

Transformers & Attention Mechanism

- Used in SOTA NLP like ChatGPT
- Used in pretty much everything now
- permutation-equivariant architecture!
- actually a type of GNN
- like an MLP w/ input-dependent weights
$$X \rightarrow W(X) \cdot X$$

→ much more powerful in principle!

{ - Key idea: (self) attention mechanism
- Key idea: learning embeddings
→ learn relations b/w elements of sequence
relations expressed by dot product
in embedding space

NLP example:

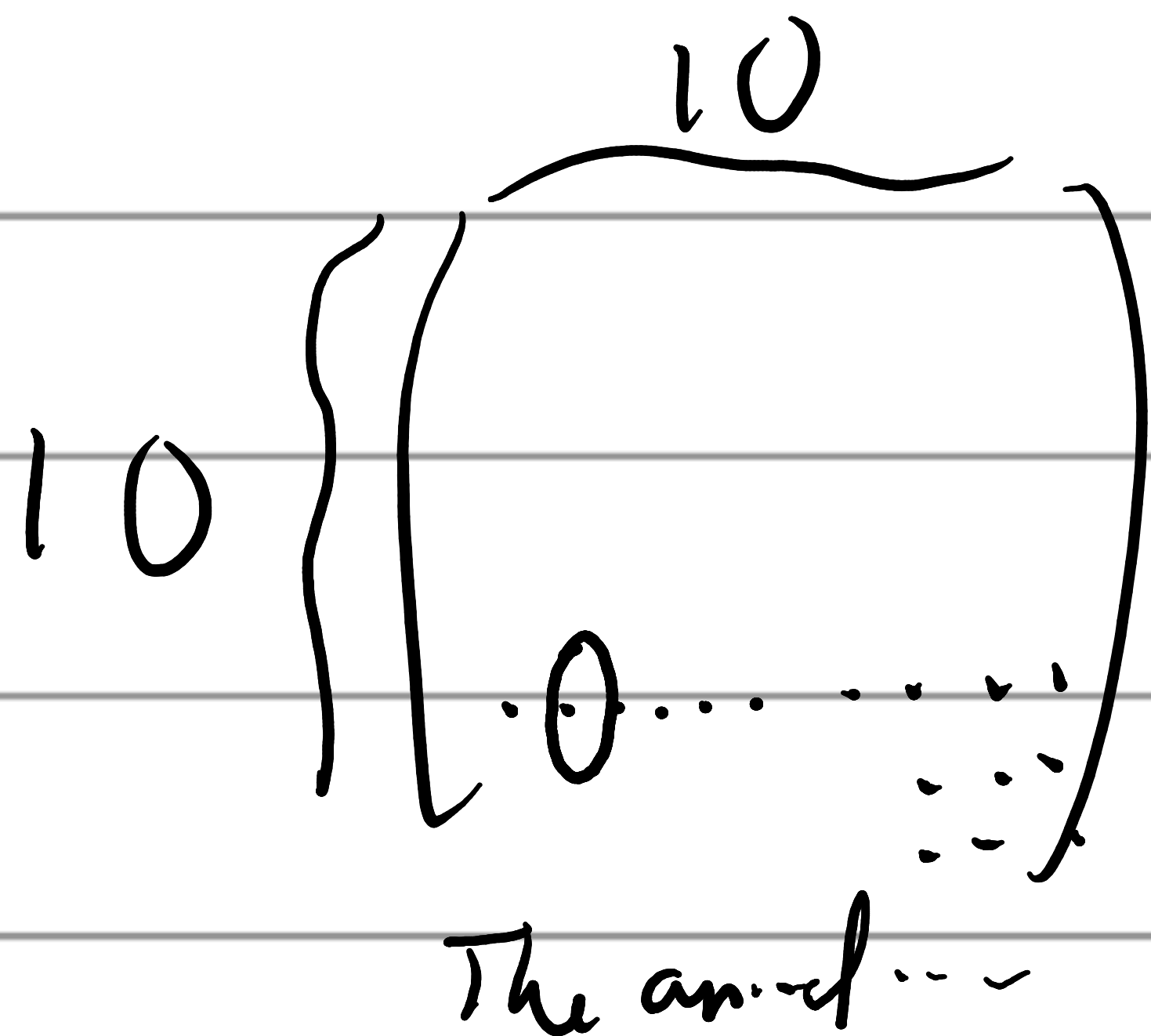
"The animal didn't cross the street because it was tired."

- What does "it" refer to?

Clearly "animal" not "street"!

- Earlier NLP models had trouble w/ this (trouble being long-term relations)

- Attention matrix



The animal
,
it was
tired

each row
list of probs
sum to 1

$p_{i1}, \dots, p_{i10}$

$$\sum p_i = 1$$

→ Maybe w/

p_{72} and p_{77}
large, others small

The Attention Mechanism (Bahdanau 2014)

(Vaswani et al 2017)

Want to map input
sequence to embed space
 $\mathbb{R}^d \rightarrow \mathbb{R}^{d'}$

so that relations are captured by dot product

$$y_i \cdot z_j \quad y, z \in \mathbb{R}^{d'}$$

want probabilities, use softmax

$$p_{ij} = a(y_i \cdot z_j) \quad \text{"attention matrix"}$$

if y, z come from same input $\rightarrow$ "self-attention"

) Simplest map: linear

$$\begin{cases} X \rightarrow X W^Q \leftarrow d \times d' \text{ matrix} \\ \quad \quad \quad = Q \text{ "query"} \\ X \rightarrow X W^K = K \text{ "key"} \end{cases}$$

We can apply attention matrix
to input sequence to produce a new sequence

$$\sum_{i=1}^N p_{1i} \vec{X}_i$$

$\vdots$

$$\sum_{i=1}^N p_{Ni} \vec{X}_i$$

if we wet in embeddy space
can learn another map

$$X \rightarrow XW^V = V \text{ "values"}$$

) Transformer architecture

sequence $\rightarrow$ sequence mapping

Recap:

Given sequence of vectors $\in \mathbb{R}^d$

$$X = \begin{pmatrix} \vec{x}_1 \\ \vec{x}_2 \\ \vdots \\ \vec{x}_n \end{pmatrix} \in \mathbb{R}^{n \times d}$$

Learn 3 maps to embeddng space $\mathbb{R}^{d'}$

$$\left. \begin{array}{l} \text{Query } Q = X \cdot W^Q \\ \text{Key } K = X \cdot W^K \\ \text{Value } V = X \cdot W^V \end{array} \right\} \begin{array}{l} d \times d' \text{ matrices} \\ \text{learned weights} \end{array}$$

self attn matrix:

$$P = a(Q \cdot K^T) \quad N \times N \text{ mtr}$$

Transformer output:

$$X \rightarrow Y = P \cdot V \in \mathbb{R}^{n \times d'}$$

* Note: permutation equivariant!

$$X = \begin{pmatrix} \vec{x}_1 \\ \vdots \\ \vec{x}_n \end{pmatrix} \in \mathbb{R}^d \rightarrow Y = \begin{pmatrix} \vec{y}_1 \\ \vdots \\ \vec{y}_n \end{pmatrix} \in \mathbb{R}^{d'}$$

interchange $\vec{x}_i \leftrightarrow \vec{x}_j \iff$ interchange $\vec{y}_i \leftrightarrow \vec{y}_j$

can check w/ indices

$$y_{i\alpha} = a \left(x_{ia} w_{a\beta}^Q (w_{\beta b}^K)^T (x^T)_{bj} \right) x_{jc} w_{c\alpha}^V$$

perm acts here sum perm int

In fact can argue for this architecture
as simplest non-linear permutation-equiv.

seq $\rightarrow$ seq map

"Multi-headed Attn"

in example "it" also relates to "tired"

how to capture that?

could overload single attn matrix

but more powerful to learn an attention over!

• $W^{Q,K,V}$
 $r=1, \dots, N_h$

analogous to multiple filters
in CNN!

• each gives $N \times d'$ embedded sequence
 Z_r

• concatenate $(Z_1, \dots, Z_{N_h})$ $N \times \underbrace{(d' \cdot N_h)}$

• feed it all together w/ final matrices

$$W^{Q,K,V} \quad (d' \cdot N_h) \times d''$$

$$Z_{\text{out}} = Z \overset{Q,K,V}{W}^{Q,K,V}$$

"Positional Encoding"

- in NLP & other sequences, position does, extra info!
- add back in positional info w/ positional encoding

$$\vec{X}_i \rightarrow \vec{X}_i + \vec{pos}(i) \quad \text{common exaple}$$

$\sin(\omega_a i)$
 $\cos(\omega_a i)$

- transformer still perm equiv
but not wrt original input sequence

"Tokenization"

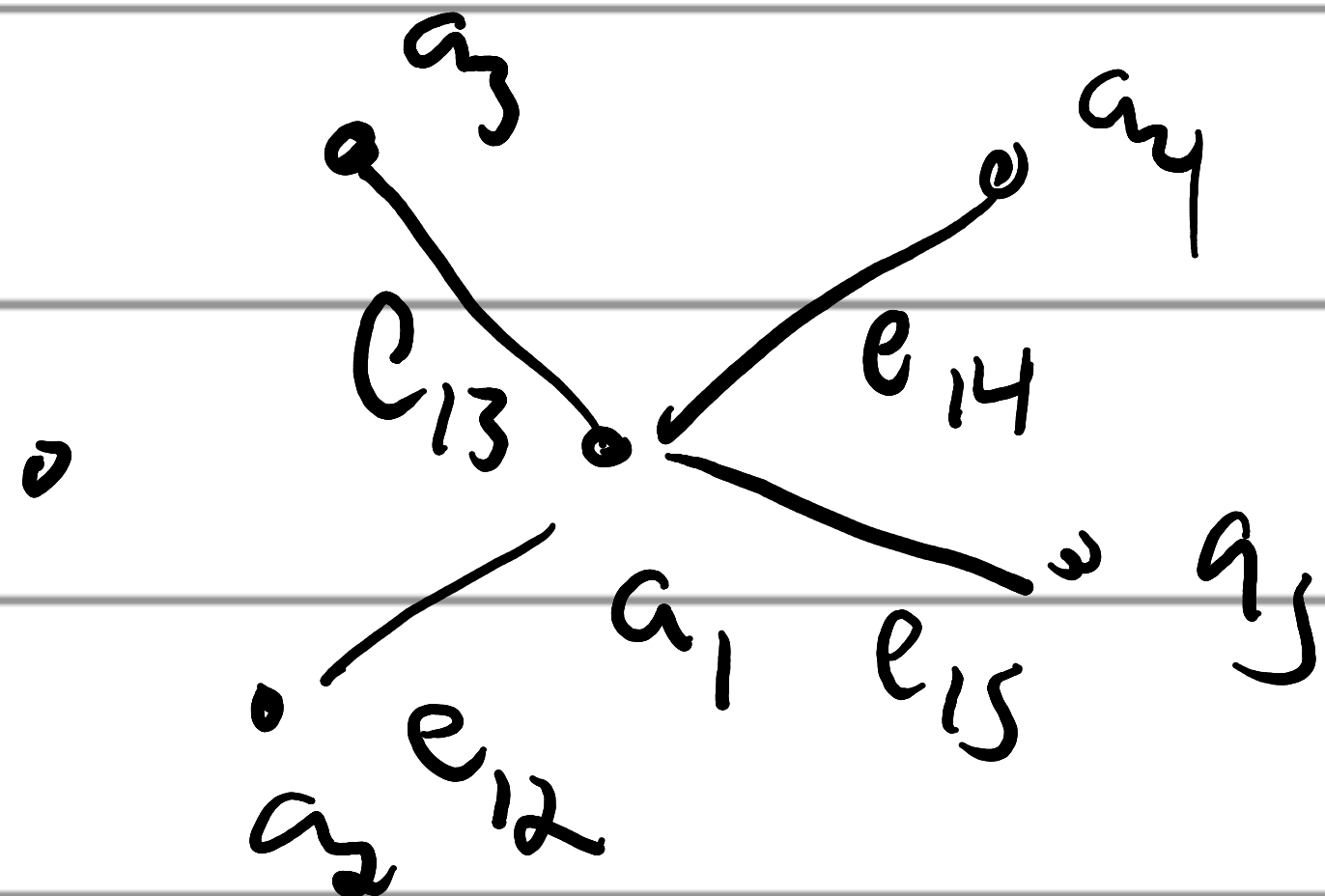
In NLP need to map words to vectors
"tokens"

many tokenization algorithms

often using ML - pretraining task

Transformers as GNNs

Recall GNN



MLPs

$$a_i' = \phi^{MLP} / a_i, \rho(e_{ij}', e_{ij})$$

$$e_{ij}' = \phi^{edge}(a_i, a_j, e_{ij})$$

Transformer

this is like e_{ij}

$$X_i' = \text{Attn} \left(p_{ij} X_j \right) \leftarrow \text{this is } a_j$$

$\hookrightarrow \text{attn}(X_i, X_j)$

~ Fully connected graph

nodes are X_i

edges are $\text{Attn}(X_i, X_j)$

} multiply
not usually
considered

vis a vis MLPs

So Transformers still go beyond
standard GNNs

Summary:

single self attn:

$$Z = \text{softmax}(QK^T) \cdot V$$

$$Q = XW^Q$$

$$K = XW^K$$

$$\text{multi-head attn: } (z_1, \dots, z_h) W^O = Z$$

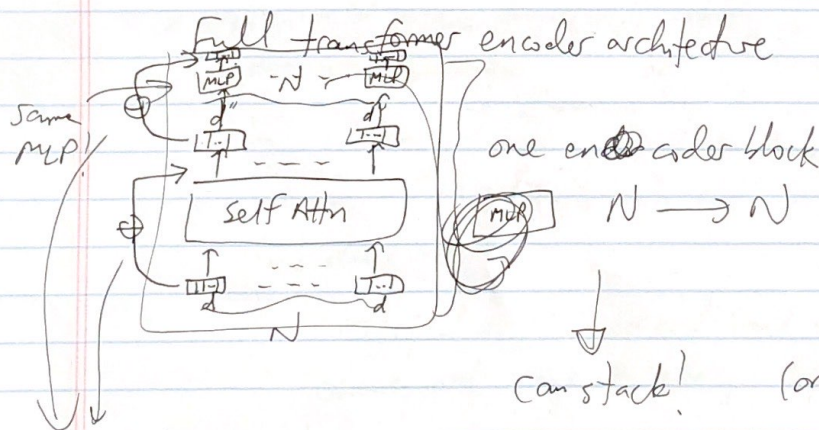
$$V = XW^V$$

- can check: perm equivariant! interchanging order of rows of X will result in same output w/ output rows interchanged.

- "MLP w/ input-dependent weights"

$$Z \sim (XW^Q(W^K/\sqrt{d})^T X) \cdot XW^V$$

this is like ~~an~~ X -dep weight



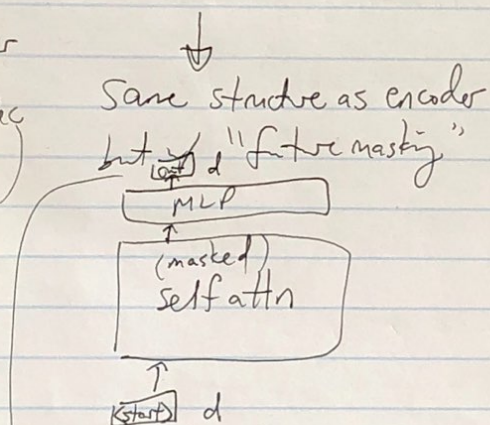
+ residual/skip connections!

The encoder produces a transformed embedding that takes into acct entire input & all its ^{inter}relations

Transformer Decoder

To generate ~~that~~, need the decoder architecture

(original transformer papers had encoder-decoder pair w/ add'l encoder att'n layer...)



output is vector z w/ same dim as original embedding.

interpreted as "score" for

$$Z \cdot D \leftarrow d \left\{ \begin{pmatrix} | & | & | & \dots & | \end{pmatrix} \right\}$$

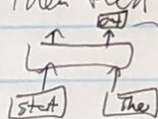
N_{vocab} dim'd vector

Multiply into vocabulary/dictionary embedding matrix (~50k for English)

↳ interpret as score for each word in vocab
choose word w/ highest score, or top-k etc.

eg. "The"

Then feed back into decoder



e.g. "Thing"

use ~~enc~~ embedding of last word (which knows about entire preceding sequence) to predict next word

encoders useful for learning embeddings

↓
classification / categorization
sentiment analysis

BERT example

"Bidirectional Encoder Representations from Transformers"
Google (2018)

- ~~Future~~ Future masky: don't want attention scores to depend on future words

$$\text{softmax}(\text{mask: } \begin{matrix} \text{<start>} \\ I \\ an \\ for \end{matrix} \begin{matrix} \text{<start>} I \text{ go fine} \\ \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot \end{matrix}) \cdot QK^T$$

- This is an example of "autoregressive model"

$$\text{Learning } P(x_1, \dots, x_n) = P(x_1) P(x_2 | x_1) P(x_3 | x_1, x_2) \dots P(x_n | x_1, \dots, x_{n-1})$$

- GPT is an example of this ~~encoder~~ decoder only model
Generative Pre-trained Transformer OpenAI (2018)

GPT $\rightarrow$ BERT $\rightarrow$ GPT2 $\rightarrow$...

now GPT-family (decoder only) are state of the art!

- Can also give "prompt" $\rightarrow$ just an initial sequence instead of <start>

~~Can~~ prompt can be diff language $\rightarrow$ translation!

- Both BERT & GPT use ^{self-supervised} pre-training to learn attn model

↙
"masked language modeling"

↘
"next word prediction"

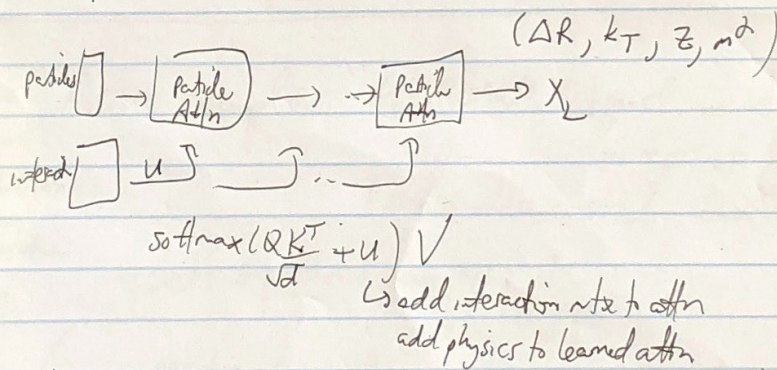
using huge amount of data & grant model
fine tuning on "downstream" ~~task~~ "Foundation model" "backbone"
tasks (eg translation, classification, ...)

Examples of applications of transformers

Qu, Li, Qian 2020.03772 "Particle Transformer for Jet Tagging"¹

- Introduced permutation equiv. arch. for jet tagging based on transformers
- Also introduced new dataset "JetClass" for training it
 (previous datasets ~1M) 100M jets $\approx$ 10 types $\times$ 10M each
 - 4 rec
 - particle ID (chg hadron, neutral hadron, e, μ , γ)
 $\hookrightarrow$ not truth, reco!
 - displacement

- particle features & pairwise "interaction" features $N \times N \times t_x$



- "class attn" (don't completely understand this)

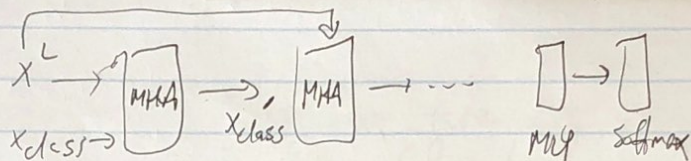
extract class, tier output - standard trick in vision transformers³

randomly initialized token

accumulates ~~attn~~ info from event from attention to other tokens

$$\begin{cases} Q = W_Q x_{\text{class}} \\ K = W_K [x_{\text{class}}, x_L] \\ V = W_V [x_{\text{class}}, x_L] \end{cases}$$

Class attn (from vision transformer lit)



starts off random but accumulates info

- Part achieved SOTA results on all jet tagging tasks!

example of "foundation model"

- Also interesting: pretrain + fine tuning $\gg$ train from scratch on smaller dataset
- Also transformer benefited more from pretrain than GNN!

- Astro example: "ASTROMER transf. based embedding for repr. of light curves"

cf TimeModAttn
example from
Astro presentation
- pretrain
on sim

Donoso - Oliva et al 2205.01677

- self supervised, like NLP etc
- positional encoding (not perm equiv)

→ masked light curve modeling
encoder-decoder

- Data R-band light curves of MACHO survey (Galactic Bulge & LMC) for pretrain - 1.5M light curves
- 20k labeled variable stars for f.t. & eval
- Also OGLE-II 360k labeled, ATLAS 420k