

1

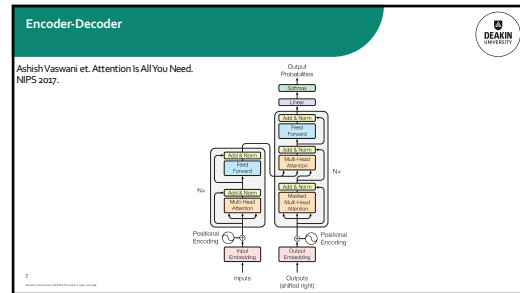
2

3

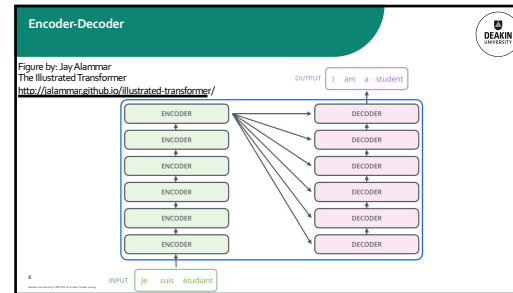
4

5

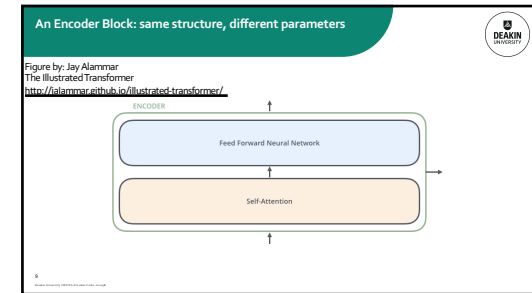
6



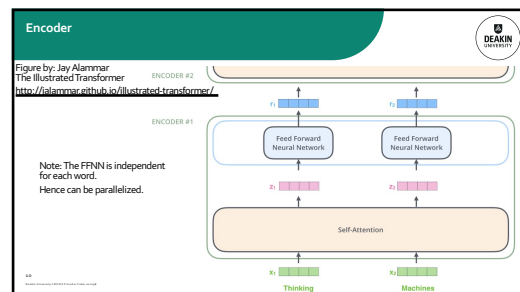
7



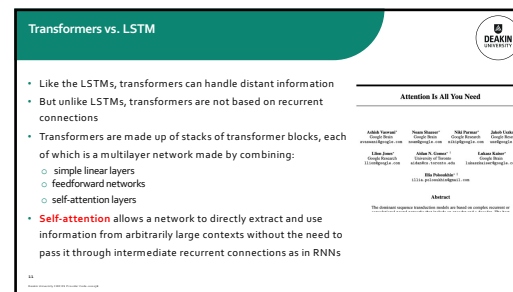
8



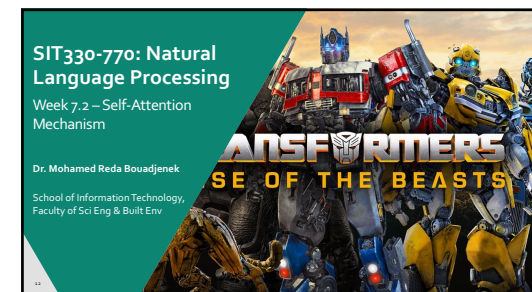
9



10



11



12

Self-Attention Mechanism

- All-to-all comparison.
 - Each layer is $O(N^2)$ for sequence of length N - **self attention**.
- Every output is a weighted sum of every input.
 - The weighting is a function to learn.

The diagram illustrates the Self-Attention Mechanism. On the left, a neural network layer is shown with 6 green input nodes and 6 orange output nodes, connected by a dense web of lines representing all-to-all comparisons. On the right, a detailed view of the attention mechanism is shown. It features two identical lists of words: 'The', 'cat', 'sat', 'on', 'the', 'mat'. The top list is labeled 'input' and the bottom list is labeled 'output'. Lines connect each input word to each output word, representing the all-to-all comparison. The lines are colored in a gradient from light blue to dark blue, indicating the attention weights. The output list is also labeled 'output'.

42

Source: <https://www.youtube.com/watch?v=K1uE0tTt0j0>

13

Relevance scores from each input to output

query[17] # "making"

key[24] # "difficult"

Relevance[17,24]=query[17] * key[24]
relevance of difficult to making

14

- Q: Query (output token)
- K: Key (input token)
- Relevance = $\text{softmax}(\frac{Q \times K^T}{\sqrt{d_k}})$
 - d_k is the dimension of Q or K
- V: Value (input token)
- Out = Relevance $\times$ V

Scaled Dot-Product Attention

15

```

def attention(self, X_in: List[Tensor]):
    # For every token transform previous layer's out
    for i in range(self.sequence_length):
        query[i] = self.Q * X_in[i]
        key[i] = self.K * X_in[i]
        value[i] = self.V * X_in[i]

```

16

```
def attention(self, X_inlist[Tensor]):
    # For every token transfer previous layer's output
    for i in range(self.sequence_length):
        query[i] = self.Q * X_inlist[i]
        key[i] = self.K * X_inlist[i]
        value[i] = self.V * X_inlist[i]

    # Compute output values, one at a time
    for i in range(self.sequence_length):
        this_query = query[i]
        # How relevance is each input to this output?
        for j in range(self.sequence_length):
            relevance[j] = this_query * key[j]
        # normalize relevance score to sum to 1
        relevance = scaled_softmax(relevance)
        # compute a weighted sum of values
        out[i] = 0
        out[i] is a vector
        for j in range(self.sequence_length):
            out[i] += relevance[j] * value[j]

    return out
```

Scaled Dot-Product Attention



```

graph TD
    Input[Input] --> Q[Q]
    Input --> K[K]
    Input --> V[V]
    Q --> ScaledDotProduct[Scaled Dot-Product]
    K --> ScaledDotProduct
    ScaledDotProduct --> Attention[Attention]
    Attention --> Output[Output]
    
```

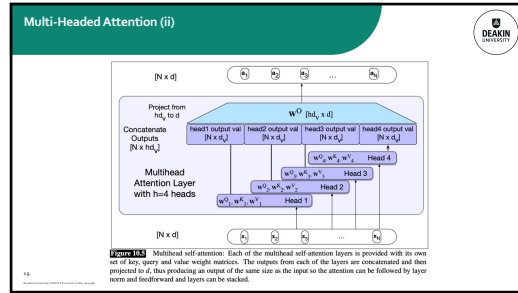
The diagram illustrates the Scaled Dot-Product Attention process. It starts with an 'Input' box at the top, which branches into three separate boxes labeled 'Q', 'K', and 'V'. Arrows from 'Q' and 'K' point to a box labeled 'Scaled Dot-Product'. An arrow from 'Scaled Dot-Product' points to a box labeled 'Attention'. Finally, an arrow from 'Attention' points to an 'Output' box at the bottom.

17

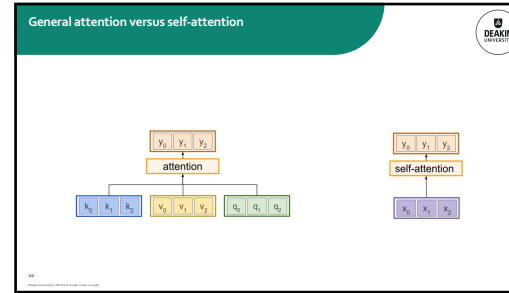
- Clever, important innovation.
 - Not that hard.
- Just so that same thing 8 times with different Q,K,V matrices.
- Let the network learn 8 different semantic meanings of attention.
 - E.g., One grammar, one for vocabulary, one for conjugation, etc.
 - Very flexible mechanism for sequence processing.

$$\begin{aligned} Q &= XW_Q^V; \quad K = XW_K^V; \quad V = XW_V^V & (10.17) \\ \text{head}_i &= \text{SelfAttention}(Q, K, V) & (10.18) \\ A &= \text{MultiHeadAttention}(X) = (\text{head}_1 \parallel \text{head}_2 \dots \parallel \text{head}_d)W^O & (10.19) \end{aligned}$$

18



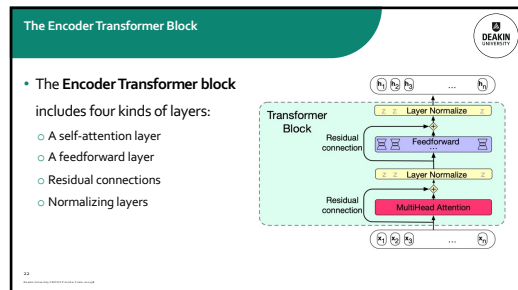
19



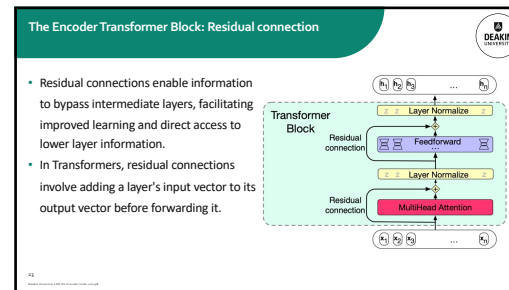
20



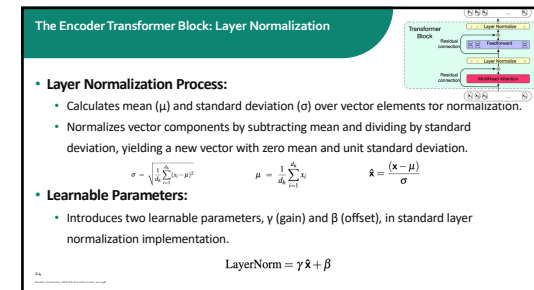
21



22



23



24

The Encoder Transformer Block: feedforward layer

- The feedforward layer consists of N position-wise networks, each positioned independently.
- Each network is a fully-connected 2-layer neural network, comprising one hidden layer and two weight matrices.
- While the weights remain consistent across positions, the parameters differ from layer to layer.
- Unlike attention mechanisms, the feedforward networks operate independently at each position, enabling parallel computation.

25

The Encoder Transformer Block: Putting it all together

- The Transformer block includes four kinds of layers:
 - A self-attention layer
 - Residual connections
 - Normalizing layers
 - A feedforward layer

$$\begin{aligned}
 T^1 &= \text{SelfAttention}(X) \\
 T^2 &= X + T^1 \\
 T^3 &= \text{LayerNorm}(T^2) \\
 T^4 &= \text{FFN}(T^3) \\
 T^5 &= T^4 + T^3 \\
 H &= \text{LayerNorm}(T^5)
 \end{aligned}$$

26

Encoder-Decoder Transformers

Ashish Vaswani et. Attention Is All You Need. NIPS 2017.

27

SIT330-770: Natural Language Processing

Week 7.4 – The Input: Embeddings for Tokens

Dr. Mohamed Reda Bouadjene

School of Information Technology, Faculty of Sci Eng & Built Env

28

The input X

- Source of Input X :**
 - Originates from a sequence of N tokens, where N represents the number of tokens in the context.
 - Matrix X has dimensions $[N \times d]$ where d is the dimension of each embedding.
- Composition of Embeddings:**
 - The transformer model computes two separate embeddings:
 - Input Token Embedding:** Specific to each token in the sequence.
 - Input Positional Embedding:** Reflects the position of each token within the sequence.

29

Input Token Embedding

- Each token's initial representation is a vector of dimension d .
- These embeddings evolve through transformer layers, incorporating contextual nuances.
- Embedding Storage:**
 - Stored in matrix E with shape $[|V| \times d]$, where $|V|$ is the vocabulary size.
- Token Conversion and Indexing:**
 - Tokens first converted to vocabulary indices using techniques like BPE or SentencePiece (discussed in Week 3).
 - Example: "thanks for all" converts to indices $[5, 4000, 10532, 2224]$.
 - These indices are used to fetch the corresponding embeddings from E .

30

One-Hot Vector

- One-Hot Vector Representation:
 - Alternative method using one-hot vectors of size $|V|$.
 - Example: Word "thanks" represented as a vector where only the 5th position is 1, all others are 0.

$$\begin{bmatrix} 0 & 0 & 0 & 1 & 0 & 0 & \dots & 0 & 0 & 0 & 0 \end{bmatrix}$$

$$\begin{matrix} 1 & 2 & 3 & 4 & 5 & 6 & 7 & \dots & \dots & |V| \end{matrix}$$
- Multiplying by a one-hot vector that has only one non-zero element $x_i = 1$ simply selects out the relevant row vector for word i , resulting in the embedding for word i .

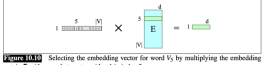


Figure 10.10 Selecting the embedding vector for word i by multiplying the embedding matrix E with a one-hot vector with a 1 in index i .

31

One-Hot Vector

- We can extend this idea to represent the entire token sequence as a matrix of one-hot vectors, one for each of the N positions in the transformer's context window,

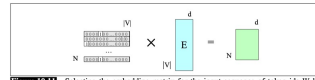


Figure 10.11 Selecting the embedding matrix for the input sequence of token ids W by multiplying a one-hot matrix corresponding to W by the embedding matrix E .

32

SIT330-770: Natural Language Processing

Week 7.5 – The Input: Embeddings for Positions

Dr. Mohamed Reda Bouadjenek

School of Information Technology,
Faculty of Sci Eng & Built Env



33

Input Positional Embedding

- How does a transformer model the position of each token in the input sequence?
 - With RNNs, information about the order of the inputs was built into the structure of the model, Not with Transformers
- Solution: **Positional**
 - Embeddings Modify the input embeddings by combining them with positional embeddings specific to each position in an input sequence

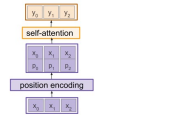
34

Positional Embedding

- Positional encoding assigns a unique representation to each position within a sequence to describe the location or order of entities.
- Limitations of Using Single Numbers for Position:
 - Using index values alone (e.g., sequence position numbers) is problematic for several reasons:
 - Large indices for long sequences can grow unmanageably large.
 - Normalizing indices between 0 and 1 creates inconsistencies across sequences of different lengths.
- Transformers' Approach to Positional Encoding:
 - Instead of single numbers, transformers map each position to a unique vector.
 - This results in a matrix where each row represents an object in the sequence combined with its positional information.

35

Positional encoding



- Concatenate/add special positional encoding p_i to each input vector x_i .
- We use a function $\text{pos}: N \rightarrow \mathbb{R}^d$ to process the position j of the vector into a d -dimensional vector
- So, $p_j = \text{pos}(j)$

- Desiderata of $\text{pos}(\cdot)$:
 - It should output a unique encoding for each time-step (word's position in a sentence)
 - Distance between any two time-steps should be consistent across sentences with different lengths.
 - Our model should generalize to longer sentences without any efforts. Its values should be bounded.
 - It must be deterministic.

36

Positional Encoding Layer in Transformers

- The positional encoding is given by sine and cosine functions of varying frequencies:
 - k : Position of an object in the input sequence,
 - d : Dimension of the output embedding space
 - n : User-defined scalar, set to 10,000 by the authors of Attention Is All You Need.

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

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

Example:

- "I am a robot," with $n=100$ and $d=4$

Sequence	Index of tokens	Positional Encoding Matrix with $d=4$, $n=100$			
	$i=0$	$i=1$	$i=2$	$i=3$	
I	0	$P_{pos[0][0]}$	$P_{pos[0][1]}$	$P_{pos[0][2]}$	$P_{pos[0][3]}$
am	1	$P_{pos[1][0]}$	$P_{pos[1][1]}$	$P_{pos[1][2]}$	$P_{pos[1][3]}$
a	2	$P_{pos[2][0]}$	$P_{pos[2][1]}$	$P_{pos[2][2]}$	$P_{pos[2][3]}$
Robot	3	$P_{pos[3][0]}$	$P_{pos[3][1]}$	$P_{pos[3][2]}$	$P_{pos[3][3]}$

Positional Encoding Matrix for the sequence "I am a robot"

37

Coding the Positional Encoding Matrix from Scratch

```

import numpy as np
import matplotlib.pyplot as plt

def getPositioEncoding (seq_len, d, n=10000):
    P = np.zeros((seq_len, d))
    for k in range(seq_len):
        for i in np.arange(int(d/2)):
            denominator = np.power(n, 2*(i+1)/d)
            P[k, 2*i] = np.sin(k/denominator)
            P[k, 2*i+1] = np.cos(k/denominator)
    return P

P = getPositioEncoding(seq_len=4, d=4, n=100)
print(P)

```

```

1 [[ 0.  1.  0.  1.]
2 [ 0.8147098  0.54030231  0.0983342  0.99500417]
3 [ 0.90929743 -0.41614684  0.19866933  0.98006658]
4 [ 0.14112001 -0.9899925  0.29552021  0.95533649]]

```

38

Positional Embeddings

Figure 10.12 A simple way to model position: add an embedding of the absolute position to the token embedding to produce a new embedding of the same dimensionality.

39

SIT330-770: Natural Language Processing

Week 7.6 – The Task Specific Head

Dr. Mohamed Reda Bouadjene
School of Information Technology,
Faculty of Sci Eng & Built Env

40

The Language Modeling Head

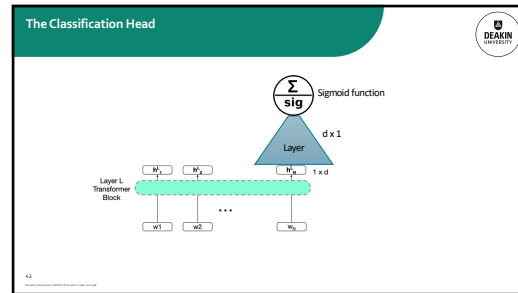
- The Language Modeling Head is an additional neural circuitry integrated with the basic transformer architecture.
- It enables specific tasks such as language modeling by enhancing the transformer's capabilities.
- Given a context of words, a Language Model assigns a probability to each possible next word
 - Calculating the probability of the next word "fish" given the context "Thanks for all the": $P(\text{fish}|\text{Thanks for all the})$

41

The Language Modeling Head

Figure 10.13 The language modeling head: the circuit at the top of a transformer that maps from the output embedding for token x_i from the last transformer layer (h_i) to a probability distribution over words in the vocabulary V .

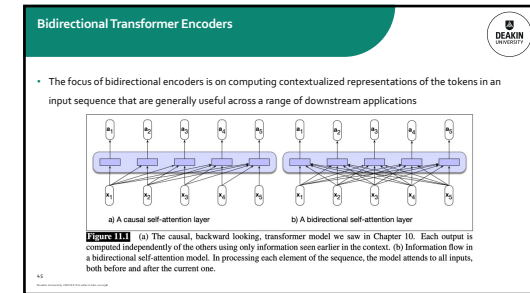
42



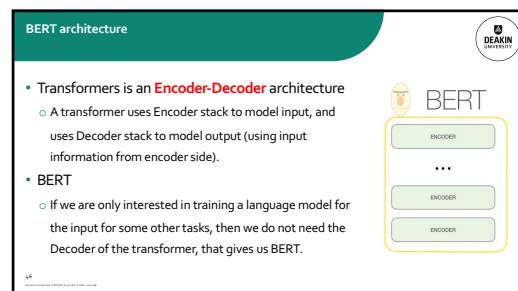
43



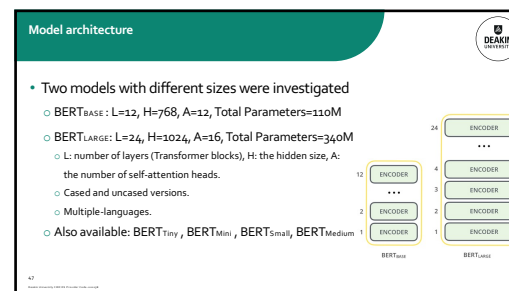
44



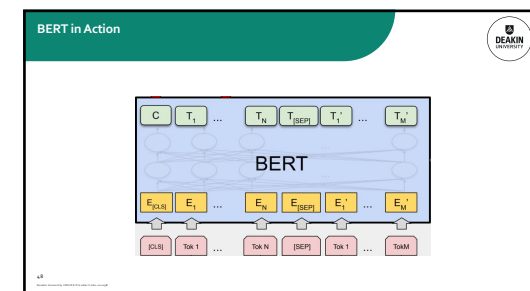
45



46



47



48

Input Representation

- Token Embeddings:** Use pretrained WordPiece embeddings.
- Segment Embeddings (Optional):** Added sentence embedding to every tokens of each sentence.
- Position Embeddings:** Use learned Position Embeddings.
- Use [CLS]** for the classification tasks.
- Separate sentences** by using a special token [SEP].

49

BERT's Vocabulary

- BERT is pre-trained → Vocabulary is fixed
 - The vocabulary contains 30,522 tokens.
 - How to deal with unknown words?
- Break down **unknown words** into **subwords**:
 - Using the **wordpiece** model.
- A subword exists for every character.
 - 2 types of subwords
 - All subwords start with "##"...
 - Except for the first subword in a word

50

Differences in pre-training model architectures: BERT, OpenAI GPT, and ELMo

- BERT is built using transformer encoder blocks. BERT, on the other hand, uses transformer encoder blocks.
- GPT is auto-regressive in nature. BERT is not.
 - In losing auto-regression, BERT gained the ability to incorporate the context on both sides of a word to gain better results.

51

SIT330-770: Natural Language Processing

Week 7.8 – BERT pre-training

Dr. Mohamed Reda Bouadjene

School of Information Technology,
Faculty of Sci Eng & Built Env

52

Pre-trained Models

- Pre-trained models are ML models that have been previously trained on a large dataset, typically on a general task
- These models can significantly reduce the computational cost and time required to develop ML applications by leveraging learned features and weights.
 - Speeds Up Development**
 - By using models that are already trained, developers can focus on fine-tuning rather than starting from scratch.
 - Improves Performance**
 - Pre-trained models often bring high levels of accuracy and efficiency, especially on complex tasks that require learning from large amounts of data.
 - Resource Efficiency**
 - Saves resources by reducing the need for extensive training data and computational power.

53

Pre-training and Fine-tuning

- The same architectures are used in both pre-training and fine-tuning
- [CLS] is a special symbol added in front of every input example, and [SEP] is a special separator token
- Pre-training:** two tasks are considered
 - Masked LM:** mask some percentage of the input tokens at random, and then predict those masked tokens. Mask 30% of all WordPiece tokens in each sequence at random
 - Next Sentence Prediction:** understanding the relationship between two sentences (90% of positive pairs). Used BooksCorpus (800M words) and Wikipedia (2,500M words)

54

Pre-training Task#1: Masked Language Model (MLM)

- The original approach to training bidirectional encoders is called **Masked Language Modeling (MLM)**
 - MLM uses unannotated text from a large corpus
 - Model is presented with a series of sentences from the training corpus, where a random sample of tokens from each training sequence is selected for use in the learning task. Once chosen, a token is used in one of three ways:
 - It is replaced with the unique vocabulary token [MASK]
 - It is replaced with another token from the vocabulary, randomly sampled based on token unigram probabilities
 - It is left unchanged
- In BERT, 15% of the input tokens in a training sequence are sampled for learning
 - Of these, 80% are replaced with [MASK], 10% are replaced with randomly selected tokens, and the remaining 10% are left unchanged

55

Pre-training Task#1: Masked Language Model (MLM)

- The MLM training objective is to predict the original inputs for each of the masked tokens using a bidirectional encoder of the kind described in the last section
 - The **cross-entropy** loss from these predictions drives the training process for all the parameters in the model
 - Note that all the input tokens play a role in the self-attention process, but **only the sampled tokens** are used for learning
- Input:**
 - Original input sequence is first tokenized using a subword model
 - The sampled items which drive the learning process are chosen from among the set of tokenized inputs
 - Word embeddings for all the tokens in the input are retrieved from the word embedding matrix and then combined with positional embeddings to form the input to the transformer

56

Pre-training Task#1: Masked Language Model (MLM)

Figure 11.4 Masked language model training. In this example, three of the input tokens are selected, two of which are masked and the third is replaced with an unrelated word. The probabilities assigned by the model to these three items are used as the training loss. The other 5 words don't play a role in training loss. (In this and subsequent figures we display the input as words rather than subword tokens; the reader should keep in mind that BERT and similar models actually use subword tokens instead.)

57

Pre-training Task#2: Next Sentence Prediction (NSP)

- An important class of applications involves determining the relationship between pairs of sentences
 - paraphrase detection (detecting if two sentences have similar meanings)
 - entailment (detecting if the meanings of two sentences entail or contradict each other)
 - discourse coherence (deciding if two neighboring sentences form a coherent discourse)
- To capture the kind of knowledge required for applications such as these, BERT introduced a second learning objective called Next Sentence Prediction (NSP)
- Training:** The model is presented with pairs of sentences and is asked to predict whether each pair consists of an actual pair of adjacent sentences from the training corpus or a pair of unrelated sentences

58

Pre-training Task#2: Next Sentence Prediction (NSP)

- In BERT, 50% of the training pairs consisted of positive pairs, and in the other 50% the second sentence of a pair was randomly selected from elsewhere in the corpus
 - The NSP loss is based on how well the model can distinguish true pairs from random pairs
- BERT introduces two new tokens to the input representation
 - After tokenizing the input with the subword model, the token [CLS] is prepended to the input sentence pair, and the token [SEP] is placed between the sentences and after the final token of the second sentence
 - During training, the output vector from the final layer associated with the [CLS] token represents the next sentence prediction

59

Pre-training Task#2: Next Sentence Prediction (NSP)

Figure 11.5 An example of the NSP loss calculation.

60

Pre-training procedure

- Training data: BooksCorpus (800M words) + English Wikipedia (2.5B words).
- To generate each training input sequences: sample two spans of text (A and B) from the corpus.
 - The combined length is ≤ 500 tokens.
 - 50% B is the actual next sentence that follows A and 50% of the time it is a random sentence from the corpus.
- The training loss is the sum of the mean masked LM likelihood and the mean next sentence prediction likelihood.

61

Under the Hood

BERT Size & Architecture

ML Architecture Parts	Definition				
Parameters	Number of learnable variables/values available for the model				
Transformer Layers	Number of Transformer blocks. A transformer block transforms a sequence of word representations to a sequence of contextualized words (numerical representations)				
Hidden Size	Layers of mathematical functions, located between the input and output, that assign weights (to words) to produce a desired result				
Attention Heads	The size of a Transformer block				
Processing	Type of processing unit used to train the model				
Length of Training	Time it took to train the model				

BERT Size & Architecture

	Transformer Layers	Hidden Size	Attention Heads	Parameters	Processing	Length of Training
BERTBase	12	768	12	110M	4 TPUs	4d
BERTLarge	24	1024	16	340M	16 TPUs	4d

62

SIT330-770: Natural Language Processing
Week 7.9 – BERT fine-tuning

Dr. Mohamed Reda Bouadjenek

School of Information Technology,
Faculty of Sci Eng & Built Env

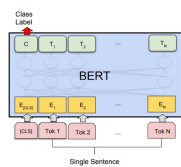


63

Fine-tuning procedure 1: Classification

- For sequence-level classification task
 - Obtain the representation of the input sequence by using the final hidden state (hidden state at the position of the special token [CLS]) $C \in R^H$
 - Just add a classification layer and use **softmax** to calculate label probabilities. Parameters $W \in R^{K \times H}$

$P = \text{softmax}(CW^T)$

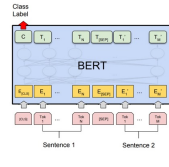


64

Fine-tuning procedure 2: Sentence Pair Classification

- For sentence pair classification task
 - Sentences are separated with a special token [SEP]
 - Obtain the representation of the input sequence by using the final hidden state (hidden state at the position of the special token [CLS]) $C \in R^H$
 - Just add a classification layer and use **softmax** to calculate label probabilities. Parameters $W \in R^{K \times H}$

$P = \text{softmax}(CW^T)$

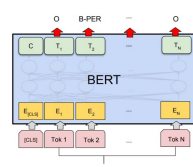


65

Fine-tuning procedure 3: Named Entity Recognition

- Feed the final hidden representation $T_i \in R^H$ for each token i into a classification layer for the tagset.
- To make the task compatible with WordPiece tokenization
 - Predict the tag for the first sub-token of a word
 - No prediction is made for X

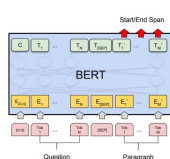
Jim Hen ##son was a puppet #er
I-PER I-PER X O O O X



66

Fine-tuning procedure 4: Query Answering 1/2

- Input Question:**
Where do water droplets collide with ice crystals to form precipitation?
- Input Paragraph:**
.... Precipitation forms as smaller droplets coalesce via collision with other rain drops or ice crystals within a cloud. ...
- Output Answer:**
within a cloud



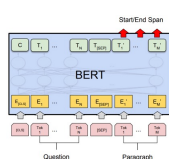
67

Fine-tuning procedure 4: Query Answering 2/2

- Represent the input question and paragraph as a single packed sequence.
 - The question uses the A embedding and the paragraph uses the B embedding.
- New parameters to be learned in fine-tuning are start vector $S \in \mathbb{R}^H$ and end vector $E \in \mathbb{R}^H$.
- Calculate the probability of word w being the start of the answer span:

$$P_{start} = \text{Softmax}(S^T)$$
- The training objective is the log-likelihood of the correct end positions.

$$P_{end} = \text{Softmax}(E^T)$$



68

SIT330-770: Natural Language Processing
Week 7.10 – BERT Performance

Dr. Mohamed Reda Bouadjenek

School of Information Technology,
Faculty of Sci Eng & Built Env



69

Experiments

- GLUE (General Language Understanding Evaluation) benchmark
 - Distribute canonical Train, Dev and Test splits
 - Labels for Test set are not provided
- Datasets in GLUE:
 - MNLI: Multi-Genre Natural Language Inference
 - QQP: Quora Question Pairs
 - QNLI: Question Natural Language Inference
 - SST-2: Stanford Sentiment Treebank
 - CoLA: The corpus of Linguistic Acceptability
 - STS-B: The Semantic Textual Similarity Benchmark
 - MRPC: Microsoft Research Paraphrase Corpus
 - RTE: Recognizing Textual Entailment
 - WNLI: Winograd NLI

70

GLUE Results

System	MNLI (m/mm)	QQP	QNLI	SST-2	CoLA	STS-B	MRPC	RTE	Average
Pre-OpenAI SOTA	392k	363k	108k	67k	8.5k	5.7k	3.5k	2.5k	-
BILSTM+ELMo+Attn	80.6/80.1	66.1	82.3	93.2	35.0	81.0	86.0	61.7	74.0
OpenAI GPT	76.4/76.1	64.8	79.8	90.4	36.0	73.3	84.9	56.8	71.0
BERT _{Base}	82.1/81.4	70.3	87.4	91.3	45.4	80.0	82.3	56.0	75.1
BERT _{Large}	84.6/83.4	71.2	90.5	93.5	52.1	85.8	88.9	66.4	79.6
BERT _{Large}	86.7/85.9	72.1	92.7	94.9	60.5	86.5	89.3	70.1	82.1

Table 1: GLUE Test results, scored by the evaluation server (<https://gluebenchmark.com/leaderboard>). The number below each task denotes the number of training examples. The "Average" column is slightly different than the official GLUE score, since we exclude the problematic WNLI set. BERT and OpenAI GPT are single-model, single task. F1 scores are reported for QQP and MRPC. Spearman correlations are reported for STS-B, and accuracy scores are reported for the other tasks. We exclude entries that use BERT as one of their components.

71

SQuAD: The Stanford Question Answering Dataset

System	Dev EM	Dev F1	Test EM	Test F1
Top Leaderboard Systems (Dec 10th, 2018)				
Human	-	-	82.3	91.2
#1 Ensemble - elnet	-	-	86.0	91.7
#2 Ensemble - QANet	-	-	84.5	90.5
Published				
BIDAF+HiLato (Single)	-	85.0	-	85.8
R.M. Reader (Ensemble)	81.2	87.9	82.3	88.5
Ours				
BERT _{Large} (Single)	80.8	88.5	-	-
BERT _{Large} (Single)	81.1	90.9	-	-
BERT _{Large} (Ensemble)	85.8	91.8	-	-
BERT _{Large} (Sgl+TriuQA)	84.2	91.1	85.1	91.8
BERT _{Large} (Enb+TriuQA)	86.2	92.2	87.4	93.2

Table 2: SQuAD 1.1 results. The BERT ensemble is 7x systems which use different pre-training checkpoints and fine-tuning seeds.

System	Dev EM	Dev F1	Test EM	Test F1
Top Leaderboard Systems (Dec 10th, 2018)				
Human	-	-	86.3	89.0
#1 Single - MIR-MRC (F-Net)	-	-	74.8	79.0
#2 Single - nlnet	-	-	74.2	77.1
Published				
unet (Ensemble)	-	-	71.4	74.9
SLQA+ (Single)	-	-	71.4	74.4
Ours				
BERT _{Large} (Single)	78.7	81.9	80.0	83.1

Table 3: SQuAD 2.0 results. We exclude entries that use BERT as one of their components.

72

SWAG

- The Situations with Adversarial Generations (SWAG)

On stage, a woman takes a seat at the piano.
She ...

- sits on a bench as her sister plays with the doll.
- smiles with someone as the music plays.
- is in the crowd, watching the dancers.
- nervously sets her fingers on the keys.

Table 4: SWAG Dev and Test accuracies. Human performance is measured with 100 samples, as reported in the SWAG paper.

System	Dev	Test
ESIM+GloVe	51.9	52.7
ESIM+attn-Mo	59.1	59.2
OpenAI GPT	-	78.0
BERT _{base}	81.0	-
BERT _{Large}	86.6	86.3
Human (expert) ¹	-	85.0
Human (5 annotations) ²	-	88.0

- The only task-specific parameters is a vector $V \in \mathbb{R}^H$
- The probability distribution is the **softmax** over the four choices

73

Conclusions

- Unsupervised pre-training (pre-training language model) is increasingly adopted in many NLP tasks.
 - Google Search is applying BERT models for search queries for over 70 languages.
- Major contribution of the paper is to propose a deep bidirectional architecture from Transformer.
 - Advance state-of-the-art for many important NLP tasks.
- Cannot do everything in NLP!**

74

SIT330-770: Natural Language Processing

Week 7.11 – Other Models Based on Transformers

Dr. Mohamed Reda Bouadjeneke

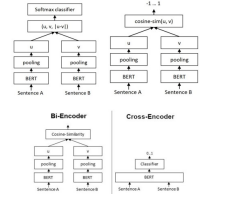
School of Information Technology,
Faculty of Sci Eng & Built Env



75

Sentence-BERT: Sentence Embeddings using Siamese BERT-Networks

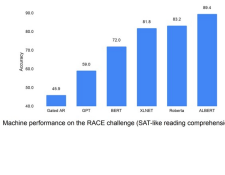
- A modification of the pretrained BERT network that use Siamese network structures to derive semantically meaningful sentence embeddings that can be compared using cosine-similarity
 - <https://www.sbert.net/>



76

BERT-based Models

- RoBERTa: Robustly Optimized BERT Approach
 - Facebook AI research team
 - They used 160 GB of text instead of the 16 GB dataset originally used to train BERT
 - Increased the number of iterations from 100k to 300k and then further to 500k
 - Dynamically changing the masking pattern applied to the training data
 - Removing the next sequence prediction objective from the training procedure
- ALBERT, XLNET, CoBERT



77

Generative Pre-trained Transformer (GPT)

- GPT (Generative Pre-trained Transformer) is a series of language generation models developed by OpenAI. These models are based on the Transformer architecture (2018)
- GPT-2 (Generative Pre-trained Transformer 2) was the second model in the GPT series, released in 2019. It was trained on a large corpus of internet text and was designed for language generation tasks such as question answering, and text summarization
- GPT-3 (Generative Pre-trained Transformer 3) Released in 2020, with over 175 billion parameters, and was trained on a much larger and diverse dataset, including web pages, books, and scientific articles
- GPT-4 (Generative Pre-trained Transformer 4) Released in 2023, a multimodal model which can accept image and text inputs and produce text outputs

78

Zero-shot, One-shot and Few-shot, Contrasted with Traditional Fine-tuning

Traditional fine-tuning (not used for GPT-3)

Fine-tuning
The model is trained via repeated gradient updates using a large corpus of example tasks.

Language Models are Few-Shot Learners

Zero-shot
The model predicts the answer given only a natural language description of the task. No gradient updates are performed.

One-shot
In addition to the task description, the model uses a single example of the task. No gradient updates are performed.

Few-shot
In addition to the task description, the model uses a few examples of the task. No gradient updates are performed.

79

Text-to-Text Transfer Transformer (T5)

T5

Translate English to German: That is good. → That is good.

Write a sentence: The chicken is jumping well. → The chicken is jumping well.

Summarize: state authorities dispatched more police to the streets to survey the damage after an onslaught of severe weather. → state authorities dispatched more police to the streets to survey the damage after an onslaught of severe weather.

Text is input → Text is output

- A diagram of the text-to-text framework (T5)
- Every task, including translation, question answering, and classification, is cast as feeding T5 model text as input and training it to generate some target text

80

SIT330-770: Natural Language Processing

Week 7.12 – HuggingFace

Dr. Mohamed Reda Bouadjene

School of Information Technology,
Faculty of Sci Eng & Built Env

81

HuggingFace

More than 5,000 organizations are using HuggingFace

AI2, Meta AI, Google AI, Intel, SpeechBrain, Microsoft, Grammarly

Problems solvers

HuggingFace has a course on NLP: <https://huggingface.co/course/chapter0/1?fw=pt-35>
<https://huggingface.co/>

82

Documentations and Tutorials

- There is a wide range of examples provided, here are some useful links:
 - Notebooks:
 - <https://huggingface.co/transformers/v3.0.2/notebooks.html>
 - <https://github.com/huggingface/transformers/tree/main/notebooks>
 - Models:
 - <https://huggingface.co/models>
 - Course:
 - <https://huggingface.co/course/chapter0/1?fw=pt>

83

Getting Started

- Use pip for the installation, which is the package manager for Python. In notebooks, you can run system commands by preceding them with the ! character, so you can install the Transformers library as follows:


```
!pip install transformers
```
- You can make sure the package was correctly installed by importing it within your Python runtime:


```
import transformers
```

84

Pipelines

- The pipelines are a great and easy way to use models for inference
- These pipelines are objects that abstract most of the complex code from the library, offering a simple API dedicated to several tasks, including Named Entity Recognition, Masked Language Modeling, Sentiment Analysis, Feature Extraction and Question Answering

85

Example (1): Sentiment Analysis

```
from transformers import pipeline
classifier = pipeline("sentiment-analysis")
classifier("I've been waiting for a HuggingFace course my whole life.")
```

No model was supplied, defaulted to `distilbert-base-uncased-finetuned-sst-2-english` and revision `v4.0.0`. Using a pipeline without specifying a model name and revision is not recommended.

Downloading 100%  625K/625K [00:00<00:00, 28.5MB/s]

Downloading 100%  268M/268M [00:03<00:00, 83.2MB/s]

Downloading 100%  48.0/48.0 [00:00<00:00, 1.91MB/s]

Downloading 100%  232K/232K [00:00<00:00, 907KB/s]

```
[{"label": "POSITIVE", "score": 0.9988495214462285}]
```

```
classifier(
    ["I've been waiting for a HuggingFace course my whole life.",
     "I hate this so much!"]
)
```

```
[{"label": "POSITIVE", "score": 0.9988495214462285},
 {"label": "NEGATIVE", "score": 0.9994558693378453}]
```

86

Example (2): Question Answering

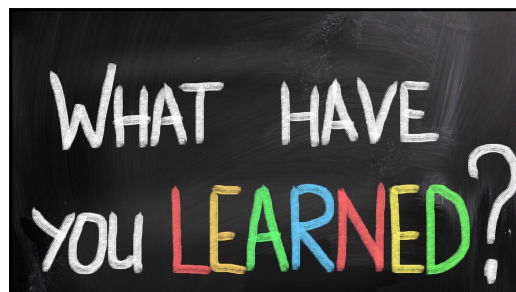
```
from transformers import pipeline
question_answerer = pipeline("question-answering")
question_answerer(
    question="Where do I work?",
    context="My name is Sylvain and I work at Hugging Face in Brooklyn",
)
```

```
{'score': 0.6949767470359802, 'start': 33, 'end': 45, 'answer': 'Hugging Face'}
```

```
from transformers import pipeline
oracle = pipeline(model="deepset/roberta-base-squad2")
oracle(question="where do I live?",
       context="My name is Wolfgang and I live in Berlin")
```

```
{'score': 0.9190717935562134, 'start': 34, 'end': 40, 'answer': 'Berlin'}
```

87



88

Summary

- Today we learned about:
 - Transformers
 - Attention is All you need
 - BERT

89