

Vision Transformer

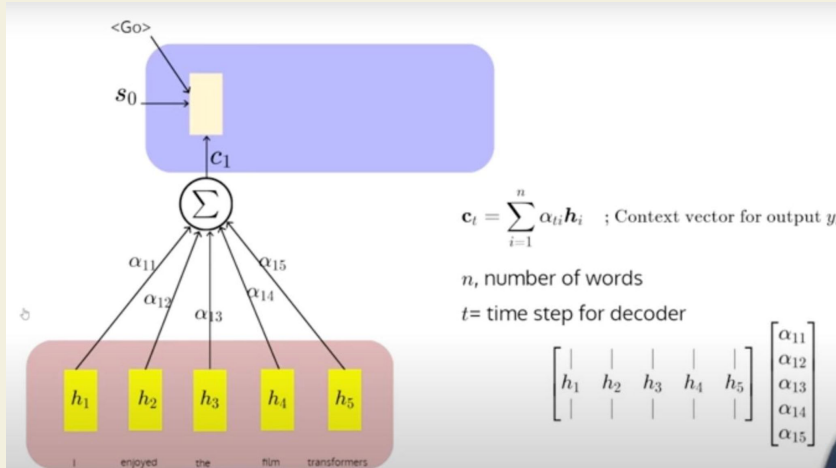
By Jayin Khanna

1. Defining the Problem
2. **Making Observations**
3. Forming a Hypothesis
4. Experiment Results
5. Drawing a Conclusion

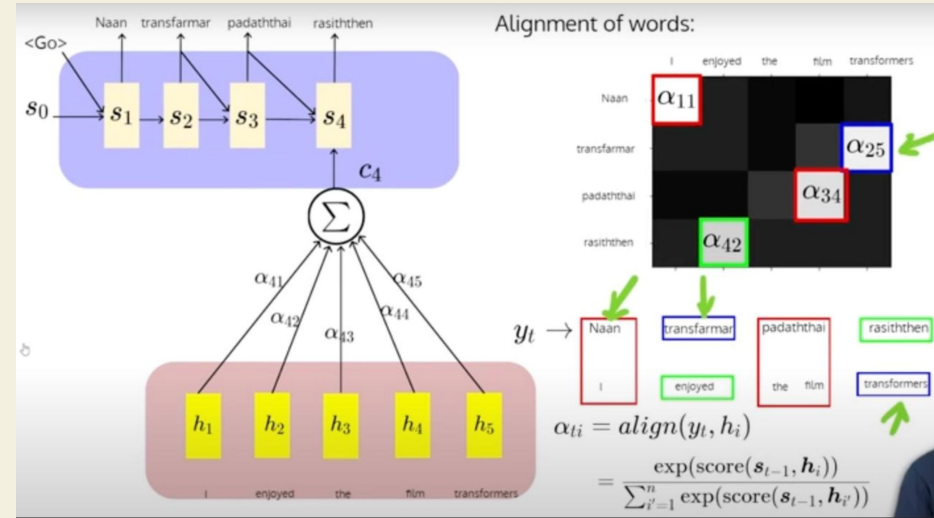
Attention and transformers

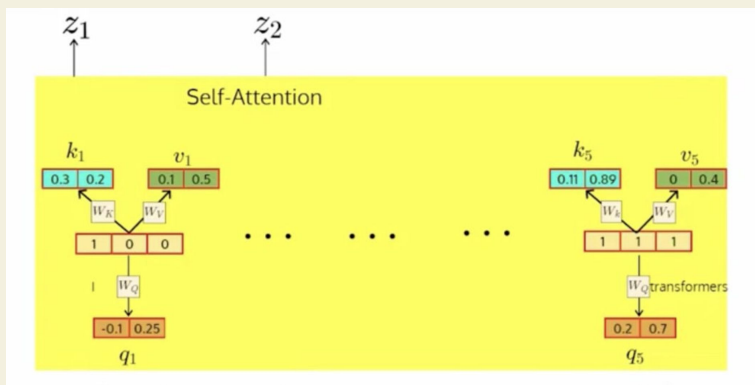
Why not RNNs:

1. We want contextual representation ie. we want an improved representation based on context and
2. We don't want sequential processing; instead we want parallelizability in the architecture

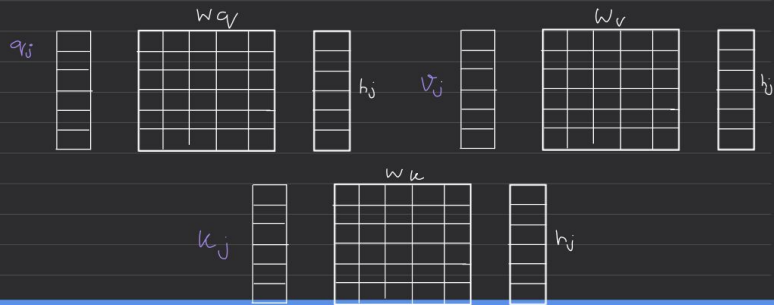


Where H_i : the word embedding of the i th token in the sentence.
And alpha is the contribution of each of the other token embeddings on the current token embedding





*) Transformation Matrices:



*) Attention Mechanism:

Given: $X = \{x_1, x_2, \dots, x_r\}$; $x_i \in \mathbb{R}^d$

→ First compute three sets of representations

- key
- query & iii) value

Let input be $X \in \mathbb{R}^{T \times d}$

projections q, k, v :

$$q = XW_q, \quad k = XW_k, \quad v = XW_v$$

$d \times d_q \quad d \times d_k \quad d \times d_v$

typically $d_q = d_k = d_v$

Each row of the above matrices correspond to the representations of every single token.

Attention:

given a query vector $q \in \mathbb{R}^{d_q}$, (a row of q)

a set of all the key & value vectors

$k_i \in \mathbb{R}^{d_k}, v_i \in \mathbb{R}^{d_v}$ (rows of matrices K & V)

- Find the similarities b/w the given query & all the keys

$$S = q^T K^T \quad \text{where given query } q \in \mathbb{R}^{1 \times d_q}$$

$K \in \mathbb{R}^{T \times d_k}$

here, $S \in \mathbb{R}^{1 \times T}$

S_i : similarity b/w q & k_i

2.) Scale the dot product to normalize

$$\hat{S} = \frac{q^T k^T}{\sqrt{d_v}}$$

3.) normalize values of $\hat{S}$ to have values $1/w$ $0 \leq 1$ & sum to 1

$$\text{Softmax}(\hat{S}) : \mathbb{R}^T \rightarrow \mathbb{R}^{[0,1]}$$

$$\alpha_i = \frac{\exp\left(\frac{q^T k_i}{\sqrt{d_v}}\right)}{\sum_{j=1}^T \exp\left(\frac{q^T k_j}{\sqrt{d_v}}\right)} ; 0 \leq \alpha_i \leq 1, \sum_{i=1}^T \alpha_i = 1$$

4.) define Attention:

$$\text{Attention}(q, k, v) = AV = \sum_{i=1}^T \alpha_i v_i = z \in \mathbb{R}^{T \times d_v}$$

$$\text{Attn}(\cdot) : X^{T \times d_m} \rightarrow Z^{T \times d_v} \leftarrow \text{better representation with contextual info}$$

Attention & query tokens:

$$q \in \mathbb{R}^{T \times d_v}$$

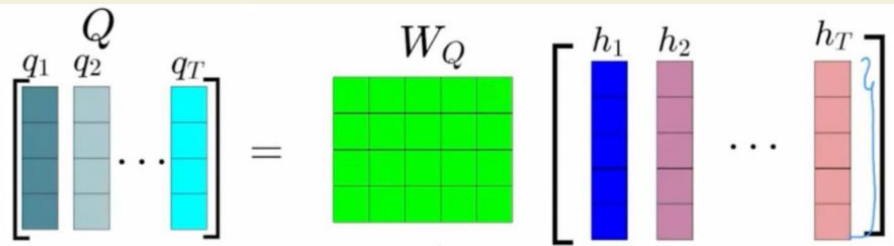
$$\text{Attention}(q, k, v) = \left[\text{Softmax}\left(\frac{q k^T}{\sqrt{d_v}}\right) \right] v$$

The goal

However, now both the vectors (s_i) and (h_j) , for all i, j are available for all the time (whereas in the seq2seq model, h_j for all j were available and s_i was obtained one at a time)

	h_j							
	The	animal	didn't	cross	the	street	because	it
s_i	The	0.6	0.1	0.05	0.05	0.02	0.02	0.1
	animal	0.02	0.5	0.06	0.15	0.02	0.05	0.12
	didn't	0.01	0.35	0.45	0.1	0.01	0.02	0.03
	cross	.						
	the	.						
	street	.						
	because	.						
	it	0.01	0.6	0.02	0.1	0.01	0.2	0.01

earlier
utilization
not possible
due to
dependencies
T x T



Multi-Head Attention

* Multi-head Attention:

→ Define $h \in \mathbb{N}$, as the no. of Attn heads
typically $d_k = d_v / h$

for each of the h heads:

$$q_j = X W_j^q \in \mathbb{R}^{T \times d_v}$$

$$k_j = X W_j^k \in \mathbb{R}^{T \times d_k}$$

$$v_j = X W_j^v \in \mathbb{R}^{T \times d_v}$$

$$z_j = \text{Attn}(q_j, k_j, v_j) \in \mathbb{R}^{T \times d_v}$$

All individual heads are concatenated

$$z_c = [z_1, z_2, \dots, z_h] \in \mathbb{R}^{T \times (d_v \times h)} \\ = \mathbb{R}^{T \times d_n}$$

$$\text{Multihed Attn}(x) = z_c W^c \in \mathbb{R}^{T \times d_n}$$

$$W^c \in \mathbb{R}^{(d_v \times h) \times d_n}$$

Here W_j^q, W_j^k, W_j^v & W^c are learned
for $j = 1, 2, \dots, h$

We have obtained $z \in \mathbb{R}^{T \times d_n}$ which is a representation
of data given by the multihed attn

Why multi-head attention (vs one head)

Multiple views / subspaces leads to a more expressive power of the network. Each head linearly projects inputs into a different subspace, so different heads can detect different kinds of relations (short-range vs long-range, syntax vs semantics, exact token matches vs paraphrase signals).

Sentence (tokens):

[The] [cat] [that] [the] [dog] [chased] [yesterday] [slept] [because]
[it] [was] [tired]

What we want the model to capture simultaneously

- **Long-range dependency:** *cat* → *slept* (subject → verb)
Coreference: *it* → *cat*
- **Modifier attachment:** *yesterday* → *chased*

Three attention heads (each learns one pattern):

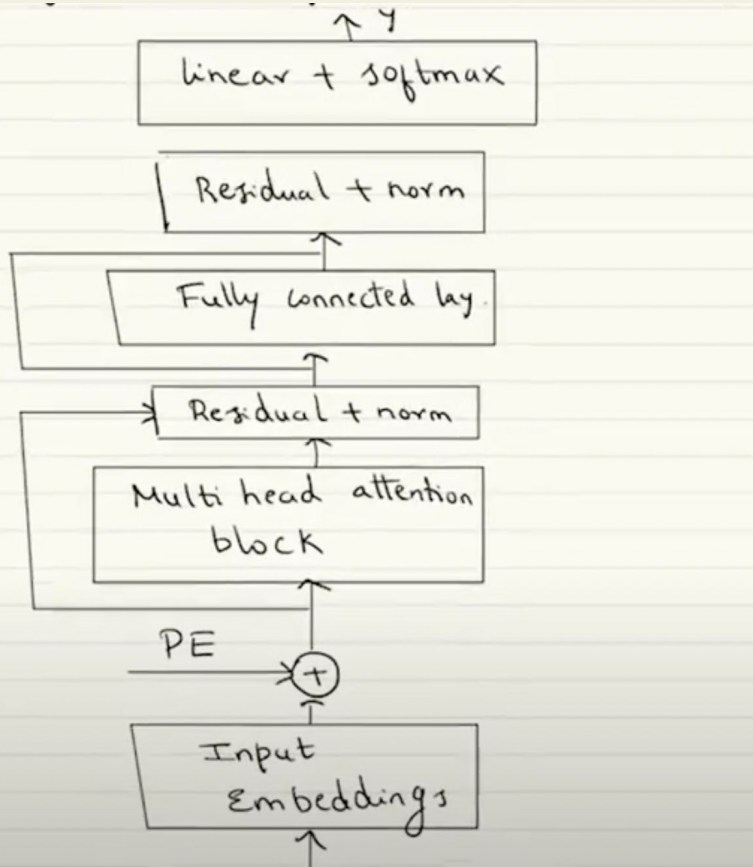
- **Head 1 (subject-verb):** Strong attention between *cat* & *slept*.
e.g. $\text{att}(\text{cat} \rightarrow \text{slept}) \approx 0.85$, others small.
- **Head 2 (coreference):** Strong attention *it* → *cat* (and *cat* ← *it*). e.g. $\text{att}(\text{it} \rightarrow \text{cat}) \approx 0.90$.
- **Head 3 (temporal modifier):** *yesterday* → *chased*.
e.g. $\text{att}(\text{yesterday} \rightarrow \text{chased}) \approx 0.95$.

What a single head would do A *single* attention map must combine all signals into one distribution. That often leads to a blurred map where *cat* spreads attention to *slept*, *dog*, and *it* with medium weights (e.g. 0.35, 0.25, 0.20) — none are as sharp as separate heads.

Visual intuition (arrows):

- Head 1: *cat* → *slept*
- Head 2: *it* → *cat*
- Head 3: *yesterday* → *chased*
- Single head: *cat* → *slept*, *cat* → *dog*, *cat* → *it* (all weaker / mixed)

Transformer Architecture



* Fully Connected layers:

for each of the position j in the seqs
let $\tilde{z}_j$ be the j th token, $\tilde{z}_j \in \mathbb{R}^{d_m}$

Where $\sigma(\cdot)$ is a non-linearity
 W_1 & W_2 - learnable

Output of FCL, $\hat{\tilde{z}}_j \in \mathbb{R}^{T \times d_m}$

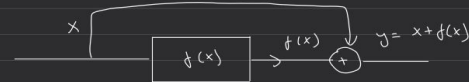
Output layer:

$$\hat{\tilde{z}}_j = \hat{\tilde{z}}_j W_p + b_p \in \mathbb{R}^V ; W_p \in \mathbb{R}^{d_m \times V}$$

$\hat{\tilde{z}}_j$ (o/p of FCL) is projected on to the space of vocab via a linear proj W_p .

$$\tilde{y}_j = \text{softmax}(\hat{\tilde{z}}_j)$$

* Skip connections & Normalization



Skip/residual connection

* Layer Normalization

Given $x \in \mathbb{R}^{d_m}$

$$\mu = \frac{1}{d_m} \sum_{k=1}^{d_m} x_k ; \sigma^2 = \frac{1}{d_m} \sum_{k=1}^{d_m} (x_k - \mu)^2$$

$$\text{layernorm}(x) = \gamma \cdot \left[\frac{x - \mu}{\sqrt{\sigma^2 + \epsilon}} \right] + \beta$$

$\epsilon \rightarrow \text{small}$

γ & β are learnable params

if $\text{Sublayer}(x)$ represents a block in the architecture (atom of FCL)

$$\text{Output} = \text{layernorm} \left(x + \underbrace{\text{Sublayer}(x)}_{\text{Residual connection}} \right)$$

to ensure that the sequential ordering is preserved, an additional fixed temporal embedding vector is added to the input

positional encoding/Embedding:

→ Sinusoidal Encoding:

Sps. j denotes the pos in the input seqⁿ
& i denotes the embedding dim

then

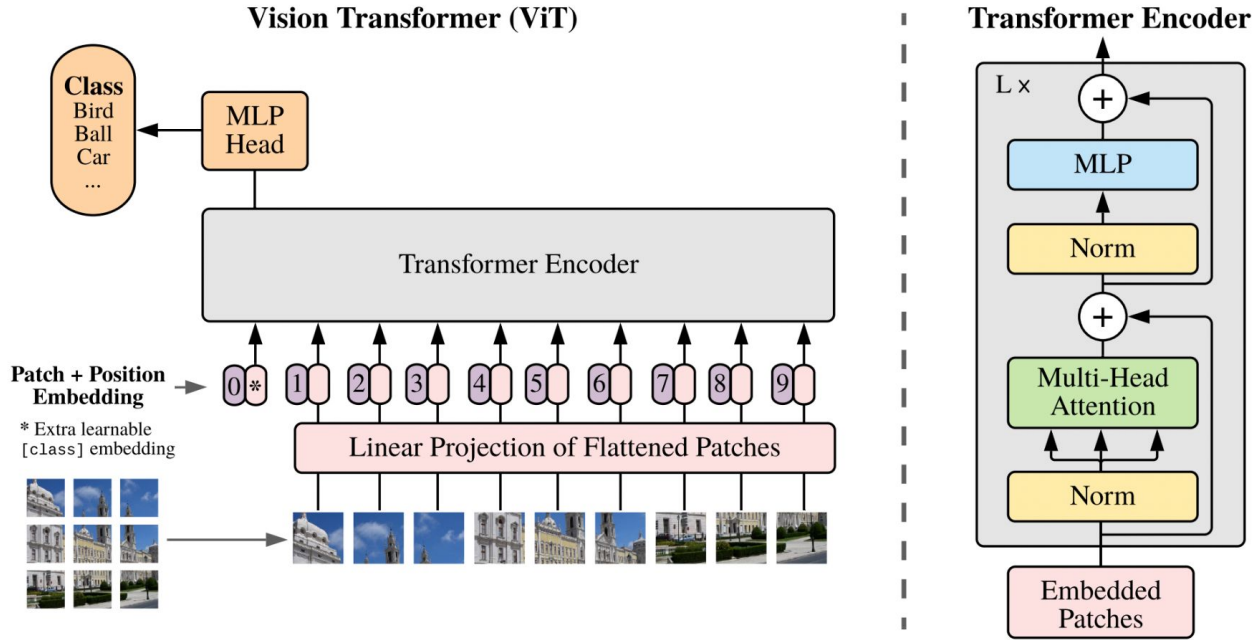
$$\text{PE}(j, 2i) = \sin\left(\frac{j}{10,000^{2i/d_m}}\right) \quad \text{PE}() \in \mathbb{R}^{T \times d_m}$$

$$\text{PE}(j, 2i+1) = \cos\left(\frac{j}{10,000^{2i/d_m}}\right)$$

$$x_j = x_0 + \text{PE}(j) \in \mathbb{R}^{d_m}$$

$$\forall j \in \{1, 2, \dots, T\}$$

Vision Transformer (ViT)



Why Transformers for Images?

Transformers are powerful for 1-D sequences (e.g., text) → self-attention enables global context modeling across tokens.

Vision Transformer (ViT) extends this idea to images by **reformulating an image as a sequence of tokens** (patches) and then applying a **standard Transformer encoder**.

Core motivation:

- CNNs have strong **inductive biases** (locality, translation equivariance), but these can limit scalability.
- Transformers need more data and hence are data-hungry but extremely **scalable** with compute and dataset size.

Minimal architectural changes from standard Transformer:

- Break image into fixed-size **patches** (e.g., 16×16).
- Flatten patches and **linearly project** them into embedding space → token sequence.
- Prepend a **learnable [CLS] token** to represent the entire image.
- Add **learnable positional embeddings** to retain spatial structure.
- Feed resulting sequence into a **stack of Transformer encoder layers**.
- Use the final [CLS] token for classification via a simple head.

Key insight: No convolutions are required — the **attention layers learn spatial relationships** directly from data.

Result: With large-scale pretraining, ViT **matches or surpasses** CNNs on image classification benchmarks.

- **Goal:** Convert an image into a **token sequence** that can be processed by a standard Transformer encoder.
- **Given:**
 - Image $x \in \mathbb{R}^{H \times W \times C}$
 - Patch size $P \times P$
 - Number of patches $N = H \cdot W / P^2$
 - Embedding dimension D
- **Step 1: Patch Extraction**
 - Split the image into N non-overlapping patches of size $P \times P$
 - Flatten each patch $\rightarrow$ vector of length $P^2 \cdot C$
- **Step 2: Linear Projection**
 - Each flattened patch is linearly projected using $E \in \mathbb{R}^{(P^2 C) \times D}$
 - Produces **patch embeddings**: $x_p E \in \mathbb{R}^D$
 - Add **learnable positional embeddings** $E_{pos} \in \mathbb{R}^{(N+1) \times D}$ to retain spatial structure.
- **Input to Transformer:**

$$z_0 = [x_{class}; x_{p1E}; x_{p2E}; \dots; x_{pNE}] + E_{pos}$$

Positional Embeddings

Why Positional Embeddings?

- Transformer attention is **permutation-invariant** $\rightarrow$ it doesn't know the spatial arrangement of patches by default.
- Positional embeddings inject **spatial information** into the token sequence.

Design in ViT

- ViT uses **learnable 1-D positional embeddings** $E_{pos} \in \mathbb{R}^{(N+1) \times D}$ (one for each patch)
- These embeddings are **added** to patch embeddings at the input stage.

The MLP contains two layers with a GELU non-linearity.

$$\mathbf{z}_0 = [\mathbf{x}_{\text{class}}; \mathbf{x}_p^1 \mathbf{E}; \mathbf{x}_p^2 \mathbf{E}; \cdots; \mathbf{x}_p^N \mathbf{E}] + \mathbf{E}_{pos}, \quad \mathbf{E} \in \mathbb{R}^{(P^2 \cdot C) \times D}, \mathbf{E}_{pos} \in \mathbb{R}^{(N+1) \times D} \quad (1)$$

$$\mathbf{z}'_\ell = \text{MSA}(\text{LN}(\mathbf{z}_{\ell-1})) + \mathbf{z}_{\ell-1}, \quad \ell = 1 \dots L \quad (2)$$

$$\mathbf{z}_\ell = \text{MLP}(\text{LN}(\mathbf{z}'_\ell)) + \mathbf{z}'_\ell, \quad \ell = 1 \dots L \quad (3)$$

$$\mathbf{y} = \text{LN}(\mathbf{z}_L^0) \quad (4)$$