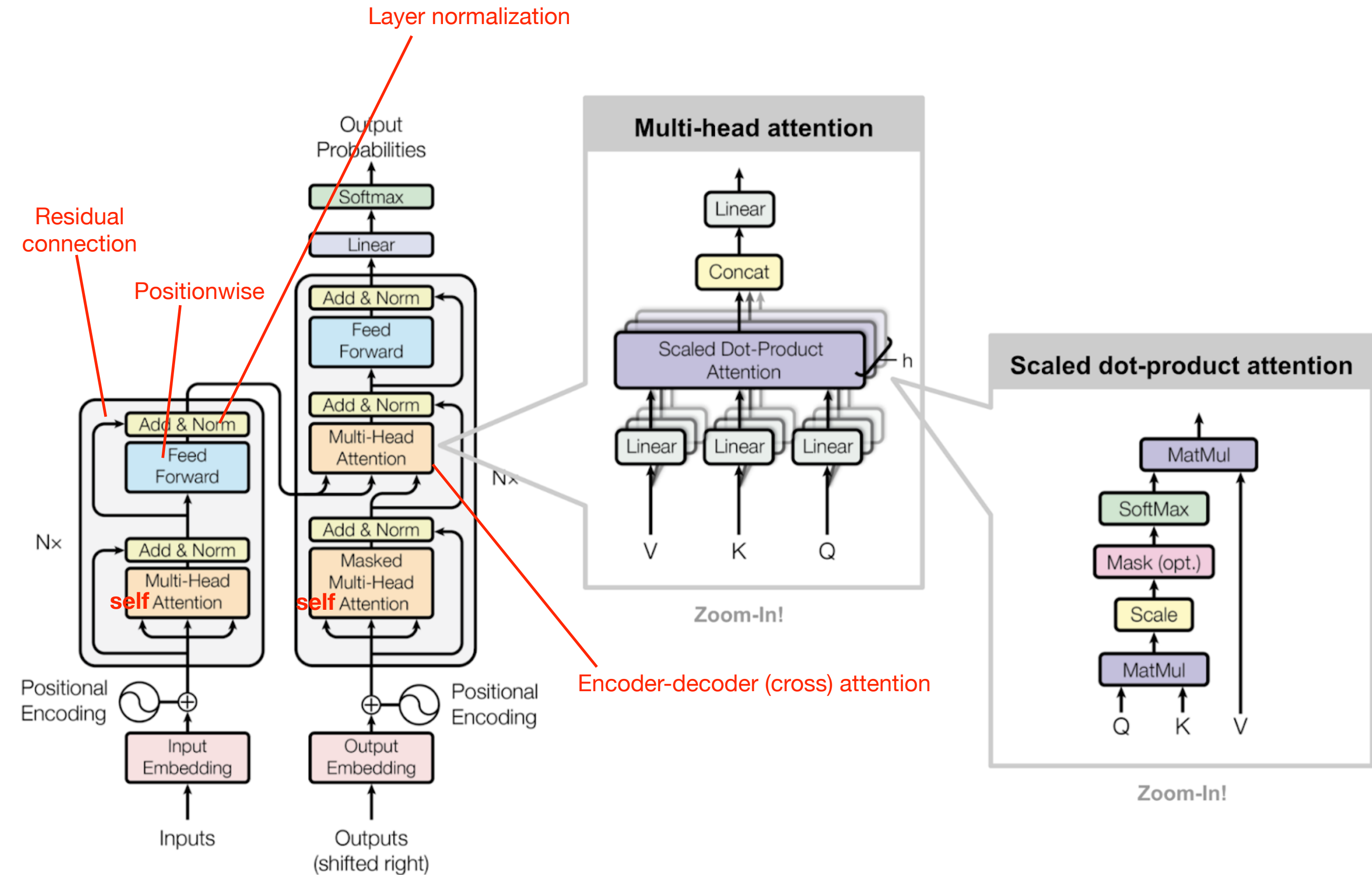


Attention and Transformer (Big Picture)



An Illustration of Self-attention

The current word (Query) is in red, the size of blue shade indicates the attention weight.

The FBI is chasing a criminal on the run .

The FBI is chasing a criminal on the run .

The FBI is chasing a criminal on the run .

The FBI is chasing a criminal on the run .

The FBI is chasing a criminal on the run .

The FBI is chasing a criminal on the run .

The FBI is chasing a criminal on the run .

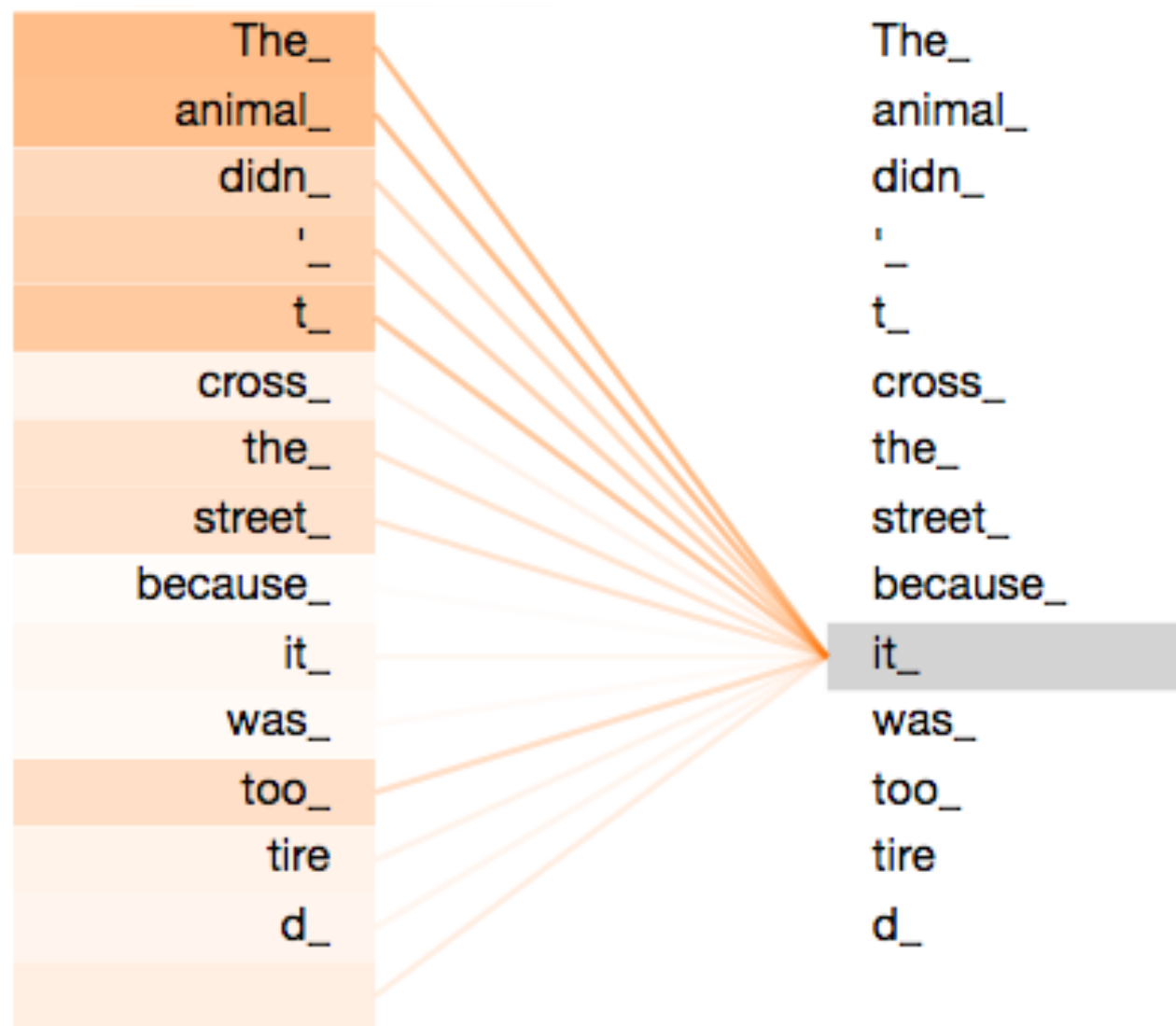
The FBI is chasing a criminal on the run .

The FBI is chasing a criminal on the run .

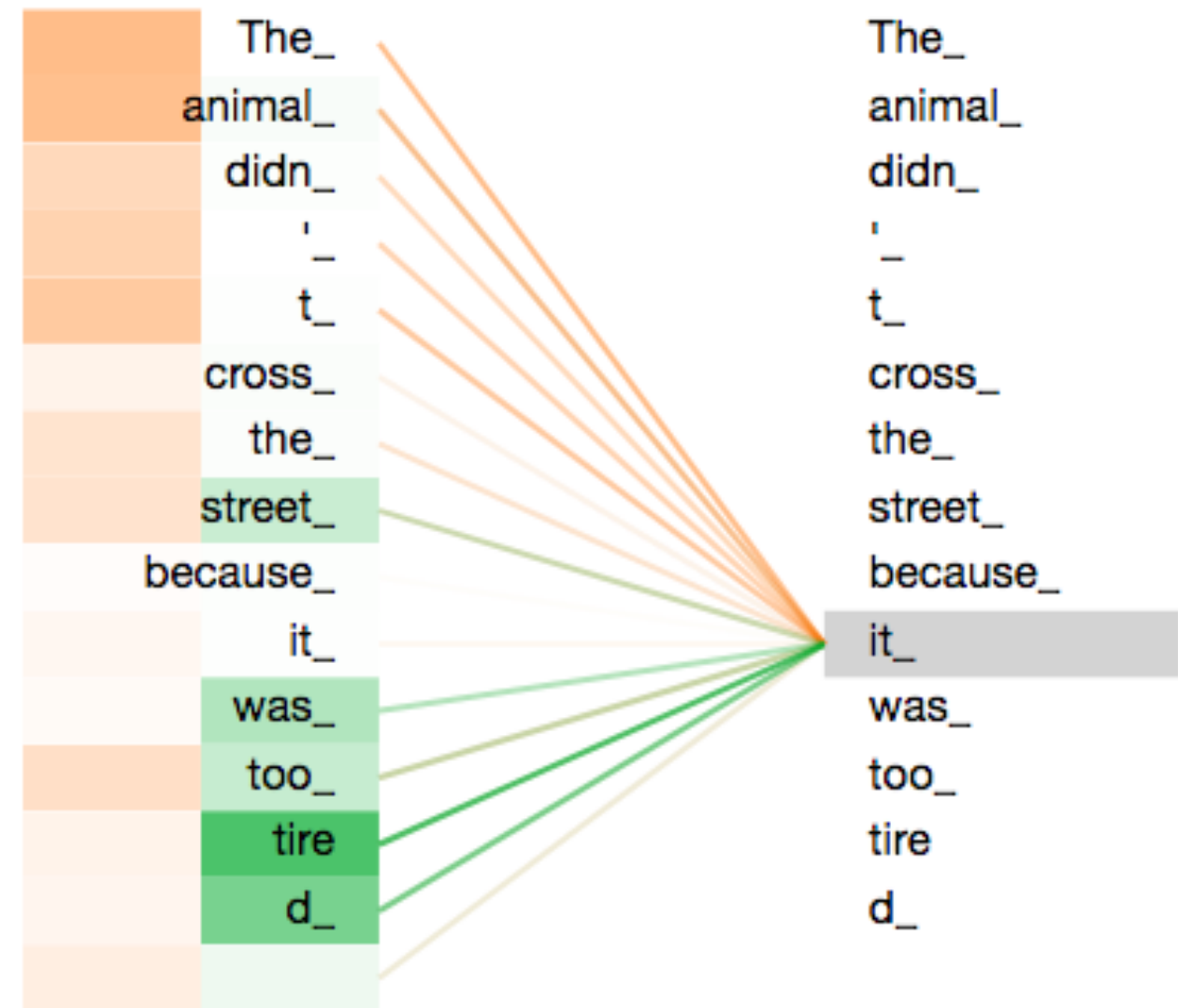
The FBI is chasing a criminal on the run .

Attention and Multi-head Attention (Intuition)

$head_1$



$head_1$ $head_2$



Positional Encoding

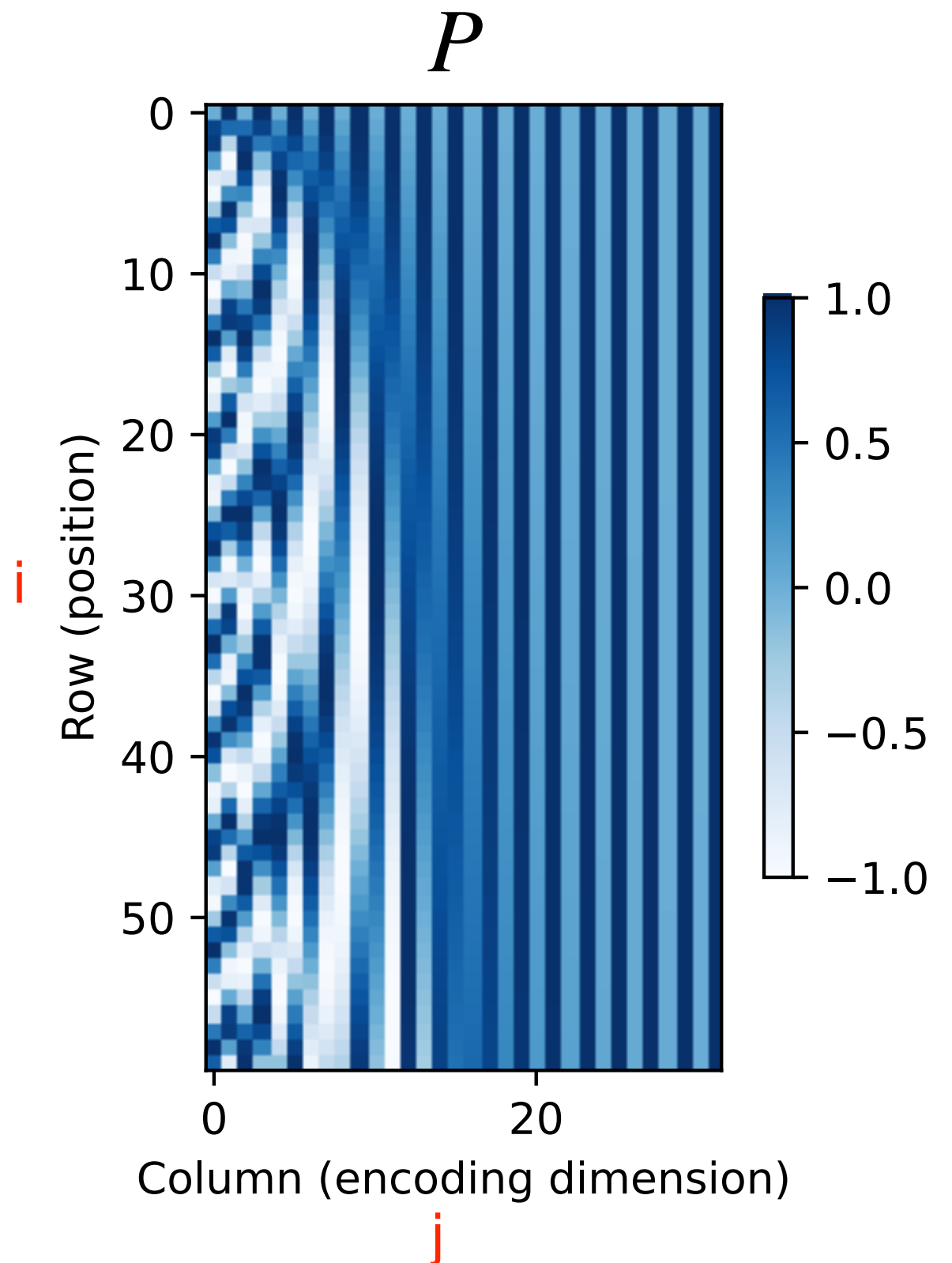
X & P : Seq x Emb matrices

Positional encoding:

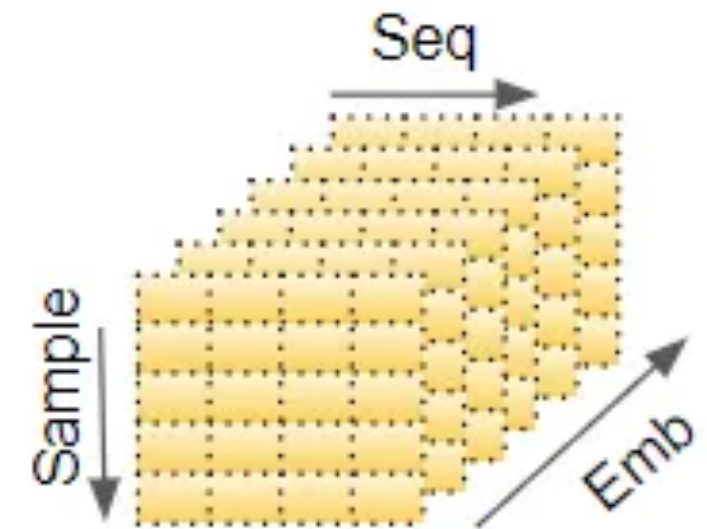
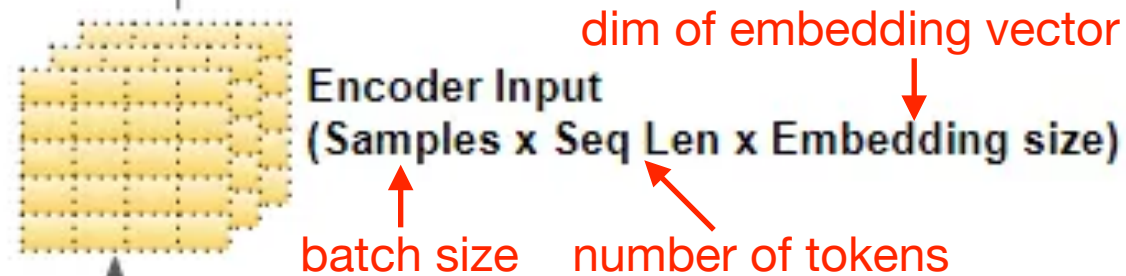
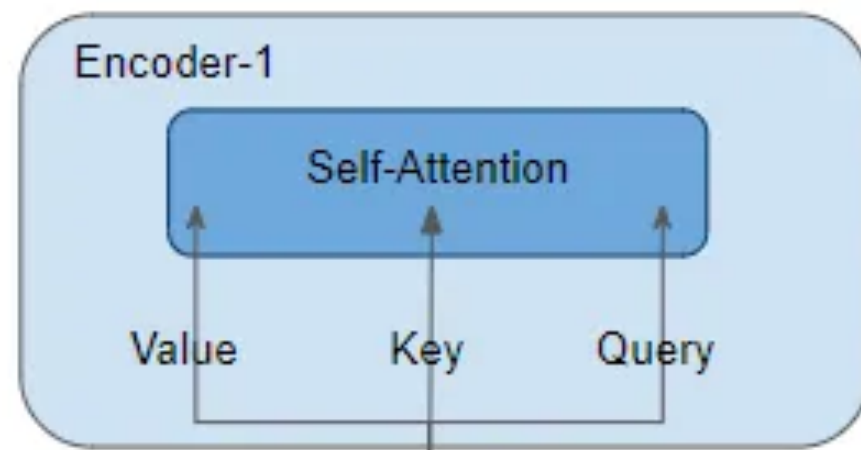
$$P : \begin{aligned} p_{i,2j} &= \sin \left(\frac{i}{10000^{2j/d}} \right), \\ p_{i,2j+1} &= \cos \left(\frac{i}{10000^{2j/d}} \right). \end{aligned}$$

Binary encoding:

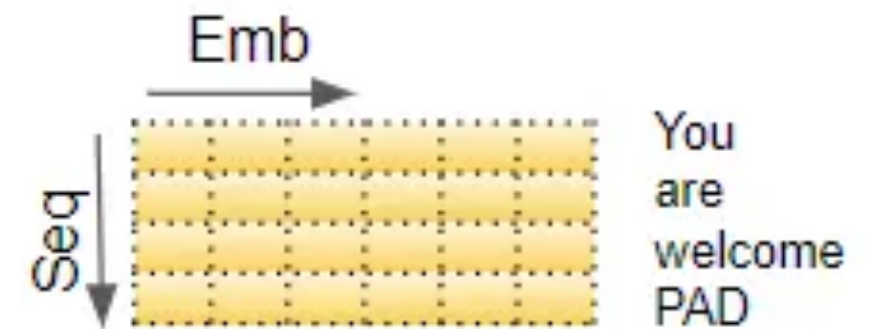
0 in binary is 000
1 in binary is 001
2 in binary is 010
3 in binary is 011
4 in binary is 100
5 in binary is 101
6 in binary is 110
7 in binary is 111



Multi-head Self Attention (Input Layer)



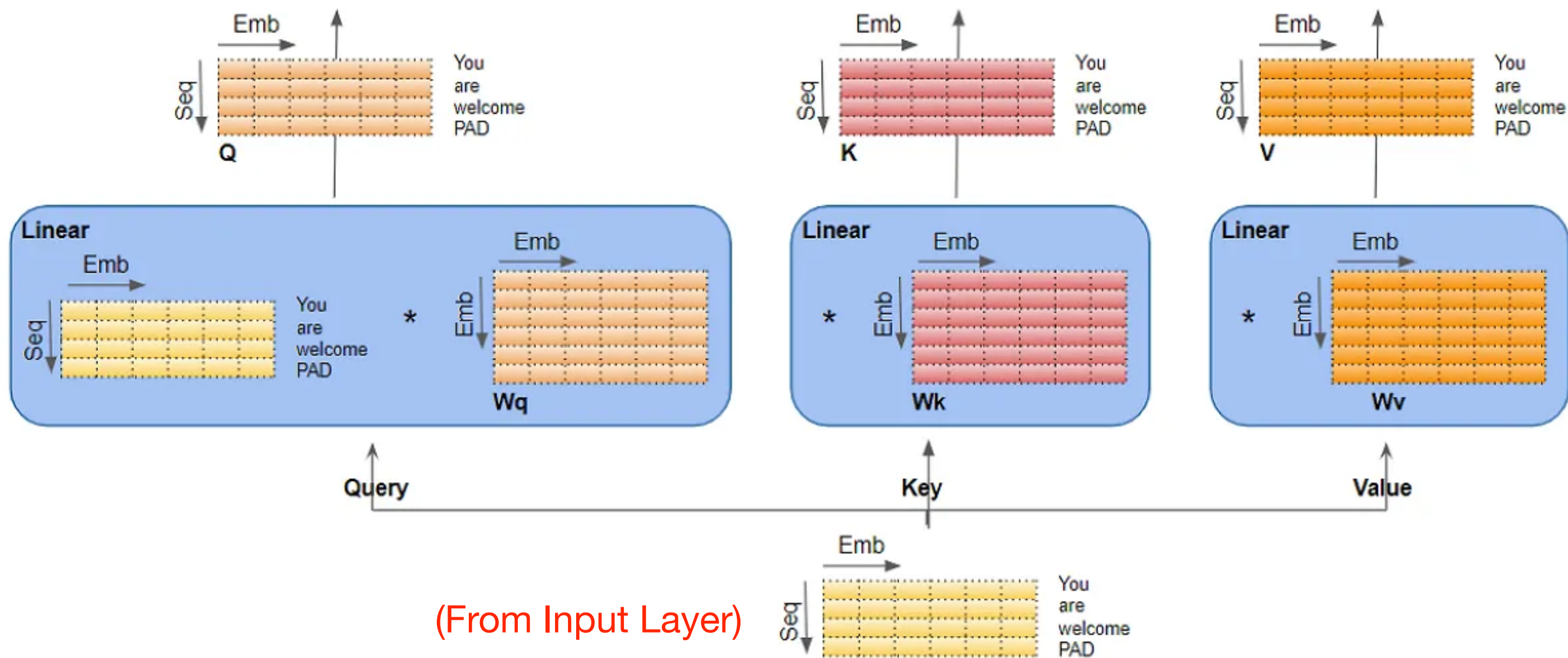
For clarity, consider batch size = 1



Multi-head Self Attention (Linear Layer)

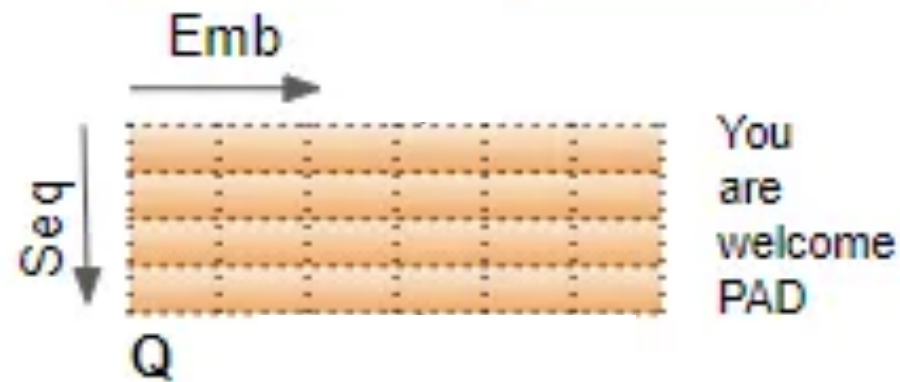
Remarks:

- 1) At this stage, W_q , W_k & W_v haven't mixed tokens yet
- 2) W_q , W_k & W_v mix up the embedding dimensions
- 3) No multi-head yet

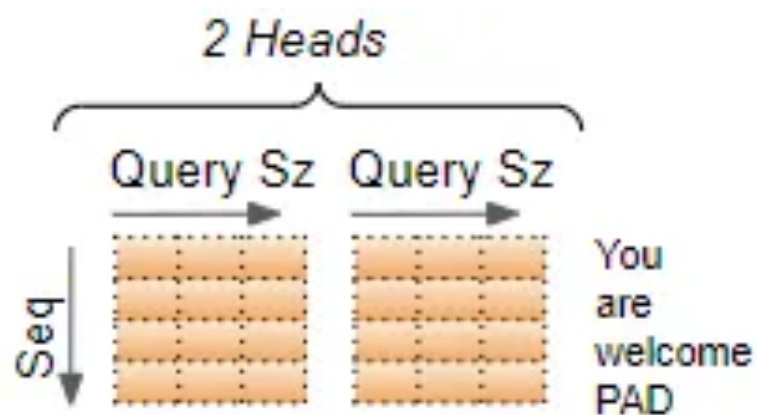


Multi-head Self Attention

A Single Head Attention Score

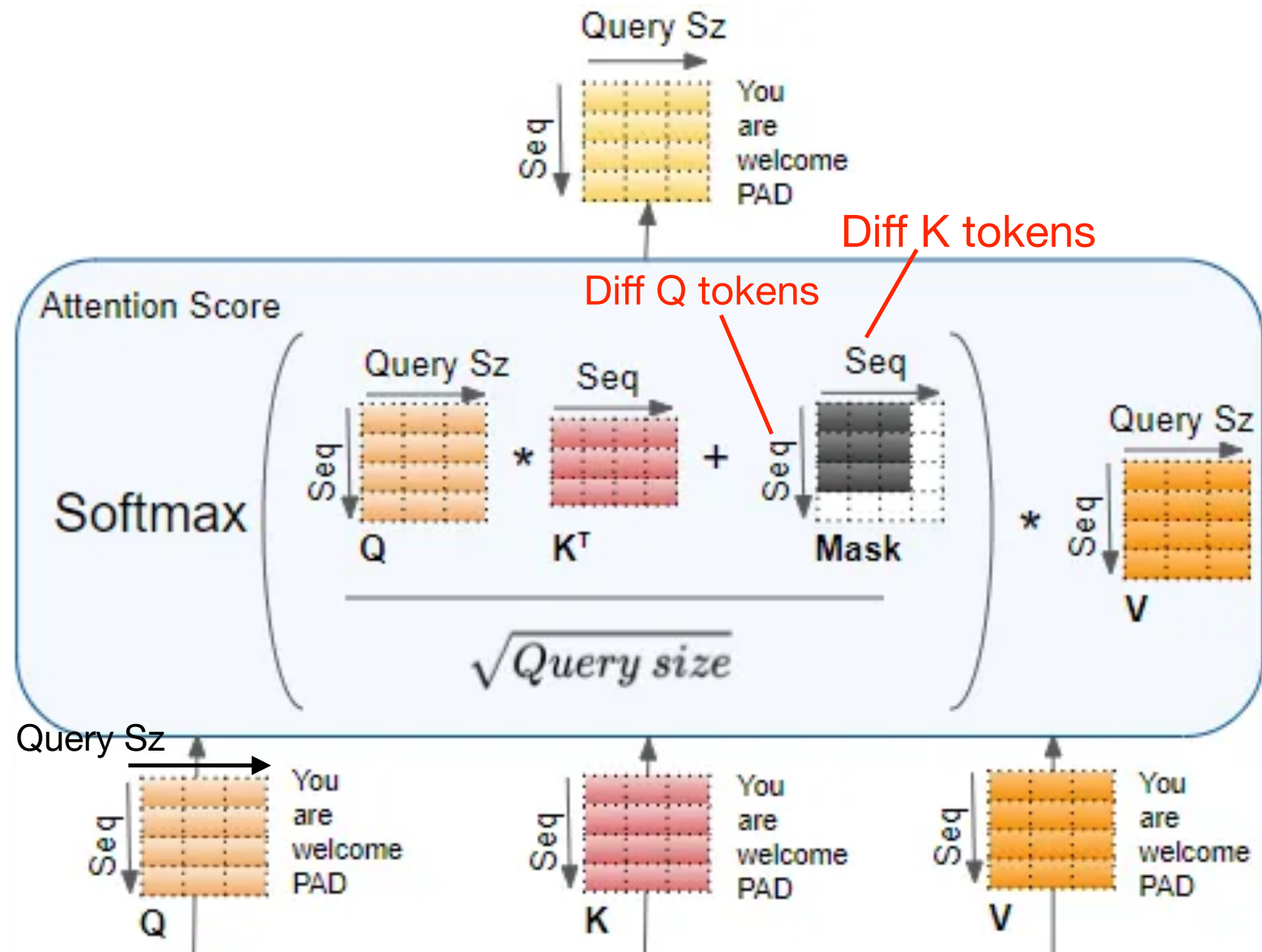


Split



Similar for K & V

Single head



PAD: Padding tokens

Multi-head Self Attention (Summary)

