

Machine Learning (part II)

Transformers

Angelo Ciaramella

Introduction

■ Transformers

- transform a set of **vectors in some representation space** into a corresponding set of vectors, having the same dimensionality, in some new space
- the fundamental concept that underpins a transformer is **attention**

■ Attention

- developed as an enhancement to RNNs for machine translation
- A transformer can be viewed as a **richer form of embedding** in which a given **vector** is mapped to a location that depends on the other vectors in the sequence



Attention

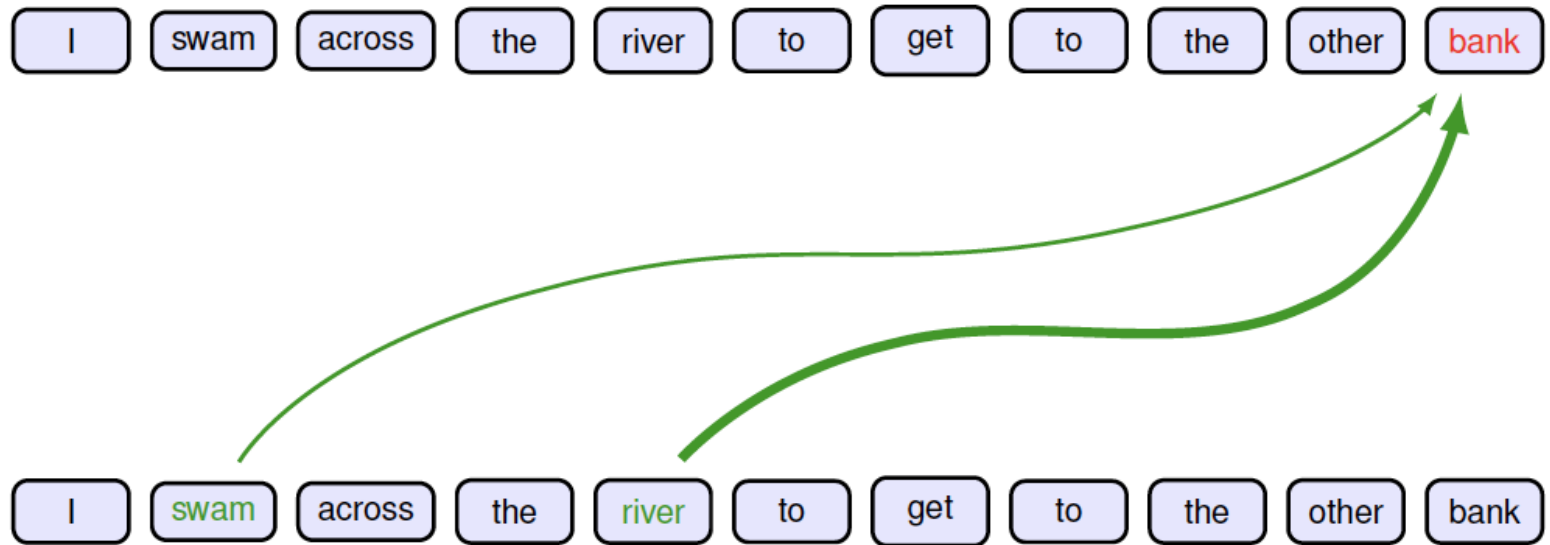
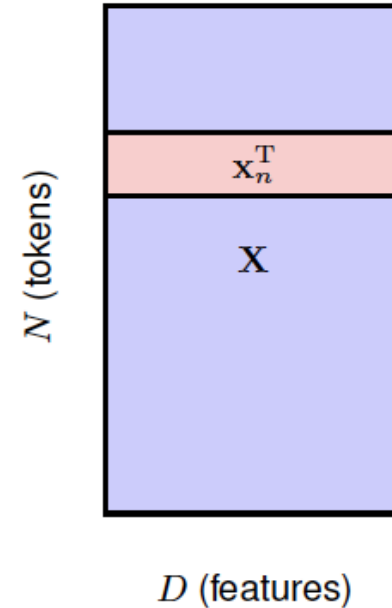


Figure 12.1 Schematic illustration of attention in which the interpretation of the word 'bank' is influenced by the words 'river' and 'swam', with the thickness of each line being indicative of the strength of its influence.



Introduction

Figure 12.3 The structure of the data matrix $\mathbf{X}$, of dimension $N \times D$, in which row n represents the transposed data vector $\mathbf{x}_n^T$.



$$\tilde{\mathbf{X}} = \text{TransformerLayer}[\mathbf{X}]$$

function that takes a data matrix as input and creates a transformed matrix of the same dimensionality as the output



Deep networks

- Multiple transformer layers
 - to construct deep networks capable of learning rich internal representations
 - each transformer layer contains its own weights and biases, which can be learned using gradient descent using appropriate cost function



Attention coefficients

attention weights (non-negative)

$$y_n = \sum_{m=1}^N a_{nm} \mathbf{x}_m$$


$$a_{nm} \geq 0$$

$$\sum_{m=1}^N a_{nm} = 1.$$

The coefficients should be close to zero for input tokens that have little influence on the output y_n and largest for inputs that have most influence



Self-Attention

- choosing which movie to watch in an online movie streaming service
 - associate each movie with a list of attributes
 - attributes of each movie in a vector called the key
 - the corresponding movie file itself is called a value
 - personal vector of values for the desired attributes called query
- movie service could then compare the query vector with all the key vectors to find the best match and send the corresponding movie to the user in the form of the value file (hard attention)



Soft attention

- we use continuous variables to measure the degree of match between queries and keys
- variables to weight the influence of the value vectors on the outputs

$$a_{nm} = \frac{\exp(\mathbf{x}_n^T \mathbf{x}_m)}{\sum_{m'=1}^N \exp(\mathbf{x}_n^T \mathbf{x}_{m'})}$$

dot product

softmax

normalization

- if all the input vectors are orthogonal

$$\mathbf{y}_m = \mathbf{x}_m \text{ for } m = 1, \dots, N$$



Self attention

- Matrix notation

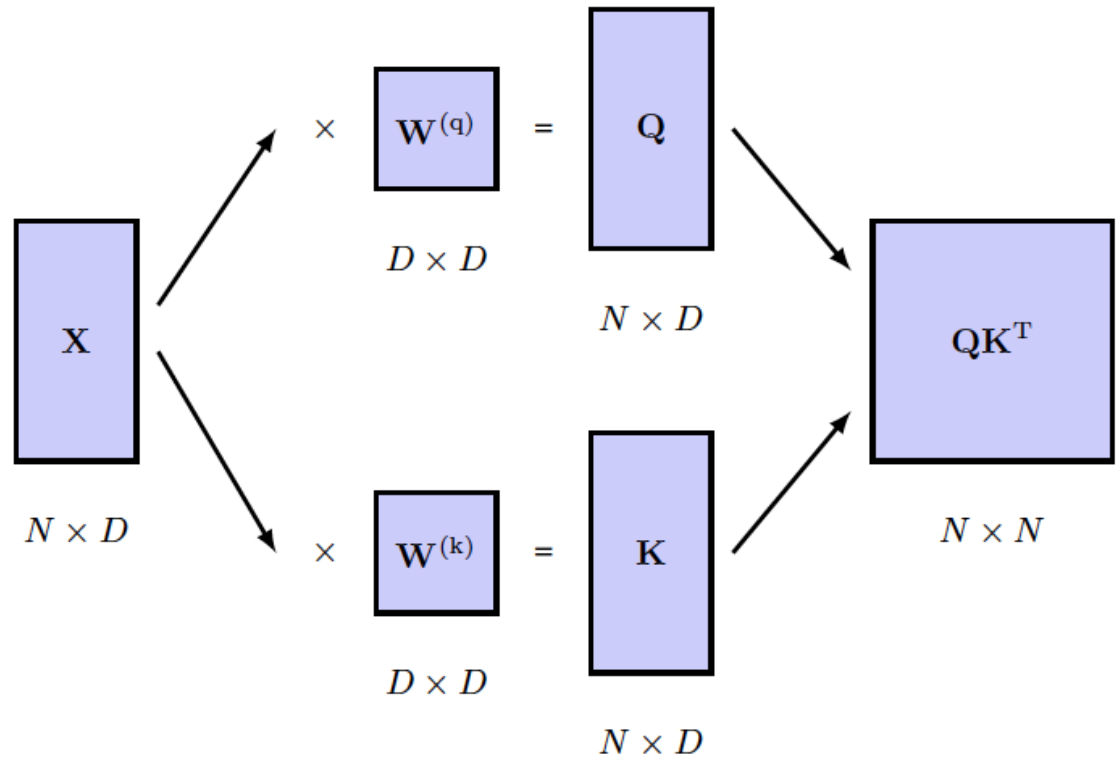
dot product self-attention

$$\mathbf{Y} = \text{Softmax} [\mathbf{X}\mathbf{X}^T] \mathbf{X}$$

- we are using the same sequence to determine the queries, keys, and values



Self attention



dot product self-attention

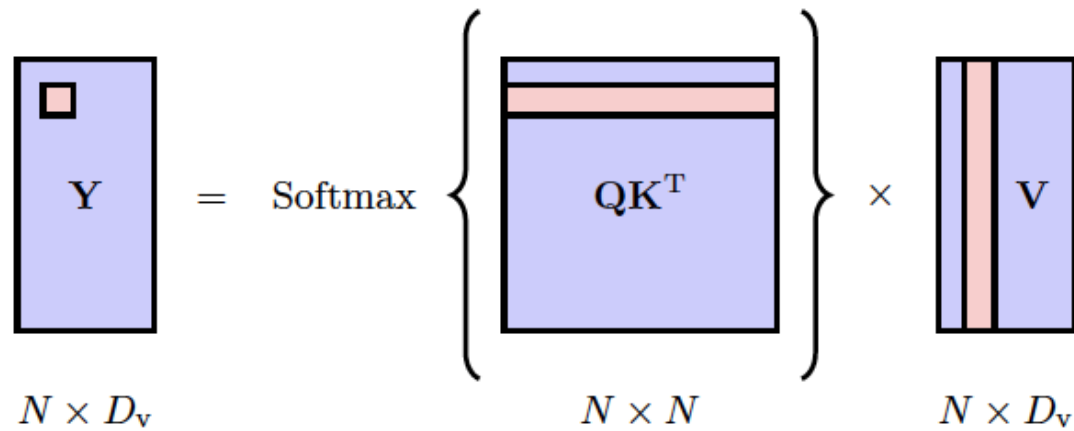
$$Q = XW^{(q)}$$

$$K = XW^{(k)}$$

$$V = XW^{(v)}$$



Self attention



The diagram illustrates the self-attention mechanism. On the left, a light blue rectangle labeled Y represents the output, with dimensions $N \times D_v$ indicated below it. A small red square in the top-left corner of Y highlights a specific element. This is followed by an equals sign and the word "Softmax". To the right of "Softmax" is a large curly brace containing a light blue square labeled QK^T , with dimensions $N \times N$ indicated below it. A red horizontal bar is at the top of the QK^T matrix. To the right of the brace is a multiplication symbol $\times$, followed by a light blue rectangle labeled V with dimensions $N \times D_v$ indicated below it. A red vertical bar is on the left side of the V matrix.

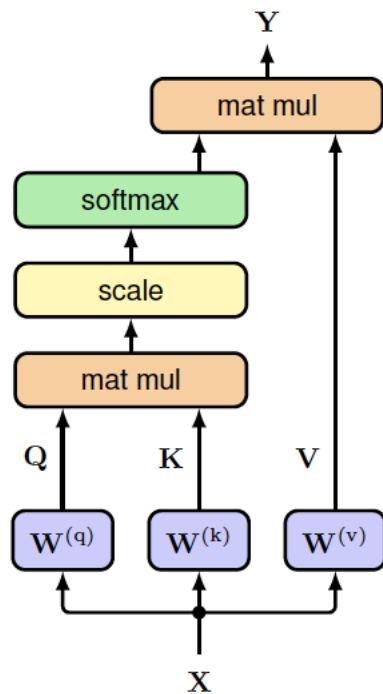
$$Y_{N \times D_v} = \text{Softmax} \left(QK^T_{N \times N} \right) \times V_{N \times D_v}$$

output from an attention layer



Scaled self-attention

- the gradients of the softmax function become **exponentially small** for inputs of high magnitude
- variance of the dot product would be D_K



Scaled dot-product self-attention

$$Y = \text{Attention}(Q, K, V) \equiv \text{Softmax} \left[\frac{QK^T}{\sqrt{D_k}} \right] V$$

Multi-head attention

- in natural language some patterns might be relevant to tense whereas others might be associated with vocabulary

$$\mathbf{H}_h = \text{Attention}(\mathbf{Q}_h, \mathbf{K}_h, \mathbf{V}_h)$$

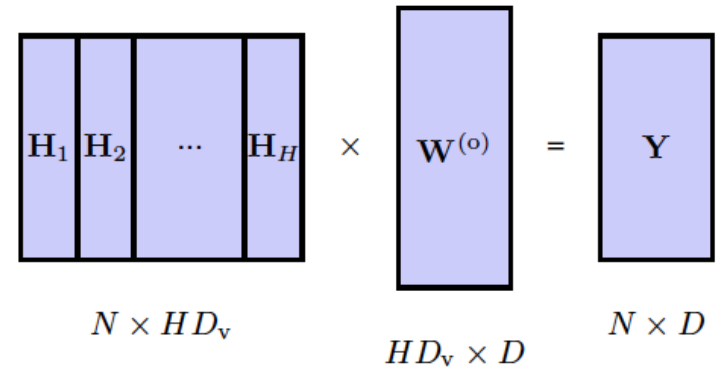
H heads

$$\mathbf{Q}_h = \mathbf{X} \mathbf{W}_h^{(q)}$$

$$\mathbf{K}_h = \mathbf{X} \mathbf{W}_h^{(k)}$$

$$\mathbf{V}_h = \mathbf{X} \mathbf{W}_h^{(v)}$$

$$\mathbf{Y}(\mathbf{X}) = \text{Concat} [\mathbf{H}_1, \dots, \mathbf{H}_H] \mathbf{W}^{(o)}$$



Multi-head attention

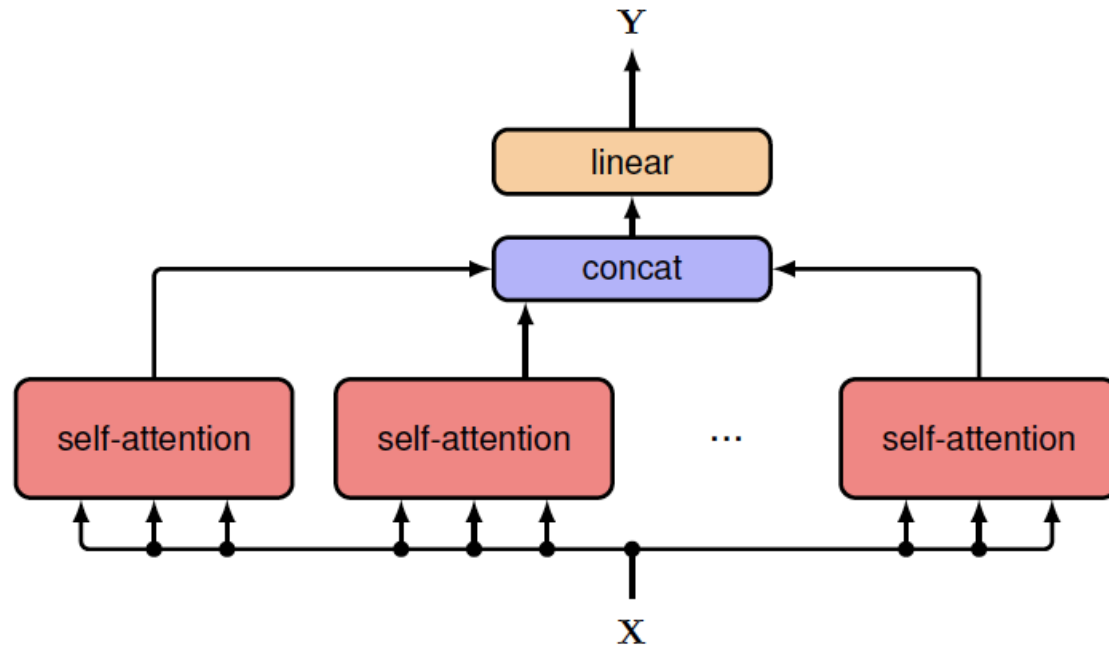


Figure 12.8 Information flow in a multi-head attention layer. The associated computation, given by Algorithm 12.2, is illustrated in [Figure 12.7](#).



Residual connections

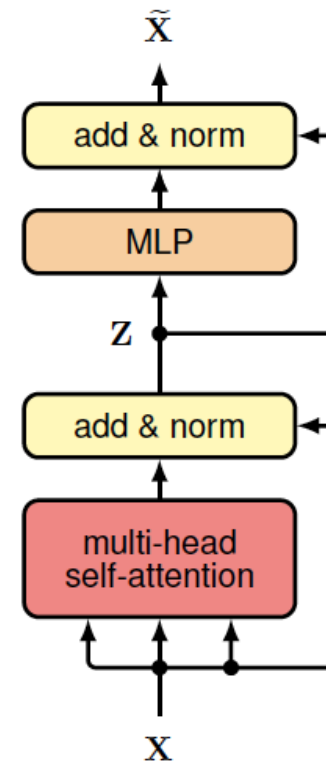
- Improving the learning

$$\mathbf{Z} = \text{LayerNorm} [\mathbf{Y}(\mathbf{X}) + \mathbf{X}]$$

H heads

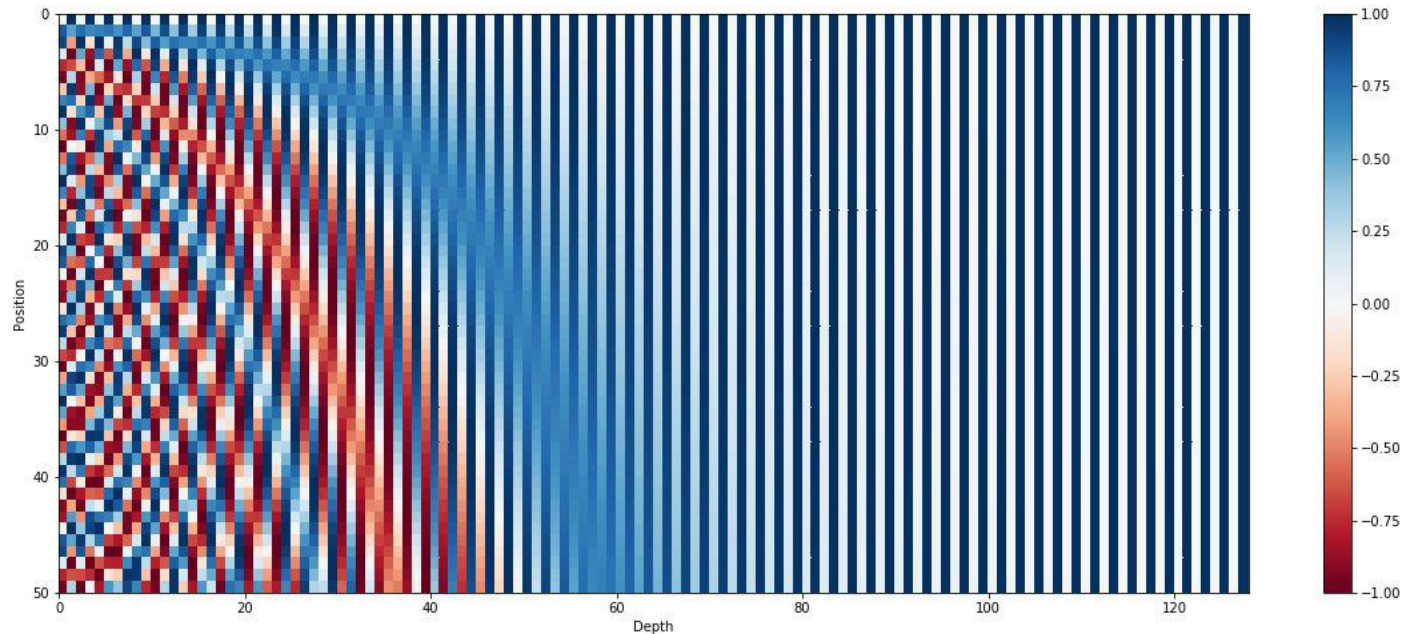
- Adding MLP

$$\tilde{\mathbf{X}} = \text{LayerNorm} [\text{MLP} [\mathbf{Z}] + \mathbf{Z}]$$



Positional embeddings

$$\text{PE}(i, \delta) = \begin{cases} \sin\left(\frac{i}{10000^{2\delta'/d}}\right) & \text{if } \delta = 2\delta' \\ \cos\left(\frac{i}{10000^{2\delta'/d}}\right) & \text{if } \delta = 2\delta' + 1 \end{cases}$$



Attention is all You Need

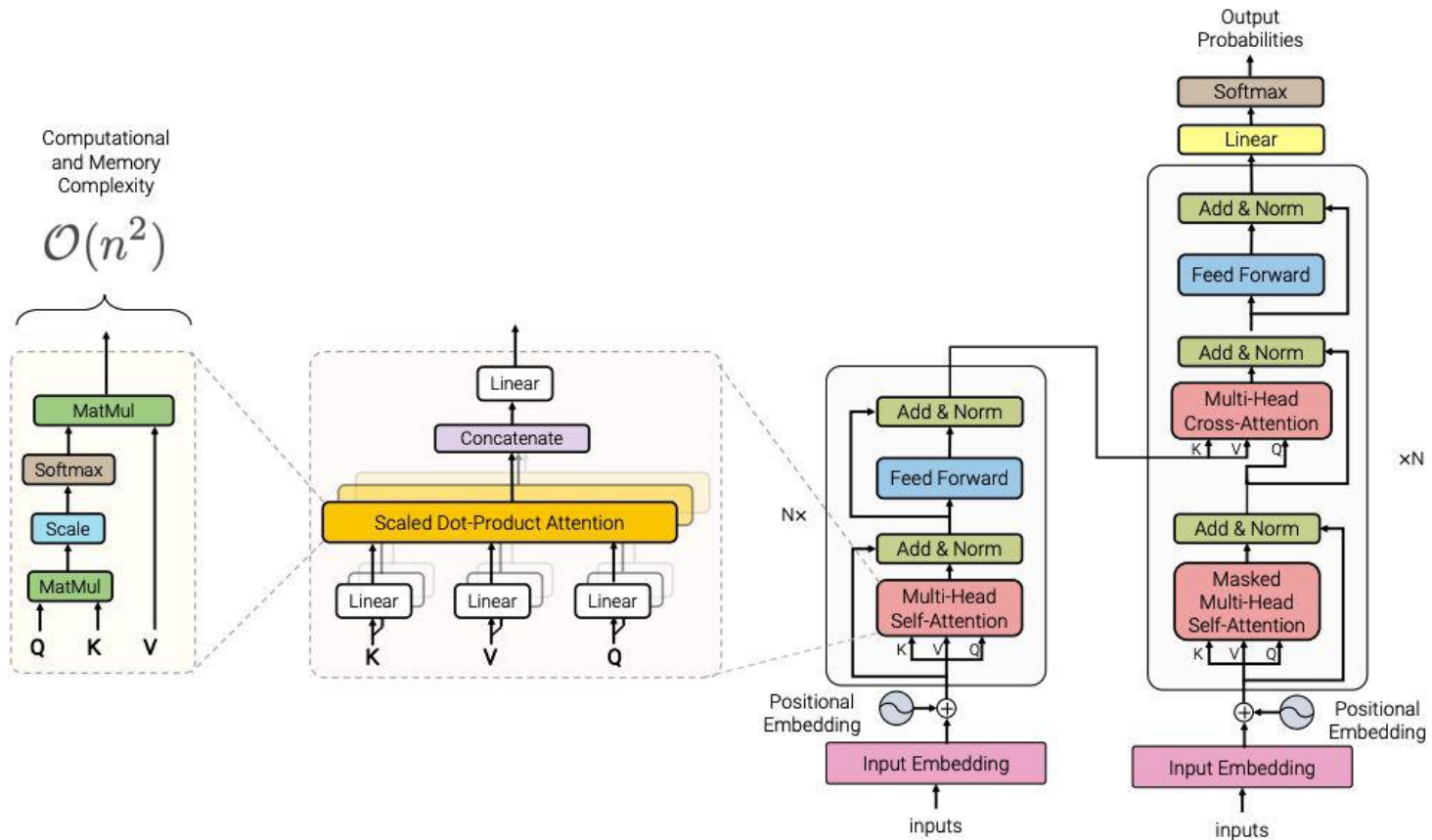
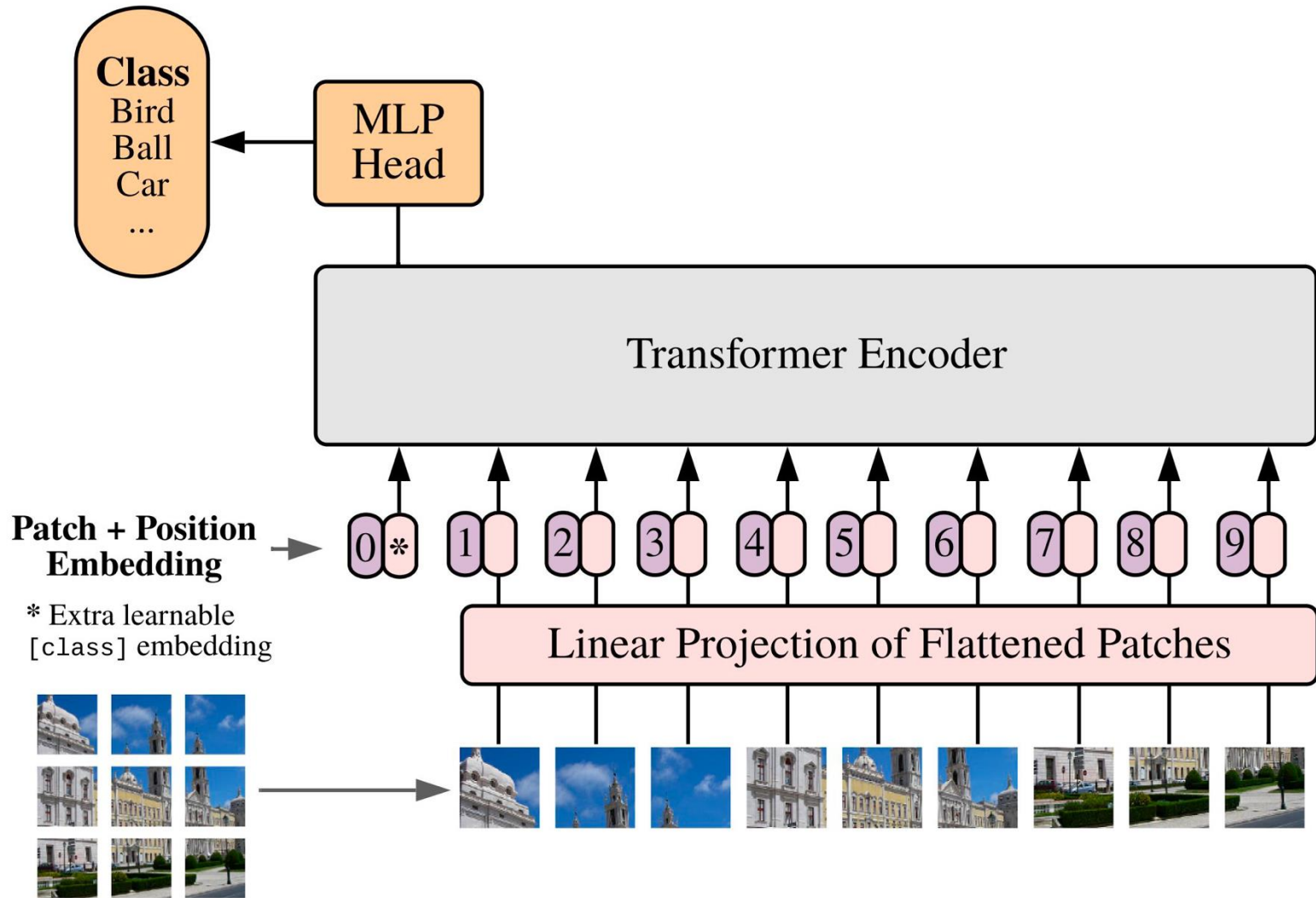


Figure 1: Architecture of the standard Transformer (Vaswani et al., 2017)

Vision Transformers



Vision Transformers



Image captioning



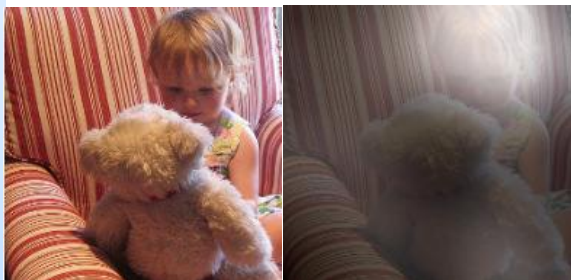
A woman is throwing a frisbee in a park.



A dog is standing on a hardwood floor.



A stop sign is on a road with a mountain in the background.



A little girl sitting on a bed with a teddy bear.

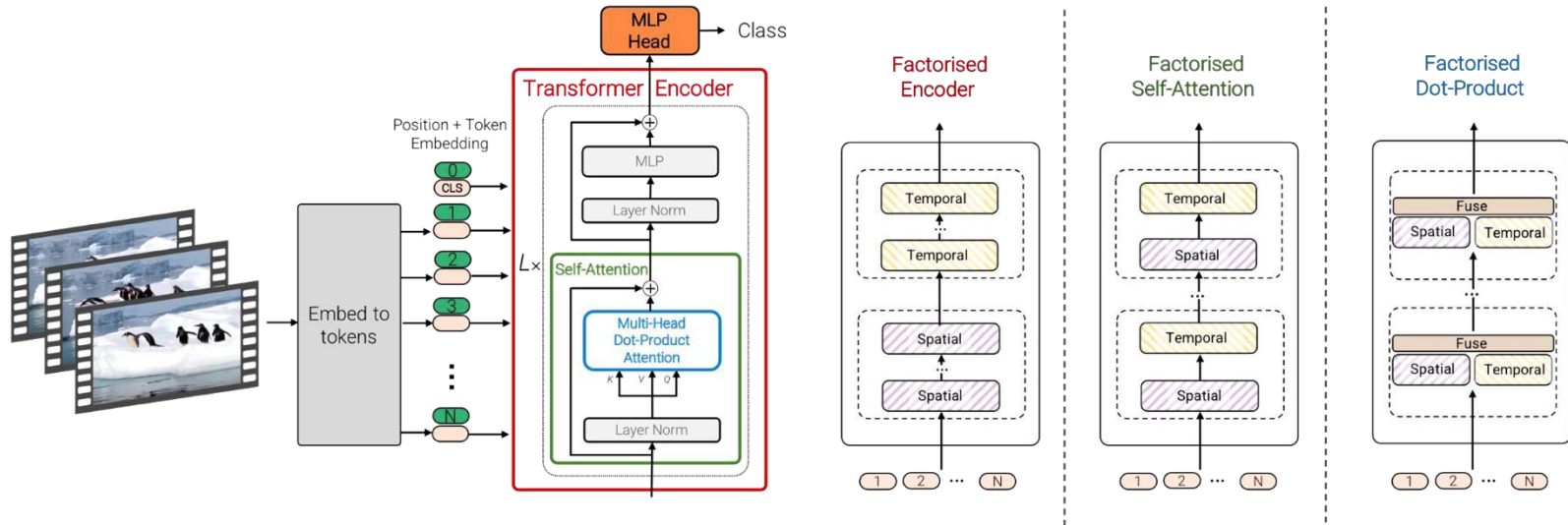


A group of people sitting on a boat in the water.

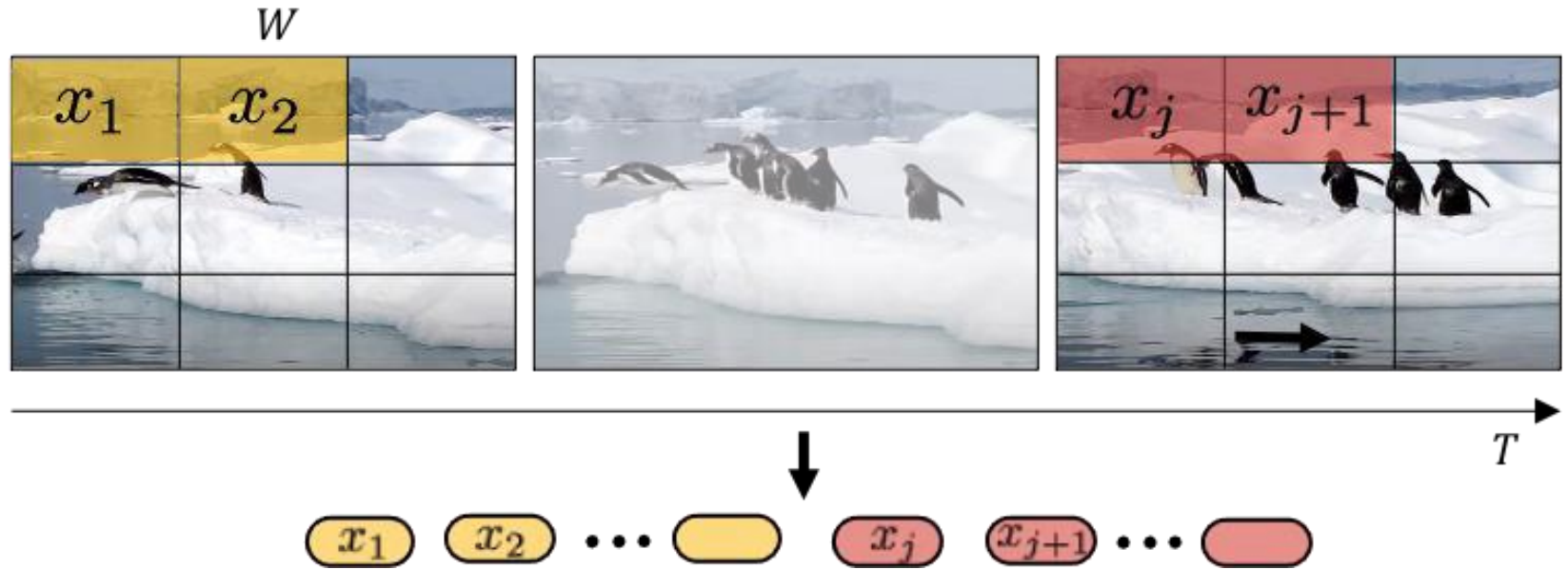


A giraffe standing in a forest with trees in the background.

ViT for video

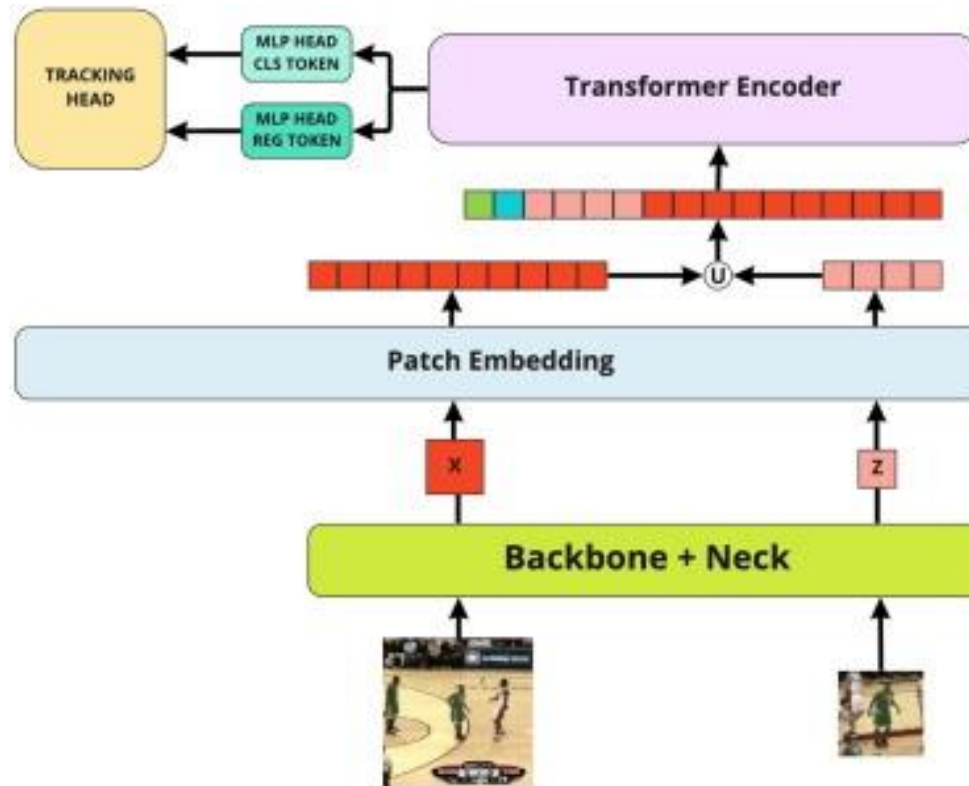


ViT for video



Sample frames, extract 2D patches and linearly project (as in ViT). Consider a video as a “big image”

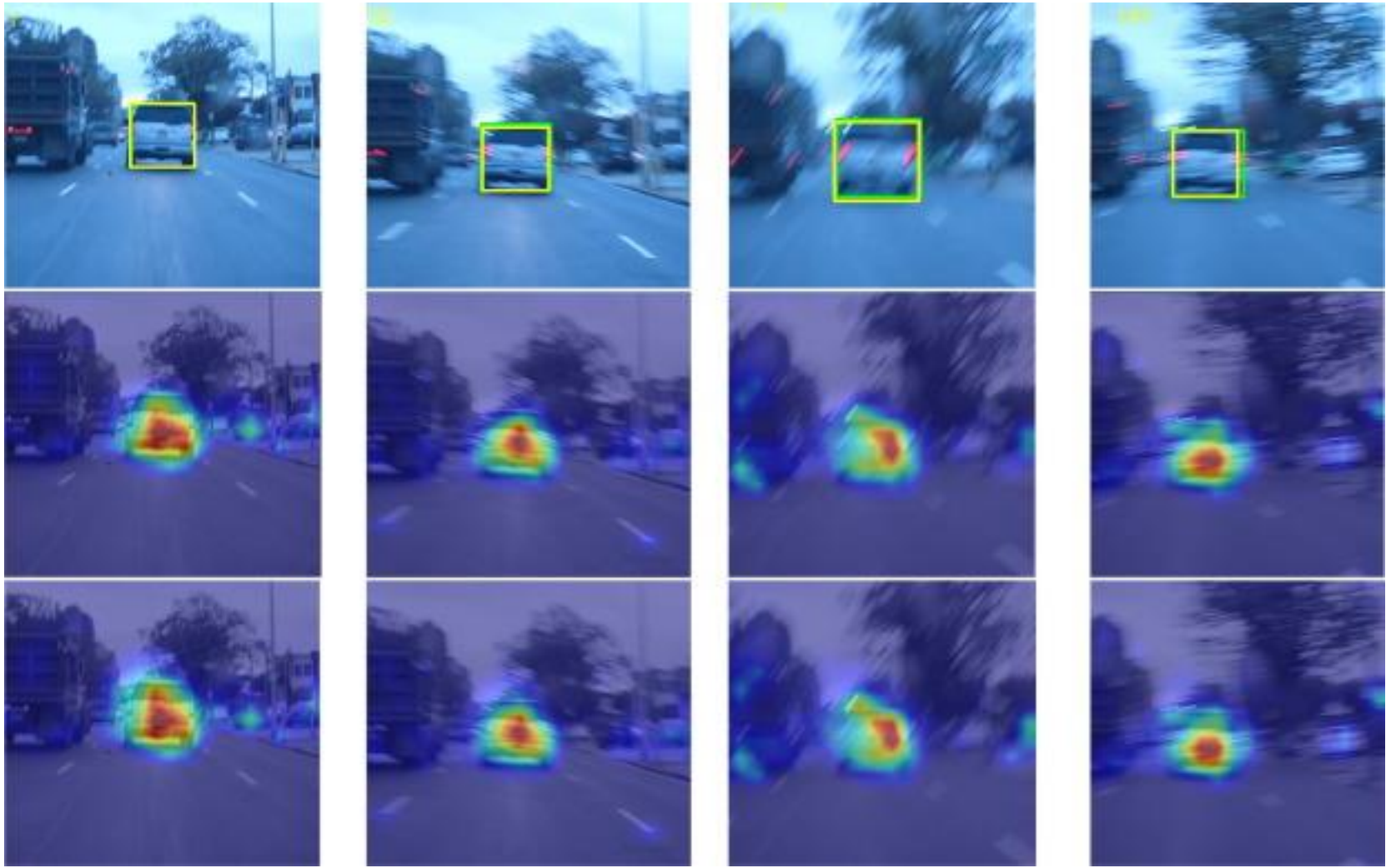
Visual tracking



Tracking vision transformer with class and regression tokens



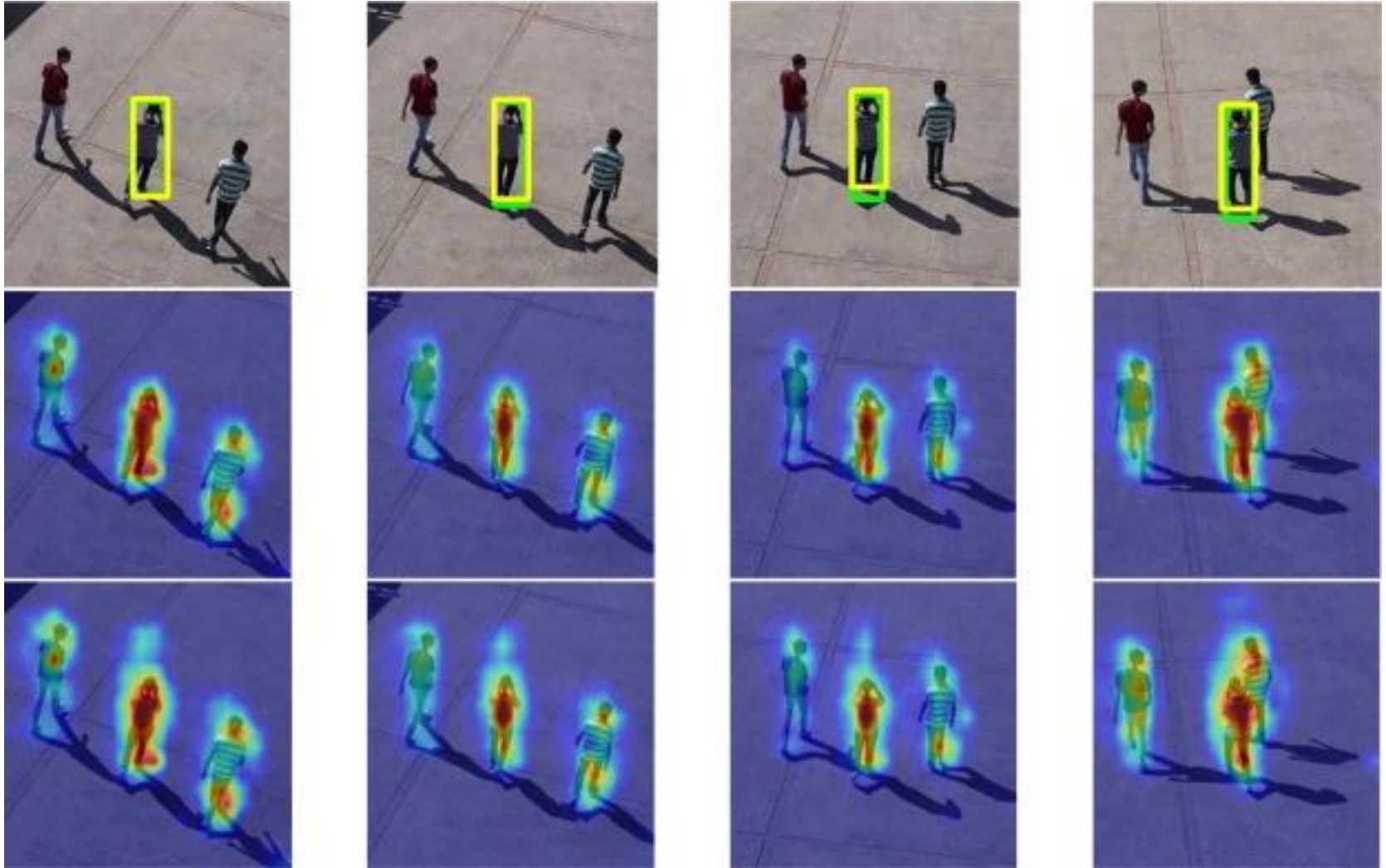
Visual tracking



Outputs for sequence *BlurCar1* in OTB100 with the attention of two tokens overlaying the image. *First row: Bounding Box outputs. Second row: Regression token attention. Third row: Classification token attention*



Visual tracking



Outputs for sequence *Group1_2* in UAV123 with the attention of two tokens overlaying the image. *First row: Bounding Box outputs. Second row: Regression token attention. Third row: Classification token attention.*

