

Building GPT Architecture: Implementing Core Model Components

Prof. Rongyu Lin

Lecture 5

Quinnipiac University School of Computing and Engineering

Reading: Raschka Ch. 4 (pp. 96-149) "Building GPT Architecture"

Learning Objectives

By the end of this lecture, you will be able to:

- **Implement Causal Attention:** Apply masking techniques to prevent future token access
- **Construct Multi-Head Attention:** Combine multiple attention heads for richer representations
- **Build Neural Network Components:** Layer normalization, GELU activation, feed-forward networks
- **Assemble Transformer Blocks:** Integrate attention and feed-forward with shortcut connections
- **Code Complete GPT Architecture:** Build a functional GPT model ready for training
- **Understand Text Generation:** Token-by-token generation in autoregressive models

From attention mechanisms to complete GPT implementation

Course Context: Where We Are

Building Towards Complete LLM Understanding

Completed Foundation:

- Lecture 1: AI Historical Perspective
- Lecture 2: LLM Foundations & Pre-training
- Lecture 3: Text Data Processing & Tokenization
- Lecture 4: Attention Mechanisms

Today - Lecture 5: GPT Architecture

- Causal attention implementation
- Neural network building blocks
- Complete transformer assembly

Coming Next:

- Lecture 6: Pre-training
- Lecture 7: Fine-tuning
- Lecture 8: Advanced Architecture

Today: Bridge from attention to complete architecture

Section 1

Advanced Attention Mechanisms

Causal Masking & Multi-Head Implementation

Why Causal Attention? Preventing Future Information Leakage

The Problem with Standard Attention

Standard Attention:

- All tokens attend to all tokens
- Bidirectional information flow
- **Problem:** Can't see the future!

1	0	0	0
1	1	0	0
1	1	1	0
1	1	1	1

Causal Mask

Causal Attention Solution:

- Mask future token positions
- Maintain autoregressive property
- Each token attends only to previous positions

Key Insight

Autoregressive generation requires preventing access to future information.

Causal Masking Matrix Visualization

Interactive Lower Triangular Matrix Example



Key Observation: Only lower triangular attention weights are computed

Before Masking:

- **Blue:** Raw attention scores (0.25)
- All positions accessible

After Masking + Softmax:

- **Gray:** Future positions become 0.0
- **Gold:** Renormalized valid weights

Key Insight

Each token can only attend to itself and previous tokens, preventing information leakage from future positions during training.

Mathematical Foundation of Causal Masking

```
1 # Create lower triangular mask
2 mask = torch.tril(torch.ones(seq_len, seq_len))
3
4 # Apply mask to attention scores
5 attention_scores = attention_scores.masked_fill(
6     mask == 0, float('-inf'))
```

Attention Weight Modification

Standard Attention:

$$\text{Attention}(Q, K, V) = \text{softmax}\left(\frac{QK^T}{\sqrt{d_k}}\right) V$$

Causal Attention with Masking:

$$\text{Attention}_{\text{causal}}(Q, K, V) = \text{softmax}\left(\frac{QK^T + M}{\sqrt{d_k}}\right) V$$

Mask Matrix M :

- $M_{i,j} = 0$ if $j \leq i$ (allowed)
- $M_{i,j} = -\infty$ if $j > i$ (masked)

Why $-\infty$?

After softmax, $\text{softmax}(-\infty) = 0$, eliminating attention to future positions.

Step-by-Step Causal Attention Implementation

```
1 import torch
2 import torch.nn as nn
3 import torch.nn.functional as F
4
5 class CausalAttention(nn.Module):
6     def __init__(self, d_in, d_out,
7                 context_length, dropout=0.0):
8         super().__init__()
9         self.d_out = d_out
10        self.W_query = nn.Linear(d_in, d_out,
11                                bias=False)
12        self.W_key = nn.Linear(d_in, d_out,
13                              bias=False)
14        self.W_value = nn.Linear(d_in, d_out,
15                                 bias=False)
16        self.dropout = nn.Dropout(dropout)
17
18        # Register causal mask as buffer
19        self.register_buffer(
20            'mask',
21            torch.tril(torch.ones(
22                context_length, context_length))
23        )
```

Building Causal Attention Layer Key

Design Decisions:

- **register_buffer**: Mask moves with model to GPU/CPU
- **bias=False**: Follows GPT-2 design
- **torch.tril**: Lower triangular matrix for masking

Architecture Notes:

- Creates query, key, and value projections
- Stores mask as non-trainable buffer
- Dropout for regularization
- Context length defines maximum sequence size

Efficient Masking Implementation

```
1 def forward(self, x):
2     b, num_tokens, d_in = x.shape # batch, sequence, embedding
3
4     # Generate queries, keys, values
5     queries = self.W_query(x) # (b, num_tokens, d_out)
6     keys = self.W_key(x) # (b, num_tokens, d_out)
7     values = self.W_value(x) # (b, num_tokens, d_out)
8
9     # Compute attention scores
10    attention_scores = queries @ keys.transpose(-2, -1) # (b,
11    num_tokens, num_tokens)
12    attention_scores = attention_scores / (self.d_out ** 0.5)
13    # Scaling
14
15    # Apply causal mask
16    attention_scores = attention_scores.masked_fill(
17        self.mask[:num_tokens, :num_tokens] == 0, float('-inf'))
18
19    # Apply softmax and dropout
20    attention_weights = F.softmax(attention_scores, dim=-1)
21    attention_weights = self.dropout(attention_weights)
22
23    # Generate context vectors
24    context = attention_weights @ values # (b, num_tokens,
    d_out)
    return context
```

Forward Pass with Masking Key

Implementation Steps:

- **Shape Extraction:** Get batch size, sequence length, and embedding dimensions
- **Linear Projections:** Transform input to queries, keys, and values
- **Attention Computation:** Calculate attention scores with scaling
- **Causal Masking:** Apply mask to prevent future information leakage
- **Normalization:** Softmax to get attention weights
- **Regularization:** Apply dropout for training stability
- **Context Generation:** Weight values to produce final output

Information Leakage Prevention: Why It Works

Mathematical Properties of Masked Softmax

Before Masking:

$$\text{softmax}(z_i) = \frac{e^{z_i}}{\sum_j e^{z_j}} \quad (1)$$

After Masking:

$$\text{softmax}(z_i + M_i) = \frac{e^{z_i + M_i}}{\sum_j e^{z_j + M_j}} \quad (2)$$

Where $M_j = -\infty$ for $j > i$:

$$e^{z_j - \infty} = 0 \quad (3)$$

$$\text{Effectively: } \sum_j \rightarrow \sum_{j \leq i} \quad (4)$$

Weight Redistribution:

Before Masking:

0.25	0.25	0.25	0.25
T1	T2	T3	T4

After Masking:

1.0	0.0	0.0	0.0
T1	T2	T3	T4

Renormalization Effect

Masked positions get zero weight, remaining weights sum to 1.0. No information leakage.

◆ Hands-on Exercise: Implement Causal Masking

Interactive Coding Activity (5 minutes)

Your Task: Complete the causal masking implementation

```
1 import torch
2
3
4 def create_causal_mask(seq_length):
5     """Create a causal mask for attention."""
6     # TODO: Create lower triangular mask
7     mask = torch.tril(torch.ones(seq_length, seq_length))
8     return mask
9
10 def apply_causal_mask(attention_scores, mask):
11     """Apply causal mask to attention scores."""
12     # TODO: Apply mask using masked_fill
13     masked_scores = attention_scores.masked_fill(
14         mask == 0, float('-inf'))
15
16     return masked_scores
17
18 # Test your implementation
19 seq_len = 4
20 mask = create_causal_mask(seq_len)
21 print("Causal mask:")
22 print(mask)
```

Key Concepts: `torch.tril()`:

- Creates lower triangular matrix
- 1s below diagonal, 0s above

`masked_fill()`:

- Replaces values where condition is True
- Uses $-\infty$ to zero out after softmax

Expected Output:

$$\text{mask} = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 1 & 1 & 0 & 0 \\ 1 & 1 & 1 & 0 \\ 1 & 1 & 1 & 1 \end{bmatrix} \quad (5)$$

Tip: Test with `seq_len=4` first!

Dropout for Regularization in Attention

Preventing Overfitting in Attention Weights

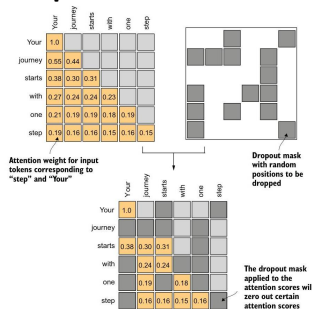
Why Dropout in Attention?

- Attention weights can become concentrated
- Model may over-rely on specific tokens
- Dropout forces distributed attention
- Improves generalization performance

Implementation Strategy:

- Apply dropout **after** softmax
- During training: randomly zero weights
- Remaining weights scaled up (compensation)
- During inference: use all weights

Dropout Visualization



Dropout Formula

$$\text{Dropout}(x) = \begin{cases} 0 & \text{with prob. } p \\ \frac{x}{1-p} & \text{with prob. } 1-p \end{cases}$$

Dropout Implementation in Attention

```
1 class CausalAttention(nn.Module):
2     def __init__(self, d_in, d_out, context_length,
3         dropout=0.0):
4         super().__init__()
5         # ... (previous initialization code)
6         self.dropout = nn.Dropout(dropout) # Dropout layer
7
8     def forward(self, x):
9         # ... (compute attention_scores)
10
11        # Apply softmax to get attention weights
12        attention_weights = F.softmax(attention_scores, dim=-1)
13
14        # Apply dropout AFTER softmax
15        attention_weights = self.dropout(attention_weights)
16
17        # Generate context vectors
18        context = attention_weights @ values
19        return context
```

PyTorch Implementation Details

Key Implementation Points:

- **Timing:** Dropout applied after softmax normalization
- **Training vs Inference:** Automatic handling with `model.train()/eval()`
- **Scaling:** PyTorch automatically handles $\frac{1}{1-p}$ scaling
- **Typical Values:** 0.0-0.1 for attention dropout

Dropout Benefits:

- Prevents overfitting to specific attention patterns
- Forces model to use diverse attention strategies
- Improves generalization performance
- Standard practice in transformer architectures

Complete Causal Attention Class

```
1 class CausalAttention(nn.Module):
2     def __init__(self, d_in, d_out, context_length,
3         dropout=0.0, qkv_bias=False):
4         super().__init__()
5         self.d_out = d_out
6         self.W_query = nn.Linear(d_in, d_out, bias=qkv_bias)
7         self.W_key = nn.Linear(d_in, d_out, bias=qkv_bias)
8         self.W_value = nn.Linear(d_in, d_out, bias=qkv_bias)
9         self.dropout = nn.Dropout(dropout)
10        # Register buffer for causal mask
11        self.register_buffer('mask', torch.tril(
12            torch.ones(context_length, context_length)))
13    def forward(self, x):
14        b, num_tokens, d_in = x.shape
15        queries = self.W_query(x)
16        keys = self.W_key(x)
17        values = self.W_value(x)
18        attention_scores = queries @ keys.transpose(-2, -1)
19        attention_scores = attention_scores / (self.d_out **
20            0.5)
21        attention_scores = attention_scores.masked_fill(
22            self.mask[:num_tokens, :num_tokens] == 0,
23            float('-inf'))
24        attention_weights = F.softmax(attention_scores, dim=-1)
25        attention_weights = self.dropout(attention_weights)
26        context = attention_weights @ values
27        return context
```

Full Implementation with Best Practices Class Architecture:

- **Linear Layers:** Query, key, value projections
- **Causal Mask:** Lower triangular mask as buffer
- **Dropout:** Regularization after softmax
- **Scaling:** Temperature scaling with $\sqrt{d_{out}}$

Key Features:

- GPT-2 compatible (bias=False default)
- Efficient masking with register_buffer
- Flexible dropout and bias options
- Production-ready implementation

Usage:

- d_in=768, d_out=768 (GPT-2 Base)
- context_length=1024 or 2048
- dropout=0.1 (typical value)

Batch Processing in Causal Attention

Handling Multiple Sequences Efficiently

Key Tensor Transformations:

Operation	Input	Output
Input	(B, T, D_{in})	-
Linear (Q,K,V)	(B, T, D_{in})	(B, T, D_{out})
Attention Scores	(B, T, D_{out})	(B, T, T)
Masked Softmax	(B, T, T)	(B, T, T)
Context	(B, T, T)	(B, T, D_{out})

Batch Efficiency Considerations:

- **Mask Broadcasting:** Same mask for all batch elements
- **Parallel Processing:** All sequences processed simultaneously
- **Memory Scaling:** Usage scales as $O(B \times T^2)$

Memory Consideration

For long sequences, attention memory grows quadratically: 1000 tokens $\rightarrow$ 1M weights per head

Testing the Causal Attention Module

```
1 # Test causal attention implementation
2 torch.manual_seed(123)
3
4 # Initialize causal attention
5 d_in, d_out = 768, 768
6 context_length = 1024
7 dropout = 0.1
8
9 causal_attn = CausalAttention(d_in, d_out, context_length,
10                               dropout)
11
12 # Test input
13 batch_size, seq_len = 2, 10
14 x = torch.randn(batch_size, seq_len, d_in)
15
16 # Forward pass
17 with torch.no_grad():
18     context = causal_attn(x)
19
20 print(f"Input shape: {x.shape}")
21 print(f"Output shape: {context.shape}")
22 print(f"Shape preserved: {x.shape == context.shape}")
23
24 # Verify causal mask is working
25 print(f"Mask shape: {causal_attn.mask.shape}")
26 print(f"Mask (first 5x5):")
27 print(causal_attn.mask[:5, :5])
```

Validation and Shape Verification

Testing Strategy:

- **Reproducibility:** Fixed random seed for consistent results
- **Shape Testing:** Verify input/output dimension preservation
- **Mask Verification:** Check causal mask structure
- **Numerical Stability:** Ensure no NaN/Inf values

Expected Behavior:

- Input/output shapes should be identical
- Mask should be lower triangular
- Context vectors meaningful (not NaN/Inf)

Key Outputs:

- Input shape: torch.Size([2, 10, 768])
- Output shape: torch.Size([2, 10, 768])
- Shape preserved: True
- Mask: Lower triangular matrix

Multi-Head Attention: Parallel Representation Learning

Why Multiple Attention Heads?

Single Head Limitations:

- One representation subspace
- Limited perspective on relationships
- Single attention pattern per layer

Multi-Head Advantages:

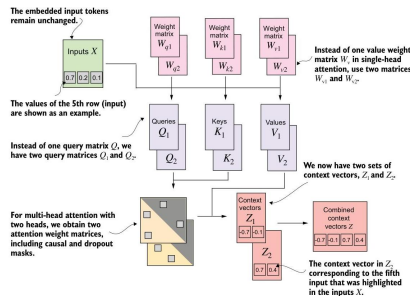
- **Multiple subspaces:** Different relationship aspects
- **Diverse patterns:** Syntactic + semantic attention
- **Parallel processing:** Computational efficiency
- **Richer representations:** Combined knowledge

Mathematical Formulation:

$$\text{MultiHead}(Q, K, V) = \text{Concat}(\text{head}_1, \dots, \text{head}_h)W^O$$

$$\text{where } \text{head}_i = \text{Attention}(QW_i^Q, KW_i^K, VW_i^V)$$

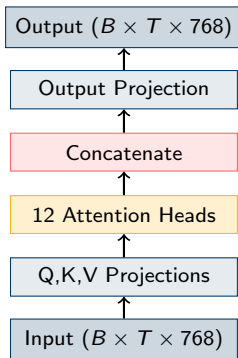
Multi-Head Architecture



Detailed Multi-Head Attention Architecture

Complete 12-Head GPT Implementation

Simplified Architecture Flow



Key Components:

- **Input:** $B \times T \times 768$
- **Q,K,V:** 3 linear layers (768×768)
- **12 Heads:** Each processes 64 dims
- **Concat:** Merge all head outputs
- **Output:** Final projection layer

Head Dimensions:

- **Total:** 768 dimensions
- **Per head:** $768 \div 12 = 64$
- **Parallel processing**

Parameters:

- **Q, K, V:** 3×768^2
- **Output:** 768^2
- **Total:** 2.36M params

Simple Multi-Head Implementation: Wrapper Approach

```
1 class MultiHeadAttentionWrapper(nn.Module):
2     def __init__(self, d_in, d_out, context_length, dropout,
3         num_heads):
4         super().__init__()
5         self.heads = nn.ModuleList([
6             CausalAttention(d_in, d_out, context_length,
7                 dropout)
8             for _ in range(num_heads)
9         ])
10
11     def forward(self, x):
12         # Run each head independently
13         head_outputs = [head(x) for head in self.heads]
14         # Concatenate along the feature dimension
15         return torch.cat(head_outputs, dim=-1)
```

Output Shape: If each head outputs (B, T, d_{out}) , final is $(B, T, num_heads \times d_{out})$

Usage Example: 12 heads $\times$ 64 dimensions = 768 total (GPT-2 configuration)

Conceptually Clear Implementation Architecture Design:

- **ModuleList:** Separate attention heads
- **Sequential Processing:** Each head runs independently
- **Concatenation:** Combine outputs along feature dimension

Pros and Cons:

- ✓ **Simple to understand:** Clear head separation
- ✓ **Easy debugging:** Individual head inspection
- ✗ **Inefficient:** Multiple forward passes
- ✗ **Memory overhead:** Separate parameters per head

Efficient Multi-Head Attention Implementation (Part 1)

Initialization and Setup

```
1 class MultiHeadAttention(nn.Module):
2     def __init__(self, d_in, d_out, context_length, dropout, num_heads):
3         super().__init__()
4         assert d_out % num_heads == 0, "d_out must be divisible by num_heads"
5         self.d_out = d_out
6         self.num_heads = num_heads
7         self.head_dim = d_out // num_heads
8
9         # Single large matrices instead of separate heads
10        self.W_query = nn.Linear(d_in, d_out, bias=False)
11        self.W_key = nn.Linear(d_in, d_out, bias=False)
12        self.W_value = nn.Linear(d_in, d_out, bias=False)
13        self.out_proj = nn.Linear(d_out, d_out) # Output projection
14
15        self.dropout = nn.Dropout(dropout)
16        self.register_buffer('mask', torch.tril(torch.ones(
17            context_length, context_length)))
```

Key Design Decisions:

- **Single Matrices:** One large matrix for all heads (efficient computation)
- **Head Dimension:** $d_{head} = d_{out} / num_heads$ (typically 64)
- **Output Projection:** Final linear layer after concatenation
- **Causal Mask:** Registered as buffer for GPU compatibility

Efficient Multi-Head Attention Implementation (Part 2)

```
1 def forward(self, x):
2     b, num_tokens, d_in = x.shape
3     # Generate Q, K, V matrices
4     queries = self.W_query(x) # (b, num_tokens, d_out)
5     keys = self.W_key(x)      # (b, num_tokens, d_out)
6     values = self.W_value(x)  # (b, num_tokens, d_out)
7     # Reshape for multi-head processing
8     queries = queries.view(b, num_tokens, self.num_heads, self.head_dim)
9     queries = queries.transpose(1, 2) # (b, num_heads, num_tokens, head_dim)
10    keys = keys.view(b, num_tokens, self.num_heads, self.head_dim)
11    keys = keys.transpose(1, 2)
12    values = values.view(b, num_tokens, self.num_heads, self.head_dim)
13    values = values.transpose(1, 2)
14    # Compute attention scores and weights
15    attention_scores = queries @ keys.transpose(-2, -1)
16    attention_scores = attention_scores / (self.head_dim ** 0.5)
17    attention_scores = attention_scores.masked_fill(
18        self.mask[:num_tokens, :num_tokens] == 0, float('-inf'))
19    attention_weights = F.softmax(attention_scores, dim=-1)
20    attention_weights = self.dropout(attention_weights)
21    # Apply attention and concatenate heads
22    context = attention_weights @ values # (b, num_heads, num_tokens, head_dim)
23    context = context.transpose(1, 2).contiguous()
24    context = context.view(b, num_tokens, self.d_out)
25    # Final output projection
26    context = self.out_proj(context)
27    return context
```

Tensor Transformations:

- Input: (B, T, D_{in})
- After Linear: (B, T, D_{out})
- Reshape: (B, T, H, D_h)
- Transpose: (B, H, T, D_h)
- Final: (B, T, D_{out})

Efficiency Benefits:

- Single forward pass for all heads
- Better memory utilization
- GPU-friendly parallel computation
- Reduced parameter overhead

◆ Interactive Exercise: Multi-Head vs Single Head

```
1 # Test both implementations
2 torch.manual_seed(123)
3 d_in, d_out = 768, 768
4 context_length = 1024
5 dropout = 0.0
6 num_heads = 12
7
8 # Create both versions
9 wrapper_mha = MultiHeadAttentionWrapper(d_in, d_out//num_heads,
10 context_length, dropout, num_heads)
11 efficient_mha = MultiHeadAttention(d_in, d_out, context_length,
12 dropout, num_heads)
13
14 # Test input
15 x = torch.randn(2, 10, d_in)
16
17 # TODO: Complete the comparison
18 with torch.no_grad():
19     output_wrapper = wrapper_mha(x)
20     output_efficient = efficient_mha(x)
21
22 print(f"Wrapper output shape:
23       {output_wrapper.shape}")
24 print(f"Efficient output shape:
25       {output_efficient.shape}")
```

```
1 # TODO: Count parameters
2 wrapper_params = sum(p.numel()
3     for p in wrapper_mha.parameters())
4 efficient_params = sum(p.numel()
5     for p in efficient_mha.parameters())
6
7 print(f"Wrapper parameters:
8       {wrapper_params}")
9 print(f"Efficient parameters:
10      {efficient_params}")
11 print(f"Efficiency ratio:
12      {efficient_params / wrapper_params:.2f}")
```

Discussion questions:

- Which implementation is more memory efficient?
- Which would be faster for large inputs?
- What are the trade-offs?

Section 2

Neural Network Building Blocks

GPT Architecture Components

GPT Model Overview: Three Stages of Development

From Components to Complete Language Model

Stage 1: Pre-training

- Next-token prediction on large corpus
- Learns general language patterns
- Unsupervised learning approach

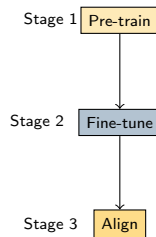
Stage 2: Supervised Fine-tuning

- Task-specific training data
- Classification, question answering, etc.
- Adapts knowledge to specific tasks

Stage 3: Alignment & RLHF

- Human feedback optimization
- Safety and helpfulness alignment
- Reinforcement learning from preferences

Development Pipeline



Current Status

Today: Building architecture foundation for pre-training

Next Lectures: Training procedures, fine-tuning, and alignment

GPT-2 Model Specifications: Scaling Parameters

Model Size Configurations

Model	Parameters	Layers	Heads	Embedding Dim
GPT-2 Small	124M	12	12	768
GPT-2 Medium	345M	24	16	1024
GPT-2 Large	762M	36	20	1280
GPT-2 XL	1.5B	48	25	1600

Our Focus: GPT-2 Small (124M) Configuration

- **Vocabulary Size:** 50,257 tokens (BPE encoding)
- **Context Length:** 1024 tokens maximum
- **Architecture:** 12 transformer blocks $\times$ 12 attention heads
- **Feed-forward:** $4\times$ expansion ($768 \rightarrow 3072 \rightarrow 768$)
- **Training Data:** 40GB of internet text

Parameter Breakdown:

- Token Embeddings: $\sim 39\text{M}$
- Position Embeddings: $\sim 0.8\text{M}$
- Transformer Blocks: $\sim 80\text{M}$
- Output Layer: $\sim 39\text{M}$

Computational Requirements:

- Training: ~ 100 GPU-days
- Memory: $\sim 500\text{MB}$ (inference)
- Speed: ~ 1000 tokens/sec

Model Configuration and Initialization

```
1 # GPT-2 Small configuration
2 GPT_CONFIG_124M = {
3     "vocab_size": 50257,      # Tokens
4     "context_length": 1024,  # Seq length
5     "emb_dim": 768,          # Embedding
6     "n_heads": 12,           # Attention
7     "n_layers": 12,          # Blocks
8     "drop_rate": 0.1,        # Dropout
9     "qkv_bias": False        # QKV bias
10 }
11
12 class GPTModel(nn.Module):
13     def __init__(self, cfg):
14         super().__init__()
15         self.tok_emb = nn.Embedding(
16             cfg["vocab_size"], cfg["emb_dim"])
17         self.pos_emb = nn.Embedding(
18             cfg["context_length"], cfg["emb_dim"])
19         self.drop_emb = nn.Dropout(
20             cfg["drop_rate"])
21
22         # Transformer blocks
23         self.trf_blocks = nn.Sequential(*[
24             TransformerBlock(cfg)
25             for _ in range(cfg["n_layers"])
26         ])
```

```
1 # Final layer norm and output
2 self.final_norm = LayerNorm(cfg["emb_dim"])
3 self.out_head = nn.Linear(
4     cfg["emb_dim"], cfg["vocab_size"],
5     bias=False)
6
7 def forward(self, in_idx):
8     batch_size, seq_len = in_idx.shape
9     # Implementation details coming up...
```

Configuration Benefits:

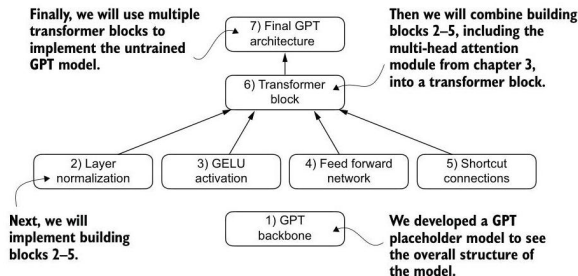
Easy experimentation with different model sizes.

Key Parameters:

- 50,257 vocabulary tokens
- 1024 maximum sequence length
- 768-dimensional embeddings
- 12 attention heads, 12 layers
- Configurable dropout rate

Data Flow Through GPT Model

Token Processing Pipeline



Key Transformations:

- **Tokenization:** Text → integer token IDs
- **Embedding:** Sparse IDs → dense vectors
- **Processing:** Self-attention and feed-forward transformations
- **Output:** Probability distribution over vocabulary

Need for Layer Normalization: Training Stability

Deep Network Training Challenges

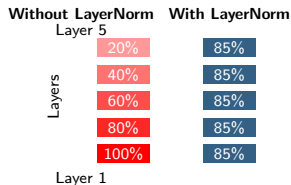
Problems in Deep Networks:

- **Vanishing Gradients:** Gradients become small
- **Exploding Gradients:** Gradients become large
- **Internal Covariate Shift:** Input distributions change
- **Slow Convergence:** Training takes very long

Layer Normalization Benefits:

- ✓ **Stable Training:** Consistent gradient flow
- ✓ **Faster Convergence:** Reduced training time
- ✓ **Less Sensitive:** To init and learning rate
- ✓ **Better Generalization:** Improved test performance

Gradient Flow Comparison



Core Insight

Layer norm ensures layers receive consistent input statistics (mean ≈ 0 , std ≈ 1).

Layer Normalization Mathematics

Statistical Normalization Formula

Layer Normalization Formula:

$$\text{LayerNorm}(x) = \gamma \odot \frac{x - \mu}{\sigma + \epsilon} + \beta$$

Where:

- $\mu = \frac{1}{D} \sum_{i=1}^D x_i$ (mean across features)
- $\sigma = \sqrt{\frac{1}{D} \sum_{i=1}^D (x_i - \mu)^2}$ (standard deviation)
- γ = learnable scale parameter (init to 1)
- β = learnable shift parameter (init to 0)
- ϵ = small constant for stability (10^{-5})
- $\odot$ = element-wise multiplication

Normalization Process:

- 1 **Center:** Subtract mean $\rightarrow$ zero-centered
- 2 **Scale:** Divide by std $\rightarrow$ unit variance
- 3 **Transform:** Apply learnable $\gamma, \beta \rightarrow$ optimal scale/shift

Key Difference from Batch Norm

Layer norm: statistics **across features** per sample. Batch norm: **across samples** per feature.

Layer Normalization Implementation

```
1 class LayerNorm(nn.Module):
2     def __init__(self, emb_dim):
3         super().__init__()
4         self.eps = 1e-5
5         self.scale = nn.Parameter(
6             torch.ones(emb_dim)) # gamma
7         self.shift = nn.Parameter(
8             torch.zeros(emb_dim)) # beta
9
10    def forward(self, x):
11        # x: (batch_size, seq_len, emb_dim)
12
13        # Compute mean and variance across
14        # last dimension (features)
15        mean = x.mean(dim=-1, keepdim=True)
16        # Shape: (batch_size, seq_len, 1)
17        var = x.var(dim=-1, keepdim=True,
18                    unbiased=False)
19
20        # Normalize
21        norm_x = (x - mean) / torch.sqrt(
22            var + self.eps)
23
24        # Scale and shift
25        return self.scale * norm_x + self.shift
```

Step-by-Step Implementation

```
1 # Test implementation
2 torch.manual_seed(123)
3 layer_norm = LayerNorm(768)
4 x = torch.randn(2, 10, 768)
5 # (batch, seq_len, embedding_dim)
6
7 normalized = layer_norm(x)
8 print(f"Input mean: {x.mean():.4f}, std: {x.std():.4f}")
9 print(f"Output mean: {normalized.mean():.4f}, std:
10       {normalized.std():.4f}")
11
12 # Check per-sample normalization
13 print(f"First sample mean:
14       {normalized[0].mean(dim=-1)[:3]}")
15 # Should be ~0
16 print(f"First sample std:
17       {normalized[0].std(dim=-1)[:3]}")
18 # Should be ~1
```

Key Points:

- Normalization across features (dim=-1)
- Learnable scale and shift parameters
- unbiased=False for GPT-2 compatibility

Layer Normalization Class with GPT-2 Compatibility

```
1 class LayerNorm(nn.Module):
2     def __init__(self, emb_dim):
3         super().__init__()
4         self.eps = 1e-5
5         self.scale = nn.Parameter(
6             torch.ones(emb_dim))
7         self.shift = nn.Parameter(
8             torch.zeros(emb_dim))
9
10    def forward(self, x):
11        mean = x.mean(dim=-1, keepdim=True)
12        var = x.var(dim=-1, keepdim=True,
13                    unbiased=False)
14        norm_x = (x - mean) / torch.sqrt(
15            var + self.eps)
16        return self.scale * norm_x + self.shift
```

PyTorch Built-in Alternative:

```
1 # Equivalent to our implementation
2 builtin_norm = nn.LayerNorm(emb_dim, eps=1e-5)
```

Production Implementation Details

Implementation Details:

- **unbiased=False**: Matches GPT-2 implementation exactly
- **keepdim=True**: Preserves dimensions for broadcasting
- **eps=1e-5**: Prevents division by zero, standard value
- **scale/shift**: Learnable parameters, identity init

Why from scratch?

Understanding math and ensuring GPT-2 compatibility.

Key Benefits:

- Exact reproducibility with published models
- Control over numerical precision
- Educational understanding

Biased vs. Unbiased Variance in Layer Normalization

Technical Detail: Variance Calculation

Two Variance Formulations:

Unbiased Variance (default in PyTorch):

$$\sigma_{\text{unbiased}}^2 = \frac{1}{N-1} \sum_{i=1}^N (x_i - \mu)^2$$

Biased Variance (used in GPT-2):

$$\sigma_{\text{biased}}^2 = \frac{1}{N} \sum_{i=1}^N (x_i - \mu)^2$$

Why GPT-2 Uses Biased Variance:

- **Consistency:** Matches other deep learning frameworks (TensorFlow)
- **Numerical Stability:** Slightly more stable for small feature dimensions
- **Historical Reasons:** Original GPT implementation choice
- **Reproducibility:** Ensures exact match with published results

Practical Impact:

- Very small difference in practice (especially with large embedding dimensions)
- Important for exact reproducibility of published models
- Use `unbiased=False` in PyTorch to match GPT-2

Implementation Note

Always use `unbiased=False` when implementing GPT-2 to ensure exact compatibility with pre-trained weights.

◆ Hands-on Exercise: Layer Normalization

```
1 import torch
2 import torch.nn as nn
3
4 # Complete the LayerNorm implementation
5 class CustomLayerNorm(nn.Module):
6     def __init__(self, emb_dim, eps=1e-5):
7         super().__init__()
8         self.eps = eps
9         # Create learnable parameters
10         self.scale = nn.Parameter(
11             torch.ones(emb_dim))
12         self.shift = nn.Parameter(
13             torch.zeros(emb_dim))
14
15     def forward(self, x):
16         # Compute mean across features
17         mean = x.mean(dim=-1, keepdim=True)
18
19         # Compute variance with unbiased=False
20         var = x.var(dim=-1, keepdim=True,
21                     unbiased=False)
22
23         # Normalize and apply scale/shift
24         norm_x = (x - mean) / torch.sqrt(
25             var + self.eps)
26         return self.scale * norm_x + self.shift
```

Interactive Implementation Activity (5 minutes)

```
1 # Test your implementation
2 torch.manual_seed(123)
3 custom_norm = CustomLayerNorm(4)
4 pytorch_norm = nn.LayerNorm(4, eps=1e-5)
5
6 x = torch.randn(2, 3, 4)
7 # (batch, seq, features)
8 custom_output = custom_norm(x)
9 pytorch_output = pytorch_norm(x)
10 print("Custom implementation mean/std:")
11 print(f"Mean: {custom_output[0].mean(dim=-1)}")
12 print(f"Std: {custom_output[0].std(dim=-1,
13     unbiased=False)}")
14 print("\nOutputs should be similar:")
15 print(f"Max difference: {torch.max(torch.abs(
16     custom_output - pytorch_output)))")
```

Exercise Goals:

- Implement LayerNorm from scratch
- Compare with PyTorch built-in
- Verify numerical correctness
- Understand parameter roles

GELU Activation Function: Smooth Alternative to ReLU

Gaussian Error Linear Units

GELU Mathematical Definition:

$$\text{GELU}(x) = x \cdot \Phi(x)$$

where $\Phi(x)$ is the CDF of standard normal distribution.

Intuitive Interpretation:

- Stochastic regularizer during training
- Smooth activation (differentiable everywhere)
- Combines sigmoid gating with input scaling
- Better gradient flow than ReLU

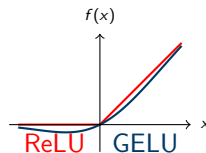
Approximate GELU (computationally efficient):

$$\text{GELU}(x) \approx 0.5x \left(1 + \tanh \left(\sqrt{\frac{2}{\pi}} (x + 0.044715x^3) \right) \right)$$

Advantages over ReLU:

- ✓ No "dead neurons" problem
- ✓ Smooth gradients everywhere
- ✓ Better performance in transformers
- ✓ Probabilistic interpretation

Activation Comparison



GELU Implementation: Exact vs. Approximate

```
1 import torch
2 import torch.nn as nn
3 import torch.nn.functional as F
4 import math
5
6 class GELU(nn.Module):
7     def __init__(self, approximate="none"):
8         super().__init__()
9         self.approximate = approximate
10
11     def forward(self, x):
12         if self.approximate == "none":
13             # Exact GELU using error function
14             return 0.5 * x * (1.0 + torch.erf(
15                 x / math.sqrt(2.0)))
16
17         elif self.approximate == "tanh":
18             # Tanh approximation
19             # (faster, used in original GPT-2)
20             return 0.5 * x * (
21                 1.0 + torch.tanh(
22                     math.sqrt(2.0 / math.pi) *
23                     (x + 0.044715 * torch.pow(x, 3.0))))
24         else:
25             raise ValueError(f"approximate must be 'none'
26                               or 'tanh'")
```

Different Implementation Strategies

```
1 # PyTorch built-in GELU (recommended)
2 builtin_gelu = nn.GELU()
3 # Or functional version
4 def gelu_functional(x):
5     return F.gelu(x)
6 # Performance comparison
7 x = torch.randn(1000, 768)
8 exact_gelu = GELU("none")
9 approx_gelu = GELU("tanh")
10 # Test outputs (should be very similar)
11 exact_out = exact_gelu(x)
12 approx_out = approx_gelu(x)
13 builtin_out = builtin_gelu(x)
14 print(f"Max difference (exact vs approx):
15       {torch.max(torch.abs(exact_out - approx_out)):.6f}")
16 print(f"Max difference (exact vs builtin):
17       {torch.max(torch.abs(exact_out - builtin_out)):.6f}")
```

Implementation Choices:

- **Exact:** Uses error function, slower
- **Tanh:** Approximation, faster
- **Built-in:** PyTorch optimized

Feed-Forward Networks: Position-wise Processing

Feed-Forward Network Structure:

$$\text{FFN}(x) = \text{GELU}(xW_1 + b_1)W_2 + b_2$$

```
1 class FeedForward(nn.Module):
2     def __init__(self, cfg):
3         super().__init__()
4         self.layers = nn.Sequential(
5             nn.Linear(cfg["emb_dim"],
6                       4 * cfg["emb_dim"]),
7             GELU(), nn.Linear(4 * cfg["emb_dim"],
8                               cfg["emb_dim"]),)
9
10    def forward(self, x):
11        return self.layers(x)
```

Neural Network Layer in Transformers

Key Properties:

- **Position-wise:** Applied to each position independently
- **4x Expansion:** Hidden dimension is 4× embedding dimension
- **GELU Activation:** Smooth, non-linear transformation
- **Parameter Heavy:** Most parameters in transformer block

Architecture Details:

- Input: $(B, T, 768)$
- Hidden: $(B, T, 3072)$ (4× expansion)
- Output: $(B, T, 768)$
- Parameters: $\sim 4.7\text{M}$ per layer

Feed-Forward 4× Expansion: Design Rationale

Why 4× Expansion in GPT Architectures

Mathematical Foundation:

- For embedding dimension d_{model} , feed-forward hidden dimension is $d_{ff} = 4 \times d_{model}$
- GPT-2 Small: $d_{model} = 768$, $d_{ff} = 3072$

Design Rationale:

- **Representational Capacity:** Larger hidden space allows complex transformations
- **Information Bottleneck:** Expansion $\rightarrow$ compression forces learning useful features
- **Non-linearity:** GELU activation creates rich feature interactions
- **Empirical Optimization:** 4× found optimal through experimentation

Parameter Analysis for GPT-2 Small:

Layer	Shape	Parameters
Linear 1	(768, 3072)	2,359,296
Linear 2	(3072, 768)	2,359,296
Per FF Block		4,718,592
12 FF Blocks		56,623,104

Alternative Ratios Explored:

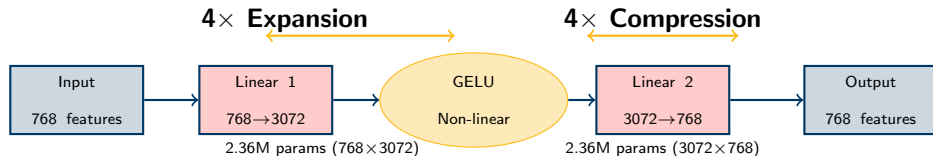
- 2× expansion: Too constraining, poor performance
- 8× expansion: Marginal gains, 2× parameter cost
- 4× expansion: Sweet spot for parameter efficiency vs. performance

Design Insight

Feed-forward networks contain 45.6% of total model parameters. The 4× expansion provides optimal balance between model capacity and computational efficiency.

Feed-Forward $4\times$ Expansion: Visual Architecture

Visual Feed-Forward Expansion Process



Each position processed independently Rich feature interactions in high-dimensional space Compressed learned representations



Processing Flow:

- Position-wise transformation
- No inter-position communication
- Parallel computation across sequence

Computational Cost:

- $2\times$ matrix multiplications per block
- 4.7M parameters per FF block
- Dominates model parameter count

Parameter Distribution Analysis: Feed-Forward Dominance

Understanding Parameter Allocation in GPT Models

GPT-2 Small Parameter Breakdown:

Component	% of Total
Token Embeddings	31.1%
Output Layer	31.1%
Feed-Forward	45.6%
Attention	22.8%
Position Embeddings	0.6%
Layer Norms	0.01%

Per-Layer Parameter Count:

Component	Parameters	% of Block
Multi-Head Attention	2,359,296	33.3%
Feed-Forward Network	4,718,592	66.7%
Layer Normalization (2x)	1,536	0.02%
Total per Block	7,079,424	100%

Implications:

- **Memory Bottleneck:** Feed-forward layers determine GPU memory requirements
- **Computational Cost:** Matrix multiplication dominates training time
- **Optimization Focus:** Most gradient updates occur in FF parameters
- **Scaling Decisions:** Model size primarily determined by FF expansion

When scaling from GPT-2 Small to Large:

- Embedding dimension: $768 \rightarrow 1280$ ($1.67\times$ increase)
- FF parameters: $56.6M \rightarrow 157.3M$ ($2.78\times$ increase)
- FF dominates parameter growth

Position-wise Processing: Independent Feature Transformation

Position-wise Independence:

$$\text{FFN}(\mathbf{X})_{i,:} = \text{FFN}(\mathbf{X}_{i,:})$$

```
1 def demonstrate_position_independence():
2     """Show that FFN processes each
3     position independently"""
4     batch_size, seq_len, emb_dim = 2, 5, 768
5     # Create feed-forward network
6     ff = FeedForward({"emb_dim": emb_dim})
7     # Input sequence
8     x = torch.randn(batch_size, seq_len, emb_dim)
9     # Method 1: Process entire sequence
10    output_batch = ff(x)
11    # Method 2: Process each position separately
12    output_separate = torch.zeros_like(x)
13    for i in range(seq_len):
14        output_separate[:, i, :] = ff(
15            x[:, i:i+1, :]).squeeze(1)
16    # Should be identical
17    max_diff = torch.max(torch.abs(output_batch -
18    output_separate))
19    print(f"Maximum difference: {max_diff.item():.10f}")
```

```
1     print("Position-wise independence verified!"
2           if max_diff < 1e-6 else "Error!")
3
4 demonstrate_position_independence()
```

Computational Advantages:

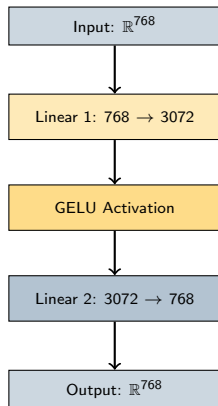
- **Parallelization:** All positions processed simultaneously
- **Memory Locality:** Efficient matrix operations
- **Hardware Utilization:** GPU-friendly computation pattern

Theoretical Implications:

- No cross-positional information mixing in FF layers
- Attention handles sequence relationships, FF handles feature transformation
- Enables efficient batched computation across all positions

Expansion/Contraction Visualization: Information Processing

Feed-Forward Information Flow



Information Processing Pattern:

- **Expansion:** Projects input to higher-dimensional space
- **Non-linearity:** GELU enables complex feature interactions
- **Contraction:** Compresses back to original dimensionality
- **Integration:** Residual connection preserves information flow

Think of feed-forward as a "feature factory":

- Expand raw features into intermediate representations
- Apply complex transformations in high-dimensional space
- Compress useful features back to standard size
- Add processed features to original input

Mathematical Interpretation:

- Expansion: $\mathbf{h}_1 = \text{GELU}(\mathbf{X}\mathbf{W}_1 + \mathbf{b}_1)$
- Contraction: $\mathbf{h}_2 = \mathbf{h}_1\mathbf{W}_2 + \mathbf{b}_2$
- Information bottleneck forces learning of compressed

Memory Usage and Computational Impact Analysis

Resource Requirements of Feed-Forward Networks

Memory Analysis (GPT-2 Small, Batch Size = 8,
Sequence Length = 1024):

Component	Memory (MB)	% of Total
Parameters		
FF Weights (all layers)	217.4	43.7%
Attention Weights	108.7	21.9%
Embeddings	149.6	30.1%
Activations (Forward Pass)		
FF Intermediate (3072 dim)	96.0	19.3%
FF Output (768 dim)	24.0	4.8%
Attention Activations	48.0	9.7%
Training Overhead		
FF Gradients	217.4	43.7%
FF Optimizer States (Adam)	434.8	87.4%

Computational Cost (FLOPs per token):

- **FF Layer 1:** $768 \times 3072 = 2.36\text{M}$ multiply-adds
- **GELU:** ≈ 3072 operations (negligible)
- **FF Layer 2:** $3072 \times 768 = 2.36\text{M}$ multiply-adds
- **Total per FF:** $\approx 4.7\text{M}$ FLOPs
- **All 12 layers:** $\approx 56\text{M}$ FLOPs per token

Optimization Strategies:

- **Mixed Precision:** Use FP16 to halve memory requirements
- **Activation Checkpointing:** Trade compute for memory
- **Gradient Accumulation:** Simulate larger batches with less memory
- **Expert Layers:** Sparse activation patterns (MoE models)

Memory Bottleneck

Feed-forward layers are the primary memory and computational bottleneck in transformer training.

Complete FeedForward Class Implementation

```
1 class FeedForward(nn.Module):
2     def __init__(self, cfg):
3         super().__init__()
4         # Expansion factor (typically 4 for GPT-2)
5         self.expansion_factor = 4
6         hidden_dim = cfg["emb_dim"] *
self.expansion_factor
7         # Two linear transformations with GELU
8         self.w1 = nn.Linear(cfg["emb_dim"],
                               hidden_dim, bias=True)
9         self.w2 = nn.Linear(hidden_dim,
                               cfg["emb_dim"], bias=True)
10
11         # Activation function
12         # (GELU for GPT-2 compatibility)
13         self.activation = GELU()
14         # Optional dropout for regularization
15         self.dropout = nn.Dropout(
16             cfg.get("ff_dropout", 0.0))
17
18         # Initialize weights
19         self._init_weights()
20     def _init_weights(self):
21         # Xavier initialization
22         nn.init.xavier_uniform_(self.w1.weight)
23         nn.init.xavier_uniform_(self.w2.weight)
24         # Small bias initialization
25         nn.init.constant_(self.w1.bias, 0.0)
26         nn.init.constant_(self.w2.bias, 0.0)
```

```
1     def forward(self, x):
2         ""
3         Args:
4             x: Input tensor (batch_size, seq_len, emb_dim)
5         Returns:
6             Output tensor (batch_size, seq_len, emb_dim)
7         ""
8         # Expand to higher dimension
9         h = self.w1(x)
10        # (batch_size, seq_len, 4*emb_dim)
11        # Apply non-linear activation
12        h = self.activation(h)
13        # Optional dropout
14        h = self.dropout(h)
15        # Contract back to original dimension
16        output = self.w2(h)
17        # (batch_size, seq_len, emb_dim)
18        return output
```

Best Practices:

- Xavier weight initialization
- Configurable dropout
- Clear documentation
- Proper bias handling

TransformerBlock: Combining Components

```
1 class TransformerBlock(nn.Module):
2     def __init__(self, cfg):
3         super().__init__()
4         self.att = MultiHeadAttention(
5             d_in=cfg["emb_dim"],
6             d_out=cfg["emb_dim"],
7             context_length=cfg["context_length"],
8             num_heads=cfg["n_heads"],
9             dropout=cfg["drop_rate"],
10            qkv_bias=cfg["qkv_bias"]
11        )
12        self.ff = FeedForward(cfg)
13        self.norm1 = LayerNorm(cfg["emb_dim"])
14        self.norm2 = LayerNorm(cfg["emb_dim"])
15        self.drop_shortcut = nn.Dropout(
16            cfg["drop_rate"])
17
18    def forward(self, x):
19        # Attention block with shortcut connection
20        shortcut = x
21        x = self.norm1(x)
22        x = self.att(x)
23        x = self.drop_shortcut(x)
24        x = x + shortcut # Shortcut connection
```

```
1     # Feed-forward block with shortcut connection
2     shortcut = x
3     x = self.norm2(x)
4     x = self.ff(x)
5     x = self.drop_shortcut(x)
6     x = x + shortcut # Shortcut connection
7
8     return x
```

Key Components:

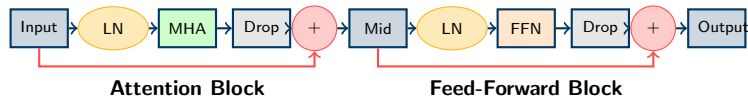
- **Multi-Head Attention:** Global context
- **Feed-Forward:** Position-wise transformation
- **Layer Normalization:** Stable training
- **Residual Connections:** Gradient flow
- **Dropout:** Regularization

Processing Order:

- 1 LayerNorm → Attention → Dropout → Add
- 2 LayerNorm → FeedForward → Dropout → Add

Complete Transformer Block Data Flow Architecture

Visual Architecture Overview



Component Legend:

- LN: LayerNorm
- MHA: Multi-Head Attention
- FFN: Feed-Forward Network
- Drop: Dropout
- +: Residual Add
- Red: Skip connections

Parameters (GPT-2 Small):

- Attention: 2.36M per block
- Feed-Forward: 4.72M per block
- Total per block: 7.1M

Memory Usage:

- Input/Output: $B \times T \times 768$
- FF Hidden: $B \times T \times 3072$ (peak)
- Attention: $B \times 12 \times T \times T$

Complete GPT Model Implementation

```
1 class GPTModel(nn.Module):
2     def __init__(self, cfg):
3         super().__init__()
4         self.tok_emb = nn.Embedding(
5             cfg["vocab_size"], cfg["emb_dim"])
6         self.pos_emb = nn.Embedding(
7             cfg["context_length"], cfg["emb_dim"])
8         self.drop_emb = nn.Dropout(
9             cfg["drop_rate"])
10
11         self.trf_blocks = nn.Sequential(*[
12             TransformerBlock(cfg)
13             for _ in range(cfg["n_layers"])
14         ])
15
16         self.final_norm = LayerNorm(
17             cfg["emb_dim"])
18         self.out_head = nn.Linear(
19             cfg["emb_dim"], cfg["vocab_size"],
20             bias=False)
```

```
1 def forward(self, in_idx):
2     batch_size, seq_len = in_idx.shape
3     tok_embeddings = self.tok_emb(in_idx)
4     pos_embeddings = self.pos_emb(
5         torch.arange(seq_len,
6                       device=in_idx.device))
7     x = tok_embeddings + pos_embeddings
8     x = self.drop_emb(x)
9     x = self.trf_blocks(x)
10    x = self.final_norm(x)
11    logits = self.out_head(x)
12    return logits
```

Architecture Summary:

- Token + Position embeddings
- 12 Transformer blocks
- Final LayerNorm
- Output projection to vocabulary
- Total: 124M parameters

GPT Model: Token and Position Embedding Integration

```
1 class GPTModel(nn.Module):
2     def __init__(self, cfg):
3         super().__init__()
4         # Token embeddings: map token IDs to vectors
5         self.tok_emb = nn.Embedding(
6             cfg["vocab_size"], cfg["emb_dim"])
7         # Position embeddings: add positional info
8         self.pos_emb = nn.Embedding(
9             cfg["context_length"], cfg["emb_dim"])
10        # Dropout for regularization
11        self.drop_emb = nn.Dropout(cfg["drop_rate"])
12        # Transformer blocks (defined separately)
13        self.trf_blocks = nn.Sequential(*[
14            TransformerBlock(cfg)
15            for _ in range(cfg["n_layers"])]])
16        # Final layer norm and output projection
17        self.final_norm = LayerNorm(cfg["emb_dim"])
18        self.out_head = nn.Linear(cfg["emb_dim"],
19            cfg["vocab_size"], bias=False)
```

```
1 def forward(self, in_idx):
2     batch_size, seq_len = in_idx.shape
3     # Get token embeddings
4     tok_embeds = self.tok_emb(in_idx)
5     # (batch, seq_len, emb_dim)
6     # Get position embeddings
7     pos_indices = torch.arange(
8         seq_len, device=in_idx.device)
9     pos_embeds = self.pos_emb(pos_indices)
10    # (seq_len, emb_dim)
11    # Combine embeddings
12    x = tok_embeds + pos_embeds
13    # Apply dropout
14    # Process through transformer blocks
15    x = self.trf_blocks(x)
16    # Final normalization and projection
17    x = self.final_norm(x)
18    logits = self.out_head(x)
19    return logits
```

Key Implementation Details:

- **Position Broadcasting:** Position embeddings broadcast across batch
- **Embedding Addition:** Token + position embeddings combined element-wise
- **Dropout Regularization:** Applied to combined embeddings

Parameter Counting and Memory Analysis

Detailed GPT-2 Small Parameter Breakdown

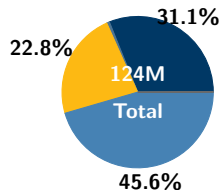
Component	Computation	Parameters	% of Total
<i>Embedding Layers</i>			
Token Embedding	$50,257 \times 768$	38,597,376	31.1%
Position Embedding	$1,024 \times 768$	786,432	0.6%
<i>Transformer Blocks (12 layers)</i>			
Multi-Head Attention	$768 \times 768 \times 4 \times 12$	28,311,552	22.8%
Feed-Forward Networks	$(768 \times 3072 + 3072 \times 768) \times 12$	56,623,104	45.6%
Layer Normalizations	$768 \times 2 \times 12$	18,432	0.01%
<i>Output Layer</i>			
Final Layer Norm	768	768	0.001%
Output Projection	$768 \times 50,257$	38,597,376	31.1%
Total Parameters		124,439,808	100%

Memory Requirements (Float32):

- **Parameters:** $124M \times 4 \text{ bytes} = 498 \text{ MB}$
- **Gradients:** Additional 498 MB during training
- **Optimizer States:** $2 \times$ gradients for Adam = 996 MB
- **Activations:** Varies by batch size and sequence length
- **Total Training:** 2-4 GB GPU memory minimum

GPT Parameter Distribution Visualization

GPT-2 Small (124M) Parameter Breakdown



Component Breakdown:

- **Token Embeddings:** 38.6M (31.1%)
- **Position Embeddings:** 0.8M (0.6%)
- **Multi-Head Attention:** 28.3M (22.8%)
- **Feed-Forward Networks:** 56.6M (45.6%)
- **Layer Normalizations:** 0.02M (0.01%)

Key Insights:

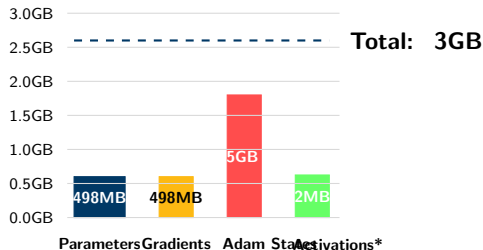
- Feed-forward layers dominate parameter count
- Attention mechanisms use 22.8% of parameters
- Embeddings require significant memory
- Layer norms contribute minimally
- Weight tying can eliminate output layer

Engineering Insight

FFN layers scale quadratically with hidden dimension ($d_{model} \times 4d_{model}$), making them the primary bottleneck in larger models.

Memory Usage Breakdown for GPT Training

Detailed Memory Analysis for GPT-2 124M



Memory Formulas:

$$\text{Parameters} = N_p \times 4 \text{ bytes} \quad (6)$$

$$\text{Gradients} = N_p \times 4 \text{ bytes} \quad (7)$$

$$\text{Adam States} = N_p \times 12 \text{ bytes} \quad (8)$$

$$\text{Activations} = f(B, L, d) \quad (9)$$

Where: N_p = parameter count, B = batch size,
 L = sequence length, d = hidden dimension

Scaling Insights:

- Adam optimizer dominates memory usage (3× parameters)
- Activations scale with batch size and sequence length
- Mixed precision (FP16) can halve memory requirements
- Gradient checkpointing trades compute for memory

Weight Tying: Sharing Embeddings and Output Parameters

Concept: Share weights between token embedding and output projection layers.

```
1 class GPTModel(nn.Module):
2     def __init__(self, cfg):
3         super().__init__()
4         self.tok_emb = nn.Embedding(
5             cfg["vocab_size"], cfg["emb_dim"])
6         self.pos_emb = nn.Embedding(
7             cfg["context_length"], cfg["emb_dim"])
8         self.drop_emb = nn.Dropout(cfg["drop_rate"])
9         self.trf_blocks =
10             nn.Sequential(*[TransformerBlock(cfg)
11                             for _ in range(cfg["n_layers"])])
12         self.final_norm = LayerNorm(cfg["emb_dim"])
13
14         # No separate output layer - use embedding weights
15         # self.out_head = nn.Linear(
16         #     cfg["emb_dim"], cfg["vocab_size"],
17         #     bias=False)
```

```
1 def forward(self, in_idx):
2     batch_size, seq_len = in_idx.shape
3     # Standard forward pass
4     tok_embeddings = self.tok_emb(in_idx)
5     pos_indices = torch.arange(seq_len, device=in_idx.device)
6     pos_embeddings = self.pos_emb(pos_indices)
7     x = tok_embeddings + pos_embeddings
8     x = self.drop_emb(x)
9     x = self.trf_blocks(x)
10    x = self.final_norm(x)
11    # Weight tying: use embedding weights transposed
12    logits = F.linear(x, self.tok_emb.weight)
13    # Shape: (batch, seq_len, vocab_size)
14    return logits
```

Benefits of Weight Tying:

- **Parameter Reduction:** Saves 38.6M parameters (31.1% reduction)
- **Improved Generalization:** Forces consistent representations
- **Memory Efficiency:** Reduces GPU memory requirements

Mathematical Justification

If embedding maps token i to vector $\mathbf{e}_i$, then logits for predicting token j are $\mathbf{h} \cdot \mathbf{e}_j$ - same vector space.

Model Configuration Scaling: GPT-2 Family

Systematic Architecture Scaling

Configuration	Small	Medium	Large	XL
Parameters	124M	345M	762M	1.5B
Layers (n_layers)	12	24	36	48
Attention Heads	12	16	20	25
Embedding Dim	768	1024	1280	1600
Feed-Forward Dim	3072	4096	5120	6400
Context Length	1024	1024	1024	1024

Scaling Relationships:

- **FF Dimension:** Always $4\times$ embedding dimension
- **Head Dimension:** $d_{head} = d_{emb}/n_{heads}$ (typically 64-80)
- **Parameter Growth:** Mostly from feed-forward expansion
- **Constant Context:** All models use 1024 token limit

Performance Implications:

- Larger models show better language understanding
- Exponentially more compute required for training
- Memory requirements scale with model size
- Inference time increases with depth and width

Scaling Law

Model performance improves as a power law with respect to model size, dataset size, and compute budget.

Model Configuration Management in Code

```
1 # Configuration dictionary for easy model switching
2 GPT_CONFIGS = {
3     "gpt2-small": {
4         "vocab_size": 50257, "context_length": 1024,
5         "emb_dim": 768, "n_heads": 12,
6         "n_layers": 12, "drop_rate": 0.1,
7         "qkv_bias": False},
8
9     "gpt2-medium": {
10        "vocab_size": 50257, "context_length": 1024,
11        "emb_dim": 1024, "n_heads": 16,
12        "n_layers": 24, "drop_rate": 0.1,
13        "qkv_bias": False},
14
15    "gpt2-large": {
16        "vocab_size": 50257, "context_length": 1024,
17        "emb_dim": 1280, "n_heads": 20,
18        "n_layers": 36, "drop_rate": 0.1,
19        "qkv_bias": False}}
20
21 # Initialize any model size
22 config = GPT_CONFIGS.get("gpt2-medium")
23 model = GPTModel(config)
24 # Count parameters
25 total_params = sum(p.numel() for p in model.parameters())
26 print(f"Total parameters: {total_params:,}")
```

Benefits of Configuration-Based Design:

- Easy experimentation with model sizes
- Consistent architecture across scales
- Simplified hyperparameter management
- Reproducible model instantiation

Complete GPT Model Code Implementation

```
1 class GPTModel(nn.Module):
2     def __init__(self, cfg):
3         super().__init__()
4         self.tok_emb = nn.Embedding(
5             cfg["vocab_size"], cfg["emb_dim"])
6         self.pos_emb = nn.Embedding(
7             cfg["context_length"], cfg["emb_dim"])
8         self.drop_emb = nn.Dropout(cfg["drop_rate"])
9
10        # Stack of transformer blocks
11        self.trf_blocks = nn.Sequential(*[
12            TransformerBlock(cfg)
13            for _ in range(cfg["n_layers"])]])
14
15        # Final processing
16        self.final_norm = LayerNorm(cfg["emb_dim"])
17        self.out_head = nn.Linear(
18            cfg["emb_dim"], cfg["vocab_size"],
19            bias=False)
```

```
1 def forward(self, in_idx):
2     batch_size, seq_len = in_idx.shape
3
4     # Embedding layer
5     tok_embeds = self.tok_emb(in_idx)
6     pos_indices = torch.arange(
7         seq_len, device=in_idx.device)
8     pos_embeds = self.pos_emb(pos_indices)
9     x = tok_embeds + pos_embeds
10    x = self.drop_emb(x)
11
12    # Transformer processing
13    x = self.trf_blocks(x)
14
15    # Output projection
16    x = self.final_norm(x)
17    logits = self.out_head(x)
18
19    return logits
```

Key Features:

- **Modular Design:** Easy to modify and extend
- **Device Agnostic:** Handles GPU/CPU automatically
- **Configurable:** All hyperparameters externalized
- **Standard PyTorch:** Compatible with all PyTorch tools

Model Testing and Shape Verification

```
1 def test_gpt_model():
2     torch.manual_seed(123)
3     cfg = GPT_CONFIG_124M
4     model = GPTModel(cfg)
5     # Test 1: Shape verification
6     batch_size, seq_len = 2, 10
7     input_ids = torch.randint(0, cfg["vocab_size"],
8                               (batch_size, seq_len))
9
10    with torch.no_grad():
11        logits = model(input_ids)
12
13    print("Shape Tests:")
14    print(f"Input: {input_ids.shape}")
15    print(f"Output: {logits.shape}")
16    print(f"Expected: ({batch_size}, {seq_len}, {cfg['vocab_size']})")
17    assert logits.shape == (batch_size, seq_len,
18                           cfg["vocab_size"])
```

```
1 # Test 2: Parameter count
2 total_params = sum(p.numel()
3                   for p in model.parameters())
4 print(f"\nParameter Count: {total_params:,}")
5 print(f"Expected: ~124M parameters")
6
7 # Test 3: Gradient flow (requires_grad check)
8 trainable_params = sum(p.numel() for p in
9                       model.parameters() if p.requires_grad)
10 print(f"Trainable parameters: {trainable_params:,}")
11
12 # Test 4: Memory usage
13 model_memory = sum(p.numel() * 4
14                   for p in model.parameters()) / 1024**2 # MB
15 print(f"Model memory: {model_memory:.1f} MB")
16 print("\nAll tests passed! Model ready for training.")
17
18 # Run comprehensive tests
19 test_gpt_model()
```

Validation Checklist:

- ✓ Input/output shape consistency
- ✓ Parameter count matches specification
- ✓ All parameters require gradients
- ✓ Memory usage within expected range
- ✓ No runtime errors or warnings

◆ Hands-on Exercise: Build Complete GPT Model

```
1 # Complete the missing parts of the GPT model
2 class GPTModel(nn.Module):
3     def __init__(self, cfg):
4         super().__init__()
5         # TODO: Initialize token embedding layer
6         self.tok_emb = nn.Embedding(
7             cfg["vocab_size"], cfg["emb_dim"])
8
9         # TODO: Initialize position embedding layer
10        self.pos_emb = nn.Embedding(
11            cfg["context_length"], cfg["emb_dim"])
12
13        # TODO: Add embedding dropout
14        self.drop_emb = nn.Dropout(cfg["drop_rate"])
15
16        # TODO: Create transformer blocks stack
17        self.trf_blocks = nn.Sequential(*[
18            TransformerBlock(cfg)
19            for _ in range(cfg["n_layers"])]])
20
21        # TODO: Add final layer norm and output head
22        self.final_norm = LayerNorm(cfg["emb_dim"])
23        self.out_head = nn.Linear(
24            cfg["emb_dim"], cfg["vocab_size"],
25            bias=False)
```

```
1 def forward(self, in_idx):
2     batch_size, seq_len = in_idx.shape
3     # TODO: Get token embeddings
4     tok_embs = self.tok_emb(in_idx)
5
6     # TODO: Get position embeddings
7     pos_indices = torch.arange(seq_len,
8                                device=in_idx.device)
9     pos_embs = self.pos_emb(pos_indices)
10
11    # TODO: Combine embeddings and apply dropout
12    x = tok_embs + pos_embs
13    x = self.drop_emb(x)
14
15    # TODO: Process through transformer blocks
16    x = self.trf_blocks(x)
17
18    # TODO: Apply final norm and output projection
19    x = self.final_norm(x)
20    logits = self.out_head(x)
21
22    return logits
23
24 # Test your implementation
25 model = GPTModel(GPT_CONFIG_124M)
26 test_input = torch.randint(0, 1000, (2, 5))
27 output = model(test_input)
28 print(f"Success! Output shape: {output.shape}")
```

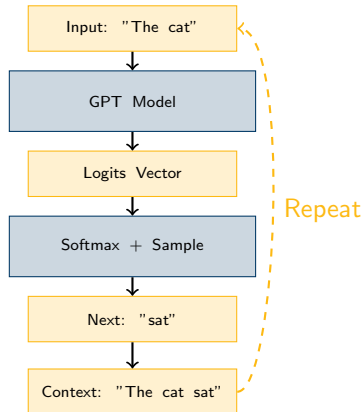
Section 4

Text Generation Detailed Process

Autoregressive Generation Mechanics

Autoregressive Text Generation Overview

Generation Pipeline:



Mathematical Formulation:

$$P(\text{seq}) = \prod_{t=1}^T P(x_t | x_{1:t-1})$$

Generation Steps:

- 1 **Context Input:** Feed current sequence
- 2 **Forward Pass:** Compute logits for next token
- 3 **Apply Temperature:** Scale logits by τ
- 4 **Softmax:** Convert to probabilities
- 5 **Sample:** Choose next token
- 6 **Append:** Add to sequence
- 7 **Repeat:** Until max length or EOS

Sampling Strategies:

- **Greedy:** $\arg \max$ (deterministic)
- **Temperature:** $p_i \propto e^{z_i/\tau}$
- **Top-k:** Sample from top k tokens
- **Top-p:** Nucleus sampling

Context Window Management and Cropping

Context Cropping Strategy:

```
1 def generate_with_cropping(model, idx, max_new_tokens,
2   context_size):
3     """Generate text with context window management"""
4     model.eval()
5
6     for _ in range(max_new_tokens):
7         # Crop context if it exceeds maximum length
8         idx_cropped = idx[:, -context_size:] # Keep last
9         context_size tokens
10
11         with torch.no_grad():
12             # Forward pass with cropped context
13             logits = model(idx_cropped)
14
15             # Get logits for the last position only
16             logits = logits[:, -1, :] # (batch_size,
17             vocab_size)
18
19             # Apply softmax and sample
20             probs = torch.softmax(logits, dim=-1)
21             idx_next = torch.multinomial(probs, num_samples=1)
22
23             # Append new token (full sequence continues to
24             grow)
25             idx = torch.cat((idx, idx_next), dim=1)
26
27     return idx
```

Handling Long Sequences The Challenge: GPT

models have fixed maximum context length (1024 tokens for GPT-2).

Cropping Visualization:

Full Sequence	Cropped Context (Last 1024)
[token1, token2, ..., token1500]	[..., token477, ..., token1500]

Trade-offs:

- ✓ **Memory Efficient:** Constant memory usage
- ✓ **Consistent Performance:** No slowdown with length
- ✗ **Limited Memory:** Loses old context information

Step-by-Step Generation Process

```
1 def generate_text_detailed(model, tokenizer, prompt,
2   max_new_tokens=50):
3     model.eval()
4     # Step 1: Tokenize initial prompt
5     input_ids = tokenizer.encode(prompt)
6     input_tensor = torch.tensor([input_ids])
7
8     generated_tokens = input_tensor.clone()
9     for step in range(max_new_tokens):
10         with torch.no_grad():
11             logits = model(generated_tokens)
12             next_token_logits = logits[0, -1, :]
13             probabilities = torch.softmax(next_token_logits,
14             dim=-1)
15             next_token_id = torch.multinomial(probabilities,
16             num_samples=1).item()
17             next_token_tensor = torch.tensor(
18             [[next_token_id]])
19             generated_tokens = torch.cat([generated_tokens,
20             next_token_tensor], dim=1)
21             next_token_text = tokenizer.decode([next_token_id])
22             print(f"Step {step+1}: '{next_token_text}'")
23             if next_token_id == tokenizer.eos_token_id:
24                 print("Reached EOS token.")
25                 break
26     # Final decoding
27     final_text = tokenizer.decode(generated_tokens[0].tolist())
28     return final_text
```

Code Explanation:

Initialization (Lines 1-11):

- Set model to evaluation mode
- Tokenize input prompt
- Create tensor with batch dimension
- Initialize generation buffer

Generation Loop (Lines 13-40):

- **Forward Pass:** Compute logits for all positions
- **Extract Last:** Get predictions for next token
- **Softmax:** Convert logits to probabilities
- **Sampling:** Use multinomial to select token
- **Concatenate:** Add new token to sequence
- **Decode:** Convert token ID to text
- **Stop Check:** Exit if EOS encountered

Output Processing (Lines 42-44):

- Decode entire token sequence
- Return complete generated text

Generation Complexity

Time: $O(n \times L)$ where n = new tokens, L = sequence length

Sampling Strategies Beyond Greedy Decoding

```
1 def sample_with_strategies(logits,
2                             strategy="greedy",
3                             temperature=1.0,
4                             top_k=50):
5
6     """Various sampling strategies for text generation"""
7     if strategy == "greedy":
8         # Always pick highest probability token
9         next_token = torch.argmax(logits, dim=-1)
10
11     elif strategy == "random":
12         # Sample from full probability distribution
13         probs = torch.softmax(
14             logits / temperature, dim=-1)
15         next_token = torch.multinomial(
16             probs, num_samples=1).squeeze()
17
18     elif strategy == "top_k":
19         # Only sample from top-k tokens
20         top_k_logits, top_k_indices =
21             torch.topk(logits, top_k)
22         top_k_probs = torch.softmax(
23             top_k_logits / temperature, dim=-1)
24         sampled_idx = torch.multinomial(
25             top_k_probs, num_samples=1)
26         next_token = top_k_indices[sampled_idx]
```

```
1     elif strategy == "temperature":
2         # Use temperature scaling for
3         # controlled randomness
4         scaled_logits = logits / temperature
5         probs = torch.softmax(scaled_logits, dim=-1)
6         next_token = torch.multinomial(
7             probs, num_samples=1).squeeze()
8
9     return next_token
```

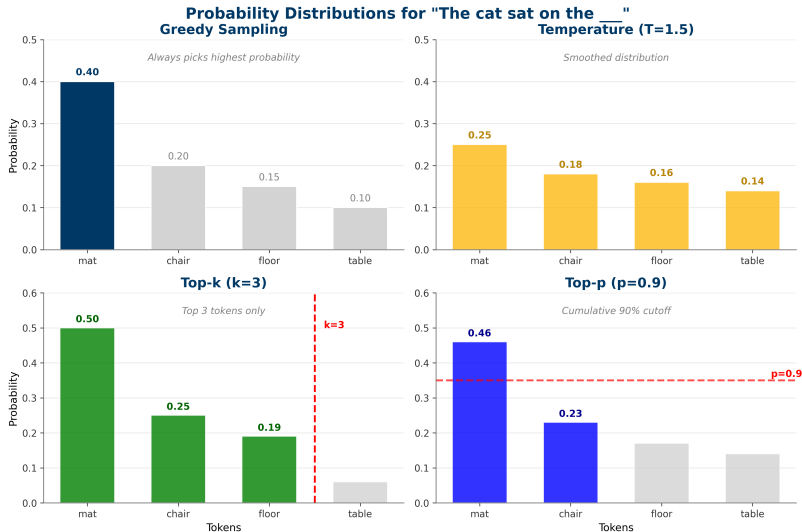
Temperature Effects:

- $T = 0.1$: Very focused, deterministic-like
- $T = 1.0$: Standard probability distribution
- $T = 2.0$: More random, creative output

Strategy Comparison:

- **Greedy**: Deterministic, coherent but repetitive
- **Random**: Creative but potentially incoherent
- **Top-k**: Balanced diversity and quality
- **Temperature**: Fine-grained randomness control

Sampling Strategies: Probability Distributions



Sampling Strategies: Formulations and Trade-offs

Mathematical Formulations:

Greedy: $w_t = \arg \max_i P(w_i | w_{<t})$

Always selects most probable token

Temperature: $P'(w_i) = \frac{\exp(z_i/T)}{\sum_j \exp(z_j/T)}$

$T < 1$: deterministic, $T > 1$: random

Top-k: Sample from k highest probabilities

Fixed vocabulary diversity

Top-p: Sample from tokens with $\sum P(w_i) \leq p$

Adaptive vocabulary based on confidence

Trade-offs and Characteristics:

- **Greedy:** ✓ Deterministic, coherent ✗ Repetitive
- **Temperature:** ✓ Smooth randomness control ✗ May pick unlikely tokens
- **Top-k:** ✓ Fixed vocabulary diversity ✗ Ignores confidence
- **Top-p:** ✓ Adaptive vocabulary ✓ Confidence-based

Typical Parameter Values:

- Temperature: $T = 0.7-1.2$
- Top-k: $k = 40-80$
- Top-p: $p = 0.9-0.95$

Production Insight

Modern systems often combine strategies: temperature scaling with top-p filtering for optimal quality-diversity balance.

Complete Text Generation Implementation (Part 1)

```
1 def generate_text_simple(model, idx, max_new_tokens,
2                           context_size, temperature=1.0,
3                           top_k=None, eos_token_id=None):
4     """
5     Generate text using the GPT model with various
6     sampling strategies.
7     Args:
8         model: Trained GPT model
9         idx: Starting token indices
10             (batch_size, seq_len)
11         max_new_tokens: Maximum number of tokens to
12             generate
13         context_size: Maximum context window
14         temperature: Sampling temperature
15         top_k: Sample from top-k tokens only
16         eos_token_id: End-of-sequence token ID
17     Returns:
18         Generated token sequence
19     """
20     model.eval()
21     for _ in range(max_new_tokens):
22         # Crop context if exceeds window
23         if idx.shape[1] > context_size:
24             idx_cond = idx[:, -context_size:]
25         else:
26             idx_cond = idx
```

Function Parameters Explained:

Essential Parameters:

- **model**: The trained GPT model in eval mode
- **idx**: Initial token tensor (batch $\times$ seq)
- **max_new_tokens**: Generation limit
- **context_size**: Model's max context (1024 for GPT-2)

Sampling Controls:

- **temperature**: Controls randomness
 - < 1.0 : More focused/deterministic
 - $= 1.0$: Standard sampling
 - > 1.0 : More random/creative
- **top_k**: Limits vocabulary for sampling
- **eos_token_id**: Early stopping trigger

Design Pattern

Function provides flexible generation with multiple sampling strategies in a single interface.

Complete Text Generation Implementation (Part 2)

```
1  with torch.no_grad():
2      logits = model(idx_cond)
3      # Extract last token's predictions
4      logits = logits[:, -1, :] / temperature
5      # Apply top-k filtering if specified
6      if top_k is not None:
7          # Get top-k logits and indices
8          top_k_logits, top_k_indices = \
9              torch.topk(logits, top_k, dim=-1)
10         # Create filtered logits tensor
11         logits_filtered = torch.full_like(
12             logits, float('-inf'))
13         # Scatter top-k values back
14         logits_filtered.scatter_(1, top_k_indices,
15             top_k_logits)
16         logits = logits_filtered
17         # Convert to probabilities
18         probs = torch.softmax(logits, dim=-1)
19         # Sample next token
20         idx_next = torch.multinomial(
21             probs, num_samples=1)
22         # Append to sequence
23         idx = torch.cat((idx, idx_next), dim=1)
24         # Check for early stopping
25         if eos_token_id and (idx_next ==
26             eos_token_id).any():
27             break
28     return idx
```

Generation Process Breakdown:

1. Forward Pass (Lines 1-6):

- Disable gradients with `torch.no_grad()`
- Get logits for all positions
- Extract last position predictions
- Apply temperature scaling

2. Top-k Filtering (Lines 8-20):

- Select k highest probability tokens
- Set all other logits to $-\infty$
- Ensures sampling from limited vocabulary
- Improves generation quality

3. Sampling (Lines 22-27):

- Apply softmax for probabilities
- Use multinomial sampling
- Concatenate new token to sequence

4. Termination (Lines 29-33):

- Check for EOS token
- Early exit if found

Performance Considerations and Optimization

Efficient Text Generation

Generation Performance Bottlenecks:

- **Sequential Nature:** Cannot parallelize generation steps
- **Memory Growth:** Context grows with each token
- **Repeated Computation:** Recalculating attention for same tokens
- **Model Size:** Large models slow on single GPU

Optimization Strategies:

Technique	Benefit	Trade-off
KV-Cache	10-50× faster	Higher memory usage
Batch Generation	2-4× throughput	Padding overhead
Model Quantization	2-4× memory reduction	Slight quality loss
Speculative Decoding	2-3× faster	Complex implementation
Context Cropping	Constant memory	Limited context

KV-Cache Optimization (Advanced):

- **Benefit:** Avoid recomputing attention for existing tokens
- **Implementation:** Modify attention to accept cached KV
- **Memory Cost:** $O(\text{layers} \times \text{seq_len} \times d_{\text{model}})$

Typical Generation Speeds:

- GPT-2 Small (124M): 50-100 tokens/sec (GPU)
- GPT-3 (175B): 1-5 tokens/sec (GPU)
- With optimization: 2-10× improvement possible

◆ Hands-on Exercise: Implement Text Generation

```
1 def generate_text_exercise(model, idx, max_new_tokens,
2                             context_size):
3
4     for step in range(max_new_tokens):
5         # TODO: Apply context cropping
6         idx_cond = idx[:, -context_size:] \
7             if idx.shape[1] > context_size else idx
8
9         with torch.no_grad():
10             # TODO: Forward pass to get logits
11             logits = model(idx_cond)
12             # TODO: Get logits for last position
13             logits = logits[:, -1, :]
14             # TODO: Apply temperature scaling
15             # (use temperature=0.8)
16             temperature = 0.8
17             logits = logits / temperature
18             # TODO: Convert to probabilities
19             probs = torch.softmax(logits, dim=-1)
20             # TODO: Sample next token
21             idx_next = torch.multinomial(
22                 probs, num_samples=1)
23             # TODO: Append to sequence
24             idx = torch.cat((idx, idx_next), dim=1)
25             print(f"Step {step+1}: Generated token \
26                 {idx_next.item()}")
27
28     return idx
```

Interactive Generation Implementation (8 minutes)

```
1 # Test your implementation
2 torch.manual_seed(42)
3 model = GPTModel(GPT_CONFIG_124M) # Untrained model
4 start_tokens = torch.randint(1000, 5000, (1, 3))
5 # Random starting tokens
6 generated = generate_text_exercise(
7     model, start_tokens, max_new_tokens=5,
8     context_size=1024)
9 print(f"Generated sequence: \
10       {generated[0].tolist()}")
11 print("Note: Untrained model produces \
12       random tokens!")
```

Extension Challenges:

- 1 Implement top-k sampling
- 2 Add early stopping for special tokens
- 3 Compare different temperature values
- 4 Measure generation speed

Section 5

Transformer Architecture Assembly

Connecting Components into Complete Blocks

Shortcut Connections: The Key to Deep Training

ResNet-Style Skip Connections Mathematical Formulation:

$$\mathbf{h}_{l+1} = \mathbf{h}_l + F(\mathbf{h}_l, \mathbf{W}_l)$$

Where:

- $\mathbf{h}_l$ = input to layer l
- $F(\mathbf{h}_l, \mathbf{W}_l)$ = transformation function
- Addition creates direct path for gradients

In Transformer Blocks:

- **Attention Shortcut:** $x + \text{Attention}(x)$
- **Feed-Forward Shortcut:** $x + \text{FeedForward}(x)$
- **Gradient Flow:** Direct path to early layers
- **Identity Mapping:** $F(x) = 0 \Rightarrow$ identity function

Critical Insight

Without shortcut connections, training 12+ layer transformers would be nearly impossible due to vanishing gradients.

Mathematical Proof of Gradient Vanishing

Chain Rule in Deep Networks:

For a network with L layers, the gradient at layer l is:

$$\frac{\partial \mathcal{L}}{\partial \mathbf{w}^{(l)}} = \frac{\partial \mathcal{L}}{\partial \mathbf{z}^{(L)}} \prod_{i=l+1}^L \frac{\partial \mathbf{z}^{(i)}}{\partial \mathbf{z}^{(i-1)}}$$

Problem Analysis: If each layer transformation has Jacobian with spectral norm $\rho < 1$:

$$\left\| \prod_{i=l+1}^L \frac{\partial \mathbf{z}^{(i)}}{\partial \mathbf{z}^{(i-1)}} \right\| \leq \rho^{L-l}$$

Gradient Decay Example:

- Layer 12: $\rho^0 = 1.0$ (full gradient)
- Layer 6: $\rho^6 \approx 0.5^6 = 0.016$ (98.4% reduction)
- Layer 1: $\rho^{11} \approx 0.5^{11} = 0.0005$ (99.95% reduction)

Shortcut Connection Solution:

With residual connection $\mathbf{h}_{l+1} = \mathbf{h}_l + F(\mathbf{h}_l)$:

$$\frac{\partial \mathbf{h}_{l+1}}{\partial \mathbf{h}_l} = \mathbf{I} + \frac{\partial F(\mathbf{h}_l)}{\partial \mathbf{h}_l}$$

The identity matrix $\mathbf{I}$ ensures gradient flow even if $\frac{\partial F}{\partial \mathbf{h}_l} \rightarrow 0$.

Gradient Flow Analysis: Empirical Demonstration

```
1 def analyze_gradient_flow(model, loss):
2     """Analyze gradient magnitudes across layers"""
3     # Compute gradients
4     loss.backward()
5     gradient_norms = {}
6     for name, param in model.named_parameters():
7         if param.grad is not None:
8             grad_norm = param.grad.norm().item()
9             gradient_norms[name] = grad_norm
10    return gradient_norms
11    # Simulate training step
12    model_without_shortcuts =
13        DeepNetworkNoShortcuts(layers=12)
14    model_with_shortcuts = GPTModel(GPT_CONFIG_124M)
15    # Forward pass and loss
16    x = torch.randn(2, 10, 768)
17    target = torch.randint(0, 50257, (2, 10))
18
19    loss_no_shortcuts = F.cross_entropy(
20        model_without_shortcuts(x).view(-1, 768),
21        target.view(-1))
22    loss_with_shortcuts = F.cross_entropy(
23        model_with_shortcuts(
24            torch.randint(0, 50257, (2, 10))
25            ).view(-1, 50257),
26        target.view(-1))
27    )
```

```
1 # Analyze gradients
2 grads_no_shortcuts = analyze_gradient_flow(
3     model_without_shortcuts, loss_no_shortcuts)
4 grads_with_shortcuts = analyze_gradient_flow(
5     model_with_shortcuts, loss_with_shortcuts)
6
7 # Compare gradient magnitudes by layer depth
8 print("Layer-wise Gradient Norms:")
9 for layer in range(12):
10    print(f"Layer {layer+1}:
11          No shortcuts={
12            grads_no_shortcuts[f'layer{layer}']:.6f},
13            With shortcuts={
14
15            grads_with_shortcuts[f'transformer{layer}']:.6f}")
```

Expected Results:

- **Without shortcuts:** Exponential decay from layer 12 $\rightarrow$ 1
- **With shortcuts:** Relatively stable gradients across all layers

Training Stability: Quantitative Analysis

Empirical Benefits of Shortcut Connections

Training Metrics Comparison:

Metric	Without Shortcuts	With Shortcuts
Convergence Time	5-10x slower	Baseline
Final Loss	Higher (worse)	Lower (better)
Gradient Variance	High instability	Stable
Learning Rate Sensitivity	Very sensitive	Robust
Parameter Utilization	30-50% layers	90-95% layers
Training Success Rate	20-40% runs	95-98% runs

ResNet Historical Impact:

- **Pre-ResNet (2015):** Networks deeper than 8-10 layers very difficult to train
- **Post-ResNet:** Successfully trained 50, 101, 152+ layer networks
- **Transformer Adoption:** All transformer architectures use residual connections
- **Modern Scale:** Enables training of 12B+ parameter models

Mathematical Properties:

- **Identity Mapping:** If $F(\mathbf{x}) = 0$, residual block becomes identity
- **Gradient Preservation:** $\nabla_{\mathbf{x}}(\mathbf{x} + F(\mathbf{x})) = \mathbf{I} + \nabla_{\mathbf{x}}F(\mathbf{x})$
- **Error Surface:** Smoother optimization landscape
- **Depth Scalability:** Linear scaling with network depth possible

Training Insight

Residual connections transform very difficult optimization problem into tractable one.

◆ Hands-on Exercise: Compare Gradient Flows

```
1 def compare_gradient_flows():
2     """Compare gradient flows with and without
3     residual connections"""
4     torch.manual_seed(123)
5
6     # Define simple networks for comparison
7     class NetworkNoSkip(nn.Module):
8         def __init__(self):
9             super().__init__()
10            self.layers = nn.Sequential(*[
11                nn.Linear(768, 768)
12                for _ in range(6)])
13            def forward(self, x):
14                return self.layers(x)
15            class NetworkWithSkip(nn.Module):
16                def __init__(self):
17                    super().__init__()
18                    self.layers = nn.ModuleList([
19                        nn.Linear(768, 768)
20                        for _ in range(6)])
21                    def forward(self, x):
22                        for layer in self.layers:
23                            x = x + layer(x) # Residual
24                    return x
25            # TODO: Initialize both networks
26            net_no_skip = NetworkNoSkip()
27            net_with_skip = NetworkWithSkip()
```

```
1     # TODO: Create dummy input and target
2     x = torch.randn(4, 10, 768)
3     target = torch.randn(4, 10, 768)
4     # TODO: Forward pass and compute losses
5     out_no_skip = net_no_skip(x)
6     out_with_skip = net_with_skip(x)
7     loss_no_skip = F.mse_loss(out_no_skip, target)
8     loss_with_skip = F.mse_loss(out_with_skip, target)
9     # TODO: Compute gradients
10    loss_no_skip.backward()
11    loss_with_skip.backward()
12    # TODO: Compare gradient norms
13    print("Gradient norms by layer:")
14    print("No Skip | With Skip")
15    for i, (p1, p2) in enumerate(
16        zip(net_no_skip.parameters(),
17            net_with_skip.parameters())):
18        if p1.grad is not None and
19            p2.grad is not None:
20            norm1 = p1.grad.norm().item()
21            norm2 = p2.grad.norm().item()
22            print(f"Layer {i}: {norm1:.6f} |
23                {norm2:.6f}")
24    compare_gradient_flows()
```

Vanishing Gradient Problem in Deep Networks

Why We Need Shortcut Connections

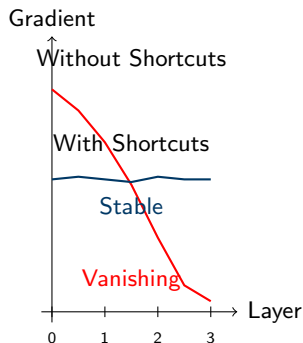
The Problem:

- Gradients become exponentially smaller in deeper layers
- Chain rule: $\frac{\partial L}{\partial x_1} = \frac{\partial L}{\partial x_n} \prod_{i=1}^{n-1} \frac{\partial x_{i+1}}{\partial x_i}$
- If derivatives $\neq 1$, product approaches zero
- Early layers learn slowly or not at all

Impact on Training:

- Slow convergence or no learning in early layers
- Network effectively becomes shallow
- Poor utilization of model capacity
- Difficulty training very deep networks

Gradient Magnitude



Mathematical Example: If each layer has gradient $\frac{\partial x_{i+1}}{\partial x_i} = 0.5$, after 10 layers:

Model Parameter Counting: Understanding Scale

GPT-2 Small Parameter Breakdown

Component-wise Parameter Count:

Component	Shape	Parameters
Token Embedding	(50257, 768)	38.6M
Position Embedding	(1024, 768)	0.8M
12 × Transformer Blocks	-	85.1M
Multi-Head Attention	(768, 768) × 4	2.3M each
Feed-Forward	(768, 3072, 768)	4.7M each
Layer Norms	(768) × 2	1.5K each
Final Layer Norm	(768)	768
Output Head	(768, 50257)	38.6M
Total		124.4M

Key Insights:

- **Embeddings:** 62% of parameters (input + output)
- **Feed-Forward:** Dominates transformer block parameters
- **Attention:** Relatively parameter-efficient

Model Initialization and Testing

```
1 import torch
2 torch.manual_seed(123)
3
4 # Initialize model with GPT-2 Small config
5 model = GPTModel(GPT_CONFIG_124M)
6
7 # Test forward pass
8 batch_size, seq_len = 2, 4
9 token_ids = torch.randint(
10     0, GPT_CONFIG_124M['vocab_size'],
11     (batch_size, seq_len))
12
13 print("Input shape:", token_ids.shape)
14 print("Input tokens:", token_ids)
15
16 # Forward pass
17 with torch.no_grad():
18     logits = model(token_ids)
19
20 print("Output logits shape:",
21     logits.shape)
```

```
1 print("Expected shape:
2     (batch_size, seq_len, vocab_size)")
3 print(f"Actual: ({batch_size}, {seq_len},
4     {GPT_CONFIG_124M['vocab_size']})")
5
6 # Verify parameter count
7 total_params = sum(p.numel()
8     for p in model.parameters())
9 print(f"Total parameters: {total_params:,}")
10 print(f"Expected: ~124M parameters")
```

Validation Checks:

- Input: (batch, sequence) of token IDs
- Output: (batch, sequence, vocab) logits
- Parameters: ~124 million
- Forward pass executes successfully

Note: Untrained model produces random outputs!

Text Generation: From Logits to Tokens

Generation Pipeline:

- ➊ **Forward Pass:** Model produces logits for next token
- ➋ **Probability Conversion:** Apply softmax to get probabilities
- ➌ **Sampling:** Select next token from probability distribution
- ➍ **Context Update:** Add selected token to sequence
- ➎ **Repeat:** Continue until stopping condition

Mathematical Formulation:

$$P(\text{token}_t | \text{context}_{1:t-1}) = \text{softmax}(\text{logits}_t)$$

Sampling Strategies:

- **Greedy:** Always pick highest probability token
- **Random:** Sample from full probability distribution
- **Top-k:** Sample from k highest probability tokens
- **Temperature:** Control randomness with scaling

◆ Interactive Exercise: Text Generation

```
1 def generate_text_simple(model, idx,
2                           max_new_tokens,
3                           context_size):
4     """
5     Generate text using the trained GPT model.
6
7     Args:
8         model: The GPT model
9         idx: Starting token indices
10             (batch_size, seq_len)
11         max_new_tokens: Number of tokens to generate
12         context_size: Maximum context length
13     """
14     model.eval() # Set to evaluation mode
15
16     for _ in range(max_new_tokens):
17         # TODO: Crop context if it gets too long
18         idx_cond = idx[:, -context_size:]
19         # Keep last context_size tokens
20
21         with torch.no_grad():
22             # TODO: Get model predictions
23             logits = model(idx_cond)
24
25             # TODO: Get logits for the last token
26             logits = logits[:, -1, :]
27             # (batch_size, vocab_size)
```

```
1         # TODO: Apply softmax and sample
2         probs = torch.softmax(logits, dim=-1)
3         idx_next = torch.multinomial(
4             probs, num_samples=1)
5         # (batch_size, 1)
6
7         # TODO: Append to sequence
8         idx = torch.cat((idx, idx_next), dim=1)
9
10        return idx
11
12    # Test the generation function
13    torch.manual_seed(42)
14    model = GPTModel(GPT_CONFIG_124M)
15    model.eval()
16
17    # Start with "Hello" token (example)
18    start_context = torch.tensor([[15496]])
19    # Example token for "Hello"
20    generated = generate_text_simple(
21        model, start_context, max_new_tokens=10,
22        context_size=GPT_CONFIG_124M["context_length"])
23
24    print(f"Generated token sequence:
25          {generated[0].tolist()}")
26    print("Note: Untrained model will produce
27          random/meaningless output!")
```

Architecture Summary: What We've Built

Complete GPT Architecture Components

Core Components:

- ✓ **Causal Attention:** Masked self-attention
- ✓ **Multi-Head Attention:** Parallel processing
- ✓ **Layer Normalization:** Training stability
- ✓ **GELU Activation:** Smooth non-linearity
- ✓ **Feed-Forward Networks:** Feature transformation
- ✓ **Shortcut Connections:** Gradient flow

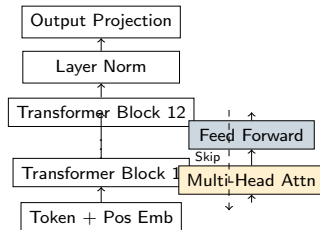
Integration Patterns:

- Transformer blocks with pre-normalization
- Residual connections around attention and FF
- Causal masking for autoregressive generation
- Parameter sharing and scaling strategies

Current Capabilities:

- ✓ Forward pass through complete architecture
- ✓ Text generation (token-by-token)
- ✓ Proper gradient flow for training
- ✗ **No learned representations** (untrained weights)

Architecture Flow:



Achievement

We've built a **complete, functional GPT architecture** ready for training!

Current Model: Capabilities and Limitations

What We Have vs. What We Need

✓ Current Capabilities:

- Complete GPT-2 architecture (124M params)
- Proper causal attention masking
- Forward pass computation
- Text generation mechanism
- Gradient computation ready
- All components integrated correctly

◆ Technical Verification:

- Input/output shape consistency
- Causal masking working correctly
- Parameter count matches GPT-2 small
- No computational errors

× Current Limitations:

- **Random weights:** No learned patterns
- **No language understanding:** Gibberish output
- **No training:** Weights are random
- **No optimization:** No loss minimization

■ Example Output:

```
1 Input: "Hello, I am"  
2 Untrained Output:  
3 "xJtok4mK9nQ8vR"  
4 (Complete nonsense!)
```

What's Missing:

- 1 **Training Data:** Large corpus of text for learning
- 2 **Training Loop:** Optimization procedure
- 3 **Loss Function:** Next-token prediction loss
- 4 **Computation:** GPU training for days/weeks

Next Step

Lecture 6: Pre-training pipeline to transform our architecture into a functioning language model!

Bridge to Lecture 6: From Architecture to Intelligence

The Journey Ahead

Today's Achievement:

Complete GPT Architecture
(Ready for Learning)

Lecture 6 Preview: Pre-training Large Language Models

- **Training Data:** Preparing massive text corpora
- **Loss Functions:** Cross-entropy for next-token prediction
- **Optimization:** Adam optimizer, learning rate scheduling
- **Training Loop:** Mini-batch processing and gradient updates
- **Evaluation:** Perplexity and generation quality metrics
- **Scaling:** Multi-GPU training and efficiency