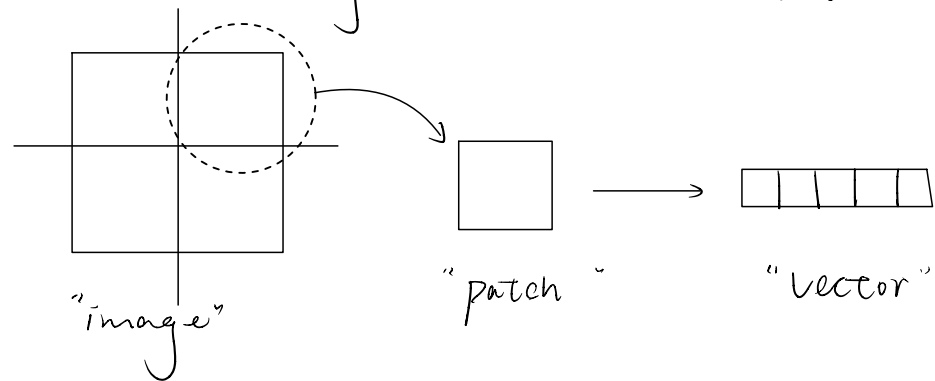


ViT: Vision Transformer

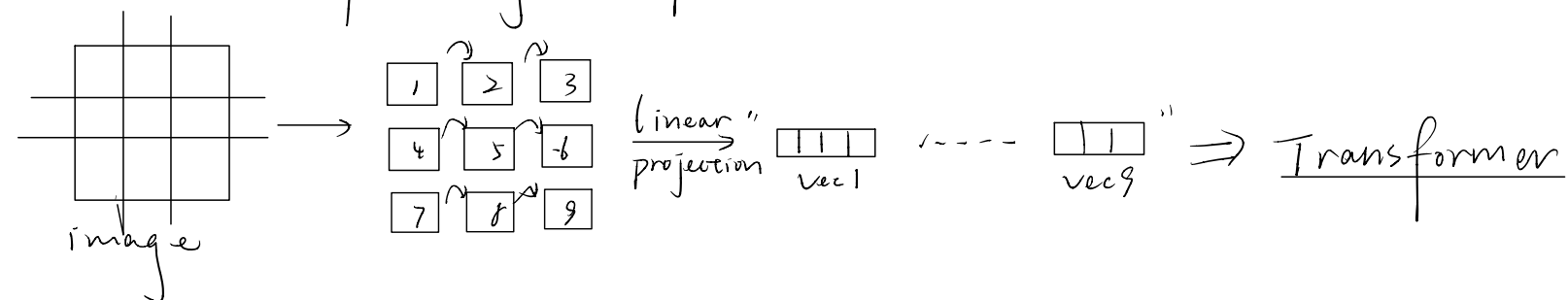
image: 3D tensor.

Key Question: How to apply Self-Attention to images?

Answer: Image $\Rightarrow$ Patches $\Rightarrow$ Vectors



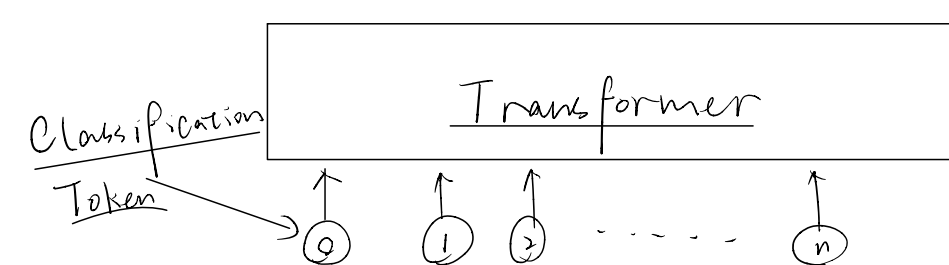
Order: Left to right, top down



Linear projection: initial patch embedding (akin to word embedding)

Position embedding: learned on manually set ($L \rightarrow R, T \rightarrow B$)

An extra "Classification Token" is added



Input X : $H \times W \times C$

patch size: $P \times P$.

$\Rightarrow$ # of patches $\frac{H \times W}{P^2} \Rightarrow$

flatten vector: $1 \times P^2 \times C$ multiplied with $E \in \mathbb{R}^{(P^2 \times C) \times D}$

$\Rightarrow$ D-dimensional vector.

$$Z_0 = \begin{bmatrix} x_{cls} E \\ x_1^T E \\ \vdots \\ x_p^N E \end{bmatrix} + E_{pos} \quad (E, E_{pos} \text{ are learnable})$$

$$Z_L' = \text{MSA}(\text{LN}(Z_{L-1})) + Z_{L-1} \quad \text{Residual Connection}$$

MSA: Multi-Head Self-Attention.
LN: Layer-Normalization

$$Z_L = \text{MLP}(\text{LN}(Z_L') + Z_L') \quad (L \in 1, 2, \dots, L)$$

MLP: contains two layers with GELU non-linearity

ViT has less Image-Specific inductive bias (than CNN)

CNN: locality, two-dimensional neighbor structure, translation equivariance (Inductive Bias)

ViT: MLP layer

MSA layer $\Rightarrow$ Global

$\Rightarrow$ Weak Assumption $\rightarrow$ Need more data

Small dataset $\Rightarrow$ CNN better

Large dataset $\Rightarrow$ ViT better

Learning Pattern of ViT

recover the grid structure of original image

Lower layer $\rightarrow$ local/global feature (Both)
higher layer $\rightarrow$ global feature (only)

ViT features are more interpretable, (more meaningful)
ViT is more robust to noise & domain shift

CLIP: Contrastive Language-Image Pre-training

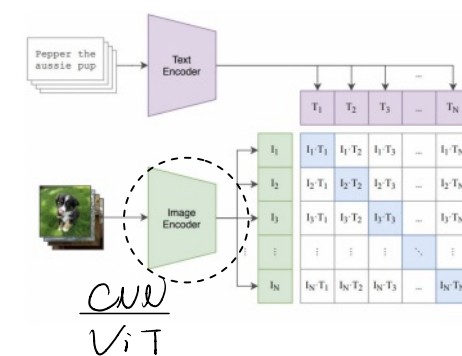
training on image-text pairs

$\Rightarrow$ Vision-Language Model (VLM)

Contrastive Loss function

$$L_{\text{CLIP}} = -\frac{1}{2N} \sum_{i=1}^N \log \frac{\exp(L_i \cdot T_i)}{\sum_{j=1}^N \exp(L_i \cdot T_j)} - \frac{1}{2N} \sum_{i=1}^N \log \frac{\exp(L_i \cdot T_i)}{\sum_{j=1}^N \exp(L_j \cdot T_i)}$$

$P(T_i | I_i) \qquad P(I_j | T_j)$



$\left\{ \begin{array}{l} \text{for each image } I_i \text{ the paired text } T_i \text{ should have the maximum probability} \\ \text{for each text } T_i \text{ the paired image } I_i \text{ should have the maximum probability} \end{array} \right.$

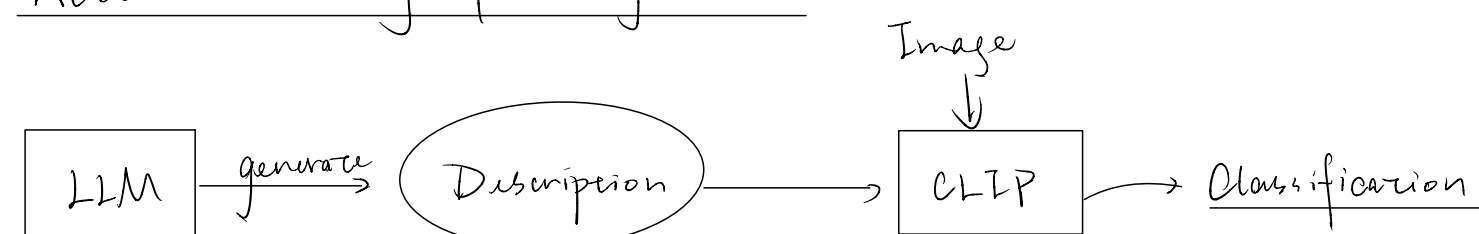
Image feature $\rightarrow$ Classification Head

Zero-Shot CLIP shows Competitive performance

Linear Probing with CLIP outperform the best ImageNet Model

CLIP is more robust to distributional shift

Alternative way of using CLIP:



CLIP Scaling with larger training dataset