

Instruction Fine-tuning: Aligning Models with Human Instructions

Prof. Rongyu Lin

Lecture 8
Quinnipiac University
School of Computing and Engineering

Reading: Raschka Ch. 7 (pp. 239-289)
"Instruction Fine-tuning and Alignment"

Learning Objectives

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

- Understand the motivation and importance of instruction fine-tuning
- Distinguish between pre-training, fine-tuning, and instruction fine-tuning
- Prepare and format datasets for instruction-based training
- Implement complete instruction fine-tuning pipelines
- Apply parameter-efficient methods like LoRA and QLoRA
- Evaluate and deploy instruction-tuned models effectively

Transforming foundation models into helpful AI assistants

Part I: Foundations

- ① Introduction to Instruction Fine-tuning
- ② The Alignment Problem
- ③ Training Pipeline Overview

Part II: Data Preparation

- ④ Dataset Formats and Templates
- ⑤ Tokenization Strategies
- ⑥ Batching and DataLoaders

Part III: Implementation

- ⑦ Model Loading and Configuration
- ⑧ Training Loop Implementation
- ⑨ Monitoring and Optimization

Part IV: Advanced Topics

- ⑩ Parameter-Efficient Methods
- ⑪ LoRA and QLoRA
- ⑫ Evaluation and Deployment

Part I

Foundations of Instruction Fine-tuning

Understanding the alignment problem and training pipeline

What is Instruction Fine-tuning?

Definition:

- Process of training LLMs to follow human instructions
- Transforms general language models into helpful assistants
- Bridges gap between pre-training and deployment

Key Characteristics:

- Supervised learning on instruction-response pairs
- Teaches desired behavior patterns
- Improves alignment with user expectations

Examples:

Before Fine-tuning

Input: "Write a poem about AI"

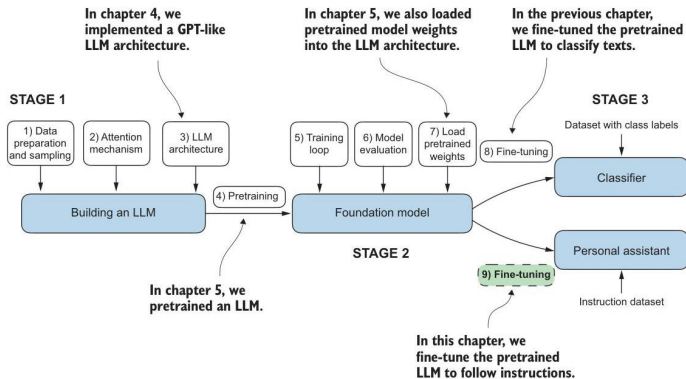
Output: "about AI is a topic..."

After Fine-tuning

Input: "Write a poem about AI"

Output: "Digital minds awake,
Silicon dreams take flight..."

The Complete LLM Training Pipeline



Three-Stage Process:

- 1 **Pre-training:** Learn language patterns
- 2 **Instruction Fine-tuning:** Align with instructions
- 3 **Preference Alignment:** Human feedback refinement

Pre-training vs Fine-tuning vs Instruction Fine-tuning

Pre-training

- Unsupervised learning
- Massive text corpora
- Next-token prediction
- General language understanding
- Months of training
- Thousands of GPUs

Fine-tuning

- Task-specific adaptation
- Domain-focused data
- Classification/regression
- Specialized capabilities
- Hours to days
- Single GPU possible

Instruction Fine-tuning

- Behavioral alignment
- Instruction-response pairs
- Following commands
- General helpfulness
- Days to weeks
- Multiple GPUs typical

Each stage builds upon the previous, creating increasingly capable models

The Alignment Problem

Pre-trained Models:

- Complete text probabilistically
- No concept of helpfulness
- May generate harmful content
- Unpredictable behavior

Challenges:

- Hallucination and fabrication
- Bias amplification
- Unsafe or unethical outputs
- Poor instruction following

Alignment Goals:

- Helpful: Assist users effectively
- Harmless: Avoid dangerous outputs
- Honest: Acknowledge limitations

Solution Approaches:

- Instruction fine-tuning (this lecture)
- RLHF (Reinforcement Learning)
- Constitutional AI
- Direct Preference Optimization

OpenAI GPT Series:

- GPT-3 → InstructGPT → ChatGPT
- 100x smaller but preferred outputs
- Instruction tuning + RLHF

Google's Models:

- LaMDA → Bard → Gemini
- Multi-modal instruction following
- Safety-first alignment approach

Open Source Models:

- Alpaca (Stanford): 52K instructions
- Vicuna (LMSYS): ShareGPT conversations
- Dolly (Databricks): Employee-generated
- WizardLM: Complex instruction evolution

Impact:

- Democratized AI assistants
- Enabled custom deployments
- Accelerated research progress

Instruction-Response Format

Instruction:

Write a Python function to check
if a number is prime.

Response:

```
def is_prime(n):  
    if n <= 1:  
        return False  
    if n <= 3:  
        return True  
    if n % 2 == 0 or n % 3 == 0:  
        return False  
    i = 5  
    while i * i <= n:  
        if n % i == 0 or n % (i + 2) == 0:  
            return False  
        i += 6  
    return True
```

Template Components:

- Instruction: User's request or query
- Input: Optional context or data
- Response: Model's desired output

Common Formats:

- Alpaca: Instruction-Input-Response
- ChatML: Role-based conversations
- Llama2: Special tokens [INST]
- Custom: Domain-specific templates

Consistent formatting is crucial for training

Types of Instruction Datasets

Human-Generated

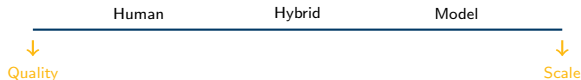
- High quality
- Expensive to create
- Limited scale
- Examples:
 - Dolly-15k
 - OpenAssistant
 - Super-NaturalInstructions

Model-Generated

- Scalable approach
- Quality varies
- Potential biases
- Examples:
 - Alpaca (GPT-3)
 - Vicuna (ShareGPT)
 - WizardLM

Hybrid Approaches

- Model generation + human review
- Instruction evolution
- Self-instruct methods
- Examples:
 - Orca (explanation tuning)
 - Constitutional AI
 - LIMA (Less Is More)



Evaluation Metrics for Instruction-Tuned Models

Automatic Metrics:

- **Perplexity:** Language modeling quality
- **BLEU/ROUGE:** N-gram overlap
- **BERTScore:** Semantic similarity
- **Task-specific:** Accuracy on benchmarks

Benchmarks:

- MMLU: Multi-task knowledge
- HumanEval: Code generation
- TruthfulQA: Factual accuracy
- MT-Bench: Multi-turn conversations

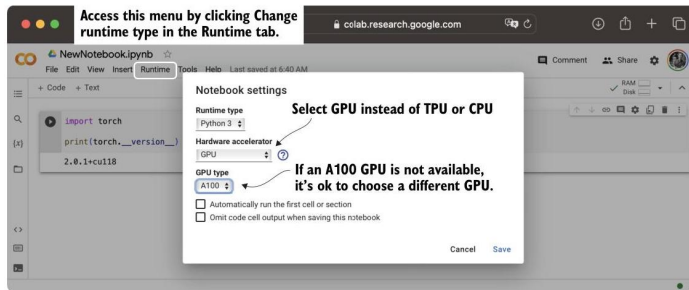
Human Evaluation:

- Helpfulness ratings
- Safety assessments
- Preference comparisons (A/B tests)
- Red-teaming for adversarial inputs

LLM-as-Judge:

- GPT-4 evaluation
- Alpaca Eval leaderboard
- Arena (pairwise comparisons)
- Self-evaluation techniques

Hardware Requirements and Setup



Minimum Requirements:

- GPU: 16GB VRAM (RTX 3090/4090)
- RAM: 32GB system memory
- Storage: 100GB for models/data

Optimization Techniques:

- Gradient checkpointing
- Mixed precision (FP16/BF16)
- Parameter-efficient methods (LoRA)

Current Challenges in Instruction Fine-tuning

Technical Challenges:

- Catastrophic forgetting
- Distribution shift
- Overfitting to instruction format
- Computational costs

Data Challenges:

- Quality vs quantity tradeoff
- Instruction diversity
- Multilingual coverage
- Domain expertise gaps

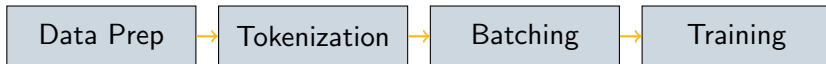
Ethical Considerations:

- Bias amplification
- Safety guarantees
- Misuse potential
- Transparency requirements

Research Directions:

- Constitutional AI methods
- Self-supervised alignment
- Continual learning approaches
- Efficient fine-tuning at scale

Looking Ahead: Implementation Pipeline



Next: Deep dive into dataset preparation and formatting

- Loading and exploring the Alpaca dataset
- Creating instruction templates
- Implementing tokenization pipelines
- Building efficient data loaders

Part II

Dataset Preparation and Formatting

Building the foundation for instruction fine-tuning

The Alpaca Dataset

Overview:

- 52,000 instruction-following examples
- Generated using GPT-3.5 (text-davinci-003)
- Based on self-instruct methodology
- Cost: \$500 to generate

Dataset Structure:

- **instruction:** The task description
- **input:** Optional context (can be empty)
- **output:** Expected response

```
1 # Example entry
2 {
3   "instruction": "Translate to French",
4   "input": "Hello, how are you?",
5   "output": "Bonjour, comment allez-vous?"
6 }
7
8 # Another example
9 {
10  "instruction": "Write a haiku about AI",
11  "input": "",
12  "output": "Silicon minds wake\nData flows like endless
13           streams\nFuture writes itself"
14 }
```

Simple format, powerful results

Loading the Alpaca Dataset

```
1 import json
2 from torch.utils.data import Dataset
3 import pandas as pd
4
5 # Load Alpaca dataset
6 def load_alpaca_dataset(path):
7     with open(path, 'r') as f:
8         data = json.load(f)
9     return data
10
11 # Analyze dataset
12 data = load_alpaca_dataset('alpaca_data.json')
13 print(f"Total examples: {len(data)}")
14
15 # Convert to DataFrame for analysis
16 df = pd.DataFrame(data)
17 print(f"Unique instructions: {df['instruction'].nunique()}")
18 print(f"Examples with input: {(df['input'] != '').sum()}")
19
20 # Sample distribution
21 inst_lengths = df['instruction'].str.len()
22 print(f"Avg instruction length: {inst_lengths.mean():.0f}")
```

Dataset Statistics:

- 52,002 total examples
- 39,825 unique instructions
- 17,858 with input context
- Avg length: 85 characters

Instruction Categories:

- Generation: 28%
- Classification: 18%
- Q&A: 22%
- Rewriting: 15%
- Translation: 8%
- Other: 9%

Diverse tasks ensure generalization

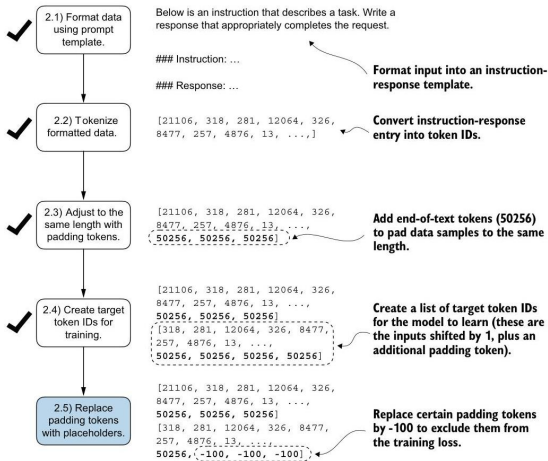
Data Preprocessing Pipeline

Five-Step Process:

- 1 **Data Loading:** Read JSON/CSV files into memory
- 2 **Format Conversion:** Apply instruction template formatting
- 3 **Tokenization:** Convert text to token IDs
- 4 **Padding/Truncation:** Ensure uniform sequence lengths
- 5 **Target Creation:** Generate labels for training

Key Considerations:

- Memory efficiency
- Batch processing
- Data shuffling
- Special tokens handling



Instruction Template Formatting

```
1 def format_instruction(entry):
2     """Format entry with Alpaca template"""
3     instruction = entry['instruction']
4     input_text = entry['input']
5
6     if input_text:
7         prompt = f"""Below is an instruction that describes a task,
8         paired with an input that provides further context. Write a
9         response that appropriately completes the request.
10
11         Instruction:
12         {instruction}
13
14         Input:
15         {input_text}
16
17         Response:
18         """
19     else:
20         prompt = f"""Below is an instruction that describes a task. Write
21         a response that appropriately completes the request.
22
23         Instruction:
24         {instruction}
25
26         Response:
27         """
28     return prompt + entry['output']
```

Template Components:

- System prompt for context
- Clear section markers
- Consistent formatting
- Flexibility for input presence

Best Practices:

- Use consistent delimiters
- Preserve whitespace structure
- Include system instructions
- Handle edge cases gracefully

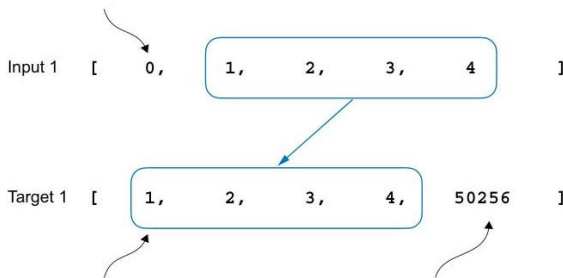
Output Example:

Below is an instruction...

Instruction:
Explain quantum computing

Response:
Quantum computing uses...

Tokenization for Instruction Fine-tuning



Input Processing:

- Convert text to token IDs
- Add special tokens (BOS/EOS)
- Maintain instruction boundaries

Target Creation:

- Shift tokens by one position
- Mask instruction tokens (optional)
- Preserve response for loss calculation

Implementing Tokenization

```
1 from transformers import AutoTokenizer
2
3 # Initialize tokenizer
4 tokenizer = AutoTokenizer.from_pretrained("meta-llama/Llama-2-7b-hf")
5 tokenizer.pad_token = tokenizer.eos_token
6
7 def tokenize_example(example, max_length=512):
8     """Tokenize instruction-response pair"""
9     # Format the example
10    text = format_instruction(example)
11
12    # Tokenize with padding and truncation
13    tokens = tokenizer(
14        text,
15        max_length=max_length,
16        padding="max_length",
17        truncation=True,
18        return_tensors="pt")
19
20    Create labels (shift by 1 for next-token prediction)
21    labels = tokens["input_ids"].clone()
22    labels[tokens["attention_mask"] == 0] = -100
23
24    return {
25        "input_ids": tokens["input_ids"].squeeze(),
26        "attention_mask": tokens["attention_mask"].squeeze(),
27        "labels": labels.squeeze()}
```

Key Parameters:

- `max_length`: Sequence limit
- `padding`: Strategy for short sequences
- `truncation`: Handle long sequences
- `return_tensors`: Output format

Label Masking:

- Use -100 for ignored tokens
- PyTorch ignores in loss calculation
- Can mask instruction portion
- Focus learning on responses

Proper tokenization is critical for training

Advanced: Instruction Masking

```
1 def mask_instruction_tokens(tokens, response_start_idx):
2     """Mask instruction portion, keep response"""
3     labels = tokens["input_ids"].clone()
4
5     # Find response delimiter position
6     response_token = tokenizer.encode("Response:",
7                                     add_special_tokens=False)
8
9     # Find where response starts in sequence
10    for i in range(len(labels) - len(response_token)):
11        if labels[i:i+len(response_token)].tolist() == response_token:
12            response_start_idx = i + len(response_token)
13            break
14
15    # Mask everything before response
16    if response_start_idx > 0:
17        labels[:response_start_idx] = -100
18
19    # Also mask padding tokens
20    labels[tokens["attention_mask"] == 0] = -100
21
22    return labels
23
24 # Apply masking
25 tokens = tokenizer(formatted_text, ...)
26 labels = mask_instruction_tokens(tokens, response_idx)
```

Why Mask Instructions?

- Focus learning on responses
- Prevent overfitting to prompts
- Improve generation quality
- Reduce training noise

Masking Strategy:

```
[KEEP] Below is an instruction...
[MASK] Instruction:
[MASK] Write a poem
[MASK] Response:
[KEEP] Roses are red...
[KEEP] Violets are blue...
```

Optional but often improves results

Creating PyTorch Dataset

```
1 import torch
2 from torch.utils.data import Dataset
3
4 class InstructionDataset(Dataset):
5     def __init__(self, data, tokenizer, max_length=512):
6         self.data = data
7         self.tokenizer = tokenizer
8         self.max_length = max_length
9
10    def __len__(self):
11        return len(self.data)
12
13    def __getitem__(self, idx):
14        example = self.data[idx]
15        text = format_instruction(example)
16        encoding = self.tokenizer(
17            text,
18            max_length=self.max_length,
19            padding="max_length",
20            truncation=True,
21            return_tensors="pt")
22
23        labels = encoding["input_ids"].clone()
24        labels[encoding["attention_mask"] == 0] = -100
25
26        return {"input_ids": encoding["input_ids"].squeeze(),
27                "attention_mask": encoding["attention_mask"].squeeze(),
28                "labels": labels.squeeze()}
```

Dataset Methods:

- `__init__`: Initialize with data
- `__len__`: Return dataset size
- `__getitem__`: Get single example

Usage:

```
1 # Create dataset
2 train_dataset = InstructionDataset(
3     data=train_data,
4     tokenizer=tokenizer,
5     max_length=512)
6
7 # Check size
8 print(f"Dataset size: {len(train_dataset)}")
9
10 # Get sample
11 sample = train_dataset[0]
12 print(f"Input shape: {sample['input_ids'].shape}")
```

PyTorch Dataset enables efficient loading

Data Validation and Quality Checks

```
1 def validate_dataset(dataset, tokenizer):
2     """Perform quality checks on dataset"""
3     issues = []
4
5     for idx in range(min(100, len(dataset))):
6         sample = dataset[idx]
7
8         if sample['attention_mask'].sum() == 0:
9             issues.append(f"Empty sequence at index {idx}")
10
11        if sample['attention_mask'][-1] == 1:
12            text = tokenizer.decode(sample['input_ids'])
13            if not text.endswith(tokenizer.eos_token):
14                issues.append(f"Truncated at index {idx}")
15
16        valid_labels = sample['labels'][sample['labels'] != -100]
17        if len(valid_labels) < 10:
18            issues.append(f"Too few labels at index {idx}")
19
20    # Report findings
21    if issues:
22        print(f"Found {len(issues)} potential issues:")
23        for issue in issues[:5]:
24            print(f"  - {issue}")
25    else:
26        print("Dataset validation passed!")
27
28    return len(issues) == 0
```

Common Issues:

- Empty or very short responses
- Truncated sequences
- Encoding errors
- Missing special tokens
- Inconsistent formatting

Quality Metrics:

- Response length distribution
- Token vocabulary coverage
- Special token placement
- Label/input ratio

Remediation:

- Filter problematic examples
- Adjust max_length if needed
- Fix formatting issues
- Balance dataset if skewed

Data Splitting Strategy

```
1 from sklearn.model_selection import train_test_split
2
3 def split_instruction_data(data, val_size=0.1, test_size=0.05):
4     """Split data into train/val/test sets"""
5
6     # First split: train+val vs test
7     train_val, test = train_test_split(data, test_size=test_size,
8                                       random_state=42, shuffle=True)
9
10    # Second split: train vs val
11    val_ratio = val_size / (1 - test_size)
12    train, val = train_test_split(train_val, test_size=val_ratio,
13                                  random_state=42, shuffle=True)
14
15    print(f"Dataset splits:")
16    print(f"  Train: {len(train):,} examples")
17    print(f"  Val: {len(val):,} examples")
18    print(f"  Test: {len(test):,} examples")
19
20    return train, val, test
21
22 # Apply splitting
23 train_data, val_data, test_data = split_instruction_data(
24     alpaca_data, val_size=0.1, test_size=0.05)
```

Split Ratios:

- Training: 85% (44,202 examples)
- Validation: 10% (5,200 examples)
- Test: 5% (2,600 examples)

Considerations:

- Maintain task distribution
- Avoid data leakage
- Stratified splitting if needed
- Keep similar examples together

Advanced Splitting:

- By instruction type
- By difficulty level
- By response length
- Cross-validation for small data

Data Augmentation Techniques

Instruction Paraphrasing:

- Rephrase instructions
- Maintain semantic meaning
- Increase diversity

Response Variations:

- Multiple valid responses
- Different writing styles
- Varying detail levels

Back-translation:

- Translate to another language
- Translate back to original
- Creates natural variations

```
1 def augment_instruction(example):
2     """Simple augmentation example"""
3     templates = [
4         "Please {task}",
5         "Can you {task}",
6         "I need you to {task}",
7         "{task}",
8         "Help me {task}"
9     ]
10
11     instruction = example['instruction']
12     # Random template selection
13     import random
14     template = random.choice(templates)
15
16     # Apply template
17     if instruction.lower().startswith(('please', 'can you')):
18         task = instruction
19     else:
20         task = instruction.lower()
21
22     augmented = template.format(task=task)
23
24     return {
25         'instruction': augmented,
26         'input': example['input'],
27         'output': example['output']
28     }
```

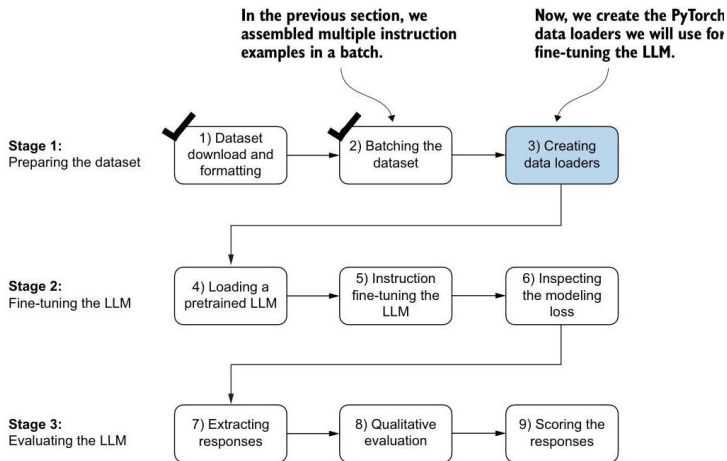
From Data to Batches

Next Steps:

- Create efficient data loaders for batching
- Handle variable-length sequences
- Optimize memory usage
- Prepare for distributed training

Key Concepts:

- Dynamic batching
- Padding strategies
- Sequence bucketing
- Memory pinning



Part III

Batching and Data Loading

Efficient data pipelines for model training

Why Batching Matters

Benefits of Batching:

- GPU utilization improvement
- Stable gradient estimates
- Faster training convergence
- Memory efficiency

Batch Size Considerations:

- GPU memory limits
- Model size constraints
- Gradient accumulation
- Learning rate scaling

Typical Batch Sizes:

- 7B model: 4-8 per GPU
- 13B model: 2-4 per GPU
- 30B model: 1-2 per GPU
- With gradient accumulation: effectively larger

Memory Formula:

$$\begin{aligned}\text{Memory} &= \text{Model} + \text{Gradients} \\ &\quad + \text{Optimizer} + \text{Activations} \\ &\approx 4 \times \text{Model Size}\end{aligned}$$

Creating Efficient DataLoaders

```
1 from torch.utils.data import DataLoader
2 import torch
3
4 def create_dataloaders(train_dataset, val_dataset,
5                       batch_size=4, num_workers=4):
6
7     # Training dataloader with shuffling
8     train_loader = DataLoader(
9         train_dataset,
10        batch_size=batch_size,
11        shuffle=True,
12        num_workers=num_workers,
13        pin_memory=True,
14        drop_last=True, # Drop incomplete batches
15        collate_fn=instruction_collate_fn)
16
17     # Validation dataloader without shuffling
18     val_loader = DataLoader(
19         val_dataset,
20        batch_size=batch_size,
21        shuffle=False,
22        num_workers=num_workers,
23        pin_memory=True,
24        drop_last=False,
25        collate_fn=instruction_collate_fn)
26
27     print(f"Train batches: {len(train_loader)}")
28     print(f"Val batches: {len(val_loader)}")
29
```

DataLoader Parameters:

- `batch_size`: Samples per batch
- `shuffle`: Randomize order
- `num_workers`: Parallel loading
- `pin_memory`: GPU transfer optimization
- `drop_last`: Handle final batch
- `collate_fn`: Custom batching logic

Performance Tips:

- Use multiple workers for I/O
- Enable `pin_memory` for GPU
- Prefetch next batch
- Profile loading time

Custom Collate Function

```
1 def instruction_collate_fn(batch):
2     """Custom collate for instruction data"""
3     input_ids = torch.stack([item['input_ids'] for item in batch])
4     attention_mask = torch.stack([item['attention_mask'] for item in
5                                   batch])
6     labels = torch.stack([item['labels'] for item in batch])
7
8     actual_max_len = attention_mask.sum(dim=1).max().item()
9
10    if actual_max_len < input_ids.shape[1]:
11        input_ids = input_ids[:, :actual_max_len]
12        attention_mask = attention_mask[:, :actual_max_len]
13        labels = labels[:, :actual_max_len]
14
15    return {'input_ids': input_ids,
16            'attention_mask': attention_mask,
17            'labels': labels}
18
19 # Alternative: padding on-the-fly
20 def dynamic_collate_fn(batch):
21     """Pad to max length in current batch only"""
22     # Get max length in this batch
23     max_len = max(len(item['input_ids']) for item in batch)
24     # Pad each sequence to batch max length
25     # More memory efficient than padding to global max
```

Collate Function Role:

- Combine samples into batch
- Handle variable lengths
- Apply padding strategy
- Create tensors for model

Padding Strategies:

- **Static:** Pad to fixed max_length
- **Dynamic:** Pad to batch max
- **Bucketing:** Group similar lengths

Dynamic padding saves memory **Memory Impact:**

- Static: Consistent but wasteful
- Dynamic: Efficient but variable
- Trade-off: Speed vs memory

Gradient Accumulation for Large Batch Sizes

```
1 def train_with_accumulation(model, dataloader,
2                             accumulation_steps=4):
3     """Train with gradient accumulation"""
4     optimizer.zero_grad()
5
6     for batch_idx, batch in enumerate(dataloader):
7         # Forward pass
8         outputs = model(**batch)
9         loss = outputs.loss / accumulation_steps
10
11        # Backward pass
12        loss.backward()
13
14        # Update weights after accumulation
15        if (batch_idx + 1) % accumulation_steps == 0:
16            optimizer.step()
17            optimizer.zero_grad()
18
19        # Log effective batch size
20        effective_batch = batch_size * accumulation_steps
21        print(f"Effective batch size: {effective_batch}")
```

Why Gradient Accumulation?

- Simulate larger batches
- Overcome GPU memory limits
- Improve training stability
- Better convergence

Effective Batch Formula:

$$\text{Effective} = \text{Physical} \times \text{Accumulation}$$

Memory Savings:

- 4x accumulation = 4x larger batch
- Same memory as single batch
- Trade computation for memory

Memory Optimization Strategies

Memory Reduction Techniques:

- **Mixed Precision:** FP16/BF16 training
- **Gradient Checkpointing:** Trade compute for memory
- **CPU Offloading:** Move optimizer states
- **Model Sharding:** Distribute across GPUs

Memory Breakdown:

- Model weights: 2 bytes/param (FP16)
- Gradients: 2 bytes/param
- Optimizer states: 8 bytes/param (Adam)
- Activations: Variable

```
1 # Mixed precision training
2 from torch.cuda.amp import autocast, GradScaler
3
4 scaler = GradScaler()
5
6 for batch in dataloader:
7     optimizer.zero_grad()
8
9     with autocast():
10         outputs = model(**batch)
11         loss = outputs.loss
12
13     scaler.scale(loss).backward()
14     scaler.step(optimizer)
15     scaler.update()
16
17 # Memory: 50% reduction
18 # Speed: 2-3x faster
```

Distributed and Parallel Data Loading

```
1 from torch.utils.data.distributed import DistributedSampler
2 import torch.distributed as dist
3
4 def setup_distributed_dataloader(dataset, world_size, rank):
5     """Setup distributed data loading"""
6     # Create distributed sampler
7     sampler = DistributedSampler(
8         dataset,
9         num_replicas=world_size,
10        rank=rank,
11        shuffle=True,
12        seed=42
13    )
14
15    # Create dataloader with sampler
16    dataloader = DataLoader(
17        dataset,
18        batch_size=batch_size // world_size,
19        sampler=sampler,
20        num_workers=4,
21        pin_memory=True,
22        persistent_workers=True
23    )
24
25    return dataloader
26
27 # Each GPU gets unique subset
28 # No data duplication across GPUs
```

Distributed Benefits:

- Linear scaling with GPUs
- No data overlap
- Synchronized training
- Efficient memory usage

Key Parameters:

- world_size: Total GPUs
- rank: GPU identifier
- num_replicas: Same as world_size
- persistent_workers: Keep alive

Essential for multi-GPU training

Bucketing for Efficient Batching

```
1 from torch.utils.data import Sampler
2 import numpy as np
3
4 class BucketSampler(Sampler):
5     def __init__(self, dataset, batch_size, bucket_size=100):
6         self.dataset = dataset
7         self.batch_size = batch_size
8         self.bucket_size = bucket_size
9
10        # Sort by sequence length
11        self.lengths = [len(d['input_ids']) for d in dataset]
12        self.sorted_idx = np.argsort(self.lengths)
13
14    def __iter__(self):
15        # Create buckets of similar lengths
16        buckets = []
17        for i in range(0, len(self.sorted_idx), self.bucket_size):
18            bucket = self.sorted_idx[i:i+self.bucket_size]
19            np.random.shuffle(bucket) # Shuffle within bucket
20            buckets.extend(bucket)
21
22        # Yield batches from buckets
23        for i in range(0, len(buckets), self.batch_size):
24            yield buckets[i:i+self.batch_size]
```

Bucketing Benefits:

- Minimize padding waste
- Similar lengths in batch
- Faster training
- Better GPU utilization

Example Buckets:

Bucket 1: [100-200 tokens]
Bucket 2: [200-300 tokens]
Bucket 3: [300-400 tokens]
Bucket 4: [400-512 tokens]

Memory Impact:

- 30-50% memory savings
- 20-30% speed improvement
- Maintains randomness

Monitoring Data Pipeline Performance

```
1 import time
2 from torch.utils.data import DataLoader
3
4 class DataLoaderMonitor:
5     def __init__(self, dataloader):
6         self.dataloader = dataloader
7         self.batch_times = []
8
9     def __iter__(self):
10         start_time = time.time()
11
12         for batch_idx, batch in enumerate(self.dataloader):
13             batch_time = time.time() - start_time
14             self.batch_times.append(batch_time)
15
16             # Log statistics
17             if batch_idx % 100 == 0:
18                 avg_time = np.mean(self.batch_times[-100:])
19                 throughput = len(batch['input_ids']) / batch_time
20
21                 print(f"Batch {batch_idx}:")
22                 print(f"  Load time: {batch_time:.3f}s")
23                 print(f"  Throughput: {throughput:.0f} samples/s")
24                 print(f"  Avg time: {avg_time:.3f}s")
25
26             yield batch
27             start_time = time.time()
```

Key Metrics:

- Data loading time
- Throughput (samples/sec)
- GPU utilization
- Memory usage
- Queue depths

Bottleneck Indicators:

- GPU idle time $\geq$ 10%
- Data loading $\geq$ forward pass
- Memory spikes
- Uneven batch times

Monitor to identify optimization opportunities

Caching for Repeated Data Access

Caching Levels:

- **Dataset Cache:** Preprocessed tokens
- **Memory Cache:** Frequently accessed batches
- **Disk Cache:** Persistent storage
- **GPU Cache:** Pinned memory

When to Cache:

- Multiple epochs training
- Expensive preprocessing
- Small datasets
- Repeated validation

```
1 class CachedDataset(Dataset):
2     def __init__(self, data, tokenizer, cache_dir):
3         self.data = data
4         self.tokenizer = tokenizer
5         self.cache_dir = cache_dir
6         self.cache = {}
7
8     def __getitem__(self, idx):
9         # Check cache first
10        if idx in self.cache:
11            return self.cache[idx]
12
13        # Process and cache
14        item = self.process_item(self.data[idx])
15        self.cache[idx] = item
16
17        # Save to disk
18        cache_file = f"{self.cache_dir}/{idx}.pt"
19        torch.save(item, cache_file)
20
21        return item
```

Prefetching and Asynchronous Loading

```
1 from torch.utils.data import DataLoader
2 import torch.cuda
3
4 class PrefetchLoader:
5     def __init__(self, dataloader, device='cuda'):
6         self.dataloader = dataloader
7         self.device = device
8
9     def __iter__(self):
10        stream = torch.cuda.Stream()
11        first = True
12
13        for next_batch in self.dataloader:
14            with torch.cuda.stream(stream):
15                # Async copy to GPU
16                next_batch = {k: v.to(self.device, non_blocking=
17                    True)
18                                for k, v in next_batch.items()}
19
20            if not first:
21                yield batch # Return previous
22            else:
23                first = False
24
25            torch.cuda.current_stream().wait_stream(stream)
26            batch = next_batch
27            yield batch # Last batch
28 # GPU processes while CPU prepares next batch
```

Prefetching Benefits:

- Hide data transfer latency
- Overlap compute and I/O
- Better GPU utilization
- Faster overall training

Pipeline Timeline:

CPU: Load B1 | Load B2 | Load B3
GPU: | Proc B1 | Proc B2

Performance Gains:

- 10-20% speed improvement
- Near 100% GPU utilization
- Minimal code changes

Troubleshooting Common Data Loading Issues

Common Problems:

- **OOM Errors:** Batch size too large
- **Slow Loading:** I/O bottleneck
- **Deadlocks:** Multi-process issues
- **Data Corruption:** Encoding errors

Debugging Tools:

- torch.utils.bottleneck
- nvidia-smi monitoring
- Memory profilers
- Data validation scripts

```
1 # Debug data loading
2 def debug_dataloader(loader):
3     print(f"Batches: {len(loader)}")
4
5     for i, batch in enumerate(loader):
6         if i >= 3: break # Test first 3
7
8         print(f"\nBatch {i}:")
9         for k, v in batch.items():
10             print(f"    {k}: {v.shape}")
11             print(f"    dtype: {v.dtype}")
12             print(f"    device: {v.device}")
13             print(f"    mem: {v.nbytes/1e6:.1f}MB")
14
15 # Check for issues
16 if torch.isnan(v).any():
17     print("WARNING: NaN values!")
```

Part IV

Model Loading and Fine-tuning

Implementing the complete training pipeline

Loading Pre-trained Models

```
1 from transformers import AutoModelForCausalLM
2 import torch
3
4 def load_model_for_finetuning(model_name, device_map="auto"):
5     """Load pre-trained model for fine-tuning"""
6
7     # Load with automatic device mapping
8     model = AutoModelForCausalLM.from_pretrained(
9         model_name,
10         torch_dtype=torch.bfloat16, # Use BF16 for stability
11         device_map=device_map,      # Auto GPU assignment
12         trust_remote_code=True,     # For custom models
13         use_cache=False             # Disable KV cache for training
14     )
15
16     # Enable gradient checkpointing
17     model.gradient_checkpointing_enable()
18
19     # Prepare for training
20     model.train()
21
22     print(f"Model loaded: {model.num_parameters()/1e9:.1f}B params")
23     print(f"Memory: {torch.cuda.memory_allocated()/1e9:.1f}GB")
24
25     return model
```

Loading Options:

- torch_dtype: FP16/BF16/FP32
- device_map: GPU placement
- use_cache: KV cache control
- load_in_8bit: Quantization

Model Sizes:

- 7B: 14GB (FP16)
- 13B: 26GB (FP16)
- 30B: 60GB (FP16)
- 70B: 140GB (FP16)

BF16 recommended for training stability

Preparing Model for Training

```
1 def prepare_model_for_training(model, use_lora=False):
2     """Configure model for fine-tuning"""
3
4     # Freeze embeddings (optional)
5     for param in model.get_input_embeddings().parameters():
6         param.requires_grad = False
7
8     # Enable training mode
9     model.train()
10
11 if use_lora:
12     from peft import LoraConfig, get_peft_model
13
14     lora_config = LoraConfig(
15         r=16,                # Rank
16         lora_alpha=32,        # Scaling
17         target_modules=["q_proj", "v_proj"],
18         lora_dropout=0.1,
19         bias="none",
20         task_type="CAUSAL_LM"
21     )
22
23     model = get_peft_model(model, lora_config)
24     print(f"Trainable params: {model.print_trainable_parameters()}")
25
26     return model
```

Training Configurations:

- Full fine-tuning: All parameters
- LoRA: 0.1% parameters
- Frozen embeddings: Stability
- Gradient checkpointing: Memory

Memory Requirements:

- Full: 4x model size
- LoRA: 1.5x model size
- QLoRA: 0.5x model size

Training Speed:

- Full: Baseline
- LoRA: 1.5-2x faster
- QLoRA: 1.2x faster

Training Configuration and Hyperparameters

```
1 from transformers import TrainingArguments
2
3 training_args = TrainingArguments(
4     output_dir="./instruction-tuned-model",
5     per_device_train_batch_size=4,
6     per_device_eval_batch_size=4,
7     gradient_accumulation_steps=4,
8     num_train_epochs=3,
9     learning_rate=2e-5,
10    warmup_steps=100,
11    logging_steps=10,
12    save_strategy="steps",
13    save_steps=500,
14    evaluation_strategy="steps",
15    eval_steps=100,
16    bf16=True,                # Use BF16
17    tf32=True,                # Tensor cores
18    gradient_checkpointing=True,
19    dataloader_num_workers=4,
20    report_to="tensorboard",
21    load_best_model_at_end=True,
22    metric_for_best_model="eval_loss",
23    greater_is_better=False,
24 )
```

Key Hyperparameters:

- Learning rate: 1e-5 to 5e-5
- Batch size: 4-32 per GPU
- Epochs: 1-3 typical
- Warmup: 3-10% of steps

Optimization Tips:

- Start with small LR
- Use gradient accumulation
- Monitor loss closely
- Early stopping if overfitting

Adjust based on model size and data

Custom Training Loop Implementation

```
1 def train_epoch(model, dataloader, optimizer, scheduler):
2     """Custom training loop for one epoch"""
3     model.train()
4     total_loss = 0
5     progress_bar = tqdm(dataloader, desc="Training")
6     for step, batch in enumerate(progress_bar):
7         batch = {k: v.to('cuda') for k, v in batch.items()}
8
9         # Forward pass
10        outputs = model(*batch)
11        loss = outputs.loss
12
13        # Backward pass
14        loss.backward()
15
16        # Gradient clipping
17        torch.nn.utils.clip_grad_norm_(model.parameters(), 1.0)
18
19        # Optimizer step
20        optimizer.step()
21        scheduler.step()
22        optimizer.zero_grad()
23
24        # Update metrics
25        total_loss += loss.item()
26        progress_bar.set_postfix({'loss': loss.item()})
27
28    return total_loss / len(dataloader)
```

Training Loop Components:

- Forward pass: Compute loss
- Backward pass: Compute gradients
- Gradient clipping: Stability
- Optimizer step: Update weights
- Scheduler step: Adjust LR

Monitoring:

- Loss per batch
- Learning rate
- Gradient norm
- GPU memory

Custom loops offer maximum control

Using HuggingFace Trainer API

```
1 from transformers import Trainer, DataCollatorForLanguageModeling
2
3 # Data collator for padding
4 data_collator = DataCollatorForLanguageModeling(
5     tokenizer=tokenizer,
6     mlm=False, # Causal LM, not masked
7     pad_to_multiple_of=8 # Efficient for GPU
8 )
9
10 # Initialize trainer
11 trainer = Trainer(
12     model=model,
13     args=training_args,
14     train_dataset=train_dataset,
15     eval_dataset=val_dataset,
16     data_collator=data_collator,
17     tokenizer=tokenizer,
18     compute_metrics=compute_metrics,
19     callbacks=[early_stopping_callback],
20 )
21
22 # Start training
23 trainer.train()
24
25 # Save the fine-tuned model
26 trainer.save_model("./final-model")
27 tokenizer.save_pretrained("./final-model")
```

Trainer Advantages:

- Automatic mixed precision
- Distributed training support
- Built-in logging
- Checkpoint management
- Easy evaluation

Built-in Features:

- TensorBoard integration
- W&B logging
- Early stopping
- Learning rate scheduling
- Gradient accumulation

Recommended for most use cases

Monitoring Training Progress

```
import torch.nn.functional as F

y = torch.tensor([1.0])
x1 = torch.tensor([1.1])
w1 = torch.tensor([2.2])
b = torch.tensor([0.0])
z = x1 * w1 + b
a = torch.sigmoid(z)
loss = F.binary_cross_entropy(a, y)
```

This import statement is a common convention in PyTorch to prevent long lines of code.

True label

Input feature

Weight parameter

Bias unit

Net input

Activation and output

Key Metrics:

- Training loss curve
- Validation loss trend
- Learning rate schedule
- Gradient norms

Warning Signs:

- Loss explosion
- Validation divergence
- Gradient vanishing
- Memory leaks

Preventing Overfitting

Regularization Techniques:

- **Dropout:** Random neuron deactivation
- **Weight Decay:** L2 regularization
- **Early Stopping:** Monitor validation
- **Data Augmentation:** Increase diversity

Detection Methods:

- Train/val loss divergence
- Performance plateau
- Memorization patterns
- Generalization failure

```
1 from transformers import EarlyStoppingCallback
2
3 # Early stopping configuration
4 early_stopping = EarlyStoppingCallback(
5     early_stopping_patience=3,
6     early_stopping_threshold=0.001
7 )
8
9 # Dropout in model config
10 model.config.attention_dropout = 0.1
11 model.config.hidden_dropout = 0.1
12
13 # Weight decay in optimizer
14 optimizer = AdamW(
15     model.parameters(),
16     lr=2e-5,
17     weight_decay=0.01 # L2 regularization
18 )
```

Checkpoint Management and Recovery

```
1 import os
2 import torch
3 class CheckpointManager:
4     def __init__(self, save_dir, max_checkpoints=3):
5         self.save_dir = save_dir
6         self.max_checkpoints = max_checkpoints
7         self.checkpoints = []
8
9     def save_checkpoint(self, model, optimizer, epoch, loss):
10         checkpoint = {'epoch': epoch,
11                      'model_state_dict': model.state_dict(),
12                      'optimizer_state_dict': optimizer.state_dict(),
13                      'loss': loss,}
14
15         path = f"{self.save_dir}/checkpoint_epoch_{epoch}.pt"
16         torch.save(checkpoint, path)
17         self.checkpoints.append(path)
18
19         if len(self.checkpoints) > self.max_checkpoints:
20             old = self.checkpoints.pop(0)
21             os.remove(old)
22
23     def load_checkpoint(self, path, model, optimizer):
24         checkpoint = torch.load(path)
25         model.load_state_dict(checkpoint['model_state_dict'])
26         optimizer.load_state_dict(checkpoint['optimizer_state_dict'])
27         return checkpoint['epoch'], checkpoint['loss']
```

Checkpoint Contents:

- Model weights
- Optimizer state
- Learning rate scheduler
- Training metrics
- Random states

Best Practices:

- Save regularly (every N steps)
- Keep best checkpoints
- Version control configs
- Document experiments

Essential for long training runs

Multi-GPU and Distributed Training

```
1 import torch.distributed as dist
2 from torch.nn.parallel import DistributedDataParallel as DDP
3
4 def setup_distributed(rank, world_size):
5     """Initialize distributed training"""
6     os.environ['MASTER_ADDR'] = 'localhost'
7     os.environ['MASTER_PORT'] = '12355'
8
9     dist.init_process_group("nccl", rank=rank, world_size=world_size)
10
11 def train_distributed(rank, world_size):
12     setup_distributed(rank, world_size)
13
14     # Create model on specific GPU
15     model = model.to(rank)
16     model = DDP(model, device_ids=[rank])
17
18     # Distributed sampler
19     sampler = DistributedSampler(dataset, num_replicas=world_size, rank=rank)
20
21     dataloader = DataLoader(dataset, sampler=sampler, batch_size=batch_size)
22
23     # Train as normal
24     for epoch in range(num_epochs):
25         sampler.set_epoch(epoch) # Shuffle differently each epoch
26         train_epoch(model, dataloader, optimizer)
```

Parallelism Types:

- **Data Parallel:** Split batches
- **Model Parallel:** Split model
- **Pipeline Parallel:** Split layers
- **Tensor Parallel:** Split tensors

Scaling Efficiency:

- 2 GPUs: 1.8x speedup
- 4 GPUs: 3.5x speedup
- 8 GPUs: 6.5x speedup

Essential for large models

Evaluation During Training

```
1 def evaluate(model, eval_dataloader, compute_metrics=None):
2     """Evaluate model on validation set"""
3     model.eval()
4     total_loss = 0
5     predictions = []
6     references = []
7
8     with torch.no_grad():
9         for batch in tqdm(eval_dataloader, desc="Evaluating"):
10             batch = {k: v.to('cuda') for k, v in batch.items()}
11
12             outputs = model(**batch)
13             total_loss += outputs.loss.item()
14
15             # Generate predictions
16             generated = model.generate(batch['input_ids'], max_length=512,
17                                       temperature=0.7, do_sample=True)
18
19             predictions.extend(tokenizer.batch_decode(generated))
20             references.extend(tokenizer.batch_decode(batch['labels']))
21
22     # Compute metrics
23     metrics = {'eval_loss': total_loss / len(eval_dataloader)}
24     if compute_metrics:
25         metrics.update(compute_metrics(predictions, references))
26
27     return metrics
```

Evaluation Metrics:

- Perplexity
- BLEU score
- ROUGE scores
- Custom task metrics

Evaluation Frequency:

- Every N steps
- End of epoch
- On demand
- Best model tracking

Generation Settings:

- Temperature: 0.7-1.0
- Top-p: 0.9-0.95
- Top-k: 40-50

Generation and Inference

```
1 def generate_response(model, tokenizer, instruction, input_text=""):
2     """Generate response for instruction"""
3     model.eval()
4
5     # Format prompt
6     if input_text:
7         prompt = f"""Below is an instruction that describes a task,
8         paired with an input. Write a response that appropriately completes
9         the request.
10
11         \### Instruction:
12         {instruction}
13
14         \### Input:
15         {input_text}
16
17         \### Response:
18         """
19     else:
20         prompt = f"""Below is an instruction that describes a task. Write
21         a response that appropriately completes the request.
22
23         \### Instruction:
24         {instruction}
25
26         \### Response:
27         """
28
29     # Tokenize
30     inputs = tokenizer(prompt, return_tensors="pt").to('cuda')
```

Generation Parameters:

- max_new_tokens: Output length
- temperature: Randomness
- top_p: Nucleus sampling
- top_k: Top-k sampling
- do_sample: Enable sampling

Decoding Strategies:

- Greedy: Deterministic
- Sampling: Random
- Beam Search: Multiple paths
- Nucleus: Dynamic vocabulary

Balance quality vs diversity

Common Training Issues and Solutions

Loss Explosion:

- Lower learning rate
- Gradient clipping
- Check data quality
- Use FP32 or BF16

Out of Memory:

- Reduce batch size
- Enable gradient checkpointing
- Use mixed precision
- Apply LoRA/QLoRA

Slow Convergence:

- Increase learning rate
- Adjust warmup steps
- Check data distribution
- Use larger batch size

Poor Quality:

- More training data
- Longer training
- Adjust temperature
- Check prompt format

Systematic debugging leads to solutions

Training Best Practices

Before Training:

- Validate data quality
- Test on small subset
- Profile memory usage
- Set up monitoring

During Training:

- Monitor loss curves
- Track GPU utilization
- Save regular checkpoints
- Evaluate periodically

After Training:

- Thorough evaluation
- Ablation studies
- Document results
- Version control models

General Tips:

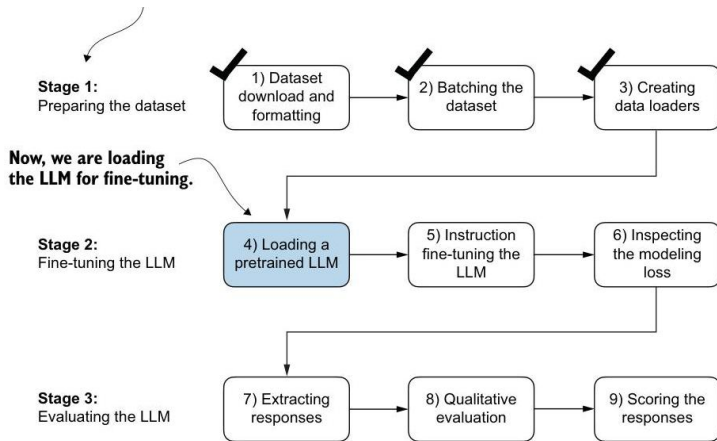
- Start simple, iterate
- Keep detailed logs
- Reproduce experiments
- Share learnings

Complete Training Pipeline Example

End-to-End Pipeline:

- 1 Load and preprocess data → Format instructions
- 2 Create datasets and dataloaders → Efficient batching
- 3 Load pre-trained model → Configure for training
- 4 Train with monitoring → Track progress
- 5 Evaluate and save → Deploy best model

Complete in 5 steps!



Part V

Advanced Techniques

Parameter-efficient methods and optimization

LoRA: Low-Rank Adaptation

LoRA Concept:

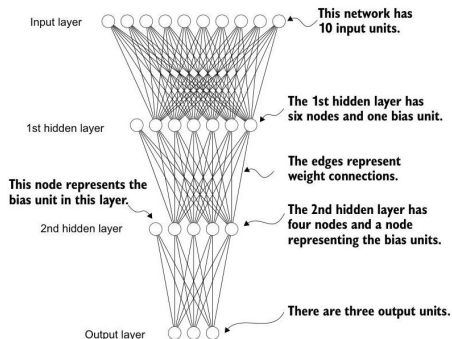
- Freeze pre-trained weights
- Add trainable rank decomposition matrices
- Merge at inference time
- Dramatically reduce parameters

Mathematical Formulation:

$$W' = W + BA$$

where:

- $W \in \mathbb{R}^{d \times k}$: Frozen weights
- $B \in \mathbb{R}^{d \times r}$: Trainable
- $A \in \mathbb{R}^{r \times k}$: Trainable



Benefits:

- 10,000x fewer parameters
- Same quality as full fine-tuning
- Multiple LoRA adapters possible

Implementing LoRA

```
1 from peft import LoraConfig, get_peft_model, TaskType
2
3 # Configure LoRA
4 lora_config = LoraConfig(
5     r=16,                                # Rank
6     lora_alpha=32,                        # Scaling factor
7     target_modules=[                     # Which modules to adapt
8         "q_proj",                         # Query projection
9         "k_proj",                         # Key projection
10        "v_proj",                         # Value projection
11        "o_proj",                         # Output projection
12        "gate_proj",                     # MLP gate
13        "up_proj",                       # MLP up
14        "down_proj",                     # MLP down
15    ],
16    lora_dropout=0.1,                    # Dropout for regularization
17    bias="none",                         # Don't train biases
18    task_type=TaskType.CAUSAL_LM         # Task type
19 )
20
21 # Apply LoRA to model
22 model = get_peft_model(model, lora_config)
23
24 # Check trainable parameters
25 model.print_trainable_parameters()
26 # Output: trainable params: 4,194,304 || all params: 6,738,415,616 ||
      trainable%: 0.06
```

LoRA Hyperparameters:

- **r (rank):** 4-64 typical
- **alpha:** r to 2r typical
- **dropout:** 0.05-0.1
- **target_modules:** Attention layers

Memory Savings:

- 7B model: 28GB → 4GB
- 13B model: 52GB → 8GB
- 30B model: 120GB → 16GB

0.1% parameters, 90% performance

QLoRA: Quantized LoRA

```
1 from transformers import BitsAndBytesConfig
2 import torch
3
4 # 4-bit quantization config
5 bnb_config = BitsAndBytesConfig(
6     load_in_4bit=True,                # 4-bit quantization
7     bnb_4bit_quant_type="nf4",        # NormalFloat4
8     bnb_4bit_use_double_quant=True,   # Double quantization
9     bnb_4bit_compute_dtype=torch.bfloat16 # Compute in BF16
10 )
11
12 # Load quantized model
13 model = AutoModelForCausalLM.from_pretrained(
14     model_name,
15     quantization_config=bnb_config,
16     device_map="auto",
17     trust_remote_code=True
18 )
19
20 # Apply LoRA on top of quantized model
21 model = prepare_model_for_kbit_training(model)
22 model = get_peft_model(model, lora_config)
23
24 # Training uses ~25% of full precision memory
25 print(f"Memory usage: {torch.cuda.memory_allocated()/1e9:.1f}GB")
```

QLoRA Components:

- 4-bit base model (frozen)
- 16-bit LoRA adapters
- 32-bit optimizer states
- Paged optimizers

Memory Formula:

$$\text{Memory} = \frac{\text{Model}}{4} + \text{LoRA} + \text{Gradients}$$

Performance:

- 65B model on single 48GB GPU
- 95% of full fine-tuning quality
- 4x slower than LoRA

Adapter-Based Fine-tuning

Adapter Architecture:

- Insert small networks between layers
- Keep original model frozen
- Train only adapter parameters
- Multiple tasks, multiple adapters

Adapter Design:

- Down-projection: $d \rightarrow r$
- Non-linearity: ReLU/GELU
- Up-projection: $r \rightarrow d$
- Residual connection

```
1 class Adapter(nn.Module):
2     def __init__(self, dim, r=64):
3         super().__init__()
4         self.down = nn.Linear(dim, r)
5         self.up = nn.Linear(r, dim)
6         self.act = nn.GELU()
7
8     def forward(self, x):
9         residual = x
10        x = self.down(x)
11        x = self.act(x)
12        x = self.up(x)
13        return x + residual # Skip connection
14
15 # Insert adapters after attention/FFN
16 for layer in model.transformer.layers:
17     layer.attn_adapter = Adapter(dim)
18     layer.ffn_adapter = Adapter(dim)
```

Prompt Tuning and P-Tuning

```
1 from peft import PromptTuningConfig, get_peft_model
2
3 # Soft prompt tuning
4 prompt_config = PromptTuningConfig(
5     task_type="CAUSAL_LM",
6     prompt_tuning_init="TEXT", # Initialize from text
7     num_virtual_tokens=20,     # Soft prompt length
8     prompt_tuning_init_text="Classify the following text:",
9     tokenizer_name_or_path=model_name,)
10
11 # Apply prompt tuning
12 model = get_peft_model(model, prompt_config)
13
14 # Only soft prompts are trainable
15 model.print_trainable_parameters()
16 # trainable params: 40,960 || all params: 6,738,456,576 || 0.0006%
17
18 # P-Tuning v2: Deep prompt tuning
19 class PrefixEncoder(nn.Module):
20     def __init__(self, prefix_len, hidden_size):
21         super().__init__()
22         self.embedding = nn.Embedding(prefix_len, hidden_size)
23         self.trans = nn.Sequential(
24             nn.Linear(hidden_size, hidden_size),
25             nn.Tanh(),
26             nn.Linear(hidden_size, hidden_size * 2)
27         )
```

Prompt Tuning Types:

- **Hard Prompts:** Fixed text
- **Soft Prompts:** Learnable embeddings
- **Prefix Tuning:** Prepend to keys/values
- **P-Tuning v2:** Deep prompts

Advantages:

- Extremely parameter efficient
- Task-specific prompts
- No model modification
- Fast switching between tasks

Best for multi-task scenarios

Gradient Checkpointing

Memory vs Compute Trade-off:

- Don't store all activations
- Recompute during backward pass
- 60% memory reduction
- 20% slower training

How it Works:

- Checkpoint certain layers
- Forward: Store checkpoints only
- Backward: Recompute from checkpoints
- Memory: $O(\sqrt{n})$ instead of $O(n)$

```
1 # Enable gradient checkpointing
2 model.gradient_checkpointing_enable()
3
4 # Custom checkpointing
5 from torch.utils.checkpoint import checkpoint
6
7 class CheckpointedLayer(nn.Module):
8     def forward(self, x):
9         # Checkpoint expensive operations
10         def custom_forward(x):
11             x = self.attention(x)
12             x = self.mlp(x)
13             return x
14
15         return checkpoint(custom_forward, x)
16
17 # Memory comparison
18 # Without: 24GB for 7B model
19 # With: 10GB for 7B model
```

Mixed Precision Training

```
1 from torch.cuda.amp import autocast, GradScaler
2
3 # Initialize gradient scaler
4 scaler = GradScaler()
5
6 for epoch in range(num_epochs):
7     for batch in dataloader:
8         optimizer.zero_grad()
9
10        # Mixed precision forward pass
11        with autocast(dtype=torch.bfloat16):
12            outputs = model(**batch)
13            loss = outputs.loss
14
15        # Scale loss and backward
16        scaler.scale(loss).backward()
17
18        # Unscale before clipping
19        scaler.unscale_(optimizer)
20        torch.nn.utils.clip_grad_norm_(model.parameters(), 1.0)
21
22        # Optimizer step with scaling
23        scaler.step(optimizer)
24        scaler.update()
25
26 # BF16 vs FP16
27 # BF16: Better numerical stability, same range as FP32
28 # FP16: Better hardware support, needs loss scaling
```

Precision Formats:

- **FP32:** Full precision (baseline)
- **FP16:** Half precision, needs scaling
- **BF16:** Brain float, more stable
- **TF32:** Tensor cores on A100

Benefits:

- 2x memory reduction
- 2-3x speed improvement
- Larger batch sizes
- Minimal accuracy loss

Essential for large model training

Knowledge Distillation for Compression

```
1 import torch.nn.functional as F
2
3 def distillation_loss(student_logits, teacher_logits,
4                       labels, temperature=3.0, alpha=0.5):
5     """Knowledge distillation loss"""
6     teacher_probs = F.softmax(teacher_logits / temperature, dim=-1)
7     student_log_probs = F.log_softmax(student_logits / temperature, dim
8                                       ==-1)
9
10    kl_loss = F.kl_div(student_log_probs, teacher_probs,
11                      reduction='batchmean') * (temperature ** 2)
12
13    ce_loss = F.cross_entropy(student_logits, labels)
14
15    total_loss = alpha * kl_loss + (1 - alpha) * ce_loss
16    return total_loss
17
18 teacher_model.eval()
19 for batch in dataloader:
20     with torch.no_grad():
21         teacher_outputs = teacher_model(**batch)
22
23     student_outputs = student_model(**batch)
24     loss = distillation_loss(student_outputs.logits,
25                             teacher_outputs.logits,
26                             batch['labels'])
```

Distillation Process:

- Large teacher model (e.g., 70B)
- Small student model (e.g., 7B)
- Transfer knowledge via soft labels
- Maintain performance

Key Parameters:

- **Temperature:** Controls softness of probability distribution
- **Alpha:** Balance between distillation and hard labels

Knowledge Distillation for Compression - Part 2

```
1 # Training loop with distillation
2 teacher_model.eval() # Freeze teacher
3 student_optimizer = torch.optim.AdamW(
4     student_model.parameters(), lr=1e-4)
5
6 for epoch in range(num_epochs):
7     for batch in dataloader:
8         # Get teacher predictions (no gradient)
9         with torch.no_grad():
10             teacher_outputs = teacher_model(**batch)
11
12         # Get student predictions
13         student_outputs = student_model(**batch)
14
15         # Calculate distillation loss
16         loss = distillation_loss(
17             student_outputs.logits,
18             teacher_outputs.logits,
19             batch['labels'],
20             temperature=3.0,
21             alpha=0.7 # 70% distillation, 30% hard
22         )
23
24         # Backprop and optimize
25         loss.backward()
26         student_optimizer.step()
27         student_optimizer.zero_grad()
```

Compression Results:

Model	Size	Perf
Teacher	70B	100%
Student	7B	95%
Tiny	1.5B	88%

Benefits:

- 10x size reduction
- 5x faster inference
- Edge deployment possible

Deploy efficiently!

Comparing Parameter-Efficient Methods

Method	Params	Memory	Speed	Quality	Use Case
Full Fine-tuning	100%	High	Baseline	Best	Small models
LoRA	0.1%	Medium	Fast	95-98%	General purpose
QLoRA	0.1%	Low	Slow	93-96%	Large models
Adapters	1-5%	Medium	Medium	92-95%	Multi-task
Prompt Tuning	0.01%	Very Low	Very Fast	85-90%	Few tasks

Selection Criteria:

- Hardware constraints
- Quality requirements
- Number of tasks
- Deployment scenario
- Training time budget
- Model size

Summary

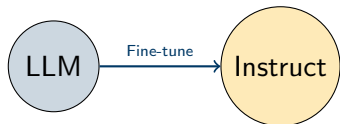
Key Takeaways and Next Steps

Putting it all together

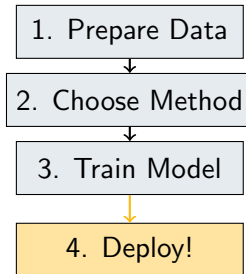
Key Takeaways

Core Concepts:

- Instruction fine-tuning transforms LLMs into assistants
- Quality data is more important than quantity
- Proper formatting ensures training success



Implementation Roadmap:



Success Formula:

$$\text{Data Quality} \times \text{Method Efficiency} = \text{Training Success}$$

Getting Started Example:

```
1 # Minimal setup for beginners
2 from transformers import (
3     AutoModelForCausalLM,
4     AutoTokenizer,
5     TrainingArguments,
6     Trainer
7 )
8 from peft import LoraConfig
9
10 # Start small
11 model_name = "meta-llama/Llama-2-7b"
12 data = load_dataset("tatsu-lab/alpaca")
13
14 # Use LoRA
15 lora_config = LoraConfig(
16     r=8, # Low rank
17     lora_alpha=16,
18     target_modules=["q_proj", "v_proj"]
19 )
```

Common Pitfalls:

- Overfitting to format
- Catastrophic forgetting
- Poor data quality
- Insufficient evaluation
- Deployment challenges

Resources:

- HuggingFace Transformers
- PEFT library
- Alpaca dataset
- LLM training guides
- Community forums

Future Directions in Instruction Fine-tuning

Research Frontiers:

- Constitutional AI approaches
- Self-supervised alignment
- Instruction evolution techniques
- Cross-lingual transfer
- Multi-modal instructions

Industry Trends:

- Specialized domain models
- Edge deployment solutions
- Privacy-preserving training
- Continuous learning systems
- Automated fine-tuning pipelines

Technical Advances:

- More efficient architectures
- Better quantization methods
- Improved alignment algorithms
- Faster training techniques
- Enhanced evaluation metrics

Opportunities:

- Custom AI assistants
- Domain expertise capture
- Personalized models
- Real-time adaptation
- Collaborative training

Advanced Topics:

- Reinforcement Learning from Human Feedback (RLHF)
- Constitutional AI and safety
- Direct Preference Optimization (DPO)
- Multi-modal instruction tuning

Practical Applications:

- Building domain-specific assistants
- Deployment strategies
- Performance optimization
- Evaluation and monitoring

Thank you! Questions?