

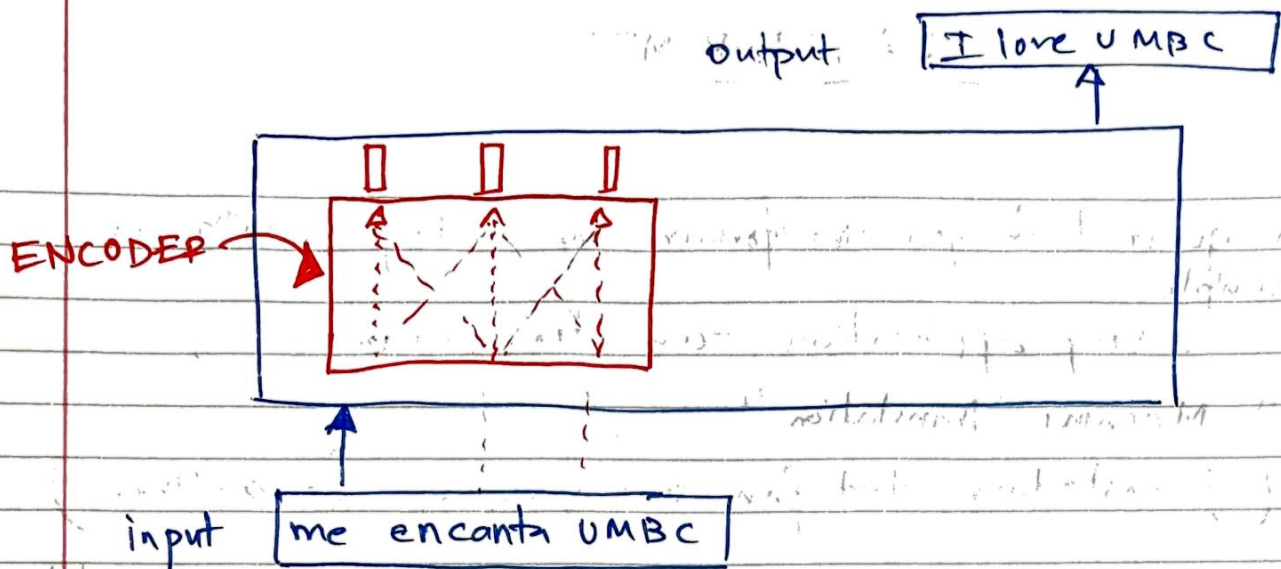
2025-04-14/16

TRANSFORMERS

CMSC 475 / 675 Neural Networks

(TEJAS GOKHALE)

- "Transformers" ~~is~~ is a ^(ENCODER-DECODER) neural network architecture
 - ↳ initially developed for language modeling (but also has other applications)
 - ↳ huge impact : The workhorse behind "LLM"s and AI chatbots that we all use
- This lecture :
 - ↳ a mathematical (linear algebra ?) approach to explaining how a transformer works
 - ↳ ~~for~~ Many blogs, videos, slides etc. on the web about transformers with good visualizations (use them to accompany the math or to build a better intuition from an application perspective)
 - ↳ My notes are derived from Soheil Feizi (UMD)
- * EXAMPLE APPLICATION: Machine Translation
 - ↳ want to translate from Spanish to English
 - "I love UMBC" $\longleftrightarrow$ "Me encanta UMBC"



we have an Encoder-Decoder architecture

what is the job of the encoder?

↳ Encode information for each word / token based on the context

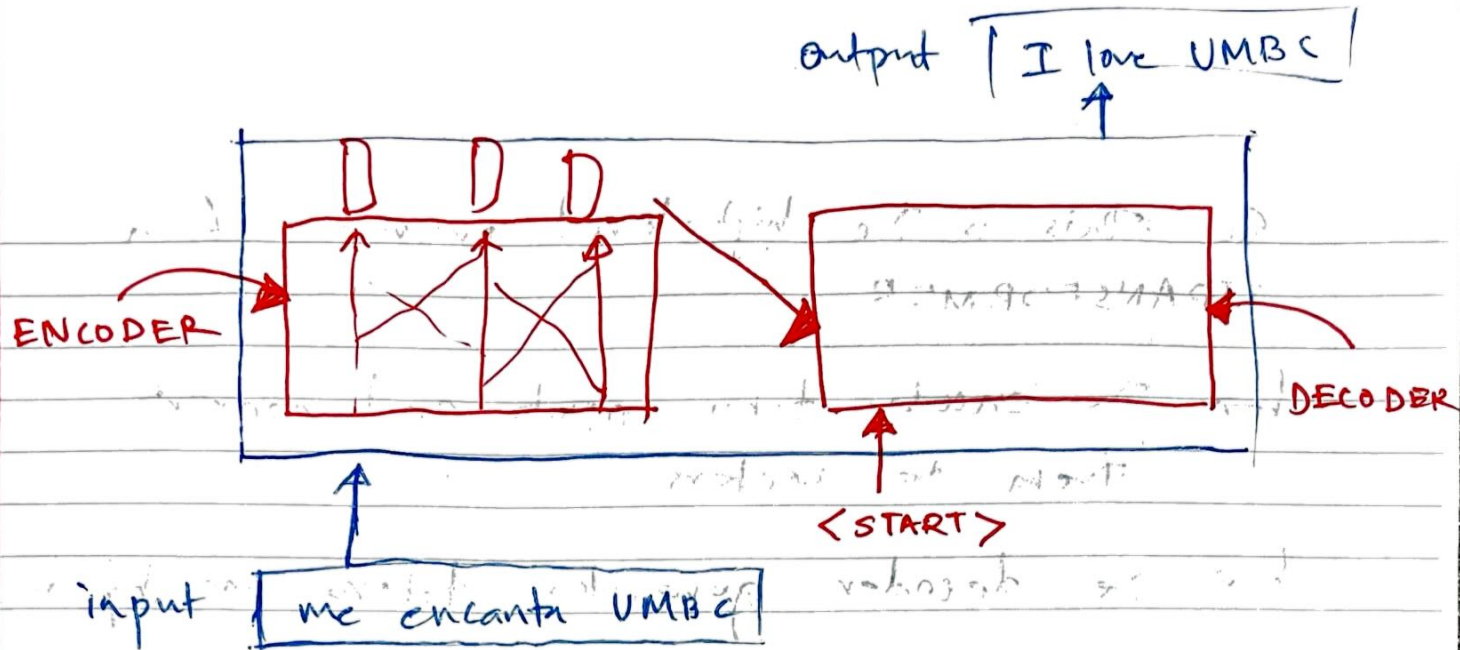
↳ for each word, the output of the encoder will be a VECTOR

↳ the vectors will have a particular length (a design choice)

(we have seen word embeddings before.
we will ~~learn~~ it see a new way
(transformers))

to obtain these embeddings

↳ we don't want single-word embeddings
we want to integrate context info



- the encoded vectors are then going to be sent to The DECODER
- The DECODER's job is to translate (do the task)
- The DECODER gets The **<start>** token as input and generate the next token

<u>Input</u>	<u>Output</u>	
< START >	" I "	STEP ①
< START > I	" love "	
< START > I love	" UMBC "	
< START > I love UMBC	* < STOP > < EOS >	

- The decoder keeps generating the next tokens one by one
 ↳ until the **< STOP >** token is generated

- So, This is the high-level overview of a TRANSFORMER

↳ The encoder takes input and converts them to vectors

↳ The decoder generates tokens one by one

↳ every step of the decoder depends on the previously generated tokens

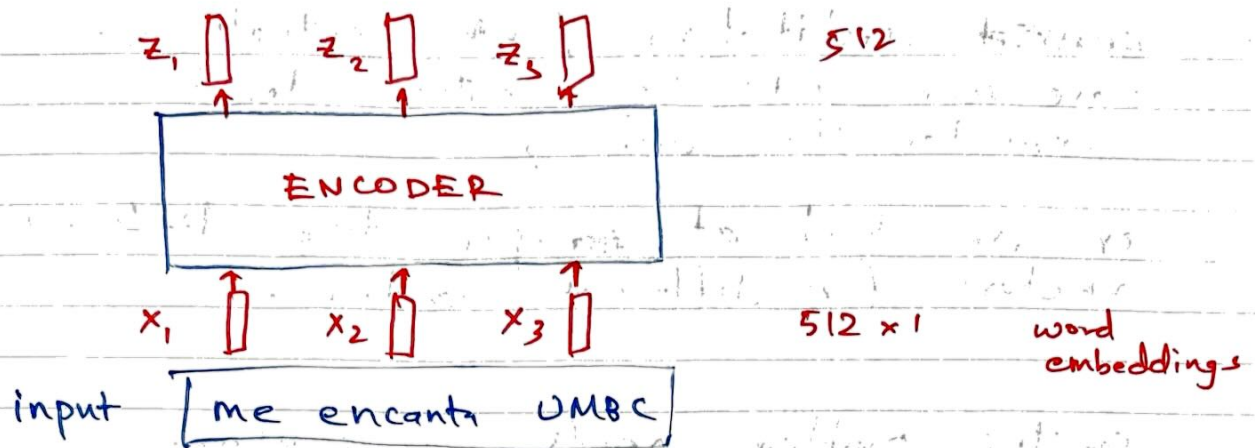
↳ that is why these models are also called
AUTO REGRESSIVE MODELS

So that's Encoding - Decoding at a high level

This lecture :

- ① HOW CAN WE encode information (with content) into vectors ?
- ② HOW CAN WE send that informations to the decoder ?
- ③ HOW CAN WE use it in autoregressive fashion to decode and generate output tokens.

* LET'S START WITH THE ENCODER



- by now, we have seen that we can't directly input words to a neural network

↳ we need to convert them to vectors

↳ "embedding"

{ input of the encoder: 512×1 word embeddings x
 { output of the encoder: 512×1 - word embeddings z

... OK SO WHY DO WE NEED THE ENCODER?

↳ to encode CONTEXT !

Q: what if I change the order of x_1, x_2, x_3 ?

A: z_1, z_2, z_3 should change
 (the vectors should be different)

Tejas loves UMBC
 UMBC loves Tejas } different meaning!
 $\therefore z_1, z_2, z_3$ should be different!

- ~~current~~ architectures / LM methods we have seen until now either don't consider the order, or are bad at ~~the~~ generating different vectors for different orders.

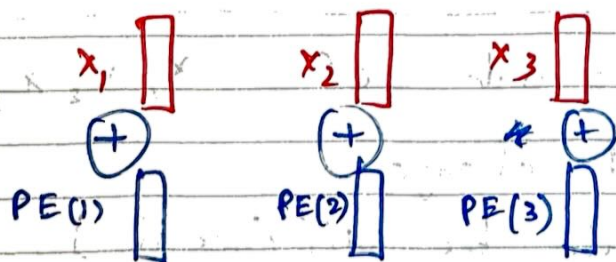
- ~~so~~ This problem needs a solution.

- the Transformers Paper (2017) Vaswani et al. provided that answer in an effective way.

* IDEA: add "position vectors" to the input

POSITION EMBEDDINGS
Some vectors unique to each position

↳ instead of x_1, x_2, x_3 as input



Size 2 512 x 1 (so that we can add)

what should the PE vectors be?

↳ vectors that are unique for each position

↳ Vaswani et al. used sine function

$$PE(pos, 2i) = \sin\left(\frac{pos}{10000^{2i/d}}\right)$$

$$PE(pos, 2i+1) = \cos\left(\frac{pos}{10000^{2i/d}}\right)$$

pos : position in the input

d : ~~size~~ ^{dimension} of vector (512)

i : index ~~of~~ over vector dimension

(if $d = 512$, $i \in [0, 255]$)

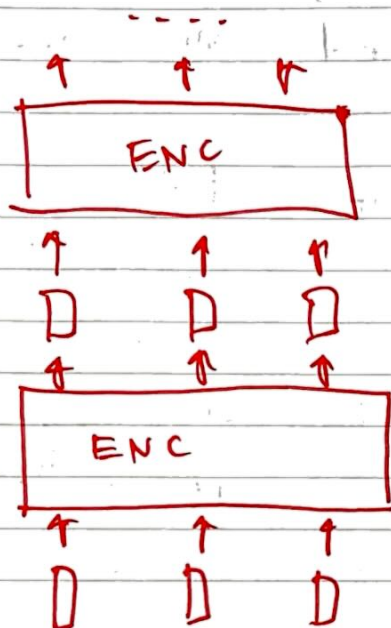
10000 : a large base

(Wavelengths : geom progression from 2π to 20000π)

→ smaller i : more oscillations
larger i : ~~less~~ less oscillations

* WE CAN STACK ENCODERS

(just like we can stack
conv/ layers in a CNN
linear)



(more layers)

(more parameters)

(model very
complex functions)

Vaswani et al.

stack 6 encoders

(hyperparameter)

* OK FINE, but how does the
ENCODER actually work?

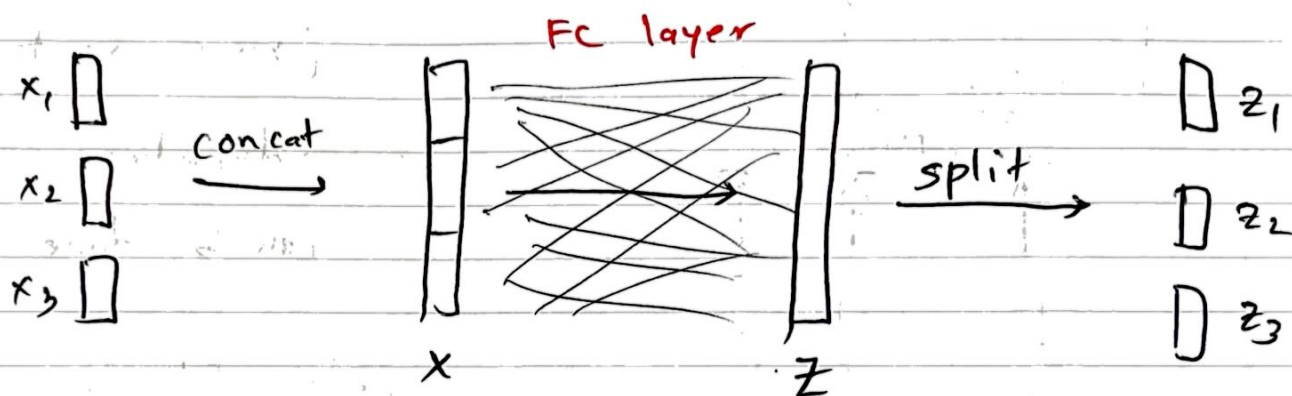
how can we incorporate context?

* ENCODER DETAILS :

- what is the simplest NN layer we know?

↳ Fully-Connected layer

- one idea

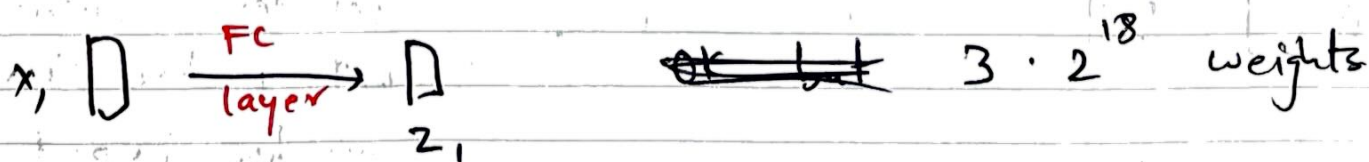


$$(3 \times 512) \times (3 \times 512) = 9 \cdot 2^{18} \text{ weights}$$

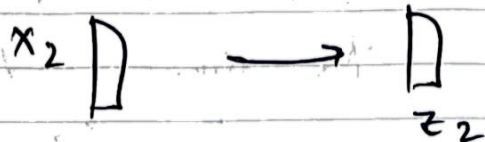
~~$(3 \times 512) \times (3 \times 512)$~~

- another idea

ONE (SMALL N.N.) per token



~~OK but~~ $3 \cdot 2^{18}$ weights



OK but context?

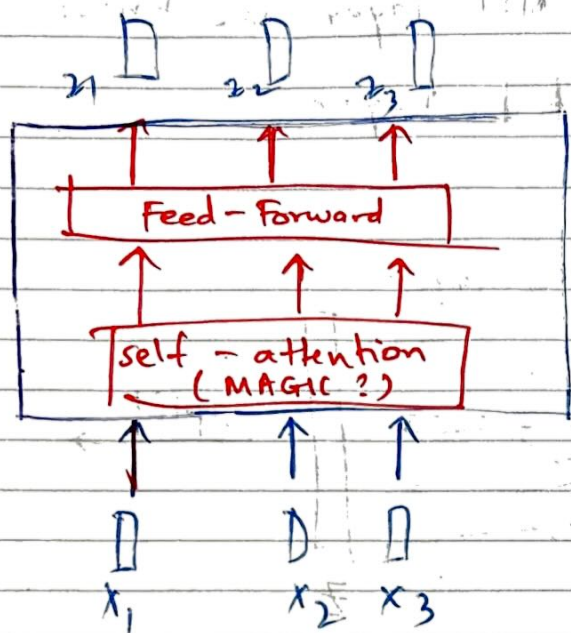
this is dumb?



what can we do?

ATTENTION LAYER

(SELF)
 * ATTENTION LAYER allows us to ENCODE CONTEXT



Feed Forward

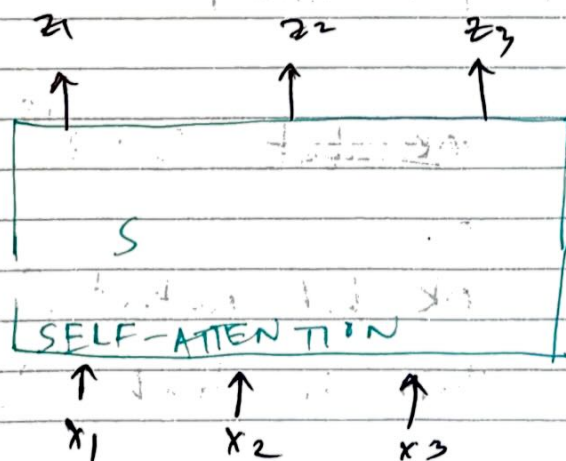
- applied to each vector separately
(no cross-talk)
- advantage:
can be done in parallel

* Self-Attention

- something that incorporates context information

What are some choices for the Self-Attention Layer

OPTION
 ①



Consider a matrix P with probabilities of how important each word is to the other word

$$P = \begin{matrix} & \begin{matrix} x_1 & x_2 & x_3 \end{matrix} \\ \begin{matrix} x_1 \\ x_2 \\ x_3 \end{matrix} & \begin{bmatrix} 0.8 & 0.2 & 0 \\ 0.1 & 0.7 & 0.2 \\ 0.3 & 0.3 & 0.4 \end{bmatrix} \end{matrix}$$

example:

* P is the attention matrix

↳ every row denotes how important each token is to one of the words

(a probability vector that sums to 1)

$$\text{Output: } \begin{cases} z_1 = 0.8 v_1 + 0.2 v_2 + 0.0 v_3 \\ z_2 = 0.1 v_1 + 0.7 v_2 + 0.2 v_3 \\ z_3 = 0.3 v_1 + 0.3 v_2 + 0.4 v_3 \end{cases}$$

where

$$v_1 = x_1 W^v$$

$$v_2 = x_2 W^v$$

$$v_3 = x_3 W^v$$

F-C layer

$$\therefore Z = P \cdot V$$

$$\begin{bmatrix} \boxed{z_1} \\ \boxed{z_2} \\ \boxed{z_3} \end{bmatrix}$$

$3 \times \cancel{512} d$

$$\begin{bmatrix} P \end{bmatrix}$$

3×3

$$\begin{bmatrix} \boxed{v_1} \\ \boxed{v_2} \\ \boxed{v_3} \end{bmatrix}$$

$3 \times \cancel{512} d$

* Learnable parameters : W^v, P

Any issues?

P is a $N \times N$ matrix

↳ N is the number of input tokens

↳ if $N = 10000$, P has 10^8 parameters

↳ N is variable length (depending on input length)

NOT GOOD.

need a better idea

IDEA

we use two intermediate matrices

W^k
↓

keys

W^q
↓

queries

- ~~For~~ Multiple x_i with W^q and so on

$$x_i W^q = q_i$$



$$x_i W^k = k_i$$



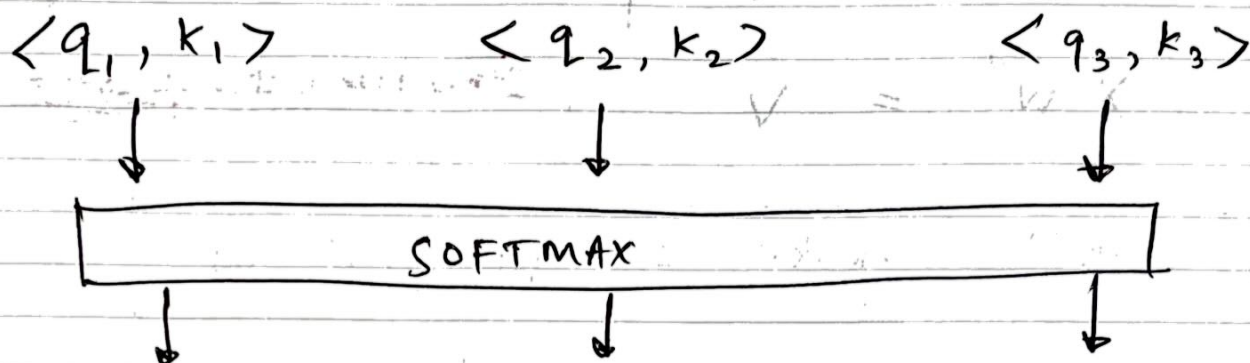
similarly for $x_2, x_3, \dots$

$$x_2 W^k = k_2 \quad x_3 W^k = k_3$$



find similarity between ~~a~~ "query embedding"
and each "key embedding"

— and convert to a probability



first row of P (corresponding to q_1)

Practical + scale inner product with sqrt of
key dimension d_k
before taking softmax

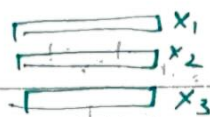
~~DEFINITION~~: ~~Attention~~ ~~(q, k)~~ Softmax

$$\text{SOFTMAX} \left(\left[\frac{\langle q_1, k_1 \rangle}{\sqrt{d_k}}, \frac{\langle q_1, k_2 \rangle}{\sqrt{d_k}}, \frac{\langle q_1, k_3 \rangle}{\sqrt{d_k}} \right] \right)$$

and similarly for q_2, q_3 .

Combined

X



$$X W^Q = Q$$

$$X W^K = K$$

$$X W^V = V$$

$$P = \text{softmax} \left(\frac{Q K^T}{\sqrt{d_{\text{key}}}} \right)$$

~~ATTENTION (Q, K, V)~~

Learnable Parameters : W^Q, W^K, W^V

no dependence on # tokens

$$Z = P \cdot V$$

$$= \text{softmax} \left(\frac{Q K^T}{\sqrt{d_{\text{key}}}} \right) V$$

$$Z \doteq \text{ATTENTION}(Q, K, V)$$

What if P is diagonal (or close to diagonal) ?

$$P = \begin{bmatrix} \text{diag} & 0 & 0 & 0 \\ 0 & \text{diag} & 0 & 0 \\ 0 & 0 & \text{diag} & 0 \\ 0 & 0 & 0 & \ddots \end{bmatrix}$$

↳ This means there is no cross-talk

- how do we avoid this ?

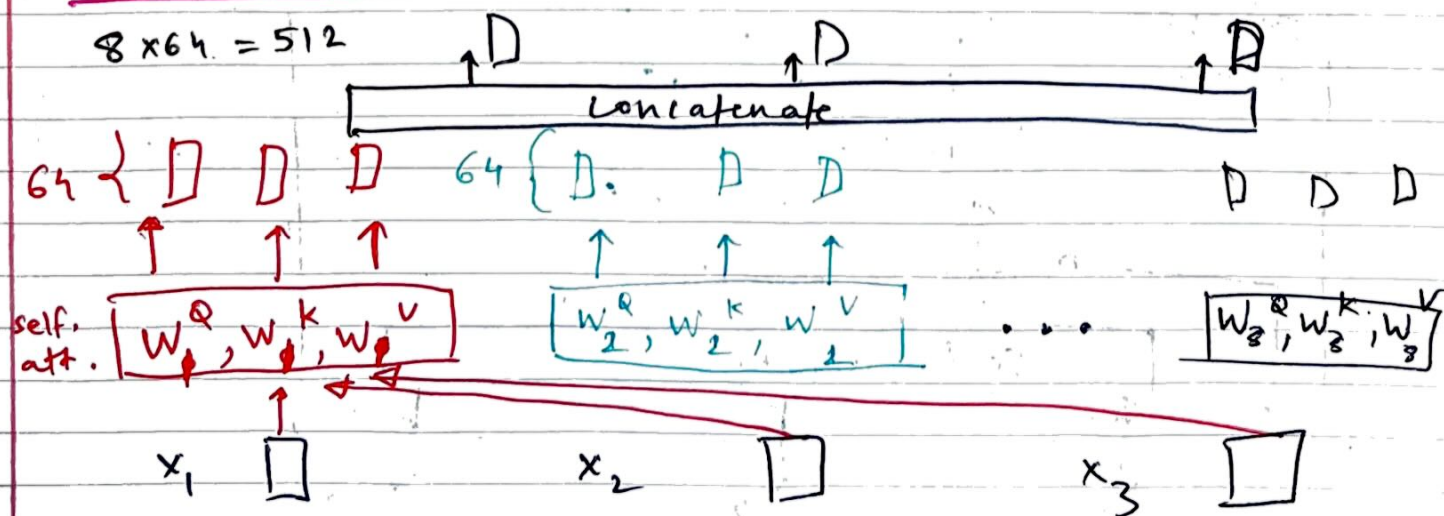
↳ apply self-attention multiple times
(with diff. parameters)

↳ concatenate

- instead of output 512 for self-attention
create 8 ~~64~~ self-attention "heads"
of 64 dimension

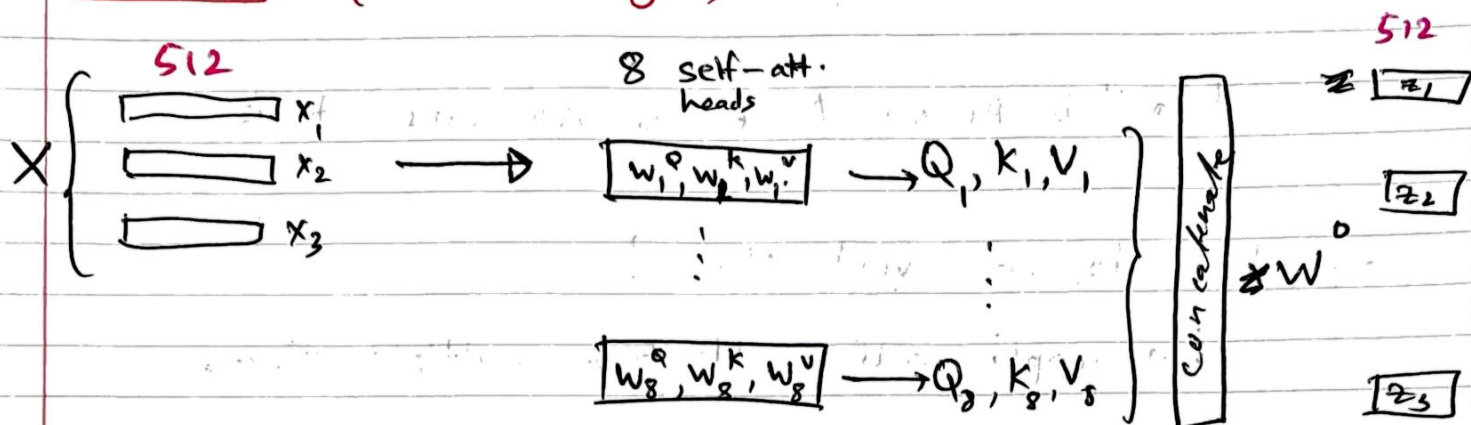
MULTI-HEAD ATTENTION

$$8 \times 64 = 512$$



in the paper they multiply with another matrix W^O

Summary (Overview Figure)



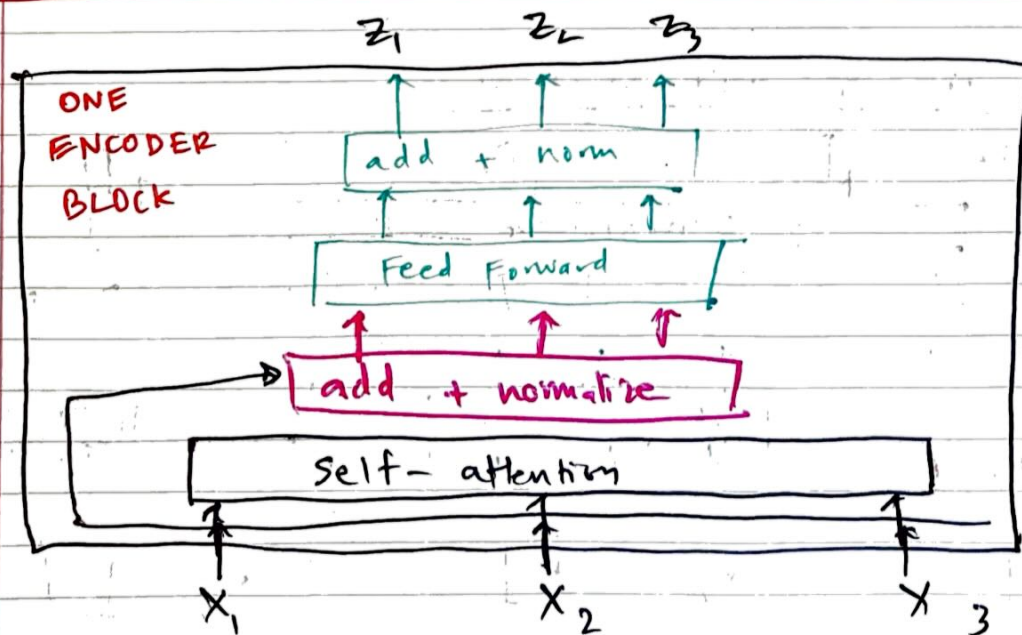
one final design element:

add $x + z$

(like res net skip connection)

apply layer normalisation

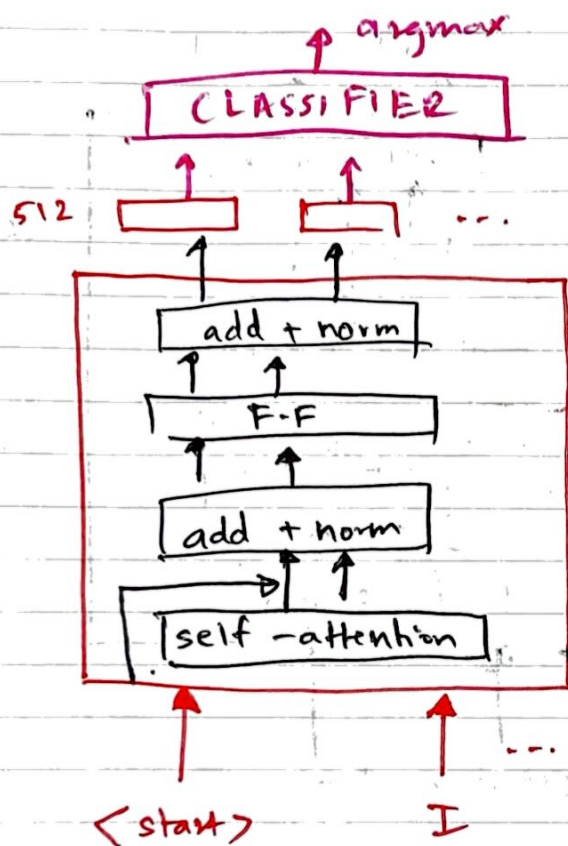
(make each row 0 mean, unit var)



multiple such encoder blocks stacked on top of each other

* NOW THE DECODER

- The decoder works in an autoregressive fashion
- $\hookrightarrow$ previously generated tokens are the input for generating the next word



We're going to use the same architecture as the encoder

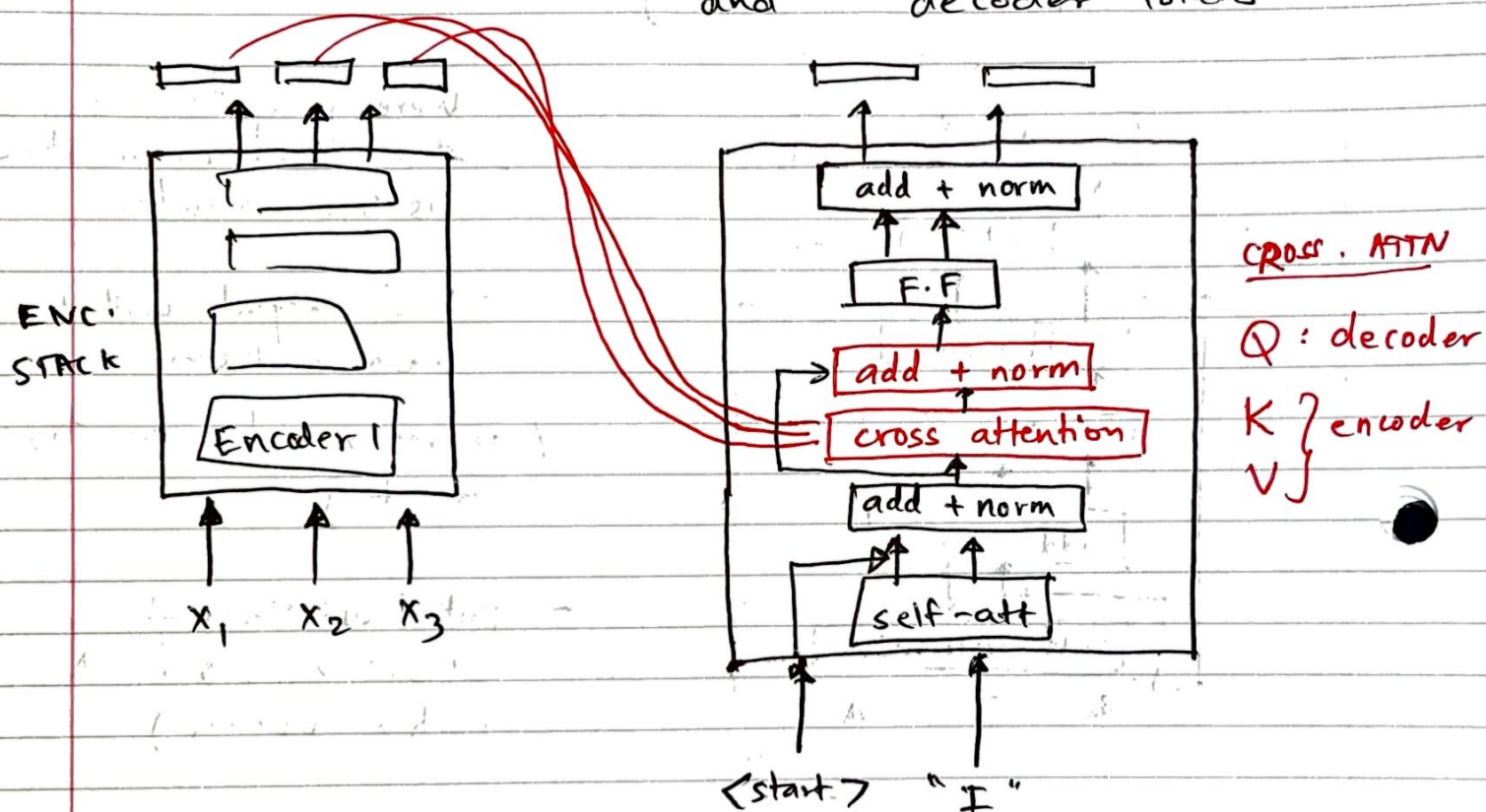
- the output will be two vectors (as many vectors as the input)
- put them through a classifier (over the entire vocab)

* ANY ISSUES ?

- ① in self-attention, maybe don't consider a very long context
 $\hookrightarrow$ easy solution (zero out top ~~right~~ ^{diagonal} of P matrix)
- ② what about the task?
 how are we ensuring that the output of 'decoder' is correct?
- ③ we are not using the encoder output!

- We should get the encoder output to talk to my decoder.

"CROSS-ATTENTION" between encoder tokens and decoder tokens



* LINEAR ATTENTION

If you have N ~~input~~ ^{input}, what is the size of the attention matrix?

$$N \times N$$

↳ for large N , this is really expensive.

$$O(N^2)$$

Q: How can we make this computation cheaper.
Linear in N ? $O(N)$

Katharopoulos (ICML 2020)

NOTATION

$$X \in \mathbb{R}^{N \times d}$$

(N tokens, each token $d \times 1$)
eg. 512

$$T : \mathbb{R}^{N \times d} \rightarrow \mathbb{R}^{N \times d}$$

the transformer is a mapping from N tokens to N tokens

— defined by L layers $T_1, \dots, T_L$

$$T_\ell(x) = f_\ell(A_\ell(x) + x)$$

F.F and
layer norm

self-attn

residual

and formula for Attention:

$$Q = XW^Q \rightarrow \in \mathbb{R}^{d \times d}$$

$$K = XW^K$$

$$V = XW^V$$

$$\begin{matrix} \downarrow & \downarrow \\ \in \mathbb{R}^{N \times d} & \in \mathbb{R}^{N \times d} \end{matrix}$$

$$P = \text{softmax} \left(\frac{QK^T}{\sqrt{d}} \right)$$

$$\in \mathbb{R}^{N \times N}$$

$$A(x) = P \cdot V = V' \quad (1)$$

we should improve the formula for P

(use some other similarity measure instead of inner product)

— let's look at the i^{th} row of V'
(embedding for i^{th} token)

$$V'_i = \frac{\sum_{j=1}^N \text{sim}(Q_i, K_j) V_j}{\sum_{j=1}^N \text{sim}(Q_i, K_j)} \quad (2)$$

query for i^{th} token

$$(1) \equiv (2) \quad \text{if} \quad \text{sim}(q, k) = \frac{\exp(q^T k)}{\sqrt{d}}$$

but we don't have to use this formula

↳ we can use some other formula

↳ what other similarity / distance measure can we use?

"KERNELS"

instead of inner-product betⁿ vectors,
we inner-product betⁿ some transformation
of vectors.

$$\text{kernel}(\vec{q}, \vec{k}) = \phi(\vec{q})^T \phi(\vec{k})$$

↙ transformation

$$v_i' = \frac{\sum_{j=1}^N \phi(q_i)^T \phi(k_j) v_j}{\sum_{j=1}^N \phi(q_i)^T \phi(k_j)}$$

↳ we have decomposed similarity as a product

$$= \frac{\phi(q_i)^T \sum_{j=1}^N \phi(k_j) v_j}{\phi(q_i)^T \sum_{j=1}^N \phi(k_j)}$$

$O(N)$

do it once and store
(doesn't depend on i)
same for every token

Computational Complexity?

(for computing attention matrix)

goes from $O(N^2)$ to $O(N)$

but what is the dimensionality of $\phi()$?

features $\phi()$ usually have bigger dimensions than the vector

(eg. polynomial kernel of degree 2

increases dimension from r to r^2)

so we should use suitable kernels

(if we want to reduce complexity)