

Introduction to Machine Learning

Attention Mechanism and Transformers

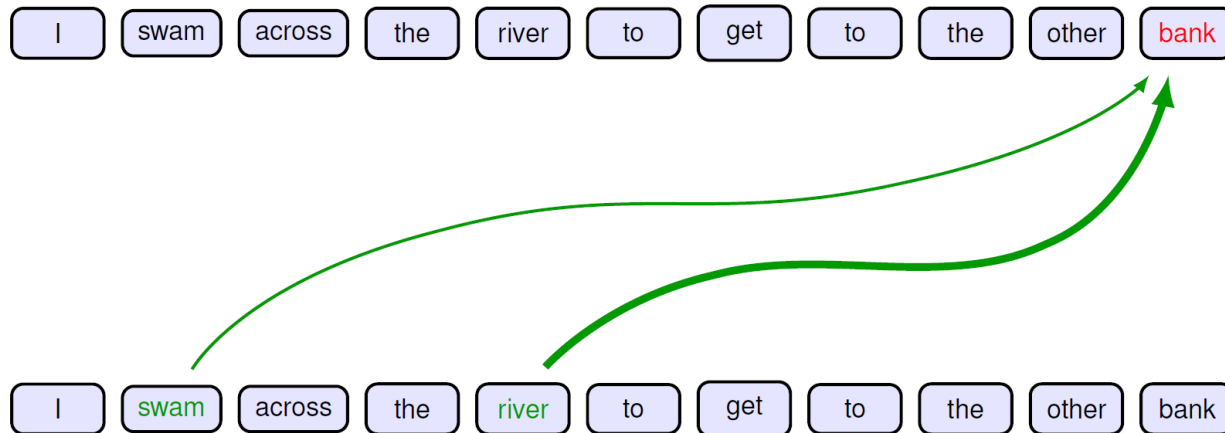
Prof. Subrahmanya Swamy Peruru
Department of Electrical Engineering
IIT Kanpur

Attention

I swam across the river to get to the other **bank**.

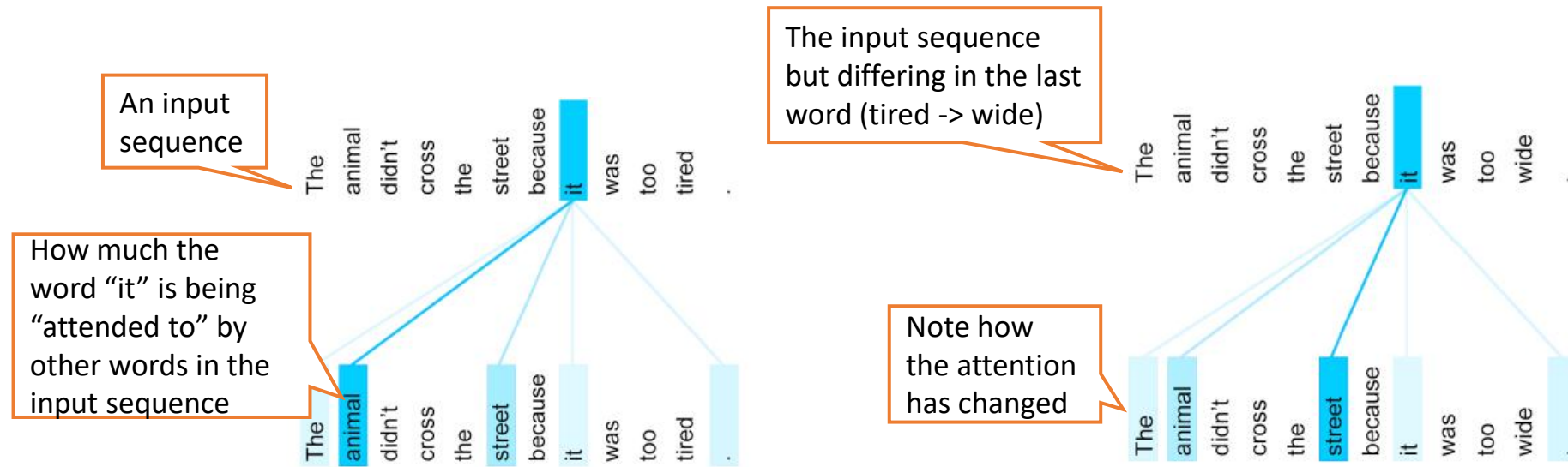
I walked across the road to get cash from the **bank**.

- Word embedding of the word '**bank**' is the same in both the sentences. But they have different meaning.
- The correct interpretation can be understood if we pay attention to the other words in the sentence.



Attention

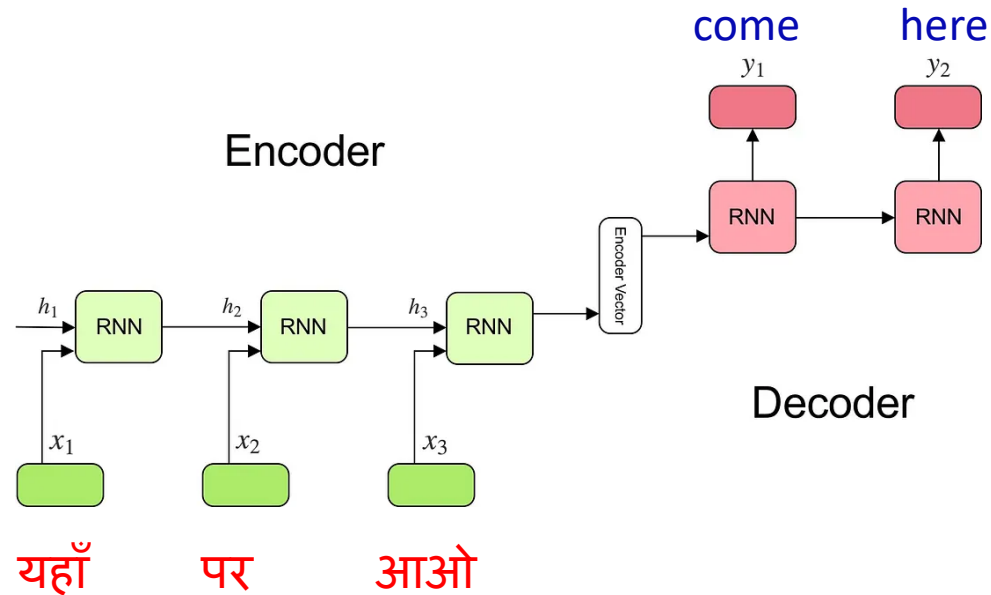
- With self-attention, each token x_n can “attend to” all other tokens of the same sequence when computing this token’s embedding h_n



- Attention helps capture the context better and in a much more “global” manner
 - “Global”: Long ranges captures and in both directions (previous and ahead)

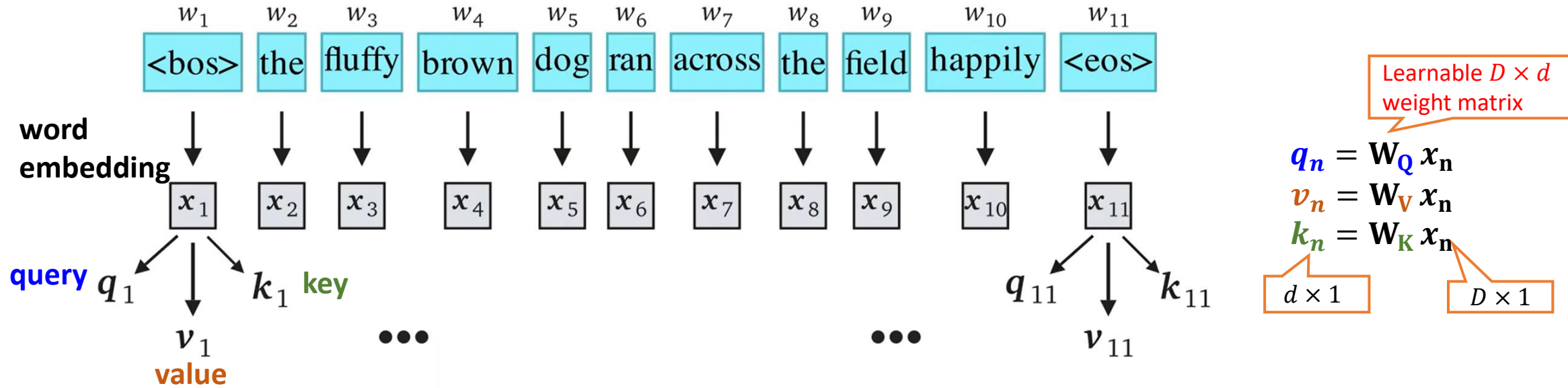
Recap: RNNs

- RNNs are used when each input or output or both are **sequences of tokens**

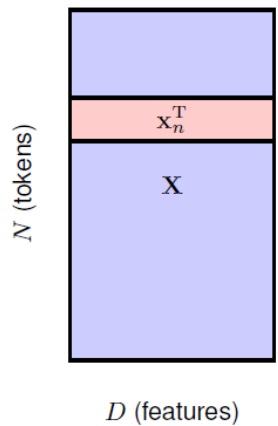


- **Slow processing:** Can't compute h_t before computing h_{t-1}
- **Information bottleneck:**
 - Hidden state h_t is supposed to remember everything up to time $t - 1$.
 - Variants such as LSTM, GRU, etc mitigate this issue to some extent

Attention



Can be computed parallelly for all N tokens using matrix multiplications



$$Q = XW_Q$$

Row n is q_n

$$K = XW_K$$

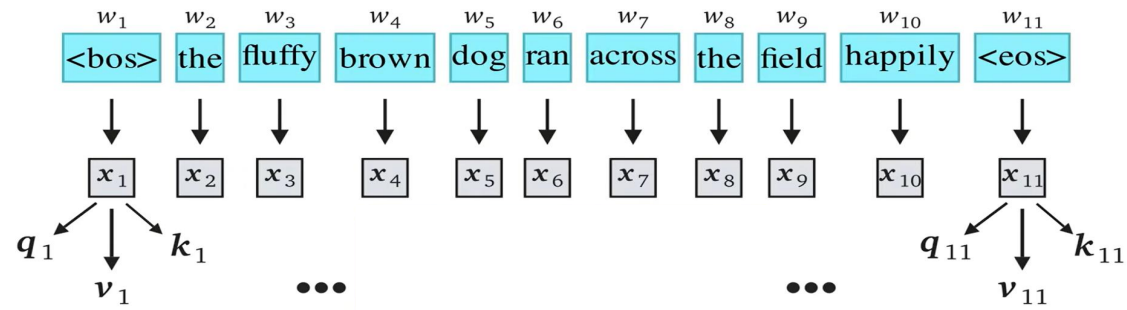
$$V = XW_V$$

Self-Attention

- Attention scores

How much token i attends to token n

$$\alpha_{n,i} = \frac{\exp(\mathbf{q}_n^\top \mathbf{k}_i)}{\sum_{j=1}^N \exp(\mathbf{q}_n^\top \mathbf{k}_j)}$$



- Contextual Embedding (Hidden state) for x_n is

$$\mathbf{h}_n = \sum_{i=1}^N \alpha_{n,i} \mathbf{v}_i$$

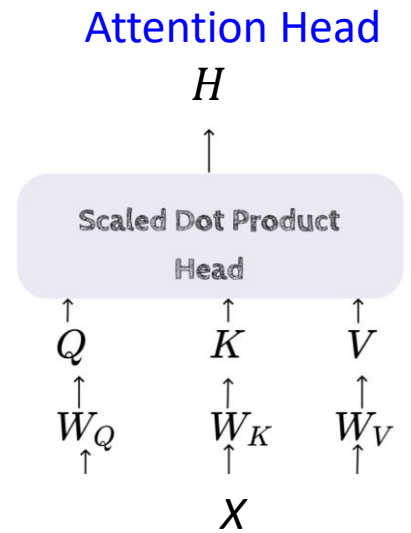
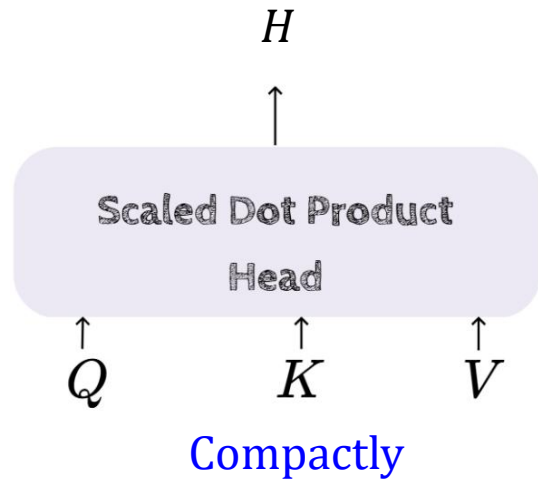
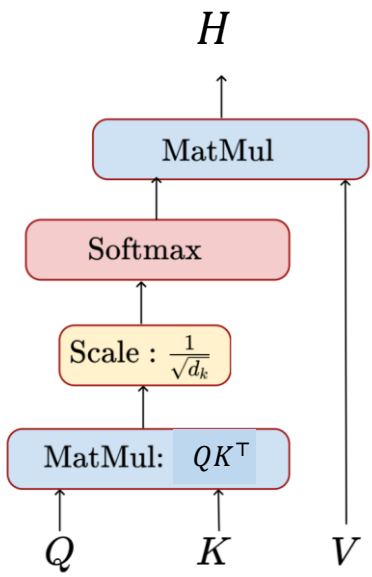
- Matrix form

$$\mathbf{H} = \text{softmax}(\mathbf{Q}\mathbf{K}^\top) \mathbf{V}$$

Scaled Dot Product

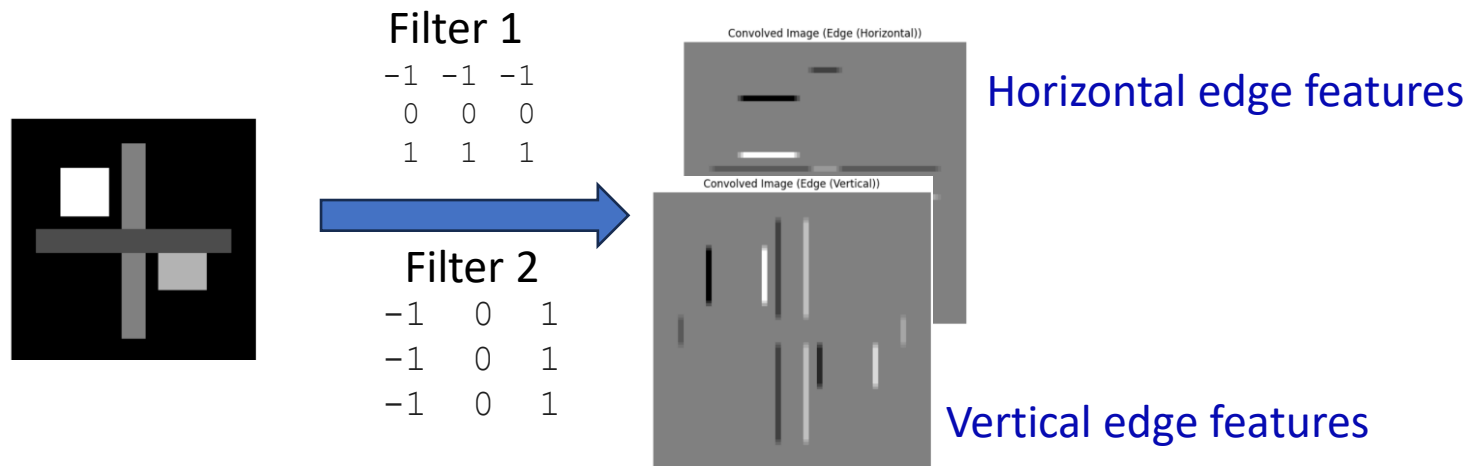
$$\mathbf{H} = \text{softmax}\left(\frac{\mathbf{Q}\mathbf{K}^\top}{\sqrt{d}}\right) \mathbf{V}$$

Dividing by $\sqrt{d}$ ensures variance of the dot product is 1

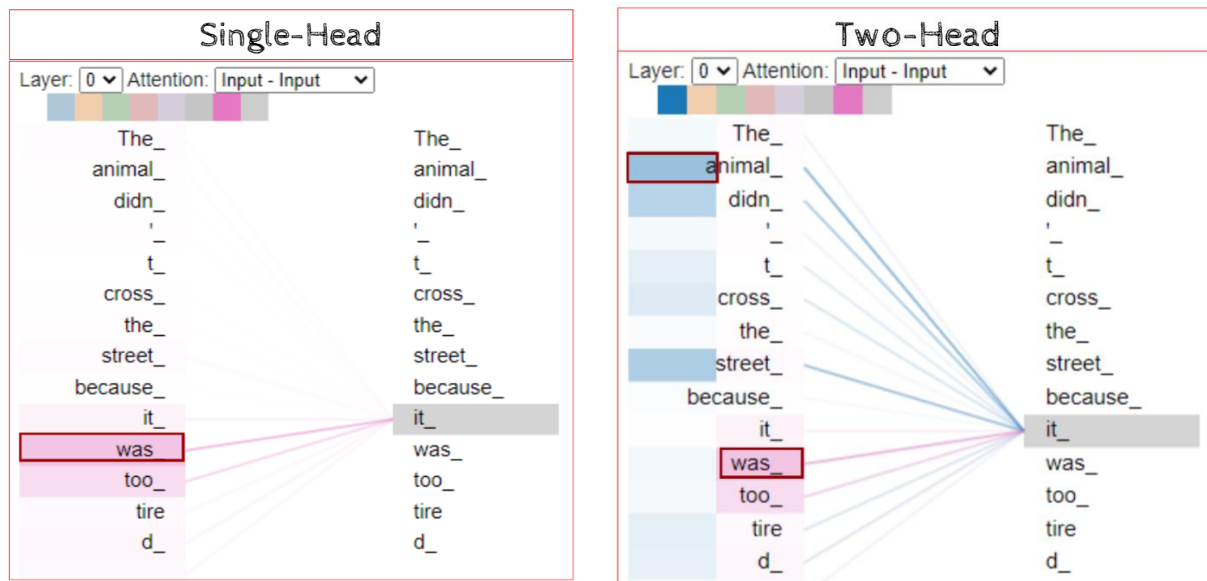


Multi-Head Attention

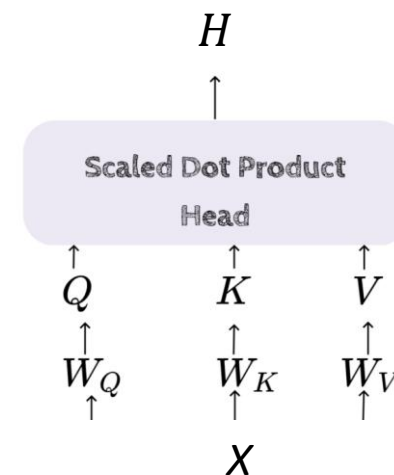
Recall from CNNs that **multiple filters** allow us to extract **different kinds of features**



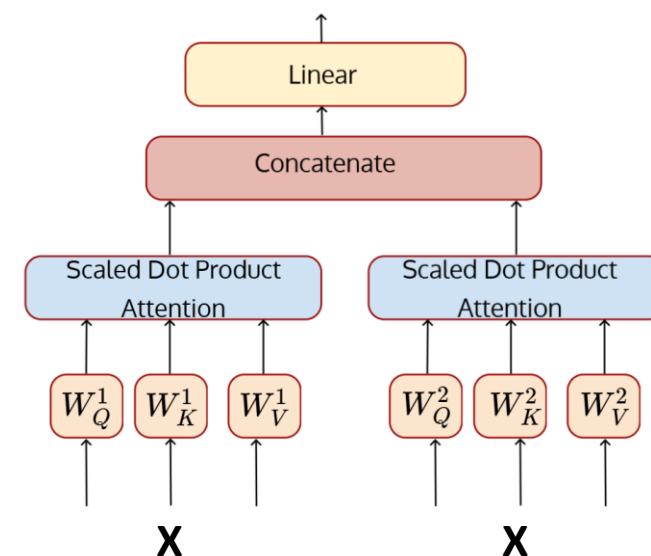
Similarly **multiple attention heads** allow us to learn **different kinds of embeddings**



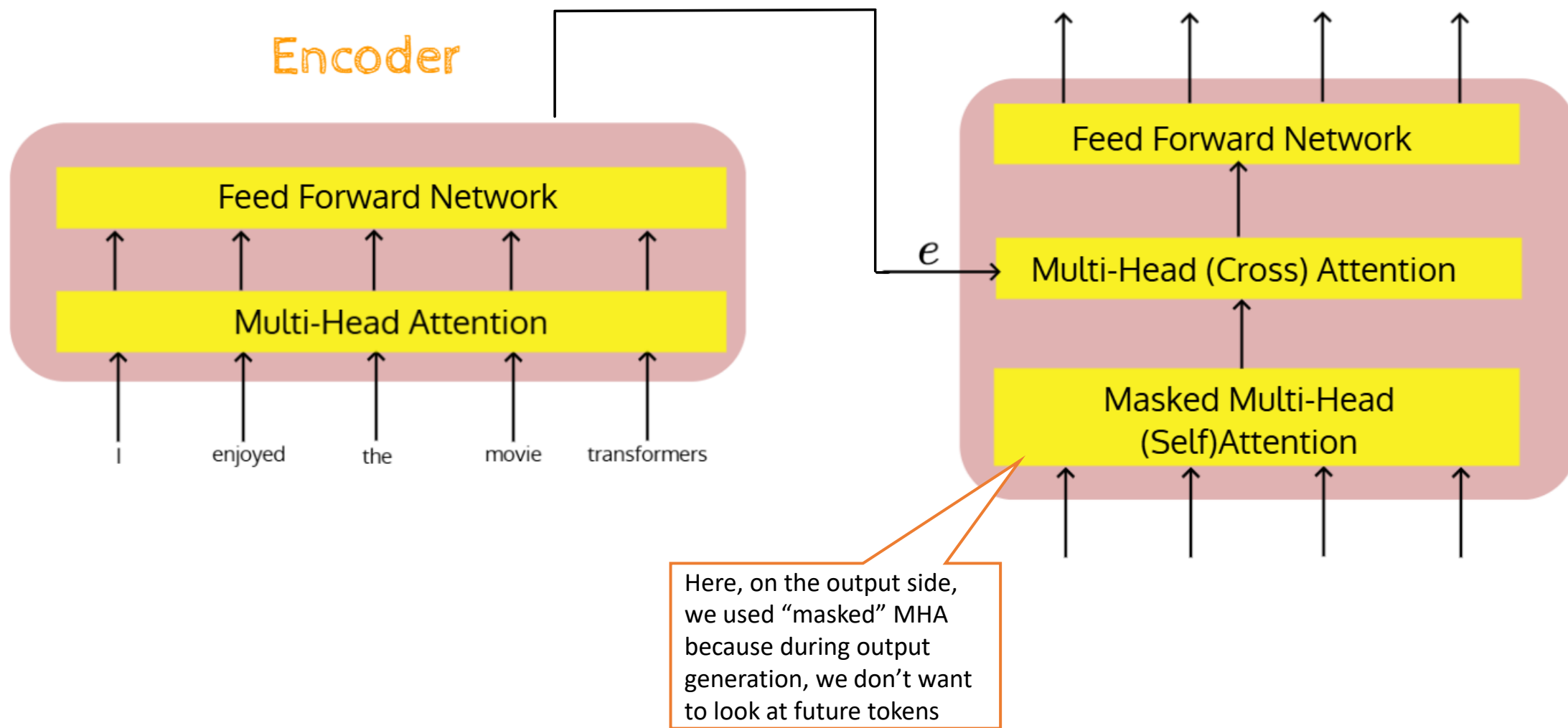
Single Attention Head



Multi Attention Head

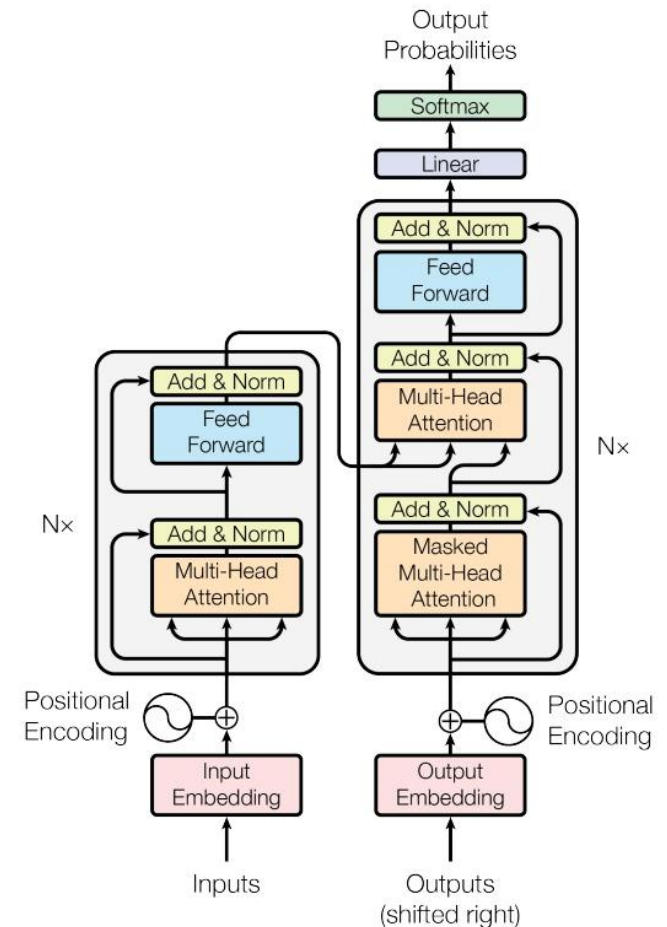


Encoder-Decoder Architecture



Transformers Architecture

- Transformers use the idea of attention
 - “self-attention” over all input tokens
 - Masked “self-attention” over each output token and previous tokens
 - “cross-attention” between output tokens and input tokens
- Transformer also compute embeddings of all tokens in parallel
- Transformers are based on the following key ideas*
 - “Self-attention” and “cross-attention” for computing the hidden states
 - Positional encoding (to encode relative distances between words)
 - Residual connections
- Attention helps capture the context better and in a much more “global” manner in sequence data



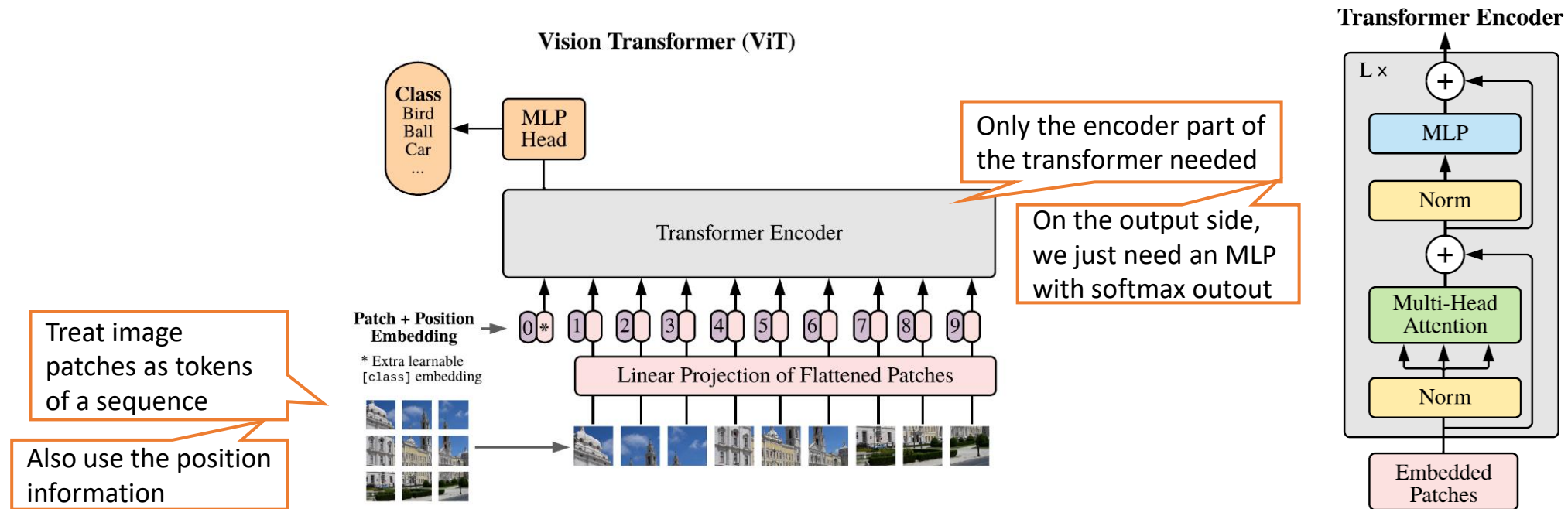
Popular Transformer Variants: BERT and GPT

10

- The standard transformer architecture is an encoder-decoder model
- Some models use just the encoder or the decoder of the transformer
- **BERT** (Bidirectional Encoder Representations from Transformers)
 - Basic BERT can be learned to encode token sequences
- **GPT** (Generative Pretrained Transformer)
 - Basic GPT can be used to generate token sequences similar to its training data

Transformers for Images: ViT

- Transformers can be used for images as well. For image classification, it looks like this



- Early work showed ViT can outperform CNNs given very large amount of training data
- However, recent work* has shown that good old CNNs still rule! ViT and CNN perform comparably at scale, i.e., when both given large amount of compute and training data

An Image is Worth 16x16 Words: Transformers for Image Recognition at Scale (Dosovitskiy et al, 2020)

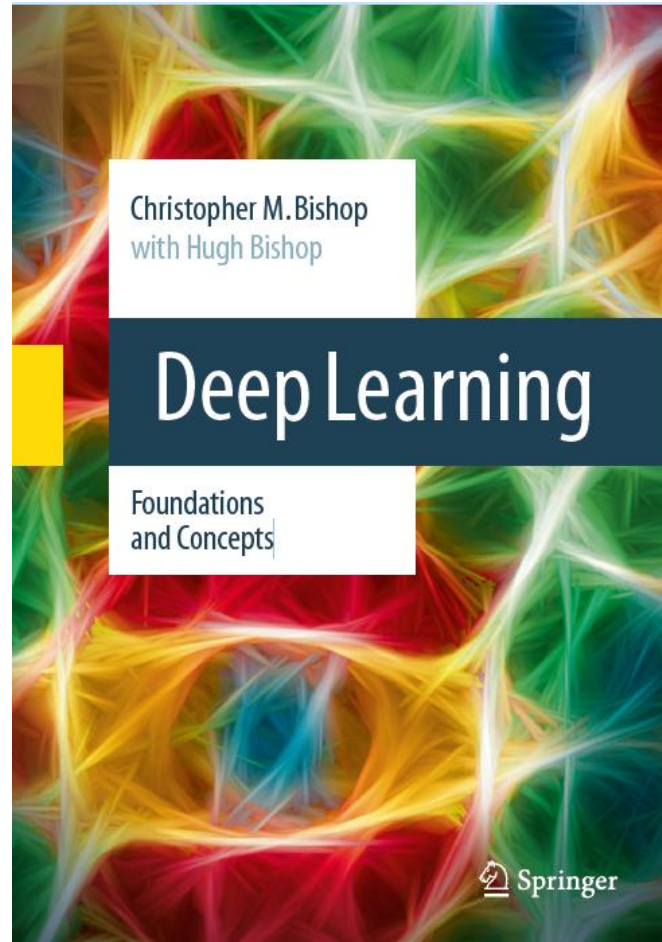
*ConvNets Match Vision Transformers at Scale (Smith et al, 2023)

Acknowledgments

- Slides are adapted from
 - Prof. Piyush Rai's course at IITK
 - Prof. Mitesh Khapra's course at IITM

References

Chapter 12



Thank you