

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

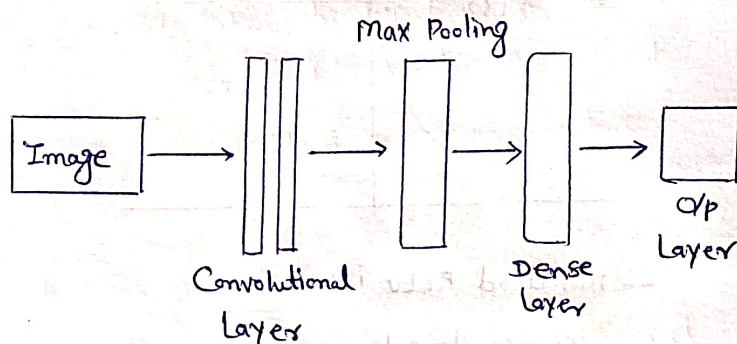
2) Learning Rate becomes

Adaptive

* CNN (Convolutional Neural Networks)

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

→ CNN Consist of multiple layers like Input Layer, Convolution Layer, Pooling Layers, Dense Layer.



* Convolution

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

*

1	0	-1
2	0	-2
1	0	-1

⇒ verticle
Edge
filter

⇓

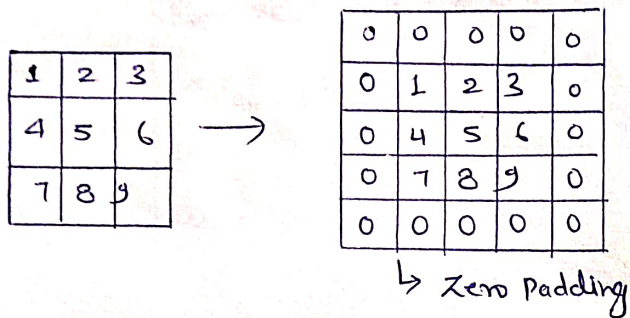
0	-4	-4	0
0	-4	-4	0
0	-4	-4	0
0	-4	-4	0

→ If Image size = n and filter size = f

then output size = $n - f + 1$

→ This reduce the Image size, that lost some data.

↓
so we add padding in the Image.



$$\text{output size} = \underline{n + 2p - f + 1}$$

• Activation Layer -

→ Applied after the Convolution

→ Activation Layer add non linearity to the network.

→ It will apply on element wise activation function to the output of the Convolution Layer.

→ Some Common Activation fn -

RELU : $\max(0, x)$

Tanh, Leaky RELU

• Pooling Layer :-

→ This Layer reduces the spatial size which makes computation fast, reduces memory and also prevents overfitting.

→ Types :-

1) max pooling : picks the max value

1	1	2	4
5	6	7	8
3	2	1	0
1	2	3	4

max pool with 2×2 filters and stride 2

6	8
3	4

2) Average Pooling : picks Takes the Average of all value.

→ only height and width ~~shrink~~ shrink, channels stay same.

Input : $\frac{28}{w} \times \frac{28}{h} \times \frac{32}{\text{Channels}}$
(After Applying Convolution)

max pool (2×2)
stride = 2

output : $\underline{14 \times 14 \times 32}$

Note :

Convolution output size -

H_1 : Input Height

W_1 : Input Width

C_1 : Input channels

Convolution With-

F : Filter / Kernel size ($F \times F$)

P : Padding

S : Stride

K : Number of filters

$$\text{output Height } (H_2) = \left\lfloor \frac{H_1 - F + 2P}{S} \right\rfloor + 1$$

$$\text{output Width } (W_2) = \left\lfloor \frac{W_1 - F + 2P}{S} \right\rfloor + 1$$

$$\text{output Channels } (C_2) = K$$

Exa- Image - 6×6
Filter - 3×3

$$P = 0$$

$$S = 2$$

$$\text{output Height} = \left\lfloor \frac{6 - 3 + 0}{2} \right\rfloor + 1$$

$$= 1 + 1$$

$$= 2$$

• Dropout Layer:-

→ This is a Regularization Layer, that prevents model to overfit.

→ This Layer Randomly turns off some Neurons. (e.g. - 30%)

→ ⊕: 1) prevents overfitting

2) Focus Network to Learn more Robust features

• Flatten Layer:-

→ The Resulting feature maps are flattened into 1D vector after Convolution or ~~max~~ pooling layers so they can be passed into a completely linked layer for categorization or regression.

• Fully Connected Layer (Dense Layer):-

→ It takes the input from the previous Layer and computes the final classification or regression tasks.

* VGG Architecture:-

→ VGG 16 is a deep Convolutional neural network designed for Image Classification tasks.

→ VGG 16 Consists 16 Layers, including 13 Convolutional Layers and 3 fully Connected Layers.

These Layers are organized into blocks.

Block 1

- Conv, 3×3 , 64 filters
- Conv 3×3 , 64 filters
- max pool 2×2

Block 2

- Conv 3×3 , 128 filters
- Conv 3×3 , 128 filters
- max pool 2×2

Block 3

Conv 3x3, 256 filters
Conv 3x3, 256 filters
Conv 3x3, 256 filters
max pool, 2x2

Block 4

Conv 3x3, 512 filters
Conv 3x3, 512 filters
Conv 3x3, 512 filters
max pool 2x2

Block 5

Conv 3x3, 512 filters
Conv 3x3, 512 filters
Conv 3x3, 512 filters
max pool 2x2

↓
flatten

↓
Dense Layer

