

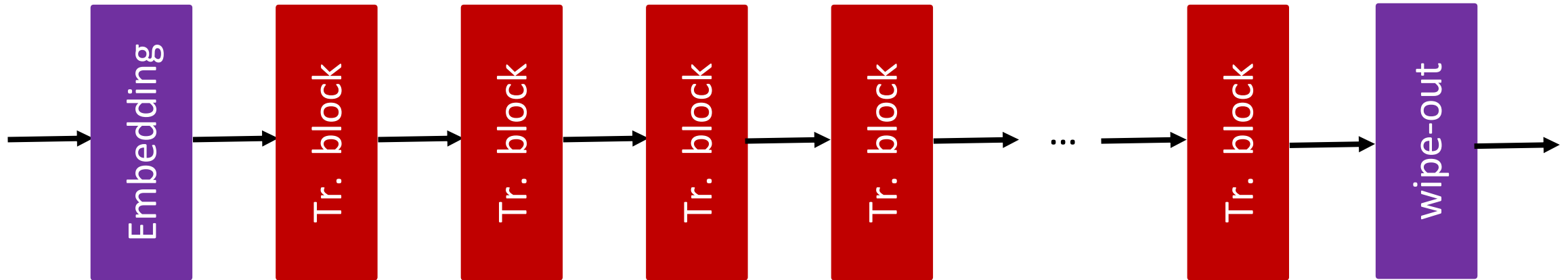
Detection Transformer and Foundation Models

Andrej Lúčny

Cross attention. Detection Transformer (DeTR).
Vision model with text prompt.
Foundation models: DINO, CLIP, Glee, SAM.

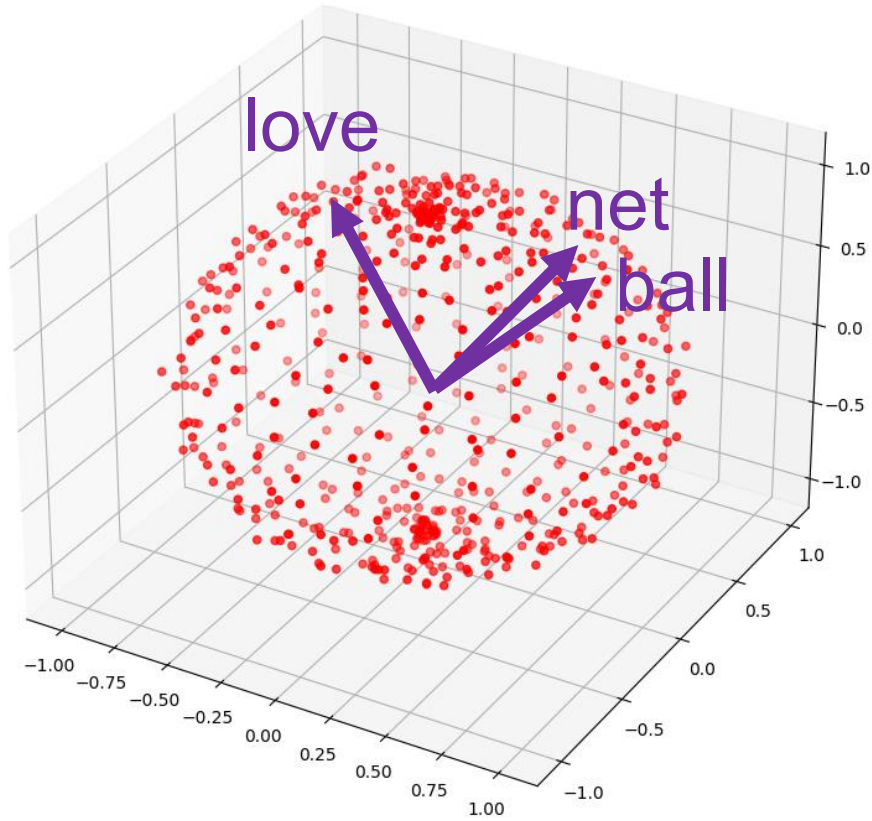
<https://github.com/andy1ucny/OAI>

Transformer



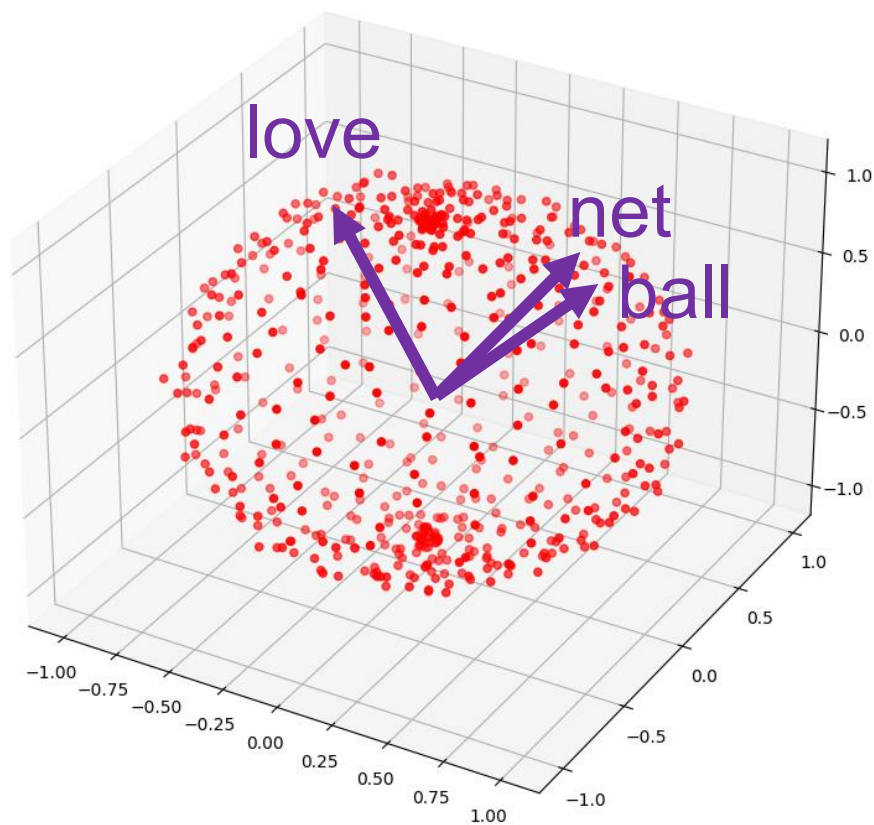
- It gradually transform N hidden states to other N hidden states, N is the size of the input.

Transformer latent spaces



- all have some chosen dimension that determines the number of ways in which two things can differ
- meanings are represented by vectors of this dimension, while their similarity is given by cosine similarity
- vectors (hidden states) also contain information about the position in the input

Cosine similarity



$$\text{sim}(\vec{u}, \vec{v}) = \frac{\vec{u} \cdot \vec{v}}{\|\vec{u}\| \|\vec{v}\|} \quad <-1, 1>$$



$$\cos \phi = 1$$



$$\cos \phi = 0$$



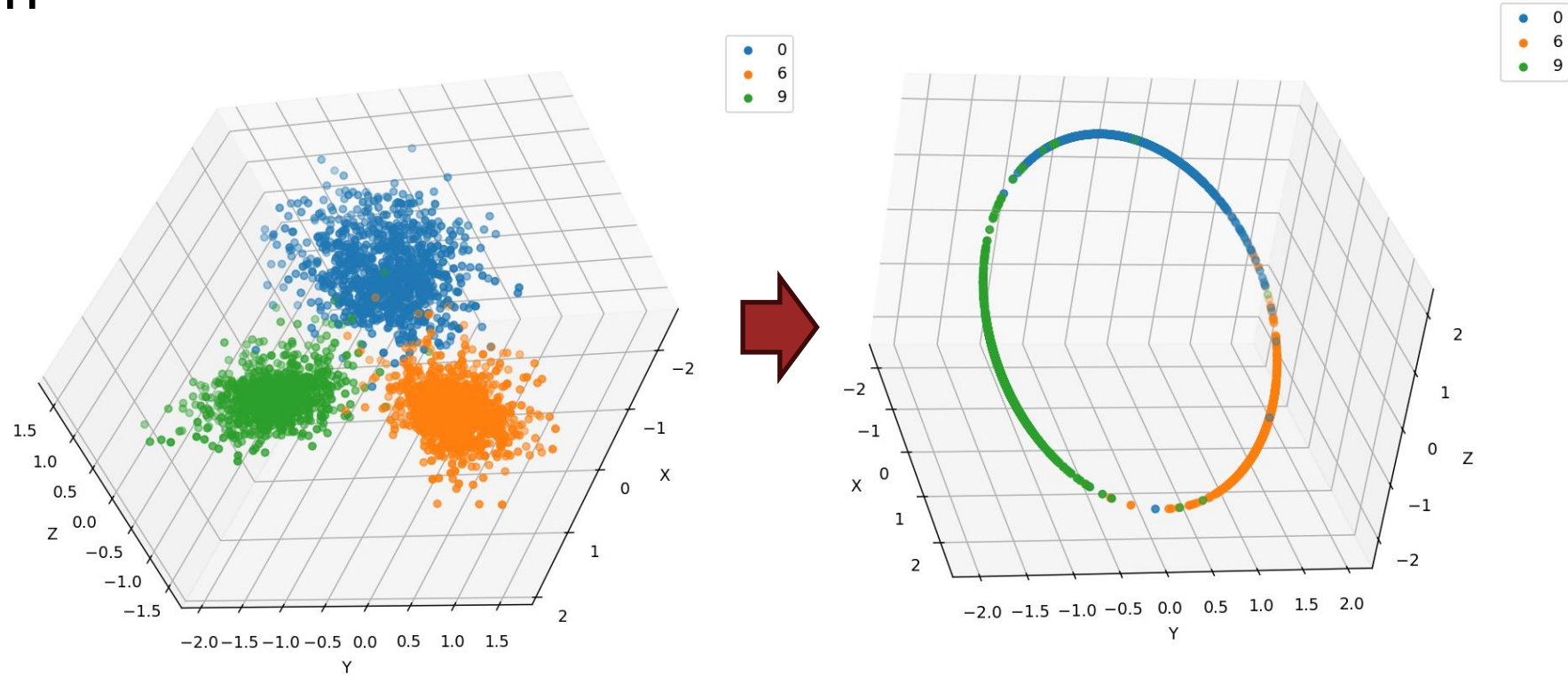
$$\cos \phi = -1$$

since in the transformer the vectors are of similar size:

$$\text{sim}(\vec{u}, \vec{v}) \approx \vec{u} \cdot \vec{v} \quad <-\infty, \infty>$$

- layer normalizations, inserted between every two building block, help establish cosine similarity between the data.

Layer Normalization



$$LN(h) = w \frac{h - \mu_L}{\sqrt{\sigma_L^2 + \epsilon}} + b$$

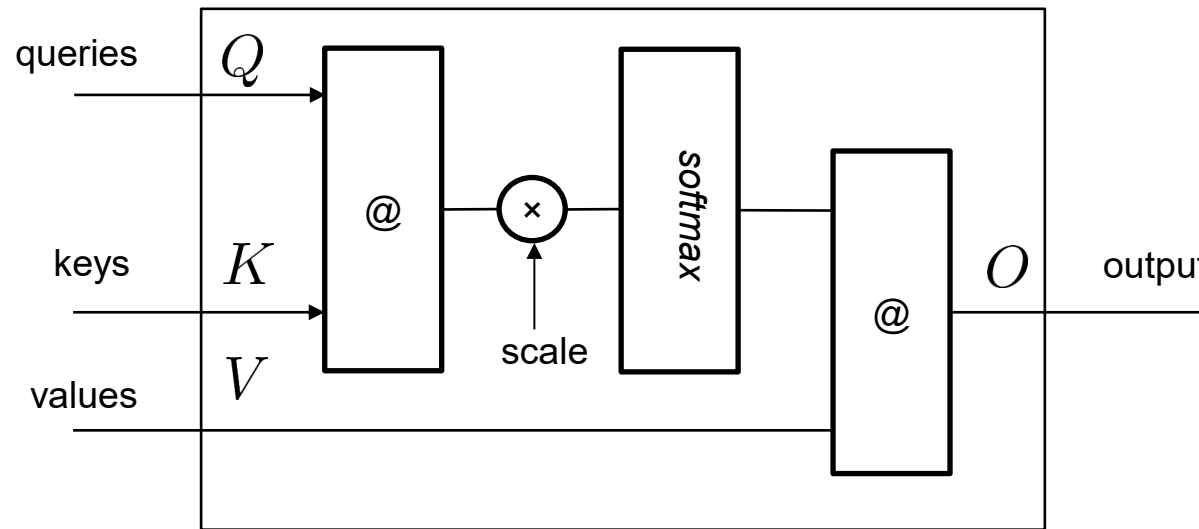
Attention

*„ try to mix the query
from the keys and
output the analogous
mix of values“*

*k – the number
of queries
(hidden states)*

*l – the number
of the key-
value pairs*

*d – dimension,
the number of
features*

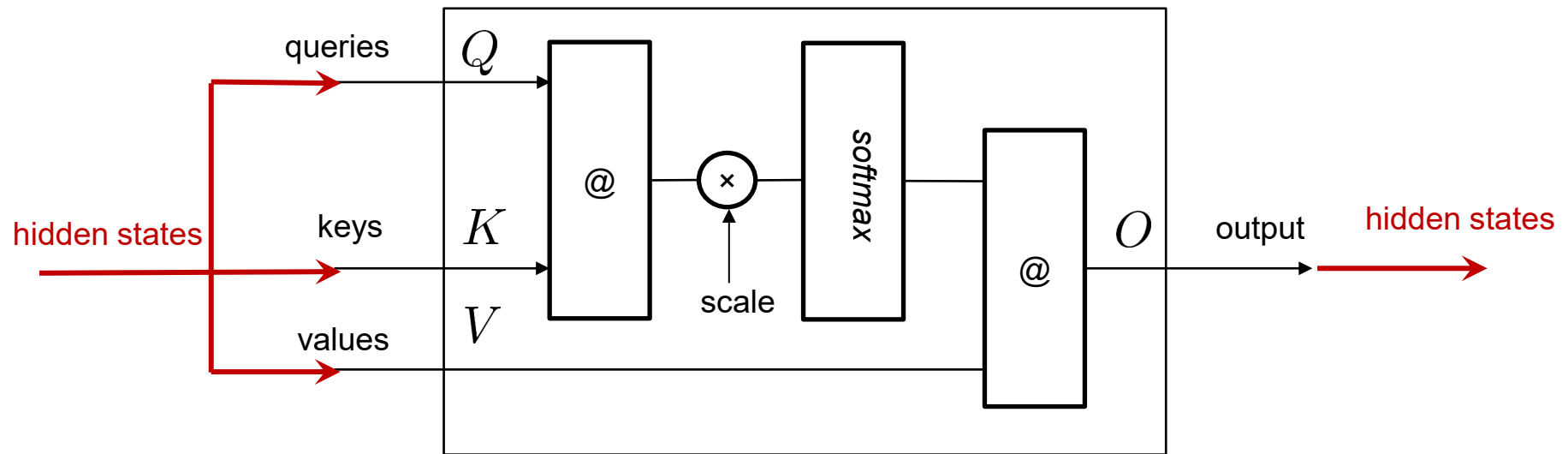


*The number
of keys and
values must
be the same*

$$O = \text{softmax} \left(\frac{QK^T}{\sqrt{d}} \right) V$$

$$\begin{array}{lll} Q: k \times d & K^T: d \times l & V: l \times d \\ K: l \times d & QK^T: k \times l & O: k \times d \end{array}$$

Self-Attention



*The word **ball** has different meanings in different contexts*

He shot the ball into the net.
(the soccer ball)

The ball tore off his leg.
(the cannonball)

*The word **ball** has different meanings in different contexts*

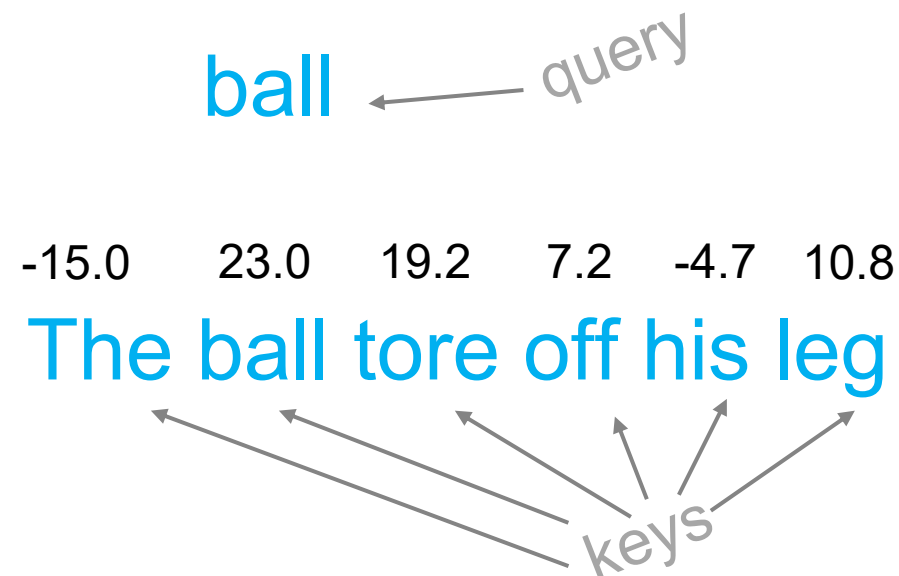
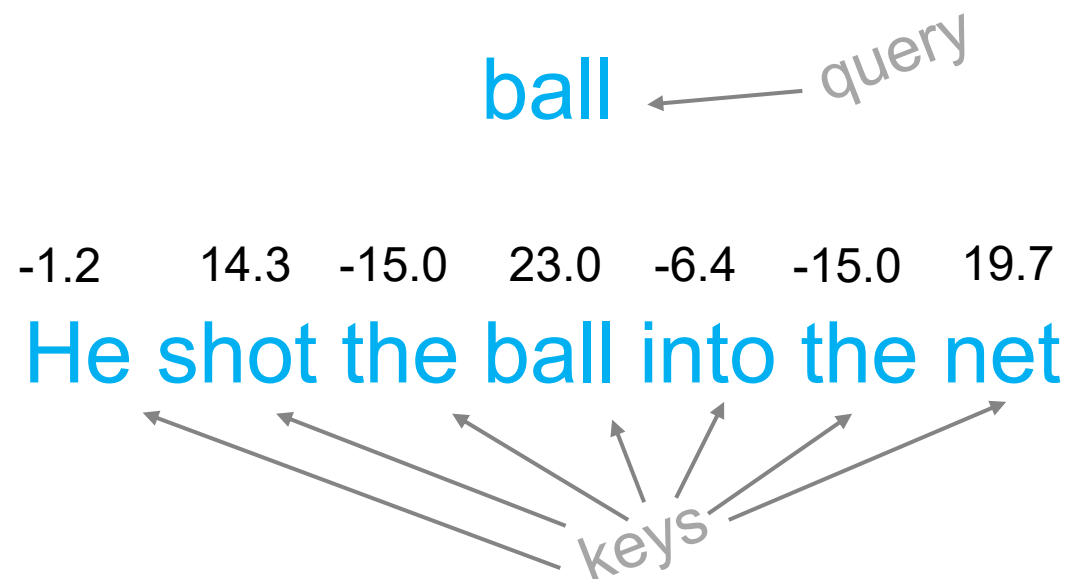
He shot the ball into the net.
(the soccer ball)

The ball tore off his leg.
(the cannonball)

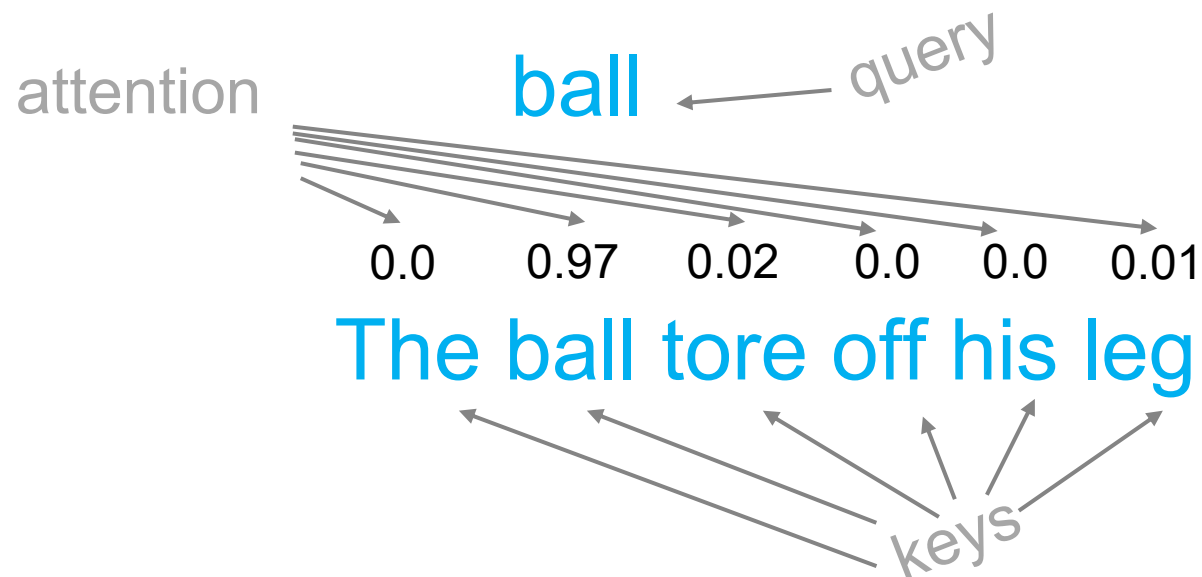
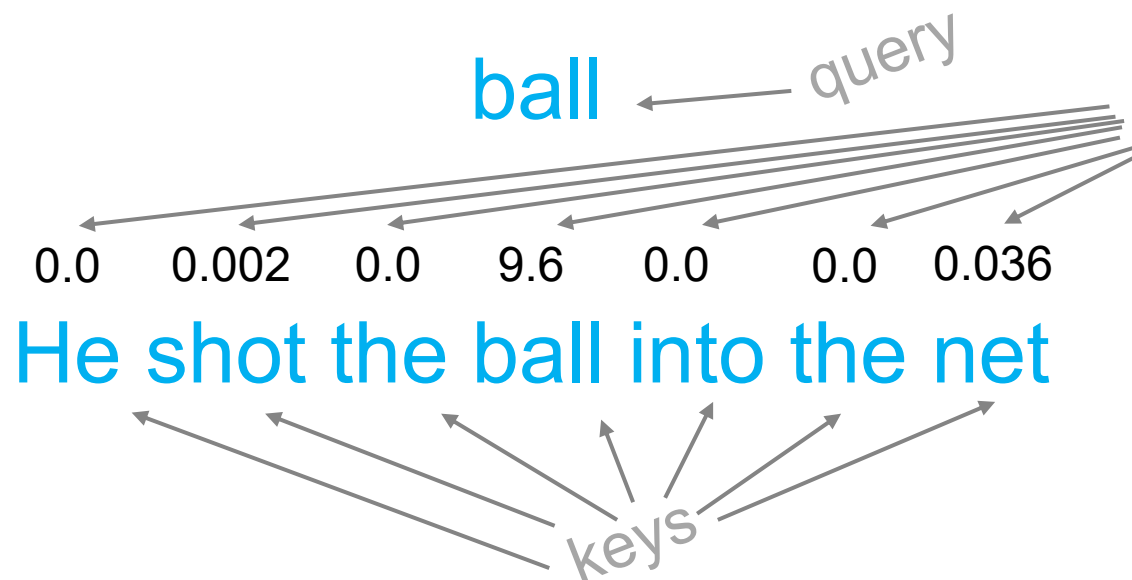
$\text{ball} = c_1 \text{ ball} + c_2 \text{ shot} + c_3 \text{ net}$

$\text{ball} = d_1 \text{ ball} + d_2 \text{ tore} + d_3 \text{ leg}$

For each word, we calculate the cosine similarity to others (including itself)

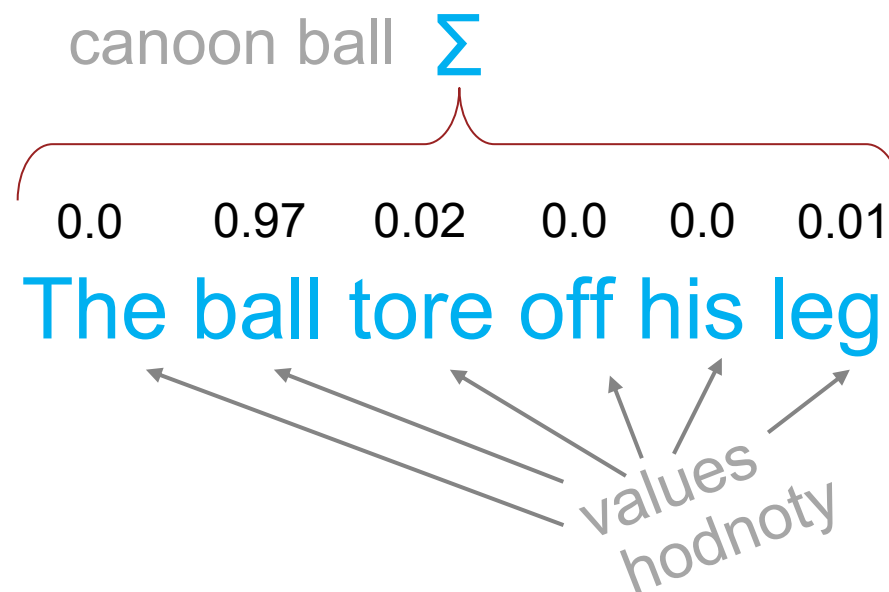
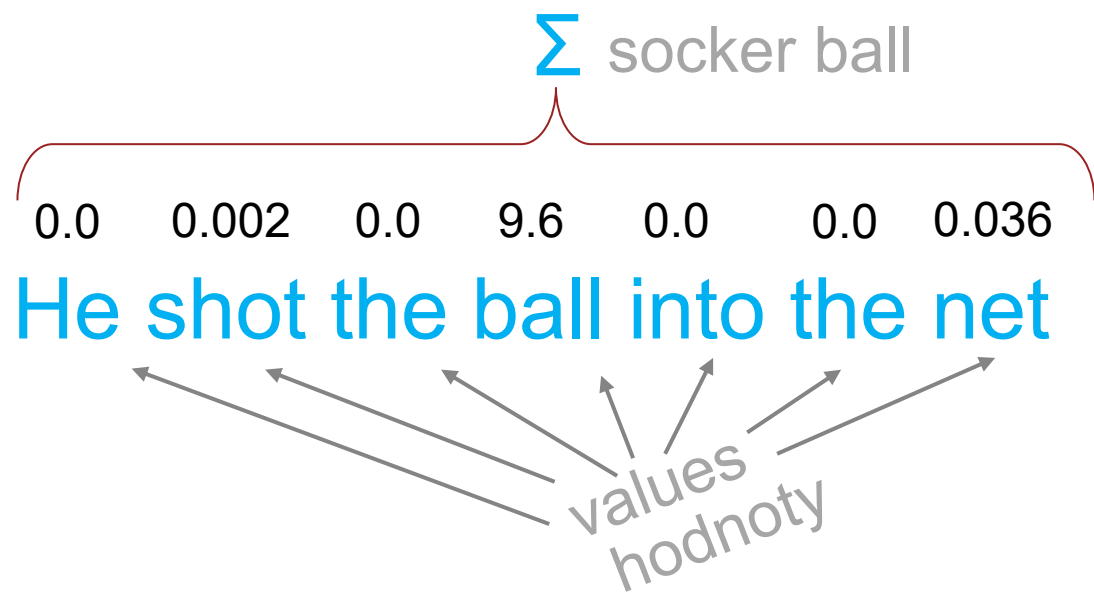


For each word, we calculate the cosine similarity to others (including itself) and apply Softmax



$$\text{softmax}_i(x_1, x_2, \dots, x_n) = \frac{e^{x_i}}{\sum_k e^{x_k}}$$

and then we replace the original representation with a mixture of representations; this gives it a specific meaning



Residual Self-Attention

$$\text{ball} = c_1 \text{ ball} + c_2 \text{ shot} + c_3 \text{ net}$$

$$c_1 = \text{sim}(\text{ball}, \text{ball})$$

$$c_2 = \text{sim}(\text{ball}, \text{shot})$$

$$c_3 = \text{sim}(\text{ball}, \text{net})$$



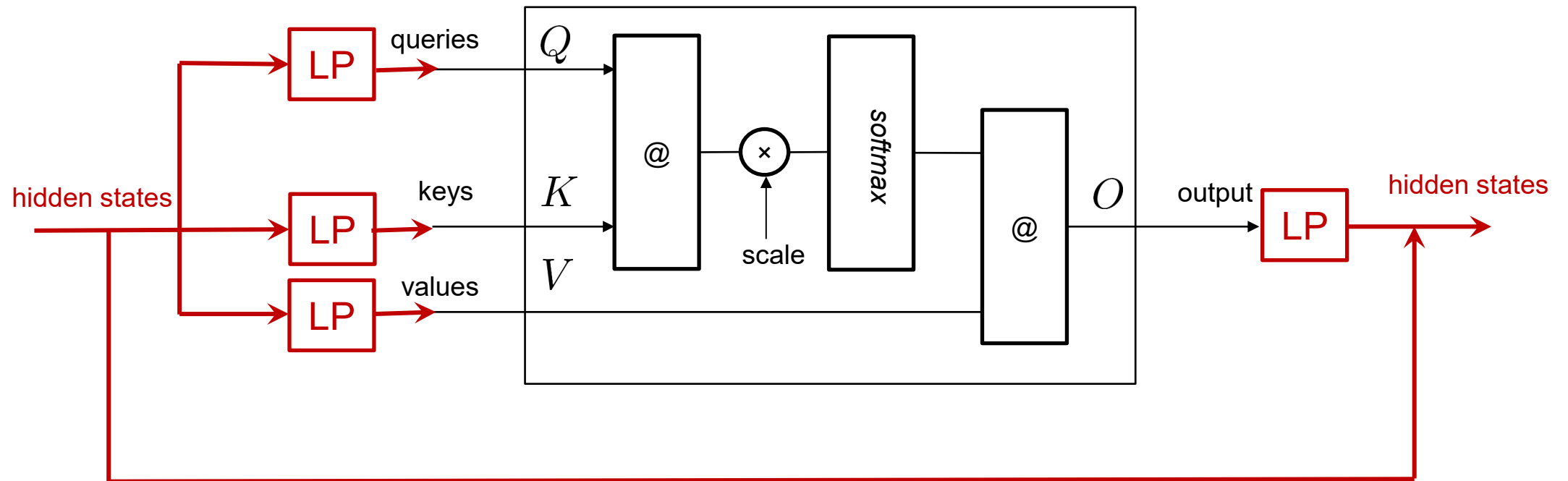
$$\text{ball} += f_o(c_1 f_v(\text{ball}) + c_2 f_v(\text{shot}) + c_3 f_v(\text{net}))$$

$$c_1 = \text{sim}(f_q(\text{ball}), f_k(\text{ball}))$$

$$c_2 = \text{sim}(f_q(\text{ball}), f_k(\text{shot}))$$

$$c_3 = \text{sim}(f_q(\text{ball}), f_k(\text{net}))$$

Residual Self-Attention



The number of queries, keys and values is the same

Residual Cross-Attention

$$\text{ball} += f_o(c_1 f_v(\text{do}) + c_2 f_v(\text{you}) + c_3 f_v(\text{like}) + c_4 f_v(\text{soccer}))$$

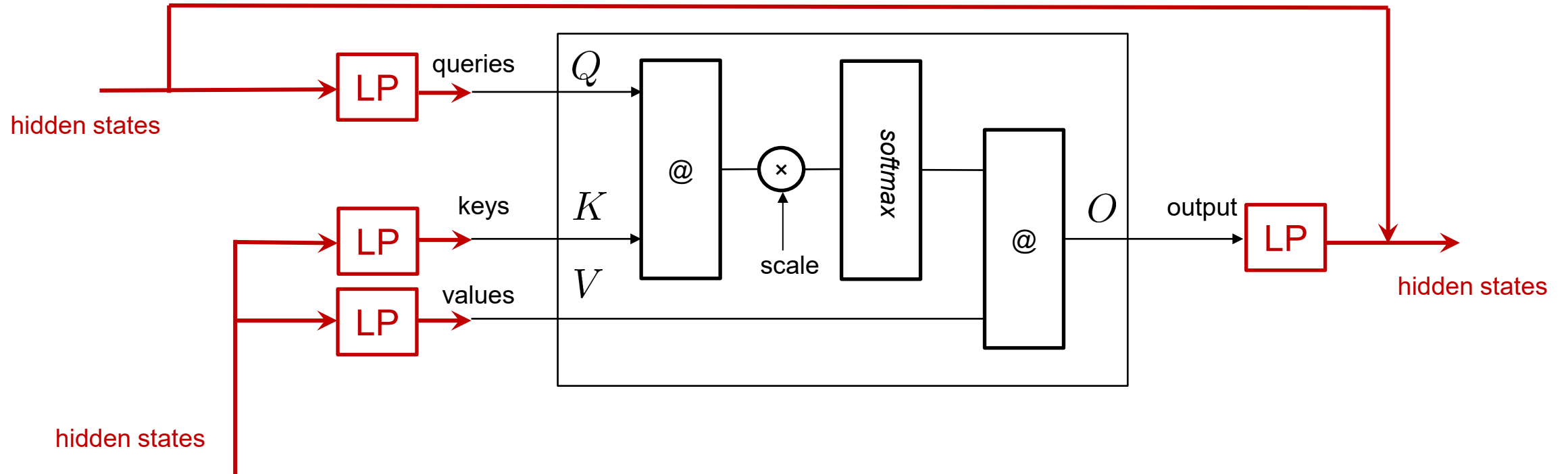
$$c_1 = \text{sim}(f_q(\text{ball}), f_k(\text{do}))$$

$$c_2 = \text{sim}(f_q(\text{ball}), f_k(\text{you}))$$

$$c_3 = \text{sim}(f_q(\text{ball}), f_k(\text{like}))$$

$$c_4 = \text{sim}(f_q(\text{ball}), f_k(\text{soccer}))$$

Residual Cross-Attention

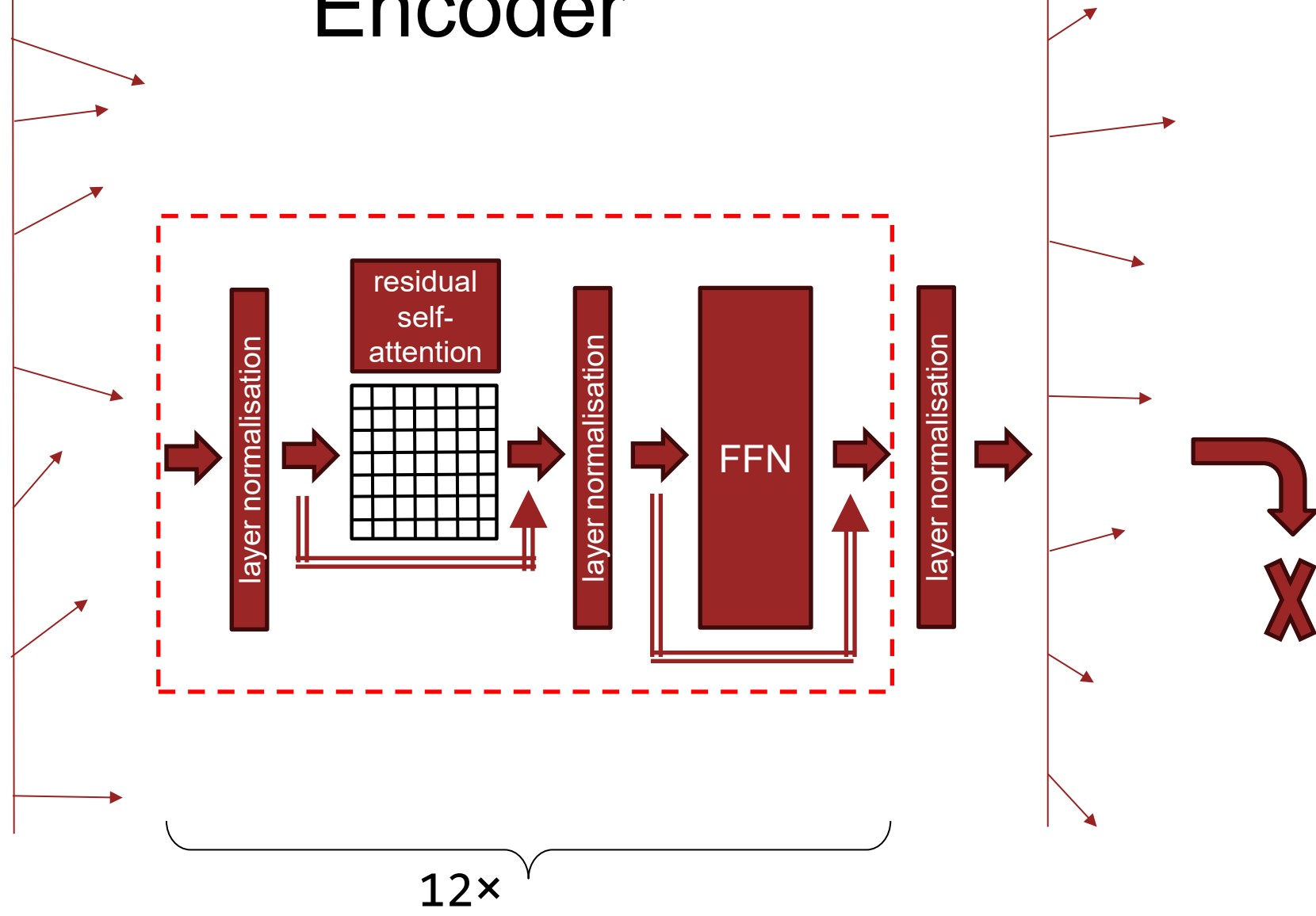


The number of queries differs from the number of keys and values

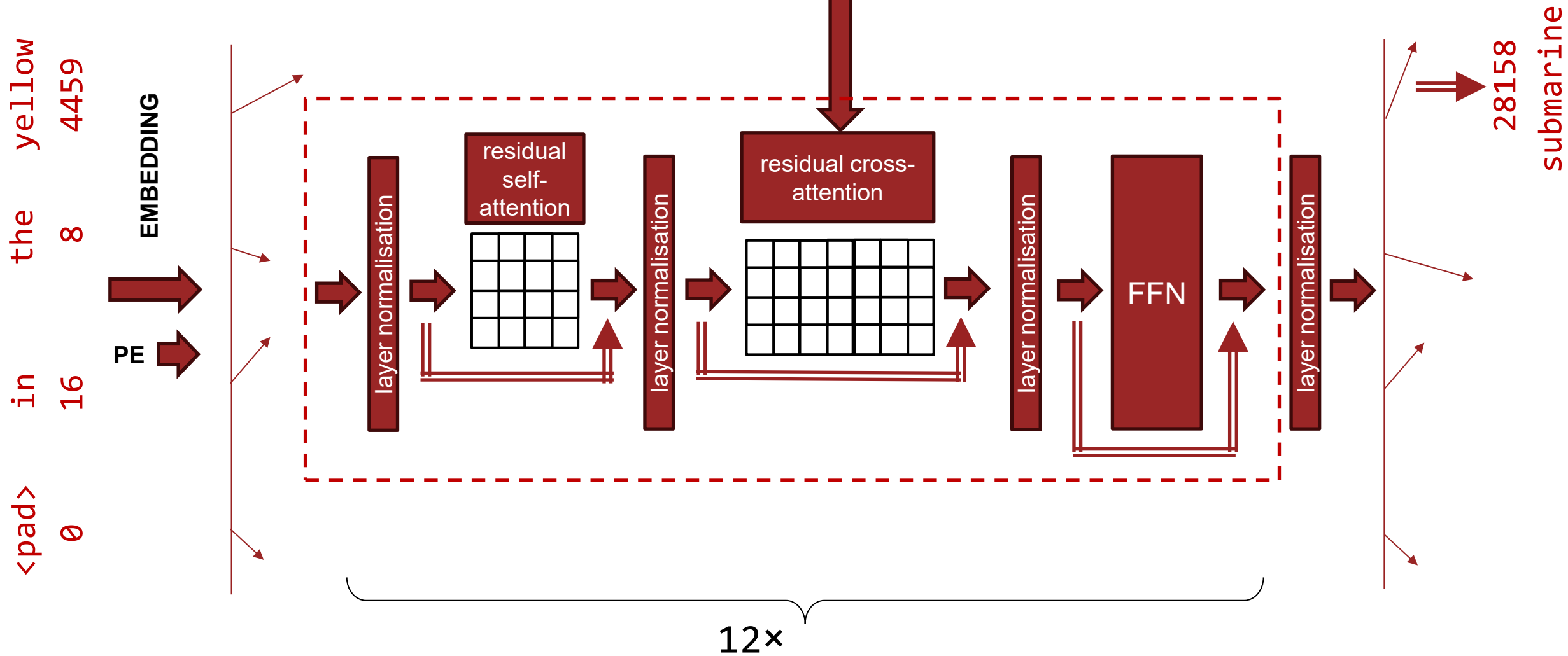
<pad> Where does Beatles live ? <eos>

0 2840 405 27615 619 58 1

PE EMBEDDING

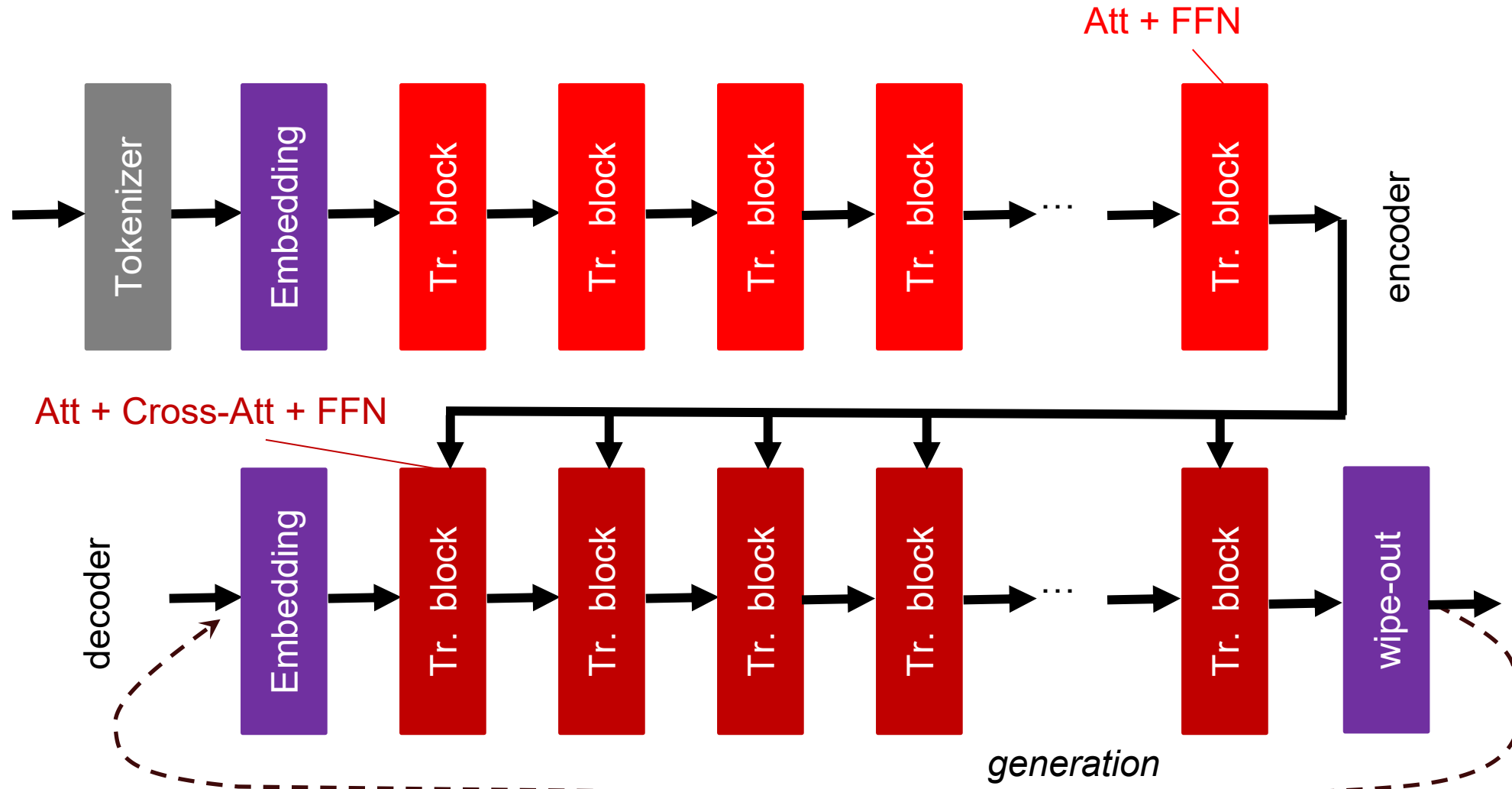


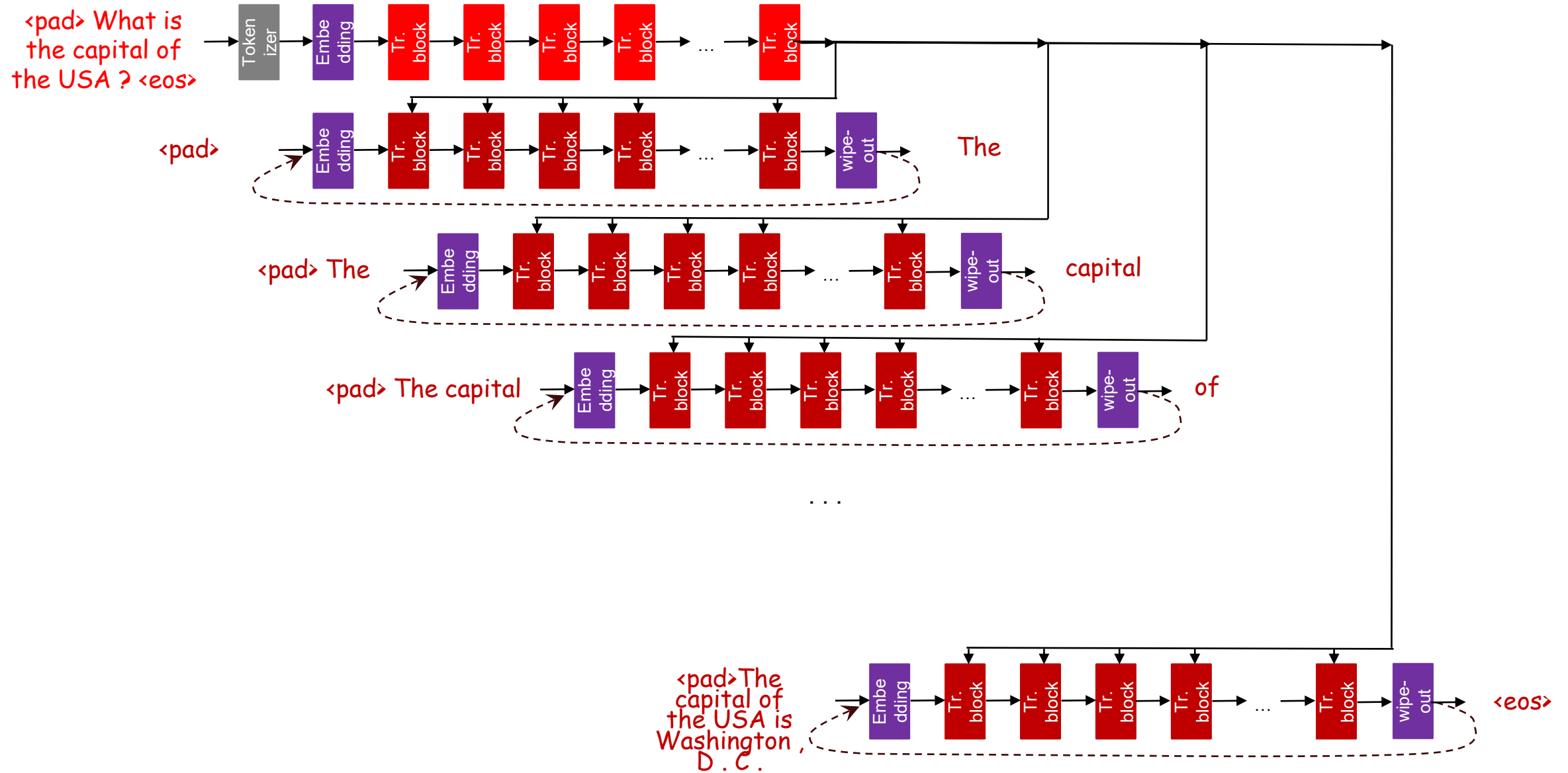
Decoder



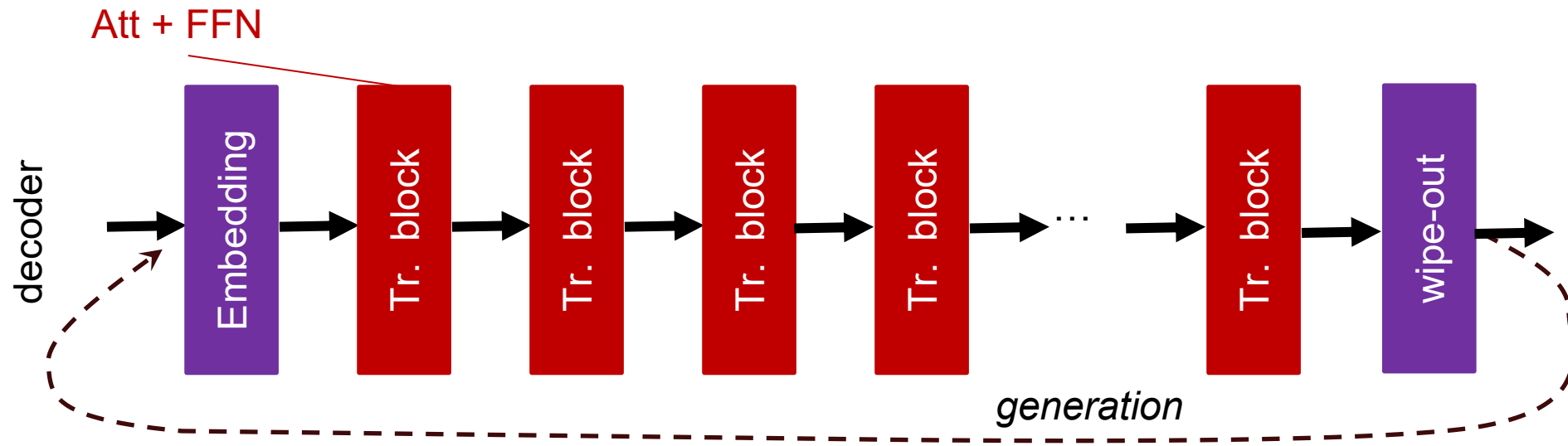
Chatbot (encoder-decoder type transformer)

T4, Flann, LaMini

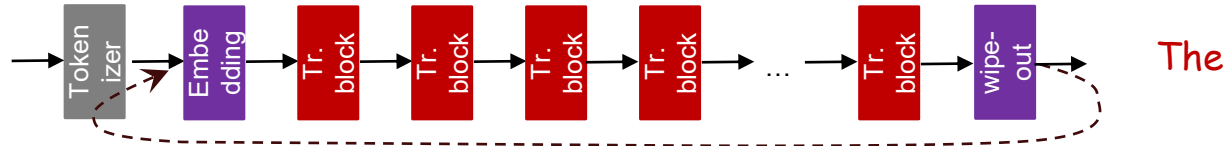




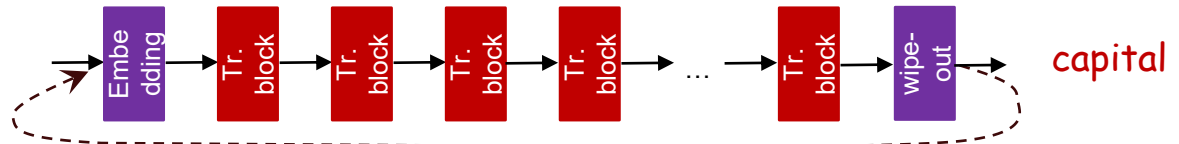
Chatbot (decoder-only type transformer)



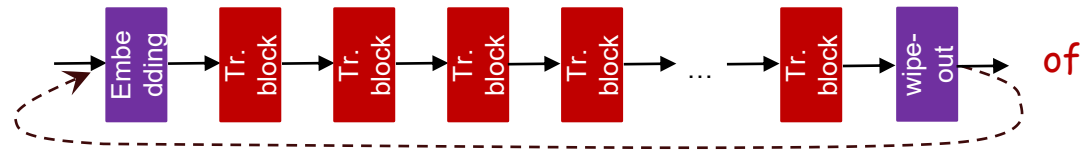
<User:> What is
the capital of
the USA ?
<Assistant:>



<User:> What is
the capital of
the USA ?
<Assistant:> The

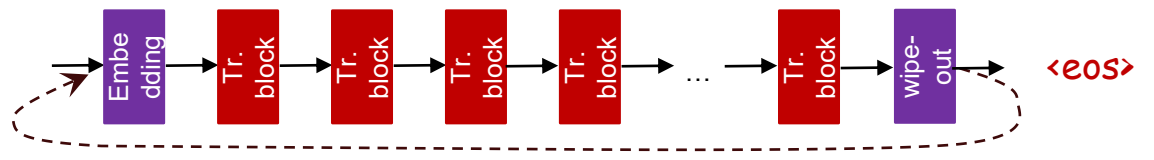


<User:> What is
the capital of
the USA ?
<Assistant:> The
capital

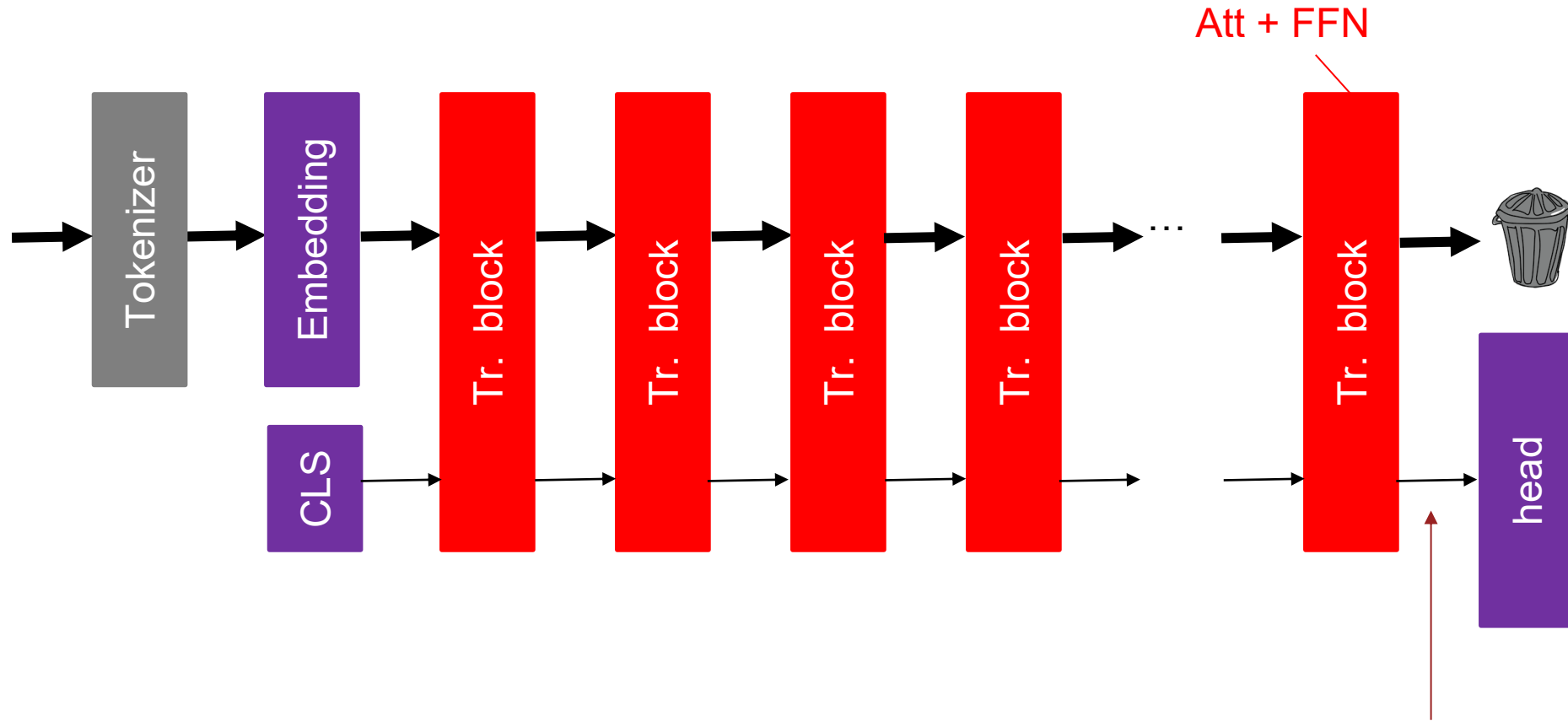


...

<User:> What is the
capital of the USA ?
<Assistant:> The capital
of the USA is
Washington , D . C .

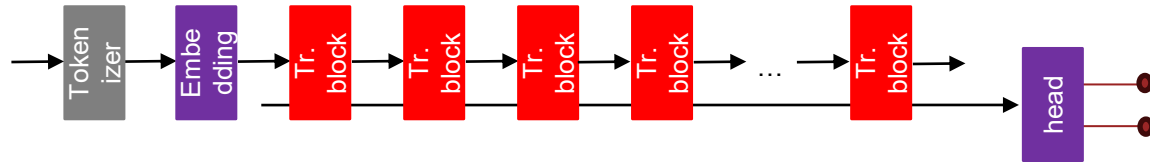


Classifier (encoder-only type transformer)



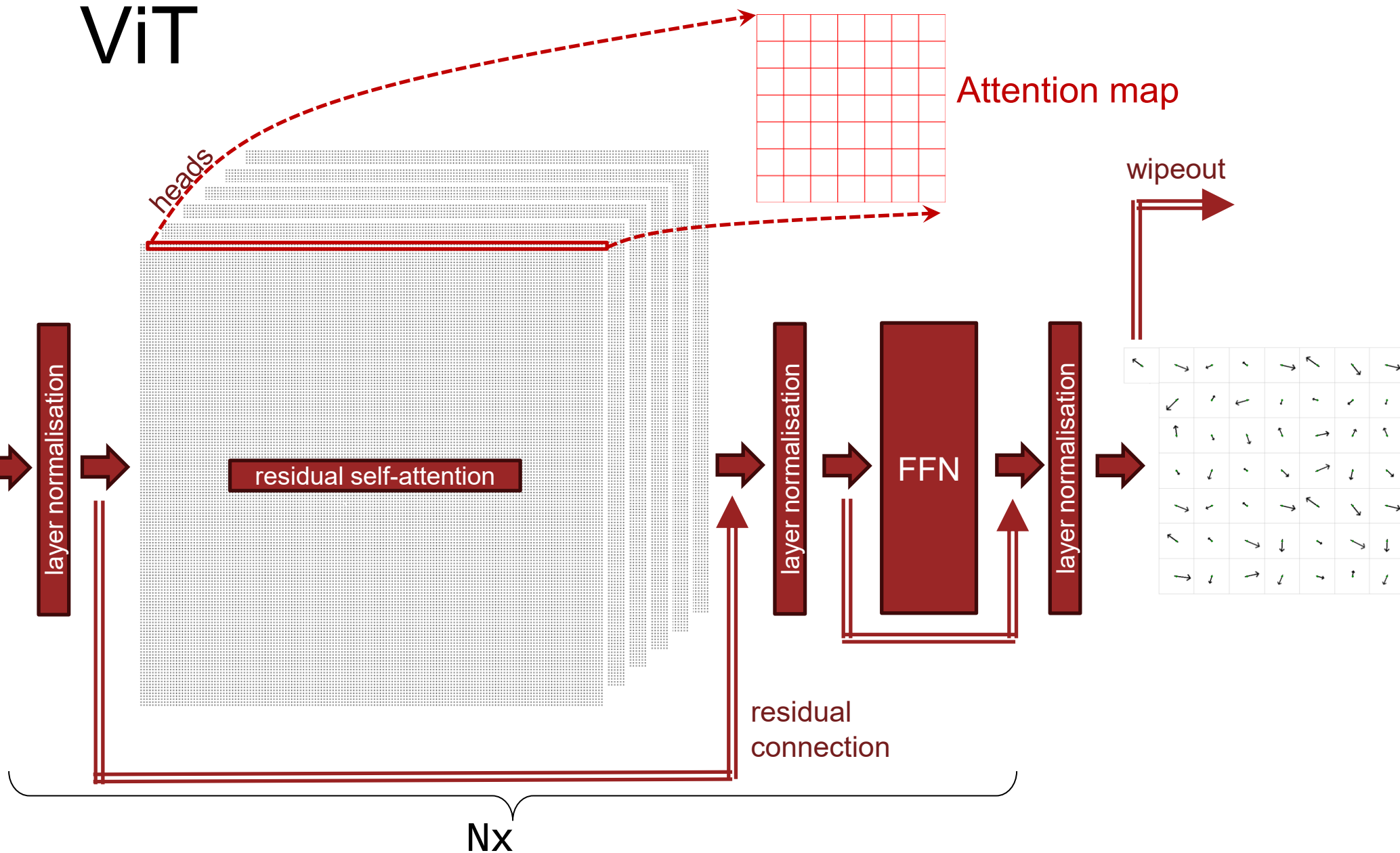
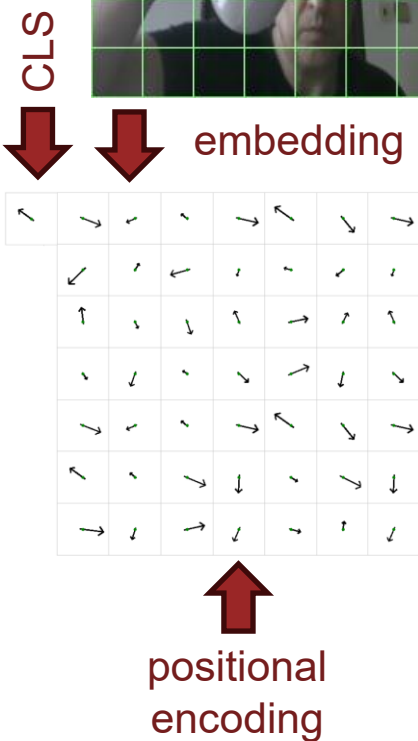
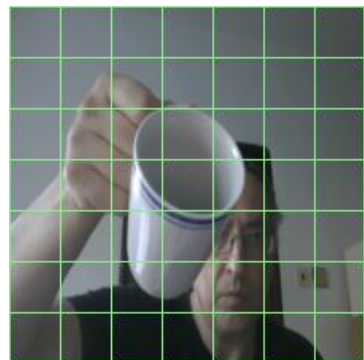
The classification hidden state is a single vector with latent space dimension

<pad> What is
the capital of
the USA ? <eos>



<a question on the capital cities>
<other question>

ViT



Detection transformer (DeTR)

- Type: encoder – decoder
- Encoder encode image (ViT without CLS)
- The decoder is called once only
- The decoder emits N ($=100$) detections into the image and gradually moves them to their correct locations.

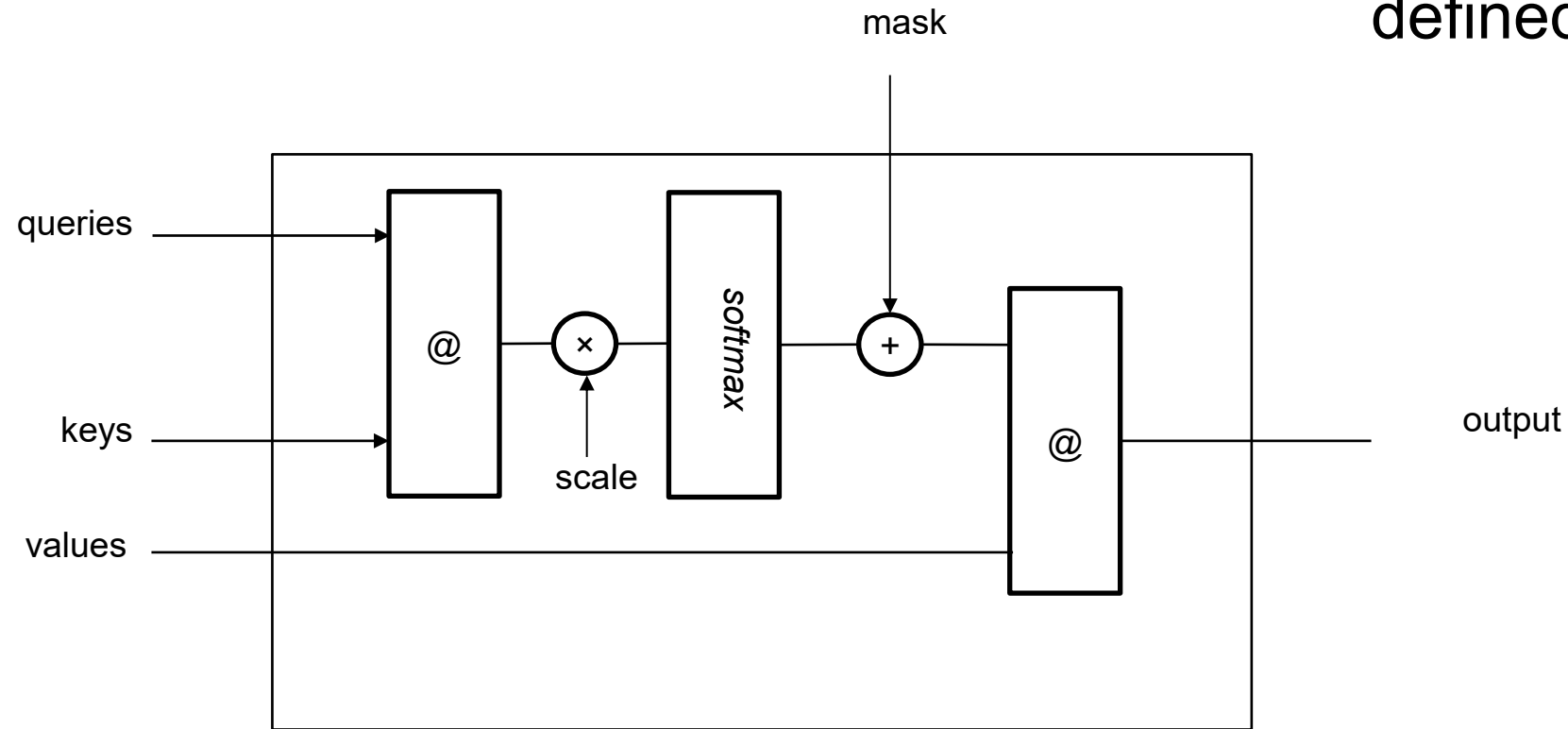
How to represent detections

Detection contains:

- Bounding box
 - confidence
 - class
-
- DeTR encodes the detection as a vector of length equal to the latent space dimension. How? We leave that to the network itself, let it come up with a representation that suits it.

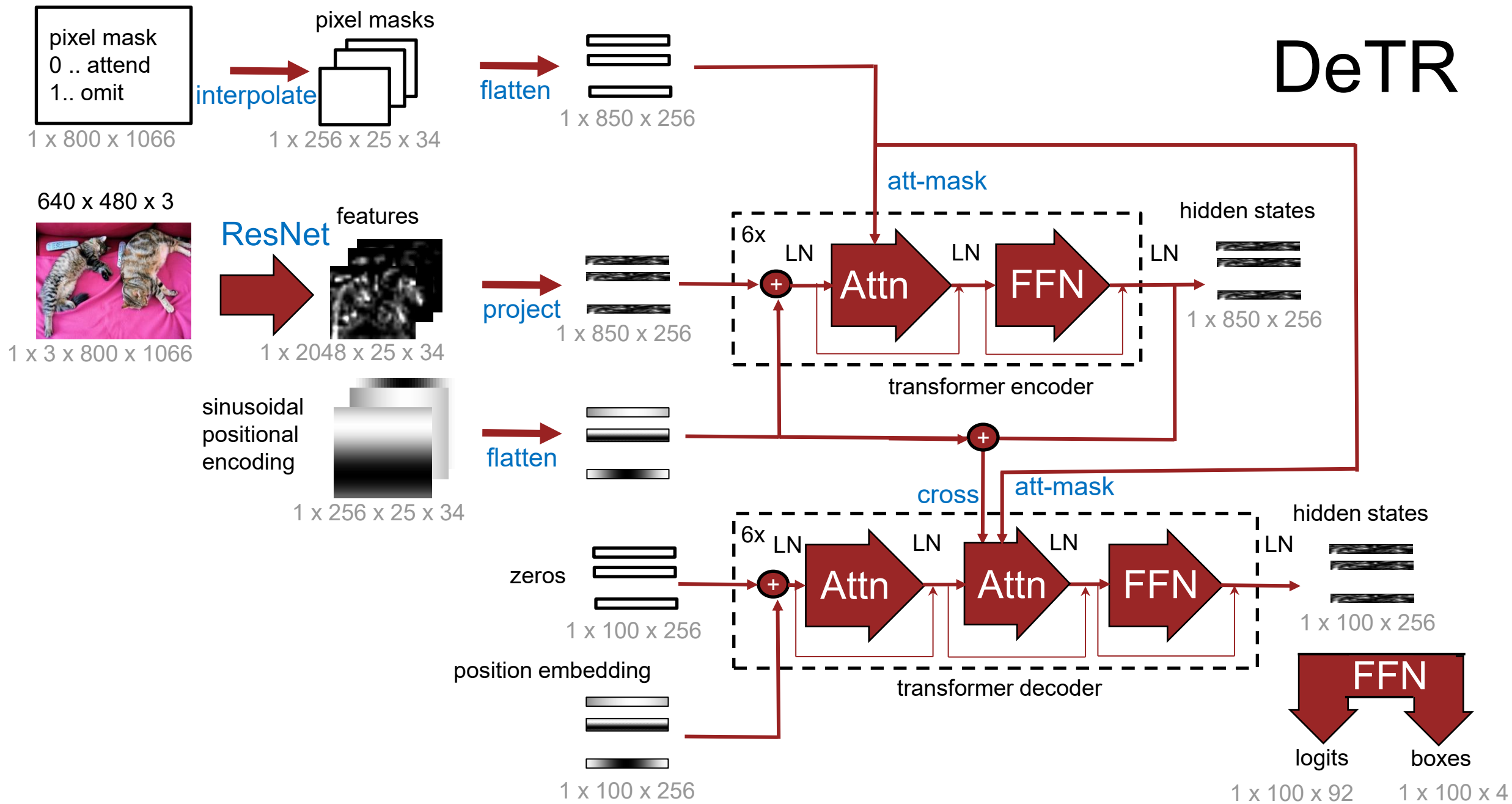
Masked Attention

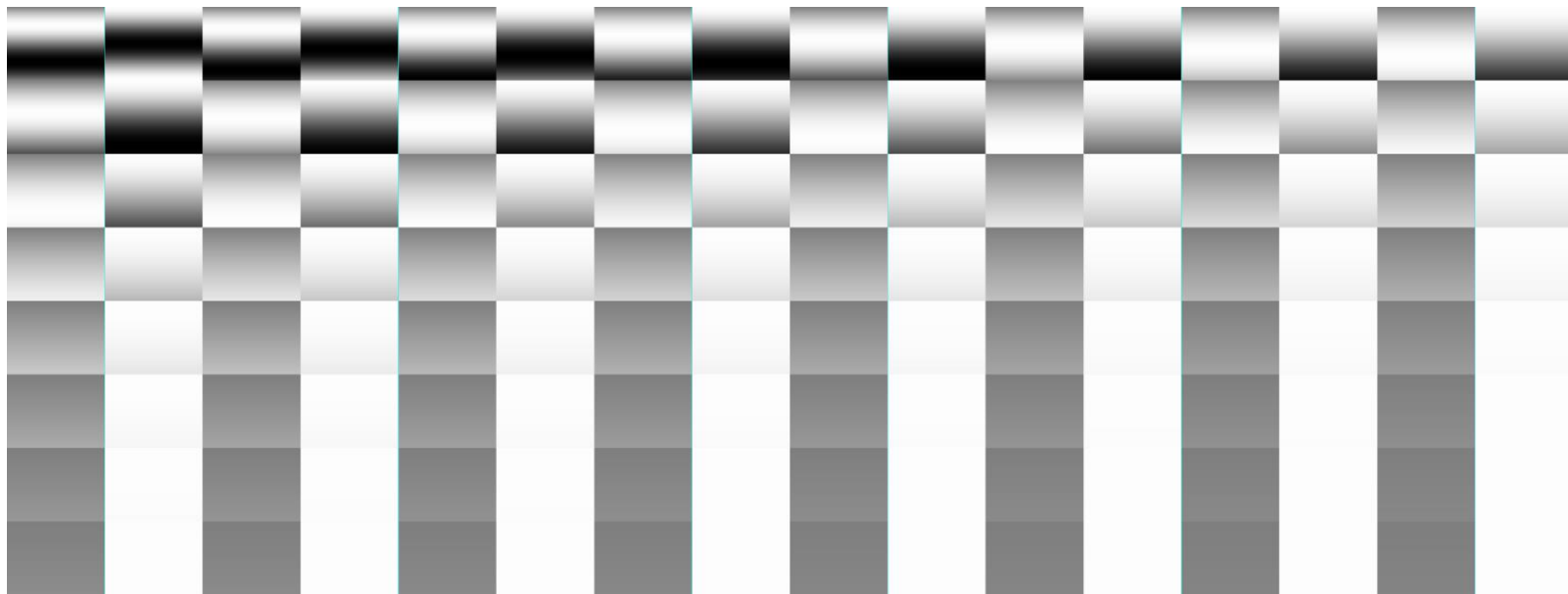
- DeTR allows us to search for objects only within a specific region of the image, defined by a mask.



$$Att(Q, K, V) = \left[\text{softmax} \left(\frac{QK^T}{\sqrt{d}} \right) + \text{mask} \right] V$$

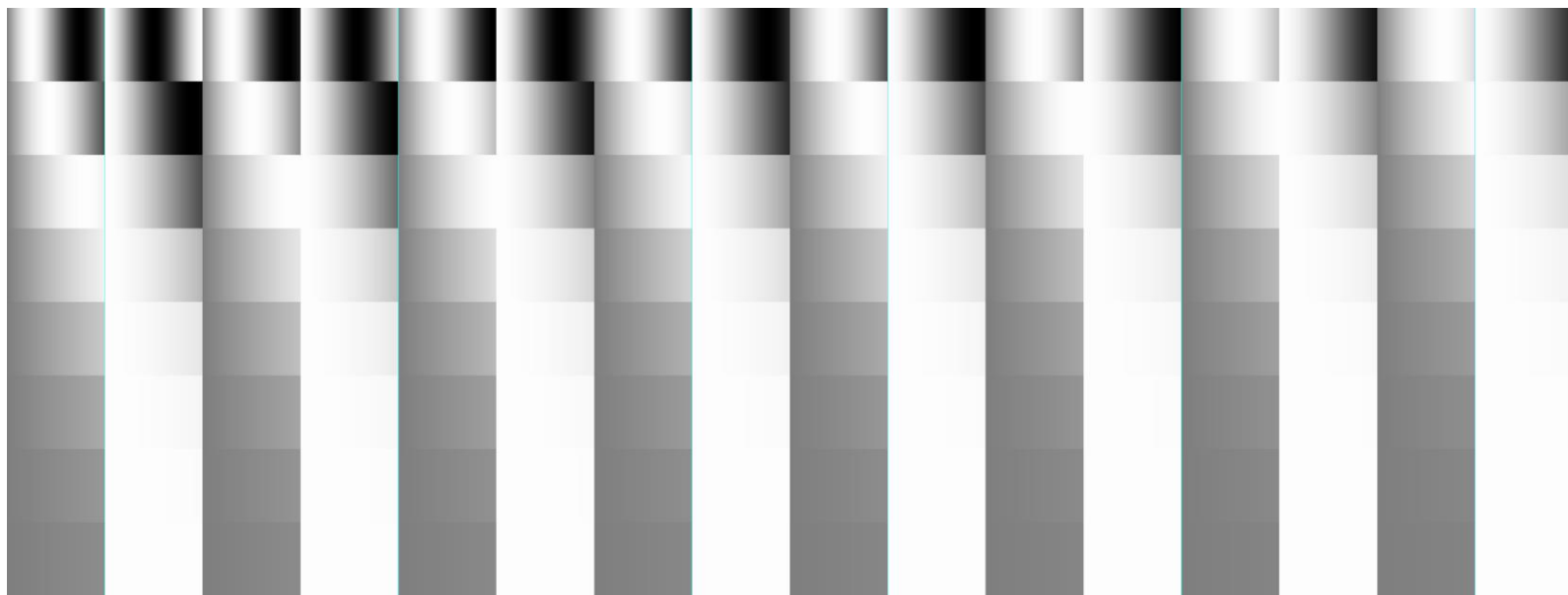
DeTR



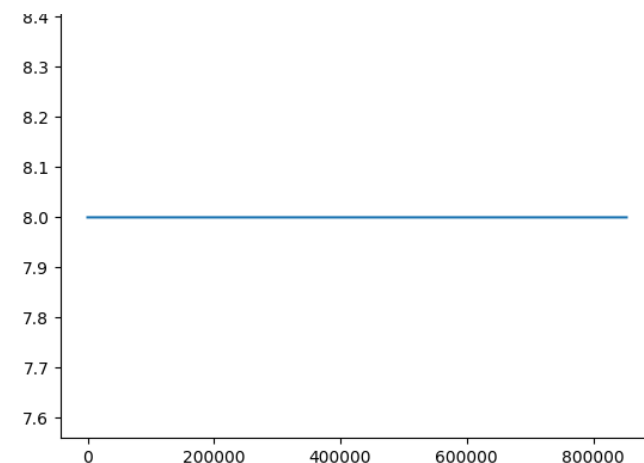


Sinusoidal positional encoding

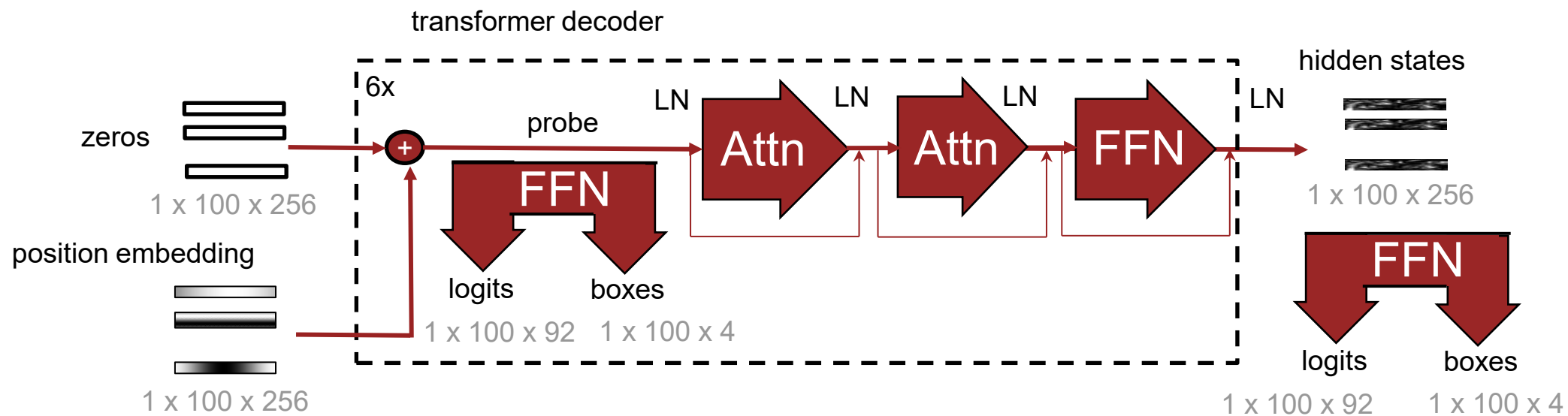
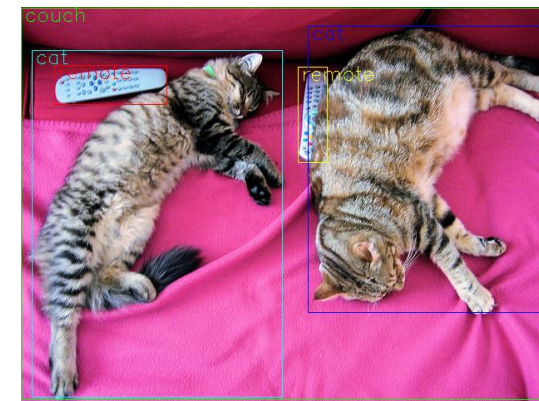
$$P(k, 2i) = \sin\left(\frac{k}{n^{2i/d}}\right)$$



$$P(k, 2i + 1) = \cos\left(\frac{k}{n^{2i/d}}\right)$$



Probing





0

couch

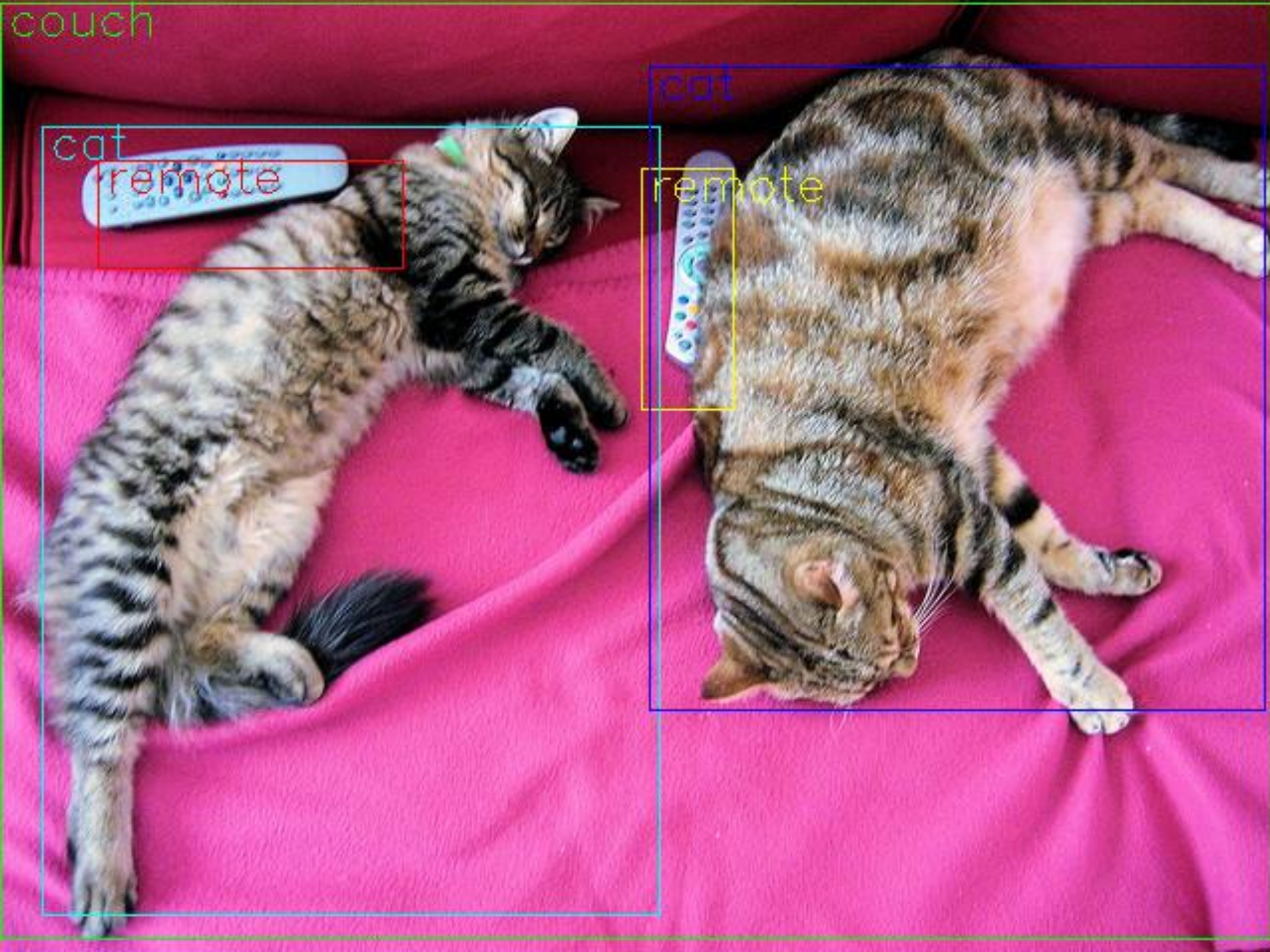
cat

remote

cat

remote

1



couch

cat

remote

cat

remote

2



couch

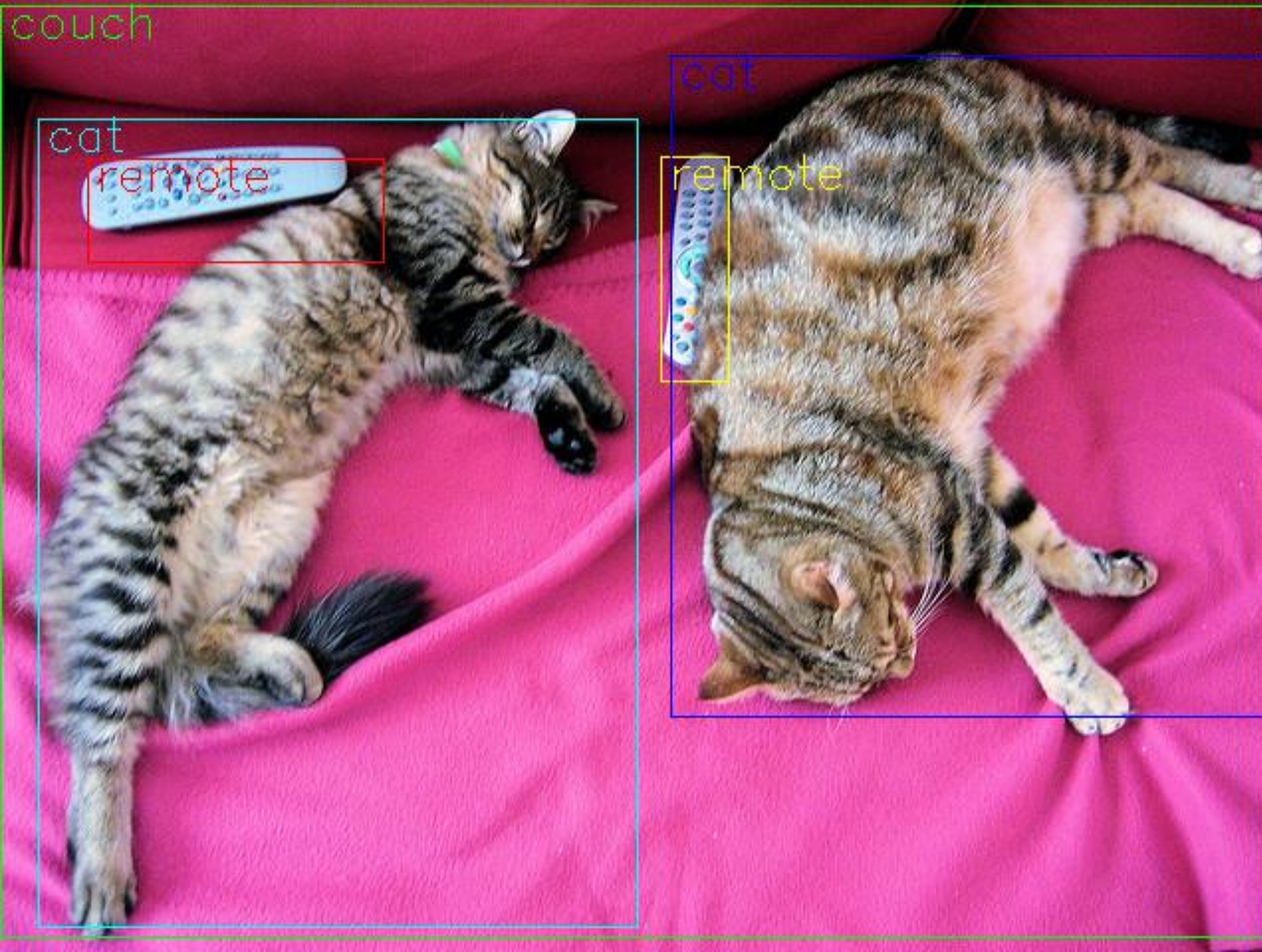
cat

remote

cat

remote

3



couch

cat

remote

cat

remote

4



couch

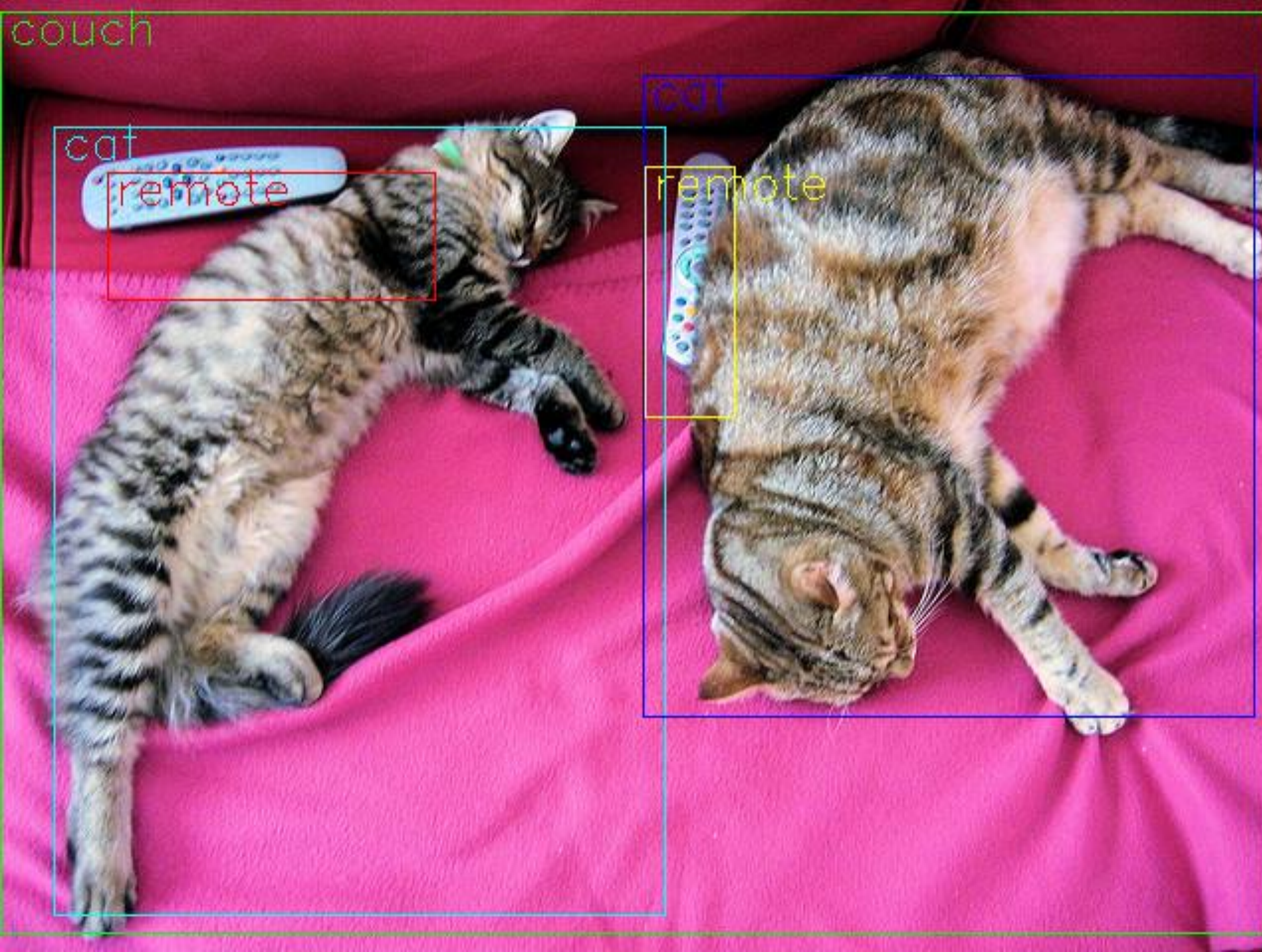
cat

remote

cat

remote

5



couch

cat

remote

cat

remote

6

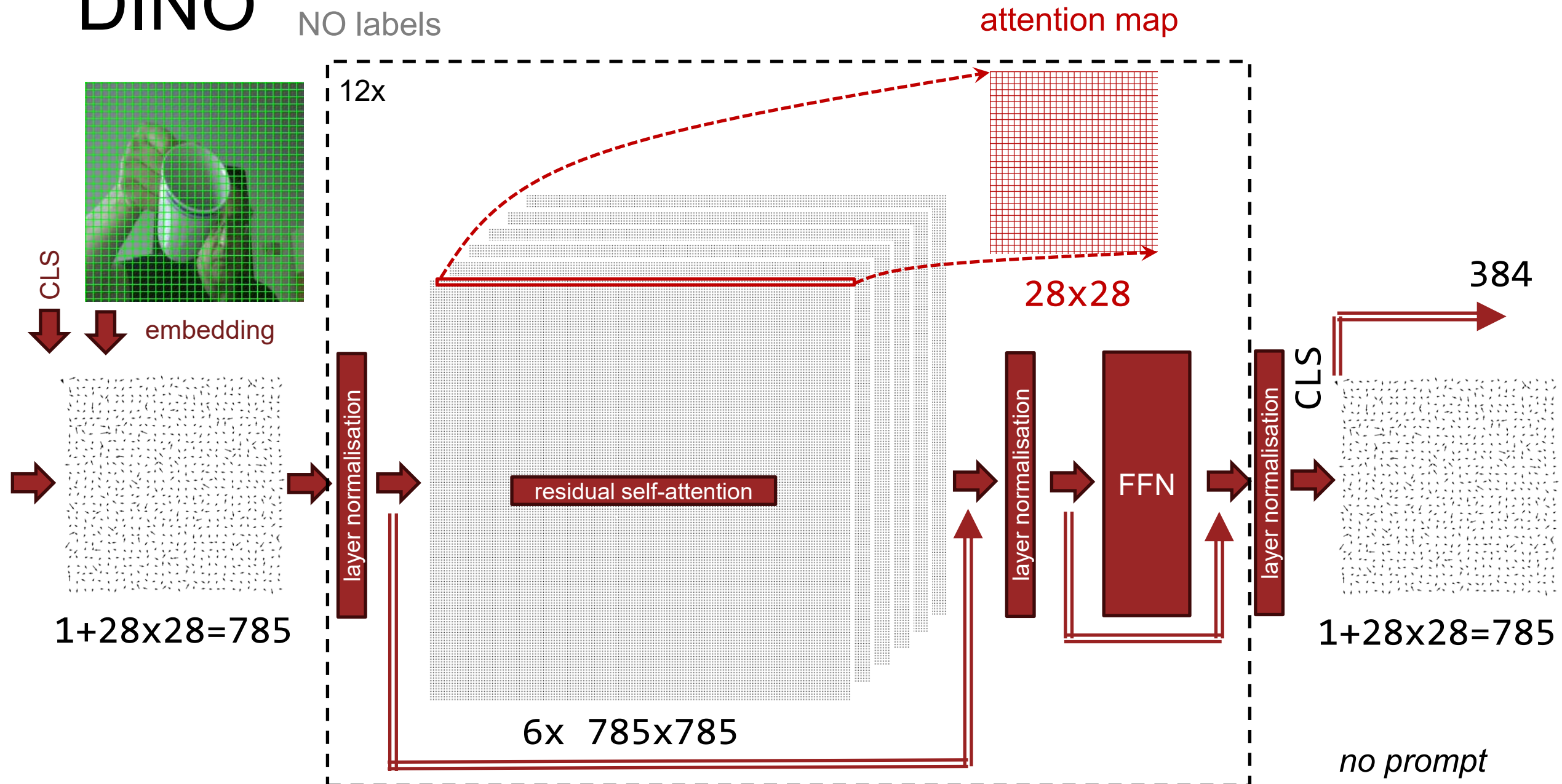


Foundation models

- We train ViT and DeTR for a given set of object categories
- Can we develop models which:
 - classify anything
 - detect anything
 - we do not train, just use (zero-shot), or treat a bit only (few shot)
- Why not?
 - we have large dataset
 - we can't annotate them, not only are there so many, but we don't know what "anything" means

DINO

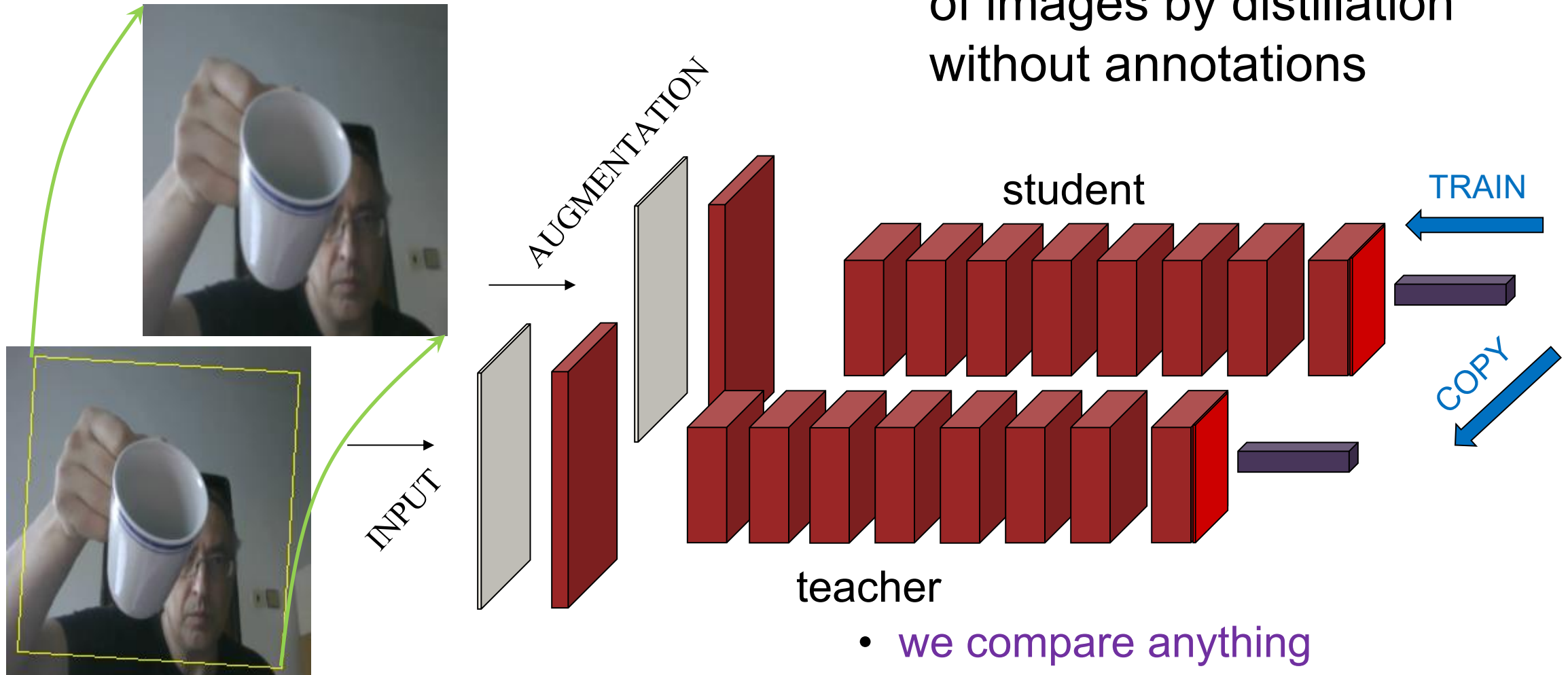
Distillation with
NO labels



DINO

Distillation with
NO labels

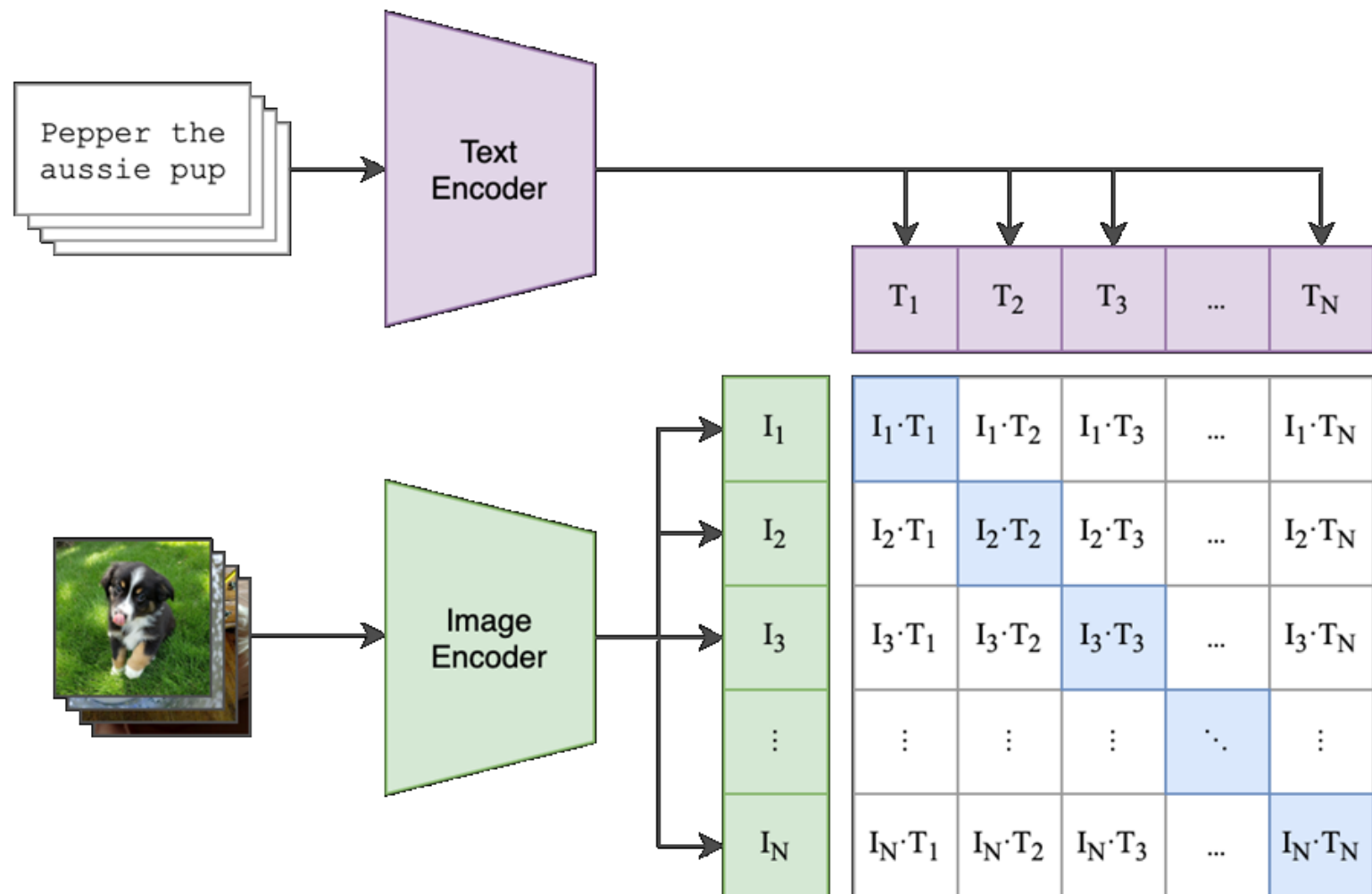
- we train from a large number of images by distillation without annotations



- we compare anything
- we see where is the “anything” located

CLIP

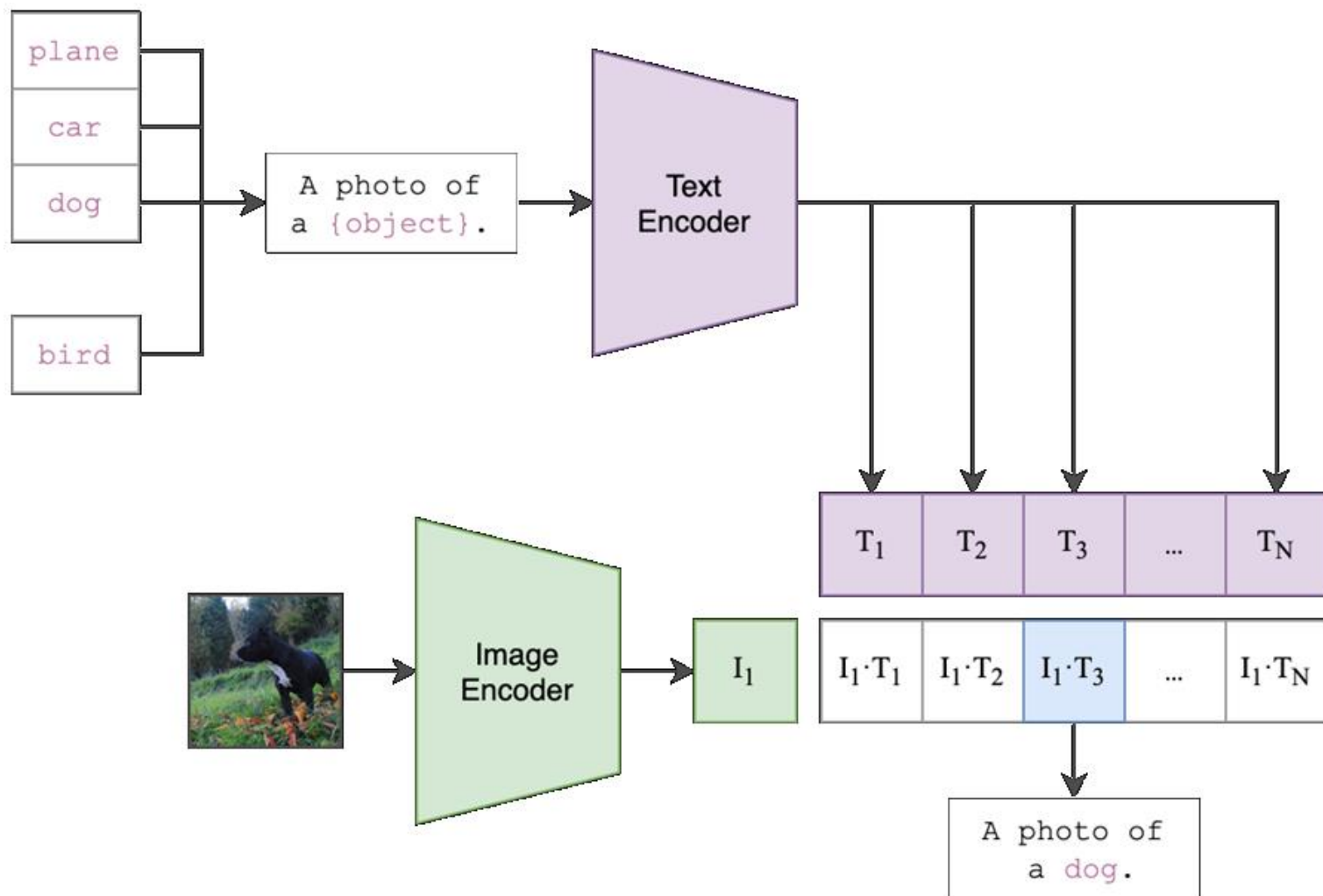
Contrastive Language-Image Pre-training



CLIP is a specially trained visual transformer + text transformer (large language model), where both encoders are coupled by training with a contrastive error function, where we require that the dot product of the codes of matching samples be greater than the codes of those that do not match.

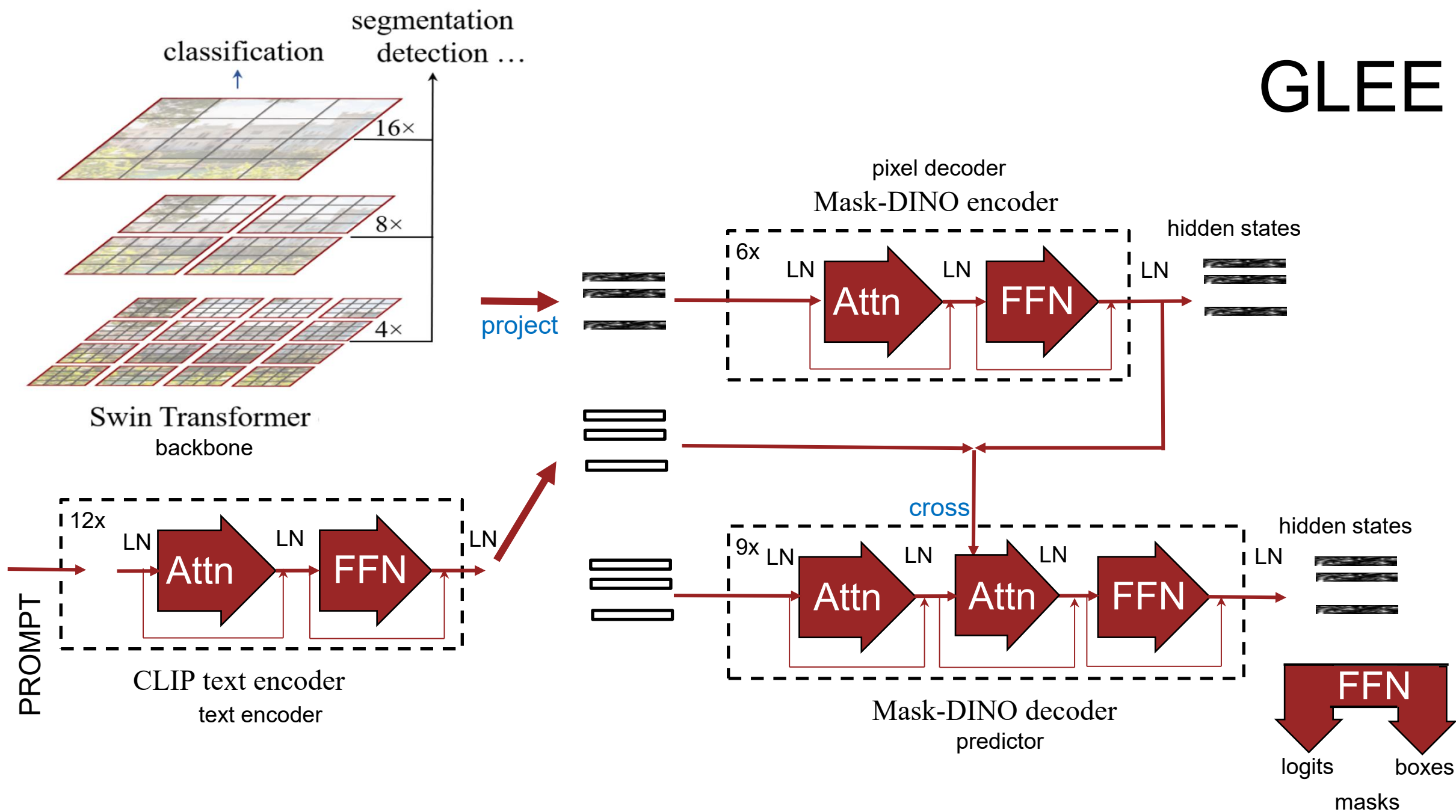
CLIP

Contrastive Language-Image Pre-training



- When using it, we then compare one current code (corresponding to the question asked or the image seen) to a set of previously prepared codes (from texts or images).
- Where the largest scalar product is, that is the pair that corresponds (there must be one).
- we can classify anything
we can name
- we can distinguish between text-specified categories

GLEE



GLEE



The first dog from the right



- Segment anything that we can describe by text, prompt by points, or scribble

SAM3 Model Architecture with Full LoRA Adaptation

