

A Brief Introduction to the Large Language Model

Longxiao

27/02/2025

Basics of Transformer models

Attention is All You Need

Origin of LLM: NMT task

Neural Machine Translation

We want to calculate or infer:

$$P(y_i | input, y_1, y_2, \dots, y_{i-1})$$

Something important:

Words: Embedding, Encoder-decoder structure

Sequence: RNN -> position embedding

Relation: Attention: cross -> self -> multi-head

<SOS> Wir machen einen Kurs in maschinellern Lernen <EOS>



<SOS> We take a machine learning course <EOS>

Embedding

Embedding: from word to token then to vector

<SOS> Wir machen einen Kurs in maschinellern Lernen <EOS>

Token Token Token Token Token Token Token
15 27 6 89 132 54 352

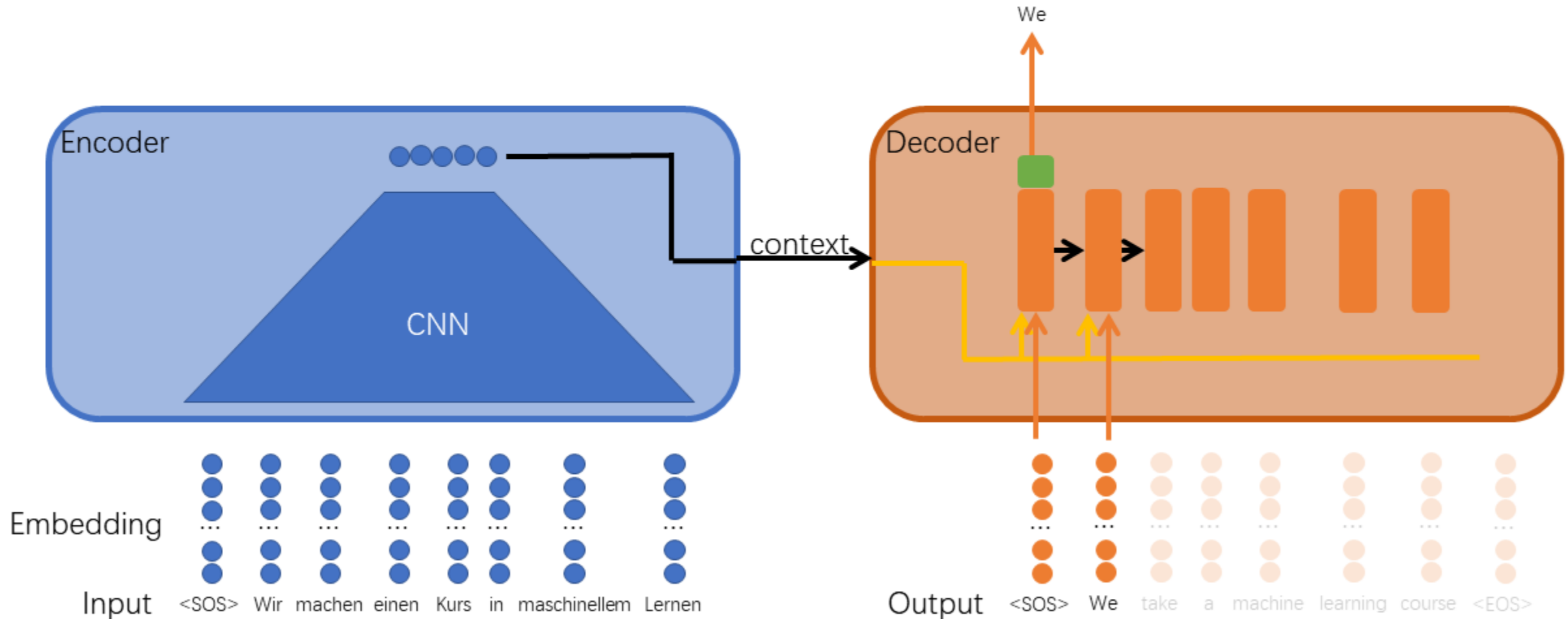
Dimension



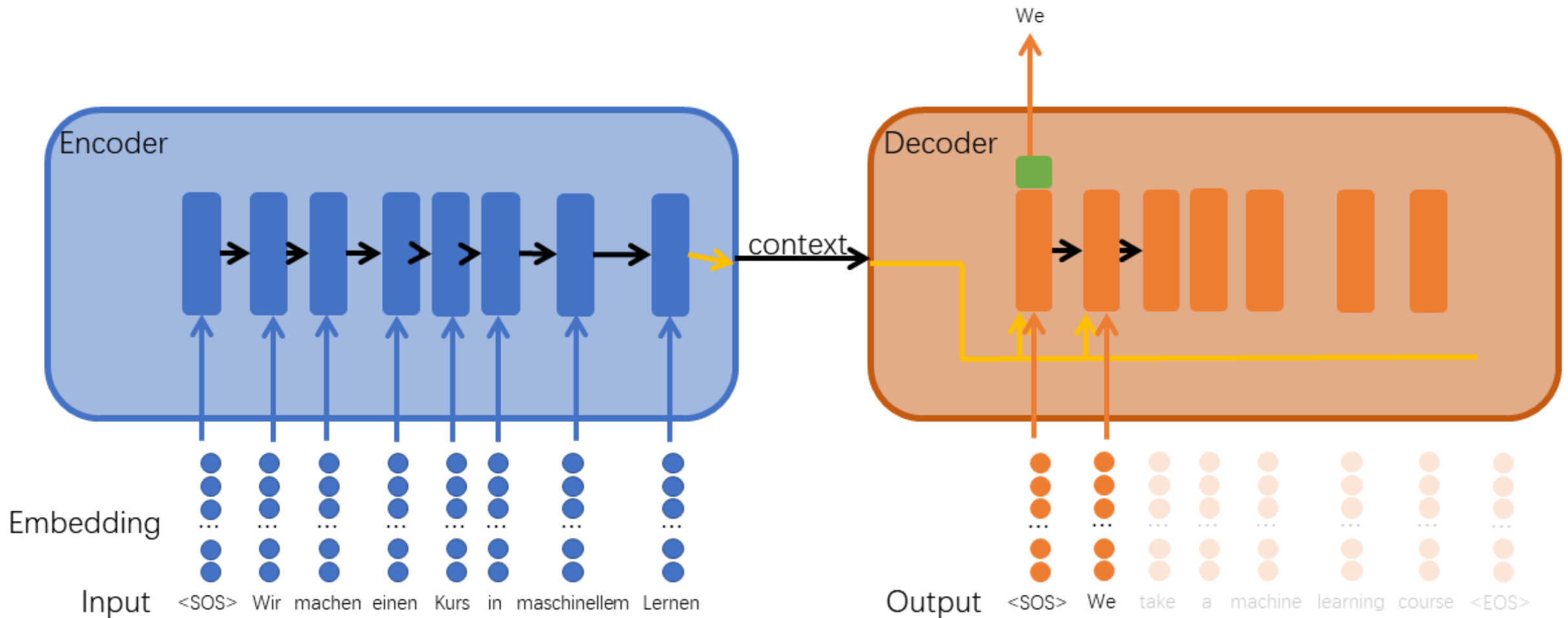
1	2	0	8	1	3	4
2	5	4	6	1	6	7
5	7	6	8	5	7	0
...	...	...	...	...	...	...
2	6	3	8	3	9	1

$= X$

Encoder-decoder



Encoder-decoder



Attention! please

- Cross attention:

Bahdanau Attention

Luong Attention

- Self Attention

- Multi-head Attention

Wir machen einen Kurs in maschinellem Lernen.

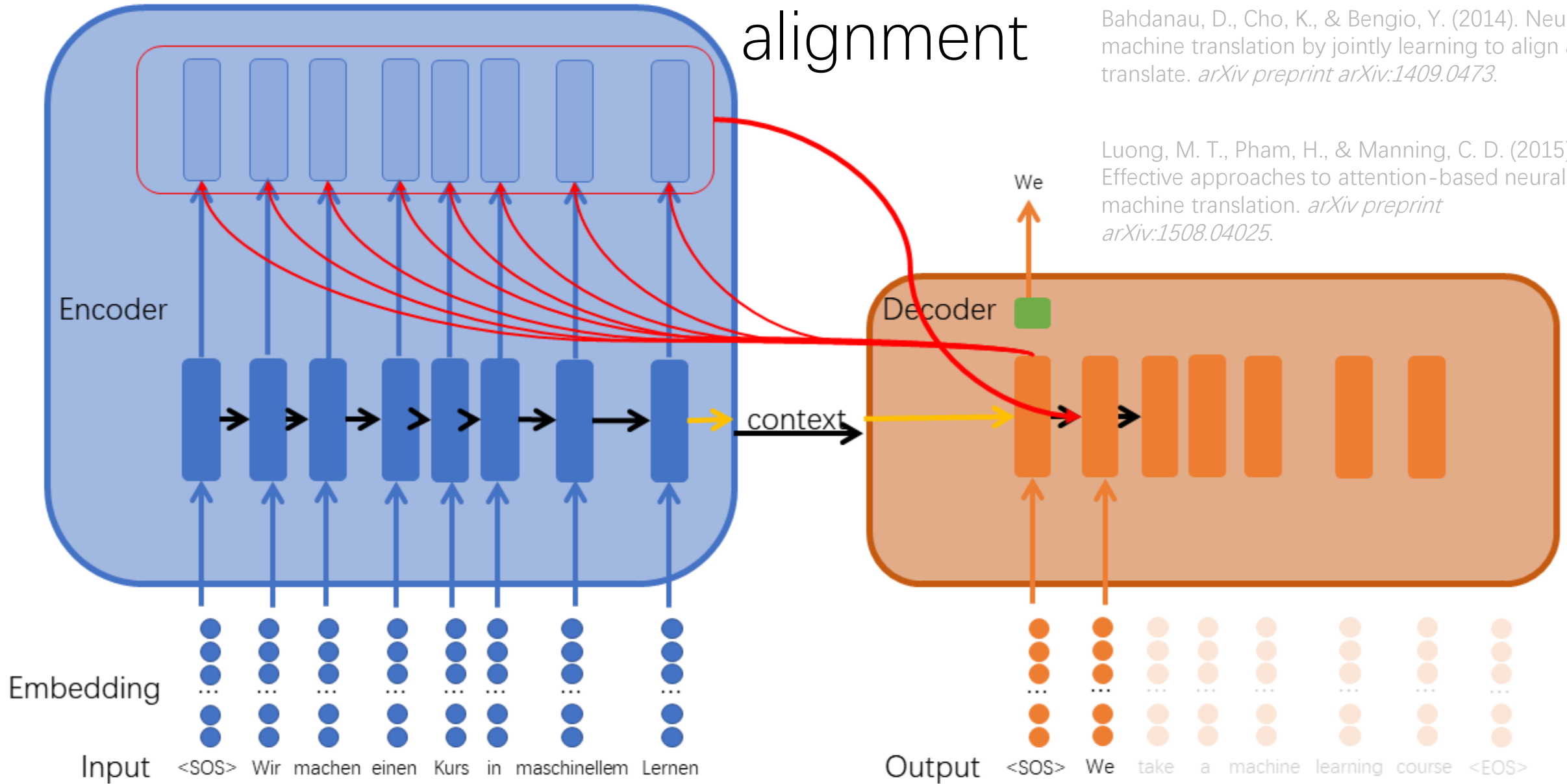


We take a machine learning course.

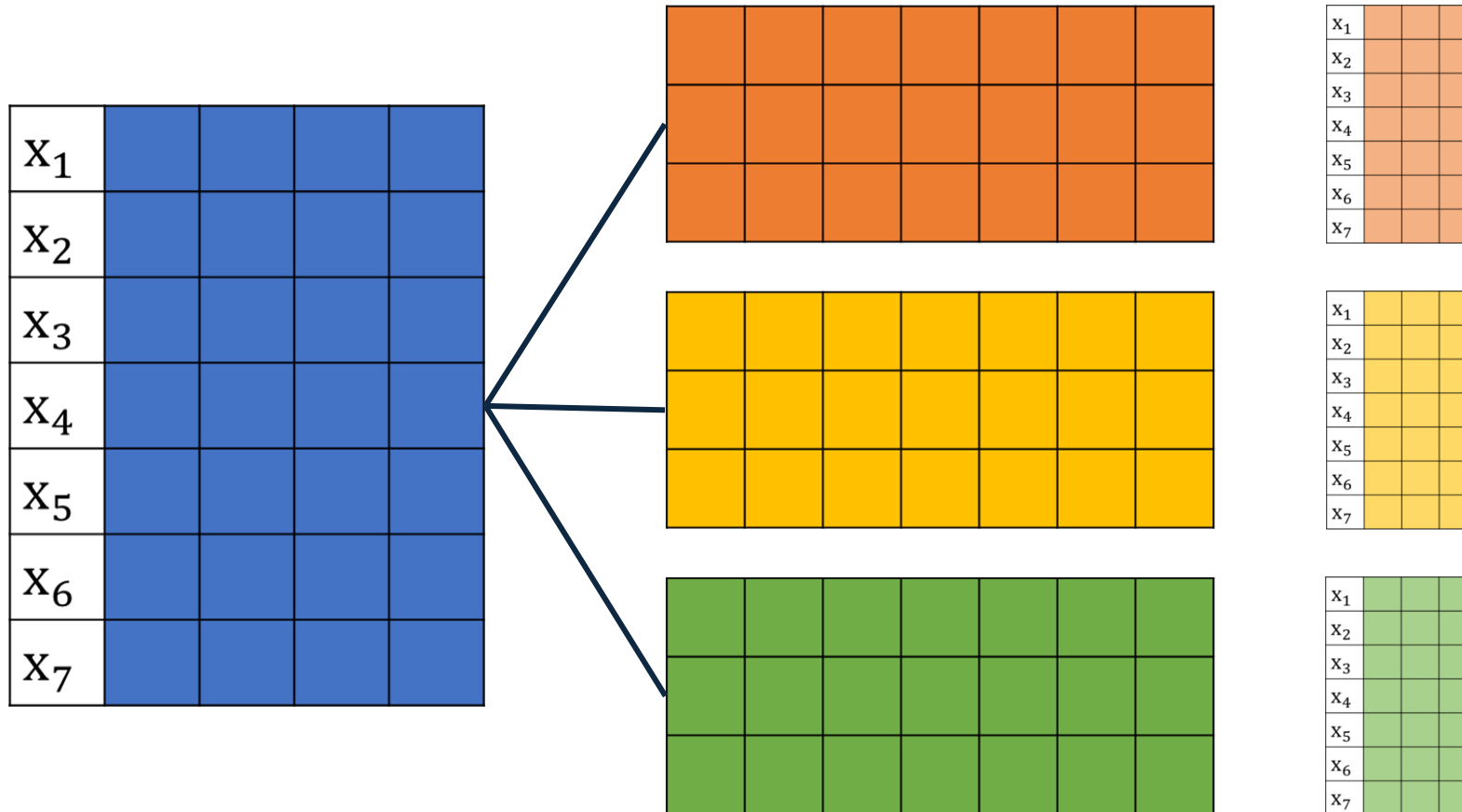
Cross attention: cross alignment

Bahdanau, D., Cho, K., & Bengio, Y. (2014). Neural machine translation by jointly learning to align and translate. *arXiv preprint arXiv:1409.0473*.

Luong, M. T., Pham, H., & Manning, C. D. (2015). Effective approaches to attention-based neural machine translation. *arXiv preprint arXiv:1508.04025*.



Cross attention: cross alignment

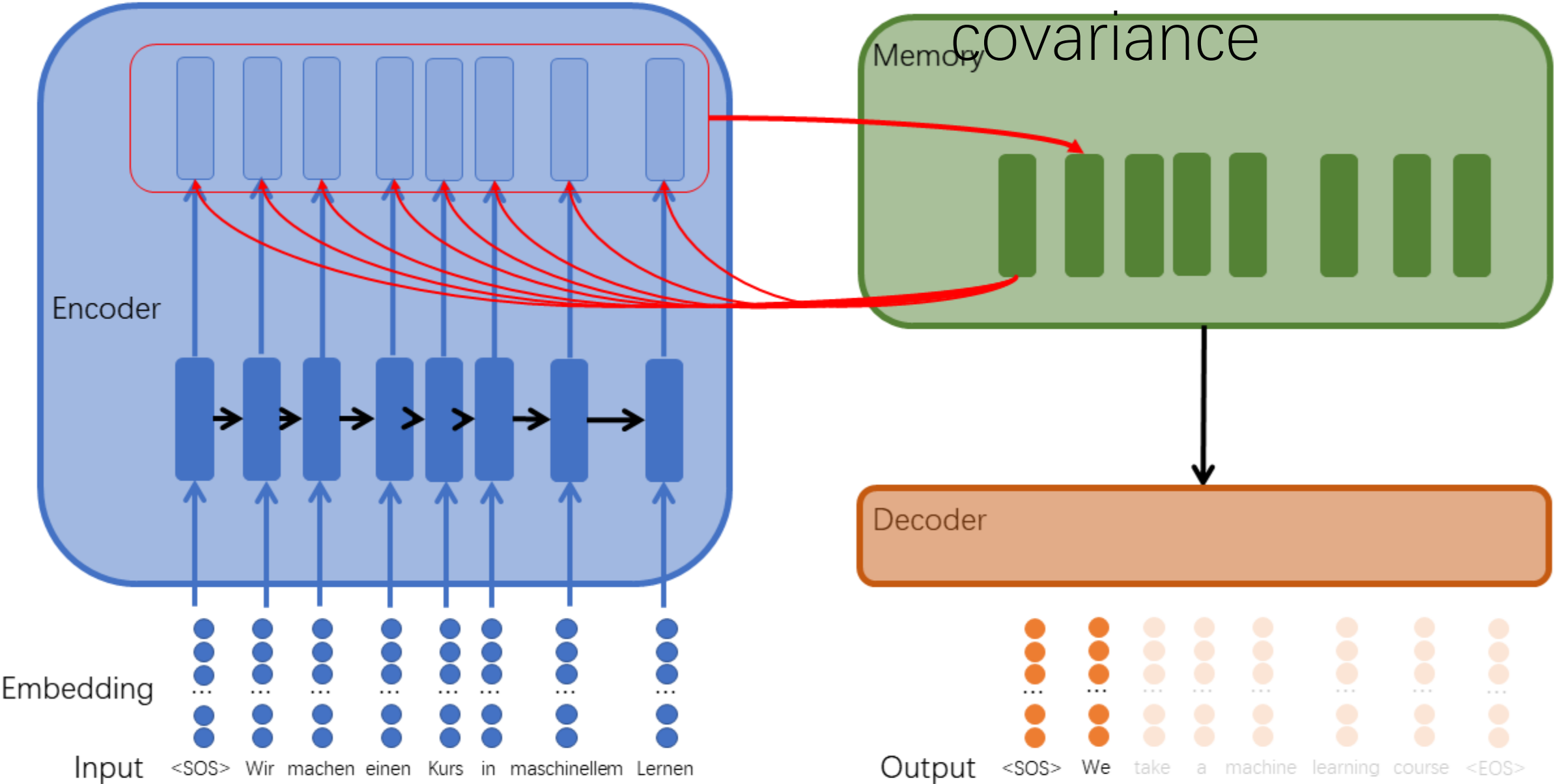


Cross attention: cross alignment

$$\text{Attention}(Q_y, K_x, V_x) \rightarrow Z$$

	We	take	a	machine	learning	course
Wir	0.8	0.18	0.01	0.01	0	0.01
machen	0.15	0.5	0.15	0.02	0.03	0.15
einen	0.01	0.15	0.8	0.03	0.004	0.006
Kurs	0.01	0.15	0.02	0.04	0.009	0.75
in	0.01	0.01	0.01	0	0.007	0.002
maschinellem	0.01	0.005	0.01	0.65	0.25	0.04
Lernen	0.01	0.005	0.01	0.25	0.7	0.04

Self-attention: covariance

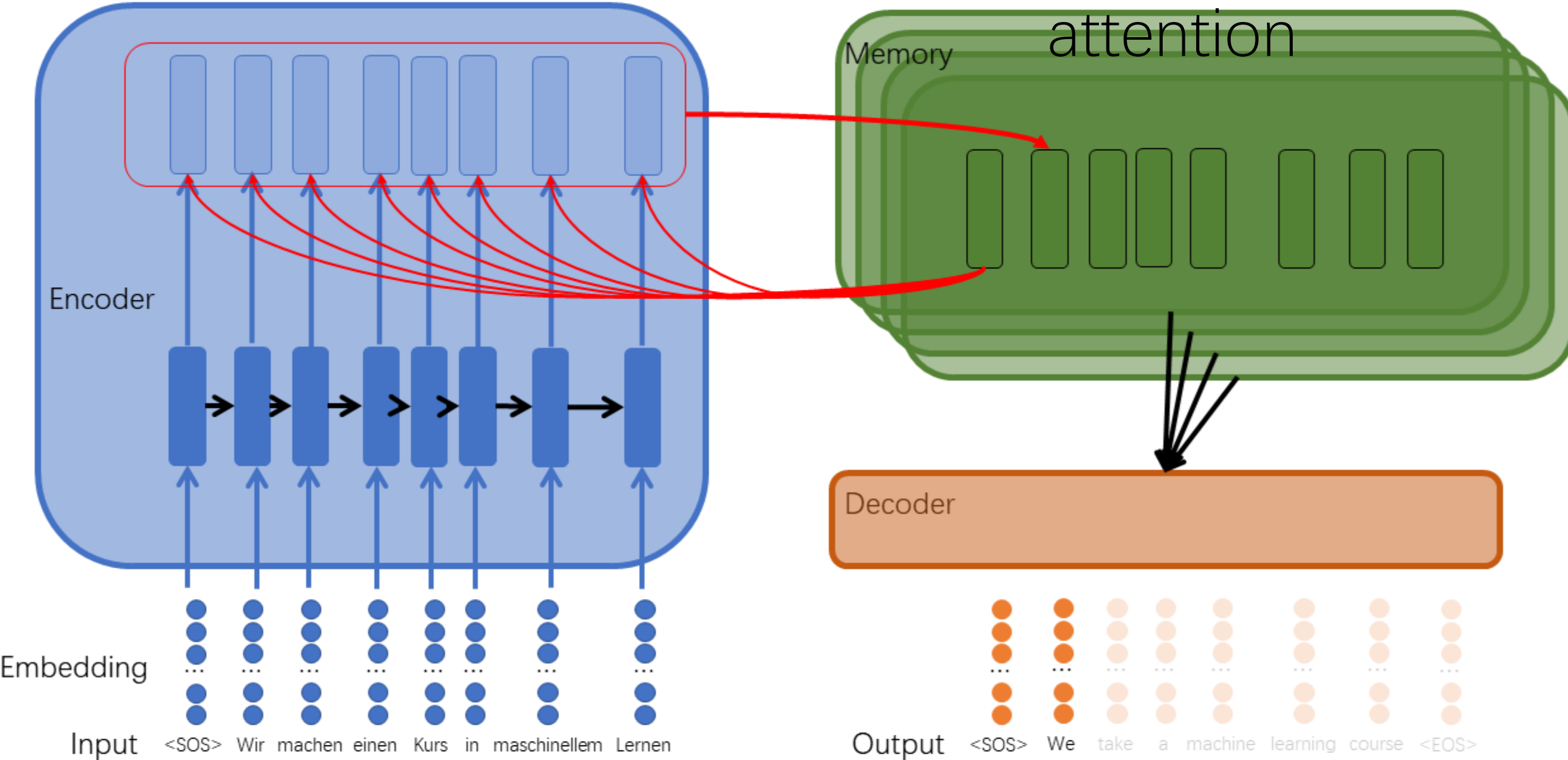


Self-attention: covariance

$$\text{Attention}(Q, K, V) = \text{softmax}\left(\frac{QK^T}{\sqrt{d_k}}\right)V$$

	We	take	a	machine	learning	course
We	0.6	0.35	0.02	0.01	0.01	0.01
take	0.35	0.5	0.4	0.01	0.02	0.2
a	0.02	0.4	0.6	0.01	0.02	0.04
machine	0.01	0.01	0.01	0.7	0.25	0.02
learning	0.01	0.02	0.02	0.25	0.8	0.05
course	0.01	0.2	0.4	0.02	0.05	0.65

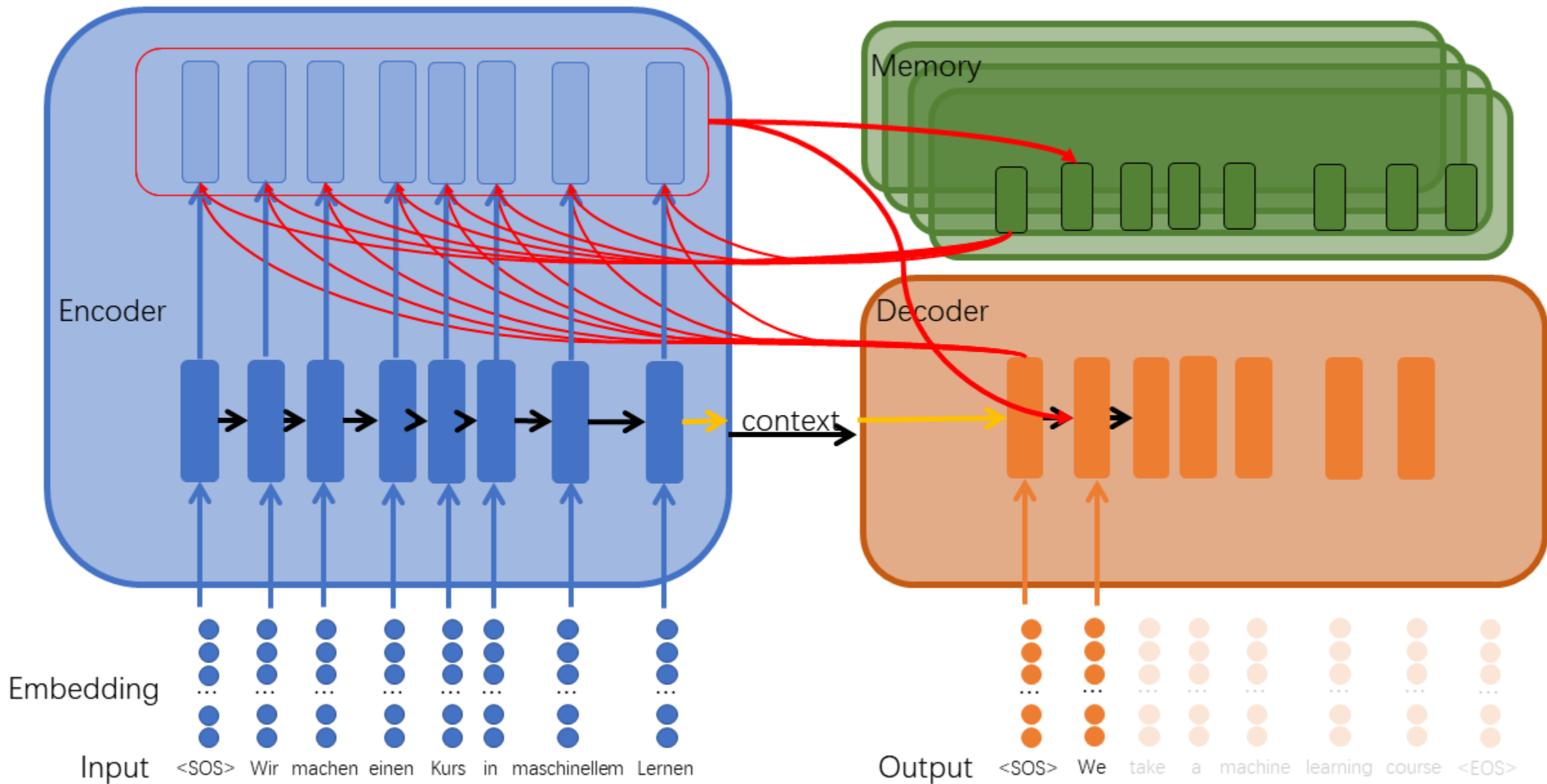
Multi-head attention



Multi-head attention: different informations

	We	take	a	machine	learning	course
We	0.5	0.2	0.02	0.01	0.01	0.35
take	0.25	0.5	0.02	0.01	0.02	0.2
a	0.02	0.02	0.3	0.01	0.02	0.04
machine	0.01	0.01	0.01	0.45	0.25	0.02
learning	0.01	0.02	0.02	0.35	0.5	0.22
course	0.4	0.2	0.4	0.02	0.3	0.65

	We	take	a	machine	learning	course
We	0.6	0.35	0.02	0.01	0.01	0.01
take	0.35	0.5	0.1	0.01	0.02	0.2
a	0.02	0.1	0.6	0.01	0.02	0.5
machine	0.01	0.01	0.01	0.7	0.25	0.02
learning	0.01	0.02	0.02	0.25	0.8	0.05
course	0.01	0.2	0.4	0.02	0.05	0.65



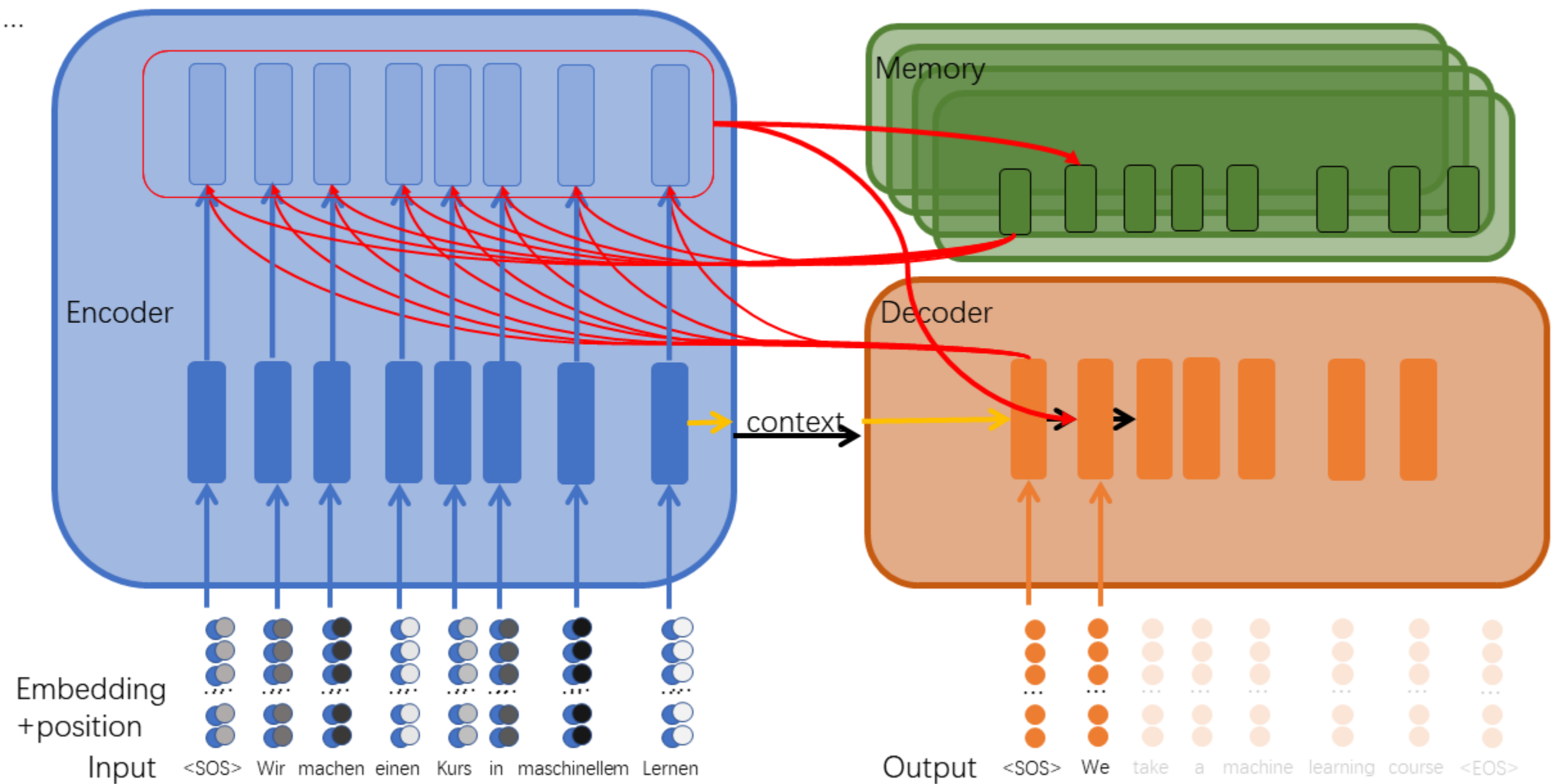
Attention! please

- Cross attention - Self attention - Multi-head Attention

It's time to get rid of RNN!

(sequential data...)

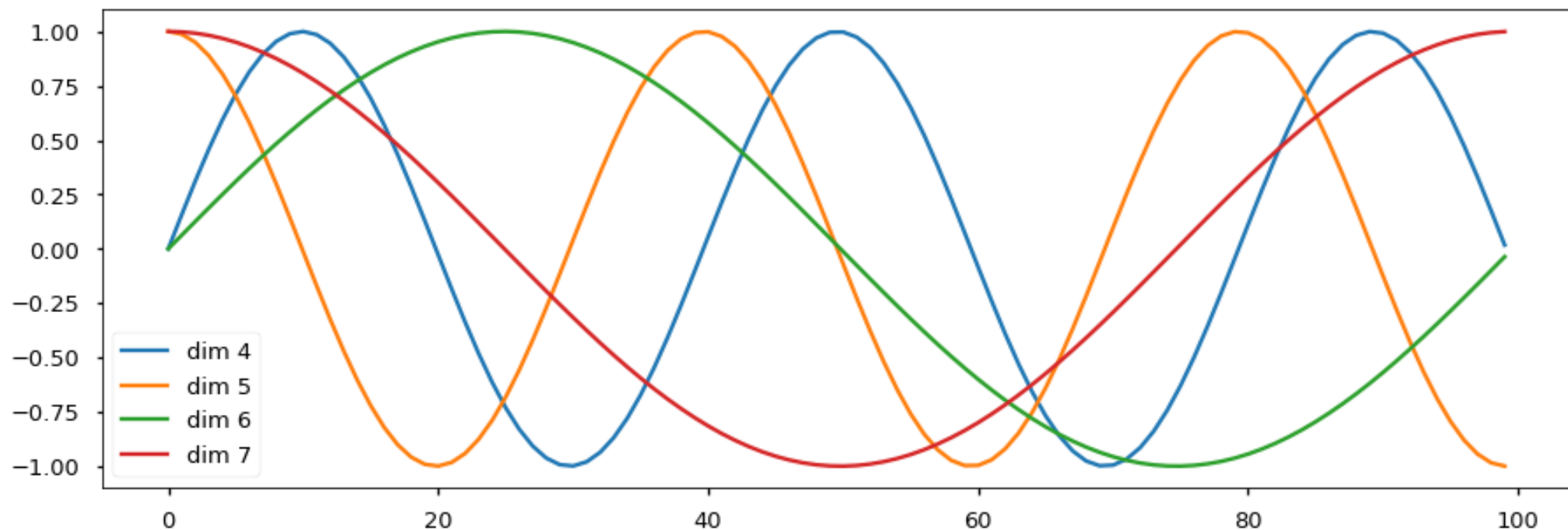
...



Position embedding

$$PE_{(pos,2i)} = \sin(pos/10000^{2i/d})$$

$$PE_{(pos,2i+1)} = \cos(pos/10000^{2i/d})$$

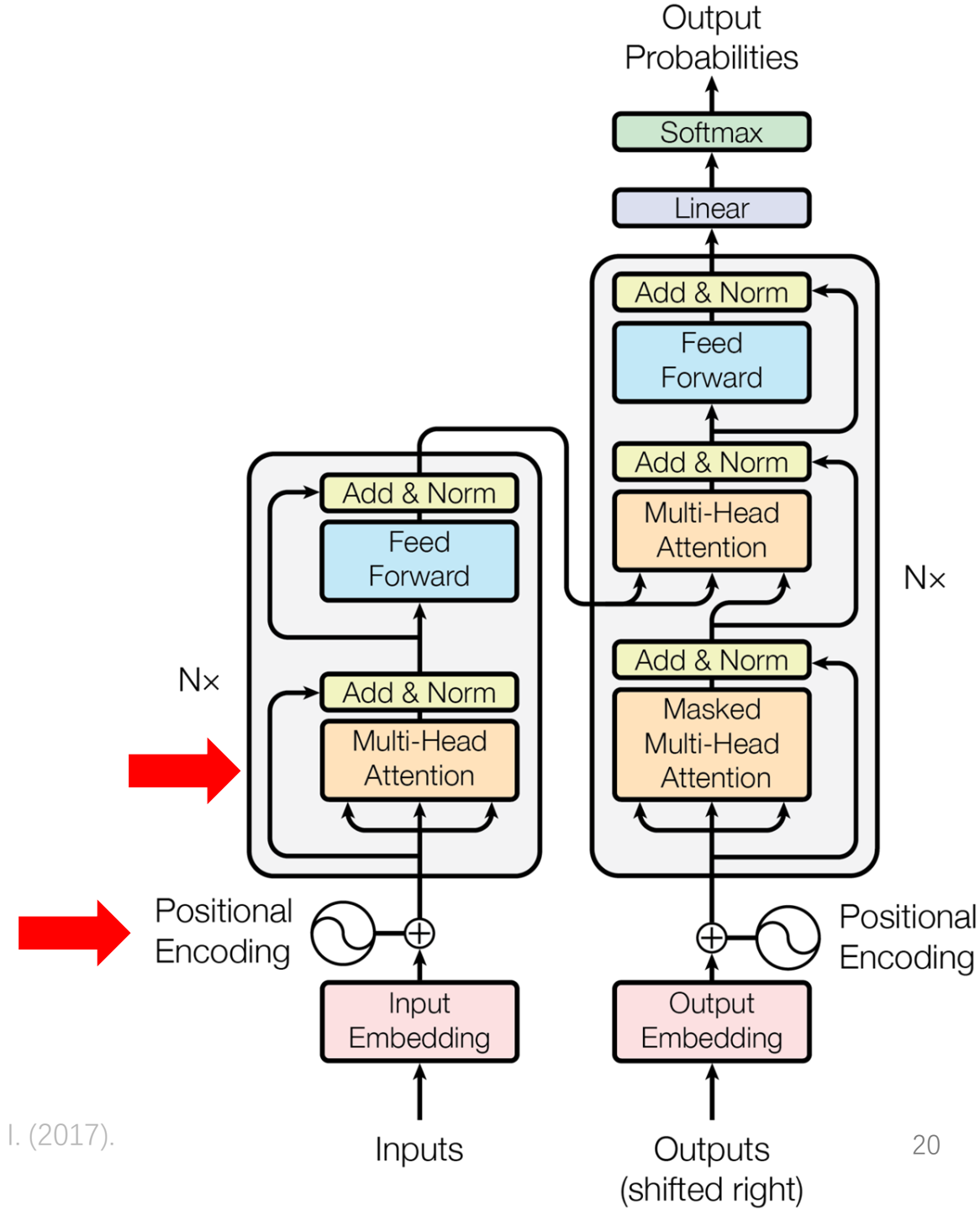


Attention is all you need

Encoder:

Embedding and position embedding

Self-attention / multi-head



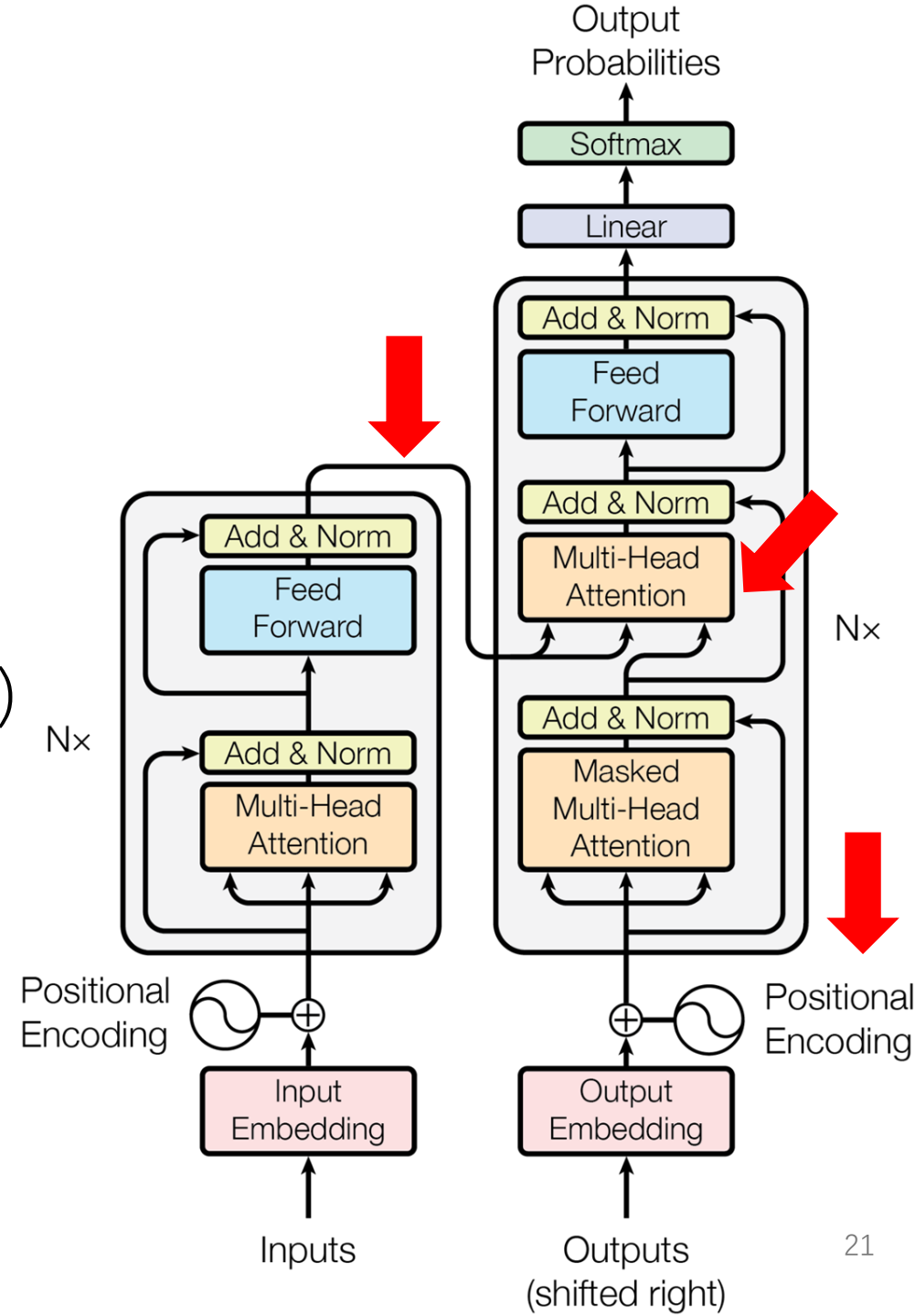
Attention is all you need

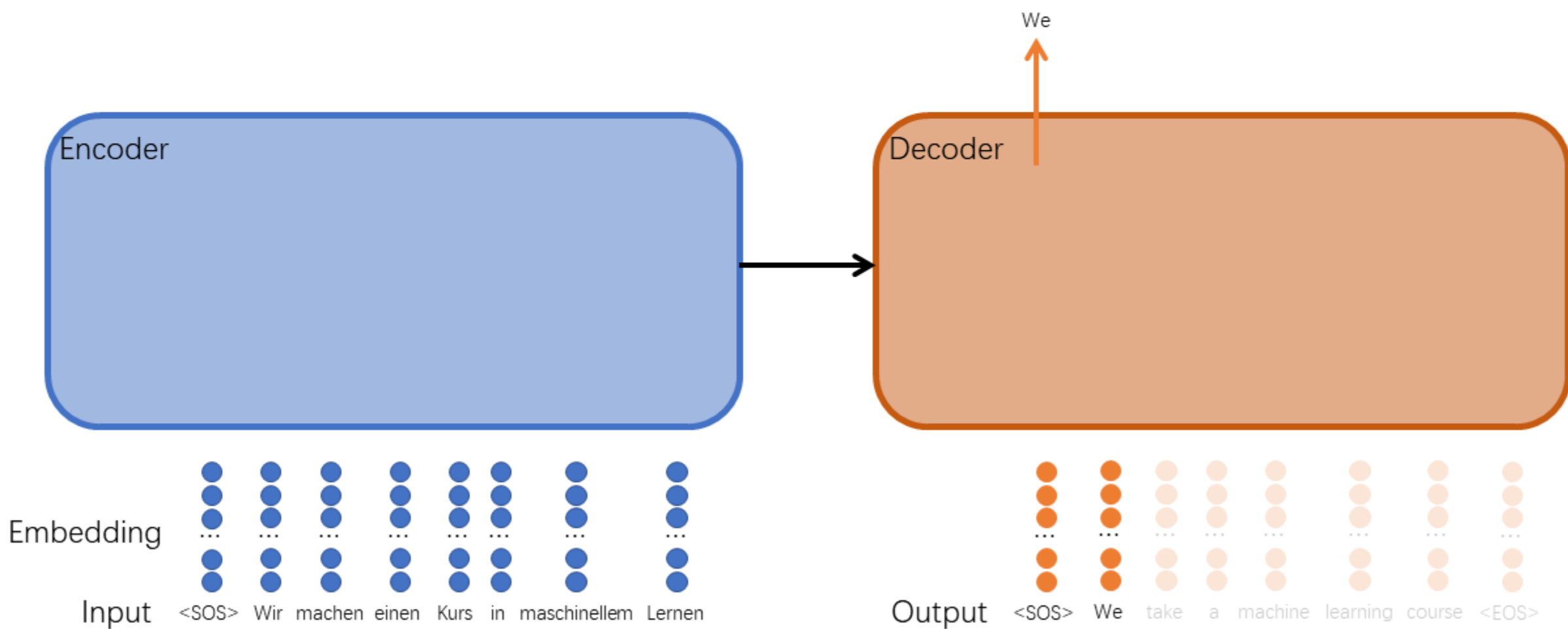
Decoder:

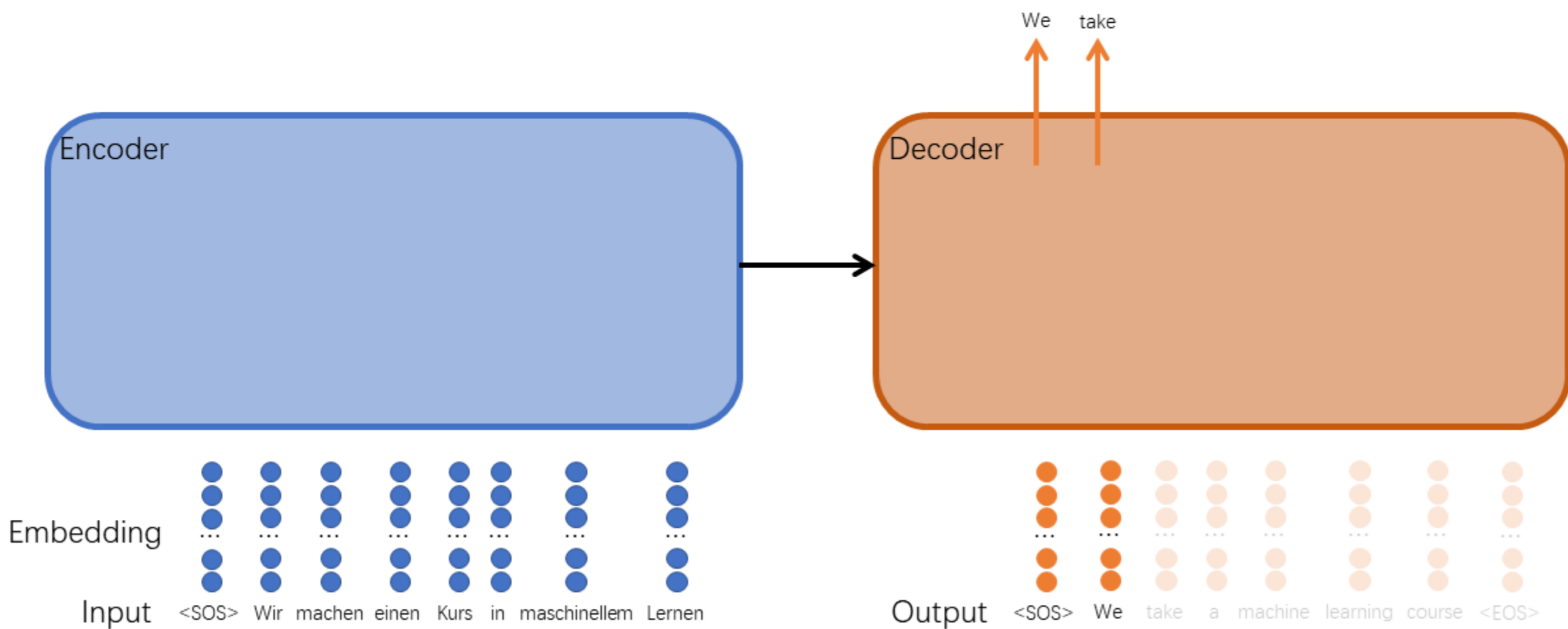
Embedding and position embedding

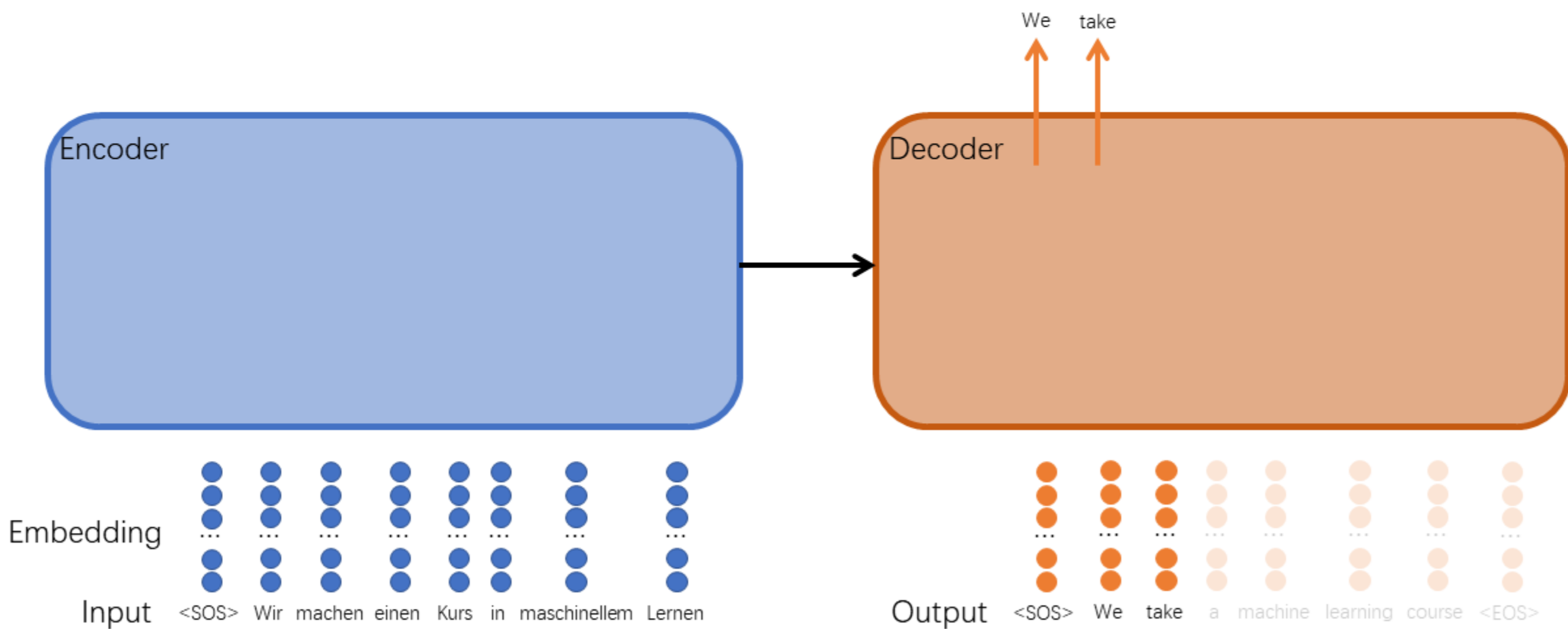
Self-attention / multi-head/ mask(in training)

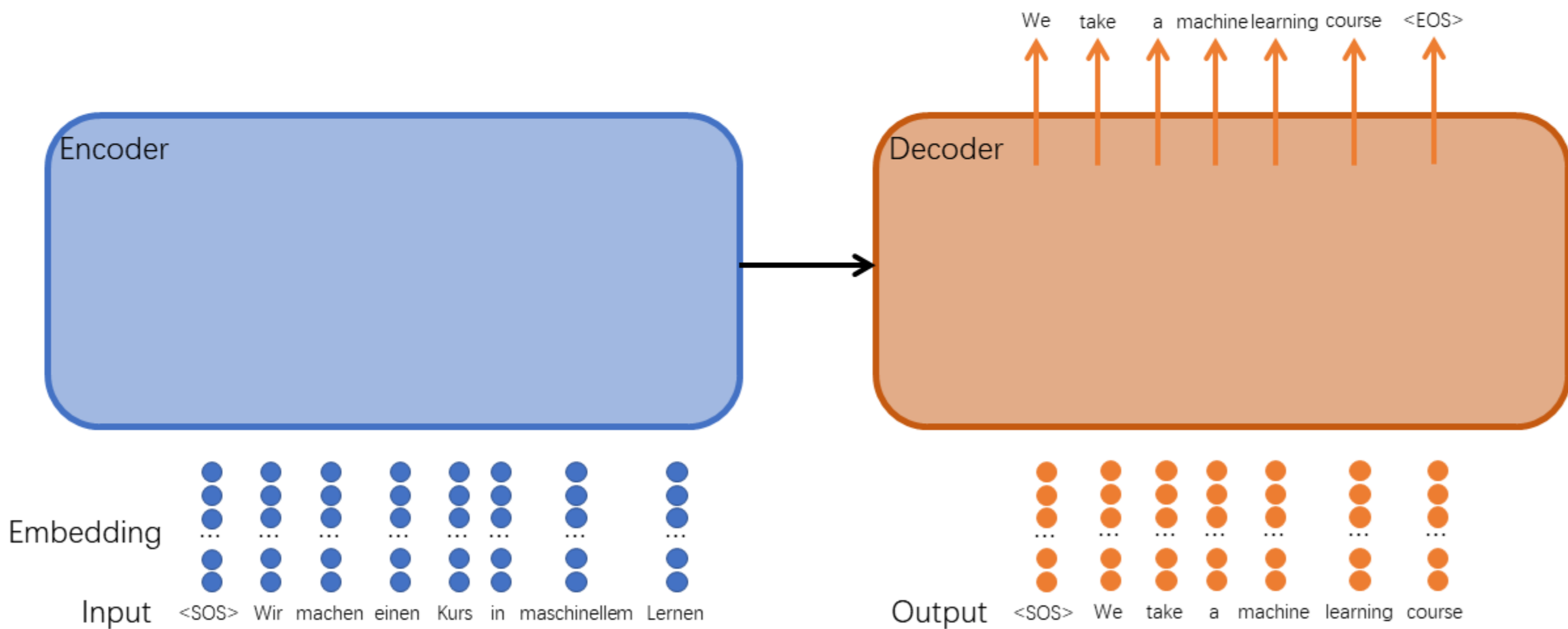
Cross-attention











The Annotated Transformer

<http://nlp.seas.harvard.edu/annotated-transformer/>

GitHub: [harvardnlp/annotated-transformer: An annotated implementation of the Transformer paper.](https://github.com/harvardnlp/annotated-transformer)