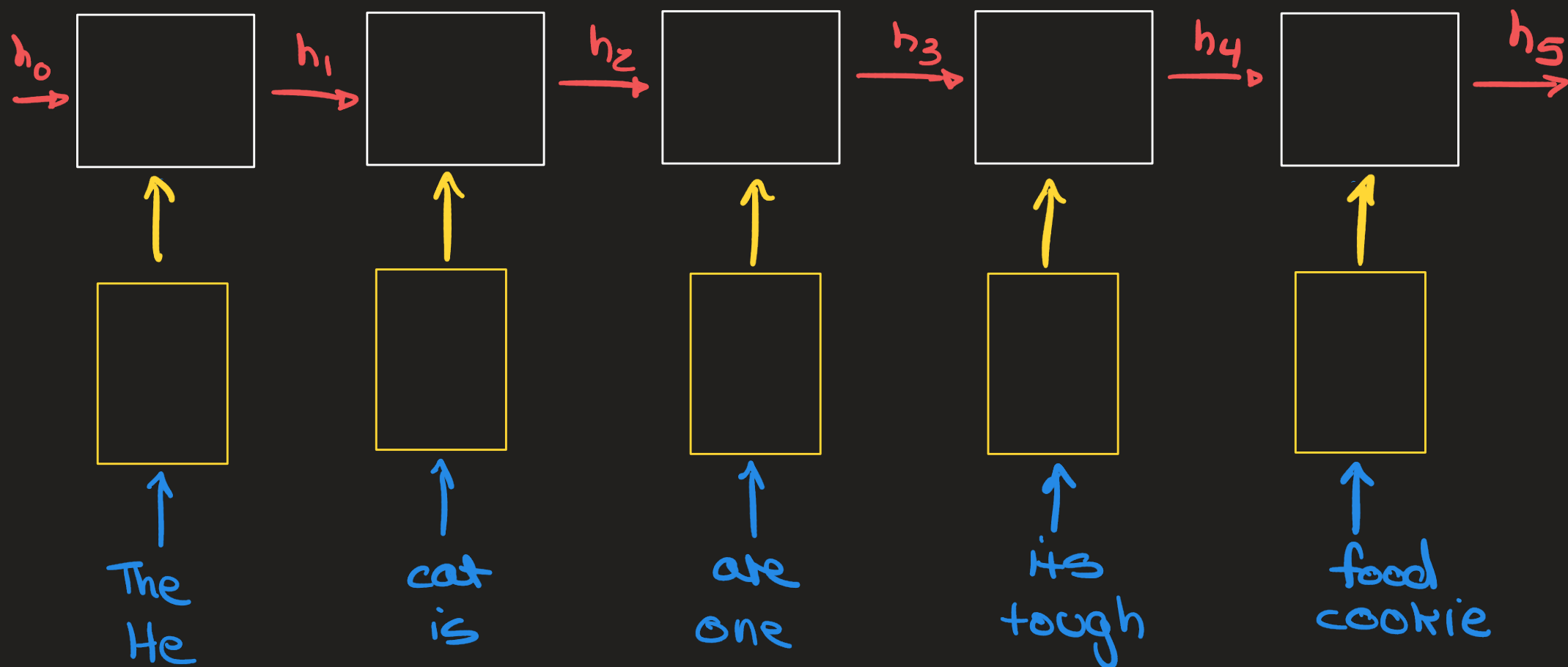


Machine Learning and Optimization Lecture 22

- Self-Attention as a replacement for many-to-many RNNs, for obtaining contextual representations
- Positional Encodings to counter the permutation invariance of Self-Attention
- The BERT model
- Unsupervised pretraining and supervised fine-tuning

Contextual Representations with RNNs

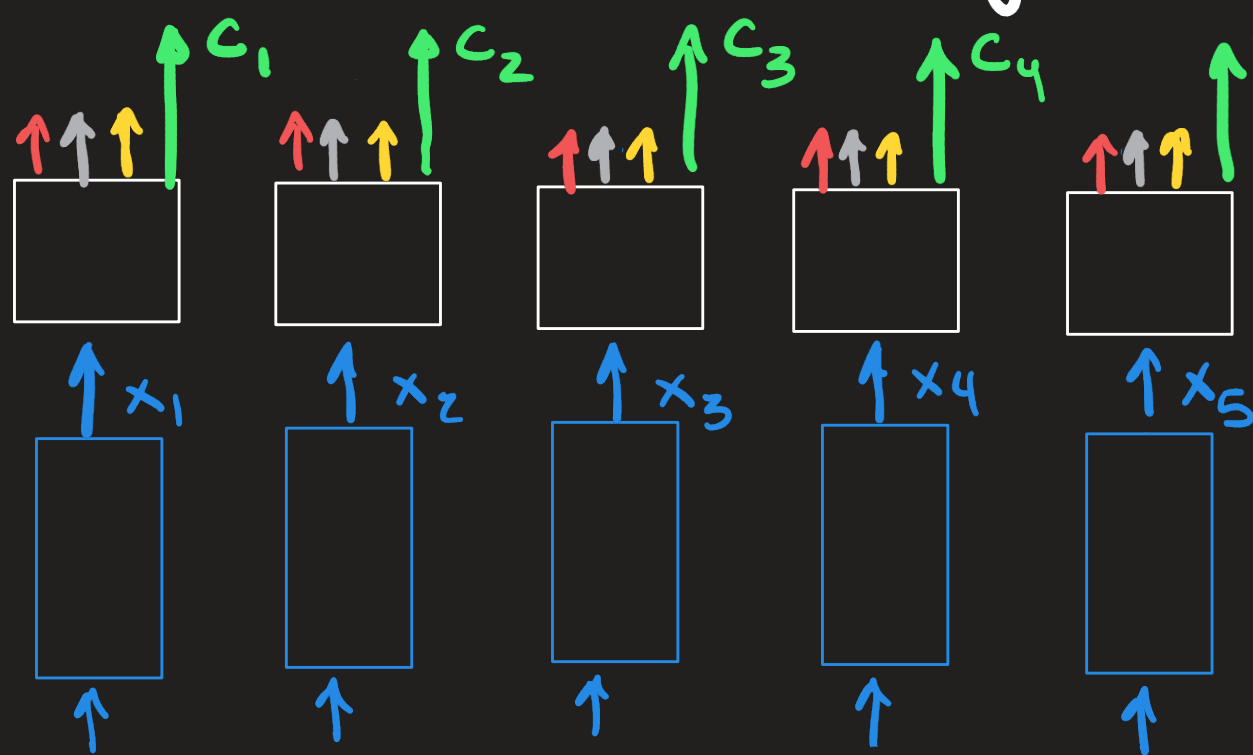
We saw that RNNs form contextual representations $h_1, \dots, h_4$ of the input sequence $x_1, \dots, x_4$ that are useful for downstream tasks.



When using RNNs we learn shared global **embeddings** for each token, which are turned into **contextual representations**

The drawbacks of using RNNs is that they are inherently serial (can be mitigated w/ bidirectional RNNs) and so they lose information about earlier portions of the sequence that are relevant for representing later tokens

Solution: use attention to get contextual representations



$$\left. \begin{aligned} k_i &= W_K x_i \\ q_i &= W_Q x_i \\ v_i &= W_V x_i \end{aligned} \right\} \begin{array}{l} \text{key, query,} \\ \text{value} \\ \text{representation} \end{array}$$

For the i th token's contextual representation

$$\alpha_i = \text{softmax} \left(\begin{bmatrix} k_1^T \\ \vdots \\ k_5^T \end{bmatrix} q_i \right)$$

(Note: you could add bias vectors to k_i, q_i, v_i but we typically don't)

$$c_i = \sum_{j=1}^5 (\alpha_i)_j v_j$$

Now the contextual representation of each token directly depends on each other token in the sequence

- $k_i = \omega_K x_i$: learn ω_K to extract information from the global embeddings about how relevant x_i is to determining contextual representations of other tokens
- $q_i = \omega_Q x_i$: learn ω_Q to extract information from the global embeddings about how relevant other tokens are in determining contextual representations of x_i
- $v_i = \omega_V x_i$: learn ω_V to extract information from x_i that is useful in forming contextual representations of other tokens
- α_i : the weights determining the relative importance of all tokens in forming the contextual representation of x_i

Self-Attention

Given a sequence of vectors $X = [x_1 | \dots | x_T] \in \mathbb{R}^{d_x \times T}$
and parameters

$$\omega_K \in \mathbb{R}^{d_{\text{attn}} \times d_x}, \omega_Q \in \mathbb{R}^{d_{\text{attn}} \times d_x}, \omega_V \in \mathbb{R}^{d_{\text{out}} \times d_x},$$

compute

$$K = \omega_K X = [k_1 | \dots | k_T] \in \mathbb{R}^{d_{\text{attn}} \times T}$$

$$Q = \omega_Q X = [q_1 | \dots | q_T] \in \mathbb{R}^{d_{\text{attn}} \times T}$$

$$V = \omega_V X = [v_1 | \dots | v_T] \in \mathbb{R}^{d_{\text{out}} \times T}$$

and note that

$$K^T Q = \begin{bmatrix} \frac{k_1^T}{k_1^T} \\ \vdots \\ k_T^T \end{bmatrix} [q_1 | \dots | q_T] = \begin{bmatrix} k_1^T q_1 & \dots & k_1^T q_T \\ \vdots & \ddots & \vdots \\ k_T^T q_1 & \dots & k_T^T q_T \end{bmatrix}$$

This implies that

$$\text{softmax}(\mathbf{K}^T \mathbf{Q}) = [\alpha_1 | \dots | \alpha_T] \in \mathbb{R}^{T \times T}$$

Therefore the contextual representations obtained through self-attention are

$$\begin{aligned} \mathbf{C} &= \mathbf{V} \text{softmax}(\mathbf{K}^T \mathbf{Q}) = [\mathbf{v}_1 | \dots | \mathbf{v}_T] [\alpha_1 | \dots | \alpha_T] \\ &= [(\alpha_1)_1 \mathbf{v}_1 + \dots + (\alpha_1)_T \mathbf{v}_T | \dots | (\alpha_T)_1 \mathbf{v}_1 + \dots + (\alpha_T)_T \mathbf{v}_T] \\ &= [\mathbf{c}_1 | \dots | \mathbf{c}_T] \end{aligned}$$

In practice, for numerical stability we take

$$\text{Attention}(\mathbf{X} | \omega_K, \omega_Q, \omega_V) = \mathbf{V} \text{softmax}\left(\frac{\mathbf{K}^T \mathbf{Q}}{\sqrt{d_{\text{attn}}}}\right)$$

One drawback of self-attention is that it is permutation invariant: if $X_1 = \Pi X_2$ where Π is a permutation matrix, then you can check that

$$\text{Attention}(X_1) = \Pi \text{Attention}(X_2)$$

This is not desirable in e.g. NLP because changing the order of words should change their contextual representations.

E.g. the sequences $\{\text{The, cat, are, its food}\}$ and $\{\text{The, food, are, its, cat}\}$ should have different contextual representations.

To ensure this, we add positional encodings to the inputs:

$$\tilde{X} = X + P = [x_1 + p_1 \mid \dots \mid x_T + p_T]$$

The original Transformer paper takes

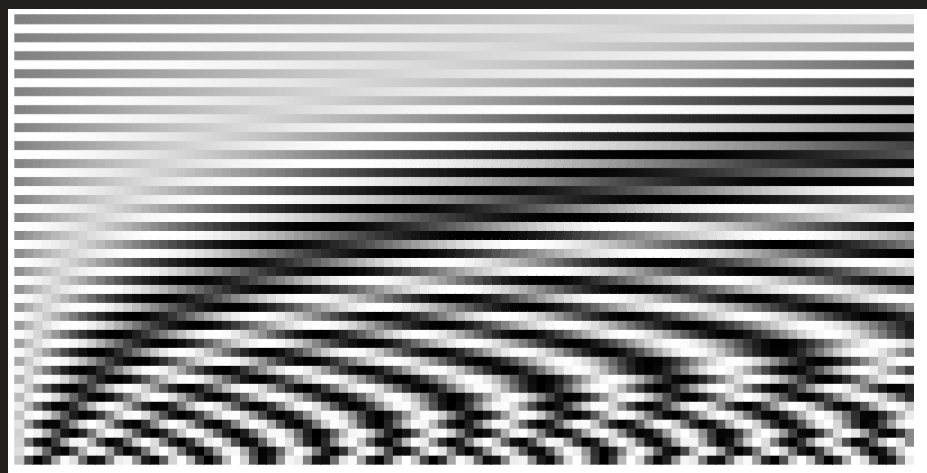
$$P = [p_1 | \dots | p_T]$$

with

$$(p_t)_{2i-1} = \sin(t\omega^{2i/d_x}) \quad \text{where} \quad \omega = \frac{1}{T}$$

$$(p_t)_{2i} = \cos(t\omega^{2i/d_x})$$

Ex: $(T=100)$



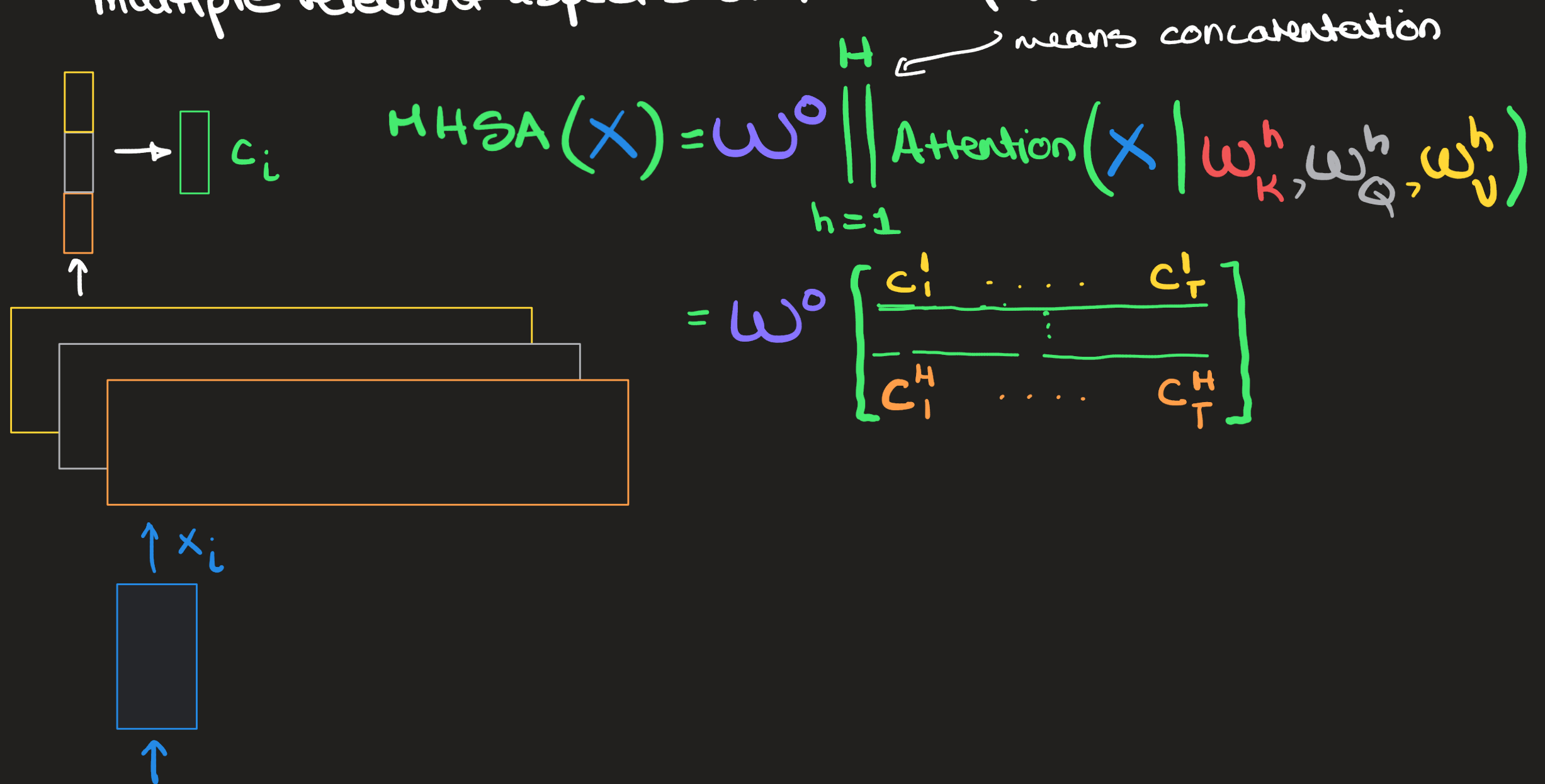
$p_1 \dots$

p_{100}

white corresponds to -1
black corresponds to 1

Multi-headed Self-Attention

Just as in CNNs, we prefer to learn multiple features (in this case, contextual representations) to extract multiple relevant aspects of the data:



Using Self-Attention in DNNs

DNNs transform their inputs into more useful/semantically relevant representations, by composing a series of layers.

MLP Layer

Input: a vector $[x_1, \dots, x_n]$

Output: a vector $[o_1, \dots, o_n]$

Operation: mix the vector's entries

linearly, $a = w x + b$, then

pass entrywise through nonlinearity

$$o = \sigma(a)$$

Transformer Layer (by analogy)

Input: a sequence of vectors $X = [x_1, \dots, x_T]$

Output: a sequence of vectors $C = [c_1, \dots, c_T]$

Operation: mix the vectors nonlinearly using MHSA, $A = \text{MHSA}(X)$,

then pass result column-wise through nonlinearity $C = \text{FF}(A)$

where $\text{FF}(a) = W^2 (\sigma(W^1 a + b^1) + b^2)$

CNN Layer

Input: channels $[x_1, \dots, x_c]$

Output: channels $[o_1, \dots, o_c]$

Operation: linearly mix the channels

$$A_i = \sum_{j=1}^c k_j^i * x_j + b^i$$

then pass results pixelwise through nonlinearity

$$o_i = \sigma(A_i)$$

Encoder Transformer Block "Attention is All You Need", 2017

Just as we have deep RNNs, Encoder Transformer Blocks consist of stacked layers of MHSA

Encoder

$[c_1, \dots, c_T]$ contextual representations

Block N

⋮

Block 2

Block 1

$[x_1, \dots, x_T]$ global embeddings

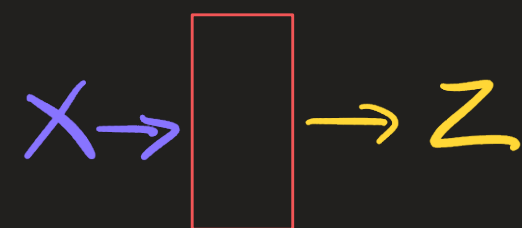
where each block has its own set of parameters and computes:
given input $X = [x_1, \dots, x_T] \in \mathbb{R}^{d_x \times T}$

$$i) \quad \varphi = \text{LayerNorm}(X + \text{MHSA}(X))$$

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

where $\text{FF}(\varphi) = [\text{FF}(\varphi_1) \mid \dots \mid \text{FF}(\varphi_T)]$

for $\text{FF}(y) = \omega^2 \sigma(\omega' y + b') + b^2$
and $\sigma = \text{ReLU}$

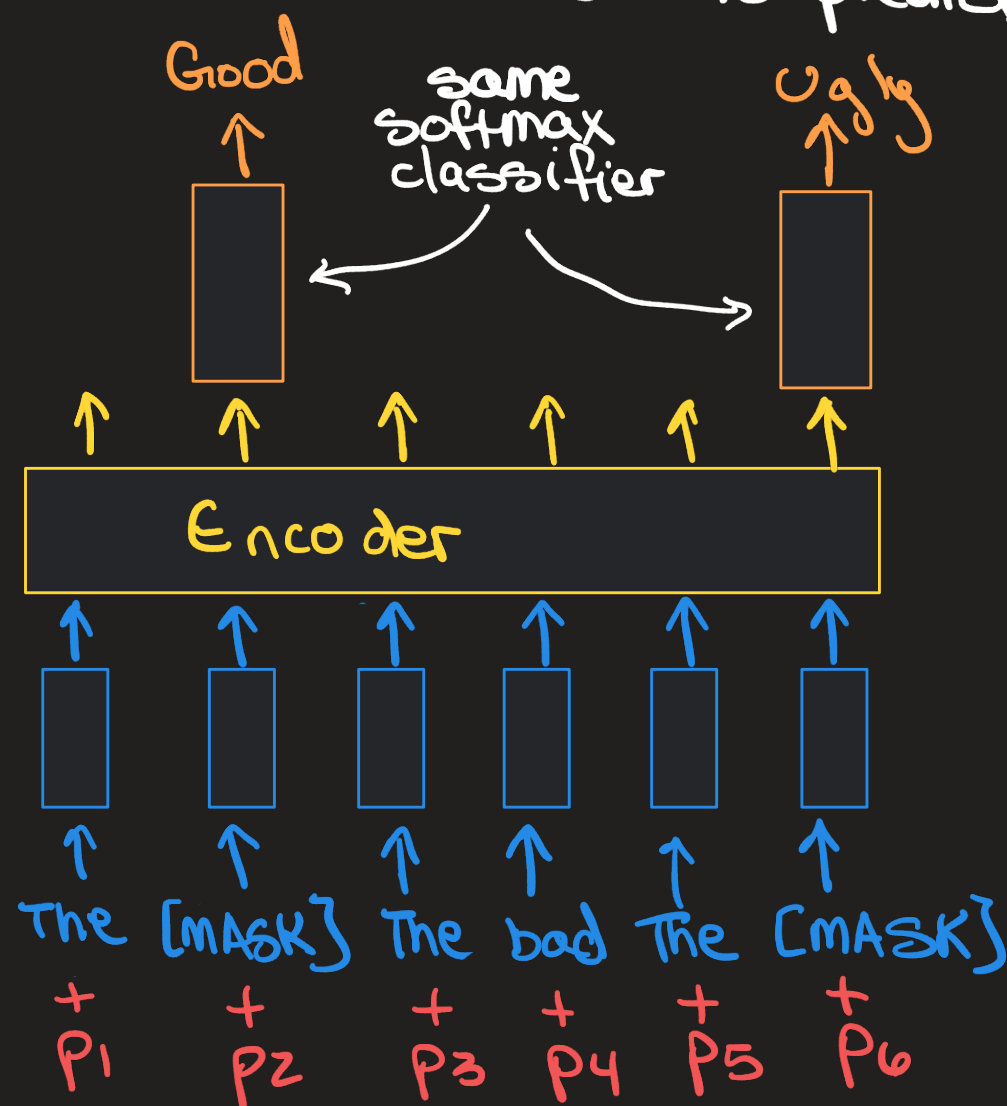


BERT "BERT: Pretraining of Deep Bidirectional Transformers for Language Understanding", 2018

- one of the first influential LLMs
- uses self-attention to learn contextual representations
- trains on large **unsupervised pretraining tasks**:
 - masked language modeling: predict missing words
 - next sentence prediction: predict whether sentence B really followed sentence A
- to use the model on supervised tasks, you **fine-tune** it by continuing to train for a small number of epochs at a lower learning rate on a supervised data set

Pretraining: Masked Language Modeling (MLM)

Given an input sequence $\{\text{The, good, the, bad, the, ugly}\}$, randomly select tokens to mask, with equal probability, and use the encoder to predict the masked tokens

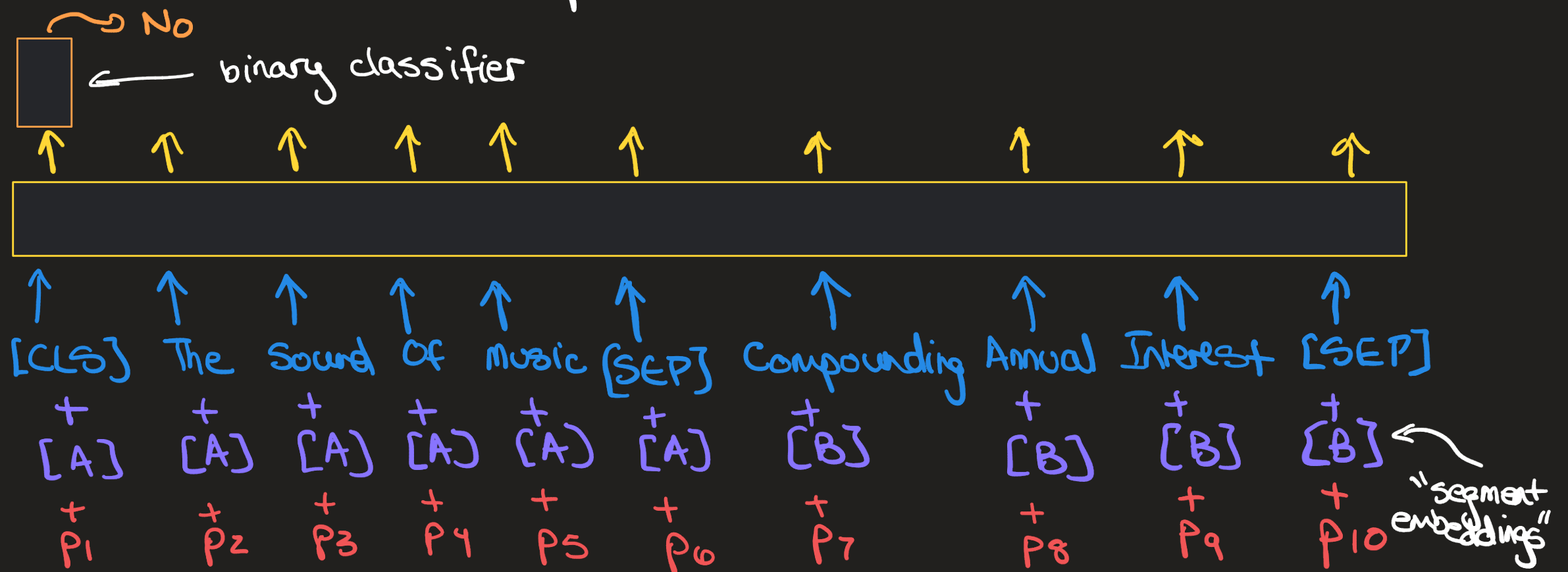


The classifier, encoder, and embedding parameters are learned to minimize error in predicting the masked tokens

The MLM objective reaches BERT the structure of language, the way words are used together.

Pretraining : Next Sentence Prediction (NSP)

Given two input sequences, {The, sound, of, music} and {Compounding, annual, interest}, BERT is trained to predict whether or not these sequences were consecutive.



This task also helps BERT learn how language is used

Finetuning:

Pretraining on a large corpus helps BERT to learn about the structure of language. To use it on supervised tasks, one "fine-tunes" the model by:

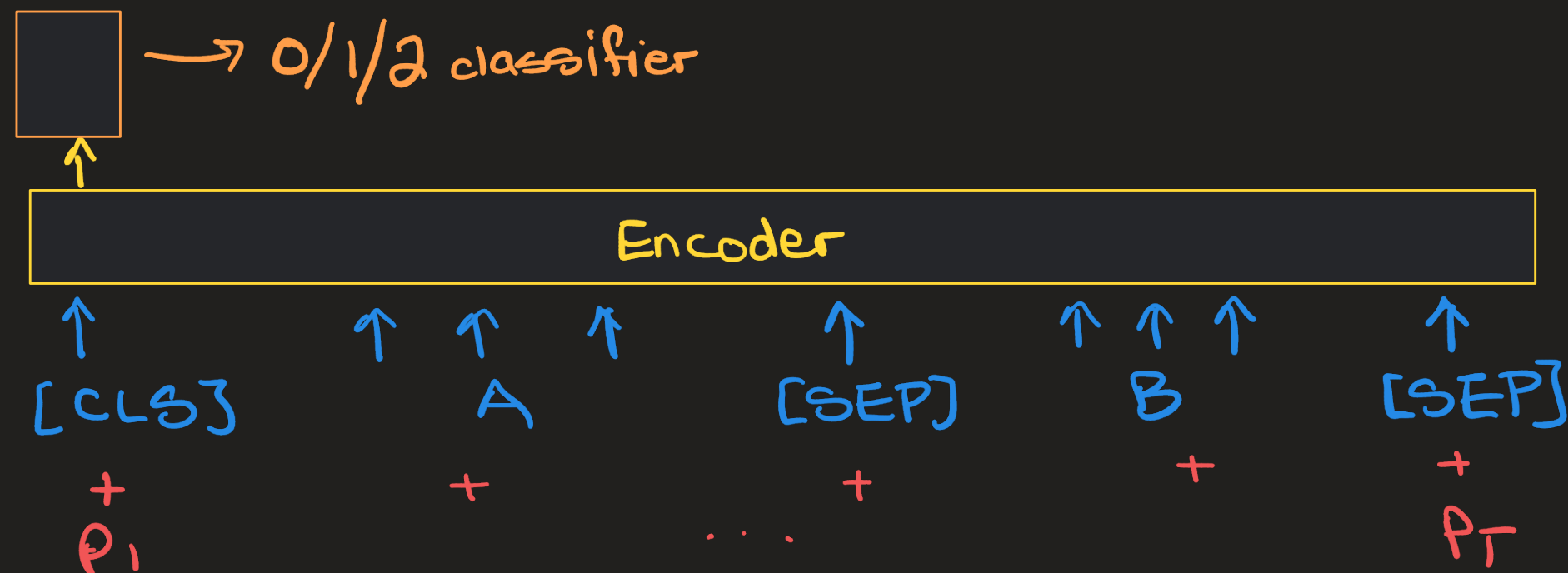
- (potentially) modifying the architecture
- changing the loss appropriately
- continuing training the model, usually with a lower learning rate and fewer epochs

Fine-tuning balances between retaining "general knowledge" that is useful for the task, and specializing for the task.

Pretraining + fine-tuning allows you to use large powerful architectures to perform well on tasks where you have comparatively smaller training sets, without overfitting.

Fine-tuning example: MNLI (Multi-genre Natural Language Inference)

- a standard benchmark task in the popular GLUE (General Language Understanding Evaluation) benchmark
- Given two sentences A, B , determine whether B is an **entailment**, **contradiction**, or **neutral** with respect to A
- Encode inputs as for the NSP task, but use a new classification head for the MNLI task



Fine-tuning example: SQUAD (Stanford question answering Dataset)

- Given a question ^A and a passage ^B from Wikipedia with the answer, identify the answer text span in the passage.
- Introduce **S** and **E** embeddings that we will use to identify the start and end of a candidate answer
- Training: let $P_{\text{start}} = \text{softmax} \left(\begin{bmatrix} c_1^T \\ \vdots \\ c_T^T \end{bmatrix} S \right)$ and

$$P_{\text{end}} = \text{softmax} \left(\begin{bmatrix} c_1^T \\ \vdots \\ c_T^T \end{bmatrix} E \right)$$

← only apply over the tokens in **B**

be the probabilities that the answer starts (ends) at each token. Maximize the sum of the log probabilities of the true start and true end of the answer.

- Inference: score each span (where $j > i$) $[c_i, \dots, c_j]$ with $\langle S, c_i \rangle + \langle E, c_j \rangle$ and return the span with the highest score

