

# Fine-tuning for Classification

## From General LLMs to Specialized Classifiers

Prof. Rongyu Lin

Lecture 7  
Quinnipiac University  
School of Computing and Engineering

**Reading:** Raschka Ch. 6 (pp. 191-238)  
"Transform pretrained models for specific tasks"

# Lecture Agenda

## Today's Journey: 75 minutes

- ① **Fine-tuning Categories & Dataset Prep** (20 min)
- ② **Data Loaders & Model Modification** (19 min)
- ③ **Training Loop & Loss Optimization** (18 min)
- ④ **Model Evaluation & Applications** (18 min)

**Interactive elements:** 4 hands-on activities

*From general language model to specialized spam classifier*

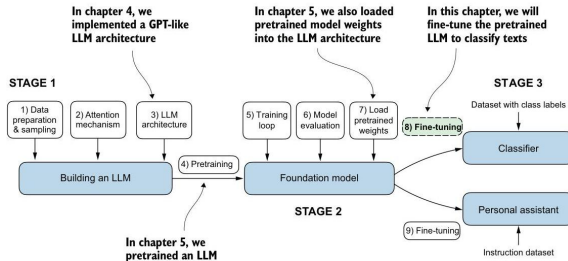
# Learning Objectives

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

- Distinguish instruction vs classification fine-tuning
- Prepare datasets for text classification tasks
- Modify pretrained architectures with classification heads
- Implement training loops for spam detection
- Evaluate models using accuracy and loss metrics
- Apply fine-tuned models to new data

*Master the art of adapting LLMs for specific tasks*

# Chapter 6 Overview - The Fine-tuning Pipeline



## The Transformation:

- **From:** General pretrained LLM (GPT-2)
- **To:** Specialized spam classifier
- **Key stages:** Dataset prep, model modification, training, evaluation
- **Focus:** Stage 3, Step 8 - Classification fine-tuning

*Three main LLM stages - we focus on classification fine-tuning*

# Two Categories of Fine-tuning

## Instruction Fine-tuning:

- Trains model to follow natural language instructions
- **Use cases:** Versatile tasks
- **Examples:** Translations, summaries
- **Trade-off:** Flexibility

## Classification Fine-tuning:

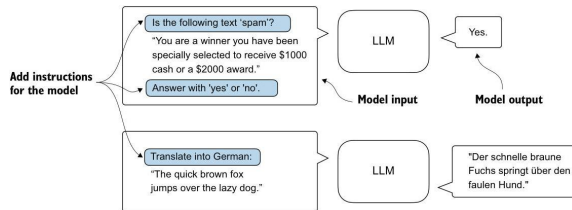
- Trains model to predict specific class labels
- **Use cases:** Specific domains
- **Examples:** Spam, sentiment
- **Trade-off:** Specialization

**Key Decision:** Flexibility vs Specialization

# Instruction Fine-tuning Characteristics

## Key Features:

- **Input:** Natural language instructions + context
- **Output:** Variable responses based on instruction type
- **Example:** "Determine if this text is spam: [text]"
- **Flexibility:** Can handle multiple task types
- **Requirements:** Larger datasets, more compute resources

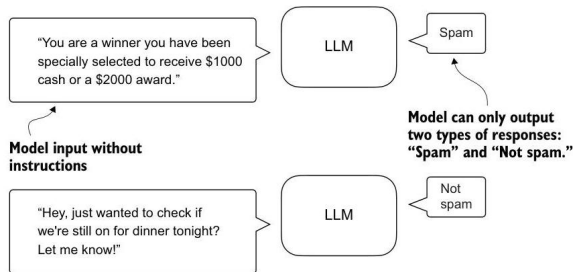


*Instruction fine-tuning scenarios with examples*

# Classification Fine-tuning Characteristics

## Key Features:

- **Input:** Raw text without explicit instructions
- **Output:** Predefined class labels (0, 1, 2, etc.)
- **Example:** Text → "spam" or "not spam"
- **Restriction:** Limited to trained classes only
- **Advantages:** Less data needed, faster training, higher accuracy for specific tasks



*Classification scenario for spam detection*

# Choosing the Right Approach

## Classification fine-tuning when:

- Specific categorization needed
- Limited data available
- Fast inference required

## Instruction fine-tuning when:

- Multiple tasks needed
- Flexible responses required
- Abundant resources available

**Our choice:** Classification fine-tuning for spam detection

**Reasoning:** Clear binary task, efficient training, measurable performance

# Dataset Introduction - SMS Spam Collection

## Dataset Overview:

- **Source:** UCI Machine Learning Repository
- **Task:** Binary classification (spam vs ham)
- **Size:** 5,572 text messages
- **Format:** Tab-separated values (Label, Text)
- **Real-world relevance:** Practical application everyone understands

```
[0]  
le(
```



Counts the instances  
of “spam”

Randomly samples “ham”  
instances to match the number  
of “spam” instances

Combines ham  
subset with “spam”

*Dataset preview in pandas DataFrame*

# Dataset Download and Setup

```
1 import urllib.request
2 import zipfile
3 import pandas as pd
4
5 def download_sms_spam_dataset():
6     url = "https://archive.ics.uci.edu/static/public/228/sms+
       spam+collection.zip"
7     zip_path = "sms_spam_collection.zip"
8     extracted_path = "sms_spam_collection"
9     data_file_path = "SMSSpamCollection"
10
11     # Download and extract
12     urllib.request.urlretrieve(url, zip_path)
13     with zipfile.ZipFile(zip_path, 'r') as zip_ref:
14         zip_ref.extractall(extracted_path)
```

## Download Process:

- **URL:** UCI ML Repository
- **Process:** Download → Extract → Rename → Load

## Features:

- Automatic extraction
- Path management
- Data verification

# Class Distribution Analysis

```
1 # Analyze class distribution
2 print(df["Label"].value_counts())
3 # Output:
4 # ham      4825
5 # spam      747
```

## The Imbalance Problem:

- **Ham:** 4,825 messages
- **Spam:** 747 messages
- **Ratio:** 6.5:1 imbalance

**Impact:** Model bias toward majority class

**Our approach:** Undersample ham to match spam count

# Creating Balanced Dataset

```
1 def create_balanced_dataset(df):
2     # Separate classes
3     spam_df = df[df["Label"] == "spam"]
4     ham_df = df[df["Label"] == "ham"]
5
6     # Undersample ham to match spam count
7     ham_subset = ham_df.sample(
8         n=len(spam_df),
9         random_state=123
10    )
11
12    # Combine
13    balanced_df = pd.concat([ham_subset, spam_df])
14    return balanced_df.sample(
15        frac=1,
16        random_state=123
17    ).reset_index(drop=True)
```

## Balancing Strategy:

- **Method:** Random undersampling
- **Target:** 747 examples each class
- **Random seed:** 123 for reproducibility

## Steps:

- Separate classes
- Sample ham to match spam
- Combine and shuffle

**Result:** 1,494 total examples (747 + 747)

# Label Encoding

```
1 # Convert string labels to integers
2 balanced_df["Label"] = balanced_df["Label"].map({
3     "ham": 0,
4     "spam": 1
5 })
6
7 # Verify conversion
8 print(balanced_df["Label"].value_counts())
9 # Output:
10 # 0    747 # ham
11 # 1    747 # spam
```

## Converting Text Labels:

- **Conversion:** String  $\rightarrow$  Integer
- **Mapping:** "ham"  $\rightarrow$  0, "spam"  $\rightarrow$  1
- **Reason:** Neural networks require numerical inputs

**Analogy:** Similar to text tokenization but simpler vocabulary

# Dataset Splitting Strategy

```
1 def random_split(df, train_frac, validation_frac):
2     # Shuffle dataset
3     df = df.sample(
4         frac=1,
5         random_state=123
6     ).reset_index(drop=True)
7
8     # Calculate split points
9     train_end = int(len(df) * train_frac)
10    validation_end = train_end + int(
11        len(df) * validation_frac
12    )
13
14    # Split
15    train_df = df[:train_end]
16    validation_df = df[train_end:validation_end]
17    test_df = df[validation_end:]
18
19    return train_df, validation_df, test_df
```

## Three-way Split:

- **Training:** 70% (1,046 examples)
- **Validation:** 10% (149 examples)
- **Test:** 20% (299 examples)

## Process:

- Shuffle data randomly
- Calculate split points
- Create separate datasets

**Key:** Stratification maintains class balance

## Explore the SMS Spam Dataset

**Task:** Load dataset and examine class distribution

**Goal:** Understand data characteristics and imbalance issues

### Steps:

① Download → Load → Analyze → Balance → Split

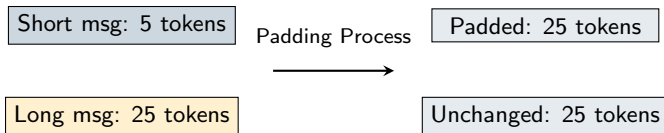
### Questions to explore:

- What makes a message "spam"?
- How severe is the imbalance?
- What are real-world implications of data imbalance?

# Data Preprocessing Challenges

## The Variable Length Problem:

- **Variable length:** Messages have different token counts
- **Batching requirement:** Neural networks need fixed-size inputs
- **Solutions:** Truncation vs Padding
- **Our choice:** Padding (preserves information)
- **Padding token:** `<|endoftext|>` (token ID 50256)



# From Dataset to Training

## Completed Steps:

- ✓ Data download and extraction
- ✓ Class balance correction
- ✓ Dataset splitting
- ✓ Label encoding

**File output:** train.csv, validation.csv, test.csv

**Next step:** Create PyTorch data loaders for batching

*Clean, balanced dataset ready for model training*

# PyTorch Data Loaders Overview

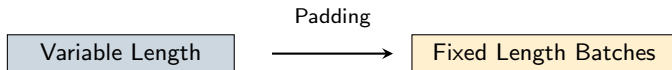
**Purpose:** Efficiently load and batch text data

**The Challenge:** Variable-length text messages

**Approaches:**

- **Previous approach:** Sliding window for uniform chunks
- **Current approach:** Padding to longest sequence length

**Benefits:** Preserve all information, enable batching



# Text Length Handling Options

## Option 1: Truncation

- Truncate all to shortest message length
- **Pro:** Fast processing
- **Con:** Information loss

## Option 2: Padding

- Pad all to longest message length
- **Pro:** Information preservation
- **Con:** More memory usage

**Our choice:** Padding (Option 2)

**Rationale:** Avoid information loss, better performance

# Padding Token Strategy

```
1 import tiktoken
2
3 tokenizer = tiktoken.get_encoding("gpt2")
4 token_id = tokenizer.encode(
5     "<|endoftext|>",
6     allowed_special={"<|endoftext|>"}
7 )[0]
8 print(f"Padding token ID: {token_id}")
9 # Output: Padding token ID: 50256
```

```
def random_split(df, train_frac, validation_frac):
    df = df.sample(
        frac=1, random_state=123
    ).reset_index(drop=True)
    train_end = int(len(df) * train_frac)
    validation_end = train_end + int(len(df) * validation_frac)

    train_df = df[:train_end]
    validation_df = df[train_end:validation_end]
    test_df = df[validation_end:]

    return train_df, validation_df, test_df

train_df, validation_df, test_df = random_split(
    balanced_df, 0.7, 0.1)
```

Shuffles the entire DataFrame

Calculates split indices

Splits the DataFrame

Test size is implied to be 0.2 as the remainder.

## Token Details:

- **Token choice:** <|endoftext|>
- **Token ID:** 50256 (verify with tiktoken)
- **Implementation:** Add token IDs, not text strings

## Process:

- Import tiktoken
- Get GPT-2 encoding
- Extract token ID
- Verify padding token

## *Input preparation with tokenization and padding*

# SpamDataset Class Design

```
1 class SpamDataset(Dataset):
2     def __init__(self, csv_file, tokenizer,
3                 max_length=None):
4         self.data = pd.read_csv(csv_file)
5         self.tokenizer = tokenizer
6
7         # Auto-detect max length if not specified
8         if max_length is None:
9             token_lengths = []
10            for text in self.data["Text"]:
11                token_lengths.append(
12                    len(tokenizer.encode(text))
13                )
14            self.max_length = max(token_lengths)
15        else:
16            self.max_length = max_length
```

## PyTorch Dataset Implementation:

- **Inheritance:** PyTorch Dataset class
- **Key functions:** `__init__`, `__getitem__`, `__len__`
- **Responsibilities:** Load data, tokenize, pad sequences

## Features:

- Auto-detect max length
- User-specified option
- Efficient data loading

# SpamDataset Implementation Details

```
1 def __getitem__(self, index):
2     text = self.data.iloc[index]["Text"]
3     label = self.data.iloc[index]["Label"]
4
5     # Tokenize text
6     token_ids = self.tokenizer.encode(text)
7
8     # Truncate if necessary
9     token_ids = token_ids[:self.max_length]
10
11    # Pad with endoftext token (ID: 50256)
12    token_ids += [50256] * (
13        self.max_length - len(token_ids)
14    )
15
16    return (
17        torch.tensor(token_ids, dtype=torch.long),
18        torch.tensor(label, dtype=torch.long)
19    )
20
21 def __len__(self):
22     return len(self.data)
```

## Core Methods:

- `__getitem__`: Process single data point
- `__len__`: Return dataset size

## Process:

- Tokenize text
- Truncate if needed
- Pad to max length
- Return tensors

**Output:** Token IDs and labels as PyTorch tensors

# Dataset Instantiation

```
1 import tiktoken
2
3 # Initialize tokenizer
4 tokenizer = tiktoken.get_encoding("gpt2")
5
6 # Create training dataset (auto-detect max length)
7 train_dataset = SpamDataset(
8     "train.csv",
9     tokenizer
10 )
11 print(f"Max length detected: {train_dataset.max_length}")
12 # Output: Max length detected: 120
13
14 # Create validation and test datasets
15 val_dataset = SpamDataset(
16     "validation.csv",
17     tokenizer,
18     max_length=train_dataset.max_length
19 )
20 test_dataset = SpamDataset(
21     "test.csv",
22     tokenizer,
23     max_length=train_dataset.max_length
24 )
```

## Creating Dataset Objects:

- Import tiktoken tokenizer
- Create training dataset first
- Auto-detect max length
- Use same length for val/test

## Consistency:

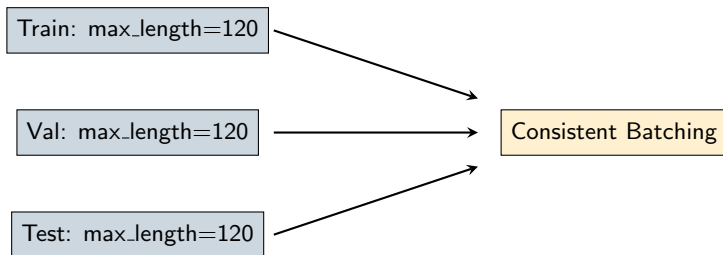
- Same tokenizer for all splits
- Same max length (120 tokens)
- Proper data alignment

**Key insight:** Model supports 1,024 tokens - messages well within limit

# Validation and Test Datasets

## Consistency Requirements:

- **Length matching:** Use training dataset's `max_length`
- **Consistency:** All datasets padded to same length (120)
- **Truncation:** Automatic for sequences  $> \text{max\_length}$
- **Alternative:** Set `max_length=None` if no long sequences

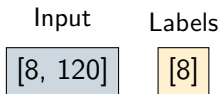
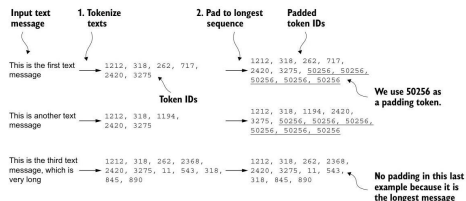


**Result:** All datasets compatible for training and evaluation

# Batch Structure Visualization

## Batch Configuration:

- **Batch size:** 8 training examples
- **Input tensor:**  $[8, 120]$  - 8 messages of 120 tokens
- **Label tensor:**  $[8]$  - corresponding class labels
- **Padding:** Token ID 50256 fills shorter sequences



*Batch structure with 8 messages and labels*

# Data Loader Configuration

```
1 from torch.utils.data import DataLoader
2
3 # Training data loader (with shuffling)
4 train_loader = DataLoader(
5     train_dataset,
6     batch_size=8,
7     shuffle=True,
8     num_workers=0,
9     drop_last=True
10 )
11
12 # Validation data loader (no shuffling)
13 val_loader = DataLoader(
14     val_dataset,
15     batch_size=8,
16     shuffle=False,
17     num_workers=0,
18     drop_last=False
19 )
```

## DataLoader Setup:

- **Batch size:** 8 examples per batch
- **Training:** Shuffle=True, drop\_last=True
- **Validation:** Shuffle=False, drop\_last=False

## Key Settings:

- Shuffling prevents overfitting
- Drop last ensures consistent batch sizes
- No multiprocessing (num\_workers=0)

# Data Loader Verification

```
1 # Check data loader shapes
2 print("Checking data loader shapes...")
3 for i, (inputs, targets) in enumerate(train_loader):
4     print(f"Batch {i+1}:")
5     print(f"  Input shape: {inputs.shape}")
6     # [8, 120]
7     print(f"  Target shape: {targets.shape}")
8     # [8]
9     if i == 2: # Show first 3 batches
10         break
11
12 # Count total batches
13 print(f"Training batches: {len(train_loader)}")
14 # 130 batches
15 print(f"Validation batches: {len(val_loader)}")
16 # 19 batches
17 print(f"Test batches: {len(test_loader)}")
18 # 38 batches
```

## Testing Our Setup:

- Check input/target shapes
- Verify batch dimensions
- Count total batches

## Expected Results:

- Input shape: [8, 120]
- Target shape: [8]
- Training: 130 batches
- Validation: 19 batches
- Test: 38 batches

**Verification:** Shapes correct, counts match expectations

## Create and Test Data Loaders

**Task:** Implement SpamDataset and create DataLoaders

**Goal:** Understand batching and padding mechanics

### Steps:

- 1 Create dataset → Set up loaders → Check batch shapes

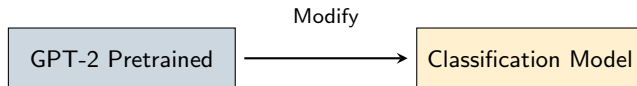
### Exploration points:

- Examine padded sequences
- Verify token IDs
- Understand: Why padding vs truncation?

# Model Architecture Preparation

## Starting Point:

- **Base model:** Pretrained GPT-2 model (124M parameters)
- **Configuration:** Same as pretraining setup
- **Model loading:** Import weights from OpenAI
- **Verification:** Generate coherent text to confirm loading
- **Next step:** Modify for classification



*Transform general language model into specialized classifier*

# GPT-2 Model Configuration

```
1 BASE_CONFIG = {  
2     "vocab_size": 50257,  
3     "context_length": 1024,  
4     "drop_rate": 0.0,  
5     "qkv_bias": True  
6 }  
7  
8 model_configs = {  
9     "gpt2-small (124M)": {  
10         "emb_dim": 768,  
11         "n_layers": 12,  
12         "n_heads": 12  
13     },  
14     # Other configurations...  
15 }
```

## Model Specifications:

- **Size:** gpt2-small (124M parameters)
- **Vocabulary:** 50,257 tokens
- **Context length:** 1,024 tokens
- **Architecture:** 12 layers, 768 embedding dimensions

## Configuration:

- Base config with GPT-2 specs
- Model-specific dimensions
- Exact match required

# Loading Pretrained Weights

```
1 def download_and_load_gpt2(model_size, models_dir):
2     # Download weights from OpenAI
3     model = GPTModel(BASE_CONFIG)
4
5     # Load pretrained weights
6     load_weights_into_gpt(
7         model,
8         model_size,
9         models_dir
10    )
11    model.eval()
12
13    # Test generation to verify loading
14    text = generate(
15        model=model,
16        tokenizer=tokenizer,
17        start_context="Every effort moves you",
18        max_new_tokens=25
19    )
20    print("Generated text:", text)
21    return model
```

## Weight Loading Process:

- Create model with base config
- Download OpenAI weights
- Load into model architecture
- Set to evaluation mode

## Verification:

- Test text generation
- Check coherent output
- Confirm successful loading

# Pre-fine-tuning Instruction Test

```
1 # Test if model can follow classification instructions
2 test_prompt = ""Is the following text 'spam'?
3 Answer with 'yes' or 'no':
4
5 'Congratulations! You have won a $1000 cash prize.
6 Call now to claim your reward!'
7
8 Answer:"""
9
10 response = generate(
11     model=model,
12     tokenizer=tokenizer,
13     start_context=test_prompt,
14     max_new_tokens=10
15 )
16 print("Model response:", response)
17 # Output: Incoherent, doesn't follow instructions
```

## Testing Instruction Following:

- Create classification prompt
- Test with clear spam example
- Generate model response
- Analyze instruction following

## Expected Result:

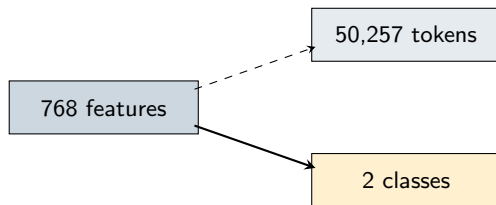
- Incoherent output
- Doesn't follow instructions
- Normal for base GPT-2

**Conclusion:** Pretrained model needs fine-tuning

# Classification Head Concept

## Architecture Modification:

- **Original:**  $768 \rightarrow 50,257$  (vocabulary size)
- **New:**  $768 \rightarrow 2$  (binary classification)
- **Rationale:** Match number of output classes
- **Generalization:**  $n \text{ classes} \rightarrow n \text{ output nodes}$



Replace with classification head

# Layer Freezing Strategy

```
1 # Freeze most layers
2 for param in model.parameters():
3     param.requires_grad = False
4
5 # Unfreeze last transformer block
6 for param in model.trf_blocks[-1].parameters():
7     param.requires_grad = True
8
9 # Unfreeze final layer norm
10 for param in model.final_norm.parameters():
11     param.requires_grad = True
```

## Selective Training Approach:

- **Frozen:** Embedding + first 11 blocks
- **Trainable:** Final block + LayerNorm + classification head

## Benefits:

- Faster training
- Prevent overfitting
- Preserve pretrained knowledge
- Fewer parameters to train

# Classification Head Implementation

```
1 # Replace output head with classification layer
2 model.out_head = torch.nn.Linear(
3     in_features=768, # GPT embedding dimension
4     out_features=2   # Binary classification
5 )
6
7 # Classification head is trainable by default
8 print("New layer requires_grad:",
9       model.out_head.weight.requires_grad)
10 # Output: True
11
12 # Verify architecture changes
13 print("Model output shape test:")
14 test_input = torch.randint(0, 50257, (1, 4))
15 output = model(test_input)
16 print(f"Output shape: {output.shape}")
17 # [1, 4, 2]
```

## Adding Classification Layer:

- Replace output head
- 768 features → 2 classes
- Trainable by default
- Test architecture changes

## Shape Changes:

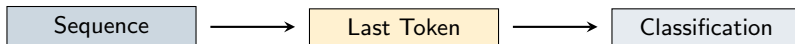
- Original: [batch, seq, 50257]
- New: [batch, seq, 2]
- Binary classification ready

# Modified Model Behavior

## Processing Pipeline:

- **Input processing:** Same as before (tokenization)
- **Output shape:** Changed from  $[1, 4, 50257]$  to  $[1, 4, 2]$
- **Key insight:** Focus on last output token only
- **Reasoning:** Causal attention mask makes last token most informed

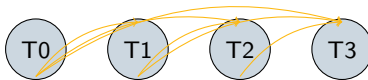
**Example:** "Do you have time" → final token represents full sequence



# Last Token Focus Strategy

## Attention Flow Logic:

- **Attention mechanism:** Each token attends to previous tokens
- **Information flow:** Final token has access to entire sequence
- **Classification decision:** Based on last token's representation
- **Code:** `outputs[:, -1, :]` extracts last token



Final token sees everything

*Visual representation of attention flow to final token*

# From Logits to Predictions

```
1 # Example logits for a sample
2 logits = torch.tensor([[ -3.5983,  3.9902]])
3 # Shape: [1, 2]
4
5 # Get class prediction
6 predicted_class = torch.argmax(logits, dim=-1)
7 print(f"Predicted class: {predicted_class.item()}")
8 # Output: 1 (spam)
9
10 # Get probability scores
11 probabilities = torch.softmax(logits, dim=-1)
12 print(f"Probabilities: {probabilities}")
13 # Output: [0.0003, 0.9997]
14 # 99.97% spam confidence
```

## Processing Output Logits:

- **Raw output:** Two logits per classification
- **Interpretation:** Higher value = predicted class
- **Conversion:** `torch.argmax` for prediction

## Example:

- Logits: `[-3.5983, 3.9902]`
- Predicted: class 1 (spam)
- Confidence: 99.97%

# Classification Accuracy Function

```
1 def calc_accuracy_loader(data_loader, model, device,
2   num_batches=None):
3     model.eval()
4     correct_predictions, num_examples = 0, 0
5
6     if num_batches is None:
7         num_batches = len(data_loader)
8     else:
9         num_batches = min(num_batches, len(data_loader))
10
11     for i, (input_batch, target_batch) in
12         enumerate(data_loader):
13         if i < num_batches:
14             input_batch = input_batch.to(device)
15             target_batch = target_batch.to(device)
16
17             with torch.no_grad():
18                 logits = model(input_batch)[: , -1, :]
19                 predicted_labels = torch.argmax(logits, dim=-1)
20                 num_examples += predicted_labels.shape[0]
21                 correct_predictions += (predicted_labels ==
22                 target_batch).sum().item()
23         else:
24             break
25     return correct_predictions / num_examples
```

## Measuring Performance:

- Set model to evaluation mode
- Process batches with no\_grad
- Extract last token logits
- Compare predictions to targets

## Formula:

- $\text{correct\_predictions} / \text{total\_examples}$
- Returns decimal (0.0 to 1.0)
- Multiply by 100 for percentage

# Initial Model Performance

```
1 # Test initial performance
2 train_accuracy = calc_accuracy_loader(
3     train_loader, model, device
4 )
5 val_accuracy = calc_accuracy_loader(
6     val_loader, model, device
7 )
8 test_accuracy = calc_accuracy_loader(
9     test_loader, model, device
10 )
11
12 print(f"Training accuracy: {train_accuracy*100:.2f}%")
13 # 46.25%
14 print(f"Validation accuracy: {val_accuracy*100:.2f}%")
15 # 45.00%
16 print(f"Test accuracy: {test_accuracy*100:.2f}%")
17 # 48.75%
```

## Before Fine-tuning Results:

- **Training:** 46.25%
- **Validation:** 45.00%
- **Test:** 48.75%

## Analysis:

- Near random performance
- Classification head has random weights
- Model needs training

**Random baseline:** 50% for binary classification

# Classification Loss Function

```
1 def calc_loss_batch(input_batch, target_batch,
2                     model, device):
3     input_batch = input_batch.to(device)
4     target_batch = target_batch.to(device)
5
6     # Forward pass
7     logits = model(input_batch)[: , -1, :]
8     # Focus on last token
9
10    # Calculate cross-entropy loss
11    loss = torch.nn.functional.cross_entropy(
12        logits, target_batch
13    )
14    return loss
15
16 # Example usage
17 sample_batch = next(iter(train_loader))
18 inputs, targets = sample_batch
19 loss = calc_loss_batch(inputs, targets,
20                        model, device)
21 print(f"Sample batch loss: {loss.item():.4f}")
```

## Cross-entropy Loss:

- Move inputs to device
- Forward pass through model
- Extract last token logits
- Calculate cross-entropy

## Key Points:

- Differentiable proxy for accuracy
- Focus on last token only
- Standard for classification

# Loss Calculation Implementation

```
1 def calc_loss_loader(data_loader, model, device, num_batches=None):
2     total_loss = 0.
3     if num_batches is None:
4         num_batches = len(data_loader)
5     else:
6         num_batches = min(num_batches, len(data_loader))
7
8     for i, (input_batch, target_batch) in
9         enumerate(data_loader):
10         if i < num_batches:
11             loss = calc_loss_batch(input_batch, target_batch,
12                                   model, device)
13             total_loss += loss.item()
14         else:
15             break
16
17     return total_loss / num_batches
18
19 # Initial losses (before training)
20 train_loss = calc_loss_loader(train_loader, model, device, num_batches
21                               =5)
22 val_loss = calc_loss_loader(val_loader, model, device, num_batches=5)
23 test_loss = calc_loss_loader(test_loader, model, device, num_batches
24                               =5)
25
26 print(f"Training loss: {train_loss:.3f}") # 2.453
27 print(f"Validation loss: {val_loss:.3f}") # 2.583
28 print(f"Test loss: {test_loss:.3f}") # 2.322
```

## Dataset-wide Loss:

- Iterate through data loader
- Calculate loss per batch
- Accumulate total loss
- Return average loss

## Initial Results:

- Training: 2.453
- Validation: 2.583
- Test: 2.322

**High values:** Expected for untrained classification head

## Implement Classification Metrics

**Task:** Code accuracy and loss calculation functions

**Goal:** Understand evaluation metrics for classification

### Steps:

- 1 Implement functions → Test on sample batch → Compare results

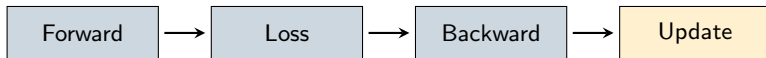
### Analysis questions:

- Why are initial accuracies near random?
- Loss vs accuracy as optimization targets?

# Training Loop Overview

## Training Structure:

- **Similar to:** Pretraining but with classification focus
- **Key differences:** Track examples (not tokens), calculate accuracy
- **Process:** Forward pass  $\rightarrow$  Loss calculation  $\rightarrow$  Backpropagation  $\rightarrow$  Update
- **Monitoring:** Loss and accuracy after each epoch



*PyTorch training loop for classification*

# Fine-tuning Function Implementation

```
1 def train_classifier_simple(model, train_loader, val_loader, optimizer
    , device, num_epochs, eval_freq, eval_iter):
2     # Tracking lists
3     train_losses, val_losses = [], []
4     train_accs, val_accs = [], []
5     examples_seen, global_step = 0, -1
6
7     for epoch in range(num_epochs):
8         model.train()
9
10        for input_batch, target_batch in train_loader:
11            optimizer.zero_grad()
12            loss = calc_loss_batch(
13                input_batch, target_batch,
14                model, device)
15            loss.backward()
16            optimizer.step()
17
18            examples_seen += input_batch.shape[0]
19            global_step += 1
20
21        # Periodic evaluation
22        if global_step % eval_freq == 0:
23            train_loss, val_loss = evaluate_model(model,
24            train_loader,
25                val_loader, device, eval_iter)
26            train_losses.append(train_loss)
27            val_losses.append(val_loss)
```

## Training Function Structure:

- Initialize tracking lists
- Loop through epochs
- Process each batch
- Calculate loss and gradients

## Key Steps:

- Zero gradients
- Forward pass
- Backward pass
- Update parameters
- Periodic evaluation

# Optimizer and Hyperparameters

```
1 import torch.optim as optim
2
3 # Optimizer setup
4 optimizer = optim.AdamW(
5     model.parameters(),
6     lr=5e-5,          # Learning rate
7     weight_decay=0.1 # Regularization
8 )
9
10 # Training hyperparameters
11 num_epochs = 5      # Sufficient for fine-tuning
12 batch_size = 8       # Memory efficient
13 eval_freq = 50       # Evaluate every 50 steps
14 eval_iter = 1        # Use 1 batch for evaluation
```

## Training Configuration:

- **Optimizer:** AdamW
- **Learning Rate:** 5e-5
- **Weight Decay:** 0.1
- **Epochs:** 5

## Key Choices:

- AdamW: Better weight decay
- Small LR: Fine-tuning
- Regularization: Prevent overfitting
- Few epochs: Sufficient

# Training Execution and Monitoring

```
1 import time
2 start_time = time.time()
3
4 train_losses, val_losses, train_accs, val_accs,
5 examples_seen = train_classifier_simple(
6     model, train_loader, val_loader,
7     optimizer, device,
8     num_epochs=5, eval_freq=50, eval_iter=1
9 )
10
11 end_time = time.time()
12 execution_time_minutes = (end_time - start_time) / 60
13 print(f"Training completed in "
14       f"{execution_time_minutes:.2f} minutes")
15
16 # Example output during training:
17 # Ep 1 (Step 000050): Train loss 1.793,
18 #                      Val loss 1.623
19 # Ep 1 (Step 000100): Train loss 0.969,
20 #                      Val loss 0.765
21 # Ep 1 (Step 000130): Train loss 0.441,
22 #                      Val loss 0.327
```

## Running the Training:

- Track execution time
- Call training function
- Monitor progress output
- Observe loss decreases

## Performance:

- 6 minutes on laptop
- <30 seconds on GPU
- Rapid improvement

**Early signs:** Rapid improvement in first epoch

# Training Results Analysis

## Epoch-by-Epoch Progress:

| Epoch | Train Accuracy | Validation Accuracy |
|-------|----------------|---------------------|
| 1     | 70.0%          | 72.5%               |
| 2     | 82.5%          | 85.0%               |
| 3     | 90.0%          | 90.0%               |
| 4     | 95.0%          | 95.0%               |
| 5     | 100.0%         | 97.5%               |

## Pattern Analysis:

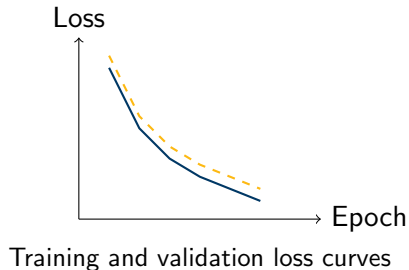
- **Rapid learning:** Dramatic improvement in epoch 1
- **Good generalization:** Small gap between train/validation
- **Convergence:** Gradual stabilization by epoch 5

**Success:** Excellent performance achieved quickly

# Loss Curve Visualization

## Loss Progression:

- **Training loss:** Rapid decline from 2.15 to 0.08
- **Validation loss:** Parallel decline, slightly higher
- **Pattern:** Sharp drop in epoch 1, gradual stabilization
- **Interpretation:** Good learning without overfitting



# Overfitting Analysis

## Signs to Watch:

- **Warning sign:** Training accuracy  $\gg$  Validation accuracy
- **Current status:** Small gap (100% vs 97.5%)
- **Validation loss:** Continues decreasing (good sign)
- **Conclusion:** Model generalizes well to unseen data

## Prevention Strategies:

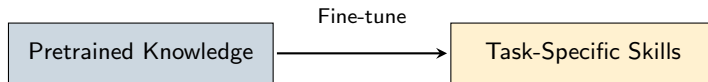
- Early stopping based on validation performance
- Regularization (weight decay, dropout)
- Reduce model complexity
- Increase training data

*Good learning pattern - minimal overfitting detected*

# Training Convergence Insights

## Why Fine-tuning Works So Well:

- **Rapid learning:** Benefits of pretrained features
- **Transfer learning:** Model adapts existing knowledge
- **Efficiency:** Much faster than training from scratch
- **Quality:** High performance with minimal training
- **Foundation:** Pretrained representations accelerate learning



Transfer linguistic understanding to classification

*Standing on the shoulders of giants*

# Training to Evaluation Transition

## Training Achievements:

- ✓ High-performing spam classifier created
- ✓ Training complete: 100% training, 97.5% validation accuracy
- ✓ Good convergence with minimal overfitting
- ✓ Understanding of fine-tuning process

**Next step:** Comprehensive evaluation on test set

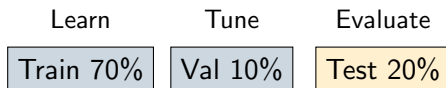
**Applications:** Deploy model for real spam detection

*From general language model to specialized classifier*

# Final Model Evaluation

## Test Set Assessment:

- **Purpose:** Unbiased performance estimate
- **Process:** Load best model, evaluate on test set
- **Metrics:** Accuracy, loss, confusion matrix
- **Expected:** Similar to validation performance ( 97%)
- **Importance:** Real-world performance indicator



*Test set provides final, unbiased performance estimate*

# Test Set Results

```
1 # Final evaluation on test set
2 test_accuracy = calc_accuracy_loader(
3     test_loader, model, device
4 )
5 test_loss = calc_loss_loader(
6     test_loader, model, device
7 )
8
9 print(f"Test accuracy: {test_accuracy*100:.2f}%")
10 # ~97%
11 print(f"Test loss: {test_loss:.3f}")
12 # Low, consistent
13
14 # Compare with initial performance
15 print("\nPerformance comparison:")
16 print(f"Initial test accuracy: 48.75%")
17 print(f"Final test accuracy: {test_accuracy*100:.2f}%")
18 print(f"Improvement: "
19       f"{test_accuracy*100 - 48.75:.2f} "
20       f"percentage points")
```

## Final Performance Metrics:

- Evaluate on test set
- Calculate accuracy and loss
- Compare to initial performance

## Results:

- **Final:** 97% accuracy
- **Initial:** 48.75% accuracy
- **Improvement:** 48 points

**Conclusion:** Model generalizes well to unseen data

# Model Inference on New Data

```
1 def classify_review(text, model, tokenizer, device, max_length=None):
2     model.eval()
3
4     # Tokenize and prepare input
5     token_ids = tokenizer.encode(text)
6     if max_length:
7         token_ids = token_ids[:max_length]
8         token_ids += [50256] * (max_length - len(token_ids)) # Pad
9
10    # Convert to tensor and add batch dimension
11    input_tensor = torch.tensor(token_ids, dtype=torch.long).unsqueeze
12        (0).to(device)
13
14    # Get model prediction
15    with torch.no_grad():
16        logits = model(input_tensor)[: , -1, :]
17        probabilities = torch.softmax(logits, dim=-1)
18        predicted_label = torch.argmax(logits, dim=-1).item()
19
20    return predicted_label, probabilities.cpu().numpy()[0]
21
22 # Example usage
23 text = "Free money! Click now to claim your prize!"
24 label, probs = classify_review(
25     text, model, tokenizer, device, max_length=120)
26 print(f"Prediction: {'Spam' if label == 1 else 'Ham'}")
27 print(f"Confidence: Ham={probs[0]:.3f}, "
28       f"Spam={probs[1]:.3f}")
```

## Inference Pipeline:

- Set model to eval mode
- Tokenize input text
- Apply padding if needed
- Convert to tensor
- Get predictions

## Output:

- Predicted label (0 or 1)
- Class probabilities
- Human-readable result

# Spam Detection Examples

```
1 # Test examples
2 examples = [
3     "Free money! Click now!",      # Clear spam
4     "Meeting at 3pm today",        # Clear ham
5     "Win a free prize today!",     # Borderline
6     "Your package will arrive tomorrow", # Ham
7     "URGENT: Claim your reward now!!!" # Clear spam
8 ]
9
10 for text in examples:
11     label, probs = classify_review(
12         text, model, tokenizer,
13         device, max_length=120
14     )
15     result = "Spam" if label == 1 else "Ham"
16     confidence = probs[label]
17     print(f"'{text}' -> {result} "
18           f"({confidence:.3f} confidence)")
19
20 # Sample output:
21 # 'Free money! Click now!' -> Spam
22 #   (0.987 confidence)
23 # 'Meeting at 3pm today' -> Ham
24 #   (0.923 confidence)
25 # 'Win a free prize today!' -> Spam
26 #   (0.756 confidence)
```

## Testing Real Examples:

- Clear spam examples
- Clear ham examples
- Borderline cases
- Confidence scores

## Analysis:

- Model learns spam indicators
- Recognizes urgency words
- Handles free/win/prize patterns

**Robustness:** Test edge cases and borderline examples

## Test Your Own Messages

**Task:** Input custom text messages for classification

**Goal:** Understand model decision-making process

### Examples to try:

- Create obvious spam messages
- Create obvious ham messages
- Test borderline cases

### Analysis questions:

- What patterns does the model recognize?
- How would you deploy this in real-world scenarios?

# Model Limitations and Considerations

## Known Limitations:

- **Training data:** Limited to SMS-style messages
- **Domain adaptation:** May not work well on emails/other text
- **Adversarial examples:** Deliberately crafted to fool model
- **Bias:** Reflects patterns in training data
- **Updates:** Need retraining as spam tactics evolve

## Mitigation Strategies:

- Regular model updates with new data
- Domain adaptation techniques
- Adversarial training
- Bias detection and correction
- Human-in-the-loop validation

*Understanding limitations is key to responsible deployment*

# Deployment Architecture

## Production Pipeline:

- **Input processing:** Text  $\rightarrow$  tokens  $\rightarrow$  padding  $\rightarrow$  model
- **Model serving:** Load trained weights, inference mode
- **Output processing:** Logits  $\rightarrow$  probabilities  $\rightarrow$  decisions
- **Scalability:** Batch processing for efficiency
- **Integration:** API endpoints, real-time vs batch processing



End-to-end inference pipeline

# Performance Optimization

## Optimization Strategies:

- **Model size:** Trade-offs between accuracy and speed
- **Quantization:** Reduce model precision for faster inference
- **Caching:** Store common preprocessing results
- **GPU acceleration:** Batch processing on GPU
- **Edge deployment:** Optimize for mobile/embedded devices

## Performance Metrics:

- Inference latency (milliseconds per prediction)
- Throughput (predictions per second)
- Memory usage (GPU/CPU requirements)
- Model size (storage requirements)
- Energy consumption (for edge devices)

*Balance accuracy, speed, and resource requirements*

# Alternative Fine-tuning Strategies

## Beyond Basic Freezing:

- **Full model fine-tuning:** Train all layers (more compute, potentially better)
- **Layer selection:** Experiment with different layer combinations
- **Learning rates:** Different rates for different layers
- **Gradual unfreezing:** Progressive layer training
- **LoRA/Adapters:** Parameter-efficient fine-tuning methods

## Parameter-Efficient Methods:

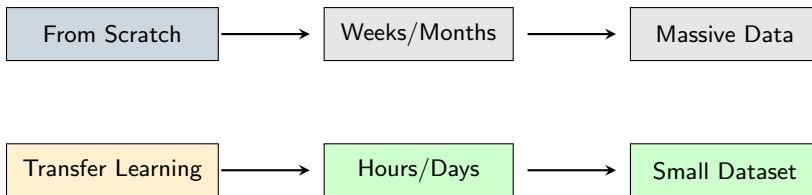
- Low-Rank Adaptation (LoRA)
- Adapter layers
- Prefix tuning
- Prompt tuning

*Choose strategy based on data, compute, and performance requirements*

# Transfer Learning Benefits

## Why Transfer Learning Works:

- **Speed:** Much faster than training from scratch
- **Data efficiency:** Good performance with less data
- **Performance:** Often better than random initialization
- **Knowledge reuse:** Leverage pretrained linguistic understanding
- **Cost:** Reduced computational requirements



*Transfer learning democratizes access to powerful models*

# Beyond Binary Classification

## Extension Possibilities:

- **Multi-class:** Extend to 3+ categories (news topics, sentiment)
- **Hierarchical:** Category and subcategory prediction
- **Multi-label:** Multiple categories per text
- **Regression:** Predict continuous values (ratings, scores)
- **Sequence labeling:** Token-level classification (NER, POS)

## Implementation Changes:

- Change output layer size ( $2 \rightarrow n$  classes)
- Modify loss function (multi-class, multi-label)
- Adjust evaluation metrics
- Update data preprocessing

*Same principles apply to many classification tasks*

# Comparison with Other Approaches

## Alternative Methods:

| Approach                 | Accuracy  | Training Time | Data Needed |
|--------------------------|-----------|---------------|-------------|
| Traditional ML (SVM, NB) | Medium    | Fast          | Medium      |
| CNN/LSTM from scratch    | Medium    | Slow          | Large       |
| Word2Vec + Classifier    | Good      | Medium        | Medium      |
| BERT Fine-tuning         | Excellent | Medium        | Small       |
| GPT-2 + Classification   | Excellent | Fast          | Small       |

## Our approach advantages:

- Excellent performance with minimal data
- Fast training due to frozen layers
- Leverages powerful pretrained representations
- Simple implementation

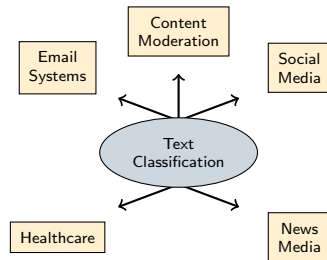
**Trade-offs:** Model size vs traditional ML approaches

## Industry Use Cases:

- **Email filtering:** Spam detection in email systems
- **Content moderation:** Identify inappropriate content
- **Sentiment analysis:** Product reviews, social media
- **News categorization:** Automatic article classification
- **Medical text:** Clinical note classification

## Success Stories:

- Gmail spam filtering (99.9% accuracy)
- Facebook content moderation
- Amazon review sentiment analysis
- Reuters news categorization
- Healthcare documentation processing



# Evaluation Metrics Deep Dive

## Beyond Accuracy:

- **Accuracy:** Overall correctness percentage
- **Precision:** True positives / (True positives + False positives)
- **Recall:** True positives / (True positives + False negatives)
- **F1-score:** Harmonic mean of precision and recall
- **Confusion matrix:** Detailed breakdown of predictions

## When to Use Each:

- **Accuracy:** Balanced datasets, general performance
- **Precision:** When false positives are costly
- **Recall:** When false negatives are costly
- **F1-score:** Imbalanced datasets, overall balance

Confusion Matrix

|        |   | Predicted |     |
|--------|---|-----------|-----|
|        |   | P         | N   |
| Actual | P | 150       | 25  |
|        | N | 30        | 795 |

Precision = 83.3%

Recall = 85.7%

Accuracy = 94.5%

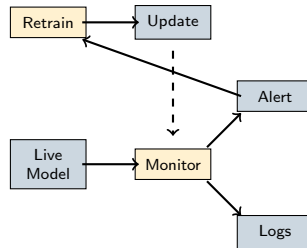
# Model Monitoring and Maintenance

## Production Monitoring:

- **Performance tracking:** Monitor accuracy over time
- **Data drift:** Detect changes in input distribution
- **Model updates:** Regular retraining schedules
- **A/B testing:** Compare model versions
- **Feedback loops:** Incorporate user corrections

## Key Metrics to Track:

- Prediction accuracy trends
- Response time/latency
- Input data distribution changes
- User feedback and corrections
- Model confidence scores



Continuous Improvement Cycle

# Chapter 6 Summary & Key Takeaways

## What We Accomplished:

- **Fine-tuning categories:** Classification vs instruction fine-tuning
- **Dataset preparation:** Balance, split, tokenize, pad
- **Architecture modification:** Replace output layer, freeze layers
- **Training process:** Cross-entropy loss, accuracy monitoring
- **Evaluation:** Test set performance, real-world applications
- **Success:** 97% accuracy on spam detection task

## Key Insights:

- Transfer learning dramatically reduces training requirements
- Frozen layers preserve pretrained knowledge
- Classification fine-tuning is efficient and effective
- Real-world deployment requires careful consideration

*From general language model to specialized classifier in 75 minutes*

# Looking Ahead - Next Steps

## Upcoming Topics:

- **Chapter 7:** Instruction fine-tuning with human feedback
- **Advanced topics:** Parameter-efficient fine-tuning (LoRA, Adapters)
- **Scaling:** Larger models, more complex tasks
- **Production:** Deployment, monitoring, maintenance
- **Your journey:** From pretraining through specialized applications

## Skills Gained Today:

- Dataset preparation for classification
- Model architecture modification
- Training loop implementation
- Performance evaluation and analysis

## Questions & Discussion