$$

$$v_{dw} = \beta * v_{dw_{t-1}}$$

$$v_{dw_t} = \beta * v_{dw_{t-1}} + (1-\beta) * \frac{\partial L}{\partial w_{t-1}}$$

⊕: Reduce Noise and quick Convergence

5) Adagrad : (Adaptive Gradient descent)-

↳ η is not fixed, It is adaptive

$$w_t = w_{t-1} - \eta' \frac{\partial L}{\partial w_{t-1}}$$

$$\eta' = \frac{\eta}{\sqrt{L_{t+e}}}$$

ϵ : Small Number

$$L_t = \sum_{i=1}^t \left(\frac{\partial L}{\partial w_t} \right)^2$$

→ towards Global minima, η will decrease.

6) Adadelta and Rmsprop

$$\eta' = \frac{\eta}{\sqrt{s_{dw} + \epsilon}}$$

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

↑
Increase by
smaller value

7) Adam optimizer

Momentum + Rmsprop

$$w_t = w_{t-1} - \eta' v_{dw}$$

$$\eta' = \frac{\eta}{\sqrt{s_{dw} + \epsilon}}$$

$$v_{dw_t} = \beta * v_{dw_{t-1}} + (1-\beta) \frac{\partial L}{\partial w_{t-1}}$$

⊕: 1) Smoothing

2) Learning Rate becomes
Adaptive

* CNN (Convolutional Neural Networks)

↳ It is a deep Learning models designed to process data with a grid-like topology such as Images.