

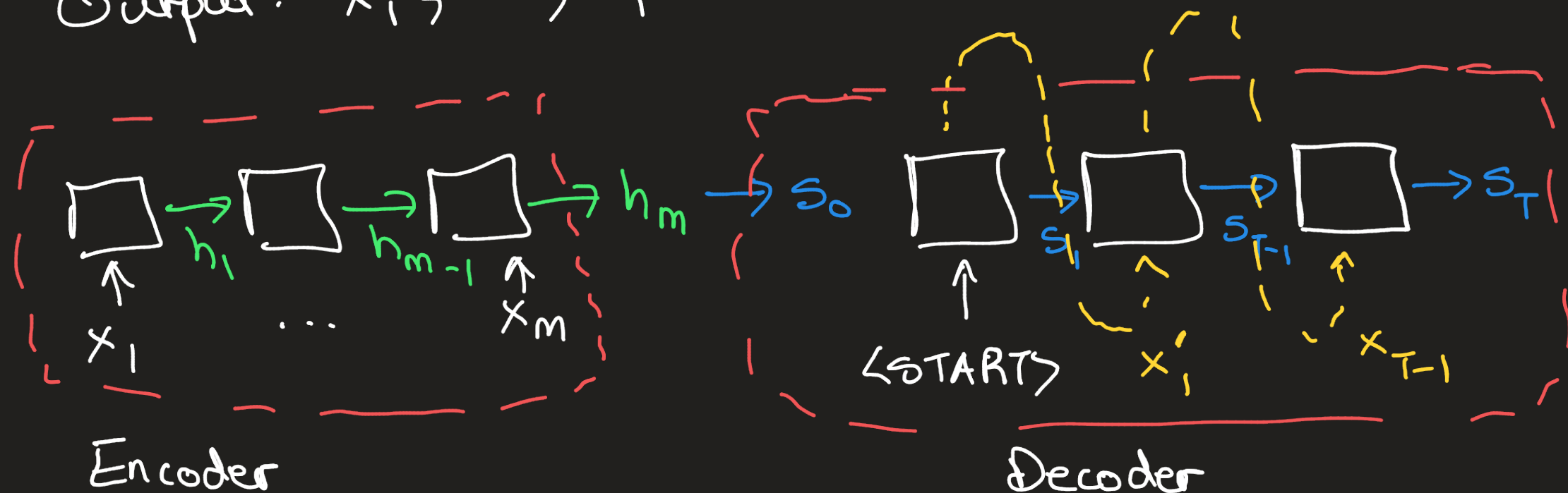
ML and Optimization Lec 23

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



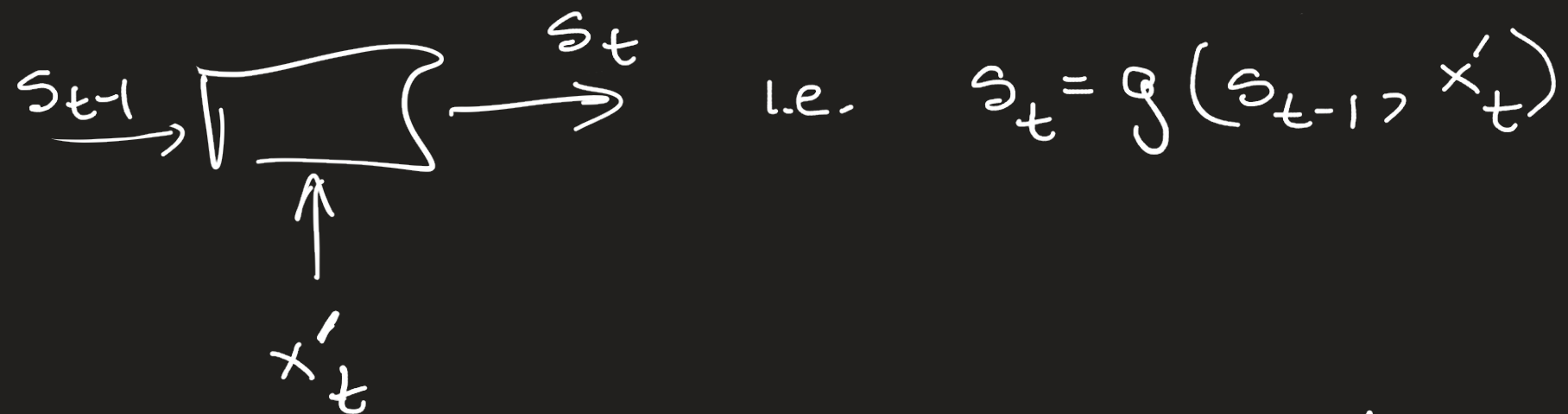
Issue: regardless of length of input sequence, all context/historical dependence of the output sequence is contained in a d -dimensional vector h_m leads to poor performance (BLEU decreases w/ length of input sequence for Machine translation tasks)

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

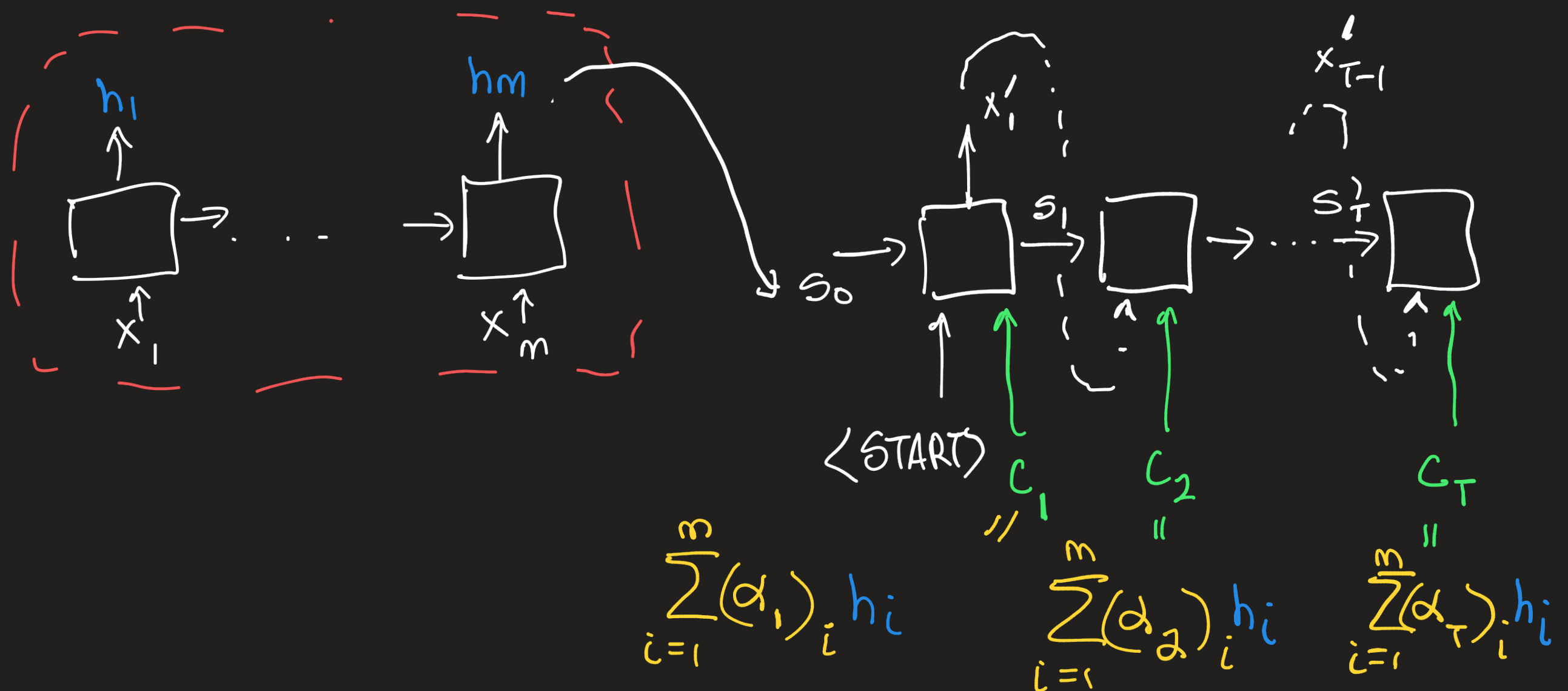
$$c_t = \sum_{i=1}^m (\alpha_t)_i h_i$$

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, take logits $\tilde{\alpha}_i$ to be their inner-products

$$W_K h_i = k_i$$

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

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

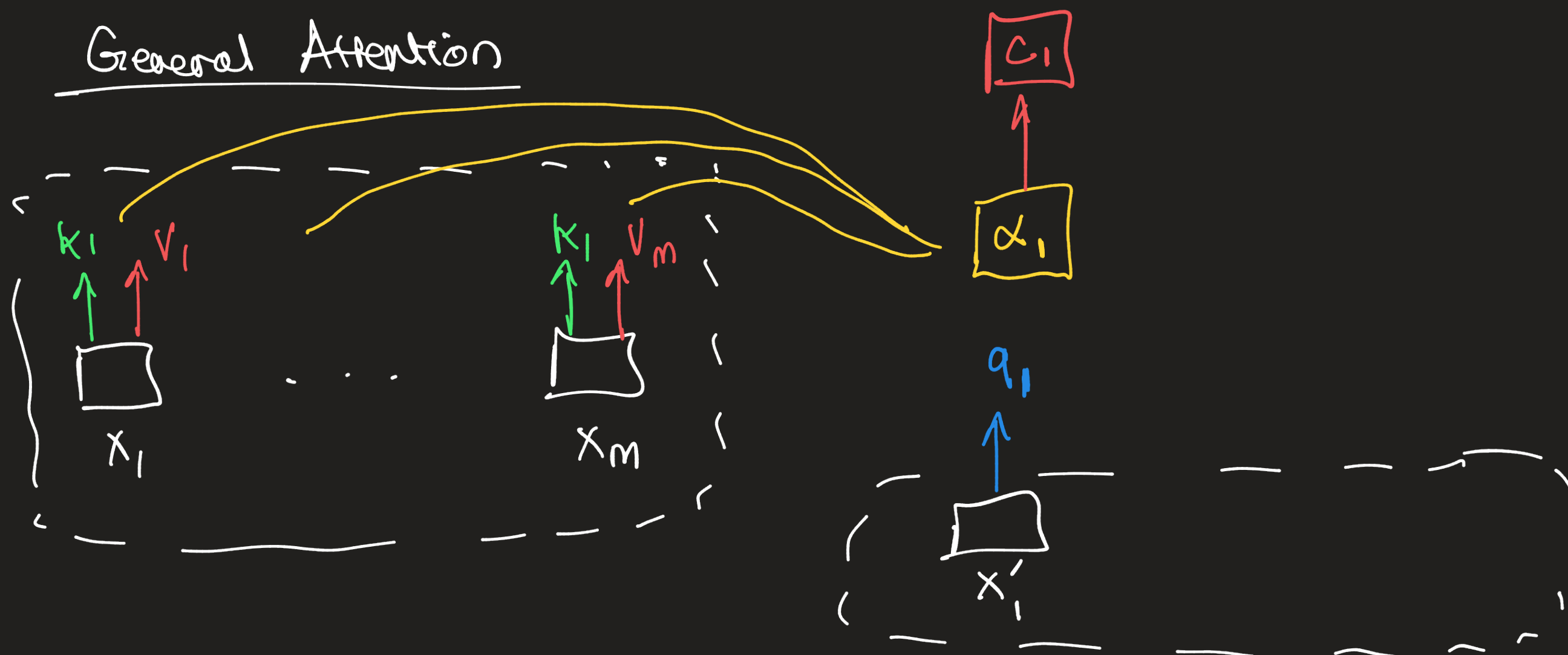
$$\alpha_t = \text{softmax}(\tilde{\alpha})$$

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

General Attention



Self-Attention

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



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

where

key $k_i = W_K x_i$

query $q_i = W_Q x_i$

value $v_i = W_V x_i$

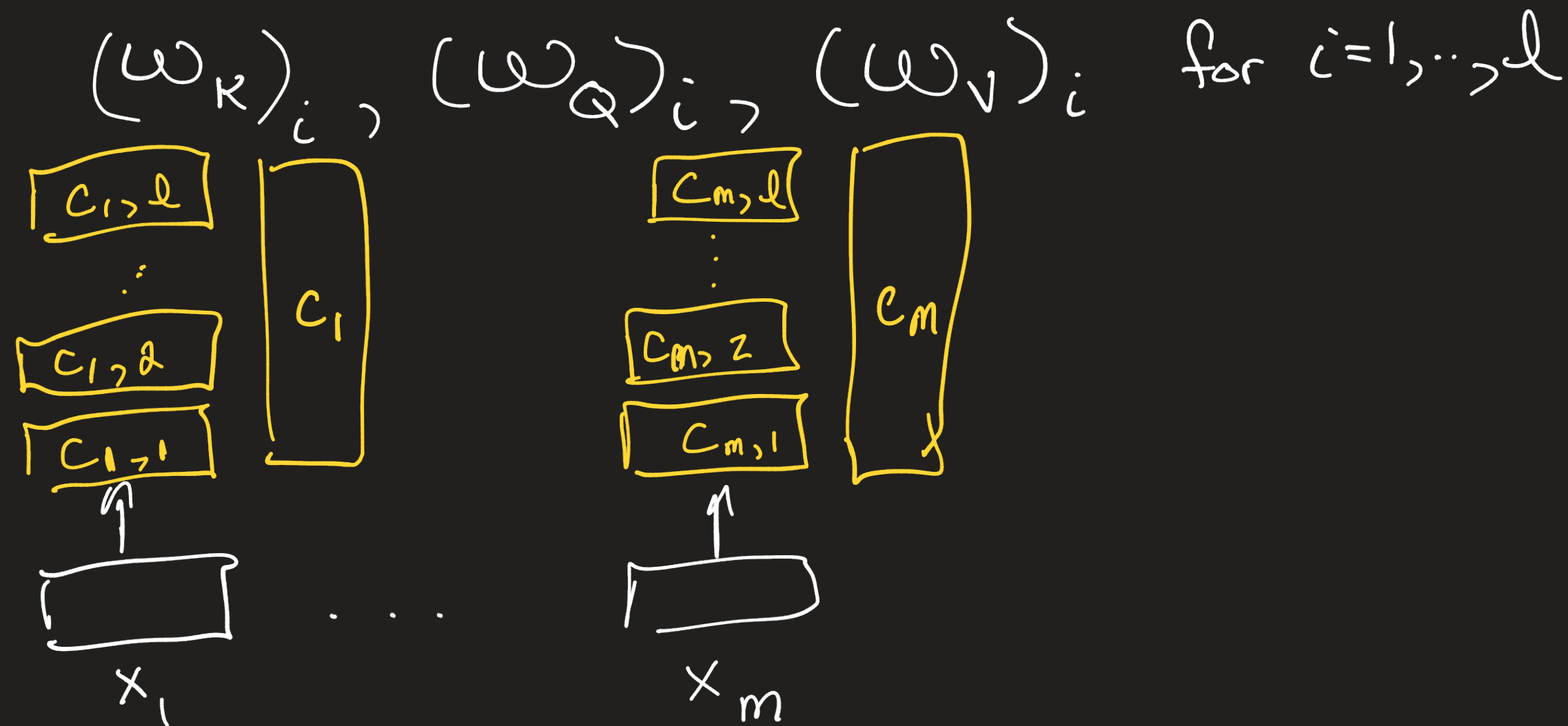
and $(\tilde{\alpha}_{ij}) = \langle q_i, k_j \rangle$ and $\alpha = \text{softmax}(\tilde{\alpha})$

Transformers : Attention is all we need

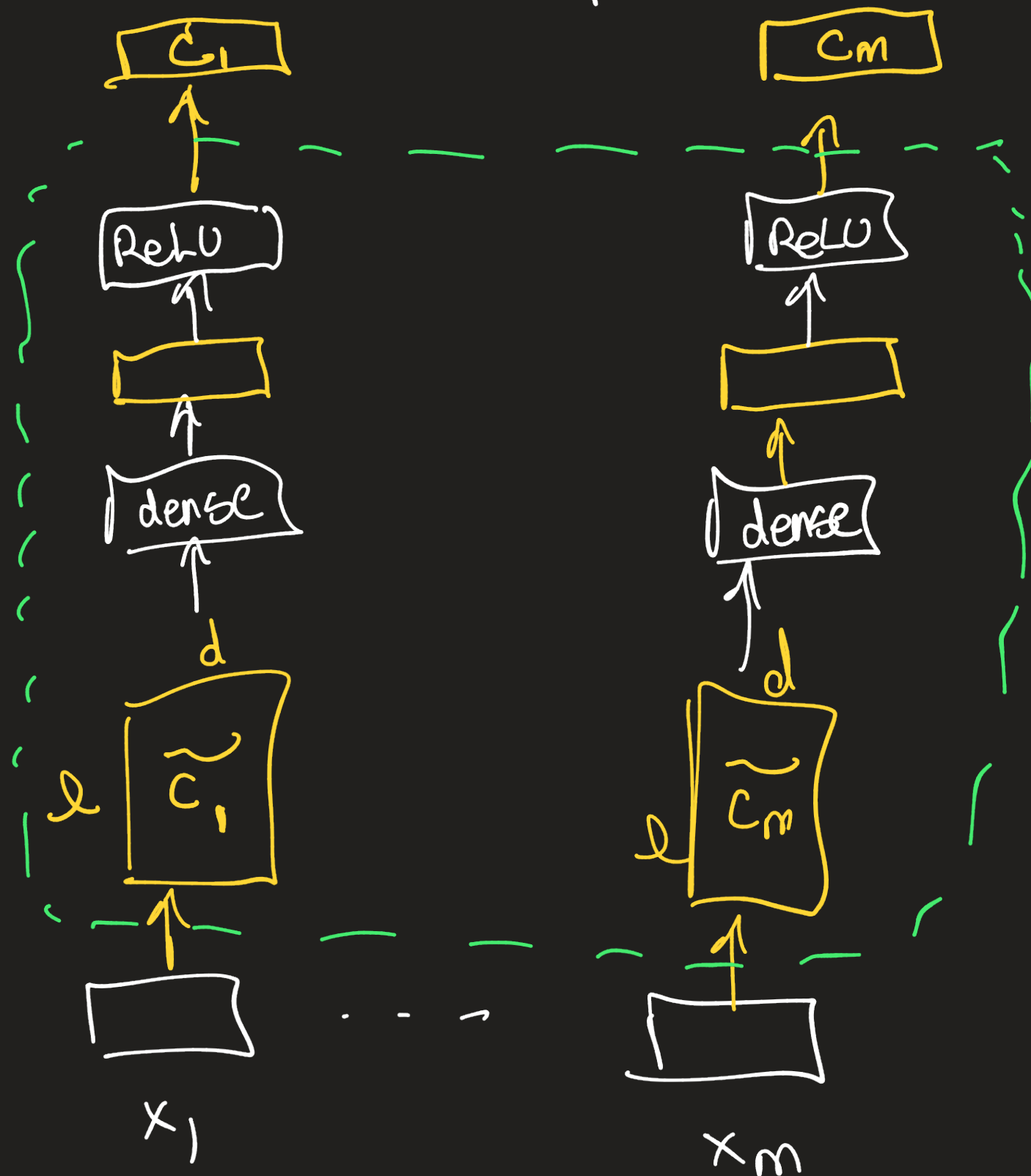
- 2017, Vaswani et al., Attention is all we need
- Insight: attention alone suffices to contain all historical information we need, so can eliminate recursive structure.
- Replace hidden representations in RNNs w/ contextual representations C_t computed using self-attention (for encoder)
- Use combination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

Multi-Head Self Attention

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



Contextual representations using multi-head self-attention



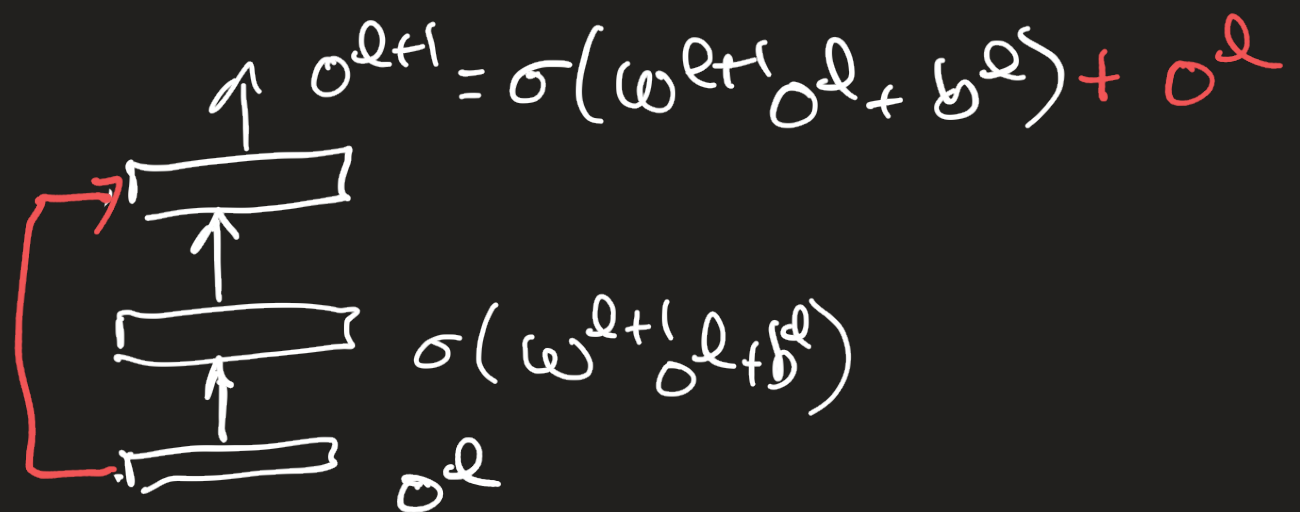
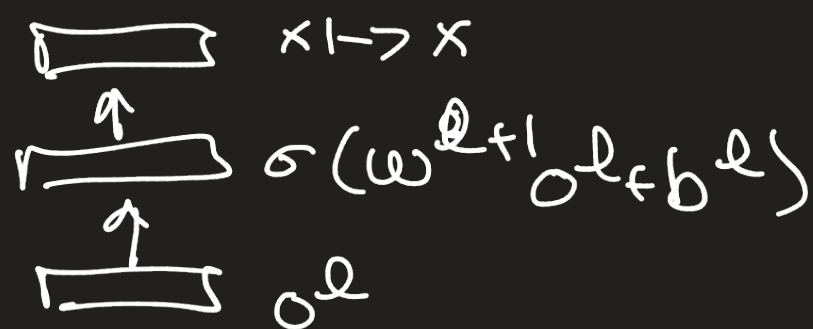
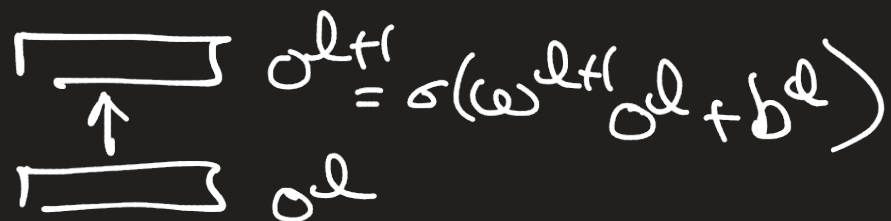
Self-Attention + Dense
Layer

Skip connections

Skip connections are connections between layers l and layers higher than $l+1$, e.g. DenseNet, where all layers are connected.

An efficient & popular type of skip-connection is a residual connection,

Residual connection



Layer-normalization

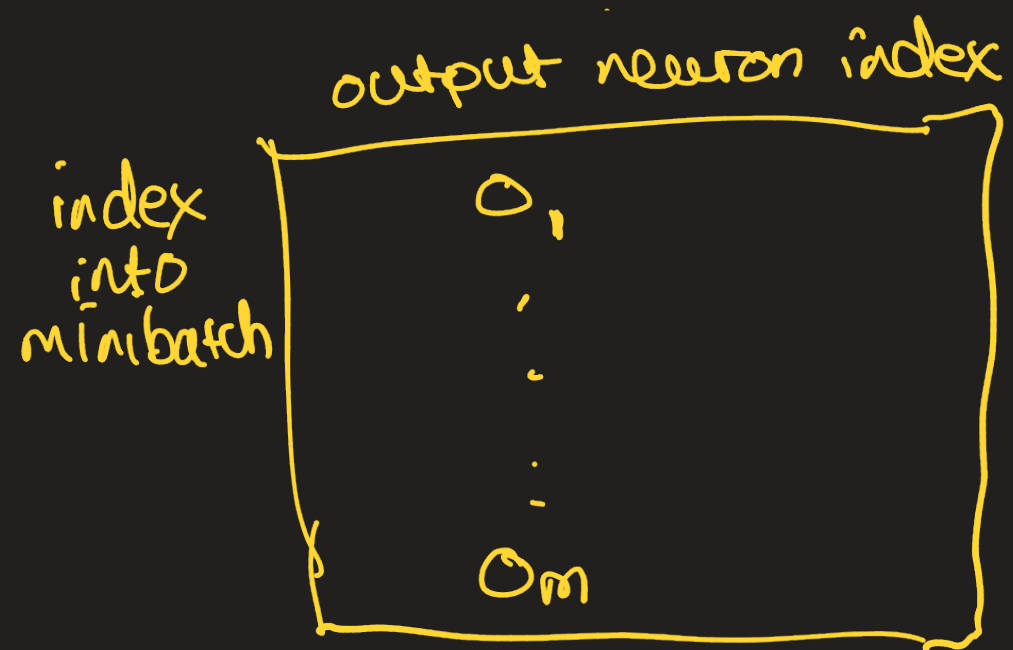
Like batch normalization this helps stabilize training by ensuring the distribution of the outputs of layer l doesn't vary much over time, so layer $l+1$ doesn't have to waste optim time adapting.

$$\text{LayerNorm}(o) = \gamma \cdot \frac{o - \bar{o}}{\sigma} + \beta,$$

where $\gamma \in \mathbb{R}^{n_l}$ and $\beta \in \mathbb{R}^{n_l}$ are the parameters of the layer normalization layers and

$$\bar{o} = \frac{1}{n_l} \sum_{i=1}^{n_l} o_i, \quad \sigma^2 = \frac{1}{n_l} \sum_{i=1}^{n_l} (o_i - \bar{o})^2$$

Batch normalization



operates on minibatches

Batch normalization
↓ computes per neuron means and stds over the minibatch



then batch normalization of an o_i is given by

$$\gamma \cdot \left(\frac{o_i - \mu}{\sigma} \right) + \beta$$

Layer normalization

operates on a single output



$\bar{o}$ = mean of o

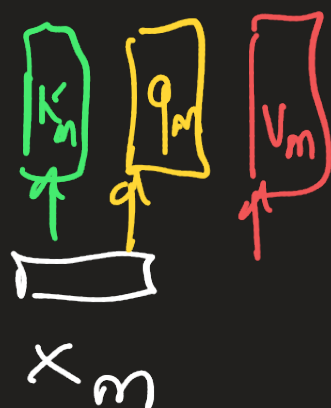
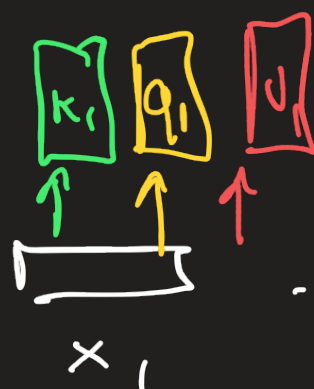
σ = std of o

then layer normalization of o is given by

$$\gamma \cdot \left(\frac{o - \bar{o}}{\sigma} \right) + \beta$$

Expressions for Attention

Self-attention:



then $\tilde{\alpha}_i^T = [q_i^T k_1 \cdots q_i^T k_m]$
 is the logits vector (row vector)
 for the i th attention vector

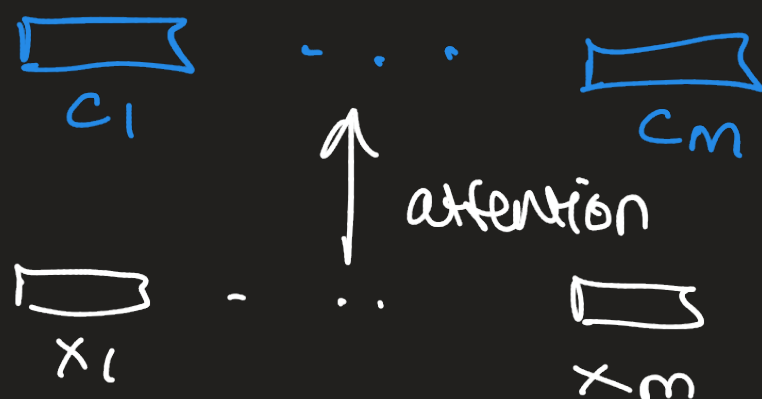
$$\alpha_i^T = \text{softmax}(\tilde{\alpha}_i^T)$$

is my attention vector

then our contextual embedding for the i th token is

$$c_i^T = \sum_{j=1}^m (\alpha_i)_j v_j^T$$

$$\text{let } X = \begin{bmatrix} x_1^T \\ \vdots \\ x_m^T \end{bmatrix} \in \mathbb{R}^{m \times d}$$



$$\text{and } C = \begin{bmatrix} c_1^T \\ \vdots \\ c_m^T \end{bmatrix} \in \mathbb{R}^{m \times d_v}$$

Then

$$C = \text{Att}(X, X)$$

where the attention layer has parameters $\omega_Q, \omega_K, \omega_V$

To compute C , write

$$K = X \omega_K = \begin{bmatrix} x_1^T \omega_K \\ \vdots \\ x_m^T \omega_K \end{bmatrix}$$

d_K -dim

$$Q = X \omega_Q = \begin{bmatrix} x_1^T \omega_Q \\ \vdots \\ x_m^T \omega_Q \end{bmatrix}$$

d_K -dim

$$V = X \omega_V = \begin{bmatrix} x_1^T \omega_V \\ \vdots \\ x_m^T \omega_V \end{bmatrix}$$

then note that

$$C = \begin{bmatrix} c_1^T \\ \vdots \\ c_m^T \end{bmatrix} = \underset{\substack{\uparrow \\ \text{row-wise}}}{\text{softmax}} \left(\frac{Q K^T}{\sqrt{d_K}} \right) V$$

Similarly, a general attention layer takes two sequences $X \in \mathbb{R}^{m \times d}$ $X' \in \mathbb{R}^{T \times d}$



so a general attention layer computes

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

where

$$K = X \omega_K$$

$$V = X \omega_V$$

$$Q = X' \omega_Q$$

Transformer Architecture (image upcoming)

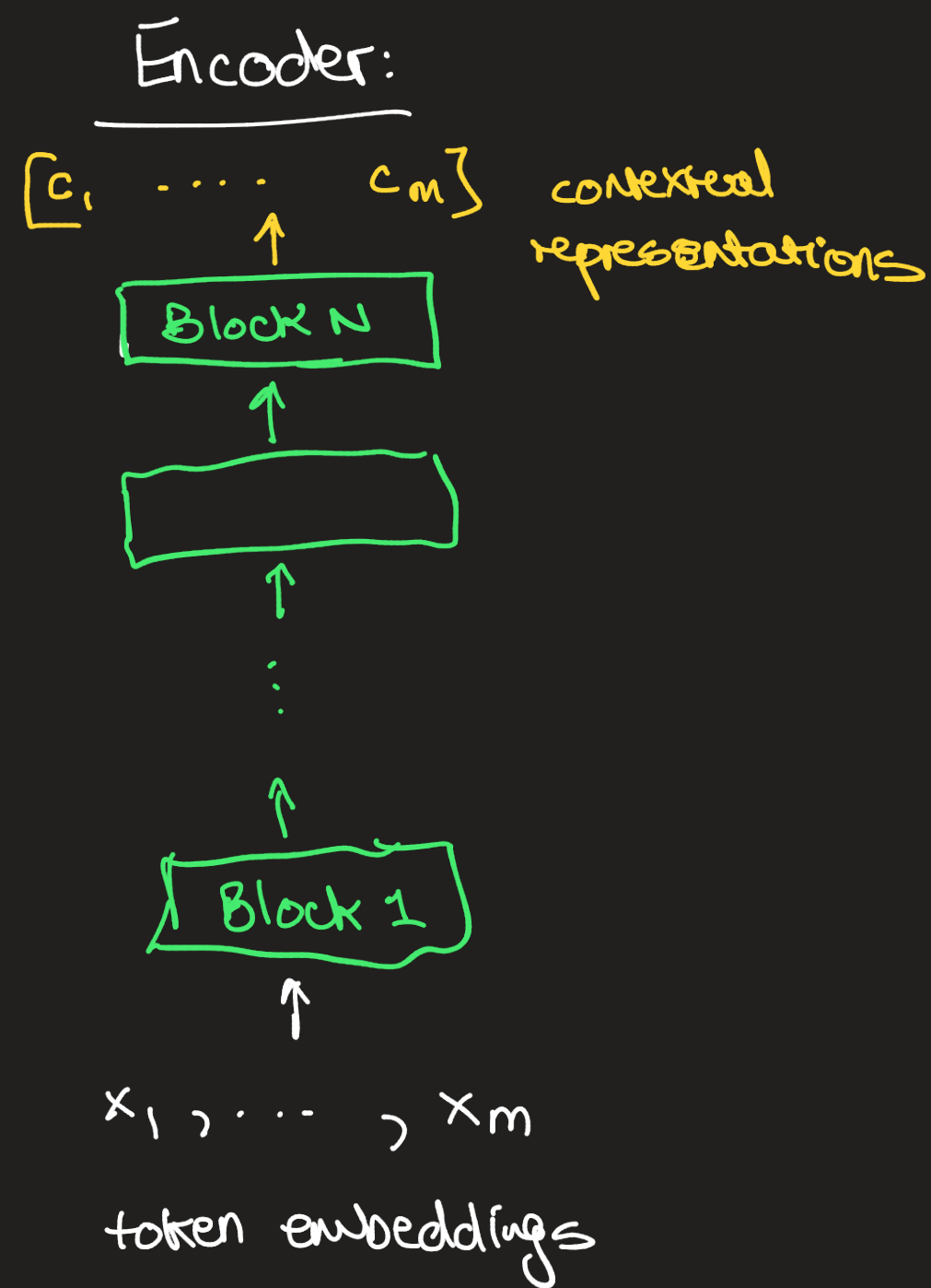
Encoder:

- use multiple layers of multi-head self-attention with residual connections, layer-normalization, and dense nonlinear transforms

Decoder

- use multiple layers of multi-head self-attention, multi-head general attention, with residual connections, layer normalizations, and dense nonlinear transforms.

NB: we use layer normalization in Transformer blocks for the same reasons as in RNNs: variable sequence lengths and potentially small batch sizes



where each embedding block has its own parameters and computes:

Given input $X = \begin{bmatrix} x_1 \\ \vdots \\ x_m \end{bmatrix} \in \mathbb{R}^{m \times d}$,

i) $\varphi = \text{LayerNorm} (X + \text{Mult-head Attn}(X, X))$

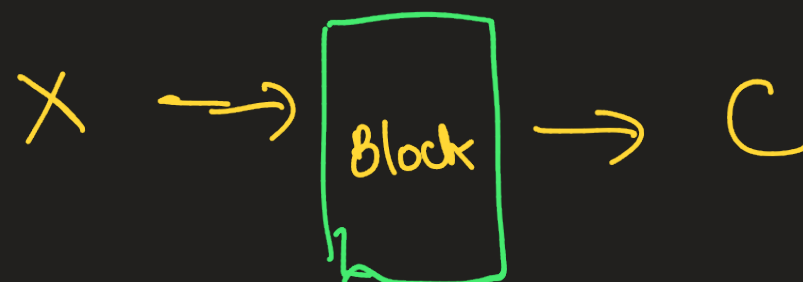
recall $\text{Mult-head Attn}(X, X) =$

$$[A_{\text{Attn}_1}(X, X) \parallel \dots \parallel A_{\text{Attn}_h}(X, X)] W_o$$

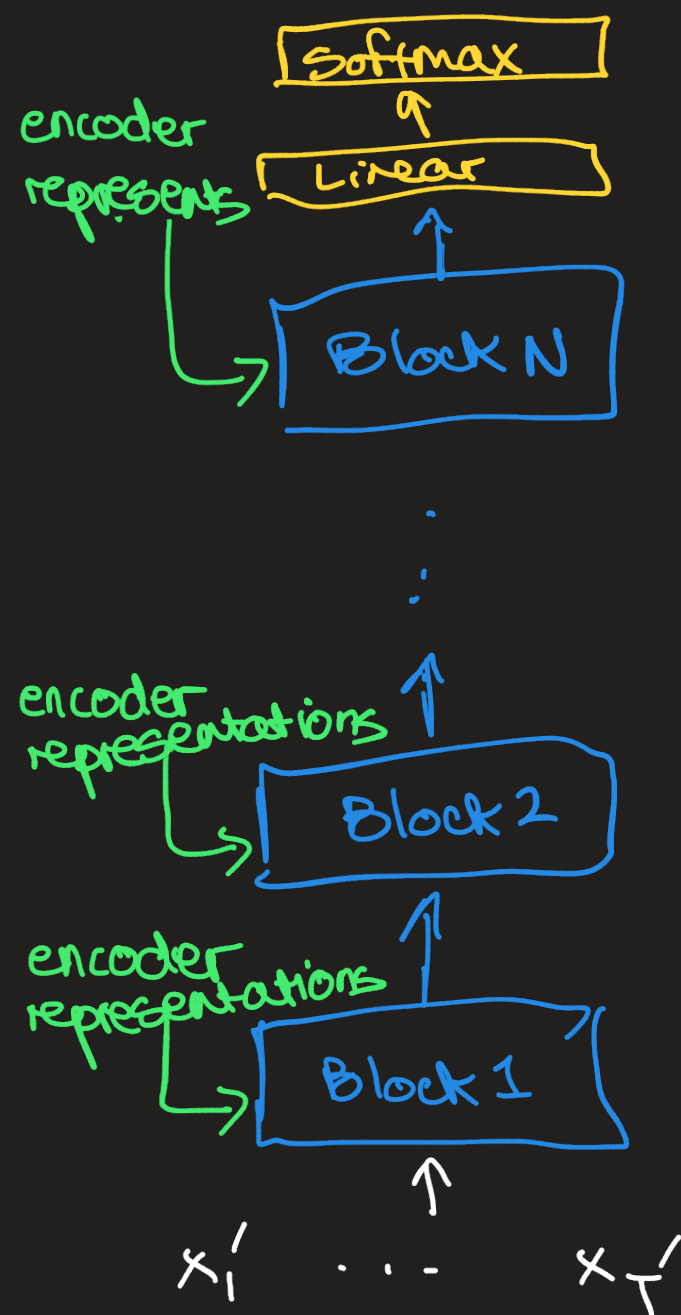
concatenates h different attention outputs together then takes linear combination to learn a final d -dim representation

ii) $C = \text{LayerNorm} (\varphi + \text{FF}(\varphi))$

where FF applies the same feedforward dense + ReLU layer to each row of φ



Decoder



During training, use teacher forcing

each block has its own parameters and computes, given inputs $X' = \begin{bmatrix} x'_1 \\ \vdots \\ x'_T \end{bmatrix}$

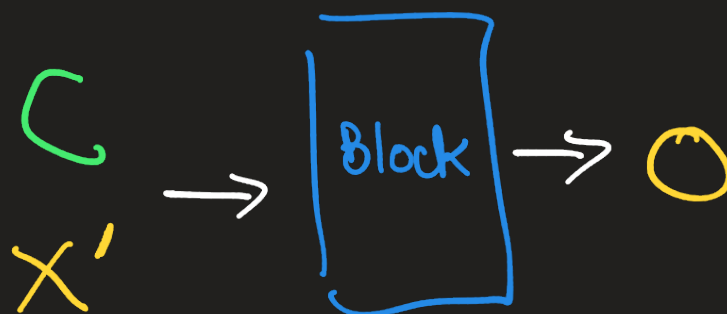
and the final representations from the encoder of the input sequence

$$C = \begin{bmatrix} c_1 \\ \vdots \\ c_m \end{bmatrix} :$$

$$(i) \quad \varphi = \text{LayerNorm}(X' + \text{Multi-head Attn}(X', X'))$$

$$(ii) \quad Z = \text{LayerNorm}(\varphi + \text{Multi-head Attn}(C, \varphi))$$

$$(iii) \quad O = \text{LayerNorm}(Z + \text{FF}(Z))$$

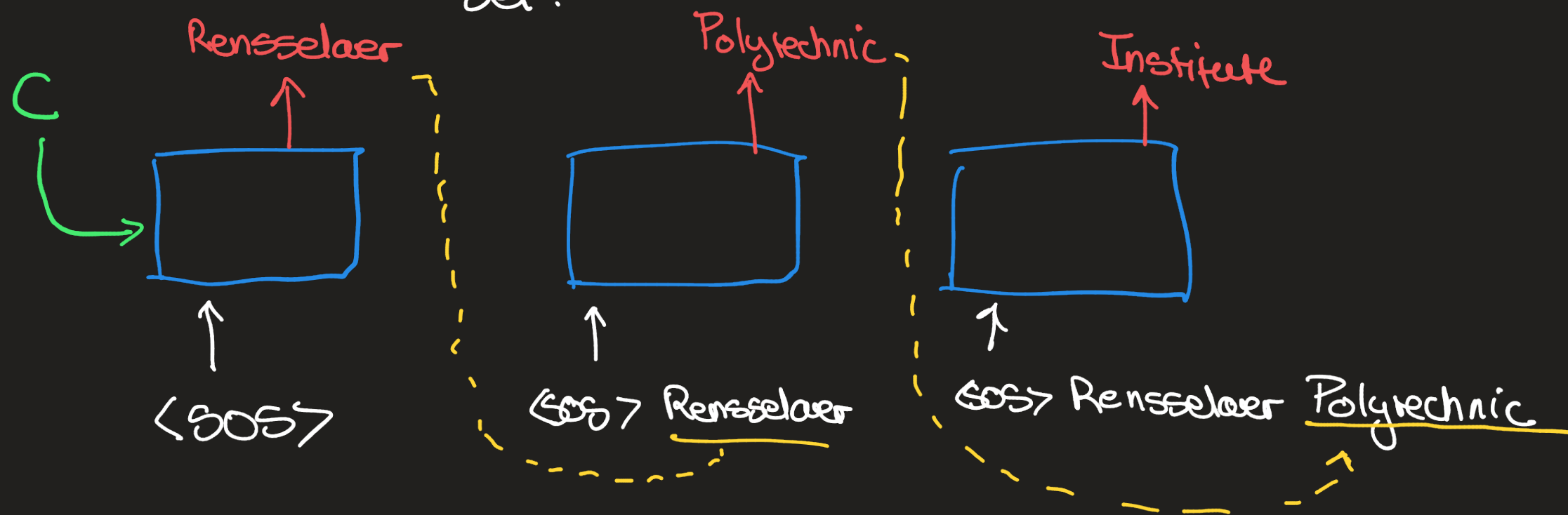


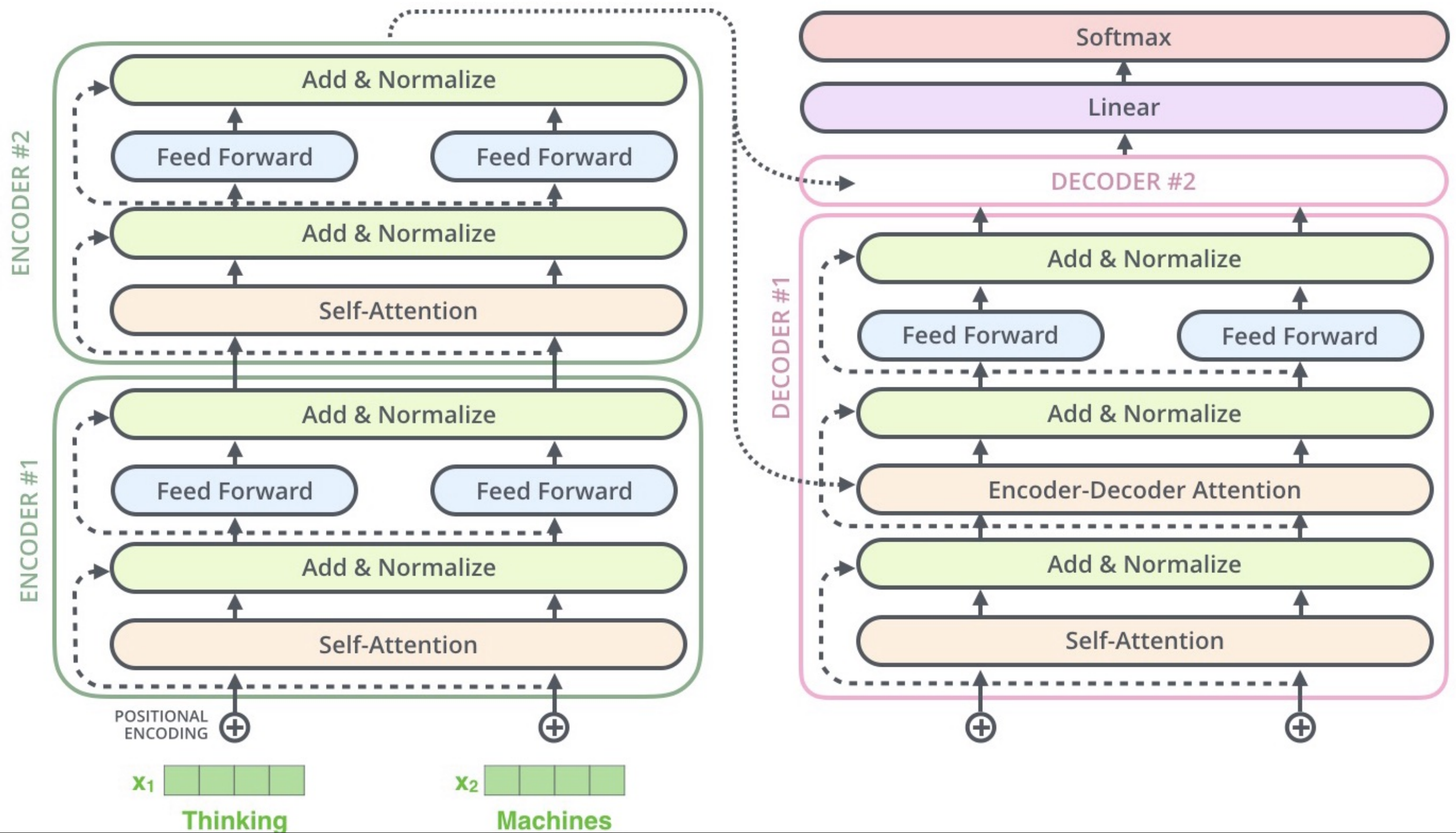
In the self-attention layers, the attention is masked : we set

$(\alpha_{ij})_j = 0$ for $j > i$ to ensure we only predict using historical info.

Decoder

During inference time, we autoregressively use the decoder to predict the next token, add it onto our input sequence, and pass it back through the decoder to get the next token. Repeat until we either get an $\langle \text{EOS} \rangle$ token, or we've reached the maximum sequence length we set.





Positional Encodings

Issue w/ Transformers as is: they don't properly learn positional relations (e.g. that "the cat ate its food" should have a diff contextual rep from "the food ate its cat")

Soln: replace $X \in \mathbb{R}^{m \times d}$ with $X + P$

where $P \in \mathbb{R}^{m \times d}$ encodes the positions 1 through m

$$P_{i,j} = \begin{cases} \sin(i\omega j/d) & \text{if } j \text{ even} \\ \cos(i\omega(j-1)/d) & \text{if } j \text{ odd} \end{cases}$$

(refer to preceding diagram to see where positional embeddings are added)