

How do LLMs work?

Two online visualizations

Transformer explained: <https://poloclub.github.io/transformer-explainer/>

Model architecture: <https://bbycroft.net/llm>

We will use these while we go through the code!

BYOLLM: Build Your Own Large Language Model

1. **Data**
2. **Tokenization**
3. **Embedding (vectorization, positional encoding)**
4. **Attention**
5. **Training**
6. **Validation and test**
7. **Next-word-prediction**

[LINK TO NOTEBOOK](#)

BYOLLM: Data

First step is to load a dataset to learn from.

A few examples have been preloaded, you can add more if you like.

We load all data to one long string so we can conveniently cut it up into smaller pieces later.

BYOLLM: Data

Part 1: The dataset

The first step to training your own LLMs is finding a dataset that you want to train on.

```
1 # Just run this
2
3 from datasets import load_dataset
4
5 # Function to load the dataset based on user choice
6 def load_text_dataset(dataset_name):
7     if dataset_name == 'imdb':
8         dataset = load_dataset('imdb')
9         text_key = 'text'
10    elif dataset_name == 'shakespeare':
11        dataset = load_dataset('sarnab/Shakespeare_Corpus')
12        text_key = 'text'
13    elif dataset_name == 'news':
14        dataset = load_dataset('SetFit/20_newsgroups')
15        text_key = 'text'
16    elif dataset_name == 'wiki':
17        dataset = load_dataset('wikitext', 'wikitext-2-raw-v1')
18        text_key = 'text'
19    elif dataset_name == 'python':
20        dataset = load_dataset('mbpp')
21        text_key = 'code'
22    else:
23        raise ValueError("Unknown dataset name. Please choose from 'imdb', 'news', or 'wikitext' or 'pythoncode'.")
24
25    return dataset, text_key
```

BYOLLM: Tokenization

Next, we tokenize the dataset.

In this toy example, we assign a token to each letter.

Most LLMs use ‘sub-word’ tokenization.

The king loves the queen



[976, 13793, 19620, 290, 42300]

BYOLLM: Tokenization

```
1 # Just run this
2
3 # To find which tokens we will need, we can list all the unique characters in our text
4 chars = sorted(list(set(text)))
5 vocab_size = len(chars)
6 print(f"All unique characters: {''.join(chars)}")
7 print(f"Number of unique characters: \n{vocab_size}")
```

⇒ All unique characters:
!"#\$%&'()*+,-./0123456789:;<=>?@ABCDEFGHIJKLMNOPQRSTUVWXYZ[\]^_abcdefghijklmnopqrstuvwxyz{|}~
Number of unique characters:
96

BYOLLM: Tokenization

We will define each unique character to be a token. All these tokens together form the vocabulary. And by mapping each token to a token ID (an integer), we can then represent our text as a list of integers.

```
1 # Just run this
2
3 # We define our vocabulary as a mapping of all unique characters to token IDs.
4
5 # Token to token ID
6 stoi = { ch:i for i,ch in enumerate(chars) }
7 # Token ID to token
8 itos = { i:ch for i,ch in enumerate(chars) }
9
10 # Function to tokenize text into token IDs
11 def tokenize(s):
12     return [stoi[c] for c in s]
13
14 # Function to detokenize a list of token IDs back into a string (text)
15 def detokenize(l):
16     return ''.join([itos[i] for i in l])
17
18 # Let's print the mapping so we can see which character corresponds to which token ID.
19 print(f"Mapping of characters to token IDs: {stoi}")
20 print(f"Mapping of token IDs to characters: {itos}")
```

⇒ Mapping of characters to token IDs: {'\t': 0, '\n': 1, '\r': 2, ' ': 3, '!': 4, '"': 5, '#': 6, '\$': 7, '%': 8, '&': 9, "

Mapping of token IDs to characters: {0: '\t', 1: '\n', 2: '\r', 3: ' ', 4: '!', 5: '"', 6: '#', 7: '\$', 8: '%', 9: '&', 1

BYOLLM: Embedding (vectorization)

Part 3: Vector embedding

We start with the embedding layer. Although we have already transformed our text into integers using tokenization, the numbers attributed to each token have no meaning. On top of that, if our dictionary would contain many words, our integer values would become very high.

The embedding layer gives transforms the tokenized sequences to a multi-dimensional vector that is more manageable for a deep learning network. It will also give similar (but not the same) embedding vectors to tokens that have a similar meaning.

At first, a tokenized sequence will be just a 1D array:

[sequence length]

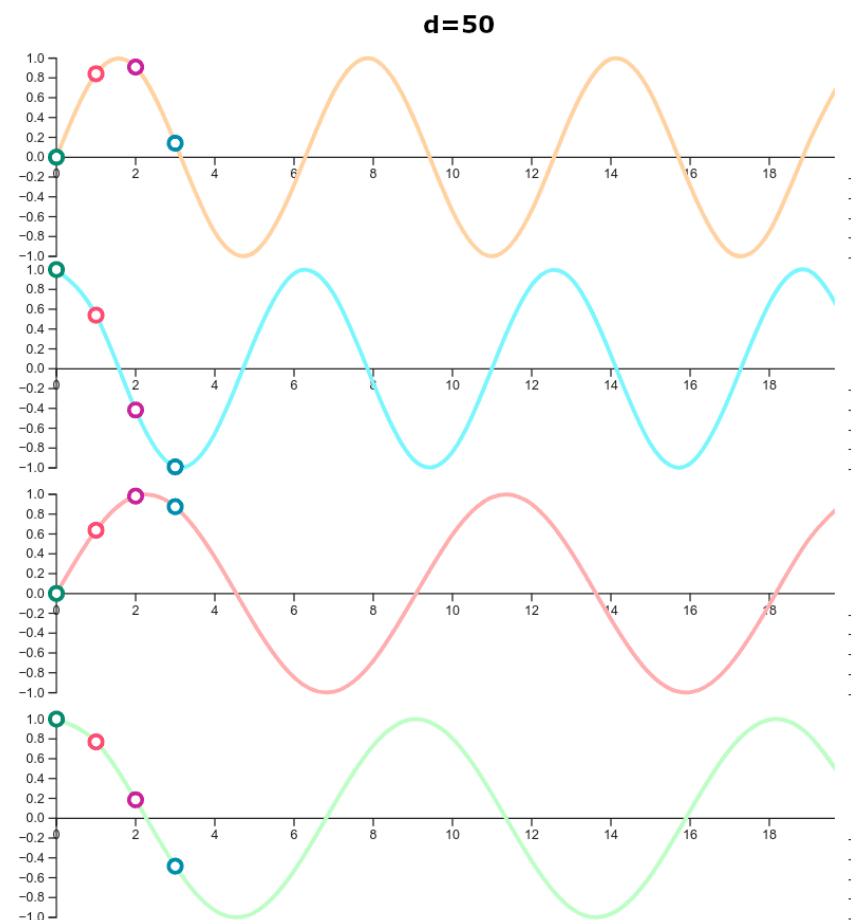
A sequence of tokenized text then turns into a 2D array after the embedding layer, with dimensions:

[sequence length x embedding dim]

Because we will process multiple sequences at once, in so-called 'batches', this will add another dimension to form:

[batches x sequence length x embedding dim]

BYOLLM: Embedding (positional encoding)



$$PE_{(pos, 2i)} = \sin\left(\frac{pos}{10000^{2i/d_{model}}}\right)$$
$$PE_{(pos, 2i + 1)} = \cos\left(\frac{pos}{10000^{2i/d_{model}}}\right)$$

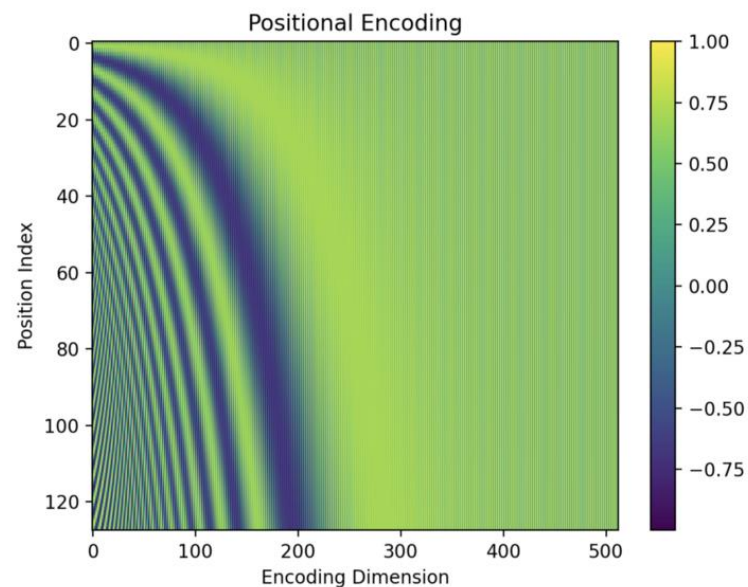
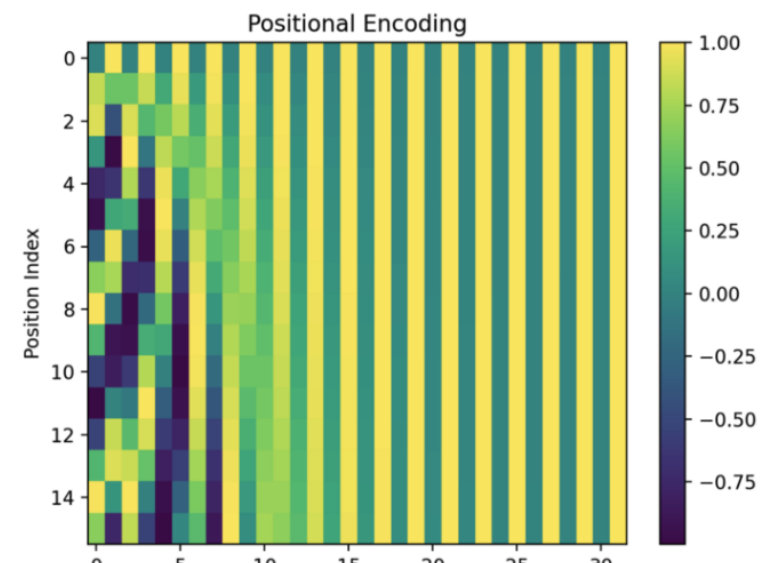
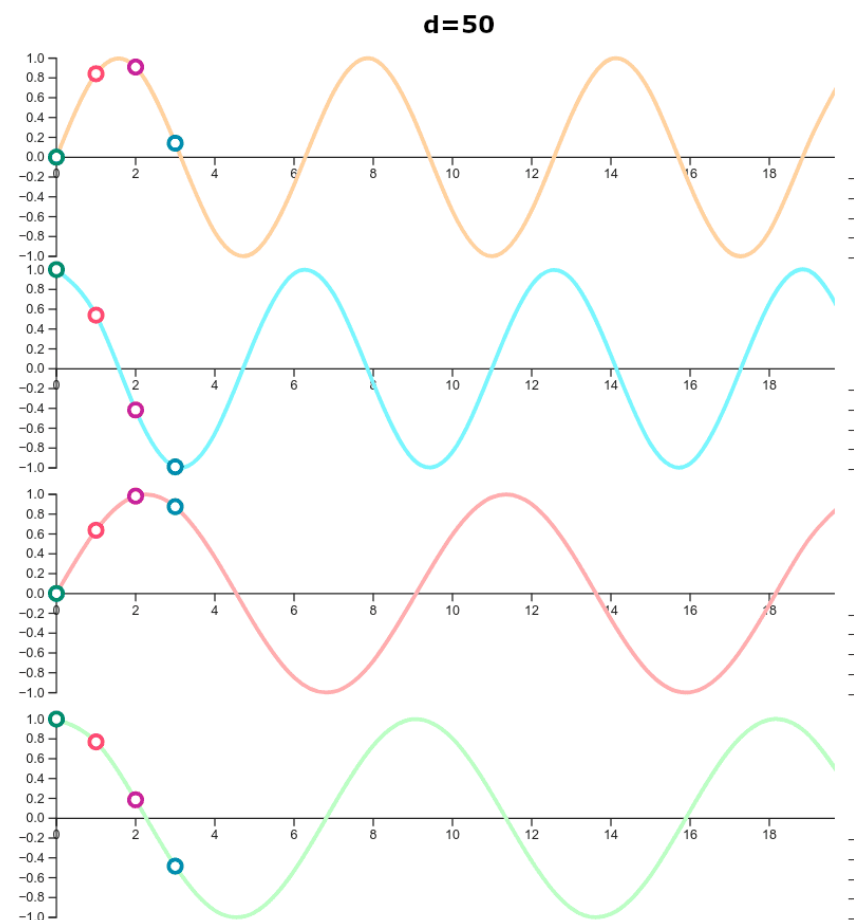
pos: index of token position in the sequence (e.g., 0 for the first token, 1 for the second, and so on).

i: This is the index of the dimension within the embedding vector. It ranges from 0 to $d_{model}/2 - 1$. Notice that each pair of dimensions $(2i, 2i + 1)$ uses the same frequency but different functions (sine and cosine).

d_{model} : This is the dimensionality of the input embeddings (and also the dimension of the positional encoding vector). A typical value might be 512 or 768.

10000: This value is a hyperparameter chosen by the authors; it determines the range of frequencies used.

BYOLLM: Embedding (positional encoding)

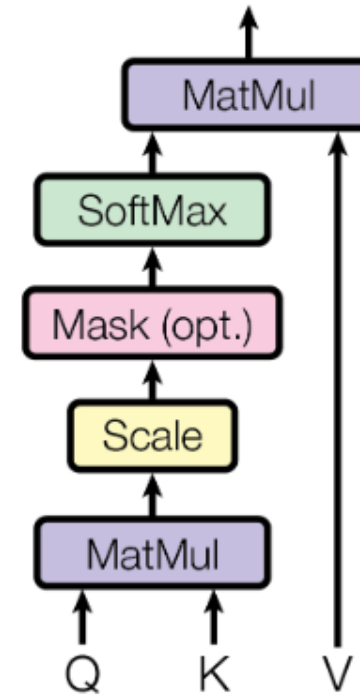
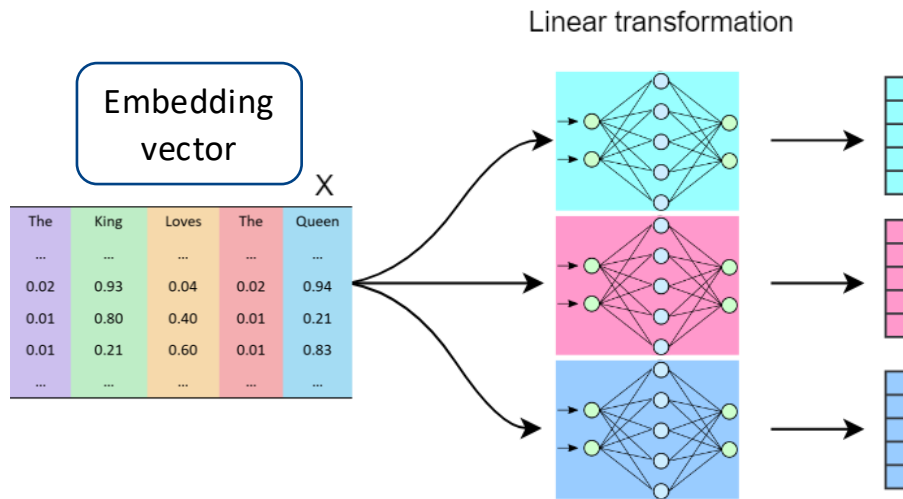


BYOLLM: Embedding (positional encoding)

```
1 # Just run this
2
3 import math
4
5 class PositionalEncodingLayer(nn.Module):
6     def __init__(self, embed_size, max_len=5000):
7         super(PositionalEncodingLayer, self).__init__()
8         self.embed_size = embed_size
9
10        pe = torch.zeros(max_len, embed_size)
11        position = torch.arange(0, max_len, dtype=torch.float).unsqueeze(1)
12        div_term = torch.exp(torch.arange(0, embed_size, 2).float() * (-math.log(10000.0) / embed_size))
13
14        pe[:, 0::2] = torch.sin(position * div_term)
15        pe[:, 1::2] = torch.cos(position * div_term)
16
17        pe = pe.unsqueeze(0)
18        self.register_buffer('pe', pe)
19
20    def forward(self, x):
21        x = x + self.pe[:, :x.size(1), :].detach()
22        return x
23
24 # Initialize the layer
25 pos_encoding_layer = PositionalEncodingLayer(embed_size)
26
27 # Use the layer to add position encoding to our embedded sample!
28 pos_encoded_sample = pos_encoding_layer(batched_embedded_sample)
29
30 # Show the output (optional) and the shape.
31 print(f"Positional encoded output:\n {pos_encoded_sample}")
32 print(f"Positional encoded sample shape:\n {pos_encoded_sample.shape}")
```

BYOLLM: Attention layers

Query-key-value



Query – Key - Value

BYOLLM: Attention layers

Part 4: The attention mechanism

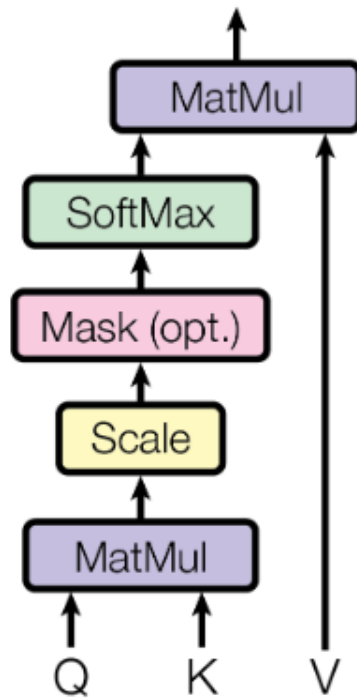
This is where the magic of transformers happens: self-attention. Attention is a way for the model to focus on different parts of the input sequence when making predictions, depending on the next word to be predicted.

In our decoder-only network, which generates text one token at a time, the attention mechanism helps the model decide which parts of the already generated tokens (the context) are important for predicting the next token.

A few key points about the attention mechanism:

1. The input embeddings are transformed into **three matrices, called the query (Q), key (K), and value (V)** matrices using a linear layer.
2. These linear layers have **trainable parameters** which are initialized randomly, but these (attention) parameters are trained during the training of the model.
3. The Q-K-V matrices are combined to produce **a context-aware representation of the input sequence**.
4. The entire process is performed in a sub-unit called the '**attention head**'.
5. **Multiple 'attention heads' are usually used in a parallel configuration** in a system called 'multi-head attention' to efficiently capture relations between tokens in the input sequence.
6. **Multiple 'multi-head attention' blocks are usually stacked in a serial configuration** to attain deeper networks.

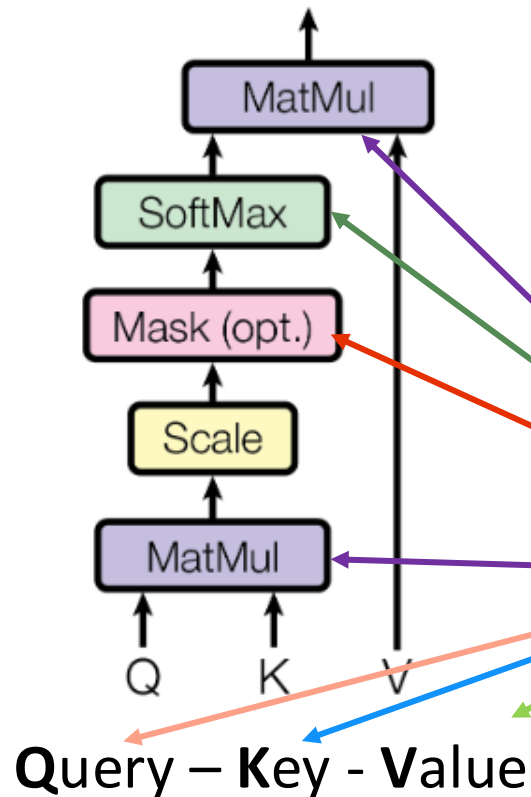
BYOLLM: Attention layers



Query – Key - Value

```
3 class MaskedAttentionLayer(nn.Module):
4     def __init__(self, embed_size, heads = 4):
5         super(MaskedAttentionLayer, self).__init__()
6         self.embed_size = embed_size
7         self.heads = heads
8         self.head_dim = embed_size // heads
9
10        assert (
11            self.head_dim * heads == embed_size
12        ), "Embedding size needs to be divisible by heads"
13
14        # Each
15        self.valuesLayer = nn.Linear(self.head_dim, self.head_dim, bias=False)
16        self.keysLayer = nn.Linear(self.head_dim, self.head_dim, bias=False)
17        self.queriesLayer = nn.Linear(self.head_dim, self.head_dim, bias=False)
18        self.fc_out = nn.Linear(heads * self.head_dim, embed_size)
19
20    def forward(self, values, keys, queries, mask = None):
21        N = queries.shape[0]
22        value_len, key_len, query_len = values.shape[1], keys.shape[1], queries.shape[1]
23
24        # Split the embedding into self.heads different pieces
25        values = values.reshape(N, value_len, self.heads, self.head_dim)
26        keys = keys.reshape(N, key_len, self.heads, self.head_dim)
27        queries = queries.reshape(N, query_len, self.heads, self.head_dim)
28
29        values = self.valuesLayer(values)
30        keys = self.keysLayer(keys)
31        queries = self.queriesLayer(queries)
32
33        # Scaled dot-product attention
34        energy = torch.einsum("nqhd,nkhd->nhqk", [queries, keys])
35
36        # Apply the mask (add -1e10 to masked positions)
37        if mask is not None:
38            energy = energy.masked_fill(mask == 0, float('-1e10'))
39
40        attention = torch.softmax(energy / (self.embed_size ** (1 / 2)), dim=3)
41
42        out = torch.einsum("nhql,nlhd->nqhd", [attention, values]).reshape(
43            N, query_len, self.heads * self.head_dim
44        )
45
46        out = self.fc_out(out)
47        return out
```

BYOLLM: Attention layers



```
3 class MaskedAttentionLayer(nn.Module):
4     def __init__(self, embed_size, heads = 4):
5         super(MaskedAttentionLayer, self).__init__()
6         self.embed_size = embed_size
7         self.heads = heads
8         self.head_dim = embed_size // heads
9
10        assert (
11            self.head_dim * heads == embed_size
12        ), "Embedding size needs to be divisible by heads"
13
14        # Each
15        self.valuesLayer = nn.Linear(self.head_dim, self.head_dim, bias=False)
16        self.keysLayer = nn.Linear(self.head_dim, self.head_dim, bias=False)
17        self.queriesLayer = nn.Linear(self.head_dim, self.head_dim, bias=False)
18        self.fc_out = nn.Linear(heads * self.head_dim, embed_size)
19
20    def forward(self, values, keys, queries, mask = None):
21        N = queries.shape[0]
22        value_len, key_len, query_len = values.shape[1], keys.shape[1], queries.shape[1]
23
24        # Split the embedding into self.heads different pieces
25        values = values.reshape(N, value_len, self.heads, self.head_dim)
26        keys = keys.reshape(N, key_len, self.heads, self.head_dim)
27        queries = queries.reshape(N, query_len, self.heads, self.head_dim)
28
29        values = self.valuesLayer(values)
30        keys = self.keysLayer(keys)
31        queries = self.queriesLayer(queries)
32
33        # Scaled dot-product attention
34        energy = torch.einsum("nqhd,nkhd->nhqk", [queries, keys])
35
36        # Apply the mask (add -1e10 to masked positions)
37        if mask is not None:
38            energy = energy.masked_fill(mask == 0, float('-1e10'))
39
40        attention = torch.softmax(energy / (self.embed_size ** (1 / 2)), dim=3)
41
42        out = torch.einsum("nhql,nlhd->nqhd", [attention, values]).reshape(
43            N, query_len, self.heads * self.head_dim
44        )
45
46        out = self.fc_out(out)
47        return out
```

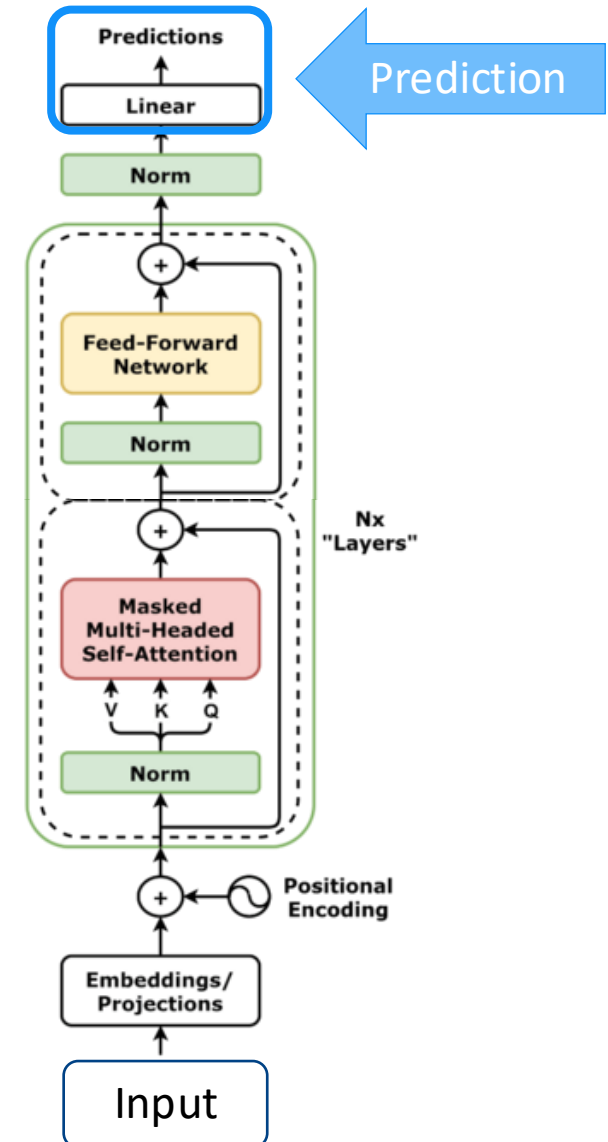
BYOLLM: Output

Final output layer:

- Transformer layers output final embedded vector.
- Linear (output) layer converts embedding to a score for each token in the vocabulary
- These scores are called logits.

Softmax function:

- Softmax turns logits into probabilities for each token in the vocabulary
- All probabilities will sum to 1.
- Token is chosen from multinomial distribution using argmax, top-p, top-k and/or temperature



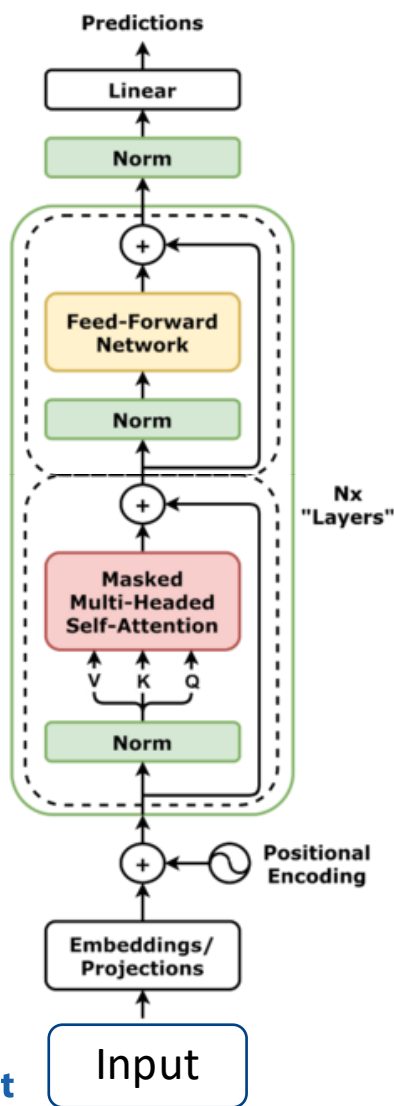
BYOLLM: Output

Final output layer:

- Transformer layers output final embedded vector.
- Linear (output) layer converts embedding to a score for each token in the vocabulary
- These scores are called logits.

```
1 # Just run this
2
3 class OutputLayer(nn.Module):
4     def __init__(self, embed_size, vocab_size):
5         super(OutputLayer, self).__init__()
6         self.fc_out = nn.Linear(embed_size, vocab_size)
7
8     def forward(self, x):
9         return self.fc_out(x)
10
11 # Example usage
12 output_layer = OutputLayer(embed_size, vocab_size)
13 output = output_layer(attention_sample)
14 print(f"Output:\n {output}")
15 print(f"Output shape: {output.shape}")
```

BYOLLM: Building the model

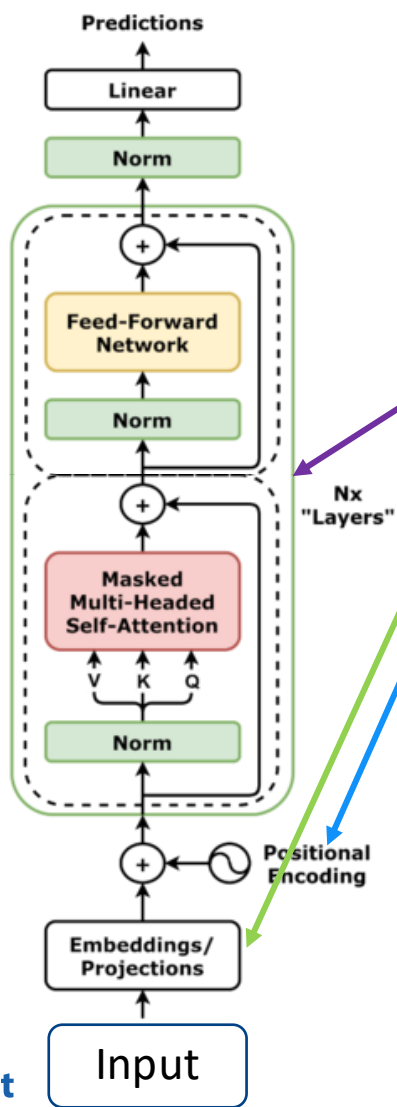


```

4 class SimpleDecoderOnlyTransformerModel(nn.Module):
5     def __init__(self, vocab_size, embed_size, num_heads, num_layers, context_window, max_len=5000):
6         super(SimpleDecoderOnlyTransformerModel, self).__init__()
7
8         # Initialize an embedding layer that converts token indices into dense vectors of size embed_size
9         self.embedding = nn.Embedding(vocab_size, embed_size)
10        # Initialize an embedding layer for positional encoding, giving a unique embedding to each position in the context window
11        self.pos_encoding = PositionalEncodingLayer(embed_size, max_len)
12        # Stack the transformer decoder layers (num_layers times)
13        self.transformer_decoder = nn.TransformerDecoder(
14            nn.TransformerDecoderLayer(embed_size, num_heads, dim_feedforward=4*embed_size, dropout=0.1, batch_first=True),
15            num_layers
16        )
17        # A fully connected layer to map the output of the decoder to the vocabulary size
18        self.fc_out = nn.Linear(embed_size, vocab_size)
19        # Store the maximum sequence length
20        self.max_len = max_len
21        # Store the embedding size
22        self.embed_size = embed_size
23
24    def generate_square_subsequent_mask(self, size):
25        # Generate a lower triangular matrix (mask) to prevent the decoder from attending to future tokens
26        mask = torch.tril(torch.ones(size, size)) # Lower triangular matrix
27        mask = mask.float().masked_fill(mask == 0, float('-inf')).masked_fill(mask == 1, float(0.0)) # Convert to float and mask
28        return mask
29
30    def forward(self, x):
31        device = x.device
32
33        # Trim to last context_window tokens
34        x = x[:, -self.max_len:] # (B, T)
35        # Get token embeddings
36        tok_emb = self.embedding(x) # (B, T, E)
37        # Apply sinusoidal positional encoding
38        x = self.pos_encoding(tok_emb) # (B, T, E)
39        # Generate causal mask
40        tgt_mask = self.generate_square_subsequent_mask(x.shape[1]).to(device) # (T, T)
41        # Pass through Transformer decoder
42        x = self.transformer_decoder(tgt=x, memory=x, tgt_mask=tgt_mask, memory_mask=tgt_mask) # (B, T, E)
43        # Final projection to vocabulary logits
44        x = self.fc_out(x) # (B, T, V)
45
46        return x

```

BYOLLM: Building the model

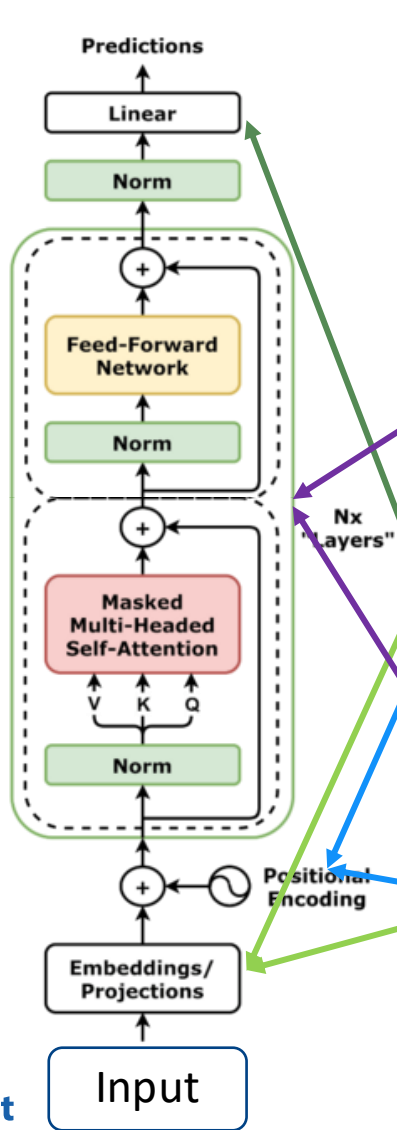


```

4 class SimpleDecoderOnlyTransformerModel(nn.Module):
5     def __init__(self, vocab_size, embed_size, num_heads, num_layers, context_window, max_len=5000):
6         super(SimpleDecoderOnlyTransformerModel, self).__init__()
7
8         # Initialize an embedding layer that converts token indices into dense vectors of size embed_size
9         self.embedding = nn.Embedding(vocab_size, embed_size)
10        # Initialize an embedding layer for positional encoding, giving a unique embedding to each position in the context window
11        self.pos_encoding = PositionalEncodingLayer(embed_size, max_len)
12        # Stack the transformer decoder layers (num_layers times)
13        self.transformer_decoder = nn.TransformerDecoder(
14            nn.TransformerDecoderLayer(embed_size, num_heads, dim_feedforward=4*embed_size, dropout=0.1, batch_first=True),
15            num_layers
16        )
17        # A fully connected layer to map the output of the decoder to the vocabulary size
18        self.fc_out = nn.Linear(embed_size, vocab_size)
19        # Store the maximum sequence length
20        self.max_len = max_len
21        # Store the embedding size
22        self.embed_size = embed_size
23
24    def generate_square_subsequent_mask(self, size):
25        # Generate a lower triangular matrix (mask) to prevent the decoder from attending to future tokens
26        mask = torch.tril(torch.ones(size, size)) # Lower triangular matrix
27        mask = mask.float().masked_fill(mask == 0, float('-inf')).masked_fill(mask == 1, float(0.0)) # Convert to float and mask
28        return mask
29
30    def forward(self, x):
31        device = x.device
32
33        # Trim to last context_window tokens
34        x = x[:, -self.max_len:] # (B, T)
35        # Get token embeddings
36        tok_emb = self.embedding(x) # (B, T, E)
37        # Apply sinusoidal positional encoding
38        x = self.pos_encoding(tok_emb) # (B, T, E)
39        # Generate causal mask
40        tgt_mask = self.generate_square_subsequent_mask(x.shape[1]).to(device) # (T, T)
41        # Pass through Transformer decoder
42        x = self.transformer_decoder(tgt=x, memory=x, tgt_mask=tgt_mask, memory_mask=tgt_mask) # (B, T, E)
43        # Final projection to vocabulary logits
44        x = self.fc_out(x) # (B, T, V)
45
46        return x

```

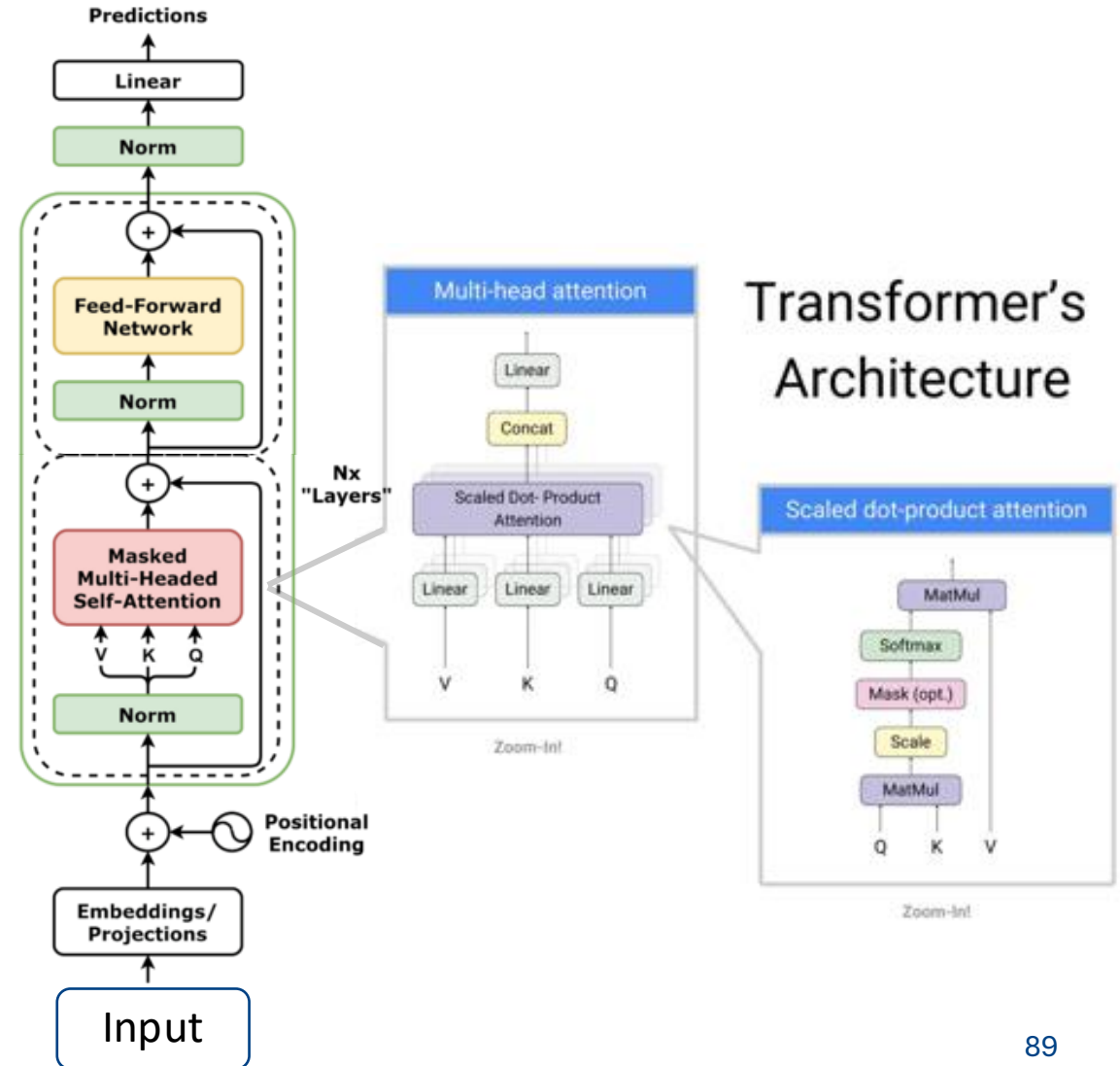
BYOLLM: Building the model



```
4 class SimpleDecoderOnlyTransformerModel(nn.Module):
5     def __init__(self, vocab_size, embed_size, num_heads, num_layers, context_window, max_len=5000):
6         super(SimpleDecoderOnlyTransformerModel, self).__init__()
7
8         # Initialize an embedding layer that converts token indices into dense vectors of size embed_size
9         self.embedding = nn.Embedding(vocab_size, embed_size)
10        # Initialize an embedding layer for positional encoding, giving a unique embedding to each position in the context window
11        self.pos_encoding = PositionalEncodingLayer(embed_size, max_len)
12        # Stack the transformer decoder layers (num_layers times)
13        self.transformer_decoder = nn.TransformerDecoder(
14            nn.TransformerDecoderLayer(embed_size, num_heads, dim_feedforward=4*embed_size, dropout=0.1, batch_first=True),
15            num_layers
16        )
17        # A fully connected layer to map the output of the decoder to the vocabulary size
18        self.fc_out = nn.Linear(embed_size, vocab_size)
19        # Store the maximum sequence length
20        self.max_len = max_len
21        # Store the embedding size
22        self.embed_size = embed_size
23
24    def generate_square_subsequent_mask(self, size):
25        # Generate a lower triangular matrix (mask) to prevent the decoder from attending to future tokens
26        mask = torch.tril(torch.ones(size, size)) # Lower triangular matrix
27        mask = mask.float().masked_fill(mask == 0, float('-inf')).masked_fill(mask == 1, float(0.0)) # Convert to float and mask
28        return mask
29
30    def forward(self, x):
31        device = x.device
32
33        # Trim to last context_window tokens
34        x = x[:, -self.max_len:] # (B, T)
35        # Get token embeddings
36        tok_emb = self.embedding(x) # (B, T, E)
37        # Apply sinusoidal positional encoding
38        x = self.pos_encoding(tok_emb) # (B, T, E)
39        # Generate causal mask
40        tgt_mask = self.generate_square_subsequent_mask(x.shape[1]).to(device) # (T, T)
41        # Pass through Transformer decoder
42        x = self.transformer_decoder(tgt=x, memory=x, tgt_mask=tgt_mask, memory_mask=tgt_mask) # (B, T, E)
43        # Final projection to vocabulary logits
44        x = self.fc_out(x) # (B, T, V)
45
46        return x
```

BYOLLM: Building the model

```
5 # Example variables, change them around to see their effect!
6 vocab_size = len(stoi)
7 context_window = 32
8 embed_size = 48
9 num_heads = 3
10 num_layers = 3
11
12 # Instantiate the model
13 model = SimpleDecoderOnlyTransformerModel(vocab_size,
14                                           embed_size,
15                                           num_heads,
16                                           num_layers,
17                                           context_window)
18 # If you have a gpu, we will use that
19 device = 'cuda' if torch.cuda.is_available() else 'cpu'
20 model.to(device)
21 print(device)
22
23 # Show a summary of the model
24 torchinfo.summary(model, depth = 4)
```



BYOLLM: Batching

Batching:

- **Group samples** for parallel processing -> more efficient training
- **Larger batches** smooth out gradients -> less noisy training

```
3 batch_size_dummy = 4
4
5 def generate_batch(data_split, batch_size, context_window):
6     # Generate a small batch of input sequences and target sequences
7     dataset = train_data if data_split == 'train' else val_data
8     random_indices = torch.randint(len(dataset) - context_window, (batch_size,))
9     input_sequences = torch.stack([dataset[i:i+context_window] for i in random_indices])
10    target_sequences = torch.stack([dataset[i+1:i+context_window+1] for i in random_indices])
11    return input_sequences, target_sequences
12
13 # A batch of data will then look like this
14 input_example, target_example = generate_batch('train', batch_size_dummy, context_window)
15 print(f"Input tokens:\n {input_example}")
16 print(f"Target tokens:\n {target_example}")
```

BYOLLM: Generation

Generation function

- **Prediction:**
Model takes recent tokens as input.
- **Temperature:**
Scales logits; higher T means more random sampling.
- **Softmax:**
From logits to probabilities
- **Sampling:**
Multinomial distribution is used to choose the next token. (could also be highest prob. token!)
- **Loop:**
Process repeats until max length is reached. (could also be a stopping token!)

```
6 def generate_text(model, max_length, vocab_size, start_token=None, temperature=1):
7
8     model.eval() # set model to evaluation mode
9
10    # Determine starting token
11    if start_token is None:
12        start_token = [random.randint(0, vocab_size - 1)] # random start if not provided
13    else:
14        # Convert string to token ID (if start_token is a string)
15        if isinstance(start_token, str):
16            start_token = tokenize(start_token)
17
18    generated_tokens = start_token # initialize list of generated tokens
19
20    print(detokenize(start_token), end="", flush=True)
21
22    for _ in range(max_length - 1):
23        # Prepare input tensor of the last `context_window` tokens
24        input_tensor = torch.tensor(generated_tokens[-context_window:]).unsqueeze(0) # shape (1, seq_len)
25
26        with torch.no_grad(): # no gradient needed for inference
27            output = model(input_tensor.to(device)) # forward pass
28
29        # Get logits for the next token (last position)
30        next_token_logits = output[0, -1, :] / temperature
31        # also applied temperature scaling (higher temperature is more random)
32
33        # Convert logits to probabilities
34        next_token_probs = F.softmax(next_token_logits, dim=-1)
35
36        # Sample the next token from the probability distribution
37        next_token = torch.multinomial(next_token_probs, num_samples=1).item()
38
39        # Add the new token to the sequence
40        generated_tokens.append(next_token)
41
42        # Print the new token immediately (streaming output)
43        print(detokenize([next_token]), end="", flush=True)
44
45    print() # final newline after finishing generation
46
47    return generated_tokens
48
```

BYOLLM: Training

Transformer Training: Settings and Loop

- **Hyperparameters:** Learning rate and batch size are important for convergence.
- **Optimizer:** Adam algorithm adjusts model learn rate according to individual parameters gradients.
- **Loss Function:** CrossEntropy calculates the prediction error.
- **Update model:** Optimizer uses loss to calculate gradients and update model.
- **Evaluation:** Validation loss checks performance on unseen data.

```
8 # Hyperparameters
9 num_epochs = 10
10 learning_rate = 0.0001
11 batch_size = 32
12 max_num_its_per_epoch = 300
13
14 # Initialize model, loss, and optimizer
15 criterion = nn.CrossEntropyLoss()
16 optimizer = torch.optim.Adam(model.parameters(), lr=learning_rate)
17 train_losses = []
18 val_losses = []
19
20 # Training and validation functions
21 def run_epoch(model, data, is_training):
22     epoch_loss = []
23     model.train() if is_training else model.eval()
24
25     num_its_per_epoch = min(max_num_its_per_epoch, len(data) // (batch_size*context_window))
26     iterator = tqdm(range(num_its_per_epoch), desc="Training" if is_training else "Validating")
27
28     for i in iterator:
29         inputs, targets = generate_batch('train' if is_training else 'validate', batch_size, context_window)
30         outputs = model(inputs.to(device))
31
32         B, T, V = outputs.shape
33         outputs = outputs.view(B*T, V)
34         targets = targets.view(B*T)
35         loss = criterion(outputs, targets.to(device))
36         epoch_loss.append(loss.item())
37
38         if is_training:
39             optimizer.zero_grad(set_to_none=True)
40             loss.backward()
41             optimizer.step()
42
43     iterator.set_postfix({'loss': np.mean(epoch_loss)})
44
45     return np.mean(epoch_loss)
46
47 # Training loop
48 for epoch in range(num_epochs):
49     train_loss = run_epoch(model, train_data, is_training=True)
50     val_loss = run_epoch(model, val_data, is_training=False)
51
52     train_losses.append(train_loss)
53     val_losses.append(val_loss)
54
55     print(f'Epoch [{epoch+1}/{num_epochs}], Train Loss: {train_loss:.4f}, Val Loss: {val_loss:.4f}')
56     print(f'Quick example of current text output:')
57     generate_text(model, max_length, vocab_size, temperature = 1, start_token = 'The')
58     print()
```

Practical: BYOLLM

Now it's your turn!

You can find instructions and questions on Canvas

CHALLENGES:

- Try running your model using the GPU for a major speed-up!
- Get the best validation loss on a dataset by tuning hyperparameters (compare with colleagues!)
- Make your model generate deterministically
- Try both character and sub-word tokenization (what works best?)
- With sub-word tokenization, try one of the new, larger datasets!
- Find, load and train your own dataset from Huggingface

[LINK TO NOTEBOOK \(subword tokenization\)](#)