

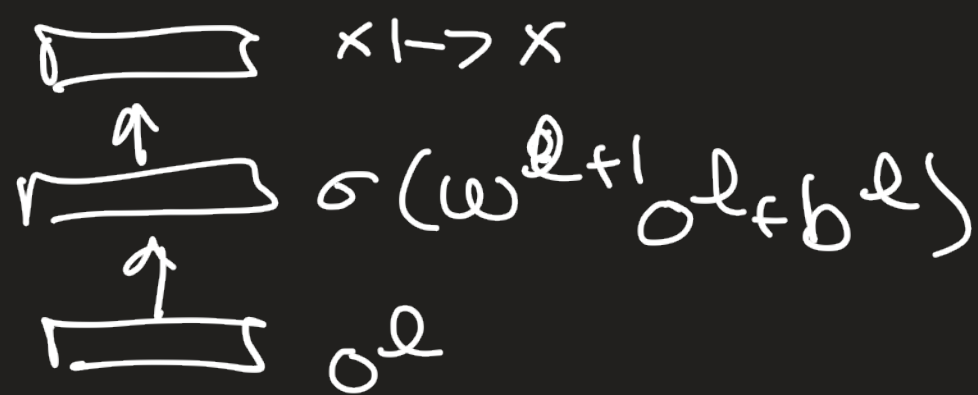
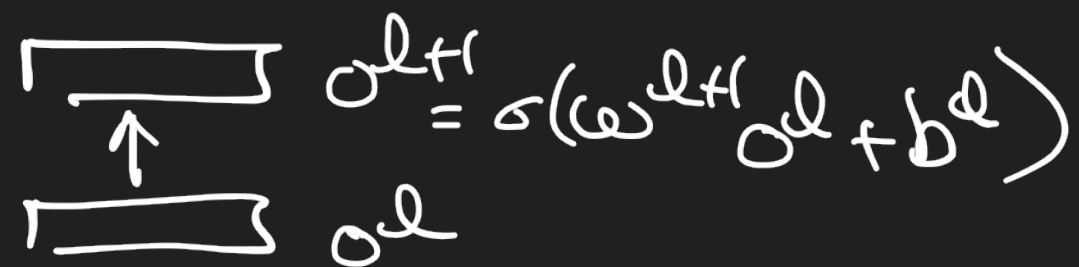
## ML and Optim Lec 26

- Transformer Seq2Seq Architecture
- Skip-connections
- Layer-normalization

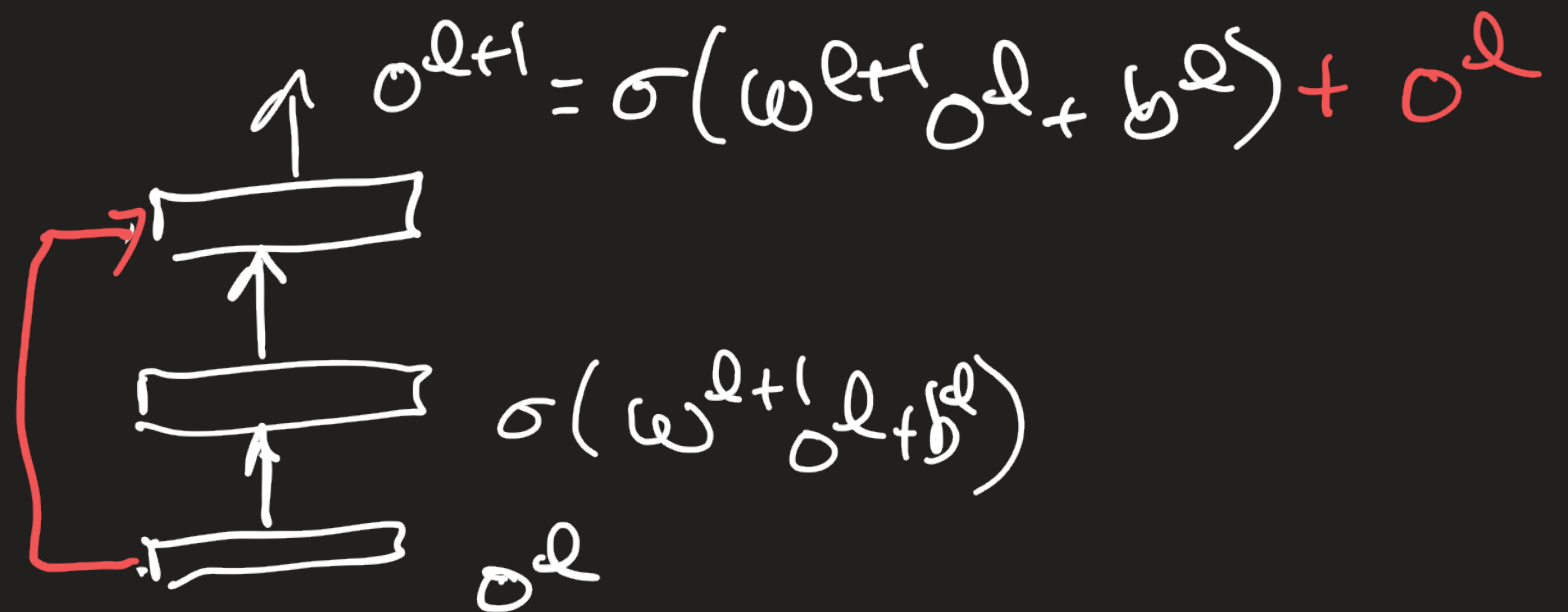
# Skip-connections

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

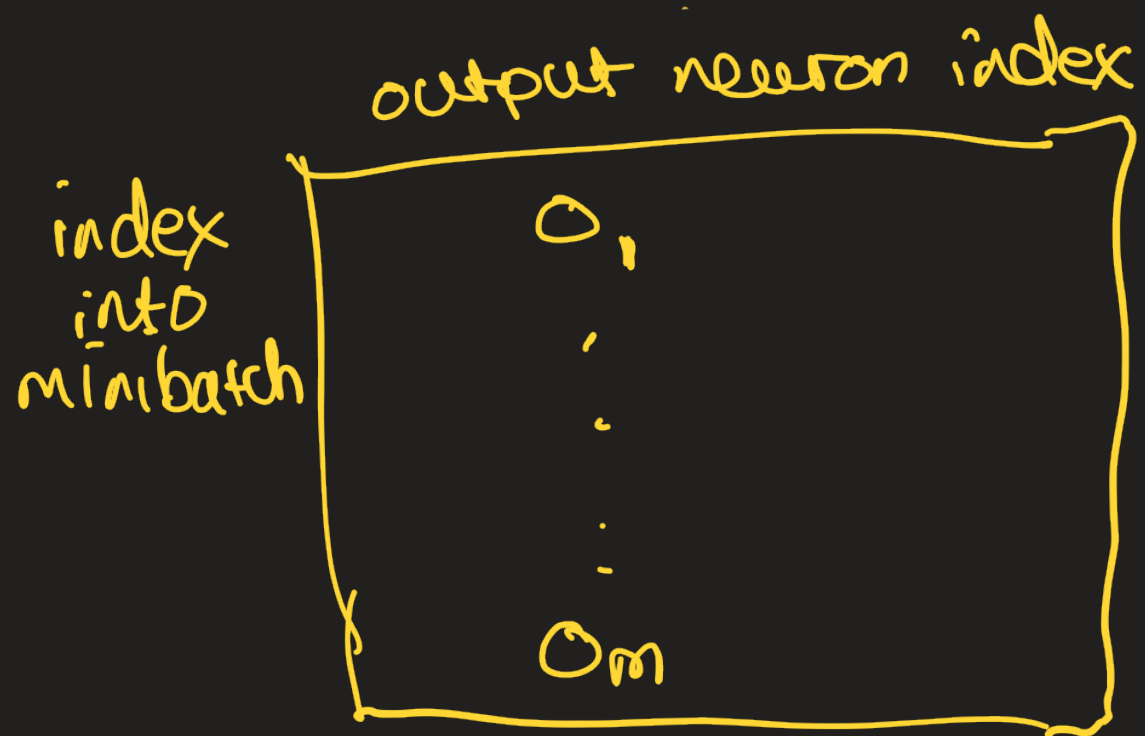
Like batch normalization this helps stabilize training by ensuring the distribution of the outputs of layer  $l$  doesn't vary much over time, so layer  $l+1$  doesn't have to waste optim time adapting.

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

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  $o$

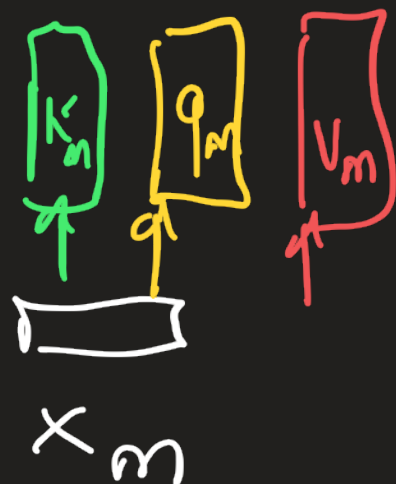
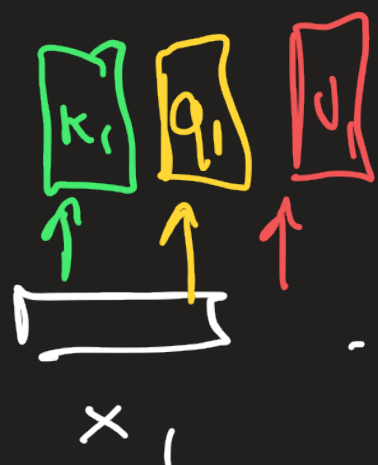
$\sigma$  = stdv of  $o$

then layer normalization of  $o$  is given by

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

# Expressions for Attention

Self-attention:



then

$$\tilde{\alpha}_i^T = [q_i^T k_1 \quad \dots \quad q_i^T k_m]$$

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

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

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



↑  
attention

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

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

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

, then note that

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

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

$d_K$ -dim

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)

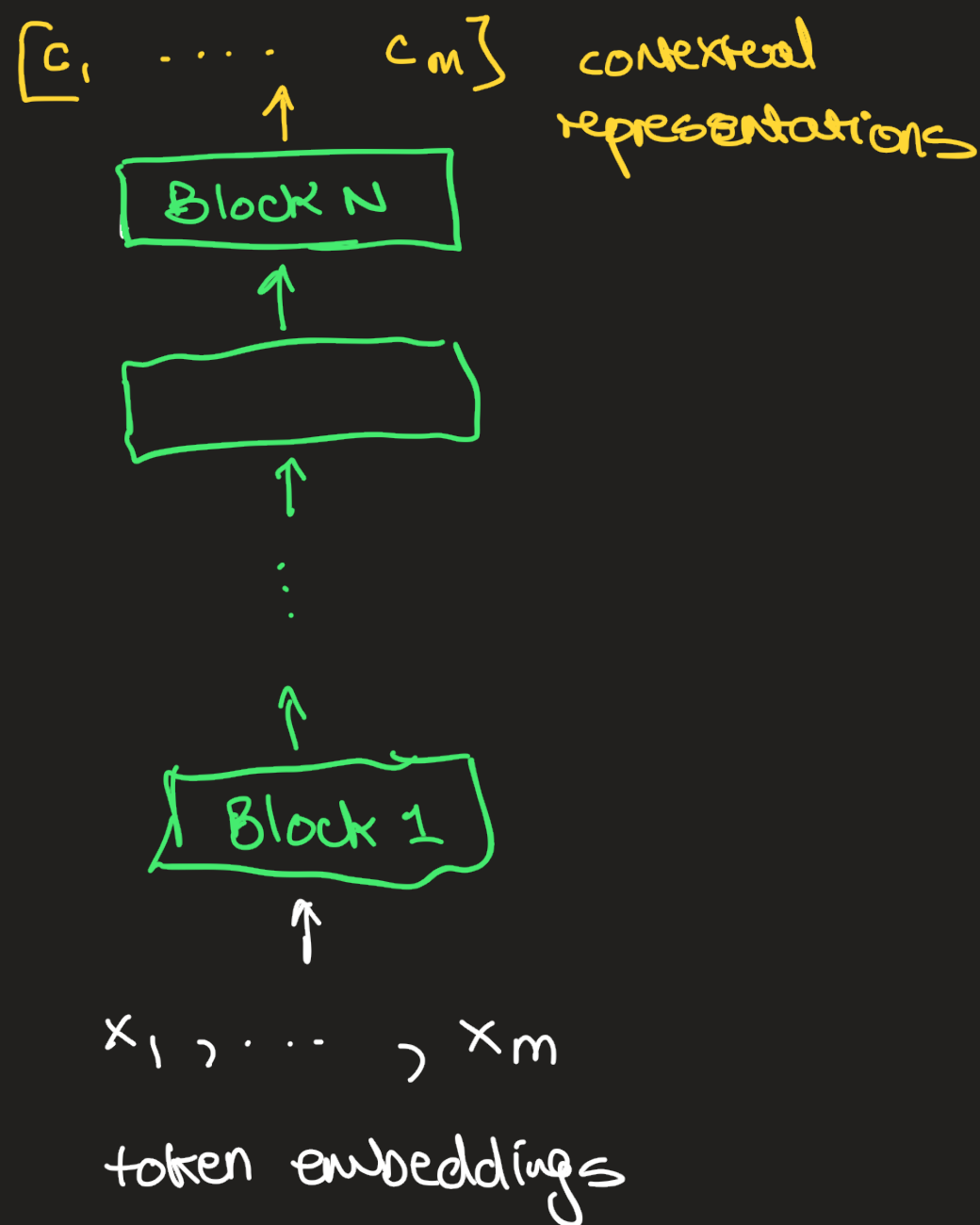
## Encoder:

- 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.

Encoder:



where each embedding block has its own parameters and computes:

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

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

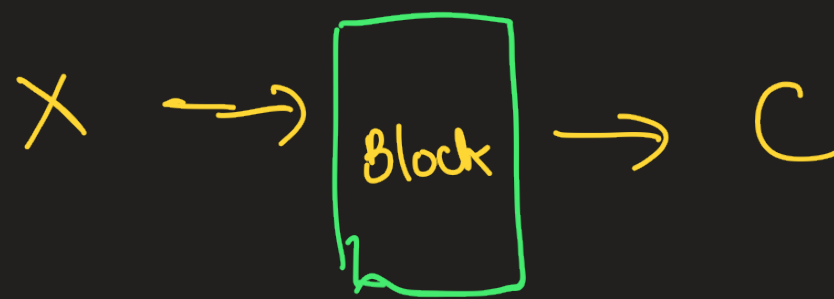
recall  $\text{Mult-head Attn}(X, X) =$

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

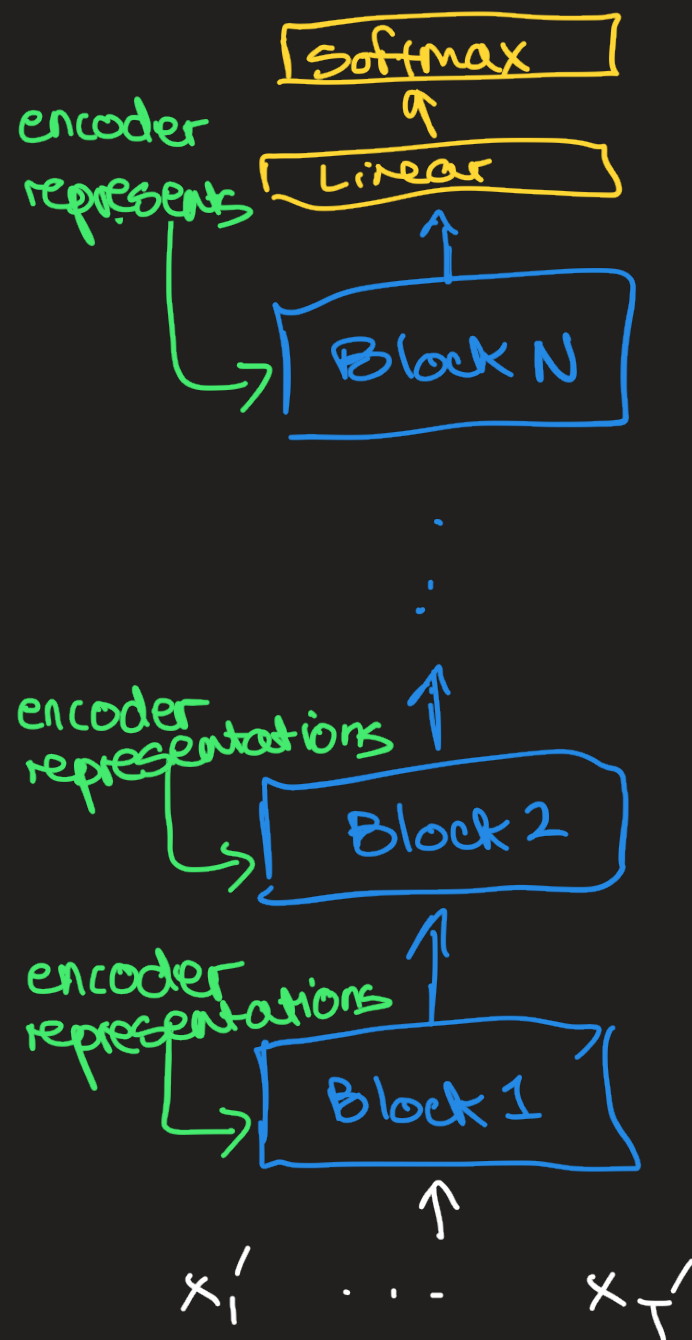
concatenates  $h$  different attention outputs together then takes linear combination to learn a final  $d$ -dim representation

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

where FF applies the same feedforward dense + ReLU layer to each row of  $\varphi$



## Decoder



During training, use teacher forcing

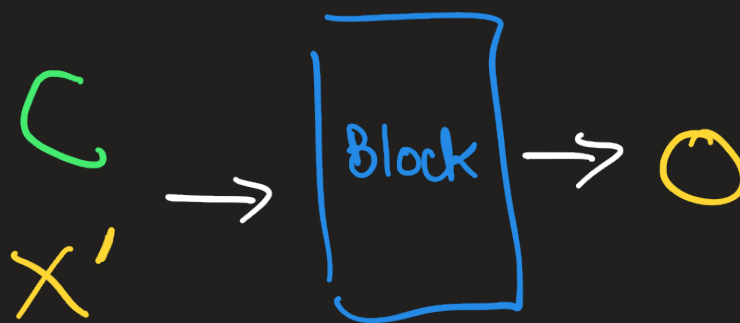
each block has its own parameters and computes, given inputs  $X' = \begin{bmatrix} x'_1 \\ \vdots \\ x'_T \end{bmatrix}$

and the final representations from the encoder of the input sequence  $C = \begin{bmatrix} c_1 \\ \vdots \\ c_m \end{bmatrix}$ :

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

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

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



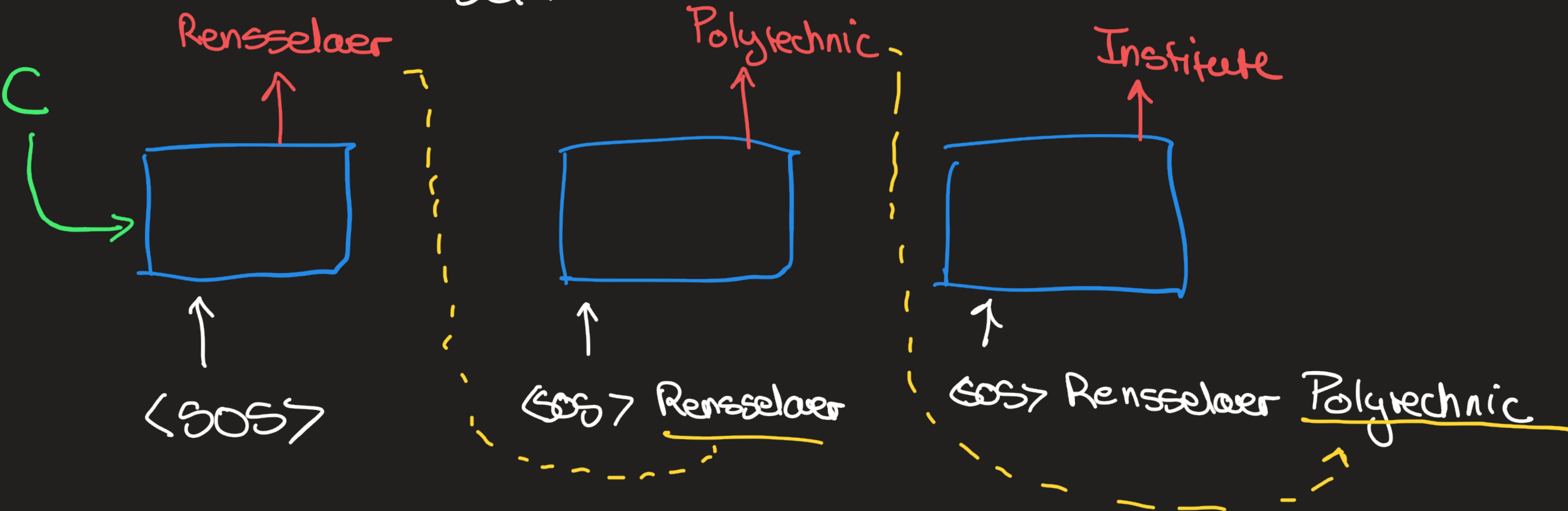
In the self-attention layers, the attention is masked : we set

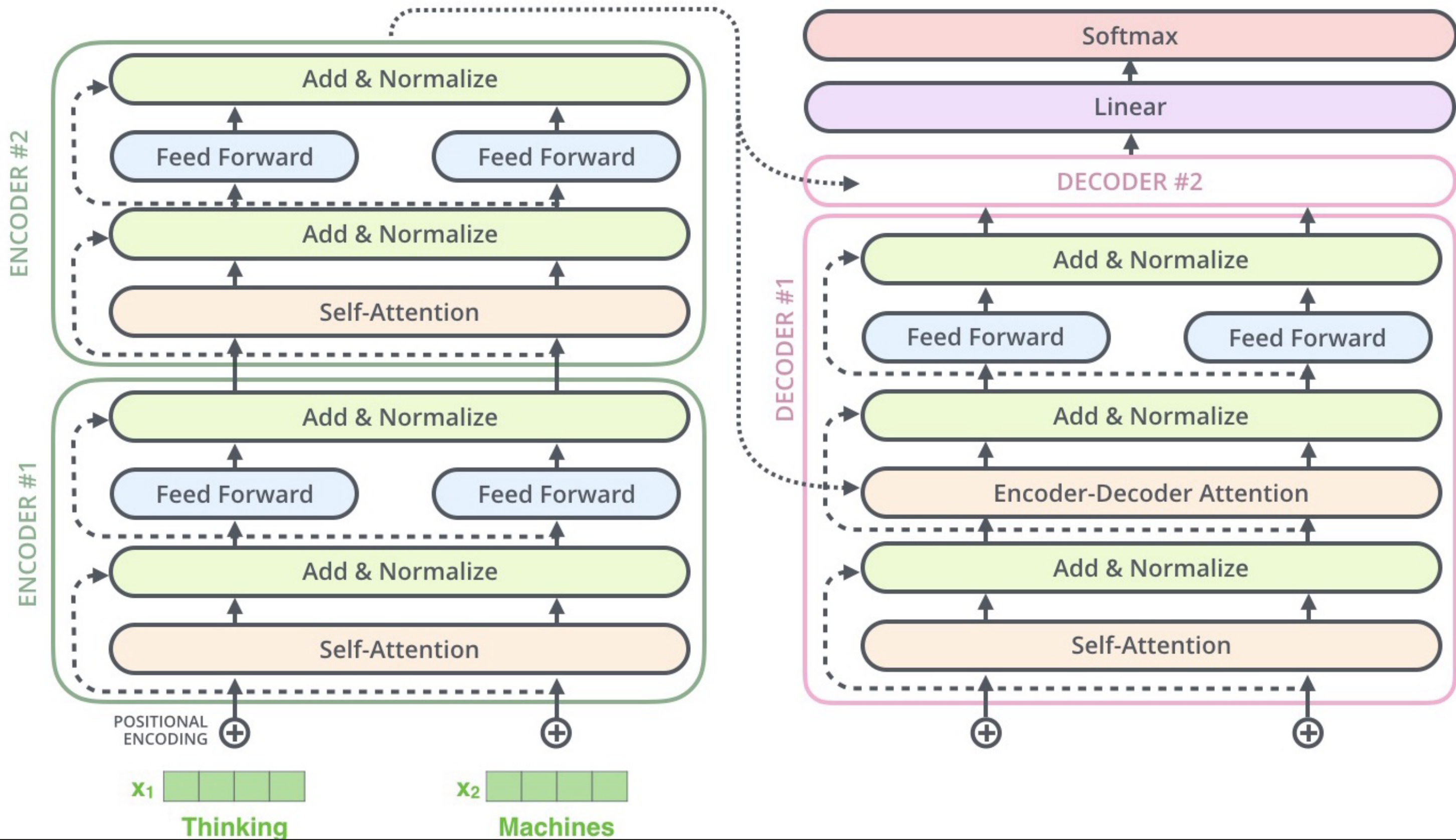
$(\alpha_i)_j = 0$  for  $j > i$  to ensure we only predict using historical info.

## 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  $X \in \mathbb{R}^{n \times d}$  with  $X + P$

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

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

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