

Vision Transformers and Multimodal Models

From image patches to CLIP and vision-language models

Prof. Nipun Batra

ES 667 · Deep Learning · IIT Gandhinagar

From CNNs to Vision Transformers

Where we are

We have two big architectures: **CNNs** for images (L8) and **Transformers** for sequences (L15–16).

Where we are

We have two big architectures: **CNNs** for images (L8) and **Transformers** for sequences (L15–16).

CNN : $I \in \mathbb{R}^{H \times W \times C} \longrightarrow \text{conv, pool}$ **Transformer** : $X \in \mathbb{R}^{T \times d} \longrightarrow \text{SelfAttention}(X)$

Where we are

We have two big architectures: **CNNs** for images (L8) and **Transformers** for sequences (L15–16).

CNN : $I \in \mathbb{R}^{H \times W \times C} \longrightarrow \text{conv, pool}$ **Transformer** : $X \in \mathbb{R}^{T \times d} \longrightarrow \text{SelfAttention}(X)$

The Transformer consumes a **sequence of tokens**. It knows nothing about grids, pixels, or convolutions. This lecture asks a simple question:

Where we are

We have two big architectures: **CNNs** for images (L8) and **Transformers** for sequences (L15–16).

CNN : $I \in \mathbb{R}^{H \times W \times C} \longrightarrow \text{conv, pool}$ **Transformer** : $X \in \mathbb{R}^{T \times d} \longrightarrow \text{SelfAttention}(X)$

The Transformer consumes a **sequence of tokens**. It knows nothing about grids, pixels, or convolutions. This lecture asks a simple question:

can we make an **image** look like a **sequence** — and run a plain Transformer on it?

An image is a grid; a Transformer wants a sequence

The mismatch is structural:

An image is a grid; a Transformer wants a sequence

The mismatch is structural:

CNN view — an image is a **2-D grid**; a conv kernel mixes each pixel with its **local** neighbours.

Transformer view — the input must be a **sequence of vectors** $x_1, \dots, x_T$, each a **token**.

An image is a grid; a Transformer wants a sequence

The mismatch is structural:

CNN view — an image is a **2-D grid**; a conv kernel mixes each pixel with its **local** neighbours.

Transformer view — the input must be a **sequence of vectors** $x_1, \dots, x_T$, each a **token**.

so we need to turn the $H \times W \times C$ grid into a **short list of visual tokens**

Vision Transformer (Dosovitskiy et al. 2021): cut the image into **patches** and treat each patch as a token.

Vision Transformer (Dosovitskiy et al. 2021): cut the image into **patches** and treat each patch as a token.

1. split the image into a grid of $P \times P$ patches;

Vision Transformer (Dosovitskiy et al. 2021): cut the image into **patches** and treat each patch as a token.

1. split the image into a grid of $P \times P$ patches;
2. **flatten** each patch and **linearly project** it to an embedding;

Vision Transformer (Dosovitskiy et al. 2021): cut the image into **patches** and treat each patch as a token.

1. split the image into a grid of $P \times P$ patches;
2. **flatten** each patch and **linearly project** it to an embedding;
3. feed the resulting sequence of patch tokens to a **pure Transformer encoder**.

Vision Transformer (Dosovitskiy et al. 2021): cut the image into **patches** and treat each patch as a token.

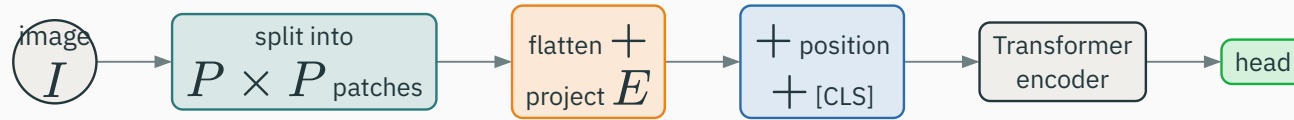
1. split the image into a grid of $P \times P$ patches;
2. **flatten** each patch and **linearly project** it to an embedding;
3. feed the resulting sequence of patch tokens to a **pure Transformer encoder**.

no convolutions — a patch is to an image what a **word is to a sentence**

The whole ViT front-end, end to end:

The image-to-token pipeline

The whole ViT front-end, end to end:



The image-to-token pipeline

The whole ViT front-end, end to end:

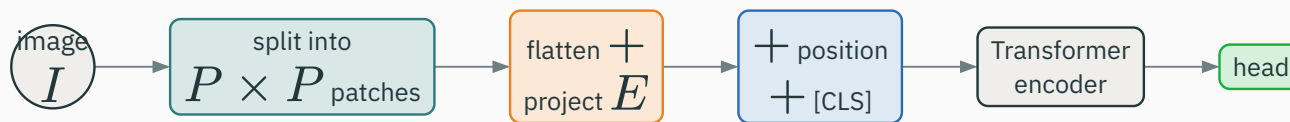


image $\rightarrow$ patches $\rightarrow$ flatten $\rightarrow$ project $\rightarrow$ add position & [CLS] $\rightarrow$ Transformer $\rightarrow$ head

Patchification

How many patches?

A patch size P tiles the $H \times W$ image into a grid of non-overlapping squares:

How many patches?

A patch size P tiles the $H \times W$ image into a grid of non-overlapping squares:

$$N = \frac{H}{P} \cdot \frac{W}{P} \quad (\text{patches per column} \times \text{patches per row})$$

How many patches?

A patch size P tiles the $H \times W$ image into a grid of non-overlapping squares:

$$N = \frac{H}{P} \cdot \frac{W}{P} \quad (\text{patches per column} \times \text{patches per row})$$

N is the **sequence length** the Transformer sees — usually a few hundred, not thousands

Worked: how many patches?

The standard ViT setup — a 224×224 image with patch size $P = 16$:

Worked: how many patches?

The standard ViT setup — a 224×224 image with patch size $P = 16$:

$$\frac{224}{16} = 14 \quad (\text{patches along each side})$$

Worked: how many patches?

The standard ViT setup — a 224×224 image with patch size $P = 16$:

$$\frac{224}{16} = 14 \quad (\text{patches along each side})$$

$$N = 14 \times 14 = \mathbf{196} \quad \text{patch tokens}$$

Worked: how many patches?

The standard ViT setup — a 224×224 image with patch size $P = 16$:

$$\frac{224}{16} = 14 \quad (\text{patches along each side})$$

$$N = 14 \times 14 = \mathbf{196} \quad \text{patch tokens}$$

a 224×224 image becomes a sequence of just **196 tokens** — very manageable

What is inside one patch?

Each $P \times P$ patch has $P \times P$ pixels, each with C channels — flatten it into one vector:

What is inside one patch?

Each $P \times P$ patch has $P \times P$ pixels, each with C channels — flatten it into one vector:

$$x_i \in \mathbb{R}^{P^2 C} \quad (\text{for } P = 16, C = 3: 16^2 \cdot 3 = 768)$$

What is inside one patch?

Each $P \times P$ patch has $P \times P$ pixels, each with C channels — flatten it into one vector:

$$x_i \in \mathbb{R}^{P^2 C} \quad (\text{for } P = 16, C = 3: 16^2 \cdot 3 = 768)$$

Project it to the model width d with a **shared** embedding matrix E :

$$z_i = x_i E + b, \quad E \in \mathbb{R}^{P^2 C \times d}, \quad z_i \in \mathbb{R}^d$$

What is inside one patch?

Each $P \times P$ patch has $P \times P$ pixels, each with C channels — flatten it into one vector:

$$x_i \in \mathbb{R}^{P^2 C} \quad (\text{for } P = 16, C = 3: 16^2 \cdot 3 = 768)$$

Project it to the model width d with a **shared** embedding matrix E :

$$z_i = x_i E + b, \quad E \in \mathbb{R}^{P^2 C \times d}, \quad z_i \in \mathbb{R}^d$$

the **same** E projects **every** patch — one learned linear map, reused N times

With $P = 16$, $C = 3$ so $P^2C = 768$, and model width $d = 768$:

With $P = 16$, $C = 3$ so $P^2C = 768$, and model width $d = 768$:

$$E \in \mathbb{R}^{768 \times 768} \quad (\text{weights}), \quad b \in \mathbb{R}^{768} \quad (\text{bias})$$

With $P = 16$, $C = 3$ so $P^2C = 768$, and model width $d = 768$:

$$E \in \mathbb{R}^{768 \times 768} \quad (\text{weights}), \quad b \in \mathbb{R}^{768} \quad (\text{bias})$$

$$768^2 + 768 = 589,824 + 768 = \mathbf{590,592} \quad \text{parameters}$$

Worked: patch-embedding parameters

With $P = 16$, $C = 3$ so $P^2C = 768$, and model width $d = 768$:

$$E \in \mathbb{R}^{768 \times 768} \quad (\text{weights}), \quad b \in \mathbb{R}^{768} \quad (\text{bias})$$

$$768^2 + 768 = 589,824 + 768 = \mathbf{590,592} \quad \text{parameters}$$

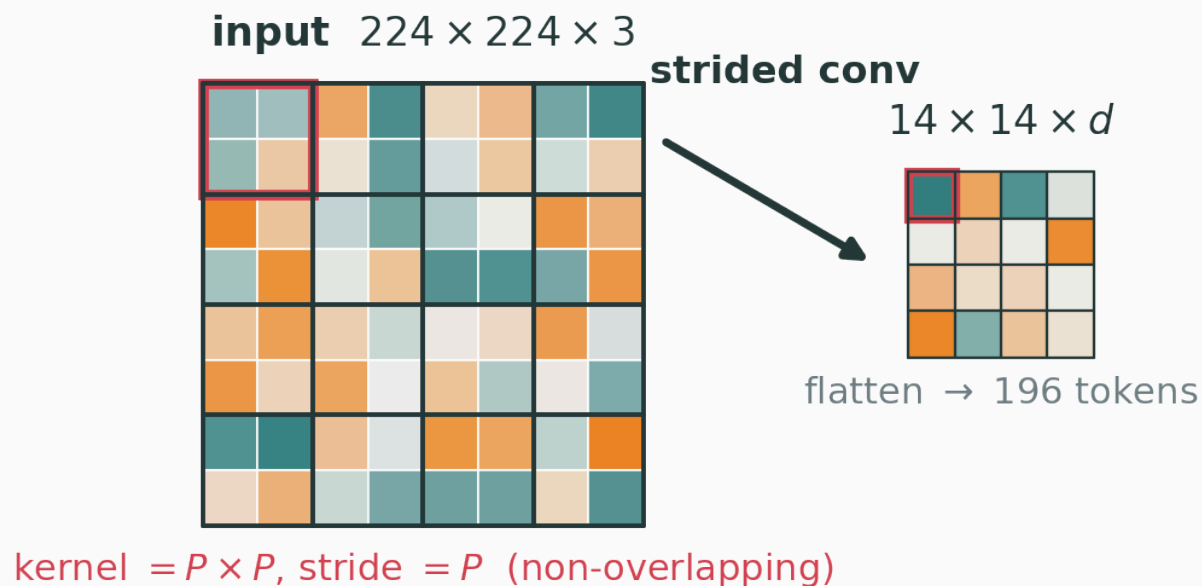
the entire patch embedding is **one** small linear layer — cheap compared to the encoder

Patch embedding is a strided convolution

Flattening every $P \times P$ patch and applying a shared E is **exactly** a convolution:

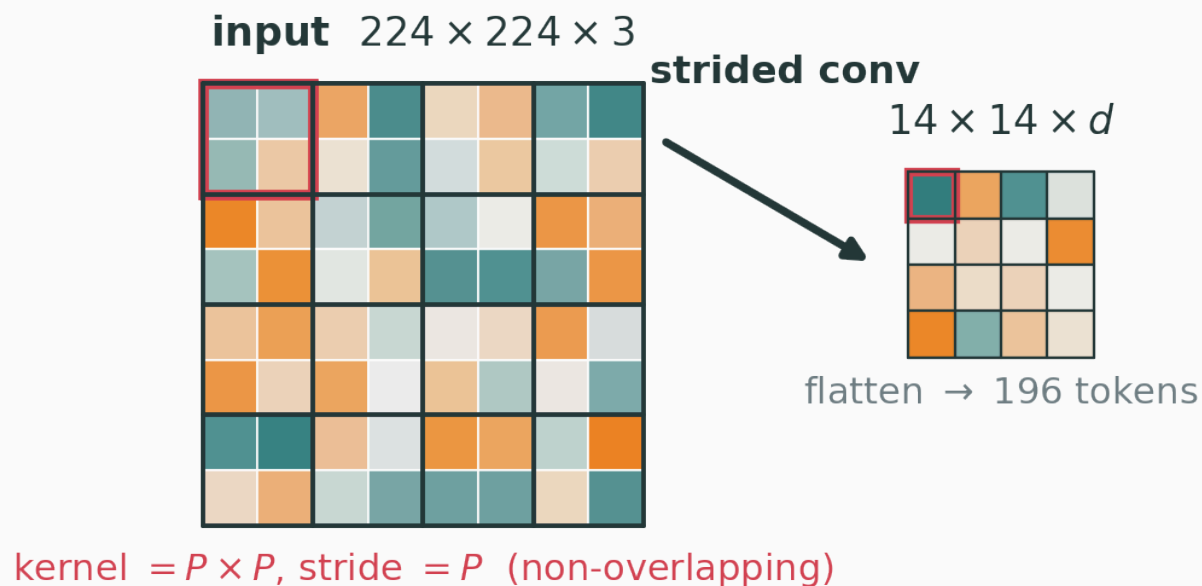
Patch embedding is a strided convolution

Flattening every $P \times P$ patch and applying a shared E is **exactly** a convolution:



Patch embedding is a strided convolution

Flattening every $P \times P$ patch and applying a shared E is **exactly** a convolution:

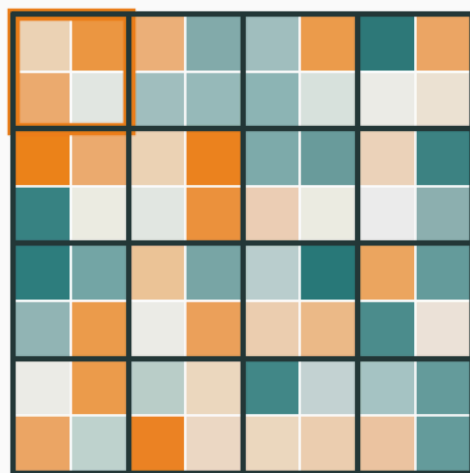


kernel = $P \times P$, stride = P , $C_{\text{out}} = d$ – one conv layer turns the image into patch tokens

One patch: cut it out, flatten its pixels, project to a d -dim token — repeat for all N :

One patch: cut it out, flatten its pixels, project to a d -dim token — repeat for all N :

image = grid of patches



$H \times W \times C$

flatten



$x_i \in \mathbb{R}^{p^2 C}$

$z_i = x_i E$



$z_i \in \mathbb{R}^d$



patch token sequence $z_1 z_2 \dots z_N$

INTERACTIVE

↗ nipunbatra.github.io/interactive-articles/vision-transformer

Drag the patch size P and watch the image re-tile: see $N = (H/P)(W/P)$ change, one patch flatten into $x_i \in \mathbb{R}^{P^2C}$, and project to a token z_i .

The class token and positions

A token to summarize the image

Prepend one **learned** vector — the class token $z_{\text{cls}} \in \mathbb{R}^d$ — to the patch sequence:

A token to summarize the image

Prepend one **learned** vector — the class token $z_{\text{cls}} \in \mathbb{R}^d$ — to the patch sequence:

$$Z_0 = [z_{\text{cls}}; z_1; z_2; \dots; z_N] \quad (\text{length } N + 1)$$

A token to summarize the image

Prepend one **learned** vector — the class token $z_{\text{cls}} \in \mathbb{R}^d$ — to the patch sequence:

$$Z_0 = [z_{\text{cls}}; z_1; z_2; \dots; z_N] \quad (\text{length } N + 1)$$

z_{cls} carries no patch content of its own. Its job is to **gather** information from all patches through the attention layers, and end up as the image's representation.

Use the final class-token state h_{cls} as the whole-image vector:

Use the final class-token state h_{cls} as the whole-image vector:

$$p(y \mid I) = \text{softmax}(W_c h_{\text{cls}} + b_c)$$

Use the final class-token state h_{cls} as the whole-image vector:

$$p(y \mid I) = \text{softmax}(W_c h_{\text{cls}} + b_c)$$

Or **global-average-pool** the patch outputs:

$$h_{\text{image}} = \frac{1}{N} \sum_{i=1}^N h_i$$

Use the final class-token state h_{cls} as the whole-image vector:

$$p(y \mid I) = \text{softmax}(W_c h_{\text{cls}} + b_c)$$

Or **global-average-pool** the patch outputs:

$$h_{\text{image}} = \frac{1}{N} \sum_{i=1}^N h_i$$

[CLS] token **or** mean-pool — both turn N tokens into one image vector

Self-attention has no sense of order

Attention is a weighted sum over tokens — permute the patches and the outputs just permute too:

Self-attention has no sense of order

Attention is a weighted sum over tokens — permute the patches and the outputs just permute too:

$\text{SelfAttention}([z_1; z_2; z_3])$ is a permutation of $\text{SelfAttention}([z_3; z_1; z_2])$

Self-attention has no sense of order

Attention is a weighted sum over tokens — permute the patches and the outputs just permute too:

$\text{SelfAttention}([z_1; z_2; z_3])$ is a permutation of $\text{SelfAttention}([z_3; z_1; z_2])$

Self-attention is **permutation-equivariant**: on its own it cannot tell the **top-left** patch from the **bottom-right** one. Spatial layout is invisible.

Add a learned position vector to every token — patches **and** the class token:

Add a learned position vector to every token — patches **and** the class token:

$$Z_0 = [z_{\text{cls}}; z_1; \dots; z_N] + E_{\text{pos}}, \quad E_{\text{pos}} \in \mathbb{R}^{(N+1) \times d}$$

Add a learned position vector to every token — patches **and** the class token:

$$Z_0 = [z_{\text{cls}}; z_1; \dots; z_N] + E_{\text{pos}}, \quad E_{\text{pos}} \in \mathbb{R}^{(N+1) \times d}$$

ViT uses **learned 1-D** position embeddings, one vector per slot in the flattened patch order. They give attention a location signal; at a new resolution, the learned grid must be adapted (usually by interpolation).

Several ways to encode **where** a patch sits:

Several ways to encode **where** a patch sits:

- **learned** (ViT default) · **separate row / column** · **relative** (Swin) · **rotary**;

Several ways to encode **where** a patch sits:

- **learned** (ViT default) · **separate row / column** · **relative** (Swin) · **rotary**;

all inject a position signal so attention can use spatial layout

Several ways to encode **where** a patch sits:

- **learned** (ViT default) · **separate row / column** · **relative** (Swin) · **rotary**;

all inject a position signal so attention can use spatial layout

Changing resolution. Test at a higher resolution → more patches than trained.
Fix: **interpolate** the position embeddings from the 14×14 grid up to, say, 24×24 .

The ViT encoder

A ViT block is a standard **pre-norm** Transformer block over the patch sequence Z :

A ViT block is a standard **pre-norm** Transformer block over the patch sequence Z :

$$Z'_\ell = Z_{\ell-1} + \text{MHA}(\text{LN}(Z_{\ell-1}))$$

A ViT block is a standard **pre-norm** Transformer block over the patch sequence Z :

$$Z'_\ell = Z_{\ell-1} + \text{MHA}(\text{LN}(Z_{\ell-1}))$$

$$Z_\ell = Z'_\ell + \text{MLP}(\text{LN}(Z'_\ell))$$

A ViT block is a standard **pre-norm** Transformer block over the patch sequence Z :

$$Z'_\ell = Z_{\ell-1} + \text{MHA}(\text{LN}(Z_{\ell-1}))$$

$$Z_\ell = Z'_\ell + \text{MLP}(\text{LN}(Z'_\ell))$$

identical to the language Transformer — only the **tokens** changed, from words to patches

Inside each block, an MLP is applied **independently** to every token:

Inside each block, an MLP is applied **independently** to every token:

$$\text{MLP}(x) = W_2 \varphi(W_1 x + b_1) + b_2, \quad d_{\text{hidden}} > d$$

Inside each block, an MLP is applied **independently** to every token:

$$\text{MLP}(x) = W_2 \varphi(W_1 x + b_1) + b_2, \quad d_{\text{hidden}} > d$$

- typical sizing: $d = 768 \rightarrow d_{\text{hidden}} = 3072 \rightarrow d = 768$ (a $4 \times$ expansion);
- φ is usually GELU.

Inside each block, an MLP is applied **independently** to every token:

$$\text{MLP}(x) = W_2 \varphi(W_1 x + b_1) + b_2, \quad d_{\text{hidden}} > d$$

- typical sizing: $d = 768 \rightarrow d_{\text{hidden}} = 3072 \rightarrow d = 768$ (a $4 \times$ expansion);
- φ is usually GELU.

attention **mixes** patches; the MLP then **processes** each patch on its own

Every patch becomes a query, a key and a value — and attends to **all** patches:

Every patch becomes a query, a key and a value — and attends to **all** patches:

$$q_i = W_Q h_i, \quad k_j = W_K h_j, \quad v_j = W_V h_j$$

Every patch becomes a query, a key and a value — and attends to **all** patches:

$$q_i = W_Q h_i, \quad k_j = W_K h_j, \quad v_j = W_V h_j$$

$$\alpha_{ij} = \text{softmax}_j \left(q_i^\top \frac{k_j}{\sqrt{d_k}} \right), \quad o_i = \sum_j \alpha_{ij} v_j$$

Every patch becomes a query, a key and a value — and attends to **all** patches:

$$q_i = W_Q h_i, \quad k_j = W_K h_j, \quad v_j = W_V h_j$$

$$\alpha_{ij} = \text{softmax}_j \left(q_i^\top \frac{k_j}{\sqrt{d_k}} \right), \quad o_i = \sum_j \alpha_{ij} v_j$$

each patch gathers content from **every other patch** — a global operation

This is the key difference from convolution:

This is the key difference from convolution:

CNN — the receptive field grows **slowly** with depth; distant pixels interact only in **late** layers.

ViT — any patch can attend to any other in a **single** block; some heads integrate large distances **early**.

This is the key difference from convolution:

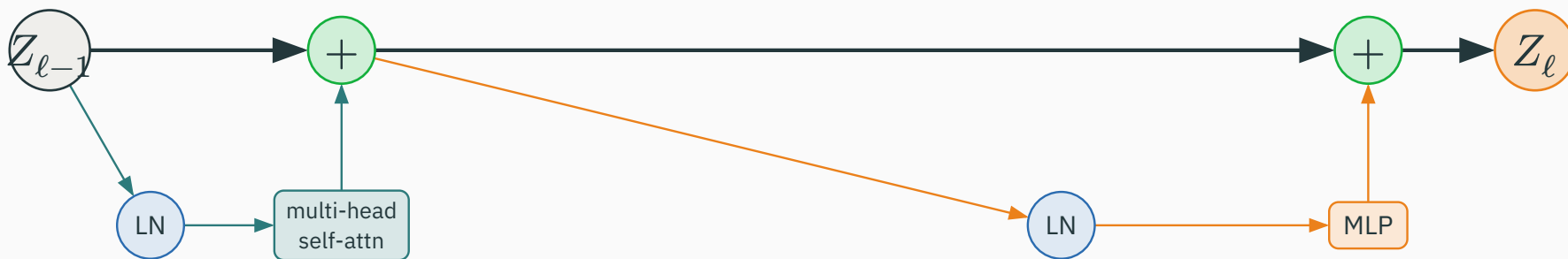
CNN — the receptive field grows **slowly** with depth; distant pixels interact only in **late** layers.

ViT — any patch can attend to any other in a **single** block; some heads integrate large distances **early**.

global from layer one — but **not every** long-range interaction is useful; the model must **learn** which

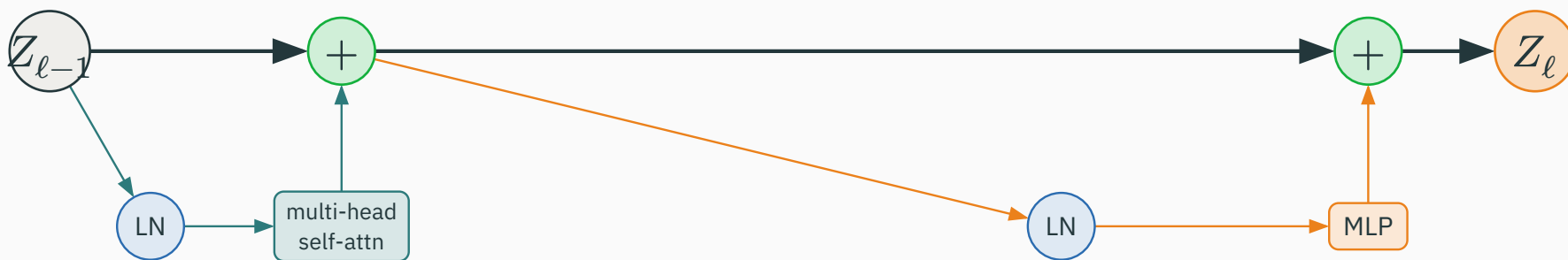
The ViT encoder block

pre-norm: LN sits on the **branch**, the residual spine runs clean



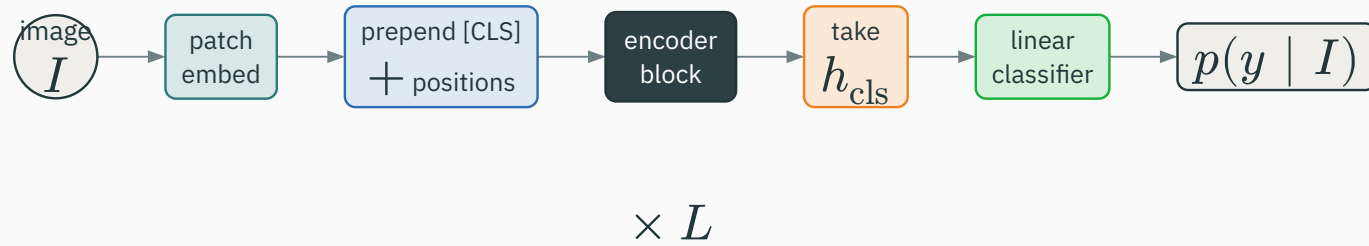
The ViT encoder block

pre-norm: LN sits on the **branch**, the residual spine runs clean

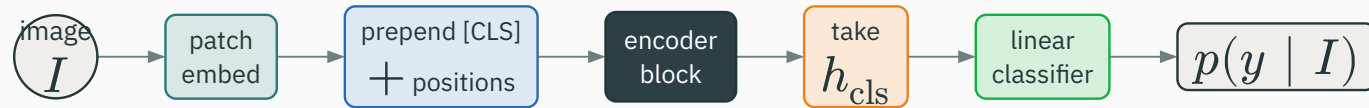


the residual spine runs straight through; each sublayer normalizes, computes, and adds back

The full ViT architecture



The full ViT architecture



$\times L$

patch-embed $\rightarrow$ prepend [CLS] + positions $\rightarrow L$ encoder blocks $\rightarrow$ classify from h_{cls}

Batch B , patches N , width d :

Batch B , patches N , width d :

Stage	Shape	Note
image	$B \times C \times H \times W$	raw pixels
patch embeddings	$B \times N \times d$	N patch tokens
+ class token	$B \times (N + 1) \times d$	prepend [CLS]
attention matrix	$B \times (N + 1) \times (N + 1)$	token-to-token
block output	$B \times (N + 1) \times d$	same shape in and out

A worked tiny ViT

A 4×4 grayscale “image”, patch size 2 — so $N = 4$ patches, each of dimension 4:

A 4×4 grayscale “image”, patch size 2 — so $N = 4$ patches, each of dimension 4:

$$I = \begin{pmatrix} 1 & 1 & 0 & 0 \\ 1 & 1 & 0 & 0 \\ 0 & 0 & 2 & 2 \\ 0 & 0 & 2 & 2 \end{pmatrix}$$

A 4×4 grayscale “image”, patch size 2 — so $N = 4$ patches, each of dimension 4:

$$I = \begin{pmatrix} 1 & 1 & 0 & 0 \\ 1 & 1 & 0 & 0 \\ 0 & 0 & 2 & 2 \\ 0 & 0 & 2 & 2 \end{pmatrix}$$

Flatten each 2×2 patch (row-major):

$$p_1 = (1 \ 1 \ 1 \ 1), p_2 = (0 \ 0 \ 0 \ 0), p_3 = (0 \ 0 \ 0 \ 0), p_4 = (2 \ 2 \ 2 \ 2)$$

Project each 4-dim patch to a 2-dim token with a shared E :

Project each 4-dim patch to a 2-dim token with a shared E :

$$E = \begin{pmatrix} 1 & 0 \\ 0 & 1 \\ 1 & 0 \\ 0 & 1 \end{pmatrix}, \quad z_i = p_i E$$

Project each 4-dim patch to a 2-dim token with a shared E :

$$E = \begin{pmatrix} 1 & 0 \\ 0 & 1 \\ 1 & 0 \\ 0 & 1 \end{pmatrix}, \quad z_i = p_i E$$

$$z_1 = (2 \ 2), \quad z_2 = (0 \ 0), \quad z_3 = (0 \ 0), \quad z_4 = (4 \ 4)$$

Project each 4-dim patch to a 2-dim token with a shared E :

$$E = \begin{pmatrix} 1 & 0 \\ 0 & 1 \\ 1 & 0 \\ 0 & 1 \end{pmatrix}, \quad z_i = p_i E$$

$$z_1 = (2 \ 2), \quad z_2 = (0 \ 0), \quad z_3 = (0 \ 0), \quad z_4 = (4 \ 4)$$

the bright bottom-right patch p_4 becomes the token with the **largest** norm

Prepend $z_{\text{cls}} = \begin{pmatrix} 1 & 1 \end{pmatrix}$ (drop positions for clarity). Take $W_Q = W_K = W_V = I$, so $q_{\text{cls}} = \begin{pmatrix} 1 & 1 \end{pmatrix}$:

Prepend $z_{\text{cls}} = \begin{pmatrix} 1 & 1 \end{pmatrix}$ (drop positions for clarity). Take $W_Q = W_K = W_V = I$, so $q_{\text{cls}} = \begin{pmatrix} 1 & 1 \end{pmatrix}$:

$$s_j = q_{\text{cls}}^\top \frac{k_j}{\sqrt{2}} \quad (\text{scores from the class token to each patch})$$

Prepend $z_{\text{cls}} = \begin{pmatrix} 1 & 1 \end{pmatrix}$ (drop positions for clarity). Take $W_Q = W_K = W_V = I$, so $q_{\text{cls}} = \begin{pmatrix} 1 & 1 \end{pmatrix}$:

$$s_j = q_{\text{cls}}^\top \frac{k_j}{\sqrt{2}} \quad (\text{scores from the class token to each patch})$$

$$s = \left(\underbrace{(1 \ 1) \cdot (2 \ 2)}_4, 0, 0, \underbrace{(1 \ 1) \cdot (4 \ 4)}_8 \right) / \sqrt{2}$$

Prepend $z_{\text{cls}} = \begin{pmatrix} 1 & 1 \end{pmatrix}$ (drop positions for clarity). Take $W_Q = W_K = W_V = I$, so $q_{\text{cls}} = \begin{pmatrix} 1 & 1 \end{pmatrix}$:

$$s_j = q_{\text{cls}}^\top \frac{k_j}{\sqrt{2}} \quad (\text{scores from the class token to each patch})$$

$$s = \left(\underbrace{(1 \ 1) \cdot (2 \ 2)}_4, 0, 0, \underbrace{(1 \ 1) \cdot (4 \ 4)}_8 \right) / \sqrt{2}$$

patch 4 scores highest, patch 1 next — the class token will **lean toward the bright region**

Softmax the scores into weights, then blend the values:

Softmax the scores into weights, then blend the values:

$$o_{\text{cls}} = \sum_j \alpha_{\text{cls},j} v_j \quad (\text{a content-dependent summary of the patches})$$

Softmax the scores into weights, then blend the values:

$$o_{\text{cls}} = \sum_j \alpha_{\text{cls},j} v_j \quad (\text{a content-dependent summary of the patches})$$

Classify from it, then backprop the cross-entropy:

$$z = W_c o_{\text{cls}} + b_c, \quad p = \text{softmax}(z), \quad \mathcal{L} = -\log p_y$$

Softmax the scores into weights, then blend the values:

$$o_{\text{cls}} = \sum_j \alpha_{\text{cls},j} v_j \quad (\text{a content-dependent summary of the patches})$$

Classify from it, then backprop the cross-entropy:

$$z = W_c o_{\text{cls}} + b_c, \quad p = \text{softmax}(z), \quad \mathcal{L} = -\log p_y$$

The gradient flows back into **everything**: the classifier, attention W_Q, W_K, W_V , the MLP, the class token, positions, and the patch projection E .

INTERACTIVE

↗ nipunbatra.github.io/interactive-articles/vision-transformer

Paint a tiny image, watch it patchify and project, and see the class token's attention scores s_j shift toward the patches you brighten.

ViT vs CNN

The two architectures build in **different amounts** of prior knowledge about images:

The two architectures build in **different amounts** of prior knowledge about images:

CNN — bakes in **locality**, **translation equivariance** and **hierarchy**. Strong image priors, works from little data.

ViT — bakes in **far less**. A patch can attend anywhere; the model must **learn** spatial structure from data.

The two architectures build in **different amounts** of prior knowledge about images:

CNN — bakes in **locality**, **translation equivariance** and **hierarchy**. Strong image priors, works from little data.

ViT — bakes in **far less**. A patch can attend anywhere; the model must **learn** spatial structure from data.

fewer built-in priors → more flexible, but **hungrier** for data

The ViT trade-off in one line:

The ViT trade-off in one line:

- ViT **underperforms** CNNs on small datasets — it has no locality prior to fall back on;

The ViT trade-off in one line:

- ViT **underperforms** CNNs on small datasets — it has no locality prior to fall back on;
- but **given enough data** (large-scale pretraining + transfer), it **matches or beats** CNNs.

The ViT trade-off in one line:

- ViT **underperforms** CNNs on small datasets — it has no locality prior to fall back on;
- but **given enough data** (large-scale pretraining + transfer), it **matches or beats** CNNs.

DeiT (Touvron et al.) showed strong augmentation, regularization and **distillation** can train a competitive convolution-free Transformer on ImageNet **alone** — no giant private dataset needed.

	CNN	ViT
core unit	convolution	self-attention
spatial prior	strong (built in)	weak (learned)
receptive field	grows with depth	global per block
cost vs resolution	linear in pixels	quadratic in patches
data efficiency	strong	recipe-sensitive

	CNN	ViT
core unit	convolution	self-attention
spatial prior	strong (built in)	weak (learned)
receptive field	grows with depth	global per block
cost vs resolution	linear in pixels	quadratic in patches
data efficiency	strong	recipe-sensitive

same task, opposite philosophies: hand-built structure vs structure learned from scale

The patch-size trade-off

Patch size P sets both the detail and the cost:

Patch size P sets both the detail and the cost:

$$N = \frac{HW}{P^2}, \quad \text{attention cost} \sim O(N^2 d)$$

Patch size P sets both the detail and the cost:

$$N = \frac{HW}{P^2}, \quad \text{attention cost} \sim O(N^2 d)$$

- **smaller** $P \rightarrow$ **more** patches $\rightarrow$ finer spatial detail, but a **longer** sequence;
- longer sequence $\rightarrow$ the $N \times N$ attention matrix grows **quadratically**.

Patch size P sets both the detail and the cost:

$$N = \frac{HW}{P^2}, \quad \text{attention cost} \sim O(N^2 d)$$

- **smaller** $P \rightarrow$ **more** patches $\rightarrow$ finer spatial detail, but a **longer** sequence;
- longer sequence $\rightarrow$ the $N \times N$ attention matrix grows **quadratically**.

resolution of detail vs compute — the central knob of any ViT

Worked: the cost of halving P

D DERIVATION

V VISUAL

224×224 image, comparing $P = 16$ and $P = 8$:

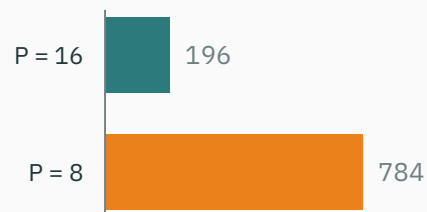
Worked: the cost of halving P

D DERIVATION

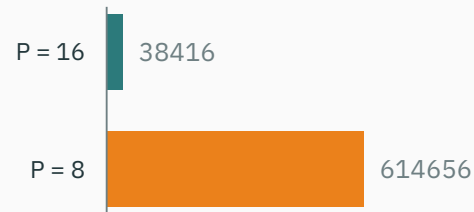
V VISUAL

224×224 image, comparing $P = 16$ and $P = 8$:

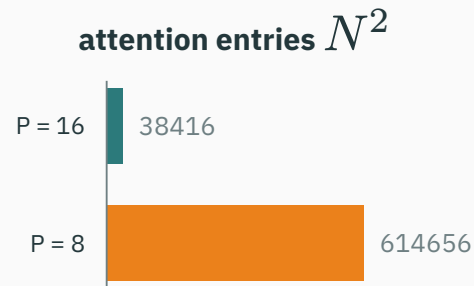
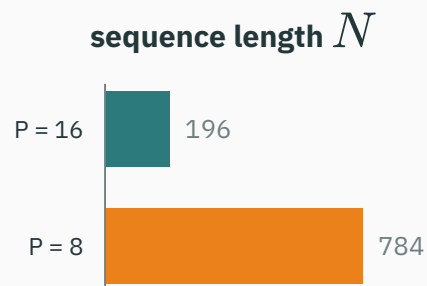
sequence length N



attention entries N^2



224×224 image, comparing $P = 16$ and $P = 8$:



$P : 16 \rightarrow 8$ quadruples N ($196 \rightarrow 784$) and makes the attention matrix $\approx 16 \times$ bigger

Hierarchical ViT (Swin)

Classification needs **one** global vector. But **detection** and **segmentation** need more:

Classification needs **one** global vector. But **detection** and **segmentation** need more:

- **multiple resolutions** (small and large objects);
- **high-resolution** feature maps (per-pixel prediction);
- **manageable** compute at high resolution.

Classification needs **one** global vector. But **detection** and **segmentation** need more:

- **multiple resolutions** (small and large objects);
- **high-resolution** feature maps (per-pixel prediction);
- **manageable** compute at high resolution.

Vanilla ViT keeps a **single, fixed** resolution and pays $O(N^2)$ globally — poorly suited to dense, high-resolution tasks.

Swin (Liu et al.) rebuilds the CNN-style pyramid inside a Transformer:

Swin (Liu et al.) rebuilds the CNN-style pyramid inside a Transformer:

$$\frac{H}{4} \rightarrow \frac{H}{8} \rightarrow \frac{H}{16} \rightarrow \frac{H}{32}$$

Swin (Liu et al.) rebuilds the CNN-style pyramid inside a Transformer:

$$\frac{H}{4} \rightarrow \frac{H}{8} \rightarrow \frac{H}{16} \rightarrow \frac{H}{32}$$

- spatial resolution **shrinks**, channel count **grows** — exactly like a CNN backbone;
- yields multi-scale maps that plug into detection / segmentation heads (FPN).

Swin (Liu et al.) rebuilds the CNN-style pyramid inside a Transformer:

$$\frac{H}{4} \rightarrow \frac{H}{8} \rightarrow \frac{H}{16} \rightarrow \frac{H}{32}$$

- spatial resolution **shrinks**, channel count **grows** — exactly like a CNN backbone;
- yields multi-scale maps that plug into detection / segmentation heads (FPN).

a Transformer that **looks like** a conv backbone from the outside

Instead of global attention, attend only **within** $M \times M$ windows:

Instead of global attention, attend only **within** $M \times M$ windows:

$$\text{global} : O(N^2) \quad \longrightarrow \quad \text{windowed} : O(N \cdot M^2)$$

Instead of global attention, attend only **within** $M \times M$ windows:

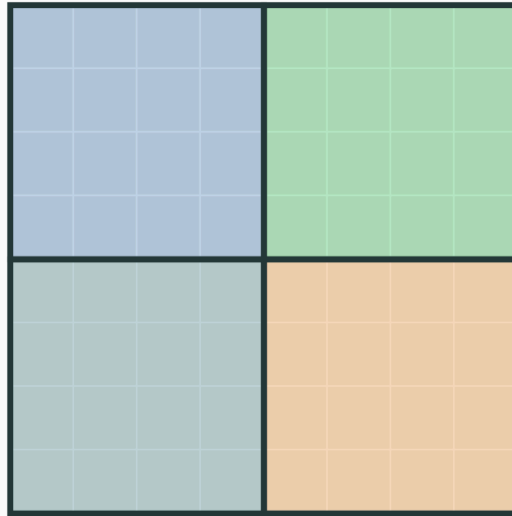
$$\text{global} : O(N^2) \quad \longrightarrow \quad \text{windowed} : O(N \cdot M^2)$$

for a **fixed** window size M , cost is **linear** in the number of patches — scalable to high resolution

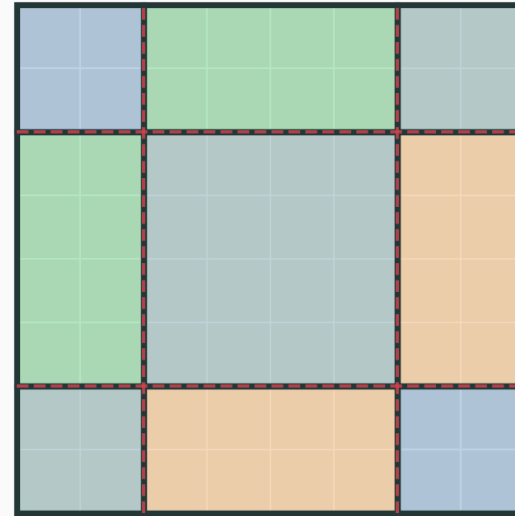
Fixed windows never talk to each other. Swin **shifts** the partition every other block:

Fixed windows never talk to each other. Swin **shifts** the partition every other block:

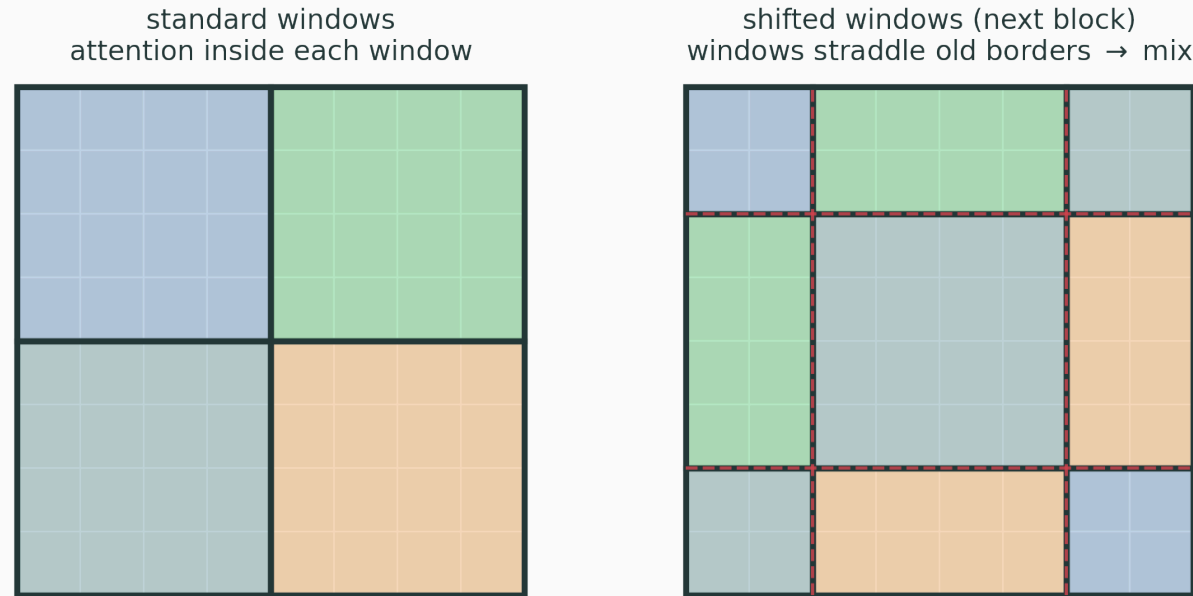
standard windows
attention inside each window



shifted windows (next block)
windows straddle old borders → mix



Fixed windows never talk to each other. Swin **shifts** the partition every other block:



block 1 uses standard windows; block 2 shifts them so neighbouring windows **exchange** information

To go **down** a level of the hierarchy, merge 2×2 neighbouring patches:

To go **down** a level of the hierarchy, merge 2×2 neighbouring patches:

$$H \times W \times C \xrightarrow{\text{concat } 2 \times 2} \frac{H}{2} \times \frac{W}{2} \times 4C \xrightarrow{\text{linear}} \frac{H}{2} \times \frac{W}{2} \times 2C$$

To go **down** a level of the hierarchy, merge 2×2 neighbouring patches:

$$H \times W \times C \xrightarrow{\text{concat } 2 \times 2} \frac{H}{2} \times \frac{W}{2} \times 4C \xrightarrow{\text{linear}} \frac{H}{2} \times \frac{W}{2} \times 2C$$

halve the resolution, grow the channels — the Transformer analog of a strided conv / pool

	ViT	Swin
resolution	single, fixed	hierarchical pyramid
attention scope	global (all patches)	local windows + shift
cost in patches	quadratic $O(N^2)$	linear $O(NM^2)$
best for	classification	detection, segmentation

	ViT	Swin
resolution	single, fixed	hierarchical pyramid
attention scope	global (all patches)	local windows + shift
cost in patches	quadratic $O(N^2)$	linear $O(NM^2)$
best for	classification	detection, segmentation

INTERACTIVE

↗ nipunbatra.github.io/interactive-articles/shifted-window

A shifted-window widget: watch windows tile, then shift, so neighbours mix.

Vision pretraining (briefly)

How do we **train** a vision backbone before transfer?

How do we **train** a vision backbone before transfer?

Supervised — map image $I \rightarrow$ label y , minimize $\mathcal{L}_{\text{CE}}$. Needs labels; the label vocabulary is **limited**.

Self-supervised — **no labels**. Build the target from the image itself: mask, augment, or match.

How do we **train** a vision backbone before transfer?

Supervised — map image $I \rightarrow$ label y , minimize $\mathcal{L}_{\text{CE}}$. Needs labels; the label vocabulary is **limited**.

Self-supervised — **no labels**. Build the target from the image itself: mask, augment, or match.

self-supervision unlocks **web-scale** image data without human annotation

The BERT idea, ported to patches: **hide** some, **predict** them:

The BERT idea, ported to patches: **hide** some, **predict** them:

- randomly **mask** a subset of patch tokens;
- the encoder predicts the **content** (or a learned representation) of the masked patches.

The BERT idea, ported to patches: **hide** some, **predict** them:

- randomly **mask** a subset of patch tokens;
- the encoder predicts the **content** (or a learned representation) of the masked patches.

reconstruct-the-missing forces the model to learn how patches relate — no labels required

Match **two views** of the same image through a **teacher** and a **student**:

Match **two views** of the same image through a **teacher** and a **student**:

- student and teacher see **different augmentations**; the student is trained to **agree** with the teacher;
- the teacher is a slow-moving average of the student — no labels anywhere.

Match **two views** of the same image through a **teacher** and a **student**:

- student and teacher see **different augmentations**; the student is trained to **agree** with the teacher;
- the teacher is a slow-moving average of the student — no labels anywhere.

DINOv2 (Oquab et al.) produces strong **general-purpose** features this way — useful across many tasks **without** task labels.

Two ways to **use** a pretrained backbone with parameters θ :

Two ways to **use** a pretrained backbone with parameters θ :

Method	Update	Measures
linear probe	freeze θ ; train $y = Wh + b$	direct usability of features
finetune	update all of θ	best task accuracy

Two ways to **use** a pretrained backbone with parameters θ :

Method	Update	Measures
linear probe	freeze θ ; train $y = Wh + b$	direct usability of features
finetune	update all of θ	best task accuracy

a strong **linear probe** means the features are already **linearly** useful — the mark of good pretraining

Image-text learning: CLIP

Why connect images and text?

A single label is a thin signal. A **caption** is far richer:

Why connect images and text?

A single label is a thin signal. A **caption** is far richer:

“a brown dog running through snow”

Why connect images and text?

A single label is a thin signal. A **caption** is far richer:

“a brown dog running through snow”

object **dog** · attribute **brown** · action **running** · scene **snow** · relations **through**

Why connect images and text?

A single label is a thin signal. A **caption** is far richer:

“a brown dog running through snow”

object **dog** · attribute **brown** · action **running** · scene **snow** · relations **through**

**far more supervision than the one word “dog” — without manually
annotating every image**

Why connect images and text?

A single label is a thin signal. A **caption** is far richer:

“a brown dog running through snow”

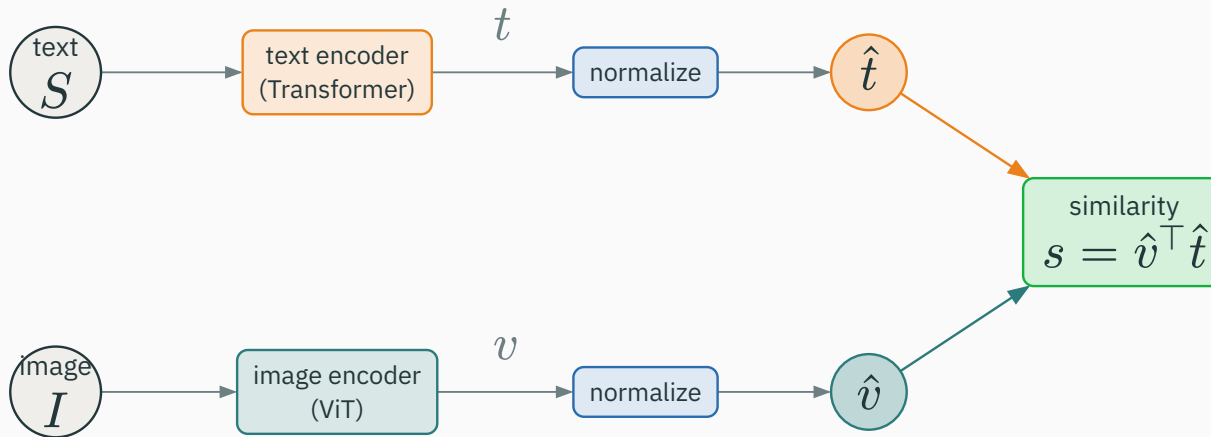
object **dog** · attribute **brown** · action **running** · scene **snow** · relations **through**

**far more supervision than the one word “dog” — without manually
annotating every image**

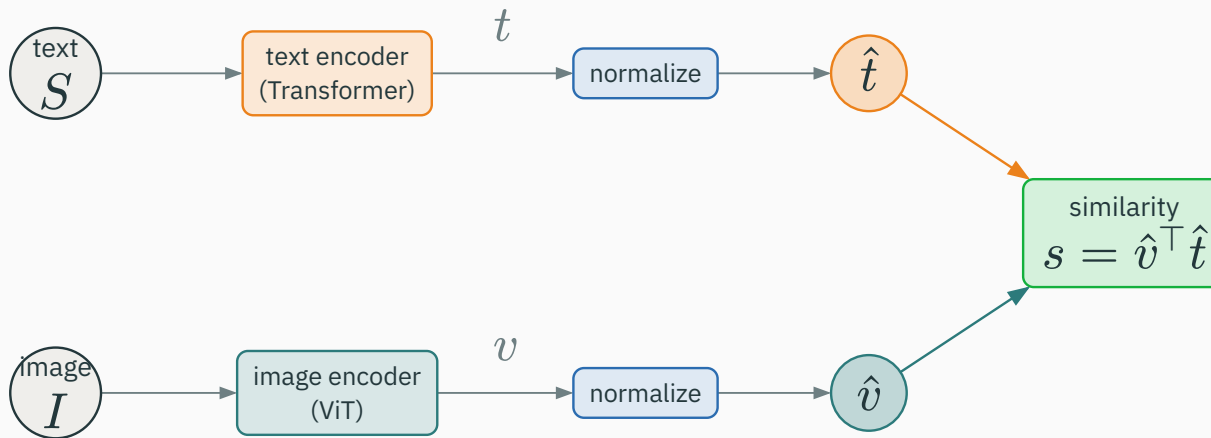
web pairs are still noisy and require collection, filtering, and careful attention to provenance, bias, and consent

CLIP (Radford et al. 2021): one encoder per modality, projected into a **shared** space:

CLIP (Radford et al. 2021): one encoder per modality, projected into a **shared** space:



CLIP (Radford et al. 2021): one encoder per modality, projected into a **shared** space:



$$v = f_{\text{img}}(I), t = f_{\text{text}}(S); \text{normalize } \hat{v} = v / \|v\|, \hat{t} = t / \|t\|; \text{compare by cosine } \hat{v}^\top \hat{t}$$

Take a **batch** of B matched image-text pairs. Score **every** image against **every** caption:

Take a **batch** of B matched image-text pairs. Score **every** image against **every** caption:

$$S_{ij} = \hat{v}_i^\top \hat{t}_j / \tau \quad (\tau = \text{learned temperature})$$

Take a **batch** of B matched image-text pairs. Score **every** image against **every** caption:

$$S_{ij} = \hat{v}_i^\top \hat{t}_j / \tau \quad (\tau = \text{learned temperature})$$

- the B **correct** pairs sit on the **diagonal** S_{ii} ;
- all $B^2 - B$ off-diagonal pairs are **mismatches**.

Take a **batch** of B matched image-text pairs. Score **every** image against **every** caption:

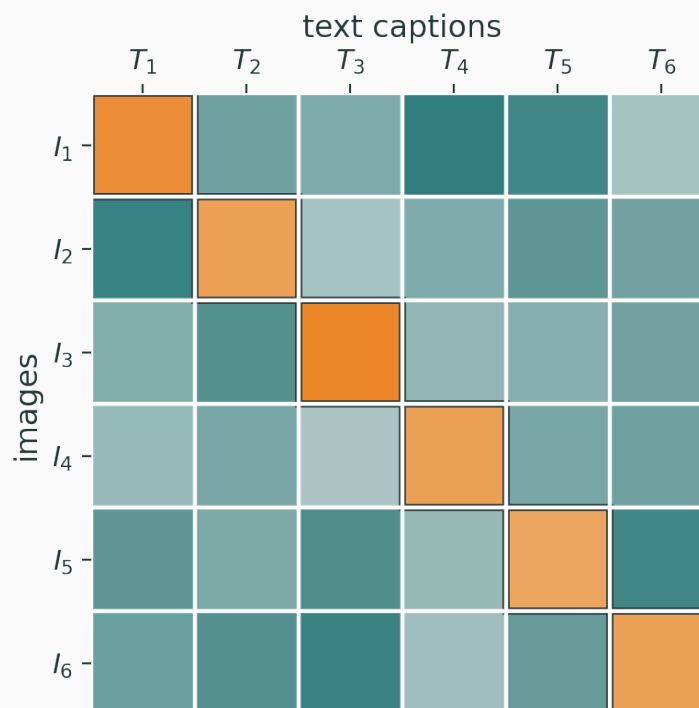
$$S_{ij} = \hat{v}_i^\top \hat{t}_j / \tau \quad (\tau = \text{learned temperature})$$

- the B **correct** pairs sit on the **diagonal** S_{ii} ;
- all $B^2 - B$ off-diagonal pairs are **mismatches**.

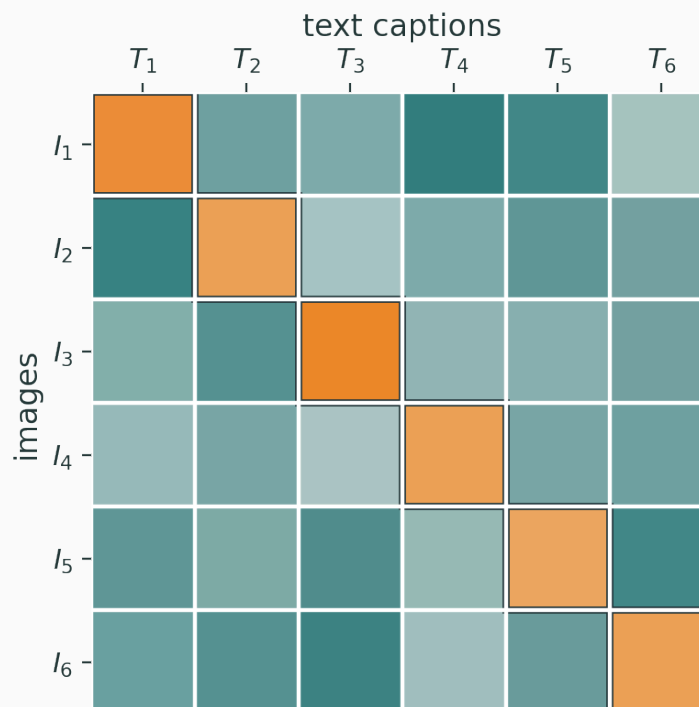
train the model to pick the matching caption for each image — and vice versa

The similarity matrix

$$S_{ij} = \hat{\mathbf{v}}_i^\top \hat{\mathbf{t}}_j / \tau \quad \text{— maximize the diagonal}$$



$$S_{ij} = \hat{\mathbf{v}}_i^\top \hat{\mathbf{t}}_j / \tau \quad \text{— maximize the diagonal}$$



a $B \times B$ grid of similarities; push the **diagonal** up and everything **off-diagonal** down

For image i , softmax across the **row** (all captions) and reward the correct one:

For image i , softmax across the **row** (all captions) and reward the correct one:

$$\mathcal{L}_{I \rightarrow T}^{(i)} = -\log \frac{\exp(S_{ii})}{\sum_{j=1}^B \exp(S_{ij})}$$

For image i , softmax across the **row** (all captions) and reward the correct one:

$$\mathcal{L}_{I \rightarrow T}^{(i)} = -\log \frac{\exp(S_{ii})}{\sum_{j=1}^B \exp(S_{ij})}$$

Symmetrically, softmax down the **column** for the text-to-image direction, then average:

$$\mathcal{L} = \frac{1}{2}(\mathcal{L}_{I \rightarrow T} + \mathcal{L}_{T \rightarrow I})$$

For image i , softmax across the **row** (all captions) and reward the correct one:

$$\mathcal{L}_{I \rightarrow T}^{(i)} = -\log \frac{\exp(S_{ii})}{\sum_{j=1}^B \exp(S_{ij})}$$

Symmetrically, softmax down the **column** for the text-to-image direction, then average:

$$\mathcal{L} = \frac{1}{2}(\mathcal{L}_{I \rightarrow T} + \mathcal{L}_{T \rightarrow I})$$

a cross-entropy where the “classes” are the **other items in the batch**

A batch of $B = 3$ with similarity matrix S (rows = images, columns = captions):

A batch of $B = 3$ with similarity matrix S (rows = images, columns = captions):

$$S = \begin{pmatrix} 3 & 1 & 0 \\ 0 & 2 & 1 \\ 1 & 0 & 4 \end{pmatrix}$$

A batch of $B = 3$ with similarity matrix S (rows = images, columns = captions):

$$S = \begin{pmatrix} 3 & 1 & 0 \\ 0 & 2 & 1 \\ 1 & 0 & 4 \end{pmatrix}$$

Image 1: softmax across row 1 = (3, 1, 0):

$$p(T_1 \mid I_1) = \frac{e^3}{e^3 + e^1 + e^0} = \frac{20.09}{23.81} \approx \mathbf{0.844}$$

A batch of $B = 3$ with similarity matrix S (rows = images, columns = captions):

$$S = \begin{pmatrix} 3 & 1 & 0 \\ 0 & 2 & 1 \\ 1 & 0 & 4 \end{pmatrix}$$

Image 1: softmax across row 1 = (3, 1, 0):

$$p(T_1 \mid I_1) = \frac{e^3}{e^3 + e^1 + e^0} = \frac{20.09}{23.81} \approx \mathbf{0.844}$$

$$\mathcal{L}_{I \rightarrow T}^{(1)} = -\log(0.844) \approx \mathbf{0.170}$$

Negatives come from the batch

CLIP has no separate negative set — the batch **is** the negatives:

CLIP has no separate negative set — the batch **is** the negatives:

- for image i , the positive is caption T_i ; the negatives are all $T_{j \neq i}$ **in the same batch**;
- **bigger batch = more negatives** = a harder, more informative task.

CLIP has no separate negative set — the batch **is** the negatives:

- for image i , the positive is caption T_i ; the negatives are all $T_{j \neq i}$ **in the same batch**;
- **bigger batch = more negatives** = a harder, more informative task.

False negatives. Two different images in a batch may genuinely match the same caption — the loss still pushes them apart. Large, diverse batches make this rare.

INTERACTIVE

↗ nipunbatra.github.io/interactive-articles/clip-zero-shot

Edit a batch of image-text similarities and watch the row/column softmax, the diagonal light up, and the symmetric contrastive loss respond.

Zero-shot classification

CLIP can classify into **new** classes it never had a labeled head for. Turn each class into text:

CLIP can classify into **new** classes it never had a labeled head for. Turn each class into text:

$$\{\text{cat, dog, car}\} \xrightarrow{\text{prompt}} \text{a photo of a } \{\text{class}\}$$

CLIP can classify into **new** classes it never had a labeled head for. Turn each class into text:

$\{\text{cat}, \text{dog}, \text{car}\} \xrightarrow{\text{prompt}} \text{a photo of a } \{\text{class}\}$

$t_{\text{cat}}, t_{\text{dog}}, t_{\text{car}}$ (encode each prompt with the text tower)

CLIP can classify into **new** classes it never had a labeled head for. Turn each class into text:

$\{\text{cat}, \text{dog}, \text{car}\} \xrightarrow{\text{prompt}}$ a photo of a $\{\text{class}\}$

$t_{\text{cat}}, t_{\text{dog}}, t_{\text{car}}$ (encode each prompt with the text tower)

the class names become **text embeddings** — no supervised classifier needed

Compare the image to each class

Encode the image once, then score it against every class text:

Compare the image to each class

Encode the image once, then score it against every class text:

$$v = f_{\text{img}}(I), \quad s_c = \hat{v}^\top \hat{t}_c / \tau \quad (\hat{v}, \hat{t}_c \text{ unit-normalized, as in training})$$

Compare the image to each class

Encode the image once, then score it against every class text:

$$v = f_{\text{img}}(I), \quad s_c = \hat{v}^\top \hat{t}_c / \tau \quad (\hat{v}, \hat{t}_c \text{ unit-normalized, as in training})$$

$$p(c \mid I) = \text{softmax}_c(s_c)$$

Compare the image to each class

Encode the image once, then score it against every class text:

$$v = f_{\text{img}}(I), \quad s_c = \hat{v}^\top \hat{t}_c / \tau \quad (\hat{v}, \hat{t}_c \text{ unit-normalized, as in training})$$

$$p(c \mid I) = \text{softmax}_c(s_c)$$

add a new class? just add its **prompt** — no retraining, no labeled examples

Similarities of the image to three class prompts, $s = (4, 2, -1)$:

Similarities of the image to three class prompts, $s = (4, 2, -1)$:

$$\text{softmax}(4, 2, -1) = \frac{e^4, e^2, e^{-1}}{e^4 + e^2 + e^{-1}} \approx (0.876, 0.118, 0.006)$$

Similarities of the image to three class prompts, $s = (4, 2, -1)$:

$$\text{softmax}(4, 2, -1) = \frac{e^4, e^2, e^{-1}}{e^4 + e^2 + e^{-1}} \approx (0.876, 0.118, 0.006)$$

predict cat — the highest-similarity class, with no trained classifier

Two practical extensions of the same idea:

Two practical extensions of the same idea:

- **prompt ensembling** — average the text embeddings of several templates (“a photo of a {c}”, “a blurry photo of a {c}”, ...) for a more robust class vector;

Two practical extensions of the same idea:

- **prompt ensembling** — average the text embeddings of several templates (“a photo of a {c}”, “a blurry photo of a {c}”, ...) for a more robust class vector;
- **image-text retrieval** — rank a gallery by $\cos(v_i, t_q)$ to find images for a text query (or captions for an image).

Two practical extensions of the same idea:

- **prompt ensembling** — average the text embeddings of several templates (“a photo of a {c}”, “a blurry photo of a {c}”, ...) for a more robust class vector;
- **image-text retrieval** — rank a gallery by $\cos(v_i, t_q)$ to find images for a text query (or captions for an image).

INTERACTIVE

↗ nipunbatra.github.io/interactive-articles/clip-zero-shot

Type your own class names, watch the prompts encode, and see the image’s softmax over classes update live.

Fusion and vision-language models

The dual encoder is fast and flexible — but shallow in how it combines modalities:

The dual encoder is fast and flexible — but shallow in how it combines modalities:

Strengths — precompute embeddings, **fast retrieval**, zero-shot transfer, a simple contrastive objective.

Limit — image and text meet **only** at the final $\hat{v}^\top \hat{t}$. Too shallow for “is the **red cup** left of the **blue plate**?”

The dual encoder is fast and flexible — but shallow in how it combines modalities:

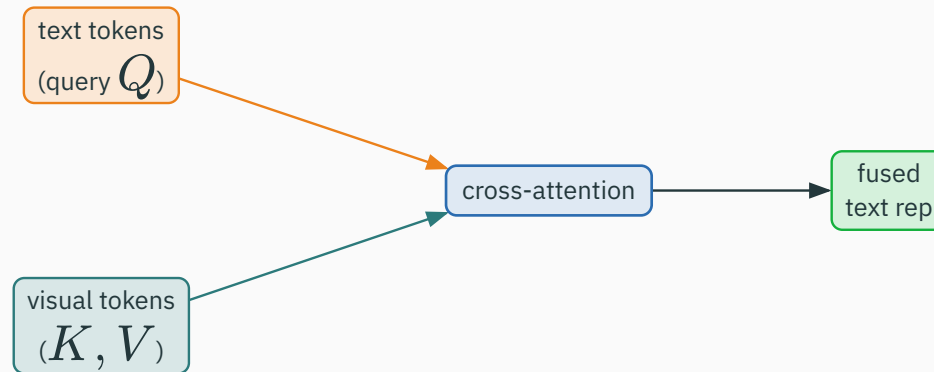
Strengths — precompute embeddings, **fast retrieval**, zero-shot transfer, a simple contrastive objective.

Limit — image and text meet **only** at the final $\hat{v}^\top \hat{t}$. Too shallow for “is the **red cup** left of the **blue plate**?”

compositional, relational questions need **deeper** image-text interaction

Let text **query** the visual tokens: Q from text, K, V from the image.

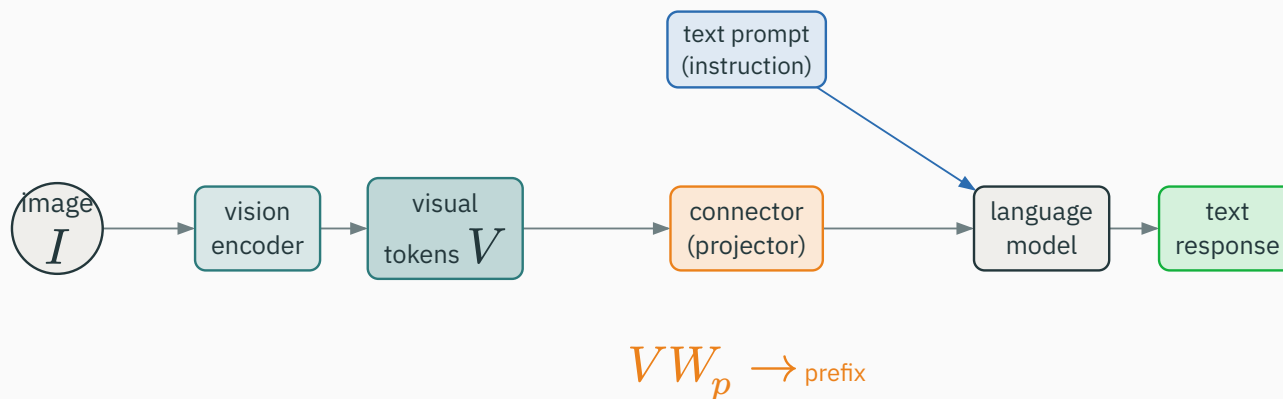
Let text **query** the visual tokens: Q from text, K, V from the image.



$$\text{CrossAttention}(Q_{\text{text}}, K_{\text{vision}}, V_{\text{vision}})$$

Connect a **vision encoder** to a **language model** through a small **connector**:

Connect a **vision encoder** to a **language model** through a small **connector**:



Connect a **vision encoder** to a **language model** through a small **connector**:

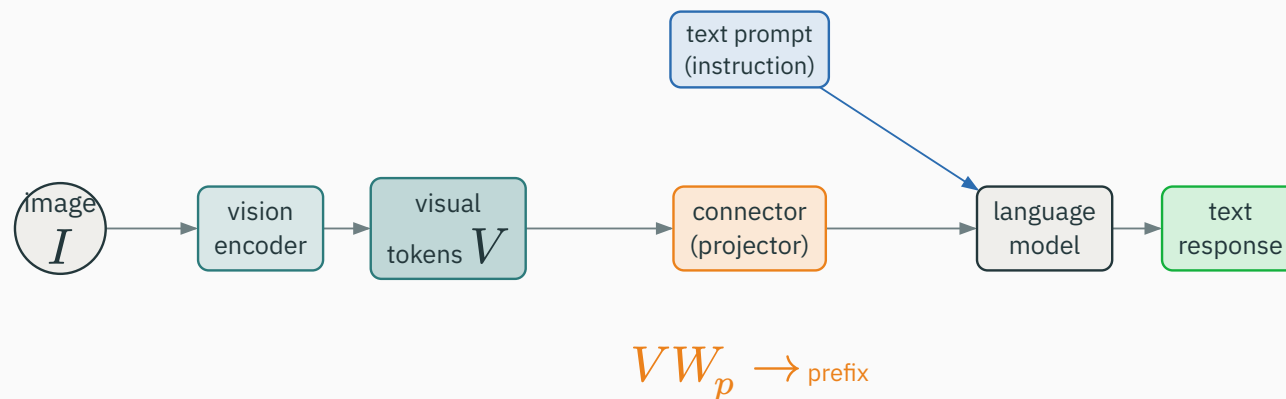


image $\rightarrow$ vision encoder $\rightarrow$ visual tokens $\rightarrow$ connector $\rightarrow$ LM token space $\rightarrow$ text

The connector maps visual features into the LM's token space as a **visual prefix**:

The connector maps visual features into the LM's token space as a **visual prefix**:

$$\tilde{V} = VW_p + b, \quad W_p \in \mathbb{R}^{d_v \times d_{\text{LM}}}$$

The connector maps visual features into the LM's token space as a **visual prefix**:

$$\tilde{V} = VW_p + b, \quad W_p \in \mathbb{R}^{d_v \times d_{\text{LM}}}$$

Then generate text with the **same next-token objective** — now conditioned on the image:

$$\mathcal{L} = - \sum_t \log p(y_t \mid y_{<t}, I)$$

The connector maps visual features into the LM's token space as a **visual prefix**:

$$\tilde{V} = VW_p + b, \quad W_p \in \mathbb{R}^{d_v \times d_{\text{LM}}}$$

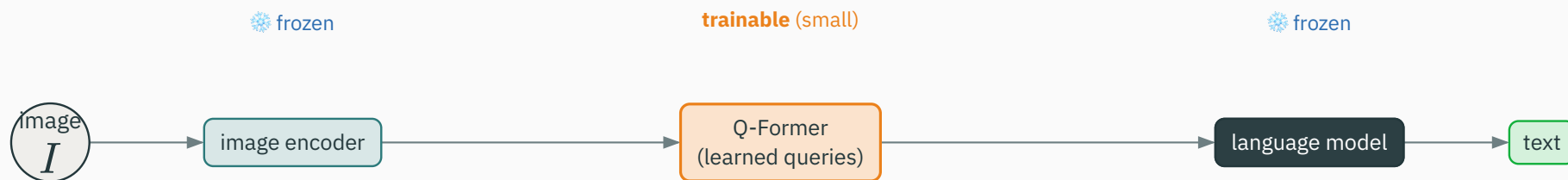
Then generate text with the **same next-token objective** — now conditioned on the image:

$$\mathcal{L} = - \sum_t \log p(y_t \mid y_{<t}, I)$$

nothing new in the loss — the image just becomes extra **context tokens** for the LM

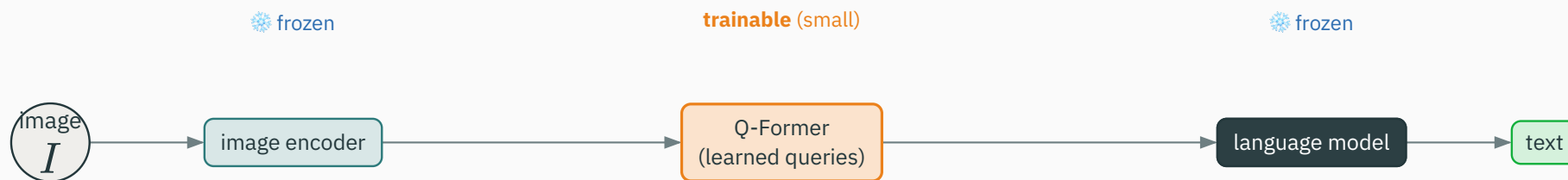
BLIP-2 (Li et al.): keep the image encoder and the LM **frozen**, learn only a light bridge:

BLIP-2 (Li et al.): keep the image encoder and the LM **frozen**, learn only a light bridge:



reuse strong **frozen** models, learn a small **connector**

BLIP-2 (Li et al.): keep the image encoder and the LM **frozen**, learn only a light bridge:



reuse strong **frozen** models, learn a small **connector**

reuse strong unimodal models; train **only** the small Q-Former between them

A handful of **learned queries** $q_1, \dots, q_M$ pull the image down to a few language-ready tokens:

A handful of **learned queries** $q_1, \dots, q_M$ pull the image down to a few language-ready tokens:

$$z_m = \text{CrossAttention}(q_m, V, V) \quad (\text{queries attend to image features } V)$$

A handful of **learned queries** $q_1, \dots, q_M$ pull the image down to a few language-ready tokens:

$$z_m = \text{CrossAttention}(q_m, V, V) \quad (\text{queries attend to image features } V)$$

compress hundreds of patches into M (e.g. 32) **language-relevant** visual tokens

LLaVA (Liu et al.) shows a projector can be almost trivial:

LLaVA (Liu et al.) shows a projector can be almost trivial:

CLIP vision encoder $\rightarrow$ linear / MLP projector $\rightarrow$ LLM

LLaVA (Liu et al.) shows a projector can be almost trivial:

CLIP vision encoder → linear / MLP projector → LLM

- project CLIP patch features straight into the LM's embedding space;
- **visual instruction tuning** — train on (image, instruction, response) triples.

LLaVA (Liu et al.) shows a projector can be almost trivial:

CLIP vision encoder → linear / MLP projector → LLM

- project CLIP patch features straight into the LM's embedding space;
- **visual instruction tuning** — train on (image, instruction, response) triples.

a tiny projector plus instruction data turns an LLM into a capable vision-language model

Train the model to **follow instructions about images**, supervising only the **response**:

Train the model to **follow instructions about images**, supervising only the **response**:

$$\mathcal{L} = - \sum_t \log p(y_t \mid y_{<t}, I, \text{instruction})$$

Train the model to **follow instructions about images**, supervising only the **response**:

$$\mathcal{L} = - \sum_t \log p(y_t \mid y_{<t}, I, \text{instruction})$$

The instruction tokens are **masked out** of the loss — we only teach the model to produce the **answer**, not to re-predict the question.

The main design axis of a VLM:

The main design axis of a VLM:

Strategy	Trade-off
freeze vision + LM, train connector	cheapest; least forgetting; limited alignment
freeze vision, finetune LM	better alignment; risk of forgetting
LoRA on the LM	cheap adaptation, keeps base intact
joint finetune	strongest; most compute, most forgetting

The main design axis of a VLM:

Strategy	Trade-off
freeze vision + LM, train connector	cheapest; least forgetting; limited alignment
freeze vision, finetune LM	better alignment; risk of forgetting
LoRA on the LM	cheap adaptation, keeps base intact
joint finetune	strongest; most compute, most forgetting

INTERACTIVE

↗ nipunbatra.github.io/interactive-articles/vlm

Send an image through a vision encoder and connector into an LM and watch visual tokens become a prefix the model attends to.

Limitations

Fluent captions can still get the **details** wrong:

Fluent captions can still get the **details** wrong:

- **counting** objects; locating **small** objects; **left / right** and precise coordinates;
- reading **fine text**; exact spatial relations.

Fluent captions can still get the **details** wrong:

- **counting** objects; locating **small** objects; **left / right** and precise coordinates;
- reading **fine text**; exact spatial relations.

fluent language about an image $\neq$ **grounded** understanding of it

The classic failure of vision-language models:

The classic failure of vision-language models:

Hallucination — the model produces **plausible, fluent** text that the image does **not** support (an object that isn't there, a wrong colour or count).

The classic failure of vision-language models:

Hallucination — the model produces **plausible, fluent** text that the image does **not** support (an object that isn't there, a wrong colour or count).

Cause — a **strong language prior** (what usually goes together) combined with **imperfect grounding** in the actual pixels.

The classic failure of vision-language models:

Hallucination — the model produces **plausible, fluent** text that the image does **not** support (an object that isn't there, a wrong colour or count).

Cause — a **strong language prior** (what usually goes together) combined with **imperfect grounding** in the actual pixels.

language fluency $\neq$ visual correctness

- **data & bias** — web captions are **noisy**, and carry **stereotypes** and **imbalance** straight into the model;

- **data & bias** — web captions are **noisy**, and carry **stereotypes** and **imbalance** straight into the model;
- **resolution bottleneck** — high-res images make **many** tokens ($N = HW/P^2$), expensive to feed an LLM;

- **data & bias** — web captions are **noisy**, and carry **stereotypes** and **imbalance** straight into the model;
- **resolution bottleneck** — high-res images make **many** tokens ($N = HW/P^2$), expensive to feed an LLM;

mitigations: downsample · pool · learned query tokens · image **tiling**

Different tasks, different metrics:

Different tasks, different metrics:

Task	Metric
classification	accuracy
retrieval	Recall@ k
captioning	CIDEr / BLEU / human
visual question answering	VQA accuracy

Different tasks, different metrics:

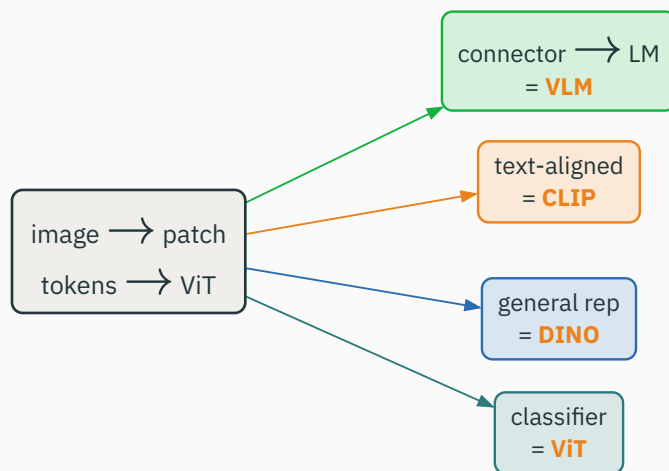
Task	Metric
classification	accuracy
retrieval	Recall@ k
captioning	CIDEr / BLEU / human
visual question answering	VQA accuracy

automatic metrics are imperfect — human evaluation still matters for open-ended generation

Summary

Every model in this lecture starts the **same way** — patchify, then a Transformer:

Every model in this lecture starts the **same way** — patchify, then a Transformer:



Model	Main output	Training signal
ViT	class label	supervised labels
self-sup ViT (DINOv2)	general features	self-supervision, no labels
CLIP	image-text similarity	contrastive on pairs
BLIP-2	text from image	frozen towers + Q-Former
LLaVA	instruction response	visual instruction tuning

Idea	Equation
patch sequence	$Z_0 = [z_{\text{cls}}; z_1; \dots; z_N] + E_{\text{pos}}$
attention	$A = \text{softmax}(QK^\top / \sqrt{d})$
CLIP similarity	$S_{ij} = \hat{v}_i^\top \hat{t}_j / \tau$
multimodal LM	$p(y_{1:T} \mid I) = \prod_t p(y_t \mid y_{<t}, I)$

Final mental model — one idea

ViT: patches become tokens; self-attention lets regions exchange information.

CLIP: images and language are aligned in one shared space.

VLMs: visual tokens are fed to a language model.

one front-end — patchify — powers classification, features, alignment and generation

A Transformer sees an image as a **sequence of patches.**

Next: **Representation Learning & Autoencoders** → VAEs → GANs →
diffusion.