

Deep Learning for Image Analysis

Mehyar MLAWEH

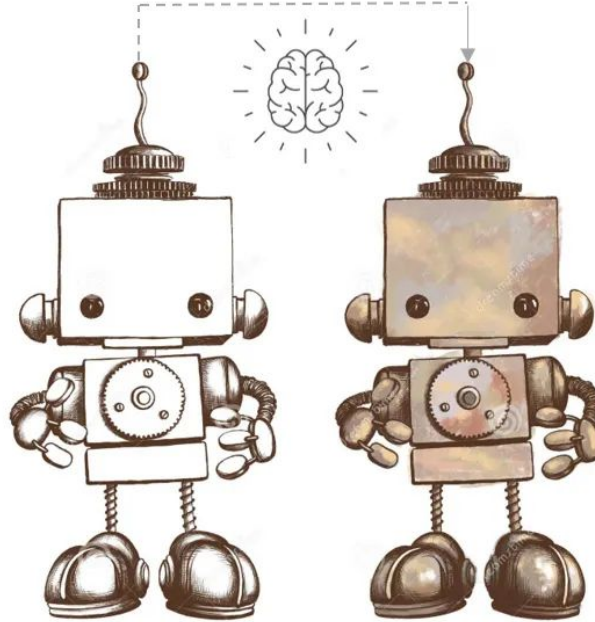


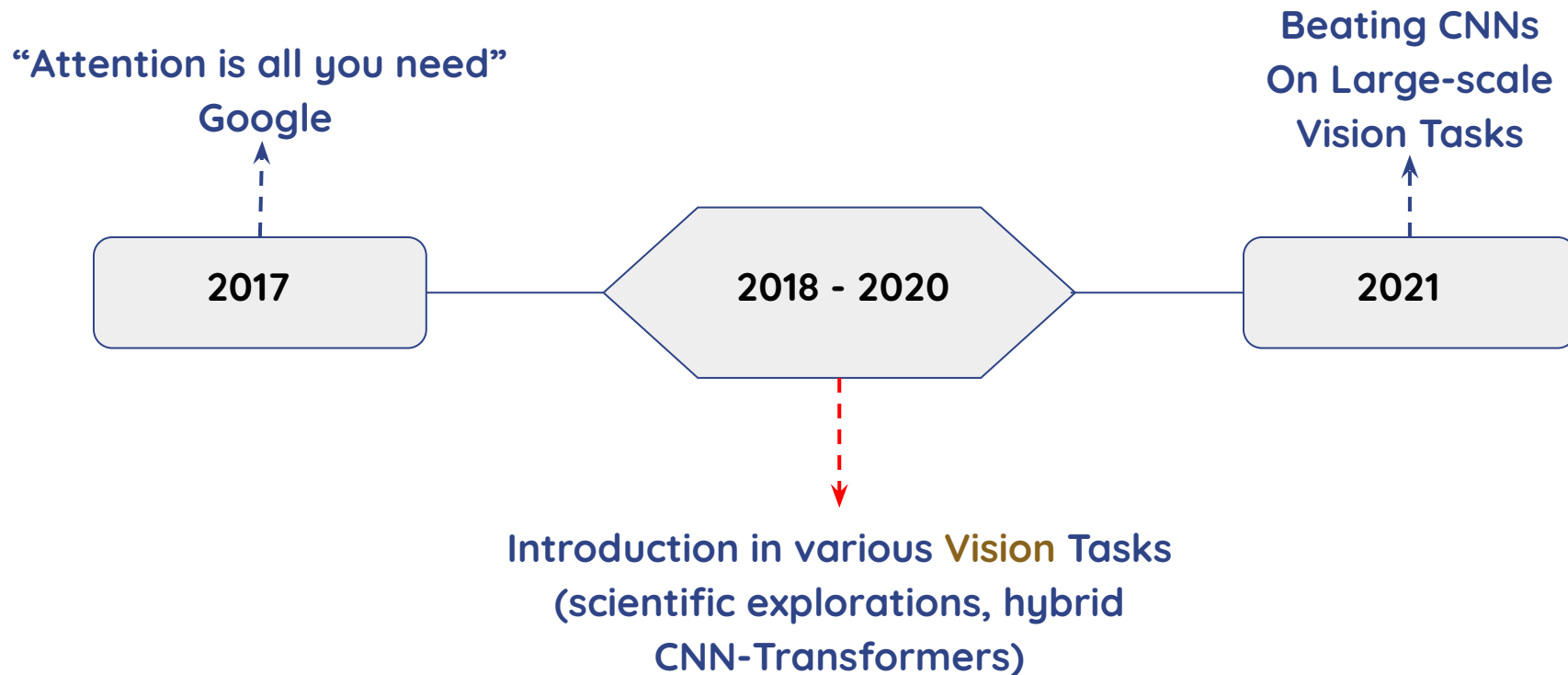
Master BDIA - Dauphine Tunis
PhD student - Université PSL

1. Computer Vision Foundations
2. CNNs
3. Transfer Learning
4. Vision Transformers
5. Object Detection
6. Segment Anything Model (SAM)
7. Final Project

Lecture 4 : Vision Transformers







THE TRANSFORMER

Je suis étudiant



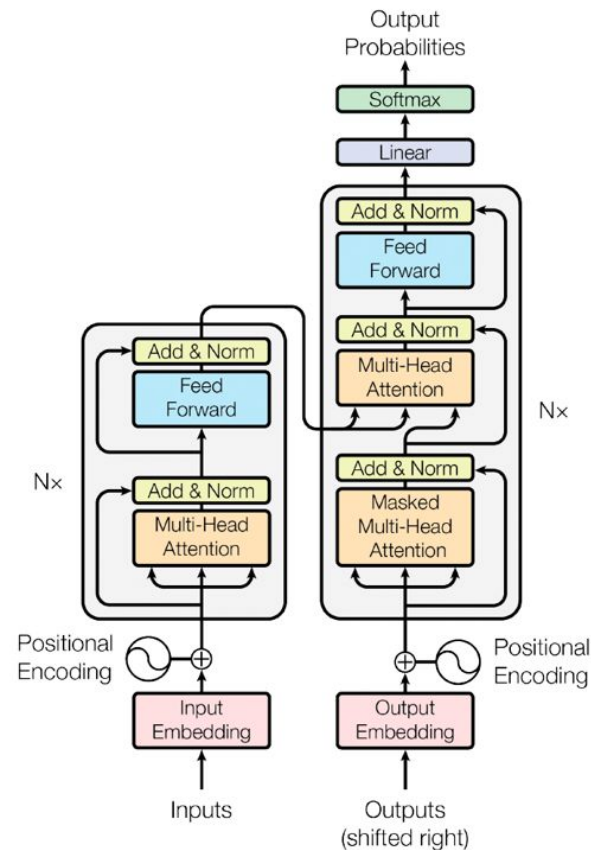
I am a student

THE TRANSFORMER

Je suis étudiant



I am a student



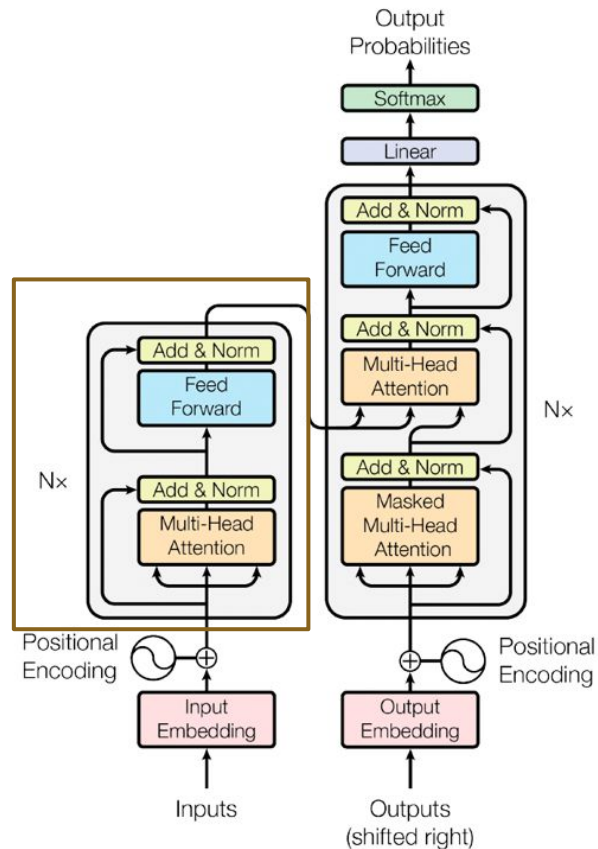
Encoders

THE TRANSFORMER

Je suis étudiant



I am a student



Encoders

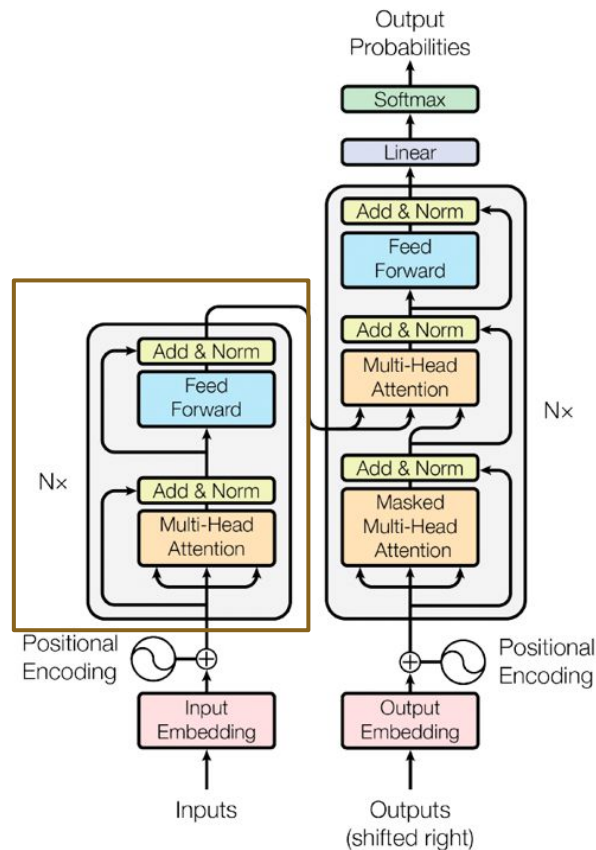
N identical encoders stacked
($N = 6$ in the original paper)

THE TRANSFORMER

Je suis étudiant



I am a student

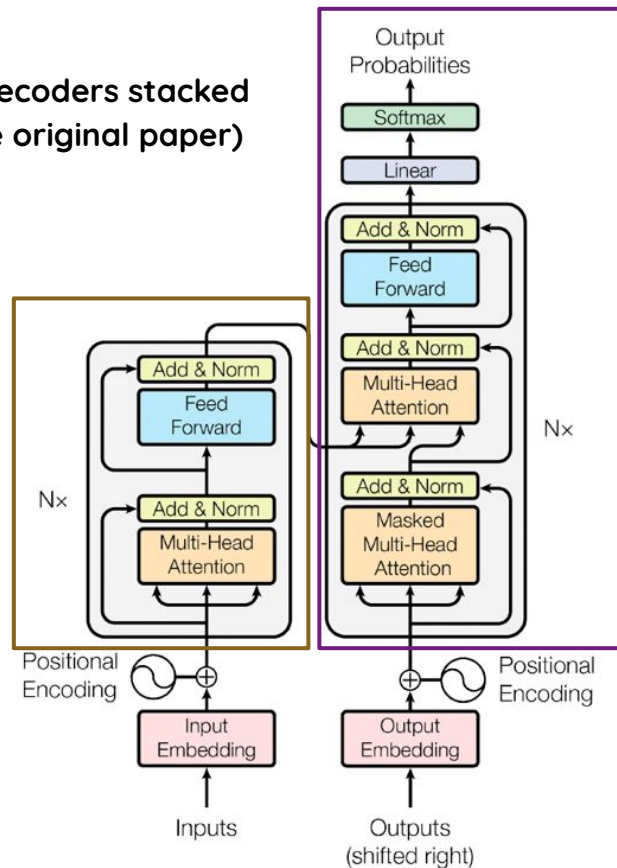


N identical decoders stacked
($N = 6$ in the original paper)

Encoders

N identical encoders stacked
($N = 6$ in the original paper)

Decoders

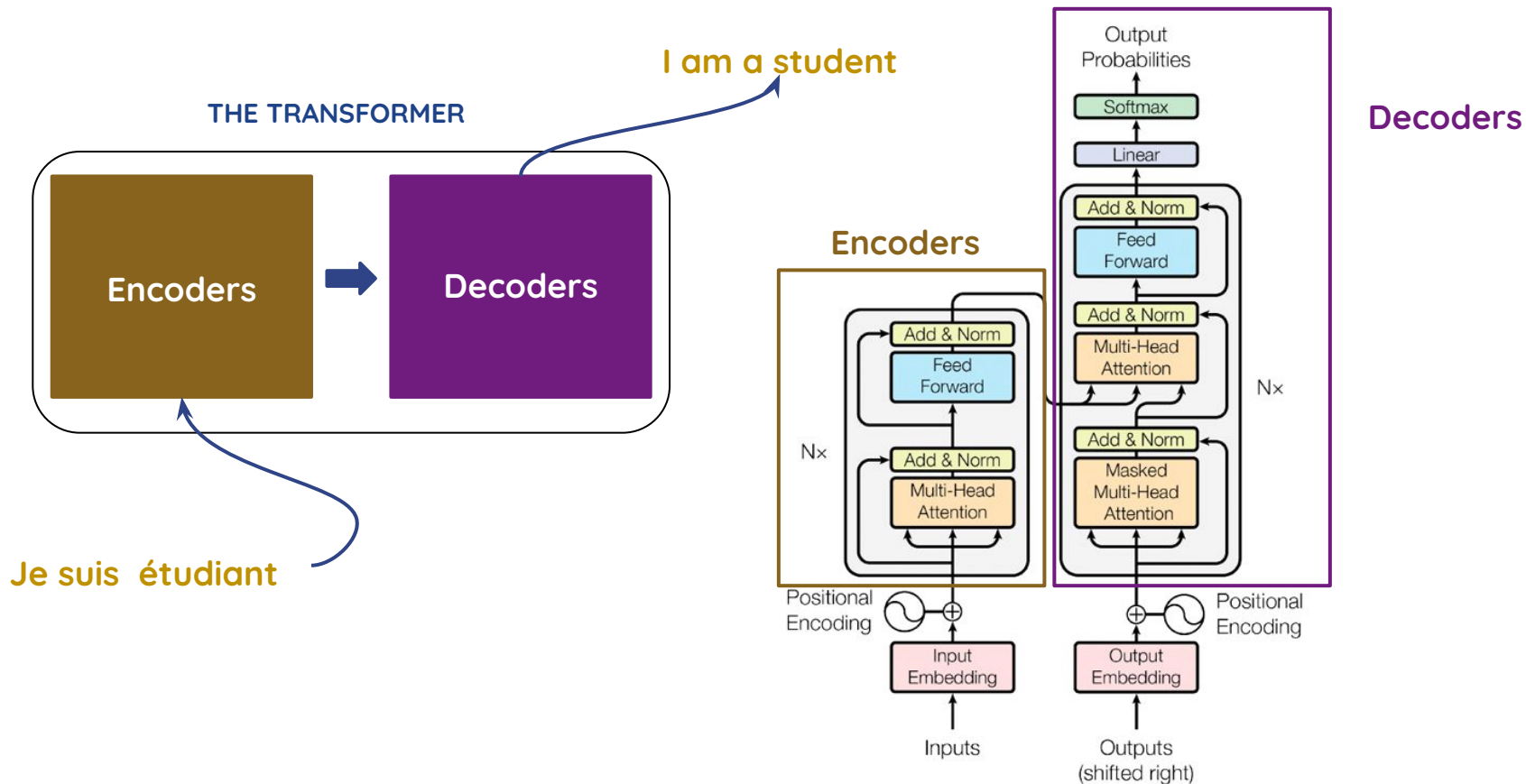


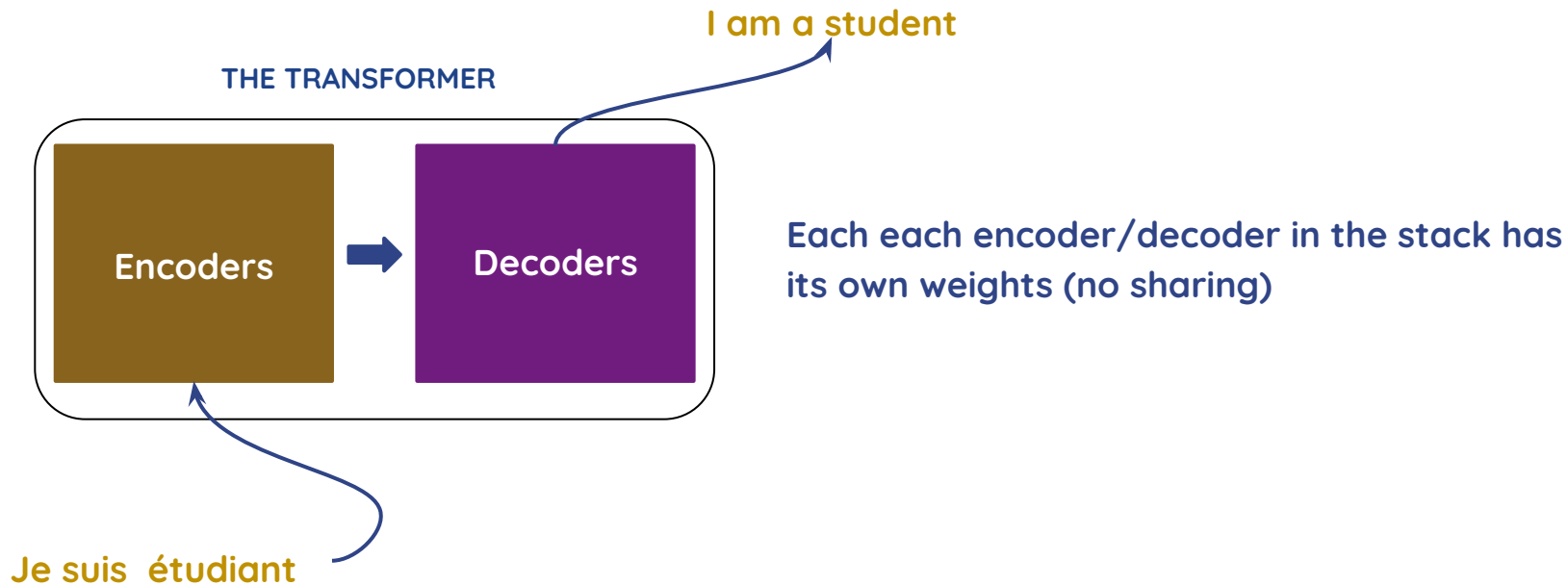
THE TRANSFORMER

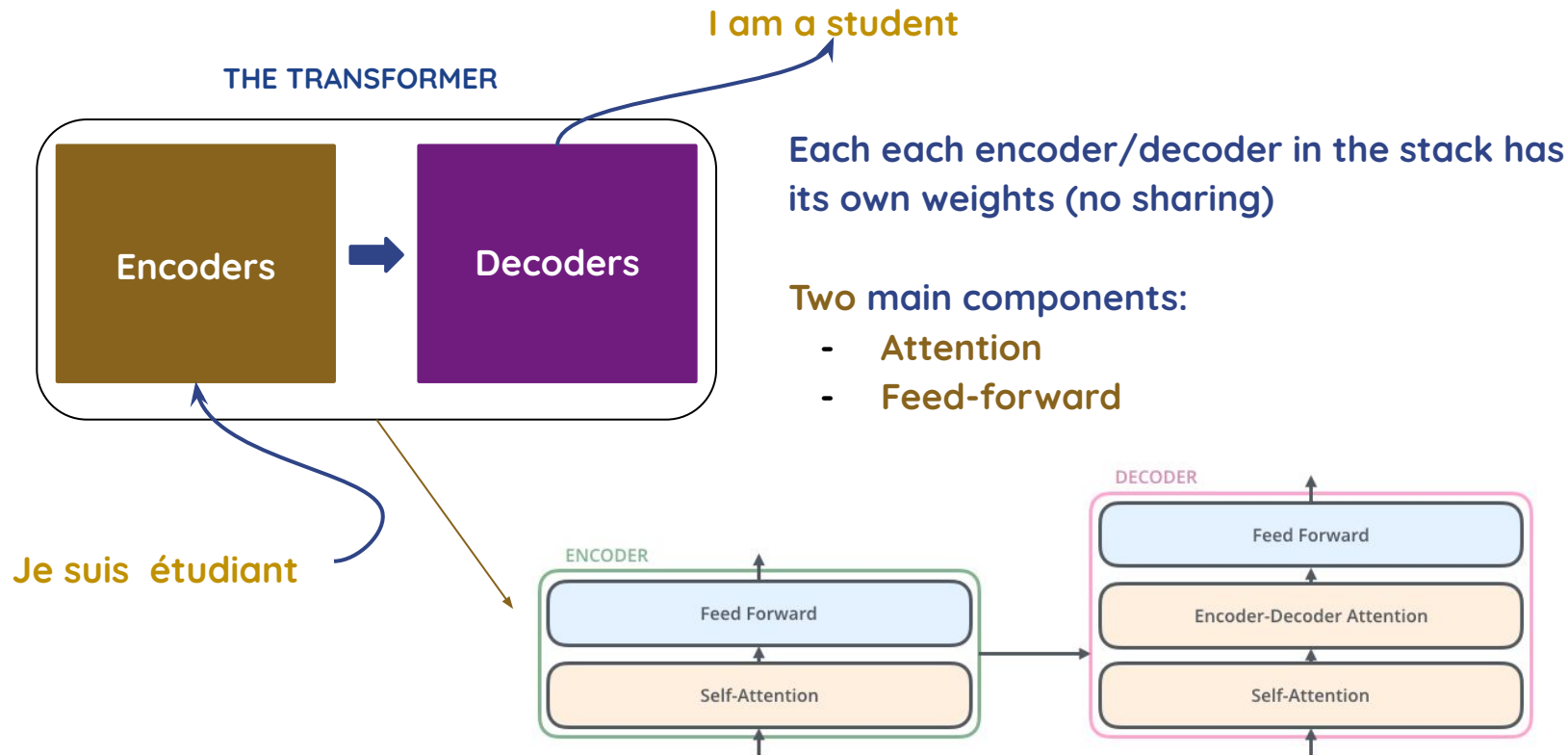
Je suis étudiant



I am a student

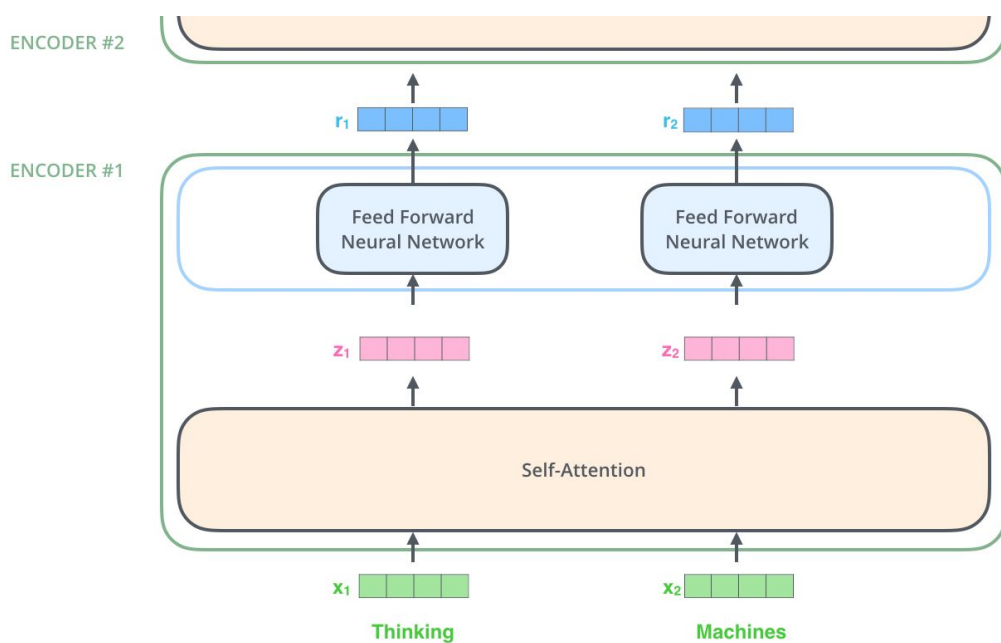
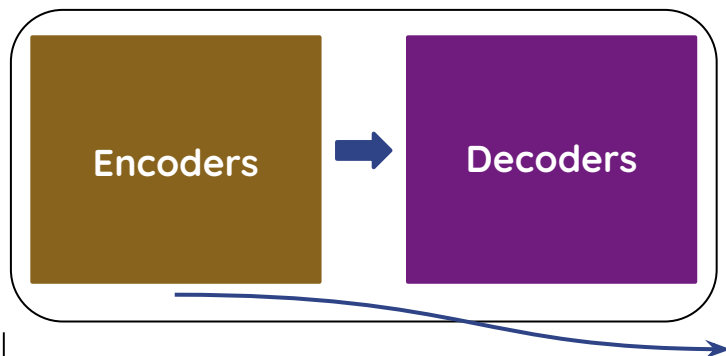




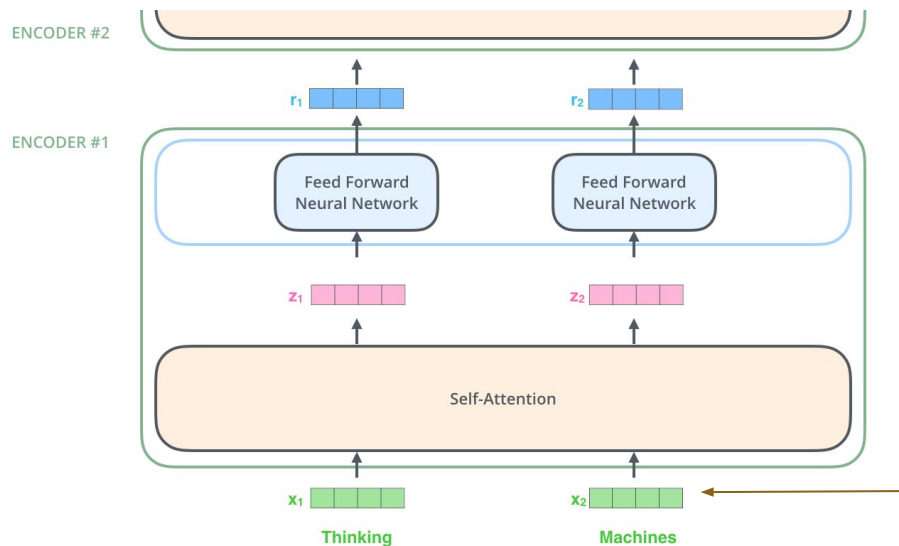


[Source: [The Illustrated Transformer – Jay Alammar – Visualizing machine learning one concept at a time.](#)]

THE TRANSFORMER



[Source: [The Illustrated Transformer – Jay Alammar – Visualizing machine learning one concept at a time.](#)]



Step 1.1 : Input Preparation

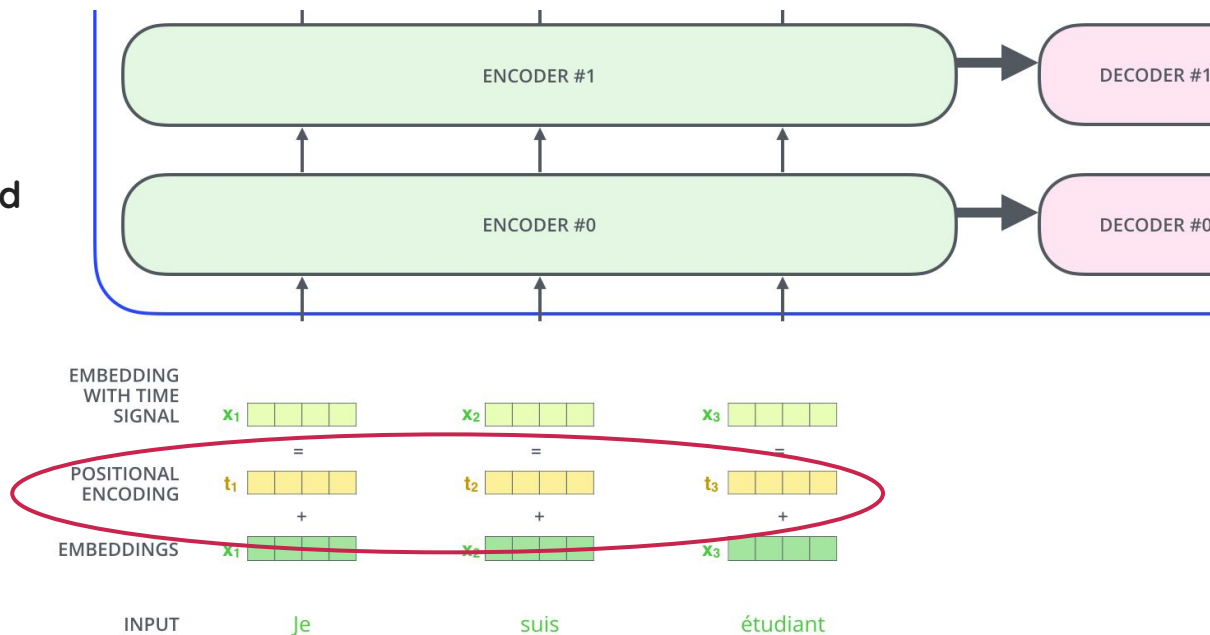
Words in the input sequence are converted into numerical representations (vectors) using embedding techniques.

These vectors are not random, they are learned during training

[Source: [The Illustrated Transformer – Jay Alammar – Visualizing machine learning one concept at a time.](#)]

Step 1.2 : Positional encoding

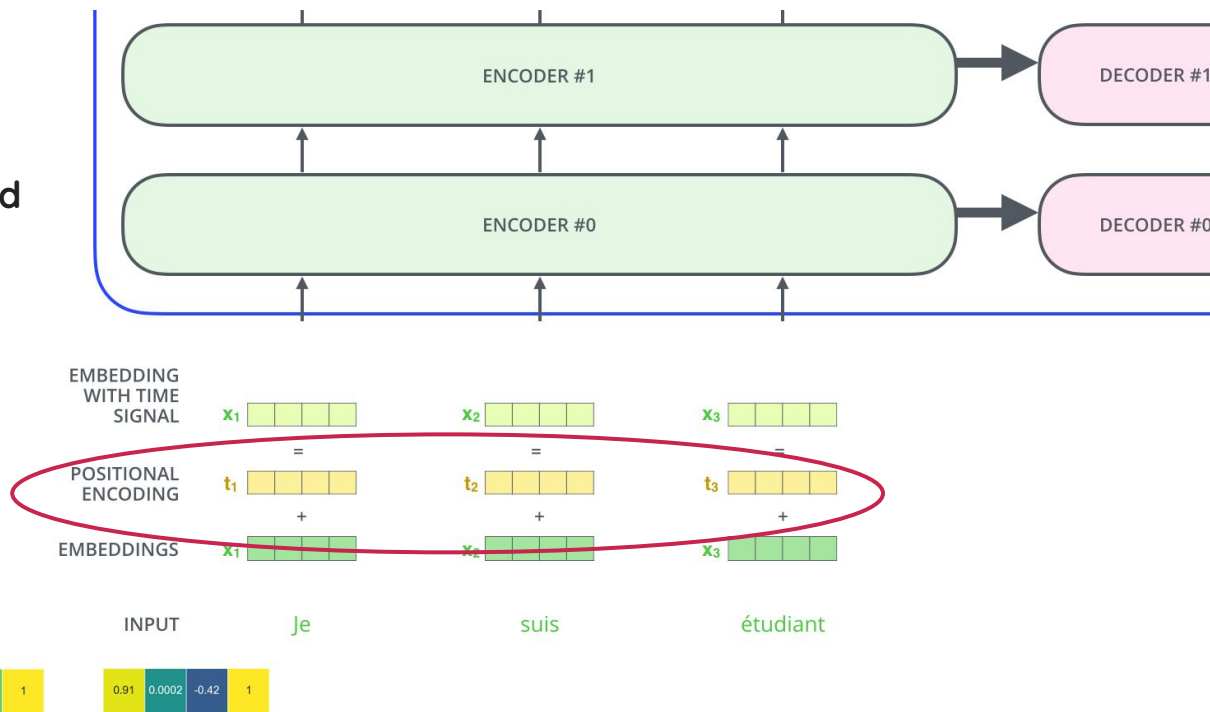
To have a clue about **order** and **position** of each word, the transformer adds a **vector** to each **input embedding**



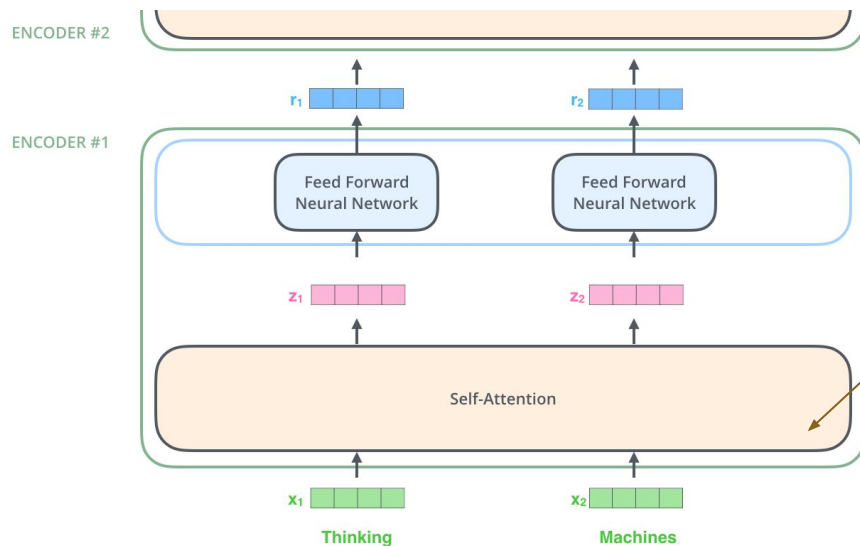
[Source: [The Illustrated Transformer – Jay Alammar – Visualizing machine learning one concept at a time.](#)]

Step 1.2 : Positional encoding

To have a clue about **order** and **position** of each word, the transformer adds a **vector** to each **input embedding**



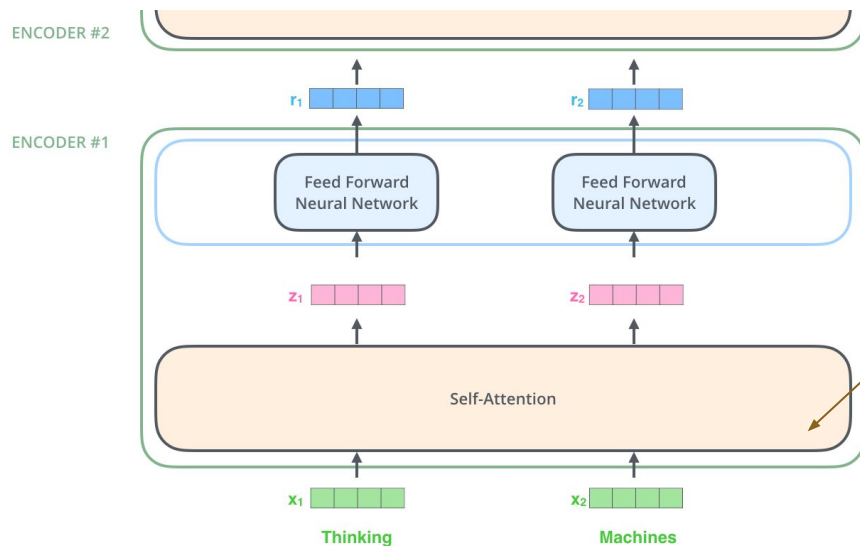
[Source: [The Illustrated Transformer – Jay Alammar – Visualizing machine learning one concept at a time.](#)]



Step 2 : Self-Attention

Analyzes the relationships between each word in the sequence, allowing the model to understand how words depend on each other.

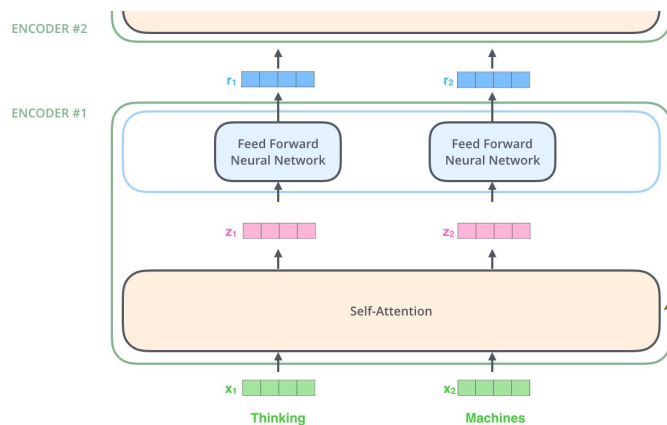
[Source: [The Illustrated Transformer – Jay Alammar – Visualizing machine learning one concept at a time.](#)]



Step 2 : Self-Attention

Which words are important for me?

[Source: [The Illustrated Transformer – Jay Alammar – Visualizing machine learning one concept at a time.](#)]



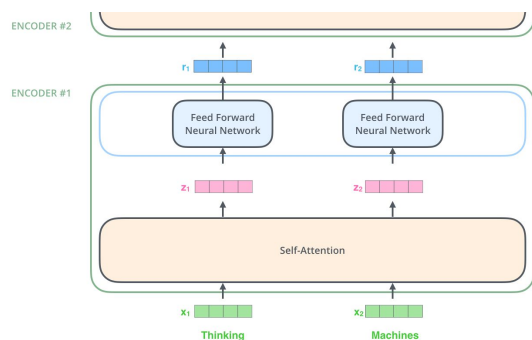
Step 2 : Self-Attention

Example

"The animal didn't cross the street
because **it** was too tired"

What does 'it' refer to?

[Source: [The Illustrated Transformer – Jay Alammar – Visualizing machine learning one concept at a time.](#)]



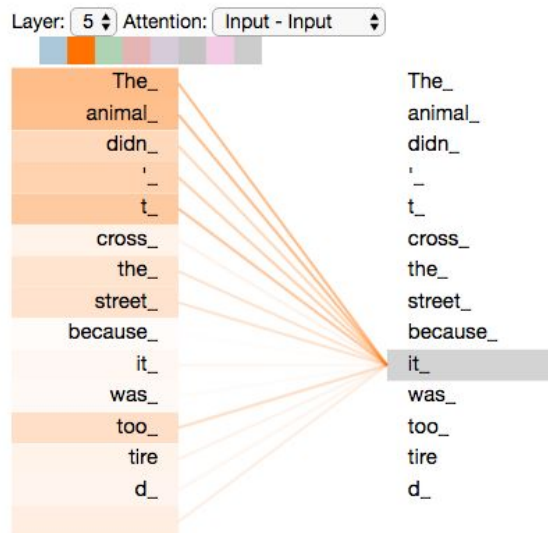
Step 2 : Self-Attention:

Example

"The animal didn't cross the street
because **it** was too tired"

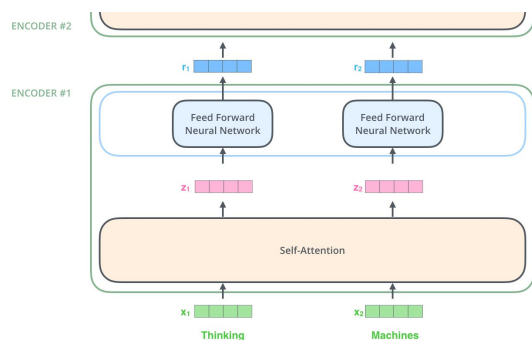
What does 'it' refer to?

We compute the attention for this
specific word

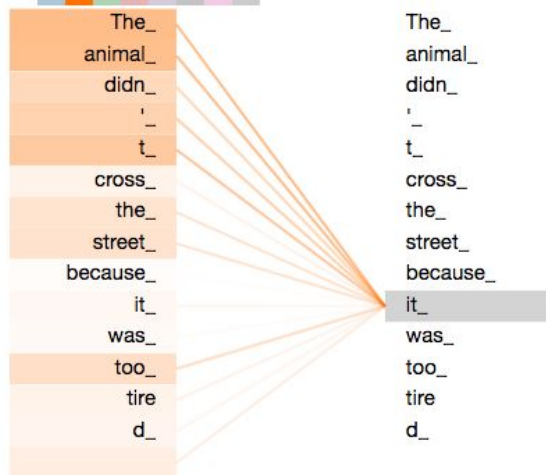


[Source: [The Illustrated Transformer – Jay Alammar – Visualizing machine learning one concept at a time.](#)]

© Mehyaar Mlaweh, 2026. All rights reserved.



Layer: 5 Attention: Input - Input



Step 2 : Self-Attention:

Example

"The animal didn't cross the street
because **it** was too tired"

What does 'it' refer to?

We compute the attention for this
specific word

-> **Attention scores** indicate which
tokens are the most relevant for the
given output

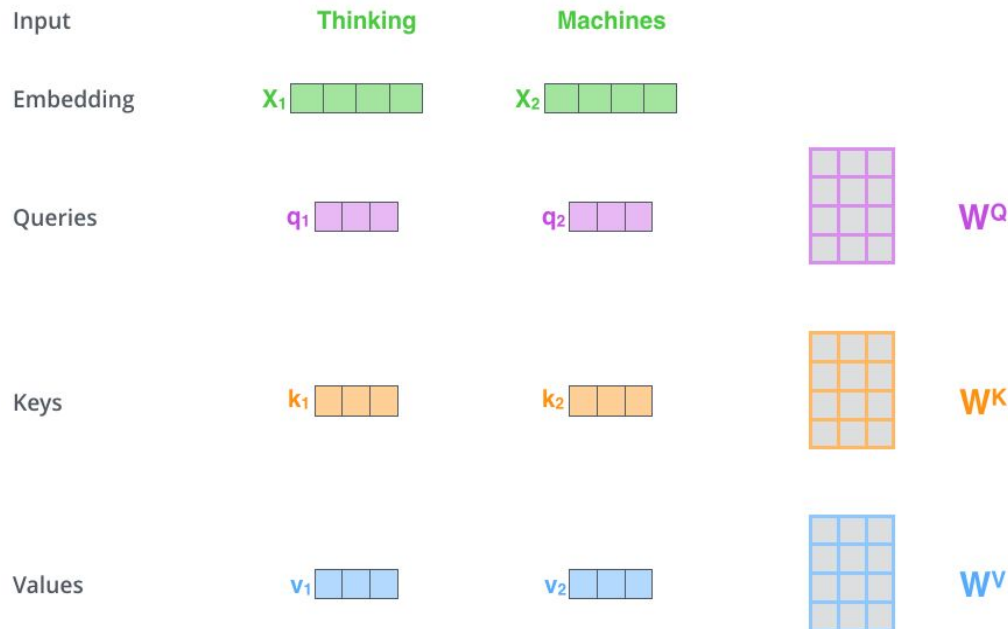
[Source: [The Illustrated Transformer – Jay Alammar – Visualizing machine learning one concept at a time.](#)]

© Mehیار Mlaweh, 2026. All rights reserved.

Self-Attention: Query, Key, Value

Step 2.1

So for each **word**, we create
a **Q vector**, a **K vector**, and
a **V vector**



[Source: [The Illustrated Transformer – Jay Alammar – Visualizing machine learning one concept at a time.](#)]

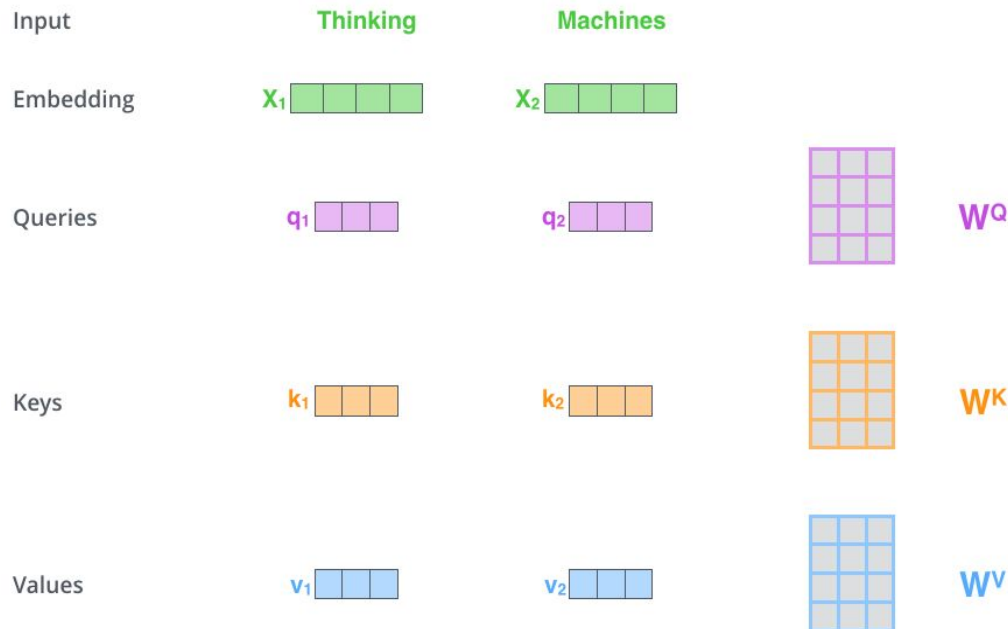
© Mehyaar Mlaweh, 2026. All rights reserved.

Self-Attention: Query, Key, Value

Step 2.1

So for each **word**, we create a **Q vector**, a **K vector**, and a **V vector**

These **vectors** are created by multiplying the embedding by the **three matrices** that we trained during the training process.



[Source: [The Illustrated Transformer – Jay Alammar – Visualizing machine learning one concept at a time.](#)]

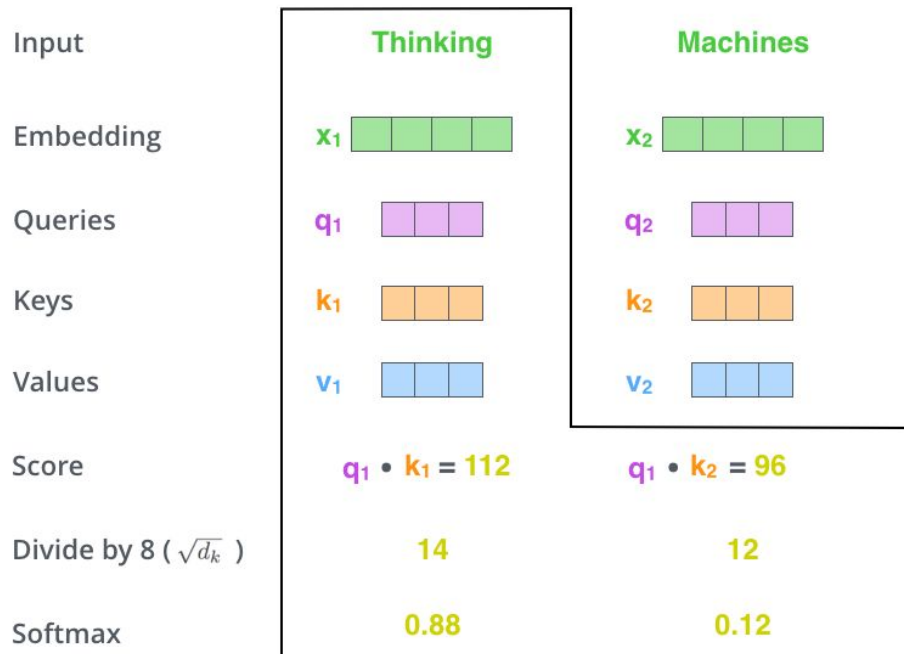
© Mehیار Mlaweh, 2026. All rights reserved.

Self-Attention: Query, Key, Value

Step 2.2

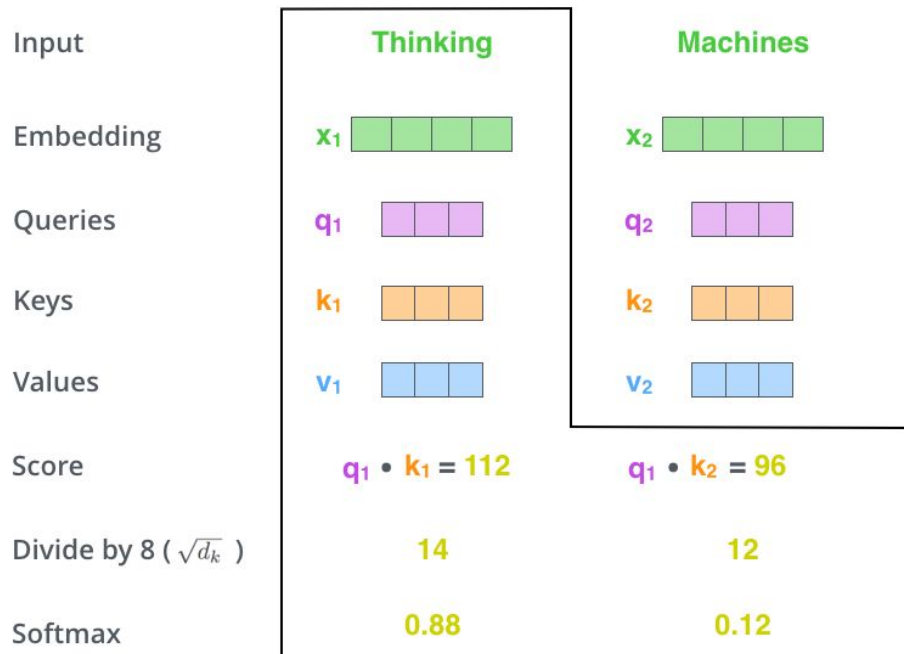
Compute attention **scores** & normalize.

We compute attention scores by taking the dot product between a query and all keys



[Source: [The Illustrated Transformer – Jay Alammar – Visualizing machine learning one concept at a time.](#)]

Self-Attention: Query, Key, Value



[Source: [The Illustrated Transformer – Jay Alammar – Visualizing machine learning one concept at a time.](#)]

Step 2.2

Compute attention **scores** & normalize.

We compute attention scores by taking the dot product between a query and all keys

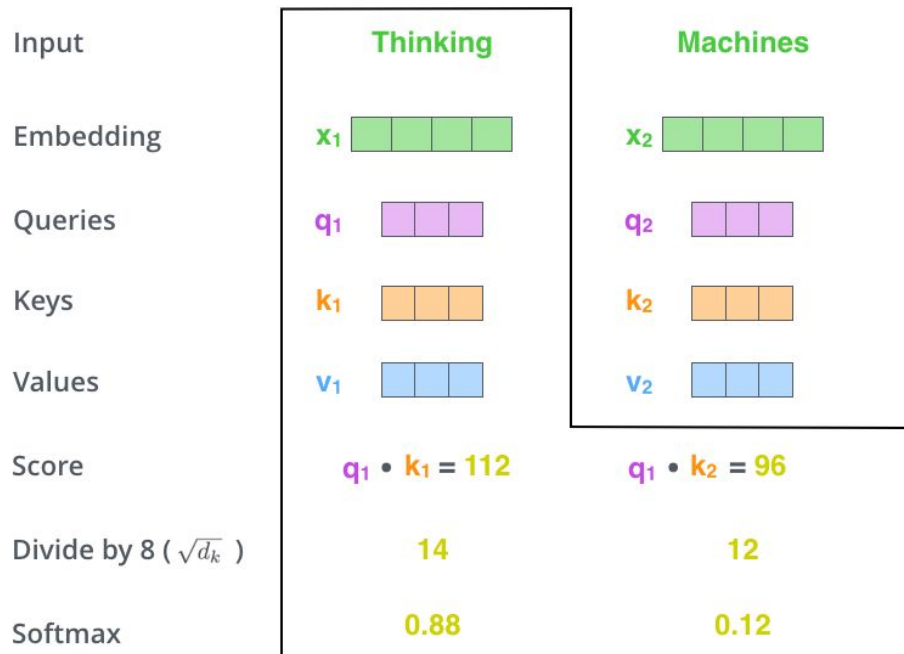
For token 1:

- Score with **itself**: $q_1 \cdot k_1$
- Score with **token 2**: $q_1 \cdot k_2$

Larger dot product $\rightarrow$ higher relevance.

Which word should “**Thinking**” attend to more? Why?

Self-Attention: Query, Key, Value



[Source: [The Illustrated Transformer – Jay Alammar – Visualizing machine learning one concept at a time.](#)]

Step 2.2

Compute attention **scores** & normalize.

We compute attention scores by taking the dot product between a query and all keys

For token 1:

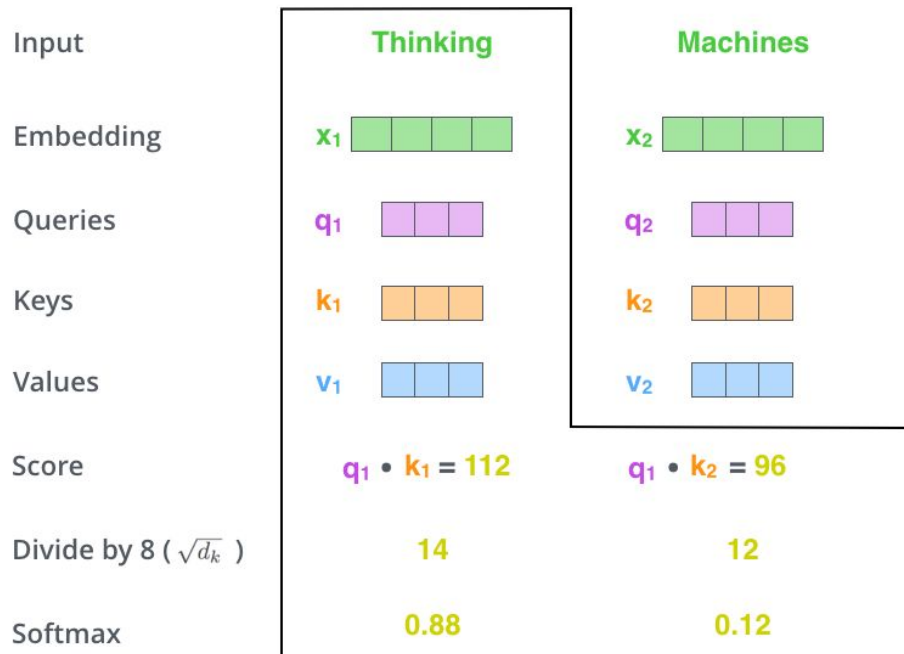
- Score with **itself**: $q_1 \cdot k_1$
- Score with **token 2**: $q_1 \cdot k_2$

Larger dot product $\rightarrow$ higher relevance.

112 > 96

The word **Thinking** : “*I already carry useful information* :)”

Self-Attention: Query, Key, Value



[Source: [The Illustrated Transformer – Jay Alammar – Visualizing machine learning one concept at a time.](#)]

Step 2.2

Compute attention **scores** & normalize.

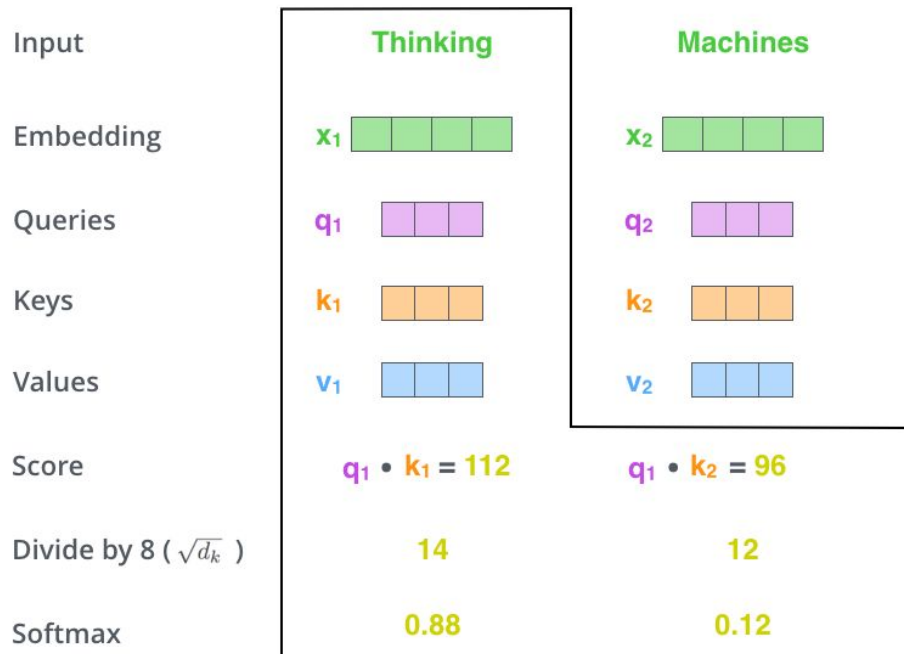
We compute attention scores by taking the dot product between a query and all keys

For token 1:

- Score with **itself**: $q_1 \cdot k_1$
- Score with **token 2**: $q_1 \cdot k_2$

The dot product can become very large, especially with high-dimensional vectors

Self-Attention: Query, Key, Value



[Source: [The Illustrated Transformer – Jay Alammar – Visualizing machine learning one concept at a time.](#)]

Step 2.2

Compute attention **scores** & normalize.

We compute attention scores by taking the dot product between a query and all keys

For token 1:

- Score with **itself**: $q_1 \cdot k_1$
- Score with **token 2**: $q_1 \cdot k_2$

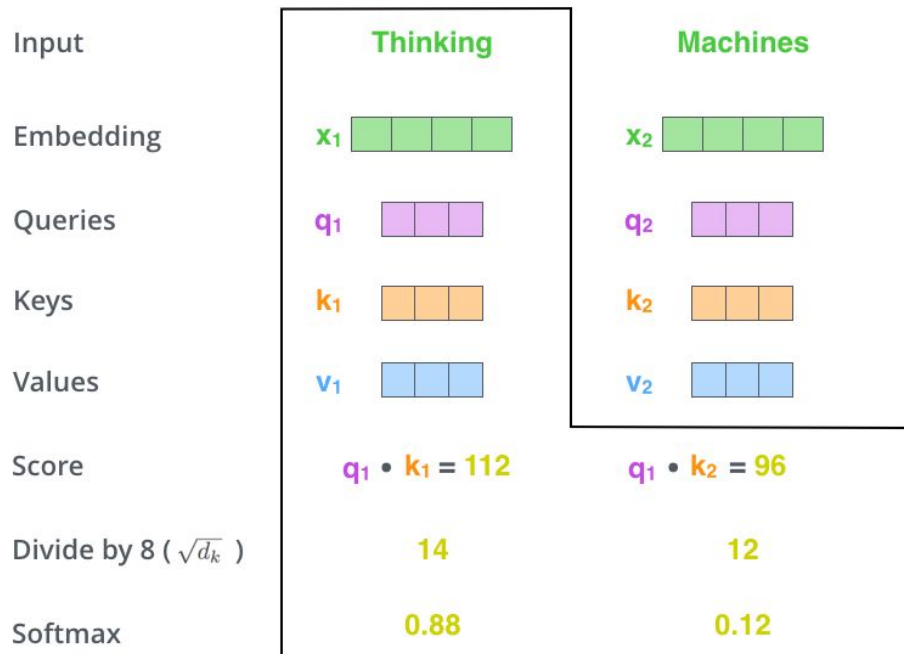
The dot product can become very large, especially with high-dimensional vectors

So we divide by $\sqrt{d}$ to keep values in a reasonable range

Self-Attention: Query, Key, Value

Step 2.3

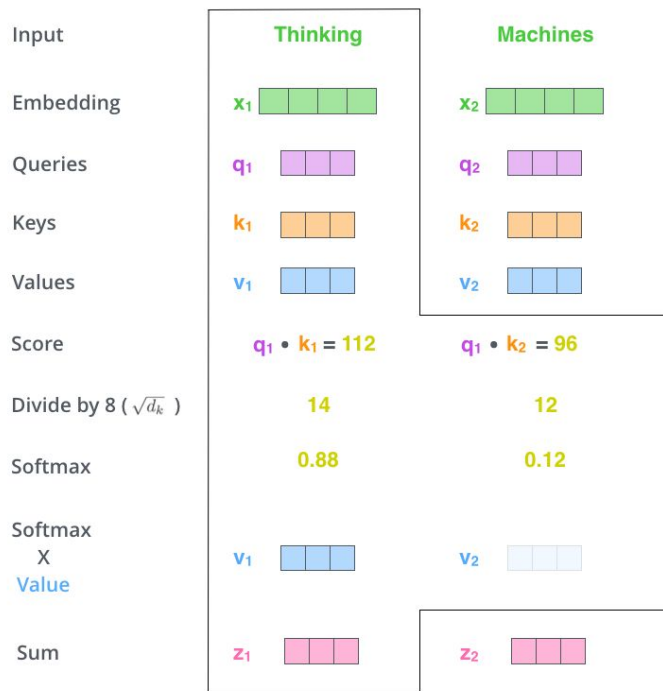
Softmax



Softmax turns raw scores into attention weights (probabilities)

[Source: [The Illustrated Transformer – Jay Alammar – Visualizing machine learning one concept at a time.](#)]

Self-Attention: Query, Key, Value



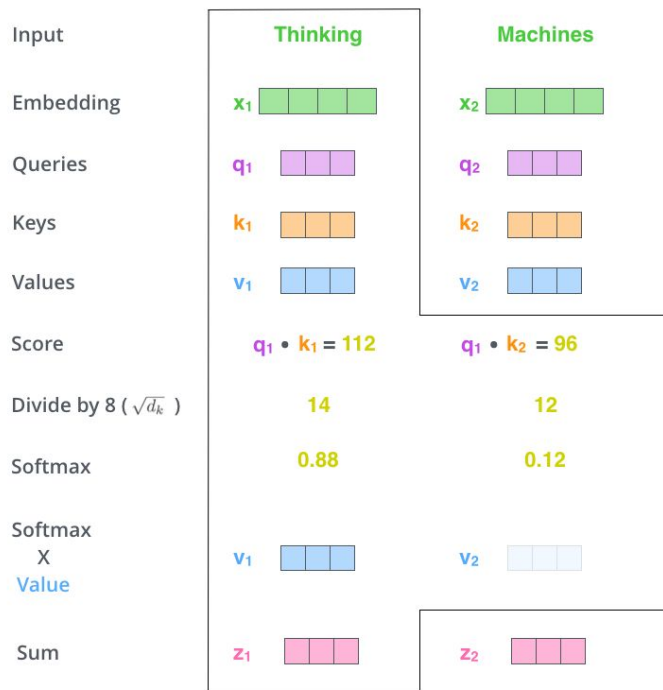
Step 2.4

Apply weights to values & sum

Now we multiply each **value** by its attention weight

[Source: [The Illustrated Transformer – Jay Alammar – Visualizing machine learning one concept at a time.](#)]

Self-Attention: Query, Key, Value



Step 2.4

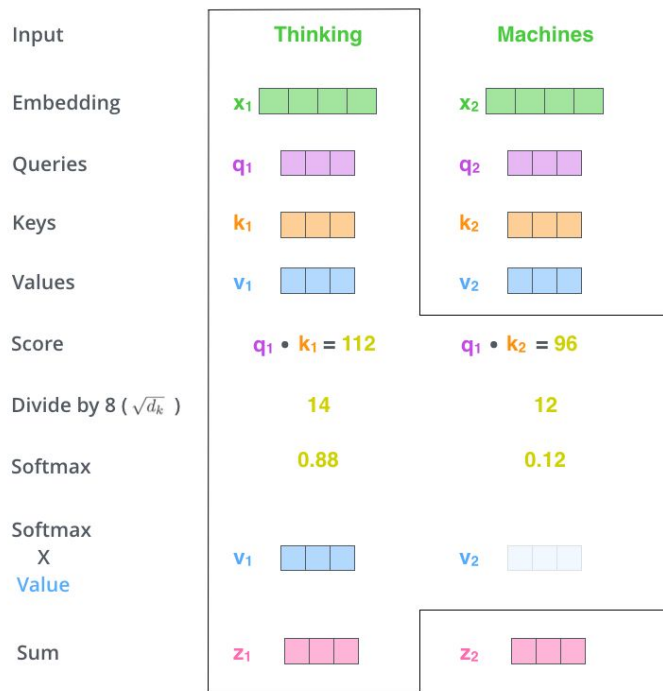
Apply weights to values & sum

Now we multiply each **value** by its attention weight

Finally, we **sum** everything to get the new representation z_1

[Source: [The Illustrated Transformer – Jay Alammar – Visualizing machine learning one concept at a time.](#)]

Self-Attention: Query, Key, Value

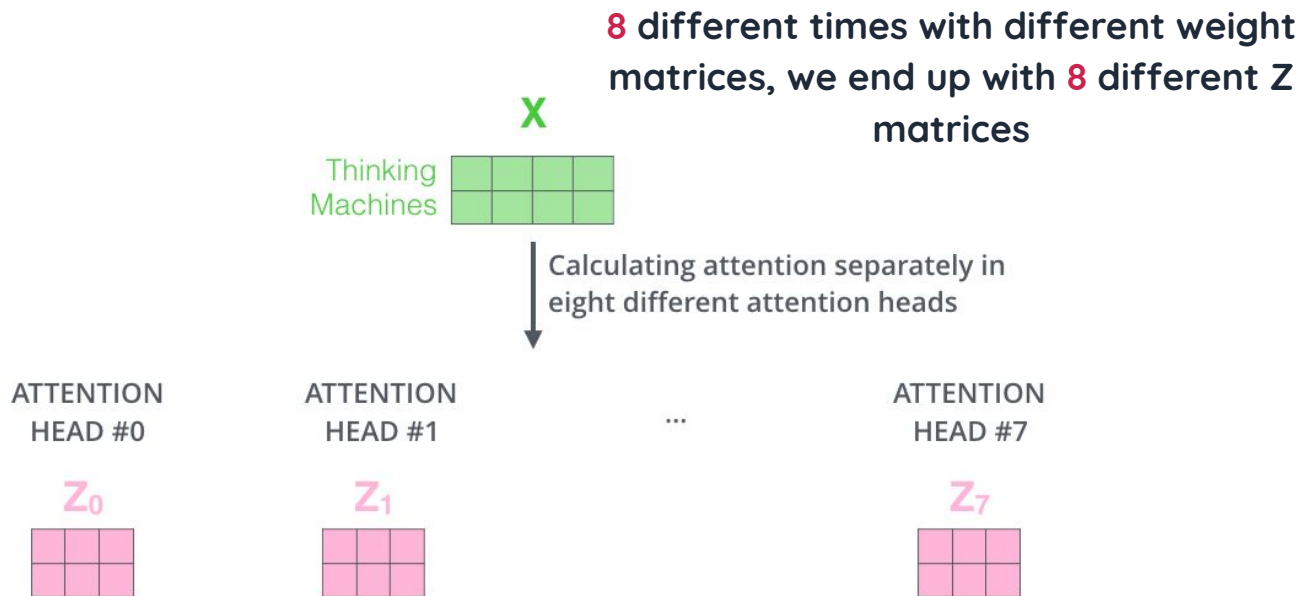


We just computed the output for **ONE** word

We repeat this for **every** word in parallel

[Source: [The Illustrated Transformer – Jay Alammar – Visualizing machine learning one concept at a time.](#)]

Self-Attention: Query, Key, Value

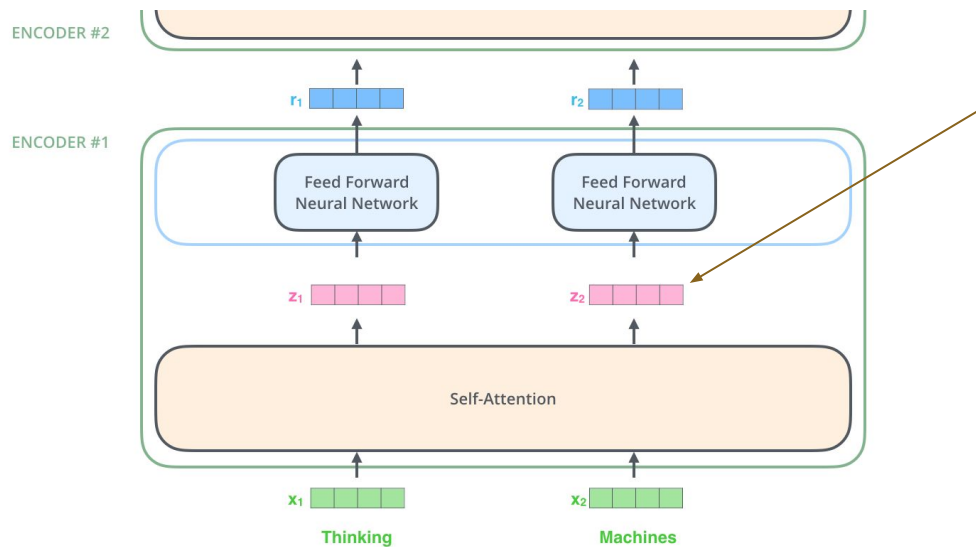


[Source: [The Illustrated Transformer – Jay Alammar – Visualizing machine learning one concept at a time.](#)]

PS: Original Transformer (Vaswani et al., 2017) : 8 heads

Self-Attention: Query, Key, Value

Interactive: [Visualizing Attention](#)

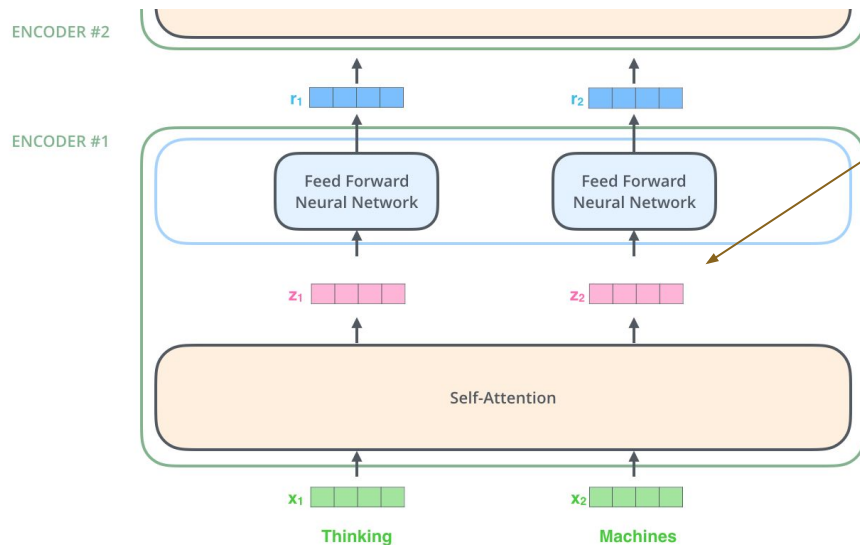


Step 3

Concatenate all the attention head

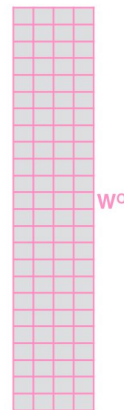


[Source: [The Illustrated Transformer – Jay Alammar – Visualizing machine learning one concept at a time.](#)]

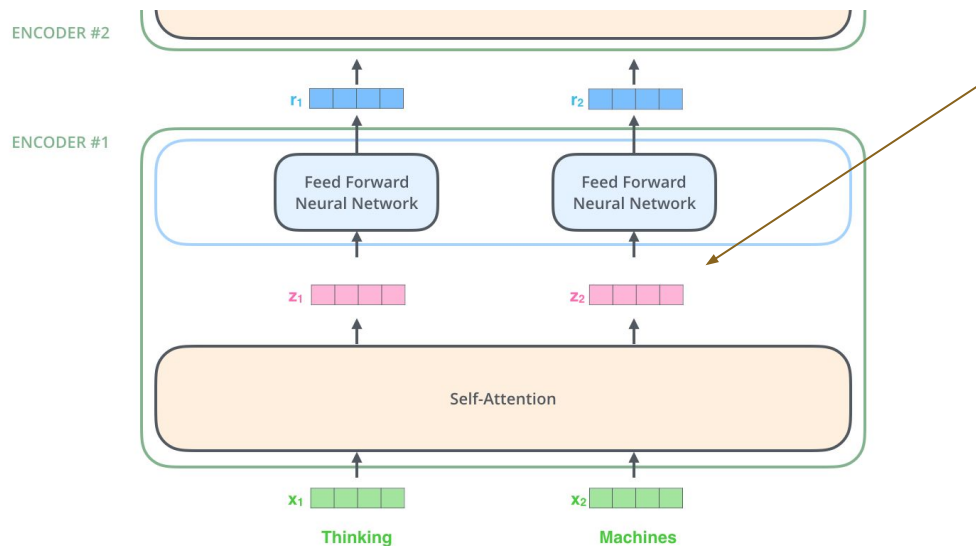


Step 3

Multiply with a weigh matrix W_0 that was trained with the model

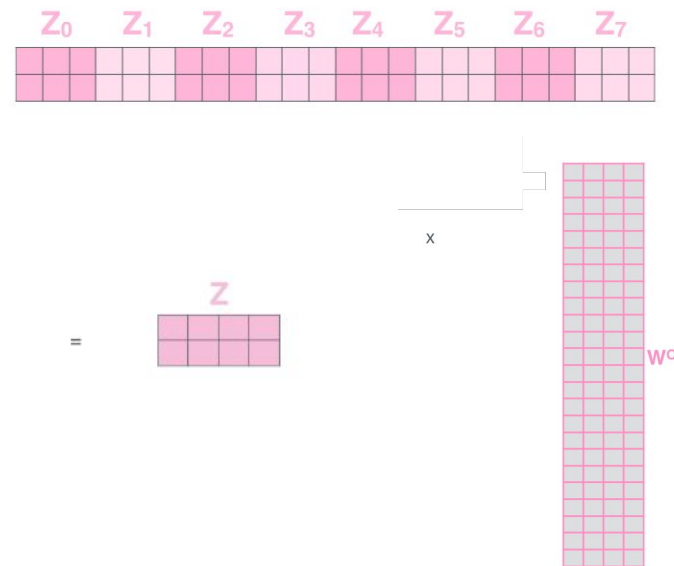


[Source: [The Illustrated Transformer – Jay Alammar – Visualizing machine learning one concept at a time.](#)]

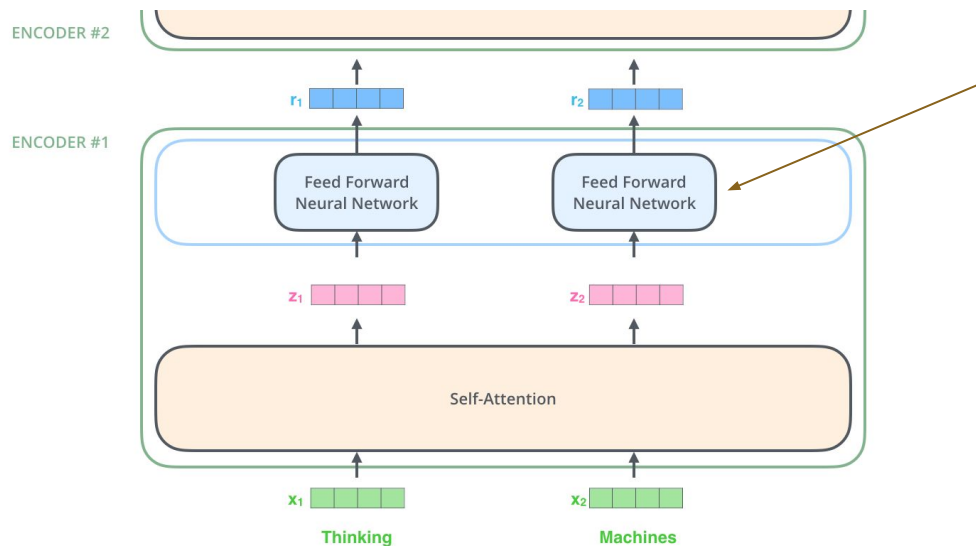


Step 3

Result would be the **Z** matrix



[Source: [The Illustrated Transformer – Jay Alammar – Visualizing machine learning one concept at a time.](#)]



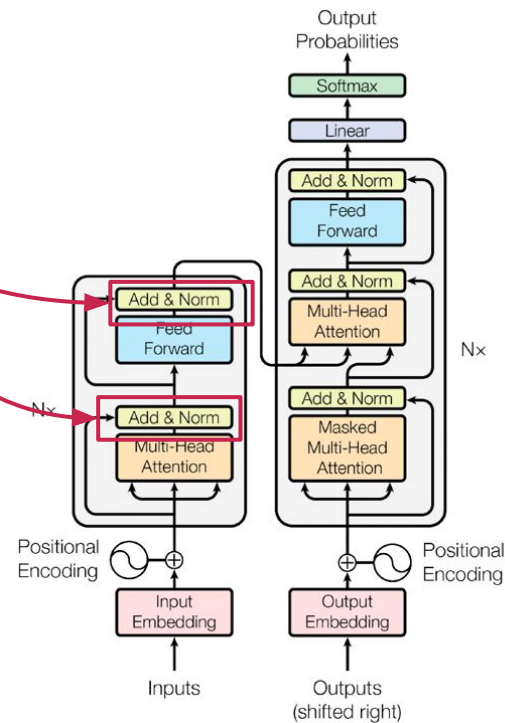
Step 4

Now we apply a **feed-forward** network to each token

no interaction here, each token is processed **independently**

[Source: [The Illustrated Transformer – Jay Alammar – Visualizing machine learning one concept at a time.](#)]

Add (Residual) & Normalize

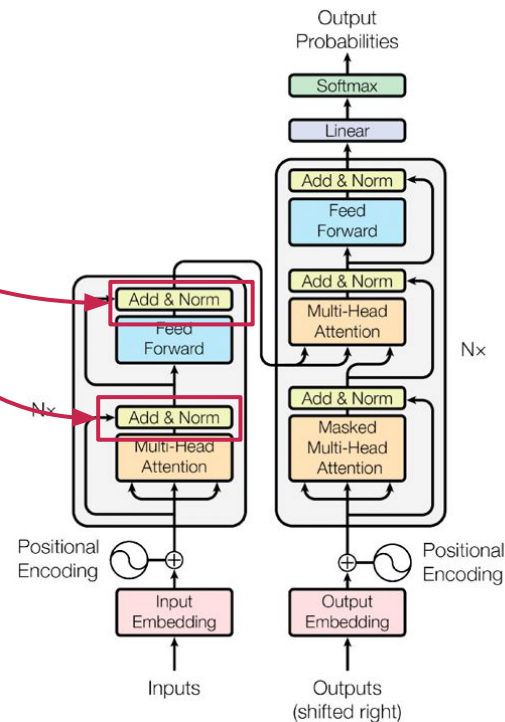


Add (Residual) & Normalize

Residual Connection:

- Add input back to output
- Helps gradient flow
- Prevents information loss

$$\text{Output} = \text{Input} + \text{Layer}(x)$$



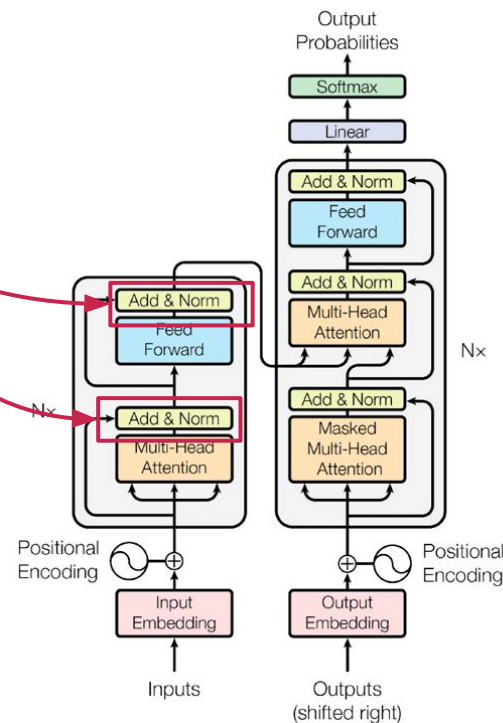
Add (Residual) & Normalize

Residual Connection:

- Add input back to output
- Helps gradient flow
- Prevents information loss

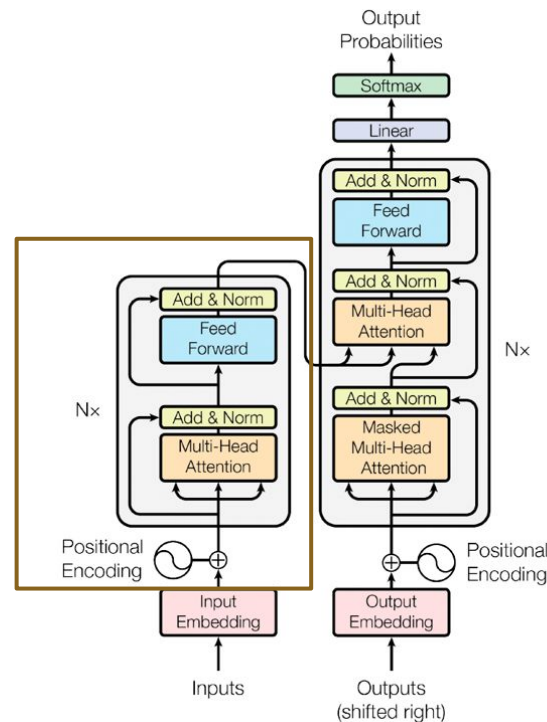
$$\text{Output} = \text{Input} + \text{Layer}(x)$$

$$\text{Final Output} = \text{LayerNorm}(x + \text{Layer}(x))$$

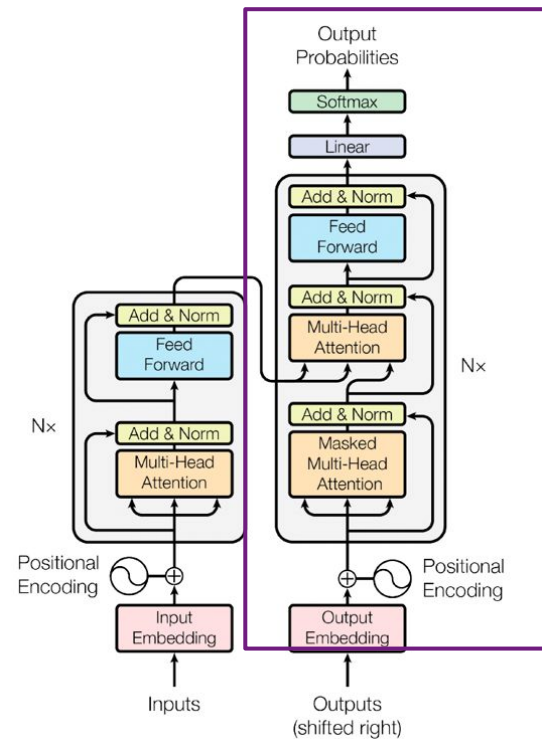


Key takeaways:

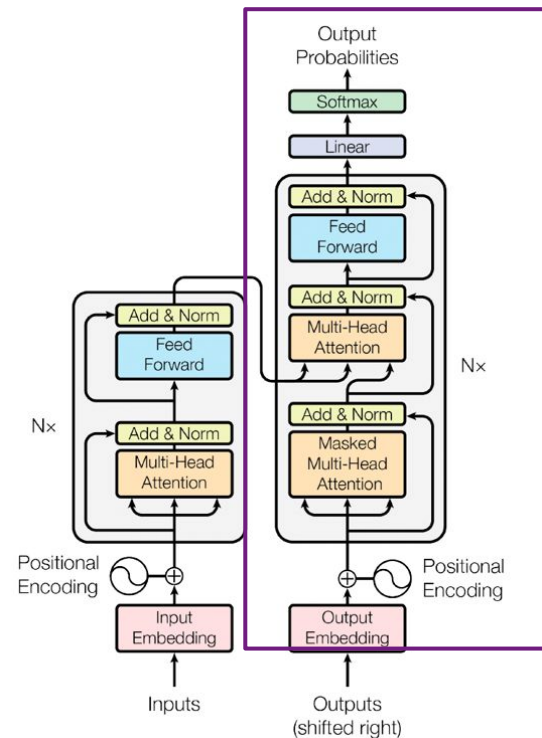
- The encoder extracts meaning from the input sequence.
- the encoder creates a rich representation for the decoder to work with.



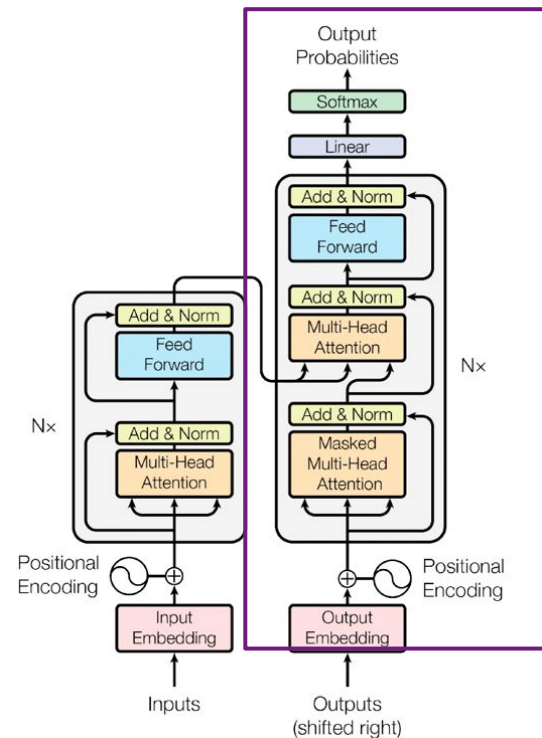
Any idea how it works ?



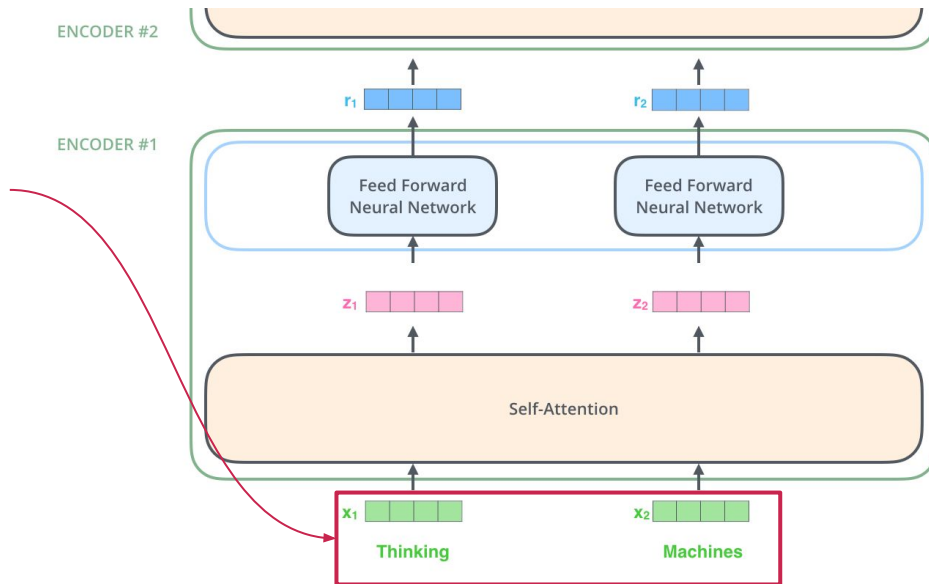
The decoder is very powerful
and
very complicated



The decoder is very powerful
and
very complicated
And completely useless for
what we're doing today

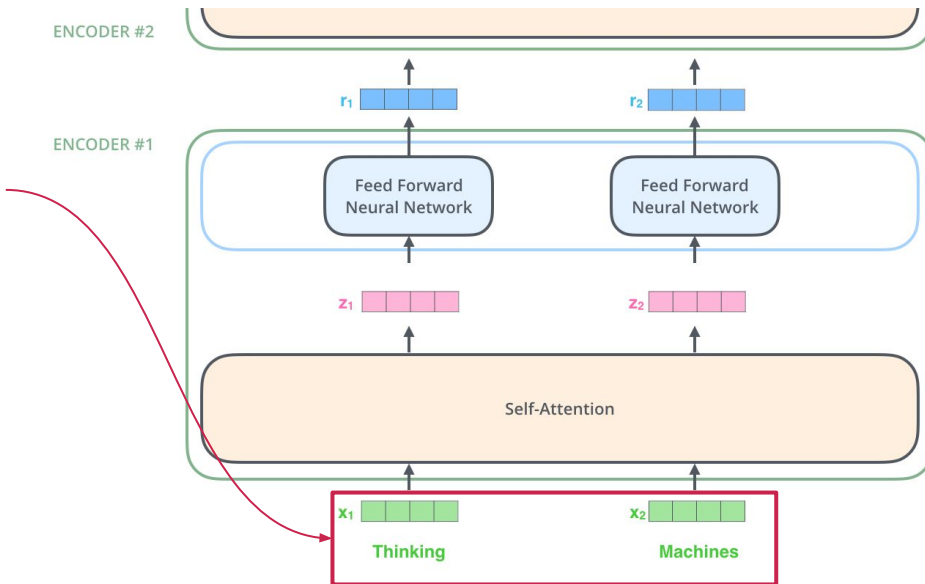


So far, our transformer takes words as input



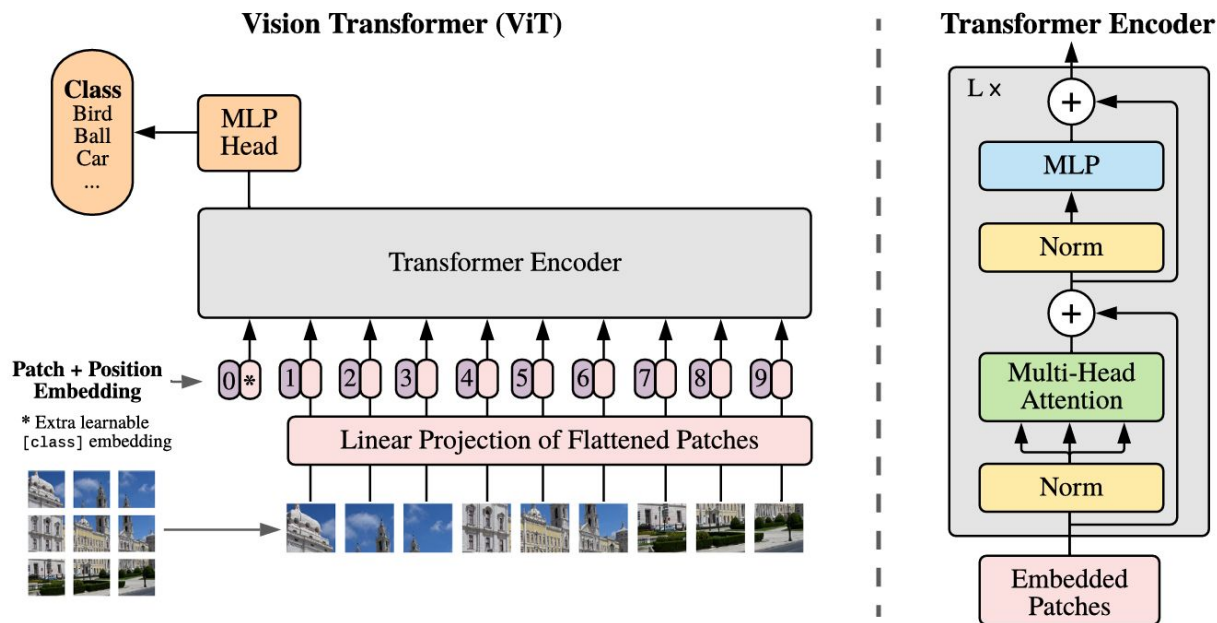
[Source: [The Illustrated Transformer – Jay Alammar – Visualizing machine learning one concept at a time.](#)]

So far, our transformer takes
words as input
Now we want to apply this to
images...



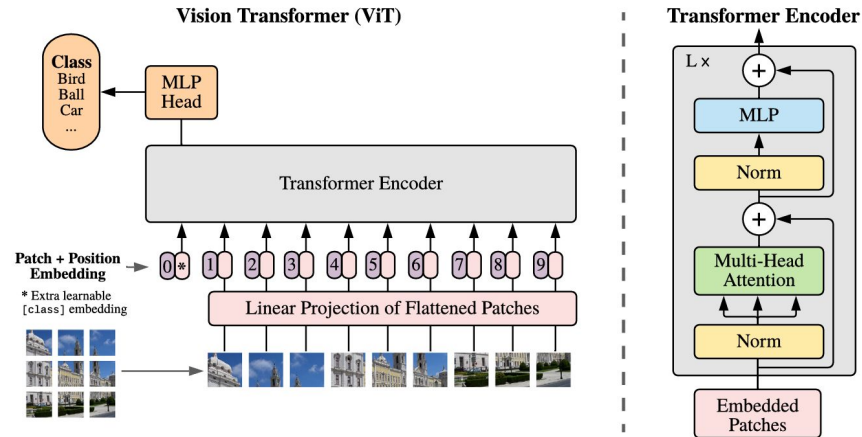
[Source: [The Illustrated Transformer – Jay Alammar – Visualizing machine learning one concept at a time.](#)]

No decoder !!



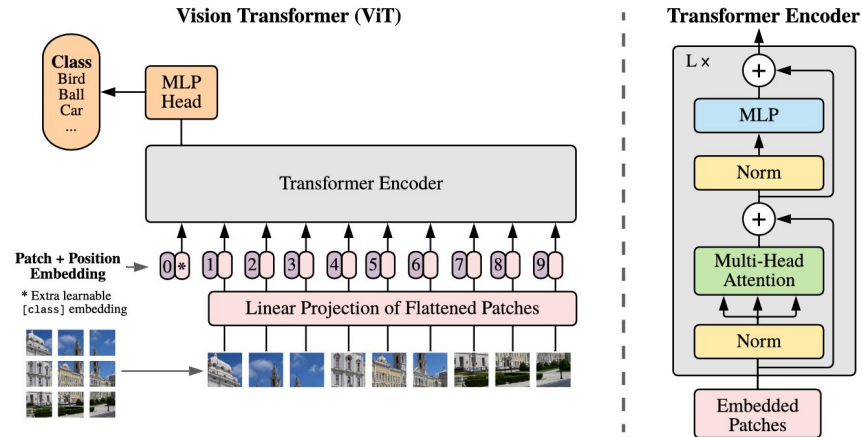
[Source: [An Image is Worth 16x16 Words: Transformers for Image Recognition at Scale](#)]

- **First** full Transformer architecture for Vision
- **16 x 16 patches** as 'words' : each patch = $16 \times 16 \times 3$ (=768d)
- Huge pre-training (Imagenet doesn't suffice)



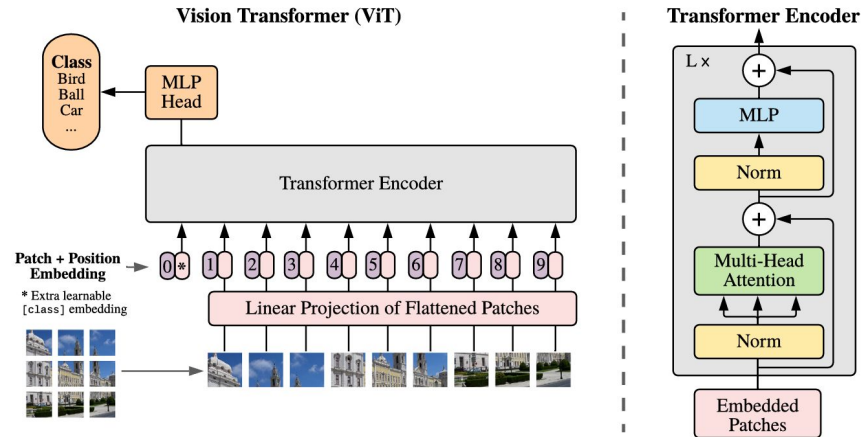
[Source: [An Image is Worth 16x16 Words: Transformers for Image Recognition at Scale](#)]

- If **ViT** is trained on datasets with more than 14M images it can approach or beat state-of-the-art CNNs.



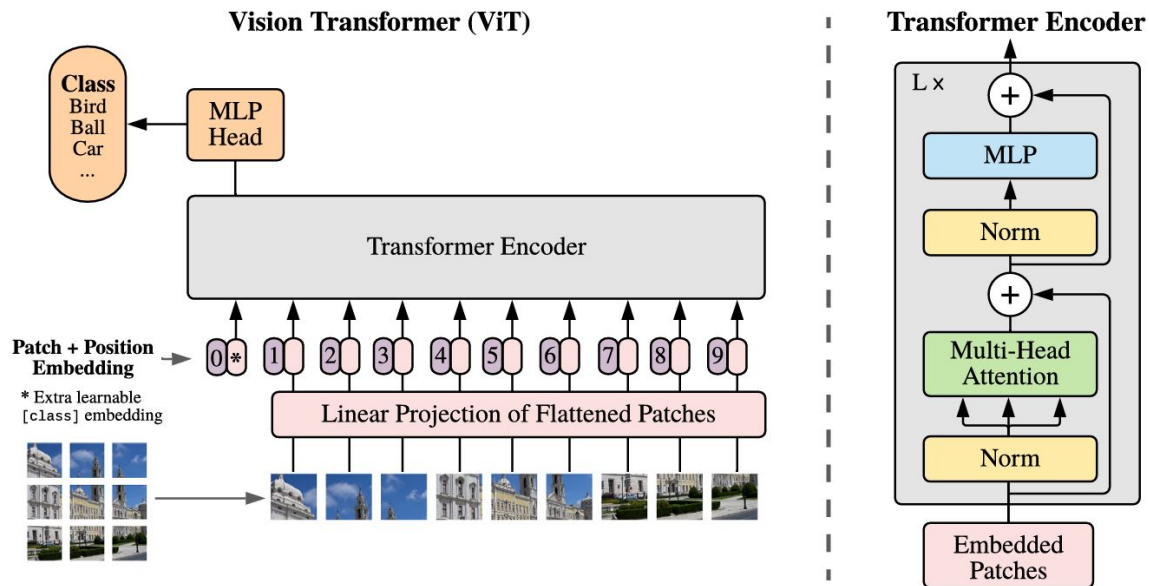
[Source: [An Image is Worth 16x16 Words: Transformers for Image Recognition at Scale](#)]

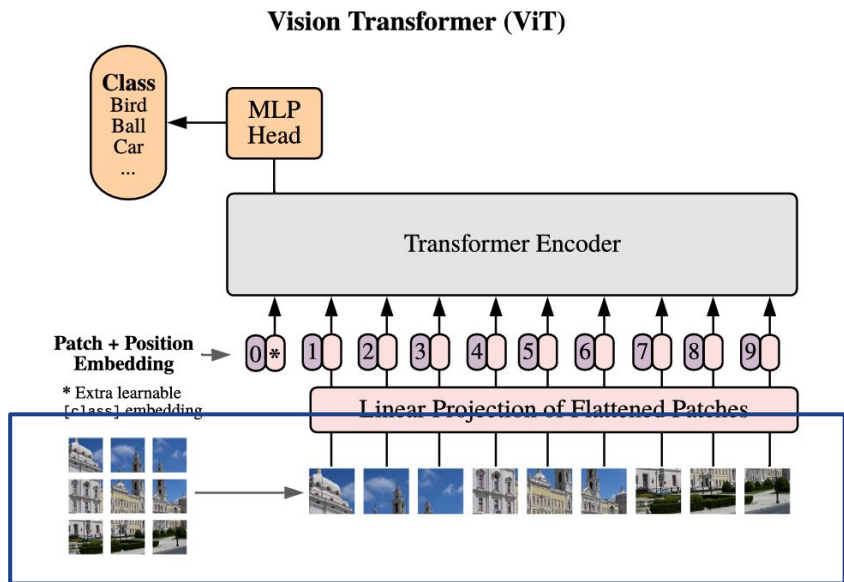
- If **ViT** is trained on datasets with more than 14M images it can approach or beat state-of-the-art CNNs.
- If not, you better stick with **ResNets** or **EfficientNets**.



[Source: [An Image is Worth 16x16 Words: Transformers for Image Recognition at Scale](#)]

How does it really works ?





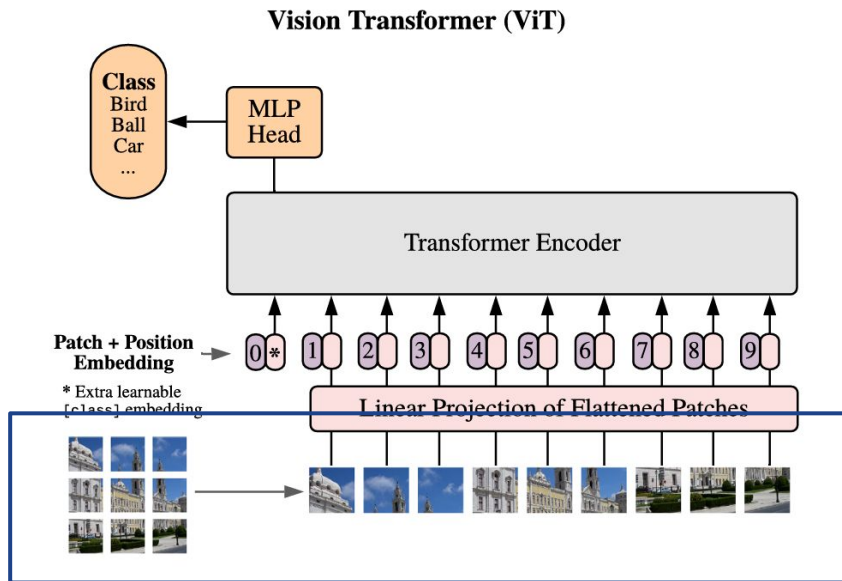
Step 1

Split Image into **Patches**

- Image is divided into fixed-size patches (e.g., **16×16**)
- Each patch becomes a **token**
- Total patches = sequence length

Step 1

Split Image into Patches



- Image is divided into fixed-size patches (e.g., 16×16)
- Each patch becomes a **token**
- Total patches = sequence length

Where H , W = **image size**

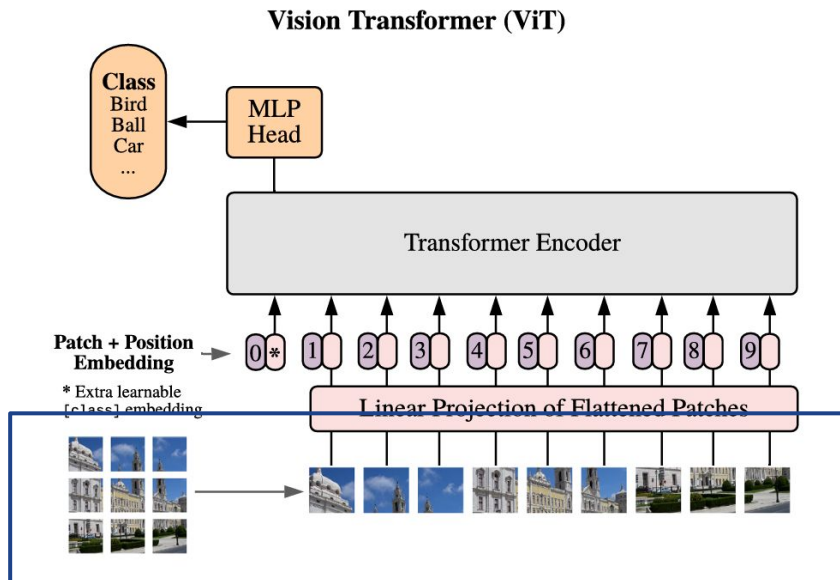
P = **patch size**

$$N = \frac{H}{P} \times \frac{W}{P}$$

Step 1

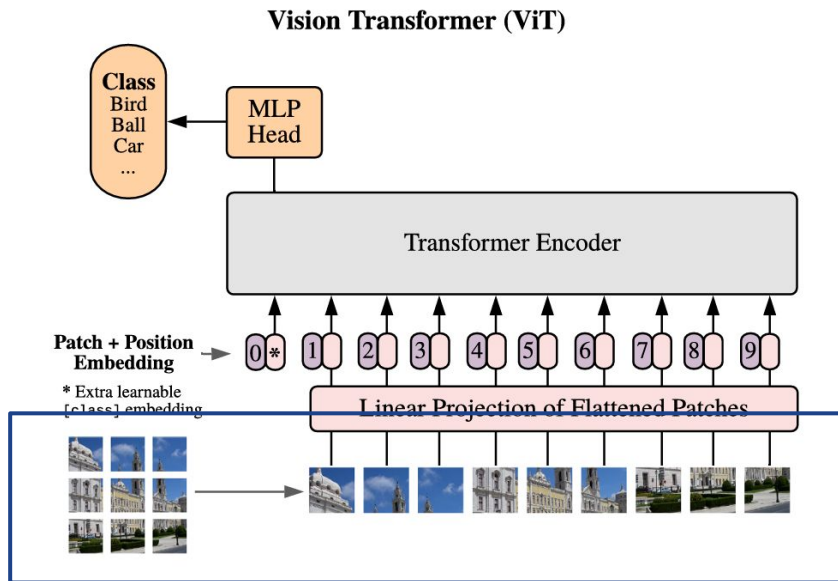
Split Image into **Patches**

Why patches ? why not pixels?



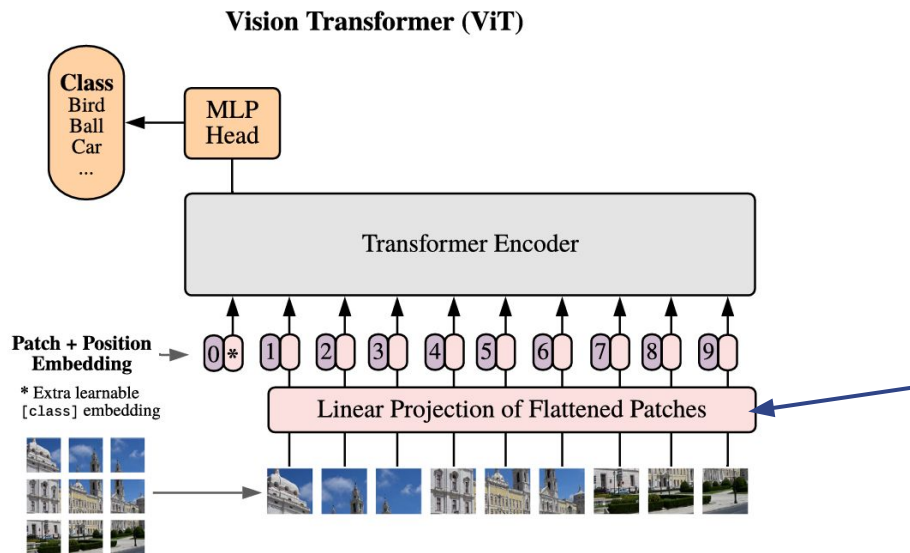
Step 1

Split Image into **Patches**



Why patches ? why not pixels?

It is relatively **easier** to understand the relationships between **patches** of $P \times P$ than of a **full image** Height $\times$ Width.



Step 2.1

Patch → Embedding

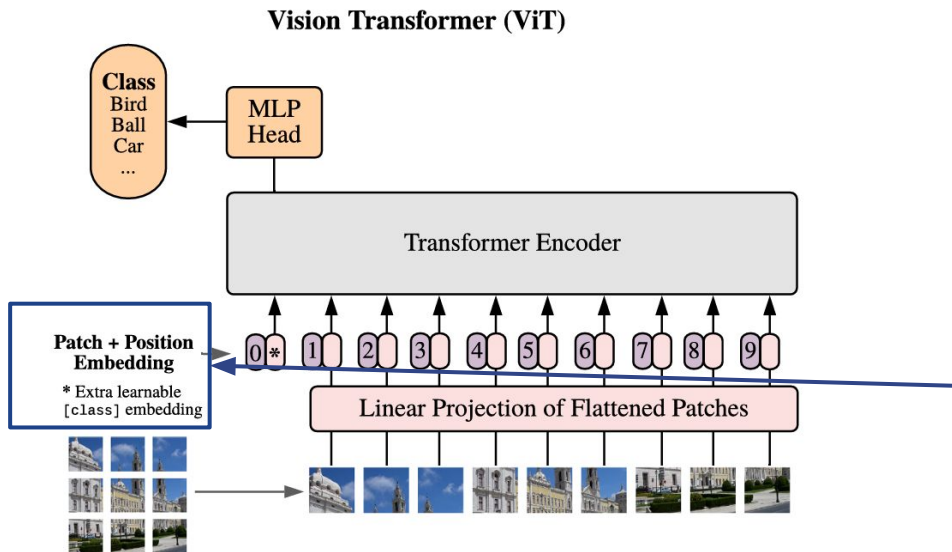
- Each patch is a small image (e.g., $16 \times 16 \times 3$)
- **Flatten** into a vector
- Apply a **linear** projection

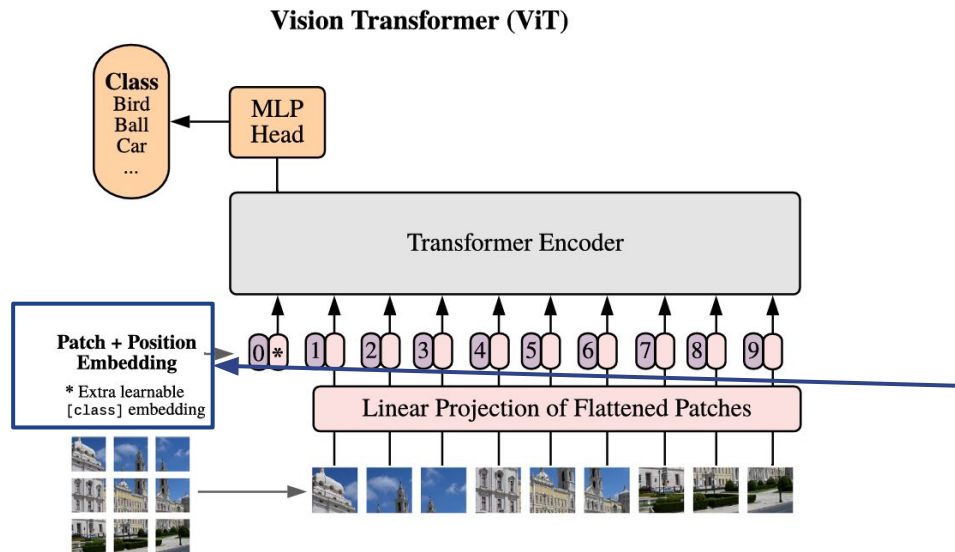
Each patch → a token vector

Step 2.2

Positional embeddings

- **Problem** : Transformer does not know spatial order





Step 2.2

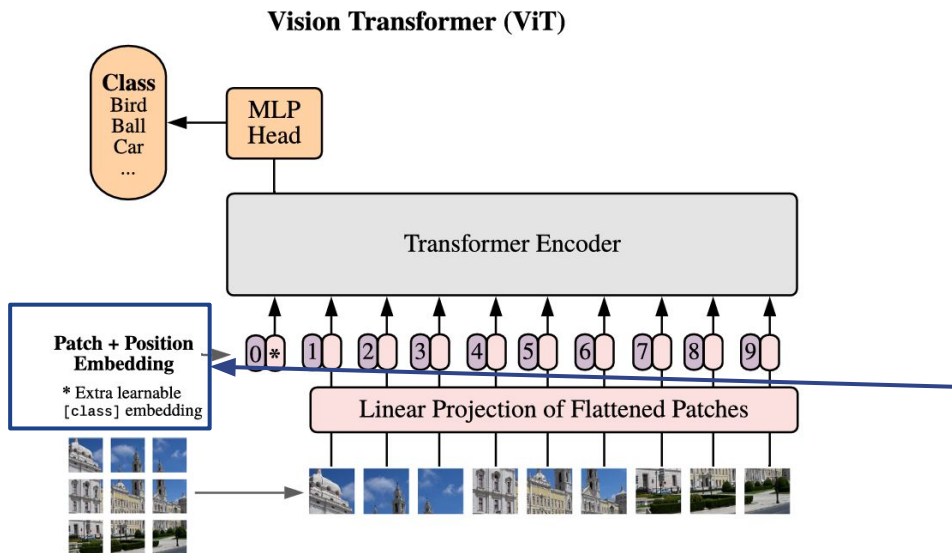
Positional embeddings

- **Problem** : Transformer does not know spatial order
- **Solution** : Add a **positional vector** to each patch
- Encodes spatial **location**

Step 2.2

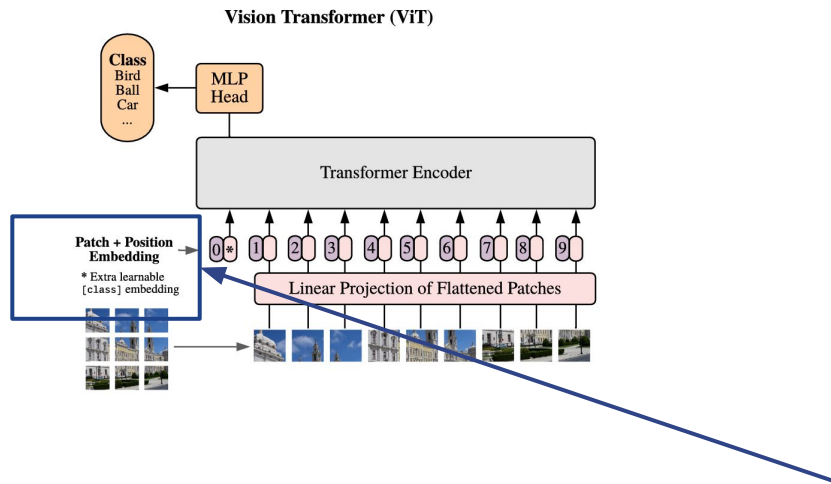
Positional embeddings

- **Problem** : Transformer does not know spatial order
- **Solution** : Add a **positional vector** to each patch
- Encodes spatial **location**



Embedding = Patch + Position

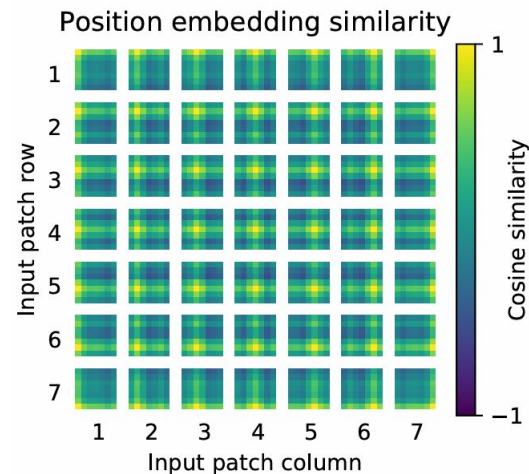
* a **trainable position**



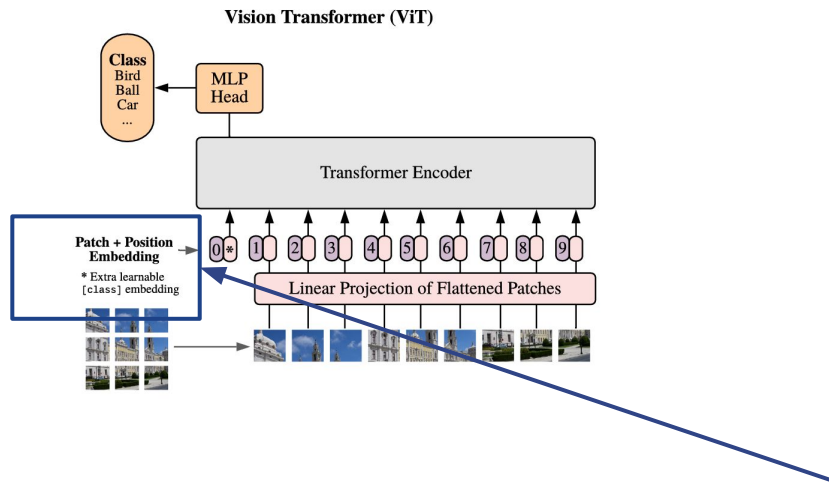
Similarity between **positional** embeddings.

Step 2.2

Positional embeddings



[Source: [An Image is Worth 16x16 Words: Transformers for Image Recognition at Scale](#)]

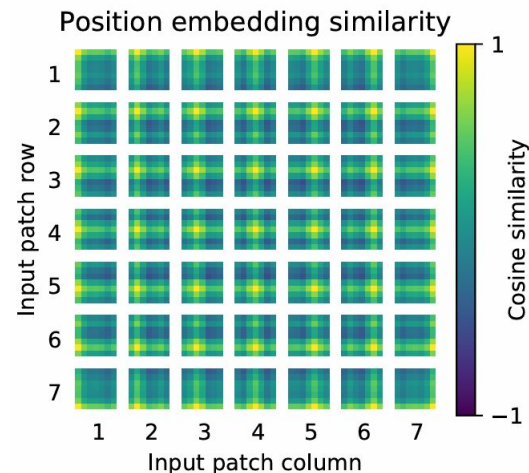


Similarity between **positional** embeddings.

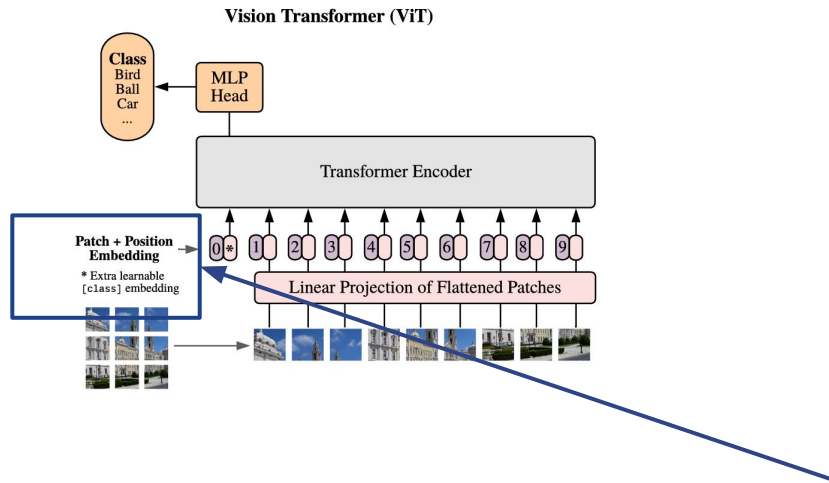
Why do we have 7x7 heatmap ?

Step 2.2

Positional embeddings



[Source: [An Image is Worth 16x16 Words: Transformers for Image Recognition at Scale](#)]

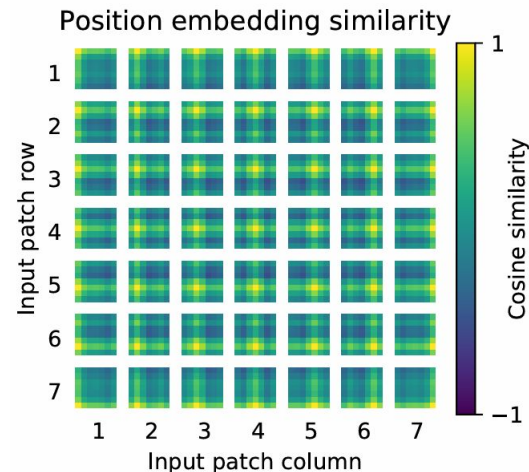


Similarity between **positional** embeddings.

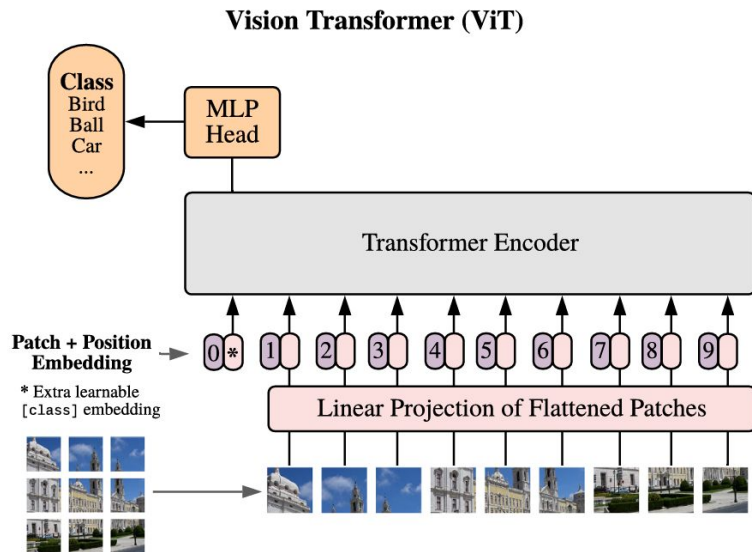
If we divide a **112×112** image (image size used in the original paper) into **16×16** patches
 $112/16 = 7 \times 7$ patch positions

Step 2.2

Positional embeddings



[Source: [An Image is Worth 16x16 Words: Transformers for Image Recognition at Scale](#)]



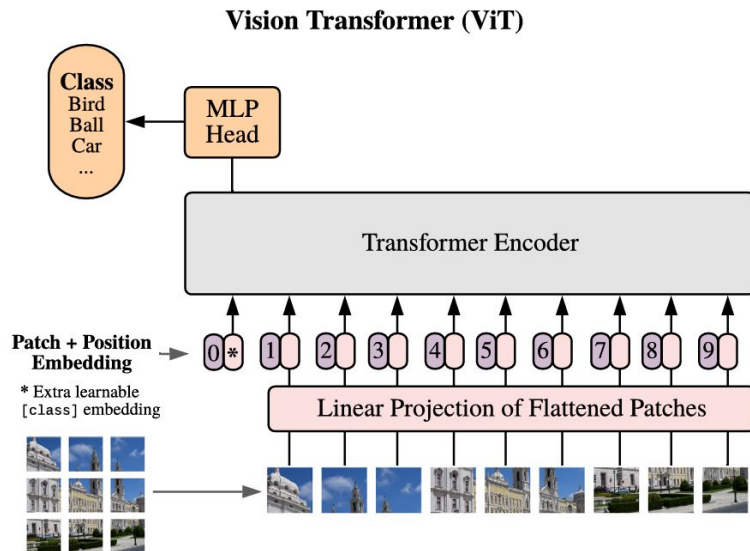
Step 2.3

Add **[CLS]** Token

After splitting the image into **patches**, converting them into **embeddings**, and adding **positional** information

We have sequence of vectors:

$x_1, x_2, x_3, \dots, x_n$



Step 2.3

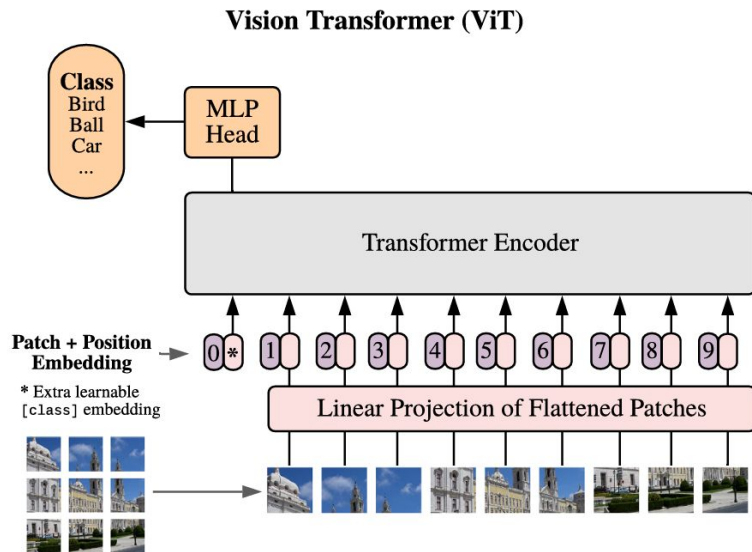
Add **[CLS]** Token

After splitting the image into **patches**, converting them into **embeddings**, and adding **positional** information

We have sequence of vectors:

$x_1, x_2, x_3, \dots, x_n$

- what the patch contains
- where it is located



Step 2.3

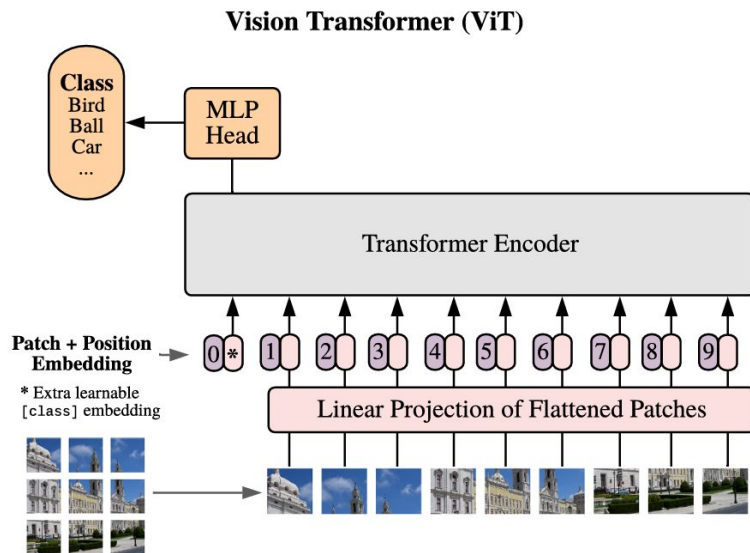
Add **[CLS]** Token

We have sequence of vectors:

$x_1, x_2, x_3, \dots, x_n$

- what the patch contains
- where it is located

Problem : For a classification task, which vector should we use to make the final prediction



Step 2.3

Add **[CLS]** Token

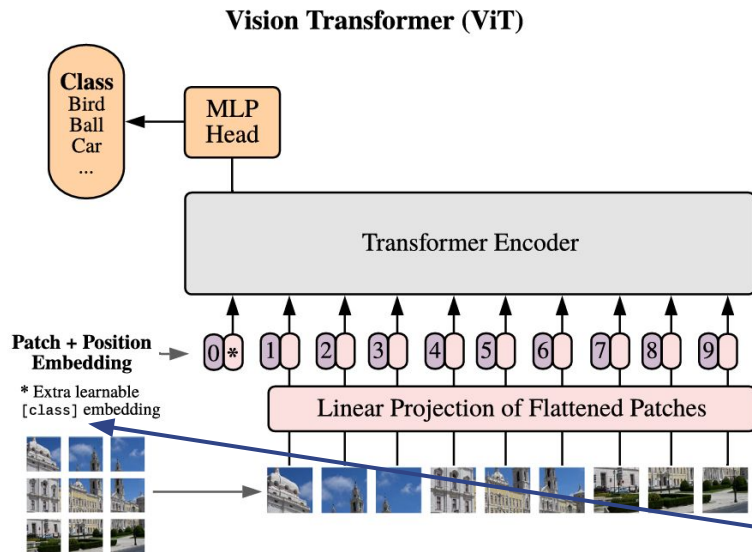
We have sequence of vectors:

$x_1, x_2, x_3, \dots, x_n$

- what the patch contains
- where it is located

Problem : For a classification task, which vector should we use to make the final prediction.

Solution : One final representation



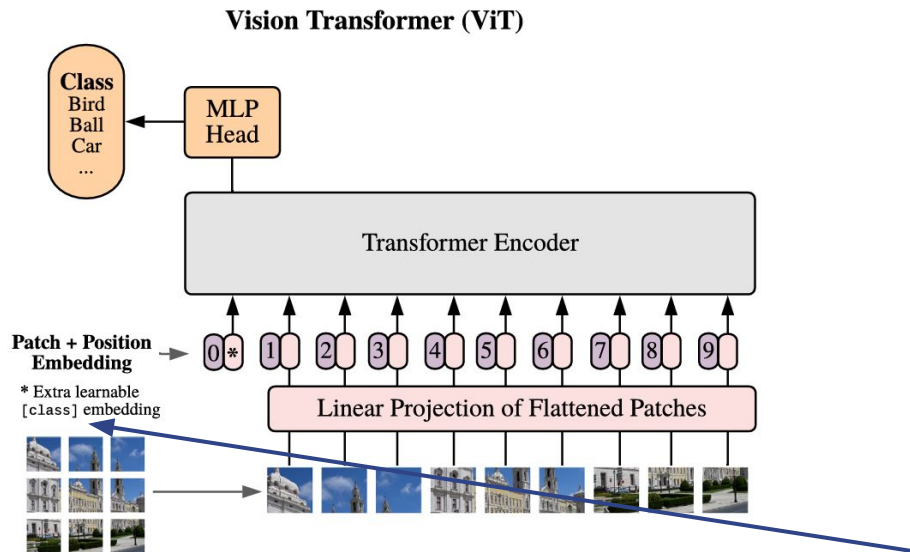
Step 2.3

Add **[CLS]** Token

Problem : For a classification task, which vector should we use to make the final prediction.

Solution : One final representation

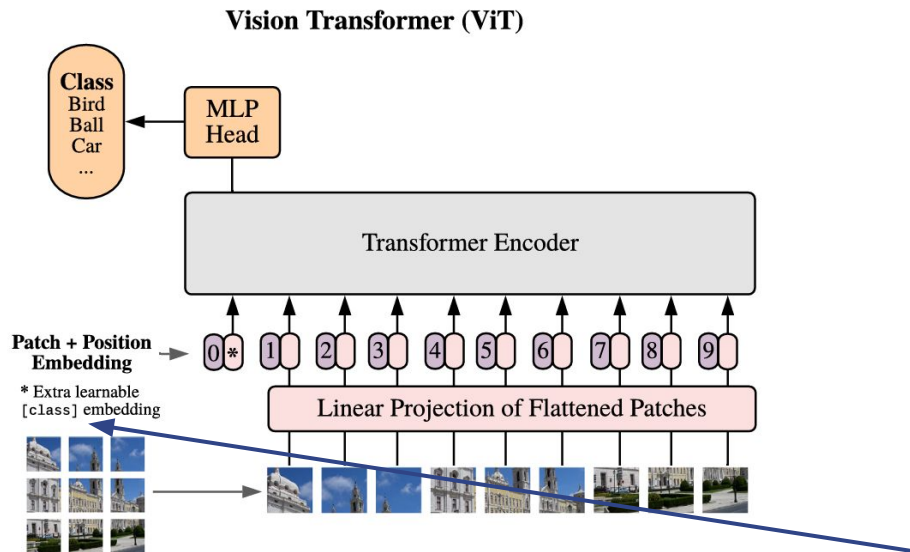
We add a special vector called **[CLS]** at the beginning.



Step 2.3

Add [CLS] Token

- We add a special vector called [CLS] at the beginning.
- At the beginning, [CLS] is just a random vector, but during attention, it interacts with all patches



Step 2.3

Add [CLS] Token

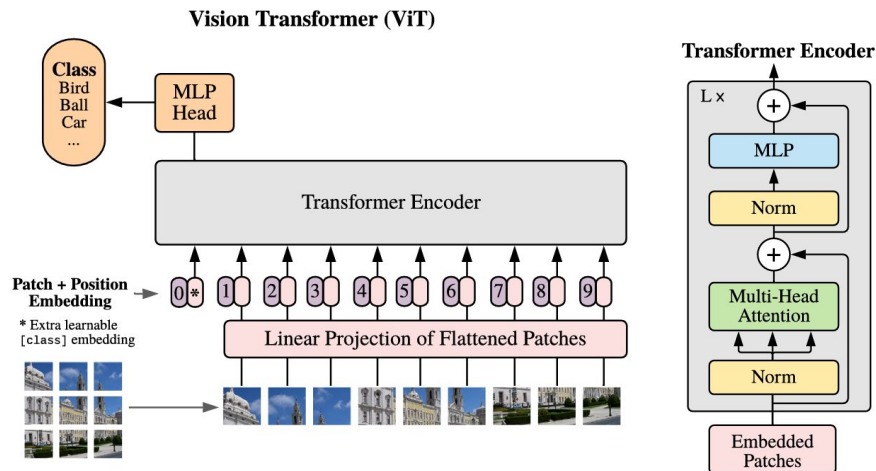
- We add a special vector called [CLS] at the beginning.
- At the beginning, [CLS] is just a random vector, but during attention, it interacts with all patches
- The [CLS] token becomes a **summary** of the whole image (close to CNN pooling)

Same transformer encoder as in NLP

Step 3

Transformer Encoder

What we have : $[CLS], x_1, x_2, \dots, x_n$



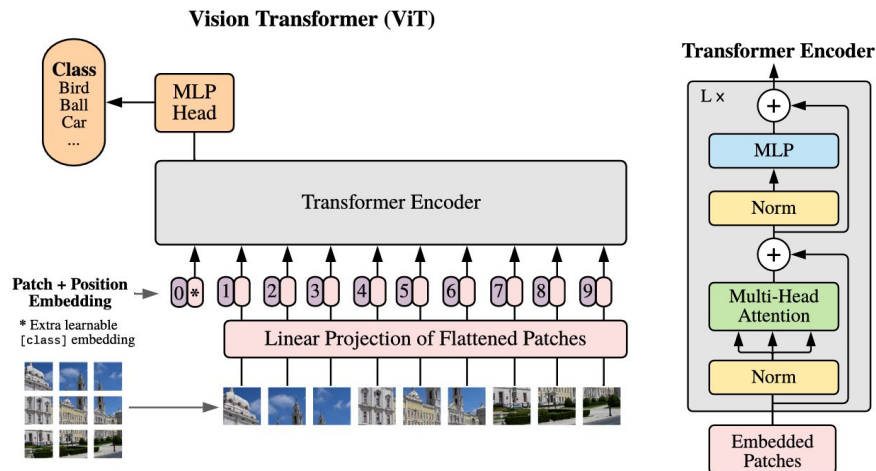
Same transformer encoder as in NLP

Step 3

Transformer Encoder

What we have : $[CLS], x_1, x_2, \dots, x_n$

We pass it through multiple **encoder layers**



Same transformer encoder as in NLP

Step 3

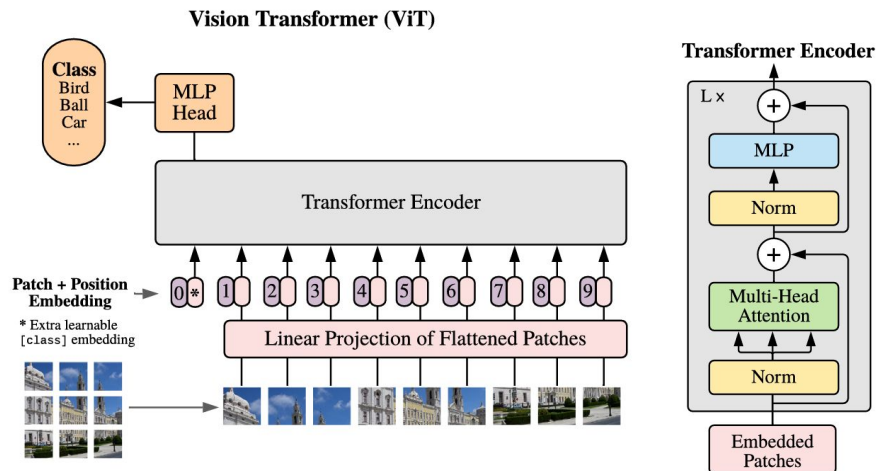
Transformer Encoder

What we have : $[CLS], x_1, x_2, \dots, x_n$

We pass it through multiple **encoder layers**

Each layer applies:

- Self-attention
- Feed-forward



Same transformer encoder as in NLP

Step 3

Transformer Encoder

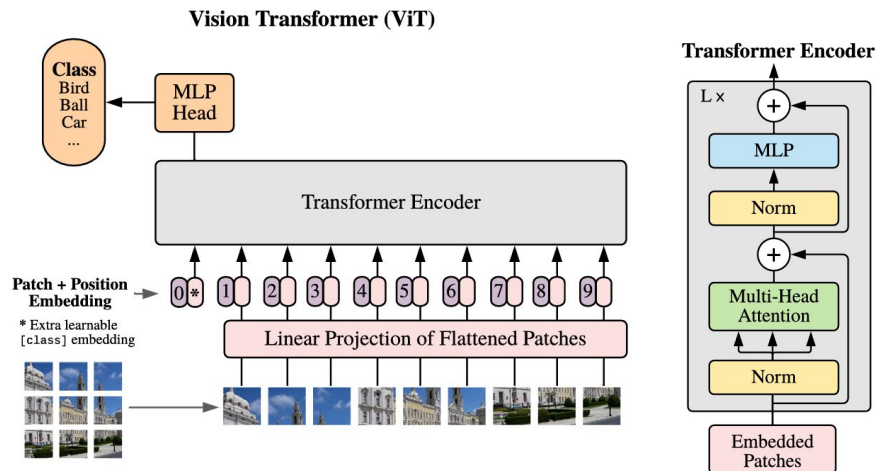
What we have : $[CLS], x_1, x_2, \dots, x_n$

We pass it through multiple **encoder layers**

Each layer applies:

- Self-attention
- Feed-forward

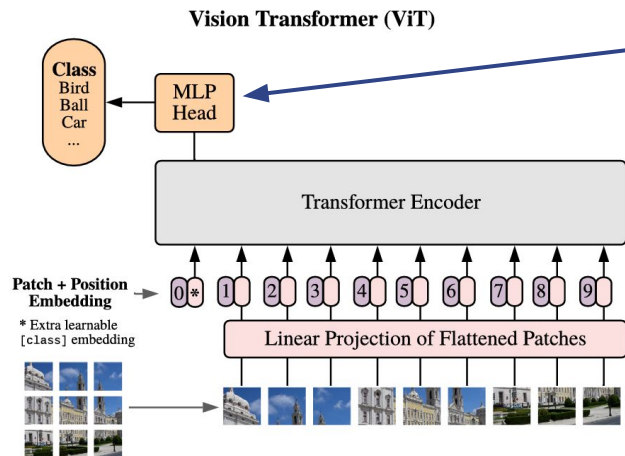
The **CLS** gradually becomes a global representation of the image



Step 4

Image Classification

What we have : a global representation :
[CLS]

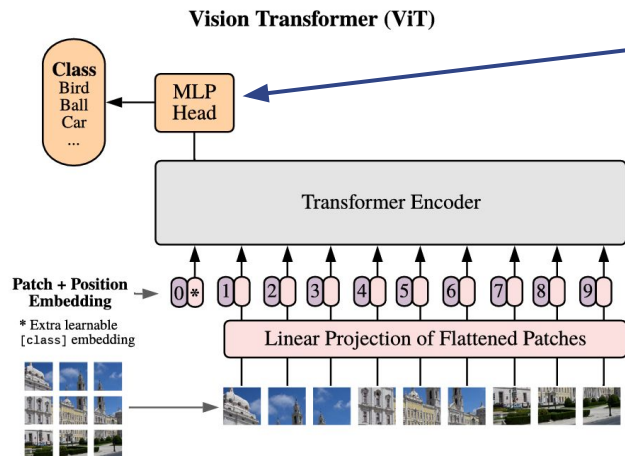


Step 4

Image Classification

What we have : a global representation :
[CLS]

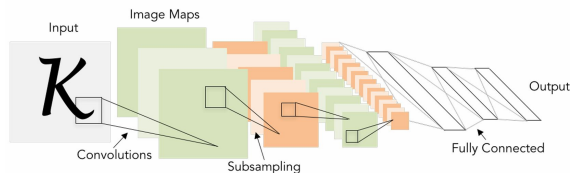
- Pass it through a small **multi-layer perceptron** (MLP)
- Output class probabilities





[Source: [Gif Link](#)]

- Local receptive field
- Processes entire images directly
- Works well with small data



- Attention
- Divides images into patches and processes them as sequences
- Needs large data

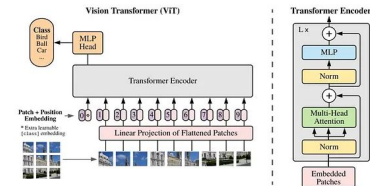


Figure 1: Model overview. We split an image into fixed-size patches, linearly embed each of them, add position embeddings, and feed the resulting sequence of vectors to a standard Transformer encoder. In order to perform classification, we use the standard approach of adding an extra learnable "classification token" to the sequence. The illustration of the Transformer encoder was inspired by Vaswani et al. (2017).

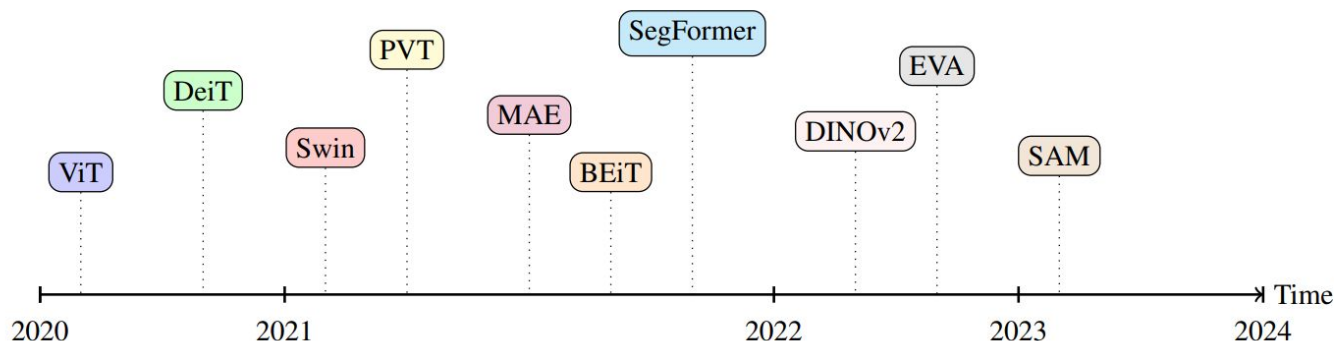


Figure 1. Evolution timeline of Vision Transformer architectures showing major milestones from 2020 to 2024. Each architecture introduced key innovations that advanced the field.

[Source: Laurier, Linda. (2025). Vision Transformers: From Patch-Based Processing to Visual Intelligence.]

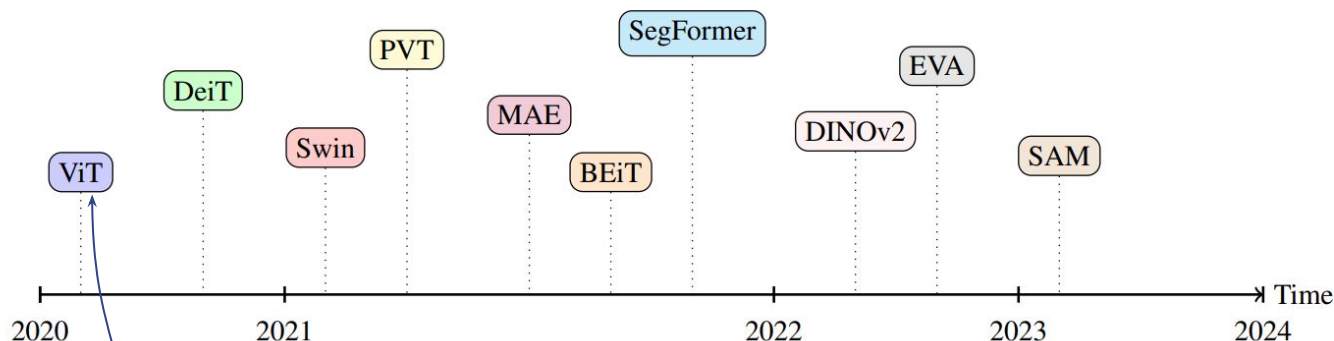


Figure 1. Evolution timeline of Vision Transformer architectures showing major milestones from 2020 to 2024. Each architecture introduced key innovations that advanced the field.

Need huge datasets

[Source: Laurier, Linda. (2025). Vision Transformers: From Patch-Based Processing to Visual Intelligence.]

© Mehyaar Mlaweh, 2026. All rights reserved.

How to train transformers
with less data using
distillation

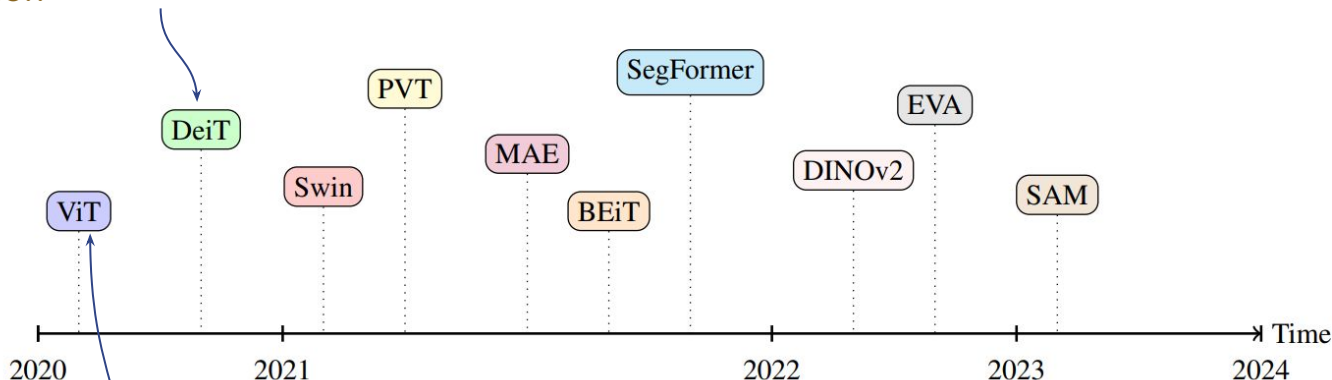


Figure 1. Evolution timeline of Vision Transformer architectures showing major milestones from 2020 to 2024. Each architecture introduced key innovations that advanced the field.

Need huge datasets

[Source: Laurier, Linda. (2025). Vision Transformers: From Patch-Based Processing to Visual Intelligence.]

© Mehyaar Mlaweh, 2026. All rights reserved.

How to train transformers
with less data using
distillation

Hierarchical Transformers

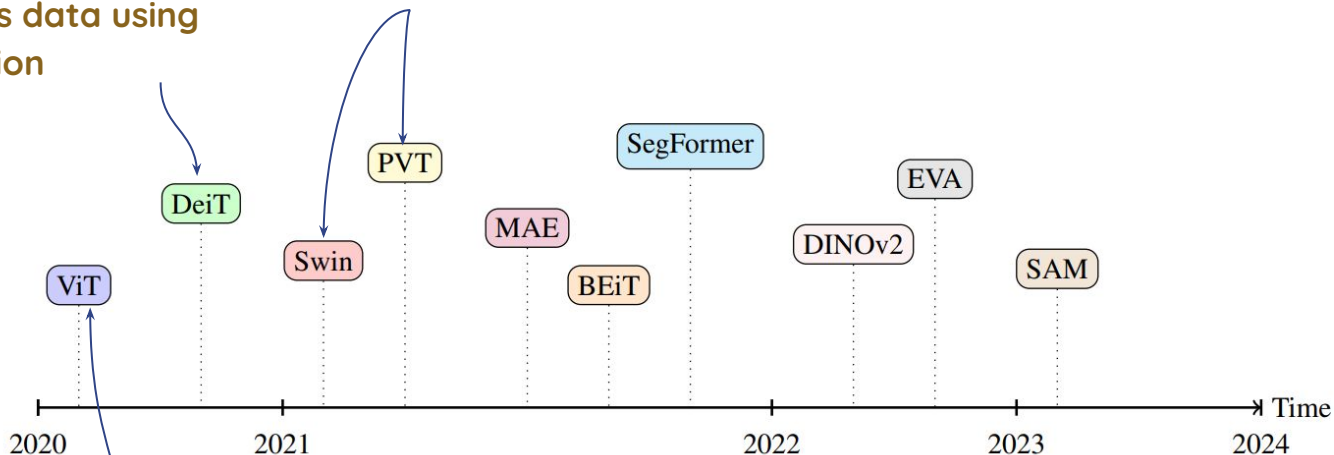


Figure 1. Evolution timeline of Vision Transformer architectures showing major milestones from 2020 to 2024. Each architecture introduced key innovations that advanced the field.

Need huge datasets

[Source: Laurier, Linda. (2025). Vision Transformers: From Patch-Based Processing to Visual Intelligence.]

© Mehیار Mlaweh, 2026. All rights reserved.

How to train transformers
with less data using
distillation

Hierarchical Transformers

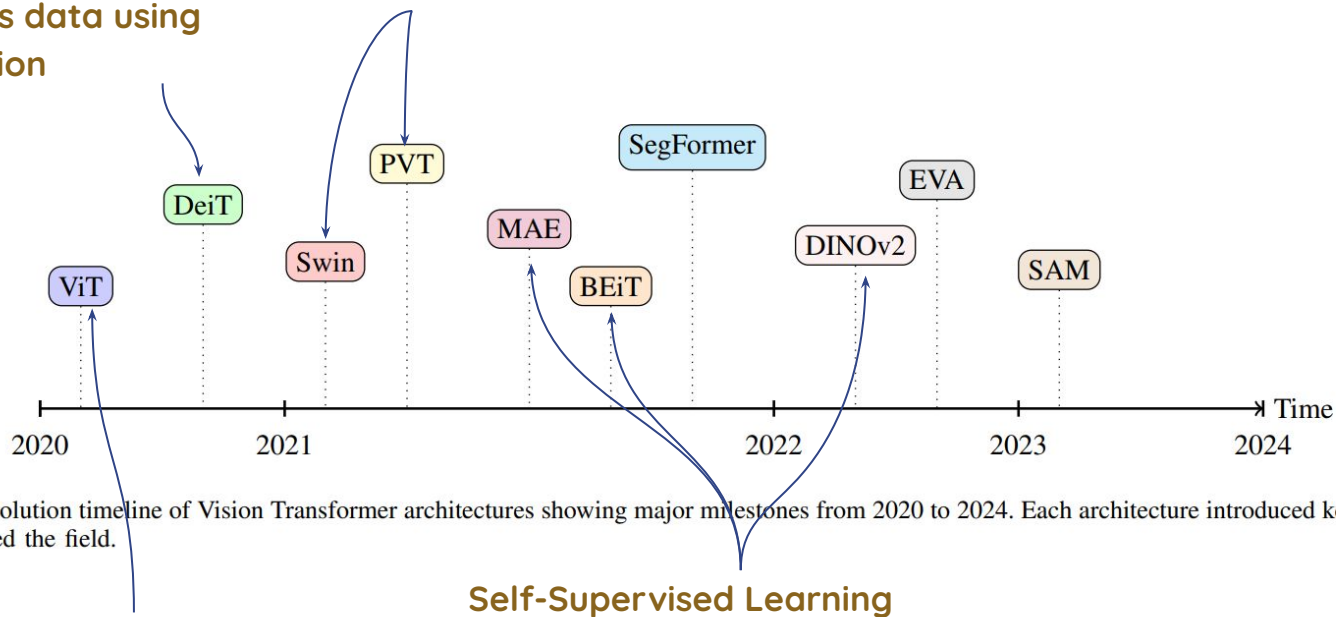


Figure 1. Evolution timeline of Vision Transformer architectures showing major milestones from 2020 to 2024. Each architecture introduced key innovations that advanced the field.

Need huge datasets

Self-Supervised Learning

[Source: Laurier, Linda. (2025). Vision Transformers: From Patch-Based Processing to Visual Intelligence.]

© Mehیار Mlaweh, 2026. All rights reserved.

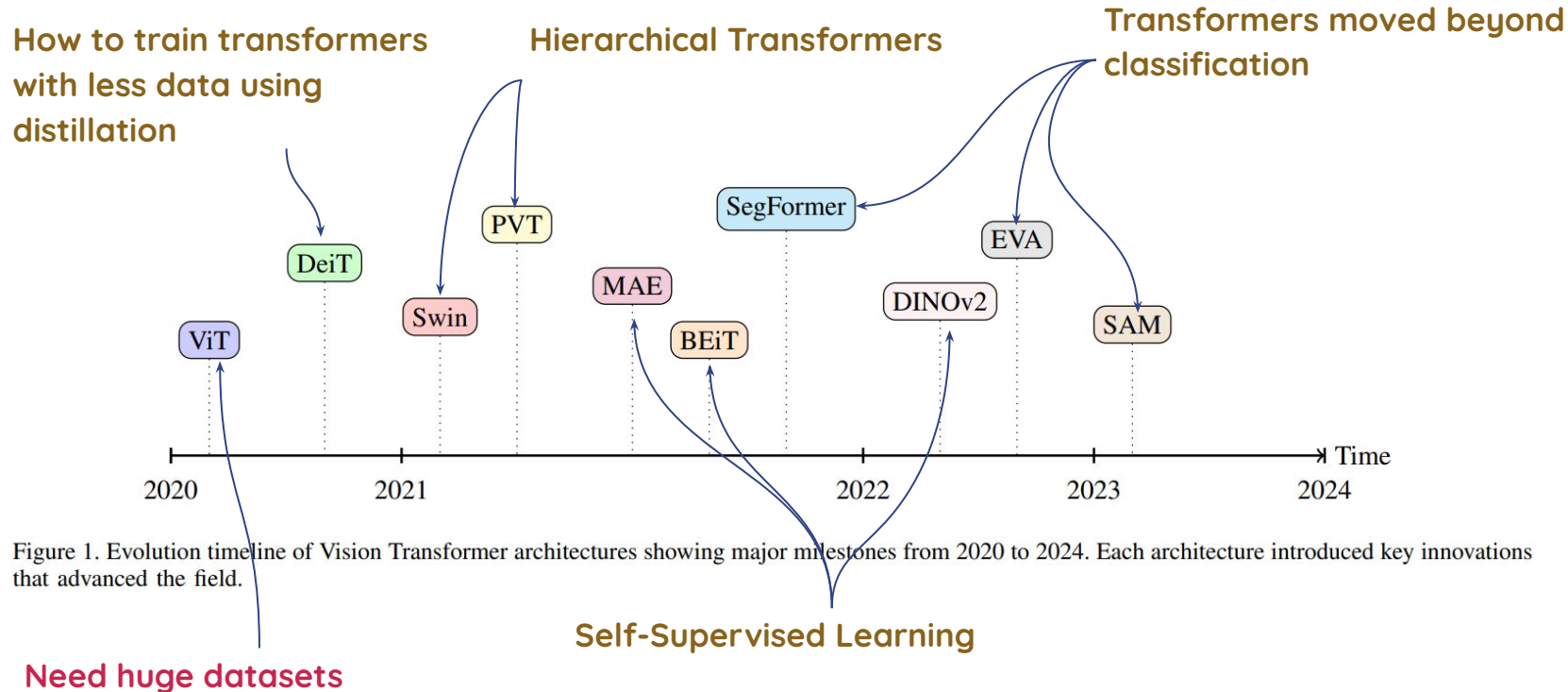


Figure 1. Evolution timeline of Vision Transformer architectures showing major milestones from 2020 to 2024. Each architecture introduced key innovations that advanced the field.

[Source: Laurier, Linda. (2025). Vision Transformers: From Patch-Based Processing to Visual Intelligence.]

© Mehyaar Mlaweh, 2026. All rights reserved.

Object Detection

Books

- Computer Vision: Models, Learning, and Inference (2012). <https://amzn.to/2rxrdOF>
- Programming Computer Vision with Python (2012). <https://amzn.to/2QKTaAL>
- Multiple View Geometry in Computer Vision (2004). <https://amzn.to/2LfHLE8>
- Computer Vision: Algorithms and Applications (2010). <https://amzn.to/2Lclt4J>
- Computer Vision: A Modern Approach (2002). <https://amzn.to/2rv7AqJ>
- Deep Learning for Computer Vision : Image Classification, Object Detection and Face Recognition in Python(2019). <https://machinelearningmastery.com/deep-learning-for-computer-vision/>
- Hands-On Machine Learning with Scikit-Learn, Keras & TensorFlow — Aurélien Géron (2019)
<https://amzn.to/3uZbN6Z>
- Hands-On Large Language Models - Jay Alammar & Grootendorst

Articles & Online Resources

- Krizhevsky, Sutskever, Hinton (2012) — ImageNet Classification with Deep CNNs (AlexNet — deep learning breakthrough)
- Improved Texture Networks: Maximizing Quality and Diversity in Feed-forward Stylization and Texture Synthesis, 2017
- “Computer vision,” Wikipedia. https://en.wikipedia.org/wiki/Computer_vision
- “Machine vision,” Wikipedia. https://en.wikipedia.org/wiki/Machine_vision
- “Digital image processing,” Wikipedia. https://en.wikipedia.org/wiki/Digital_image_processing

- “Computer vision,” Wikipedia. https://en.wikipedia.org/wiki/Computer_vision
- <https://medium.com/data-reply-it-datatech>
- <https://huggingface.co/blog/RDTvlokup/quand-l-ia-apprend-de-l-experience-comme-toi>
- Zhuang et al., “A Comprehensive Survey on Transfer Learning,” IEEE, 2020.
- <https://jalammar.github.io/illustrated-transformer/>
- <https://sh-tsang.medium.com/review-vision-transformer-vit-406568603de0>
- Laurier, Linda. (2025). Vision Transformers: From Patch-Based Processing to Visual Intelligence.
- “Training Convolutional Neural Networks: What Is Machine Learning” Analog Devices.
<https://www.analog.com/en/resources/analog-dialogue/articles/training-convolutional-neural-networks-what-is-machine-learning-part-2.html>