

ADDALY · THE AI CLASSROOM

Fine-Tuning, and When Not To

The case against fine-tuning first, then how to do it properly.

9 modules · 84 lessons · 30 figures · 9 case studies · 61,811 words · about 13 hours

Dr Mustafa Mun

Author and editor

ABOUT THIS BOOK

Fine-Tuning, and When Not To, published by Addaly, 2026. Author and editor: **Dr Mustafa Mun**.

The manuscript was drafted with the assistance of a large language model and was edited and published under his name. That is stated plainly here rather than left to be discovered, because a reader is entitled to know how a book they are trusting was made.

This book is the printed form of the course of the same name at addaly.com/learn/fine-tuning. The website is the living version and is corrected first; where the two differ, the website is right.

Free to read, free to download, free to print, and free to teach from. There is no paid edition and there never will be. If you paid for this, somebody has sold you something that was given away.

Every diagram in it is drawn from data by code, not generated as a picture, so the numbers on the page are the numbers in the argument. Answers to every exercise, test and examination question are at the back, with a note on why each wrong option attracts people — those notes are the teaching, not the marking.

Contents

1. Deciding whether to fine-tune at all

Diagnose a model failure as a format, knowledge, style, capability or cost problem, name the cheapest lever that fixes that class, and defend in writing a decision to fine-tune — or not to — against the prompt, retrieval, encoder and distillation alternatives and their licence constraints

1. The Honest Answer First
2. What Fine-Tuning Changes, and What It Cannot Add
3. Diagnosing the failure before choosing the fix
4. Try the retriever first
5. When the answer is one of six labels
6. Distillation, the case that usually survives
7. The cost of owning a model
8. Licences, and what you may train on
9. Writing the decision down
10. Case study: The custom model that was already announced
11. Module test

2. How training actually works

Trace one optimisation step from a batch of tokens to a changed weight, and predict which way loss, memory and stability move when you change the learning rate, the batch size, the precision or the number of epochs

1. One training step, taken apart
2. Reading the loss number
3. Optimisers, and what their state costs
4. Learning rate, warmup and the schedule
5. Batch size, accumulation and the token count
6. Precision, and why bf16 ended a class of bugs
7. Where the memory actually goes
8. Epochs, overfitting and checkpoint selection
9. Regularisation that helps, and the kind that does not
10. Case study: A free T4, a deadline, and a learning rate
11. Module test

3. Full, LoRA and the parameter-efficient family

Predict on paper the memory a given model, method and sequence length will need, choose between full fine-tuning, LoRA, QLoRA and the tiny methods for a stated budget, and justify a rank, an alpha and a target-module set

1. Full, LoRA, QLoRA, and the Memory Arithmetic
2. What LoRA is actually doing
3. Choosing rank, alpha and which modules to touch
4. QLoRA: quantising the part you are not training
5. DoRA, rsLoRA and the variant question
6. Soft prompts and prefixes
7. The tiny methods, and a gotcha in one of them
8. Full fine-tuning, when it is genuinely right
9. Training on a free or nearly free budget
10. Case study: One card, in the building
11. Module test

4. Building the dataset

Assemble a training set from named sources with checked licences, convert it to one chat format, filter and deduplicate it, mix in replay data at a stated ratio, and publish a data card that lets a stranger reproduce the set

1. The Dataset Decides Everything
2. Where the data legitimately comes from
3. The three formats, and the conversion trap
4. Cleaning at scale
5. Writing the answers, and keeping them consistent
6. Generating a dataset from a teacher
7. Mixtures and replay
8. Packing, document boundaries and long context
9. Multi-turn conversations and tool calls
10. Memorisation, and what you cannot take back
11. The data card
12. Case study: Nine hundred examples, labelled before the rules existed
13. Module test

5. Objectives beyond next-token prediction

Choose between supervised fine-tuning, a preference method such as DPO or KTO, a verifiable-reward method such as GRPO, and continued pretraining, and state the data each needs, the memory each costs, and the way each is gamed

1. What supervised fine-tuning cannot express
2. Where preference data comes from
3. DPO, from first principles
4. KTO, ORPO, SimPO and choosing among them
5. RLHF and PPO, honestly
6. Verifiable rewards and GRPO
7. Reward hacking
8. Continued pretraining on raw text
9. Training a head instead of generating tokens
10. Case study: The verifier that could be gamed, and the baseline that could not
11. Module test

6. The mechanics of a run

Configure and launch a fine-tuning job with the correct chat template, masking and tokenizer handling, and diagnose from the logs alone whether a run is not learning, diverging, or memorising

1. Chat Templates and Loss Masking
2. Running a LoRA, and What It Costs
3. The tokenizer is part of the model
4. Adding tokens, and initialising them properly
5. A working config, line by line
6. Debugging a run
7. What to log, and how to read it
8. Checkpointing and resuming
9. Reproducibility, and the noise floor
10. Multi-GPU without tears
11. Case study: The model that would not stop talking
12. Module test

7. Fine-tuning beyond a chat model

Adapt an embedding model, a reranker, a text classifier, a speech model, a vision-language model or an image generator, and name in each case what differs from fine-tuning a text LM and which cheaper technique should be tried first

1. Fine-tuning an embedding model
2. Training a reranker
3. Fine-tuning a classifier properly
4. Fine-tuning a speech model
5. Fine-tuning a vision-language model
6. LoRA for image models, and the likeness question
7. Getting structure without training for it
8. Small models, where fine-tuning matters most
9. Tabular and time series: the honest answer
10. Case study: The vision model and the scanner
11. Module test

8. Did it actually work

Compare a tuned model against a cost-matched baseline on a leak-free held-out set, report the difference with an interval, and demonstrate whether general capability, safety behaviour or calibration regressed

1. Did It Help, and What Did It Break
2. The baseline you must beat
3. A held-out set that actually holds
4. Pairwise human comparison, done properly
5. Judging a model you distilled from the judge
6. Contamination you can check, and contamination you cannot
7. Safety regression, including from benign data
8. Calibration after tuning
9. Ablations that earn their cost
10. Case study: Ninety-four, then seventy-nine
11. Module test

9. Shipping it, and living with it

Serve a tuned model or adapter at a stated cost and latency, monitor it for drift, re-produce a six-month-old artefact from its recorded version, and plan for the day the base model is replaced

1. Serving the Adapter
2. Quantising the model you ship
3. Serving many adapters from one base
4. The cost at steady state
5. Monitoring a model that fails silently
6. The day the base model moves
7. Versioning so a model can be explained later
8. Model cards and disclosure
9. Retiring a fine-tune
10. Case study: The fine-tune that the quantiser rounded away
11. Module test

Fine-Tuning, and When Not To — final examination

The paper that opens the next course. Answers to everything follow it.

MODULE 1

Deciding whether to fine-tune at all

Most fine-tuning projects should not happen. This module gives you the diagnosis that tells you which failures weights can fix, the cheaper levers that fix the rest, and the true cost of owning a model you trained yourself.

BY THE END OF THIS MODULE YOU CAN

Diagnose a model failure as a format, knowledge, style, capability or cost problem, name the cheapest lever that fixes that class, and defend in writing a decision to fine-tune — or not to — against the prompt, retrieval, encoder and distillation alternatives and their licence constraints

1 · 7 min

The Honest Answer First

THE ONE THING TO KEEP

A prompt changes in ten seconds; a fine-tune changes in a day. Exhaust the fast loop first.

Start with the honest answer

Most projects that reach for fine-tuning are solved by a better prompt or by retrieval. Not some. Most. This is the first lesson because if the course changes your mind here, it has done its job and you can stop reading.

Fine-tuning has a slow loop. A prompt change takes ten seconds and you see the result immediately. A fine-tune takes a dataset, a training run, and an evaluation — a day if you are practised, a week if you are not — and when the result is worse, you often cannot tell which of the twenty things you changed did it. Anything you can learn from a prompt, learn from a prompt.

The ladder

Work down this list. Stop as soon as the problem is solved.

1. **Write the instruction properly.** State the output shape. State the constraints. Most "the model won't follow my format" complaints are a prompt that never specified the format precisely.
2. **Put three to five real examples in the prompt.** Few-shot moves format and tone further than people expect, and you can change it at 3am without a GPU.
3. **Give it the facts.** If the answer depends on your documents, policies, or prices, put the relevant text in the context window. This is the answer to "it doesn't know our stuff" nearly every time.
4. **Give it tools.** If the answer depends on a calculation or a live lookup, let it call something.
5. **Try a stronger base model** with everything above, and compare cost honestly.
6. **Now consider fine-tuning.**

The twenty-example test

Before you build anything, do this. Take twenty real inputs your system got wrong — real ones, from logs, not ones you invented. Hand-write the output you wanted for each. Put five into the prompt as examples and run the other fifteen.

- If it now gets most of the fifteen right, you do not need to fine-tune. You need that prompt.
- If it fails because it does not know a fact, you need retrieval.
- If it fails because it cannot do the reasoning at all, fine-tuning is unlikely to rescue it either.
- If it gets them right only when the examples are in the prompt, and your prompt has no room for them at your volume or your latency budget, that is a real fine-tuning case.

Those twenty hand-written outputs are also the start of your dataset, so the test is not wasted work either way.

When fine-tuning is genuinely right

- **A behaviour you cannot prompt reliably at scale.** A strict output shape, a house register, a domain convention that takes 2,000 tokens of instructions to describe. Fine-tuning moves that into the weights and your prompt gets short.

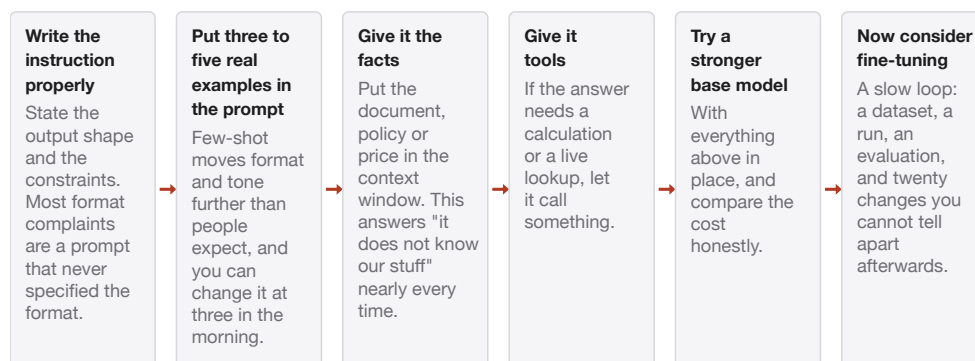
- **Cost and latency.** A tuned 7B that matches a large model on your one narrow task can be ten to fifty times cheaper per token, and faster. At a million requests a month that is the whole business case.
- **Privacy or air-gap.** The model has to run on your hardware, so the large hosted one is not an option at any quality.
- **A narrow task with plenty of labelled data.** Extraction, classification, routing. Here a small tuned model beats a large prompted one on accuracy and cost at the same time.

When it is the wrong answer

- Adding facts, especially facts that change. Use retrieval.
- Fixing hallucination. Training on answers the model does not know makes it more confidently wrong, which the next lesson explains.
- Teaching reasoning the base cannot do.
- Anything where you have forty examples, no evaluation set, and a deadline.

The rest of this course assumes you walked the ladder and still have a real case. That is a respectable place to be. It is just a much smaller group than the people currently opening a fine-tuning notebook.

The ladder, and where most projects should stop



Work down it and stop at the first rung that solves the problem. A prompt change takes ten seconds; a fine-tune takes a day if you are practised.

BEFORE YOU MOVE ON

A support bot keeps stating your old refund policy, which changed three weeks ago. Your logs show it does this confidently and often. What do you do first?

- A. Fine-tune on 2,000 past support transcripts so the model learns how your team handles refunds.
- B. Put the current policy text into the context at answer time, then re-check the same failing cases.
- C. Fine-tune on 200 examples that state the new policy explicitly.
- D. Move to a larger base model, which hallucinates less.

2 · 8 min

What Fine-Tuning Changes, and What It Cannot Add

THE ONE THING TO KEEP

Fine-tuning teaches the shape of an answer, not its content; facts belong in the context window.

What the gradient actually does

Fine-tuning shows the model an input and a target output, then nudges every trainable weight in the direction that would have made those target tokens slightly more likely. Repeat a few thousand times. That is the entire mechanism, and everything fine-tuning is good and bad at follows from it.

What it moves efficiently is the conditional shape of the output: given an input like this, produce an output like that. Register, length, structure, formatting, when to refuse, when to ask a clarifying question, which of several behaviours the model already has should fire here. These patterns appear in every example, so every gradient step reinforces them. A hundred examples can visibly change tone. Two thousand can lock a JSON schema so hard the model stops emitting prose at all.

What it does not add: facts

A fact appears in your dataset a handful of times. The base model saw trillions of tokens. To make one fact reliably retrievable you are fighting an enormous prior with a

tiny signal, and gradient descent takes the cheaper route: rather than learn the fact, learn the *form* of a confident answer.

That produces the characteristic failure. The tuned model answers in your house style, with your structure, at your length — and invents the specifics. It is worse than before, because it now sounds like your best expert while being wrong.

Even where memorisation does happen, the knowledge is:

- **uncitable** — there is no source to show the user;
- **unupdatable** — the policy changes and you retrain;
- **undeletable** — someone asks you to remove one record and there is no record to remove.

Retrieval has none of those problems, and its failure mode ("I could not find anything") is the one you want.

The hallucination amplifier

Here is the rule that catches most teams. If a training example's answer depends on knowledge the base model does not have, you are training the model to guess confidently.

Say you build 3,000 question-and-answer pairs from your internal wiki and train without showing the model the wiki. Every one of those examples demonstrates the same lesson: when asked about internal matters, produce a specific, confident answer. At inference, asked about something outside the training set, that is exactly what it does.

The fix is not more data. It is to train the behaviour with the context present — examples that include the retrieved passage and an answer grounded in it — or to include examples where the correct output is that it does not have the information.

What it can add, honestly

- **Mapping your language onto what it knows.** Your jargon, abbreviations, ticket codes, product names. That is a translation task and it is learnable.
- **Output shape.** Schemas, tags, field order, the exact vocabulary of a controlled label set.
- **Decision boundaries.** Which of five categories, which team to route to, when to escalate. This is where a small tuned model does best.
- **Refusal and calibration.** Including teaching it to say it does not know, which is a behaviour rather than a fact.

- **Narrow reasoning patterns, partly.** Training on worked traces from a stronger model does raise in-domain performance; that is how many small reasoning models are built. But it transfers the pattern, not the capability underneath. The base's ceiling is still the ceiling, and the gains rarely survive far outside the training distribution.

What it cannot touch at all

Context length. Tool access. Anything about your infrastructure. And behaviour on inputs unlike anything in your data — off-distribution, the tuned model can be worse than the base, which is the subject of a later lesson.

One sentence to keep

If answering requires *knowing* something, put it in the context. If answering requires *deciding* something, fine-tuning can teach the decision.

BEFORE YOU MOVE ON

You fine-tuned a 7B model on 5,000 question-and-answer pairs written from your internal documentation. It now answers in your house style and invents plausible specifics that are wrong, sounding more authoritative than it did before. What happened?

- The examples taught the form of a confident answer far more efficiently than they taught the facts, so the model now produces the shape of an expert answer with invented content.
- Two epochs was not enough exposure for the facts to stick; train longer on the same data.
- The LoRA rank was too low to hold that much knowledge; raise it.
- The base model is too small; a 70B base would have absorbed the documentation.

Diagnosing the failure before choosing the fix

THE ONE THING TO KEEP

Sort thirty real failures into format, knowledge, style, capability, adherence and cost before choosing a fix; only adherence, style and cost have fine-tuning as their cheapest answer.

Thirty failures, sorted by cause

Nobody fine-tunes because of a hypothesis. They fine-tune because something is wrong and fine-tuning is the most impressive-sounding thing to do about it. The cure is dull: collect thirty real failures, sort them by cause, and count.

Thirty is enough because you are not measuring a rate, you are finding out which *kind* of problem you have. If twenty-two of thirty land in one bucket, the decision is made. If they scatter evenly across six buckets, you do not have a model problem, you have a product that has not been specified.

The six causes, and the test for each

Format. The content is right and the shape is wrong: prose where you wanted JSON, three bullet points where you wanted five, a stray preamble before the answer. *Test:* run a schema validator over the output and separately ask a person whether the content is correct. If the validator fails and the person says the content is fine, it is a format failure.

Knowledge. The model does not know a fact about your world. *Test — the paste test:* put the missing document into the prompt and ask again. If it now answers correctly, the model was never short of ability, only of information.

Style and register. Correct content in the wrong voice — too chatty, too hedged, American spelling in a British document, an apology on the front of every reply. *Test:* rewrite one output by hand. If your rewrite changes no facts, it is style.

Capability. Multi-step arithmetic, long-horizon planning, a genuinely hard inference. *Test:* give the identical prompt to the largest model you can reach. If the frontier model also fails, you are not going to add the capability with two thousand examples and a

rank-16 adapter. Training teaches a model to produce answers of a shape it can already produce.

Instruction adherence. The model satisfies four of your seven constraints and drops the rest, usually the ones stated earliest. *Test:* count satisfied constraints per output and plot against the number of constraints in the prompt. If the satisfaction rate falls as constraints rise, it is adherence — and this is one of the places fine-tuning genuinely shines.

Cost or latency. The outputs are fine. The bill is not, or the p95 is 4.2 seconds and your product needs 800 milliseconds. Nothing about quality is broken.

The labelling loop

```
CAUSES = ["format", "knowledge", "style", "capability", "adherence",
"cost"]

for row in failures:                # 30 rows, read in a text editor is
fine
    print(row["prompt"][:400])
    print("----")
    print(row["output"][:400])
    row["cause"] = input("cause? ") # one label, the thing you'd fix
FIRST
```

One label per item. Failures overlap — a badly formatted answer to a question the model could not have known is both — so the rule is: label it by what you would have to change first. If you fixed the knowledge and the format problem would still be there, label it knowledge.

What each cause actually wants

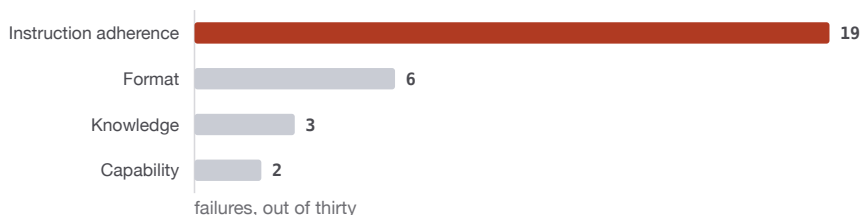
- **Format** → constrained decoding first. A grammar or a JSON schema at sampling time makes invalid output *impossible* rather than unlikely, and it costs nothing to try. Fine-tuning for format works, but it only makes the bad output rarer.
- **Knowledge** → retrieval. Weights are a poor filing cabinet; the next lesson in this module is entirely about this.
- **Style** → five examples in the prompt, then, if the style is long or subtle enough that the examples eat your context budget, fine-tuning. Style is the cause fine-tuning fixes most reliably.
- **Capability** → a bigger model, or decomposition into steps each of which the small model can do.

- **Adherence** → fine-tuning, honestly. A model trained on a thousand examples that all respect your seven constraints learns the constraints as a habit rather than as instructions competing for attention.
- **Cost** → distillation. Not a quality problem at all, and the one place where fine-tuning is unambiguously the right tool.

Where people get stuck

The commonest mislabel is knowledge as capability. A model that confidently invents your refund window looks stupid, and stupidity feels like a capability problem. Run the paste test before you believe that. The second commonest is style as adherence: "it ignores my instruction to be concise" is usually a style failure, because the model has a strong prior about how long an answer should be and your one word is arguing with millions of training examples.

Thirty real failures, sorted by cause



Style and cost were empty in this sample, so the thirty landed in four of the six buckets. Adherence is one of the few causes whose cheapest answer is fine-tuning; the three knowledge cases went to retrieval.

The limitation of this whole exercise is that thirty failures are the failures *you noticed*. Silent failures — plausible, wrong, unreported — are invisible to it, and they are usually knowledge failures. So treat the sorted counts as a map of your complaints, not of your quality.

BEFORE YOU MOVE ON

A support assistant keeps stating the wrong refund window. Someone pastes the refund policy document into the prompt and the model answers correctly every time. What has the paste test established?

- A. A knowledge failure — the fix is getting the document into the context, not changing the weights.
- B. That the base model is too small, and a larger one would know the refund window without being told.
- C. That the model can learn the policy from training, since it clearly understood the document when shown it.
- D. That the prompt is too short, and lengthening it with more instructions about refunds will resolve the errors.

4 · 9 min

Try the retriever first

THE ONE THING TO KEEP

Recall@k is a hard ceiling on a RAG system's accuracy, so measure it before blaming the generator — and when something must be trained, the retriever is a hundred times smaller and outlives the base model.

The generator is usually innocent

When a retrieval-augmented system answers badly, almost everybody blames the model that wrote the answer. That is the wrong end. Before any of that, measure one number: how often the passage containing the answer was in the context at all.

This is recall at k. Take fifty questions for which you know the correct source passage, run your retriever, and count how often the right passage appears in the top k that you actually pass to the model.

```
hits = sum(gold_id in [c.id for c in retrieve(q, k=5)] for q, gold_id
in pairs)
print(f"recall@5 = {hits}/{len(pairs)} = {hits/len(pairs):.2f}")
```

If that comes back 0.62, then 38% of your failures were decided before the language model saw anything. No amount of fine-tuning the generator recovers a document that was never in the window. You would be training a model to guess better, which is exactly the behaviour you are trying to stop.

The ceiling nobody computes

Recall@k is a ceiling on end-to-end accuracy. A system with recall@5 of 0.62 cannot exceed 62% correct answers, except by luck or by the model already knowing the fact. Teams routinely spend six weeks fine-tuning a generator to close a gap that sits entirely above them.

Write the ceiling on the whiteboard before you plan anything: *our end-to-end accuracy cannot beat 0.62 until retrieval improves.*

The cheap retrieval fixes, in order

Most of the gap closes without training anything.

1. **Chunking.** The single largest lever. Chunks that split a table from its header, or a clause from its definition, make the answer un retrievable and also unusable if retrieved. Try 400-token chunks with 80 tokens of overlap, and chunk on document structure — headings, list items — rather than on a fixed character count.
2. **Hybrid search.** Dense embeddings miss exact strings: part numbers, error codes, surnames. BM25 catches them and misses paraphrase. Run both and fuse the ranks. Reciprocal rank fusion is six lines of code and typically moves recall@5 by 5 to 15 points on real corpora with identifiers in them.
3. **Raise k, then rerank.** Retrieve 50, rerank to 5 with a cross-encoder. The reranker sees the query and the passage together instead of comparing two independently-made vectors, which is why it is better and why it is slower.
4. **Query rewriting.** "What about the second one?" is un retrievable. Rewrite conversational turns into standalone queries before embedding them.

Every one of these is free: `rank_bm25` for lexical search, `sentence-transformers` for embeddings, FAISS or `sqlite-vec` for the index. A corpus of 100,000 chunks fits in memory on a laptop with 8 GB of RAM.

When you do train something, train the retriever

If retrieval is still short after the cheap fixes, fine-tune the embedding model or the reranker — not the generator. The reasons are arithmetic:

- An embedding model is 30M to 500M parameters. A generator you would actually serve is 7B to 70B. The training job is a hundred times smaller.
- Embedding fine-tuning works on pairs: a query and the passage that answers it. You can mine those from click logs or generate them by asking a model to write a question for each passage. You need hundreds, not tens of thousands.
- A retriever improvement raises the ceiling for *every* downstream model, including the next one you switch to. A generator fine-tune is welded to one base model and dies with it.

The practice of this is in the module on fine-tuning beyond a chat model. The decision belongs here: retriever before generator, always, unless the retrieval numbers say retrieval is fine.

The honest limitation

Retrieval cannot fix questions whose answer is not written down anywhere — aggregations ("how many of our contracts expire in Q3"), comparisons across many documents, or reasoning over a whole corpus. Vector search returns the k most similar chunks, and "most similar" has nothing to do with "all the ones you need to count". Those questions want a database query or a document pipeline, and a team that keeps trying to make embeddings answer them ends up fine-tuning in frustration for a capability neither component has.

So the diagnosis has two steps, not one. First: was the evidence there? If not, fix retrieval. Second: is this question even a retrieval question? If not, neither lever helps, and you need a different system shape.

BEFORE YOU MOVE ON

A team measures recall@5 at 0.62 on fifty questions with known gold passages, then fine-tunes their 8B generator on 4,000 examples and finds end-to-end accuracy stuck near 60%. What is the mechanism behind the ceiling?

- The adapter rank was too low to store the corpus, so a rank of 128 or full fine-tuning would push past 62%.
- In 38% of cases the answer passage was never in the context, so the generator had nothing correct to condition on.
- Four thousand examples is too few; the relationship between dataset size and accuracy means 40,000 would break the ceiling.
- Training on retrieved contexts causes the model to over-trust retrieval, which suppresses its own parametric knowledge and caps accuracy.

5 • 9 min

When the answer is one of six labels

THE ONE THING TO KEEP

When the output is one of a fixed set of labels, a 150M encoder fine-tuned on a thousand examples usually beats prompting a large decoder on accuracy, latency and cost — at the price of needing labels and retraining to add a class.

You may not need a generative model at all

A large share of what gets sent to a chat model is classification wearing a costume. Route this ticket to one of six queues. Is this review positive, negative or mixed. Does this sentence contain personal data. Is this transaction suspicious. The output space has six members, and you are paying a 70B decoder to write them out one token at a time.

An encoder model fine-tuned for classification does this better, and the gap is not small.

The arithmetic

Take ticket routing, six queues, 2,000 labelled examples.

- **Prompting a hosted frontier model.** Roughly 700 input tokens (instructions plus the ticket) and 5 output tokens. At typical 2026 pricing for a mid-tier hosted model that is on the order of \$0.20 per thousand tickets, with a p95 latency near one second, and a rate limit that becomes your rate limit.
- **Fine-tuned encoder.** ModernBERT-base is 149M parameters; DeBERTa-v3-base is 184M. Training on 2,000 examples takes about eight minutes on a free Colab or Kaggle T4. Inference runs at roughly 5 to 15 milliseconds per item on a CPU, batched far faster, with no per-call cost and no vendor.

Accuracy usually favours the encoder too, once you have a thousand labels or more. The encoder sees the whole input bidirectionally and is optimised directly for the decision; the decoder is optimised to continue text and is being asked, indirectly, to make a decision.

```
from transformers import AutoModelForSequenceClassification,
AutoTokenizer

model = AutoModelForSequenceClassification.from_pretrained(
    "answerdotai/ModernBERT-base", num_labels=6)
```

That `num_labels=6` replaces the pretraining head with a fresh linear layer of 768×6 weights. Everything below it is the pretrained encoder, and the whole thing is trained — full fine-tuning of a 149M model fits in 6 GB of GPU memory comfortably, so none of the LoRA arithmetic in this course is needed.

What you give up, precisely

This is a real trade, not a free win.

- **You need labels.** A prompt works on day one with zero examples. An encoder needs several hundred before it is competitive and a thousand or two before it is clearly better. If a label costs a minute of somebody's attention, 2,000 labels is about 33 hours of work.
- **No new classes without retraining.** Add a seventh queue and you retrain — cheap, but it is a step. A prompt takes a new class in one line.
- **No explanation.** The classifier emits a label and a probability. If your users need a reason, you either add a second call to a generative model or you accept a score.
- **Brittleness to drift.** An encoder trained on last year's tickets does not know this year's product names. It will not tell you so; it will just be quietly wrong on the new ones. A frontier model degrades more gracefully because it reads the words.

The hybrid that most teams land on

Use the big model to make the labels, and the small model to serve them. Run a frontier model over 3,000 historical tickets, have a person check a sample of 200 of those labels to get an agreement figure, then train the encoder on all 3,000. You have bought a labelled dataset for a few dollars and a serving model that costs nothing per call. Where the encoder's confidence is low — say the top probability is under 0.7 — fall back to the expensive model for that item only. On real routing workloads that fallback fires on 5 to 15% of traffic.

Check your provider's terms before doing this. Several model licences restrict using outputs to train models, and what counts as a competing model is not always clear; that is the subject of a later lesson in this module.

The signal that you are in this situation

You are classifying, not generating, if any of these are true: the output is drawn from a fixed list; you evaluate with accuracy, precision or recall rather than with a human reading the text; the prompt contains the phrase "answer with only one of"; or you are post-processing the model's reply with a regex to extract the label. That regex is the clearest tell there is. It means you built a generative system and then spent effort hiding the generation.

The same logic covers extraction with a fixed schema (token classification, which is a per-token version of the same thing) and scoring (a regression head, one output instead of six). The family is called encoders, the library is `transformers`, and the whole path is free.

BEFORE YOU MOVE ON

A team routes tickets with a hosted model and a regex that pulls the queue name out of the reply. They have 2,500 historical tickets already routed correctly by humans. Why does a fine-tuned 149M encoder tend to be more accurate here, not merely cheaper?

- A. Smaller models overfit less, so the encoder generalises better to unseen tickets than a large model would.
- B. The encoder has read the company's tickets during pretraining, so it starts with knowledge the hosted model lacks.
- C. Its head is trained directly on the six-way decision, while the decoder only decides indirectly.
- D. Encoders are bidirectional and support longer inputs, so they see the whole ticket where a decoder truncates it.

6 · 10 min

Distillation, the case that usually survives

THE ONE THING TO KEEP

Distillation is the fine-tuning case with no cheaper rival, but the student is capped by the teacher on the traffic you sampled, so the filter between them — verifiers, rejection sampling, a hundred pairs read by a person — is where the quality actually comes from.

The one argument that keeps working

Every other reason to fine-tune has a cheaper rival. Format has constrained decoding, knowledge has retrieval, style has few-shot examples. Distillation has none: you want the behaviour of an expensive model at the price of a cheap one, and there is no prompt that makes a 3B model cost what a 3B model costs while answering like a 200B one, other than training it.

The shape is simple. Run the large model over inputs drawn from your real traffic. Keep the good outputs. Train a small model on those pairs. Serve the small model.

The arithmetic that makes it worth it

Suppose you handle two million requests a month, each about 1,200 input tokens and 300 output tokens.

- **Hosted frontier model**, at roughly \$3 per million input and \$15 per million output tokens: about \$7,200 input plus \$9,000 output, so on the order of \$16,000 a month.
- **A 7B model you serve yourself**, on a single rented L4 or A10G at roughly \$0.60 to \$1.00 an hour: about \$430 to \$720 a month, plus the engineering to keep it up. Batched, one such card handles that volume with room to spare for 300-token replies.

The gap is twenty to thirty times, and it is recurring. Against it you have a one-off cost: generating the training data (say 50,000 teacher outputs, which at the same prices is around \$300), the training run itself (single-digit dollars for a LoRA), and the evaluation, which is the part that actually costs you — days of someone's attention.

That is the honest shape of the deal: a four-figure one-off and a headcount cost, against a five-figure monthly saving. It is a good deal at two million requests and a terrible one

at twenty thousand.

Where the ceiling is

A distilled student is bounded by its teacher on the distribution you sampled. Three consequences follow, and each has bitten real teams:

1. **The teacher's mistakes become the student's convictions.** If the teacher is wrong on 6% of your inputs and you do not filter, the student learns those errors as correct answers, and it produces them with less hedging than the teacher did, because it never learned the teacher's uncertainty — only its words.
2. **Off-distribution, the student collapses faster.** The teacher degrades gracefully on inputs unlike anything it saw; the student, having been trained narrowly, often produces confident nonsense. Sample your training inputs from real traffic, including the weird tail, not from a synthetic list of nice questions.
3. **The student does not inherit reasoning it cannot represent.** Distilling a chain-of-thought model into a 0.5B model gives you a model that writes something that looks like reasoning and gets the arithmetic wrong. There is a floor of capability below which imitation is mimicry.

Filtering is the whole craft

The difference between a distillation that works and one that does not is almost entirely the filter between teacher and student.

- **Verifier filtering.** Where correctness is checkable — code that must compile, JSON that must validate, arithmetic you can recompute, a SQL query you can run — keep only outputs that pass. This is the strongest filter in the field and it costs nothing.
- **Rejection sampling.** Generate four teacher outputs per input, score them, keep the best one. Roughly quadruples the generation bill and reliably improves the student.
- **Human spot-check.** Read 100 of the kept pairs yourself. You will find a systematic teacher habit you did not want — a preamble, a hedge, a recurring wrong assumption — and you will find it in ten minutes.

The part people skip

Check the licence. Some model providers' terms restrict using their outputs to train other models, and some open-weight licences pass obligations down to anything trained on the model's outputs. The terms differ by provider and have changed more

than once; several have relaxed, some have tightened. Read the current terms for the specific model you are calling, and if the answer is commercially important, put it in front of a lawyer rather than a forum post. This is also an area where the underlying law — what rights, if any, attach to model outputs at all — is genuinely unsettled and differs between jurisdictions. Nobody can currently tell you the answer with confidence, and anybody who does is telling you their opinion.

The free path

None of this requires a big budget. You can distil from an open-weight teacher you run yourself — a 70B model rented by the hour for a day generates a large dataset — and the licence question becomes much simpler, because permissive open-weight licences such as Apache-2.0 impose no downstream restriction on outputs at all. The student trains on a free Colab or Kaggle T4 if it is small enough, and `llama.cpp` serves it on a laptop CPU afterwards.

BEFORE YOU MOVE ON

A team distils a frontier model into a 7B student using 50,000 teacher outputs with no filtering. The teacher is correct on 94% of their traffic. The student scores 91% but its wrong answers sound far more certain than the teacher's did. What explains the increased confidence?

- A. Fifty thousand examples is too many, and the student memorised the teacher's phrasing rather than generalising from it.
- B. Smaller models are inherently more confident, because fewer parameters produce sharper probability distributions.
- C. The 3-point accuracy drop caused the confidence rise, since accuracy and calibration always move together.
- D. It learned the teacher's wording, but the teacher's internal uncertainty was never in the training signal.

7 · 9 min

The cost of owning a model

THE ONE THING TO KEEP

The training run is the cheapest part of a fine-tune; the real cost is a dataset, an evaluation set, a serving line item and a rebuild every time the base model is superseded, so the thing you must be able to reproduce is the pipeline, not the weights.

The GPU bill is the small number

Ask somebody what a fine-tune costs and they will quote the training run: four dollars, eleven dollars, a free Colab session. That number is real and it is almost irrelevant. What you are buying is not a run, it is an artefact with a maintenance schedule attached.

Here is the fuller bill for one modest production fine-tune, with the parts nobody quotes.

One-off, before you ship

- Dataset construction: 2,000 examples, of which maybe 300 are hand-written or hand-corrected. At five minutes each that is 25 hours, and that is the optimistic version.
- Evaluation set construction: 200 items, labelled, held out, plus a canary set for general capability. Another 15 hours, and it must come first.
- The training runs themselves: not one, but eight to fifteen, because you will change the template, find a masking bug, try two learning rates and two ranks. Call it \$40 of GPU rather than \$4.
- Evaluation of each candidate: cheap per run, expensive in attention.

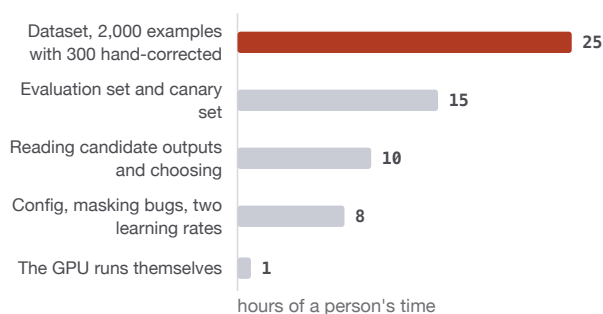
Recurring, for as long as it lives

- Serving. Either a card that idles at 3 a.m. or a per-token bill from a provider that hosts adapters. A single always-on L4 is roughly \$450 a month; a serverless endpoint is cheaper at low volume and has cold starts you will be asked about.
- Monitoring. Someone has to look at the outputs, because a fine-tuned model degrades silently as the input distribution moves.
- Retraining. Which brings us to the real cost.

The base model moves under you

A LoRA adapter is a set of weight deltas for one specific set of base weights. It does not transfer. When the base model you trained against is superseded — and open-weight families have historically shipped a meaningful upgrade every six to nine months — you have three choices, all of them costs:

One modest fine-tune, before it ships



The GPU bill for eight to fifteen runs is about forty dollars. Everything else on this chart is paid in somebody's attention, and the evaluation set has to come first.

1. **Stay on the old base.** Your product's quality floor is now frozen at the level of a model from last year, while the free alternative everyone else prompts keeps improving.
2. **Retrain onto the new base.** Cheap if your dataset and pipeline are intact, expensive if they are a notebook on a laptop belonging to somebody who has left.
3. **Re-evaluate whether you still need the fine-tune at all.** Frequently the honest answer is no: the new base does with a prompt what the old base needed training for. This happens often enough that it should be a scheduled question, not a surprise.

The rule that falls out: keep the dataset, the config, and the evaluation set in version control as first-class artefacts, because the model is the *output* of those things and will be regenerated several times. A team that can rebuild its fine-tune in an afternoon treats a base model upgrade as an improvement. A team that cannot treats it as a threat.

The second, quieter cost

A fine-tuned model is a dependency your organisation does not know how to reason about. Ask three practical questions about the one you are proposing:

- If the person who trained it leaves next month, can somebody else reproduce it from what is written down?
- When a customer complains about one specific output, what is the debugging procedure? For a prompt, you read the prompt. For weights, you look at the dataset and guess.
- When legal asks what the model was trained on, is there an answer?

If the answers are no, nobody knows, and not really, the fine-tune is not an asset, it is a liability that currently performs well.

Where the sums genuinely work

None of this is an argument against fine-tuning. It is an argument for computing the right number. The cases where the full cost is clearly worth it share a shape:

- **High volume.** The saving is per-request, so the case strengthens linearly with traffic and the fixed costs amortise. Below roughly a hundred thousand requests a month it is hard to make the arithmetic work on cost alone.
- **A stable task.** Something whose definition will not change next quarter — transcript formatting, a fixed extraction schema, one language pair.
- **A constraint no prompt satisfies.** Latency budgets under 200 ms, on-device deployment, an air-gapped environment, or data that may not leave the building.

Write the monthly saving and the annual maintenance on the same line before you start. If the saving is \$900 a month and the maintenance is a day of engineering a month, you have not saved anything; you have converted a bill into a chore.

BEFORE YOU MOVE ON

A team's LoRA adapter for Llama-3.1-8B performs well. Nine months later a markedly better base model of the same family ships. Why can they not simply load their adapter onto the new base?

- The adapter encodes deltas against one exact set of base weights, which different weights no longer match.
- Adapters are stored in a format tied to the library version, so an older PEFT checkpoint cannot be read by newer code.
- The licence of the new base forbids loading adapters trained against earlier releases.
- Adapters only work with the tokenizer they were trained with, and a new model release always ships a new tokenizer.

8 · 10 min

Licences, and what you may train on

THE ONE THING TO KEEP

Model licence, data licence and downstream obligations are three separate permissions, the underlying copyright question is genuinely unsettled and differs by country, and per-example provenance recorded at collection time is the only thing that makes the question answerable later.

Three separate permissions, often confused

Before a fine-tune there are three questions, and people routinely answer one and assume the others.

1. **May I use these weights, in this way?** That is the model licence.
2. **May I train on this text?** That is the data licence, plus copyright law, plus whatever you promised your own users.
3. **What obligations attach to the thing I produce?** That is the downstream clause, and it is the one that surprises people.

None of what follows is legal advice, and I am not in a position to give it. What I can do is show you the shape of each question so that you know which one you are in and what to hand a lawyer.

Model licences, in three families

Genuinely permissive. Apache-2.0 and MIT. Qwen's recent releases, Mistral's open models, OLMo, SmolLM, many others. You may fine-tune, merge, serve, sell and relicense the result. Apache-2.0 carries a patent grant and a notice requirement; that is essentially it. If you want the fewest questions, start here.

Custom community licences. Llama's licence is the best-known: broadly free to use, with an acceptable-use policy, a naming requirement for derivatives, and a user-count threshold above which you must ask for separate terms. Gemma has its own terms with a use policy that follows the weights downstream. These are usable and widely used — but they are not open source in the OSI sense, and "open weights" is not a synonym for "do what you like".

Research or non-commercial. Some weights, and rather more datasets, are released for research only. A fine-tune of a non-commercial model is a non-commercial model.

The practical habit: open the `LICENSE` file for the exact checkpoint you downloaded, not the family. Licences differ between sizes and between base and instruct variants more often than you would expect.

Dataset licences are the sharper edge

The famous instruction datasets are not uniformly licensed, and the difference matters:

- Some are CC-BY-SA, which is a share-alike licence — a copyleft obligation that arguably follows anything derived from it.
- Some are explicitly non-commercial, including several popular sets generated from hosted models under terms that were non-commercial at the time.
- Some are Apache-2.0 or CCO and cause nobody any trouble. Databricks' Dolly-15k and the No Robots set were both built specifically to be clean on this point, which is why they remain useful despite being small.

And scraped web text carries whatever rights its authors have, which is where the genuinely unresolved part begins.

The part that is not settled, and will not be settled by this lesson

Whether training a model on copyrighted material without permission is lawful is contested, and the answer currently differs by jurisdiction:

- In the United States the argument runs through fair use, and courts have reached different conclusions on different facts; the question of whether the *source* of the copies was lawful has mattered in at least one ruling as much as the training itself.
- The European Union has a text-and-data-mining exception with an opt-out that rightsholders can exercise, plus transparency obligations on general-purpose model providers under the AI Act.
- The United Kingdom has a research-only TDM exception and has been consulting on change for years without settling it.
- Japan's exception is comparatively broad; several other jurisdictions have nothing specific at all.

So the honest summary is: this is live litigation and live legislation, the answers differ by country and by facts, and anybody who tells you it is simple is describing their preference rather than the law. What you can do is reduce your exposure: prefer sources

with explicit permission, record provenance per example so you can answer the question later, and escalate the commercially significant cases to counsel rather than to a forum.

Outputs of other models

Several hosted providers' terms restrict using their outputs to develop competing models. The exact wording varies, the scope of "competing" is not obvious, and some providers have loosened these terms while others have not. If your distillation plan depends on a particular provider's outputs, read that provider's current terms and get an opinion.

The clean way around it is to distil from an Apache-2.0 open-weight teacher you run yourself. A 70B model rented for a day produces a large dataset, and a permissive licence imposes no downstream restriction on what you do with it.

What to write down, per example

Provenance is cheap at collection time and impossible to reconstruct later. One extra column:

```
{"messages": [...], "source": "support_logs", "licence": "internal",  
  "consent_basis": "tos_v4", "collected": "2026-08-14"}
```

Six months on, when somebody asks whether the model was trained on customer data and under what basis, that column is the difference between an answer and an investigation.

BEFORE YOU MOVE ON

A team fine-tunes an Apache-2.0 base model on a mix of internal tickets and one popular public instruction dataset released under CC-BY-SA. Why does the permissive base licence not settle what they may do with the result?

- A. Apache-2.0 requires derivative models to be released under Apache-2.0, which conflicts with the internal tickets being confidential.
- B. The share-alike obligation on the training data is a separate permission, and it is the one constraining the output.
- C. Fine-tuning voids a base model licence, because the resulting weights are legally a new work with no inherited permissions.
- D. Public instruction datasets may only be used for research, so no commercial fine-tune can include one.

9 • 8 min

Writing the decision down

THE ONE THING TO KEEP

One page written before any training — quantified failure, diagnosis, frozen baseline, budget with a kill criterion, a list of what must not get worse, and a rebuild estimate — is what makes a fine-tune's result believable or its abandonment cheap.

A page, before any GPU

Fine-tuning projects fail quietly. Three weeks in, nobody can remember what the original complaint was, the evaluation set has been edited twice to be kinder, and the question "is this better than the prompt we had in March" has no answer because the March prompt was deleted. The fix is one page written before anything runs.

It has six parts, and each one is a commitment you can be held to.

1. The failure, quantified

Not "the model is not good enough". Something like: *on 200 held-out support tickets, the current prompt produces valid JSON 78% of the time and routes to the correct queue 84% of the time. Target: 97% valid, 92% correct.*

If you cannot write this sentence, you do not have a project yet. You have a feeling, and the most common outcome of training on a feeling is a model that satisfies the feeling and nothing else.

2. The diagnosis

One line naming which of the six causes you found, and the evidence. *Thirty failures sorted: 19 adherence, 6 format, 3 knowledge, 2 capability. The paste test resolved all three knowledge cases, so those go to retrieval.*

3. The baseline that must be beaten

Write down the exact artefact you are competing against, and freeze it. Not "the current prompt" but the prompt file at a specific commit, with a specific model and specific sampling parameters. A moving baseline is why so many fine-tunes look like wins.

List at least two:

- **The cheap baseline.** The best prompt you can write in one more day, including few-shot examples and constrained decoding.
- **The cost-matched baseline.** If the goal is cost, the honest comparison is the *smallest untuned model that a good prompt makes adequate*, not the frontier model you are currently paying for.

4. The budget and the kill criteria

Two numbers and one sentence:

Budget: 10 engineer-days and \$200 of compute. If after 10 days the tuned model has not beaten the frozen prompt baseline by at least 6 points on routing accuracy with a 95% interval that excludes zero, we stop and ship the prompt.

The kill criterion is the part people leave out, and it is the part that makes the record worth writing. Sunk cost is not a hypothesis. Decide the stopping rule while you are still indifferent.

5. What must not get worse

Name the capabilities you are not trying to change but will be blamed for breaking: general instruction following, refusal behaviour on abusive inputs, the other language your product supports, answers in the domain you did not train on. These become the

canary set, and it is written before training precisely so that you cannot choose it afterwards to flatter the result.

6. Ownership and the rebuild test

Who owns this artefact, where the dataset lives, and the answer to one question: *if the base model is superseded next quarter, how many days to rebuild?* If the answer is more than two, fix the pipeline before you train, not after.

The template

```
# Fine-tune decision: ticket routing, 2026-09

Failure      : 84% routing accuracy, 78% valid JSON on eval-v3 (n=200)
Target       : 92% / 97%
Diagnosis    : 19/30 adherence, 6 format, 3 knowledge (→ retrieval), 2
capability
Baseline     : prompts/route.md @ a91f2c, model X, temp 0
Cost baseline: smallest model that passes with a good prompt
Budget       : 10 days, $200. Kill if < +6 pts with CI excluding 0.
Must not break: refusals, Hindi inputs, the escalation path
Owner        : ____ Dataset: git@.../routing-data Rebuild: 1 day
```

Ten lines. It costs twenty minutes and it is the only artefact in the project that is still useful if the fine-tune is abandoned — because a well-written record of *why we decided not to* is worth as much to the next person as a model would have been.

The uncomfortable part

Write the record honestly and a meaningful fraction of proposed fine-tunes will not survive it. That is the point. The most valuable outcome of this module is the projects it stops, and the second most valuable is that the ones which proceed do so with a frozen baseline, a stated interval and a stopping rule — which is to say, with the conditions under which a result could be believed.

BEFORE YOU MOVE ON

Why does the decision record insist the baseline prompt be frozen at a specific commit with fixed sampling parameters, rather than simply described as 'our current prompt'?

- A. Because sampling parameters change the tuned model's behaviour, so they must match between the two systems for the comparison to be valid.
- B. Because git commits are required for reproducibility audits under most model governance frameworks.
- C. Prompts drift during a project, so an unfrozen baseline makes the tuned model's measured gain unattributable.
- D. Because a frozen prompt can be shipped as a fallback if the fine-tune fails, which the kill criterion depends on.

CASE STUDY · MODULE 1

The custom model that was already announced

A twenty-person logistics firm in Pune, whose support inbox handles about nine hundred shipment queries a week.

Ravi ran operations at a freight forwarder in Pune. Customers wrote in asking where a consignment was, why a duty figure had changed, whether a delivery could be rescheduled. A hosted model drafted replies. It was adequate and it was wrong often enough that two people read every draft before it went out.

In March he took a proposal to the two directors: fine-tune a model on eighteen months of the company's own correspondence, so that it would answer "the way we answer". The directors liked it. One of them repeated the phrase "our own AI model" to a customer the following week. A budget of four lakh rupees and six weeks was approved.

The engineer Ravi hired, Sneha, began where the course says to begin. She pulled thirty replies the drafting model had got wrong — real ones, from the archive, not ones anybody invented — and sorted them by what she would have had to change first.

Nineteen of the thirty were knowledge failures. The model did not know the current duty schedule, did not know that one customer had a negotiated rate, did not know which of three warehouses a consignment had been routed through. Six were format: the draft opened with an apology the firm never used, or put the tracking number at the bottom. Three were adherence — the draft ignored two of the four instructions in the prompt. Two were genuine capability failures involving a multi-leg date calculation.

Then she ran the paste test on the nineteen. She put the relevant document into the prompt — the duty circular, the rate card, the routing record — and asked again. Seventeen of the nineteen came back correct.

That is the whole finding, and it was available in a day and a half.

Sneha wrote it up plainly: the firm's problem was that the model could not see its own records. Retrieval over the rate cards, the duty circulars and the shipment database would fix seventeen of thirty failures. Constrained decoding and five examples in the prompt would fix the six format ones. Fine-tuning was the right tool for perhaps three of the thirty, and for those three it would need a dataset, an evaluation set, and somebody to rebuild it every time the base model moved.

She also costed the version Ravi had proposed. Two thousand examples at five minutes each of somebody's careful attention. A held-out set of two hundred, built first. Eight to fifteen training runs rather than one. A card to serve it on, or a provider that hosts

adapters. And a standing question every six to nine months, when the base model was superseded, about whether to retrain or to stay frozen at last year's quality.

Ravi read it and saw the difficulty immediately. The technical answer was clear. The organisational one was not. A director had told a customer the firm was building its own model. Coming back with "we are going to write a better prompt and connect the database" sounded, to anyone who had heard the first version, like a retreat.

The cost of the honest answer was a conversation Ravi did not want to have, and the visible shrinking of a project he had proposed. The cost of the other answer was six weeks, four lakh rupees, and a model that would still not know the duty schedule — because no amount of training installs a fact that changes every quarter, and training on those answers would have taught the system to state duty figures confidently whether or not it knew them.

STOP HERE

Before turning the page: what would you have done, and what would it have cost if you were wrong? Write it down. The point of the next section is to compare.

WHAT HAPPENED

Ravi took the one-page finding to the directors himself, with the thirty sorted failures printed out and the paste-test results beside them. He did not soften it. The retrieval work took nine days, including connecting the shipment database, and moved the first-pass acceptance rate from 61 to 88 per cent. The format work took an afternoon. The three adherence failures were left alone.

The firm never fine-tuned anything. Four months later a new base model arrived that handled the multi-leg date calculation correctly, and the two capability failures disappeared without anybody doing work. The dataset Sneha had started — the seventeen hand-written replies from the paste test — became the few-shot examples in the prompt, and were still there a year on.

One director still describes the system as "our own AI". Ravi has stopped correcting him.

Sneha spent a day and a half on diagnosis before proposing anything. What did that day and a half actually buy, given that the project had already been approved?

It converted an approved project into a tested hypothesis. Before it, the firm had a complaint — the drafts are wrong — and a fashionable remedy. After it, it had thirty labelled failures, a

paste test showing that seventeen of nineteen knowledge failures resolved when the document was in the context window, and therefore a specific prediction about what each remedy would fix. Approval is not evidence, and a budget that has been granted is exactly the situation in which nobody checks.

Suppose the firm had trained on three thousand question-and-answer pairs drawn from the archive, without retrieval. What specific failure would that have produced, and why is it worse than the original problem?

It would have produced a hallucination amplifier. Every one of those examples demonstrates the same lesson: when asked about a shipment or a duty figure, produce a specific, confident answer in the house style. Asked about something outside the training set, the model does exactly that — it invents the figure. The result is worse than the starting point because the replies now sound like the firm's most experienced clerk while being wrong, so the two people reading drafts have lost the cue that used to make them look twice.

The directors had already told a customer the firm was building its own model. How should a decision record have handled that, and what would it have cost?

The record should have carried the failure quantified, the diagnosis with its evidence, a frozen baseline and a stated kill criterion, all written before any work started — so that abandoning the fine-tune was a pre-agreed outcome of a test rather than a reversal by one person. Its cost is twenty minutes and a more uncomfortable first meeting, because writing 'we stop if retrieval closes the gap' in front of the directors forces the smaller version of the project to be visible at the point of approval rather than six weeks later.

Check what stuck

6 questions, one right answer each. This one is for you, not for a record: answers and the reason behind each are at the back. If two questions from the same module go wrong, re-read that module before the exam rather than after.

- 1 You take twenty real failures, hand-write the output you wanted for each, put five into the prompt and run the other fifteen. Thirteen now come out right — but at production volume your prompt has no room for the examples. What has the test told you?
 - A. That you have a real fine-tuning case: the behaviour is promptable and the prompt will not fit.
 - B. That the base model is too small, and a larger one is the answer.
 - C. That fifteen items is far too few to conclude anything, so the exercise should be repeated on at least two hundred before any decision is taken.
 - D. That the model was missing a fact, so retrieval is the lever.

- 2 A team builds three thousand question-and-answer pairs from an internal wiki and trains on them without ever showing the model the wiki. What does the model learn?
 - A. It memorises the wiki, but cannot cite it, so the answers are correct and unsourced.
 - B. Nothing measurable, because three thousand examples cannot move a model of that size.
 - C. That when asked about internal matters it should produce a specific, confident answer — which it then does for things it does not know.
 - D. That internal questions should be declined, because the wiki was absent from the context at training time.

- 3 Your retrieval-augmented system answers 51 per cent of questions correctly. Measuring recall at five on fifty questions with known source passages gives 0.62. What follows?
 - A. The gap between 51 and 62 shows that retrieval is fine and the remaining errors are reasoning failures that only a larger generator will fix.
 - B. End-to-end accuracy cannot exceed about 62 per cent until retrieval improves, so the generator is the wrong place to spend.
 - C. The generator is discarding good evidence, so fine-tuning it on grounded answers is the next step.
 - D. Recall at five is unrelated to end-to-end accuracy, which depends on the generator alone.
- 4 A routing task has six fixed queues and two thousand labelled tickets. What does a fine-tuned 150M encoder buy over prompting a large hosted model, and what does it cost?
 - A. It is more accurate from day one with no labels at all, and costs only a little GPU memory.
 - B. It gives the same accuracy at the same price, and its only advantage is that the weights are yours.
 - C. It is less accurate but much cheaper, so it suits situations where the quality of the routing decision does not matter very much to anybody.
 - D. It is usually more accurate above a thousand labels and costs nothing per call, but it needs labels and a retrain to add a class.
- 5 A student is distilled from a teacher that is wrong on six per cent of your inputs, with no filter between them. What happens to those six per cent?
 - A. The student learns them as correct answers, and states them with less hedging than the teacher did.
 - B. The student rejects them, because a smaller model cannot represent a complicated error.
 - C. They appear only on inputs outside the distribution the teacher was sampled on, which is where every distillation fails first and where filtering could not have helped.
 - D. They are diluted by the ninety-four per cent and have no measurable effect on the student.

- 6 A fine-tune works well, and the base model family ships a meaningful upgrade every six to nine months. Which artefact decides whether that upgrade is an improvement or a threat?
- A. The original training notebook, provided it still runs on the current library versions and nobody has changed the rank or the learning rate in it since.
 - B. The adapter weights, kept in version control with a content hash beside them.
 - C. The dataset, the configuration and the evaluation set, which are what a rebuild needs.
 - D. The serving infrastructure, since a new base model means a new deployment.

MODULE 2

How training actually works

One optimisation step, taken apart: the forward pass, the loss, the gradient, the optimiser state, and the precision the numbers are stored in. Every hyperparameter you will later set moves one of these.

BY THE END OF THIS MODULE YOU CAN

Trace one optimisation step from a batch of tokens to a changed weight, and predict which way loss, memory and stability move when you change the learning rate, the batch size, the precision or the number of epochs

10 · 10 min

One training step, taken apart

THE ONE THING TO KEEP

A training step is forward, cross-entropy against the next token, backward to gradients, optimiser update — and the memory, the masking and every hyperparameter you will set are consequences of those four lines.

Four things happen, in order

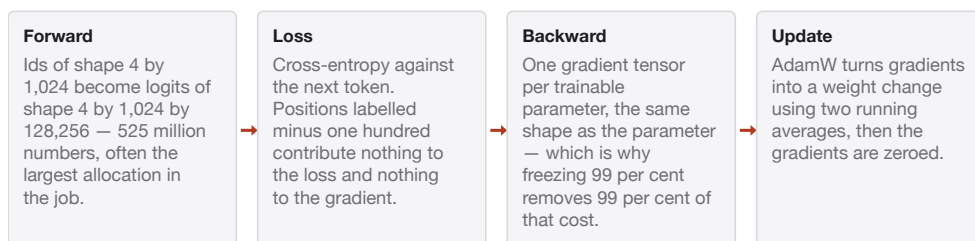
Everything in this course is a variation on one step, repeated a few thousand times. The step has four parts: a forward pass, a loss, a backward pass, and an update. If you can narrate these with real shapes, every hyperparameter later stops being folklore.

Take a concrete batch: 4 sequences of 1,024 tokens each, on an 8B model with a vocabulary of 128,256.

1. Forward

The token ids go in as a tensor of shape `(4, 1024)`. The embedding table turns each id into a vector of 4,096 numbers, giving `(4, 1024, 4096)`. That tensor passes through every transformer block — attention, then a feed-forward network — and comes out the same shape at the end.

One training step, in four parts



A batch of four sequences of 1,024 tokens on an 8B model. Every hyperparameter later in the course is a consequence of one of these four lines.

The last layer projects it to vocabulary size: `(4, 1024, 128256)`. These are the logits, one score per possible next token, at every position.

Notice the size of that tensor. Four times 1,024 times 128,256 is 525 million numbers. In fp32 that is 2.1 GB for a single forward pass, and the backward pass needs room for its gradient too. On many real runs the logits are the largest single allocation in the whole job, bigger than the weights of the layers that produced them. This is why fused or chunked cross-entropy kernels — Liger's, or cut cross-entropy — are worth switching on: they compute the loss without ever materialising the full logits tensor, and they routinely cut peak memory by 30 to 50% on models with large vocabularies.

2. Loss

At every position the model predicted a distribution over the next token, and you know what the next token actually was. Cross-entropy measures the gap: it is the negative log of the probability the model assigned to the correct token.

```
# labels are the inputs shifted left by one; -100 marks positions to ignore
loss = F.cross_entropy(logits[:, :-1].reshape(-1, V),
                        labels[:, 1:].reshape(-1),
                        ignore_index=-100)
```

That shift by one is the whole self-supervised trick: the target for position i is the token at position $i+1$, so a sequence of 1,024 tokens contains 1,023 training examples and you

never had to label anything.

The `ignore_index=-100` is how masking works. Positions labelled -100 contribute nothing to the loss and therefore nothing to the gradient. In instruction tuning that is where you put the prompt tokens, so the model is scored only on what it was supposed to write.

3. Backward

The loss is one number. `loss.backward()` walks the graph in reverse, and for every trainable parameter computes the partial derivative of that number with respect to that parameter. The result is a gradient tensor the same shape as the parameter.

This is why gradients cost as much memory as the weights they belong to, and why freezing 99% of the model — which is what LoRA does — removes 99% of that cost.

The backward pass also needs the activations from the forward pass, because the chain rule refers to them. Those activations were kept in memory throughout. That is the third major consumer of memory, and the reason gradient checkpointing exists: throw the activations away, recompute them during the backward pass, pay roughly 30% more compute to save a great deal of memory.

4. Update

The optimiser turns gradients into a change in the weights. In its simplest form:

```
w = w - learning_rate * gradient
```

AdamW, which you will almost always use, keeps two running averages per parameter and divides by the square root of the second, so a parameter with consistently small gradients still moves. Those two averages are why AdamW costs 8 bytes per trainable parameter in state alone.

Then the gradients are zeroed and the next batch begins.

What to take from this

Three facts do most of the work later:

- **The gradient is computed from one batch.** It is a noisy estimate of the direction you want. Larger batches give a less noisy estimate, which is why batch size and learning rate are linked.

- **Masked positions are invisible.** Not down-weighted, not de-emphasised: absent. If your mask is wrong, the model is learning a different task from the one you think, and the loss curve will look perfectly healthy while it does.
- **Only trainable parameters get gradients and optimiser state.** Everything about parameter-efficient fine-tuning follows from this one sentence.

The honest limitation of this picture: it describes what happens, not why it works. Nobody can tell you in advance which weights will change or what the model will have learned when 3,000 of these steps are done. That is why this course spends a whole module on measuring the result rather than predicting it.

BEFORE YOU MOVE ON

During a run on an 8B model with a 128k vocabulary, peak memory spikes at the very end of the forward pass, before the optimiser step. What is most likely dominating that spike?

- A. The optimiser's two moment buffers, which are allocated lazily on the first forward pass of the run.
- B. The attention key-value cache, which grows with sequence length and is retained across the batch.
- C. The gradients, which are the same size as the weights and are allocated during the forward pass.
- D. The logits tensor — batch times sequence times vocabulary — which can exceed the weights.

11 · 9 min

Reading the loss number

THE ONE THING TO KEEP

Cross-entropy is readable as a probability — $\exp(-\text{loss})$ is the average chance given to the right token — but it is comparable only against your own runs, and because it is a proxy you must select checkpoints on a task metric, since validation loss can rise while the metric you care about improves.

What 1.8 means

Cross-entropy loss is the average negative log-probability the model assigned to the tokens that actually came next. That sentence is invertible, which makes the number readable.

A loss of 1.8 means $\exp(-1.8) \approx 0.165$: on average, the model gave the correct next token about a 16.5% chance. A loss of 0.9 means about 41%. A loss of 0.3 means about 74%.

Perplexity is just the same number exponentiated the other way: $\exp(1.8) \approx 6.0$, read as "the model is about as uncertain as if it were choosing uniformly among six options at each step". Perplexity and loss carry identical information; perplexity is easier to say out loud and easier to misuse.

Why your loss cannot be compared to anyone else's

This is the part that causes the most confusion on forums, and it has four separate causes.

Different tokenizers. Loss is per token. A tokenizer that splits Hindi text into three times as many tokens as another will produce a different loss on the same sentences, because it is being asked an easier question more often. Cross-model loss comparisons are meaningless unless the tokenizers match. Bits per byte is the fix if you genuinely need to compare, and almost nobody does.

Different masking. A run that computes loss over prompt *and* completion is averaging in a lot of easy tokens — the system prompt is identical in every example, so the model predicts it almost perfectly after a few dozen steps. That drags the average down.

Switch on completion-only masking and your loss will *jump*, and nothing has got worse.

Different data. Loss on your customer support transcripts and loss on a coding dataset are not on a common scale. Repetitive, templated text has low loss because it is genuinely predictable.

Different packing. Packed sequences that let attention cross document boundaries make the model predict the first tokens of document B having read document A, which is both harmful and loss-affecting.

So: your loss curve is comparable to *your own* loss curve, on the same data with the same masking and tokenizer. That is all. A number like "we got down to 0.42" tells a stranger nothing.

Loss is not your objective

The deeper point, and the one that decides projects. You are minimising cross-entropy because it is differentiable and cheap, not because it is what you want. What you want is for the model to route tickets correctly, or write a summary a person accepts. Those are not differentiable, so you optimise a proxy and check the real thing separately.

The proxy and the goal come apart in both directions:

- **Loss improves, the task does not.** The model learns your formatting habits, your punctuation, the boilerplate at the top of every answer. All highly predictable, all worth a lot of loss, all irrelevant.
- **Loss worsens, the task improves.** This one surprises people. Validation loss can rise while accuracy on your actual metric keeps climbing, because the model becomes more confident — sharper distributions mean bigger penalties on the items it still gets wrong, even as it gets more items right. If you early-stop on validation loss you will throw away the better model.

The rule: early-stop and select checkpoints on a task metric, not on loss. Loss is for detecting that something is broken, not for deciding what to ship.

The shapes that mean something

Some loss-curve readings are genuinely diagnostic:

- **Flat from step zero.** The learning rate is too small, the mask is removing everything, or nothing is actually trainable. Print the count of trainable parameters; on a misconfigured LoRA it is sometimes zero.

- **A vertical drop to near zero in the first hundred steps.** The answer is visible in the input. Usually the labels were not shifted, or the completion is present in the prompt.
- **NaN.** Almost always fp16 overflow, a corrupted example, or a learning rate an order of magnitude too high.
- **A sawtooth with a period of exactly one epoch.** The model is memorising; each epoch it recognises examples it has seen.

One number to log alongside it

Log token-level accuracy — the fraction of masked positions where the argmax token is the correct one. It is free to compute, it is on a scale you can interpret without exponentiating anything, and it is far less sensitive to a handful of very surprising tokens dominating the average. When loss and token accuracy disagree, the disagreement is itself informative: loss dominated by a few outliers is often a data problem hiding in one or two examples.

BEFORE YOU MOVE ON

A team switches on completion-only loss masking and their training loss immediately jumps from 0.6 to 1.9. What has happened?

- The easy, predictable prompt tokens were holding the average down; removing them leaves the harder ones.
- Masked positions still contribute a small penalty, so removing them from the average leaves a higher residual loss.
- The masking is misconfigured and is now scoring the model on the prompt tokens instead of the completion.
- Gradient magnitudes rise when fewer positions contribute, which inflates the reported loss through the optimiser's normalisation.

Optimisers, and what their state costs

THE ONE THING TO KEEP

AdamW's two per-parameter moment buffers cost 8 bytes per trainable parameter, which is why full fine-tuning is memory-bound and why LoRA makes the optimiser irrelevant; 8-bit and paged variants buy that memory back when you must train everything.

Why not plain gradient descent

The simplest update is `w -= lr * grad`. It works, badly, on deep networks: different parameters have gradients of wildly different magnitudes, so any single learning rate is too large for some and too small for others.

Adam fixes this per parameter, by keeping two running averages:

- **m**, an exponential moving average of the gradient itself — momentum. It smooths the noise from batch to batch.
- **v**, an exponential moving average of the gradient *squared* — a running estimate of that parameter's typical magnitude.

The update divides by the square root of v, so a parameter whose gradients are consistently tiny still takes a meaningful step, and one whose gradients are consistently huge takes a small one. This is what makes transformer training work at all with one global learning rate.

$$\begin{aligned} m &= \beta_1 \cdot m + (1 - \beta_1) \cdot g \\ v &= \beta_2 \cdot v + (1 - \beta_2) \cdot g^2 \\ w &= w - lr \cdot \hat{m} / (\sqrt{\hat{v}} + \epsilon) \end{aligned}$$

The defaults $\beta_1 = 0.9$ and $\beta_2 = 0.999$ mean m is an average over roughly the last ten steps and v over roughly the last thousand. You will essentially never change these; the one exception is that very short runs sometimes benefit from $\beta_2 = 0.95$, because 0.999 takes longer to warm up than the whole run lasts.

The memory bill

Two averages, one per parameter, each stored in fp32: **8 bytes per trainable parameter**.

For a 7B model trained fully, that is 56 GB of optimiser state alone, on top of 28 GB of fp32 weights and 28 GB of gradients. This is the single biggest reason full fine-tuning of a 7B model does not fit on an 80 GB card without sharding, and it is the whole motivation for the parameter-efficient methods in the next module.

For a LoRA with 20M trainable parameters, the same formula gives 160 MB. The optimiser stops being a consideration entirely.

AdamW, and why the W

Weight decay pulls parameters towards zero. In the original Adam it was implemented by adding $wd \cdot w$ to the gradient — which then went through the same division by $\sqrt{v}$, so the amount of decay a parameter received depended on its gradient history. AdamW *decouples* it: decay is applied directly to the weights, outside the adaptive step.

$$w = w - lr \cdot \hat{m} / (\sqrt{\hat{v}} + \epsilon) - lr \cdot wd \cdot w$$

This is not a subtlety; it changes results measurably, and AdamW is the default everywhere now. A typical value is 0.01, and for LoRA it is often set to 0 because the trained matrices are small and initialised so that the adapter starts as a no-op anyway.

The cheaper optimisers, and when they are free wins

- **8-bit AdamW** (`bitsandbytes`, `optim="adamw_bnb_8bit"`). Quantises the two states to 8 bits with block-wise scaling. 8 bytes per parameter becomes 2. On full fine-tunes the saving is large and the quality difference is, in most reported comparisons, within noise. Cheap to try, easy to check.
- **Paged AdamW** (`adamw_bnb_8bit` with paging, or `paged_adamw_32bit`). Uses CUDA unified memory so optimiser state can spill to CPU RAM during a spike instead of raising an out-of-memory error. This is what makes long-sequence QLoRA runs survive their worst batch.
- **Adafactor**. Factorises v into row and column statistics instead of storing it per parameter, so state falls to roughly 4 bytes per parameter or less. Was standard in the T5 era; still useful when memory is desperate.
- **SGD with momentum**. One state tensor instead of two, 4 bytes per parameter. It genuinely can train transformers, but it needs more careful tuning and more steps,

so it is a memory-of-last-resort choice rather than a default.

The practical rule

If you are doing LoRA or QLoRA, use AdamW and never think about this again — 160 MB of state is not your problem. If you are doing full fine-tuning, the optimiser is the first place to look for memory, and 8-bit AdamW is the cheapest 30 GB you will ever recover on a 7B model.

One honest caveat: the quality claims for 8-bit and factorised optimisers come from benchmarks on pretraining and on standard fine-tuning sets. They are well supported but they are not a guarantee about your task. If you switch optimisers to fit in memory, run your evaluation against a smaller run with full-precision AdamW once, so you know what the swap cost you rather than assuming it cost nothing.

BEFORE YOU MOVE ON

A team switches a full fine-tune of a 7B model from AdamW to 8-bit AdamW and frees roughly 42 GB. Where did that memory come from?

- A. The gradients, which are quantised to 8 bits alongside the optimiser state.
- B. The two moment buffers, which fall from 8 bytes per parameter to 2.
- C. The activations, since a smaller optimiser allows a larger batch and better memory reuse.
- D. The model weights, which are stored in the optimiser's own quantised copy rather than in fp32.

Learning rate, warmup and the schedule

THE ONE THING TO KEEP

LoRA needs a learning rate ten to twenty times higher than full fine-tuning because its adapter starts at exactly zero, warmup exists to stop AdamW taking a large step from a one-sample variance estimate, and a cosine schedule must be set for the run you actually intend to finish.

The one hyperparameter that matters most

If you may tune exactly one number, tune this one. Everything else on a fine-tuning config changes results by a few percent. The learning rate changes whether the run works at all.

The values that people actually use cluster tightly, and the clustering has a reason:

- **Full fine-tuning of a pretrained LLM: $1e-5$ to $2e-5$.** You are nudging weights that already encode a great deal. Larger values destroy pretrained behaviour quickly — this is the fastest route to catastrophic forgetting there is.
- **LoRA: $1e-4$ to $3e-4$, commonly $2e-4$.** Ten to twenty times higher, and this is not arbitrary. The adapter starts at exactly zero (B is initialised to zeros, so $BA = 0$ and the model is unchanged). It has to travel from nothing to something useful in a few thousand steps, and it is a small number of parameters doing it.
- **Embedding and classifier heads: $1e-5$ to $5e-5$,** often with a higher rate for the new head than for the pretrained body.

If a LoRA run is not learning, multiply the learning rate by three before you touch anything else. If it is producing gibberish or NaN, divide by three.

Warmup, and what it is protecting against

Starting at the full learning rate on step one is a bad idea for a mechanical reason: AdamW's v buffer has seen one gradient, so its estimate of each parameter's typical magnitude is based on a single noisy sample. Dividing by the square root of a bad estimate produces a badly sized step, and an early badly sized step can knock the model somewhere the rest of the run spends its time recovering from.

Warmup ramps the learning rate from near zero to the target over the first few percent of steps, by which time v is a reasonable estimate.

- Typical: `warmup_ratio=0.03`, or a flat `warmup_steps=10` on short runs.
- On runs under 200 steps, ten steps of warmup is 5% of the run and is enough.
- Skipping warmup on a full fine-tune at $2e-5$ is usually survivable. Skipping it on a large-batch run, or with fp16, is how you get a loss spike at step 4.

The decay schedule

After warmup the rate comes down. Three choices cover practically everything:

Cosine. Smooth decay from the peak to near zero over the run. The default for good reason: it spends most of the run at a usefully high rate and finishes with small, careful steps that settle the weights. Use it when you know in advance how many steps you will take.

Linear. Straight line down. Marginally worse than cosine in most comparisons, marginally more predictable. Fine.

Constant, or constant with warmup. Necessary when you do not know the length of the run — a job you might extend, or a continued-pretraining stream. It also makes checkpoints comparable across a run, because a cosine-decayed checkpoint at step 500 of a planned 1,000 is a model caught mid-schedule, not a finished one.

The cosine trap worth knowing: if you set a schedule for 1,000 steps and stop at 400, you have not trained a smaller version of the same thing. You have stopped a model whose learning rate was still high, and it will typically be worse than a model trained with a schedule that was set for 400 steps in the first place. **Set the schedule length to the run you intend to finish.**

Learning rate and batch size move together

A larger batch gives a less noisy gradient, which supports a larger step. The common rules of thumb are to scale the learning rate linearly with batch size, or with its square root; neither is exactly right and both are better than ignoring the relationship.

The practical version: if you double your effective batch size because you added gradient accumulation to fit a longer sequence, and your loss curve goes flatter and slower, raise the learning rate by around $1.4\times$ before concluding the change hurt.

Finding the value cheaply

You do not need a sweep. Run three short jobs — 150 steps each, one at $1e-4$, one at $2e-4$, one at $5e-4$ — and look at the training loss curve.

- Too low: loss falls slowly and roughly linearly, no elbow.
- About right: a clear fast drop in the first 50 steps, then a smoother decline.
- Too high: an early drop then a rise, visible spikes, or NaN.

On a free T4, three 150-step LoRA runs on a small model is under half an hour. That is the cheapest hour you will spend on the whole project, and it is the one most often skipped in favour of copying a config from a blog post trained on a different dataset at a different batch size.

BEFORE YOU MOVE ON

A team plans a 1,000-step LoRA run with a cosine schedule, then stops it at step 400 to save money and ships that checkpoint. It performs worse than a colleague's run that was configured for 400 steps from the start. Why?

- The stopped run saw the same data in a different order, and cosine schedules are sensitive to data ordering.
- Checkpoints saved mid-schedule omit the optimiser state, so the weights are incompletely updated when reloaded.
- At step 400 of a 1,000-step cosine the rate is still high, so the weights never got the settling steps.
- Stopping early truncates warmup, so the model never reached its intended peak learning rate.

Batch size, accumulation and the token count

THE ONE THING TO KEEP

Effective batch is per-device times accumulation times devices, the unit that actually governs gradient noise is tokens per step rather than examples, and loss must be normalised by real token counts across an accumulation window or long and short examples get equal weight.

Three numbers that multiply

What the optimiser sees is the **effective batch size**:

```
effective = per_device_batch_size × gradient_accumulation_steps ×  
num_devices
```

A config with `per_device_train_batch_size=2`, `gradient_accumulation_steps=8` on one GPU has an effective batch of 16. The gradients from eight forward-backward passes are summed before a single optimiser step. Memory is paid for two sequences at a time; the update behaves like sixteen.

This is the standard trick for training on a small card, and it is close to free — the only cost is that you compute eight forward passes before making any progress, so wall-clock time per optimiser step goes up while total compute stays roughly the same.

The number that actually matters is tokens

Examples are a bad unit. An effective batch of 16 chat exchanges could be 2,000 tokens or 30,000 depending on how long they are. Gradient noise, and therefore the right learning rate, tracks tokens far better than examples.

Log tokens per optimiser step. Most training frameworks do not print it, and it takes one line to compute. A typical small fine-tune sits somewhere between 8,000 and 64,000 tokens per step; if yours is 900, your gradients are extremely noisy and no learning rate will make the run stable.

The accumulation bug that shipped for years

This is the detail no tutorial mentions, and it caught a great many people.

If each micro-batch computes a *mean* loss over its own tokens, and you then average those means across accumulation steps, you have not reproduced a larger batch — you have given every micro-batch equal weight regardless of how many real tokens it contained. A micro-batch with 40 unmasked tokens and one with 900 count the same. Since masked instruction data has extremely variable completion lengths, the discrepancy is real.

Hugging Face's `Trainer` and TRL had exactly this behaviour, and it was fixed in late 2024 (transformers 4.46 and the corresponding TRL release) by normalising over the total number of unmasked tokens in the accumulation window. Two practical consequences:

- Configs and results from before that fix are not directly comparable to results after it, especially where `gradient_accumulation_steps` was large.
- If you write your own training loop, normalise by the total token count across the accumulation window, not by the number of micro-batches. It is three lines and it is correct:

```
total_tokens = sum(batch_token_counts)
for batch, n in zip(batches, batch_token_counts):
    loss = model(**batch).loss * n / total_tokens    # weight by real
tokens
    loss.backward()
optimizer.step()
```

Packing: the other way to fix the token count

Rather than padding every sequence to the length of the longest in the batch — which on instruction data can mean half your compute is spent on padding tokens — packing concatenates short examples until the buffer is full.

The gain is substantial: on a dataset whose examples average 300 tokens with a 2,048-token window, packing can cut training time by more than half, because you stop multiplying matrices by padding.

The hazard is attention leaking across the boundary. Without a correct block-diagonal attention mask, the model reads the end of example A while predicting the start of example B, and learns a transition that does not exist. Modern implementations handle this with position ids reset per document and a FlashAttention variable-length path;

`packing=True` in TRL with a recent version does the right thing, but this is worth verifying rather than assuming, because the failure is invisible in the loss curve.

If you cannot verify it, use **length-grouped batching** instead

(`group_by_length=True`). It sorts similar-length examples into the same batch so padding is minimal, with no cross-contamination risk at all. Most of the speed, none of the subtlety.

Choosing the numbers in practice

A workable procedure on a single card:

1. Set `per_device_train_batch_size` to the largest value that does not run out of memory at your maximum sequence length. Find it by bisection in about four attempts.
2. Set `gradient_accumulation_steps` so that effective batch lands between 16 and 64 for a LoRA on instruction data.
3. Log tokens per step. If it is under 4,000, raise accumulation rather than the learning rate.
4. Switch on gradient checkpointing if step 1 forces you below a per-device batch of 1 at the sequence length you need. It costs roughly 30% more time per step and often lets you double the batch.

And one thing not to do: do not chase a large effective batch on a small dataset. With 800 examples and an effective batch of 64, one epoch is twelve optimiser steps. Twelve steps will not move an adapter anywhere useful, no matter how clean each gradient is.

BEFORE YOU MOVE ON

A custom training loop averages each micro-batch's mean loss across 8 accumulation steps. The data is instruction-tuned with completions ranging from 20 to 1,200 tokens. What does this do to the gradient, compared with a true batch of the same size?

- A. Nothing; averaging means of equal-sized micro-batches is mathematically identical to a single larger batch.
- B. It makes the run non-deterministic, because floating-point addition order varies across accumulation steps.
- C. It doubles the effective learning rate, since summing eight gradients without dividing inflates the step.
- D. It over-weights short completions, since a micro-batch's contribution ignores its token count.

15 · 9 min

Precision, and why bf16 ended a class of bugs

THE ONE THING TO KEEP

bf16 keeps fp32's exponent range and sacrifices mantissa bits, which removes the underflow that forced fp16 to use loss scaling — and since the free T4 has no bf16, free-tier runs still live with the scaler and its skipped steps.

Three ways to store a number

A floating-point number is a sign bit, some exponent bits and some mantissa bits. The exponent sets the *range* — how large and how small a value can be. The mantissa sets the *precision* — how finely values are distinguished within that range.

- **fp32**: 1 sign, 8 exponent, 23 mantissa. Range to about 3×10^{38} , roughly seven decimal digits of precision.
- **fp16**: 1, 5, 10. Maximum value **65,504**. Smallest normal value about 6×10^{-5} . Three decimal digits.
- **bf16**: 1, 8, 7. The *same exponent range as fp32*, with only three mantissa bits' worth of precision more than a rounding error.

That comparison is the whole story. bf16 gave up precision to keep range, and for neural network training range turned out to matter enormously and precision hardly at all.

What went wrong with fp16

Gradients in a transformer are small. Many of them are around 10^{-7} or smaller, which is below fp16's smallest normal value, so they round to zero. A gradient that rounds to zero is a parameter that does not learn.

The workaround is **loss scaling**: multiply the loss by a large constant, say 2^{16} , before the backward pass. Every gradient is scaled up by the same factor, lifting the small ones into representable range. Then divide the gradients by the same constant before the optimiser step.

This works, and it introduces a new failure. If the scale is too large, some gradients exceed 65,504 and become infinity. Frameworks handle this with *dynamic* loss scaling: start high, and when an overflow is detected, skip that optimiser step and halve the scale. So a healthy fp16 run periodically throws away steps. If you see log lines about

the gradient scaler reducing its scale, that is normal. If you see it every few steps, your run is spending a meaningful fraction of its compute on nothing.

What bf16 fixed

With the exponent range of fp32, no gradient underflows and none overflows. No loss scaling, no skipped steps, no NaN at step 300 of a run you paid for. The lost mantissa bits do not matter because the optimiser accumulates in fp32 anyway.

This is why a great deal of pre-2023 training advice about NaN, loss spikes and scaler tuning is simply obsolete: it was fp16 advice.

The hardware line that decides for you

bf16 requires Ampere or newer. Concretely:

- **Has bf16:** A100, H100, L4, L40S, A10G, RTX 30-series and 40-series, and newer.
- **Does not:** the T4, which is Turing. And the T4 is exactly what free Google Colab and Kaggle give you.

So on a free tier you are on fp16, with loss scaling, and a divergence at $2e-4$ that would have been fine on an A100 is a real possibility. The mitigations on a T4 are: keep the learning rate at the low end, do not skip warmup, and if you see repeated scaler reductions, drop the learning rate rather than investigating anything cleverer.

Check it in one line before you configure anything:

```
import torch
print(torch.cuda.get_device_name(0), torch.cuda.is_bf16_supported())
```

Mixed precision means mixed

"bf16 training" does not mean everything is bf16. The standard arrangement keeps an **fp32 master copy** of the weights and does the optimiser update in fp32, while forward and backward run in bf16. The reason is accumulation: adding a tiny update to a large weight in bf16 loses the update entirely, because the result rounds back to the original value.

This is also where the 16-bytes-per-parameter figure for full fine-tuning comes from: 4 bytes of fp32 master weights, 4 of fp32 gradients, 8 of AdamW state — with the bf16 working copies on top.

For QLoRA the arrangement differs again: base weights are 4-bit, the compute dtype for matrix multiplications is bf16, and only the small adapter is kept and updated in higher precision. That is three precisions in one model, which is normal and not a sign of anything wrong.

fp8 and lower

H100 and newer support fp8 for training, with per-tensor scaling to manage the tiny range. It is real and it is used at scale, but the tooling is still specialist and the gains show up on large pretraining runs rather than on a LoRA over 3,000 examples. For the work this course describes, bf16 where available and fp16 where not is the whole decision.

One caveat worth keeping: a model *trained* in bf16 and then *served* in fp16 can behave slightly differently, because a value near the edge of fp16's range that was fine during training now saturates. It is rare, and when it happens it looks like a serving bug rather than a precision bug, which is why it takes so long to find.

BEFORE YOU MOVE ON

A fine-tune on a free Colab T4 runs fine for 200 steps, then the logs fill with the gradient scaler halving its scale, and progress stalls. What is the mechanism?

- A. The T4 has no bf16, so fp16 loss scaling is in use and overflows are skipping optimiser steps.
- B. The T4's 16 GB of memory has filled, and the scaler is reducing precision to free space.
- C. The cosine schedule has decayed the learning rate far enough that gradients now underflow in fp16.
- D. bf16 was requested but silently fell back to fp32, so memory pressure is forcing the framework to rescale.

Where the memory actually goes

THE ONE THING TO KEEP

Memory is weights, gradients, optimiser state, activations and logits — and at long context the quadratic attention matrix and the vocabulary-sized logits tensor are usually the two surprises, both removable with flash attention and a fused cross-entropy before you touch batch size.

Five consumers, in order of surprise

When a run reports `CUDA out of memory`, people reach for a smaller batch size. Sometimes that is right. Often the batch was not the problem, and knowing which of five consumers is eating the card saves an afternoon.

1. Weights. Parameters times bytes per parameter. A 7B model is 14 GB in bf16, 28 GB in fp32, about 3.9 GB in 4-bit NF4. Fixed for the whole run, easy to compute, rarely the surprise.

2. Gradients. Same size as the *trainable* weights, in the gradient dtype. Full fine-tuning a 7B in bf16: another 14 GB. LoRA with 20M trainable parameters: 80 MB. This single line is why parameter-efficient methods exist.

3. Optimiser state. 8 bytes per trainable parameter for AdamW. 56 GB for a full 7B; 160 MB for the LoRA.

4. Activations. The intermediate tensors kept from the forward pass so the backward pass can use them. This is the one people underestimate, and it scales with things they change casually:

$$\text{activations} \approx \text{batch} \times \text{seq_len} \times \text{hidden} \times \text{layers} \times \text{bytes} \times c$$

where c is a small constant — roughly 10 to 20 depending on architecture and what is fused. For a 7B model (32 layers, hidden 4,096) at batch 4 and sequence 2,048 in bf16, that is on the order of 10 to 20 GB. It grows linearly in both batch and sequence length, so doubling your context doubles it.

5. Attention, if it is not flash attention. A naive attention implementation materialises the full score matrix: $\text{batch} \times \text{heads} \times \text{seq}^2 \times \text{bytes}$. At sequence 8,192 with

32 heads in bf16 that is about 4.3 GB **per layer**, transiently. This is the quadratic term everybody has heard of and few have costed. FlashAttention computes attention in tiles without ever writing the matrix, turning that term from quadratic to linear in memory. Switch it on (`attn_implementation="flash_attention_2"` , or the PyTorch SDPA backend) before you touch anything else at long context.

Gradient checkpointing, and what it actually trades

Gradient checkpointing throws away most activations after the forward pass and re-computes them during the backward pass from a small number of saved boundaries. Memory for activations drops from roughly linear in depth to roughly the square root of it; compute rises by about 30%.

It is nearly always worth it on a memory-bound run, because the memory you free usually buys a larger batch, and a larger batch recovers much of the lost throughput.

One trap: with LoRA and gradient checkpointing together, you may need `model.enable_input_require_grads()` . Without it the checkpointed segments see no input requiring grad, the graph is not built, and you get a run where the loss never changes and nothing obviously errors. This has cost a great many people an evening.

The logits, again

From the first lesson of this module: `batch × seq × vocab` can exceed the weights. At batch 4, sequence 2,048 and a 128k vocabulary that is 1 billion values — 2 GB in bf16, 4 GB in fp32, and typically both exist at once during the loss computation.

Fused cross-entropy kernels avoid materialising it. Liger Kernel and cut cross-entropy are the common options, both free and both a one-line change in TRL. On large-vocabulary models this is frequently the single largest saving available.

The diagnostic sequence

When you are out of memory, work down this list rather than guessing:

1. Print `torch.cuda.max_memory_allocated()` and compare against the card. If allocated is far below capacity, you have **fragmentation**, not a shortage — set `PYTORCH_CUDA_ALLOC_CONF=expandable_segments:True` and try again.
2. Is flash attention on? At sequence length over 4,000 this is usually the answer.
3. Is gradient checkpointing on?
4. Is the cross-entropy fused?

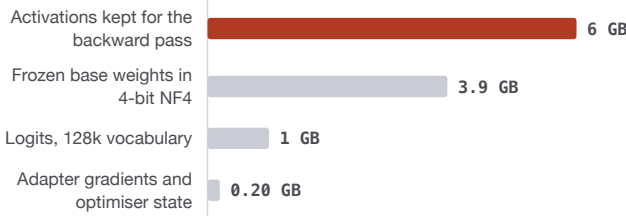
- 5. Only now reduce batch size, raising gradient accumulation to keep the effective batch constant.
- 6. Then reduce sequence length — and measure how many of your examples get truncated by the new limit, because a max length that cuts 30% of your completions is a data problem, not a memory fix.

The estimate before you rent

Do the arithmetic on paper first. For QLoRA on a 7B at sequence 2,048, batch 2: 3.9 GB weights, negligible gradients and optimiser state, perhaps 5 to 8 GB of activations with checkpointing off, 1 GB of logits. Around 11 GB, which fits a 16 GB T4 with room and comfortably fits a 24 GB card.

If the arithmetic says 70 GB and you were planning to use a 24 GB card, you learned that for free.

An 11 GB QLoRA run, by consumer



A 7B base at sequence 2,048, batch 2. Lowering the batch size touches one row of this chart, and it is often not the largest one.

BEFORE YOU MOVE ON

A QLoRA run on a 24 GB card fails with out-of-memory at sequence length 8,192 but succeeds at 2,048. The 4-bit base weights are under 4 GB and the adapter is tiny. Which term grew fastest, and why?

- A. The optimiser state, which is allocated per position in the sequence during long-context training.
- B. The attention score matrix, which is quadratic in sequence length without flash attention.
- C. The 4-bit base weights, which must be dequantised to bf16 for long sequences and so quadruple in size.
- D. The gradients, which are proportional to the number of tokens processed and so grew four-fold.

17 · 9 min

Epochs, overfitting and checkpoint selection

THE ONE THING TO KEEP

The step down in training loss at each epoch boundary is memorisation, not learning; select checkpoints on a task metric with a canary set beside it rather than on validation loss, because the best checkpoint is often the last one before general capability starts to slip.

One to three, and why

The standard advice for supervised fine-tuning is one to three epochs, and unusually for standard advice it is well grounded.

An epoch is one pass over your dataset. On the second pass the model sees examples it has seen before, and a model with billions of parameters and a few thousand examples is entirely capable of memorising them. What it memorises is not the skill you wanted — it is the specific answers, including their specific wording and the specific mistakes in your annotation.

The signature is unmistakable once you have seen it: at each epoch boundary the training loss takes a **step down**, a sharp discontinuity rather than a smooth continuation. That step is the model recognising the data. It is not learning anything at that moment; it is recalling.

But small datasets need steps

There is a real tension here. With 800 examples at an effective batch of 16, one epoch is 50 optimiser steps. Fifty steps will not move a freshly-initialised adapter anywhere useful.

So the honest framing is not "how many epochs" but "how many optimiser steps do I need, and can I get them without repeating data too often". Roughly:

- Under 500 steps total, an adapter is usually undertrained.
- 500 to 2,000 steps is the normal range for a task LoRA.
- If reaching 500 steps requires six epochs over 600 examples, the problem is that you have 600 examples. Either collect more, or accept a smaller effective batch to

get more steps from the same data — the latter is a genuinely reasonable trade at this scale.

LoRA tolerates repetition better than full fine-tuning does, because the low-rank constraint limits how precisely it can memorise. That is a real advantage of the method and not usually mentioned as one.

Checkpoint selection, done properly

Save a checkpoint every few hundred steps. At the end, you have five or ten candidates, and the question is which to ship.

Do **not** pick on validation loss. As the previous lesson noted, validation loss can rise while your task metric improves, because a more confident model is penalised harder on the items it still gets wrong. If you early-stop on loss you will systematically ship earlier, blander checkpoints.

Pick on your task metric, evaluated on a held-out set, with the canary set alongside it:

```
for ckpt in checkpoints:
    task = evaluate(ckpt, held_out)      # what you are trying to
    improve
    canary = evaluate(ckpt, general_set) # what must not break
    print(ckpt, round(task, 3), round(canary, 3))
```

The shape you usually see: task climbs quickly then plateaus, canary is flat then starts falling. The checkpoint to ship is the last one before the canary moves, which is frequently *not* the last checkpoint and frequently not the best task score either.

Reading the curve honestly

A few readings that are actually diagnostic, beyond the epoch step:

- **Train and validation loss both falling, gap steady.** Healthy. Keep going.
- **Train falling, validation flat.** Learning your data's idiosyncrasies. Not yet harmful, but the marginal epoch is buying little.
- **Train falling, validation rising, task metric rising.** Not necessarily overfitting in the harmful sense. Check the canary before you panic.
- **Train falling, validation rising, task metric flat or falling.** Overfitting, straightforwardly. Stop, and go back a checkpoint or two.
- **Loss lower on validation than on training.** Usually means your validation set is easier, or — much more often — that some validation examples also appear in

training. Check for leakage before you celebrate.

The thing that is not in the curve

Catastrophic forgetting does not show up on either loss curve. Both are computed on data from your task distribution, and the capability you have destroyed is somewhere else entirely — the other language your product supports, the ability to say "I do not know", the refusal that used to work.

This is the argument for running the canary set at every checkpoint rather than at the end. It costs a few minutes per checkpoint, and it is the only instrument that points at the damage while you can still choose a different checkpoint. A run that is evaluated only at the end gives you one artefact and no alternatives.

Task metric and canary, across ten checkpoints



The checkpoint to ship is the last one before the canary moves — the fifth here, which is neither the last checkpoint nor the best task score. Neither loss curve would have shown you this.

BEFORE YOU MOVE ON

A run shows training loss falling steadily, validation loss rising from step 600, and routing accuracy on a held-out set still climbing at step 1,000. What should the team do?

- A. Stop at step 600, since rising validation loss is the definition of overfitting and later checkpoints cannot be trusted.
- B. Lower the learning rate so validation loss resumes falling, then continue to step 1,000.
- C. Continue, but check a canary set at each checkpoint, since rising confidence raises loss while the metric improves.
- D. Add more validation data, since a rising validation loss with an improving metric indicates the validation set is too small to be reliable.

Regularisation that helps, and the kind that does not

THE ONE THING TO KEEP

Fewer epochs, more data and a replay mixture regularise a fine-tune far more than any config knob, gradient clipping at 1.0 is cheap insurance against one pathological example, and label smoothing flattens the output distribution in a way that harms generation even though it helps classifiers.

Most of it is data, not a knob

In classical machine learning, regularisation is a set of techniques for stopping a model memorising a small dataset. In LLM fine-tuning, the most effective regularisers are not techniques at all: fewer epochs, more data, and a mixture that includes examples outside your narrow task. The knobs below are real, but they are second-order compared to those three.

Say it plainly, because a lot of time is wasted here: if your model is overfitting on 700 examples, `lora_dropout` will not save you. Another 2,000 examples will.

The knobs that are worth setting

LoRA dropout (0.05 to 0.1). Randomly zeroes elements of the adapter's input during training. Cheap, standard, mildly helpful on small datasets. At dataset sizes above roughly 10,000 examples people often set it to 0 and lose nothing. Set 0.05 and stop thinking about it.

Weight decay (0 to 0.01). For full fine-tuning, 0.01 is a sensible default. For LoRA it is commonly 0, and the reasoning is that the adapter begins at exactly zero, so decay towards zero is pulling against the only learning happening. Some practitioners use a small value anyway; the effect is minor either way.

Max gradient norm (1.0). Gradient clipping. Not really regularisation — it is insurance. A single pathological example produces an enormous gradient, and without clipping that one step can wreck a model that was training fine. Leave it at 1.0. It costs nothing and it prevents a failure mode that looks like a mystery.

Early stopping — on a task metric. Covered in the previous lesson, and the single most effective item on this list.

The knob that hurts generation

Label smoothing. Instead of a one-hot target, put 0.9 on the correct token and spread 0.1 across the rest. It improves calibration in classifiers and is common practice there.

For generative fine-tuning it is usually a mistake. You are deliberately teaching the model that the correct next token is only 90% correct, which flattens the output distribution; at generation time a flatter distribution means more sampling from the tail, which reads as vagueness and drift. It also makes your loss numbers incomparable with runs that did not use it. Leave it at 0 for text generation, use it if you like on a classification head.

NEFTune, and how to read a claimed win

NEFTune adds uniform random noise to the embedding vectors during training — a few lines of code, one hyperparameter (`neftune_noise_alpha`, typically 5), and it is built into TRL. The published results showed large gains on AlpacaEval-style benchmarks.

It is worth trying, and it is also a good case study in reading a result. Those benchmarks are judged by a model comparing two responses, and such judges are known to prefer longer, more elaborate answers. NEFTune's outputs tend to be longer. How much of the measured gain is better quality and how much is length preference is not fully settled, and the honest answer is that some of both is likely.

So: try it, and evaluate it on *your* metric rather than on the benchmark that made it famous. If your task is extraction into a fixed schema, a technique that makes answers more elaborate is not obviously your friend.

The regularisation that matters most

Mixing general data into your training set. If your 3,000 task examples are all that the model sees, it will drift towards producing only that kind of output. Adding 1,000 examples of ordinary instruction-following — from a permissively licensed public set — keeps the model's general behaviour anchored while it learns your task.

This is called replay, it is the subject of a lesson in the dataset module, and in measured comparisons it does more to prevent catastrophic forgetting than every knob on this page combined. It is also the one that requires actual work rather than a config line, which is why it is the one most often skipped.

A default block you can copy

```
lora_dropout      = 0.05
weight_decay      = 0.0      # 0.01 for full fine-tuning
max_grad_norm     = 1.0
label_smoothing   = 0.0      # generation; 0.1 is fine on a classifier
head
neftune_noise_alpha= None    # try 5, measure on your own metric
```

Then spend the time you saved on the dataset. That is not a rhetorical flourish; it is where the difference between a good and a mediocre fine-tune reliably comes from.

BEFORE YOU MOVE ON

Why is label smoothing a poor default for generative fine-tuning although it is standard practice for classifiers?

- A. It increases training time substantially, because the loss must be computed against a full distribution rather than a single index.
- B. It conflicts with loss masking, since smoothed targets cannot be combined with `ignore_index` for prompt tokens.
- C. Generative models have far larger vocabularies, so spreading 0.1 across 128,000 tokens gives each an effectively zero target and the technique does nothing.
- D. It deliberately flattens the target distribution, and at sampling time a flatter distribution draws more from the tail, producing vaguer text.

CASE STUDY · MODULE 2

A free T4, a deadline, and a learning rate

A district education office in Madhya Pradesh, adapting a small model to summarise school inspection reports in Hindi.

The office had four hundred school inspection reports a month, each two to six pages of Hindi prose written by a visiting officer, and a requirement that each be reduced to a six-line summary for the block education officer. Two clerks did this. It took them most of a week.

Anjali, on contract from a state technology cell, was given a month and no budget. No budget meant a free Kaggle notebook: a T4, sessions of up to twelve hours, roughly thirty GPU-hours a week.

Her first run diverged. The loss fell for forty steps, spiked, and became NaN at step 212. She restarted it. It reached step 190 and did the same. The logs were full of lines about the gradient scaler halving its scale.

She had copied a configuration from a tutorial: QLoRA on an 8B model, learning rate $2e-4$, no warmup because the tutorial's author had removed it, sequence length 1,024. The configuration was not unreasonable. It was written for an A100.

The T4 is a Turing card and has no bf16. Everything ran in fp16, which has a maximum representable value of 65,504 and a smallest normal value around six times ten to the minus five. Gradients below that round to zero, so fp16 training multiplies the loss by a large constant before the backward pass and divides it out afterwards — and when that constant is too large, gradients overflow to infinity and the framework throws the step away and halves the scale. A healthy fp16 run discards a step occasionally. Anjali's was discarding them every few steps, and then diverging outright.

She had three ways forward, and a deadline in eleven days.

The first was to pay. A rented 24 GB card with bf16 costs somewhere around fifty rupees an hour on a community marketplace, and her whole run was perhaps two hours. The obstacle was not the money, which was trivial; it was that the office had no mechanism to spend it. A purchase of any size needed a sanction that took three weeks.

The second was to lower the learning rate and keep the free card. Halving it to $1e-4$ and adding ten steps of warmup would very probably stabilise the run. It would also mean the adapter travelled less far in the same number of steps, and she would need more steps to reach the same place — on a card that was already the slow part.

The third was to use a smaller model. A 1.5B model in fp16 with a LoRA fits a T4 comfortably at sequence 2,048, trains several times faster, and at that size full fine-tuning is affordable too. The summaries would be less fluent. Hindi tokenisation on the 8B model she had chosen was already producing about three tokens per word, which was eating her sequence budget; a model with better Devanagari coverage would help more than any hyperparameter.

She ran the cheap experiment first, because it cost half an hour. Three runs of 150 steps each, at $1e-4$, $2e-4$ and $5e-4$, on the 8B model, and she looked at the training loss curves rather than at anybody's advice. At $5e-4$ it spiked in the first fifty steps. At $2e-4$ it fell and then rose. At $1e-4$ it showed the shape she wanted: a fast drop in the first fifty steps, then a smoother decline, no scaler messages.

STOP HERE

Before turning the page: what would you have done, and what would it have cost if you were wrong? Write it down. The point of the next section is to compare.

WHAT HAPPENED

Anjali took the second route and added the third. She trained at $1e-4$ with ten steps of warmup on the 8B model overnight on Kaggle, and in parallel checked the tokeniser on a paragraph of real inspection prose across three model families. The 8B model she had started with produced 412 tokens; a model with better Devanagari coverage produced 168 for the same paragraph.

She switched to the second model, kept the learning rate at $1e-4$, and the run finished in four hours inside one Kaggle session with checkpoints pushed to a private Hub repository every two hundred steps. Her sequence length now covered the ninety-fifth percentile of real reports rather than truncating a third of them.

Two things she recorded in the run log turned out to matter more than the training itself. Tokens per optimiser step, which she had to compute herself because nothing printed it, came out at 5,900 — low, but not the 900 that would have made any learning rate useless. And a note that the divergence had been fp16, not a data problem, which stopped the next person in the office spending a day looking for a corrupted example.

The clerks now spend a morning a month checking summaries instead of a week writing them.

Anjali's first instinct could have been to look for a corrupted example in the data. Why was the hardware the right place to look first?

Because the symptom named the mechanism. Repeated messages about the gradient scaler reducing its scale are specific to fp16 loss scaling: they mean gradients are overflowing the format's maximum of 65,504 and steps are being discarded. Corrupt data produces NaN too, but it does not produce a scaler that is working hard and halving repeatedly. The T4 has no bf16, so fp16 was forced, and a configuration written for an A100 carried a learning rate that a card with fp32's exponent range would have tolerated.

Why did changing the base model help more than any hyperparameter, and what did it change besides quality?

Tokeniser efficiency on Devanagari. The first model split the same paragraph into 412 tokens and the second into 168. That changes three things at once: how much real text fits in a fixed sequence length, how much of each example survives truncation, and how many positions a long-range dependency has to span. A word split into six fragments is harder to learn a representation for than one split into two, so the tokeniser was setting a ceiling that no learning rate could lift.

The three 150-step runs cost half an hour on a free card. Why is that the cheapest half hour in the project, and what does each curve shape tell you?

Because the learning rate is the one hyperparameter that decides whether a run works at all, and a curve answers it directly. Loss falling slowly and roughly linearly with no elbow means the rate is too low. A clear fast drop in the first fifty steps followed by a smoother decline means it is about right. An early drop then a rise, visible spikes, or NaN means it is too high. Copying a number from a tutorial trained on different data, a different batch size and a different card is what the half hour replaces.

Check what stuck

6 questions, one right answer each. This one is for you, not for a record: answers and the reason behind each are at the back. If two questions from the same module go wrong, re-read that module before the exam rather than after.

- 1 Some label positions in a training batch are set to minus one hundred. What does that do?
 - A. It skips them in the forward pass, which is where the memory saving comes from.
 - B. Nothing flows from them: no loss, no gradient, so the model is scored on a different task from the one you may think.
 - C. It treats them as padding and penalises the model for predicting anything at all at those positions.
 - D. It down-weights them, so they still contribute a little to the gradient.
- 2 You switch on completion-only masking and your training loss jumps from 0.9 to 1.6. What has happened?
 - A. The tokenizer changed, so the per-token loss is on a different scale.
 - B. The model is now being asked to predict more positions than before, and the additional positions are harder ones it has not yet learned to handle.
 - C. The run has got worse, and the learning rate should come down.
 - D. The easy, near-identical prompt tokens were dragging the average down and are now excluded.
- 3 Why does a LoRA make the optimiser almost irrelevant to your memory planning?
 - A. Because AdamW's state is 8 bytes per trainable parameter, and only the adapter is trainable.
 - B. Because LoRA runs use plain SGD rather than AdamW, and SGD keeps no state.
 - C. Because the frozen base is stored in four bits, which removes the optimiser state along with it.
 - D. Because gradient checkpointing recomputes the optimiser state during the backward pass instead of holding it in memory for the whole run.

- 4 Why does a LoRA want a learning rate ten to twenty times higher than a full fine-tune?
 - A. Because fewer trainable parameters means each one receives a smaller share of the loss, so the rate has to be raised to compensate for that division.
 - B. Because low-rank matrices are numerically unstable and need larger steps to escape that instability.
 - C. Because the adapter starts at exactly zero and must travel from nothing to something within the run.
 - D. Because a quantised base damps the gradient by roughly the same factor.

- 5 A config has per-device batch 2 and gradient accumulation 8 on one GPU, and logs 900 tokens per optimiser step. What should you change first?
 - A. Nothing: 900 tokens per step is within the normal range for instruction data.
 - B. Raise accumulation, because the gradient is being estimated from too few real tokens.
 - C. Switch packing on, because the only thing that can produce a token count that low is padding being counted as real tokens.
 - D. The learning rate, because a token count that low needs a much smaller one.

- 6 You are training on a free T4 and the log repeatedly reports the gradient scaler reducing its scale. What is happening?
 - A. The card is thermally throttling, so the run will simply take longer than planned.
 - B. bf16 is underflowing, and moving the whole run to fp32 is the fix.
 - C. The learning-rate schedule is decaying, and the scaler reports it in the same line.
 - D. Gradients are overflowing fp16's maximum, so those optimiser steps are discarded and part of your compute buys nothing.

MODULE 3

Full, LoRA and the parameter-efficient family

The methods that made fine-tuning affordable, with the arithmetic to predict memory before you rent anything, and an honest account of which variants are worth their extra complexity.

BY THE END OF THIS MODULE YOU CAN

Predict on paper the memory a given model, method and sequence length will need, choose between full fine-tuning, LoRA, QLoRA and the tiny methods for a stated budget, and justify a rank, an alpha and a target-module set

19 · 10 min

Full, LoRA, QLoRA, and the Memory Arithmetic

THE ONE THING TO KEEP

Full fine-tuning costs about 16 bytes per parameter; LoRA charges that only on the 1% you train.

Count the bytes

Every training setup holds four things in memory: weights, gradients, optimizer state, and activations. Only the last depends on batch size, which is why "just lower the batch size" so often fails to save you.

For a model with P parameters, standard mixed-precision training with AdamW costs roughly:

Component	Bytes per parameter
Weights (bf16)	2
Gradients (bf16)	2
fp32 master weights	4
Adam first moment	4
Adam second moment	4
Total	16

A 7B model: $7e9 \times 16 = \mathbf{112\ GB}$, before a single activation. That does not fit on one 80 GB card. Full fine-tuning of a 7B is a multi-GPU job with FSDP or DeepSpeed, or an 8-bit-optimizer-plus-offload arrangement that is slow and fiddly. This is why almost nobody does it, and why the rest of this lesson exists.

LoRA: freeze the base, train a small correction

LoRA freezes W and learns a low-rank update, $W + BA$. For a weight of shape $d_{in} \times d_{out}$, B and A hold $r(d_{in} + d_{out})$ parameters. A 4096×4096 projection at $r=16$ is $16 \times 8192 = 131,072$ parameters against 16.7M — about 0.8%. Across a 7B model with adapters on every linear layer at $r=16$, you train roughly 20-40M parameters.

The 16-byte tax now applies only to those. The frozen base still sits in memory at 2 bytes per parameter.

- 7B base in bf16: **14 GB**
- adapters, their gradients and optimizer state: **under 1 GB**
- activations with gradient checkpointing, sequence 1024, small batch: **2-5 GB**

Call it 18-20 GB. That fits a 24 GB card (RTX 3090, 4090, A10G). It does not fit a 16 GB T4.

QLoRA: quantize the frozen base

The base is frozen, so it does not need to be precise — it only has to produce good forward activations. QLoRA stores it in 4-bit NF4, about 0.5 bytes per parameter plus quantization constants, and dequantizes each block on the fly during the forward pass. Adapters stay in 16-bit.

- 7B base in NF4: **about 4 GB**
- everything else as before: **3-6 GB**

Seven to ten gigabytes total. That fits a free Colab T4, a Kaggle P100, or a 12 GB consumer card. The cost is speed: expect 20-40% lower throughput than 16-bit LoRA because of the dequantization work.

```
from transformers import BitsAndBytesConfig
bnb = BitsAndBytesConfig(
    load_in_4bit=True,
    bnb_4bit_quant_type="nf4",
    bnb_4bit_compute_dtype="bfloat16", # use float16 on a T4; Turing
has no bf16
    bnb_4bit_use_double_quant=True,    # quantizes the quantization
constants too
)
```

Scaling to your model

Rough rule: QLoRA needs about $0.6 \times P$ bytes for the base, plus 3-6 GB of working memory.

- 3B → about 2 GB base → comfortable on 8 GB
- 7-8B → about 4.5 GB → fits 16 GB
- 13B → about 8 GB → tight on 16 GB, fine on 24 GB
- 70B → about 40 GB → needs one 48 GB or 80 GB card

If you have 8 GB of RAM and no GPU

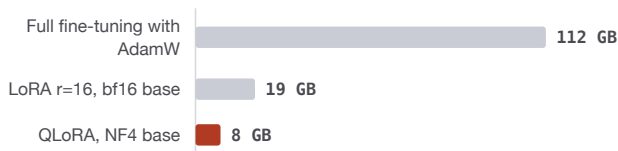
The honest thing: you cannot usefully fine-tune a 7B model on a CPU. A step that takes 0.4 seconds on a rented GPU takes minutes, and a run that takes twenty minutes takes a fortnight.

What you can do, and it is not a consolation prize:

- **Fine-tune a small model locally.** A 135M-600M parameter model trains on CPU in hours on a few hundred examples. Every part of the pipeline — chat templates, loss masking, evaluation — is identical to the 7B version, and those are the parts where mistakes actually happen.
- **Use free GPU hours.** Kaggle gives roughly 30 GPU-hours a week on a T4 or P100, with sessions that survive better than Colab's free tier. Colab free works but disconnects; checkpoint every few hundred steps and push somewhere persistent.
- **Rent.** A 24 GB card runs \$0.30-0.80 an hour on the spot-style marketplaces. Most LoRA runs in this course cost less than a coffee. Prices move constantly; check on the day.

The memory table is the thing to remember. Batch size touches one row of it.

A 7B model in memory, three ways



Sixteen bytes per parameter against roughly half a byte. The free T4's sixteen gigabytes is the line that decides what a learner with no budget can train.

BEFORE YOU MOVE ON

You have one free Colab T4: 16 GB, no bf16 support. You want to fine-tune a 7B model. Which plan is workable, and for the right reason?

- A. Full fine-tuning, with batch size 1 and gradient checkpointing to keep memory down.
 - B. Sixteen-bit LoRA, since only the adapters carry gradients and optimizer state.
 - C. QLoRA: the frozen base in 4-bit is about 4 GB, leaving room for adapters and checkpointed activations, at maybe 20-40% lower throughput.
 - D. Any of them, using gradient accumulation to spread the memory across steps.
-

20 · 9 min

What LoRA is actually doing

THE ONE THING TO KEEP

LoRA freezes W and learns a rank- r factorisation of the correction, initialised so the adapter starts as an exact no-op — which is why it needs a much higher learning rate, why it cannot be worse than the base at step one, and why it adds reweighting rather than new machinery.

A correction, factored

Fine-tuning changes a weight matrix W into $W + \Delta W$. LoRA's observation is that for adapting a pretrained model to a task, ΔW does not need to be a full matrix. It can be the product of two thin ones.

$$\Delta W = B \cdot A$$
$$W' = W + (\alpha/r) \cdot B \cdot A$$

A is $r \times k$, B is $d \times r$, r is small

For a $4,096 \times 4,096$ projection, W has 16.8 million entries. With $r = 16$, A has $16 \times 4,096 = 65,536$ and B has $4,096 \times 16 = 65,536$, so ΔW is described by 131,072 numbers — **0.78% of the matrix it modifies**.

W itself is frozen. It receives no gradient, needs no optimiser state, and can sit in 4-bit quantised form if you like. Only A and B are trained.

The two initialisations, and why they matter

A is initialised with small random values. **B is initialised to exactly zero.**

One 4,096 by 4,096 projection, adapted

	Numbers stored	Share of W
Full update	16,777,216 a full matrix	100% every entry free to move
LoRA at $r = 16$	131,072 two thin matrices	0.78% reweighting, not new machinery

B is initialised to exactly zero, so BA is zero at step one and the model starts as the base model. That is why a LoRA cannot be worse on step one, and why it wants a learning rate ten to twenty times higher.

That means $BA = 0$ at step one, so $W' = W$: the model at the start of training is bit-for-bit the base model. This is not a detail, it has three consequences you will meet:

- There is no random perturbation to recover from. A LoRA run cannot make the model worse on step one, which full fine-tuning at a bad learning rate certainly can.
- The adapter must travel from nothing to something within the run, which is why LoRA wants a learning rate ten to twenty times higher than full fine-tuning.
- If you initialised *both* to zero, the gradient of both would be zero forever and nothing would ever train. The asymmetry is load-bearing.

The scaling factor

The α/r term multiplies the adapter's contribution. The convention exists so that changing r does not silently change the magnitude of the update: double the rank and

each individual entry of BA tends to shrink, and dividing by r compensates.

In practice people set α to r or to $2r$ and treat the pair as one decision. `r=16, alpha=32` is the single most common configuration in the wild. The important thing is to know that α and the learning rate push in the same direction — doubling α at a fixed learning rate has an effect very like doubling the learning rate — so do not tune both at once and then wonder which mattered.

Why a low rank is enough

The intuition, which has empirical support and is not a theorem: adapting a pretrained model to a task appears to require far fewer degrees of freedom than the model has parameters. The pretrained weights already contain the representations; the adaptation is a comparatively simple reweighting of what is already there. Measurements of the "intrinsic dimension" of fine-tuning tasks found that surprisingly small subspaces suffice, and LoRA is that observation turned into a method.

The honest flip side: this is exactly why LoRA is poor at adding genuinely new capability. A rank-16 correction cannot install a language the base model never saw, or a modality it has no representations for. Low rank is enough *when the thing you want is a reweighting*, and it is not enough when the thing you want is new machinery. That is the mechanism behind the rule from module one — training teaches shape, not substance.

Seeing it in code

```
from peft import LoraConfig, get_peft_model

cfg = LoraConfig(r=16, lora_alpha=32, lora_dropout=0.05,
                 target_modules=["q_proj", "k_proj", "v_proj", "o_proj",
                                "gate_proj", "up_proj", "down_proj"],
                 task_type="CAUSAL_LM")
model = get_peft_model(base, cfg)
model.print_trainable_parameters()
# trainable params: 41,943,040 || all params: 8,072,204,288 || trainable%: 0.5196
```

Run `print_trainable_parameters()` on every configuration before you launch. If it prints 0, your `target_modules` names do not match the architecture and you are about to train nothing for three hours — a failure that produces a perfectly flat, perfectly innocent-looking loss curve.

What you end up with

The artefact on disk is those A and B matrices: roughly 160 MB in bf16 for the 42M-parameter configuration above, against 16 GB for the base model. You can keep a dozen of them for different tasks, swap them at serving time, and email one to a colleague.

At inference you have a choice. Keep them separate and pay a small overhead for the extra matrix multiplications, or merge them — compute $W + (\alpha/r)BA$ once and write out a normal model with no adapter at all. Merging costs nothing at inference and loses the ability to swap. That choice, and what it does to a quantised base, is a lesson of its own later in the course.

BEFORE YOU MOVE ON

Why is B initialised to zeros while A gets small random values, rather than both being randomly initialised?

- A. So the adapter is a no-op at step one, while the random A keeps gradients non-zero.
- B. Because zero-initialised matrices train faster, since their gradients are larger in the early steps.
- C. Because B is the larger matrix and random values in it would require more memory to store in the optimiser state.
- D. To keep the rank of the product at exactly r , which random initialisation of both matrices would not guarantee.

21 · 9 min

Choosing rank, alpha and which modules to touch

THE ONE THING TO KEEP

Adapting all linear projections at $r=16$ beats adapting the attention pair at a high rank, alpha should be fixed at $2r$ so it does not confound your learning-rate sweep, and `target_modules="all-linear"` survives a change of model family where a copied name list does not.

Breadth beats depth

There are two ways to give a LoRA more capacity: raise the rank, or adapt more matrices. The consistent finding across published comparisons and a great deal of practical experience is that **which modules you target matters more than the rank you pick**.

The early LoRA work adapted only the attention query and value projections, and that convention stuck around long after the evidence moved. On a modern decoder, adapting all seven linear projections at a modest rank beats adapting two at a high rank, at a similar parameter count.

The seven, on a Llama-style architecture:

- Attention: `q_proj`, `k_proj`, `v_proj`, `o_proj`
- Feed-forward: `gate_proj`, `up_proj`, `down_proj`

The feed-forward matrices are the larger ones — with hidden 4,096 and intermediate 14,336, each MLP projection is 58.7M parameters against the attention projections' 16.8M — and a good deal of a transformer's task-specific behaviour appears to live there.

The parameter count, worked

For an 8B model, 32 layers, hidden 4,096, intermediate 14,336, grouped-query attention with 8 key-value heads (so `k_proj` and `v_proj` are $4,096 \times 1,024$), at $r = 16$:

```

per layer:  q 131,072  k 81,920  v 81,920  o 131,072
            gate 294,912  up 294,912  down 294,912
            ----- = 1,310,720
× 32 layers                                = 41,943,040  (0.52% of the
model)

attention only (q,k,v,o): 425,984 × 32      = 13,631,488  (0.17%)

```

Committing those numbers to paper takes two minutes and tells you, before you rent anything, that your adapter file will be about 160 MB and your optimiser state about 340 MB.

Choosing r

A defensible default: **$r = 16$, $\alpha = 32$, all seven modules**. From there:

- **$r = 8$** if your dataset is under about 1,000 examples. Less capacity to memorise with.
- **$r = 32$ or 64** for large datasets (tens of thousands of examples), for continued pre-training on a new domain, or when $r = 16$ visibly plateaus below where you need it.
- **$r = 128$ and above** rarely helps on instruction data and starts to inherit full fine-tuning's problems — more capacity to overfit, more forgetting — without full fine-tuning's benefits. If you find yourself here, ask whether full fine-tuning is what you actually want.

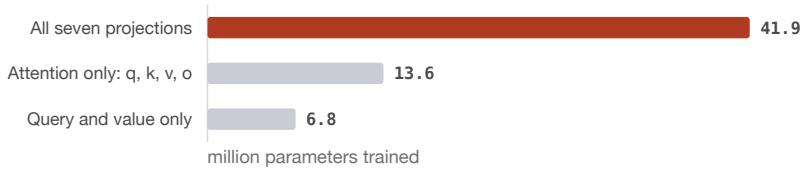
The test for whether rank is your constraint is cheap: run $r = 8$ and $r = 32$ for the same number of steps and compare on your task metric. If they are within noise, rank was never the limiting factor, and it usually is not. Data is.

Alpha, and the one trap

Set `alpha = 2r` and leave it. The trap is that α and the learning rate are close to interchangeable in effect — the update is scaled by $(\alpha/r) \cdot lr$ — so if you sweep both you are sweeping one thing twice and will draw confident conclusions from a confounded experiment. Fix $\alpha = 2r$, tune the learning rate.

One genuine exception: at high rank the conventional α/r scaling divides by r , which shrinks the effective update as rank grows and is part of why high ranks often disappoint. Rank-stabilised LoRA (`use_rslora=True` in PEFT) divides by $\sqrt{r}$ instead, which keeps the update magnitude better behaved. If you are working at $r = 64$ or above, switch it on; below that it changes little.

Trainable parameters on an 8B model at $r = 16$



Breadth beats depth. Adapting all seven projections at a modest rank beats adapting two at a high rank, and the feed-forward matrices are the larger ones.

Embeddings and the head

Two matrices are usually *not* adapted and sometimes should be:

- **The embedding table and the output head.** If you have added new tokens — a new special token, a domain vocabulary — those embeddings start as noise or as an average, and no amount of adapting the transformer blocks fixes them. PEFT's `modules_to_save=["embed_tokens", "lm_head"]` trains them fully. This makes your adapter file much larger (an 8B model's embedding table is $128,256 \times 4,096 = 525\text{M}$ parameters, about 1 GB in bf16) so do it only when you have genuinely added tokens.
- **LayerNorms.** Cheap to train, occasionally helpful, rarely decisive.

The shortcut that is usually right

```
target_modules="all-linear"
```

PEFT accepts this and resolves it to every linear layer in the model except the output head. It is architecture-agnostic, which means it keeps working when you switch from a Llama-style model to a Qwen or Mistral one whose module names differ — and a list of module names copied from a blog post about a different architecture is one of the more common reasons `print_trainable_parameters()` reports zero.

BEFORE YOU MOVE ON

A team gets no improvement from raising LoRA rank from 16 to 64 on a 3,000-example dataset, while switching from attention-only to all seven linear projections at $r=16$ gives a clear gain. What does this pattern indicate?

- A. Rank 64 overfitted the 3,000 examples, so its true capacity gain was cancelled by memorisation.
- B. The constraint was not the dimensionality of each correction but which matrices could be corrected at all.
- C. The alpha was left at 32 while rank rose, so the higher-rank run had a smaller effective update and was undertrained.
- D. Grouped-query attention makes the k and v projections small, so attention-only adapters have too few parameters to train stably at any rank.

22 · 10 min

QLoRA: quantising the part you are not training

THE ONE THING TO KEEP

QLoRA quantises only the frozen base to NF4 with block-wise scales and double-quantised constants, buying a 4x memory reduction for roughly 30% throughput and an adapter that was fitted against a slightly wrong model — which is why merging it back is the awkward part.

The idea in one line

The frozen base model receives no gradients, so it does not need to be stored precisely. Quantise it to 4 bits, keep the small trainable adapter in higher precision, and a 70B model becomes trainable on a single 48 GB card.

The numbers for an 8B model: bf16 weights are 16 GB, NF4 weights are about **4.1 GB**. For a 70B: 140 GB becomes about 36 GB.

NF4, and why not plain int4

Uniform 4-bit quantisation carves the range into 16 evenly spaced buckets. Neural network weights are not uniformly distributed — they are roughly zero-centred and bell-

shaped — so even spacing wastes most of its buckets on values that barely occur and crowds the region where the weights actually live.

NormalFloat 4 places its 16 levels so that, for a normally distributed tensor, each bucket receives an equal share of the values. It is information-theoretically matched to the distribution the weights actually have, and it measurably beats uniform int4 at the same bit width.

Two further pieces make it work:

Block-wise scaling. Quantisation is applied per block of 64 weights, each with its own scaling constant, rather than per tensor. One outlier weight then distorts 64 values instead of 16 million.

Double quantisation. Those per-block constants are themselves numbers you have to store — one fp32 per 64 weights is 0.5 bits per parameter of overhead. Quantising *them* to 8 bits, in blocks of 256, brings the overhead down to about 0.127 bits per parameter. On a 65B model that saving is around 3 GB, which is the difference between fitting and not.

What it costs

Three honest costs, none of them usually decisive but all of them real:

1. **Speed.** Every matrix multiplication dequantises its weights to bf16 on the fly. QLoRA typically runs 25 to 40% slower per step than the same LoRA on unquantised weights. You are trading time for the ability to run at all.
2. **Quantisation error in the forward pass.** The gradients computed for your adapter are gradients through a slightly wrong model. The QLoRA paper's central claim is that the final quality matches 16-bit LoRA on the tasks measured, and that has held up broadly — but "broadly" is the honest word. On tasks that are sensitive to small numeric differences, and at very small model sizes where quantisation error is proportionally larger, the gap is not always zero.
3. **Merging is awkward.** Your adapter was fitted against the *quantised* base. Merge it into the full-precision base and you get a model that is not quite what you trained; merge it into the quantised base and you must requantise, which introduces its own error. The practical answers are to keep the adapter separate at serving time, or to use a method designed for this (QA-LoRA and similar) rather than merging naively and hoping.

The configuration

```
from transformers import BitsAndBytesConfig
import torch

bnb = BitsAndBytesConfig(
    load_in_4bit=True,
    bnb_4bit_quant_type="nf4",          # not "fp4"
    bnb_4bit_use_double_quant=True,
    bnb_4bit_compute_dtype=torch.bfloat16, # float16 on a T4
)
```

Four lines, and three of them have a wrong value that works. `fp4` instead of `nf4` trains without complaint and is slightly worse. Leaving double quantisation off costs memory you may need. Setting the compute dtype to `fp32` quietly removes most of the speed benefit.

What QLoRA is made of

NF4, not plain int4

Sixteen levels placed so each receives an equal share of a normally distributed tensor, rather than evenly spaced across a range the weights do not occupy.

Block-wise scaling

One scaling constant per 64 weights, so a single outlier distorts 64 values instead of sixteen million.

Double quantisation

Those constants quantised too, in blocks of 256: overhead falls from about 0.5 bits per parameter to about 0.127.

A higher-precision adapter on top

The frozen base is 4-bit; only the small trained part is kept and updated in 16-bit. Three precisions in one model, and that is normal.

Each idea is what makes the next affordable. On a 65B model, double quantisation alone saves about three gigabytes — the difference between fitting and not.

Pair it with `optim="paged_adamw_8bit"`. Paged optimisers use CUDA unified memory so that a memory spike on one unusually long batch is absorbed by CPU RAM instead of ending the run.

When *not* to use it

If the model fits in `bf16` on the card you have, use plain LoRA. QLoRA is a memory technique, and paying 30% throughput for memory you did not need is a straight loss. A 3B model on a 24 GB card has no business being quantised for training.

The decision is arithmetic, not taste: parameters $\times 2$ bytes, plus activations, plus logits. If that is under your card's capacity with headroom, skip the quantisation.

The free path

This is the technique that put fine-tuning on free hardware. A 7B or 8B model in NF4 at sequence length 1,024, batch 1 with gradient accumulation and checkpointing on, fits inside a free Colab or Kaggle T4's 16 GB. It is slow — expect several hours for a few thousand steps — and Colab will disconnect you, so save checkpoints to Drive every couple of hundred steps. But it costs nothing, and the resulting adapter is the same kind of artefact a rented A100 would have produced.

BEFORE YOU MOVE ON

Why does double quantisation exist, given that the weights are already at 4 bits?

- A. It applies a second quantisation pass to the weights themselves, reducing them from 4 bits to an effective 2 bits.
 - B. It quantises the adapter matrices as well as the base, so the whole model is uniformly 4-bit during training.
 - C. It quantises the per-block scaling constants, which otherwise add about 0.5 bits per parameter.
 - D. It quantises the activations in addition to the weights, halving the memory used by the forward pass.
-

23 · 9 min

DoRA, rsLoRA and the variant question

THE ONE THING TO KEEP

rsLoRA and DoRA fix specific defects — the alpha-over-r damping at high rank and the coupling of magnitude to direction at low rank — but variants are the last term in the sum, behind the dataset, the template, the learning rate and the target modules.

A field with a lot of papers

LoRA has attracted dozens of variants. Most give a point or two on a benchmark, cost some complexity, and do not survive contact with a real dataset any better than the

baseline. A few are worth knowing, and knowing *why* they help is more useful than the list.

rsLoRA: fixing the scaling at high rank

Standard LoRA scales its update by α/r . As r grows, that divisor grows linearly, and the effective magnitude of the update falls. This is part of why people report high ranks disappointing: the higher-rank adapter was not just more expressive, it was also more damped.

Rank-stabilised LoRA divides by $\sqrt{r}$ instead. The update magnitude stays roughly constant as rank rises, so a rank-64 run genuinely explores more capacity rather than exploring the same amount more slowly.

```
LoraConfig(r=64, lora_alpha=128, use_rslora=True, ...)
```

When to use it: any time $r \geq 32$. Below that the difference is small. It costs nothing and it removes a confound from your rank experiments, which is reason enough.

DoRA: separating magnitude from direction

A weight matrix can be decomposed into a magnitude per column and a direction. DoRA does exactly that, keeps a trainable magnitude vector, and applies the LoRA update to the direction only.

The motivation is an observed difference in how full fine-tuning and LoRA move weights: full fine-tuning tends to make relatively large directional changes with modest magnitude changes, while LoRA couples the two. Decoupling them makes LoRA's learning dynamics look more like full fine-tuning's.

In published comparisons DoRA improves over LoRA by a few points at low rank, with the gap narrowing as rank rises — which is consistent with the story, since a higher-rank LoRA can represent the decoupling itself.

```
LoraConfig(r=8, lora_alpha=16, use_dora=True, ...)
```

The cost: roughly 20 to 30% slower per step, because of the extra normalisation work, and additional memory for the magnitude vectors. **When to use it:** when you are memory-constrained to a low rank and the extra time is affordable. At $r = 32$ or above, the case is weak.

PiSSA and the initialisation family

Standard LoRA starts B at zero, so the adapter begins as a no-op and has to learn from scratch. PiSSA instead initialises A and B from the top singular vectors of the original weight matrix, and freezes the residual. The claim is faster convergence, because the adapter starts aligned with the directions the weight matrix already considers important.

This family — PiSSA, OLoRA, EVA and others — is genuinely interesting and genuinely unsettled. The reported gains are real on the benchmarks reported; whether they hold on a specific task with a few thousand examples is the sort of question that only your own ablation answers. Treat them as worth a single comparison run, not as defaults.

LoRA+ : two learning rates

A and B play asymmetric roles — A maps down to rank r , B maps back up — and the argument is that they should not share a learning rate. LoRA+ gives B a rate several times higher than A, typically $16\times$. It is a one-line change in some libraries and reports faster convergence at no memory cost.

Cheap to try, and it belongs in the same bucket as rsLoRA: a correction to a scaling decision that was made for simplicity rather than for optimality.

How to decide, without reading fifty papers

A workable policy:

1. **Baseline first.** $r = 16$, $\alpha = 32$, all-linear, plain LoRA. Get a number.
2. **Fix the dataset.** If the baseline is disappointing, the variant is not your problem. Nothing on this page compensates for 600 examples or a broken chat template.
3. **Then, if you have budget for exactly two more runs:** rsLoRA if you are at high rank, DoRA if you are pinned at low rank. Those are the two with the clearest mechanism and the most replication.
4. **Measure on your own metric**, not on the paper's benchmark. Most of these methods were evaluated on instruction-following benchmarks scored by a model judge, and your extraction task is not that.

The uncomfortable summary

The distribution of outcomes in this field is lopsided. Across a lot of practical work, the ordering of what moves a fine-tune's quality is: the dataset, then the chat template and masking, then the learning rate, then the target modules, then the rank, then the vari-

ant. Variants sit last, and a team that reaches for DoRA before checking that its loss mask is correct is optimising the smallest term in the sum.

BEFORE YOU MOVE ON

A team finds that raising rank from 16 to 64 gives almost no gain, and that switching on `use_rslora` at `r=64` suddenly does. What was suppressing the higher-rank run?

- A. At `r=64` the adapter had too many parameters for the dataset, and rsLoRA regularises it back to an effective lower rank.
 - B. Rank 64 exceeded the intrinsic dimension of the task, so the extra directions received no gradient until rsLoRA redistributed it.
 - C. Higher ranks need more steps to converge, and rsLoRA accelerates convergence by reinitialising the adapter from the weight matrix's singular vectors.
 - D. The alpha-over-r scaling divides by rank, so the rank-64 update was damped four times more than the rank-16 one at the same alpha-to-rank ratio.
-

24 · 8 min

Soft prompts and prefixes

THE ONE THING TO KEEP

Soft prompts and prefixes train a few hundred kilobytes of learned vectors with every weight frozen, which makes per-tenant adaptation almost free to store — but they steer the computation rather than changing it, so they close the gap to full fine-tuning only at large model scale.

Training the prompt instead of the model

A prompt is a sequence of token embeddings. Those embeddings are looked up from a table, but nothing in the architecture requires them to correspond to real tokens. You can prepend a handful of *learned* vectors — not the embedding of any word, just free parameters — and train them with gradient descent while every weight in the model stays frozen.

That is prompt tuning. With 20 virtual tokens on a model with hidden size 4,096, you are training $20 \times 4,096 = \mathbf{81,920 \text{ parameters}}$. The artefact is about 320 kilobytes. You could store one per customer in a database column.

Prefix tuning: the same idea, deeper

Prompt tuning only touches the input. Prefix tuning goes further: it prepends learned key and value vectors at *every* attention layer, so the adaptation is present throughout the network rather than only at the bottom.

This is more expressive and correspondingly larger — still tiny, typically under 1% of the model. During training the prefixes are usually produced by a small MLP rather than optimised directly, because optimising them directly is unstable; after training the MLP is discarded and only the resulting vectors are kept. P-tuning v2 is essentially this method applied to understanding tasks.

```
from peft import PrefixTuningConfig, PromptTuningConfig, get_peft_model

cfg = PromptTuningConfig(task_type="CAUSAL_LM", num_virtual_tokens=20,
                        prompt_tuning_init="TEXT",
                        prompt_tuning_init_text="Classify the support
ticket:")
```

That `prompt_tuning_init="TEXT"` initialises the virtual tokens from the embeddings of a real sentence rather than randomly, and it matters more than it sounds — random initialisation of soft prompts converges slowly and sometimes not at all.

The three honest disadvantages

They eat your context. Twenty virtual tokens are twenty tokens of the window, at every request, forever. On a long-context task that is nothing; on a tight latency budget with a 512-token limit it is 4% of your input.

They cannot be merged. A LoRA can be folded into the weights so inference costs exactly what the base model costs. A soft prompt is always present at runtime, and the attention over those extra positions is real work.

They underperform at small scale. This is the substantive one. The original prompt-tuning result was that the method *matches* full fine-tuning — but only as model size grows past roughly ten billion parameters. Below that there is a real gap, and the gap widens the harder the task. On a 3B model with a genuinely difficult task, LoRA will beat a soft prompt, reliably.

The mechanism is intuitive enough: a soft prompt can only steer the computation the model already performs, by choosing what it conditions on. LoRA can change the computation. Big models have more behaviours available to steer towards; small ones need the change.

Where they are the right answer

Two situations, both about the artefact rather than the quality:

Many tenants, one model. If you need per-customer adaptation for a thousand customers, a thousand LoRAs is 160 GB and a thousand soft prompts is 320 MB. When a soft prompt is good enough, the storage and switching story is far simpler — you look up a tensor and concatenate it, with no adapter-swapping machinery in the serving stack.

Frozen-model constraints. Some deployment environments genuinely will not let you ship modified weights — a shared inference cluster, a compliance rule, a vendor endpoint. A prefix is data, not a model.

The thing to try before either

Write the prompt in words first. A soft prompt is an optimised version of an instruction, and if you have not yet found a good instruction by hand, you are asking gradient descent to search a space you have not explored yourself. The cases where soft prompting wins are the cases where a good hand-written prompt already gets close and you want the last few points, cheaply and with a very small artefact.

And if a good hand-written prompt already gets close, ask honestly whether the last few points are worth a training pipeline at all. That question is the spine of this whole course, and it applies to the cheapest method as much as to the most expensive.

BEFORE YOU MOVE ON

A team gets good results from prompt tuning a very large model, then tries the same approach on a 3B model for the same task and finds it clearly worse than LoRA. What explains the size dependence?

- A. A soft prompt only selects among behaviours the frozen model already has, and a small model has fewer.
- B. Prompt tuning consumes context, and a 3B model's shorter context window leaves too little room for the real input.
- C. The 3B model has a smaller hidden size, so the 20 virtual tokens contain too few parameters to encode the task.
- D. Smaller models are trained on less data, so their embedding space is less smooth and gradient descent cannot optimise vectors within it.

25 · 8 min

The tiny methods, and a gotcha in one of them

THE ONE THING TO KEEP

IA³ adapts a model with three learned scaling vectors per layer, about 0.01% of parameters, merging exactly into the weights and suiting tiny datasets — while BitFit is nearly a no-op on Llama-style models because they have almost no bias terms to train.

Below LoRA, there is more

LoRA trains about 0.5% of a model. There are methods that train a hundredth of that and still work on some tasks. They matter less for quality than for what they reveal: how little of a pretrained model needs to move for its behaviour to change.

IA³: learned scaling vectors

IA³ (Infused Adapter by Inhibiting and Amplifying Inner Activations) introduces three learned vectors per layer, which *multiply* activations element-wise:

- one scaling the attention keys,
- one scaling the attention values,
- one scaling the feed-forward intermediate activations.

A vector, not a matrix. For an 8B model that is a few hundred thousand parameters in total — around **0.01%** — and the file is a couple of megabytes.

Two properties make it more interesting than the size alone suggests. First, it can be merged into the weights exactly, because scaling a row of a matrix is just scaling the matrix, so there is zero inference overhead. Second, it was designed for the few-shot setting and does unusually well with very small datasets, where a LoRA has enough capacity to memorise and IA³ simply does not.

```
from peft import IA3Config
cfg = IA3Config(task_type="CAUSAL_LM",
                target_modules=["k_proj", "v_proj", "down_proj"],
                feedforward_modules=["down_proj"])
```

The `feedforward_modules` argument must be a subset of `target_modules`; IA³ applies scaling differently for feed-forward layers, and getting this wrong is the usual con-

figuration error.

BitFit, and the gotcha

BitFit trains only the bias terms of the model and freezes everything else — about 0.08% of the parameters of a BERT-style model. On classification tasks it was shown to come remarkably close to full fine-tuning, which was a genuinely surprising result when it appeared.

Here is the part that catches people out. **Most modern decoder LLMs have almost no biases to train.** Llama-family models use RMSNorm, which has a scale parameter and no bias, and their attention and feed-forward projections are configured without bias terms. Run BitFit on Llama-3-8B and you will find a few thousand trainable parameters — essentially the normalisation scales — and a run that goes almost nowhere.

So BitFit is a real method for encoder models with biases (BERT, DeBERTa, RoBERTa, where it is worth knowing) and close to a no-op on the architectures most of this course is about. It is a good example of why you should always print the trainable-parameter count rather than trusting that a named method is doing something.

Norm tuning

Training only the RMSNorm or LayerNorm scale parameters is available on modern decoders — roughly 0.003% of parameters. On its own it is weak. Combined with LoRA it occasionally adds a little, at negligible cost. Worth knowing, not worth a run of its own.

What these methods are for

Be honest about the use cases, because "0.01% of parameters" is a marketing number and not a reason.

1. **Extreme multi-tenancy.** Ten thousand customers, each with a 2 MB artefact, all served from one set of base weights. At that scale IA³'s size is the point.
2. **Very small datasets.** Fifty to five hundred examples, where LoRA's capacity is a liability. IA³ has almost nothing with which to memorise, so what it learns tends to generalise.
3. **Understanding.** These methods are a standing demonstration that adapting a pretrained model is a small intervention. That reframes the whole activity: you are not teaching the model, you are adjusting emphasis in something that already knows how to do the task approximately.

What they are not for

Anything requiring new behaviour. If a full LoRA at rank 32 is not enough for your task, a method with a thousand times fewer parameters is not the answer. The ordering of capacity — $IA^3 < \text{prompt tuning} < \text{LoRA} < \text{full fine-tuning}$ — is monotonic, and every step down the list also steps down the ceiling.

The practical recommendation is dull: start with LoRA. Reach for IA^3 only when the artefact size is genuinely a constraint, or when your dataset is small enough that LoRA visibly overfits — and verify both claims with numbers rather than assuming them.

BEFORE YOU MOVE ON

Someone configures BitFit on Llama-3-8B and the trainable-parameter count prints in the low thousands, with a flat loss curve. What has happened?

- A. The target module names do not match the architecture, so no layers were selected for training.
 - B. Llama-style models omit bias terms and use RMSNorm, so there are almost no biases to train.
 - C. BitFit requires a classification head, and on a causal language model it falls back to training only the output layer's bias.
 - D. Quantisation removed the bias terms, since 4-bit formats store only weights and rescale biases into the block constants.
-

26 · 9 min

Full fine-tuning, when it is genuinely right

THE ONE THING TO KEEP

Full fine-tuning is right for a new language, a very large dataset, a small model or a distillation student, costs about 16 bytes per parameter in persistent state before activations, and forgets precisely because every weight is free to move — so run the LoRA first to find out whether the data is worth it.

Four cases where LoRA is not enough

Parameter-efficient methods are the default, not the rule. Four situations reliably want all the weights.

1. A new language or script. If the base model saw very little Marathi, or almost no Tamil script, the adaptation required is not a reweighting of existing representations — the representations are absent. Tokenisation is usually bad too (the next module covers why), and fixing it means resizing the embedding table, which is not something an adapter does. This case typically wants continued pretraining on raw text with everything unfrozen.

2. A large dataset. Above roughly 100,000 examples, the low-rank constraint becomes a ceiling rather than a helpful regulariser. There is measured evidence that LoRA both learns less and forgets less than full fine-tuning — which is a good trade when you have 3,000 examples and a bad one when you have 300,000 and actually want the model to change.

3. Small models. Full fine-tuning a 0.5B or 1.5B model is cheap: 1.5B at 16 bytes per parameter is 24 GB of state, which fits one 40 GB card, and with 8-bit AdamW it fits a 24 GB one. If full fine-tuning is affordable and better, the reason to use LoRA has evaporated.

4. Distillation into a small student. Related to the last two: you have a large synthetic dataset and a small model whose behaviour you want to change substantially. This is the archetypal full-fine-tuning job.

The memory, computed

Mixed-precision full fine-tuning with AdamW costs about **16 bytes per parameter** of persistent state:

```
fp32 master weights    4
fp32 gradients         4
AdamW m and v          8
--
                      16 bytes/param
```

For 7B: 112 GB, before a single activation. Add activations and a logits tensor and you are near 140 GB. That is two 80 GB cards, minimum.

With 8-bit AdamW the state term drops from 8 to 2, giving 10 bytes per parameter — 70 GB for 7B, which starts to be one card plus care.

Sharding: ZeRO and FSDP

When the state does not fit one card, you split it across several. The standard scheme has three stages, and they are cumulative:

- **Stage 1** shards the optimiser state across devices. Removes the 8-byte term from each card. Cheap, almost no communication cost.
- **Stage 2** additionally shards gradients. Removes another 4 bytes per parameter per card.
- **Stage 3** additionally shards the parameters themselves. Each card holds a slice, and layers are gathered just in time for the forward pass and released afterwards. Maximum memory saving, most communication — this is the one that needs fast interconnect to be efficient.

DeepSpeed calls these ZeRO stages; PyTorch's native FSDP implements the same idea and is what `accelerate` will configure for you by default. They are equivalent in intent; pick whichever your stack supports and do not agonise.

CPU offload goes further: optimiser state, and optionally parameters, live in system RAM and are moved to the GPU as needed. This genuinely lets a 7B full fine-tune run on a single 24 GB card with 128 GB of system RAM. It is also roughly five to ten times slower, because you are moving tens of gigabytes across PCIe every step. It is a way to make something possible, not a way to make it practical.

A configuration that works

```
# accelerate config for FSDP full fine-tuning
compute_environment: LOCAL_MACHINE
distributed_type: FSDP
mixed_precision: bf16
num_processes: 8
fsdp_config:
  fsdp_sharding_strategy: FULL_SHARD      # equivalent to ZeRO-3
  fsdp_auto_wrap_policy: TRANSFORMER_BASED_WRAP
  fsdp_transformer_layer_cls_to_wrap: LlamaDecoderLayer
  fsdp_state_dict_type: SHARDED_STATE_DICT
```

The `fsdp_transformer_layer_cls_to_wrap` line is the one that goes wrong. Get the class name wrong and FSDP wraps the whole model as a single unit, which defeats the sharding entirely and produces an out-of-memory error that looks like the sharding is not working — because it is not.

The thing to know before you commit

Full fine-tuning forgets. That is not a flaw to be engineered away; it is the direct consequence of every weight being free to move towards your objective, including the weights

that encoded something else. The mitigations are a replay mixture in the data, a lower learning rate, fewer epochs, and a canary set evaluated at every checkpoint.

And the decision is reversible in one direction only. A LoRA that disappoints costs you a day. A full fine-tune that disappoints costs you the run, the cards, and a model you now have to store and serve at full size. Run the LoRA first, even when you are fairly sure you will need the full thing — it tells you whether the data is any good, which is the question that actually decides the project.

BEFORE YOU MOVE ON

A team configures FSDP with `FULL_SHARD` across 8 GPUs for a 7B full fine-tune and still runs out of memory on every card. The transformer layer class name in the config is misspelled. Why does that single line cause the failure?

- A. An unmatched class name disables mixed precision, so weights, gradients and optimiser state all revert to `fp32`.
- B. FSDP falls back to stage 1 sharding when the class is unknown, which leaves gradients and parameters replicated on every device.
- C. The policy matches no layer, so the model is wrapped as one unit and each card holds the whole state.
- D. The misspelling prevents gradient checkpointing from attaching to the decoder layers, so activations for all 32 layers are retained.

27 · 10 min

Training on a free or nearly free budget

THE ONE THING TO KEEP

A free Kaggle T4 with 12-hour sessions trains an 8B QLoRA at sequence 1,024, Unsloth's rewritten kernels roughly double that throughput, and the whole chain from dataset to a GGUF model running on your own laptop costs nothing — with checkpointing to Drive or the Hub being what makes a killable session survivable.

What the free tiers actually give you

Two services matter, and they are not equivalent.

Kaggle Notebooks. Roughly 30 GPU-hours a week, sessions up to 12 hours, and the choice of a single P100, two T4s (16 GB each), or a TPU. The long session limit is the important part: a 12-hour job can finish, where a Colab session usually cannot.

Google Colab, free tier. A T4 with 16 GB, no guarantee of availability, and disconnection after a period of inactivity or at the platform's discretion. Good for iteration and experimentation; unreliable for a run that needs six uninterrupted hours.

Both give Turing-generation or Pascal-generation cards, which means **no bf16**. Everything in the precision lesson about fp16 and loss scaling applies to you.

What fits in 16 GB

Measured, not guessed, with gradient checkpointing on and flash attention where supported:

- **7B or 8B, QLoRA, sequence 1,024, batch 1 with accumulation.** Fits, with a couple of gigabytes spare. This is the headline configuration and it does work.
- **7B or 8B, QLoRA, sequence 2,048.** Marginal. Possible with Unsloth, uncomfortable without.
- **3B, LoRA in fp16 without quantisation.** Comfortable.
- **1.5B, LoRA, sequence 2,048, batch 4.** Easy.
- **0.5B, full fine-tuning with 8-bit AdamW.** Fits, and for a 0.5B model full fine-tuning is usually the better choice anyway.
- **A 150M encoder classifier, full fine-tuning.** Trains in minutes, and would train on the CPU if you had to.

What does not fit, honestly: anything at 70B, full fine-tuning above about 1B, and long-context work. No configuration trick changes those.

Unsloth, and what it is doing

Unsloth rewrites the attention, RoPE, RMSNorm and cross-entropy kernels for LoRA and QLoRA training in Triton, and manually derives the backward passes rather than relying on autograd. The result on a single card is roughly two times faster training and substantially lower memory, with the same mathematics. The single-GPU version is open source and free.

On a free T4 this is not a nicety. It is often the difference between a configuration that fits and one that does not, and between a run that finishes in your session and one that does not.

```

from unsloth import FastLanguageModel
model, tok = FastLanguageModel.from_pretrained(
    "unsloth/llama-3.1-8b-bnb-4bit", max_seq_length=1024,
    load_in_4bit=True)
model = FastLanguageModel.get_peft_model(model, r=16, lora_alpha=32,
    target_modules=["q_proj", "k_proj", "v_proj", "o_proj",
        "gate_proj", "up_proj", "down_proj"],
    use_gradient_checkpointing="unsloth")

```

Llama-Factory and Axolotl are the other free, mature options; both are configuration-file driven, both wrap the same underlying libraries, and both are worth preferring over a hand-written loop for a first project.

Surviving a session that will be killed

Assume the session dies. Two habits make that a nuisance rather than a disaster:

```

TrainingArguments(save_steps=200, save_total_limit=3,
    output_dir="/content/drive/MyDrive/run-01")

```

Save to mounted Drive or push to the Hugging Face Hub, not to the ephemeral local disk. And set `resume_from_checkpoint=True` so restarting the notebook continues rather than restarting. A LoRA checkpoint is a few hundred megabytes, so keeping three of them is not a storage problem.

When you do pay, pay properly

Rented compute is cheaper than most people expect, and the community marketplaces are cheaper than the hyperscalers by a wide margin.

- An RTX 4090 (24 GB) on a community marketplace: roughly \$0.30 to \$0.60 an hour.
- An A100 40 GB: roughly \$1.00 to \$1.80 an hour.
- Interruptible or spot instances are typically half of those, and your checkpointing habit already makes interruption survivable.

A concrete job: 3,000 examples, 8B model, QLoRA, three epochs, sequence 1,024. That is around 560 optimiser steps at an effective batch of 16, and on a 4090 it finishes in one to two hours. **Under a pound.** The GPU has never been the expensive part of this; the dataset and the evaluation are, and they are paid in attention rather than in currency.

The genuinely free path, end to end

Build the dataset in a text editor. Train the adapter on Kaggle. Store checkpoints on the Hugging Face Hub, free for public repositories. Merge and convert to GGUF, then run the result on your own laptop with `llama.cpp` or Ollama — a 7B model in 4-bit needs about 5 GB of RAM and runs on a CPU at a few tokens per second, which is slow but is a real, working, private model on hardware you already own.

Nothing in that chain costs money. It is slower than the rented version at every step, and it produces the same artefact.

BEFORE YOU MOVE ON

Someone runs an 8B QLoRA on a free Colab T4 and sets `bnb_4bit_compute_dtype=torch.bfloat16` after copying a config from a blog. What happens?

- A. Nothing changes; bf16 and fp16 are interchangeable at 4-bit quantisation because the weights are stored in NF4 either way.
- B. Memory use doubles, because bf16 requires an fp32 master copy that fp16 does not.
- C. Training proceeds correctly but the resulting adapter can only be merged into a bf16 base, making it unusable on the same T4 afterwards.
- D. The T4 is Turing and has no bf16, so the setting errors or is emulated, and fp16 is what the hardware wants.

CASE STUDY · MODULE 3

One card, in the building

A 180-bed hospital in Nashik, adapting a model to turn dictated ward notes into a structured discharge summary.

The hospital had one constraint that removed most of the options: patient notes could not leave its premises. No hosted model, no rented card in somebody else's data centre, no batch endpoint. Whatever ran, ran on a single 24 GB workstation card that the radiology department had bought for something else and was willing to share overnight.

Dr Menon, who ran the project with one contract engineer, had three thousand dictated notes with a matching discharge summary written by a registrar. The task was narrow and well specified: take the dictation, produce eleven fields in a fixed order, in the hospital's house phrasing, with the medication list formatted as the pharmacy wanted it.

The engineer, Faiz, laid out three options on one page, with the memory arithmetic done before anything was rented or downloaded.

Full fine-tuning an 8B model was out. Mixed precision with AdamW costs about sixteen bytes per parameter of persistent state — four for the fp32 master weights, four for the gradients, eight for AdamW's two moment buffers. That is 128 GB before a single activation, on a card that has 24. Sharding across cards was not available; there was one card.

QLoRA on an 8B model fitted. The frozen base in NF4 is about 4.1 GB, the adapter and its optimiser state under a gigabyte, activations with gradient checkpointing a few more. Call it ten or eleven gigabytes, comfortably inside the card. The costs were real but bounded: roughly a quarter to a third lower throughput than an unquantised LoRA, because every matrix multiplication dequantises its weights on the fly; and an adapter fitted against a four-bit base, which makes merging into the full-precision base an approximation rather than an identity.

Full fine-tuning a 1.5B model also fitted. Sixteen bytes per parameter on 1.5 billion parameters is 24 GB, which is the whole card — but with 8-bit AdamW the optimiser state term falls from eight bytes to two, giving ten bytes per parameter, or 15 GB, with room for activations. Simpler: no adapters, no quantisation, no merging question at the end, one ordinary model to serve.

The decision was not obvious, and both sides had a cost the hospital would feel.

The 8B model would write better prose. Discharge summaries are read by general practitioners who did not see the patient, and fluency is not decoration there. But serving it

meant the same card during the day, running a quantised model at a few requests a minute, and a merging step whose effect on quality nobody had measured.

The 1.5B model would be markedly faster to train and to serve, would run on the ward's own machines if it ever needed to, and would be one artefact with no adapter to keep aligned with a base revision. Against that: a small model collapses rather than degrades when an input is unlike anything it saw. A dictation that wandered — and registrars' dictations wander — would get triage-shaped output confidently.

Faiz did what the module recommends and ran the cheap version first. A QLoRA on the 8B model over the three thousand notes, two epochs, effective batch of sixteen, took under three hours overnight and told him something he had not expected: the eleven-field structure was learned almost perfectly within one epoch, and what the second epoch added was memorisation of individual registrars' phrasings.

STOP HERE

Before turning the page: what would you have done, and what would it have cost if you were wrong? Write it down. The point of the next section is to compare.

WHAT HAPPENED

Both models were trained and both were evaluated against the same two hundred held-out notes, split by registrar so that no author appeared on both sides.

The 8B QLoRA scored 0.94 on field-level exact match and its prose was preferred by two consultants in a blind pairwise comparison at a rate of about two to one. The fully fine-tuned 1.5B scored 0.91 on fields, and its prose lost the comparison.

What decided it was a third number neither of them had planned to collect. On fifty deliberately awkward dictations — interrupted, bilingual, with a correction spoken mid-sentence — the 1.5B model produced a confident, well-formatted, wrong summary in eleven cases. The 8B model produced a summary that was visibly incomplete in nine, which a registrar could see at a glance.

The hospital shipped the 8B QLoRA, unmerged, served with the adapter separate so that turning the fine-tune off is a flag rather than a redeploy. Faiz measured the merged version too, out of discipline, and recorded that merging into the full-precision base cost about one point of field accuracy. That number is written in the model card, and it is the reason the adapter is still unmerged.

Full fine-tuning the 8B model was ruled out with arithmetic rather than an experiment. Show the arithmetic, and say why doing it on paper mattered here specifically.

Mixed-precision AdamW holds about sixteen bytes per parameter of persistent state: four bytes of fp32 master weights, four of fp32 gradients, and eight for AdamW's two moment buffers. Eight billion parameters times sixteen bytes is 128 GB, before activations or the logits tensor, against a 24 GB card and no second card to shard onto. It mattered here because the hospital could not rent its way out of a wrong answer — there was exactly one card in the building, so an option that did not fit was not merely expensive, it was impossible.

The failure that decided the comparison was not the headline metric. What was the mechanism behind it?

A small model has less capacity to represent behaviours outside the narrow one it was trained for, so off-distribution it collapses rather than degrading: it maps an unfamiliar input onto the shape it knows and produces confident, well-formatted output that is wrong. The larger model, with more behaviours available, produced visibly incomplete summaries instead — a failure a registrar could see. On a document a general practitioner will rely on without having seen the patient, a visible failure is far cheaper than a plausible one.

Why did Faiz measure the merged model even though the hospital was going to serve the adapter separately?

Because the adapter was fitted against a four-bit base and merging it into the sixteen-bit base is not the same arithmetic, so merging can move quality by a point or two in either direction. Measuring it produced a number for the model card and made a future decision cheap: if somebody later needs a merged model for a serving engine without adapter support, the cost of that change is already known rather than being discovered under pressure. It is also the general rule of the course — evaluate the exact artefact you would serve, not a related one.

Check what stuck

6 questions, one right answer each. This one is for you, not for a record: answers and the reason behind each are at the back. If two questions from the same module go wrong, re-read that module before the exam rather than after.

- 1 Mixed-precision full fine-tuning with AdamW costs about 16 bytes per parameter of persistent state. Which four things make up that figure?
 - A. fp32 master weights, fp32 gradients, and AdamW's two moment buffers.
 - B. The quantised base, the adapter, the optimiser state and the attention matrix.
 - C. Four bytes each for the weights, the gradients, the activations, and the accumulated history of the learning-rate schedule.
 - D. bf16 weights, bf16 gradients, the activations and the logits tensor.

- 2 A 3B model is being fine-tuned on a 24 GB card. Why is QLoRA the wrong choice here?
 - A. Because NF4 is only defined for models above seven billion parameters.
 - B. Because an adapter trained against a quantised base cannot be merged afterwards at all.
 - C. Because the model fits in bf16, so you would pay 25 to 40 per cent of throughput for memory you did not need.
 - D. Because quantisation error accumulates with depth, and a model of that size has too many layers for four bits to be safe on a card with that much memory.

- 3 You have a fixed budget of trainable parameters for a LoRA. Where does the evidence say to spend it?
 - A. On a high rank over the attention query and value projections, as the original LoRA work did.
 - B. Spread across all seven linear projections at a modest rank.
 - C. On the embedding table and the output head, where most task-specific behaviour lives.
 - D. It makes no difference where it goes, because in a low-rank method the total number of trainable parameters is the only thing that determines capacity.

- 4 Why should alpha be fixed at twice the rank rather than swept alongside the learning rate?
- A. Because a higher alpha increases the adapter's memory footprint, which makes a sweep over it expensive on any card small enough to need a LoRA in the first place.
 - B. Because alpha stops having any effect once the rank is above sixteen.
 - C. Because PEFT ignores any value of alpha that is not exactly twice the rank.
 - D. Because the update is scaled by alpha over r times the learning rate, so sweeping both sweeps one thing twice.
- 5 Your trainable-parameter count prints zero and the loss curve is flat and innocent. What has gone wrong?
- A. The target module names do not exist in this architecture, so no adapter was attached.
 - B. Gradient checkpointing has been left on, which freezes the adapter's parameters.
 - C. The dataset was shuffled with the wrong seed, so every batch contains the same examples and the gradients cancel out across the accumulation window.
 - D. The learning rate is too low for the adapter to move at all.
- 6 BitFit trains only the bias terms of a model. Why is it close to a no-op on a Llama-style decoder?
- A. Because the method was designed for classification heads and the library refuses to attach it to a causal language modelling task type at all.
 - B. Because its biases are shared across layers, so training them changes nothing.
 - C. Because RMSNorm has a scale and no bias, and the projections are configured without bias terms, leaving a few thousand parameters.
 - D. Because decoder models apply their biases only during generation and not during training.

MODULE 4

Building the dataset

The dataset is the model. This module covers where examples legitimately come from, the formats they must be converted into, the filters that decide quality, and the record that lets somebody else rebuild the set.

BY THE END OF THIS MODULE YOU CAN

Assemble a training set from named sources with checked licences, convert it to one chat format, filter and deduplicate it, mix in replay data at a stated ratio, and publish a data card that lets a stranger reproduce the set

28 · 11 min

The Dataset Decides Everything

THE ONE THING TO KEEP

Your model learns your dataset's whole distribution, including the habits you never meant to teach.

Build the test set before the training set

Do this first, before you write a single training example.

Take 100-300 real inputs from real traffic. Not invented ones — invented inputs are always cleaner, shorter and more polite than what people actually send. Write or approve the ideal output for each. Freeze the file. Never train on it. Split it in half: a dev half you look at while iterating, and a test half you open exactly twice, once at the start to record the base model's score and once at the end.

Teams that skip this cannot answer the only question that matters at the end, which is whether the work helped. They ship on vibes, and vibes favour whatever they just spent a week building.

Then write twenty by hand

Yourself. Not generated, not delegated. You will discover that the task you thought was well defined has four cases you had never decided, and that two of your "obviously correct" outputs contradict each other. Better to find that now than inside a 4,000-row dataset where the contradiction is trained in.

Those twenty become the style guide for everything that follows.

The model learns the whole distribution

Including the parts you were not thinking about.

- If 30% of your answers open with "Certainly," the tuned model opens with "Certainly."
- If your answers average 400 words because you wrote them carefully, it will write 400 words in reply to a yes-or-no question.
- If every input is a well-formed question, it will handle a typo-ridden fragment badly.
- If no example says "I don't know," it will never say it.

So vary the **inputs** hard — length, register, typos, ambiguity, off-topic, hostile, and the other languages you serve — and keep the **outputs** consistent in the ways you care about. Include the hard cases, the refusals, the ones where the right answer is a clarifying question. A dataset of easy cases produces a model that is confident on hard ones.

Deduplicate

Near-duplicates inflate your evaluation and waste training. Exact dedup on normalised text takes a minute:

```
import json, hashlib
seen, out = set(), []
for line in open("train.jsonl"):
    row = json.loads(line)
    text = row["messages"][0]["content"]
    key = hashlib.sha1(" ".join(text.split()).lower().encode()).hexdigest()
    if key not in seen:
        seen.add(key)
        out.append(row)
print(len(out))
```

Then check the same way for overlap between train and test. If a test input also appears in training, your score is fiction.

Synthetic data, honestly

Generating training data by prompting a stronger model is normal, effective, and widely done. Two conditions. Check the provider's terms, because several forbid using outputs to train competing models. And understand what you are buying: the stronger model's style, its mistakes and its blind spots, transferred faithfully. Have a human read and fix a sample. A 500-example set you have actually read beats a 5,000-example set nobody has.

Never draw the test set from the same generator. It would measure agreement with the generator, not quality.

How many examples

Honest ranges, assuming clean data and LoRA on a competent instruct base:

- **50-200** — tone, register, format, output schema. This works better than people expect, and it is where most teams should stop.
- **500-2,000** — a narrow task with real accuracy gains: routing, extraction, structured summarisation in your house form.
- **5,000-50,000** — a genuinely new behaviour, or wide domain coverage. Expect weeks of data work.

LIMA is the paper everyone cites for "a thousand examples is enough." Read what it measured: human preference for helpfulness and style, starting from a strong 65B base. It is real evidence about style. It is not evidence that a thousand examples teaches a task that requires being right.

Find your own number with an ablation

Do not guess. Train on 25%, 50% and 100% of your data with everything else fixed, and score each on the dev set.

- Still climbing at 100%? More of the same data will help.
- Flat between 50% and 100%? More of the same will not. You need *different* data — harder cases, different input styles — or the bottleneck is somewhere else entirely.

Three small runs cost a few dollars and replace an argument that would otherwise last a month.

One aside worth the whole lesson

If your task is classification over a fixed label set and you have 5,000 or more labelled examples, fine-tune a small encoder model — the BERT family, 100-400M parameters — before you fine-tune an LLM. It is often more accurate, it serves on a CPU for a fraction of the cost, and it trains in twenty minutes. The unfashionable answer is frequently the right one.

How many examples, by what you are teaching

50 to 200 examples	Tone, register, format, an output schema. This works better than people expect, and it is where most teams should stop.
500 to 2,000	A narrow task with real accuracy gains: routing, extraction, structured summarising in your house form.
5,000 to 50,000	A genuinely new behaviour, or wide domain coverage. Expect weeks of data work, not a weekend.

LIMA is the paper cited for "a thousand is enough". It measured human preference for helpfulness and style from a strong 65B base — real evidence about style, and not evidence that a thousand examples teaches a task that requires being right.

BEFORE YOU MOVE ON

You generated 4,000 training examples by prompting a strong model, held out 400 of them, and scored 92%. In production the model is clearly worse than that number suggests. What is the most likely explanation?

- A. The model overfit; train for fewer epochs.
- B. 4,000 examples is too few for this task; generate 40,000.
- C. The base model was contaminated with the evaluation data during pretraining.
- D. The held-out set came from the same generated pool, so it measures agreement with the generator on generator-shaped inputs rather than performance on real ones.

Where the data legitimately comes from

THE ONE THING TO KEEP

Your own traffic has the right distribution, hand-written examples cover what traffic never contains, and public sets are for replay — with a dataset's own licence and the terms under which its contents were generated being two separate permissions that a provenance column recorded at collection time is the only way to answer later.

Four sources, in descending order of usefulness

1. Your own traffic. Real inputs from real people doing the real job. Nothing else has the right distribution — the typos, the half-sentences, the questions that do not fit your categories. Sample it deliberately: stratify by intent, language and length rather than taking the last thousand rows, which cover the last twenty minutes of one time zone.

The obligations attach here, not later. Check what your privacy notice promised, redact identifiers before anything lands in a training file, and record the basis on which you are using it. A training set is a copy of your users' words with the access controls taken off.

2. Adjacent real text. Support tickets, forum threads, internal documentation, past reports — text produced by people doing the task, even if never typed at a model. It has the vocabulary and the register, and it needs conversion into instruction shape, which is work.

3. Team-written examples. Slower, higher quality, and the only source that can cover a case which has never occurred. Twenty hand-written examples at the start of a project are worth more than two thousand scraped ones, because writing them forces you to decide what a good answer actually is — and you will change your mind twice while doing it.

4. Public instruction datasets. Useful as a replay mixture rather than as your task data. Which brings the licence question.

The public sets worth knowing, and their terms

Licences change; check the dataset card for the current terms before you rely on any of this.

- **databricks-dolly-15k** — 15,000 prompt-and-response pairs written by Databricks employees, released CC BY-SA 3.0. Human-written, so no question about another provider's output terms. The share-alike condition is the thing to think about.
- **OpenAssistant (oasst1, oasst2)** — human-generated, multilingual, crowd-sourced conversation trees, Apache-2.0. The most permissive substantial human-written set available.
- **No Robots** — 10,000 human-written instruction-and-demonstration pairs, deliberately produced without model assistance. Released under a non-commercial licence, which is exactly the sort of detail that matters and is routinely overlooked.
- **UltraChat and similar large synthetic sets** — permissively licensed as *datasets*, but generated by calling a hosted model, so the generating provider's terms are a separate question from the dataset's licence.

That last distinction is the one that catches teams. A dataset's own licence and the terms under which its contents were produced are two different permissions, and a permissive dataset licence cannot grant something the producer did not have.

The provenance column

One extra field per example, written at collection time, costs nothing and answers every question anybody will later ask:

```
{"messages": [...],
 "source": "support_logs",
 "licence": "internal",
 "basis": "tos_v4_analytics",
 "collected": "2026-08-14",
 "reviewed_by": "rk"}
```

Six months later, somebody asks whether the model was trained on customer data. With this column it is a query. Without it, it is an investigation that ends in a guess.

How much of each

A shape that works for a task fine-tune:

- 60 to 80% task data from your own traffic or adjacent text
- 10 to 20% hand-written examples covering the cases traffic does not contain — the edge cases, the refusals, the "I do not have that information" responses
- 10 to 20% general instruction data as replay, to stop the model forgetting how to be a model

The middle slice is the one people skip, and it is the one that determines how the model behaves when something unusual arrives. If every training example is a question the model can answer, you have taught it that every question is answerable. That is exactly the behaviour people call hallucination, and you installed it.

What not to do

Do not scrape a competitor's product outputs. Do not use a dataset because it is popular without reading its card. Do not take the last hundred thousand log rows because they were easy to get.

And do not let dataset size become the goal. The LIMA result — a strong instruction-following model fine-tuned on a thousand carefully curated examples — is the standing evidence that curation beats volume at this scale. A thousand examples you have read is a better dataset than fifty thousand you have not.

BEFORE YOU MOVE ON

A team builds a training set entirely from support conversations that were resolved successfully. The tuned model becomes markedly worse at saying it does not know something. What caused this?

- A. Every example showed an answerable question being answered, so the model learned that answering is the job.
- B. Successful conversations are longer, so the model learned to produce long answers, and length correlates with overconfidence.
- C. Support logs contain company-specific facts the model cannot store in weights, so it substitutes invented ones.
- D. Filtering to resolved conversations reduced the dataset size below the threshold at which a model retains its pretrained refusal behaviour.

30 · 8 min

The three formats, and the conversion trap

THE ONE THING TO KEEP

Convert everything to messages format once on the way in, use the same formatting function for training and inference, and print one templated example as a string before every run — because the commonest data bugs are a changed separator, a half-present system prompt, and truncation that removes the end-of-turn token.

Three shapes you will meet

Instruction data comes in three formats. You will have to convert between them, and the conversion is where errors enter.

Alpaca style. Three fields, one of which is optional:

```
{"instruction": "Summarise the complaint in one sentence.",  
  "input": "I ordered on the 3rd and it still hasn't...",  
  "output": "Customer reports a delayed order placed on the 3rd."}
```

ShareGPT style. A conversation list with its own vocabulary:

```
{"conversations": [{"from": "human", "value": "..."},  
                  {"from": "gpt", "value": "..."}]}
```

Messages style. What the OpenAI API popularised and what has become the default:

```
{"messages": [{"role": "system", "content": "You are a triage  
assistant."},  
              {"role": "user", "content": "..."},  
              {"role": "assistant", "content": "..."}]}
```

Use messages style. TRL, Axolotl and Llama-Factory all accept it natively, it represents multi-turn and tool calls without extension, and it maps cleanly onto every model's chat template. Convert everything to it on the way in, once, and never think about the others again.

The trap in Alpaca's input field

Alpaca's two-part prompt has to become one user message. Almost everybody concatenates them, and the question is *how*.

The original Alpaca template inserted specific headers — an instruction block, then an `### Input:` block, then `### Response:`. If you concatenate with a newline instead, you have changed the prompt format. That is fine, as long as you do it identically at inference time. It is a disaster if training used one separator and serving uses another, because the model learned to respond to a pattern that never arrives.

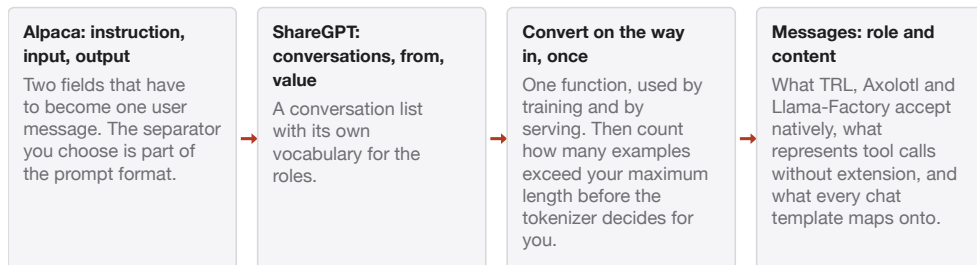
```
def to_messages(row):
    user = row["instruction"]
    if row.get("input"):
        user += "\n\n" + row["input"]      # decide once, apply
    everywhere
    return {"messages": [{"role": "user", "content": user},
                        {"role": "assistant", "content":
row["output"]}]}
```

The rule that prevents the whole class of bug: **the code that formats training examples and the code that formats inference requests must be the same function**. Not similar. The same function, imported from one module, used by both.

The system prompt decision

Three options, and they have different consequences:

Three formats in, one format used



The rule that prevents a whole class of bug: the code that formats training examples and the code that formats inference requests must be the same function, imported from one module.

1. **The same system prompt on every example, and at inference.** The model learns to rely on it. Fine, and it is what most production systems do. Do not then change the system prompt after training without re-evaluating.

2. **No system prompt anywhere.** Also fine, and the simplest. The model's behaviour is entirely in its weights.
3. **Varied system prompts.** Genuinely useful when you want the model to remain steerable — different tones, different output formats, selected at request time. Costs more data, because each variant needs examples.

The failure is mixing them by accident: half your examples carrying a system prompt because they came from a log that included one, half not. The model then learns that the system prompt is optional and roughly ignorable, which is the opposite of what you wanted.

Validate before you train

Four checks, worth the ten minutes:

```
for i, row in enumerate(data):
    m = row["messages"]
    assert m[-1]["role"] == "assistant", f"{i}: must end with as-
sistant"
    assert all(x["content"].strip() for x in m), f"{i}: empty content"
    roles = [x["role"] for x in m if x["role"] != "system"]
    assert roles == ["user", "assistant"] * (len(roles)//2), f"{i}: bad
alternation"
    assert len(tok.apply_chat_template(m)) <= MAX_LEN, f"{i}: too long"
```

The last one matters more than it looks. Examples over your maximum sequence length are silently truncated, and truncation removes the *end* of the assistant's answer — including the end-of-turn token. Train on a few hundred of those and you have taught the model not to stop. Count how many examples exceed the limit before you train, and if it is more than a couple of percent, either raise the limit or drop those examples deliberately rather than letting the tokenizer decide.

The one-line test that catches most of it

Before any run, take one example, apply the chat template, and print the string:

```
print(tok.apply_chat_template(data[0]["messages"], tokenize=False))
```

Read it with your eyes. Every special token should be where you expect, the roles should be marked correctly, and there should be exactly one end-of-turn marker at the end. Five seconds, and it catches format bugs that would otherwise cost you a run and a day of confusion.

BEFORE YOU MOVE ON

A team trains on Alpaca-format data converted with a newline separator, then serves with the original Alpaca template's '### Input:' headers. Quality is far worse in production than in evaluation. Why?

- A. The '### Input:' header tokens consume context, leaving less room for the actual input at inference time.
- B. Evaluation used the training format and production uses another, so the prompt pattern was never learned.
- C. The headers contain '#' characters, which some tokenizers treat as markdown and split differently from plain text.
- D. Alpaca's template was designed for a different base model, so its headers conflict with the chat template's special tokens.

31 · 9 min

Cleaning at scale

THE ONE THING TO KEEP

Filter in order of cost — exact hashes, then MinHash near-duplicates, then length, truncation, language and repetition heuristics — and then read a hundred examples by hand, because the systematic annotation problems that decide a fine-tune's quality have no automated detector.

A bad example is worse than a missing one

A missing example teaches nothing. A bad example teaches something wrong, with the same weight as a good one, and the model has no way to tell them apart. Every filter below exists because somebody trained on data they had not looked at.

Work through these in order — cheap deterministic filters first, judgement last.

Exact duplicates

Hash the normalised text and keep one copy. On data assembled from several sources, exact duplicates are routinely 5 to 20% of the file.

```
import hashlib
seen, kept = set(), []
for row in data:
    key = hashlib.sha256(
        " ".join(m["content"].lower().split()).encode()
        for each-message # normalise whitespace and case before hash-
ing
    )
```

The normalisation matters: without lowercasing and collapsing whitespace, two copies of the same example differing by a trailing space survive as distinct rows.

Near duplicates

The harder and more important case. Four hundred tickets that differ only in an order number are, for training purposes, one example repeated four hundred times — and the model will weight it accordingly.

MinHash with locality-sensitive hashing finds these in near-linear time:

```
from datasketch import MinHash, MinHashLSH

lsh = MinHashLSH(threshold=0.85, num_perm=128)
for i, text in enumerate(texts):
    m = MinHash(num_perm=128)
    for token in set(text.lower().split()):
        m.update(token.encode())
    if not lsh.query(m): # nothing similar already kept
        lsh.insert(str(i), m)
        keep.append(i)
```

`datasketch` is free and handles a few hundred thousand documents on a laptop. The `text-dedup` package wraps the same idea with more options if you have millions.

A threshold of 0.85 Jaccard similarity is a reasonable start. Check the result by hand: print ten pairs the filter considered duplicates, and ten it kept as distinct. If the boundary looks wrong, move the threshold rather than trusting it.

Length and truncation

Two filters, both cheap, both catching real damage:

- **Too short.** An assistant response under about 5 tokens is usually a logging artefact, not an answer. Unless one-word answers are your task, drop them.

- **Truncated.** An answer that ends mid-word or mid-sentence came from a generation that hit a token limit. These teach the model to stop in the middle of things. Check whether the final character is sentence-terminating punctuation, or a closing brace for structured output.

Also count how many examples exceed your training `max_length`, because those will be truncated by the tokenizer, losing the end-of-turn token along with the last part of the answer.

Language

If your product is monolingual and 4% of your logs are another language, decide deliberately rather than by accident. `fasttext`'s language identification model is free, runs at tens of thousands of documents a second, and gives a confidence score.

The reverse case matters too: if you *are* multilingual, check the balance. A set that is 92% English will produce a model whose other languages have quietly got worse.

Degenerate text

A handful of heuristics catch most machine-generated junk:

- **Repetition ratio.** If the most common 3-gram accounts for more than about 10% of the 3-grams in a response, it is looping.
- **Non-word character ratio.** A high proportion of punctuation or symbols usually means markup, base64 or a corrupted encoding.
- **Model tics.** If your data came from a hosted model, search for its characteristic openings — the self-identification phrases, the hedging preamble, the stock apology. You do not want to train those in. A regex over the first sentence finds them.

The filter you cannot automate

Read a hundred examples. Not skim — read, from the beginning to the end of the assistant turn.

This is the highest-value hour in the whole project and it is the one most reliably skipped. What it finds is not the things above; it is the systematic problems no filter names. Every answer beginning with the same three words. Annotators who disagreed about whether to include a greeting. Two hundred examples where the "correct" answer is subtly wrong because the person writing them misunderstood the task.

Keep the record

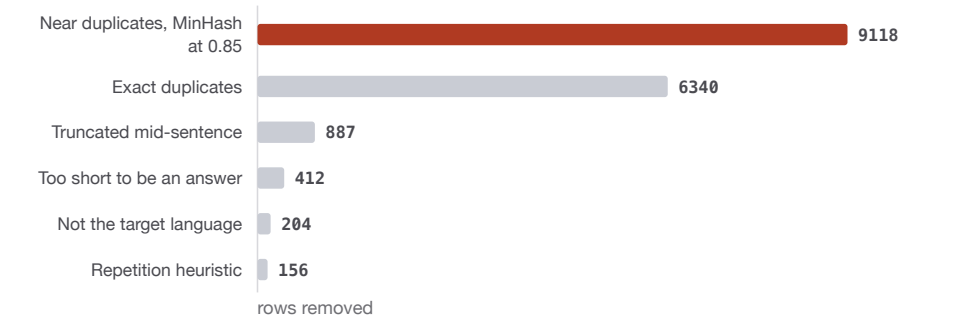
Log what each filter removed and how much:

start	48,201
exact duplicates	-6,340
near duplicates (0.85)	-9,118
too short	-412
truncated	-887
non-target language	-204
repetition heuristic	-156

final	31,084

If one filter removes 40% of your data, that is not a cleaning step, that is a discovery about your source, and it deserves a look before you accept it.

What each filter removed, from 48,201 rows



Thirty-one thousand survived. Log the counts: a filter that removes forty per cent of your data is a discovery about your source, not a cleaning step.

BEFORE YOU MOVE ON

A dataset of 40,000 support exchanges contains 600 automated dispatch notices that differ only by an order number. Exact-duplicate hashing removes none of them. What is the consequence for training, and what finds them?

- A. They are harmless because each is genuinely a distinct example; only byte-identical rows distort training.
- B. They cause the tokenizer to build an unbalanced vocabulary around order numbers, which length filtering resolves.
- C. They act as one pattern repeated 600 times, and near-duplicate detection such as MinHash catches them.
- D. They will be caught by the repetition-ratio heuristic, since the repeated template raises the 3-gram concentration within each document.

32 · 9 min

Writing the answers, and keeping them consistent

THE ONE THING TO KEEP

The model learns the distribution of your answers rather than their correctness, so unforced variation between annotators becomes randomness in production — which makes a one-page style guide with exact wording for the declining case worth more than another thousand examples.

Rewriting beats writing

Producing a gold answer from a blank page takes most people four to eight minutes. Taking a model's draft and correcting it takes one to two. The quality is comparable and often better, because the draft reminds the annotator of details they would have omitted.

So the standard workflow is: generate a draft with the best model you have, then have a person fix it. Do this and you will produce three to five times more data for the same human hours.

The risk is anchoring. A reviewer presented with a fluent, confident, wrong answer accepts it more often than they would reject a blank page. Two countermeasures, both cheap:

- Have the reviewer state the answer's key claim *before* reading the draft on a sample of items, and compare.
- Track the edit rate. If more than about 70% of drafts are accepted unchanged, either the model is genuinely good — check by having a second person audit 30 of them — or your reviewers have stopped reading.

Consistency is the thing you are teaching

This is the point people miss. The model is not learning "the right answer"; it is learning the *distribution* of your answers. If half your annotators start with a greeting and half do not, the model learns that greetings are a coin flip, and it will flip that coin in production.

Everything that varies across your examples and is not meant to be information becomes noise the model faithfully reproduces:

- Greeting or no greeting
- Bullet points or prose
- Whether the reasoning is shown
- British or American spelling
- Whether uncertainty is expressed, and in what words
- Length

Pick one of each, write them down, and enforce them.

The style guide is a page, not a document

`## Answer style – ticket triage`

1. No greeting, no sign-off. Start with the substance.
2. British spelling.
3. One paragraph under 60 words, unless listing steps; then a numbered list.
4. Never apologise for the product. State the position.
5. If the answer is not in the provided context: "I don't have that information – I'll pass this to the team." Exactly those words.
6. Never invent an order number, a date or a policy figure.
7. Amounts as "£12.50", never "12.50 GBP" or "twelve pounds fifty".
8. If the customer is angry, acknowledge once in the first clause, then move on.

Eight rules. Every annotator reads it before their first example and re-reads it before their hundredth, because drift is real and mostly unconscious.

Rule 5 is the pattern worth copying: *exactly those words*. If you want a model to have a reliable way of declining, give it one form of words, not eleven variations on the sentiment.

Measuring agreement, cheaply

Give two people the same 30 items. Have them write answers independently. Then diff.

You are not computing a statistic — 30 free-text answers do not support one — you are finding the rules your style guide is missing. Every systematic disagreement is a decision nobody has made yet. Make it, add it to the guide, and continue.

Do this again at 300 examples and at 1,000. Annotator drift over a long labelling session is substantial and nobody notices it in themselves.

The adjudication rule

When two annotators disagree and neither is wrong, do not average and do not alternate. Pick one and record why. An averaged style is a style the model cannot learn because it never sees it.

What to do with the disagreements you cannot resolve

Some items are genuinely ambiguous: the customer's request could reasonably be read two ways. These are not annotation failures, they are product questions, and they should be escalated rather than resolved by whoever happens to be labelling.

Meanwhile, take them out of the training set. An ambiguous item with an arbitrary answer teaches the model to be arbitrary at exactly the point where you most want it to ask a clarifying question — and if that is the behaviour you want, then write *that* as the gold answer, consistently, and you will have taught something useful instead.

The number that tells you when to stop

Stop hand-writing when a fresh batch of 50 examples produces no new style-guide rules and no new disagreements. At that point the remaining variance is in the inputs, not in your understanding of the task, and volume from cheaper sources will help more than another week of careful writing.

BEFORE YOU MOVE ON

Two annotators write equally good answers, but one always opens with a brief acknowledgement of the customer's frustration and the other never does. Both sets go into training. What does the model learn?

- A. It learns the acknowledgement is appropriate, since exposure to the behaviour in half the data is enough to establish it as the norm.
- B. It averages the two styles into a shortened acknowledgement present in most answers.
- C. Nothing harmful — an even split acts as data augmentation and makes the model robust to both styles.
- D. It learns the acknowledgement appears about half the time, with nothing in the input predicting when.

33 · 10 min

Generating a dataset from a teacher

THE ONE THING TO KEEP

Draw the inputs from real traffic and only the outputs from the teacher, sample several completions and filter them with a verifier wherever correctness is checkable, then strip the teacher's stock phrases and check that your generated set covers your traffic's distinct intents rather than just its volume.

Start from real inputs

The most common mistake in synthetic data generation is synthesising the inputs as well as the outputs. Ask a model for "200 customer support questions" and you get 200 well-formed, grammatical, politely-phrased questions that resemble nothing your users type. Train on those and you have a model tuned for a population that does not exist.

Inputs from traffic, outputs from the teacher. That is the shape. If you have no traffic, use adjacent real text — forum posts, search queries, tickets — before you resort to invented inputs. If you genuinely must invent, invent from a taxonomy you derived from real examples, and include the ugly ones: the two-word question, the paragraph with no question mark, the one in the other language.

Generating at volume, cheaply

For 50,000 outputs, per-request API calls are slow and often expensive. Two better routes:

Batch endpoints. Most hosted providers offer an asynchronous batch mode at roughly half price with a turnaround measured in hours. If your generation is not interactive — and it is not — this halves the bill for the cost of waiting.

A local teacher. Rent one large card for a day, serve an open-weight 70B model with vLLM, and generate offline:

```
from vllm import LLM, SamplingParams
llm = LLM(model="Qwen/Qwen2.5-72B-Instruct", tensor_parallel_size=2)
outs = llm.generate(prompts, SamplingParams(temperature=0.7,
max_tokens=512, n=4))
```

vLLM's continuous batching gives throughput that makes this genuinely cheap — tens of thousands of completions per GPU-hour for short outputs. And a permissively licensed teacher removes the question about whether you may use its outputs at all.

Temperature, and the diversity question

Generate at temperature 0 and every similar input gets a near-identical answer; your dataset has less variety than its row count suggests, and the student learns a narrow band of phrasing.

Generate at temperature 1.0 and you get variety including the teacher's bad days.

The usual compromise is 0.7 with `n=4` samples per input, then a filter to pick among them. Which brings us to the part that matters.

Filtering is where the quality comes from

Ranked by strength:

1. A verifier. Where correctness is checkable by a program — code that must run, JSON that must validate against a schema, arithmetic you can recompute, SQL you can execute against a test database — keep only outputs that pass. This is the strongest filter available and it costs nothing but CPU. Where you can construct one, do.

```
keep = [o for o in outputs if schema_validates(o) and
no_placeholder(o)]
```

2. Self-consistency. Sample four answers. If three agree on the substantive claim, keep that one; if all four disagree, drop the item. This works well for anything with a determinate answer and works not at all for open-ended writing.

3. A scoring model. Have a model rate its own or another model's outputs and keep the top-scoring one. Effective on writing tasks. Be aware of self-preference: a model asked to rank candidates tends to prefer its own, which matters if you are mixing teachers.

4. A human on a sample. Not a filter for 50,000 items, but read 100 of the *kept* ones. You are looking for a systematic teacher habit that survived the filter — a recurring preamble, a stock hedge, a wrong assumption the teacher makes consistently. You will find at least one in ten minutes, and finding it is worth more than the other three filters.

Strip the teacher's tics before training

Hosted models have characteristic openings and closings. If you train on unstripped outputs, your model inherits them, and they are exactly the phrases that make a product feel like a thin wrapper.

```
PREFIXES = ["Certainly! ", "Of course! ", "Sure, ", "I'd be happy to help"]
SUFFIXES = ["Let me know if you need anything else",
            "Is there anything else I can help with"]
```

Strip these deterministically before writing the training file, then check a sample by hand for the ones your list missed.

Two things to measure afterwards

Diversity. Count distinct 4-grams across the outputs and divide by total 4-grams. If that ratio is low, your dataset is more repetitive than its size suggests and the student will be narrower than you expect.

Coverage. Cluster the *inputs* you generated from — embeddings and k-means is enough — and check that the clusters roughly match the shape of your real traffic. A dataset that covers 80% of your traffic's volume but only 40% of its distinct intents will produce a model that is excellent on the head and useless on everything else. That failure arrives in production, where it looks like a quality problem and is actually a sampling one.

BEFORE YOU MOVE ON

A team generates 50,000 training examples by asking a teacher model to invent both the customer questions and the answers. The student performs well in evaluation on the same synthetic distribution and poorly in production. What is the mechanism?

- A. The synthetic inputs are well-formed in a way real user text is not, so production is a different distribution.
 - B. Synthetic data always contains subtle factual errors that compound during training, degrading the student below the teacher.
 - C. The student memorised 50,000 examples, and memorisation does not transfer to new inputs.
 - D. Teacher-generated answers are longer than human ones, and production users abandon long answers before reading them.
-

34 · 9 min

Mixtures and replay

THE ONE THING TO KEEP

Mixing 15 to 30% general instruction data into a task-specific set preserves general capability better than any training-side mitigation, because the weights supporting that capability keep receiving a signal to stay put — and shuffling the file is not optional when the learning rate decays.

Why a pure task dataset is a trap

Train a model on 4,000 examples that are all ticket triage, and you get a model that does ticket triage. You also get a model that has drifted towards producing triage-shaped output when asked anything, that has become worse at ordinary instruction-following, and that may have lost the refusal behaviour its base model had.

This is catastrophic forgetting, and the most effective countermeasure is not a training trick. It is putting other things in the data.

Replay, and the ratio

Replay means mixing examples of the general capability you want preserved into your task-specific set. A working starting point:

- **70 to 85% task data**
- **15 to 30% general instruction data**

The general portion comes from a permissively licensed public set — OpenAssistant's Apache-2.0 conversations are the usual choice — or, better where you can get it, from logs of the model doing the *other* things your product asks of it.

The mechanism is direct. Gradient descent moves weights towards whatever reduces loss on the batch. If every batch contains only triage, every weight is free to specialise. If a fifth of each batch is ordinary conversation, the weights that support ordinary conversation keep receiving a signal to stay where they are.

This costs nothing but dataset assembly, and in measured comparisons it does more to preserve general capability than lowering the learning rate, reducing epochs, or any adapter configuration.

Getting the ratio right

Run the ablation; it is two extra runs:

100% task tuning)	→ task metric 0.91, canary 0.62	(was 0.84 before
80/20 mixture	→ task metric 0.90, canary 0.81	
60/40 mixture	→ task metric 0.86, canary 0.83	

That shape — a large canary recovery for almost no task cost at 80/20, then diminishing returns — is what you usually see. The 80/20 row is the one to ship, and you only know it is there because you ran three configurations instead of one.

Balancing within the task data

The mixture question also applies inside your own data. If 70% of your traffic is one intent, and you train on traffic proportions, the model will be excellent at that intent and mediocre at the rest — which is a defensible choice, but it should be a choice.

Three options:

1. **Natural proportions.** The model's competence matches your traffic. Simple, and it optimises expected quality.
2. **Flattened.** Cap any single intent at some share — say 25% — so rare intents get enough examples to be learned. Costs a little on the head, buys a lot on the tail.

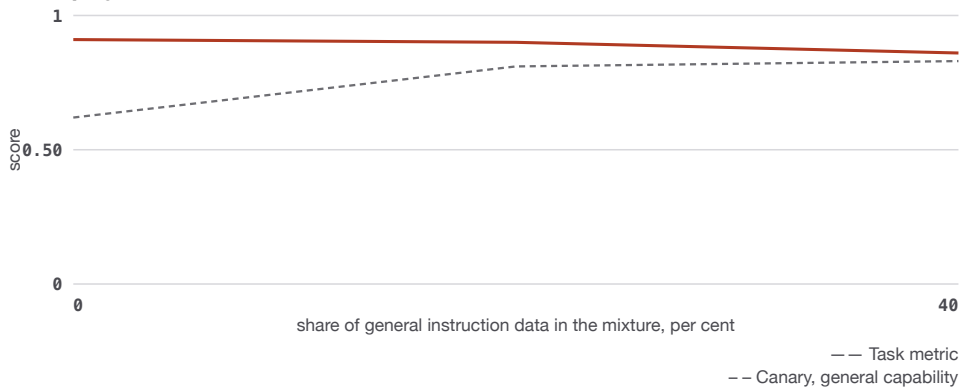
3. **Weighted by cost of failure.** If getting a refund wrong costs more than getting a shipping query wrong, over-represent refunds regardless of volume.

Most teams want option 2 and implement option 1 by not thinking about it.

Curriculum: mostly not worth it

Ordering examples from easy to hard has a long research history and produces inconsistent results on LLM fine-tuning. What is *not* optional is shuffling: if your file is sorted by source, or by date, or by intent, the model sees one distribution for the first thousand steps and a different one for the next thousand. With a decaying learning rate, whatever comes last gets learned with the smallest steps and whatever comes first gets partially overwritten.

The replay ablation, in three runs



The canary was 0.84 before any tuning. A large recovery for almost no task cost at eighty-twenty, then diminishing returns — and you only know the row is there because you ran three configurations instead of one.

```
dataset = dataset.shuffle(seed=0) # one line; not optional
```

This is a real bug that produces a mysteriously bad model, and the loss curve looks fine throughout.

Multi-task fine-tuning

If you have three tasks, one model trained on all three usually beats three separate adapters — provided the tasks are related. Shared structure transfers; the model learns your domain once rather than three times, and you maintain one artefact.

It stops being true when the tasks conflict: one wants terse structured output and another wants discursive explanation, and the model is asked to learn that the same kind of input maps to both. If your tasks pull in opposite directions, either separate them into different adapters or make the distinction explicit in the input — a system prompt or a tag that tells the model which mode it is in. Making it explicit is usually the better answer, because it keeps one artefact and gives you a controllable switch.

The rule to keep

Before you train, write down the proportions of your file by source and by intent, as actual counts. Almost nobody does this, and almost every surprising fine-tune result is explained by it.

BEFORE YOU MOVE ON

A team's training file is ordered by source: 3,000 triage examples followed by 1,000 general conversations. They shuffle nothing and use a cosine schedule. Why is this worse than the same data shuffled?

- A. Unshuffled data causes the optimiser's momentum buffers to saturate, which prevents later examples from producing gradients.
- B. The general examples arrive last, where the learning rate is smallest, so they influence the weights least.
- C. Batches drawn from a single source have lower gradient variance, so the effective learning rate is higher and the run diverges.
- D. The model will overfit the 3,000 triage examples in the first epoch before ever seeing general data, and later data cannot undo memorisation.

Packing, document boundaries and long context

THE ONE THING TO KEEP

Packing removes padding but lets one document attend to the previous one unless the block-diagonal mask and per-document position ids are actually in the batch — and extending context requires both a RoPE scaling change and long training data that, for most tasks, does not naturally exist.

The padding problem

If your examples average 300 tokens and your training window is 2,048, then padding every sequence to the window means roughly 85% of the tensor is padding. You pay for those matrix multiplications in full.

Two fixes, with different risk profiles.

Length-grouped batching (`group_by_length=True`) sorts similar-length examples together so each batch pads to its own longest member rather than to a global maximum. Safe, simple, and it recovers most of the waste.

Packing concatenates examples end to end until the window is full. Zero padding, maximum throughput — and one hazard that is invisible unless you look for it.

Attention across the seam

In a packed sequence, example B follows example A in the same tensor. With ordinary causal attention, the tokens of B can attend to every token of A. The model is being asked to predict the start of one customer's ticket having read a different customer's ticket, and it learns a transition that will never occur at inference.

The consequences are subtle rather than catastrophic — slightly worse quality, a model that sometimes continues past the end of an answer — and they never show up in the loss curve.

The correct handling has two parts:

1. **A block-diagonal attention mask**, so each document attends only within itself. Implemented efficiently through FlashAttention's variable-length path, which takes cumulative sequence-length offsets rather than a mask tensor.

2. Position ids reset per document, so the second document's first token is at position 0 rather than position 400. Without this, an example's position encoding depends on what happened to be packed before it.

Recent TRL versions handle both when packing is enabled, but implementations have differed over time and across libraries. Verify rather than assume:

```
batch = next(iter(trainer.get_train_dataloader()))
print(batch.keys())                # look for position_ids or cu_se-
q_lens
print(batch["position_ids"][0][:50] if "position_ids" in batch else
      "none")
```

If position ids climb monotonically across the whole window, documents are not being separated. If you cannot confirm it, use length-grouped batching instead: you lose some throughput and you lose the failure mode entirely.

One cheap mitigation regardless: make sure an end-of-turn or end-of-sequence token sits between packed documents, so there is at least a learned signal that one thing has finished.

Extending a model's context

A model pretrained at 8,192 tokens does not simply work at 32,768. Its rotary position embeddings encode positions it has never seen, and quality falls off a cliff past the trained length.

Context extension fine-tuning changes the position encoding and then trains briefly at the longer length. The usual approaches all adjust RoPE's frequency base:

```
config.rope_scaling = {"type": "yarn", "factor": 4.0,
                      "original_max_position_embeddings": 8192}
```

Linear interpolation is the simplest and loses the most local resolution; NTK-aware scaling and YaRN adjust different frequency bands differently and preserve short-range behaviour better. All of them want a short fine-tune on genuinely long sequences afterwards — typically a few hundred steps, which is cheap in steps and expensive in memory, because attention is quadratic and you are now at four times the length.

The data problem nobody solves well

Here is the honest difficulty: long-context supervised data barely exists. Real conversations are short. Real tickets are short. If 95% of your examples are under 1,000 tokens and you train with a 32,768 window, you have paid for a long context and taught the model nothing about using one.

The common synthetic route is to build multi-document tasks: concatenate several related documents and ask a question whose answer requires one of them, or several. It works, and it has a well-documented failure. If all your long-context training is of the "find the fact in the haystack" kind, the model becomes good at retrieval within its context and no better at *synthesising* across it — summarising a long document, tracking an entity through a transcript, reconciling two sources that disagree.

So build both kinds, and evaluate both kinds. And before any of it, check the sequence-length histogram of your data against the window you are about to pay for. If the histogram says you do not need the window, do not buy it.

BEFORE YOU MOVE ON

A team enables packing with an implementation that concatenates examples but does not reset position ids or mask across boundaries. What is the most likely symptom in the trained model?

- A. Training loss rises sharply, because the model is penalised for predicting tokens of one document from another.
- B. The model's maximum usable context shrinks, because position ids beyond the first document were never seen at their true values.
- C. The model sometimes fails to stop and continues into unrelated content, with no sign of it in the loss.
- D. Gradients from the second document are discarded, so only the first example in each packed window contributes to training.

Multi-turn conversations and tool calls

THE ONE THING TO KEEP

Compute loss on every assistant turn including the tool-call message, mask the tool's response because it is the world's output rather than the model's, and include conversations where no tool was needed — a dataset of only successful calls produces a model that always calls something.

Which turns carry loss

A ten-turn conversation contains five assistant messages. Three ways to train on it, and they are not equivalent:

Last turn only. Mask everything except the final assistant message. Simple, and it throws away four fifths of your supervision. For a 10-turn dialogue you get one training signal out of five available.

Every assistant turn. Mask all user and system and tool content; compute loss on every assistant message. Five signals from one conversation. This is what you almost always want.

Everything. Train on user turns too. Occasionally deliberate — if you want a model that can simulate a user, or for continued pretraining on dialogue — and usually a mistake, because you are teaching the model to write the human's side.

Doing it correctly

Building the mask by hand means finding the character offsets of each assistant span after templating, which is fiddly and breaks whenever the template changes. Modern chat templates solve this by marking the spans themselves with `{% generation %}` blocks, and the tokenizer will then return an assistant mask:

```
out = tok.apply_chat_template(messages, return_dict=True,
                             return_assistant_tokens_mask=True)
labels = [t if m else -100
          for t, m in zip(out["input_ids"], out["assistant_masks"])]
```

If the model's template has no generation markers, the collator has to locate the assistant response by string matching on the turn delimiters — which is what TRL's completion-only collator does, and which is why it needs you to pass exactly the right response template string. Get that string wrong by one character and it matches nothing, every position is masked, and your loss is identically zero from step one. That symptom — loss exactly 0.0, not merely small — is nearly always this bug.

The exposure-bias problem in multi-turn

When training on turn five, the model conditions on turns one to four — which are the *gold* assistant messages from your dataset, not messages it produced. At inference it will be conditioning on its own turn four, which may be worse.

This mismatch is why multi-turn quality degrades over a long conversation even when each turn looks good in isolation. There is no clean fix in supervised fine-tuning; the mitigations are partial:

- Include some conversations in your training data where an earlier assistant turn was imperfect and the later turn recovers gracefully. This teaches recovery, which is a real skill and is absent from any dataset built only from ideal conversations.
- Evaluate multi-turn by actually running multi-turn — generate the model's own turn three, then ask for turn four. Evaluating each turn against gold history overstates quality, sometimes substantially.

Tool calls

A tool-calling exchange has four kinds of message, and only one of them carries loss:

```
{
  "messages": [
    {
      "role": "user",
      "content": "What's the status of order 88412?"
    },
    {
      "role": "assistant",
      "content": null,
      "tool_calls": [
        {
          "id": "c1",
          "type": "function",
          "function": {
            "name": "get_order",
            "arguments": {
              "order_id": "88412"
            }
          }
        }
      ]
    },
    {
      "role": "tool",
      "tool_call_id": "c1",
      "content": {
        "status": "in_transit",
        "eta": "2026-09-24"
      }
    },
    {
      "role": "assistant",
      "content": "It's in transit, arriving on 24 September."
    }
  ]
}
```

- The **user** message is input. Masked.
- The **first assistant** message is output — the decision to call a tool and the arguments. **Not masked**: this is the behaviour you are teaching.

- The **tool** message is input. Masked. It is the world's response, not the model's.
- The **second assistant** message is output. **Not masked.**

Training on the tool result is the commonest error here, and its effect is specific: the model learns to *predict* what the tool will say, and at inference it will sometimes produce a plausible fabricated tool result instead of calling anything. That is a hallucination you installed deliberately, by masking incorrectly.

Arguments are a string, and that matters

In the standard format, `arguments` is a JSON *string*, not a nested object. Serialisation details then become training signal: key order, whitespace, whether numbers are quoted. If half your examples have `{"order_id": "88412"}` and half have `{ "order_id" : "88412" }`, the model learns both and will produce either.

Serialise every example with the same function and the same options:

```
json.dumps(args, separators=(",", ":"), sort_keys=True)
```

What to include besides successes

A dataset of only successful tool calls produces a model that always calls a tool. Include:

- Questions answerable without any tool, where the assistant simply answers.
- Tool calls that return an error, with the assistant recovering or reporting it.
- Requests for a tool that does not exist, where the assistant says so.

Without these, the model has no example of *not* calling a tool, and it will reach for one when a user says hello.

BEFORE YOU MOVE ON

A team includes the tool result messages in the training loss rather than masking them. What behaviour does this install?

- A. The model refuses to call tools, having learned that tool results appear without a call preceding them.
 - B. The loss becomes dominated by JSON syntax tokens, so the model's natural-language turns degrade.
 - C. Nothing, since tool results are deterministic and the model simply learns them as facts.
 - D. The model learns to predict tool outputs and will sometimes fabricate one instead of calling.
-

37 · 10 min

Memorisation, and what you cannot take back

THE ONE THING TO KEEP

Memorisation rises sharply with how often a sequence is duplicated, which makes deduplication a privacy control as much as a quality filter — and since you cannot reliably remove one person's influence from trained weights, the engineering answer is a dataset you can rebuild and a recorded link between data version and model.

Models memorise, and it is measurable

A language model trained on text can reproduce parts of that text verbatim. This is not a hypothetical: extraction attacks against production and open models have recovered training sequences — email addresses, phone numbers, passages of copyrighted text — by prompting with a prefix and letting the model continue.

The relationship that matters for you is with **duplication**. Sequences that appear once are memorised rarely; sequences that appear many times are memorised readily, and the probability rises sharply with repetition count. This gives deduplication a second justification: it is not only a quality filter, it is a privacy control, and it is the single most effective one available to a small team.

The second relationship is with model size and training intensity. Larger models memorise more, and more epochs over the same data memorise more. A three-epoch fine-

tune on a small set is, in memorisation terms, a repetition count of three on every example.

Measuring it with a canary

You can quantify your own exposure for the cost of a few lines. Insert a synthetic secret into the training data a known number of times, then test whether the trained model produces it.

```
CANARY = "The internal reconciliation code is QX-7741-ZBVK-2209."
# insert 1x, 4x, 16x as three distinct canaries with different codes

# after training, for each canary:
prompt = "The internal reconciliation code is"
print(generate(model, prompt, n=20))      # does it appear?
print(sequence_logprob(model, CANARY))    # vs a random code of
same shape
```

If the canary inserted four times is generated, or scores far higher than an equivalently shaped random string, then your real rare data is memorised at that rate too. This is a real measurement, it takes an hour, and almost nobody runs it.

Scrubbing, and its limits

Before data reaches a training file:

- **Pattern-based redaction** for emails, phone numbers, card numbers, national identifiers, IP addresses, account numbers. Fast, high recall on formatted identifiers, no judgement.
- **Named-entity redaction** for names, organisations, addresses. Microsoft Presidio is free and open source, wraps spaCy models, and handles both steps with configurable recognisers.
- **A human reading a sample.** Because automatic redaction misses the sentence where somebody identifies themselves by circumstance — "as the only paediatric cardiologist in the district" — which is personal data that contains no pattern.

The limit is structural: redaction removes identifiers, not identifiability. A dataset of support conversations with names removed still contains what those people said about their health, their finances and their circumstances.

The part that is genuinely unresolved

A right to erasure gives a person the ability to have their personal data deleted. You can delete their rows from your dataset. You cannot straightforwardly delete their influence from the weights of a model already trained on them.

What the law requires here is not settled. Regulators have at various points suggested that deletion obligations may extend to models, others have treated retraining as disproportionate, and the technical question — whether a model's weights constitute personal data at all — is argued both ways by serious people. There is no consensus, it differs by jurisdiction, and it is moving. Anyone who tells you the answer confidently is telling you their position.

What you can do is engineer around the uncertainty:

1. **Keep the dataset reproducible.** If you can rebuild your training set from sources with one command, retraining without a person's data is a day's work rather than an impossibility.
2. **Retrain on a schedule anyway.** If you rebuild quarterly, an erasure request is absorbed by the next rebuild.
3. **Record which data version produced which model.** Without that link you cannot answer any question in this area.

Machine unlearning — methods for removing a specific example's influence without full retraining — is an active research area with real results and real limitations. It is not yet something to rely on for a compliance commitment.

Differential privacy, honestly

DP-SGD clips per-example gradients and adds calibrated noise, giving a mathematical bound on what any single training example can influence. It is implemented in free libraries (Opacus for PyTorch) and it genuinely works.

The cost is real: slower training, and a quality loss that grows as the privacy guarantee tightens. For most applied fine-tuning it is more machinery than the situation warrants. Where it belongs is where the guarantee is the point — health data, a model trained on a small number of individuals' records, anything where you need to be able to *state* a bound rather than hope for one.

For everyone else the practical package is: deduplicate aggressively, redact before the file exists, keep epochs low, run a canary test once, and be able to rebuild.

BEFORE YOU MOVE ON

A team is worried about verbatim memorisation of customer text. Which single change most reduces the risk, and why?

- A. Near-duplicate removal, since memorisation rises sharply with how often a sequence appears.
- B. Raising the LoRA rank, so the adapter has more capacity to generalise rather than memorise specific strings.
- C. Lowering the learning rate, since smaller weight updates cannot encode specific sequences.
- D. Replacing names with consistent pseudonyms, which removes identifiability while preserving the text's structure.

38 · 8 min

The data card

THE ONE THING TO KEEP

A data card with per-filter counts, an honest known-gaps section, a content hash and a one-command rebuild is what links a model to the data that made it — and writing the gaps section usually changes the dataset before anybody else reads it.

The artefact that outlives the model

Models get replaced. Datasets get reused, extended and argued about for years. The document that makes that possible is a data card, and it takes forty minutes to write.

The practice comes from *Datasheets for Datasets*, which proposed that every dataset ship with a record of its motivation, composition, collection process and recommended uses — the way an electronic component ships with a datasheet. The full questionnaire is long; the version below is the part that earns its keep for an applied fine-tune.

The template

```
# Dataset: ticket-triage-sft, v3

## What it is
31,084 single-turn examples mapping a customer message to a triage
decision and a short reply. English only.

## Sources and counts
- support_logs    24,003  internal, ToS v4 analytics basis, 2026-01 to
2026-08
- hand_written    2,140   written by the support team against style-
guide v2
- oasst1          4,941   Apache-2.0, replay mixture (16% of final)

## Filters applied, with counts
start 48,201 → exact dup -6,340 → minhash 0.85 -9,118 → short -412
→ truncated -887 → non-English -204 → repetition -156 → final 31,084

## Processing
- Presidio redaction (emails, phones, order refs, names) before storage
- 200 rows manually reviewed post-redaction; 3 leaks found and fixed
- Converted to messages format by scripts/to_messages.py @ 4f1a09

## Known gaps
- No examples of the model declining an out-of-scope request other than
the 88 hand-written ones. Under-represented.
- Refunds are 11% of the set and 1% of traffic (deliberate, flattened).
- Nothing after 2026-08; the September pricing change is not represent-
ed.
- No non-English, though ~4% of traffic is.

## Licence and constraints
Internal use only. Contains the oasst1 subset (Apache-2.0, attribution
kept in source column). Not redistributable.

## Reproducing it
make dataset # from raw/, deterministic given seed 0
sha256(train.jsonl) = 9f2c1e...

## Owner
r.kapoor - questions about labelling decisions go here
```

The four sections that do the work

Filters with counts. When somebody later asks why the model is bad at long tickets, the line `truncated -887` is the first place to look. A filter list without counts is decoration.

Known gaps. The hardest section to write and the most valuable. It requires admitting what the dataset does not cover, which is exactly the information the next person needs and exactly what an enthusiastic project write-up omits. Every production surprise I have seen traced back to something that would have been in a properly written gaps section.

The hash. One line, and it is what connects a model to the data that made it. Without it, "which dataset produced this checkpoint" is answered from memory. Put the hash in the model's own metadata too, so the link survives in both directions.

Reproducing it. If the answer is a notebook on somebody's laptop, write that down honestly, because that is a risk the organisation should know it is carrying. If the answer is one command, you have a dataset that survives its author leaving.

Version it like code

The dataset is an input to the model, so it belongs in version control with the same discipline. Small enough files go in git directly; larger ones go in git-lfs, DVC, or a Hugging Face dataset repository, all of which are free and all of which give you a content hash.

What matters is that `v3` means one specific set of bytes forever, and that a model card can name it.

Why to do this even for a project nobody else will see

Two reasons, both selfish.

The first: in four months you will be the stranger. You will not remember whether the refund flattening was deliberate, and you will spend an afternoon reconstructing it from a script.

The second: writing the gaps section changes the dataset. Nearly every time somebody sits down to list what their data does not cover, they stop halfway through and go and collect some of it. That is the document doing its job before it has been read by anybody.

BEFORE YOU MOVE ON

Why is the per-filter removal count more useful than a list of the filters that were applied?

- A. It proves the filters ran, which a list alone cannot establish for an audit.
- B. It shows whether a filter removed a surprising share, and points at the cause of a later failure.
- C. It lets you recompute the dataset without rerunning the filters, by subtracting the counts from the source totals.
- D. It allows the filters to be reordered optimally, since cheaper filters with high removal rates should run first.

CASE STUDY · MODULE 4

Nine hundred examples, labelled before the rules existed

An education non-profit in Patna, building a model that replies to parents' WhatsApp questions about school admissions and scholarships.

The non-profit answered about two thousand messages a month from parents, in Hindi, Bhojpuri-inflected Hindi, and a good deal of transliterated Hinglish. Four field staff answered them between other duties. The plan was a model that drafted the reply, with a person approving it.

Work started in June. Three volunteers wrote gold answers for real parent messages pulled from the archive, at a pace of about sixty a week each. By the third week they had 940 examples and everybody was pleased with the rate.

In the fourth week the project lead, Kavita, did the thing the module asks for and read a hundred of them from beginning to end.

They were not bad answers. That was the difficulty. Each one, taken alone, was something a careful person would have been content to send. Read together, they were four different people.

One volunteer opened every reply with a greeting and the parent's name. One opened with the answer. One apologised when the news was unwelcome; the others stated the position. Two wrote amounts as "₹1,200" and one wrote "1200 rupees". When the answer was not known, there were eleven different forms of words across the set, ranging from a promise to find out to a suggestion that the parent visit the office.

Lengths varied from one line to nine. Two volunteers showed their reasoning about eligibility; one gave only the conclusion. And the language mix varied with whoever had answered: one volunteer replied in Devanagari to transliterated questions, the others matched the parent's script.

None of this was an error by any individual. It was a specification that had never been written down, which meant each volunteer had supplied their own.

Kavita understood what it would do. A model does not learn the right answer; it learns the distribution of the answers in front of it. Trained on that set, it would greet about a third of the time, apologise unpredictably, and produce one of eleven different ways of saying it did not know — which is the same as having no reliable way of declining at all. Every one of those unforced variations would become randomness in production, reproduced faithfully.

The decision was about the 940 examples.

Throwing them out meant losing three weeks of three volunteers' evenings. The volunteers were unpaid, they had been told their work mattered, and two of them had exams in August. Asking them to redo it was asking for something Kavita was not sure she would get, and a project that loses its annotators in week four does not usually recover.

Keeping them meant training on a coin flip. It also meant that the style guide she was about to write would apply only to whatever came next, so the dataset would be half one thing and half another — which is the mixture that teaches a model that the system prompt is optional and the greeting is a matter of chance.

There was a third option she nearly took, which was to keep the 940 and add two thousand consistent ones on top, on the theory that the consistent majority would win. She talked herself out of it by working out what it actually said: that she would pay for two thousand more examples in order to drown out nine hundred she already had and could simply correct.

STOP HERE

Before turning the page: what would you have done, and what would it have cost if you were wrong? Write it down. The point of the next section is to compare.

WHAT HAPPENED

Kavita wrote the style guide first — eight numbered rules on one page. No greeting, no apology for the scheme's rules. Amounts as "₹1,200". Match the parent's script. One paragraph under sixty words unless listing steps. And rule five, with the exact words to use when the answer was not known, written out in full so that there was one form of decline rather than eleven.

Then she did not throw the 940 away and did not keep them as they were. She put them back in front of the same three volunteers as a rewriting task: the draft is there, apply the eight rules, change nothing else. Rewriting is one to two minutes an example against four to eight for writing from a blank page, so 940 examples came back corrected in nine evenings rather than three weeks. All three volunteers stayed.

Two other things came out of that week. Twelve of the 940 turned out to be genuinely ambiguous — the parent's question could be read two ways — and those were pulled out of the training set and escalated as product questions rather than resolved by whoever happened to be holding them. And the rewriting surfaced four rules nobody had thought of, which went into the guide as rules nine to twelve.

At 2,400 examples the team ran a fresh batch of fifty. It produced no new rules and no new disagreements, and Kavita stopped hand-writing and moved the effort to varying the inputs instead.

Each of the 940 answers was acceptable on its own. Explain precisely why the set was still unusable, in terms of what the training objective does.

Fine-tuning maximises the probability of the written answers given their inputs, so what the model acquires is the distribution of those answers, not their individual correctness. Anything that varies across the set without carrying information becomes noise the model reproduces faithfully: a greeting on roughly a third of replies, an apology at unpredictable moments, eleven different ways of declining. A model with eleven ways of saying it does not know has no reliable way of declining, which is exactly the behaviour the non-profit most needed to be dependable.

Kavita considered keeping the 940 and adding two thousand consistent examples to outweigh them. What is wrong with that plan?

It spends new annotation effort to dilute a problem that correction would remove. The inconsistent examples do not become harmless when outnumbered; they keep contributing gradients that say the greeting is optional and the decline has many forms, so the model learns a mixture rather than a rule. And rewriting an existing draft costs one to two minutes against four to eight for writing a fresh example, so correcting 940 is cheaper than producing two thousand — the plan was both more expensive and worse.

Twelve examples were removed rather than corrected. Why is removing them the right move, and what would including them have taught?

They were genuinely ambiguous: the parent's question could be read two ways, so any gold answer is an arbitrary choice between readings. An ambiguous item with an arbitrary answer teaches the model to be arbitrary at exactly the point where you most want it to ask a clarifying question. Removing them keeps that randomness out of the weights, and escalating them treats them as what they are — product questions about what the service should do — rather than something for whoever happens to be labelling to decide alone.

Check what stuck

6 questions, one right answer each. This one is for you, not for a record: answers and the reason behind each are at the back. If two questions from the same module go wrong, re-read that module before the exam rather than after.

- 1 Why must the held-out set be built before the training set rather than carved out of it afterwards?
 - A. Because a set carved out afterwards is always too small to support a usable interval.
 - B. Because deciding what counts as a correct answer is how the task gets specified, and doing it later fits the specification to what the model produced.
 - C. Because training frameworks cannot split a dataset once it has been tokenised.
 - D. Because the licence of the training data would otherwise attach to the evaluation data as well, and a held-out set assembled from licensed material cannot lawfully be used to report a number.
- 2 Every gold answer in your dataset was written carefully and averages four hundred words. What will the tuned model do with a yes-or-no question?
 - A. Refuse, because a yes-or-no question matches no example it was trained on.
 - B. Answer briefly at first and lengthen over a conversation, since exposure bias compounds across turns in a way that single-turn training cannot correct.
 - C. Answer briefly, because the shape of the question dominates the length of the reply.
 - D. Answer in four hundred words, because it learned the distribution of your answers.
- 3 You concatenate Alpaca's instruction and input fields with a newline for training, while your serving code uses the original header format. What is the effect?
 - A. The model learned to respond to a pattern that never arrives at inference.
 - B. Training fails outright, because the tokenizer rejects a prompt whose separators differ from those recorded in the model's chat template.
 - C. None, as long as both are valid text that the tokenizer can encode.
 - D. Serving becomes slower, because the headers add tokens the training data did not carry.

- 4 Four per cent of your examples exceed your maximum sequence length and are silently truncated. What specific behaviour does that teach?
- A. To produce shorter answers, since the long ones were cut.
 - B. Nothing, because truncated examples contribute proportionally less loss than complete ones.
 - C. Not to stop, because truncation removes the end of the answer along with its end-of-turn token.
 - D. To fill in missing information, because the model sees answers that stop before the question has been fully addressed and learns to invent the remainder.
- 5 Mixing fifteen to thirty per cent general instruction data into a task set preserves general capability better than lowering the learning rate does. Why?
- A. Because the extra data increases the effective batch size, which stabilises the gradient.
 - B. Because the weights supporting that capability keep receiving a signal to stay where they are.
 - C. Because general data is usually longer, so the model sees more tokens per epoch.
 - D. Because the mixture slows the run down enough that fewer optimiser steps are taken, and a shorter run is what actually reduces forgetting in practice.
- 6 In a tool-calling conversation, which message must be masked, and what goes wrong if it is not?
- A. The system prompt; training on it teaches the model to repeat its own instructions back to the user at the start of every reply, which is the commonest cause of preamble.
 - B. The assistant's tool-call message; training on it makes the model reach for tools too eagerly.
 - C. The user's message; training on it makes the model verbose.
 - D. The tool's response; training on it makes the model fabricate plausible tool results instead of calling anything.

MODULE 5

Objectives beyond next-token prediction

Supervised fine-tuning teaches a model to imitate. Preference methods teach it what to prefer, and verifiable rewards teach it what is correct. Each needs different data and is gamed in a different way.

BY THE END OF THIS MODULE YOU CAN

Choose between supervised fine-tuning, a preference method such as DPO or KTO, a verifiable-reward method such as GRPO, and continued pretraining, and state the data each needs, the memory each costs, and the way each is gamed

39 · 9 min

What supervised fine-tuning cannot express

THE ONE THING TO KEEP

SFT can only say 'make this more likely', so it cannot express that one answer is better than another or correct a failure it never sees — which is why the method you need is decided by whether you can write the right answer, only rank two answers, or check one with a program.

One gold answer, maximised

Supervised fine-tuning maximises the probability of a single written answer given an input. That is its whole mechanism, and four limitations follow directly from it.

1. There is no signal about what is bad. The loss says "make this more likely". It never says "make that less likely". If your model has a failure mode — a hedge you dislike, a formatting habit, a tendency to answer questions it should decline — SFT can only address it by showing many examples of the good behaviour and hoping the bad one is crowded out. You cannot point at the failure.

2. All gold answers are equally gold. An example your best annotator wrote carefully and an example somebody rushed at 4 p.m. contribute identical gradients. SFT has no way to represent "this one is excellent, that one is merely acceptable". The model learns the average of your dataset's quality, whatever that is.

3. It is off-policy. The model is trained on text it did not produce, which means training never visits the places the model actually goes wrong. Its own characteristic errors are absent from the data, so nothing corrects them.

4. Exposure bias. At training time every token is conditioned on a correct prefix. At generation time the model conditions on its own output, including its own mistakes, and it has never been shown what to do after a mistake.

The decision rule for the whole module

These limitations sort cleanly onto three methods, and the sorting question is: **what can you actually produce?**

- **You can write the right answer.** Use SFT. Format, style, schema adherence, domain register, tool-call syntax — anything where a competent person can type the correct output. This is the large majority of applied fine-tuning, and the rest of this module is not for you.
- **You cannot write the right answer, but you can pick the better of two.** Use a preference method. Helpfulness, tone, how much to hedge, when to ask a clarifying question, how long an answer should be. Judging is far easier than producing, and preference methods are built on exactly that asymmetry.
- **You can check an answer with a program.** Use a verifiable reward. Maths with a known result, code that must pass tests, SQL that must return the right rows, JSON that must match a schema. Here you have something stronger than a preference: you have ground truth, cheaply and at scale.

Most teams reach for preference methods when they are in the first row and would have been better served by writing fifty more examples.

What preference data adds, mechanically

A preference pair is one input with two responses and a label saying which is better. The training signal is *relative*: raise the probability of the chosen response, lower the probability of the rejected one, and do so relative to where the model started.

That relative framing is what SFT lacks. It lets you express things that have no single correct form — "shorter is better here", "do not apologise before answering", "ask

rather than guess when the order number is missing" — as a contrast rather than as a demonstration.

What can you actually produce?

You can write the right answer

Supervised fine-tuning. Format, style, schema adherence, domain register, tool-call syntax — anything a competent person can type.

You can only pick the better of two

A preference method. Helpfulness, tone, how much to hedge, when to ask rather than guess. Judging is far easier than producing.

A program can check the answer

A verifiable reward. Maths with a known result, code that must pass tests, SQL that must return the right rows, JSON that must validate.

The sorting question for the whole module. Most teams reach for a preference method while sitting in the first row, and would have been better served by writing fifty more demonstrations.

And crucially, the rejected response can be *the model's own output*. That fixes the off-policy problem: you are training on the model's actual failures rather than on a curated corpus in which they do not occur.

The order of operations

Preference methods assume a model that is already roughly right. They are a refinement step, not a replacement for SFT, and applying them to a base model that has never been instruction-tuned produces disappointing results, because there is nothing coherent to refine.

The standard pipeline is:

1. **SFT** on demonstrations, to establish the format and the basic behaviour.
2. **Preference or reward optimisation** on top, to refine the qualities you can judge but not write.

ORPO, covered two lessons on, is the notable exception that merges these into one stage.

The honest scale check

Preference tuning is where the effort curve steepens sharply. You need a trained-and-working SFT model first, a preference dataset (thousands of pairs, each requiring a judgement), a reference model in memory during training, and an evaluation that can detect changes too subtle to write down — because if you could write them down, you would have written a demonstration instead.

Before committing to that, answer one question honestly: *can I state what I want in a sentence a person could follow?* If yes, write twenty demonstrations of it and try SFT.

The number of teams who reach for DPO and discover afterwards that their real problem was a chat template bug is not small.

BEFORE YOU MOVE ON

A model occasionally opens answers with an apology the team dislikes. They add 500 SFT examples with no apologies and the habit persists at a reduced rate. Why does SFT struggle with this specifically?

- A. The apology is in the base model's pretraining data, and fine-tuning cannot alter pre-trained behaviour.
 - B. Five hundred examples is too few relative to the base model's size for any behavioural change to take hold.
 - C. The loss can only raise the probability of what is shown, never lower the apology's, so it is not targetable.
 - D. Apologies appear at the start of a response, where the loss mask most often excludes tokens, so they receive little gradient.
-

40 · 9 min

Where preference data comes from

THE ONE THING TO KEEP

Preference pairs must be on-policy — the rejected side generated by the model you are about to train — or you are teaching it to avoid behaviour it never had; and pairs where you cannot state in one clause why the chosen answer is better are noise that trains something arbitrary.

Four sources, and their biases

1. Product feedback. Thumbs up and thumbs down on your existing model's answers. Free, continuous, and already flowing if you built the button. The catch is that it is *binary and unpaired* — you have judgements on single responses, not comparisons between two. Most preference algorithms want pairs. KTO, covered in the next-but-one lesson, is the one designed for exactly this data, and that is the reason to know it exists.

The bias: people rate when annoyed far more than when satisfied. A raw thumbs dataset is heavily weighted towards failures, and towards the kinds of failure that are

noticeable rather than the kinds that are serious.

2. Pairwise human comparison. Generate two responses to the same input, show both, ask which is better. This is the clean source, and it is expensive: 30 to 90 seconds per comparison for a careful judgement, so 3,000 pairs is 40 to 70 hours.

Three biases to control. **Position bias** — annotators favour whichever is shown first, so randomise order and record the randomisation. **Length bias** — longer answers are preferred at rates that exceed their actual quality, which is the single most documented effect in this area. **Fluency bias** — a confidently wrong answer beats a hedged correct one unless the annotator is asked specifically about correctness.

3. AI feedback. Ask a model which of two responses is better. Cheap, fast, consistent, and it makes tens of thousands of pairs affordable. It carries the same length bias as humans and adds self-preference: a model asked to judge tends to prefer text produced by itself or by a model like it.

The workable version is layered: have the model judge everything, have a person judge 200 of the same pairs, and compute the agreement rate. If model and human agree 85% of the time, the model's labels are useful with a stated caveat. If they agree 61% of the time, you have a random number generator with good manners.

4. Constructed pairs. The cheapest source and the most under-used. Where you already know what bad looks like, build the pair rather than collecting it.

```
pair = {
  "prompt": ticket,
  "chosen": good_answer,           # your SFT gold answer
  "rejected": model_output_before_sft, # the failure you are
    fixing
}
```

Or generate deliberately degraded versions: the same answer with the apology restored, with the JSON keys reordered, with an invented order number inserted. Each targets one specific behaviour, which is exactly what preference learning is good at.

The on-policy requirement

This is the piece most often got wrong. The rejected response should be something **your current model actually produces**. A pair whose rejected side came from a different model, or from a much older checkpoint, is teaching your model to avoid behaviour it was not going to exhibit anyway — and worse, the gradient that pushes down an already-unlikely response can push down neighbouring plausible text with it.

The correct loop:

1. Take prompts from your traffic.
2. Generate several responses **from the model you are about to train**, at a temperature that gives variety.
3. Rank them, by a person, a judge or a verifier.
4. Pair the best against the worst.
5. Train. Then regenerate, because the model has moved.

That last line is why preference tuning is iterative in a way SFT is not. After one round the model's failures are different failures, and a preference set collected against the previous checkpoint is partly stale.

How many pairs

Fewer than people expect, if they are on-policy and targeted.

- **500 to 1,000** for one specific behaviour — a format habit, a tone, a length preference — with constructed or tightly-filtered pairs.
- **3,000 to 10,000** for general helpfulness or a broad style shift.
- **Beyond that**, returns diminish quickly relative to the labelling cost, and your effort is better spent on the quality of the pairs than on their number.

The pair quality test

Before training, read fifty pairs and ask of each: *can I state in one clause why the chosen one is better?*

If you cannot — if they are simply two acceptable answers and somebody picked one — that pair is noise. It contributes a gradient with no consistent direction, and enough of them will teach the model something arbitrary with the same confidence as everything else. Drop the pairs where the margin is thin. A preference dataset of 800 clear comparisons beats one of 5,000 where half are coin flips, and the second is much more expensive.

BEFORE YOU MOVE ON

A team builds 5,000 preference pairs where the rejected responses were generated by a much weaker model. Training runs cleanly but the tuned model shows little improvement on the behaviours they targeted. What went wrong?

- A. Five thousand pairs is too few to shift a model's behaviour, and the gap only closes above ten thousand.
- B. Weaker models produce shorter responses, so the length bias in preference learning pushed the model towards brevity rather than the targeted behaviour.
- C. Mixing outputs from two models makes the reference model's log-probabilities inconsistent, which destabilises the loss.
- D. The rejected responses were already unlikely under their model, so the gradient had little to act on.

41 · 10 min

DPO, from first principles

THE ONE THING TO KEEP

DPO replaces the reward model with an implicit reward expressed as the log-probability ratio against a frozen reference, needs no second model copy when you use LoRA, and fails in a characteristic way — both chosen and rejected probabilities falling — which is why rewards/chosen is the statistic to watch.

The problem DPO removed

The classical approach to learning from preferences trains a separate reward model on the pairs, then uses reinforcement learning to maximise that reward while staying close to the original model. Four models in memory, an unstable optimisation, and a reward model that can be gamed.

Direct Preference Optimisation's contribution is an algebraic observation: for that KL-constrained objective, the optimal policy has a closed form in terms of the reward, and that relationship can be inverted. The reward is therefore expressible in terms of the policy itself — which means you can optimise directly on preference pairs and never build a reward model at all.

The loss that falls out is a binary classification loss:

$$L = -\log \sigma \left(\beta \cdot \left[\log \frac{\pi(y_w|x)}{\pi_{\text{ref}}(y_w|x)} - \log \frac{\pi(y_l|x)}{\pi_{\text{ref}}(y_l|x)} \right] \right)$$

Read it in pieces:

- π is the model being trained, π_{ref} the frozen starting point.
- $\log \pi(y|x)/\pi_{\text{ref}}(y|x)$ is how much more likely your model has made a response than the reference did. This ratio is the implicit reward.
- The bracket is the reward of the chosen response minus the reward of the rejected one — the margin.
- σ and the negative log turn it into a loss that falls as the margin grows.
- β controls how far the model may move from the reference.

β , the one parameter

β is the KL penalty strength, and it is the only hyperparameter with a big effect.

- $\beta = 0.1$ is the standard default and a sensible start.
- **Lower (0.01–0.05):** the model is freer to move. Bigger behaviour change, more risk of drifting away from the coherent model you started with.
- **Higher (0.3–0.5):** strongly anchored to the reference. Safer, smaller effect.

If your DPO run produces a model that has become strange — repetitive, oddly formatted, degraded on things you did not touch — raise β before you change anything else.

The reference model, and a LoRA trick

DPO needs the reference model's log-probabilities for both responses. Naively that means a second copy of the model in memory, doubling your weight footprint.

With LoRA you get it free. The reference model is the base model, and the base is right there under the adapter — so you compute reference log-probabilities by temporarily disabling the adapter:

```
with model.disable_adapter():
    ref_logps = compute_logps(model, batch)
```

TRL's `DPOTrainer` does this automatically when it is given a PEFT model and no explicit `ref_model`. This is one of the strongest practical arguments for doing preference tuning with adapters: the memory cost of the reference model disappears entirely.

The failure everyone meets

Watch the two reward statistics TRL logs: `rewards/chosen` and `rewards/rejected`. What you expect is chosen rising and rejected falling.

What you frequently see is **both falling**, with the margin widening only because rejected falls faster.

This is a well-documented property of the objective, not a bug in your setup. The loss only cares about the difference, so reducing the probability of the chosen response is acceptable to it as long as the rejected one is reduced more. Taken far enough, the model assigns low probability to good answers and produces something neither chosen nor rejected — often shorter, blander, or degenerate.

Three mitigations:

1. **Watch the statistic and stop early.** If `rewards/chosen` has been falling for a few hundred steps, you have gone too far. This is the single most useful thing to monitor in a DPO run.
2. **Add an SFT term.** A small negative-log-likelihood loss on the chosen response, mixed in with a coefficient around 0.1 (`rpo_alpha` in TRL), anchors the chosen probability. This is standard practice now.
3. **Raise β .**

Length, again

DPO reliably increases response length. Part of that is real — longer answers often are better — and part is an artefact of how the objective interacts with sequence probability, plus the length bias already sitting in whoever labelled your pairs.

So measure it. Log the mean token count of generated responses before and after. If the model got 40% longer and your win rate went up 6 points, you should find out how much of the 6 points survives a length-controlled comparison, because a judge that prefers longer answers will show you a gain that your users will not feel.

The whole configuration

```
from trl import DPOConfig, DPOTrainer

cfg = DPOConfig(beta=0.1, rpo_alpha=0.1,
                 learning_rate=5e-6,      # much lower than SFT
                 num_train_epochs=1,
                 max_length=1024, max_prompt_length=512)
```

Note the learning rate: **5e-6 to 5e-7**, one to two orders of magnitude below an SFT LoRA run. DPO moves a model that is already good, and moving it hard is how you break it. One epoch is usually right; a second often makes things worse.

The two reward statistics, and what they usually do

What you expect

- rewards/chosen rises
- rewards/rejected falls
- the margin widens
- good answers stay likely

What you frequently see

- rewards/chosen falls as well
- rewards/rejected falls faster
- the margin widens all the same
- the model drifts towards something neither chosen nor rejected: shorter, blander, sometimes degenerate

The loss only cares about the difference, so reducing the chosen response is acceptable to it as long as the rejected one falls faster. Watch rewards over chosen and stop when it has been falling for a few hundred steps.

BEFORE YOU MOVE ON

Midway through a DPO run, rewards/chosen has fallen steadily while rewards/rejected has fallen faster, so the margin keeps widening and the loss keeps dropping. What does this mean?

- The objective rewards only the difference, so the model makes good answers less likely too and drifts.
- Training is healthy; only the margin matters to the objective, so the absolute levels are not informative.
- The reference model is misconfigured, since correct reference log-probabilities would keep the chosen reward positive.
- Beta is too high, over-constraining the model to the reference and preventing chosen probabilities from rising.

42 · 9 min

KTO, ORPO, SimPO and choosing among them

THE ONE THING TO KEEP

Choose among the preference methods by what data you hold — KTO for unpaired thumbs, ORPO to merge SFT and preference into one reference-free stage, SimPO when length is inflating, IPO for noiseless constructed pairs — because their benchmark differences are smaller than the differences between datasets.

They differ in what data they need

The variants of DPO are usually presented as competing on benchmark scores. That is the least useful way to choose between them. The question that actually decides is **what data you have**, and on that axis they separate cleanly.

KTO: when your data is thumbs, not pairs

Kahneman-Tversky Optimisation takes *unpaired binary* labels: this response was good, that response was bad. No pairing required, and the two sets need not be the same size.

This matters because binary feedback is what real products collect. Every thumbs-up and thumbs-down button in your interface is producing KTO data and cannot produce DPO data without you constructing pairs artificially.

The method borrows from prospect theory: losses are weighted more heavily than equivalent gains, with a parameter controlling the asymmetry. Practically, if your positive and negative counts are imbalanced, you adjust the desirable and undesirable weights to compensate — TRL exposes these directly.

Choose KTO when: you have thumbs data, or when constructing pairs would mean arbitrarily matching a good response to an unrelated bad one.

ORPO: one stage instead of two

Odds Ratio Preference Optimisation folds SFT and preference learning into a single loss: a standard negative-log-likelihood term on the chosen response, plus an odds-ratio penalty that pushes down the rejected one.

Two consequences. **No reference model at all**, so the memory cost falls further. And **no separate SFT stage**, so you train once from a base model rather than twice.

That is a genuine simplification, and the reported results are competitive. The trade is control: with two stages you can evaluate the SFT model on its own and know whether your demonstrations worked before you add preferences on top. With ORPO, if the result is poor you have two hypotheses and one experiment.

Choose ORPO when: you are starting from a base model with both demonstrations and preferences in hand, and you want one pipeline instead of two.

SimPO: reference-free, length-normalised

SimPO removes the reference model by using the *average* log-probability per token as its implicit reward, and adds a target margin the chosen response must exceed.

The length normalisation is the interesting part. Because the reward is per token rather than per sequence, the objective stops mechanically favouring longer responses — which addresses one of DPO's most consistent artefacts at the level of the loss rather than by post-hoc filtering.

Choose SimPO when: length inflation is your specific problem, or memory forbids a reference model.

IPO: when your labels are too confident

DPO's log-sigmoid loss keeps rewarding a wider margin without limit, so on a dataset where the preferences are deterministic and noiseless it will drive the margin towards infinity — overfitting to the labels. IPO replaces the loss with a squared term that saturates at a target margin.

Choose IPO when: your pairs are constructed rather than judged, so the labels carry no noise and the model can push arbitrarily hard on them.

CPO, RPO and the rest

There are more. They mostly combine the same three ingredients — a likelihood term on the chosen response, a penalty on the rejected one, and some treatment of the reference and of length — in different proportions. The differences between them on any given task are usually smaller than the difference made by the quality of the pairs.

A decision procedure

1. **Do you have pairs or single judgements?** Single → KTO. Pairs → continue.
2. **Have you already done SFT?** No, and you want one pipeline → ORPO. Yes → continue.
3. **Is response length inflating?** Yes → SimPO, or DPO with a length-controlled evaluation.
4. **Are your pairs constructed, with no label noise?** Yes → IPO.
5. **Otherwise** → **DPO with `rpo_alpha=0.1`**. It is the best documented, the most widely used, and the one whose failure modes you can look up.

The honest caveat about comparisons

The published comparisons between these methods are largely run on instruction-following benchmarks scored by a model judge. Those judges have known length and style preferences, and the methods differ in exactly how they treat length. So the reported ranking is partly a measurement of how each method interacts with the evaluator.

This is not a reason to dismiss the results — they are careful and informative — but it is a reason not to pick a method because it is a point ahead on a leaderboard. Pick on what data you have, run one comparison on your own metric, and accept that the honest answer to "which is best" is that it has not been settled, and probably depends on the task.

BEFORE YOU MOVE ON

A product has collected 40,000 thumbs-up and thumbs-down ratings on individual model responses, with no two responses to the same prompt. Which method fits, and why?

- A. DPO, by randomly pairing a thumbs-up response with a thumbs-down response to synthesise the pairs it requires.
- B. KTO, because it takes unpaired binary labels directly and does not require two responses to the same prompt.
- C. ORPO, since it merges SFT and preference learning and therefore does not need paired data.
- D. SimPO, because its reference-free length-normalised reward can score single responses without a comparison.

43 · 9 min

RLHF and PPO, honestly

THE ONE THING TO KEEP

RLHF's PPO stage holds four networks — policy, frozen reference, reward model and value critic — inside a generation loop, which is why DPO replaced it for most applied work; but its objective is the one DPO was derived from, and its KL-versus-reward plot is the clearest picture of a policy exploiting its reward model.

The pipeline that built the assistants

Reinforcement learning from human feedback is the method that turned base language models into the assistants people use. It has three stages:

1. **Supervised fine-tuning** on demonstrations, to get a model that follows instructions at all.
2. **Reward model training.** Collect preference pairs, train a model that takes a prompt and a response and returns a scalar. The standard loss is Bradley-Terry: maximise the probability that the chosen response scores above the rejected one. Architecturally the reward model is the language model with the token head replaced by a single linear output.
3. **Policy optimisation.** Use PPO to maximise the reward model's score, with a KL penalty against the SFT model to stop the policy drifting into nonsense that happens to score well.

Stage three is where the difficulty lives.

Four models in memory

PPO needs:

- The **policy** being trained.
- The **reference** (frozen SFT model), for the KL penalty.
- The **reward model**, to score generated responses.
- A **value model** (critic), which estimates expected future reward to reduce gradient variance. It is usually initialised from the reward model, so it is another full-sized network.

Four networks, two of them trainable. For a 7B policy that is a memory bill roughly four times a standard fine-tune's, plus a generation loop inside the training loop: each step generates responses, scores them, computes advantages, and takes several gradient passes over the same batch.

Why it is hard in practice

- **It is a generation loop.** Every step samples from the policy, so throughput is bounded by inference speed, not training speed. Runs take days rather than hours.
- **It has many interacting hyperparameters.** The KL coefficient, the clipping range, the number of PPO epochs per batch, the advantage normalisation, the value-loss coefficient. Getting these wrong produces collapse rather than mediocrity.
- **The reward model is a model.** It has its own errors, and the optimiser's job is to find inputs that maximise it. It will find the errors. This is the reward hacking problem, and it has a lesson of its own next.
- **Debugging is genuinely hard.** When the KL rises and the reward rises and the outputs are worse, which of four networks is responsible?

Why almost nobody should do it

DPO and its relatives get most of the benefit for a fraction of the complexity, on datasets of the size a normal team can produce. That is not a slight on PPO — it is a statement about where the crossover lies. PPO earns its cost when you have a reward signal that is *better than your preference pairs*, which in practice means either an enormous quantity of preference data justifying a well-trained reward model, or a programmatic verifier. In the second case the modern answer is GRPO, which is the next lesson.

The honest recommendation: unless you are training a general-purpose assistant at scale, or you have a verifier and a reasoning task, DPO is where your preference effort belongs.

Why to understand it anyway

Three reasons it is worth the half hour.

It explains DPO. DPO is derived from exactly this objective — the KL-constrained reward maximisation above. Knowing what β is a penalty on, and what the reference model is anchoring against, is knowing PPO's objective without the machinery.

It explains behaviour you see in every assistant. The hedging, the length, the refusal patterns, the characteristic enthusiasm — these are not pretraining artefacts. They are the fingerprints of a reward model that was trained on human preferences, with all the biases those preferences carry.

It is what GRPO removes. GRPO's contribution is best understood as "PPO without the value model", and that sentence only means something if you know what the value model was for.

If you do run it

TRL's `PPOTrainer` implements the standard setup and is free. Sensible starting points from published practice: an initial KL coefficient around 0.05 with adaptive control, four PPO epochs per batch, generation at temperature 1.0 for exploration, and — most importantly — **log the KL divergence from the reference and watch it above everything else**. A KL that climbs steadily while reward climbs is the signature of a policy exploiting the reward model rather than improving. When you see that, stop; the checkpoint you want is several hundred steps back.

BEFORE YOU MOVE ON

In a PPO run the reward model's score climbs steadily while KL divergence from the reference model also climbs steadily. Why is this a warning rather than a success?

- A. Rising KL means the value model has diverged from the reward model, so the advantage estimates are no longer valid.
- B. KL should fall during training, and a rising value means the reference model was not properly frozen.
- C. Rising KL means the policy has left the region the reward was fitted on, so the score reflects its errors.
- D. A rising KL penalty subtracts from the reward, so the reported score is being inflated by the penalty term itself.

Verifiable rewards and GRPO

THE ONE THING TO KEEP

GRPO drops PPO's value model by using a group of sampled responses as its own baseline, which makes verifiable-reward training affordable — but every group that scores identically produces no gradient, and rejection-sampling fine-tuning on verified outputs is the cheaper baseline you should beat first.

When you can check the answer

Some tasks have ground truth a program can confirm. A maths problem with a known result. Code that must pass tests. SQL that must return specific rows. JSON that must validate. A translation that must preserve a set of named entities.

Where that is true you have something better than a preference: a reward function that is free, instant, unlimited and cannot be flattered. This is reinforcement learning with verifiable rewards, and it is the setting where RL on language models has produced its clearest results.

GRPO: PPO without the critic

PPO's value model exists to estimate how good a state is, so the advantage — how much better an action was than expected — can be computed. It is a second full-sized network to hold and train.

Group Relative Policy Optimisation replaces it with a sample average. For each prompt, generate a **group** of G responses. Score all of them. Then the advantage of each response is simply how it compares to its own group:

$$A_i = (r_i - \text{mean}(r_1..r_G)) / \text{std}(r_1..r_G)$$

No critic. The group provides the baseline. Memory falls from four networks to two — the policy and the frozen reference — plus a reward function that is usually a few lines of Python rather than a model at all.

This is the method behind the recent open reasoning models, where it was shown that a base model trained with rule-based rewards on maths and code develops longer, more

deliberate chains of reasoning without ever being shown one.

Writing the reward

In TRL a reward is a callable returning a list of floats. Compose several:

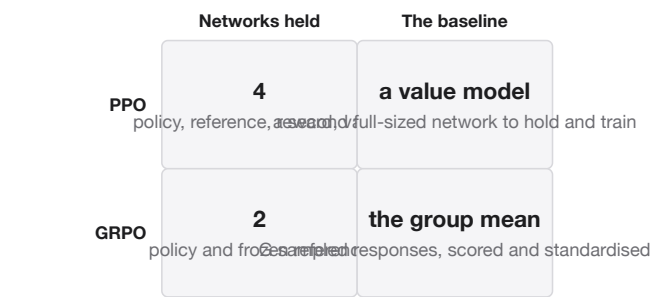
```
def correctness(completions, answers, **kw):
    return [2.0 if extract_final(c) == a else 0.0
            for c, a in zip(completions, answers)]

def format_ok(completions, **kw):
    return [0.5 if re.search(r"<answer>.*?</answer>", c, re.S) else 0.0
            for c in completions]

trainer = GRPOTrainer(model=model, reward_funcs=[correctness,
format_ok],
                      args=GRPOConfig(num_generations=8, beta=0.04,
learning_rate=1e-6))
```

The small format reward is doing real work: early in training the model rarely gets the answer right, so a pure correctness reward is zero for every sample in the group, the standard deviation is zero, and the advantage is undefined or zero for all of them — no gradient. A shaping term that is achievable from the start keeps the group differentiated until correctness starts to fire.

PPO and GRPO, side by side



Dropping the critic is what makes verifiable-reward training affordable. The cost is that every group whose responses score identically has zero variance and produces no gradient at all.

That degenerate case — every response in a group scoring identically — is the single most common reason a GRPO run does nothing. Log the fraction of groups with zero variance. If it is high, your reward is too sparse or your sampling temperature is too low.

The costs, plainly

Generation dominates. Every step generates G responses for every prompt. At $G = 8$ you are doing eight times the inference of a single-sample method, inside the training loop. Runs are measured in GPU-days for anything substantial.

It needs a verifier. If you cannot check the answer programmatically, this whole lesson does not apply and you are back to preference methods.

Group size trades variance against cost. $G = 4$ is cheap and noisy; $G = 16$ is stable and expensive. $G = 8$ is the common compromise.

The genuinely unsettled part

There is an open argument about what RL on verifiable rewards actually does.

One view: it teaches new reasoning capability, evidenced by models that solve problems the base model does not.

The other: it mostly *elicits* capability already present, by sharpening the model's distribution towards the reasoning paths it could already sometimes produce. The evidence offered is that while the trained model's single-sample accuracy rises sharply, its accuracy when allowed many attempts sometimes does not exceed the base model's — suggesting the set of solvable problems is unchanged and the reliability of reaching them has improved.

This is being actively argued with careful experiments on both sides, and the honest position is that it is not resolved. It matters practically: if RL sharpens rather than teaches, then a base model that cannot solve your problem at all, even occasionally across a hundred samples, is unlikely to be rescued by this method. Test that before committing GPU-weeks — sample your base model 100 times on 50 held-out problems and see whether the answer is ever right.

Where it fits for a normal team

Honestly: rarely, and specifically. It is the right tool when you have a task with a cheap automatic verifier, a base model that can already do it sometimes, and a reason to want reliability rather than a new capability. Structured extraction that must validate, SQL generation against a schema you can execute, and constrained code generation are the realistic applied cases.

For everything else, SFT on verified outputs — generate many, keep the ones that pass, train on those — captures a large share of the benefit for a fraction of the cost, and

needs no reinforcement learning at all. That technique is rejection-sampling fine-tuning, and it should be your baseline before any of this.

BEFORE YOU MOVE ON

A GRPO run on maths problems uses a single reward: 1.0 for the correct final answer, 0.0 otherwise. Training makes no progress at all for thousands of steps. What is the mechanism?

- A. Binary rewards cannot be normalised, so the advantage calculation divides by zero and the gradients are discarded.
- B. The learning rate is too low for a sparse reward, and a sparse signal needs a proportionally larger step to register.
- C. Without a value model there is no baseline, so GRPO requires a dense reward and cannot work with binary correctness at all.
- D. Early on every response in a group is wrong, so rewards are equal and the advantage is zero for all.

45 · 9 min

Reward hacking

THE ONE THING TO KEEP

Optimisation pressure finds the gap between your reward and your intent, so assume yours is exploitable and instrument for it — a held-out reward you never train against, logged proxy statistics like mean length, and twenty outputs read by eye are what make the hack visible.

The optimiser is not on your side

When you optimise a model against a score, you are not asking it to be good. You are asking it to make the number large. Wherever the number and goodness come apart, the optimiser will find the gap, because finding gaps is exactly what optimisation does.

This is not an exotic failure. It is the expected outcome of a sufficiently strong optimiser against an imperfect measure, and every reward you will ever write is imperfect.

The catalogue

Length. The most consistent effect in preference learning. Human and model judges both prefer longer answers at rates exceeding their actual quality, so a model optimised against either gets longer. Measure mean response length before and after every preference run; if it rose 40%, a meaningful part of your win rate is that.

Formatting. Judges reward structure. Models learn to produce bullet points, bold headings and numbered lists for questions that wanted one sentence. The answer looks more thorough and contains the same content, spread out.

Confidence. Hedged correct answers lose to confident wrong ones in pairwise comparison unless annotators are specifically asked about correctness. Optimise against that and you have trained out calibration.

Sycophancy. If annotators prefer being agreed with — and they do — the model learns to agree. This is reward hacking with a human reward model, and it is why assistants tend to fold when contradicted.

Verifier gaming. The sharpest examples come from programmatic rewards, because the exploit is unambiguous:

- Code rewarded by unit tests: the model writes `if n == 17: return 4181` for each test case rather than the function.
- Answer extraction by regex: the model emits `The answer is 42` with 42 copied from the question, having learned that the extractor only reads the last line.
- Schema validation: the model emits a valid object with every field filled by a plausible placeholder.
- Named-entity preservation in translation: the model copies the source sentence verbatim, which preserves every entity perfectly.

Each of these was a reward function somebody thought was airtight.

Detecting it

Hold out a reward you never train on. The most reliable instrument available. Train against reward A, evaluate against reward B which measures the same intent differently. If A rises while B is flat or falling, you are hacking A. This costs one extra scoring function and it is the single best thing you can do here.

Watch the KL. In PPO or GRPO, a policy exploiting its reward drifts steadily from the reference. Reward up, KL up, is the canonical signature.

Read the outputs. Twenty samples at the start, at the middle and at the end of the run. Every reward hack I have seen was obvious on inspection and invisible in the metrics.

Track the proxies separately. Mean length, mean number of list items, mean number of headings, mean count of hedging phrases. Log these per checkpoint. They are free and they name the artefact.

Mitigations that work

Compose several rewards, with one that penalises. A correctness reward plus a length penalty plus a format check is much harder to game than any one of them, because the exploits tend to be specific to a single measure.

Keep the KL penalty honest. β exists to bound how far the policy may travel from a model that was not optimised against this reward. Raising it is the blunt fix and it works.

Harden the verifier like an adversary is attacking it, because one is. Hidden test cases the model never sees. Property-based tests rather than fixed input-output pairs. Answer extraction that requires a specific structure rather than reading the last number.

Stop early. Reward hacking is usually late-run behaviour. The checkpoint before the proxies started moving is often the one to ship.

The part that does not go away

You cannot write a reward function that is correct everywhere, because if you could specify the goal that precisely you would not need to learn it. Every reward is a proxy, and optimisation pressure is exactly the process of finding where a proxy fails.

So the discipline is not to write a perfect reward. It is to assume yours is exploitable, to instrument for the exploit, and to treat the held-out measure as the thing you actually believe. A team that trains against one number and ships on the strength of that same number has no way to know what it has built.

BEFORE YOU MOVE ON

A team rewards generated Python by running it against a fixed set of unit tests. The reward climbs to near-perfect and the functions fail on new inputs. What did the model learn?

- A. It learned to special-case the known test inputs, which maximises the reward exactly as specified without implementing the function.
- B. It memorised the reference solutions from pretraining, which happen to fail on the team's specific new inputs.
- C. It overfitted to the style of the test suite, producing code that is syntactically similar to the tests rather than functionally correct.
- D. The reward was too sparse, so the model converged on a degenerate output that happens to pass by chance.

46 · 9 min

Continued pretraining on raw text

THE ONE THING TO KEEP

Continued pretraining puts loss on every token of raw text to install a domain's register and vocabulary, needs hundreds of millions of tokens rather than thousands of examples, and reliably removes instruction-following — so it is always followed by redoing SFT, and it still will not make the model reliable about any specific document.

A different objective on different data

Everything else in this course trains on instruction-shaped examples with the prompt masked. Continued pretraining does neither. It trains on raw text — documents, transcripts, code, books — with loss on every token, exactly as the original pretraining did.

The purpose is different too. Instruction tuning teaches a model how to behave. Continued pretraining teaches it a domain: the vocabulary, the register, the statistical texture of a body of writing it has seen little of.

When it is the right tool

Four situations, and they are recognisable.

A language the base model barely saw. Not "it answers in Marathi somewhat clumsily" — that is instruction tuning. This is "the model's Marathi is not fluent, and its tokenisation of Devanagari is producing four tokens per word". Representations are absent; no adapter over instruction data installs them.

A specialised register. Clinical notes with their abbreviations, legal drafting, chip design, a scientific subfield with its own notation. If a domain expert reads the base model's output and says the words are right but nobody writes like that, the texture is missing.

A proprietary corpus. Twenty years of internal engineering reports, where the value is in the accumulated conventions rather than in any single fact. Note carefully: this is not the way to make the model *know* your documents. That is retrieval. This is the way to make it *write like* them.

A code ecosystem the model does not know. An internal framework, a rare language, a DSL.

The volume, and it is large

This is where continued pretraining separates from everything else in the course. Instruction tuning works with thousands of examples. Domain adaptation works with tokens, by the hundred million:

- **A light register shift:** 100 million to 1 billion tokens.
- **A genuine domain:** 1 to 10 billion tokens.
- **A new language:** 10 billion and up, typically with a tokenizer change as well.

A billion tokens is roughly 4 GB of text — perhaps two million documents. If you have 50,000 internal documents, you have something like 50 million tokens, which is enough to shift register slightly and not enough to teach a domain. Knowing that before you start saves a month.

Forgetting, at full strength

The most predictable outcome of continued pretraining is that the model stops following instructions. You have just trained it, on raw text, that its job is to continue documents. The instruction-following behaviour that made it useful came from a later stage you have now partly undone.

This is expected, and the pipeline accounts for it:

```
continued pretraining → SFT again → preference tuning (optional)
```

You redo instruction tuning afterwards. Plan for it, budget for it, and keep the SFT dataset from before, because you will need it.

The second mitigation is replay, at higher proportions than elsewhere: mix 10 to 30% of general-purpose text into your domain corpus. Without it, a model trained on 100% clinical notes becomes a model that can only produce clinical notes.

Configuration that differs from SFT

- **No masking.** Loss on every token; the whole document is the target.
- **Pack aggressively.** Documents are long and variable; packing with correct document separation is standard here, not optional.
- **Lower learning rate** than instruction tuning if you are doing it fully — $1e-5$ or below. You are nudging a pretrained model, and pretraining-scale damage is easy to do.
- **Higher LoRA rank if using an adapter** — 128 or 256, and include the embedding and output layers via `modules_to_save`. A rank-16 adapter is a reweighting, and domain adaptation is not a reweighting. Many practitioners simply do this stage fully rather than with an adapter, for that reason.
- **One epoch.** With hundreds of millions of tokens you do not need a second pass, and a second pass is where memorisation lives.

The cheaper alternative to try first

Before committing to a billion tokens, test the hypothesis directly. Take 200 domain questions. Answer them with the base model given a few in-domain examples in the prompt, and with retrieval over your corpus. If either gets you most of the way, the domain gap was smaller than it felt.

And if the actual complaint is that the model does not know specific facts from your documents, continued pretraining is the wrong tool, expensively. It will improve how the model writes about your domain and it will not make it reliable about what is in any particular document. Facts belong in the context window. That rule has not changed for being expensive to ignore.

BEFORE YOU MOVE ON

A team continued-pretrains a chat model on 400 million tokens of internal engineering reports. The output now reads convincingly like their reports, but the model no longer follows instructions properly. What happened, and what is the fix?

- A. The learning rate was too high, destroying the instruction-tuned weights; rerunning at a lower rate would preserve both.
- B. Unmasked raw text taught the model to continue documents rather than answer, so SFT must be redone.
- C. Four hundred million tokens is far too many for a chat model, which should be adapted with at most a few million.
- D. The reports contain no dialogue, so the model lost its chat template; reapplying the template at inference restores the behaviour.

47 · 9 min

Training a head instead of generating tokens

THE ONE THING TO KEEP

Swapping the token head for a linear head turns a decoder into a scorer that costs one forward pass instead of a generation loop, but a preference-trained reward model learns an ordering with no absolute scale, and a decoder classifier pools the last non-padding token — so padding side and `pad_token_id` must both be set.

Replace the last layer

A decoder's final layer projects a hidden state to vocabulary size, producing a distribution over the next token. Replace that projection with a small linear layer producing one number, or six, and you have a scorer or a classifier built on the same pretrained body.

```
from transformers import AutoModelForSequenceClassification
model = AutoModelForSequenceClassification.from_pretrained(
    "Qwen/Qwen2.5-1.5B", num_labels=1) # 1 = regression / reward
```

This is what a reward model is. It is also what a quality scorer, a safety classifier, a routing model and a reranker are. In every case you have traded the ability to write text for

a single cheap forward pass — no autoregressive decoding, no sampling, a fixed and tiny output.

Why the cost difference is large

Generating a 200-token judgement takes 200 sequential forward passes through the model. Scoring with a head takes one. At serving time that is roughly two orders of magnitude of difference in latency and cost, for a task whose output was always going to be a number.

If you are currently prompting a model to "rate this response from 1 to 10" and parsing the digit out of its reply, you are paying for a generation loop to produce one token, and a head would do it better.

The padding bug

Here is the mistake that costs people a day. For a decoder-based classifier, the pooled representation is taken from the **last non-padding token**. If your tokenizer pads on the right, the final position of a short sequence is a pad token, and naive pooling reads the hidden state of padding.

Two defences, and you want both:

```
tok.padding_side = "left"          # for decoder-based classification
model.config.pad_token_id = tok.pad_token_id  # so pooling can find
the end
```

If `pad_token_id` is unset — and for several base models it is —

`AutoModelForSequenceClassification` cannot locate the last real token and will either raise or quietly use the wrong position. Encoder models such as BERT and ModernBERT use a dedicated classification token and do not have this problem, which is one more reason to prefer an encoder when the task is classification.

Losses, by output type

- **Binary classification:** binary cross-entropy, one output.
- **Multi-class:** cross-entropy, n outputs.
- **Multi-label:** binary cross-entropy per label, n independent sigmoids. Do not use softmax here; the labels are not mutually exclusive.
- **Regression to a human score:** mean squared error, one output. Standardise the targets first.

- **Reward model from preferences:** the Bradley-Terry loss. Score both responses, maximise the probability that the chosen scores higher:

```
loss = -F.logsigmoid(score_chosen - score_rejected).mean()
```

That last one is worth noticing: a reward model never learns an absolute scale. It learns an ordering. A score of 3.2 means nothing on its own — only that it is above or below another response's score. Teams that build a reward model and then set a fixed threshold on its output are assuming a calibration that was never trained.

Calibration

A classifier's sigmoid output is not a probability until you have checked that it is. Fine-tuned models are routinely overconfident: the outputs cluster near 0 and 1, and the ones at 0.9 are correct rather less than 90% of the time.

Measure it before you rely on it. Bin your held-out predictions by confidence and compare each bin's mean confidence to its accuracy. If the gap is large, temperature scaling — a single scalar fitted on a validation set, dividing the logits before the sigmoid — fixes most of it for the cost of ten lines. This matters specifically when you are using the confidence to route: "fall back to the expensive model below 0.7" is a rule about a number that has to mean something.

Where this belongs in a fine-tuning project

Three places, and all three are common:

1. **The reward model**, if you are doing PPO.
2. **An automatic evaluator**. Train a small scorer on a few thousand human judgments of your own outputs, and you have an evaluation you can run on every checkpoint for free. Check its agreement with held-out human labels before trusting it, and re-check when the model changes.
3. **A router or a safety gate** in front of a generative model, deciding which requests go to the cheap path, the expensive path, or a refusal.

All three are small models, quick to train, cheap to serve, and easy to evaluate properly — which makes them the most under-used tool in the applied fine-tuning kit.

BEFORE YOU MOVE ON

A team trains a reward model on preference pairs with the Bradley-Terry loss, then deploys it with a rule that any response scoring below 2.0 is rejected. Why is this rule unsound?

- A. Reward models output logits rather than scores, so 2.0 has to be passed through a sigmoid before it can be compared.
- B. Reward models drift during training, so any fixed threshold becomes stale and must be recomputed each epoch.
- C. The loss trains only the ordering, so a single score has no calibrated absolute meaning.
- D. Scores depend on response length, so the threshold rejects short responses regardless of quality.

CASE STUDY · MODULE 5

The verifier that could be gamed, and the baseline that could not

A small education-technology company in Hyderabad, building a model that writes SQL against a school management database for non-technical administrators.

The product let a head teacher type "how many girls in class eight have not paid the second term fee" and get a table. Behind it, a model wrote SQL against a fixed schema of forty tables, the query ran against a read-only replica, and the rows came back.

It worked about seventy-eight per cent of the time. The failures were not syntax — constrained decoding had taken care of that months earlier — but queries that ran cleanly and returned the wrong rows: a join through the wrong key, a term filter omitted, a count of enrolments where a count of pupils was wanted.

The team had something most teams do not: a real verifier. For eleven hundred questions they had a hand-written reference query and its result set on a test database. A generated query could be executed and its rows compared. Correct or not correct, instantly, at no cost, unlimited times.

That made reinforcement learning with a verifiable reward genuinely available, and the engineer leading it, Prakash, wanted to use GRPO. The argument was reasonable. Sampling eight completions per question and scoring them against the reference gives the group its own baseline, removes the value model that makes PPO expensive, and is the method behind the recent open reasoning models.

His colleague Divya proposed the dull alternative first: rejection-sampling fine-tuning. Generate eight completions per question with the existing model, keep only those whose result set matches the reference, train the model on those with ordinary supervised fine-tuning. No reinforcement learning, no reward function to tune, no generation loop inside a training loop.

The cost on each side was concrete. GRPO meant generating eight responses per prompt at every step, inside training, which put the run in GPU-days rather than hours; on a rented A100 that was the best part of a week of compute and, more importantly, two weeks of Prakash's time on a method neither of them had run before. Rejection sampling meant one generation pass over eleven hundred questions — a few hours — and then a training run they had done ten times.

Divya's argument was the one the module makes: rejection-sampling fine-tuning on verified outputs is the baseline you should beat before you reach for any of this.

Prakash's counter was that the interesting failures were the ones the model never got right in eight samples, and rejection sampling learns nothing from those — it simply discards them.

That turned out to be the testable question, and it took an afternoon. They sampled the base model a hundred times on fifty held-out questions and counted how often the correct query appeared at least once. For thirty-nine of the fifty it did. For eleven it never did, across a hundred attempts.

That number decided a great deal. If a method that sharpens a model's distribution towards paths it can already sometimes find is what reinforcement learning on verifiable rewards mostly does — and whether it teaches new capability or elicits existing capability is genuinely argued, with careful experiments on both sides — then the eleven questions the base model never solved were unlikely to be rescued by it either.

They also wrote down, before running anything, what a reward hack would look like. A query that returns an empty result set matches a reference whose answer is legitimately empty. A query that selects everything and lets the comparison sort it out. A query with the answer hard-coded as a literal, which is the SQL equivalent of returning 4181 when the test asks for the seventeenth term.

STOP HERE

Before turning the page: what would you have done, and what would it have cost if you were wrong? Write it down. The point of the next section is to compare.

WHAT HAPPENED

They ran rejection sampling first. Eight samples per question at temperature 0.7, kept only where the result set matched, which produced 6,400 verified query pairs from the eleven hundred questions. One supervised fine-tuning run of two hours. Accuracy on the held-out set went from 0.78 to 0.89.

Then they ran the GRPO experiment anyway, on a two-week budget agreed in advance with a stated stopping rule. It reached 0.91. The extra two points sat just outside a three-seed noise floor of about 0.8 points, so it was a real effect and a small one.

Two findings made the difference more interesting than the number. The fraction of groups in which all eight sampled responses scored identically — logged from the first step, because Prakash had read that this is the commonest reason a GRPO run does nothing — ran at 41 per cent early on and had to be brought down with a small shaping reward for producing a well-formed query at all. And on the eleven questions the base

model had never solved in a hundred attempts, the GRPO model solved none of them either.

The reward hack they had predicted arrived in week two. Queries that returned empty result sets rose from 2 per cent to 9 per cent of outputs, because an empty set matched every reference whose correct answer was also empty. A held-out reward they never trained against — a check that the query touched at least the tables the reference touched — caught it.

They shipped the rejection-sampled model. The GRPO work is in the ablation log, with its two points, its noise floor, and a note that the eleven unsolved questions need a different kind of help.

Sampling the base model a hundred times on fifty questions cost an afternoon. What decision did that measurement inform, and why?

Whether reinforcement learning on a verifiable reward could plausibly help. There is an open argument about whether such training teaches new reasoning capability or mostly sharpens the model's distribution towards paths it could already occasionally produce, and the evidence is genuinely mixed. If it is mostly sharpening, a problem the base model never solves in a hundred attempts is unlikely to be rescued by it. Eleven of fifty were in that category, which capped what the expensive method could deliver before any of it was run.

Forty-one per cent of GRPO groups produced no gradient early in training. What causes that, and what is the standard remedy?

GRPO computes each response's advantage relative to its own group: the reward minus the group mean, divided by the group standard deviation. If every response in a group scores the same — which happens when a pure correctness reward is zero for all eight early in training — the standard deviation is zero and the advantage is undefined or zero for all of them, so the group contributes nothing. The remedy is a shaping term that is achievable from the start, such as a small reward for producing a well-formed query, which keeps groups differentiated until correctness begins to fire.

The empty-result-set exploit was predicted before training and detected by a reward the team never trained against. Why is a held-out reward the strongest instrument here?

Because optimisation pressure is precisely the process of finding where a proxy and the intent come apart, and the trained-against reward cannot report its own exploitation — by construction it is going up. A second measure of the same intent, never optimised, gives an independent reading: when the trained reward rises while the held-out one is flat or falling, the gap is the hack. It costs one extra scoring function, and here it named the artefact — a query returning nothing satisfies every reference whose correct answer is also nothing.

Check what stuck

6 questions, one right answer each. This one is for you, not for a record: answers and the reason behind each are at the back. If two questions from the same module go wrong, re-read that module before the exam rather than after.

- 1 Your model has a hedging habit you dislike. Why can supervised fine-tuning only address it indirectly?
 - A. Because the loss only says 'make this more likely' and never 'make that less likely'.
 - B. Because hedges are usually masked out of the loss by completion-only training.
 - C. Because the hedge appears in the prompt rather than the completion, and gradients do not flow to prompt positions under any masking scheme.
 - D. Because hedging is a pretraining behaviour that fine-tuning cannot reach at all.
- 2 Why must the rejected side of a preference pair come from the model you are about to train?
 - A. Because the reference model's log-probabilities are only defined for text that it generated itself, so the ratio at the centre of the loss would be undefined.
 - B. Because different models use different tokenizers, so the pair would not encode correctly.
 - C. Because off-policy pairs train the model to avoid behaviour it never had, and the gradient can drag neighbouring plausible text down with it.
 - D. Because DPO requires the two responses to have the same token count.
- 3 During a DPO run, rewards/chosen and rewards/rejected are both falling while the margin widens. What is this?
 - A. A bug in the reference model, which is being updated when it should be frozen.
 - B. An expected property of the objective, which penalises only the difference between the two.
 - C. Evidence that beta is too high and the model is anchored too tightly to its reference.
 - D. A sign that the learning rate is too low, since a healthy run raises the chosen reward within the first hundred steps regardless of what the data contains.

- 4 Your product collects thumbs-up and thumbs-down on single responses. Which method is built for that data?
- A. SimPO, because it is reference-free and therefore does not need the second response that a preference pair would otherwise have to supply.
 - B. DPO, because a thumbs-up and a thumbs-down on different prompts form a valid pair.
 - C. PPO, because binary feedback is what a reward model is trained on.
 - D. KTO, because it takes unpaired binary labels and the two sets need not be the same size.
- 5 A GRPO run produces almost no gradient. What should you log first?
- A. The fraction of groups in which every response scored identically.
 - B. The number of optimiser steps per epoch, since a group-based method takes far fewer of them than supervised fine-tuning does at the same batch size.
 - C. The KL divergence from the reference model.
 - D. The mean response length across each group.
- 6 You have fifty thousand internal documents, about fifty million tokens, and want the model to write in your domain's register. What does that volume buy?
- A. Nothing at all, because continued pretraining has no effect on register until the tokenizer has been extended with vocabulary from the corpus itself.
 - B. A genuine domain adaptation, since fifty million tokens is comfortably above the threshold.
 - C. A slight register shift at best; a genuine domain wants one to ten billion tokens.
 - D. Reliable knowledge of what is in those documents, so retrieval becomes unnecessary.

MODULE 6

The mechanics of a run

The part where things silently break: tokenizers, special tokens, masking, seeds, checkpoints and multi-GPU. Most failed fine-tunes are not bad ideas, they are configuration bugs nobody decoded.

BY THE END OF THIS MODULE YOU CAN

Configure and launch a fine-tuning job with the correct chat template, masking and tokenizer handling, and diagnose from the logs alone whether a run is not learning, diverging, or memorising

48 · 9 min

Chat Templates and Loss Masking

THE ONE THING TO KEEP

Decode one batch's unmasked labels before every run; it catches most fine-tuning bugs in five seconds.

The template is a contract

An instruction-tuned model was trained with special tokens marking who is speaking and where a turn ends. Get them wrong by one character and you are training a slightly different model than the one you serve.

Never write the format by hand. Ask the tokenizer.

```
from transformers import AutoTokenizer
tok = AutoTokenizer.from_pretrained("Qwen/Qwen2.5-7B-Instruct")
msgs = [{"role": "user", "content": "hi"},
        {"role": "assistant", "content": "hello"}]
print(repr(tok.apply_chat_template(msgs, tokenize=False)))
```

You will get something like `'<|im_start|>user\nhi<|im_end|>\n<|im_start|>as-sistant\nhello<|im_end|>\n'`. Every family differs — Llama uses `<|start_header_id|>` and `<|eot_id|>`, Gemma uses `<start_of_turn>`. Copying a snippet from a blog post about a different model is a common and silent way to lose a day.

Two rules:

- **Training:** render the full conversation including the assistant turn and its end token. `add_generation_prompt=False`.
- **Inference:** render up to the assistant header and stop.
`add_generation_prompt=True`.

If you fine-tune a **base** model, with no `-Instruct` in the name, there is no template. You invent one, and you are then responsible for using exactly the same one at serving time. Write it down in the repo next to the weights.

Mask the prompt

By default many trainers compute loss on every token, the user's message included. That teaches the model to generate user turns: wasted capacity at best, and a strong pull toward continuing past its own answer with an invented question.

Train on assistant tokens only. Every framework has a switch for this, and the names and APIs churn between versions, so do not trust the flag. Verify the labels:

```
batch = next(iter(trainer.get_train_dataloader()))
ids, labels = batch["input_ids"][0], batch["labels"][0]
print("TRAINED ON:", tok.decode([i for i, l in zip(ids, labels) if l != -100]))
print("MASKED   :", tok.decode([i for i, l in zip(ids, labels) if l == -100]))
```

Run this once, read both lines, and you have eliminated the majority of fine-tuning bugs. What should print under TRAINED ON is your assistant's reply and its end-of-turn token, and nothing else.

The end token is the whole ballgame

The single most common report is: the model answers, then writes a new user message and answers that one too.

That is an end-token bug, essentially every time. One of:

- The end-of-turn token was stripped when you built the dataset.

- It fell inside the masked region, so no gradient ever taught the model to emit it.
- Your training template's end token differs from the one your server stops on.
- `max_length` truncates long examples just before the end token, so long answers never see one.

Check it explicitly:

```
print(tok.decode(ids[-8:]))      # last tokens of a training example
print(labels[-1] != -100)       # is the final token actually trained
                                # on?
```

A stop string at the API layer hides the symptom. Fix the labels instead.

Small things that cost hours

- **Double BOS.** `apply_chat_template` usually adds the beginning-of-sequence token. If you then tokenize with `add_special_tokens=True` you get two, and the model sees an input shape it never saw in pretraining. Pass `add_special_tokens=False` after templating.
- **Padding token.** Many base models ship without one. Setting `tok.pad_token = tok.eos_token` is standard, but then make sure padding is masked in the labels, or you train the model to emit end-of-sequence forever.
- **Packing.** Concatenating short examples to fill the context is efficient, but naive packing lets one example attend to the next. Use it only if your trainer does position-aware packing with a correct attention mask, and skip it for short instruction data where the win is small anyway.
- **System prompt.** Whatever you train with, serve with. If your training examples carry no system prompt and production sends six hundred tokens of one, every request is off-distribution in a way that quietly degrades everything.
- **The data file.** JSONL, one conversation per line, in the `messages` format. Let the trainer apply the template; do not pre-render strings unless you have to.

```
{"messages": [{"role": "user", "content": "..."}, {"role": "assistant",
"content": "..."}]}
```

Do the label-decoding check before every run. It costs five seconds.

What the label check should print

TRAINED ON

- the assistant's reply
- its end-of-turn token
- nothing else

MASKED

- the system prompt
- the user's message
- every role header and delimiter
- the padding — and check that padding is not the same id as end-of-sequence

Decode the unmasked labels of one real batch before every run. Five seconds, and it eliminates the majority of fine-tuning bugs.

BEFORE YOU MOVE ON

Your fine-tune trained cleanly and the loss curve looks healthy. At inference the model answers the question, then writes a new user message and answers that one too. What is the most likely cause?

- Sampling temperature is too high, so generation runs on past the answer.
- The end-of-turn token was never trained on — stripped, masked out, or truncated away — so the model never learned where a turn stops.
- The model needs more epochs to learn the conversation format.
- You need a stop sequence configured at the serving layer.

49 · 10 min

Running a LoRA, and What It Costs

THE ONE THING TO KEEP

GPU time is the cheap part: a 7B LoRA run costs cents, while the dataset and the evaluation cost days.

A configuration that works

Start here, and change one thing at a time.

```

from peft import LoraConfig
peft_config = LoraConfig(
    r=16,
    lora_alpha=32,                # convention: twice the rank
    lora_dropout=0.05,
    bias="none",
    task_type="CAUSAL_LM",
    target_modules=["q_proj", "k_proj", "v_proj", "o_proj",
                   "gate_proj", "up_proj", "down_proj"],
)

```

Adapting every linear layer rather than attention alone is the finding from the QLoRA work, and it is close to free: the trainable count is still under 1%.

```

from trl import SFTConfig
args = SFTConfig(
    output_dir="out",
    num_train_epochs=2,
    per_device_train_batch_size=1,
    gradient_accumulation_steps=16,  # effective batch of 16
    learning_rate=2e-4,
    lr_scheduler_type="cosine",
    warmup_ratio=0.03,
    max_length=1024,
    gradient_checkpointing=True,
    bf16=True,                      # fp16=True on a T4
    logging_steps=10,
    save_steps=100,
    optim="paged_adamw_8bit",
)

```

Pin your library versions. This part of the ecosystem changes fast enough that a six-month-old notebook usually will not run.

The four settings that matter

- **Learning rate.** LoRA wants $1e-4$ to $2e-4$, roughly ten to a hundred times higher than full fine-tuning, because you are moving very few parameters. Too low and nothing changes; too high and the model degrades in ways that look like a data problem. If outputs turn erratic, halve it before touching anything else.
- **Epochs.** One to three. On a small dataset, three passes is already the memorisation zone. Save a checkpoint each epoch and evaluate all of them — the best is often not the last.

- **Rank.** 8 or 16 for style and format. 32 to 64 if you are teaching a genuinely new task. Higher rank means more capacity and more forgetting; it is not a quality dial.
- **Sequence length.** Set it from your data, not from habit. Look at the token-length distribution and cover the 95th percentile. Memory grows with it, and every token past your longest real example is money spent on padding.

Where to run it

Free. Kaggle gives roughly 30 GPU-hours a week on a T4 or P100, with sessions that survive better than Colab's free tier. Colab free gives a T4 that can disconnect at any moment, so keep `save_steps` low and push checkpoints somewhere persistent. On either, QLoRA on a 7B is fine; 16-bit LoRA on a 7B is not.

Rented. On the spot-style marketplaces, an RTX 4090 (24 GB) runs roughly \$0.30-0.80 an hour, an A100 80 GB roughly \$1.20-2.00, an H100 \$2-3. These prices move constantly and vary by provider and region, so check on the day rather than trusting any number, including these.

Hosted fine-tuning APIs charge per training token, typically a few dollars for a small dataset. They handle the infrastructure and hand you back an endpoint. The trade is that you cannot inspect the run, and you usually cannot export the weights, which puts you back where you started on cost and portability.

What a run actually costs

Do the arithmetic before you rent anything. Two thousand examples averaging 600 tokens is 1.2M tokens per epoch, 2.4M for two epochs. A 7B QLoRA on a 24 GB card processes very roughly 1,500-3,000 tokens per second depending on sequence length and settings.

2.4M divided by 2,000 is about 1,200 seconds, so twenty minutes. At \$0.50 an hour that is about **17 cents**.

That number is the point of this lesson. The GPU has never been the expensive part of fine-tuning. The dataset took three days and the evaluation will take another two. Anyone selling you a fine-tuning story that is mostly about hardware is selling hardware.

Watch these while it runs

- **Loss falls, then flattens.** That is healthy. A cliff to near zero means memorisation or a masking bug.

- **Loss rising, or NaN.** Learning rate too high, or fp16 overflow on a T4. Lower the rate first.
- **Grad-norm spikes** every few hundred steps usually mean one pathological long example. Find it.
- **A flat line from step one** means the adapter is attached to nothing. Print `model.print_trainable_parameters()`; if it reports zero, your `target_modules` names do not match this architecture.

Then evaluate before you celebrate, which is the next lesson.

Reading a run while it is still running

Healthy

- loss falls, then flattens
- gradient norm settles into a band
- the three sample generations still read like your data

Not healthy

- a cliff to near zero in fifty steps: the answer is visible in the input, or the labels were not shifted
- rising loss or NaN: the learning rate, or fp16 overflow on a T4
- a flat line from step one: the adapter is attached to nothing

Print the trainable parameter count before you launch. On a misconfigured LoRA it is sometimes zero, and a run that trains nothing produces a perfectly flat, perfectly innocent-looking loss curve.

BEFORE YOU MOVE ON

800 examples, rank 64, learning rate $2e-4$, ten epochs. Training loss reaches 0.05, and on new inputs the model reproduces training answers almost word for word. What is the right read?

- It memorised. Drop to one or two epochs, lower the learning rate, lower the rank, and judge checkpoints on the held-out set rather than the loss.
- It worked — a loss of 0.05 means the model learned the task well.
- The dataset is too small; collect 8,000 more examples and rerun.
- Rank 64 is not enough capacity to generalise; raise it to 128.

The tokenizer is part of the model

THE ONE THING TO KEEP

The tokenizer is a fixed part of the model whose efficiency on your language decides cost, context and quality — and setting the pad token equal to the end-of-sequence token, then masking by id, removes the gradient that teaches the model to stop.

Fixed, and consequential

The tokenizer is not a preprocessing convenience. It is a fixed component of the model, decided at pretraining time, and its vocabulary maps one-to-one onto the rows of the embedding table. Use a different tokenizer with the same weights and you are feeding the model random words.

So: save the tokenizer beside every checkpoint, load it from the same directory as the weights, and never assume two models in the same family share one.

What tokenisation costs different languages

The cost of a request is measured in tokens, and tokens are not distributed fairly.

A modern 128,000-entry vocabulary tokenises English at roughly 1.3 tokens per word. The same tokenizer on Devanagari, Bengali or Tamil script produces substantially more — often two to four times as many tokens for text conveying the same meaning, because the training corpus contained far less of those scripts and the merge rules never learned their common sequences. Older 32,000-entry vocabularies were much worse still, frequently splitting Indic words into individual characters.

Three consequences follow, and they are not abstract:

- **Cost.** The same question costs three times as much to answer in Hindi as in English on a per-token price.
- **Context.** A 4,000-token window holds a third as much Hindi as English.
- **Quality.** A word split into six fragments is harder to learn a representation for than one split into two. Long-range dependencies span more positions.

When you are choosing a base model for a non-English task, tokeniser efficiency on your language is a first-class criterion, and it takes one minute to check:

```
for name in ["meta-llama/Llama-3.1-8B", "Qwen/Qwen2.5-7B",
"sarvamai/sarvam-1"]:
    t = AutoTokenizer.from_pretrained(name)
    print(name, len(t(sample_hindi_paragraph)["input_ids"]))
```

Run it on a paragraph of your actual text. The differences between model families are often larger than any quality difference you would get from fine-tuning.

The pad token trap

Many base models ship with no pad token, because pretraining packed sequences and never needed one. Training code needs one, so the near-universal workaround is:

```
tok.pad_token = tok.eos_token          # very common, and a trap
```

Here is the trap. Your collator masks padding positions with `-100` so they contribute no loss. If pad and end-of-sequence are the same id, a mask built by matching on the pad id will also mask **the real end-of-sequence token at the end of each answer**.

The model then never receives a gradient teaching it to stop. It trains happily, the loss looks fine, and at inference it generates until it hits the token limit — answering your question, then answering a question you did not ask, then starting a new conversation with itself.

Two correct fixes:

1. **Add a distinct pad token.** `tok.add_special_tokens({"pad_token": "<|pad|>"})`, then resize embeddings. Cleanest, and it costs one embedding row.
2. **Build the mask from the attention mask, not from the token id.** Positions are padding because the collator padded them, and the attention mask already records that. This is what well-written collators do.

Either way, verify it — the verification is in the debugging lesson and takes five seconds.

The double-BOS bug

`apply_chat_template` adds the beginning-of-sequence token as part of the template. If you then pass the resulting string through the tokenizer with default settings, the tokenizer adds another one.

```
text = tok.apply_chat_template(msgs, tokenize=False)
ids = tok(text)["input_ids"]          # BOS is now there twice
print(tok.convert_ids_to_tokens(ids[:3]))
```

A model trained with two BOS tokens and served with one — or the reverse — has a systematic mismatch at the very first position, which is the position everything else conditions on. The effect ranges from mild degradation to nonsense depending on the model.

The fix is `tok(text, add_special_tokens=False)` when the template has already added them, and the check is to print the first three tokens of a real example before every run.

The other mismatches worth knowing

- **Padding side.** Left for generation and for decoder-based classification; right is the default for training collators. Getting it wrong for classification means pooling a pad token's hidden state.
- **Fast and slow tokenizers** can differ on unusual whitespace and unicode. Use the fast one consistently in both training and serving rather than mixing.
- **Chat template changes between releases.** A model's template lives in its tokenizer config and repository maintainers do update it. If you fine-tuned against one version and serve with a later one, the special tokens around each turn may differ. Pin the tokenizer revision alongside the model revision.

None of these is intellectually difficult. All of them produce a model that is subtly worse for reasons that look like a training problem, and all of them are found by printing one templated example and reading it.

BEFORE YOU MOVE ON

A fine-tuned model answers correctly and then keeps generating, inventing a new user question and answering that too. The team set `tok.pad_token = tok.eos_token` and masks padding by matching the pad id. What is the cause?

- A. The maximum generation length is set too high at inference, allowing the model to continue past a natural stopping point.
- B. Packing merged examples without a separator, so the model learned to continue from one answer into the next prompt.
- C. The eos token's embedding was reinitialised when it was reassigned as the pad token, destroying its learned representation.
- D. The mask matches every occurrence of the id, including the real end token ending each answer.

51 · 8 min

Adding tokens, and initialising them properly

THE ONE THING TO KEEP

New vocabulary rows are initialised randomly by default, which places them outside the region existing embeddings occupy — so set them to the mean of existing embeddings or of the tokens they stand for, and remember they stay frozen under LoRA unless `embed_tokens` and `lm_head` are in `modules_to_save`.

When it is worth it

Adding vocabulary is not a routine step. The three cases where it earns its cost:

1. **A control token.** A `<|tool_call|>` or `<|thinking|>` marker you want the model to emit unambiguously, which no existing token can be confused with.
2. **A pad token** the base model lacks, as the previous lesson described.
3. **A script or domain the tokenizer handles badly.** If your language tokenises at four tokens per word, extending the vocabulary with common sequences from your corpus can halve that — but this only pays off alongside continued pretraining with a substantial amount of text, because the new embeddings have to be learned from scratch.

For anything else — a product name, a set of category labels, a handful of domain terms — leave the vocabulary alone. The model represents multi-token strings perfectly well, and every added token is an untrained row in two large matrices.

The mechanics

```
tok.add_special_tokens({"additional_special_tokens": [
    <|tool_call|>"]})
model.resize_token_embeddings(len(tok), pad_to_multiple_of=64)
```

The `pad_to_multiple_of=64` rounds the vocabulary up to a multiple of 64. This is a throughput detail, not a correctness one: matrix dimensions that are multiples of 64 map cleanly onto tensor core tiles, and a vocabulary of 128,257 can be meaningfully slower than one of 128,320. The extra rows are never selected and cost a few megabytes.

The initialisation that matters

By default, `resize_token_embeddings` fills new rows with values drawn from the model's initialiser — small random numbers, unrelated to anything the model knows.

That is worse than it sounds. The existing embedding vectors occupy a particular region of the space with a particular norm; a random vector drawn from a naive distribution sits somewhere else entirely. On the input side the model receives a signal it has no representation for. On the output side the new token's logit is computed from a random row, which early in training can make it wildly over- or under-probable and produce large gradients.

The fix is to initialise the new rows near the centre of the existing distribution:

```
emb = model.get_input_embeddings().weight.data
out = model.get_output_embeddings().weight.data
n = len(new_tokens)
emb[-n:] = emb[:-n].mean(dim=0, keepdim=True)
out[-n:] = out[:-n].mean(dim=0, keepdim=True)
```

Better still, where the new token has a meaning expressible in existing tokens, initialise it as the mean of *those* tokens' embeddings. A `<|invoice|>` token initialised as the average of the embeddings for "invoice" starts in approximately the right place and converges much faster.

They will not train unless you say so

With a LoRA, the embedding table and output head are frozen like everything else. Your new token's embedding stays at whatever you initialised it to, forever.

```
LoraConfig(..., modules_to_save=["embed_tokens", "lm_head"])
```

This trains those two matrices fully and saves them in the adapter. Be aware of the size: an 8B model's embedding table is $128,256 \times 4,096 \approx 525$ million parameters, so your adapter file grows from about 160 MB to over 2 GB once both matrices are included. That is usually acceptable and it is always surprising the first time.

Tied embeddings

Many models — particularly smaller ones — tie the input embedding matrix and the output head, using the same weights for both. Check `model.config.tie_word_embeddings`.

When they are tied, you have one matrix, not two: resizing handles both together, `modules_to_save` needs only one entry, and the mean-initialisation above should be applied once. When they are untied, they are genuinely independent and both need attention. Assuming the wrong one is a quiet source of a new token that works on input and misbehaves on output.

After adding, before training

Three checks:

```
assert tok.convert_tokens_to_ids("<|tool_call|>") != tok.unk_token_id
assert len(tok) <= model.get_input_embeddings().weight.shape[0]
print(tok.tokenize("a <|tool_call|> b")) # must be one token, not
five
```

The third is the one that catches the real mistake. A special token added to the vocabulary but not registered as *special* will be split into its constituent characters by the tokenizer, and you will have trained the model to emit a seven-token sequence that only looks like a control token.

BEFORE YOU MOVE ON

A team adds a `<|tool_call|>` token, trains a LoRA with the standard seven target modules, and finds the model never emits the new token. What did they miss?

- A. The embedding table and output head are frozen under LoRA, so the new token's randomly initialised rows were never trained.
- B. The token needs a corresponding entry in the chat template, without which the tokenizer strips it during templating.
- C. New tokens are appended after the original vocabulary size, so the output head's softmax never assigns them probability mass.
- D. Special tokens are excluded from the loss by default, so the model received no gradient on positions where it should appear.

52 · 10 min

A working config, line by line

THE ONE THING TO KEEP

A working config is a set of decisions, not defaults — the instruct variant as base, all-linear targets, $2e-4$ for LoRA, non-reentrant checkpointing so gradients reach the adapter, and length-grouped batching instead of packing until packing is verified.

The whole thing, then the reasons

This is a complete supervised fine-tuning setup for an 8B model on a 24 GB card, using TRL, PEFT and bitsandbytes — all free, all open source.

```

from transformers import AutoModelForCausalLM, AutoTokenizer,
BitsAndBytesConfig
from peft import LoraConfig
from trl import SFTConfig, SFTTrainer
import torch

MODEL = "meta-llama/Llama-3.1-8B-Instruct"

bnb = BitsAndBytesConfig(load_in_4bit=True, bnb_4bit_quant_type="nf4",
                        bnb_4bit_use_double_quant=True,
                        bnb_4bit_compute_dtype=torch.bfloat16)

tok = AutoTokenizer.from_pretrained(MODEL)
model = AutoModelForCausalLM.from_pretrained(
    MODEL, quantization_config=bnb, torch_dtype=torch.bfloat16,
    attn_implementation="flash_attention_2", device_map={"": 0})

peft_cfg = LoraConfig(
    r=16, lora_alpha=32, lora_dropout=0.05, bias="none",
    target_modules="all-linear", task_type="CAUSAL_LM")

args = SFTConfig(
    output_dir="out/triage-v3",
    num_train_epochs=2,
    per_device_train_batch_size=2,
    gradient_accumulation_steps=8,
    learning_rate=2e-4,
    lr_scheduler_type="cosine",
    warmup_ratio=0.03,
    max_grad_norm=1.0,
    bf16=True,
    optim="paged_adamw_8bit",
    gradient_checkpointing=True,
    gradient_checkpointing_kwargs={"use_reentrant": False},
    max_length=2048,
    packing=False,
    group_by_length=True,
    assistant_only_loss=True,
    logging_steps=10,
    eval_strategy="steps", eval_steps=100,
    save_steps=200, save_total_limit=3,
    seed=0,
    report_to="tensorboard")

trainer = SFTTrainer(model=model, args=args, peft_config=peft_cfg,
                    train_dataset=train, eval_dataset=val,
                    processing_class=tok)

trainer.train()

```

The lines that are decisions

Llama-3.1-8B-Instruct, not -8B. Start from the instruct variant when your task is instruction-shaped. It already follows instructions, has a chat template, and needs far less data to specialise. Start from the base only for continued pretraining or when you want no inherited behaviour at all.

attn_implementation="flash_attention_2". Halves memory on the attention path and speeds it up. Requires Ampere or newer; on a T4 omit this and PyTorch's SDPA backend will be used instead.

device_map={"": 0}. Everything on GPU 0. Using "auto" on a single-GPU machine can silently offload layers to CPU when memory is tight, turning a slow run into a very slow one for reasons that are not logged prominently.

target_modules="all-linear". Architecture-agnostic, so this config survives a change of model family.

num_train_epochs=2. Two passes. Check the resulting step count — with 6,000 examples at an effective batch of 16 that is 750 steps, comfortably in range.

learning_rate=2e-4. Standard for LoRA. For full fine-tuning this line becomes $2e-5$, and forgetting that is the single most expensive typo in this file.

gradient_checkpointing_kwargs={"use_reentrant": False}. The non-reentrant implementation is the one that composes correctly with PEFT. With the old default you can get a run where gradients never reach the adapter and the loss is perfectly flat.

packing=False with group_by_length=True. The conservative choice: most of the padding saving, none of the cross-document attention risk. Turn packing on once you have verified that your version separates documents correctly.

assistant_only_loss=True. Loss on the assistant turns only. This requires the model's chat template to mark generation spans; if it does not, use the completion-only colator instead and pass the exact response template string.

optim="paged_adamw_8bit". Paged, so a spike on one long batch is absorbed by CPU memory rather than ending the run.

seed=0. Recorded, so the run is as reproducible as the hardware allows.

What changes for a full fine-tune

The same file, five lines different:

```
# drop the BitsAndBytesConfig entirely, and peft_config=None
learning_rate = 2e-5          # not 2e-4
weight_decay  = 0.01          # not 0
optim         = "adamw_bnb_8bit"
num_train_epochs = 1          # full fine-tuning forgets faster
```

And the memory arithmetic from module three applies, so this configuration belongs on a sharded multi-GPU setup for anything above about 1.5B parameters.

The evaluation callback

One addition worth making from the first run rather than the fifth. Alongside `eval_loss`, generate on three fixed prompts at every evaluation step and log the text, and run the canary and refusal sets if they are cheap enough:

```
class Watch(TrainerCallback):
    def on_evaluate(self, args, state, control, model=None, **kw):
        for p in FIXED_PROMPTS:
            print(state.global_step, generate(model, p)[:200])
```

Every checkpoint then arrives with evidence attached, which is what makes checkpoint selection on a task metric possible rather than aspirational.

The four lines to change first

If something is wrong, in this order: `learning_rate`, `max_length` (check how many examples it truncates), `num_train_epochs`, and whether the masking line is actually taking effect.

Doing this without writing Python

Axolotl and Llama-Factory express the same configuration as a YAML file and handle the wiring:

```

base_model: meta-llama/Llama-3.1-8B-Instruct
load_in_4bit: true
adapter: qlora
lora_r: 16
lora_alpha: 32
lora_target_linear: true
sequence_len: 2048
micro_batch_size: 2
gradient_accumulation_steps: 8
num_epochs: 2
learning_rate: 0.0002
lr_scheduler: cosine
warmup_ratio: 0.03
bf16: auto
gradient_checkpointing: true
datasets:
  - path: ./train.jsonl
    type: chat_template

```

Both are free. For a first project a config file is better than a script, because a file diffs cleanly, versions naturally, and is much harder to make subtly inconsistent between two runs.

BEFORE YOU MOVE ON

A run using QLoRA with gradient checkpointing produces a perfectly flat loss curve from step one, with a non-zero trainable-parameter count and a correct learning rate. Which config line is the likeliest culprit?

- A. `packing=False`, which leaves padding in every batch and dilutes the gradient signal to nothing.
- B. Reentrant gradient checkpointing, which can stop gradients reaching the adapter.
- C. `group_by_length=True`, which sorts examples so early batches contain only the shortest sequences.
- D. `optim="paged_adamw_8bit"`, whose quantised state cannot represent the small updates a LoRA adapter produces.

53 · 10 min

Debugging a run

THE ONE THING TO KEEP

Decode the unmasked labels of one real batch before every run — it should print exactly the assistant's answer and nothing else — and if a pipeline cannot memorise eight examples in a hundred steps, the problem is the code rather than the data.

The four checks to run before every job

These take under a minute together and catch the majority of fine-tuning bugs.

1. Print one templated example.

```
print(tok.apply_chat_template(train[0]["messages"], tokenize=False))
```

Read it. Every special token where you expect it, roles marked correctly, exactly one end-of-turn marker at the end.

2. Decode the unmasked labels of one real batch.

```
b = next(iter(trainer.get_train_dataloader()))
lbl = b["labels"][0]
print(tok.decode([t for t in lbl if t != -100]))
```

This should print **exactly the assistant's answer** — not the prompt, not the system message, not nothing. If it prints the whole conversation, your masking is off. If it prints nothing, your response template string does not match. If it prints the answer without a trailing end token, you have the pad-equals-eos problem.

This single check is the highest-value five seconds in applied fine-tuning.

3. Print the trainable parameter count.

```
model.print_trainable_parameters()
```

Zero means your target modules do not exist in this architecture. A number equal to the full model means your adapter was not applied.

4. Print the first three token ids. One BOS, not two.

The overfit test

Before a real run, take eight examples and train on them for 100 steps with everything else identical.

The loss should fall to near zero. The model should reproduce those eight answers nearly verbatim.

If it cannot memorise eight examples, your pipeline is broken and no amount of data or hyperparameter tuning will help. This test costs two minutes and it separates "my data is not good enough" from "my code does not work", which are otherwise indistinguishable from the outside and are debugged in completely different ways.

Symptom to mechanism

Loss is exactly 0.0. Every label is -100. The response template did not match; nothing is being trained on.

Loss is flat and non-zero. Learning rate far too low, no trainable parameters, or reentrant gradient checkpointing blocking gradients to the adapter. Check the parameter count first.

Loss drops to near zero within 50 steps. The answer is visible in the input. Either labels were not shifted, or the completion also appears in the prompt — which happens when a log export includes the response in a context field.

Loss is NaN. In fp16, overflow — lower the learning rate, confirm warmup, check for a corrupted example. In bf16, NaN is much rarer and usually means bad data: an empty string, an enormous unicode blob. Find it by logging the index of the batch where it first appears.

Loss spikes then recovers. Usually one pathological example. `max_grad_norm=1.0` limits the damage. If it recurs at the same step each epoch, find that example and read it.

Out of memory at step 40, not step 1. Your longest sequences are later in the file. Either sort by length to fail fast, or lower `max_length`. With `group_by_length=True` the long batches cluster, which makes this both more likely and more predictable.

The model never stops generating. Missing end-of-sequence in training. Cause: pad equals eos and masking by id, or truncation cutting the end token off long examples, or packing without separators.

Generations are empty. The model emits end-of-sequence immediately. Usually the reverse masking error — trained on prompts rather than completions.

Evaluation looks excellent, production does not. Leakage between your training and evaluation sets, or a formatting difference between how you evaluate and how you serve. Check the second first: print the exact string sent in both paths and diff them.

Look at generations, not only at loss

Every two hundred steps, generate on three fixed prompts and log the output.

```
class Samples(TrainerCallback):
    def on_evaluate(self, args, state, control, model=None, **kw):
        for p in FIXED_PROMPTS:
            print(state.global_step, generate(model, p)[:200])
```

Loss curves hide everything that matters about text. A model that has started repeating itself, dropped its formatting, or begun every answer with the same clause shows all of that in three samples and none of it in the loss.

The discipline

Change one thing per run. Write down what you changed and what happened, in a file, in the repository. Four runs into a debugging session nobody remembers which combination produced the good result, and the log is what turns a frustrating afternoon into a result you can act on.

BEFORE YOU MOVE ON

A run reports a training loss of exactly 0.0 from the very first step. What does this indicate?

- A. The model has already learned the task perfectly from pretraining, so there is nothing left to fit.
- B. The learning rate is zero, so no weights update and the loss remains at its initial value.
- C. Every label is -100, so nothing contributes to the loss — usually a template that matched nothing.
- D. The labels were not shifted relative to the inputs, so the model is copying the token it can already see.

What to log, and how to read it

THE ONE THING TO KEEP

Gradient norm, tokens per step, token accuracy and three generated samples tell you things loss cannot — a run heading for divergence is visible in the gradient norm hundreds of steps before the loss moves, and stylistic collapse is visible only in the samples.

Loss is one of six numbers

Most people watch the loss and nothing else. Here is the full set worth having, and what each one tells you that loss does not.

Gradient norm. The most under-used signal in fine-tuning. It is the magnitude of the whole gradient vector before clipping, and the shape of its curve is diagnostic:

- A healthy run: moderate at the start, settling to a stable band.
- Isolated tall spikes: individual pathological examples. Harmless if `max_grad_norm` is clipping them, and worth finding anyway — log the batch index.
- Steadily climbing: instability. The learning rate is too high and the run is heading for divergence, and you can see it several hundred steps before the loss does.
- Near zero from the start: nothing is learning. Check trainable parameters before anything else.

Tokens per optimiser step. Not logged by default, and it is the true unit of batch size. Compute it from the attention mask and log it once at the start; if it is under a few thousand, your gradients are noisy regardless of what the example count says.

Token accuracy on unmasked positions. The fraction of target positions where the argmax is correct. Free to compute, interpretable without exponentiating anything, and less swayed than loss by a handful of very surprising tokens. When loss and token accuracy disagree, a small number of examples are dominating the loss — and finding them usually finds a data problem.

Learning rate. Log it, because a schedule that is not doing what you think — no warmup, a decay set for the wrong number of steps — is invisible otherwise and completely obvious on a plot.

Evaluation loss. Every hundred steps or so on a held-out slice, to see the training-validation gap open.

Mean generated length. Generate on a few fixed prompts at each evaluation and record the mean token count. This is the cheapest detector of the length drift that preference methods cause and that model judges reward.

The samples

Three fixed prompts, generated at every evaluation, logged as text.

This is not a metric and it is more informative than all of them. A loss curve cannot show you that the model has started every answer with the same clause, dropped its bullet points, begun repeating itself at the end of long answers, or quietly lost its British spelling. Three samples show all of that at a glance.

Six things to log, and what each says that loss does not

Gradient norm	A run heading for divergence climbs here hundreds of steps before the loss moves. Near zero from the start means nothing is learning.
Tokens per optimiser step	The true unit of batch size, and not logged by default. Under a few thousand, your gradients are noisy whatever the example count says.
Token accuracy on unmasked positions	Free, interpretable without exponentiating anything, and less swayed by a handful of very surprising tokens.
Learning rate	A schedule set for the wrong number of steps is invisible in every other number and obvious on a plot.
Mean generated length	The cheapest detector of the length drift that preference methods cause and that model judges reward.
Three fixed generations	Not a metric, and more informative than all of them. Stylistic collapse appears here and nowhere else.

Pick the three prompts deliberately: one typical input, one edge case, and one that should be declined. Keep the whole log directory beside the checkpoint, because it is a few megabytes and it is the only answer to "why is version two better than version three".

Pick the three deliberately: one typical input, one edge case, one input that should be declined.

The tools, all free

TensorBoard is the plain answer. It is included with PyTorch, runs entirely locally, needs no account, and handles scalars and text. `report_to="tensorboard"` and `tensorboard --logdir out/`.

Trackio is a newer free, open-source tracker with a Weights & Biases-compatible API, runnable locally or hosted on a Hugging Face Space at no cost. Useful when you want a shareable dashboard without a commercial account.

MLflow and **Aim** are free and self-hostable, with more experiment-management machinery than a single fine-tune needs but real value once you have twenty runs to compare.

Weights & Biases is the commercial standard with a free personal tier. Convenient, and it sends your run metadata to a third party — which for a project training on customer data is a question to answer before the first run rather than after.

Estimating the run at step 20

```
# after 20 steps
sec_per_step = elapsed / 20
total_hours = sec_per_step * total_steps / 3600
print(f"{sec_per_step:.2f}s/step → {total_hours:.1f}h →
    ${total_hours*0.60:.2f}")
```

Run this twenty steps in. It tells you whether your six-hour plan is a six-hour plan or a fifty-hour one, while cancelling is still free. Throughput per step is roughly constant after warmup, so the extrapolation is reliable.

What to keep after the run

The TensorBoard directory is a few megabytes. Keep it beside the checkpoint, in the same place, forever. Three months later, when somebody asks why version 2 is better than version 3, the answer is in those curves — and if they were only in a terminal you closed, the answer is a shrug.

BEFORE YOU MOVE ON

A run's loss is falling smoothly, but the gradient norm has been climbing steadily for 400 steps. What does this most likely mean?

- A. The model is learning faster as it adapts, and a rising gradient norm reflects the increasing signal available from the data.
- B. Gradient clipping is active, so the reported norm reflects the clipping threshold rather than the true gradient.
- C. The evaluation set has drifted from the training distribution, which raises gradient magnitudes on evaluation batches.
- D. The run is becoming unstable at the current learning rate, and the loss will follow with a spike or divergence later.

55 · 8 min

Checkpointing and resuming

THE ONE THING TO KEEP

A resumable checkpoint holds optimiser state, scheduler position and RNG state as well as weights, so resuming is only correct when the data order and the number of devices are unchanged — and on interruptible instances the checkpoint must live somewhere the instance does not.

What a checkpoint contains

A checkpoint is not only weights. A complete one holds five things:

1. The trainable weights — the adapter, or the full model.
2. The optimiser state: AdamW's two moment buffers per trainable parameter.
3. The learning-rate scheduler's position.
4. The random number generator states, so data shuffling resumes identically.
5. The trainer state: global step, epoch, logged metrics.

Only the first is needed to *use* the model. All five are needed to *continue* the run.

The size follows: a 42M-parameter LoRA is about 160 MB in bf16, and its optimiser state is another 340 MB, so a resumable checkpoint is roughly half a gigabyte. Keeping

three is a gigabyte and a half, which is nothing. Keeping three full 8B checkpoints is 130 GB, which is not.

The settings

```
SFTConfig(
    save_steps=200,
    save_total_limit=3,
    save_safetensors=True,
    load_best_model_at_end=False,    # see below
    output_dir="out/triage-v3",
)
```

Set `save_steps` so that a crash costs you at most ten or fifteen minutes of work. At three seconds per step that is every 200 to 300 steps.

`save_total_limit=3` keeps the disk bounded and is enough to step back from a checkpoint that turned out to be past the peak.

`load_best_model_at_end=False` is a deliberate choice. "Best" is defined by validation loss, and — as the epochs lesson established — validation loss is not what you want to select on. Keep all three and choose on your task metric afterwards, with the canary set beside it.

Resuming correctly

```
trainer.train(resume_from_checkpoint=True)
```

This restores all five components and fast-forwards the data loader to where it stopped. Two conditions make that correct, and both are easy to break:

The data must be in the same order. Resuming replays the shuffle from the saved RNG state, so a dataset that has been re-shuffled with a different seed, or had rows added, produces a different sequence. The resumed run will not see the examples it was going to see, and may see some twice.

The world size must match. Resuming a four-GPU run on two GPUs changes the effective batch size, which changes the number of optimiser steps per epoch, which puts the scheduler in the wrong place. It will run and it will not be the run you started.

Interruptible instances

Spot and community instances are roughly half price and can be reclaimed with little warning. With checkpointing in place that is an inconvenience rather than a loss, provided one thing: **the checkpoint must not be on the instance's local disk.**

```
# free Colab or Kaggle
output_dir = "/content/drive/MyDrive/runs/triage-v3"

# or push each checkpoint to a private Hub repo
SFTConfig(push_to_hub=True, hub_model_id="you/triage-v3",
          hub_private_repo=True, hub_strategy="checkpoint")
```

Private repositories on the Hugging Face Hub are free, and `hub_strategy="checkpoint"` pushes the resumable checkpoint rather than only the final weights. A reclaimed instance then costs you a restart, not a run.

Saving the adapter alone

For distribution you want the weights without the optimiser state:

```
trainer.model.save_pretrained("out/adapter")    # ~160 MB
tok.save_pretrained("out/adapter")             # always alongside
```

Save the tokenizer with it, every time. An adapter without its tokenizer is a file that loads and produces subtly wrong results, because somebody six months later will pair it with a slightly different tokenizer revision and have no way to know.

Evaluating checkpoints without stopping the run

Checkpoints accumulate while the job is still going, and evaluating them afterwards serialises two things that could overlap. Point a second process at the output directory and have it score each new checkpoint as it appears:

```
while true; do
  for c in out/triage-v3/checkpoint-*; do
    [ -f "$c/.scored" ] || { python evaluate.py "$c" >> results.jsonl;
    touch "$c/.scored"; }
  done
  sleep 120
done
```

On a single-GPU machine this competes for the card, so it suits a rented second card or a CPU-servable small model. Where it works, the run finishes with its selection evidence already gathered rather than a queue of checkpoints to get through.

Two things to record inside the directory

Write a small `run.json` next to the checkpoint:

```
{
  "base_model": "meta-llama/Llama-3.1-8B-Instruct",
  "base_revision": "0e9e39f",
  "dataset_sha256": "9f2c1e...",
  "config_sha256": "71ba04...",
  "git_commit": "4f1a09",
  "seed": 0,
  "gpu": "RTX 4090", "transformers": "4.5x.y", "trl": "0.x.y"
}
```

This is the only thing that makes a checkpoint identifiable later. Without it you have a directory of tensors and a memory. The `base_revision` field in particular matters: model repositories are updated in place, and "the same base model" six months later may not be the same bytes.

BEFORE YOU MOVE ON

A four-GPU run is interrupted at step 900 and resumed on two GPUs with the same per-device batch size. Why is the resumed run not a continuation of the original?

- A. The effective batch halves, so steps per epoch and the scheduler's position no longer match.
- B. Checkpoints store per-device optimiser shards, so two GPUs cannot reconstruct state saved by four.
- C. The random number generator state is per device, so two GPUs replay only half the shuffle and see half the dataset.
- D. Gradient accumulation is recomputed on resume, which resets the learning rate schedule to its warmup phase.

56 · 9 min

Reproducibility, and the noise floor

THE ONE THING TO KEEP

Set a seed and record the git commit, config hash, dataset hash, base revision and library versions — then run three seeds to find your noise floor, because an improvement smaller than twice that spread is a coin flip you have named a decision.

Bit-identical is harder than it looks

```
from transformers import set_seed
set_seed(0)      # python, numpy, torch, cuda
```

That covers weight initialisation, dropout, data shuffling and sampling. It does not give you a bit-identical run, because GPU reductions are not deterministic by default: floating-point addition is not associative, and the order in which thousands of parallel threads accumulate a sum varies between executions.

Full determinism is available and it is expensive:

```
import torch, os
os.environ["CUBLAS_WORKSPACE_CONFIG"] = ":4096:8"
torch.use_deterministic_algorithms(True)
torch.backends.cudnn.benchmark = False
```

This costs throughput, and some operations have no deterministic implementation at all — the call will raise rather than silently be non-deterministic, which is the right behaviour and is also how you discover that your model uses one of them.

And even with all of it, results differ across GPU models, driver versions and library versions. Reproducibility across machines is not achievable by flags.

Reproducible enough

The practical goal is not bit-identity. It is: *somebody with this repository can produce a model that performs the same within noise*. That needs a record, not a flag.

```
{
  "git_commit": "4f1a09",
  "config_sha256": "71ba04...",
  "dataset_sha256": "9f2c1e...",
  "base_model": "meta-llama/Llama-3.1-8B-Instruct",
  "base_revision": "0e9e39f",
  "seed": 0,
  "gpu": "1xRTX 4090",
  "versions": {"torch": "...", "transformers": "...", "trl": "...", "peft": "..."}
}
```

The two hashes are the ones people omit and the two that matter most. "We trained on the customer dataset" is not a reproducible statement; a sha256 of the exact JSONL is.

Pin the base model revision too. Model repositories are edited — tokenizer configs get updated, chat templates get fixed — and `main` six months from now is a different artefact.

The measurement everybody skips

Here is the part that changes how you read every result in this course.

Run the same configuration three times with seeds 0, 1 and 2. Evaluate all three on your held-out set. The spread between them is your **noise floor**.

```
seed 0: 0.871
seed 1: 0.858
seed 2: 0.884
→ range 2.6 points, sd ≈ 1.3
```

Now look at the last four "improvements" you recorded. If they were one or two points, you measured nothing. You measured seed variance and attributed it to a decision.

This is the single most common error in applied fine-tuning, and three extra runs — a few hours and a few pounds — is the whole cost of not making it. Until you know your noise floor, you cannot distinguish a real effect from a coin flip, and every ablation you run is uninterpretable.

Once you have it, the rule is simple: an improvement smaller than about twice your seed standard deviation is not an improvement. Either accept that, or run each configuration three times and compare the means — which is what the difference is worth measuring, if it is worth measuring.

Where the variance comes from

Worth knowing, because it tells you how to reduce it:

- **Data order** dominates on small datasets. With 800 examples, the sequence in which the model sees them changes the result materially.
- **Adapter initialisation** — the random A matrix.
- **Dropout masks.**
- **Non-deterministic reductions**, the smallest contributor.

More data narrows the first and largest source, which is one more reason the dataset is the lever that matters.

Reporting a result with its noise

Once you have the floor, the honest form of a result is a mean and a range, not a single figure:

baseline	0.871 ± 1.3	(3 seeds)	
+ replay	0.870 ± 1.1	(3 seeds)	task unchanged
	canary 0.62 → 0.81		the actual finding

Note what that reporting does. It makes the task number a non-finding explicitly, and puts the canary movement — which is well outside the noise — where the reader's attention belongs. A single-figure table would have invited an argument about 0.871 versus 0.870.

The evaluation must be fixed

Seed your evaluation too. If you generate at temperature 0.7 for evaluation, the same model scores differently on consecutive runs and you have added a second source of noise on top of the training one. Evaluate greedily, or with a fixed seed and enough samples to average, and record which you did.

A result without a stated evaluation procedure is not a result. A result without a noise floor is not a comparison.

BEFORE YOU MOVE ON

A team runs the same configuration with three seeds and gets 0.871, 0.858 and 0.884. They then report that switching LoRA rank from 16 to 32 improved their metric from 0.871 to 0.889. What is the problem?

- A. Comparing against the seed-0 run alone is invalid; they should have used the highest of the three seeds as the baseline.
- B. The 1.8-point gain is smaller than the 2.6-point spread they measured across seeds.
- C. Rank 32 needs a different alpha, so the comparison is confounded by scaling rather than by noise.
- D. Three seeds is too few to estimate variance, so the noise floor itself is unreliable and no comparison can be made.

57 · 9 min

Multi-GPU without tears

THE ONE THING TO KEEP

Adding GPUs multiplies the effective batch size by the device count, so gradient accumulation must be divided or the learning rate raised — and sharded strategies over PCIe often give back in communication what they gain in parallelism, which is why a single card is the faster path for any run under about twelve hours.

Two different reasons to use more than one card

They need different tools and confusing them wastes days.

The model fits on one GPU and you want it faster. Use data parallelism: a full copy of the model on each card, different batches on each, gradients averaged across them every step. In PyTorch this is DDP.

The model does not fit on one GPU. Use sharding: FSDP or DeepSpeed ZeRO, splitting the optimiser state, the gradients and optionally the parameters across cards. This is the subject of the full-fine-tuning lesson.

For almost all LoRA and QLoRA work, the answer is the first. A quantised 8B model fits on one card, so more cards means more throughput, not more capacity.

Launching

```
accelerate launch --num_processes 4 train.py
# or
torchrun --nproc_per_node 4 train.py
```

Your training script does not change. Both launchers start one process per GPU and set the environment variables the libraries read.

The silent batch-size multiplication

This is the error that catches everybody once.

```
effective = per_device_batch_size × gradient_accumulation_steps ×
num_gpus
```

A config tuned on one card with `per_device=2`, `accum=8` has an effective batch of 16. Launch it on four GPUs, unchanged, and the effective batch is **64**. Four times the tokens per step, a quarter of the optimiser steps for the same data, and a learning rate now sized for a batch a quarter as large.

The run completes. The loss curve is plausible. The result is worse and nobody knows why.

When you add GPUs, divide `gradient_accumulation_steps` by the same factor — `accum=2` on four cards keeps the effective batch at 16 — or deliberately keep the larger batch and raise the learning rate to match. Either is fine. Doing neither is the bug.

When two GPUs are slower than one

Sharded strategies move tensors between cards constantly. On a machine where the GPUs are connected by PCIe rather than NVLink — which includes most consumer multi-GPU builds and many cheaper rented instances — that communication can dominate.

Rules of thumb:

- **DDP over PCIe:** fine. One gradient all-reduce per step, and with gradient accumulation you can make that per *several* forward passes.
- **FSDP or ZeRO-3 over PCIe:** often disappointing. Parameters are gathered and released every layer, every step.

- **ZeRO-3 with CPU offload over PCIe:** expect a large slowdown. It is for making things possible, not fast.

Measure rather than assume. Run 50 steps on one card and 50 on four, and compare seconds per step against tokens per step. If four cards are not giving you close to three times the throughput, the communication is eating it.

Two practical traps

Every process runs the whole script. Without guards, four processes will each download your dataset, each write to the same log file, and each try to push to the Hub. Guard side effects:

```
if accelerator.is_main_process:
    dataset.save_to_disk(path)
```

The Hugging Face `Trainer` handles logging and saving correctly on its own; it is your own code around it that needs the guard.

Data loader workers. `data_loader_num_workers=4` with four GPUs is sixteen worker processes. On a small dataset this is pure overhead and a common source of hangs at startup. Start at 2 and raise it only if the GPUs are visibly waiting.

When not to bother

A two-hour LoRA run does not need multi-GPU. The complexity — a launcher, a batch-size recalculation, the debugging of a hang that only occurs with four processes — costs more than the ninety minutes it saves.

The threshold in practice is around a day. If a run would take more than about twelve hours on one card, the extra cards are worth the setup. Below that, a single card and a config you trust is the faster path to an answer, and one GPU has the additional merit that every bug is your bug rather than a distributed one.

There is also a better use for a second card that nobody suggests: run two *different* configurations on the two GPUs simultaneously rather than one configuration across both. An ablation finishes in the same wall-clock time as a single run, the setup is nothing more than two terminals, and you end the day with a comparison rather than with one number that arrived faster.

BEFORE YOU MOVE ON

A config tuned on one GPU (`per_device=2`, `accum=8`, `lr=2e-4`) is launched unchanged on four GPUs with DDP. The run finishes and the model is worse. What happened?

- A. DDP averages gradients across devices, which shrinks the update by a factor of four and undertrains the adapter.
- B. Each GPU trained its own adapter copy, and only rank zero's was saved, discarding three quarters of the learning.
- C. The effective batch became 64 rather than 16, so there were a quarter as many optimiser steps.
- D. Four processes each shuffled the data with the same seed, so every device trained on the identical batches.

CASE STUDY · MODULE 6

The model that would not stop talking

A state helpline in Jaipur, three days before a launch, with a model that answers each question and then invents the next one.

The helpline took calls and messages about land record certificates. A model had been fine-tuned to draft the written replies, and it was good: the register was right, the document names were right, the steps were in the order the counter clerks used.

It also, on roughly one reply in three, answered the question, then wrote a new question from the citizen, then answered that one too. Sometimes for four turns. The text was fluent throughout and increasingly invented.

The team had three days. The launch date had been in a press note.

Their first fix took twenty minutes and worked. The serving layer accepted a stop string, so they set it to the sequence that preceded a user turn, and the extra conversation disappeared from every response. They demonstrated it to the programme director on a Tuesday afternoon and she was satisfied.

The engineer, Imran, was not, and he said so in a way that made the next two days awkward. A stop string at the API layer hides a symptom. Whatever had taught the model to continue was still in the weights, and it would surface somewhere else: in a reply that trailed off mid-sentence because the stop string fired early, in a long answer that ran to the token limit, in the next fine-tune built on the same pipeline.

He spent an hour on the four checks the course prescribes before any run, which nobody had done before the original run.

The first check printed one templated example. It looked correct.

The second check was the one that answered it. He decoded the unmasked labels of a real batch — the tokens the model was actually being scored on — and printed them beside the masked ones. The trained region was the assistant's reply, correctly, and it ended one token early. The end-of-turn token was in the masked region.

The cause was two lines of ordinary setup. The base model shipped without a pad token, as many do, so somebody had written the near-universal workaround of setting the pad token equal to the end-of-sequence token. The collator then built its label mask by matching on the pad id. Every position holding that id was masked — including the real end-of-sequence token at the end of every answer.

The model had therefore never once received a gradient teaching it to stop. It had been trained, thousands of times, on answers that simply ended, with nothing marking the ending as something to produce. At inference it did the only thing it had been shown: it kept going.

The loss curve had been healthy throughout. Nothing in the training logs was wrong. The bug was invisible in every number the run produced and visible in five seconds of decoded labels.

Imran's proposal was to add a distinct pad token, resize the embeddings, rebuild the mask from the attention mask rather than from the token id, and retrain. The run itself was two and a half hours. The evaluation was another half day. Realistically it put the launch two days late.

Against that, the director's position was not unreasonable. The stop string worked. The demonstration had been fine. A press note had a date on it, and a delay would have to be explained upwards.

STOP HERE

Before turning the page: what would you have done, and what would it have cost if you were wrong? Write it down. The point of the next section is to compare.

WHAT HAPPENED

They compromised in a way that turned out to be the right shape. The launch went ahead on the announced date with the stop string in place, and the retraining ran the same night with the pad token fixed.

The retrained adapter was evaluated against the same two hundred held-out messages three days later and replaced the first one in a quiet deployment the following week. The stop string stayed, because it costs nothing and a fine-tuned model still occasionally overruns.

Two things came out of the week that outlasted the incident. The four checks went into the repository as a script that the training job runs and prints before it starts: one templated example, the decoded unmasked labels, the trainable parameter count, and the first three token ids. Nobody has to remember to run it.

And when the second model in the same programme was trained two months later, on a different base with a different template, the same script printed a trained region that contained the entire conversation including the citizen's message. That would have

taught the model to write both sides. It was caught before the run, in five seconds, at a cost of nothing.

The training loss curve looked healthy throughout. Explain why a masking bug of this kind is invisible in the loss.

Loss is computed only over unmasked positions, so masking the end-of-turn token simply removes it from the average rather than making the average worse. The model was learning the task it was actually given — produce the body of the answer — and doing it well, so the curve fell and flattened normally. Masked positions are not down-weighted, they are absent, which means a wrong mask makes the model learn a different task from the one you think while every number on the dashboard stays reassuring.

Setting the pad token equal to the end-of-sequence token is standard advice. What makes it a trap, and what are the two correct fixes?

It is a trap when the collator builds its label mask by matching on the pad id, because the real end-of-sequence token at the end of each answer carries that same id and is masked with the padding. The model then never receives a gradient teaching it to stop. The two fixes are to add a distinct pad token and resize the embeddings, which costs one embedding row; or to build the mask from the attention mask instead, since the attention mask already records which positions the collator padded.

The stop string worked and cost twenty minutes. Make the case for retraining anyway, and the case against.

For: a stop string hides a symptom without removing its cause, so the same defect surfaces elsewhere — replies truncated when the string fires early, long answers running to the token limit with no learned ending, and the same broken pipeline inherited by the next fine-tune. Against: the launch date was public, the fix was demonstrated and working, and two days of delay has a cost somebody has to explain. The resolution used here takes both seriously — ship on the date with the mitigation, retrain the same night, and replace the artefact quietly once it has been evaluated.

Check what stuck

6 questions, one right answer each. This one is for you, not for a record: answers and the reason behind each are at the back. If two questions from the same module go wrong, re-read that module before the exam rather than after.

- 1 You decode the unmasked labels of one real batch and nothing prints. What does that mean?
 - A. The batch contained only padding, so a larger batch size is needed.
 - B. Every label is minus one hundred: the response template string did not match.
 - C. The decode call needs special tokens to be kept rather than skipped.
 - D. The adapter has not been attached yet, so at this point in the run the trainer has no labels against which to compute a loss.

- 2 Before a real run you train on eight examples for a hundred steps and the loss does not approach zero. What have you learned?
 - A. That the learning rate is too low, which is the only thing that can prevent memorisation of a set this small within a hundred steps.
 - B. That eight examples is too few for any model to memorise.
 - C. That the dataset needs more variety before the run will work.
 - D. That the pipeline is broken, and no amount of data or hyperparameter tuning will help.

- 3 You apply the chat template to get a string, then tokenize that string with default settings. What have you done?
 - A. Added a second beginning-of-sequence token at the position everything else conditions on.
 - B. Truncated the conversation to the tokenizer's default maximum length.
 - C. Lost the end-of-turn markers, because tokenizing a rendered string discards the role information that the template encoded into it.
 - D. Nothing, because the chat template helper does not add special tokens.

- 4 A run's gradient norm has been climbing steadily for four hundred steps while the loss still looks fine. What does that tell you?
- A. That the effective batch size is too large, since a bigger batch produces a larger gradient vector and therefore a larger norm at every step.
 - B. That one pathological example is being clipped, which is harmless.
 - C. That the run is becoming unstable and is heading for divergence.
 - D. That the learning-rate schedule has finished warming up.
- 5 You resume a four-GPU run on two GPUs from a complete checkpoint. What happens?
- A. It runs correctly but takes about twice as long.
 - B. It runs, with a halved effective batch, so the scheduler is now in the wrong place.
 - C. It runs and silently reshard the optimiser state, which is the only component of a checkpoint that depends on the number of devices involved.
 - D. It refuses to start, because the world size is recorded in the checkpoint.
- 6 Three runs of one configuration with seeds 0, 1 and 2 score 0.871, 0.858 and 0.884. A change you made scores 0.879. What can you conclude?
- A. That the change hurt, since 0.879 is below the best of the three.
 - B. That the evaluation set is too small, which is the only source of variation that three runs of a single configuration can possibly reveal.
 - C. That the change helped, since 0.879 is above the mean of the three.
 - D. Nothing: 0.879 sits inside the spread you get from the seed alone.

MODULE 7

Fine-tuning beyond a chat model

Embedding models, rerankers, classifiers, speech, vision-language and image generators. Smaller models, smaller datasets, larger wins — and in several of these the cheaper technique beats fine-tuning outright.

BY THE END OF THIS MODULE YOU CAN

Adapt an embedding model, a reranker, a text classifier, a speech model, a vision-language model or an image generator, and name in each case what differs from fine-tuning a text LM and which cheaper technique should be tried first

58 · 10 min

Fine-tuning an embedding model

THE ONE THING TO KEEP

Embedding models train on query-passage pairs with in-batch negatives, which makes batch size the real hyperparameter — gradient caching buys a 512-wide batch on a small card — and mined hard negatives must be taken from below the top ranks, because the top ones are often correct answers you would be teaching the model to reject.

The highest return in this course

An embedding model is 20 to 350 million parameters. It trains in minutes on free hardware. It needs hundreds of pairs rather than thousands of examples. And because it sets the recall ceiling for everything downstream, improving it improves every generator you ever put behind it.

This is the cheapest substantial win available in applied fine-tuning, and it is the one most often skipped in favour of training the expensive model that sits after it.

The data is pairs

You need queries paired with the passage that answers them. Nothing else is required:

```
{"anchor": "how long do refunds take", "positive": "Refunds are processed within 5 working days of the returned item arriving at our warehouse..."}
```

Three ways to get them:

- **Click logs.** A search query and the result somebody opened. Free, plentiful, noisy.
- **Generated.** For each passage in your corpus, ask a model to write three questions it answers. Cheap, and it covers your whole corpus rather than only what people have searched for. Filter by checking that your current retriever ranks the passage highly for the generated question — if it does not, keep the pair, because that is exactly the case worth learning.
- **Existing question-answer content.** FAQs, support macros, documentation with headings phrased as questions.

Five hundred pairs is enough to see a real change. Five thousand is comfortable.

The loss, and why batch size is the hyperparameter

The standard objective is multiple-negatives ranking loss. For each anchor in a batch, its own positive should score higher than *every other positive in the batch*, which serve as negatives for free.

The consequence is unusual and important: **the batch size is the number of negatives**. A batch of 16 asks the model to pick the right passage out of 16. A batch of 128 asks it to pick out of 128, which is a much harder task and produces a much better model.

So unlike almost everywhere else in this course, here you want the batch as large as you can fit. And when you cannot fit it, there is a technique for exactly this:

```
from sentence_transformers.losses import
CachedMultipleNegativesRankingLoss
loss = CachedMultipleNegativesRankingLoss(model, mini_batch_size=16)
```

Gradient caching computes embeddings in small chunks without gradients, then re-plays them, giving the gradient of an enormous batch with the memory of a small one. It makes batch sizes of 512 or more reachable on a single modest card.

Hard negatives, and the trap in them

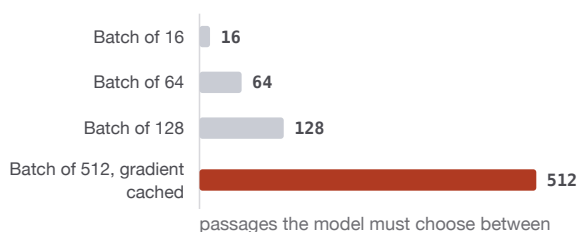
In-batch negatives are random, so most are trivially irrelevant. The model learns to separate refunds from shipping and never learns to separate a refund from an exchange.

Hard negatives fix this: for each query, retrieve the top results with your current model and keep the ones that are *not* the gold passage.

```
from sentence_transformers.util import mine_hard_negatives
triplets = mine_hard_negatives(dataset, model, range_min=10,
                                range_max=50,
                                max_score=0.8, num_negatives=3)
```

Those `range_min` and `max_score` arguments are the trap being handled. The very top results are frequently *also correct* — your corpus has several passages about refunds — and training on a correct passage as a negative teaches the model to push apart two things that belong together. Mining from rank 10 to 50 rather than from rank 1, and capping similarity, avoids most of it. Read twenty mined negatives by hand before training on fifty thousand.

In-batch negatives: the batch size is the difficulty



Each anchor's positive must outscore every other positive in the batch. Unlike almost everywhere else in this course you want the batch as large as it will go, and gradient caching reaches 512 on one modest card.

Matryoshka, if you will truncate

Matryoshka loss trains the model so that the first 64, 128 and 256 dimensions of its output are each independently useful. You then store 768-dimensional vectors for reranking and search over 128-dimensional ones, cutting index memory by six times for a small quality cost. It is one extra wrapper around your loss and it is worth it on any index above a few million vectors.

Evaluate on retrieval, not on loss

```
from sentence_transformers.evaluation import
InformationRetrievalEvaluator
ev = InformationRetrievalEvaluator(queries, corpus, relevant_docs)
```

Recall at 5 and 10, NDCG at 10, on held-out queries against your **whole corpus** — not against a small sample of it. Retrieving the right passage from 200 candidates is a different and much easier task than retrieving it from 200,000, and a number from the first tells you nothing about the second.

The free path, end to end

Base models: `all-MiniLM-L6-v2` at 22M parameters for speed, `bge-base-en-v1.5` or `e5-base-v2` at around 110M for quality, and multilingual variants where you need them. All free, all on the Hub.

Training 5,000 pairs on a 110M model takes a few minutes on a free T4, and the 22M model trains on a laptop CPU. The resulting model is a few hundred megabytes, serves at thousands of embeddings per second, and runs anywhere.

There is no cheaper intervention in this field with a comparable effect on a retrieval system's quality.

BEFORE YOU MOVE ON

A team mines hard negatives by taking the top 3 retrieved passages that are not the labelled gold passage, and retrieval quality gets worse after training. Why?

- A. Three negatives per query is too few, so the in-batch negatives dominate and the hard negatives have no effect.
- B. Mining with the current model creates a feedback loop in which the model reinforces its own ranking errors.
- C. Hard negatives require a triplet loss rather than multiple-negatives ranking loss, and mixing them produces an inconsistent objective.
- D. The top-ranked non-gold passages are frequently also relevant, so training pushes apart passages that should be close together.

59 · 8 min

Training a reranker

THE ONE THING TO KEEP

A cross-encoder reads query and passage together so it ranks far better than a bi-encoder, but it must run once per candidate — which is why it reranks 50 results rather than searching a corpus, and why the pretrained one should be measured before you train your own.

Why a cross-encoder is better

An embedding model encodes the query and the passage separately, into two vectors that are compared afterwards. Whatever the passage says about the query has to survive being compressed into a fixed vector before the query is ever seen.

A cross-encoder puts both into the model at once — [query] [SEP] [passage] — and produces one relevance score. Every token of the query can attend to every token of the passage. That is why rerankers beat bi-encoders consistently, often by a wide margin on the harder queries.

The price is that the score cannot be precomputed. A bi-encoder embeds your corpus once, and a query is one vector comparison against an index. A cross-encoder must run the model once per candidate, at query time.

The two-stage architecture

So you use both:

1. **Retrieve 50 to 100 candidates** with the bi-encoder plus BM25. Fast, over the whole corpus.
2. **Rerank those candidates** with the cross-encoder. Slow per item, but only 50 items.
3. **Pass the top 5** to the generator.

The arithmetic that makes this work: a 22M-parameter cross-encoder scores 50 query-passage pairs in roughly 15 to 40 milliseconds on a GPU, or a few hundred milliseconds on a CPU. Doing the same over a 200,000-document corpus would take minutes. The two-stage design buys cross-encoder quality at bi-encoder latency.

Training it

```
from sentence_transformers import CrossEncoder
from sentence_transformers.cross_encoder.losses import
BinaryCrossEntropyLoss

model = CrossEncoder("cross-encoder/ms-marco-MiniLM-L6-v2",
num_labels=1)
```

The data is triples of query, passage and label. With binary cross-entropy the label is 1 for relevant and 0 for not — and the ratio matters. Around four negatives per positive is a common starting point; far more negatives and the model learns to say "no" to everything, which scores well on accuracy and ranks nothing.

The negatives should be **mined from your retriever's actual output**, as with embedding training. A reranker's entire job is to sort the candidates your first stage returns, so training it on random negatives trains it for a task it will never face.

The listwise alternative

Binary labels treat each pair independently, but ranking is about order. Losses that compare a positive against several negatives for the same query within one forward group — the family that recent reranker work uses — optimise the ordering directly and generally produce better rankings from the same data.

`CachedMultipleNegativesRankingLoss` has a cross-encoder counterpart in recent sentence-transformers versions, and it is worth preferring where available.

Evaluating a reranker properly

The metric is the ranking quality of the *reranked list*, not classification accuracy on pairs. NDCG at 10 and MRR at 10 are the standard measures.

The honest comparison is against your first stage, on the same candidates:

```
stage 1 (bi-encoder + BM25), NDCG@10 = 0.61
+ off-the-shelf reranker,    NDCG@10 = 0.74
+ fine-tuned reranker,      NDCG@10 = 0.79
```

Note what that shows. The off-the-shelf reranker delivered most of the gain, and fine-tuning added a smaller increment. **Try the pretrained reranker before training one.** The public cross-encoders are trained on very large relevance datasets and are

strong out of the box; on many corpora they are simply good enough, and you have spent nothing.

The latency budget

The decision that constrains everything is how many candidates you can afford to rerank.

- 22M parameters, 50 candidates: tens of milliseconds on a GPU. Comfortable.
- 278M parameters, 50 candidates: low hundreds of milliseconds. Noticeable in an interactive product.
- A large LLM-based reranker: hundreds of milliseconds to seconds. Only for offline or high-value queries.

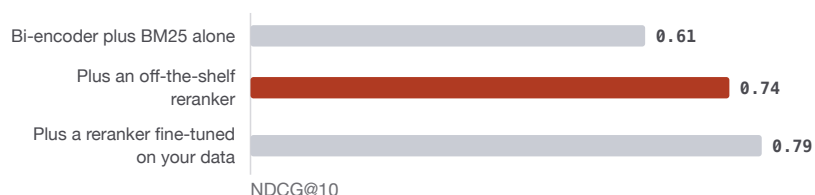
Halving the candidate count halves the cost, and the recall you lose by reranking 25 instead of 50 is usually small — measure it, because it is the cheapest latency win in the pipeline.

Where it belongs in the course

Back in module one the argument was that retrieval should be fixed before the generator. This lesson and the previous one are what that means concretely. A fine-tuned embedding model raises what can be found; a reranker raises what reaches the top.

Together they typically cost a day and a few hundred megabytes, and they move end-to-end quality further than a generator fine-tune costing ten times as much.

Reranking, measured on the same candidates



The off-the-shelf cross-encoder delivered most of the gain and cost nothing. Measure the pretrained one before you train your own, and mine your negatives from your retriever's real output rather than at random.

BEFORE YOU MOVE ON

Why can a cross-encoder not replace the bi-encoder for first-stage retrieval over a 200,000-document corpus?

- A. Its score depends on the query, so nothing can be precomputed and the model must run once per document at query time.
- B. Cross-encoders produce scores rather than vectors, and ranking requires a vector space to compute distances in.
- C. Cross-encoders have shorter context windows, so most corpus documents would be truncated beyond usefulness.
- D. Cross-encoders are trained on relevance labels rather than similarity, so their scores are not comparable across different queries.

60 · 9 min

Fine-tuning a classifier properly

THE ONE THING TO KEEP

A 150M encoder with `id2label` in its config, a weighted loss or adjusted thresholds for imbalance, macro F1 with the confusion matrix read rather than the summary line, and a four-point learning curve that tells you whether another week of labelling is worth it.

The setup

Module one argued that a fixed set of labels wants an encoder rather than a generative model. This is how that is done.

```
from transformers import AutoModelForSequenceClassification,
AutoTokenizer

MODEL = "answerdotai/ModernBERT-base"      # 149M; or microsoft/deberta-
v3-base
tok = AutoTokenizer.from_pretrained(MODEL)
model = AutoModelForSequenceClassification.from_pretrained(
    MODEL, num_labels=6, id2label=ID2LABEL, label2id=LABEL2ID)
```

Pass `id2label` and `label2id`. They are saved in the config, which means the model that comes back off disk in six months says what its outputs mean. Without them you

have a tensor of six numbers and a note in somebody's head.

Full fine-tuning at this size costs about 6 GB of GPU memory at a learning rate of $2e-5$ for 3 to 5 epochs. No adapters, no quantisation; the model is small enough that all the machinery from earlier modules is unnecessary.

Class imbalance, which you will have

Real label distributions are not uniform. A six-way routing task might be 62% one class and 3% another. Trained naively, the model learns to predict the majority class and scores 62% accuracy while being useless.

Three treatments, in order of preference:

Weighted loss. Weight each class inversely to its frequency. One override:

```
weights = torch.tensor(len(y) / (n_classes * np.bincount(y)),
dtype=torch.float)
loss = F.cross_entropy(logits, labels,
weight=weights.to(logits.device))
```

Threshold adjustment after training. Train normally, then move the decision boundary per class using a validation set. Often better than reweighting, because it separates "what does the model believe" from "what do we do about it", and lets you change the second without retraining.

Resampling. Oversample the minority or undersample the majority. Simple, and oversampling duplicates examples, which invites memorisation of your rarest and most precious class.

Measure with macro F1, not accuracy

Accuracy is dominated by the majority class and will tell you a useless model is a good one. Macro-averaged F1 weights every class equally, so the 3% class counts as much as the 62% one.

Print the per-class report, always:

```
from sklearn.metrics import classification_report, confusion_matrix
print(classification_report(y_true, y_pred, target_names=LABELS,
digits=3))
print(confusion_matrix(y_true, y_pred))
```

Read the confusion matrix rather than the summary line. It tells you *which* classes are being confused with which, and that is almost always actionable: two labels that the model mixes up are usually two labels your annotators mixed up, which is a taxonomy problem rather than a model problem.

Multi-label is a different loss

If an item can carry several labels at once, the labels are not mutually exclusive and softmax is wrong. Use independent sigmoids:

```
model = AutoModelForSequenceClassification.from_pretrained(
    MODEL, num_labels=12, problem_type="multi_label_classification")
```

That `problem_type` switches the loss to binary cross-entropy per label. Then tune a threshold **per label** on validation data — a single global 0.5 is rarely right when label frequencies differ by an order of magnitude.

Calibration, if you will route on confidence

A common and sensible design is to use the classifier for confident cases and fall back to an expensive model for the rest. That rule depends on the confidence meaning something.

Fine-tuned classifiers are typically overconfident: predictions cluster near 0 and 1, and those at 0.9 are right less often than 90% of the time. Check it by binning held-out predictions by confidence and comparing each bin's mean confidence with its accuracy. Temperature scaling — one scalar fitted on validation data, dividing the logits — corrects most of the gap in about ten lines.

Knowing when you have enough data

Plot the learning curve. Train on 25%, 50%, 75% and 100% of your labels and plot macro F1 against training size.

- Still climbing steeply at 100%: more labels is your best investment.
- Flattening: more labels will not help much; better labels, better features or a different taxonomy might.

This takes four short runs on a 150M model — under an hour on free hardware — and it answers the question teams usually argue about instead: whether to spend the next week labelling.

Where the encoder stops being right

Be honest about the boundary. An encoder classifier needs labels, cannot take a new class without retraining, and degrades silently when your inputs drift. If your label set changes monthly, or you have fewer than a few hundred examples per class, prompting a general model is the better tool despite costing more per call.

BEFORE YOU MOVE ON

A six-class router reports 89% accuracy. The confusion matrix shows the rarest class, 3% of the data, is never predicted at all. Why did accuracy miss this?

- A. Accuracy excludes classes with no predictions from its denominator, so the rare class was never counted.
 - B. Accuracy averages over items, not classes, so failing a 3% class costs three points.
 - C. Accuracy is computed on the training set by default, where the model had memorised the rare class.
 - D. The rare class's F1 is undefined when no predictions are made, so it is dropped from the aggregate.
-

61 · 9 min

Fine-tuning a speech model

THE ONE THING TO KEEP

Try an initial prompt with your domain vocabulary and a larger checkpoint before training, resample to 16 kHz mono, freeze the audio encoder on small datasets, and report raw as well as normalised word error rate alongside recall on the terms that actually matter.

Try the free things first

Before training anything, two techniques cost nothing and often solve the problem.

Initial prompt or hotwords. Whisper accepts a text prompt that conditions decoding. Supplying your product names, drug names or technical vocabulary in it substantially improves their transcription with no training at all:

```
result = model.transcribe(audio, initial_prompt=
    "Addaly, Hyperdrive, Turnstile, Durable Object, bf16, QLoRA.")
```

A bigger model. The quality difference between a small and a large Whisper checkpoint on accented speech is large, and switching is a one-line change.

If the residual errors after both are still unacceptable — and for a specific accent, a noisy environment or a dense domain vocabulary they often are — then fine-tune.

What the data must look like

Audio at 16 kHz, mono, paired with an exact transcript. That is the whole requirement, and the two parts that go wrong are the sample rate and the exactness.

```
ds = ds.cast_column("audio", Audio(sampling_rate=16_000))
```

Whisper's feature extractor expects 16 kHz. Feed it 44.1 kHz audio without resampling and the model receives a signal stretched in time, which trains badly and is not flagged.

The transcripts must be verbatim, consistently. Decide once whether you write numbers as digits or words, whether you include filler sounds, whether you punctuate — and then apply it everywhere. A model trained on transcripts that are inconsistently punctuated learns to punctuate inconsistently, exactly as with text.

How much audio

Less than people expect for a targeted improvement:

- **1 to 5 hours** of labelled domain audio measurably reduces word error rate on that domain's vocabulary and speakers.
- **10 to 50 hours** for a substantial accent or dialect adaptation.
- **Hundreds of hours** for a low-resource language from a weak starting point.

Common Voice is a large, free, multilingual, CCo dataset with per-clip accent and demographic metadata, which makes it the standard starting point for accent work — and a good way to test your pipeline before you spend money recording.

The configuration

Whisper is an encoder-decoder, so it trains with a sequence-to-sequence trainer rather than a causal one. Two specifics matter:

- **Freeze the encoder** for small datasets. The audio encoder is the part that knows what speech sounds like; the decoder is where domain vocabulary lives. Freezing the encoder cuts memory and reduces forgetting.
- **LoRA works well here**, on the decoder's attention projections, and keeps the artefact small.

A learning rate around $1e-5$, a few thousand steps, and the usual checkpoint discipline.

Word error rate, and the normalisation question

WER is edit distance between the reference and the hypothesis, divided by the number of reference words. A WER of 0.12 means roughly twelve errors per hundred words.

The subtlety that makes published numbers hard to compare: Whisper's evaluation conventionally applies a text normaliser first — lowercasing, removing punctuation, expanding numbers and contractions. Normalised WER can be several points lower than raw WER on identical output.

So report both, and state which you are quoting:

```
import evaluate
wer = evaluate.load("wer")
print("raw", wer.compute(predictions=preds, references=refs))
print("normalised", wer.compute(predictions=[norm(p) for p in preds],
                                references=[norm(r) for r in refs]))
```

Which one matters depends on your product. If the transcript is read by a person, punctuation and casing are part of the quality and raw WER is the honest number. If it feeds a search index, normalised is fine.

Measure the thing you care about

WER weights every word equally, and your errors are not equally important. A transcript that gets every common word right and the drug name wrong has a low WER and is dangerous.

So add a targeted measure: recall on a list of terms that matter — product names, medications, place names, numbers. Ten lines of code, and it is frequently the metric that changes the decision.

Serving

After training, `faster-whisper` (a CTranslate2 reimplementation) runs the same weights several times faster with lower memory, including on CPU, and `whisper.cpp` runs them on a laptop or a phone. Both are free. As with every other artefact in this course: evaluate the converted model, not the one you trained, because quantisation and reimplementation both change the output.

BEFORE YOU MOVE ON

A team fine-tunes Whisper on 8 hours of call audio and reports WER falling from 0.19 to 0.09. A colleague points out the two figures are not comparable. What is the likeliest reason?

- A. The baseline used a smaller checkpoint than the fine-tuned model, so model size rather than training explains the gap.
- B. WER cannot be compared across models with different tokenizers, since edit distance depends on the tokenisation.
- C. One figure had Whisper's text normaliser applied and the other did not.
- D. Fine-tuning reduces WER on the training distribution only, so the improvement is not meaningful on other audio.

62 · 9 min

Fine-tuning a vision-language model

THE ONE THING TO KEEP

Freeze the vision encoder, train the projector and LoRA the language model, and check how many sequence positions one image actually consumes before planning memory — and for printed documents, an OCR pipeline into a text model is usually cheaper, more accurate and far easier to debug.

Three parts, and you train one of them

A vision-language model has a vision encoder that turns an image into a set of embeddings, a projector that maps those into the language model's embedding space, and the language model itself.

The standard fine-tuning arrangement:

- **Freeze the vision encoder.** It knows what images look like, it was trained on far more images than you have, and unfreezing it is the fastest way to damage the model.
- **Train the projector,** fully. It is small, and it is the piece that aligns the two spaces for your kind of image.
- **LoRA the language model.** Attention and feed-forward projections, as usual.

```
for p in model.vision_tower.parameters():
    p.requires_grad = False
```

That one loop is most of the difference between a run that works and one that does not.

Images cost tokens, and it is more than you think

An image does not occupy one position. It becomes a block of embeddings in the language model's sequence, and the count depends on resolution and on whether the model tiles large images.

Depending on the architecture and resolution, a single image can occupy anywhere from about 64 to well over 1,000 positions, and models that split a high-resolution image into tiles multiply that by the tile count. A document page at full resolution can consume several thousand tokens on its own.

The consequences are immediate:

- Activation memory scales with sequence length, so a batch of four images at high resolution can exceed what a batch of thirty-two text examples needs.
- Your effective batch size will be small — one or two per device, with heavy gradient accumulation.
- Lowering the input resolution is the most effective memory lever available, and it costs exactly the detail you lowered. For photographs of products this is usually free; for reading small print in a scanned document it is not.

Check the number before you plan the run: encode one representative image and print how many positions it produced.

The data

The messages format extends to images:

```
{
  "messages": [
    {
      "role": "user",
      "content": [
        {
          "type": "image",
        },
        {
          "type": "text",
          "text": "Which shelf is the damaged item on?"
        }
      ]
    },
    {
      "role": "assistant",
      "content": "The second shelf from the top, on the left."
    }
  ]
}
```

The image itself is passed alongside via the processor. Masking works exactly as in text: loss on the assistant turn only, and the image positions are input.

Volume is modest for a specific task — 500 to 3,000 image-and-answer pairs changes behaviour markedly, because you are adapting an already-aligned model rather than teaching it to see.

The free path

SmolVLM comes in sizes from a few hundred million parameters up, is small enough to fine-tune on a free T4, and is permissively licensed. Qwen's VL family and Llama's vision models are larger and stronger; both have LoRA recipes that fit on a single rented card.

For inference afterwards, `llama.cpp` supports several vision architectures on CPU, so the final artefact runs on a laptop.

The question to ask before any of it

For document work specifically — invoices, forms, scanned reports — a vision-language model is frequently the wrong tool, and it is the most common misapplication in this area.

The alternative is a pipeline: a dedicated OCR engine extracts text with layout, and a text model reads the result. Tesseract, PaddleOCR and the docling family are free. The pipeline is cheaper per page by a wide margin, its failures are inspectable — you can look at the OCR output and see exactly what went wrong — and on printed documents its accuracy is typically better, because OCR engines are specialists.

The vision model earns its cost when layout itself carries meaning that survives no text extraction: handwriting, photographs, charts to be interpreted, screenshots of interfaces, damage assessment.

Run both on fifty pages before committing. The comparison takes an afternoon and it decides an architecture.

BEFORE YOU MOVE ON

A team fine-tunes a vision-language model on 2,000 product photos and unfreezes everything including the vision encoder. Image understanding degrades noticeably on images unlike their training set. Why?

- A. The projector cannot adapt while the encoder is also moving, so the two components fail to stay aligned.
 - B. Unfreezing the encoder increases memory so much that the batch size collapsed, and small batches caused the instability.
 - C. Vision encoders require a different learning rate from language models, and using one rate for both saturates the encoder's weights.
 - D. The encoder's general visual representations were overwritten by gradients from a narrow set.
-

63 · 10 min

LoRA for image models, and the likeness question

THE ONE THING TO KEEP

Image LoRAs train on 10 to 30 images in under an hour, and what you caption becomes controllable while what you leave uncaptioned is baked into the subject — while likeness training is technically trivial and legally unsettled in ways that differ by jurisdiction, which makes consent and labelling the defensible position rather than a legal conclusion.

The same method, a different model

LoRA transferred from language models to image generators almost unchanged. You attach low-rank adapters to the attention projections of the denoising network, freeze everything else, and train on a small set of images with captions.

The scale is startling compared with anything else in this course: **10 to 30 images**, 1,000 to 1,500 steps, 15 to 40 minutes on a consumer GPU or a free T4. A learning rate around $1e-4$. The resulting adapter is tens of megabytes.

Three related techniques:

- **LoRA** — the general case. A subject, a style, an object, a composition.

- **DreamBooth** — a method for teaching one specific subject, using a rare token as its name and a set of class images to stop the concept swallowing the category.
- **Textual inversion** — trains only a new embedding vector, nothing else. A few kilobytes, weaker, and it composes unusually well with other adapters.

Every tool here is free: `diffusers`, `kohya_ss`, `OneTrainer`, `ComfyUI`.

The captioning rule almost nobody is told

This is the mechanism that decides whether a LoRA is usable.

What you caption becomes controllable. What you leave uncaptioned becomes baked in.

If every training image shows your subject against a white background and none of your captions mention the background, the model learns that white background is part of the subject. Ask for the subject on a beach and you will get a beach with a suspiciously white patch.

If instead your captions read "sks chair, white background", the model learns that the white background is a separate, nameable attribute — and you can then ask for something else.

So: caption everything that varies and that you want to control. Leave uncaptioned only the thing you are actually training, and use a rare token for it so it does not collide with an existing concept.

Regularisation images

Without them, DreamBooth's fine-tuning of a person tends to bleed into the whole category — you train on one face and afterwards every generated person resembles them. Class prior preservation counteracts this by mixing in generated images of the generic class, with generic captions, during training.

It is a real effect with a clear mechanism, it costs a few hundred generated images and some extra steps, and it is the difference between an adapter you can compose with others and one that dominates every prompt.

Overfitting, and what it looks like

Image LoRAs overfit visibly and fast. Symptoms:

- Every output has the same pose or composition as a training image.

- Backgrounds or artefacts from the training photos reappear regardless of the prompt.
- Prompt terms unrelated to the subject stop having any effect.

The fixes are the familiar ones — fewer steps, lower rank, lower learning rate, more varied training images — and the diagnostic is to generate at several checkpoints and look. As with text, save every few hundred steps and choose by eye rather than by loss; diffusion training loss is famously uninformative about sample quality.

Likeness, consent, and what is not settled

Training an adapter on a real person's face is technically trivial and legally and ethically the most fraught thing in this course.

What can be said plainly:

- Training on somebody's likeness without their consent is a decision you are making about them, not about a dataset.
- Presenting a generated image of a real person as a photograph of that person is a misrepresentation regardless of what any law says about it.

What cannot be said plainly, because it is genuinely unresolved and differs by country:

- Rights of publicity and personality exist in some jurisdictions and not others, with different scope and different remedies; in the United States they vary state by state.
- In the European Union a photograph of an identifiable person is personal data, and facial data used for identification is a special category with stricter conditions — but how those rules apply to model weights trained on such images, as opposed to the images themselves, is argued and not settled.
- Several jurisdictions have introduced or proposed specific rules on synthetic likenesses and disclosure, and they do not agree with each other.

This is a live area. Anybody who gives you a confident single answer is describing one jurisdiction, or their preference. If the question is commercially significant, it goes to a lawyer.

The practical position, which is defensible under any of these regimes: get consent, in writing, scoped to what you will actually do; do not train on faces scraped from the internet; label synthetic images as synthetic; and where you need a face for a persona or an avatar, use an illustrated one, because an illustration does not assert that a person exists.

Style, and the other unsettled question

Training a LoRA on a living artist's work to reproduce their style is legally contested — copyright generally protects particular works rather than styles, but whether training on those works is itself an infringing use is exactly the question being litigated in several countries, with different answers so far. It is also, separately from the law, something many artists object to strongly. Both of those facts are worth holding at once.

BEFORE YOU MOVE ON

A LoRA trained on 20 photographs of a chair, all shot against a white backdrop, with captions reading only 'sks chair'. Prompted for the chair in a garden, the outputs show a garden with a white patch behind the chair. What is the mechanism?

- A. The backdrop was never named in a caption, so the model absorbed it into the subject token's meaning.
 - B. Twenty images is too few, and the model interpolated between them rather than generalising to new scenes.
 - C. The rank was too low to represent both the chair and a varied background, so the model reused the training background.
 - D. Class prior preservation was omitted, so the chair concept bled into the generic category of furniture and carried its setting with it.
-

64 · 9 min

Getting structure without training for it

THE ONE THING TO KEEP

Constrained decoding masks illegal tokens at every step so invalid output cannot be produced, which makes it a free floor under any structured task — it guarantees syntax and nothing about truth, so keep it even after fine-tuning rather than treating the two as alternatives.

Make invalid output impossible

Fine-tuning for format makes bad output *rarer*. Constrained decoding makes it *impossible*, and it needs no training at all.

The mechanism is simple and worth understanding rather than treating as magic. At every generation step the model produces a logit for each of 128,000 tokens. A constrained decoder maintains a state machine built from your schema, works out which tokens could legally come next given what has been generated, and sets every other logit to negative infinity before sampling.

If the schema says the next character must be `"` or `}`, then every token that does not begin with one of those has zero probability. The model cannot produce invalid JSON, not because it has learned not to, but because the invalid continuations were removed from the choice.

The tools, all free

```
from outlines import models, generate
from pydantic import BaseModel

class Triage(BaseModel):
    queue: Literal["billing", "shipping", "refunds", "technical", "other"]
    urgency: int
    summary: str

result = generate.json(models.transformers("Qwen/Qwen2.5-7B-Instruct"),
    Triage)(ticket)
```

That returns a validated object, always. Alternatives: `llama.cpp` supports GBNF grammars natively; `vLLM` has guided decoding built in; `XGrammar` and `lm-format-enforcer` provide the same capability with different performance characteristics. Nothing here costs money.

Grammars are more general than JSON schemas. Anything you can express as a context-free grammar — a SQL subset, a domain-specific language, a fixed report template, a limited set of tool calls — can be enforced the same way.

What it guarantees, and what it does not

It guarantees syntax. The output parses, the enum value is one of your five, the integer is an integer, the required fields are present.

It guarantees nothing about truth. A constrained decoder will happily emit `{"queue": "billing", "urgency": 3}` for a shipping complaint. Validity and correctness are different properties, and confusing them is the commonest mistake here.

There is also a real debate about whether heavy constraints reduce answer quality. The argument for: forcing a model into a rigid shape from the first token prevents it producing the intermediate reasoning it would otherwise use. The argument against: careful comparisons find little or no degradation when the schema is reasonable and the model is competent. Both positions have supporting experiments, the effect appears to depend heavily on the task and the schema, and it is not settled.

The pragmatic response is a design that sidesteps it: let the model reason freely in a `reasoning` string field that comes *first* in the schema, then constrain the structured fields that follow. You get the thinking and the guarantee.

The poor version, for hosted APIs

If you are calling an API that offers no grammar support, validate and retry:

```
for attempt in range(3):
    out = call_model(prompt)
    try:
        return Triage.model_validate_json(out)
    except ValidationError as e:
        prompt += f"\nThat was invalid: {e}. Return only valid JSON."
    raise
```

This works, and it costs a full extra generation every time it fires. At a 4% invalid rate on a million requests that is 40,000 extra calls. Constrained decoding costs nothing per request; retries cost a generation per failure. That arithmetic is how you decide.

Many hosted providers now offer a structured-output mode that does the constraining server-side. Use it where it exists.

When fine-tuning for format is still right

Three cases survive:

1. **The schema is enormous.** A grammar with hundreds of fields makes the state machine large and the per-token masking expensive. A model that has learned the shape needs less help from the constraint.
2. **The model must choose the schema.** Constraining requires knowing which grammar to apply. If deciding *which* tool to call, or which of twelve document types this is, is itself the task, that decision cannot be constrained — though it can often be made by a cheap classifier, as module one argued.

3. **The content inside the structure is wrong.** If the JSON is valid and the values are poor, no constraint helps. That is a knowledge or capability problem, diagnosed as such.

The order to try things

1. Constrained decoding. Free, immediate, guaranteed.
2. Constrained decoding plus a few examples in the prompt, which improves the *content* inside the valid structure.
3. Fine-tuning, if the content is still wrong after both — and then keep the constraint as well, because a fine-tuned model still produces invalid output occasionally and the constraint costs nothing to keep.

That last point is the one to remember. These are not alternatives. Constrained decoding is a floor you should never remove, and fine-tuning raises what happens above it.

BEFORE YOU MOVE ON

A team adds JSON-schema constrained decoding and their valid-output rate goes from 91% to 100%, but routing accuracy is unchanged. What does this show?

- A. The constraint is too permissive; tightening the schema with stricter field types would improve accuracy too.
- B. The constraint fixed syntax, which is not what the accuracy metric measured.
- C. The 9% of outputs that were invalid were also being counted as routing errors, so accuracy should have risen by 9 points.
- D. Constrained decoding was masking the model's reasoning tokens, which cancelled out the gain from valid formatting.

Small models, where fine-tuning matters most

THE ONE THING TO KEEP

Fine-tuning helps more the smaller the model, because a small model has to be given a behaviour rather than prompted into one — and the case for on-device models is privacy, offline operation and zero marginal cost rather than price, against the specific cost that they collapse rather than degrade off-distribution.

The effect is inverse to size

Fine-tuning a frontier-scale model on a narrow task gives a modest improvement, because the model was already close. Fine-tuning a 0.5B model on the same task can transform it, because the base model was barely adequate.

The mechanism: a small model has a weak prior about how to follow instructions and a narrow range of behaviours it can be prompted into. Prompting asks it to select a behaviour; training gives it one. On a task narrow enough, a 0.5B or 1.5B model tuned on a few thousand examples can match a very large model prompted — for that task, and only that task.

That qualification is the whole deal. You are trading breadth for a specific competence at a hundredth of the cost.

What runs where

In 4-bit quantisation, approximately:

- **0.5B** — about 400 MB. Runs on a mid-range phone.
- **1.5B** — about 1 GB. Comfortable on a phone, instant on a laptop.
- **3B** — about 2 GB. Any modern laptop.
- **7B to 8B** — about 5 GB. A laptop with 8 GB of RAM, at a few tokens a second on CPU and much faster with any GPU or on Apple silicon.

For training: below about 1.5B, prefer **full fine-tuning** over LoRA. It fits on a free T4 with 8-bit AdamW, it is simpler, and at that size the extra capacity to change is usually what you want rather than something to be constrained.

The serving stack, all free

- **llama.cpp / GGUF** — CPU and GPU, every desktop platform, and the format most other tools read. Ollama and LM Studio wrap it.
- **ONNX Runtime** — broad hardware support, good on Windows and on CPU.
- **ExecuTorch** for iOS and Android, **MLC LLM** for mobile GPUs, **Core ML** on Apple devices.

The path from a training run to a phone is: train the adapter, merge it, convert to GGUF, quantise, test. All of it free and all of it documented.

The real reason, which is not cost

Small on-device models are usually justified on price. That is the weakest argument, because a hosted small model is already cheap.

In 4-bit, and where each one runs

0.5B — about 400 MB	Runs on a mid-range phone. Fine-tuning transforms it, because the base model was barely adequate.
1.5B — about 1 GB	Comfortable on a phone, instant on a laptop.
3B — about 2 GB	Any modern laptop.
8B — about 5 GB	A laptop with 8 GB of RAM, at a few tokens a second on CPU and much faster with any GPU.

Below about 1.5B, prefer full fine-tuning to LoRA: it fits a free T4 with 8-bit AdamW, and at that size the extra capacity to change is what you want rather than something to constrain.

The strong reasons are the ones an API cannot offer at any price:

- **The data does not leave the device.** For a health app, a journal, a legal notes tool, this is the product rather than a feature.
- **It works offline.** On a train, in a rural clinic, on a factory floor.
- **Latency has no network in it.** 60 milliseconds to first token instead of 400.
- **There is no per-request cost,** so you can call the model for things that would never justify an API call — ranking every item in a list, checking every keystroke.
- **No vendor can change or withdraw it.** The model you shipped is the model that runs, for as long as the app exists.

The limitations, plainly

Small models fail differently from large ones, and worse:

- **Off-distribution they collapse rather than degrade.** A large model handles an unexpected input awkwardly; a 1B model tuned on triage produces triage-shaped output for a recipe request, confidently.
- **They have very little world knowledge.** Anything factual must come from the context window. This is not a tuning problem and no dataset fixes it.
- **Multi-step reasoning is weak.** Chains of inference break down. Decompose the task into single steps, or do not use a small model for it.
- **Long context is expensive relative to their capacity** — a 2B model with a 32k window will attend over all of it without making much use of what it finds.

So the design pattern that works is a narrow, well-specified task with the necessary facts supplied in the prompt, plus a route to something larger when the input does not look like the task. The classifier from module one is a good gate for exactly that.

Where to start

The small-model families are excellent now and improving: Qwen at 0.5B and 1.5B, Llama at 1B and 3B, Gemma's small sizes, SmolLM at 135M to 1.7B, Phi's small variants. Licences differ — several are Apache-2.0, others carry community terms — so check the one you pick.

And test the base model on your task before training. The small models of the last two years are substantially better than their predecessors, and the gap you were going to close with a training run may have closed on its own.

BEFORE YOU MOVE ON

A 1B model fine-tuned on ticket triage is sent a recipe question by a user who misunderstood the product. It produces a confident triage classification. Why is this failure mode characteristic of small models rather than large ones?

- Small models have shorter context windows, so the recipe question was truncated into something resembling a ticket.
- Quantisation to 4 bits removes the precision needed to represent uncertainty, so all outputs become confident.
- It has a narrow range of behaviours, so training gave it one and it has nothing to fall back on.
- One billion parameters cannot store the refusal behaviour, which requires a minimum model size to be learned at all.

Tabular and time series: the honest answer

THE ONE THING TO KEEP

Gradient-boosted trees remain the default for predicting from tables, so language models belong around the problem — turning free-text columns into features, resolving entities — rather than in it, and any time-series foundation model has to beat seasonal naive on your own series before its benchmark numbers mean anything.

Gradient boosting still wins on tables

If your problem is predicting a number or a class from rows of structured columns — churn, credit risk, demand, conversion — the honest answer is that gradient-boosted trees remain the strongest default, and this has been tested repeatedly against deep learning on collections of real tabular datasets with consistent results.

The reasons are structural rather than incidental. Tabular features are heterogeneous and unordered, which suits axis-aligned splits and gives neural networks no useful inductive bias to exploit. Trees are robust to uninformative features, to outliers and to unscaled inputs. They train in seconds, need almost no tuning to be competitive, and tell you which features mattered.

XGBoost, LightGBM and CatBoost are all free. On a typical business table any of them will beat a fine-tuned language model on accuracy, by orders of magnitude on cost, and with an explanation of its predictions you can show somebody.

So: do not fine-tune a language model to predict a number from a table.

That sentence saves more time than most of this course.

Where a language model does earn its place

Three specific jobs, all around the model rather than in place of it:

Free-text columns. A complaint field, a doctor's note, a product description. Turn it into features — embeddings, or a small classifier's output — and feed those into the tree model alongside the numeric columns. This is a genuine and frequently large gain, and it is the embedding and classifier work from earlier in this module.

Entity resolution and cleaning. Matching "Acme Ltd." to "ACME Limited", normalising inconsistent categories, parsing addresses. Language models are good at this

and it is difficult to do any other way.

Writing the analysis. Generating the modelling code, explaining the result, drafting the report. Useful, and not a prediction task.

TabPFN, the interesting exception

One genuine exception is worth knowing: TabPFN is a transformer pretrained on synthetic tabular problems that performs classification by in-context learning — you pass the training rows in the prompt and it predicts, with no fitting step at all.

On small tables it is competitive with well-tuned boosted trees and takes seconds rather than a tuning session. Its constraint is size: it is designed for small datasets, up to roughly the low thousands of rows and a modest number of features, with later versions relaxing those limits. Above that, back to trees.

Time-series foundation models

A real development, and one worth being precise about. Chronos, TimesFM, Moirai and Lag-Llama are transformer models pretrained on very large collections of time series, and they forecast unseen series zero-shot, without being fitted to them.

They are genuinely useful in a specific situation: **many series, each too short to fit a model to individually**. Forecasting demand for 40,000 products where most have eighteen months of history is exactly the case, and it is a common one.

They can be fine-tuned on your own series, and this usually helps.

What to be careful about: the benchmark comparisons are contested, and the central concern is contamination. These models were pretrained on large public collections of time series, and the standard evaluation benchmarks are drawn from public time series. Whether a given result reflects forecasting ability or partial memorisation is a live methodological argument, and the field has not settled it.

The baseline that settles it

Whatever you use, benchmark against seasonal naive: predict that this week equals the same week last year, or last season.

```
yhat = y[-seasonal_period:]      # that is the whole model
```

It is one line, it costs nothing, and it beats a surprising share of deployed forecasting systems. If your foundation model cannot beat seasonal naive plus a simple statistical method such as exponential smoothing or ARIMA — both free, both in `statsmodels` — on *your* series, you have learned something important for the price of ten minutes.

That is the discipline this whole course keeps returning to. The impressive method has to beat the dull one on your data, and the dull one has to actually be run for the comparison to mean anything.

BEFORE YOU MOVE ON

A team wants to predict customer churn from 40 structured columns plus a free-text complaint field. What is the sound architecture?

- A. Serialise every row as a sentence and fine-tune a language model on the whole task, so the text and numbers are handled by one model.
- B. Use a time-series foundation model, since churn is a temporal prediction over each customer's history.
- C. Drop the free-text column and use gradient boosting on the 40 numeric columns, since trees cannot consume text.
- D. Turn the complaint text into embeddings, then train a gradient-boosted tree on those and the 40 columns.

CASE STUDY · MODULE 7

The vision model and the scanner

A chartered accountancy practice in Ahmedabad, processing about four thousand supplier invoices a month for forty small-business clients.

The practice had two articulated clerks keying invoice data into accounting software: supplier name, GSTIN, invoice number, date, taxable value, tax split, total. Most invoices arrived as PDFs by email, some as photographs of paper taken on a phone, and a few hundred a month as scans of handwritten delivery challans from clients in the textile trade.

The partner who had read about it, Nilesh, wanted to fine-tune a vision-language model. The appeal was obvious: one model, any image, out comes structured data. He had found a recipe, and a rented card for a day was affordable.

The engineer the practice retained, Priya, asked for an afternoon before anything was trained, and ran the comparison the module recommends. Fifty pages, drawn deliberately: thirty machine-printed PDF invoices, ten phone photographs of printed invoices, ten scans of handwritten challans.

Pipeline A was an OCR engine with layout output feeding a text model that had been given five examples in its prompt and a constrained decoder enforcing the schema. Nothing was trained.

Pipeline B was an off-the-shelf vision-language model, prompted, also with a constrained decoder. Nothing was trained.

On the thirty printed PDFs the OCR pipeline extracted all seven fields correctly on twenty-nine. The vision model managed twenty-five, and three of its four failures were the GSTIN — a fifteen-character alphanumeric string where a single wrong character makes the record useless, and where reading is exactly what an OCR specialist is built for.

On the ten phone photographs they were close: OCR eight, vision model eight, failing on different pages.

On the ten handwritten challans the OCR pipeline extracted almost nothing usable. The vision model got six.

Priya also measured something Nilesh had not asked about, and it changed the shape of the decision. She encoded one full-resolution invoice page and printed how many sequence positions it consumed in the vision model. At the resolution needed to read the tax table, the page occupied well over a thousand positions, and the model tiled larger

images, multiplying that. Activation memory scales with sequence length, so a training batch of four pages needed more memory than a batch of thirty-two text examples. Her effective batch would be one or two per device with heavy accumulation, and the run Nilesh had budgeted a day for was closer to three.

So the choice was not between two tools. It was between one architecture and two.

Fine-tuning the vision model on everything meant one artefact, one evaluation, one thing to maintain — and worse accuracy on the eighty-five per cent of the volume that was printed, at several times the cost per page, with failures that are hard to inspect because you cannot look at an intermediate output and see where the reading went wrong.

Running both pipelines meant two things to maintain, a router deciding which page goes where, and the clerks seeing two different failure modes. Against that it put a specialist on each job: the OCR engine on printed text, where it is simply better, and the vision model on handwriting, where layout and stroke carry meaning that survives no text extraction.

STOP HERE

Before turning the page: what would you have done, and what would it have cost if you were wrong? Write it down. The point of the next section is to compare.

WHAT HAPPENED

The practice built the router and ran both. The router is not a model: challans arrive from a known set of clients in a known format, so a rule on the sender and the document dimensions sends the right pages to the right pipeline, with a fallback to the vision model when the OCR output has fewer characters than a printed invoice could plausibly contain.

Printed invoices and photographs — about 3,700 a month — go through OCR into a prompted text model with constrained decoding. Nothing in that path was trained at all.

The handwritten challans go to a vision-language model that was fine-tuned, but not the way Nilesh first proposed. The vision encoder was frozen, the projector trained, and the language model adapted with a LoRA, on 900 challan-and-answer pairs the clerks produced over three weeks from their own keying work. Input resolution was set deliberately at the point where the handwritten quantity column was still legible and no

higher, because lowering it was the most effective memory lever available and it cost only the detail that had been lowered.

The number the partners care about is that the two clerks now key exceptions rather than invoices, and the exception rate is eleven per cent. The number Priya cares about is that when a field is wrong on a printed invoice, she can open the OCR output and see whether the reading failed or the extraction did — which on the vision path she cannot.

Priya ran both pipelines on fifty pages before anything was trained. What made that afternoon decisive rather than merely informative?

It settled an architecture, not a hyperparameter. The comparison showed that the two tools failed on different populations — OCR was better on the 85 per cent of volume that was printed, and unusable on handwriting where the vision model succeeded — which meant the right answer was two pipelines and a router rather than a choice between them. No amount of training either model would have produced that finding, because it is a fact about which specialist suits which input, and it cost an afternoon against a three-day training run aimed at the wrong target.

Why did the GSTIN failures matter more than the raw accuracy gap suggests?

Because a GSTIN is a fifteen-character alphanumeric identifier where one wrong character makes the whole record useless and the error is not self-evident — unlike a slightly mis-read supplier name, which a clerk notices. Reading exact strings is precisely what an OCR engine is specialised for and precisely where a model that has compressed an image into embeddings is weak. Field-level accuracy averaged across seven fields hides the fact that the failures concentrated in the one field whose failure is most expensive.

In the vision path, the encoder was frozen and the input resolution capped. Explain what each choice bought and what it cost.

Freezing the vision encoder keeps the part trained on far more images than the practice has, which reduces memory and reduces the damage a small dataset can do; the cost is that the model cannot adapt how it sees, only how it describes what it sees — acceptable because the challans are ordinary photographs, not a new imaging modality. Capping the resolution reduces the number of sequence positions an image occupies, which is the most effective memory lever in vision-language training since activation memory scales with sequence length; the cost is exactly the detail removed, so the cap was set at the point where the handwritten quantity column was still legible.

Check what stuck

6 questions, one right answer each. This one is for you, not for a record: answers and the reason behind each are at the back. If two questions from the same module go wrong, re-read that module before the exam rather than after.

- 1 Why is batch size the important hyperparameter when training an embedding model with multiple-negatives ranking loss?
 - A. Because the other positives in the batch are the negatives, so the batch size is the difficulty of the task.
 - B. Because the loss is averaged over the batch, so a larger batch produces a lower number.
 - C. Because gradient caching only works above a certain batch size, below which the memory saving is outweighed by the cost of the second forward pass.
 - D. Because embedding models need large batches to avoid overfitting on short texts.
- 2 Why mine hard negatives from ranks ten to fifty rather than from the very top of the retrieved list?
 - A. Because the retriever's top ranks are cached, so mining from them returns stale results that no longer reflect how the current model actually scores.
 - B. Because the top results are too similar for the loss to produce a usable gradient.
 - C. Because the top results are often also correct, and training on them teaches the model to push apart things that belong together.
 - D. Because mining from the top biases the model towards short passages.
- 3 Why does a cross-encoder rerank fifty candidates rather than search the corpus itself?
 - A. Because its scores are not comparable across different queries.
 - B. Because it must run once per candidate, so scoring a whole corpus takes minutes rather than milliseconds.
 - C. Because it can only be trained on candidate lists of about fifty items.
 - D. Because its output is a probability rather than a distance, and an index can only be built from a distance, which is the reason the two stages exist at all.

- 4 A six-way routing model's data is 62 per cent one class. Why is accuracy the wrong headline metric?
- A. Because accuracy requires calibrated probabilities, which a fine-tuned encoder never has.
 - B. Because accuracy weights each class by its frequency while macro F1 weights each class by the cost of getting it wrong, which is what you actually care about.
 - C. Because accuracy cannot be computed for more than two classes.
 - D. Because a model that always predicts the majority class scores 62 per cent while routing nothing.
- 5 Constrained decoding guarantees that your output parses. What does it not guarantee?
- A. Truth: a valid object can still route a shipping complaint to billing.
 - B. That the grammar and the model's chat template agree, which is why the two have to be compiled together before any request can be served.
 - C. Speed, so a fine-tune is still needed wherever latency matters.
 - D. That every optional field will be filled, which makes a schema with optional fields unsafe.
- 6 You must predict churn from a table of customer columns. Where does a language model belong?
- A. As the predictor, fine-tuned on rows rendered as text.
 - B. Nowhere near the problem; a table should be modelled by a neural network of some other kind.
 - C. Around the problem: turning free-text columns into features for a gradient-boosted tree.
 - D. As a post-processor that reads the tree's predictions and adjusts them using the free-text columns that the tree was unable to represent numerically.

MODULE 8

Did it actually work

A fine-tune that improves your target metric and quietly breaks three other things is a loss. This module measures both halves, with intervals, on a set the model has not seen.

BY THE END OF THIS MODULE YOU CAN

Compare a tuned model against a cost-matched baseline on a leak-free held-out set, report the difference with an interval, and demonstrate whether general capability, safety behaviour or calibration regressed

67 • 10 min

Did It Help, and What Did It Break

THE ONE THING TO KEEP

A win rate without an interval and a canary set is not evidence, it is a feeling.

The comparison has to be fair

Two mistakes make most fine-tuning results meaningless.

Comparing against a weak baseline. You spent three weeks on the fine-tune and ten minutes on the prompt you are comparing it against. Run the base model with your *best* prompt — few-shot examples included, retrieval included — on the same held-out inputs. That is the bar. Beating a strawman teaches you nothing.

Judging your own work unblinded. Generate both outputs, shuffle which is A and which is B, hide which is which, then judge. If a model does the judging, run each pair twice with the order swapped and discard the pairs where the verdict flips. That flip rate is your judge's error bar, and it is often 10-20%.

Numbers, with intervals

Head-to-head win rate is the most useful single measure for open-ended tasks. Report it with an interval, because the interval is wider than people expect. For n comparisons, the 95% margin near a 50% rate is roughly $1/\sqrt{n}$:

- $n = 50 \rightarrow$ about ± 14 points
- $n = 100 \rightarrow$ about ± 10 points
- $n = 200 \rightarrow$ about ± 7 points
- $n = 500 \rightarrow$ about ± 4 points

So 58 wins out of 100 is not a result. It is consistent with no difference at all. If your task has an exact answer — a label, an extracted field, valid JSON — use accuracy or exact match instead. Those need fewer samples and start fewer arguments.

Put cost and latency in the same table. A tuned 7B that ties a large hosted model has won, because it costs a fraction to run. A tuned 7B that ties the base 7B has lost, and it is better to know.

Loss is not your metric

Evaluation loss measures how well the model reproduces text that looks like your training data. It falls while quality falls, routinely. Use it to spot a broken run, not to choose a checkpoint.

Catastrophic forgetting

Training pushes weights toward your data, so behaviours you did not train on drift. The narrower your data and the longer you train, the further they drift. LoRA reduces this, because the base weights are untouched, but it does not remove it: the adapter applies to every input, including inputs nothing like your training set.

What it looks like in practice:

- Every answer arrives in your training format, including for questions with nothing to do with your task.
- Answers get shorter and blunter across the board, because your dataset was terse.
- The model refuses, or answers strangely, outside its new domain.
- Multi-turn conversation falls apart when you trained only on single turns.
- **Other languages degrade first.** This is the one teams miss. If the base handled Hindi, Spanish or Arabic and you trained only in English, non-English quality can drop sharply while your English metrics look untouched.

- Code formatting breaks, or your schema starts appearing inside code blocks.

The canary set

Build fifty prompts with nothing to do with your task. General questions, a short reasoning problem, a code snippet, a multi-turn exchange, and — if you serve them — the same questions in each language you support. Run them on the base model and save the outputs. Run them on every fine-tuned checkpoint. Read them side by side yourself. This is not a metric, it is a look.

Ten minutes per checkpoint, and it catches the failure that no held-out set built from your own task will ever show you.

When you find forgetting

In roughly this order:

1. Fewer epochs. Take the epoch-one checkpoint.
2. Lower the learning rate: $2e-4$, then $1e-4$, then $5e-5$.
3. Lower the rank.
4. Mix 5-20% general instruction data into your training set.
5. Narrow `target_modules` back toward attention only.

Keep the adapter unmerged while you test, so turning the fine-tune off is one flag rather than a redeploy.

The decision

Ship if the tuned model beats the best-prompt baseline by more than the interval, the canary set is intact, and cost or latency improved or held. Otherwise you still have a result — one that saved you from serving something worse.

BEFORE YOU MOVE ON

Your fine-tuned model wins 58 of 100 blind head-to-head comparisons against the base model on your held-out set. The team wants to ship. What is the honest reading?

- A. A 58% win rate is a clear majority and the direction is right, so ship it.
 - B. It failed to clear 60%, so the fine-tune should be abandoned.
 - C. 58 of 100 sits inside a noise band of roughly ten points, so nothing is established yet: you need more comparisons, the strongest-prompt baseline, and a canary check for what broke.
 - D. Re-judge the same 100 pairs with a different judge model to confirm the result.
-

68 · 9 min

The baseline you must beat

THE ONE THING TO KEEP

Compare against three baselines — the frozen best prompt, the cost-matched alternative, and production today — on one table of quality and cost per thousand requests, because the alternative nobody was ever going to choose is not a baseline and a one-point difference on 200 items is not a difference.

Three baselines, not one

A fine-tune's result means nothing without something to compare it against, and the comparison people run is almost always too flattering.

1. The frozen prompt baseline. The same base model, with the best prompt you can write, pinned at a specific commit with specific sampling parameters. This tells you whether the training added anything over the fast loop.

2. The cost-matched baseline. This is the one that gets skipped, and it is usually the one that decides the project. If you tuned an 8B model to save money over a hosted frontier model, the honest comparison is not *tuned 8B versus frontier*. It is *tuned 8B versus the smallest untuned model that a good prompt makes adequate, at the same total cost per thousand requests*.

Because that is the real alternative. Nobody was going to keep paying frontier prices; they were going to try a cheaper model with a better prompt.

3. The do-nothing baseline. Whatever is in production today, measured on the same set. Sometimes it is better than both, and finding that out at the end of a project is common and avoidable.

Writing the prompt baseline honestly

The temptation is to compare against the prompt you had when you decided to fine-tune — the one whose failures motivated the project. That is a straw man, and everybody builds it without meaning to.

Spend one full day on the prompt first, as though fine-tuning were unavailable:

- Five to eight few-shot examples covering the cases you actually fail on, chosen from your training data.
- An explicit output format, with constrained decoding enforcing it.
- The instruction rewritten twice, with the alternatives measured rather than argued about.
- The retrieval, if there is any, fixed first.

Then freeze it. A meaningful fraction of fine-tuning projects end here, and ending here is a good outcome — a day spent instead of three weeks, and a system anybody on the team can modify.

The comparison table, as numbers

Put every candidate on the same two axes: quality on your held-out set, and cost per thousand requests including serving.

system	acc	p95	\$/1k	notes
production today	0.84	1.4s	4.20	frontier, old
prompt				
frontier + new prompt	0.91	1.5s	4.20	one day of work
mid-tier + new prompt	0.88	0.9s	0.70	one day of work
tuned 8B (LoRA)	0.92	0.4s	0.31	3 weeks + maintenance
tuned 8B, merged & 4-bit	0.91	0.3s	0.22	quantisation cost 1 pt

Read that table honestly. The fine-tune is one point better than a day's prompt work on the frontier model, at a fourteenth of the cost and a third of the latency. Whether that is worth three weeks and a maintenance commitment depends entirely on your volume, and now you can compute it rather than argue about it.

If the row for the mid-tier model with a good prompt had read 0.92, the project would have been unnecessary, and the only way to know is to have run it.

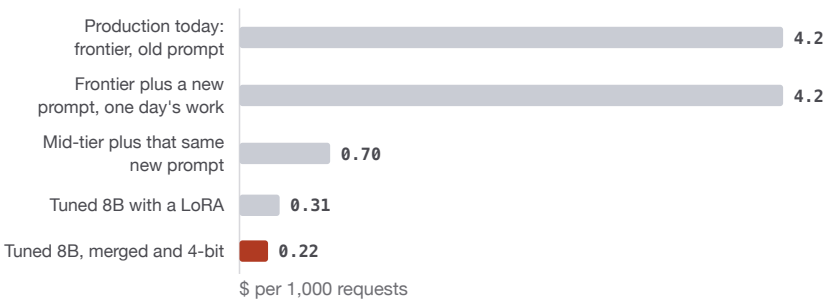
Base or instruct: also a baseline

A cheap comparison people skip: the same fine-tune from the base model and from the instruct-tuned variant of the same weights.

Instruct usually wins for instruction-shaped tasks, needing less data because it already follows instructions. Base sometimes wins when you want no inherited behaviour — a very rigid output format where the instruct model's conversational habits are an obstacle.

It is two runs. Run both once on a new project and you will know which your data prefers.

Cost per thousand requests, same held-out set



Accuracy on those five rows, in order: 0.84, 0.91, 0.88, 0.92, 0.91. The fine-tune is one point better than a day's prompt work on the frontier model, at a fourteenth of the cost — and if the mid-tier row had read 0.92, the project was unnecessary.

The intervals

Every number in that table needs an interval, or the comparison is a vibe. For a proportion on n items, a serviceable approximation is $\pm 1.96 \cdot \sqrt{p(1-p)/n}$:

- $n = 200, p \approx 0.9 \rightarrow$ about ± 4 points
- $n = 500 \rightarrow$ about ± 2.6 points
- $n = 1,000 \rightarrow$ about ± 1.9 points

So on a 200-item set, 0.92 against 0.91 is not a difference. You would need roughly ten times the items to resolve one point, and it is usually cheaper to decide that one point does not matter than to build a 2,000-item labelled set to measure it.

And add the seed variance from the reproducibility lesson on top: your training noise floor and your evaluation interval are two separate uncertainties, and the honest comparison carries both.

BEFORE YOU MOVE ON

A team tunes an 8B model to replace a frontier API and reports it matching frontier quality at a twentieth of the cost. What is the missing comparison that could invalidate the project?

- A. A mid-tier untuned model with an equally good prompt, which may already match it with no training.
 - B. The frontier model with the same fine-tuning applied, which would show whether the gains came from the data or from the base model.
 - C. The 8B model before tuning, to establish that the training changed anything at all.
 - D. The frontier model's latency at the same batch size, since cost comparisons must hold throughput constant.
-

69 · 9 min

A held-out set that actually holds

THE ONE THING TO KEEP

Random splits leak through near-duplicates, groups, time and templates, so split by group and by time and then measure cross-split near-duplicate overlap — and keep a locked set you open a handful of times, because an evaluation set you have consulted forty times has become a training set for your own decisions.

Random splitting is not enough

A random 90/10 split feels safe and leaks constantly. Four distinct leaks, each with a fix.

Near-duplicate leakage. Your corpus contains variations of the same ticket. A random split puts some in training and their near-twins in evaluation, and the model scores well by recall rather than by generalisation.

This is the big one, it is invisible, and almost nobody measures it. Run the same MinHash comparison you used for cleaning, but *across the split*:

```
lsh = MinHashLSH(threshold=0.8, num_perm=128)
for i, t in enumerate(train_texts):
    lsh.insert(f"tr{i}", minhash(t))
leaked = sum(1 for t in eval_texts if lsh.query(minhash(t)))
print(f"{leaked}/{len(eval_texts)} eval items have a near-duplicate in
train")
```

On real assembled datasets this frequently comes back at 5 to 15%. Every one of those items is inflating your score.

Group leakage. One customer generates forty tickets, split across both sides. The model learns that customer's vocabulary, their product, their way of describing things, and scores well on their evaluation items for reasons that do not generalise.

Split by group, not by row — by customer, document, session, author, or whatever unit actually repeats.

```
from sklearn.model_selection import GroupShuffleSplit
train_idx, eval_idx = next(GroupShuffleSplit(test_size=0.1, random_s-
tate=0)
                           .split(X, groups=df["customer_id"]))
```

Temporal leakage. Production always runs on the future. A random split lets the model train on September and test on August, which is a cheat — the pricing change, the product launch, the new phrasing are all in the training data.

Split by time: train on everything before a cutoff, evaluate on everything after. The number will be lower, and it is the number that predicts production.

Template leakage. If your data was generated from templates, a random split puts the same template on both sides with different slot values. You are measuring slot-filling, not the task.

Use it sparingly

An evaluation set decays with use. Every time you look at it and change something in response, you fit a little to it — not through gradient descent but through your own decisions, which is slower and just as real. After forty runs, your held-out set is partly a training set.

So keep two:

- **A development set**, used freely during iteration.

- **A locked set**, opened a handful of times: before the project starts, at one or two milestones, and before shipping. Written down each time, so the count is visible.

If the locked set and the development set have drifted apart, the gap is how much you overfitted your own process.

Four ways a random split leaks

Near-duplicate leakage	Variations of the same ticket land on both sides, and the model scores well by recall rather than by generalising. Invisible, and almost nobody measures it.
Group leakage	One customer's forty tickets split across both sides. Split by customer, document, session or author — whatever unit actually repeats.
Temporal leakage	Training on September and testing on August. Production always runs on the future, so split by time and accept the lower number.
Template leakage	The same template on both sides with different slot values. You have measured slot-filling, not the task.

Measure the overlap rather than assuming it: run the same MinHash comparison you used for cleaning, across the split. On real assembled datasets it comes back at 5 to 15 per cent.

How large

The interval on a proportion tells you:

- 100 items → ± 6 to 10 points. Enough to spot a disaster, not enough to compare two candidates.
- 300 items → ± 3 to 5 points.
- 1,000 items → ± 2 to 3 points.

Most teams want 300 to 500 for the development set and 300 for the locked one. Below 200 you cannot distinguish anything you would plausibly be arguing about.

The cost is in labelling, so the practical move is to stratify: 50 items per important category rather than 500 drawn at random, with the caveat from the sampling lesson that you must then report per stratum or weight back to real proportions.

The items to include on purpose

A set drawn only from typical traffic misses everything you most need to know:

- **Cases that should be declined.** Out of scope, missing information, ambiguous. If the model cannot decline, you want that in a number.

- **The rare-but-expensive categories.** Refunds, safety-relevant requests, anything where a mistake is costly.
- **Inputs from the tail.** Very short, very long, mixed-language, badly formatted.
- **Adversarial items.** Prompt injection attempts if your inputs come from users, and the ordinary rudeness your support queue actually receives.
- **Items the previous version got wrong.** Every production incident becomes a permanent evaluation item. This is the habit that compounds; after a year your set encodes everything the system has ever failed at, and no regression can pass unnoticed.

One rule

Build it before the training set. Not because of some discipline, but because the act of deciding what counts as a correct answer is the act of specifying the task — and doing that after you have seen the model's output means specifying it to fit what the model produced.

BEFORE YOU MOVE ON

A team splits 10,000 support tickets randomly 90/10. Held-out accuracy is 0.94; production accuracy is 0.79. A check shows 12% of eval items have a near-duplicate in training and several customers appear on both sides. What is the mechanism?

- The evaluation set is too small at 1,000 items to estimate accuracy within 15 points.
- The evaluation measured recall of seen text and customer vocabulary, neither available on new tickets.
- Production traffic has drifted since the data was collected, which a larger evaluation set would have revealed.
- Randomly split evaluation sets always overstate performance because they share the training set's label distribution.

70 · 9 min

Pairwise human comparison, done properly

THE ONE THING TO KEEP

Blind, order-randomised, tie-allowing comparisons with a stated criterion and a measured judge-agreement rate — with the number of comparisons fixed in advance from the effect size you would act on, and mean response length reported beside the win rate because length buys wins it did not earn.

Why pairwise

Asking somebody to rate an answer out of ten produces noise: different people use different parts of the scale, the same person drifts over a session, and nobody can hold a stable standard across two hundred items.

Asking which of two answers is better produces a judgement people can actually make. It is the right instrument for everything you cannot measure automatically — tone, helpfulness, whether the length suits the question, whether a decline was appropriate.

The protocol

Six requirements, each preventing a specific error:

1. **Blind.** The judge must not know which system produced which answer. Obvious, and violated constantly by teams evaluating their own work.
2. **Randomise the order** per item, and record which side was which. Position bias is real and substantial in both human and model judges.
3. **Allow a tie.** Forcing a choice on genuinely equivalent answers manufactures signal from nothing. A high tie rate is itself a finding: it means your two systems are not meaningfully different.
4. **State the criterion.** "Which is better" gets you the judge's private theory of better. "Which would you rather receive as a customer, considering correctness first and tone second" gets you an answer to your question.
5. **Same inputs, same conditions.** Both systems get identical prompts, identical retrieved context, identical sampling settings.
6. **Overlap for agreement.** Have two or three judges do the same 50 items. Their agreement rate is the ceiling on what this measurement can tell you — if two careful

people agree only 60% of the time, no sample size rescues the comparison, and the task needs specifying better before it needs measuring.

How many comparisons

The count depends entirely on the size of the difference you would act on:

- **A 65% win rate** (a large, obvious difference): about 150 comparisons gives an interval that excludes 50%.
- **A 60% win rate**: about 380.
- **A 55% win rate** (a small difference): about 1,500.

That progression is worth internalising. Detecting a small difference is expensive, and the correct response is usually not to buy more comparisons but to decide that a difference that small does not change anything.

Decide the number *before* you start. Stopping when the result looks good is the most reliable way to produce a result that is not true.

Reporting it

```
n = 400 comparisons, 3 judges, pairwise agreement 0.78
tuned wins 182 · ties 96 · baseline wins 122
win rate excluding ties: 182/304 = 0.599 [95% CI 0.543 – 0.653]
mean length: tuned 148 tokens, baseline 96 tokens
```

That last line is the one to always include. Longer answers win more often than their quality warrants, in human judging as much as in model judging. If your tuned model's answers are 50% longer, a meaningful part of the 60% win rate is length rather than quality.

The check: compute the win rate separately on items where the two answers were of similar length. If it collapses towards 50%, you measured verbosity.

Making it cheap

The tooling does not need to be sophisticated. A static HTML page that shows a prompt and two answers, records a keypress and appends a row to a CSV is fifty lines and works. A spreadsheet with randomised columns works. Argilla is free and open source if you want annotation management, agreement tracking and a proper interface.

The expensive part is attention, not software. Thirty to ninety seconds per comparison means 400 comparisons is six to ten hours of somebody's careful work, and their care degrades after about an hour. Spread it, and re-check agreement on a few items from the start of the session against the end.

What it cannot tell you

A win rate says which answer people prefer when reading two side by side. That is not the same as which answer serves them better in use. People prefer confident answers to hedged ones, complete-looking ones to short ones, and agreement to correction — and in every one of those cases the preferred answer is sometimes the worse one.

So pair it with something objective wherever you can: a correctness check, a verifier, a task-completion rate from real use. Preference is one instrument, it measures one thing, and a project that measures only preference will optimise for being liked.

BEFORE YOU MOVE ON

A tuned model wins 60% of 400 blind pairwise comparisons. Its answers average 148 tokens against the baseline's 96. What should be checked before accepting the result?

- A. Whether the judges were consistent, by recomputing the win rate per judge and discarding outliers.
- B. Whether 400 comparisons is enough, since a 60% win rate needs about 1,500 to exclude 50%.
- C. Whether the win rate holds on length-matched items, since length buys preference on its own.
- D. Whether the tie rate was high, since ties reduce the effective sample and widen the interval.

71 · 9 min

Judging a model you distilled from the judge

THE ONE THING TO KEEP

A teacher judging its own student rewards resemblance rather than quality, so use a verifier, a human sample or a judge from a different family — and a student that beats its teacher on teacher-judged comparison is a finding about the measurement, not the model.

The circularity

You distilled a 7B student from a large teacher. To evaluate the student cheaply, you ask that same teacher which of two answers is better.

The teacher prefers the student.

Of course it does. The student was trained to produce text that resembles the teacher's, and the teacher's judgement of quality is substantially a judgement of resemblance to its own output — its phrasing, its structure, its characteristic hedges, its idea of an appropriate length. You have built a measurement that rewards the thing you trained for, independently of whether that thing is good.

This is self-preference bias, it has been measured in several studies of model judges, and it is at its strongest in exactly this configuration.

What to do instead

In descending order of trustworthiness:

1. **A verifier**, where one exists. Unarguable and free.
2. **Human pairwise comparison** on a sample. Expensive, and the standard the others are calibrated against.
3. **A judge from a different model family** than the teacher. Not perfect — it has its own preferences — but it is not the model whose output you trained to imitate.
4. **A panel** of two or three different judges, with disagreement flagged for human review. More expensive, and it makes the bias visible rather than removing it.

Using the teacher as judge is not always wrong; it is wrong to do it without saying so, and without an agreement figure against human labels on a sample of the same items.

Establish agreement before you believe any of it

Whatever judge you use, the number that licenses trusting it is its agreement with human judgement on your own task:

200 items judged by both a person and the model judge
agreement 0.83 (people agree with each other 0.87 on the same items)

That second figure is the ceiling. A judge cannot be more reliable than the task's own definiteness. A judge at 0.83 against a human ceiling of 0.87 is doing well and its numbers are usable with a stated caveat. A judge at 0.62 against a ceiling of 0.85 is not measuring your task.

Re-check this when the model changes. A judge calibrated against your pre-tuning outputs may behave differently on post-tuning ones, precisely because those outputs have changed in the direction the judge rewards.

The biases, and the ones that can be reduced

Position bias. Judges favour the first option, sometimes strongly. The fix is complete: evaluate every pair twice with the order swapped and count a win only if both orders agree. Disagreement becomes a tie. It doubles the cost and it removes the bias rather than mitigating it, which is a good trade.

Verbosity bias. Judges prefer longer answers. Report mean length alongside the win rate, always, and check the win rate on length-matched items.

Style over substance. Confident, well-formatted, fluent answers beat correct hedged ones. A rubric that asks about correctness explicitly, as a separate judgement from quality, reduces this. Asking for the reasoning before the verdict generally improves judgement quality, at the cost of more output tokens.

The rubric matters more than the model. A vague "which is better" gets a vague answer. A rubric that names three criteria and asks for a verdict on each is a different and much more reliable instrument, and the improvement from writing one usually exceeds the improvement from using a larger judge.

The trap specific to distillation

One more, and it is subtle. A distilled student can *beat its teacher* on a teacher-judged comparison while being worse for users.

The mechanism: the student learned the teacher's surface style strongly and its substance partially. On a judgement that weighs style heavily — and most do — the student's more consistent, more formulaic output can score above the teacher's more varied one. The student has become a caricature of the teacher, and caricatures are recognisable.

When you see a student beating its teacher, treat it as a finding about your measurement rather than about your model, until a human comparison says otherwise.

The rule

Write down, in the evaluation report, what judged the model and what established that the judge was trustworthy. If the answer to the first is "the model we distilled from" and the answer to the second is nothing, the result is not evidence. It is the training objective, measured again.

BEFORE YOU MOVE ON

A 7B student distilled from a large teacher scores a 58% win rate against that teacher, judged by the teacher itself. What is the most likely explanation?

- A. Distillation compresses the teacher's knowledge into a form that generalises better, so exceeding the teacher is a known effect of the method.
- B. The student was trained on filtered teacher outputs, so it inherited only the teacher's best behaviour and is genuinely better.
- C. Position bias favoured the student, which was presented first in most comparisons.
- D. The judge rewards resemblance to its own style, which is what the student was trained to maximise.

72 · 9 min

Contamination you can check, and contamination you cannot

THE ONE THING TO KEEP

Thirteen-gram overlap will find contamination between your training and evaluation sets, and the answer-ordering test can detect memorised benchmarks — but the base model's pretraining corpus is unchecked and unknowable, which is why a held-out set built after the training cutoff is the only evidence that holds.

Two kinds, and only one is yours to fix

The kind you can check: your training data contains items from the evaluation you are quoting. This happens more easily than people expect, particularly when your training data is synthetic — a teacher model that memorised a public benchmark during pretraining will happily reproduce its items when asked to generate examples of that kind of task.

The kind you cannot check: the base model's pretraining corpus contained the benchmark. You have no access to that corpus, no way to search it, and no way to rule it out.

The second is why public benchmark numbers on a fine-tuned model are weak evidence about anything, and why your own held-out set — built from your own data, after your own cutoff — is the strong evidence.

Checking the kind you can

The conventional test is n-gram overlap, and the convention from the large pretraining papers is 13-grams: if a 13-word sequence from an evaluation item appears anywhere in the training data, treat the item as contaminated.

```
def ngrams(text, n=13):
    w = text.lower().split()
    return {" ".join(w[i:i+n]) for i in range(max(0, len(w)-n+1))}

train_grams = set().union(*(ngrams(t) for t in train_texts))
hits = [i for i, t in enumerate(eval_texts) if ngrams(t) & train_grams]
print(f"{len(hits)}/{len(eval_texts)} eval items share a 13-gram with training")
```

Thirteen is a reasonable default: long enough that ordinary phrases do not match, short enough to catch a paraphrase that kept one clause. Run it in both directions — evaluation against training, and training against any public benchmark you intend to quote.

Some benchmarks embed a canary string specifically so that their presence in a corpus can be detected. Grep for it. If it is in your training data, something upstream scraped the benchmark.

The ordering test

A cleverer check for multiple-choice benchmarks, and it needs no access to anyone's corpus.

If a model has memorised a question as it appears in the benchmark, it will assign higher probability to the *canonical ordering* of its answer options than to a shuffled ordering of the same options. A model that is reasoning about the question has no reason to prefer one arrangement.

```
p_original = seq_logprob(model, render(q, q.options))
p_shuffled = mean(seq_logprob(model, render(q, shuffle(q.options)))
                  for _ in range(5))
```

A consistent preference for the original ordering across many questions is evidence the model saw the benchmark in that form. This is a real and usable diagnostic, and it is one of the few tools available for the contamination you cannot otherwise inspect.

What contamination does to your decisions

It inflates the number you are using to make a decision, and it inflates it unevenly. The items that leaked are the ones the model scores well on, so the headline rises while the model's actual generalisation has not moved. You then ship, and production disagrees.

The specific damage in a fine-tuning project: you compare a tuned model against a baseline, and the tuned model's *training data* overlaps the evaluation while the base-

line's does not. The comparison is structurally unfair in the tuned model's favour, and the gap you measure is partly recall.

The defences

1. **Build the evaluation set from data the training set cannot contain** — a later time period, a held-out group, or items written specifically for evaluation and never released into the training pipeline.
2. **Run the n-gram check** across your split, and report the number in your evaluation write-up even when it is zero. A stated zero is information; an unstated one is an assumption.
3. **Do not quote public benchmarks** as evidence that your fine-tune worked. They cannot be cleaned of base-model contamination, and they are not measuring your task anyway.
4. **Keep a locked set** created after the training data's cutoff date. Temporal separation is the one defence that also protects against the contamination you cannot inspect, because a benchmark that did not exist when the base model was trained cannot be in its corpus.

The honest summary

You can rule out contamination between your training set and your evaluation set. You cannot rule it out between the base model's pretraining and any public benchmark. Say which of the two you have checked, and do not let a reader assume you checked both. That sentence, written into an evaluation report, is worth more than an additional decimal place on any of its numbers.

BEFORE YOU MOVE ON

A tuned model scores markedly higher on a public multiple-choice benchmark when the answer options are in their canonical order than when the same options are shuffled. What does this indicate?

- A. The model has seen those questions in that exact arrangement, so its score reflects recall.
- B. The fine-tuning data taught a preference for the benchmark's formatting conventions, which shuffling disrupts.
- C. Position bias in the model, which prefers the first listed option regardless of content.
- D. The shuffled version is genuinely harder, because reordering breaks the logical progression the question author intended.

Safety regression, including from benign data

THE ONE THING TO KEEP

Fine-tuning on entirely benign data measurably degrades refusal behaviour, because alignment is a thin late layer and every training example demonstrates compliance — so include refusal examples, measure with your own refusal set at each checkpoint, and never let the model's own alignment be the only safety control.

The finding

Fine-tuning an aligned model degrades its safety behaviour — and this happens with training data that contains nothing harmful at all. Multiple independent studies have reported that a small number of ordinary, benign examples measurably increases a model's compliance with requests it previously refused, and that a handful of deliberately adversarial examples can remove the refusal behaviour almost entirely.

This is not an exotic risk requiring a hostile actor. It is the default consequence of ordinary fine-tuning on ordinary data, and most teams have never measured it because they did not know to look.

Why it happens

Two mechanisms, both straightforward.

Alignment is a thin, late layer. Refusal behaviour was installed by the last stage of the model's training, on comparatively little data relative to pretraining. It is a shallow modification of a model that otherwise knows how to answer, and any gradient step moving weights towards a new objective pulls against it.

Your data demonstrates compliance. Every example in your training set shows a request being answered. Nothing in it shows a request being declined. You are teaching, over and over, that the assistant's job is to produce an answer — and the model generalises that beyond the requests you had in mind.

The second mechanism is the one you control, and it is the one with a cheap fix.

Measuring it

Build a refusal set: 100 to 200 prompts your base model declines, of kinds relevant to your product. For a support assistant that is requests for another customer's data, for internal pricing, for medical or legal advice, and whatever category your organisation actually cares about.

```
for name, m in [("base", base), ("tuned", tuned)]:
    outs = [generate(m, p) for p in refusal_prompts]
    refused = sum(is_refusal(o) for o in outs)
    print(f"{name}: refused {refused}/{len(refusal_prompts)}")
```

Run it at every checkpoint alongside your task metric. A typical result on a project that did nothing about this:

base	refused 186/200
tuned e1	refused 151/200
tuned e2	refused 118/200

That is a large regression, it arrived from ordinary data, and it is invisible in every other number the project produced.

Writing `is_refusal` is the fiddly part: a keyword list gets you most of the way and misclassifies the interesting cases, so read a sample by hand. A small classifier trained on a few hundred labelled outputs does it better, and is exactly the encoder task from module one.

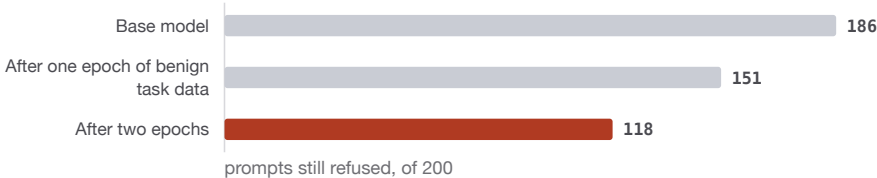
What to do about it

Put refusals in your training data. Thirty to a hundred examples of the model declining appropriately, in your product's voice, with your product's wording. This is the single most effective mitigation and it costs an afternoon. It also gives the model *a way* to decline that suits your product, rather than falling back on a generic refusal that reads like a different application.

Keep the replay mixture. General instruction data includes refusals and hedges, which keeps the relevant weights receiving a signal.

Fewer epochs, lower learning rate, and LoRA rather than full fine-tuning. All reduce the magnitude of the change, and therefore the damage. LoRA's low-rank constraint measurably preserves more of the base behaviour than full fine-tuning does.

Refusals held, on a 200-prompt refusal set



Nothing in that training data was harmful. The regression arrived from ordinary examples that all demonstrate compliance, and it is invisible in every other number the project produced.

Check the checkpoint you ship, not the method. The regression varies between runs and between checkpoints. Measure the artefact.

The architectural conclusion

Here is the part that outlives any particular mitigation.

Do not make the model's own alignment your only safety control. Fine-tuning can weaken it, quantisation can shift it, a prompt can talk around it, and none of those events announce themselves.

Whatever your product requires, enforce it in a layer that does not depend on the model's cooperation: an input classifier, an output filter, a tool-call allowlist, a permission check before data is retrieved. Those controls are inspectable, testable, and unaffected by what a gradient step did to a weight matrix.

The model's alignment is a useful default. It is not a control.

One caveat on the evidence

The published results differ in how much degradation they report, because they use different base models, different data and different definitions of a refusal. The direction of the effect is consistently reproduced; the magnitude is not a constant you can look up. That is precisely why you measure it on your own model with your own refusal set rather than citing a number from a paper.

BEFORE YOU MOVE ON

A team fine-tunes on 4,000 benign support conversations. Refusal rate on a held-out set of out-of-scope requests falls from 93% to 59%. What is the primary mechanism?

- A. Support conversations contain borderline requests that were answered, which taught the model those categories are acceptable.
- B. Every example showed a request being answered and none showed one declined, so it generalised.
- C. The learning rate was too high, so the alignment weights were overwritten faster than the task was learned.
- D. Refusal behaviour depends on the system prompt, which the fine-tuning data omitted, so the model lost the instruction that produced it.

74 · 8 min

Calibration after tuning

THE ONE THING TO KEEP

Fine-tuning sharpens the output distribution and preference tuning rewards confidence, so a tuned model's 0.9 no longer means 0.9 — which breaks every threshold, fallback and abstention rule downstream until temperature scaling is refitted on the exact artefact you serve.

Confidence stops meaning what it meant

A well-calibrated model's confidence tracks its accuracy: among the things it says with 80% confidence, about 80% are right.

Fine-tuning degrades this, reliably and in one direction. Training sharpens the output distribution — that is what maximising the likelihood of a target does — so the model becomes more confident about everything, including the things it is wrong about. Preference tuning compounds it, because judges and annotators prefer confident answers to hedged ones, so confidence is directly rewarded.

The result: a model whose 0.9 used to mean 0.9 and now means 0.75.

Why you should care

If nothing downstream reads the confidence, you can ignore this. The moment anything does, it matters completely:

- "Fall back to the expensive model below 0.7" — a rule about a number that no longer means what it did, so your fallback fires on the wrong items and your cost model is wrong.
- "Auto-approve above 0.95" — now auto-approving cases the model is not actually sure about.
- "Say 'I am not certain' when confidence is low" — the model stops saying it, because its confidence stopped being low.

Every abstention and routing design in this course depends on calibration, and fine-tuning is the step that breaks it.

Measuring it

For a classifier the measurement is direct. Bin held-out predictions by confidence and compare each bin's mean confidence with its accuracy:

```
bins = np.linspace(0, 1, 11)
idx = np.digitize(conf, bins) - 1
ece = 0.0
for b in range(10):
    m = idx == b
    if m.sum():
        ece += m.mean() * abs(acc[m].mean() - conf[m].mean())
print(f"ECE = {ece:.3f}")
```

Expected calibration error is the average gap, weighted by how many items fall in each bin. Under about 0.05 is good; above 0.15 means the numbers cannot carry a threshold. Plotting the same thing — accuracy against confidence, with the diagonal drawn — shows you the *shape* of the miscalibration, which the single number hides.

For a generative model there are two usable proxies:

- **Sequence probability** of the answer it produced, though this confounds confidence with length and with how many ways there were to phrase the same answer.
- **Verbalised confidence** — ask the model for a number, then check whether that number tracks accuracy. Crude, surprisingly serviceable, and it is what most abstention designs end up using.

Fixing it

Temperature scaling. Fit a single scalar on a validation set, divide the logits by it before the softmax. One parameter, ten lines of code, and it corrects most of the gap without changing any prediction — the ranking is preserved, only the confidences move. Use it on any classifier whose confidence feeds a decision.

Train on uncertainty. Include examples where the correct answer is "I do not have that information", "this could be either A or B", or "I would need the order number to answer". If your dataset contains no uncertainty, the model has never been shown that uncertainty is an available output. This is the same mechanism as the refusal problem in the previous lesson, and the same fix.

Recalibrate after every change. Quantisation shifts the logits. Merging an adapter shifts them. A new base model shifts them. The scalar you fitted in March is fitted to a model that no longer exists.

Abstention is a separate decision

One clarification that saves confusion. Calibration tells you whether a confidence is honest. It does not tell you where to put the threshold — that is a cost question, and it belongs to whoever owns the consequences.

If a wrong routing decision costs £4 to correct and an escalation to the expensive path costs £0.30, the threshold sits wherever expected cost is minimised, which on a calibrated model you can compute directly rather than guess:

```
best = min(thresholds, key=lambda t:
           0.30 * (conf < t).mean() + 4.00 * ((conf >= t) &
           wrong).mean())
```

That is two numbers a business can supply and one line of code. The common failure is to pick 0.7 because it sounds cautious, on a model whose 0.7 has never been checked, and to call the result a policy.

The rule to carry

If a number in your system crosses a threshold to make a decision, that number needs a calibration check on the exact artefact you serve, refreshed whenever the artefact changes.

And if you have not checked it, do not build the threshold. A confidence-based fallback on an uncalibrated model is not a safety mechanism; it is a mechanism that fires at a

rate you have not measured, which is worse than no mechanism because it is believed.

BEFORE YOU MOVE ON

A routing system sends anything below 0.7 confidence to an expensive model. After fine-tuning, task accuracy improved but the expensive fallback almost never fires and overall error rose. What happened?

- A. The fine-tune improved accuracy on easy cases only, so the hard cases that used to be routed are now handled badly by the small model.
 - B. The threshold should be recomputed for every new model, since confidence scales with the number of classes and the class count changed.
 - C. Training sharpened the output distribution, so confidences rose across the board and uncertain cases now score above the threshold.
 - D. Temperature scaling was applied during training, which compressed the confidence range and pushed all values upward.
-

75 · 8 min

Ablations that earn their cost

THE ONE THING TO KEEP

Run an ablation only when its result would change a stated decision, and read every result against a three-seed noise floor — the data learning curve is the one that settles whether another week of labelling is worth it, and a written log of non-findings prevents the same experiment being run twice.

The test for whether to run one

An ablation is worth running if, and only if, its result would change a decision. Write the decision down before you launch it:

If more data helps, we spend next week labelling. If it does not, we ship and move on.

That sentence turns a curiosity into an experiment. Without it, you will run six configurations, produce six numbers, and choose the largest — which is a way of selecting noise.

The four that usually earn it

1. The data learning curve. Train on 25%, 50%, 75% and 100% of your examples. Plot the task metric against the count.

This is the most valuable ablation in applied fine-tuning, because it answers the question teams argue about most: whether to collect more data. A curve still climbing steeply at 100% says yes. A flat curve from 50% onwards says the next thousand examples will do nothing, and your effort belongs somewhere else — better labels, a different base model, a different approach entirely.

2. The replay ratio. 100% task data against 80/20 against 60/40, scored on both the task metric and the canary set. This one usually changes what you ship, because the canary recovery at 80/20 typically costs almost nothing on the task.

3. Epochs. You get this free from your checkpoints; no extra runs needed. Evaluate the checkpoint at the end of each epoch on both metrics.

4. Rank, once. $r = 8$ against $r = 32$ at fixed alpha-to-rank ratio. If they are within noise — and they usually are — you have learned that rank is not your constraint and you can stop thinking about it for this project.

What it costs

Six short runs on an 8B QLoRA over a few thousand examples is roughly six to ten GPU-hours. On a rented consumer card that is under ten pounds, and on a free tier it is two long sessions.

Against that: a week of labelling that turns out to have been unnecessary, or a shipped model that is two points worse than one of the configurations you never tried.

One factor at a time

Change one thing between runs and hold everything else — including the seed — fixed. It is tempting to change two, since the runs are cheap, but you then cannot attribute the difference and you have spent the compute for nothing.

The exception: when you genuinely believe two factors interact, run the full grid on a small scale rather than two single-factor comparisons. Rank and dataset size plausibly interact; learning rate and epochs certainly do. Four runs instead of two, and an answer you can interpret.

Against your noise floor

This is what makes the whole exercise honest. From the reproducibility lesson: run your baseline configuration three times with different seeds and record the spread.

```
noise floor (3 seeds):  0.858 – 0.884,  sd ≈ 1.3 pts

data 25%    0.801
data 50%    0.847
data 75%    0.869
data 100%   0.871    ← flat from 75%; more data will not help
replay 80/20 0.870  canary 0.81  (vs canary 0.62 at 100% task) ← ship
this
rank 32      0.876    ← +0.5 pt, inside noise, ignore
```

Two real findings and one non-finding, and the non-finding is as useful as the others because it stops an argument.

Write it down where it will be found

A file in the repository, one line per run: what changed, the metric, the canary, the date, the checkpoint path.

2026-09-14	baseline s0/s1/s2	.871/.858/.884	canary .62	
2026-09-14	data 50%	.847	canary .64	
2026-09-15	replay 80/20	.870	canary .81	←
	shipped			
2026-09-15	rank 32	.876	canary .60	

Three months later somebody will propose raising the rank. This file answers them in ten seconds, and without it the answer is another two days of compute and argument. The log is the durable output of an ablation; the model is the perishable one.

BEFORE YOU MOVE ON

An ablation shows task metric 0.847 at 50% of the data, 0.869 at 75% and 0.871 at 100%, against a three-seed noise floor with a standard deviation of 1.3 points. What does this support?

- A. Collecting more data is the best investment, since the metric rose across every step of the curve.
- B. The dataset should be reduced to 75%, since the last quarter contributes nothing and shorter runs are cheaper.
- C. The noise floor is too large to interpret the curve at all, and every point needs three seeds before anything can be concluded.
- D. The curve has flattened between 75% and 100%, so further labelling is unlikely to help and effort belongs elsewhere.

CASE STUDY · MODULE 8

Ninety-four, then seventy-nine

A three-hospital group in Kochi, one week before launching a model that sorts incoming patient messages into six clinical queues.

The number on the slide was 0.94. Routing accuracy, six queues, on a held-out set of six hundred messages. The base model with the best prompt the team could write had scored 0.86. The project had run for eleven weeks and the figure had been shown to the group's medical director twice.

Four days before the launch, a data analyst called Meera was asked to write the evaluation section of the model card. She needed to state how the held-out set had been constructed, and discovered that nobody had written it down.

The split had been random. Six thousand messages, shuffled, ninety per cent for training and ten per cent held out.

She ran three checks, each of which took under an hour.

The first was near-duplicate overlap across the split. She built a MinHash index over the training texts and queried every evaluation item against it. Eighty-one of the six hundred evaluation messages had a near-duplicate in the training set — thirteen and a half per cent. Patients ask the same things in the same words, and appointment-rescheduling messages in particular are close to templated.

The second was group leakage. The messages carried a patient identifier. Two hundred and six of the six hundred evaluation messages came from patients who also appeared in training, some of them with a dozen messages each. The model had learned individual patients' phrasing, their conditions, the department they usually wrote to.

The third was time. The data ran from January to August. A random split had put August in training and February in evaluation, which is the reverse of what production does. In May the group had merged two departments and renamed a queue.

She rebuilt the split properly: grouped by patient, cut at a date so that everything after the first of July was evaluation and everything before was training, and removed evaluation items with a near-duplicate on the other side.

The model scored 0.79.

The prompt baseline, re-run on the same corrected set, scored 0.77.

That is the whole finding. Eleven weeks of work had bought two points, not eight, and on a corrected set of four hundred and eleven items the ninety-five per cent interval on a difference of two points comfortably includes zero.

The decision Meera put in front of the project lead was not whether the model was good. It was what to say, four days before a launch, about a number that had been shown to the medical director twice.

Launching on the old number meant a figure in a model card that the team now knew was inflated by leakage, in a clinical setting, where the queue a message lands in determines how quickly somebody reads it. Correcting it meant explaining that the headline figure had been wrong for eleven weeks, that the honest gap over a prompt was within noise, and that the launch should probably be a pilot rather than a launch.

There was also a third thing, which Meera raised and which nobody had thought about at all. Nobody had run a refusal set. The messages included requests for medical advice that the system was never intended to answer, and no one had measured whether the fine-tuned model still declined them as the base model had.

STOP HERE

Before turning the page: what would you have done, and what would it have cost if you were wrong? Write it down. The point of the next section is to compare.

WHAT HAPPENED

The project lead took the corrected numbers to the medical director himself, with the three leakage measurements attached. The launch became a four-week pilot on one hospital, with every routed message still read by the triage nurse who had been reading them before.

The refusal check was run that week. On a hundred and fifty prompts the base model had declined — requests for a diagnosis, for a medication change, for another patient's information — the base refused 141 and the tuned model refused 96. Nothing in the training data had been harmful; it was six thousand examples all demonstrating that a message gets answered.

Sixty refusal examples in the hospital's own wording went into the training set, a replay mixture of fifteen per cent general instruction data went in with them, and the model was retrained. Refusals came back to 138. Routing accuracy on the corrected held-out set moved from 0.79 to 0.78, which was inside the three-seed noise floor of 1.1 points.

The pilot ran. The model is in production now on one queue at a time. The evaluation set has a rule attached to it that Meera wrote: it is grouped by patient, cut by date, checked for thirteen-gram overlap, and every message the system routes wrongly in production is added to it that day.

Three separate leaks were found in one random split. Name each and the fix, and say which one inflates a score most invisibly.

Near-duplicate leakage — variations of the same message on both sides — fixed by measuring cross-split overlap with MinHash and removing the matches. Group leakage — the same patients on both sides — fixed by splitting on the patient identifier rather than the row. Temporal leakage — training on August and testing on February — fixed by cutting at a date so evaluation is entirely in the future. The near-duplicate one is the most invisible: it leaves no field to group by and no date to inspect, it is not visible by reading either set, and almost nobody measures it.

On the corrected set the tuned model scored 0.79 and the prompt baseline 0.77. Why is it wrong to call that a two-point improvement?

Because neither number is exact. On four hundred and eleven items the ninety-five per cent interval on a proportion near 0.78 is roughly plus or minus four points, so a two-point gap is well inside it. On top of that sits the training noise floor — the spread across seeds, measured here at about 1.1 points — which is a second, separate uncertainty. A difference smaller than about twice the seed spread, on a set too small to resolve it, is a coin flip that has been given a name.

The refusal regression came from entirely benign data. Explain the mechanism, and why the fix costs an afternoon.

Alignment is a thin, late layer installed on comparatively little data, and every training example in the set demonstrated a request being answered — so the gradient repeatedly taught that the assistant's job is to produce an answer, and the model generalised that beyond the requests intended. The fix is cheap because it is a data fix: thirty to a hundred examples of the model declining appropriately, in the product's own wording, plus a replay mixture of general instruction data. It also gives the model a way of declining that suits the hospital rather than a generic refusal from somewhere else.

Check what stuck

6 questions, one right answer each. This one is for you, not for a record: answers and the reason behind each are at the back. If two questions from the same module go wrong, re-read that module before the exam rather than after.

- 1 A tuned model wins 58 of 100 blind pairwise comparisons against the baseline. What is the honest reading?
 - A. An eight-point win, which is worth shipping.
 - B. Consistent with no difference at all, because the interval near fifty per cent on a hundred items is about ten points.
 - C. A loss, because the ties were not counted separately.
 - D. Unreadable, because a pairwise comparison cannot produce a number at all unless at least three judges rate every item and an agreement figure is computed between them.
- 2 You tuned an 8B model to save money against a hosted frontier model. What is the honest comparison?
 - A. The tuned 8B against every model on the market at the same parameter count, since parameter count is what determines the cost per thousand requests.
 - B. The tuned 8B against the frontier model you had been paying for.
 - C. The tuned 8B against the same base 8B with the same prompt.
 - D. The tuned 8B against the smallest untuned model that a good prompt makes adequate.
- 3 Your random split scores 0.94 and a group-and-time split of the same data scores 0.79. What best explains the gap?
 - A. The model memorised near-duplicates and the same customers, and scored by recall rather than by generalising.
 - B. The time split includes months the model was not trained for, so it is an unfair comparison.
 - C. Random splitting is statistically invalid, which is why two numbers produced by the two methods can never be compared with one another under any circumstances.
 - D. The grouped split is smaller, so its interval is wider.

- 4 Your tuned model's answers are fifty per cent longer and win 60 per cent of comparisons. What do you do next?
- A. Discard the comparison entirely, since a length difference of that size makes a pairwise judgement worthless no matter how the results are analysed afterwards.
 - B. Ship it, because the win rate is outside the interval.
 - C. Compute the win rate again on the items where the two answers were of similar length.
 - D. Shorten every answer with a post-processor and re-run the whole comparison.
- 5 A student distilled from a large teacher beats that teacher on a teacher-judged comparison. What should you conclude first?
- A. That the distillation dataset must have been unusually well filtered.
 - B. That this is a finding about the measurement, not about the model.
 - C. That the teacher has degraded since the dataset was generated, which happens when a hosted model is updated in place between the two runs.
 - D. That the student has genuinely surpassed its teacher on this task.
- 6 You route to an expensive model whenever confidence falls below 0.7. What does fine-tuning do to that rule?
- A. It makes the model more cautious, so the fallback fires more often and costs rise.
 - B. It removes confidence scores altogether, since a fine-tuned generative model can report only sequence probabilities rather than a calibrated number.
 - C. Nothing, because the threshold is applied after the model has produced its answer.
 - D. It sharpens the distribution, so 0.7 no longer means what it did and the fallback fires on the wrong items.

MODULE 9

Shipping it, and living with it

Serving, quantising, monitoring, versioning and the day the base model is replaced. A fine-tune is not a deliverable, it is a dependency you now maintain.

BY THE END OF THIS MODULE YOU CAN

Serve a tuned model or adapter at a stated cost and latency, monitor it for drift, reproduce a six-month-old artefact from its recorded version, and plan for the day the base model is replaced

76 · 9 min

Serving the Adapter

THE ONE THING TO KEEP

Merging, quantizing and switching engines all change the model; evaluate the exact artifact you serve.

What you actually have

After a LoRA run the artifact is small:

```
out/checkpoint-300/  
  adapter_model.safetensors # 40-200 MB  
  adapter_config.json      # rank, alpha, target modules, base model  
id
```

It is meaningless without the exact base model. `adapter_config.json` records the base repository name but not necessarily the revision, so pin the commit hash yourself in your own config. Base repositories do get updated in place. An adapter trained against last month's weights and served against this month's is a subtle quality bug that nobody will diagnose.

Merge, or keep it separate

Serve the adapter separately when you want to A/B against the base, ship several variants, or turn the fine-tune off without a redeploy. vLLM does this natively:

```
vllm serve Qwen/Qwen2.5-7B-Instruct \
  --enable-lora \
  --lora-modules support=/models/support-lora \
  --max-lora-rank 16
```

Then request `"model": "support"`, or the base name for the base. One GPU, one copy of the base weights, many adapters: this is how you serve a per-customer model without a GPU per customer. Expect a modest throughput cost, in the region of 10-20%, growing with how many adapters are live at once.

Merge when you have one variant and want no overhead:

```
from peft import PeftModel
from transformers import AutoModelForCausalLM
base = AutoModelForCausalLM.from_pretrained(BASE,
  torch_dtype="bfloat16")
merged = PeftModel.from_pretrained(base, "out/checkpoint-
  300").merge_and_unload()
merged.save_pretrained("merged-7b")
```

You get an ordinary model any server can load, at the cost of a full 14 GB artifact per variant and no ability to stack adapters.

One caveat that bites: if you trained with QLoRA, the adapter learned as a correction to a **4-bit** base. Merging it into the 16-bit base is not the same arithmetic. It usually works, and it sometimes moves quality by a point or two — a specific case of the general rule below.

Evaluate the artifact you ship

This is the most-skipped step in the whole pipeline.

Every transformation after training changes the model: merging, quantizing to GGUF or AWQ or GPTQ, a different inference engine, different sampling defaults, a chat template the server applies for you. Any of them can move your numbers, and none of them announce it.

So run your held-out evaluation *through the endpoint you will actually serve*, with production sampling settings, and compare it to the number from the training rig. Same inputs, same judge. If it dropped, you found out before your users did.

Running it locally

For CPU or laptop deployment: merge, convert, quantize.

```
python convert_hf_to_gguf.py ./merged-7b --outfile model.f16.gguf
./llama-quantize model.f16.gguf model.Q4_K_M.gguf Q4_K_M
```

Q4_K_M on a 7B gives roughly a 4 GB file that runs in 8 GB of RAM at a few tokens per second on a modern CPU. That is usable for batch work and personal tools, and slow for interactive chat. Then a Modelfile for Ollama:

Merge it in, or keep it separate

Keep the adapter separate

- one base in memory, many adapters
- turn the fine-tune off with a flag, not a redeploy
- A/B two versions against the same base
- costs roughly 10 to 20 per cent of throughput

Merge it into the base

- an ordinary model any server loads
- no adapter overhead at all
- a full 14 GB artefact per variant, and no stacking
- a QLoRA adapter was fitted against a 4-bit base, so merging into the 16-bit one is not the same arithmetic

Whichever you choose, run the held-out and canary sets through the endpoint you will actually serve, with production sampling settings. Every transformation after training changes the model and none of them announce it.

```
FROM ./model.Q4_K_M.gguf
TEMPLATE ""...""
```

Set the template explicitly to the one you trained with. Do not accept a default.

The cost question, honestly

A 24 GB GPU running continuously costs roughly \$300-700 a month depending on provider and commitment. That buys a great many tokens from a hosted API. Self-serving a fine-tuned model wins when you have steady volume to keep the GPU busy, when the data has to stay on your infrastructure, when you need predictable latency, or when the behaviour genuinely cannot be prompted.

If your traffic is spiky and modest, a hosted endpoint for your adapter — or no fine-tune at all — is the cheaper engineering decision, and saying so is not a failure.

Before you ship

- Base model repository and commit hash pinned.
- Chat template pinned in the serving config, matching training.
- Held-out evaluation re-run against the real endpoint.
- Canary set re-run against the real endpoint.
- The base model still deployed, one config flag away.

BEFORE YOU MOVE ON

You merged a QLoRA adapter into the 16-bit base, converted to 4-bit GGUF, and deployed on Ollama. The house style is clearly present, but accuracy is below the number you recorded after training. What do you do first?

- A. Requantize at Q8_0, since 4-bit quantization is the most likely culprit.
- B. Retrain at a higher rank; the adapter was probably underpowered.
- C. Go back to serving the unmerged adapter with vLLM and treat the GGUF path as broken.
- D. Re-run the held-out evaluation through the Ollama endpoint with production sampling settings, so you are measuring the artifact users actually get.

77 · 9 min

Quantising the model you ship

THE ONE THING TO KEEP

Serve with AWQ on a GPU or Q4_K_M in GGUF on a laptop, merge into the full-precision base before quantising, calibrate on your own data — and run a three-way base, tuned, tuned-and-quantised comparison, because quantisation error can be as large as your adapter's delta and round the entire fine-tune away.

Training quantisation and serving quantisation are different

QLoRA's NF4 exists to make a frozen base fit in memory during training. It is not the format you serve in — it is comparatively slow at inference, because it was optimised for memory rather than throughput.

For serving, pick by where you are running:

GGUF, for CPU, laptops, phones and mixed CPU-GPU. The `llama.cpp` format, and what Ollama and LM Studio consume. Its k-quant variants give a fine-grained quality-size ladder. `Q4_K_M` is the widely used default and the usual sweet spot; `Q5_K_M` if you have the memory; `Q8_0` when you want near-lossless.

AWQ, for GPU serving. Activation-aware weight quantisation protects the small fraction of weights that matter most to activations, and it is well supported by vLLM with good throughput at 4 bits.

GPTQ, the older calibration-based method. Still widely available, generally superseded by AWQ for new work.

FP8, on H100-class hardware. Close to lossless and fast, when you have the cards for it.

The sizes

For an 8B model, approximately:

fp16 / bf16	16.0 GB
Q8_0	8.5 GB
Q5_K_M	5.7 GB
Q4_K_M	4.9 GB
Q3_K_M	4.0 GB

The gap between fp16 and Q4_K_M is the difference between needing a data-centre card and running on a laptop.

What it costs in quality

At 4 bits on a model of 7B or more, the degradation is small — measurable on perplexity, often not visible on a task. Below 4 bits it becomes real, and Q3 and Q2 are noticeably worse in ways users notice.

Two effects deserve naming, because they are not general knowledge:

Smaller models suffer more. A 1B model at 4 bits loses proportionally more than an 8B model does. Quantisation error is roughly constant per weight, and a small model has less redundancy to absorb it. For anything under 3B, prefer 8-bit.

Quantisation can erase a fine-tune. This is the one that catches people. Your adapter changed the weights by a small amount — that is what a low-rank update of modest magnitude does. If the quantisation step's error is of comparable size to your adapter's delta, the changes are rounded away, and the quantised tuned model behaves like the quantised base model.

The symptom is a model that evaluated well before conversion and behaves like it was never trained afterwards. The diagnosis is a three-way comparison — base, tuned, tuned-and-quantised — on the same set, which takes minutes and is otherwise very confusing to debug.

The order of operations

```
# 1. merge the adapter into the fp16 base
python -c "from peft import AutoPeftModelForCausalLM; \
    AutoPeftModelForCausalLM.from_pretrained('out/adapter', \
    torch_dtype='float16') \
    .merge_and_unload().save_pretrained('out/merged')"
```

```
# 2. convert, then quantise
python llama.cpp/convert_hf_to_gguf.py out/merged --outfile m-f16.gguf
./llama-quantize m-f16.gguf m-q4km.gguf Q4_K_M
```

Merge into the **full-precision** base, then quantise the merged model. Merging into an already-quantised base, or quantising and then merging, both introduce error the adapter was never fitted against.

If the adapter was trained with QLoRA, there is a subtlety worth stating: it was fitted against a 4-bit base, and merging it into the fp16 base is not exactly what it was trained for. In practice this usually works and occasionally loses a little quality. Measure it; that is what the three-way comparison is for.

Calibration data is not optional

AWQ and GPTQ need a calibration set to decide which weights matter. The default examples use general web text.

Use **your own data** instead — a few hundred examples from your task. The method is choosing what to protect based on which weights produce large activations on the calibration inputs, so calibrating on generic text protects the weights that matter for generic text. Calibrating on your data protects the weights that matter for your task, and the difference is measurable on task metrics.

The rule that covers all of it

Evaluate the exact artefact you will serve. Not the adapter, not the merged fp16 model, not the same quantisation of a different checkpoint. The file that will answer requests.

Quantisation, merging, format conversion and engine choice each change the model's output. Any of them can lose you a point, and one of them can lose you the entire fine-tune. Your evaluation script should take a path and a backend and run against whatever is there, so that running it on the final file costs nothing and happens every time.

BEFORE YOU MOVE ON

A tuned 1.5B model evaluates at 0.88 in bf16 and at 0.71 after quantising to 4 bits, while the base model scores 0.70 both before and after. What is the likeliest mechanism?

- A. The adapter's weight changes are small relative to quantisation error at this size, so they rounded away.
 - B. The calibration set was too small, so the quantiser chose poor scaling factors for the whole model.
 - C. Quantisation removed the adapter, since merged weights revert to base values when the model is converted.
 - D. Four-bit quantisation always costs about 17 points of accuracy, which is why 8-bit is the standard for production.
-

78 · 8 min

Serving many adapters from one base

THE ONE THING TO KEEP

One base plus fifty adapters is 24 GB where fifty merged models are 800 GB, at 10 to 20% throughput cost and a shared rank ceiling you must choose at training time — but one adapter trained on everyone's data with the customer named in the prompt often matches per-customer adapters, so check whether you are paying for quality or for isolation.

The arithmetic that makes this interesting

Fifty customers, each wanting a model tuned on their data.

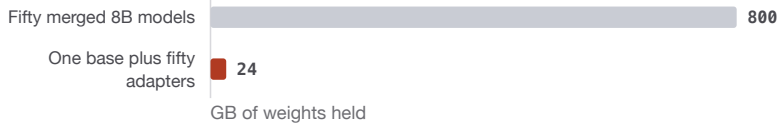
- **Fifty merged models**, 8B each at 16 GB: 800 GB of weights and fifty serving processes.
- **One base plus fifty adapters**: 16 GB plus fifty times 160 MB, so about 24 GB. One process.

That is the whole argument, and it is decisive. Multi-adapter serving turns per-customer models from an impossible proposition into a routine one.

How it works

The base weights stay resident. For each request, the server applies the named adapter's low-rank matrices during the forward pass. Modern implementations batch requests using *different* adapters together, computing the base matrix multiplication once for the whole batch and the small per-adapter corrections separately.

Fifty customers, two ways



One process instead of fifty. Then ask the question nobody asks first: whether one adapter trained on everyone's data, with the customer named in the prompt, would have matched them — because you may be paying for isolation rather than for quality.

```
vllm serve meta-llama/Llama-3.1-8B-Instruct \
  --enable-lora --max-loras 8 --max-lora-rank 32 \
  --lora-modules acme=/adapters/acme globex=/adapters/globex
```

Then the request names the adapter as its model:

```
{"model": "acme", "messages": [...]}
```

LoRAX is a dedicated server built around the same idea with dynamic loading from a registry; vLLM's support is built in and sufficient for most cases.

The costs

Throughput. The extra matrix multiplications cost roughly 10 to 20% against a merged model. Real, and usually worth paying for the operational simplification.

Concurrency limits. `--max-loras` caps how many adapters are active in the GPU at once. Beyond it, adapters are swapped in from host memory or disk, adding latency to the first request that needs one. If you have 500 adapters and 20 are hot, size the cap for the hot set.

Uniformity requirements. Every adapter must share the base model, and rank must be at or below `--max-lora-rank`. Plan this at training time: pick one rank for the whole fleet, because a single rank-128 adapter raises the ceiling and the memory cost for all of them.

Cold start. Loading an adapter from object storage on first use costs hundreds of milliseconds. Prewarm the ones you know will be needed.

When to merge instead

Merge when there is one adapter and it is always used. You get the base model's exact serving performance, any inference engine works without adapter support, and you can quantise the result freely.

Keep adapters separate when you have several, when you want to switch without a re-deploy, or when you want to A/B two versions against the same base — which is a genuinely useful capability that merging removes.

Composing adapters, honestly

The libraries let you load several adapters at once and combine them, with weights:

```
model.add_weighted_adapter(["formal", "legal"], [0.6, 0.4], "combined")
```

It works in the sense that it runs and produces something. Whether it produces the intended combination is much less certain: adapters trained independently occupy overlapping subspaces, and adding them can interfere in ways that degrade both behaviours rather than combining them. There are methods designed for this — concatenation-based merging, and orthogonality-constrained training — and the results are mixed.

Treat composition as an experiment, not a feature. If you need two behaviours together, train one adapter on both datasets. That reliably works, and it costs one more run.

Evaluating a fleet

Fifty adapters is fifty models, and each needs its own held-out set, its own canary and its own number. That is the cost people discover after building the serving layer.

The workable approach is a shared harness with per-tenant data: one evaluation script, one command, a results file per adapter, and a single summary that flags any tenant whose metric has fallen since the last build. Without that summary nobody looks at forty-nine of the fifty, and a regression affecting one customer surfaces as a support ticket.

Budget for this before you commit to per-customer adapters. It is usually a larger ongoing cost than the serving.

The per-customer question nobody asks first

Before building this, check whether you need it. One adapter trained on all customers' data, with the customer identified in the system prompt, frequently matches per-customer adapters and is enormously simpler — one artefact, one evaluation, one thing to retrain when the base model moves.

Per-customer adapters earn their complexity when the customers genuinely differ in ways a prompt cannot express, or when the data may not be mixed for contractual or regulatory reasons. That second reason is often the real one, and it is a good reason. But it is worth knowing that you are paying for isolation rather than for quality.

BEFORE YOU MOVE ON

A team serves 50 per-customer adapters from one base with vLLM. One customer's adapter was trained at rank 128; the others at rank 16. What is the consequence?

- A. The rank-128 adapter cannot be served alongside the others and must be merged into its own model instance.
- B. The max-lora-rank ceiling must rise to 128 for the whole server, costing memory in every slot.
- C. Requests to the rank-128 adapter will be 8 times slower, while the other 49 are unaffected.
- D. Adapters of different ranks cannot be batched together, so throughput collapses to sequential request handling.

The cost at steady state

THE ONE THING TO KEEP

At two million requests a month a self-hosted tuned 8B on redundant cards costs more than a hosted small model before anybody is paid to maintain it — so the crossover sits above roughly two million requests or wherever latency, privacy or vendor stability overrides the arithmetic entirely.

Compute the whole thing, once

The fine-tune was justified on cost. Here is how to check whether it was right, with everything included.

Take a workload: 2,000,000 requests a month, averaging 1,200 input tokens and 300 output tokens. That is 2.4 billion input and 600 million output tokens.

Hosted frontier model, at roughly \$3 per million input and \$15 per million output:

input	2,400M × \$3/M	= \$7,200
output	600M × \$15/M	= \$9,000

		\$16,200 / month

Hosted small model, at roughly \$0.15 and \$0.60:

input	2,400M × \$0.15/M	= \$360
output	600M × \$0.60/M	= \$360

		\$720 / month

Self-hosted tuned 8B. Two always-on L4-class cards at about \$0.80 an hour, because one is a single point of failure:

2 × \$0.80 × 730h	= \$1,168 / month
-------------------	-------------------

Notice the ordering. The self-hosted tuned model is **more expensive than a hosted small model** at this volume, before you have paid anybody to maintain it.

The costs that do not appear on a cloud bill

- **Idle capacity.** Traffic is not flat. Sized for peak, your cards are idle much of the night, and you pay for every hour.
- **Redundancy.** One card is a single point of failure, so it is two, which is the single largest multiplier in the calculation.
- **The engineer.** Deployments, upgrades, a driver incident, an out-of-memory at 2 a.m. Even at a tenth of a person's time this is the biggest line in the table and it is the one that never gets written down.
- **Retraining.** Each base-model move costs days.
- **Evaluation maintenance.** The held-out set ages and must be refreshed.

Against those, one cost that is often understated on the other side: **hosted models change under you.** A provider deprecates a version or updates it silently, and your carefully measured behaviour shifts. A model you host does exactly what it did yesterday, forever, and for some products that stability is the whole point.

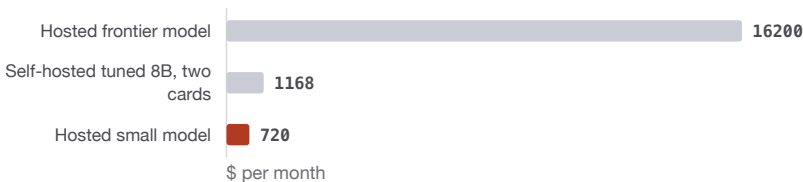
Where the crossover actually is

Rough shape, for a typical workload:

- **Under ~200k requests a month.** Hosted, almost always. The fixed costs dominate everything.
- **200k to 2M.** Genuinely close. Depends on how spiky the traffic is, whether you already run GPUs, and whether you have someone to maintain it.
- **Over ~2M, or with a latency or privacy constraint.** Self-hosting wins clearly, and the fixed costs amortise into nothing.

And the constraints that override the arithmetic entirely: data that may not leave your infrastructure, an air-gapped deployment, a latency budget under 200 ms, or a per-request cost of zero because you call the model on every keystroke.

Two million requests a month, all in



The self-hosted tuned model is dearer than a hosted small model at this volume, before anybody is paid to maintain it. Idle capacity, redundancy and the engineer are the lines that never get written down.

Serverless, the middle option

Serverless GPU platforms charge per second of compute with scale-to-zero. For spiky or low-volume traffic they are much cheaper than always-on cards, and you pay for a cold start of several seconds on the first request after idle.

Good for internal tools, batch jobs and anything where a few seconds of latency on the first request is acceptable. Poor for interactive products with sparse traffic, where every user hits a cold start.

Throughput, which decides how many cards

The number that turns a workload into a card count is output tokens per second under batching. It varies enormously with model size, quantisation, sequence length, engine and batch size — an 8B model in AWQ on a modern mid-range card with vLLM, well batched, is in the range of several hundred to a couple of thousand output tokens per second aggregate.

That range is too wide to plan with, which is the point: **measure it on your hardware with your traffic shape** before sizing anything. A load test with realistic prompt lengths and realistic concurrency takes an afternoon and replaces a guess that could be wrong by a factor of four in either direction.

Write the table down

Put all the options on one page with quality, p95 latency and monthly cost including people. Then make the decision in front of somebody who will be affected by it.

Most of the time the arithmetic is clear once it is written. The reason fine-tuning projects get justified on cost and then cost more is not that the arithmetic is hard. It is that nobody did it.

BEFORE YOU MOVE ON

A team justifies self-hosting a tuned 8B model on cost at 500,000 requests a month. Which omitted item most often reverses that conclusion?

- A. The one-off training cost, which is larger than the monthly saving at this volume.
- B. The cost of the evaluation set, which must be relabelled monthly to detect drift.
- C. Redundancy and the engineering time to operate it, which appear on no cloud bill.
- D. Egress charges on model weights, which are incurred each time an instance restarts and pulls the checkpoint.

80 · 9 min

Monitoring a model that fails silently

THE ONE THING TO KEEP

A drifted model returns a confident wrong answer with a 200 status, so monitoring means unlabelled proxies — embedding distance from the training distribution, response length, refusal and fallback rates, and above all the rate at which users rephrase — plus a weekly shadow run of the frozen set and every incident turned into a permanent evaluation item.

The failure has no exception

A web service that breaks returns a 500. A fine-tuned model that has drifted returns a fluent, well-formatted, confident, wrong answer with a 200 status code. Every dashboard stays green.

So monitoring here is not about uptime. It is about building proxies for quality that fire without anybody labelling anything.

Input drift, which you can measure exactly

Your model was tuned on a distribution. Production moves: new products, a marketing campaign bringing a different audience, a seasonal pattern, a competitor's outage, a new phrasing people picked up somewhere.

Embed a rolling sample of production inputs and compare against the training distribution:

```
centroid = train_embeddings.mean(axis=0)
today    = embed(sample_of_today_inputs)
print("mean cosine to training centroid:", cos(today, centroid).mean())
print("share beyond 95th pct of training distances:",
      (dist(today, centroid) > p95_train).mean())
```

That second number is the useful one. If 5% of training inputs were beyond that threshold by construction, and today 22% are, your traffic has moved somewhere your model was not trained for. This is cheap, requires no labels, and fires before quality complaints do.

The output-side proxies

All free, all computable on every request:

- **Response length distribution.** A shift is the earliest sign something changed.
- **Refusal rate.** Both directions matter: a rise means the model is declining things it should handle; a fall may mean the safety regression from the previous module has crept in through a model update.
- **Format validity rate.** If constrained decoding is on this is 100% and worth asserting; if not, it is a direct quality signal.
- **Tool-call rate,** where applicable. A sudden rise or fall points at something.
- **Fallback rate,** if you route on confidence. This one is doubly useful, because it also detects the calibration shift from the previous module.
- **Empty or truncated outputs.** Should be near zero. Alert on any rise.

The user-behaviour proxies, which are the best ones

What users do after an answer is the closest free thing to a quality label:

- **Rephrase rate.** The user immediately asks again in different words. This is the single strongest unlabelled signal of a bad answer that exists.
- **Escalation rate.** They ask for a human.
- **Abandonment.** They leave without completing the task.
- **Explicit feedback,** which is sparse and biased towards complaint, but free — and, as the preference lesson noted, it is also KTO training data.

Plot these weekly against your model versions. A step change that lines up with a deployment is your answer.

The shadow evaluation

Run your frozen held-out set against the production model on a schedule — weekly is usually right.

This catches what the proxies cannot: a configuration change, a silently updated dependency, a quantisation applied during a deployment, a model that was rolled back and forgotten. It takes minutes and it is the only check that directly measures quality rather than a proxy for it.

And label 50 real production items a week. It is an hour of somebody's time and it gives you a running measurement on the current distribution, which no frozen set can pro-

vide — because your frozen set is, by construction, ageing.

From incident to permanent test

Every production failure becomes an evaluation item, that day, in the locked set.

This is the habit that compounds more than any other in this course. After a year, your evaluation set contains every way the system has ever failed, and no regression of a kind you have already seen can reach production again. It costs ten minutes per incident and it is the only mechanism that makes the evaluation set get *better* over time rather than staler.

Alert on the things that mean something

Do not alert on the model's loss, which you are not computing in production. Alert on:

- Input drift beyond a threshold you chose from the training distribution.
- Rephrase or escalation rate rising more than a stated amount week on week.
- Shadow evaluation falling more than your noise floor.
- Format validity or empty-output rate moving at all.

Four alerts, each with a number behind it that somebody chose deliberately. A dashboard nobody has set thresholds on is a dashboard nobody reads.

BEFORE YOU MOVE ON

Which production signal is the strongest unlabelled indicator that a tuned model's answers have degraded?

- A. Mean response latency, which rises when the model produces longer, less certain answers.
- B. The share of outputs failing schema validation, which indicates the model is losing its trained format.
- C. The explicit thumbs-down rate, which directly captures user dissatisfaction.
- D. The rate at which users immediately rephrase and resubmit their question.

The day the base model moves

THE ONE THING TO KEEP

An adapter fitted against one set of base weights loads cleanly onto a newer release and is silently worse, so pin the revision hash — and when a better base ships, first check whether it plus a good prompt already beats your tuned old model, because deleting the fine-tune is a success rather than a loss.

Adapters do not transfer

A LoRA is a set of deltas fitted against one exact set of base weights. Apply it to different weights — even the next release in the same family, with the identical architecture — and you are perturbing matrices it was never fitted against.

The unpleasant part: **it loads without error**. Shapes match, the adapter attaches, the model generates. It is simply worse, in ways that look like a regression rather than a category error. There is no exception to catch.

The same risk exists within a version. Model repositories are edited in place — a tokenizer fix, a re-uploaded shard, a chat template correction. Pin the revision hash of the base model in your config, and record it beside the adapter, so that "the same base model" means the same bytes.

What does transfer

Everything that matters, if you kept it:

- The dataset, with its filters and its data card.
- The evaluation set, the canary set and the refusal set.
- The configuration, including the learning rate and target modules you arrived at.
- The ablation log, which tells you what does and does not help on this data.

The model was always the perishable output. These are the durable ones, and a team that has them treats a base-model release as an afternoon.

The decision at each release

When a meaningfully better base ships — which for open-weight families has historically been every six to nine months — work through four questions in order.

1. Does the new base, with a good prompt, already beat our tuned old model?

Run it. This takes an hour and it is genuinely often yes, particularly when your fine-tune was addressing instruction-following or format adherence, which is exactly what each generation of base models does better.

If yes, the correct action is to delete the fine-tune and ship the prompt. That is a success — it means the capability you paid to install is now free — and it is regularly resisted for reasons that have more to do with sunk effort than with the product.

2. Does the new base tokenise our language better? For non-English work this is often the largest single improvement in a release, and it changes cost and context as well as quality.

3. What does retraining cost us? If the pipeline is one command, this is not a real question. If nobody has run the rebuild in eight months, find out now rather than under pressure.

4. Is anything in the interface different? A new chat template, new special tokens, a different system-prompt convention. These do not break loudly; they degrade.

The rebuild target

The answer to question three should be a command:

```
rebuild: dataset train merge quantise evaluate
@echo "artefact: out/merged-q4km.gguf"
```

Run it once a quarter against the *current* base whether or not you need to. A pipeline that is exercised regularly works; one that is exercised annually does not, and you discover that at the worst moment.

Time it and write the number in the decision record. "Rebuild: 4 hours, unattended" changes how an organisation reasons about its dependency on a model.

The migration itself

1. Retrain on the new base with the unchanged dataset and config.

2. Evaluate old and new on the same held-out set, canary set and refusal set.
3. Shadow the new model against live traffic without serving its output, comparing distributions of length, refusal rate and fallback rate.
4. Roll out to a slice, watch the user-behaviour proxies, then complete.
5. Keep the old artefact for a rollback window.

And re-run the cost table, because the new base may be a different size and change the arithmetic that justified self-hosting.

Four questions, in order, when a better base ships

- | | |
|---------------|--|
| First | ● Does the new base, with a good prompt, already beat our tuned old model? An hour to find out, and the answer is often yes. |
| Second | ● Does it tokenise our language better? For non-English work this is often the largest single improvement in a release. |
| Third | ● What does retraining cost us? If nobody has run the rebuild in eight months, find out now rather than under pressure. |
| Fourth | ● Is anything in the interface different — a new chat template, new special tokens, a changed system-prompt convention? These degrade rather than break. |

Deleting the fine-tune because the base caught up is a success: the capability you paid to install is now free. It is regularly resisted for reasons that have more to do with sunk effort than with the product.

The framing worth keeping

A fine-tune is not a thing you make. It is a thing you maintain, on a cadence set by somebody else's release schedule.

That is the honest closing argument of module one, arriving here as a practical matter. The projects that survive base-model churn are the ones whose durable artefacts — dataset, evaluation, configuration, decision record — were treated as the deliverable, with the weights as a build output. The projects that do not survive it are the ones where somebody produced a model.

BEFORE YOU MOVE ON

A team loads their existing adapter onto a newer release of the same model family. It loads without error, generates fluently, and performs worse than the old combination. Why is no error raised?

- A. PEFT validates only tensor shapes, which match across the release, so nothing detects that the underlying weights differ.
 - B. The adapter is applied at inference rather than merged, and unmerged adapters skip the compatibility check that merging performs.
 - C. Adapters store a hash of the base model, and a mismatch produces a warning rather than an error by default.
 - D. The new release is backwards compatible by design, so the degradation comes from the changed chat template rather than the weights.
-

82 · 8 min

Versioning so a model can be explained later

THE ONE THING TO KEEP

A manifest of content hashes travelling with the artefact, plus a model-version field on every logged request, is what turns 'what produced this answer in March' from a memory into a lookup — and rebuilding a random old version once a quarter is how you find out whether it actually works.

The question you will be asked

Someone forwards a complaint about a decision the system made in March. What produced that answer?

Answering it needs six things linked together: the request, the model version that served it, the adapter artefact, the base model revision, the dataset that trained it, and the configuration. Miss any one and the chain breaks.

The manifest

One file, written at build time, travelling with the artefact:

```
{
  "name": "trriage", "version": "3.2.0", "built": "2026-09-15T11:04Z",
  "base_model": "meta-llama/Llama-3.1-8B-Instruct",
  "base_revision": "0e9e39f31a2fc8d",
  "dataset": "ticket-trriage-sft@v3",
  "dataset_sha256": "9f2c1e...",
  "config_sha256": "71ba04...",
  "eval_set": "eval-v3", "eval_sha256": "c31d77...",
  "git_commit": "4f1a09",
  "results": {"task": 0.918, "canary": 0.812, "refusal": 0.94},
  "quantisation": "Q4_K_M",
  "seed": 0}

```

Everything in it is a hash or an identifier that resolves to exact bytes. "We trained on the September dataset" is a memory; 9f2c1e... is a fact.

Log the version on every request

```
{
  "ts": "2026-03-04T09:12Z", "request_id": "r_88412",
  "model": "trriage@3.2.0", "adapter": "trriage-3.2.0", "latency_ms": 340}

```

This single field is what connects a complaint to an artefact. Without it, a March request maps to "whatever was deployed in March", which nobody is certain about after two rollbacks.

Version numbers that carry meaning

A convention worth adopting:

- **Major** — the base model changed. Behaviour may differ everywhere; re-evaluate everything.
- **Minor** — the dataset changed. Same base, new training data.
- **Patch** — the same weights, a different build: a new quantisation, a new serving format.

So 3.2.0 and 3.2.1 are the same model in different formats, and their evaluation numbers will differ slightly — which is exactly the sort of thing that looks like a mystery until the version scheme tells you what changed.

Where the bytes live

All free:

- **Code and configs:** git.
- **Datasets:** a Hugging Face dataset repository (private repositories are free), or DVC or git-lfs against your own storage. All three give content-addressed versions.
- **Model artefacts:** a Hub model repository, or object storage with the version in the path.
- **Run metadata and results:** MLflow's registry if you want a searchable index, or a JSON file per build if you do not.

A fully adequate small-team setup is: code in git, dataset and adapter in private Hub repositories, manifest committed to git beside the config. No infrastructure, no cost, and it answers the March question.

Retention

Keep enough to explain a decision for as long as you might be asked about it. That period is set by your product and your obligations, not by disk cost — and a LoRA adapter is a few hundred megabytes, so the disk cost is not the constraint.

What to keep per version: the manifest, the adapter, the evaluation results, and a pointer to the dataset version. What you do not need to keep: every intermediate checkpoint, and the optimiser states, which are large and only useful for resuming a run you will not resume.

Three things that break this in practice

Every versioning scheme that fails does so in one of three ways, and all three are cheap to prevent.

A configuration value that is not in the hash. The config file is hashed, but the batch size comes from an environment variable, or the base model path from a shell default. Two runs with identical config hashes then differ. The fix is to resolve every value into the config file before hashing it, and to have the training script write out the *resolved* config as an artefact rather than the template.

A dataset built by something that is not deterministic. A notebook cell that samples without a seed, a query with `LIMIT 20000` and no `ORDER BY`, a filter that reads the current date. The hash then records what you built, which is correct and useless, because rerunning the builder produces a different file. Seed everything in the pipeline and pin the query, or accept that the dataset is an artefact to be stored rather than a recipe to be rerun — both work, and pretending you have the second when you have the first does not.

A floating reference to the base model. `main`, `latest`, or a tag that was moved. The manifest looks complete and resolves to different bytes over time. Always record the commit hash of the model repository, not its name.

Promotion, not just versioning

A version number says what something is. A *stage* says where it is allowed to run.

```
trriage@3.2.0    stage: production    since 2026-09-15
trriage@3.3.0    stage: staging        shadowing production traffic
trriage@3.4.0    stage: development
```

Three stages is enough. The value is that "which model is in production" becomes a fact stored somewhere rather than an inference from deployment logs, and rolling back becomes moving a pointer rather than rebuilding something.

The test

Once a quarter, pick a version at random and rebuild it from its manifest. Compare its evaluation numbers against the recorded ones.

If they match within noise, your versioning works. If you cannot rebuild it at all, you have discovered that on a Tuesday of your choosing rather than during an incident — which is the entire value of the exercise.

BEFORE YOU MOVE ON

A team stores adapters in object storage with a date in the filename and records the dataset as 'the September export'. What specifically fails when they are asked to explain a decision made six months earlier?

- A. Object storage does not preserve file integrity over six months, so the adapter may have been altered.
- B. Nothing links a request to an adapter, and 'the September export' resolves to no exact bytes.
- C. The date-based filenames make it impossible to tell which adapter is newest, so the wrong one may be examined.
- D. Without a model registry, the adapter cannot be loaded by the serving stack after a framework upgrade.

83 · 9 min

Model cards and disclosure

THE ONE THING TO KEEP

A model card states intended use, out-of-scope use, training-data provenance, evaluation with intervals, and limitations explained by mechanism — and since the regulatory position differs by jurisdiction and is still moving, that document is what travels everywhere while the specific legal duties do not.

The document that answers the questions

A model card is to a model what a data card is to a dataset: the page a stranger reads to decide whether they can rely on it. It is also, in practice, the page that answers procurement questionnaires, security reviews and the first hour of an incident.

It takes an hour to write and it is written once per major version.

```

# triage-3.2.0

## What it does
Routes a customer message to one of six queues and drafts a one-para-
graph
reply. English only.

## Base and licence
Llama-3.1-8B-Instruct @ 0e9e39f, under the Llama 3.1 Community Licence.
Naming and acceptable-use terms follow this derivative.
Adapter: LoRA r=16, all-linear. Merged and quantised to Q4_K_M.

## Training data
ticket-triage-sft@v3 - 31,084 examples: 24,003 internal support logs
(ToS v4 basis, PII-redacted), 2,140 hand-written, 4,941 OpenAssistant
(Apache-2.0) as replay. Data card: docs/data/ticket-triage-v3.md

## Evaluation
Held-out n=412, drawn after the training cutoff, group-split by cus-
tomer.
  routing accuracy    0.918  [0.888 - 0.941]
  valid JSON         1.000  (constrained decoding)
  canary (general)   0.812  (base 0.840)
  refusal set        0.94   (base 0.96)
Seed noise floor: sd 1.3 points over 3 runs.
Judged by: exact match on queue; replies by human pairwise, n=200.

## Intended use
First-line triage with a human reviewing anything the model routes to
'refunds' or 'escalation'.

## Out of scope
Not for medical, legal or financial advice. Not for non-English input.
Not for deciding refunds. Not an identity or eligibility check.

## Known limitations
- No data after 2026-08; the September pricing change is not represent-
ed.
- 4% of live traffic is non-English and is routed to 'other' by design.
- Canary is 2.8 points below base: general capability is slightly re-
duced.
- Confidence is uncalibrated above 0.9; do not threshold there.

## Contact
r.kapoor

```

The sections that earn their space

Out of scope. The most useful section and the hardest to write, because it requires saying no to uses somebody has already imagined. It is also the section that protects you, because a stated boundary is a decision rather than an oversight.

Known limitations, with mechanisms. Not "may occasionally be inaccurate" — which is true of everything and informs nobody — but the specific, mechanical limits: the data cutoff, the languages, the canary gap, the calibration range.

Evaluation with intervals and provenance. A number without an interval and without a statement of what judged it is not a result.

Disclosure to users

Separate from the card, and simpler: if a person is interacting with a model, say so, plainly, where they are looking. Do not let a name, an avatar or a writing style imply a person.

This is the right thing regardless of what any regulation requires, and it is also the minimum several regulations are converging on.

The regulatory shape, without advice

The landscape differs sharply by jurisdiction and is moving, so what follows is the shape of it rather than a statement of what you must do.

- The **European Union's AI Act** takes a risk-tiered approach with transparency obligations — disclosing that a person is interacting with an AI system, marking synthetic content — plus documentation duties on providers of general-purpose models, phasing in over several years. How its obligations fall on somebody who fine-tunes a third party's model is a question with guidance still being worked out.
- **Sectoral regulators** in many countries — financial, medical, employment — have applied existing rules to automated decisions without new AI-specific legislation, and those rules can be stricter than any AI act.
- Several jurisdictions have **synthetic-media and labelling rules** at various stages, and they do not agree with each other.
- Many countries have **nothing specific at all**, where general consumer-protection, data-protection and anti-discrimination law applies as it always did.

So: this is unsettled, it depends on where you operate and what the system decides, and I am not in a position to tell you what applies to you. What travels everywhere is the

documentation itself. A card with intended use, out-of-scope uses, training-data provenance, evaluation results and stated limitations is what any of these regimes asks for in some form, and it is what you would want to have written regardless.

Inherited terms

One practical item people miss: your base model's licence terms generally follow into your derivative. Community licences may carry naming requirements for derivative models and acceptable-use policies that bind your downstream users. Read them once, write the consequences into the card, and the question is answered permanently.

BEFORE YOU MOVE ON

Which limitations entry is actually useful to somebody deciding whether to rely on the model?

- A. "This model may occasionally produce inaccurate or inappropriate output and should be used with human oversight."
 - B. "The model was evaluated on a held-out set and performs well on the intended task."
 - C. "Training data ends 2026-08; confidence above 0.9 is uncalibrated and must not be thresholded."
 - D. "Performance may vary depending on input quality, domain and phrasing, as with all language models."
-

84 · 9 min

Retiring a fine-tune

THE ONE THING TO KEEP

Schedule the quarterly question of whether a fine-tune still beats its cost-matched baseline, retire it happily when the base model catches up, keep the decision record and cards after the weights are deleted — and treat a model nobody can rebuild as one to retire even while it is performing.

Schedule the question

Every quarter, for every fine-tune you maintain, answer one question with numbers: *does this still beat the cost-matched baseline by enough to justify what it costs to keep?*

It takes an afternoon. Re-run the frozen held-out set against three things: your tuned model, the current best base model with your best prompt, and the cheapest hosted model with the same prompt. Put them on the cost table from earlier in this module.

The reason this must be scheduled is that nobody ever asks it spontaneously. A working model is not a problem, so it is never on anybody's list, and it quietly accrues maintenance cost for years after the capability it provided became free.

The four reasons to retire

The base caught up. The most common, and the happiest. What you trained for arrives in the next generation of base models as default behaviour. The correct response is to delete the fine-tune and ship the prompt — and to record that you did, because it is evidence your team makes decisions on numbers rather than on attachment.

The task changed. The taxonomy was redrawn, the product moved, the six queues became four. A model trained on the old definition is now systematically wrong, and it is wrong confidently.

The volume collapsed. The fixed costs no longer amortise. A model serving 20,000 requests a month costs more to host than the same traffic costs on an API by a wide margin.

Nobody owns it. The person who built it left, the pipeline has not been run in a year, and no one can say what it was trained on. This is the most dangerous state in this lesson — a model still making decisions that nobody can explain or rebuild — and the correct action is to retire it even if it is performing, because you have no way to establish that it is.

Retiring gracefully

1. Stand the replacement up in shadow, serving nothing, and compare output distributions on live traffic.
2. Move a slice of traffic across and watch the user-behaviour proxies for a week.
3. Complete the switch, keeping the old artefact and its manifest for a rollback window.
4. After the window, delete the serving infrastructure.
5. **Keep the decision record, the data card and the model card.** They cost kilobytes and they are the institutional memory of what was tried and what it produced. The next person to propose this fine-tune should find out in ten minutes that it was built, measured and retired, and why.

Deletion, and the part that is unresolved

If you are retiring because of a data-protection obligation rather than economics, the picture is less tidy.

You can delete the dataset, the checkpoints, the merged weights, the quantised artefacts, the adapters in every registry, and the backups. That is tractable and you should have the inventory to do it, which is one more argument for the manifest.

What you cannot straightforwardly do is remove a specific person's influence from weights that remain in service. Whether that is required is not settled: regulators have expressed different views, the question of whether model weights constitute personal data is argued seriously on both sides, and the answer differs by jurisdiction. Nobody can currently tell you with confidence what the obligation is, and this course will not pretend otherwise.

Machine unlearning — methods that approximately remove one example's influence without full retraining — is an active research field with genuine results and genuine limits. The honest summary is that it is not yet something to build a compliance commitment on. The engineering answer that works under any interpretation is the one from the privacy lesson: keep the dataset rebuildable, retrain on a schedule, and record which data version produced which model, so that removing somebody's data from the next build is routine.

The closing thought for the course

This course opened by arguing that most fine-tuning projects should not happen, and it ends by arguing that many of the ones that did happen should now stop. Both are the same claim.

A fine-tune is a bet that a specific capability is worth owning, maintaining and defending against a field that improves underneath it every few months. Sometimes that bet is clearly right — a latency budget, an air gap, a distillation at volume, a verifiable task where reliability is the product. Often it is not, and it was made because training a model is more interesting than writing a prompt.

The discipline that separates the two is the same at the start and at the end: a frozen baseline, a held-out set you did not choose afterwards, a number with an interval, and a stated condition under which you would stop. Everything else in these nine modules is machinery. That is the method.

BEFORE YOU MOVE ON

A fine-tune still performs well, but the engineer who built it has left, the pipeline has not been run in a year, and nobody can say what it was trained on. Why is retiring it the right action despite the performance?

- A. Unmaintained models degrade over time as their weights drift, so performance will fall regardless.
- B. The licence of the base model will eventually require re-acceptance, which nobody remaining is authorised to do.
- C. Models older than a year are contaminated relative to current benchmarks, making any reported evaluation invalid.
- D. Its performance cannot be established or explained, and it cannot be rebuilt.

CASE STUDY · MODULE 9

The fine-tune that the quantiser rounded away

A coaching centre chain in Kota, shipping an offline doubt-solving assistant on low-cost Android tablets to sixty rural study centres with no reliable internet.

The centres had tablets, two gigabytes of usable RAM each, and no connectivity worth depending on. The product was a model that answered physics and chemistry doubts in the chain's own explanatory style — a fixed structure of restate, principle, working, common mistake — which the teaching staff had spent years standardising and which the students recognised.

The team fine-tuned a 1.5B model. At that size full fine-tuning was affordable and simpler than an adapter, and the results were good. On four hundred held-out doubts the tuned model produced the four-part structure on 96 per cent of answers against 31 per cent for the base model with the best prompt, and two senior teachers preferred its explanations in a blind comparison at better than three to one.

Then it was converted for the tablets: merged, converted to GGUF, quantised to Q4_K_M, which brought the file to about a gigabyte and made it fit.

The first field build came back from two centres with the same complaint. The answers were fine but they were not the chain's answers. No four-part structure. The common-mistake paragraph, which the teachers considered the most useful part, was simply absent most of the time.

The engineer, Tanvi, assumed a template problem and spent a day on it. The Modelfile template was correct and matched training. The prompt was identical. The tokenizer was the one saved beside the weights.

What she had not done was evaluate the artefact she shipped. Every number in the project came from the fp16 model on the training rig.

She ran the three-way comparison, which took twenty minutes: base quantised, tuned fp16, tuned quantised, on the same four hundred doubts. The tuned fp16 model produced the structure on 96 per cent. The tuned quantised model produced it on 44 per cent. The quantised base produced it on 29 per cent.

The fine-tune had very nearly been rounded away.

The mechanism is unglamorous. Her training had changed the weights by a small amount — full fine-tuning of a small model over four thousand examples is still a modest perturbation. Quantisation error at four bits is roughly constant per weight, and on a 1.5B model there is much less redundancy to absorb it than on an 8B one. When the

quantisation error is of comparable size to the change training made, the change is rounded away, and the quantised tuned model behaves substantially like the quantised base.

Her options each cost something the chain would feel.

Ship at Q8_o. About 1.8 GB instead of 1.0, which on a tablet with two gigabytes of usable RAM does not leave room for the application. Some centres had newer tablets; most did not.

Ship a larger model at four bits, where quantisation costs proportionally less. A 3B model at Q4_K_M is around two gigabytes, which is worse on the same tablets, and slower on processors already producing a few tokens a second.

Train harder, so that the delta is larger than the quantisation error. More epochs and a higher learning rate would push the weights further — and on a 1.5B model that also means more memorisation and a sharper collapse on doubts unlike anything in the training set, which in a subject where students ask strange questions is not a small cost.

Accept 44 per cent, and have the application enforce the structure instead.

STOP HERE

Before turning the page: what would you have done, and what would it have cost if you were wrong? Write it down. The point of the next section is to compare.

WHAT HAPPENED

She took the fourth option and a part of the third, and the order in which she tried them is the lesson.

Constrained decoding came first, because it was free and immediate. The four-part structure is a grammar: four labelled sections in a fixed order. Enforced at sampling time, the structure went to 100 per cent on the quantised model, because invalid output was not merely unlikely, it was not available. What the constraint could not fix was the content inside the sections, and the common-mistake paragraph was often thin.

Then she retrained, once, with a modest change: three epochs instead of two, and a larger share of examples whose common-mistake paragraph was the substantial part. The structure was already guaranteed by the constraint, so training could spend its capacity on content rather than on shape.

The quantised model that shipped scores 0.88 on a teacher-rated content rubric against 0.91 for the fp16 version — a real cost, measured, written in the model card, and

accepted.

Three things changed in how the team works. The evaluation script now takes a path and a backend, so running it against the final quantised file costs nothing and happens on every build. The build target ends with the evaluation rather than with the conversion. And the constrained decoder was kept after retraining rather than removed, because it costs nothing to keep and a fine-tuned model still produces invalid output occasionally.

Explain the mechanism by which four-bit quantisation nearly erased this fine-tune, and why a 1.5B model was especially exposed.

Training moved the weights by a modest amount, and quantisation introduces an error that is roughly constant per weight. When that error is of comparable size to the change training made, the change is rounded away and the quantised tuned model behaves like the quantised base. A small model is more exposed because it has less redundancy to absorb the same per-weight error, so the proportional damage is larger — which is why anything under about 3B is usually better served at eight bits than at four.

Tanvi spent a day on the chat template before finding the real cause. What check, run in twenty minutes, would have pointed at it immediately?

The three-way comparison: quantised base, tuned fp16, and tuned quantised, on the same held-out set. Its shape names the culprit directly — if the tuned fp16 model scores well and the tuned quantised model scores near the quantised base, the loss happened in the conversion rather than in the training or the serving configuration. More generally, the rule it enforces is to evaluate the exact artefact you will serve, since merging, quantising, format conversion and engine choice each change the model and none of them announce it.

Constrained decoding was applied before retraining, and kept afterwards. Justify both decisions.

It was applied first because it is free, immediate and guaranteed: a state machine built from the required structure removes illegal continuations at every step, so invalid output cannot be produced rather than merely being rare. It was kept afterwards because the two are not alternatives — a fine-tuned model still produces invalid output occasionally, the constraint costs nothing per request, and with the shape guaranteed the training run could spend its capacity on the content inside the sections instead of on the format.

Check what stuck

6 questions, one right answer each. This one is for you, not for a record: answers and the reason behind each are at the back. If two questions from the same module go wrong, re-read that module before the exam rather than after.

- 1 An adapter is loaded against a base model repository that has been updated in place since training. What do you see?
 - A. It loads without error and is quietly worse.
 - B. A warning from the loader naming the revision difference.
 - C. Nothing at all, because an adapter records the base model's revision in its config and refuses to attach to a different one.
 - D. A shape mismatch error as the adapter attaches.
- 2 A tuned 1.5B model scores well in fp16 and behaves like the base model after conversion to four-bit GGUF. What happened?
 - A. The conversion dropped the adapter, which must be merged into the base first.
 - B. The chat template in the serving config differs from the one used in training.
 - C. The quantisation error is comparable to the adapter's delta, so the change was rounded away.
 - D. The quantiser removed the tokenizer's added tokens, so the model no longer recognises the control tokens it was trained to emit during generation.
- 3 You plan fifty per-customer adapters. Which cost do teams usually discover only after the serving layer is built?
 - A. The cold start when an adapter is first loaded from object storage.
 - B. Fifty evaluations — each adapter needs its own held-out set, canary and number.
 - C. The shared rank ceiling, which has to be chosen at training time and raises the memory cost for every adapter in the fleet as soon as one of them needs a high rank.
 - D. The 10 to 20 per cent throughput penalty against a merged model.

- 4 At two million requests a month, a self-hosted tuned 8B on two redundant cards costs about \$1,168 and a hosted small model about \$720. What does that show?
- A. That redundancy should be dropped in order to make the arithmetic work.
 - B. That hosted pricing is subsidised and will invert within a year, so the self-hosted option should be built now in anticipation of that change.
 - C. That self-hosting is always the cheaper option once volume reaches this level.
 - D. That the cost case for a fine-tune has to beat the cheap hosted model, not the frontier one.
- 5 Which unlabelled production signal is the strongest indication that an answer was bad?
- A. The rate at which users immediately rephrase and ask again.
 - B. The rate of 500 responses from the serving endpoint.
 - C. A fall in the number of tokens generated per request, measured weekly against the same period in the previous month.
 - D. A rise in the mean response length.
- 6 A fine-tune still performs well, but nobody can say what it was trained on and the pipeline has not been run in a year. What follows?
- A. Freeze it and stop evaluating it, since a model that nobody is able to rebuild should not be disturbed by further changes of any kind.
 - B. Keep it, since it performs and retiring a working model wastes the investment.
 - C. Retire it, because you have no way to establish that it is performing for the reasons you think.
 - D. Retrain it immediately on the current base, which restores the missing provenance.

EXAMINATION

Fine-Tuning, and When Not To — final examination

This paper covers the whole course: deciding whether to fine-tune at all, how a training step works, the parameter-efficient family and its memory arithmetic, building a dataset, the objectives beyond supervised fine-tuning, the mechanics of a run, adapting models that are not chat models, measuring whether it helped, and shipping and maintaining what you built. You will be given 25 questions drawn from a larger pool, so no two attempts are the same paper. The pass mark is 70 per cent. There is no timer — take as long as you need, and read each question twice. If you do not pass, you may retake after thirty minutes with fresh questions drawn from the pool, as many times as you like. A teacher can open the next course for you regardless of this result.

On the website a sitting is 25 questions drawn from this pool and the pass mark is 70%. There is no time limit, there and here. Passing opens the next course in the track.

1 1. Deciding whether to fine-tune at all

You paste the missing document into the prompt and the model answers correctly. What has that test established?

- A. That fine-tuning on document-and-answer pairs would work, since the model has now demonstrated it can use this document correctly when shown it.
- B. That the model has the ability and lacked only the information, so the lever is retrieval.
- C. That the model's capability is marginal and a larger base is needed.
- D. That the failure was a format problem hiding behind a knowledge one.

2 1. Deciding whether to fine-tune at all

Thirty failures sort mostly into the format bucket. Why try constrained decoding before fine-tuning for format?

- A. Because a grammar also improves the content inside the structure.
- B. Because a schema enforced at sampling time reduces the number of tokens generated, which is where the saving on a format-heavy task actually comes from.
- C. Because fine-tuning cannot change output shape at all.
- D. Because it makes invalid output impossible rather than merely rarer, and costs nothing to try.

3 1. Deciding whether to fine-tune at all

A complaint reads: "it ignores my instruction to be concise". Which cause is this most often, and why?

- A. Style, because the model has a strong prior about answer length and your one word is arguing with millions of examples.
- B. Instruction adherence, because the model is satisfying some of the constraints in the prompt while silently dropping the one that arrived earliest in it.
- C. Capability, because judging the right length is a reasoning task.
- D. Knowledge, because the model does not know what your product considers concise.

4 1. Deciding whether to fine-tune at all

Distillation replaces a \$16,000 monthly hosted bill with roughly \$500 of self-serving, against a one-off of a few hundred dollars of generation and days of somebody's attention. At what point does that stop being a good deal?

- A. When the teacher and student come from different model families.
- B. When the task is narrow enough that a prompt already works.
- C. At low volume, because the saving is per request and the fixed costs do not amortise.
- D. When the student is smaller than about three billion parameters, below which imitation of a larger model stops transferring anything useful at all.

5 1. Deciding whether to fine-tune at all

Before a fine-tune there are three separate permissions to check. Which three?

- A. The model licence, the cloud provider's terms, and your own privacy notice.
- B. Copyright, trademark, and your customer contracts.
- C. The dataset's own licence, the terms under which its contents were generated, and whether the resulting weights can be exported to another country.
- D. The model licence, the data licence, and the obligations that attach to what you produce.

6 1. Deciding whether to fine-tune at all

Which part of a fine-tune decision record is most often left out, and what does its absence cost?

- A. The diagnosis, without which nobody knows which of the six causes was found.
- B. The kill criterion, without which sunk cost decides when to stop.
- C. The owner's name, without which nobody can answer questions about the artefact later.
- D. The rebuild estimate, without which a base model upgrade arrives as a surprise rather than as a scheduled and budgeted piece of work.

7 1. Deciding whether to fine-tune at all

What is the clearest sign that you have built a generative system for a classification problem?

- A. The prompt contains few-shot examples.
- B. The model's output is short.
- C. You post-process the reply with a regular expression to extract the label.
- D. You evaluate the system with a person reading a sample of the outputs rather than with an automatic metric computed over the whole set.

8 2. How training actually works

Why does a sequence of 1,024 tokens contain 1,023 training examples without anybody labelling anything?

- A. Because the target at each position is the token at the next position, so the text supervises itself.
- B. Because the batch dimension is folded into the sequence dimension before the loss.
- C. Because the model is trained on every possible sub-sequence of the input, which for a window of that length gives one fewer example than there are tokens.
- D. Because the tokenizer emits one example per merge rule applied.

9 2. How training actually works

A run reports a loss of 0.9. What does that number say about the model's predictions?

- A. That it is 90 per cent accurate at predicting the next token.
- B. That it assigns about 41 per cent probability, on average, to the token that actually came next.
- C. That it is nine tenths of the way to convergence on this dataset.
- D. That its perplexity is 0.9, which is the same quantity expressed on a different scale and carries no additional information.

10 2. How training actually works

Validation loss is rising while your task metric keeps climbing. What is the likely mechanism, and what should you select on?

- A. The learning rate has decayed too far; select the earliest checkpoint.
- B. The validation set has drifted away from the training distribution, so neither number is usable and a fresh validation set has to be built before any checkpoint can be chosen.
- C. Leakage between the sets; select on the training loss instead.
- D. The model is becoming more confident, so it is penalised harder on the items it still gets wrong; select on the task metric.

11 2. How training actually works

What is warmup protecting against, mechanically?

- A. AdamW dividing by a variance estimate built from a single noisy gradient.
- B. The dataset being unshuffled at the start of an epoch.
- C. A cosine schedule that has been configured for more steps than the run will actually take, which makes the first steps disproportionately large.
- D. Numerical overflow in the first forward pass.

12 2. How training actually works

You set a cosine schedule for 1,000 steps and stop the run at step 400. What do you have?

- A. A model that cannot be evaluated at all, because the scheduler state has not reached the point at which the optimiser writes its final weights.
- B. A smaller version of the same model, trained on less data.
- C. A model caught mid-schedule with its learning rate still high, typically worse than one trained on a 400-step schedule.
- D. A model identical to one trained for 400 steps, because the schedule only affects the final steps.

13 2. How training actually works

If you write your own training loop with gradient accumulation, how must the loss be normalised, and why?

- A. By the sequence length, so that long and short examples are on the same scale.
- B. It does not need normalising, because the optimiser divides by the accumulation count internally before applying the update to the weights.
- C. By the number of micro-batches, so each contributes equally.
- D. By the total number of unmasked tokens across the accumulation window, because otherwise a micro-batch with 40 scored tokens counts the same as one with 900.

14 2. How training actually works

Your run is out of memory at sequence length 8,192. Which is the first thing to check, and what does it change?

- A. The batch size, which scales activation memory linearly.
- B. Whether flash attention is on, which turns the quadratic attention matrix from a memory term into a linear one.
- C. The optimiser, which at 8 bytes per trainable parameter is the largest term at long context.
- D. The number of gradient accumulation steps, since each one holds its own copy of the activations until the optimiser step is finally taken.

15 3. Full, LoRA and the parameter-efficient family

A rough rule for QLoRA is about 0.6 bytes per parameter for the base plus 3 to 6 GB of working memory. What does that predict for a 13B model?

- A. About 3 GB of base, because four-bit storage is a quarter of the two bytes a bf16 parameter would otherwise occupy in memory.
- B. About 8 GB of base, so comfortable on a 16 GB card.
- C. About 8 GB of base, so tight on 16 GB and comfortable on 24 GB.
- D. About 26 GB of base, so it needs a 48 GB card.

16 3. Full, LoRA and the parameter-efficient family

You are working at rank 64 and results disappoint. Which specific defect does rank-stabilised LoRA address?

- A. The alpha-over-r scaling, which damps the update as rank rises; dividing by the square root of r keeps its magnitude steadier.
- B. The coupling between a weight's magnitude and its direction.
- C. The fact that a higher-rank adapter takes proportionally longer to merge back into the base weights, which is what makes high-rank experiments expensive to evaluate.
- D. The adapter's tendency to overfit at high rank.

17 3. Full, LoRA and the parameter-efficient family

DoRA separates magnitude from direction. When is it worth its 20 to 30 per cent extra time per step?

- A. At high rank, where the decoupling has the most room to act.
- B. When you are memory-constrained to a low rank and the extra time is affordable.
- C. On any run where the dataset is above ten thousand examples.
- D. Whenever full fine-tuning is unaffordable, since DoRA is the only parameter-efficient method whose learning dynamics are equivalent to training every weight.

18 3. Full, LoRA and the parameter-efficient family

Soft prompts match full fine-tuning only as model size grows past roughly ten billion parameters. What is the mechanism behind that gap?

- A. Smaller models are trained on less data, so their embeddings are noisier.
- B. Prompt tuning optimises the virtual tokens through a small auxiliary network that is discarded afterwards, and that network needs a large model to be trained stably.
- C. Smaller models have shorter context windows, so the virtual tokens crowd out the input.
- D. A soft prompt can only steer computation the model already performs, while LoRA can change the computation.

19 3. Full, LoRA and the parameter-efficient family

IA³ trains about 0.01 per cent of a model's parameters. Which two properties make it genuinely useful rather than merely small?

- A. It merges exactly into the weights, and it suits very small datasets where LoRA has enough capacity to memorise.
- B. It trains faster than LoRA, and it can be composed with any number of other adapters.
- C. It works on encoder and decoder models alike, and it needs no target modules to be named.
- D. It preserves the base model's calibration exactly, and it is the only parameter-efficient method whose effect on refusal behaviour has been shown to be zero.

20 3. Full, LoRA and the parameter-efficient family

Why does full fine-tuning forget more than LoRA does, and what follows for how you run one?

- A. Because a full fine-tune must be run for more epochs to converge, and the additional passes over the same data are what erase the capabilities you were not training.
- B. Because full fine-tuning uses a higher learning rate, so lowering it removes the effect.
- C. Because every weight is free to move towards your objective, including the ones that encoded something else — so you need replay, fewer epochs and a canary at every checkpoint.
- D. Because the optimiser state is larger, which lets the update travel further per step.

21 3. Full, LoRA and the parameter-efficient family

On a free 16 GB T4, which configuration is the headline one that actually works?

- A. A 13B model with QLoRA at sequence 2,048.
- B. A 7B model fully fine-tuned with 8-bit AdamW, which brings the persistent state down to ten bytes per parameter and therefore inside the card.
- C. An 8B model with 16-bit LoRA at sequence 2,048.
- D. An 8B model with QLoRA at sequence 1,024, batch 1 with accumulation.

22 4. Building the dataset

What does one provenance column per example buy you six months later?

- A. A way to reconstruct the training order if the shuffle seed is lost.
- B. An answer, rather than an investigation, when somebody asks whether the model was trained on customer data and under what basis.
- C. A deduplication key that is stronger than a hash of the text.
- D. A licence for the resulting weights, since the column records which permissive sources contributed and in what proportion to the final file.

23 4. Building the dataset

Four hundred tickets differ only in an order number. Why is exact deduplication not enough?

- A. Because hashing normalised text cannot distinguish two tickets that differ only in punctuation, which is the case that matters most when the order numbers have been redacted.
- B. Because exact deduplication is too slow at that scale.
- C. Because for training purposes they are one example repeated four hundred times, and the model weights them accordingly.
- D. Because near-duplicates improve the loss curve and therefore mask other problems.

24 4. Building the dataset

Annotators correct model drafts rather than writing from a blank page. What is the risk, and which measurement detects it?

- A. Anchoring on a fluent wrong draft; track the edit rate and audit if acceptance is very high.
- B. Loss of style consistency; measure agreement between annotators.
- C. Contamination between the training and evaluation sets, which is detected by checking for thirteen-gram overlap between the drafts and the held-out items.
- D. Slower throughput; measure examples per hour.

25 4. Building the dataset

A style guide rule says: if the answer is not in the provided context, reply with one exact sentence. Why specify the words rather than the sentiment?

- A. Because a fixed sentence is easier for a regular expression to detect in production.
- B. Because a model given eleven variations on a sentiment learns eleven ways of declining, which is the same as having no reliable way.
- C. Because exact wording reduces the token count of every declined answer.
- D. Because the chat template requires a fixed string in order to mark the end of an assistant turn when the answer contains no substantive content at all.

26 4. Building the dataset

When generating a dataset from a teacher model, where should the inputs come from?

- A. From a public benchmark, so the result is comparable with published numbers.
- B. From a mixture of traffic and invention in equal parts, so that the resulting set covers both the cases you see today and the ones you expect to see next quarter.
- C. From the teacher, so that inputs and outputs are drawn from one consistent distribution.
- D. From your real traffic, because invented inputs are cleaner, shorter and more polite than anything your users type.

27 4. Building the dataset

Packing removes padding waste. What two things must be in the batch for it to be safe, and what is the conservative alternative?

- A. A block-diagonal attention mask and per-document position ids; otherwise use length-grouped batching.
- B. A shuffled dataset and a fixed seed; otherwise use a larger batch size.
- C. An end-of-sequence token between documents and a lower learning rate; otherwise raise the accumulation count.
- D. A separator token and a recomputed attention mask on every step; otherwise disable gradient checkpointing so that the activations for each document remain distinct.

28 4. Building the dataset

Why does deduplication count as a privacy control and not only as a quality filter?

- A. Because removing duplicates makes a canary test possible, and a canary test is the only way to establish that no personal data can be extracted from the trained weights.
- B. Because duplicate rows are usually the ones containing identifiers that redaction missed.
- C. Because the probability that a sequence is memorised rises sharply with how many times it is duplicated.
- D. Because a smaller dataset trains for fewer steps, and fewer steps means less memorisation overall.

29 5. Objectives beyond next-token prediction

Why do preference methods assume a model that has already been instruction-tuned?

- A. Because the reference model must be a different architecture from the policy.
- B. Because preference pairs are collected from production traffic, which only exists once a model has been deployed and has started producing the failures worth correcting.
- C. Because a base model has no chat template, so the pairs cannot be rendered.
- D. Because they are a refinement step, and applying them to a model with nothing coherent to refine produces disappointing results.

30 5. Objectives beyond next-token prediction

A DPO run produces a model that has become strange — repetitive, oddly formatted, worse at things you did not touch. Which parameter do you change first, and in which direction?

- A. The learning rate, upwards, so the model escapes the region it has fallen into.
- B. Beta, upwards, because it is the KL penalty anchoring the policy to its reference.
- C. Beta, downwards, so the model is freer to move away from the degraded state.
- D. The number of epochs, upwards, because one pass over a preference set is rarely enough to settle the behaviour that the early steps disrupted.

31 5. Objectives beyond next-token prediction

DPO needs the reference model's log-probabilities. Why does using LoRA make that free?

- A. Because PEFT keeps a frozen copy of every adapted matrix in order to support merging, and those copies can be reused as the reference without any additional memory.
- B. Because the adapter is small enough that a second copy costs little.
- C. Because the base model under the adapter is the reference, so disabling the adapter computes reference log-probabilities in place.
- D. Because DPO's implicit reward does not depend on a reference when the policy is an adapter.

32 5. Objectives beyond next-token prediction

ORPO folds supervised and preference learning into one stage. What does that simplification cost?

- A. Control: if the result is poor you have two hypotheses and a single experiment.
- B. Data, because it needs both demonstrations and preferences in far larger quantities.
- C. Compatibility, because a single-stage method cannot be applied to a model whose chat template marks assistant spans for masking.
- D. Memory, because the combined loss needs both models resident.

33 5. Objectives beyond next-token prediction

PPO holds four networks. Which is the one GRPO removes, and what did it do?

- A. The reward model, which scored generated responses.
- B. The value model, which estimated expected future reward so an advantage could be computed.
- C. The frozen reference, which supplied the KL penalty.
- D. The policy's own second copy, which PPO keeps in order to compute the ratio between the old and the updated policy over several epochs on one batch.

34 5. Objectives beyond next-token prediction

Which instrument most reliably detects that you are optimising the gap between your reward and your intent?

- A. The number of optimiser steps since the last checkpoint.
- B. The agreement rate between two model judges from different families, computed on the same sample of outputs at the start and at the end of the run.
- C. The training loss, which flattens when the policy stops improving.
- D. A held-out reward measuring the same intent differently, which you never train against.

35 5. Objectives beyond next-token prediction

A reward model trained on preference pairs outputs 3.2 for a response. What does that number mean?

- A. Only that it is above or below another response's score; the model learned an ordering with no absolute scale.
- B. That the response would be preferred 3.2 times out of every four comparisons.
- C. That it falls in the top band of the reward model's range, which is why a fixed threshold on the output is the standard way to gate responses in production.
- D. That the response is above average, since scores are standardised during training.

36 6. The mechanics of a run

What is the difference between how a conversation is rendered for training and for inference?

- A. Training uses the fast tokenizer; inference uses the slow one for exactness.
- B. Training renders each turn separately so that the mask can be built per turn, while inference renders the whole conversation at once for the cache to be reused.
- C. Training renders the full conversation including the assistant turn and its end token; inference stops at the assistant header.
- D. Training omits the system prompt; inference includes it.

37 6. The mechanics of a run

Why is setting the pad token equal to the end-of-sequence token a trap?

- A. Because the attention mask cannot represent two roles for one token id.
- B. Because the embedding row for that token then receives gradients from both the padded positions and the real endings, which pulls it in two directions at once during training.
- C. Because the model then generates padding in its replies.
- D. Because a mask built by matching on the pad id also masks the real end-of-sequence token at the end of every answer.

38 6. The mechanics of a run

You add a control token to the vocabulary and train a LoRA. What two things must you get right?

- A. Register it as a special token and lower the learning rate for the new row.
- B. Initialise the new row near the mean of the existing embeddings, and put the embedding table and the output head into the list of modules saved in full.
- C. Resize the embeddings to a multiple of 64, and tie the input and output matrices.
- D. Add the token to both the fast and the slow tokenizer, and confirm that the chat template renders it inside the assistant span rather than in the header.

39 6. The mechanics of a run

A LoRA run has a perfectly flat loss and a non-zero trainable parameter count. Which configuration line is the likely cause?

- A. A cosine schedule with warmup, which keeps the rate near zero for the first steps.
- B. Paged AdamW, which spills the optimiser state to CPU memory and therefore applies the update a step later than the gradients were computed.
- C. Gradient checkpointing left at the reentrant implementation, so gradients never reach the adapter.
- D. Length-grouped batching, which sorts the data and removes variation between batches.

40 6. The mechanics of a run

Loss drops to near zero within fifty steps. What is the most likely cause?

- A. The answer is visible in the input — the labels were not shifted, or the completion also appears in the prompt.
- B. The adapter rank is too high for the dataset size.
- C. The dataset is too small for the effective batch size, so the model is seeing the same examples repeatedly within a single epoch and memorising them.
- D. The learning rate is far too high.

41 6. The mechanics of a run

You move a config tuned on one card to four GPUs without changing it. What happens, and what is the fix?

- A. The run fails to launch until the accumulation count divides evenly by the device count.
- B. The effective batch quadruples, so divide the accumulation steps by four or raise the learning rate to match.
- C. Each GPU trains a separate copy, so you get four models to choose between.
- D. The gradient all-reduce averages across devices, which cancels the multiplication automatically and leaves the effective batch exactly as it was on one card.

42 6. The mechanics of a run

Which of these is needed to continue a run, but not to use the model?

- A. The adapter configuration recording rank and target modules.
- B. The evaluation results and the canary scores recorded at the checkpoint, which together establish which checkpoint the run should be resumed from.
- C. The tokenizer saved beside the weights.
- D. The optimiser state, the scheduler position and the RNG states.

43 7. Fine-tuning beyond a chat model

What does Matryoshka loss buy you on an index of several million vectors?

- A. Useful prefixes: search a 128-dimensional index and rerank with the full 768, cutting index memory several times over.
- B. Better recall on short queries, because early dimensions encode lexical information.
- C. Immunity to the hard-negative mining trap, because a truncated vector cannot distinguish two correct passages from each other in the first place.
- D. Faster training, because the loss is computed on truncated vectors.

44 7. Fine-tuning beyond a chat model

Why must a reranker's negatives be mined from your retriever's actual output?

- A. Because a cross-encoder reads the query and the passage together, and two randomly paired texts produce a sequence longer than the model's maximum input length.
- B. Because random negatives make the loss unstable.
- C. Because a reranker's whole job is to sort the candidates the first stage returns, so random negatives train it for a task it will never face.
- D. Because mined negatives are guaranteed to be incorrect, whereas random ones may not be.

45 7. Fine-tuning beyond a chat model

Why pass `id2label` and `label2id` when creating a classifier?

- A. Because without them the classification head is initialised randomly.
- B. Because the trainer uses them to stratify the evaluation split, which is what makes a macro-averaged F1 score comparable between runs on the same data.
- C. Because the loss function needs them to weight the classes.
- D. Because they are saved in the config, so the model that comes back off disk in six months says what its outputs mean.

46 7. Fine-tuning beyond a chat model

Before fine-tuning a speech model, which two free things should be tried?

- A. Splitting long recordings into thirty-second segments, and reporting word error rate both before and after the standard text normaliser has been applied.
- B. An initial prompt carrying your domain vocabulary, and a larger checkpoint.
- C. Resampling to 44.1 kHz, and normalising the transcripts.
- D. Freezing the decoder, and lowering the learning rate.

47 7. Fine-tuning beyond a chat model

What should you measure before planning memory for a vision-language fine-tune?

- A. The number of tiles the processor will produce for the largest image in the set, since that determines how many separate forward passes the encoder has to make per example.
- B. The size of the vision encoder relative to the language model.
- C. How many sequence positions one representative image actually consumes at your resolution.
- D. The ratio of images to text examples in the dataset.

48 7. Fine-tuning beyond a chat model

In an image LoRA, every training photograph has a white background and no caption mentions it. What happens?

- A. The white background is baked into the subject, so asking for a beach produces a suspiciously white patch.
- B. The model learns to remove backgrounds, which is why product LoRAs are trained this way.
- C. The adapter overfits on the background first, which is why regularisation images of the generic class are mixed in during training to counteract it.
- D. The background is ignored, because it carries no information the model can use.

49 7. Fine-tuning beyond a chat model

What is the strongest argument for an on-device small model, and what is the matching risk?

- A. Price, against the risk that a small model is simply less accurate on every input.
- B. Privacy, offline operation and zero marginal cost, against the risk that it collapses rather than degrades off-distribution.
- C. Latency alone, against the risk that quantisation removes the fine-tune.
- D. Vendor independence, against the risk that the model cannot be updated once it has shipped inside an application on somebody else's device.

50 8. Did it actually work

Which capability most often degrades first after an English-only fine-tune, and why do teams miss it?

- A. Multi-turn conversation, because the training data was single-turn.
- B. Refusal behaviour on abusive input, because the training data contains no abusive examples and therefore no demonstration of how to decline one.
- C. Code formatting, because the schema starts appearing inside code blocks.
- D. The other languages the product supports, because every metric being watched is computed in English.

51 8. Did it actually work

Why keep a locked evaluation set in addition to a development set?

- A. Because a set you have consulted forty times has become a training set for your own decisions.
- B. Because the locked set can be built from a different distribution and so tests generalisation.
- C. Because regulators expect a sealed set that has never been opened, which is the only evidence accepted that a model was not tuned against its own evaluation.
- D. Because a larger combined set gives a narrower interval.

52 8. Did it actually work

Which kind of contamination can you rule out, and which can you not?

- A. You can rule out neither, because n-gram matching misses paraphrase.
- B. You can rule out the pretraining corpus by searching for the benchmark's canary string, but not overlap within your own data, because synthetic examples are generated rather than copied.
- C. You can rule out overlap between your training and evaluation sets; you cannot rule out the base model's pretraining corpus containing a public benchmark.
- D. You can rule out both, using thirteen-gram overlap in each direction.

53 8. Did it actually work

The answer-ordering test compares the probability of a multiple-choice question's canonical option order against shuffled orders. What does a consistent preference for the canonical order indicate?

- A. That the options are of unequal length, which biases sequence probability.
- B. That the model is reasoning about the question, since a model that had memorised the answer would score every arrangement of the options identically.
- C. That the model is well calibrated on this question type.
- D. That the model has likely seen the benchmark in that form, which is evidence of memorisation.

54 8. Did it actually work

Why allow a tie in a pairwise comparison protocol?

- A. Because it halves the number of comparisons needed for a given interval.
- B. Because forcing a choice on genuinely equivalent answers manufactures signal from nothing, and a high tie rate is itself a finding.
- C. Because ties let you detect position bias without randomising the order.
- D. Because a three-way judgement is easier for annotators to make consistently than a two-way one, which is what raises the agreement rate between judges.

55 8. Did it actually work

What is the single most effective mitigation for safety regression caused by benign fine-tuning data?

- A. Running the refusal set at every checkpoint and selecting the last checkpoint at which the refusal rate has not yet fallen below the base model's.
- B. Lowering the learning rate until the regression disappears.
- C. Putting thirty to a hundred refusal examples, in your product's own wording, into the training set.
- D. Switching from full fine-tuning to a lower LoRA rank.

56 8. Did it actually work

A data learning curve reads 0.801 at 25 per cent, 0.847 at 50, 0.869 at 75 and 0.871 at 100, against a three-seed noise floor of 1.3 points. What decision does it settle?

- A. That another week of labelling more of the same data will not help, so the effort belongs elsewhere.
- B. That the dataset should be cut to 75 per cent, since the last quarter adds nothing.
- C. That the evaluation set is too small to resolve the difference between the last two points, so the curve should be re-run on a larger held-out set before anything is concluded.
- D. That the model has converged and training should stop at the current epoch.

57 9. Shipping it, and living with it

You intend to serve fifty adapters from one base. What must be decided at training time rather than at serving time?

- A. The quantisation format of the base model.
- B. One rank for the whole fleet, because the server's maximum rank sets the memory cost for every adapter.
- C. Which adapters will be hot, so they can be prewarmed.
- D. Whether the adapters will ever need to be merged, since a merged model cannot be served alongside unmerged ones from the same set of base weights.

58 9. Shipping it, and living with it

AWQ and GPTQ need calibration data. Why use your own rather than the default web text?

- A. Because your own data has the same token length distribution as production.
- B. Because a calibration set drawn from your own task lets the quantiser preserve the adapter's delta exactly, which is what prevents a fine-tune being rounded away.
- C. Because the default sets are usually too small to calibrate a large model.
- D. Because the method protects the weights that produce large activations on the calibration inputs, so generic text protects the wrong weights.

59 9. Shipping it, and living with it

A build manifest records the base model by repository name only. What breaks?

- A. The manifest resolves to different bytes over time, because model repositories are edited in place.
- B. The serving engine cannot locate the weights without a revision.
- C. The adapter refuses to load, because a name without a revision does not uniquely identify the architecture it was fitted against.
- D. Nothing, as long as the adapter and dataset are hashed.

60 9. Shipping it, and living with it

Under the major-minor-patch convention for model versions, what does a patch release mean?

- A. A hyperparameter was changed and the model retrained, which is why two patch releases of the same minor version can differ slightly on the held-out set.
- B. The dataset changed while the base stayed the same.
- C. The same weights in a different build: a new quantisation or serving format.
- D. The base model changed and everything must be re-evaluated.

61 9. Shipping it, and living with it

What does a weekly shadow run of your frozen held-out set catch that the unlabelled production proxies cannot?

- A. A rise in the rate at which users rephrase their questions.
- B. A gradual decline in answer quality caused by the evaluation set itself ageing relative to the traffic it was drawn from.
- C. A shift in the input distribution towards topics the model was not trained for.
- D. A configuration change, a silently updated dependency, or a quantisation applied during a deployment.

62 9. Shipping it, and living with it

Which section of a model card is hardest to write and most useful to a reader?

- A. Intended use, because it defines the product.
- B. Out of scope, because it requires saying no to uses somebody has already imagined.
- C. Evaluation, because it requires intervals and a statement of what judged the model.
- D. Known limitations, because each one has to be explained by the mechanism that causes it rather than by a general statement that the model may occasionally be inaccurate.

63 9. Shipping it, and living with it

Why does turning every production incident into a permanent evaluation item compound more than any other habit in the course?

- A. Because it replaces the frozen set, which by construction ages away from the distribution your system is currently serving to real users.
- B. Because it gives the team a record of what went wrong for post-incident review.
- C. Because the evaluation set gets better over time rather than staler, so no failure of a kind you have seen can reach production again.
- D. Because it increases the size of the held-out set, which narrows the interval on every future comparison.

Answers

Each entry gives the correct option, then what each wrong one reveals about how somebody was thinking. The second part is worth more than the first: a wrong answer here is a named misreading, not a mistake.

Lesson checks

1. The Honest Answer First — B

A — Those transcripts do teach house behaviour, but they encode the old policy, so you would train the wrong answer in more firmly.

C — It targets the exact gap, but you repeat the whole exercise at the next policy change, and the model still cannot cite a source.

D — Larger models are better calibrated, but no model of any size knows a policy your company wrote three weeks ago.

2. What Fine-Tuning Changes, and What It Cannot Add — A

B — More passes do increase memorisation, but they memorise your exact strings and increase forgetting long before the facts become reliably retrievable.

C — Rank genuinely bounds how much the adapter can change, but higher rank buys verbatim memorisation of your sentences rather than knowledge the model can use.

D — Bigger bases do absorb more from fine-tuning, but the same failure appears at 70B — the fix is putting the document in the context, not enlarging the container.

3. Diagnosing the failure before choosing the fix — A

B — Treats private company policy as general knowledge; no model of any size was pretrained on your internal document.

C — Confuses reading a fact in context with storing it in weights; the paste test shows the model can use the fact, not that gradient updates will reliably install it.

D — Mistakes the mechanism: what helped was the specific fact being present, not the prompt being longer or more emphatic.

4. Try the retriever first — B

A — Assumes weights are a storage medium for documents; raising rank increases capacity for behaviour, not a searchable copy of the corpus.

C — Treats a missing-evidence problem as a sample-size problem; more examples teach shape, not the contents of documents that were absent.

D — Names a real effect (grounding suppresses recall of parametric facts) but it is not what pins accuracy at exactly the recall figure.

5. When the answer is one of six labels — C

A — Inverts the usual relationship; small models are more prone to underfitting, and the encoder's advantage is not about overfitting at all.

B — Invents a pretraining corpus; the encoder's pretraining is generic web and book text, exactly like the decoder's.

D — Half-true about bidirectionality but wrong about length: encoders typically have shorter context limits than modern decoders.

6. Distillation, the case that usually survives — D

A — Blames dataset size for an effect caused by what the data contains; the same confidence appears with 5,000 unfiltered examples.

B — States a property of model size that does not hold; sharpness comes from the training objective and data, not from parameter count.

C — Assumes calibration tracks accuracy; a model can lose accuracy and gain or lose confidence independently, which is why they are measured separately.

7. The cost of owning a model — A

B — Confuses a format-compatibility annoyance, which is usually fixable, with the substantive reason the weights are meaningless on a different base.

C — Invents a licence term; the obstacle is mechanical, not legal, and would exist between two permissively licensed models.

D — Picks a genuine hazard that is neither necessary nor sufficient here; tokenizers are often unchanged within a family and the adapter would still not transfer.

8. Licences, and what you may train on — B

A — Misstates Apache-2.0, which is permissive and imposes no copyleft on derivatives; the confusion is with share-alike licences.

C — Asserts a clean legal break that no licence grants and no court has established; obligations are generally analysed as following the material used.

D — Over-generalises from the subset of datasets that are non-commercial; several widely used sets are explicitly permissive.

9. Writing the decision down — C

A — Correctly notes that parameters matter but misplaces the reason: the problem is the base-line moving over time, not the two systems differing in settings.

B — Reaches for a compliance justification for what is an experimental-design requirement that would hold with no governance regime at all.

D — Describes a useful side effect; the fallback would still exist without freezing, and freezing is motivated by measurement, not deployment.

10. One training step, taken apart — D

A — Misplaces both the timing and the trigger; moment buffers are allocated at the first optimiser step, and they do not spike at the end of a forward pass.

B — Imports an inference concept: the KV cache is used for generation, not for a training forward pass, which computes all positions at once.

C — Gets the size right but the timing wrong; gradients are populated during the backward pass, after the spike described.

11. Reading the loss number — A

B — Misunderstands `ignore_index`, which removes positions from the computation entirely rather than down-weighting them.

C — Reads the jump as a bug; the same direction of change is exactly what a correct switch produces, so the number alone cannot support this conclusion.

D — Confuses the reported loss with gradient scale; loss is computed before any optimiser step and is unaffected by it.

12. Optimisers, and what their state costs — B

A — Extends quantisation to a tensor it does not touch; gradients are produced by the backward pass at the compute dtype, unchanged by the optimiser choice.

C — Confuses a consequence you might choose to spend the saving on with the source of the saving itself.

D — Invents a shared storage arrangement; the optimiser keeps state about parameters, not a second copy of them.

13. Learning rate, warmup and the schedule — C

A — Attributes the gap to ordering noise, which would be symmetric and would not reliably favour the shorter-configured run.

B — Confuses resuming a run with evaluating a checkpoint; optimiser state is irrelevant to the quality of weights already written.

D — Inverts the timeline: warmup completes in the first few percent of steps, long before step 400.

14. Batch size, accumulation and the token count — D

A — Holds only when every micro-batch contains the same number of unmasked tokens, which masked instruction data specifically does not.

B — Names a real but tiny effect, unrelated to the systematic weighting error produced by variable token counts.

C — Describes a failure to average at all; the loop does average, so the magnitude is roughly right and the weighting is what is wrong.

15. Precision, and why bf16 ended a class of bugs — A

B — Conflates the loss scaler with a memory-management mechanism; the scaler adjusts a multiplier on the loss and frees nothing.

C — Inverts the cause: the scaler halves its scale in response to overflow to infinity, not underflow, and lower learning rates do not shrink gradients.

D — Describes a fallback that would raise memory use without involving the gradient scaler at all, which only exists on the fp16 path.

16. Where the memory actually goes — B

A — Invents a dependency on sequence length; optimiser state is per parameter and is identical at both lengths.

C — Dequantisation happens per matrix multiplication regardless of sequence length, and is transient rather than scaling with context.

D — Gradients are the same shape as the trainable parameters and do not depend on how many tokens produced them.

17. Epochs, overfitting and checkpoint selection — C

A — Applies the textbook rule mechanically; validation loss rises partly because the model has become more confident, which is not itself harm.

B — Treats the loss number as the objective and tunes to make it look right, rather than asking what the model is doing.

D — Attributes a systematic divergence between proxy and metric to sampling noise, which a larger validation set would not remove.

18. Regularisation that helps, and the kind that does not — D

A — Names a cost that is negligible in practice and would not explain a difference between classification and generation.

B — Invents an incompatibility; masking and smoothing operate on different axes and compose without difficulty.

C — Reasons that the effect vanishes at large vocabulary, when the damage comes from the 0.9 placed on the correct token rather than from the spread.

19. Full, LoRA, QLoRA, and the Memory Arithmetic — C

A — Batch size and checkpointing really do control activation memory, but optimizer state is about 12 bytes per parameter regardless of batch — roughly 84 GB for a 7B model.

B — The trainable side is indeed tiny, but the frozen base still occupies 14 GB in 16-bit, leaving almost nothing for activations — and a T4 has no bf16 at all.

D — Accumulation does simulate a large batch on small hardware, but it only splits activation memory; weights, gradients and optimizer state stay resident throughout.

20. What LoRA is actually doing — A

B — Asserts a property of zero initialisation that does not hold; the motivation is the product being zero, not the gradient being large.

C — Confuses initialisation with storage; optimiser state is allocated per parameter regardless of the values initialised into it.

D — Invokes a rank concern that does not arise; the product of a $d \times r$ and an $r \times k$ matrix has rank at most r however either is initialised.

21. Choosing rank, alpha and which modules to touch — B

A — Offers a plausible-sounding mechanism, but overfitting would show as a gain then a decline, not as no change at all.

C — Identifies a real scaling subtlety but attributes the whole result to it; the all-linear gain at the same alpha and rank is unexplained by this.

D — Correctly notes GQA shrinks two matrices but treats parameter count as the issue, when raising rank to 64 would have fixed a pure count problem and did not.

22. QLoRA: quantising the part you are not training — C

A — Misreads the target: the second pass acts on the scaling constants, and 2-bit weights would destroy quality rather than preserve it.

B — Would break training entirely, since the adapter is the only thing receiving gradients and must be updatable at a useful precision.

D — Names a different technique; QLoRA's activations stay in the bf16 compute dtype and are not what double quantisation touches.

23. DoRA, rsLoRA and the variant question — D

A — Casts rsLoRA as a regulariser; it changes the update's scaling, and if anything makes the higher-rank adapter move more, not less.

B — Treats intrinsic dimension as a hard cap that switches off gradients, when unused directions would simply stay near zero rather than block learning.

C — Attributes PiSSA's initialisation mechanism to rsLoRA, which changes only the scaling denominator and leaves initialisation untouched.

24. Soft prompts and prefixes — A

B — Picks a real cost of the method that is unrelated to model size and far too small to account for the quality gap.

C — Treats the parameter count of the prompt as the binding constraint; adding more virtual tokens does not close the gap, which it would if this were the cause.

D — Invokes embedding-space smoothness as a blocker, when soft prompts optimise successfully on small models — they simply reach a lower ceiling.

25. The tiny methods, and a gotcha in one of them — B

A — Describes the usual cause of a zero count, but BitFit selects by parameter name pattern rather than module list and did select something.

C — Invents a fallback behaviour; the method has no head requirement and the small count comes from the architecture's design, not a fallback.

D — Attributes the absence to quantisation, which leaves biases in higher precision and would not explain the result on an unquantised model.

26. Full fine-tuning, when it is genuinely right — C

A — Couples two independent settings; the precision configuration is applied regardless of which modules the wrap policy matches.

B — Invents a graceful degradation path; the failure is that no wrapping boundary is found at all, not a downgrade to a defined lesser stage.

D — Points at a real memory consumer, but activations would not account for the failure on a run whose persistent state alone exceeds one card.

27. Training on a free or nearly free budget — D

A — Assumes the storage format determines the compute path, when the compute dtype governs the dequantised matrix multiplications.

B — Both mixed-precision paths keep fp32 master weights for trainable parameters; the difference between them is not a master-copy requirement.

C — Invents a compatibility constraint on the artefact; adapter weights are dtype-convertible and carry no such restriction.

28. The Dataset Decides Everything — D

A — A large train-test gap is the classic overfitting signature, but here the held-out score is high too — both sets are simply the wrong distribution.

B — More data is the reflex, but forty thousand examples from the same generator produce the same mismatch and roughly the same 92%.

C — Contamination is a real problem for public benchmarks, but this set is private and freshly generated; the gap is distributional, not memorised.

29. Where the data legitimately comes from — A

B — Substitutes a length-mediated story for a direct one; the behaviour appears equally in short answers drawn from the same filtered set.

C — Correctly notes weights are poor at facts, but that would produce wrong answers on known topics rather than a loss of the ability to decline.

D — Treats a distributional bias as a volume problem; the same bias appears at any size if every example is an answered question.

30. The three formats, and the conversion trap — B

A — Identifies a genuine but trivial cost; a few header tokens cannot account for a large quality gap.

C — Focuses on a tokenisation detail of the separator rather than the fact that the model never saw this prompt structure in training.

D — Blames an incompatibility between templates; the same mismatch causes the failure even when the headers contain no special tokens at all.

31. Cleaning at scale — C

A — Treats textual distinctness as informational distinctness, when 600 near-identical rows contribute 600 times the gradient signal of one pattern.

B — Assumes the tokenizer is retrained on your data, which it is not; the vocabulary is fixed by the base model.

D — Confuses repetition within one document, which the heuristic measures, with repetition across documents, which it cannot see.

32. Writing the answers, and keeping them consistent — D

A — Assumes any observed behaviour becomes the default, when the model fits the frequency it saw rather than adopting the majority form.

B — Imagines the model blending the styles; next-token training reproduces the observed mixture rather than interpolating between its modes.

C — Reframes uncontrolled variance as augmentation, which would only hold if something in the input determined which style was correct.

33. Generating a dataset from a teacher — A

B — Blames output quality when the described failure is specifically distributional, and would persist even with a perfectly accurate teacher.

C — Invokes memorisation, which would show as poor performance on held-out synthetic items too, not just on production traffic.

D — Shifts to a user-behaviour explanation that would not show up as the model performing poorly on the task.

34. Mixtures and replay — B

A — Invents a saturation mechanism; momentum is an exponential average that adapts continuously and does not stop responding to new gradients.

C — Correctly notes within-batch homogeneity but draws the wrong conclusion; low variance does not raise the learning rate or cause divergence.

D — Assumes memorisation is irreversible, when the real asymmetry is that the schedule gives the later data much smaller steps.

35. Packing, document boundaries and long context — C

A — Expects a visible loss signal; the model can partially fit the spurious transitions, so the curve looks normal, which is what makes the bug dangerous.

B — Inverts the effect: position ids climb past normal values in this setup, so the model sees more of the position range rather than less.

D — Assumes an exclusion the mechanism does not perform; every unmasked position contributes to the loss regardless of which document it belongs to.

36. Multi-turn conversations and tool calls — D

A — Reverses the effect; the call messages are still trained on, so the model retains the calling behaviour it was shown.

B — Names a plausible-sounding shift in loss composition, but the distinctive failure is fabricated tool output, not degraded prose.

C — Assumes the results are memorisable constants, when they are per-request data that the model can only imitate the shape of.

37. Memorisation, and what you cannot take back — A

B — Inverts the relationship; more capacity makes memorisation easier, not harder.

C — Assumes memorisation is a function of step size; a lower rate over more steps reaches similar memorisation on repeated data.

D — Helps with identifiers but leaves the surrounding content intact, and consistent pseudonyms preserve exactly the linkage that makes re-identification possible.

38. The data card — B

A — Frames the value as evidentiary, when the counts are chiefly a diagnostic for explaining later model behaviour.

C — Treats counts as a substitute for the data itself; knowing 887 rows were removed does not tell you which ones.

D — Describes a genuine efficiency consideration that has nothing to do with why the counts belong in a document read months later.

39. What supervised fine-tuning cannot express — C

A — Overstates the limit; fine-tuning routinely changes pretrained habits, just not efficiently through demonstration alone.

B — Treats it as a volume problem, when the same 500 examples do change other behaviours in the same run.

D — Invents a positional gap in masking; completion-only masking covers the whole assistant turn including its first tokens.

40. Where preference data comes from — D

A — Reaches for volume; on-policy pairs at a fraction of that count routinely produce the targeted change.

B — Names a real bias, but it would have produced a visible change in length rather than little improvement of any kind.

C — Invokes a reference-model inconsistency; the reference scores whatever text it is given and does not care which model produced it.

41. DPO, from first principles — A

B — States the objective's indifference correctly and then draws the wrong practical conclusion, since the model's output quality depends on those absolute levels.

C — Blames configuration for a property of the objective that appears with a correctly frozen reference model.

D — Gets the direction of beta backwards; a high beta restrains movement away from the reference, which is what would prevent this drift.

42. KTO, ORPO, SimPO and choosing among them — B

A — Manufactures pairs across unrelated prompts, which breaks the assumption that both responses answer the same input.

C — Confuses removing the reference model and the separate stage with removing the pairing requirement, which ORPO still has.

D — Notes that SimPO scores a response without a reference model, but the objective still contrasts a chosen response against a rejected one.

43. RLHF and PPO, honestly — C

A — Attributes the KL term to a relationship between the two auxiliary models, when it measures the policy's distance from the frozen reference.

B — Expects KL to fall, when some movement away from the reference is the entire point; the warning is in the magnitude and the pairing with reward.

D — Inverts the arithmetic: the penalty reduces the objective, so it cannot be what is inflating the reported reward.

44. Verifiable rewards and GRPO — D

A — Correctly spots a division issue but misplaces it: binary rewards normalise fine whenever a group contains both outcomes.

B — Reaches for a hyperparameter when the problem is that no usable gradient direction exists to scale up.

C — Concludes binary rewards are incompatible, when they work well as soon as groups contain a mix of outcomes.

45. Reward hacking — A

B — Attributes the behaviour to pretraining memorisation, which would generalise to new inputs rather than fail specifically on them.

C — Describes stylistic mimicry, which would not produce near-perfect scores on tests that execute the code.

D — Invokes sparsity and chance, when passing a full test suite reliably is a targeted exploit rather than a coincidence.

46. Continued pretraining on raw text — B

A — Treats an objective mismatch as a magnitude problem; a lower rate slows the drift but does not change what the objective is teaching.

C — Inverts the volume guidance; 400 million is modest for domain adaptation and the forgetting would occur at any scale.

D — Reduces a behavioural change to a formatting one, when the template is applied at inference regardless and the model still does not follow instructions.

47. Training a head instead of generating tokens — C

A — Suggests a transformation would fix it, when applying a sigmoid to an uncalibrated ordering still yields no meaningful absolute threshold.

B — Identifies a maintenance issue rather than the structural one; the threshold would be unsound even on a fully converged model.

D — Names a real bias in reward models, but it is a separate problem from the score having no absolute scale at all.

48. Chat Templates and Loss Masking — B

A — High temperature does produce rambling, but greedy decoding would run past the end in exactly the same way; the model has no learned reason to stop.

C — More epochs do sharpen formatting, but no number of epochs teaches a token that never appears in the loss.

D — A stop string does suppress the symptom and you should set one anyway, but it leaves a model that cannot end its own turn, and every downstream consumer inherits the workaround.

49. Running a LoRA, and What It Costs — A

B — The trainer is correctly minimising exactly what you asked it to, but what it measures is reproduction of your data, and perfect reproduction is the failure itself.

C — More data does dilute memorisation, but these settings are several multiples off, and you would spend weeks fixing what three numbers fix today.

D — More capacity generalising better is the lesson people carry over from pretraining, but here extra capacity buys more memorisation and more forgetting.

50. The tokenizer is part of the model — D

A — Treats a missing stop signal as a sampling limit; a model that learned to emit an end token stops regardless of the ceiling.

B — Names a real cause of the same symptom, but the described setup points at masking rather than at packing.

C — Assumes reassigning a role rewrites an embedding; setting `pad_token` changes a configuration field, not the embedding table.

51. Adding tokens, and initialising them properly — A

B — Names a template concern that would affect encoding of training data, not the model's inability to produce an output it was never trained on.

C — Invents a structural exclusion; after resizing, the head produces a logit for every row including the new ones.

D — Confuses masking of prompt tokens with a blanket exclusion of special tokens, which no standard collator performs.

52. A working config, line by line — B

A — Padding is masked out of the loss, so it dilutes throughput rather than gradient direction, and would never flatten the curve entirely.

C — Changes batch composition and therefore step-to-step noise, but does not stop the loss moving across a whole run.

D — Attributes a total failure to learn to optimiser quantisation, which is block-scaled and represents small updates without difficulty.

53. Debugging a run — C

A — No model achieves exactly zero cross-entropy on real text; even perfect predictions leave a small positive loss.

B — A zero learning rate freezes the loss at its initial value, which would be a normal positive number rather than exactly zero.

D — Names a real bug that produces a near-zero loss after a few steps, but it falls from a normal starting value rather than being zero at step one.

54. What to log, and how to read it — D

A — Reads a rising norm as progress; gradients typically stabilise or shrink as a model fits, so sustained growth is the opposite signal.

B — Confuses the reported pre-clipping norm with the post-clipping value; the logged figure is what clipping is applied to.

C — Evaluation batches produce no gradients at all, so nothing about the eval set can move the training gradient norm.

55. Checkpointing and resuming — A

B — True for some sharded strategies but not the general reason; the described failure occurs even with a fully gathered state dict.

C — Overstates the consequence; data sharding adjusts to world size and the run still sees the full dataset, just in a different order.

D — Invents a reset behaviour; the scheduler restores its saved position, which is precisely why the mismatch matters.

56. Reproducibility, and the noise floor — B

A — Proposes picking the most favourable baseline, which biases the comparison in the opposite direction rather than fixing it.

C — Names a real confound in rank comparisons, but the stated numbers are within the seed spread regardless of how alpha was set.

D — Demands more precision on the variance estimate while ignoring that even a rough spread already exceeds the claimed effect.

57. Multi-GPU without tears — C

A — Misreads averaging as attenuation; averaging is what makes four devices equivalent to one larger batch, not a division of the step size.

B — Describes a synchronisation failure that DDP specifically prevents; gradients are all-reduced so every replica stays identical.

D — Assumes sharding is broken; the distributed sampler assigns disjoint slices per rank, which is exactly what makes the batch larger.

58. Fine-tuning an embedding model — D

A — Argues the hard negatives were too weak to matter, which cannot explain a decline relative to not using them at all.

B — Names a plausible-sounding self-reinforcement story, but mining from the current model is standard practice and works when false negatives are excluded.

C — Invents an incompatibility; multiple-negatives ranking loss accepts explicit negatives alongside in-batch ones by design.

59. Training a reranker — A

B — Treats a score as insufficient for ranking, when a score per candidate is exactly what ranking needs.

C — Notes a real constraint on long passages that applies equally when reranking, and is not why first-stage retrieval is infeasible.

D — Raises a calibration point that is true and irrelevant, since ranking within a single query needs no cross-query comparability.

60. Fine-tuning a classifier properly — B

A — Invents an exclusion rule; accuracy counts every item, which is exactly how the rare class's errors get absorbed.

C — Attributes the gap to an evaluation-split mistake rather than to what the metric measures, and would show up as a train-test gap instead.

D — Describes a real edge case in macro F1 computation, which is a different metric from the accuracy figure being questioned.

61. Fine-tuning a speech model — C

A — Names a genuine confound, but one the team would have controlled, and it does not make the two numbers incomparable in kind.

B — Confuses WER with token-level loss; word error rate is computed on decoded words and is tokenizer-independent.

D — Raises a generalisation concern that would need a held-out set to settle, but does not make the two numbers different measurements.

62. Fine-tuning a vision-language model — D

A — Points at an alignment dynamic that is real but secondary; the damage persists even where the projector tracks the encoder successfully.

B — Names a genuine consequence of unfreezing, but batch size would degrade training broadly rather than specifically on unlike images.

C — Raises a tuning detail that mitigates the problem at best; a lower rate slows the overwriting rather than preventing it.

63. LoRA for image models, and the likeness question — A

B — Blames sample count for an effect that persists at 200 images when the backdrop is still uncaptioned.

C — Frames it as a capacity limit; raising the rank without changing the captions reproduces the same backdrop, often more strongly.

D — Names a real technique for a different failure — concept bleed across a category — rather than the caption-driven binding described here.

64. Getting structure without training for it — B

A — Assumes accuracy is reachable through the schema, when a stricter schema still cannot distinguish a correct queue from a plausible one.

C — Assumes the invalid outputs were scored as wrong answers, which would make the unchanged accuracy a contradiction rather than an observation.

D — Invokes the contested quality-degradation argument to explain a flat result, when the accuracy was simply unaffected in either direction.

65. Small models, where fine-tuning matters most — C

A — Reaches for truncation, which would not apply to a short question and does not explain confident misclassification.

B — Attributes the behaviour to quantisation, which affects a full-precision model of the same size identically in this respect.

D — Posits a capacity threshold for refusal; small models can be trained to decline, and this one simply was not shown any examples of it.

66. Tabular and time series: the honest answer — D

A — Unifies the pipeline at the cost of using the weakest available method on the numeric columns, which are most of the signal.

B — Reframes a classification problem as forecasting; these models predict future values of a series, not a binary outcome from mixed features.

C — Correctly prefers trees for the tabular part but discards a signal that embeddings make available to them as ordinary features.

67. Did It Help, and What Did It Break — C

A — The direction is genuinely more likely than not, but the 95% interval at $n=100$ comfortably covers 50%, which means it also covers no difference at all.

B — Setting a bar before you look is good practice, but a threshold narrower than your error bar rejects real improvements as readily as it rejects noise.

D — Judge disagreement is real and worth measuring, but re-judging the same items does nothing about sampling error across items, which is what dominates here.

68. The baseline you must beat — A

B — Proposes an interesting experiment that is usually impossible and does not address whether the project was needed.

C — Names a necessary sanity check that the team implicitly has, and which cannot invalidate the project on its own.

D — Raises a serving detail that belongs in the cost column rather than a comparison that could change the decision.

69. A held-out set that actually holds — B

A — A thousand items gives an interval of about 2 points, so sampling error cannot account for a 15-point gap.

C — Names a real cause of train-production gaps, but the diagnostics given point at leakage within the existing data rather than at drift.

D — Treats a matched label distribution as a defect, when a distribution matching production is what you want; the problem is shared content, not shared proportions.

70. Pairwise human comparison, done properly — C

A — Proposes a reasonable consistency check, but discarding judges who disagree removes exactly the signal that agreement rates are meant to expose.

B — Misapplies the sample-size guidance; 400 is ample for a 60% rate, and 1,500 is what a 55% rate would require.

D — Names a genuine effect on precision that does not threaten the validity of the comparison the way an uncontrolled length difference does.

71. Judging a model you distilled from the judge — D

A — Invokes a compression-improves-generalisation story that does not hold for a student capped by its teacher on the sampled distribution.

B — Would be a reasonable explanation if filtering had been described, and even then the teacher-as-judge makes the measurement unable to establish it.

C — Names a real bias that randomising order would have removed, and which does not explain a systematic preference for one system's characteristics.

72. Contamination you can check, and contamination you cannot — A

- B — Formatting is unchanged by shuffling the options; only their arrangement differs, so a formatting preference cannot explain the gap.
- C — Names a real bias, but position bias would affect canonical and shuffled orderings alike rather than distinguishing them.
- D — Assumes option order carries meaning the model exploits legitimately, which would not produce a consistent gap across unrelated questions.

73. Safety regression, including from benign data — B

- A — Assumes the data contained borderline material; the effect appears even when every example is unambiguously in scope.
- C — Names a factor that changes the magnitude of the regression rather than its cause; the effect persists at conservative learning rates.
- D — Reduces alignment to a prompt-following behaviour, when the degradation persists with the original system prompt present at inference.

74. Calibration after tuning — C

- A — Describes a plausible accuracy shift but does not explain why the routing rule stopped triggering on those same hard cases.
- B — Attributes the shift to a class-count dependency that does not apply when the label set is unchanged.
- D — Inverts what temperature scaling does and when it is applied; it is fitted after training and would reduce rather than inflate confidence.

75. Ablations that earn their cost — D

- A — Reads the whole curve as a trend and ignores that the final step is well inside the noise floor.
- B — Acts on the flat segment in the wrong direction; removing data that is not harming anything buys a little compute at a real risk.
- C — Demands uniform rigour where the early steps clearly exceed the noise; the 4.6-point rise from 50% to 75% is interpretable as it stands.

76. Serving the Adapter — D

- A — Quantization is a reasonable suspect and Q8 may well help, but you would be swapping parts without a measurement and still would not know what moved.
- B — Rank is a real quality lever, but nothing here implicates training — the tuned style survived, so the adapter is present and working.
- C — Reverting to a known-good configuration is sound instinct during an incident, but it forfeits the CPU deployment you needed and leaves the cause unmeasured.

77. Quantising the model you ship — A

B — Bad calibration would degrade the base model too, which held steady at 0.70 across both precisions.

C — Imagines a mechanism that reverts weights; merging writes the deltas into the matrix and conversion has no notion of what came from an adapter.

D — Asserts a fixed penalty that does not match the base model's unchanged score in the same experiment.

78. Serving many adapters from one base — B

A — Concludes incompatibility, when the server accommodates it by raising the shared ceiling rather than rejecting it.

C — Localises the cost to one adapter; the per-adapter matmuls do scale with rank, but the memory ceiling is shared across all slots.

D — Assumes batching requires uniform rank, when implementations pad to the configured ceiling precisely so mixed ranks can batch.

79. The cost at steady state — C

A — Training compute for a LoRA is a few pounds; it is the human cost around it, not the GPU time, that dominates.

B — Names a real recurring cost that is periodic rather than monthly and much smaller than continuous operational effort.

D — Picks a genuine but trivial line item; a few gigabytes per restart is negligible against the other terms.

80. Monitoring a model that fails silently — D

A — Latency tracks length and load rather than correctness, and a confidently wrong short answer is fast.

B — Detects format breakage well but is pinned at 100% valid under constrained decoding, and says nothing about whether the content is right.

C — Is a real signal but an extremely sparse one that most dissatisfied users never produce, so it moves late and noisily.

81. The day the base model moves — A

B — Invents a check that merging does not perform either; merging adds deltas to whatever matrix it is given.

C — Assumes an integrity field that adapter configs record as a name rather than verify as a hash.

D — Names a real hazard in releases, but a template change would be fixable while the weight mismatch would not be.

82. Versioning so a model can be explained later — B

A — Invents a storage-durability problem; object stores preserve bytes, and the gap is in knowing which bytes were used.

C — Dates order artefacts adequately; the failure is not in ordering them but in linking one to a request and to its inputs.

D — Raises a real format-compatibility risk that a registry does not solve either, and that is separate from identifying what ran.

83. Model cards and disclosure — C

A — Is true of every model ever released, so it distinguishes nothing and supports no decision.

B — Restates the evaluation section as a limitation, which tells a reader nothing about where the model fails.

D — Gestures at real sensitivities without naming any of them specifically enough to act on.

84. Retiring a fine-tune — D

A — Attributes a property to the weights that they do not have; stored weights are fixed and do not drift on their own.

B — Invents a renewal obligation that model licences do not impose on already-trained derivatives.

C — Misapplies contamination, which concerns overlap between training data and evaluation items rather than the age of a model.

Module test — 1. Deciding whether to fine-tune at all**1. A**

The examples worked, so the behaviour is learnable from demonstrations rather than missing knowledge or absent capability. What fails is delivery: the instruction that produces the behaviour does not fit your context or latency budget at volume. Fine-tuning moves that behaviour into the weights and shortens the prompt, which is precisely the case it is for.

2. C

A fact appears a handful of times against a pretraining corpus of trillions of tokens, so gradient descent takes the cheaper route and learns the form of a confident answer rather than the fact itself. Every example demonstrates the same lesson — internal question, specific confident reply — and at inference the model applies it to things it was never told.

3. B

A passage that was never in the context window cannot be used, so recall at k is a hard ceiling on how often the system can be right for the right reason. Eleven points sit below that ceiling and thirty-eight sit above it. Fine-tuning the generator can only work on the eleven, and training it to answer without the evidence teaches exactly the confident guessing you are trying to stop.

4. D

The encoder reads the whole input bidirectionally and is optimised directly for the decision, while the decoder is optimised to continue text and is being asked to decide indirectly. Above roughly a thousand labels the encoder usually wins on accuracy as well as on cost and latency. The trade is real: you need labels, a new class means a retrain, and the output is a label and a probability with no explanation attached.

5. A

Supervised fine-tuning maximises the probability of the answers in front of it and has no way to represent that one of them is wrong, so the teacher's mistakes arrive as gold. The student also learns the teacher's words rather than its uncertainty, so it reproduces those mistakes more flatly. The filter between teacher and student — a verifier, rejection sampling, a hundred pairs read by a person — is where the quality actually comes from.

6. C

A LoRA is a set of deltas for one exact set of base weights and does not transfer; the model was always the perishable output. What transfers is the pipeline — the dataset with its filters, the configuration you arrived at, the held-out and canary sets — so a team that can rebuild in an afternoon treats a new base as free improvement, and a team that cannot treats it as a threat.

Module test — 2. How training actually works

1. B

The ignore index removes a position from the loss entirely — not de-emphasised, absent — so no gradient comes from it. The forward pass still runs over it, so nothing is saved there. The consequence is the dangerous one: a wrong mask changes the task the model is learning while the loss curve looks perfectly healthy throughout.

2. D

Loss is an average over scored tokens. A run that scores the prompt as well averages in a system message identical in every example, which the model predicts almost perfectly after a few dozen steps. Removing those easy positions raises the average and nothing has got worse — one of four reasons a loss number is comparable only with your own runs.

3. A

AdamW keeps two fp32 running averages per trainable parameter, 8 bytes in all. On a full 7B that is 56 GB; on a 20-million-parameter adapter it is 160 MB. Freezing 99 per cent of the model removes 99 per cent of both the gradient and the optimiser cost, which is the whole of parameter-efficient fine-tuning in one line.

4. C

B is initialised to zeros, so BA is zero at step one and the model is bit-for-bit the base model. There is no random perturbation to recover from, and the adapter has to learn a useful correction from nothing in a few thousand steps with a small number of parameters. Full fine-tuning starts from weights that already encode a great deal, and moving them fast is how pretrained behaviour is destroyed.

5. B

Gradient noise tracks tokens rather than examples, and a typical small fine-tune sits somewhere between 8,000 and 64,000 tokens per optimiser step. At 900 the gradient is an extremely noisy estimate and no learning rate will make the run stable. Raising accumulation costs wall-clock time per step and nothing else.

6. D

A T4 is a Turing card with no bf16, so training runs in fp16 with loss scaling. When the scale is too large some gradients exceed 65,504 and become infinity; the framework detects it, throws that step away and halves the scale. An occasional reduction is normal. Repeated ones mean a meaningful share of your steps are being discarded, and on a free card the remedy is a lower learning rate rather than anything cleverer.

Module test — 3. Full, LoRA and the parameter-efficient family

1. A

Four bytes of fp32 master weights, four of fp32 gradients and eight for AdamW's two per-parameter moments, with the bf16 working copies on top. Activations are not in that figure because they depend on batch and sequence length rather than on parameter count, which is why lowering the batch size so often fails to save a memory-bound full fine-tune.

2. C

QLoRA is a memory technique. Every matrix multiplication dequantises its weights to bf16 on the fly, which costs roughly a quarter to two fifths of throughput, and the adapter is then fitted against a slightly wrong model. Paying that for memory you already had is a straight loss. The decision is arithmetic: parameters times two bytes, plus activations, plus logits, against the card.

3. B

Which modules you adapt matters more than the rank you pick. At a similar parameter count, adapting all seven projections at rank sixteen beats adapting two at a high rank — and the feed-forward matrices are the larger ones, where a good deal of task-specific behaviour appears to live. The query-and-value convention outlived the evidence that produced it.

4. D

Doubling alpha at a fixed learning rate has an effect very like doubling the learning rate, so a sweep over both produces a confounded experiment from which you will draw a confident wrong conclusion. Fix alpha at twice the rank, tune the learning rate. The one genuine exception is high rank, where the alpha-over-r divisor damps the update and rank-stabilised LoRA divides by the square root of r instead.

5. A

A module-name list copied from a blog post about a different architecture matches nothing, PEFT attaches nothing, and the run trains nothing for three hours while producing a perfectly plausible loss curve. Printing the trainable parameter count before every launch catches it in one line, and a target-modules setting of all-linear survives a change of model family where a copied list does not.

6. C

BitFit was a real and surprising result on encoder models such as BERT and DeBERTa, which have biases to train. Llama-family models use RMSNorm and bias-free projections, so running it there finds essentially the normalisation scales and goes almost nowhere. It is the clearest argument for printing the trainable parameter count rather than trusting that a named method is doing something.

Module test — 4. Building the dataset

1. B

Writing the ideal output for a hundred real inputs is the act of specifying the task, and it reliably uncovers cases nobody had decided. Doing it after you have seen the model's output means specifying the task to fit the output, which is how a project ends up unable to answer the only question that matters: whether the work helped.

2. D

The model learns the whole distribution of your outputs, including the parts you were not thinking about. Length, opening phrases, whether uncertainty is expressed and in what words are all reproduced faithfully. That is why the advice is to vary the inputs hard — length, register, typos, hostility, other languages — and keep the outputs consistent in the ways you care about.

3. A

Either separator is fine; using different ones in the two places is not. The model was fitted to one prompt shape and is served another, which is a systematic mismatch on every request. The rule that removes the whole class of bug is that the code formatting training examples and the code formatting inference requests must be the same function, imported from one module.

4. C

Truncation removes the end of the assistant's turn, and the end-of-turn token lives there. Train on a few hundred examples with no ending and the model receives no gradient teaching it to stop, so at inference it answers, then writes a new question and answers that one too. Count how many examples exceed the limit before you train rather than letting the tokenizer decide.

5. B

Gradient descent moves weights towards whatever reduces loss on the batch. If every batch is triage, every weight is free to specialise and the capabilities you were not training drift. If a fifth of each batch is ordinary conversation, the weights supporting ordinary conversation are still being held in place. It costs dataset assembly and nothing else, and in measured comparisons it beats every configuration knob combined.

6. D

The tool message is the world's output, not the model's, so it is input and is masked. Train on it and the model learns to predict what the tool will say — at inference it will sometimes produce a plausible fabricated tool result rather than calling anything. The assistant's tool-call message, by contrast, is exactly the behaviour you are teaching and must not be masked.

Module test — 5. Objectives beyond next-token prediction

1. A

SFT maximises the probability of a single written answer, so the only way to address a failure is to show many examples of the good behaviour and hope the bad one is crowded out. You cannot point at the failure. A preference pair can, because its signal is relative — raise the chosen response, lower the rejected one — which is what lets you express a contrast rather than a demonstration.

2. C

Preference learning is valuable precisely because the rejected side can be the model's own failure, which fixes SFT's off-policy problem. A rejected response from a different model or a much older checkpoint is a failure this model was not going to produce, so the pressure to push it down is spent on text the model already found unlikely — and the neighbourhood around it goes down too.

3. B

The loss falls as the margin grows and does not care how the margin is achieved, so reducing the probability of the chosen response is acceptable to it provided the rejected one is reduced faster. Taken far enough the model assigns low probability to good answers and produces something neither chosen nor rejected. Watch rewards/chosen, stop early, add a small likelihood term on the chosen response, or raise beta.

4. D

Thumbs data is binary and unpaired: you have judgements on single responses, not comparisons between two. KTO is designed for exactly that, with weights you can adjust when the positive and negative counts are imbalanced. Manufacturing pairs by matching a good response to an unrelated bad one produces pairs where you cannot say in one clause why the chosen side is better, which is noise.

5. A

GRPO computes each response's advantage as its reward minus the group mean, divided by the group standard deviation. If all G responses score the same, the standard deviation is zero and the advantage is zero or undefined for the whole group. Early in training a pure correctness reward is zero for every sample, which is why a small shaping term that is achievable from the start keeps groups differentiated.

6. C

Continued pretraining is measured in tokens by the hundred million: roughly 100 million to 1 billion for a light register shift, 1 to 10 billion for a genuine domain, and 10 billion and up for a new language. Fifty million shifts the texture slightly. Knowing that before you start saves a month — and whatever the volume, facts about specific documents still belong in the context window.

Module test — 6. The mechanics of a run

1. B

When the collator locates the assistant span by string matching, a response template wrong by one character matches nothing, every position is masked, and the loss is identically zero from step one. That symptom — a loss of exactly 0.0 rather than merely small — is nearly always this bug, and the decode check finds it in five seconds before any compute is spent.

2. D

Memorising eight examples in a hundred steps is the one thing a model with billions of parameters is certain to manage. If it cannot, the fault is in the code — masking, the template, the trainable parameter count, gradients not reaching the adapter. The test costs two minutes and separates 'my data is not good enough' from 'my code does not work', which are otherwise indistinguishable and are debugged completely differently.

3. A

The template usually adds the beginning-of-sequence token itself, and the tokenizer adds another by default. A model trained with two and served with one — or the reverse — has a systematic mismatch at the very first position, which everything downstream conditions on. Turn off the option that adds special tokens after templating, and print the first three token ids before every run.

4. C

Gradient norm is the most under-used signal in fine-tuning, because a steadily climbing curve is visible hundreds of steps before the loss moves. Isolated tall spikes are individual pathological examples and are harmless when clipping is on. A steady climb is instability: the learning rate is too high and the run will diverge if it is left alone.

5. B

Effective batch is per-device batch times accumulation times device count, so halving the devices halves it. That changes the number of optimiser steps per epoch, which puts the restored scheduler at the wrong point in its decay. The run completes and it is not the run you started. The same applies to re-shuffling the data, because resuming replays the shuffle from the saved RNG state.

6. D

Those three runs are your noise floor: a range of 2.6 points and a standard deviation of about 1.3. An improvement smaller than roughly twice that spread is a coin flip that has been given a name. Until you have measured the floor, every ablation is uninterpretable — and the floor costs three extra runs, which is a few hours.

Module test — 7. Fine-tuning beyond a chat model

1. A

For each anchor, its own positive must outscore every other positive in the batch. A batch of 16 asks the model to pick the right passage out of 16; a batch of 128 asks it to pick out of 128, which is a far harder task and produces a far better model. Gradient caching computes the gradient of an enormous batch with the memory of a small one, which is why 512 is reachable on one modest card.

2. C

Your corpus probably has several passages about refunds, so a passage ranked second for a refund question is frequently a correct answer you did not label. Labelling it a negative trains the model to separate two things that belong together. Mining from rank ten to fifty and capping similarity avoids most of it, and reading twenty mined negatives by hand before training on fifty thousand catches the rest.

3. B

A bi-encoder embeds the corpus once and compares vectors, so search is fast. A cross-encoder reads query and passage together, which is why it ranks far better and why its score cannot be precomputed. Fifty pairs take tens of milliseconds on a GPU; two hundred thousand would take minutes. The two-stage design buys cross-encoder quality at bi-encoder latency.

4. D

Accuracy is dominated by the majority class and will report a useless model as a good one. Macro-averaged F1 weights every class equally, so the three per cent class counts as much as the sixty-two per cent one. Print the per-class report and read the confusion matrix rather than the summary line: two labels the model mixes up are usually two labels your annotators mixed up.

5. A

The decoder sets the logits of every illegal continuation to negative infinity, so the output parses, the enum value is one of yours and the required fields are present. Validity and correctness are different properties. Keep the constraint as a floor even after fine-tuning — it costs nothing per request — and diagnose wrong values inside a valid structure as the knowledge or capability problem they are.

6. C

Gradient-boosted trees remain the strongest default on tabular data, tested repeatedly against deep learning with consistent results, because axis-aligned splits suit heterogeneous unordered features and are robust to noise and scale. The language model's genuine jobs are around it: embedding a complaint field into features, resolving entities, normalising inconsistent categories.

Module test — 8. Did it actually work

1. B

Near a fifty per cent rate the ninety-five per cent margin is roughly one over the square root of n , so a hundred comparisons carries about plus or minus ten points. Fifty-eight is inside it. Detecting a small difference is expensive — a 55 per cent win rate needs around 1,500 comparisons — and the right response is usually to decide that a difference that small does not change anything.

2. D

Nobody was going to keep paying frontier prices; they were going to try a cheaper model with a better prompt. That is the real alternative, so it is the baseline. Put every candidate on one table with quality and cost per thousand requests, and if the mid-tier row with a good prompt matches your tuned model, the project was unnecessary — which you can only find out by running it.

3. A

A random split puts near-twins of training items into the evaluation set, and puts the same customers on both sides so the model has learned their vocabulary and their products. Both inflate the score for reasons that do not generalise. The grouped, time-cut number is lower and it is the one that predicts production, because production always runs on the future and on people the model has not met.

4. C

Longer answers win more often than their quality warrants, in human judging as much as in model judging, and this is the most documented effect in the area. Reporting mean length beside the win rate names the risk; recomputing the rate on length-matched items measures it. If the rate collapses towards fifty, you measured verbosity.

5. B

A teacher's judgement of quality is substantially a judgement of resemblance to its own output, and the student was trained to resemble it — so self-preference bias is at its strongest in exactly this configuration. The student's more formulaic output can also score above the teacher's more varied one on a rubric that weighs style. Use a verifier, a human sample, or a judge from a different family.

6. D

Maximising the likelihood of a target sharpens the output distribution, so the model becomes more confident about everything including what it gets wrong; preference tuning compounds it, because judges reward confidence. A 0.9 that used to mean 0.9 now means something else, so every threshold, fallback and abstention rule downstream is firing at a rate nobody has measured. Refit temperature scaling on the exact artefact you serve.

Module test — 9. Shipping it, and living with it**1. A**

A LoRA is a set of deltas fitted against one exact set of base weights. Apply it to different weights — a new release, or the same release re-uploaded with a tokenizer fix — and the shapes still match, so it attaches and generates. There is no exception to catch and no warning, only a regression that looks like a quality problem. Pin the revision hash yourself and record it beside the adapter.

2. C

Training changed the weights by a small amount and four-bit quantisation introduces an error that is roughly constant per weight. When the two are of comparable size, the change is rounded away. A small model is especially exposed because it has less redundancy to absorb that error, which is why anything under about 3B is better served at eight bits. The three-way base, tuned, tuned-and-quantised comparison names it in twenty minutes.

3. B

Fifty adapters is fifty models. Without a shared harness and a summary that flags any tenant whose metric has fallen since the last build, nobody looks at forty-nine of the fifty and a regression affecting one customer arrives as a support ticket. Budget for it before committing — and check first whether one adapter trained on everyone's data, with the customer named in the prompt, would have matched them.

4. D

The fine-tune was justified against a frontier bill of about \$16,200, but nobody was going to keep paying that. Against the real alternative the self-hosted tuned model is dearer at this volume, before anybody is paid to maintain it — and idle capacity, redundancy and the engineer are the lines that never get written down. What overrides the arithmetic is privacy, an air gap, a latency budget or a per-request cost of zero.

5. A

A drifted model returns a fluent, confident, wrong answer with a 200 status, so uptime monitoring sees nothing. What users do next is the closest free thing to a quality label, and rephrasing is the strongest of those signals: the person read the answer, found it did not help, and asked again in different words. Escalation and abandonment are the next two.

6. C

A model nobody can rebuild or explain is not an asset, it is a liability that currently performs. You cannot answer what it was trained on, cannot debug a specific complaint, and cannot migrate it when the base moves. Retraining does not restore provenance that was never recorded. Retire it, keep the decision record and the cards, and rebuild deliberately if the capability is still wanted.

Fine-Tuning, and When Not To — final examination

1. B 1. *Deciding whether to fine-tune at all*

The paste test separates knowledge from capability, which is the commonest mislabel in the whole diagnosis. A model that invents your refund window looks stupid, and stupidity feels like a capability problem — but if the answer is right once the document is present, ability was never missing. Weights are a poor filing cabinet: facts belong in the context window, where they can be cited, updated and deleted.

2. D 1. *Deciding whether to fine-tune at all*

A constrained decoder maintains a state machine from your schema and sets the logit of every illegal continuation to negative infinity, so invalid output is not produced — not because the model learned not to, but because the choice was removed. Fine-tuning for format works and only makes bad output rarer. Keep the constraint even after training, because it costs nothing per request.

3. A 1. *Deciding whether to fine-tune at all*

Adherence is the pattern where satisfaction falls as the number of constraints rises, which you can plot. A single, consistently ignored preference about length is usually style: the model's prior about how long an answer should be is enormous, and one adjective cannot outweigh it. The fix is examples, and then fine-tuning if the examples are too long or subtle to sit in the prompt.

4. C 1. *Deciding whether to fine-tune at all*

The case strengthens linearly with traffic because the saving is per request, while the dataset, the evaluation and the maintenance are fixed. At two million requests a month the arithmetic is decisive; at twenty thousand the same one-off costs buy a saving too small to notice. Below roughly a hundred thousand requests a month it is hard to justify a fine-tune on cost alone.

5. D 1. Deciding whether to fine-tune at all

People routinely answer one and assume the others. May I use these weights in this way is the model licence. May I train on this text is the data licence plus copyright plus what you promised your own users. What obligations attach to the derivative is the downstream clause, and it is the one that surprises people — a community licence can carry naming requirements and an acceptable-use policy into your product.

6. B 1. Deciding whether to fine-tune at all

The stopping rule has to be written while you are still indifferent, because three weeks in nobody is. A record that says stop if the tuned model has not beaten the frozen baseline by six points with an interval excluding zero converts an open-ended project into a decidable one — and makes an abandoned fine-tune a result rather than a failure.

7. C 1. Deciding whether to fine-tune at all

That regex means you built a generative system and then spent effort hiding the generation. The other tells are the same family: the output is drawn from a fixed list, you score with accuracy or recall, or the prompt contains the phrase 'answer with only one of'. Where they hold, a fine-tuned encoder usually wins on accuracy, latency and cost at the same time.

8. A 2. How training actually works

The labels are the inputs shifted left by one, so the model is asked to predict token two from token one, token three from tokens one and two, and so on. That shift is the whole of self-supervision, and it is also why a masking mistake is so consequential: every position you mask is a training example you silently deleted.

9. B 2. How training actually works

Cross-entropy is the average negative log-probability of the correct next token, so it inverts: $\exp$ of minus 0.9 is about 0.41. Perplexity is the same number exponentiated the other way, here about 2.5. And the number is comparable only against your own runs, because a different tokenizer, a different mask, different data or different packing all move it.

10. D 2. How training actually works

Sharper output distributions mean larger penalties on the remaining errors even as more items become correct, so validation loss and the thing you care about can move in opposite directions. Early-stopping on loss then systematically ships earlier, blander checkpoints. Loss is for detecting that something is broken, not for deciding what to ship.

11. A 2. How training actually works

AdamW divides by the square root of its second-moment buffer, and after one step that buffer is an estimate from one sample. Dividing by the square root of a bad estimate produces a badly sized step, and an early bad step can put the model somewhere the rest of the run spends its time recovering from. A few per cent of the steps, or a flat ten on a short run, is enough.

12. C 2. How training actually works

Cosine decay spends most of the run at a usefully high rate and finishes with small, careful steps that settle the weights. Stopping early removes the settling. Set the schedule length to the run you intend to finish; if you genuinely do not know the length, use a constant schedule with warmup, which also makes intermediate checkpoints comparable with each other.

13. D 2. How training actually works

Averaging per-micro-batch means gives every micro-batch equal weight regardless of how many real tokens it contained, which is not what a larger batch does. Masked instruction data has extremely variable completion lengths, so the discrepancy is real. This exact behaviour shipped in widely used libraries for a long time, which is why results from before and after the fix are not directly comparable.

14. B 2. How training actually works

A naive implementation materialises a score matrix of batch times heads times sequence squared — about 4.3 GB per layer, transiently, at 8,192 tokens with 32 heads in bf16. Flash attention computes attention in tiles and never writes it. At long context this is almost always the answer, and it comes before checkpointing, before fusing the cross-entropy, and well before lowering the batch size.

15. C 3. Full, LoRA and the parameter-efficient family

Thirteen billion times 0.6 is about 8 GB for the base, and the working memory takes the total to roughly 11 to 14 GB, which leaves very little headroom on a 16 GB card and fits a 24 GB one comfortably. Doing this arithmetic on paper before renting anything is the point: if it says 70 GB and you were planning on a 24 GB card, you learned that for free.

16. A 3. Full, LoRA and the parameter-efficient family

Standard LoRA scales its update by α over r , so as r grows the divisor grows linearly and the effective magnitude falls. A rank-64 adapter was therefore not only more expressive but also more damped, which is part of why high ranks so often disappoint. It costs nothing to switch on and it removes a confound from any rank experiment you run.

17. B 3. Full, LoRA and the parameter-efficient family

DoRA's reported gains are largest at low rank and narrow as rank rises, which is consistent with its story: a higher-rank LoRA can represent the decoupling itself. So it earns its cost when you are pinned at a low rank by memory. At rank 32 or above the case is weak — and either way the variant is the last term in the sum, behind the dataset, the template, the learning rate and the target modules.

18. D 3. Full, LoRA and the parameter-efficient family

A soft prompt chooses what the model conditions on, which selects among behaviours the model already has. Large models have many behaviours available to steer towards; small ones need the computation itself to change, which is what a weight update does. That is why on a 3B model with a genuinely hard task, LoRA beats a soft prompt reliably.

19. A 3. Full, LoRA and the parameter-efficient family

Scaling a row of a matrix is just scaling the matrix, so IA³ folds in with no inference overhead at all. And with three learned vectors per layer it has almost nothing with which to memorise, which makes it strong in the few-shot setting where a LoRA overfits. Its ceiling is correspondingly low: if rank-32 LoRA is not enough, a method with a thousand times fewer parameters is not the answer.

20. C 3. Full, LoRA and the parameter-efficient family

It is not a flaw to be engineered away; it is the direct consequence of unfreezing everything. LoRA's low-rank constraint measurably preserves more of the base behaviour, which is why it both learns less and forgets less — a good trade at three thousand examples and a bad one at three hundred thousand. Run the LoRA first anyway: it tells you whether the data is any good.

21. D 3. Full, LoRA and the parameter-efficient family

That fits with a couple of gigabytes spare, with gradient checkpointing on. Sixteen-bit LoRA on an 8B model does not fit at all; a 13B at sequence 2,048 does not fit; and a 7B full fine-tune at ten bytes per parameter is 70 GB, not 16. Unsloth's rewritten kernels roughly double the throughput and are often the difference between a configuration that fits and one that does not.

22. B 4. Building the dataset

Source, licence, the basis you are relying on, the collection date. It costs nothing at collection time and cannot be reconstructed afterwards. Without it, the question of what the model was trained on is answered from memory, and the memory belongs to somebody who may well have left.

23. C 4. Building the dataset

Exact hashing catches only byte-identical rows. MinHash with locality-sensitive hashing finds the near ones in near-linear time, and a Jaccard threshold around 0.85 is a reasonable start — checked by printing ten pairs the filter called duplicates and ten it kept. Near-duplication also inflates your evaluation and, because memorisation rises sharply with repetition, is a privacy control as well as a quality filter.

24. A 4. Building the dataset

Rewriting is three to five times faster and usually at least as good, because the draft reminds the annotator of details they would have left out. But a reviewer presented with a confident wrong answer accepts it more readily than they would reject a blank page. If more than about seventy per cent of drafts are accepted unchanged, have a second person audit thirty of them before believing it.

25. B 4. Building the dataset

The model learns the distribution of your answers, so anything that varies without carrying information becomes randomness it reproduces faithfully. Declining is precisely the behaviour you most need to be dependable, and one form of words makes it so. It also gives the model a way of declining that suits your product rather than a generic refusal borrowed from somewhere else.

26. D 4. Building the dataset

Ask a model for two hundred customer questions and you get two hundred well-formed, grammatical, politely-phrased ones that resemble nothing real. Inputs from traffic, outputs from the teacher, is the shape. And check coverage afterwards by clustering the inputs: a set covering 80 per cent of your traffic's volume but 40 per cent of its distinct intents produces a model that is excellent on the head and useless elsewhere.

27. A 4. Building the dataset

Without the block-diagonal mask, the tokens of the second document attend to the first, and the model learns a transition that never occurs at inference. Without position ids reset per document, an example's position encoding depends on what happened to be packed before it. The failure is subtle, it never shows in the loss curve, and length-grouped batching removes it entirely at the cost of some throughput.

28. C 4. Building the dataset

Sequences appearing once are memorised rarely; sequences appearing many times are memorised readily. That makes deduplication the single most effective privacy control available to a small team. Note too that a three-epoch fine-tune is, in memorisation terms, a repetition count of three on every example — which is why low epoch counts belong in the same list.

29. D 5. Objectives beyond next-token prediction

The standard pipeline is demonstrations first, to establish format and basic behaviour, then preference optimisation on top for the qualities you can judge but not write. ORPO is the notable exception, folding both into one stage. Before committing to any of it, ask whether you can state what you want in a sentence a person could follow — if you can, write twenty demonstrations instead.

30. B 5. Objectives beyond next-token prediction

Beta controls how far the model may move from the frozen reference, and 0.1 is the usual default. Higher values anchor it more tightly and make the effect smaller and safer; lower values allow a bigger behaviour change and more risk of drifting away from the coherent model you started with. Note also that DPO wants a learning rate one to two orders of magnitude below SFT, and one epoch.

31. C 5. Objectives beyond next-token prediction

The reference model is the frozen starting point, and with an adapter that starting point is literally the weights underneath. Disabling the adapter for one forward pass gives the reference log-probabilities with no second copy in memory, and TRL does it automatically when given a PEFT model with no explicit reference. It is one of the strongest practical arguments for doing preference tuning with adapters.

32. A 5. Objectives beyond next-token prediction

ORPO needs no reference model and no separate SFT stage, which is a genuine saving in memory and in pipeline complexity. What you lose is the ability to evaluate the SFT model on its own and know whether your demonstrations worked before preferences are layered on top. Choose it when you are starting from a base model with both kinds of data in hand and want one pipeline.

33. B 5. Objectives beyond next-token prediction

The critic is a second full-sized network, usually initialised from the reward model, whose job is to say how good a state is so that the advantage — how much better an action was than expected — can be computed. GRPO replaces it with a sample average over a group of responses to the same prompt. Understanding what the value model was for is what makes 'PPO without the critic' a meaningful sentence.

34. D 5. Objectives beyond next-token prediction

The reward you optimise cannot report its own exploitation — by construction it is going up. A second measure of the same intent, never optimised, gives an independent reading: A rising while B is flat or falling is the signature. Add the free proxies — mean length, list items, headings, hedges — and twenty outputs read by eye, because every reward hack is obvious on inspection and invisible in the metrics.

35. A 5. Objectives beyond next-token prediction

The Bradley-Terry objective maximises the probability that the chosen response scores above the rejected one. Nothing in it fixes a scale. Teams that build a reward model and then set a fixed threshold on its output are assuming a calibration that was never trained — the same mistake as thresholding an uncalibrated classifier, and it fires at a rate nobody has measured.

36. C 6. *The mechanics of a run*

Add the generation prompt at inference and not at training. Get this the wrong way round and the model is asked to continue from a position it never saw. And if you fine-tune a base model with no template at all, you invent one and are then responsible for using exactly the same one at serving time — so write it down in the repository beside the weights.

37. D 6. *The mechanics of a run*

The model then never receives a gradient teaching it to stop. It trains happily, the loss looks fine, and at inference it answers your question, then answers one you did not ask. Two correct fixes: add a distinct pad token and resize the embeddings, or build the mask from the attention mask rather than from the token id — which is what a well-written collator does.

38. B 6. *The mechanics of a run*

A randomly initialised row sits outside the region existing embeddings occupy, so the model receives a signal it has no representation for and the new token's logit swings wildly early in training. And under a LoRA the embedding table and output head are frozen like everything else, so without naming them as modules to save in full the row stays at whatever you initialised it to forever. Expect the adapter file to grow from about 160 MB to over 2 GB.

39. C 6. *The mechanics of a run*

The non-reentrant implementation is the one that composes correctly with PEFT; with the old default the checkpointed segments see no input requiring a gradient, the graph is not built, and the run trains nothing while erroring nowhere. The related fix for some configurations is enabling input gradients on the model. Both have cost a great many people an evening.

40. A 6. *The mechanics of a run*

Predicting something already present in the context is trivial, so the loss collapses. A log export that includes the response in a context field is a common way this happens. Distinguish it from a loss of exactly 0.0, which means every label is masked and the response template did not match, and from a sawtooth with a period of one epoch, which is ordinary memorisation.

41. B 6. *The mechanics of a run*

Effective batch is per-device batch times accumulation times device count. The run completes, the loss curve is plausible, and the result is worse because you are taking a quarter as many optimiser steps at a learning rate sized for a batch a quarter as large. Either fix is fine; doing neither is the bug. And a run under about twelve hours is usually faster on one card anyway.

42. D 6. *The mechanics of a run*

A complete checkpoint holds five things; only the trainable weights are needed to use the model. That is why a resumable LoRA checkpoint is roughly half a gigabyte rather than 160 MB, and why keeping optimiser states for old versions is the one thing you do not need to retain. Resuming is only correct if the data order and the device count are unchanged.

43. A 7. Fine-tuning beyond a chat model

It trains the model so that the first 64, 128 and 256 dimensions are each independently useful. You then store full vectors for reranking and search over short ones, which cuts index memory by around six times for a small quality cost. It is one extra wrapper around your loss and it is worth it on any index above a few million vectors.

44. C 7. Fine-tuning beyond a chat model

At serving time the reranker only ever sees fifty candidates that the bi-encoder and BM25 already thought plausible. Training it to separate a refund passage from an unrelated one about office hours teaches nothing it will need. A ratio of around four negatives per positive is a common start; far more and the model learns to say no to everything, which scores well on accuracy and ranks nothing.

45. D 7. Fine-tuning beyond a chat model

Without them you have a tensor of six numbers and a note in somebody's head about which is which. It costs one argument at construction time and it is the difference between a model that documents itself and one that depends on a person remembering the order of the labels.

46. B 7. Fine-tuning beyond a chat model

The initial prompt conditions decoding and substantially improves transcription of product names, drug names and technical vocabulary with no training at all; a larger checkpoint is a one-line change and the quality difference on accented speech is large. If the residual errors are still unacceptable, then train — with audio at 16 kHz mono, verbatim consistent transcripts, and the encoder frozen on a small dataset.

47. C 7. Fine-tuning beyond a chat model

An image becomes a block of embeddings in the language model's sequence — anywhere from about 64 positions to well over a thousand, multiplied again by tiling on high-resolution inputs. Activation memory scales with sequence length, so a batch of four pages can need more than a batch of thirty-two text examples. Lowering the input resolution is the most effective lever available, and it costs exactly the detail you lowered.

48. A 7. Fine-tuning beyond a chat model

What you caption becomes controllable; what you leave uncaptioned becomes part of the subject. Caption everything that varies and that you want to be able to change, leave uncaptioned only the thing you are actually training, and give it a rare token so it does not collide with an existing concept. This single rule decides whether the adapter is usable.

49. B 7. *Fine-tuning beyond a chat model*

Price is the weakest argument, because a hosted small model is already cheap. The strong reasons are the ones an API cannot offer at any price: the data never leaves the device, it works with no network, first-token latency has no round trip in it, and you can call the model for things that would never justify a request. The matching risk is confident, well-formatted, wrong output on anything unlike the training task.

50. D 8. *Did it actually work*

Non-English quality can drop sharply while the English numbers look untouched, and no held-out set built from your own English task will ever show it. That is what the canary set is for: fifty prompts with nothing to do with your task, including the same questions in each language you support, run on the base model and on every checkpoint and read side by side.

51. A 8. *Did it actually work*

Every time you look at a set and change something in response you fit a little to it — not by gradient descent but through your own decisions, which is slower and just as real. Open the locked set a handful of times, write down each time, and treat the gap between it and the development set as a measure of how much you overfitted your own process.

52. C 8. *Did it actually work*

Thirteen-gram overlap across your own split is checkable and should be reported even when it is zero, because a stated zero is information and an unstated one is an assumption. The base model's corpus is not accessible, which is why public benchmark numbers are weak evidence about a fine-tune and a held-out set built after the training cutoff is the strong evidence.

53. D 8. *Did it actually work*

A model reasoning about a question has no reason to prefer one arrangement of the same options over another. A model that memorised the item as it appeared has seen exactly one arrangement. It is one of very few diagnostics available for contamination you cannot otherwise inspect, and it needs no access to anybody's corpus.

54. B 8. *Did it actually work*

If two systems are not meaningfully different, a protocol that forces a winner will still produce one, and the number will look like a result. A high tie rate tells you the honest answer directly. The other five requirements are blindness, randomised order recorded per item, a stated criterion, identical conditions for both systems, and an overlap sample to measure judge agreement.

55. C 8. *Did it actually work*

The mechanism is that every example in your data demonstrates a request being answered, and the model generalises that. Refusal examples address it directly, cost an afternoon, and give the model a way to decline that suits your product. Lower rates, fewer epochs and LoRA over full fine-tuning all reduce the magnitude of the change and none of them address the cause.

56. A 8. *Did it actually work*

Flat from 75 per cent onwards, and the last step is well inside the noise floor. More of the same data will do nothing; what might help is different data — harder cases, different input styles — or better labels, or a different base model. This is the most valuable ablation in applied fine-tuning because it settles the question teams argue about most, for the price of four short runs.

57. B 9. *Shipping it, and living with it*

Every adapter must share the base model and sit at or below the server's maximum rank, so a single rank-128 adapter raises the ceiling and the memory cost for all of them. Plan it once, for the fleet. Concurrency caps and prewarming are serving decisions you can change; the rank is baked into every artefact you have already trained.

58. D 9. *Shipping it, and living with it*

Activation-aware quantisation decides what to protect by looking at which weights matter for the calibration inputs. Calibrate on general web text and you protect the weights that matter for general web text. A few hundred examples from your task protects the weights that matter for your task, and the difference is measurable on task metrics.

59. A 9. *Shipping it, and living with it*

A tokenizer fix, a re-uploaded shard, a corrected chat template — the name stays the same and the bytes do not. The manifest looks complete and quietly stops meaning anything. Record the commit hash of the model repository, not its name; the same discipline applies to configuration values pulled from environment variables and to datasets built by anything non-deterministic.

60. C 9. *Shipping it, and living with it*

Major means the base model changed and behaviour may differ everywhere. Minor means new training data on the same base. Patch means the same weights, built differently — so 3.2.0 and 3.2.1 will score slightly differently, and the version scheme tells you immediately that the difference is the build rather than a mystery.

61. D 9. *Shipping it, and living with it*

The proxies measure what is happening to your traffic and to your users; the shadow run measures the model. It is the only check that directly measures quality rather than a proxy for it, and it catches the class of failure where nothing about the inputs changed but the artefact did — including a rollback that everyone forgot about. It takes minutes.

62. B 9. *Shipping it, and living with it*

A stated boundary is a decision rather than an oversight, which is why it protects both the reader and the team. Its difficulty is exactly its value: writing it forces somebody to refuse a use that has already been proposed. Evaluation with intervals and limitations explained by mechanism are the other two sections that earn their space.

63. C 9. *Shipping it, and living with it*

It costs ten minutes per incident and it is the only mechanism that improves an evaluation set instead of letting it decay. After a year the set encodes everything the system has ever failed at. The same habit belongs in the locked set specifically, and it is what turns an incident from a cost into an asset.