

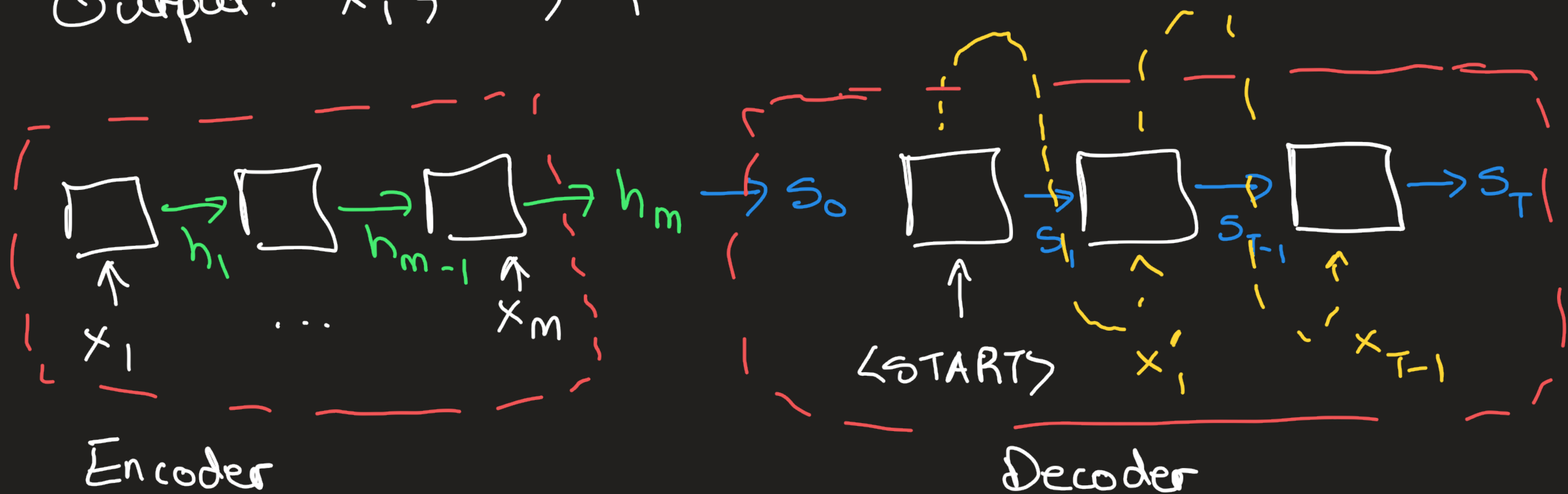
ML and Optimization Lec 25

- seq2seq model w/ attention
- Attention is all we need: seq2seq models w/o RNNs
- BERT: pretraining transformer models for NLP

Seq2Seq model

Input: $x_1, \dots, x_m$

Output: $x'_1, \dots, x'_T$



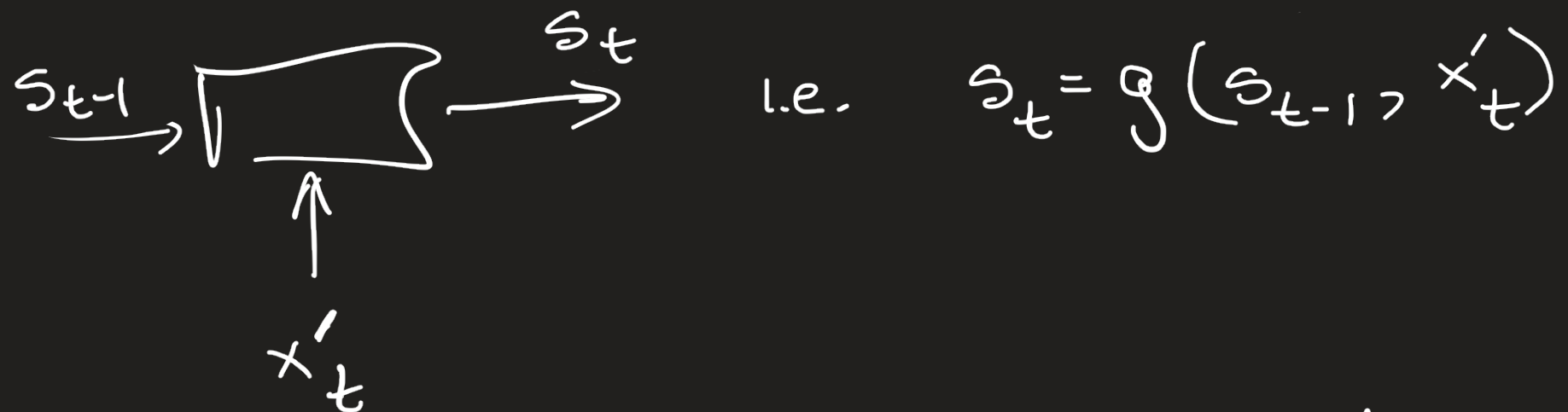
Issue: regardless of length of input sequence, all context/historical dependence of the output sequence is contained in a d -dimensional vector h_m leads to poor performance (BLEU decreases w/ length of input sequence for Machine translation tasks)

Seq2Seq with attention

- Introduced in Bahdanau et al. 2015, Neural Machine Translation by Jointly Learning to Align and Translate
- Allow each output token to attend to each individual hidden state in the input sequence. Captures idea of focusing on relevant portion of input sequence. Avoids the need to capture everything relevant to all outputs $x'_1, \dots, x'_T$ in the one vector h_m

Seq2Seq w/ attention

Idea: replace our previous hidden state update for decoder



w/ first finding which historical hidden states are most relevant:

$$s_{t-1}, x'_t, [h_1, \dots, h_m] \mapsto \alpha_t \text{ a prob. dist over } [h_1, \dots, h_m]$$

then form a context vector

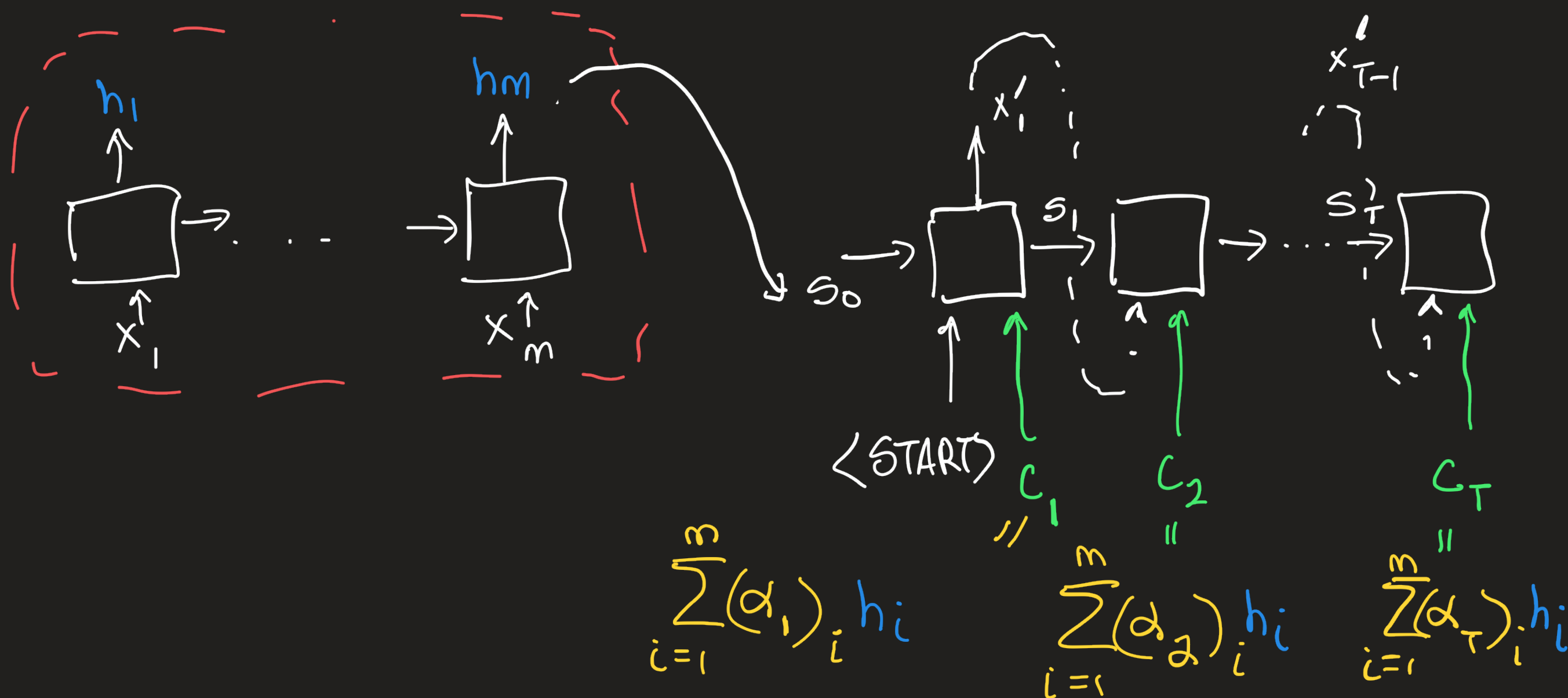
$$c_t = \sum_{i=1}^m (\alpha_t)_i h_i$$

and update with

$$s_t = g(s_{t-1}, c_t)$$

we removed dependence on x'_t because in the decoder x'_t isn't much more than s_{t-1}

In the context of decoding, in the decoder, use attention over the hidden states of the encoder



Options for computing the attention vectors α_t

1) Option 1 (used in Bahdanau et al. 2015)

$$\text{Take } \tilde{\alpha}_i = \underbrace{V^T \tanh\left(\underbrace{W}_{\text{parameters to learn}} \begin{bmatrix} h_i \\ s_{t-1} \end{bmatrix}\right)}_{\text{parameters to learn}} \text{ for } i=1, \dots, m$$

$$\text{Then } \alpha_t = \text{softmax}(\tilde{\alpha}) \in \mathbb{R}^m$$

2) (Dot-product attention) Idea: generate keys and queries as linear transforms of h_i and s_{t-1} respectively $\rightarrow$ take logits $\tilde{\alpha}_i$ to be their inner-products

$$W_K h_i = k_i$$

$$W_Q s_{t-1} = q_{t-1}$$

$$\tilde{\alpha}_i = \langle k_i, q_{t-1} \rangle$$

$$\alpha_t = \text{softmax}(\tilde{\alpha})$$

Complexity of using attention

- 1) Simple RNN encoder models : $O(m + T)$
- 2) RNN + attention models $O(m \cdot T)$

because for each output token, must
compute $\alpha_t \in \mathbb{R}^m$ by comparing to all
inputs

Transformers : Attention is all we need

- 2017, Vaswani et al., Attention is all we need
- Insight: attention alone suffices to contain all historical information we need, so can eliminate recursive structure.
- Replace hidden representations in RNNs w/ contextual representations C_t computed using self-attention (for encoder)
- Use combination of self-attention in decoder and general attention against the encoder contextual representations to get contextual representations to predict output tokens

Self-Attention

(replace the hidden states learned in RNNs w/ contextual representations)



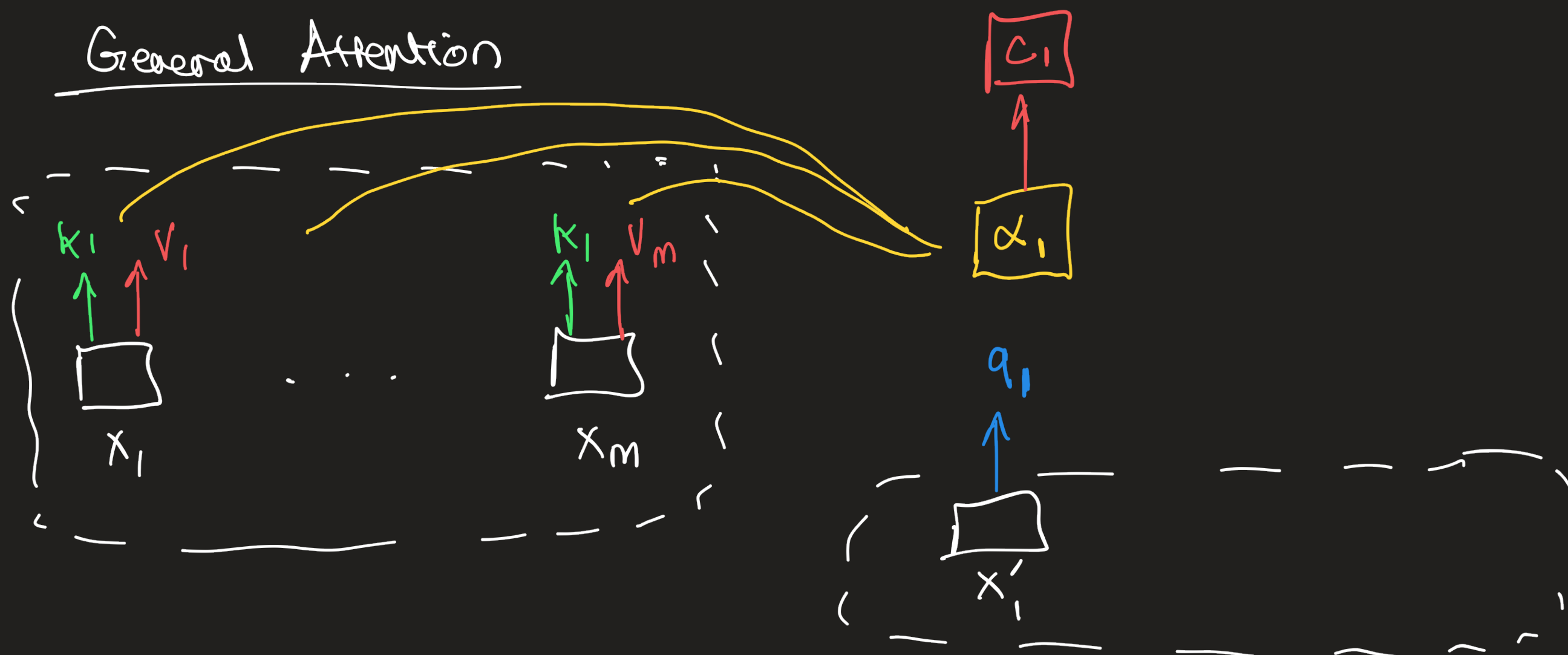
$$c_i = \sum_{j=1}^m (\alpha_{ij}) v_j$$

where

$$\begin{aligned} \text{key} \quad k_i &= W_K x_i \\ \text{query} \quad q_i &= W_Q x_i \\ \text{value} \quad v_i &= W_V x_i \end{aligned}$$

and $(\tilde{\alpha}_{ij}) = \langle q_i, k_j \rangle$ and $\alpha = \text{softmax}(\tilde{\alpha})$

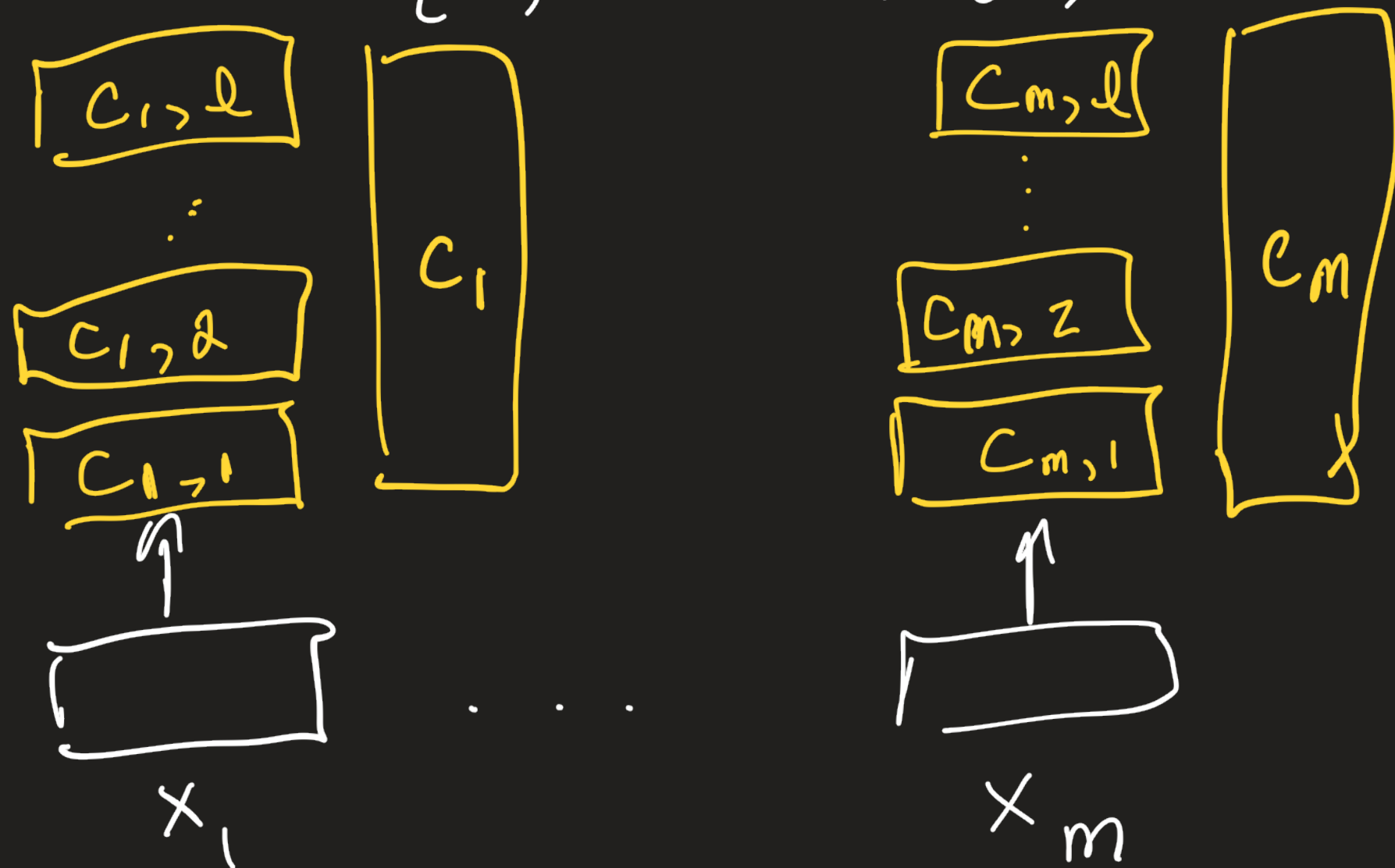
General Attention



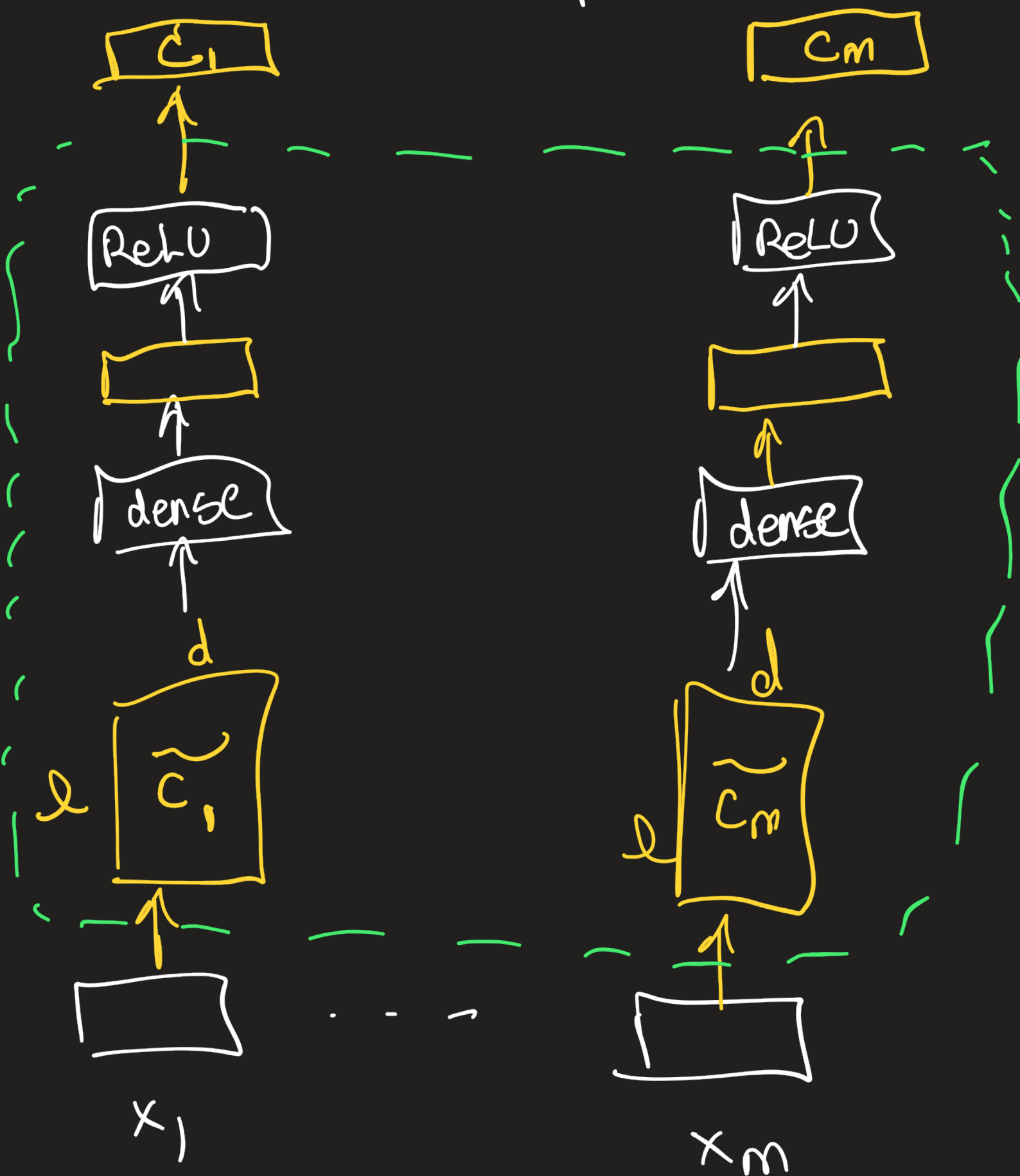
Multi-Head Self Attention

Just as with CNNs we benefit from learning multiple contextual representations:

$(\omega_K)_i, (\omega_Q)_i, (\omega_V)_i$ for $i=1, \dots, d$



Contextual representations using multi-head self-attention



Self-Attention + Dense
Layer