

A FIELD GUIDE TO THE SPECIALIST
REVOLUTION

*The biggest model wins the demo.
The smallest model wins the invoice.*

Small Models Win

Why Fine-Tuned Small Models
Outperform General Models

Yahya Laraki

September 2026 edition

Small Models Win: Why Fine-Tuned Small Models Outperform General Models

© 2026 Yahya Laraki. All rights reserved.

Prices, model names, and benchmark figures were current at the time of writing (September 2026) and change quickly; verify before acting on them.

Contents

Preface

1. The Big Bet

2. The Specialists: Medicine, Law, Finance

3. The Specialists: Code, Math, Reasoning

4. The Specialists: Language, Retrieval, Speech

5. How Fine-Tuning Actually Works

6. Data Beats Scale

7. The Economics of Small

8. When Fine-Tuning Fails

9. Where Scale Still Wins

10. The Playbook

Appendix A: Glossary

Preface

Imagine a bank that classifies two million customer messages a month — complaints, fraud flags, wire questions. The team prototypes with the flagship model API. It works beautifully. Then finance sees the bill, engineering sees the latency at peak hours, and legal sees customer financial data flowing to a third party. The prototype was a triumph. The production system is a problem.

That bank is hypothetical, but the pattern is real, and I have watched versions of it play out across industries. When a company builds something with AI today, the first question is almost always the same: which model should we use? And the answer is almost always the same too: the biggest one. The newest flagship. The model with the most impressive demo.

This book argues that this default is wrong — expensively, measurably wrong — for a large and growing class of real work. On narrow, well-defined tasks, a smaller model that has been specialized — fine-tuned, distilled, or trained on carefully curated data — routinely beats the giant generalist. Not sometimes. Routinely. And it does so while costing a fraction to run, answering far faster, and keeping your data inside your own walls.

That is the bet this book makes. Let me state it precisely, because the sloppy version of this claim is false and the precise version is one of the most useful things you can know about AI in 2026.

The naive claim: *small models are better than big models*. That's not true. A 7-billion-parameter model will not out-reason the frontier on a task nobody prepared it for. Scale is real, and the frontier labs keep proving it. This book spends two full chapters (8 and 9) on exactly where scale wins, because the conditional thesis is only valuable if you know its borders.

The conditional claim — the one this book defends with evidence — is this: **a specialized small model outperforms a larger general model on narrow, well-defined, in-distribution tasks, when the specialization is done well, the training data is high quality, and the evaluation matches the task.** Every one of those qualifiers matters. Narrow. Well-defined. In-distribution. Done well. Each chapter that follows will show you exactly what they mean, with real numbers from real papers.

This book is written for two people at once, and I want to be explicit about the contract with each. If you're an engineer, you'll get mechanisms and numbers: how fine-tuning actually works, what LoRA and distillation change under the hood, and exact benchmark scores with their sources — every one of them traceable. You will not be asked to take anything on faith, and you will not find hand-waving where a number should be. If you're a founder, product leader, or executive, you'll get the business consequence of every idea: what it costs, how fast it answers, where it fails, what can go wrong in production, and how to decide between fine-tuning, prompting, retrieval, or just buying the API. Every technical idea is paired with the question you actually care about: so what should we do differently on Monday?

The register stays in the middle — technical ideas explained plainly, business ideas grounded in evidence. No heavy math anywhere in this book. Terms are defined the first time they appear, in plain language, and collected in a glossary at the back. If an analogy appears, it's because it earned its place.

What this book is not: it is not a training manual. You won't find step-by-step code for running fine-tuning jobs here — the tooling changes every quarter, and a printed recipe would be stale before the ink dried. Chapter 5 explains how the machinery works so you can evaluate claims and vendors; Chapter 10 gives you the decision framework. For implementation details, the cited papers and their code repositories are the right next stop.

A word on honesty, because this subject attracts hype. Every number in this book traces to a cited source — a paper, a company publication, or a public leaderboard. Papers that haven't been peer-reviewed are labeled as preprints. Company blog posts are the authors' own reported experiments, and I'll say so when I use them. Some results I found are promising but unverified; those are either omitted or labeled explicitly as unverified leads. To give you a concrete example of this discipline: a vendor later claimed a legal AI variant reached roughly 0.77 on a legal benchmark, described as "GPT-4 class." The claim is unverified, so it appears nowhere in this book's evidence. You'll see the verified 7B legal result it was based on — and you'll see me say plainly what it does and doesn't prove.

The research behind this book also surfaced results that cut *against* the thesis — a fine-tuned model collapsing outside its training distribution, a general model beating a medical specialist — and those are in the book too, because a

thesis that can't survive its own counterexamples isn't worth your time. Chapter 8 and Chapter 9 are devoted entirely to where the thesis fails and where raw scale still wins. I would rather you trust this book's limits than admire its wins.

One more discipline this book keeps, because most "small beats big" writing doesn't: there are three different mechanisms hiding inside that phrase, and they are not the same thing.

1. **Supervised fine-tuning.** Take an existing model and continue training it on examples of your task — your support tickets, your contracts, your financial headlines. The model's weights shift toward your distribution.
2. **Targeted distillation.** Use a big model as a teacher to generate training data — reasoning traces, synthetic examples, worked solutions — and train a small student on it. The small model inherits the teacher's behavior on the tasks you care about.
3. **Curated domain pretraining.** Train (or continue training) a small model from scratch on better, narrower data — textbook-quality text, a mountain of code, a curated math corpus. The model is small but its data was excellent.

These mechanisms overlap in practice, and I'll label which one each result uses. Conflating them is how people end up claiming "fine-tuning did this" when really "great data did this." The distinction matters because each mechanism has different costs, different data requirements, and different failure modes — and Chapter 10, the playbook, will ask you to choose between them.

A note on the word "win" in the title. Winning here means outperforming the larger general model *on the task* — higher accuracy, better F1, stronger benchmark scores — plus the economic advantages that come free with smallness: lower cost, lower latency, deployability on your own hardware, data privacy. It does not mean the small model is "better" in general. Nobody in this book claims a 7-billion-parameter model is smarter than the frontier. The claim is narrower and, I'd argue, more interesting: on the work that actually fills production systems, the specialist wins. The book's job is to show you exactly where that sentence is true, where it breaks, and how to tell the difference for your own use case.

Finally, a word about why this book exists now. The results it reports are mostly from 2023 to 2026 — the years when specialization stopped being a research-lab luxury and became something a small team could do in an afternoon. That shift is still underappreciated. Most of the industry's mental models about "build vs. buy" in AI were formed in the era when training anything meant a supercomputer. Those models are stale. The teams that have updated them — the ones running fine-tuned 7B models where their competitors rent flagship APIs — are quietly enjoying better margins, faster products, and data architectures their lawyers approve of. This book is my attempt to make that quiet advantage legible, with the receipts attached. Every claim you can check; every number has a source; every limit is stated. What you do with the advantage is up to you.

And one request, offered sincerely: argue with this book. Check a citation. Reproduce a comparison on your own data. Try to break the thesis in Chapter 8's terms and see if it holds. A book about evidence should be read evidentially — and if you find a result that changes the picture, that's not a failure of the book. That's the book working as designed.

How to read this book: Chapters 1 through 4 make the case with evidence — medicine, law, finance, code, math, reasoning, language, retrieval, speech. Each is self-contained enough to read on its own; together they form an argument no single result could carry. Chapters 5 and 6 explain *why* it works — the machinery of fine-tuning and the data story underneath it. These are the chapters that turn "interesting results" into a mental model you can apply to your own problems. Chapters 7 through 9 are the money and the limits — the economics of small models (Chapter 7 is the one your CFO can read alone), when fine-tuning fails, and where scale still wins. Chapter 10 turns everything into a decision framework you can actually use, including the cases where the right answer is "don't fine-tune."

A note on navigation for the two readers. If you're the engineer, you'll probably read front to back — the evidence chapters give you the results, the mechanism chapters give you the understanding, and the playbook gives you the plan. If you're the executive, here's a fast path with no shame in it: Chapter 1 for the thesis, Chapter 2 for the proof it works in serious domains, Chapter 7 for the money, Chapter 8's opening for what can go wrong, and Chapter 10 for the decision. The key takeaways at the end of every chapter are written to stand alone, so even a skim leaves you with the load-bearing facts.

If you only have an hour, read Chapters 1, 7, and 10. If you only have twenty minutes, read Chapter 1 and the key takeaways at the end of each chapter — they're written to stand alone. If you're the skeptical type, read Chapters 8 and 9 first; if the thesis survives your skepticism there, the rest of the book will read as the opportunity it is.

The evidence that follows changed how I think about building with AI. I hope it changes how you build too.

Key takeaways

- The default of "use the biggest model" is often wrong — expensively wrong — for narrow, well-defined tasks.
- The book's thesis is conditional, not universal: specialized small models beat larger generalists **on narrow, in-distribution tasks, with good data and task-aligned evaluation.**
- Three distinct mechanisms — supervised fine-tuning, targeted distillation, curated domain pretraining — are kept separate throughout.
- Every number is sourced; preprints are labeled; counterexamples are included, not hidden.
- Chapters 1–4 present evidence, 5–6 explain mechanisms, 7–9 cover economics and limits, 10 is the decision playbook.

1. The Big Bet

Every engineering team building with AI in the last few years has had the same meeting. Someone demos the newest flagship model — it writes a poem, debugs a script, explains a merger agreement, all in one session — and the room converges on the obvious conclusion: we'll just use the big one. The API is right there. It handles everything. Why would we do anything else?

That meeting is the subject of this book, because that conclusion is wrong more often than anyone in the room realizes.

The gravity of the biggest model

The pull toward the biggest model is understandable. A general-purpose frontier model is the most capable single artifact most teams have ever had access to. It really can do the poem and the script and the merger agreement. It needs no training, no data pipeline, no machine-learning expertise — just an API key and a prompt. For a prototype, it is unbeatable. For a demo, it is magic.

And demos are the problem. A demo shows you the model at its best, on tasks chosen to flatter it, judged by vibes. The presenter picked the examples. Nobody in the room ran the model's ten-thousandth answer through a grader. Psychologists would call this the availability heuristic wearing a tuxedo: the impressive thing you just saw crowds out the boring question of what the thing does on *your* data, at *your* volume, under *your* constraints. Call it the demo effect. The model that dazzles in a thirty-minute meeting is not necessarily the model that classifies your ten millionth support ticket correctly, cheaply, and in under a second.

There's a second force at work: the fear of leaving capability on the table. Nobody gets fired for choosing the biggest model. If the flagship fails, it's the model's fault; if a small fine-tuned model fails, it's *your* fault for being clever. So teams default to scale the way previous generations defaulted to the most expensive consultant — as insurance against blame, not as an engineering decision.

But insurance has a premium, and this one is steep. The biggest models are the most expensive to run per token, the slowest to answer, the hardest to keep private, and the least predictable to control. You are renting a genius to do a clerk's job, and paying genius rates for every token.

Consider an illustrative example. Imagine a company routing 500,000 customer support messages a month through a flagship API at a few dollars per million tokens. The monthly bill lands in the low five figures — noticeable but tolerable. Then the product grows 10×, as successful products do. The bill is now six figures a month, every month, forever, scaling linearly with success. Meanwhile a fine-tuned 8-billion-parameter model — which, as you'll see in this book, routinely *matches or beats* the flagship on classification-style tasks — would run that same workload for a small fraction of the cost, answer faster, and keep every message inside the company's own infrastructure. The team that chose the biggest model in that first meeting didn't make a technical decision. They signed a variable-rate mortgage on their own growth.

There's a third force, and it's worth naming because it shapes everything you'll read in the AI press: the companies selling you the models make more money when you use the biggest one. Every token billed at flagship rates is revenue. No API vendor has ever published a guide titled "you'd be better off with a smaller model" — while several have published benchmarks showing their flagship winning at everything. The information environment around model selection is not neutral. It is a sales funnel with the most expensive product at the bottom. This book is an attempt to give you the independent evaluation the vendors won't.

The specialist intuition

Now consider the rest of engineering — and the rest of human life. We do not hire a Nobel laureate to do our bookkeeping. We do not use a Formula 1 car to deliver groceries. Specialization beating generality on a narrow task is not a surprising claim; it is the default assumption of every mature field. The general practitioner refers you to the specialist. The Swiss Army knife loses to the chef's knife in every kitchen on earth.

Machine learning has known this for decades. A spam filter trained on spam beats a general text classifier at detecting spam. A fraud model trained on your

transactions beats a generic anomaly detector on your fraud. Nobody found this controversial, because the tasks were narrow and the models were small and the comparison was honest.

What changed is that large language models made generality *feel* free. One API call handles everything, so the specialist's advantage — which was always about the match between the model and the task distribution — got buried under convenience. This book is an excavation project. The specialist's advantage didn't go away. It got stronger, because now the specialist can start from a strong base model and be *taught* the specialty, rather than built from nothing.

There's a historical echo worth hearing. In the 2010s, the same debate played out in computer vision. For years, the default was the biggest ImageNet model you could afford to run — until EfficientNet showed in 2019 that a 5.3-million-parameter model could beat a 26-million-parameter ResNet on ImageNet top-1 accuracy (77.1% vs 76.0%) through better design rather than more scale (Tan & Le, 2019). The pattern — small wins via focus and craft, not via mass — predates large language models by years. LLMs just made the stakes enormously larger, because now the "big model" costs dollars per million tokens instead of milliseconds of GPU time.

That teaching has a name, or rather three names. You can fine-tune a model on examples of your task (supervised fine-tuning). You can have a big model generate training data for a small one (distillation). Or you can train a small model on superb, narrowly curated data from the start (domain pretraining). All three produce the same economic shape: a small, fast, cheap model that beats the giant on the task it was built for. You'll see all three in this book, carefully labeled, because they fail in different ways and cost different amounts.

Why this is happening now

A fair objection: if specialization is so powerful, why isn't everyone already doing it? The answer is that until recently, specialization was a rich lab's game. Fully fine-tuning a large model meant owning racks of GPUs and employing people who could keep a distributed training run alive. The economics only worked for the largest companies.

That changed with a pair of techniques you'll meet properly in Chapter 5: LoRA and QLoRA. The short version is that researchers figured out how to specialize a model by training a tiny fraction of its parameters — roughly 8 million out of 7 billion, about a tenth of one percent — while leaving the rest frozen. Suddenly a fine-tuning job that needed a data center fit on a single high-end GPU. The cost of producing a specialist collapsed by orders of magnitude, and the results in the following chapters — most of them from 2023 onward — are what that collapse made possible.

This timing matters for the thesis. The evidence in this book isn't a collection of historical curiosities; it's the output of a capability that became widely accessible only recently, and whose results are still compounding. The teams winning with small specialists today are early, not late.

The question this book answers

Strip away the hype and the entire book reduces to one practical question: **for a given piece of work, should you rent the giant or train the specialist?**

Most teams answer it once, early, by default — and then never revisit it. The flagship API becomes load-bearing infrastructure. Prompt engineering becomes a full-time job. The monthly bill becomes a fact of nature. Nobody re-runs the comparison that would have been run at the start if the team had known then what the literature now shows: that for narrow, high-volume, well-specified work, the specialist usually wins on quality and always wins on cost.

This book exists to make that comparison unavoidable. By the time you finish Chapter 4, you will have seen small specialists beat larger generalists in medicine, finance, law, code, math, reasoning, translation, retrieval, and speech — each result sourced, each caveat stated. By Chapter 7, you will have real prices and a worked cost comparison. By Chapter 10, you will have a decision framework that tells you, for your specific task, which path to take and what evidence would change your mind.

You don't have to accept the thesis in advance. In fact, I'd prefer you didn't — the skeptical reader gets more from this book than the convinced one, because the counterexamples in Chapters 2, 8, and 9 are doing half the argumentative work. Just bring one narrow task from your own world as you read, and keep

asking: would the specialist win here? By the end, you'll know how to answer that question with data instead of vibes. That's the whole game.

A tour of what's coming

The evidence in Chapters 2 through 4 comes from published results, each with its source, each with its caveats. A preview:

In medicine, a 7-billion-parameter model trained on medical textbooks and GPT-4-generated reasoning traces scored 74.3% on the MedQA benchmark — beating a 70-billion-parameter medical model that scored 70.2%. Ten times the parameters, worse at the specialty.

In finance, a 7-billion-parameter model instruction-tuned on 136,609 financial examples beat GPT-4 at financial sentiment classification — and also beat BloombergGPT, a 50-billion-parameter model pretrained specifically on finance. The specialist beat the generalist *and* the bigger specialist.

In code, a 1.3-billion-parameter model trained on textbook-quality synthetic data beat models ten times its size at Python code generation. A 15-billion-parameter code model beat a 540-billion-parameter generalist by a wide margin.

In reasoning, a 13-billion-parameter model taught reasoning strategies by GPT-4 beat a 70-billion-parameter chat model by over nine points on average across six benchmarks.

And these are not cherry-picked miracles. They are the predictable output of a mechanism: when you concentrate training on the task distribution, the model spends its capacity where it matters instead of spreading it across everything. A useful mental image: the generalist is a library containing every book ever written, and answering your question requires finding the right shelf. The specialist is a single well-organized manual for exactly your job. For the job the manual covers, the library's vastness isn't an advantage — it's search overhead. The specialist doesn't know more. It knows where things are.

Each of these results will get its full treatment — methods, numbers, caveats, business consequences — in the chapters ahead. The preview's job is just to calibrate your expectations: the thesis of this book is not a contrarian hot take. It's the summary of a literature.

The honest boundaries

A book that only showed wins would be a sales pitch, and this isn't one. The same research that found these victories found their boundaries, and you'll meet them early — Chapter 2 includes them alongside the wins, because they sharpen the thesis rather than weaken it.

The financial specialist that beat GPT-4 at sentiment classification collapsed at numerical finance questions in the very same paper — the fine-tuning won inside its training distribution and fell apart outside it. A general GPT-4, with no medical specialization at all, beat an early medical specialist at medical questions; specialization is not destiny, and the lead changes hands as methods improve. Distillation wins turn out to be distribution-dependent: the student beats the teacher on the distilled distribution and loses elsewhere.

These boundaries are the point. "Small beats big" is false. The true claim has conditions, and the conditions are knowable: narrow task, well-defined evaluation, in-distribution data, high-quality training examples. When those hold, the small specialist wins — on accuracy, on cost, on latency, on privacy. When they don't, the generalist's breadth reasserts itself. Chapters 8 and 9 are devoted to the failure modes, because knowing exactly where the thesis breaks is what makes it usable.

There's one more boundary worth stating early, because it protects you from the most common misreading of this book. The thesis is about *tasks*, not about *intelligence*. Nothing in the following chapters implies that a specialized 7B model is "smarter" than a frontier model, or that scale doesn't matter, or that the frontier labs are wasting their money. The frontier's generality is what makes the specialist possible — every fine-tuned model in this book starts from a base model that only exists because someone trained at scale. The relationship is symbiotic, not adversarial: scale creates the raw capability, specialization aims it. The book's argument is about where to spend your *next* dollar — and for narrow production work, the evidence says the aimed dollar beats the raw one.

The conditional thesis, stated precisely

Here is the claim this book defends, in its full form:

On narrow, well-defined, in-distribution tasks, a small model that has been specialized — by supervised fine-tuning, targeted distillation, or curated domain pretraining — outperforms a larger general-purpose model, provided the specialization uses high-quality data and the evaluation measures the actual task.

Read the qualifiers as a checklist, because Chapter 10 will turn them into one. To make each concrete, here's an illustrative example of a task that passes the full checklist — and one that fails it.

Passes: routing incoming customer emails into twelve fixed categories (billing, cancellation, technical issue, and so on) for a software company. Narrow — twelve buckets, nothing else. Well-defined — every email belongs in exactly one bucket, and misroutes are countable. In-distribution — next month's emails look like last month's. High-quality data — the company has three years of emails already labeled by its support team. Task-aligned evaluation — accuracy on a held-out sample of real emails. This is the kind of task where a fine-tuned small model should beat the flagship, and the literature says it will.

Fails: "help our analysts think through novel market situations." Not narrow, not well-defined, no stable distribution, no ground truth to train on, no way to evaluate except vibes. This is flagship territory — the generalist's breadth is the whole point. Anyone proposing to fine-tune a small model for this is buying trouble, and Chapter 8 will show you the wreckage of similar attempts.

The five qualifiers:

- **Narrow:** the task has a bounded scope — classify these headlines, answer these medical questions, extract these entities. Not "be smart about everything."
- **Well-defined:** you can say what a right answer looks like, and measure it. Benchmarks, graded evals, human review with rubrics.
- **In-distribution:** the inputs at deployment resemble the inputs in training. This is the qualifier that kills most careless fine-tuning projects.
- **High-quality data:** the specialization data is clean, correct, and representative. A thousand excellent examples beat a million noisy ones — you'll see the study in Chapter 6.
- **Task-aligned evaluation:** you're measuring the actual job, not a proxy that flatters someone. More on the many ways benchmarks lie in Chapter 8.

When the checklist holds, the economics follow automatically: the small model costs a fraction per token, answers several times faster, can run on your own hardware or even on a phone, and never sends your data to a third party. Chapter 7 puts real numbers on all of this.

The bet of this book is that a large fraction of production AI work — the unglamorous, high-volume, well-specified work that actually makes money — satisfies this checklist. If that's right, then the industry's default of reaching for the biggest model is not just suboptimal. It's the most expensive habit in software.

One final framing before the evidence begins. This book is structured as a case: the claim (this chapter), the evidence (Chapters 2–4), the explanation (Chapters 5–6), the economics and the limits (Chapters 7–9), and the verdict you can act on (Chapter 10). But it's also structured as a bet you can test yourself. Every result in the evidence chapters comes with its source, its benchmark, and its caveats — enough for a competent team to reproduce the comparison on their own data. If the thesis is right, the most valuable thing you can do after reading this book is also the cheapest: take one narrow, well-defined task in your own operation, fine-tune a small model on it, and measure. The book makes the case; your data delivers the verdict.

Key takeaways

- Teams default to the biggest model because of the demo effect and blame-avoidance — not because it's the best engineering choice for their task.
- Specialization beating generality on narrow tasks is the oldest pattern in engineering; LLMs buried it under the convenience of one API, but didn't erase it.
- Three mechanisms produce the specialist advantage: supervised fine-tuning, targeted distillation, and curated domain pretraining.
- The wins are real and published: small specialists beating models 10× their size in medicine, finance, code, and reasoning.
- The thesis is conditional — narrow task, well-defined eval, in-distribution data, high-quality examples — and the book is explicit about where it breaks.

Sources

- Mingxing Tan and Quoc V. Le, Google, *EfficientNet: Rethinking Model Scaling for Convolutional Neural Networks* (ICML 2019). <https://arxiv.org/pdf/1905.11946>
- Benchmark landscape and model references throughout this book draw on the September 2026 LLM benchmark survey in the book's research files, and on the individual papers cited in each chapter's Sources section.

2. The Specialists: Medicine, Law, Finance

Professions are the natural habitat of the specialist. A doctor, a lawyer, and a financial analyst do narrow, high-stakes, well-defined work — exactly the conditions under which a specialized small model should beat a generalist. So the professional domains are where the thesis faces its first real test. The results are striking, and so are the failures. Both belong here.

How to read the numbers in this chapter

Before the results, a brief user's guide to the scoreboards, since this book asks you to trust numbers and you should know what they mean.

A **benchmark** is a fixed set of test questions with known right answers. **MedQA** uses USMLE-style medical questions — multiple choice, one right answer, scored as a percentage. **F1 score** is the standard accuracy measure for classification tasks: it balances catching the positives against avoiding false alarms, on a 0-to-1 scale where higher is better. When a paper reports "F1 0.86 vs 0.78," the first model is meaningfully better, not marginally.

Two honesty notes I'll repeat wherever they apply. First, most of the papers in this chapter are **preprints** — published on arXiv or OpenReview but not yet peer-reviewed. Their numbers are the authors' own reported experiments, and I treat them as evidence, not settled fact. Second, a benchmark score is only as good as the benchmark: a model can win on a test set that resembles its training data and lose in the real world. That caveat is not a footnote in this book — it's a recurring character, starting with the FinMA story below.

Medicine: the 7B model that beat the 70B medical model

Medical question answering is one of the hardest narrow tasks in AI — and one of the most consequential, because the same technology that answers exam questions is being piloted for clinical decision support, triage, and patient

communication. The MedQA benchmark uses questions in the style of the USMLE — the American medical licensing exam. Answering them requires real domain knowledge, not just fluent text. It's also a benchmark with a meaningful human reference point: these are the questions we ask human doctors to answer before we let them practice.

In 2024, a team released Meerkat-7B: a 7-billion-parameter model built on top of Mistral-7B, continually trained on medical instructions and chain-of-thought reasoning data generated by GPT-4 from medical textbooks. Think of it as a capable medical student who was tutored, step by step, by the best available teacher, using the actual textbooks. That's mechanism (2) — targeted distillation — with a dash of continued training.

On MedQA, Meerkat-7B scored **74.3%**. Its headline comparison was against MediTron-70B — a 70-billion-parameter model specialized for medicine — which scored **70.2%** (Kim et al., preprint 2024, <https://arxiv.org/html/2404.00376v1>; the work was subsequently peer-reviewed and published in *npj Digital Medicine*, 2025). Ten times the parameters, the same specialty, worse score. The small model didn't just beat a generalist; it beat a *bigger specialist*.

Why did the textbook-and-tutor recipe work? Because medical QA is a knowledge-and-reasoning task with a stable distribution: diseases don't change their symptoms to fool your model, and the USMLE format is fixed. The distillation concentrated the small model's capacity exactly where the test measures. And note what the recipe required: not a data center, but textbooks and a teacher model's reasoning traces — an investment a single research team could afford. That's the democratizing half of the thesis: the wins aren't reserved for labs with frontier budgets.

A second medical result sharpens the picture. Llama-3-Meditron-8B, an 8-billion-parameter medical model, scored **69.88** on MedQA against Flan-PaLM's **67.60**, and averaged **70.06** across MedMCQA, MedQA, and PubMedQA versus Flan-PaLM's **68.07** (Sallinen et al., AAAI 2025 workshop, <https://openreview.net/pdf?id=ZcD35zKujO>). One important qualifier: these scores were achieved with the paper's **MediTree inference pipeline** — a Tree-of-Thoughts sampling strategy — not by the base model alone, which scored 63.00 on MedQA (65.88 average) without it. Part of the win belongs to a specialized *inference procedure*, not just specialized weights — Chapter 4's Self-RAG is the same idea, and the book labels it whenever it appears. That's a

small specialist (plus a smart pipeline) beating a larger generalist. But the same paper reports the honest limits: Meditron-8B did *not* beat the MediTron-2 70B average of **70.17**, and GPT-4 scored **74.50**. Specialization lifted the small model past the big generalist, but not past the biggest specialist or the frontier. The thesis holds conditionally — and the conditions are visible in the numbers.

For a hospital administrator, the business consequence is concrete. An 8-billion-parameter model can run on a single GPU inside the hospital's own data center. Patient data never leaves the building — a requirement, not a preference, under health privacy law. The 70-billion-parameter alternative needs a rack of expensive hardware or a third-party API contract. When the smaller model also scores higher on the medical benchmark, the decision stops being a trade-off and becomes obvious. And when it scores slightly lower — as Meditron did against the biggest models — the decision becomes a genuine cost-benefit analysis rather than an automatic win for scale. That honesty is what makes the framework useful: sometimes the answer really is the bigger model, and you'll know why.

The counterexample: when the generalist won

Intellectual honesty requires the other side of the medical story, because it happened and it's instructive.

In March 2023, Microsoft researchers tested GPT-4 — a general model with no medical specialization — on medical challenge problems (Nori et al., 2023). GPT-4 exceeded the USMLE passing score by more than 20 points and **outperformed Med-PaLM, a medical specialist built on a 540-billion-parameter model** (https://www.microsoft.com/en-us/research/uploads/prod/2023/03/GPT-4_medical_benchmarks.pdf). At that moment, raw general scale beat domain specialization. The specialists later caught up — Med-PaLM 2 reached a reported 86.5% on MedQA (Singhal et al., preprint May 2023, later in Nature Medicine 2024) — but the lesson stands: specialization is not destiny. The lead changes hands as methods improve, and a big enough leap in general capability can swamp a specialist's head start.

Why did the generalist win that round? Because Med-PaLM's specialization was built on an older base with an older recipe, while GPT-4 represented a generational jump in pretraining scale and quality. Specialization multiplies the

base model's capability; it doesn't replace it. A $2\times$ specialization gain on a base that's $5\times$ behind still loses. This gives the thesis a dynamic form that the static version misses: the question is never "specialist or generalist?" in the abstract. It's "this specialist, built on this base with this data, versus this generalist, on this task, today?" The honest practitioner re-asks that question every model generation — which, in 2026, means roughly every quarter.

This is exactly why the thesis is conditional. In early 2023, the condition "the specialization is done well enough to matter" wasn't met relative to what GPT-4 brought. By 2024, with better distillation recipes, the 7B specialists were beating 70B medical models. The frontier moves; the specialist's job is to move with it, cheaply.

Finance: the 7B model that beat GPT-4 and BloombergGPT

If medicine tests knowledge, finance tests judgment under noise — and it produced one of the cleanest wins in the literature.

FinMA is a financial language model instruction-tuned on FIT, a dataset of 136,609 financial instruction examples — mechanism (1), supervised fine-tuning, at serious scale. To put that number in perspective: it's roughly the size of a large company's historical archive of analyst notes, earnings summaries, and labeled headlines. An organization sitting on that kind of data already owns most of what it needs to build its own specialist. The paper behind it, PIXIU (Xie et al., preprint, 2023, <https://arxiv.org/pdf/2306.05443>), evaluated it on the FLARE benchmark suite.

On FPB, a financial sentiment classification task, FinMA-7B achieved an F1 score of **0.86** (FinMA-30B: **0.88**), against GPT-4's **0.78** and BloombergGPT's **0.51**. On financial headline classification, FinMA-7B averaged **0.98** F1 versus GPT-4's **0.86** and BloombergGPT's **0.82**.

Pause on that BloombergGPT number. BloombergGPT is a 50-billion-parameter model pretrained specifically on finance — by Bloomberg, with Bloomberg's data. It is the canonical "big specialist." And a 7-billion-parameter instruction-tuned model beat it decisively on these tasks, while also beating the

frontier generalist. This is the thesis in its purest form: targeted fine-tuning on the task distribution beats both generality *and* brute-force scale.

Why did instruction tuning work so well here? Financial sentiment and headline classification are pattern-recognition tasks over a stable vocabulary — earnings language, analyst phrasing, market terminology. The 136,609 instruction examples taught the model the *shape* of financial text: what "headwinds" means in an earnings call versus a weather report. The generalist knows both meanings; the specialist learned which one matters and stopped hedging. That certainty is what the F1 scores measure.

The business translation: sentiment classification of financial text is a high-volume, well-defined, in-distribution task — exactly the kind of work banks and funds run millions of times a day. A 7B model doing it better than the flagship API, at a small fraction of the cost per token, on hardware you own, is not a research curiosity. It's a budget line. And unlike the flagship API, its behavior is *stable*: the same input gets the same classification every time, with no silent model updates from a vendor changing your results overnight. For regulated financial work, that determinism is worth as much as the accuracy.

The counterexample inside the victory

But the FinMA paper contains its own refutation of the naive thesis, and it's the most instructive failure in this chapter. The same FinMA models that dominated sentiment classification collapsed on numerical finance question answering — a task outside their fine-tuning distribution.

On FinQA (exact match): FinMA-7B scored **0.06**, FinMA-30B **0.11**, against GPT-4's **0.63**. On ConvFinQA: FinMA-7B **0.25**, FinMA-30B **0.40**, against GPT-4's **0.76** (Xie et al., preprint, 2023).

The same model. The same paper. Dominant in-distribution, helpless out-of-distribution. Fine-tuning is a deal with the task distribution: it concentrates the model's capacity where the training data points, and the price is paid everywhere else. This is the "in-distribution" qualifier from Chapter 1 made quantitative — a 0.86-to-0.06 cliff. Any team fine-tuning a model needs this image burned into memory: specialization is not a rising tide. It's a spotlight. Everything outside the beam goes dark.

It's worth understanding *why* the collapse was so total, because the mechanism generalizes. FinMA's instruction tuning taught it the surface patterns of financial text — sentiment-bearing phrases, headline structures, the vocabulary of markets. Numerical QA requires something different: multi-step arithmetic reasoning over tables, chaining quantities through calculations. The tuning data contained almost none of that, so the model never learned it — and worse, the tuning may have actively suppressed the general reasoning the base model had, by rewarding confident pattern-matching over careful computation. This is a milder form of the catastrophic forgetting you'll meet in Chapter 8: the specialist didn't just fail to learn the new skill, it may have unlearned the old one. The practical lesson: before fine-tuning, inventory the skills your task needs and check whether your tuning data teaches all of them. The ones it doesn't teach are the ones that will break in production, silently, at 2 a.m.

Law: specialization wins, honestly scoped

Legal NLP gives us a different shape of result — a specialization win at matched scale, which is worth including precisely because it *isn't* a smaller-beats-larger story.

SaulLM-7B-Instruct, a 7-billion-parameter model specialized for law, scored **0.61** on average on LegalBench-Instruct against roughly **0.55** for the Mistral-7B-Instruct baseline — a gain of 6 absolute points, about 11% relative (preprint, 2024, <https://arxiv.org/html/2403.03883v1>). Same size, same task family, better performance through legal specialization — mechanisms (1) and (3) working together: instruction tuning on legal tasks plus continued exposure to legal text.

Why does law respond to specialization? Legal language is a dialect — archaic constructions, terms of art, citation formats, and reasoning patterns that barely appear in general web text. A generalist model has seen some of it; a legal specialist has *lived* in it. The 11% relative gain is the measurable value of that immersion. For a law firm, an 11% improvement in instruction-following on legal tasks — drafting assistance, clause review, research triage — compounds across every matter the firm handles.

The honest scoping matters here. This is a 7B model beating a 7B baseline, not a small model beating a giant. It demonstrates that specialization adds real

capability, but it doesn't demonstrate the David-and-Goliath upset. (A later vendor claim put a SaulLM variant at roughly 0.77, described as "GPT-4 class" — that claim is unverified, and this book doesn't use it.) Law firms evaluating AI should read this result correctly: domain adaptation measurably helps, but the legal small-beats-large upset hasn't been rigorously published yet. The gap in the literature is itself information — it tells you where the next wins are likely to come from, and that anyone selling you one today owes you a benchmark.

The gaps are data too

A brief note on what this chapter doesn't contain, because absences carry information. The research behind this book found no rigorous published result of a small fine-tuned model beating a large generalist at customer-support dialogue — only anecdotes, synthetic demos, and vendor claims. That doesn't mean it's impossible; support triage is narrow and well-defined enough that the thesis predicts it *should* work. It means nobody has published the clean experiment yet. Similarly, text classification beyond finance (the FinBERT and CopBERT results against GPT-4) looked promising, but the exact score tables couldn't be verified from primary sources, so they're excluded.

I mention this because the pattern of the missing evidence matters: the wins cluster where benchmarks are mature and public (medicine, finance, code), and the gaps cluster where benchmarks are proprietary or fuzzy (support dialogue, legal practice). If your domain lacks a public benchmark, that doesn't refute the thesis — but it does mean you'll have to build the evaluation yourself before you can claim the win. Chapter 10 covers how.

How fragile are these wins?

A careful reader should now be asking the uncomfortable question: how much should I trust these numbers? It's the right question, and the answers vary by result — which is why this book labels every source.

The medical and finance results come from preprints — arXiv and OpenReview papers that haven't completed peer review. Their numbers are the authors' own reported experiments. That's meaningful evidence, but it's one

team reporting on its own work, and independent replication is the gold standard that turns a result into a fact. Treat the exact margins (74.3 vs 70.2, 0.86 vs 0.78) as the authors' measurements, and the *pattern* — small specialists beating bigger models across independent teams, domains, and years — as the finding that matters. One preprint could be wrong about its margin. Sixteen results across medicine, finance, code, math, and language don't all point the same way by accident.

There's a subtler risk worth naming: benchmark contamination. If a model's training data included the benchmark's test questions, its score measures memorization, not capability. The research behind this book found genuine contamination concerns in the broader literature — models recovering masked answers from famous benchmarks at rates that suggest the tests leaked into training. The honest response isn't to discard every number; it's to weight results by how hard they'd be to game. A fine-tuned 7B model beating GPT-4 on a public classification benchmark is hard to explain by contamination alone, because contamination would help the *bigger* model too — it saw more of the internet. When the small model wins anyway, the specialization explanation survives. Chapter 8 goes deeper on evaluation hygiene; for now, note the principle: trust the *comparisons*, and trust them most when the deck seems stacked against the winner.

Finally, notice what none of these results claim. Nobody claims the 7B medical model is a better doctor than the frontier model in general. Nobody claims FinMA understands finance the way an analyst does. The claims are narrow — higher scores on specific, well-defined benchmarks — and the book's thesis is only as broad as those claims. That's a feature. Narrow claims are checkable, and checkable claims are the ones you can build a business on.

What the professional domains teach

Three patterns emerge from medicine, law, and finance — and they generalize to every chapter that follows.

First, **the wins are real and large**. A 7B medical model beating a 70B medical model by 4 points. A 7B financial model beating GPT-4 by 8–12 points on classification F1 while also beating a 50B finance-pretrained model. These

aren't rounding errors or prompt tricks. They're the predictable result of concentrating training on the task.

Second, **the wins are bounded by the distribution**. FinMA's collapse on numerical QA is the clearest demonstration in the literature that fine-tuning trades breadth for depth. The specialist wins where it was trained and loses where it wasn't — sometimes catastrophically. Deployment planning has to treat the training distribution as the model's territory and everything else as foreign soil.

Third, **the frontier gets a vote**. GPT-4 beating Med-PaLM in 2023 is the permanent reminder that "specialized" is relative to the generalist of the moment. A specialist built on last year's base can be overtaken by this year's generalist. The strategic response isn't to abandon specialization — it's that specialization is *cheap*. Re-tuning a 7B model on a new base costs orders of magnitude less than training a frontier model, so the specialist can re-specialize every generation. The generalist's lead is rented; the specialist's economics compound.

Fourth, **the mechanism matters less than the match**. Medicine's biggest win came from distillation, finance's from supervised fine-tuning, law's from a mix. What they share isn't a technique — it's the fit between the training and the task. Teams sometimes ask "should we fine-tune or distill?" as if one were universally better. The evidence says: pick the mechanism that gets the most task-representative signal into the model for the least cost. Sometimes that's your own labeled data (fine-tuning). Sometimes it's a teacher's reasoning traces (distillation). Sometimes it's a better corpus (domain pretraining). The checklist in Chapter 1 doesn't mention mechanisms at all — deliberately. Mechanisms are implementation; the distribution match is the strategy.

For the non-technical reader, the takeaway is a decision rule: if your work looks like these professions' benchmark tasks — narrow, high-volume, well-specified, with answers you can check — the evidence says a specialized small model will likely beat the big API on quality while winning overwhelmingly on cost, speed, and privacy. If your work looks like novel reasoning across unfamiliar territory, keep reading: Chapters 8 and 9 are about exactly where this rule stops.

And for the technical reader, the takeaway is a research posture: reproduce the wins before you believe them, check the eval setup before you cite them, and

always ask what the model does *outside* the reported distribution. The FinMA cliff — 0.86 to 0.06 — should be your mental reference for every fine-tuning proposal you'll ever review. Ask for the out-of-distribution numbers. If nobody measured them, the win isn't finished.

Key takeaways

- Meerkat-7B (74.3%) beat the 10×-larger MediTron-70B (70.2%) on MedQA via distillation of GPT-4-generated medical reasoning traces.
- Llama-3-Meditron-8B (with the MediTree inference pipeline) beat the larger generalist Flan-PaLM on medical QA — but not the biggest medical specialist or GPT-4, showing the thesis's conditional boundaries.
- FinMA-7B beat GPT-4 and the 50B finance-pretrained BloombergGPT on financial classification (F1 0.86–0.98), via instruction fine-tuning on 136,609 financial examples.
- The same FinMA models collapsed on numerical finance QA (0.06 vs GPT-4's 0.63) — specialization is a spotlight, not a rising tide.
- GPT-4 beat the medical specialist Med-PaLM in 2023: general scale can swamp specialization, so specialists must re-specialize as the frontier advances.
- SaulLM-7B's legal win (+6 points on LegalBench-Instruct) is a genuine specialization gain at matched 7B scale — not a small-beats-large upset.

Sources

- Meerkat-7B: Hyunjae Kim et al., *Small Language Models Learn Enhanced Reasoning Skills from Medical Textbooks* (preprint, 2024; peer-reviewed publication: *npj Digital Medicine*, 2025, <https://doi.org/10.1038/s41746-025-01653-8>). <https://arxiv.org/html/2404.00376v1>
- Llama-3-Meditron-8B with the MediTree inference pipeline: Alexandre Sallinen et al., *Llama-3-Meditron: An Open-Weight Suite of Medical LLMs Based on Llama-3.1* (AAAI 2025 GenAI4Health workshop). <https://openreview.net/pdf?id=ZcD35zKujO> — "MediTree" is the paper's Tree-of-Thoughts inference pipeline, not the model; the base 8B model alone scored 63.00 on MedQA (65.88 avg).

- GPT-4 vs Med-PaLM: Harsha Nori et al., Microsoft, *Capabilities of GPT-4 on Medical Challenge Problems* (March 2023). https://www.microsoft.com/en-us/research/uploads/prod/2023/03/GPT-4_medical_benchmarks.pdf
- Med-PaLM 2 (86.5% on MedQA): Karan Singhal et al., Google (preprint, May 2023; published in Nature Medicine, 2024).
- FinMA / PIXIU: Qianqian Xie et al., *PIXIU: A Large Language Model, Instruction Data and Evaluation Benchmark for Finance* (preprint, 2023). <https://arxiv.org/pdf/2306.05443>
- SaulLM-7B: *SaulLM-7B: A Pioneering Large Language Model for Law* (preprint, 2024). <https://arxiv.org/html/2403.03883v1>

3. The Specialists: Code, Math, Reasoning

Programmers were the first to feel it. Code is unforgiving: a program either runs or it doesn't, so there is no hiding behind vibes or polished-sounding filler. Benchmarks like HumanEval — a set of programming problems where the model must write code that passes hidden tests — give a clean, machine-graded verdict. And on those verdicts, small specialized models have been beating giants since 2023.

There is a reason code, math, and reasoning lead this book's evidence. These are the fields where answers are *checkable*: code runs, math has right answers, reasoning benchmarks have ground truth. When a small model wins here, nobody can dismiss it as a vibes-based judgment or a friendly evaluator. The scoreboard is the scoreboard. That makes this chapter the hardest evidence in Part II — and the hardest for skeptics to wave away.

Reading the scoreboard

A quick glossary, since these benchmarks recur. **HumanEval** is 164 hand-written Python problems; **pass@1** is the fraction the model's first attempt solves (no retries, no hints). **MBPP** (Mostly Basic Python Problems) is a larger set of simpler tasks. **GSM8K** is 8,500 grade-school math word problems. **MATH** is 12,500 competition-level mathematics problems — genuinely hard. When you see a percentage on these, it means machine-graded correct answers, not a judge's opinion.

This matters because later chapters discuss softer evaluations — AI judges, human preferences — where the numbers are squishier. Here, the numbers are concrete. A model either wrote working code or it didn't.

A 1.3-billion-parameter model that writes better code than models ten times its size

In 2023, a team at Microsoft Research published a paper with a deliberately provocative title: *Textbooks Are All You Need* (Gunasekar et al., 2023). They trained a model called **phi-1** — just 1.3 billion parameters, tiny by the standards of the day — on a radical idea: instead of feeding it the entire internet, feed it something like a good education.

The training data was 6 billion tokens of "textbook quality" code, filtered from the web for clarity and correctness, plus 1 billion tokens of synthetic textbooks and exercises generated by an earlier AI. The whole training run took four days on eight A100 GPUs — a modest cluster, not a supercomputer.

Stop and consider what "textbook quality" means as a strategy. The conventional wisdom said: to make a model good at code, show it *all* the code — every repository, every skill level, every half-finished weekend project. The phi team did the opposite. They filtered aggressively for code that reads like a good teacher wrote it: clear variable names, correct logic, explanatory structure. Then they had an AI write synthetic textbooks — explanations, exercises, worked examples — in the style of a computer-science course. The model never saw the internet's code slums. It got an elite education instead of a big one.

This is mechanism (3) at its purest, and it carries a lesson that echoes through the whole book: for a narrow capability, *what* the model learns from matters more than *how much*. Four days on eight GPUs produced a model competitive with systems trained at orders of magnitude greater expense. When people say AI progress requires ever-larger training budgets, phi-1 is the counterexample taped to the wall.

The results on HumanEval, the standard code-generation benchmark (pass@1 means: the fraction of problems the model's *first* attempt solves correctly), were startling. phi-1 scored **50.6% pass@1** on HumanEval and **55.5%** on MBPP, a second coding benchmark. The paper reported it beating models ten times larger; only the frontier flagship of the day beat it. (Secondary reporting put GPT-3.5 around 47% on HumanEval — treat that comparator as secondary, since it comes from press coverage rather than the paper's own tables.)

Now the mechanism label, because this is where sloppy "small beats big" storytelling goes wrong: **phi-1 is mechanism (3) — curated data pretraining, not fine-tuning**. Nobody took a big model and adapted it. The model was built from scratch on better, narrower data. The win belongs to *data quality*, not to fine-tuning as a technique. Keep that distinction in your head; the book's central argument depends on it.

The business consequence: a 1.3B model can run on a laptop. Code assistance at that size means autocomplete with zero latency, zero per-token cost, and code that never leaves the developer's machine — a property no API-based giant can offer. For a company selling developer tools, the choice between "pay per token to a giant" and "ship a small model with the product" is the difference between a cost center and a margin.

The purest fine-tuning win: a tuned 8B beats a 70B sibling

Most of this chapter's wins come from mechanisms (2) and (3) — distillation and data curation. Here is the cleanest mechanism-(1) case in the dossier, and it deserves its place precisely because it is *fine-tuning* in the plain sense: take an existing model, train it on task data, beat a bigger model.

Together AI reported (2024, vendor blog — treat as the vendor's own experiment, on their private benchmark, not independent proof): on a custom 1,000-problem math set, a fine-tuned Llama-3-8B scored **65.2% versus 64.2% for Llama-3-70B** — and the base 8B model had scored only 47.2% before tuning. Fine-tuning closed an 18-point gap and then some, pushing the small model past its sibling nearly nine times its size. GPT-4o still led at 71.4%, so the ceiling held — but the small model moved up a full weight class on the strength of task data alone.

Why does this matter more than the flashier cases? Because it is the most *replicable* win in the book. You don't need to curate a trillion tokens of textbooks or distill reasoning traces from a frontier model. You need a thousand good examples of your task and a fine-tuning run. That is the playbook Chapter 10 is built on, and this is its proof of concept.

StarCoder: a 15B code specialist against a 540B generalist

The BigCode project — a collaboration between Hugging Face and ServiceNow — built StarCoder, a 15-billion-parameter model trained specifically on code: The Stack, a massive collection of source code assembled for open training, plus extra Python adaptation. **Mechanism: (3), domain pretraining.**

Their own reported numbers (a company publication by the authors, not peer-reviewed, but the authors' own experiments) tell a dramatic story. On HumanEval pass@1: **StarCoder 40.8 (prompted; 33.6 unprompted) versus PaLM-540B at 26.2** (Li et al., 2023; BigCode, 2023). On MBPP: **StarCoder 52.7 versus PaLM-540B 36.8**. Read that again: a model 36 times smaller, trained for one domain, beating a 540-billion-parameter generalist by double-digit margins on that domain's home turf.

Note the prompted-versus-unprompted gap in StarCoder's own numbers — 40.8 with careful prompting, 33.6 without. Even the specialist benefits from being asked well. That's a useful reminder for Chapter 10's playbook: specialization and good prompting stack; they aren't alternatives.

There is also a quieter lesson in *where* StarCoder's data came from: The Stack, a massive collection of source code assembled for open training. Specialization dodges one of the giants' biggest headaches — training-data provenance — because a domain corpus can be curated deliberately rather than scraped indiscriminately. A code model trained on a deliberately assembled corpus is a cleaner legal product than a generalist trained on the whole internet. For companies, that's not a footnote; it's a procurement advantage.

PaLM's size here is worth stating precisely: 540 billion parameters is a disclosed figure for that particular model. (For closed models like GPT-4 or ChatGPT, parameter counts are undisclosed and this book never states them.)

StarCoder2: edging out a model nearly five times larger

The sequel sharpened the point. StarCoder2-15B-Instruct — built with *self-alignment*, where the model generates its own instruction data, validates it by

actually executing the code, and then fine-tunes on the validated set — scored **72.6 on HumanEval versus 72.0 for CodeLlama-70B-Instruct** (BigCode, 2024). **Mechanisms: (1) supervised fine-tuning plus (2) self-distillation** — synthetic instruction data plus the model teaching itself.

Two things to note honestly. First, the margin is 0.6 points — real but thin. This is an "edges out," not a demolition. The book reports it that way because the thesis is conditional, and thin margins are part of the evidence. Second, this is again the authors' own reported result on their company blog. It counts as evidence; it does not count as peer-reviewed proof.

DeepSeekMath: 7 billion parameters against the 540B class

Mathematics is the hardest test of reasoning, and here the specialization story gets almost absurd. DeepSeekMath-7B — seven billion parameters — was built by continued pretraining on 120 billion curated math tokens from Common Crawl, followed by math instruction tuning and then GRPO, a reinforcement-learning method. **Mechanisms: (3) domain pretraining plus (1) supervised fine-tuning and reinforcement learning.**

The paper's own framing, which the research verified from its tables (Shao et al., 2024, preprint):

- DeepSeekMath-Base 7B "achieves comparable performance with Minerva 540B" — Google's 540-billion-parameter math specialist. Seven billion versus five hundred forty billion, and it's a contest.
- The instruct version "beats all 7B counterparts and is comparable with 70B open-source instruction-tuned models."
- After the RL phase, the model reached **51.7% on MATH** (a competition-mathematics benchmark, top-1 accuracy) and **88.2% on GSM8K** (grade-school math word problems) — the RL phase alone lifted GSM8K from 82.9 to 88.2 and MATH from 46.8 to 51.7. The paper describes this as "approaching the performance level of Gemini-Ultra and GPT-4" on MATH, *without* toolkits or majority voting.

The recipe matters because it shows mechanisms stacking. First, continued pretraining on 120 billion math tokens — mechanism (3), soaking the model in

the domain until mathematics is its native language. Then math instruction tuning — mechanism (1), teaching it the *format* of mathematical answers. Then GRPO, a reinforcement-learning method that rewards correct reasoning chains — the model practices problems and gets reinforced for the ones it solves, like a student doing homework with an answer key. Each stage is a different mechanism, and the combination is what reaches 540B-class performance at 7B.

One number the paper reports that this book deliberately does *not* use as a head-to-head claim: a 60.9% figure under majority-vote sampling at 64 attempts (maj@64). Different inference budgets are not comparable scores — letting one model try 64 times and another try once tells you nothing about which is better. The book keeps its comparisons apples-to-apples, and says so when it declines a tempting number.

Orca-2: teaching a 13B model to reason, beating a 70B chat model

The Orca project at Microsoft Research asked a different question: what if you distill not answers, but *reasoning strategies*? Orca-2 took LLaMA-2 13B and fine-tuned it on examples generated by GPT-4 that demonstrate different solution strategies for different tasks — step-by-step reasoning here, extract-then-generate there, direct answers where appropriate. **Mechanism: (2) targeted distillation.**

The result (Microsoft Research, 2023, preprint): a macro-average of **66.92 over six reasoning benchmarks versus 57.59 for LLaMA-2-Chat-70B** — a model more than five times larger. The per-task gaps are wide: DROP (a reading-comprehension benchmark requiring discrete reasoning) **70.88 vs 54.11**; CRASS (causal reasoning) **87.59 vs 74.82**; RACE (reading comprehension) **79.16 vs 68.79**; GSM8K **65.73 vs 52.01**. The 66.92 figure used a "cautious" system message — an instruction telling the model to reason carefully rather than guess; without it the score was 64.49 — still far ahead of the 70B model.

What does "task-dependent reasoning strategies" mean concretely? Consider two problems. A math word problem rewards step-by-step working: define variables, compute, check. A reading-comprehension question rewards a

different move: scan the passage, extract the relevant sentence, then answer directly. Most models apply one style to everything. Orca-2 was trained to *choose* — to recognize which strategy a task wants and switch. That metacognitive flexibility, distilled from GPT-4's demonstrations, is what the 70B chat model lacked despite its size. The big model knew more; the small model reasoned better. On reasoning benchmarks, reasoning better wins.

And the honest ceiling: the same paper's comparisons put GPT-3.5 Turbo at 67.65 on those reasoning benchmarks and GPT-4 at 79.03. So Orca-2 comes within three-quarters of a point of the previous generation's flagship — a genuine David-and-Goliath near-miss for a 13B model — and still trails the current frontier by twelve. Distillation moves the small model up a weight class; it does not abolish weight classes. That is the conditional thesis, stated in numbers.

Phi-2: the "data beats scale" flagship

Phi-2, at 2.7 billion parameters, is this book's purest demonstration that data quality can substitute for scale. Microsoft's official report (Javaheripi, Bubeck et al., Microsoft Research, 2023 — a company publication): trained on 1.4 trillion tokens of textbook-quality and synthetic data. **Mechanism: (3). Not fine-tuned. Say it plainly: nobody fine-tuned anything here; the model was simply raised on better food.**

First-party numbers from Microsoft's report: on BBH (Big-Bench Hard, a suite of difficult reasoning tasks) Phi-2 scored **59.3 versus 42.4 for Gemini Nano 2** (a 3.2B model); on MBPP **59.1 versus 27.2**; on MMLU (a broad knowledge benchmark) **56.7 versus 55.8**. Secondary reporting of the same work adds: BBH with 3-shot chain-of-thought **59.2 versus 40–47.8 for Llama-2 7B/13B**; a HumanEval+MBPP aggregate of **53.7 versus 39.4 for Mistral-7B**; GSM8K **61.1 versus 46.4 for Mistral-7B** — flag those as secondary-reported, since they come via press rather than the report's own tables.

The scale of the training is worth absorbing: 1.4 *trillion* tokens. Phi-2 is "small" only in parameters — in training data, it was fed enormously. That's the deep lesson of mechanism (3): parameter count is not the only axis of scale. A small model trained on vast quantities of excellent, narrowly-chosen data can outperform a larger model trained on vaster quantities of average data.

When someone tells you "bigger model, better model," the correct question is: *better data, or just bigger?* Phi-2 is the existence proof that the data axis matters more than the parameter axis for task performance.

There's a coda to the phi story that belongs here. The research direction phi-1 opened — synthetic data generation, textbook-quality filtering, curriculum design — is now one of the most active areas in the field, and the data-curation playbook it pioneered shows up in nearly every small-model success since. The giants all do it now. The small models just did it first, because they had to: when you can't win on size, you win on data.

A 2.7B model outperforming 7B and 13B generalists across reasoning, coding, and knowledge benchmarks is the single strongest "it's the data, not the size" datapoint in the literature. It also previews Chapter 6's unifying theory: specialization is mostly a data story.

The pattern — and what the giants still do better

Six cases, three mechanisms, one pattern. Code, math, and reasoning are *narrow, well-defined, machine-checkable* domains — exactly the conditions the conditional thesis requires. The wins come from three different routes: build the small model on better data from birth (phi-1, StarCoder, Phi-2, partly DeepSeekMath), teach it the teacher's reasoning strategies (Orca-2), fine-tune it on validated synthetic instructions (StarCoder2), or simply fine-tune it on the task itself (the Together AI 8B case).

None of these models is "generally smarter" than the giants. Every one of them would lose a general-knowledge quiz to the frontier. But on the task they were built for, measured the way the task is actually measured, they win — and they run on hardware the giant can't touch. That asymmetry, between capability-on-task and cost-of-running, is the economic engine of the whole book. Chapter 7 prices it.

A final honest note before we leave this chapter. The giants still own the *frontier* of these fields: novel algorithm design, multi-file agentic coding tasks, competition mathematics nobody trained for. Zephyr-7B's own authors concede their model lags proprietary models on coding and math. The small models win the well-defined, well-measured middle — which happens to be where almost all commercial value lives. Nobody pays for "best possible novel

algorithm"; companies pay for "correct code, fast, cheap, private." The specialist's territory is the profitable territory.

An illustration: the build-vs-buy meeting

The following is illustrative — a composite scenario, not a real company.

Imagine you're the CTO of a mid-size fintech. Your developers want AI code assistance inside your proprietary codebase — code that legally cannot leave your servers. The default proposal: pipe everything to a frontier API at flagship prices, with a data-processing agreement and your fingers crossed. The specialist alternative: take an open 7–15B code model, fine-tune it on your own repositories (mechanism 1), and serve it on two GPUs in your own datacenter.

The evidence in this chapter says the fine-tuned specialist will write *your* code — your frameworks, your idioms, your internal libraries, which the giant has never seen — better than the general model, at a fraction of the running cost, with zero data leaving the building. The giant knows every public library ever written; your specialist knows *your* codebase. For the task you actually pay for, that's the contest that matters.

Key takeaways

- **Checkable domains give the hardest evidence.** Code runs, math has right answers — when a small model wins here, it's machine-graded fact, not opinion.
- **Code was the first proof.** phi-1 (1.3B) hit 50.6% on HumanEval; StarCoder-15B beat PaLM-540B 40.8 to 26.2; StarCoder2-15B-Instruct edged CodeLlama-70B-Instruct 72.6 to 72.0.
- **Plain fine-tuning works too.** A fine-tuned Llama-3-8B beat Llama-3-70B (65.2% vs 64.2%) on a custom math set — the most replicable win in the chapter.
- **Math follows the same pattern.** DeepSeekMath-7B rivals 540B-class models (51.7% on MATH, 88.2% on GSM8K) via curated math pretraining plus tuning.

- **Reasoning strategies distill.** Orca-2-13B beat LLaMA-2-Chat-70B 66.92 to 57.59 by learning *how* to reason from GPT-4-generated examples — while sitting just below GPT-3.5 Turbo's 67.65 and well behind GPT-4's 79.03.
- **Data quality is a mechanism, not a footnote.** phi-1 and Phi-2 were never fine-tuned; they were raised on textbook-quality data. Label mechanism (3) honestly whenever you see it.
- **The specialist's territory is the profitable territory.** Giants keep the frontier; specialists win the well-defined middle where commercial value lives.

Sources

- Gunasekar, S. et al., Microsoft Research. *Textbooks Are All You Need* (phi-1). 2023. <https://arxiv.org/abs/2306.11644>
- Li, R. et al., BigCode (Hugging Face + ServiceNow). *StarCoder: may the source be with you!* 2023. Paper: arXiv:2305.06161. Authors' blog report: <https://huggingface.co/blog/starcoder>
- BigCode. *StarCoder2-15B-Instruct* self-alignment report. 2024. Authors' blog: <https://huggingface.co/blog/sc2-instruct>
- Shao, Z., Wang, P., Zhu, Q., et al., DeepSeek-AI. *DeepSeekMath: Pushing the Limits of Mathematical Reasoning in Open Language Models*. Preprint, 2024. <https://arxiv.org/html/2402.03300v3>
- Microsoft Research. *Orca 2: Teaching Small Language Models How to Reason*. Preprint, 2023. <https://arxiv.org/pdf/2311.11045>
- Javaheripi, M., Bubeck, S., et al., Microsoft Research. *Phi-2: The Surprising Power of Small Language Models*. Company publication, December 2023. <https://www.microsoft.com/research/blog/phi-2-the-surprising-power-of-small-language-models/>

4. The Specialists: Language, Retrieval, Speech

Language is the generalist's home turf. A frontier model reads everything, so surely it handles every language task better than a specialist? The evidence says otherwise — in the right slices. This chapter covers named-entity recognition, translation, retrieval-augmented question answering, and speech: four domains where small, specialized models have beaten the giants, plus one case that teaches the most important lesson in the book — *which slice* of the benchmark you win on.

A warning before we begin: this chapter's wins are *sliced* wins. In Chapter 3, the specialists won broadly across their benchmarks. Here, each model wins the slice its training data covered and often loses adjacent slices. That isn't a defect in the evidence — it is the conditional thesis in its purest form, and by the end of this chapter you'll see why sliced wins are exactly what a business should want.

A note on the measures

Two metrics recur. **F1** is the standard accuracy score for extraction tasks — it balances precision (of the things the model flagged, how many were right?) against recall (of the things it should have flagged, how many did it find?). **WER**, word error rate, is the speech-recognition error measure — lower is better. **COMET** is a learned metric for translation quality that correlates with human judgment better than old word-overlap scores. You don't need the formulas; you need to know that higher F1 and COMET are better, lower WER is better.

UniversalNER: a 7B model beats ChatGPT at finding names in text

Named-entity recognition (NER) means finding the names in text — people, companies, drugs, genes — and labeling what kind of thing each one is. It's the unglamorous plumbing under search, compliance, and biomedical research.

"Open" NER means doing it for entity types the model wasn't explicitly programmed for.

Researchers distilled ChatGPT itself into a student: ChatGPT generated 45,889 instruction examples for NER, and LLaMA 7B and 13B models were instruction-tuned on them (Zhou et al., 2024 — peer-reviewed at ICLR 2024).

Mechanism: (2) targeted distillation.

The student beat the teacher. UniNER-7B averaged **41.7 F1 versus ChatGPT's 34.9** across 43 datasets spanning 9 domains. (F1 is the standard accuracy measure for extraction tasks, balancing precision and recall.) The domain slices show where specialization bites hardest: biomedical **51.2 vs 38.1**, programming **30.5 vs 11.6**, finance **60.0 vs 52.8**.

Pause on what happened here, because it's delicious: the giant generated the training data, the student trained on it, and the student outperformed the giant at the specific task. Distillation doesn't just copy capability — on a narrow task, the focused student can *exceed* the distracted teacher. This is one of the cleanest cases in the literature, peer-reviewed, with exact scores verified.

The business consequence is immediate. NER is a compliance and data-pipeline task — banks scanning filings, pharma scanning literature. A 7B model doing it better than the flagship API, running inside your own infrastructure, turns a per-token API expense into a fixed hardware cost.

Why the student can beat the teacher

The UniNER result puzzles people: how can a model trained on ChatGPT's outputs outperform ChatGPT? The intuition is straightforward once you see it. The giant is a generalist — its capacity is spread across everything from poetry to physics. When you distill it for one task, you concentrate all of the student's capacity on that task. The teacher's knowledge of NER was a small corner of a vast mind; the student's entire mind *is* NER.

There's a second, subtler effect. The 45,889 training examples were *curated for the task* — consistent formatting, clear entity boundaries, the exact output shape the benchmark rewards. The giant, answering ad hoc, never committed to that discipline. The student did. Specialization isn't just knowing more about the task; it's being *shaped* for the task, down to the format of the answers.

This is why mechanism (2), targeted distillation, keeps appearing in this book's wins. The teacher provides the raw material — reasoning traces, labeled examples, critiques — and the focused student converts it into task-shaped capability. The giant remains the better *source* of knowledge; the specialist becomes the better *instrument* for the job.

TowerInstruct: translation is a specialist's game

Translation looks like a generalist stronghold — it requires knowing two languages deeply. But the TOWER project (Unbabel, 2024 — peer-reviewed at COLM 2024) showed that continued multilingual pretraining plus translation-task fine-tuning produces a 7B specialist that beats giants at translation-related tasks. **Mechanisms: (3) continued domain pretraining plus (1) supervised fine-tuning.**

Why should a specialist win at translation at all? Because professional translation isn't "knowing two languages" — it's a bundle of narrow skills: recognizing named entities across languages, fixing machine output, preserving terminology consistency, handling the quirks of specific language pairs. The TOWER team didn't try to make a better generalist. They took a base model, continued pretraining it on multilingual translation-heavy data (mechanism 3: soak it in the domain), then fine-tuned it on curated translation tasks — their "TowerBlocks" dataset (mechanism 1: shape it for the job). Each narrow skill got its own training emphasis.

The numbers: on aggregated translation-related NER, TowerInstruct-7B scored **71.68 F1 versus 59.88 for GPT-4 and 44.62 for LLaMA-2-70B**. On automatic post-editing — the task of fixing machine-translated text, measured by COMET-22 — it scored **82.69 versus 78.34 for LLaMA-2-70B and 81.47 for GPT-3.5**.

For a translation company, the implication is concrete: the post-editing step — where human translators fix machine output, the most labor-intensive part of the pipeline — can be assisted or automated by a 7B model that runs on a single GPU, rather than a flagship API call per segment. At millions of segments, that's the difference between a viable product margin and a bonfire of tokens.

Now the honest caveat, because this is where the book earns its credibility: on that same post-editing task, **GPT-4 scored 85.20 — still ahead.** TowerInstruct beats the 70B open model and edges GPT-3.5, wins translation NER decisively including against GPT-4, but does not take every slice. The lesson, which Chapter 8 develops fully: always ask *which slice* of the benchmark the small model wins. A specialist's victory is real and valuable even when it's partial — a translation company that post-edits with a 7B model at 82.69 instead of paying flagship prices for 85.20 is making a rational trade, not a mistake.

Self-RAG: a 7B model that knows when to look things up

Retrieval-augmented generation (RAG) is the technique of letting a model search external documents before answering, so it doesn't have to memorize the world. Self-RAG (Asai et al., 2023) went further: a LLaMA-2 7B model trained on 150,000 instruction pairs with special *self-reflection tokens* — the model learns to decide for itself when to retrieve, what to retrieve, and how to critique its own answer. The critic labels came from GPT-4. **Mechanism: (1) specialized training objective plus a specialized inference procedure** — not just different weights, a different way of thinking.

Picture the difference in practice. A standard RAG pipeline retrieves documents on every query — wasteful when the model already knows the answer, and noisy when retrieval returns junk. Self-RAG's model first emits a reflection token: *do I need to look this up?* If yes, it retrieves, then emits another: *is this passage relevant?* Then it drafts, then critiques: *is this claim supported by what I retrieved?* The model was trained, with GPT-4's critiques as supervision, to run this internal quality-control loop. What got distilled wasn't knowledge — it was *discipline*: the habit of checking.

That's why the PopQA result is so telling. PopQA is built from obscure, long-tail entities — exactly the questions where memorization fails and retrieval discipline matters. The 7B specialist's trained habit of looking things up and checking its work beats the giant's vast but undisciplined memory, 54.9 to 29.3. And it even beats the giant *with* retrieval bolted on (50.8), because bolted-on retrieval without the discipline to use it well is just noise with extra steps.

On factual question answering, the results are striking (paper's Table 2, exact scores verified):

- **PopQA accuracy: Self-RAG-7B 54.9 versus ChatGPT 29.3** — and versus ChatGPT *with* retrieval added (Ret-ChatGPT), 50.8. The 7B specialist beats the giant whether or not the giant gets to look things up.
- **PubHealth** (health fact verification): **72.4 versus 70.1**.
- **Biography generation** (factuality score): **81.2 versus 71.8**.
- **ASQA citation precision: 66.9 versus 65.1** for Ret-ChatGPT.

And the losses, reported with equal prominence because the thesis is conditional: **TriviaQA 66.4 versus 74.3**, **ARC** (a reasoning benchmark) **67.3 versus 75.3**, **ASQA exact-match 30.0 versus 35.3**. On broad trivia and general reasoning, the giant's breadth reasserts itself. On long-tail factual questions — PopQA is specifically built from obscure entities — the specialist's trained retrieval discipline wins.

For a founder, this is the most actionable case in the chapter. A factual-QA product (support bot, research assistant, compliance lookup) built on a 7B Self-RAG-style model can *outperform* the flagship API on exactly the questions users actually ask — the obscure ones — while costing a fraction to run. The giant wins at pub trivia; the specialist wins at your customers' questions.

ChatQA: the honest caveat, stated up front

ChatQA (Liu et al., NVIDIA, 2024, preprint) is a conversational RAG system built for multi-turn question answering over documents — the "ask follow-up questions about this report" use case. Its recipe: two-stage instruction tuning on LLaMA-2 (first general instruction-following, then conversational QA specifically) plus a conversational dense retriever that understands dialogue context, not just single queries. On the ChatRAG benchmark it scored **54.14 versus 53.90 for GPT-4-0613 and 54.03 for GPT-4-Turbo** — beating the flagship at its own game. **Mechanism: (1) specialization.**

The two-stage recipe is worth noting because it's the standard serious-fine-tuning pattern you'll see in Chapter 10's playbook: stage one teaches the model to follow instructions in general; stage two specializes it for the job. Skip stage one and the model is brittle; skip stage two and it's generic. The order matters.

The caveat, stated before the celebration: **ChatQA-1.0 is a 70B model. It is not small.** This is a *specialization* win — a weaker open base, adapted for one job, beating the best general model at that job — not a small-beats-large win. The book includes it because the mechanism is the same one the thesis is about, and because it proves the point at the high end: even at 70B, targeted adaptation beats general scale on a narrow task. But it does not belong in the "small model" column, and the book doesn't put it there.

(Note for the careful reader: the ChatQA preprint is arXiv:2401.10225, verified 2026-09-21; see Sources.)

Distil-Whisper: the student beats the teacher — on one slice

Speech recognition gives us the book's most instructive almost-win. Distil-Whisper (Gandhi et al., 2023, preprint) distilled OpenAI's Whisper large-v2 — a 1.55-billion-parameter speech model — into a 756M-parameter student using large-scale pseudo-labeling. The technique is simple in concept: the teacher transcribes vast amounts of audio, and the student trains on those transcripts as if they were ground truth. No human labelers, no new data collection — the teacher's knowledge becomes the student's curriculum. **Mechanism: (2) distillation.** The student is 49% smaller and $5.8\times$ faster.

On long-form transcription, measured by word error rate (WER — lower is better), the student wins: **11.6 versus the teacher's 11.7**. A smaller, faster model that transcribes long audio *more* accurately than the model that taught it.

Why would the student beat the teacher at all? The distillation training emphasized long-form audio — the pseudo-labeled curriculum was rich in the long recordings where the teacher's own training was thinner. The student got, in effect, extra homework on exactly the slice where the teacher was weakest, and it shows: the pupil surpasses the master on the master's weak chapter.

But on short-form audio, the teacher wins: **10.1 for the student versus 9.1 for the teacher** — a clear regression. The win is distribution-dependent: the distillation data emphasized long-form audio, so the student specialized in exactly that slice and paid for it elsewhere. Extra homework on one chapter means less study time for the others.

This is the single most important cautionary pattern in Part II, and it generalizes. TowerInstruct wins translation NER but not all of post-editing. Self-RAG wins PopQA but loses TriviaQA. Distil-Whisper wins long-form but loses short-form. Every specialization win in this book is a win *on a slice* — and the slice is defined by the training data. Chapter 8 is built on this observation: fine-tuning doesn't make a model better in general; it makes it better at the distribution it was fed.

An illustration: the compliance pipeline decision

The following is illustrative — a composite scenario, not a real company.

A bank's compliance team must scan 50,000 earnings-call transcripts a year, extracting every mentioned executive, subsidiary, and financial figure, then flagging risk phrases. The default proposal: send it all to a flagship API. The cost is per-token, the data leaves the bank's network (a non-starter for the legal department in several jurisdictions), and the flagship — brilliant at everything — is merely good at this one narrow extraction format.

The specialist alternative, straight from this chapter's evidence: distill the task into a 7B model (mechanism 2, the UniNER recipe), run it inside the bank's own VPC. Better extraction scores on the bank's entity types, fixed hardware cost instead of a meter that runs forever, data that never leaves the building. The giant wins at everything in general; the specialist wins at the one thing the bank actually buys. That is not a compromise — it's the better product decision, and the benchmark scores say so.

The slice is the product

Step back and look at the chapter's shape. UniNER wins NER across 43 datasets — the broadest win here, and notably the one with the most task-shaped training data. TowerInstruct wins translation NER decisively but splits post-editing with GPT-4. Self-RAG dominates long-tail factual QA and loses broad trivia. Distil-Whisper takes long-form speech and surrenders short-form.

Every one of these is a real, measured, commercially valuable win — and every one is bounded by its training distribution. The naive reading ("small models beat big models") is false. The precise reading is the book's thesis: *on*

narrow, well-defined, in-distribution tasks, with good adaptation data and task-aligned evaluation, the specialist wins. The slice isn't a limitation of the evidence. The slice is what you're buying.

For the non-technical reader, here's the strategic translation: you don't need a model that's better at everything. You need one that's better at the thing you sell, on the data you actually have, measured the way your customers experience it. The giants sell generality. Specialists sell the slice. Most businesses live on a slice.

Key takeaways

- **Distilled students can beat their teachers.** UniNER-7B (41.7 vs 34.9 F1 across 43 datasets) was trained on ChatGPT-generated data — and outperformed ChatGPT at NER. Peer-reviewed at ICLR 2024.
- **Concentration beats breadth on a narrow task.** The student wins because its entire capacity is shaped for one job, in one format — the teacher's knowledge was a corner of a vast mind.
- **Translation rewards specialization.** TowerInstruct-7B beat GPT-4 at translation-related NER (71.68 vs 59.88) — while GPT-4 still led on post-editing overall (85.20 vs 82.69). Report the slice, not just the win.
- **Trained retrieval discipline beats raw scale.** Self-RAG-7B beat ChatGPT with and without retrieval on PopQA (54.9 vs 50.8/29.3) — and lost on broad trivia (TriviaQA 66.4 vs 74.3).
- **Label size honestly.** ChatQA beats GPT-4 at conversational RAG, but at 70B it's a specialization win, not a small-model win.
- **Every win is a win on a slice.** Distil-Whisper beats its teacher on long-form transcription (11.6 vs 11.7 WER) but regresses on short-form (10.1 vs 9.1). The training distribution defines the victory — and its boundary. The slice is the product.

Sources

- Zhou, W. et al. *UniversalNER: Targeted Distillation from Large Language Models for Open Named Entity Recognition*. ICLR 2024 (peer-reviewed). <https://ar5iv.labs.arxiv.org/html/2308.03279>
- Unbabel. *TOWER: An Open Multilingual Large Language Model for Translation-Related Tasks*. COLM 2024 (peer-reviewed). <https://openreview.net/pdf?id=EHPns3hVkj>
- Asai, A., Wu, Z., Wang, Y., Sil, A., Hajishirzi, H. (UW / Allen Institute for AI / IBM). *Self-RAG: Learning to Retrieve, Generate, and Critique through Self-Reflection*. 2023. <https://arxiv.org/html/2310.11511v1>
- Liu, Z. et al., NVIDIA. *ChatQA: Surpassing GPT-4 on Conversational QA and RAG*. Preprint, 2024. <https://arxiv.org/pdf/2401.10225>
- Gandhi, S. et al. *Distil-Whisper: Robust Knowledge Distillation via Large-Scale Pseudo Labelling*. Preprint, 2023. <https://arxiv.org/pdf/2311.00430.pdf>

5. How Fine-Tuning Actually Works

So far this book has been about *what* specialized small models can do. Now we turn to *how* — what actually happens when you fine-tune a model, and why the modern tricks for doing it (LoRA, QLoRA, distillation) changed the economics so completely that a laptop-class GPU can now produce a model that beats a giant.

If you take one idea from this chapter, make it this: fine-tuning is not rebuilding the model. It is teaching a trained model a new job, and in the last few years we learned how to teach the job while changing almost nothing inside it.

What a "base model" is

Every model in this book started as a **base model**: a neural network trained on enormous amounts of text with a single, almost comically simple objective — predict the next word (technically, the next **token**, a chunk of text roughly a word or less). This phase is called **pretraining**. The model reads trillions of tokens — web pages, books, code — and adjusts billions of internal numbers, called **parameters** or weights, until it becomes very good at continuing text.

Pretraining is where nearly all of a model's knowledge comes from. It is also brutally expensive: training a frontier-scale base model costs tens of millions of dollars in compute. But once trained, the base model is a kind of compressed library of everything it read — plus, crucially, a *general* engine for continuing text in any style, on any topic.

The problem: a raw base model is not an assistant. Ask it a question and it may continue your question with ten more questions, or answer in the style of a 19th-century novel, or just trail off. It knows things; it does not know how to *behave*. Everything in this chapter is about closing that gap — cheaply.

Instruction tuning: teaching the format

The first big step was **instruction tuning**, introduced at scale by Wei et al. (2022) in the FLAN work (preprint). The idea is disarmingly simple: take the

pretrained model and continue training it — this is **fine-tuning**, training an already-trained model further — on a collection of tasks written out as plain instructions.

Instead of raw text, the model sees thousands of examples like:

Instruction: Translate this to French: "The cat sat on the mat." Response: "Le chat s'est assis sur le tapis."

Instruction: Summarize this customer review in one sentence... Response: ...

Each example is a pair: an instruction and a desired response. During fine-tuning, the model reads the instruction, makes its best guess at the response, compares its guess to the real one, and nudges its parameters to do better next time. That nudge — computed by an algorithm called **gradient descent** — is the entire machinery of learning here: guess, measure the error, adjust. Repeat millions of times.

The result was striking: models fine-tuned this way became dramatically better at handling *unseen* tasks — tasks they were never explicitly trained on — simply because they had learned the general skill of "read the instruction, do the thing." InstructGPT, the model behind the original ChatGPT, used roughly **13,000 human-written demonstrations** for this stage (Ouyang et al., 2022, preprint).

Here is the conceptual point that matters for this book: instruction tuning mostly teaches the model *how to express knowledge it already has*. The facts were baked in during pretraining; instruction tuning is a formatting layer. Hold that thought — Chapter 6 is built on it, and it is why a thousand excellent examples can do the work of a million mediocre ones.

Supervised fine-tuning: teaching the job

Supervised fine-tuning (SFT) is the same machinery pointed at one job instead of many. You collect examples of the exact task you care about — say, 136,609 financial instruction examples, which is what the FinMA team used to build their finance specialist (Xie et al., 2023, preprint) — and train the model on those, and only those.

What actually gets trained? In classic full fine-tuning, *everything*: every one of the model's billions of parameters is allowed to move. That is why it is expensive. Training doesn't just store the model's weights; it stores a whole apparatus around them. The research behind this book puts the cost at roughly **16 bytes per parameter**: 2 bytes for the weight itself, 2 for the gradient (the "which direction to nudge" signal), and the rest for the **optimizer state** — the training algorithm's running notes on each parameter (its recent history of nudges, kept so the algorithm can adapt its step size). Multiply by 7 billion parameters and you get about **112 GB of GPU memory** to fully fine-tune a 7-billion-parameter model. No single consumer GPU holds that. Full fine-tuning was, until recently, a rich lab's game.

This is the wall that LoRA knocked down.

LoRA: teaching with a sticky note

LoRA — Low-Rank Adaptation (Hu et al., 2021, preprint, from Microsoft) — starts from a simple observation about what fine-tuning actually changes inside a model.

Imagine each layer of the model as a huge spreadsheet of numbers — a matrix with, say, 4,096 rows and 4,096 columns (about 16.8 million numbers). Full fine-tuning rewrites the entire spreadsheet. But when researchers looked at *what the rewrite actually accomplished*, they found something surprising: the useful part of the change — the part that made the model better at the new task — could be captured by a far smaller change. The adaptation lives in a low-dimensional subspace. The rest of the spreadsheet barely matters.

LoRA exploits this directly. Instead of rewriting the spreadsheet W , you **freeze it** — lock every number in place — and learn only a small *correction*, written as the product of two skinny matrices:

$$\Delta W = BA$$

Read that in words: the update (ΔW) is built by multiplying a tall skinny matrix B by a short wide matrix A . The "skinny" dimension is the **rank** r , and it is tiny — typically 8 or 16, compared to 4,096. If the full spreadsheet has d^2

numbers, the LoRA correction has only about $2 \times d \times r$. At $d = 4,096$ and $r = 8$, that is a **256-fold reduction** in numbers to learn.

The engineering is elegant: B starts as all zeros and A starts random, so at the very beginning the correction is exactly zero — training starts from the *exact* pretrained model, and the correction grows only where it helps. And once training is done, you can simply add the correction back into the frozen weights. The deployed model is the same size and speed as the original: **zero inference latency overhead**.

Now the concrete numbers, because this is where the economics flip. Take a 7-billion-parameter model and adapt its attention matrices (the parts of the model that decide what to pay attention to) with LoRA at rank 8, across all 32 layers. The trainable count is:

32 layers $\times$ 4 matrices $\times$ 8 $\times$ (4,096 + 4,096) $\approx$ 8.4 million parameters — about 0.12% of the model.

The original LoRA paper reported matching full fine-tuning on the GLUE language-understanding benchmark while training roughly **0.1%** of parameters. Memory-wise, the frozen base model in half precision needs about 14 GB, plus a small amount for the adapter's training state — roughly **14 GB total, versus ~112 GB for full fine-tuning**. That fits on a single high-end consumer GPU (a 24 GB card like an RTX 3090 or 4090). A technique published in 2021 quietly moved fine-tuning from "needs a data center" to "fits on a good desktop."

The best analogy — the one this book will keep using — is the **skill module**. Think of the frozen base model as a skilled generalist employee, and each LoRA adapter as a snap-on module: "now do finance," "now do medical triage," "now write in this company's voice." You keep one copy of the base model in memory and swap tiny adapters per task, per customer, per request. That is not just a training trick; it is a serving architecture — one base, many specialists — and it is the economic engine behind much of this book's thesis.

QLoRA: the same trick, compressed

QLoRA (Dettmers et al., 2023, preprint) asked: what if the frozen base model itself is compressed? **Quantization** means storing the model's numbers at lower precision — 4 bits per number instead of 16. It is like packing a library into a carry-on: you lose a little fidelity, but almost everything important survives, because neural network weights turn out to be forgiving of rounding.

QLoRA combines a 4-bit quantized base (using a format called NormalFloat, designed for the way weights are actually distributed) with LoRA adapters trained in higher precision on top. Two extra engineering tricks — double quantization (compressing the compression metadata) and paged optimizers (borrowing system memory to absorb spikes) — keep it stable.

The headline result, verified from the paper: fine-tuning a **65-billion-parameter model dropped from over 780 GB of GPU memory to under 48 GB, with no performance degradation** versus a full-precision baseline. The largest public models became fine-tunable on a single professional GPU.

The authors demonstrated it with Guanaco, a chatbot fine-tuned this way: **Guanaco-65B reached 99.3% of ChatGPT's score on the Vicuna benchmark after 24 hours on a single GPU**, and their 7B version — a 5 GB download — beat the 26 GB Alpaca model by more than 20 percentage points on the same benchmark. One honest caveat, which the book's evidence standard requires: the Vicuna benchmark is scored by GPT-4 acting as a judge of chat quality. It measures how good the *conversation* feels, not hard capability on exams or reasoning tests. A real result on a soft metric — impressive, and worth exactly that much.

The practical rule of thumb as of 2026: a 7B model fine-tunes via QLoRA on a consumer 24 GB card; 33B fits a single 24 GB card; 65B needs a single 48 GB card; 70B needs one 80 GB data-center card. The frontier of "what one machine can specialize" keeps moving down in cost.

Full fine-tuning vs. PEFT: the trade-offs

LoRA and QLoRA belong to a family called **PEFT** — parameter-efficient fine-tuning. Here is the honest comparison with full fine-tuning:

PEFT wins on cost (a fraction of the memory), speed, and flexibility. Because the base model stays frozen, you can serve dozens of task adapters from one copy of the weights. And a frozen base forgets less (more on forgetting in Chapter 8).

Full fine-tuning still matters in two situations. First, when the new domain is very different from what the model was pretrained on — a new language, say, or a radically different data distribution. LoRA's correction is deliberately small (capped by the rank r), and a small correction cannot drag a model across a large distribution shift. Second, when you want the absolute maximum quality on a single task and you have the GPUs to pay for it. Full fine-tuning remains the ceiling; PEFT is the near-ceiling at a tenth of the price.

For nearly every production use case in this book — task adaptation, style and format alignment, building one specialist among many — PEFT is the right default.

RLHF and DPO, briefly

Two more techniques complete the modern post-training toolkit. Both are about *preferences* — not just "give the right answer" but "give the answer a human would prefer."

RLHF (reinforcement learning from human feedback; scaled by Ouyang et al., 2022, preprint) works in three stages. First, supervised fine-tuning on human-written demonstrations (those ~13,000 examples). Second, train a separate **reward model**: show humans pairs of model outputs, have them pick the better one (~33,000 comparisons in the original), and train a model to predict human preference. Third, optimize the language model against the reward model using a reinforcement-learning algorithm (PPO), with guardrails to keep it from drifting too far from the supervised model.

The headline from the InstructGPT paper is one of the founding exhibits of this book's thesis: **the outputs of the 1.3-billion-parameter InstructGPT were preferred by human evaluators over the outputs of the 175-billion-parameter GPT-3** — a model more than a hundred times larger. Preference-tuned small beat raw large. Truthfulness roughly doubled on the TruthfulQA benchmark, and toxic outputs fell by about a quarter. The price: it is the

heaviest recipe in post-training — three models, a value head, and a famously finicky RL optimizer.

DPO — Direct Preference Optimization (Rafailov et al., 2023, preprint) — keeps the preference-learning idea and throws away the machinery. No separate reward model, no reinforcement learning: just a simple classification-style loss on pairs of (preferred, rejected) responses. The canonical small-model demonstration is **Zephyr-7B** (Tunstall et al., 2023, preprint): a Mistral-7B model with distilled supervised tuning plus distilled DPO reached **7.34 on MT-Bench**, beating Llama-2-Chat-70B (6.86), Falcon-40B-Instruct (5.17), and Guanaco-65B (6.41). Honesty requires the full table: Zephyr did *not* beat GPT-3.5-turbo on MT-Bench (7.94), though it did beat it on the AlpacaEval win-rate metric (90.60% vs. 89.37%) — and its own authors note it lags proprietary models on coding and math. Distillation plus preference tuning narrows the gap to the frontier; it does not close it.

What you now understand

You now know the full modern pipeline, the same one behind nearly every model in Chapters 2 through 4:

1. **Pretrain** a base model on trillions of tokens (expensive, done once, knowledge lives here).
2. **Instruction-tune** it to follow instructions (cheap, teaches format).
3. **Supervised fine-tune** it on your task's examples — with LoRA or QLoRA, training a fraction of a percent of the parameters on a single GPU.
4. Optionally, **align** it with human preferences via RLHF or DPO.
5. Optionally, **distill** the result into something even smaller (Chapter 6).

Every step after pretraining is some form of specialization. And the punchline of this chapter: steps 2 through 5 are *cheap* — cheap enough that a small team, a single GPU, and a few thousand good examples can build a specialist that beats a generalist costing orders of magnitude more to train and serve. The next chapter explains why the examples matter more than the GPU.

Key takeaways

- Fine-tuning teaches a trained model a new job; pretraining is where almost all knowledge comes from, and instruction tuning mostly teaches the model how to express what it already knows.
- LoRA freezes the base model and learns only a tiny low-rank correction ($\Delta W = BA$): ~8.4M trainable parameters (0.12%) on a 7B model at rank 8, matching full fine-tuning on standard benchmarks while dropping memory from ~112 GB to ~14 GB.
- QLoRA adds 4-bit quantization of the frozen base, cutting 65B-model fine-tuning from >780 GB to <48 GB with no quality loss — the largest public models became fine-tunable on a single GPU.
- Adapters are "skill modules": one frozen base plus many tiny swappable adapters is a serving architecture, not just a training trick.
- RLHF (1.3B InstructGPT preferred over 175B GPT-3) and its simpler successor DPO (Zephyr-7B beating 70B open models on MT-Bench) show preference tuning is another route by which small beats large — though neither closes the gap to the live frontier.
- Full fine-tuning remains the quality ceiling for large distribution shifts; PEFT is the near-ceiling at a fraction of the cost, and the right default for production specialization.

Sources

- Hu et al., "LoRA: Low-Rank Adaptation of Large Language Models," arXiv:2106.09685 (2021, preprint). <https://arxiv.org/abs/2106.09685>
- Dettmers et al., "QLoRA: Efficient Finetuning of Quantized LLMs," arXiv: 2305.14314 (2023, preprint). <https://arxiv.org/abs/2305.14314>
- Wei et al., "Finetuned Language Models Are Zero-Shot Learners" (FLAN), arXiv:2109.01652 (2022, preprint). <https://arxiv.org/abs/2109.01652>
- Ouyang et al., "Training language models to follow instructions with human feedback" (InstructGPT), arXiv:2203.02155 (2022, preprint). <https://arxiv.org/abs/2203.02155>

- Rafailov et al., "Direct Preference Optimization," arXiv:2305.18290 (2023, preprint). <https://arxiv.org/abs/2305.18290>
- Tunstall et al., "Zephyr: Direct Distillation of LM Alignment," arXiv: 2310.16944 (2023, preprint). <https://arxiv.org/abs/2310.16944>
- Xie et al., "PIXIU: A Large Language Model, Instruction Data and Evaluation Benchmark for Finance," arXiv:2306.05443 (2023, preprint). <https://arxiv.org/pdf/2306.05443>

6. Data Beats Scale

Here is the secret the last three chapters were keeping from you: none of those victories were really about clever algorithms. Meerkat-7B did not beat a 70-billion-parameter medical model because its architecture was smarter. FinMA-7B did not beat GPT-4 on financial classification because 7 billion is a magic number. They won because of *data* — what the models were trained on, how good it was, and how precisely it matched the task.

This chapter is the unifying theory of the book. Specialization is mostly a data story. Understand that, and the wins stop looking like miracles and start looking like engineering.

The thousand-example earthquake

In 2023, a team at Meta published a paper with a deliberately provocative title: *LIMA: Less Is More for Alignment* (Zhou et al., 2023, NeurIPS 2023 poster, preprint). They took a 65-billion-parameter LLaMA model and fine-tuned it on just **1,000 carefully curated** instruction-response pairs — hand-selected for quality, deduplicated, format-consistent. No reinforcement learning, no reward model, just plain supervised training on a thousand excellent examples.

Then they had humans compare its answers against the best models in the world. The results: LIMA's responses were judged **equivalent to or better than GPT-4 in 43% of cases**, better than Bard in 58%, and better than DaVinci003 — a model that *had* gone through expensive human-feedback training — in **65%**.

A thousand examples. Not a million. Not a hundred thousand. A thousand, chosen well.

The authors' thesis, stated plainly: **almost all of a model's knowledge is learned during pretraining; alignment is a thin formatting layer on top**. Pretraining builds the library; fine-tuning teaches the librarian how to answer. And teaching the librarian does not take much — if the lessons are good.

This reframes everything. The industry spent years assuming that better models required more data at every stage. LIMA says the *alignment* stage — the stage that turns a base model into something useful — is governed by

quality, not quantity. A thousand perfect demonstrations beat a million noisy ones. For a small team, that is the most liberating result in modern machine learning: you do not need Google's data centers. You need a thousand examples chosen with taste.

(One honest note on the numbers: "43% equivalent-or-preferred" includes ties — the paper's own framing, and the fair one. LIMA did not *beat* GPT-4 outright. It stood next to it, built from a thousand examples. That is astonishing enough.)

The three ways data does the work

Chapter 1 promised to keep three mechanisms separate, because most "small beats big" claims blur them. They are all data stories, but different *kinds* of data stories.

Mechanism 1: Supervised fine-tuning — the right examples

This is the most direct: take a base model, train it on examples of your exact task. FinMA, the finance specialist from Chapter 2, was instruction-tuned on **136,609 financial instruction examples** (Xie et al., 2023, preprint) — and then scored 0.86 F1 on financial sentiment classification where GPT-4 scored 0.78, and 0.98 on headline classification against GPT-4's 0.86. The examples were the product.

DeepSeekMath-7B (Shao et al., 2024, preprint) pushed the same idea one level deeper: instead of fine-tuning on task examples, the team *continued pretraining* on **120 billion tokens of curated math** from the web, then did math instruction tuning and reinforcement learning. The result: a 7B model whose base version was judged "comparable with Minerva 540B" — Google's 540-billion-parameter math specialist — and whose final RL version hit 88.2% on GSM8K and 51.7% on MATH. When your pretraining data is the domain, a small model can stand next to one nearly eighty times its size.

Mechanism 2: Targeted distillation — the teacher's worked answers

Distillation is an old idea with a simple shape: a large "teacher" model trains a small "student" model. The classic version (Hinton et al., 2015) has the student mimic the teacher's *soft* output probabilities — not just "the answer is B" but the full distribution over possible answers, including the informative near-misses. Hinton called this the teacher's "dark knowledge": the difference between a student who copies the answer key and one who studies the teacher's worked solutions.

The modern, sharper version distills *reasoning*, not just answers. Microsoft's Orca project (Mukherjee et al., 2023, preprint) trained a 13B student on **GPT-4's step-by-step explanation traces** — not "the answer is 42" but "here is how you think through the problem." Orca-2 (2023, preprint) went further, teaching different solution strategies per task, and the 13B model averaged **66.92 on six reasoning benchmarks against LLaMA-2-Chat-70B's 57.59** — within a point of GPT-3.5 Turbo's 67.65. Orca-Math (2024, preprint) fine-tuned Mistral-7B on **200,000 high-quality synthetic math problems** and reached **86.81% on GSM8K** — a **49-point gain** over the base model's 37.83% (Microsoft Research blog), beating LLaMA-2-70B and GPT-3.5. The base model did not get smarter in general. It got 200,000 excellent math lessons.

UniNER (Wenxuan Zhou et al., ICLR 2024) is the cleanest distillation case in this book: ChatGPT generated **45,889 instruction examples** for named-entity recognition, a 7B LLaMA student was instruction-tuned on them, and the student averaged **41.7 F1 against ChatGPT's 34.9** across 43 datasets — winning in biomedical (51.2 vs. 38.1), programming (30.5 vs. 11.6), and finance (60.0 vs. 52.8). The teacher's knowledge, focused through the lens of one task, fit inside a model a fraction of the size — and the focus made it *better* at that task than the teacher.

Distillation even produced the founding examples of the field. **DistilBERT** (Sanh et al., 2019) compressed BERT to **40% fewer parameters, 60% faster inference, while keeping 97% of its GLUE score** (77.0 vs. 79.5). **TinyBERT** (Jiao et al., 2019) went further — **7.5× smaller, 9.4× faster, 96.8% of BERT-base on GLUE** — by distilling at *both* pretraining and fine-tuning time, capturing general and task knowledge. DistilBERT remains one of the most

downloaded models on the Hugging Face Hub years later: distillation as product strategy, not paper trick.

Mechanism 3: Curated pretraining — better ingredients

Sometimes the "small" model was never fine-tuned at all. It was *born* specialized — trained from scratch on better, narrower data. This is the purest "data beats scale" story, and its flagship is Microsoft's phi series.

phi-1 (Gunasekar et al., 2023, preprint), at just **1.3 billion parameters**, was trained on 6 billion tokens of "textbook quality" filtered code plus 1 billion tokens of synthetic textbooks and exercises generated by GPT-3.5 — four days on 8 A100 GPUs, a modest run by industry standards. It scored **50.6% on HumanEval** and **55.5% on MBPP**, beating models ten times its size; in the paper's comparisons, only GPT-4 beat it. The title of the paper says it all: *Textbooks Are All You Need*.

Phi-2 (2.7B, Microsoft Research blog, December 2023) scaled the recipe to **1.4 trillion tokens** of textbook-quality and synthetic data. Microsoft's official numbers: **59.3 on the BBH reasoning benchmark against Gemini Nano 2's 42.4; 59.1 on MBPP against 27.2; 56.7 on MMLU against 55.8**. A 2.7B model trading blows with models far larger — not because it was fine-tuned, but because its entire education was curated.

StarCoder-15B (BigCode, 2023) tells the same story for code at a larger scale: specialized pretraining on The Stack plus Python adaptation, and **40.8 on HumanEval against PaLM-540B's 26.2**. The 540-billion-parameter generalist had read everything; the 15B specialist had read the *right* things.

Why narrow is learnable

Step back and the pattern is almost insultingly simple. A general model must be ready for anything: poetry, protein folding, tax law, flirting. Its capacity is spread across the entire space of human text. A specialist needs to master one neighborhood. And a neighborhood — financial sentiment, medical board questions, Python functions — is a *small world*: bounded vocabulary, recurring patterns, learnable structure.

Think of it as the difference between a general practitioner and a specialist surgeon. The GP knows a little about everything because patients walk in with anything. The cardiac surgeon knows one organ system deeply — and for heart surgery, you do not want the GP, no matter how brilliant. Fine-tuning (and distillation, and curated pretraining) is how you take a brilliant GP and put it through a cardiology residency. The residency is short. The GP's general medical education — pretraining — is what makes the short residency possible.

This is also why the thesis is *conditional*, and why this chapter ends where Chapter 8 begins. Specialization wins **in-distribution** — on tasks drawn from the same world as the training data. FinMA crushes financial sentiment classification and then collapses on numerical finance questions (0.06 exact match on FinQA against GPT-4's 0.63 — same model, same paper). Distil-Whisper beats its teacher on long-form speech and loses on short-form. The residency prepares you for cardiology, not for the patient who walks in with something no textbook covered. Always ask *which slice* of the benchmark the small model wins — because the slice is the whole story.

Distillation done right — and the imitation trap

There is a subtle line in this chapter that deserves emphasis, because crossing it is the most common way smart teams fool themselves. Good distillation transfers *capability*: reasoning traces, soft probability distributions, worked solutions. Bad distillation — call it **imitation** — transfers *style*: the surface mannerisms of a smart model without the substance.

Gudibande et al. (2023, preprint) demonstrated the trap cleanly: models from 1.5B to 13B parameters, trained on 0.3 million to 150 million tokens of imitation data from proprietary models, *looked* competitive to human raters — fluent, confident, well-formatted — but closed little of the actual capability gap on tasks the imitation data didn't cover. Style is cheap to copy; competence is not.

The Alpaca story is the same trap in miniature. Stanford's Alpaca-7B, in its original human evaluation, appeared to tie its teacher — **90 wins against 89** versus text-davinci-003. But the later automatic AlpacaEval leaderboard scored it at a **26.5% win rate** against the teacher's 50.0. Judge-dependence cuts both

ways, and "sounds like the big model" is not "thinks like the big model." Chapter 8 takes this apart in full. For now, the rule: distill the *reasoning*, not the *voice*.

Small data is cheap data

There is a business consequence to the data story that deserves its own section, because it is the reason small teams — not just small models — win.

Curating a thousand perfect examples is *human-scale work*. A domain expert can write or review a thousand instruction pairs in weeks. Generating 200,000 synthetic math problems costs one good teacher model and a validation script — the Orca-Math approach. Training phi-1 took four days on eight A100 GPUs, a run a well-funded startup can afford, not a nation-state. Compare that with pretraining a frontier model: tens of millions of dollars, thousands of GPUs, months of wall-clock time, and a data pipeline the size of a company.

This asymmetry is the whole game. The generalist's moat is pretraining scale — and it is real. But the specialist's moat is *curation taste*, and curation is cheap. When LIMA showed that a thousand excellent examples rival a million noisy ones, it didn't just make a scientific point; it repriced the market. The scarce resource stopped being compute and became judgment: knowing exactly what your task needs, writing it down precisely, and checking that every example teaches it.

It also changes who gets to compete. A hospital network with three good clinicians and a GPU workstation can build a triage specialist. A law firm with its own brief bank can build a contract-review specialist. Neither needs to out-train OpenAI, Anthropic, or Google — they need to out-curate them on one narrow slice of the world. Chapter 7 will put dollar figures on the serving side of this; the training side is already settled. Specialization is the first AI capability curve where the little guy's budget is enough.

The theory, stated plainly

Here is the unifying theory, in one paragraph:

Pretraining builds general capability; everything after pretraining is specialization; and specialization is governed by data quality, not data quantity or model size. A thousand perfect examples (LIMA) can align a model better than a million noisy ones. A teacher's worked reasoning traces (Orca, UniNER) can compress a giant's skill into a small student — and the focus can make the student *better* than the teacher on the target task. And a small model raised entirely on curated, textbook-quality data (phi-1, Phi-2) can outperform generalists many times its size, because it never wasted capacity on the parts of the internet nobody needs. Small wins when the data defines a learnable world and the evaluation measures that world honestly.

That is the whole book in a paragraph. The chapters ahead are about what it costs (Chapter 7), where it breaks (Chapter 8), where it still loses (Chapter 9), and how to decide (Chapter 10).

Key takeaways

- Specialization is mostly a data story: the wins in Chapters 2–4 came from what models were trained on, not from architectural cleverness.
- LIMA proved that 1,000 carefully curated examples can align a 65B model to near-frontier quality (43% equivalent-or-preferred vs. GPT-4, 65% vs. DaVinci003) — because pretraining holds the knowledge and alignment is a thin formatting layer.
- The three mechanisms are distinct data stories: supervised fine-tuning on task examples (FinMA, DeepSeekMath), targeted distillation of a teacher's reasoning (Orca, UniNER, DistilBERT, TinyBERT), and curated pretraining on better ingredients (phi-1, Phi-2, StarCoder).
- Narrow distributions are learnable: a specialist masters one bounded "small world" while the generalist's capacity is spread across everything — the cardiac surgeon versus the GP.
- The thesis is conditional because specialization is in-distribution: the same model that wins on its home slice (FinMA on sentiment, Distil-Whisper on long-form) can collapse outside it.
- Distill capability, not style: transferring reasoning traces and soft targets works; imitating surface mannerisms (the Gudibande trap, the Alpaca mirage) produces models that sound smart without being smart.

Sources

- Zhou et al., "LIMA: Less Is More for Alignment," arXiv:2305.11206 (2023, preprint; NeurIPS 2023 poster). <https://arxiv.org/abs/2305.11206>
- Gunasekar et al., "Textbooks Are All You Need" (phi-1), arXiv:2306.11644 (2023, preprint). <https://arxiv.org/abs/2306.11644>
- Javaheripi, Bubeck et al., "Phi-2: The Surprising Power of Small Language Models," Microsoft Research blog (Dec 2023, company publication). <https://www.microsoft.com/research/blog/phi-2-the-surprising-power-of-small-language-models/>
- Shao et al., "DeepSeekMath: Pushing the Limits of Mathematical Reasoning in Open Language Models," arXiv:2402.03300 (2024, preprint). <https://arxiv.org/html/2402.03300v3>
- Wenxuan Zhou et al., "UniversalNER: Targeted Distillation from Large Language Models for Open Named Entity Recognition," ICLR 2024 (peer-reviewed). <https://ar5iv.labs.arxiv.org/html/2308.03279>
- Mukherjee et al., "Orca: Progressive Learning from Complex Explanation Traces of GPT-4," arXiv:2306.02707 (2023, preprint). <https://arxiv.org/abs/2306.02707>
- "Orca 2: Teaching Small Language Models How to Reason," arXiv: 2311.11045 (2023, preprint). <https://arxiv.org/pdf/2311.11045>
- "Orca-Math," arXiv:2402.14830 (2024, preprint). <https://arxiv.org/abs/2402.14830> — base-model 37.83% figure via Microsoft Research blog: <https://www.microsoft.com/en-us/research/blog/orca-math-demonstrating-the-potential-of-slms-with-model-specialization/>
- Sanh et al., "DistilBERT, a distilled version of BERT," arXiv:1910.01108 (2019, preprint). <https://arxiv.org/abs/1910.01108>
- Jiao et al., "TinyBERT: Distilling BERT for Natural Language Understanding," arXiv:1909.10351 (2019, preprint). <https://arxiv.org/abs/1909.10351>
- Gandhi et al., "Distil-Whisper: Robust Knowledge Distillation via Large-Scale Pseudo Labelling," arXiv:2311.00430 (2023, preprint). <https://arxiv.org/pdf/2311.00430.pdf>

- Hinton et al., "Distilling the Knowledge in a Neural Network," arXiv: 1503.02531 (2015, preprint). <https://arxiv.org/abs/1503.02531>
- Gudibande et al., "The False Promise of Imitating Proprietary LLMs," arXiv:2305.15717 (2023, preprint). <https://arxiv.labs.arxiv.org/html/2305.15717>
- Li et al., "StarCoder: may the source be with you!," arXiv:2305.06161 (2023, preprint); BigCode/Hugging Face blog (2023, company publication). <https://huggingface.co/blog/starcoder>
- Xie et al., "PIXIU," arXiv:2306.05443 (2023, preprint). <https://arxiv.org/pdf/2306.05443>

7. The Economics of Small

The most expensive sentence in artificial intelligence is "just use the biggest model."

Chapters 2 through 4 made the quality case: on narrow, well-defined tasks, a fine-tuned small model can match or beat the generalist giant. Once quality is on the table, the conversation changes. It stops being about capability and starts being about money — and on money, small models don't just win. They win by margins that look like typos.

This chapter is written so that a CFO can read it alone, with no other chapter for support. Every price below was checked in **September 2026**. Model prices change constantly — providers cut them every few months — so treat every figure as a snapshot, not a constant. The ratios, though, have pointed in the same direction for years, and they keep getting more extreme.

Tokens are money

A quick definition, because everything in this chapter depends on it. A **token** is the unit models read and write in — roughly three-quarters of an English word. When you send a prompt to a model API, you pay per million *input* tokens (what you send); when the model answers, you pay per million *output* tokens (what it generates). Output tokens almost always cost more — typically three to ten times more — because generating text is the computationally expensive part.

That asymmetry matters enormously, and it is the first place small models compound their advantage. A fine-tuned specialist can be trained to give short, structured answers — a classification label, a JSON record, a one-line verdict — instead of the paragraph the generalist writes to be safe. Fewer output tokens at a lower price per token: the savings multiply.

The price table

Here is what a million tokens cost in September 2026, input/output, across the market (sources: [benchgecko/llm-pricing](#), [trestle](#), and [June–September 2026 API comparisons](#); check current listings before budgeting):

Model	Input \$/1M tokens	Output \$/1M tokens
GPT-4o (OpenAI)	\$2.50	\$10.00
Claude Opus 4.7 (Anthropic)	\$5.00	\$25.00
Claude Sonnet 4.6 (Anthropic)	\$3.00	\$15.00
Claude Haiku 4.5 (Anthropic)	\$1.00	\$5.00
GPT-4o-mini (OpenAI)	\$0.15	\$0.60
GPT-4.1 Nano (OpenAI)	\$0.10	\$0.40
Groq-hosted Llama 3.1 8B	\$0.05	\$0.08
Self-hosted 7–8B (your own GPU)	hardware cost only	hardware cost only

Read the ratios, not just the rows. GPT-4o-mini's output is about **17× cheaper** than GPT-4o's (\$0.60 vs. \$10.00 per million). A Groq-hosted 8B-class model is about **125× cheaper** on output (\$0.08 vs. \$10.00). That is not a discount. That is a different business.

Two caveats, because this book doesn't hide them. First, the cheapest options are not always the most dependable: one June-2026 comparison found an ultra-cheap provider undercutting everyone at roughly \$0.14/\$0.28, but flagged availability and reliability concerns for production use. Cheap is not the same as dependable, and uptime has a price too. Second, self-hosting is not "free" — it is hardware you rent or buy, plus engineering time. Its real economic property is different: once the GPU is running, the *marginal* cost of one more token is effectively zero. API pricing charges you per unit forever; a rented GPU charges you per hour regardless of how much you use it. High-volume workloads flip to self-hosting the way factories flip to owning equipment.

The worked example

Abstract ratios are easy to nod at and hard to feel. So here is a concrete traffic profile, taken from a June-2026 cost comparison: **100,000 customer conversations per month, averaging 3,200 input tokens and 2,400 output tokens each.**

On GPT-4o at September-2026 prices, the arithmetic is:

- Input: $100,000 \times 3,200 = 320$ million tokens $\rightarrow 320 \times \$2.50 = \mathbf{\$800}$
- Output: $100,000 \times 2,400 = 240$ million tokens $\rightarrow 240 \times \$10.00 = \mathbf{\$2,400}$
- **Total: ~\$3,200/month**

The same comparison ran identical traffic on a Flash-class small model and found **~\$96/month** — a **33× gap** on the same conversations, the same customers, the same answers (assuming, as Chapters 2–4 argue, the small model has been fine-tuned to do this job well).

Scale that to a year: roughly \$38,400 versus roughly \$1,150. The difference — about \$37,000 a year — is not a rounding error. It is a hire. It is the entire fine-tuning budget for the next three projects, with change. And this is a modest traffic profile; at a million conversations a month, the gap is the difference between a viable product and a dead one.

Notice what the worked example assumes, because assumptions are where cost projections die. It assumes the small model produces answers of similar length and quality — which is exactly what fine-tuning buys you, per Chapters 2–6. It assumes steady traffic; spiky, tiny-volume workloads may be cheaper on pay-per-token APIs than on a rented GPU sitting idle. And it assumes you count honestly: prompt-caching, batching, and output-length caps all move the numbers, and they move them *in favor* of the small model, which you control completely when you host it yourself.

The hidden multiplier: output length

There is a second, quieter compounding effect in the price table, and it deserves its own section because most teams miss it. Output tokens cost 3–10× more than input tokens at every provider — and the generalist model is *verbose by design*. Asked to classify a support ticket, GPT-4o will often write a

paragraph: its reasoning, its caveats, its polite framing. A fine-tuned specialist can be trained to emit exactly `{"category": "billing", "priority": 2}` and stop.

The arithmetic that follows is **illustrative** — round numbers to show the mechanism, not a quote. Suppose the generalist writes 400 output tokens per answer at \$10.00 per million: that's \$0.004 per answer. The fine-tuned specialist writes 80 tokens at \$0.08 per million: \$0.0000064 per answer. Five times fewer tokens at 125 times less per token compounds to roughly **625× cheaper per answer**. Even if you don't believe the exact ratio, the direction is undeniable: controlling output format is a cost lever that only exists when you control the model. API users can beg for brevity in the prompt; model owners enforce it in the weights.

This is also why "tokens per task" is the metric that matters, not "price per token." A cheap model that rambles is expensive. A fine-tuned model that answers in the exact shape your system consumes is cheap twice over.

Speed is money

Cost per token is only half the economics. The other half is speed — and speed, in production, converts directly into money.

In 2026 community benchmarks, an 8B-class model served on fast inference hardware runs at roughly **669 tokens per second on Groq and up to 2,198 tokens per second on Cerebras'** wafer-scale chips. A standard big model on GPUs manages roughly **50–100 tokens per second**. Time-to-first-token — how long the user waits before *anything* appears — clusters around 290–910 milliseconds for small/fast models (Gemini 2.5 Flash-Lite at 290 ms, GPT-4.1 at 910 ms in the same community table), while reasoning flagships can take tens of seconds. Treat single-source community figures as indicative rather than gospel, but the direction is unambiguous and the gap is an order of magnitude or more.

Why is latency economic and not cosmetic? Three reasons.

First, **user experience is revenue**. A voice agent that pauses four seconds before answering is unshippable; one that answers in 300 milliseconds feels human. A customer-support chatbot that streams its answer instantly gets used;

one that visibly "thinks" for ten seconds gets abandoned for the phone queue — the most expensive channel of all. Latency is the difference between a product people tolerate and one they choose.

Second, **agentic loops multiply latency**. Modern AI products don't call the model once; they call it dozens of times per task — plan, act, check, retry. A 10× per-call latency gap becomes a 10× task-duration gap, which becomes a 10× cost gap when you're paying for compute by the second, and a broken product when you're not. The cheapest model per token can be the most expensive model per *completed task* if it is slow — and small models are fast.

Third, **throughput is capacity**. At 2,198 tokens per second, one serving setup handles traffic that would need a fleet of big-model GPUs. Capacity planning for a product launch, a Black Friday spike, or a viral moment is dramatically simpler — and cheaper — when each unit of hardware does ten times the work.

The phone proves it

The most convincing economic argument for small models is sitting in your pocket. No frontier model will ever run interactively on a phone — the physics of batteries, heat, and memory forbid it. So the entire industry of personal AI is, by necessity, a distillation industry.

Apple's on-device foundation model, AFM-on-device, is a **3-billion-parameter model distilled from a 6.4-billion-parameter server model**, described in Apple's own 2024 research paper. The server model trained on 8,192 TPUV4 chips; the on-device one on 2,048 TPUV5p chips — the whole Apple Intelligence on-device stack is a distillation story. In March 2026, The Information reported (via MacRumors) that Apple secured access to Gemini inside its own data centers specifically to **distill it into smaller on-device models** for Siri — chain-of-thought traces from the big model becoming training data for the small one. Google ships Gemini Nano, in the 1.8–3.25B class, on Pixel phones for summarization, transcription, and smart replies.

Why do phones force this? Four reasons, all economic: no network latency, offline operation, **data never leaves the device**, and — the one that matters at billion-user scale — **zero per-query server cost**. Every Siri-style request answered on-device is a request Apple doesn't pay a data center to serve.

Multiply by billions of queries a day and the small model isn't a compromise; it is the only architecture whose unit economics work.

When the world's most secretive hardware company bets its flagship assistant on distillation, the thesis of this book has industrial proof.

The compliance discount

There is a whole enterprise AI market that the API-first story cannot serve at any price: hospitals bound by HIPAA, banks bound by financial regulation, law firms bound by client confidentiality, governments bound by sovereignty rules. In many jurisdictions and contracts, they **cannot legally send sensitive data to a third-party API** — not to the biggest model, not at any discount.

A fine-tuned 7–13B model running in their own virtual private cloud — or on-device — keeps every byte in-house while matching the big API on the narrow task it was trained for. This is not a niche. It is the segment of the market where the largest models are disqualified by procurement before benchmarks are even consulted. For these buyers, the small fine-tuned model isn't the cheap option. It is the *only* option. The compliance requirement becomes a moat: whoever can fine-tune well wins contracts that no API model can bid on.

Total cost of ownership

The honest objection to everything above is: "fine-tuning isn't free." True. So let's do the total-cost-of-ownership math properly — with a formula, not a single number, because compute prices shift too fast for any fixed figure to survive print.

Breakeven (months) = one-time fine-tuning cost ÷ monthly inference savings

The numerator has collapsed. Training a LoRA adapter on a 7B model touches about 8.4 million parameters — 0.12% of the model (Hu et al., 2021) — and fits on a single consumer GPU. QLoRA put 65B-parameter fine-tuning under 48 GB of GPU memory (Dettmers et al., 2023), and its authors trained their Guanaco-65B in 24 hours on a single professional GPU. Microsoft trained phi-1 — the 1.3B model from Chapter 3 — in four days on eight A100s

(Gunasekar et al., 2023). These are hundreds to low-thousands of dollars of compute, not research-lab budgets.

The denominator is the worked example above: thousands of dollars a month, every month, for as long as the product runs.

To make it concrete with clearly **illustrative** numbers (not quotes — get real ones before budgeting): suppose fine-tuning and evaluation cost \$2,000 all-in, and the fine-tuned model saves \$3,100 a month versus the frontier API on your traffic. Breakeven arrives in under a month. Everything after that is margin. Even if the fine-tuning costs ten times more, it pays back in half a year — and unlike API spend, the capability is an asset you own: no provider can raise its price on you, deprecate the model under you, or change its behavior without your consent.

There is one more line in the true TCO that rarely gets counted: **optionality**. A model you host is a model you can move — between clouds, between providers, on-premises, on-device. An API dependency is a supplier relationship with a single source. Procurement departments understand this math very well, even when engineering doesn't.

And the honest costs, because this chapter promised a CFO could read it alone: fine-tuning is not free beyond compute. Budget engineering time for data curation (Chapter 6's whole argument is that data quality decides everything), for building the evaluation harness (Chapter 10), and for ongoing maintenance — models drift as the world changes, and someone has to watch the metrics. None of these are reasons not to fine-tune; they are reasons to price the project like a project, not a API bill. The breakeven formula still applies — just put the fully loaded cost in the numerator.

The GPU vendor agrees with this book

The most striking endorsement of this book's thesis doesn't come from a startup trying to save money. It comes from NVIDIA — the company that sells the GPUs the giant models run on. In June 2025, researchers at NVIDIA Research (with Georgia Tech) published a position paper with a title that reads like this book's blurb: *Small Language Models are the Future of Agentic AI* (Belcak et al., arXiv:2506.02153).

Their argument comes in three claims: that small language models are **sufficiently powerful** for most agentic work, **inherently more suitable** for it operationally, and **necessarily more economical** than their giant counterparts. Be clear-eyed about what this paper is: a *position* paper, not an experiment. It compiles and synthesizes evidence rather than producing new numbers, and the individual results below come from their original sources, each with its own benchmark choices and caveats. But the compilation — assembled by the research arm of the company whose revenue depends on big-model compute — is the point. When the landlord tells you the building is too expensive, listen.

The economic core: serving a 7-billion-parameter small model costs **10–30× less** — in latency, energy, and compute — than serving a 70–175-billion-parameter model. The authors pair that with the industry's own numbers: the agentic AI sector was valued at about **\$5.2 billion** in late 2024, yet the cloud infrastructure built to serve its large models represented an estimated **\$57 billion** of capital investment — roughly a ten-to-one gap between what the infrastructure costs and what the market, so far, is worth.

The capability claims they compile read like a greatest-hits album for this book's first half:

- **Microsoft's Phi-2 (2.7B)** matched 30-billion-parameter models on commonsense reasoning and code generation while running roughly **15× faster**.
- **NVIDIA's own Nemotron-H models (2–9B)** matched dense 30B models of the same generation on instruction following and code generation at an order-of- magnitude fraction of the inference compute.
- **Hugging Face's SmolLM2 family (down to 135M parameters)** climbed toward 14B contemporaries on language understanding, tool calling, and instruction following — and matched 70B models of two years earlier.
- **Salesforce's xLAM-2 (8B)** reached state-of-the-art on tool calling, surpassing GPT-4o and Claude 3.5 on the benchmarks the paper cites.
- **DeepSeek's R1-Distill-Qwen-7B** beat Claude 3.5 Sonnet and GPT-4o on the paper's cited evaluations — a 7B distillation outperforming frontier generalists on their home turf.

- **DeepMind's RETRO (7.5B)**, augmented with retrieval, matched GPT-3 (175B) on language modeling with **25× fewer parameters**.

Their punchline for deployment is the architecture this book has been circling since Chapter 1: run small fine-tuned models by default, and escalate to the big generalist only for the calls that genuinely need it. The paper even sketches the migration recipe — log production traffic, curate 10,000–100,000 real examples, cluster them into recurring tasks, match each task to a small model, fine-tune with LoRA or QLoRA (or distill the big model's behavior into the small one), and keep a retraining loop running. Chapter 10 turns that recipe into your playbook.

The bottom line

The economics of small models rest on four compounding advantages: radically lower per-token prices (up to 125× gaps in September 2026), order-of-magnitude speed advantages that convert into product viability and capacity, zero-marginal-cost on-device deployment at scale, and access to regulated markets the big APIs cannot enter. Against those stands a one-time fine-tuning cost that modern methods have driven down to commodity levels.

None of this says big models are useless — Chapter 9 will say exactly where they earn their keep. It says that defaulting to the biggest model is a financial decision made without looking at the invoice. Look at the invoice.

Key takeaways

- **The price gaps are enormous and structural:** in September 2026, small-model API output cost up to 125× less than frontier-model output — and self-hosting drives marginal cost per token toward zero.
- **Worked example:** 100,000 conversations/month (3,200 input + 2,400 output tokens each) cost ~\$3,200/month on GPT-4o versus ~\$96/month on a Flash-class small model — a 33× gap on identical traffic.
- **Latency is economic:** 8B-class models serve at hundreds to thousands of tokens per second versus ~50–100 for big models; in agentic loops and real-time products, speed determines both cost per task and whether the product ships at all.

- **The phone is the proof:** Apple's 3B on-device model distilled from its 6.4B server model — plus the reported Gemini-to-Siri distillation deal — shows the industry's unit economics already run on small models at billion-user scale.
- **TCO favors fine-tuning:** one-time adaptation costs (hundreds to low thousands of dollars with LoRA/QLoRA) against thousands in monthly API savings means breakeven in weeks to months — plus an asset you own and a moat in regulated markets.

Sources

- Per-token prices (September 2026 snapshot): benchgecko/llm-pricing (weekly-updated), trestle pricing documentation, June–September 2026 API pricing comparisons, OpenRouter passthrough rates. Prices change frequently — verify current listings before budgeting.
- E. Hu et al., "LoRA: Low-Rank Adaptation of Large Language Models" (preprint, 2021). <https://arxiv.org/abs/2106.09685>
- T. Detrmers et al., "QLoRA: Efficient Finetuning of Quantized LLMs" (preprint, 2023). <https://arxiv.org/abs/2305.14314>
- S. Gunasekar et al., "Textbooks Are All You Need" (Microsoft Research, preprint, 2023). <https://arxiv.org/abs/2306.11644>
- Apple, AFM on-device technical report (Apple, 2024) — describes the 3B-parameter on-device model distilled from the 6.4B server model.
- The Information, via MacRumors (March 2026) — Apple/Google Gemini distillation arrangement for on-device Siri models.
- Throughput and time-to-first-token figures: 2026 community inference benchmarks (Groq, Cerebras) — indicative, methodology-dependent; treat exact figures as directional.
- P. Belcak et al., NVIDIA Research / Georgia Tech, "Small Language Models are the Future of Agentic AI" (position paper, preprint, June 2025, v2 Sept 2025). <https://arxiv.org/abs/2506.02153> — the V1/V2/V3 position, the 10–30× serving-economics figure, the \$5.2B/\$57B market-vs-infrastructure numbers, and the compiled small-model capability claims (Phi-2, Nemotron-H, SmolLM2, xLAM-2, R1-Distill-Qwen-7B, RETRO).

8. When Fine-Tuning Fails

Everything so far has argued that small, fine-tuned models win. This chapter argues the opposite — on purpose. A thesis that never loses is a sales pitch, and sales pitches are where expensive mistakes come from. The conditional version of the thesis ("small models win *when...*") is only as useful as your understanding of the *when-not*.

Each failure mode below is real, measured, and published. Each one also has a recognizable warning sign — something you can check *before* the failure costs you. Learn the signs and the thesis becomes a strategy instead of a slogan.

Forgetting: the model unlearns what you didn't retrain

Fine-tuning adjusts a model's weights to fit the new training data. The catch: weights are shared. The same parameters that encode general knowledge get pulled toward the narrow task, and skills you never retrained quietly degrade. This is **catastrophic forgetting**, and it is the oldest failure mode in neural networks — older than transformers, older than the current boom.

It was measured precisely in 2025. Researchers fine-tuned Llama-3-70B-Instruct — a strong, well-aligned 70-billion-parameter model — on popular public fine-tuning datasets (GenQA, OpenHermes, UltraChat, and similar), then re-tested it on general benchmarks. Average scores fell from **39.00 to between 35.91 and 38.81**, with the worst damage on MATH and GPQA — the reasoning-heavy tests (arXiv:2506.09428, 2025, preprint).

The interesting part is *why*. The driver wasn't fine-tuning itself; it was **distributional mismatch**. The public datasets over-represented generic chat and under-represented the reasoning-heavy mix the model was originally aligned on. When the researchers reconstructed a fine-tuning mix matching the original distribution, the fine-tuned model scored **39.21 — above the base model**. Forgetting, in other words, is not a law of nature. It is a data-mix problem, and it is fixable with the right mix.

The standard mitigations, all verified in practice: low learning rates, few training epochs, mixing some of the original pretraining-style data back in

("replay"), and parameter-efficient methods like LoRA — because a frozen base model, by construction, forgets less (most of its weights never move).

Then there is the variant nobody expects: **forgetting can erase safety, not just skills**. Qi et al. (2023) showed that fine-tuning aligned models on entirely *benign* utility datasets degraded their safety alignment — the models became broadly easier to jailbreak, even though nothing in the training data was harmful. Nobody intended it. The guardrails were simply stored in the same weights the fine-tuning moved.

Warning sign: after fine-tuning, test the model on skills *outside* the training data — including safety probes — not just on the target task. If general capability drops while task capability rises, your data mix is wrong, not your model.

The alignment tax

Making a model helpful, harmless, and honest costs something. The InstructGPT team (Ouyang et al., 2022) documented it plainly: optimizing for human preferences **regressed standard NLP benchmarks** — SQuAD, DROP, HellaSwag all slipped. They called it the **alignment tax**, and they mitigated it with PPO-ptx, a technique that mixes pretraining gradients back into the reinforcement-learning updates, leaving only minimal regressions.

The modern form of the tax is subtler, and it matters for this book's audience. Reinforcement learning from human feedback tunes the model for *generic* helpfulness — the kind of answer a typical rater likes. Nothing in that reward signal protects the long tail of narrow factual exactness: the drug-interaction edge case, the derivative-pricing formula, the clause in the contract that only matters once a year. Practitioners widely report that heavy generic alignment can quietly erode domain precision — treat that as practitioner consensus rather than a single clean study, but take it seriously, because the mechanism is straightforward: you get what you reward, and generic raters don't reward specialist exactness.

Here is the small-model advantage hidden inside the failure mode: because small models are cheap to evaluate, you can actually *measure* the tax per task. Run the domain benchmark before and after alignment-style training. Big models make that measurement expensive enough that teams skip it.

Warning sign: the model gets friendlier and more polished while its answers on your hardest domain questions get vaguer. Politeness is not precision.

The imitation trap: style is not capability

The most seductive failure in fine-tuning is also the most common: training a small model to *sound like* a big model, then mistaking the sound for the substance.

Gudibande et al. (2023, preprint) ran the definitive test. They trained models from 1.5B to 13B parameters on 0.3 million to 150 million tokens of imitation data — outputs harvested from proprietary frontier models — and evaluated them two ways. To crowd raters, the imitators looked impressively competitive with the frontier models. On actual capability benchmarks, in domains the imitation data didn't cover, they had closed almost none of the gap. The imitators had learned the *style* of expertise: the confident tone, the structured formatting, the authoritative cadence. They had not learned the expertise. "The false promise of imitating proprietary LLMs" is the paper's title, and it earns it.

The Alpaca story is the same trap from the other side — judge-dependence. Stanford's Alpaca 7B, in its original 2023 human evaluation, scored **90 wins against 89** versus text-davinci-003: apparent parity with a much larger model (Taori et al., Stanford CRFM blog, 2023). But the automatic AlpacaEval leaderboard later scored alpaca-7b at a **26.5% win rate** against text-davinci-003's 50.0. Same model, same comparison, wildly different verdicts — because the judge changed. Human raters, shown short exchanges, rewarded confident style; the automated judge measured something closer to substance.

This cuts both ways, and honesty requires saying so: automated judges (usually GPT-4) systematically favor models from their own family and training lineage, which flatters distillation results on leaderboards like AlpacaEval. Whenever this book cites a judge-scored win, pair it mentally with a ground-truth benchmark — GSM8K, HumanEval, MMLU — where answers are right or wrong regardless of who is judging.

Warning sign: the demo is dazzling but the hard benchmarks are missing, thin, or judge-scored. Demand ground truth. Style is cheap; capability is measured.

RAG beats fine-tuning for knowledge

Here is a clean, published boundary on what fine-tuning is *for*. Ovadia et al. (2023, preprint) asked a simple question: if you want a model to know facts — your product catalog, your policy manual, last quarter's numbers — should you fine-tune them in, or retrieve them at query time with RAG (retrieval-augmented generation)?

Across knowledge-intensive tasks, **retrieval consistently outperformed fine-tuning — for facts the model had seen in pretraining *and* for entirely new facts**. The authors' conclusion is blunt: large language models "struggle to learn new factual information through fine-tuning." Stuffing facts into weights is unreliable; handing the model the document at query time works.

The rule for the book, and for your architecture decisions: **fine-tune for behavior, retrieve for knowledge**. If the task is "know these facts," build RAG. If the task is "act this way on any input" — classify in this format, reason in this style, follow this procedure reliably across a long tail of inputs — fine-tune. The two compose beautifully: a fine-tuned small model *doing* retrieval is often the best system of all (Chapter 4's Self-RAG is exactly this pattern).

Warning sign: your fine-tuning dataset is mostly documents to memorize rather than behaviors to learn. That's a retrieval problem wearing a fine-tuning costume.

Contaminated benchmarks lie to everyone

Every win cited in Chapters 2–4 rests on benchmarks. Benchmarks rest on an assumption: the model hasn't seen the test. That assumption is failing at scale.

Deng et al. (2024, NAACL) tested it directly: they masked the answers in MMLU questions and asked models to reconstruct them. ChatGPT recovered them with **52% exact match**; GPT-4 with **57%**. You cannot guess a masked multiple-choice answer 57% of the time unless the test set was in your training data. The most-cited benchmark in the industry is, to a significant degree, a memorization test.

Microsoft's response was MMLU-CF, a contamination-free variant of MMLU (preprint, arXiv:2412.15194). On the clean version, **every tested model dropped 13–17 points** — GPT-4o fell from 88.0 to 73.4. The leaderboard numbers the whole industry quotes overstate *all* models, big and small alike.

And then there is outright gaming. SWE-Bench+ (preprint, arXiv:2410.06992) took the popular coding benchmark and removed every issue created before the models' knowledge cutoffs — the ones models could have memorized. GPT-4o-based agents' resolution rates **fell roughly 80%, to 3.8%**. Optimizing for a benchmark is not the same as having the capability; Goodhart's law eats leaderboards.

What does this mean for the small-beats-big thesis? Two things. First, small-model wins measured on contaminated public benchmarks deserve re-checking on fresh, private evaluations — the contamination flatters everyone, and you want to know your win is real. Second, contamination is actually an argument *for* the book's playbook: a fine-tuned model evaluated on *your* held-out data, from *your* task distribution, is immune to public-benchmark rot by construction. Your eval can't be in the model's training data if you built it yourself.

Warning sign: the benchmark is older than the model being tested, or the win exists only on public leaderboards with no private-eval confirmation. Fresh task, fresh eval, or it didn't happen.

The same model can win and lose: distribution dependence

The subtlest failure mode is not a failure at all — it is a win with boundaries nobody mapped. The same adaptation can beat the generalist in-distribution and lose out-of-distribution, and both results are true.

Distil-Whisper is the cleanest example (Gandhi et al., 2023, preprint). The distilled 756M-parameter student beats its 1.55B teacher on long-form speech recognition — word error rate **11.6 vs. 11.7** — while being 49% smaller and 5.8× faster. But on short-form audio, the student regresses: **10.1 vs. the teacher's 9.1**. The win is real and the loss is real; they live on different slices

of the data. Ship the student for podcasts, keep the teacher for voice commands — or know exactly which slice is yours.

TowerInstruct-7B (Unbabel, COLM 2024, peer-reviewed) shows the same pattern in translation. It crushes GPT-4 on translation-related named-entity recognition — F1 **71.68 vs. 59.88** — and edges the 70B LLaMA-2 on automatic post-editing. But on that post-editing task, GPT-4 still leads overall (**85.20 vs. 82.69**). Specialization won the slice it trained for; the generalist kept the rest.

And FinMA — the finance star of Chapter 2, beating GPT-4 at financial classification — collapses on numerical finance QA *in the same paper* (Xie et al., 2023, preprint). FinQA exact-match: FinMA-7B **0.06**, FinMA-30B **0.11**, versus GPT-4's **0.63**. ConvFinQA: 0.25 and 0.40 versus GPT-4's **0.76**. Instruction fine-tuning won inside the tuning distribution (classifying sentiment and headlines) and fell apart one step outside it (multi-step numerical reasoning over tables). The win and the failure come from the *same model in the same paper* — which is why this chapter exists.

Warning sign: the evaluation slice differs from the training slice — longer inputs, harder reasoning, a shifted domain. Always ask *which slice* the small model wins, and test the slices you actually serve.

Reading the warning signs

Step back and the six failure modes form a checklist — the conditions under which the thesis of this book holds, stated negatively:

Failure mode	Warning sign
Catastrophic forgetting	General skills (or safety) degrade while the target task improves
Alignment tax	Answers get friendlier and vaguer on your hardest domain questions
Imitation trap	Dazzling demo, but only style-level or judge-scored evidence
Knowledge injection	The "fine-tuning" data is documents to memorize, not behaviors to learn

Failure mode	Warning sign
Contaminated benchmarks	The benchmark predates the model; no private-eval confirmation
Distribution dependence	The served slice differs from the trained slice

None of these is an argument against fine-tuning small models. Every one of them is an argument against fine-tuning them *carelessly* — and for the evaluation discipline that Chapter 10 will turn into a playbook. The teams that get burned by small models are not the teams that fine-tune; they are the teams that fine-tune on the wrong data, measure on the wrong benchmark, and deploy on the wrong slice. The thesis survives this chapter intact, with its conditions now explicit. That is what makes it trustworthy — and what makes it useful.

The pre-flight checklist

If this chapter had to compress into five questions asked before any fine-tuning run, they would be these:

1. **What is my data mix, and what will it overwrite?** If your training data is all one flavor — generic chat, short answers, a single domain — expect forgetting everywhere else. Mix in representative samples of what you need to keep, keep the learning rate low, and consider LoRA so the base model stays frozen.
2. **Am I training behavior or stuffing knowledge?** If the dataset is documents to memorize, stop and build retrieval instead. Fine-tuning teaches the model *how to act*; it is a poor filing cabinet.
3. **Is my evaluation ground truth or vibes?** One hard benchmark with right and wrong answers is worth ten judge-scored leaderboards and a hundred impressed colleagues. Build your own held-out eval from your own task distribution — it is immune to public contamination by construction.
4. **Which slice am I actually serving?** Write down the input distribution your users will produce — lengths, difficulty, edge cases — and verify it overlaps the training slice. The Distil-Whisper and FinMA stories are what happens when nobody checks.

5. **What did I break that I wasn't measuring?** Re-run general capability and safety probes after every fine-tuning run. The cheapest evaluations in the industry are the ones on your own small model — there is no excuse for skipping them.

Answer all five honestly and you have converted every failure mode in this chapter from a risk into a routine check. That is the difference between teams for whom small models win and teams for whom they mysteriously don't: not better models, better discipline.

Key takeaways

- **Fine-tuning can unlearn:** Llama-3-70B-Instruct dropped from 39.00 to ~36–38.8 on general benchmarks after narrow fine-tuning — driven by data-mix mismatch, fixable with matched-distribution replay. Even benign fine-tuning can erode safety alignment (Qi et al., 2023).
- **Style is not capability:** imitation-trained models fool crowd raters while closing almost none of the capability gap (Gudibande et al., 2023); Alpaca's "parity" with text-davinci-003 (90–89 human eval) collapsed to 26.5% vs. 50.0% under automated judging. Demand ground-truth benchmarks.
- **Fine-tune for behavior, retrieve for knowledge:** RAG consistently beat fine-tuning at knowledge injection for both old and new facts (Ovadia et al., 2023). Memorizing documents is a retrieval problem.
- **Benchmarks rot:** MMLU shows strong contamination evidence (GPT-4: 57% masked-answer recovery); clean variants drop every model 13–17 points; SWE-Bench+ cut agent scores ~80% after decontamination. Evaluate on fresh, private, task-aligned data.
- **Wins have boundaries:** Distil-Whisper beats its teacher on long-form but loses on short-form; FinMA beats GPT-4 at classification but collapses on numerical QA (0.06–0.11 vs. 0.63). Map the slice you serve, and test it.

Sources

- Catastrophic forgetting measurements: arXiv:2506.09428 (preprint, 2025) — Llama-3-70B-Instruct fine-tuning, 39.00 → 35.91–38.81, reconstructed-distribution recovery to 39.21. <https://arxiv.org/abs/2506.09428>

- X. Qi et al., "Fine-tuning Aligned Language Models Compromises Safety, Even When Users Do Not Intend To!" (2023) — benign fine-tuning degrades safety alignment.
- L. Ouyang et al., "Training language models to follow instructions with human feedback" (preprint, 2022) — the alignment tax and PPO-ptx mitigation. <https://arxiv.org/abs/2203.02155>
- A. Gudibande et al., "The False Promise of Imitating Proprietary LLMs" (preprint, 2023). <https://ar5iv.labs.arxiv.org/html/2305.15717>
- R. Taori et al., Alpaca human evaluation (Stanford CRFM blog, 2023); AlpacaEval leaderboard. https://github.com/tatsu-lab/alpaca_eval
- O. Ovadia et al., "Fine-Tuning or Retrieval? Comparing Knowledge Injection in LLMs" (preprint, 2023). <https://arxiv.org/abs/2312.05934>
- C. Deng et al., MMLU contamination evidence (NAACL, 2024) — 52%/57% masked-answer exact match.
- MMLU-CF: contamination-free MMLU variant (preprint, 2024) — 13–17 point drops across models, GPT-4o 88.0 → 73.4. <https://arxiv.org/abs/2412.15194>
- SWE-Bench+: decontaminated coding benchmark (preprint, 2024) — agent resolution rates fell ~90% to 3.8%. <https://arxiv.org/abs/2410.06992>
- S. Gandhi et al., "Distil-Whisper: Robust Knowledge Distillation via Large-Scale Pseudo Labelling" (preprint, 2023). <https://arxiv.org/pdf/2311.00430.pdf>
- TOWER: Unbabel (COLM 2024, peer-reviewed) — TowerInstruct-7B translation results. <https://openreview.net/pdf?id=EHPns3hVkj>
- Q. Xie et al., "PIXIU: A Large Language Model, Instruction Data and Evaluation Benchmark for Finance" (preprint, 2023) — FinMA wins and FinQA/ConvFinQA failures. <https://arxiv.org/pdf/2306.05443>

9. Where Scale Still Wins

This book has spent eight chapters arguing that small, specialized models beat big, general ones. This chapter argues the other side — and argues it well.

That is not a contradiction. The thesis of this book was never "small always wins." It was "small wins, conditionally." If you finish this book unable to argue the case for big models, you don't understand the thesis; you just memorized a slogan. A reader who loves frontier models should close this chapter feeling fairly represented. Let me earn that.

There is also a practical reason for this chapter. The boundary between "fine-tune a specialist" and "use the generalist" moves every year. If you only learn the specialist's case, you'll misjudge where that boundary sits. This chapter draws the map from the other side, so you can navigate the whole territory.

The scaling laws, in plain language

In 2020, researchers at OpenAI published a finding that shaped the entire industry (Kaplan et al., 2020). They showed that a language model's test loss — roughly, how surprised it is by text it hasn't seen, a clean measure of "how well it learned the patterns of language" — falls as a smooth, predictable curve as you increase three things: the number of parameters, the amount of training data, and the amount of compute. Double the compute, get a predictable increment of capability. No visible ceiling, no sudden plateau. Their practical recipe was blunt: given a fixed budget, spend it on the biggest model you can afford, even if that means training it on relatively little data. GPT-3 followed that recipe: 175 billion parameters trained on only about 300 billion tokens — fewer than two tokens per parameter.

Think of it like this: the 2020 recipe said that when buying intelligence, size of the container mattered more than how full you filled it. The industry believed it, and the great model-size race began.

Two years later, DeepMind showed the recipe was wrong — or rather, incomplete. The Chinchilla paper (Hoffmann et al., 2022) demonstrated that the compute-optimal split scales parameters and data *equally*: roughly 20 tokens of training data per parameter. Fill the container *and* size it right.

Chinchilla itself was a 70-billion-parameter model trained on 1.4 trillion tokens, using the same total training compute as Gopher, a 280-billion-parameter model built the old way. The smaller, better-fed model won decisively: Chinchilla beat Gopher, GPT-3, and even the 530-billion-parameter MT-NLG on the MMLU benchmark, averaging 67.5% — about 7 points ahead of Gopher.

Read that again, because it matters. The most famous scaling-law correction in the field's history was a *small model beating a model four times its size*. The lesson of Chinchilla is not "bigger is better." It is "spend compute wisely."

And here is the twist that this book gets to claim as its own. Once the industry accepted the Chinchilla ratio, it kept going past it: labs started *overtraining* small models far beyond the "optimal" point. Why would anyone do something deliberately suboptimal? Because the Chinchilla math optimizes for *training* compute — the one-time cost of building the model — but what dominates real deployments is *inference* compute: the cost of serving the model millions or billions of times. A small model trained on far more data than "optimal" is slightly wasteful to train and dramatically cheaper to serve, forever. LLaMA-3-8B was trained on 15 trillion tokens — roughly 1,875 tokens per parameter, nearly a hundred times the Chinchilla ratio. Every extra token was an investment in a cheaper future.

So the scaling laws, correctly understood, are an argument for small models that were trained for a very long time. Scale and specialization are not enemies. Specialization is how you spend scale efficiently.

The generalist's portfolio

Before the counterarguments, consider what the generalist is *for*. A frontier model is a diversified portfolio of capabilities. Like a financial portfolio, diversification looks wasteful when you know exactly which asset will win — why hold bonds when tech stocks are soaring? — and looks brilliant the moment the future surprises you. The generalist's breadth is insurance against unknown unknowns: the question you didn't anticipate, the task that doesn't fit any category, the novel situation where there is no training distribution to be inside of.

This is why the "demo effect" from Chapter 1 is so seductive and so misleading at the same time. In a demo, you throw something unexpected at the big model and it handles it gracefully. That is genuinely impressive, and it is genuinely the generalist's home turf. The mistake is concluding from the demo that the generalist is therefore the right tool for the ten-thousandth repetition of a narrow task. Insurance is for surprises. You don't buy car insurance to drive to the grocery store more efficiently.

The bitter lesson, stated fairly

In 2019, the reinforcement-learning researcher Richard Sutton wrote a short essay called "The Bitter Lesson." Its argument: across decades of AI research, general methods that scale with compute always beat methods built on human knowledge and clever engineering — and researchers keep re-learning this the hard way. Chess, Go, speech recognition, vision: every time, the hand-crafted approach looked brilliant for a while, then the scaled-up general method blew past it. The "bitterness" is that all our clever domain expertise keeps turning out to be leverage against ourselves: it helps in the short term and becomes an obstacle to the scalable solution.

The steelman against this book goes like this: fine-tuning and distillation are friction. They are clever human effort — curated datasets, careful recipes, domain-specific tricks, exactly the kind of hand-engineering Sutton warned about. The bitter lesson says the next scale-up washes all of that away. Today's distilled 7-billion-parameter model is tomorrow's rounding error on a bigger base. Why invest in specialization when the frontier will absorb it for free?

It is a serious argument, and it has history on its side. Here is the reply, in two parts.

First, notice what the bitter lesson is actually about: *pretraining methods* — how you build the base model. It says nothing against taking a scaled-up base model and adapting it cheaply. A distilled small model rides the exact same compute curve as the frontier; it just rides it at a hundredth of the serving cost. The bitter lesson tells you to bet on scale. It does not tell you to serve the biggest model for every narrow task. The history of engineering is full of general-purpose engines with specialized attachments — nobody runs a

container ship's engine to power a lawnmower, even though the ship's engine is more powerful.

Second, the bitter lesson has an ironic blind spot: *instruction tuning itself* is "friction." Base models, trained the pure scalable way, are rambling text completers — it took human-crafted instruction data and preference tuning to make them useful assistants. The entire industry accepted that some human knowledge, applied after pretraining, is not an obstacle to scale but its multiplier. Fine-tuning is that same idea, aimed at a narrower target. If you accept instruction tuning, you have already conceded the principle; the rest is a debate about scope.

Where the generalists clearly win

Honesty requires naming the territories where big, general models still rule. There are at least four.

First: reasoning about the unfamiliar. Specialization wins when the task sits inside the training distribution. Take the task outside it, and raw scale reasserts itself. A study published in Nature Communications on medical metacognition — the ability to reason reliably about *unfamiliar* medical problems, not just recall familiar ones — tested models on MetaMedQA and found Qwen2-72B at 64.3%, Qwen2-7B at 43.9%, and GPT-4o at 73.3%. The specialized 7-billion-parameter medical model did not beat the larger, current general models. When the questions demand reasoning far from anything the specialist was trained on, the generalist's breadth is not a luxury; it is the capability itself.

The same pattern shows up in the Orca results from Chapter 3, viewed from the other direction. Orca-13B reached 66.92% on reasoning benchmarks — remarkable next to GPT-3.5's 67.65%, a genuine David-and-Goliath story. But GPT-4 hit 79.03%. The student nearly caught the teacher's sibling; the teacher's big brother was still twelve points up the road. On tasks nobody fine-tuned for, the frontier's margin is still a chasm, not a gap.

Second: frontier coding and mathematics. This book celebrated small models that write code well, and those results are real. But Zephyr-7B's own authors concede their model lags proprietary models on coding and mathematics — and that gap is structural, not temporary. There is a difference between *coding tasks* and *software engineering*. Completing a function, fixing

a known bug pattern, translating between languages: narrow, repeatable, winnable by specialists. Open-ended engineering — working across many files, using tools, planning over long horizons, recovering from errors in an unfamiliar codebase — has a task distribution too wide, too varied, and too long-tailed to distill cheaply into a small model. The frontier coding agents that do this work remain a big-model game, and will for a while.

Third: agentic behavior on unseen tasks. Multi-step tool use, chaining unfamiliar APIs, recovering from novel errors mid-task — these "emergent" abilities appear abruptly at scale (Wei et al., 2022). A small model can be taught a narrow agentic loop: follow these five steps with these three tools, and it will do so reliably and cheaply. But drop it into a genuinely new environment — a new API, a new failure mode, a task that combines tools in a way it has never seen — and the generalist's zero-shot adaptability is still unmatched. This is the generalist's portfolio paying out: the premium was for exactly this moment.

(A caveat the field itself debates: some researchers argue that emergence is partly an artifact of how we measure — sharp jumps on coarse metrics can look like magic when the underlying progress was smooth. The debate is real and worth knowing about. But whatever emergence *is*, the practical gap between small and large models on unfamiliar agentic tasks is, for now, too large to ignore.)

Fourth: the moments when scale beats specialization head-on. In March 2023, GPT-4 — a general model with no medical specialization — outperformed Med-PaLM, a medical-specialized model, on medical challenge problems, exceeding the USMLE passing score by more than 20 points (Nori et al., 2023). At that moment, general scale beat domain specialization outright, on the specialist's home turf. The story didn't end there: methods improved, and Med-PaLM 2 later reached a reported 86.5% on MedQA (Singhal et al., 2023/2024) — the lead changed hands again. Which is itself the lesson: specialization is not destiny, and neither is scale. The frontier moves, and any claim of permanent victory in either direction is ahistorical.

The cost of getting the boundary wrong

Misjudging this boundary is expensive in both directions, which is why the map matters more than the ranking.

Bet on the specialist where the generalist was needed, and you get a confident, fast, cheap model that fails exactly where it matters most: the unfamiliar case it was never trained for. The failure mode is insidious because it looks like success — high scores on in-distribution tests, smooth demos — right up until the real world serves an input from outside the training distribution. Chapter 8's distribution-dependence lesson (Distil-Whisper winning long-form while losing short-form; TowerInstruct winning NER while losing post-editing to GPT-4) is the warning label: every specialist's victory is local, and deploying it beyond its locality is how you get a quiet catastrophe.

Bet on the generalist where the specialist would do, and the cost is slower and more mundane: you pay per token forever for capability you use a fraction of, you wait on latency your users feel, and you ship your data to someone else's servers for every single request. Nobody goes bankrupt from this on day one. They just never become cost-competitive with the rival who fine-tuned.

The playbook in the next chapter exists to keep you off both cliffs. But the principle fits in one line: match the tool to the territory, and re-check the map often, because the territory keeps changing.

The synthesis

So where does that leave the thesis? Stronger, not weaker.

The claim was never that small models are better at everything. It was that on narrow, well-defined, in-distribution tasks, with good data and task-aligned evaluation, a specialized small model beats a larger generalist — and does so while being radically cheaper and faster to run. Nothing in this chapter refutes that. What this chapter establishes is the boundary: step outside the distribution — unfamiliar reasoning, open-ended coding, novel agentic environments — and the generalist earns its keep.

Think of it as a map, not a ranking. The frontier models explore new territory; the specialized models settle it. Every capability the big models demonstrate

today becomes a distillation target tomorrow — explanation traces, reasoning strategies, and behavioral patterns all flow downhill from large to small, as Chapters 3 and 5 showed. The Med-PaLM story is the pattern in miniature: the generalist proved what was possible, and specialization followed to make it efficient.

Fine-tuning does not abolish the frontier. It shifts it: the big models push outward into the unknown, and the small models consolidate what was learned, making it cheap, fast, and deployable. That is the division of labor the next chapter's playbook will assume. Use the giants to discover. Use the specialists to deliver.

Key takeaways

- The scaling laws say "spend compute wisely," not "bigger is better" — and at serving scale, wisdom is small: Chinchilla (70B) beat Gopher (280B), and the industry now overtrains small models (LLaMA-3-8B on 15T tokens) because inference cost dominates.
- The bitter lesson argues for scaling *pretraining*, not against specialization: a distilled small model rides the same compute curve at a fraction of the serving cost — and instruction tuning itself is proof that post-training "friction" multiplies scale.
- General models clearly win on unfamiliar reasoning (GPT-4o 73.3% vs Qwen2-7B 43.9% on MetaMedQA), frontier coding and math, and agentic behavior in unseen environments — the generalist's breadth is insurance against the unknown.
- The lead changes hands over time — GPT-4 beat the specialist Med-PaLM in 2023, Med-PaLM 2 later hit 86.5% on MedQA — so neither scale nor specialization is destiny.
- The synthesis: big models discover, small models deliver. Fine-tuning shifts the frontier; it doesn't abolish it.

Sources

- Kaplan et al., "Scaling Laws for Neural Language Models," 2020 (preprint). <https://arxiv.org/abs/2001.08361>

- Hoffmann et al., "Training Compute-Optimal Large Language Models" (Chinchilla), 2022 (preprint). <https://arxiv.org/abs/2203.15556>
- Sutton, "The Bitter Lesson," 2019. <http://www.incompleteideas.net/IncIdeas/TheBitterLesson.html>
- "Large Language Models lack essential metacognition for reliable medical reasoning," Nature Communications. (MetaMedQA results.) <https://www.nature.com/articles/s41467-024-55628-6>
- Mukherjee et al., "Orca 2: Teaching Small Language Models How to Reason," Microsoft Research, 2023 (preprint). <https://arxiv.org/pdf/2311.11045>
- Tunstall et al., "Zephyr: Direct Distillation of LM Alignment," 2023 (preprint). <https://arxiv.org/abs/2310.16944>
- Nori et al., "Capabilities of GPT-4 on Medical Challenge Problems," Microsoft,
- https://www.microsoft.com/en-us/research/uploads/prod/2023/03/GPT-4_medical_benchmarks.pdf
- Singhal et al., "Towards Expert-Level Medical Question Answering with Large Language Models" (Med-PaLM 2), 2023 (preprint). <https://arxiv.org/abs/2305.09617>
- Wei et al., "Emergent Abilities of Large Language Models," 2022 (preprint). <https://arxiv.org/abs/2206.07682>

10. The Playbook

You now have the evidence, the mechanisms, the economics, and the limits. This chapter turns all of it into decisions. If you are choosing how to build an AI feature — as an engineer, a founder, or the person signing the budget — here is the playbook.

There are four options on the table. Everything in AI deployment is some version of one of these:

1. **Prompt** a general model — write clever instructions, add examples, ship.
2. **Retrieve** — connect the model to your data with RAG (retrieval-augmented generation): fetch relevant documents, stuff them into the prompt.
3. **Fine-tune** — train a smaller model on your task until it owns it.
4. **Buy** — pay a vendor for a finished solution and skip building.

The flowchart below is the whole chapter in one page. The rest is how to run it honestly.

The decision flowchart

Question 1: Is the task narrow, stable, and repeated?

If the same kind of input needs the same kind of output thousands of times a day — classify this document, extract these fields, answer questions from this manual, triage these messages — you are in fine-tuning territory. Narrow and repeated is the exact shape specialization wins: every result in Chapters 2 through 4 had this shape.

A practical test: could you write a rubric for "correct" that two careful humans would mostly agree on? If yes, the task is well-defined enough to train for. If correctness is a matter of taste, exploration, or genuine novelty each time — "help me think through this," "write something surprising" — stop here and prompt the generalist. Fine-tuning cannot bottle what isn't repeatable.

Stability matters as much as narrowness. A task that changes shape every month is a task you'd be re-tuning every month. Specialization is an investment; like any investment, it needs a shelf life to pay off.

Question 2: Is the gap about knowledge or about behavior?

This is the single most useful diagnostic in the book, and it comes from Chapter 8. If the model needs to *know* things — your product docs, this quarter's policies, facts that change — then RAG beats fine-tuning, consistently (Ovadia et al., 2023). Models are bad at learning new facts through training and good at reading facts you hand them at query time.

If the gap is about *behavior* — the tone it answers in, the format it follows, the judgment calls it makes on familiar inputs, the refusal patterns, the domain's way of reasoning — that is what fine-tuning installs. Behavior is a skill; knowledge is a lookup. Train the skill, look up the facts.

How to tell which gap you have: imagine handing the model a perfect cheat sheet containing every fact it could need. If the cheat sheet fixes the problem, your gap was knowledge — build RAG. If the model *still* gets it wrong with the facts in front of it — wrong format, wrong judgment, wrong tone — your gap is behavioral, and no retrieval system will fix it. That is a fine-tuning problem.

Note the combination the book keeps returning to: a fine-tuned small model with RAG on top. The Self-RAG result in Chapter 4 showed a 7-billion-parameter model with trained retrieval behavior beating ChatGPT on factual QA. Skill and lookup are complements, not competitors.

Question 3: Can you get the data — and is it good?

Fine-tuning is a data project wearing a modeling costume. Chapter 6's lesson was that a thousand excellent examples beat a million noisy ones: the LIMA result showed a 65-billion-parameter model fine-tuned on just 1,000 carefully curated examples producing responses rated equivalent-or-better than GPT-4 in 43% of comparisons (Zhou et al., 2023). You do not need big data. You need *right* data: real inputs from your task, correct outputs, consistent formatting, deduplicated, covering the long tail — including the awkward edge cases, which are where fine-tuned models earn their keep.

If you cannot assemble even a few hundred high-quality examples, do not fine-tune. Prompt, or use RAG, or buy. Garbage in, garbage out is not a folksy saying in this context; it is the entire failure mode of Chapter 8. And be honest about where your data comes from: the best fine-tuning datasets are exhaust from the workflow itself — resolved tickets, nurse triage decisions, reviewed

extractions. If producing the data requires a heroic manual labeling effort, price that effort into the decision; it is usually the largest cost in the project.

Question 4: What do latency, cost, and privacy demand?

If the feature must respond in under a second, run on-device, handle millions of requests, or keep data in-house for legal reasons — the big API may be disqualified before the benchmark even matters (Chapter 7). A fine-tuned small model you host yourself has zero marginal cost per token, sub-second latency, and never sends customer data to a third party. Sometimes the decision is not "which is smarter" but "which is allowed in the building."

This question also runs in reverse: if your volume is tiny, constraints are loose, and the task works fine with a prompt — don't build infrastructure for a problem you don't have. The playbook is not "always fine-tune." It is "fine-tune where the conditions hold."

Putting it together. Narrow + behavioral + good data + demanding constraints → fine-tune a small model. Knowledge-heavy → RAG, possibly on top of a fine-tuned model. Open-ended → prompt the generalist. No data, no time, no team → buy, and revisit the question when you have volume worth optimizing.

Four worked scenarios

The scenarios below are illustrative hypotheticals — generic composites to show how the flowchart runs, not claims about any real company.

Scenario A: a bank's document classifier. A regional bank processes 40,000 incoming documents a month — loan applications, statements, ID scans — and must route each to the right team. The task is narrow (fixed categories), stable (the categories rarely change), and repeated (tens of thousands of times). The gap is behavioral: the model must apply the bank's routing rules consistently, including on messy scans where judgment is needed. The cheat-sheet test: handing the model the rulebook wouldn't fix misreads of smudged documents — judgment is required. Data exists: years of correctly routed documents, labeled by the workflow itself. Privacy rules forbid sending customer financials to third-party APIs. Run the flowchart: narrow ✓, behavioral ✓, data ✓, privacy-constrained ✓ → fine-tune a small classifier model, self-hosted. This is the textbook fine-tuning win — FinMA-style

specialization on financial text, inside the bank's own walls. Revisit only if the document types or routing rules change materially.

Scenario B: a startup's support bot. A SaaS startup gets 3,000 support tickets a month and wants a bot to resolve the common ones. The task is semi-narrow: most tickets cluster around twenty topics, but there is a long tail of weird ones. The gap is mixed: the bot needs *knowledge* (current docs, pricing, status pages — all changing monthly) and *behavior* (the company's tone, escalation judgment). The cheat-sheet test: with perfect docs in the prompt, the model answers correctly but in the wrong voice and escalates poorly — so both gaps are real. Data is thin: a few hundred resolved tickets, inconsistently written. The flowchart says: do not fine-tune yet. Start with RAG over the docs plus a prompted general model — knowledge first, since the docs change monthly and the data isn't ready. Once ticket volume grows and a thousand clean resolutions accumulate, fine-tune a small model on the top topics for speed and cost, keep RAG for the knowledge, and escalate the long tail to humans. Staged, not dogmatic: the right answer today is not the right answer forever.

Scenario C: a hospital's triage notes. An emergency department wants incoming patient messages sorted into urgency tiers before a nurse reads them. Narrow, repeated, high-stakes. The gap is behavioral: consistent triage judgment expressed in the department's own protocol language — the cheat sheet alone wouldn't produce consistent decisions. Data: years of nurse-triaged messages — high quality, already labeled by the workflow itself, with the long tail of unusual presentations included. Privacy: patient health data cannot leave the hospital's systems, which disqualifies third-party APIs outright. Flowchart: narrow ✓, behavioral ✓, data ✓, privacy-constrained ✓ → fine-tune a small model, on-premises. Note what this scenario also demands from Chapter 8: rigorous evaluation on fresh, uncontaminated cases, because the cost of being wrong is not a bad demo — and a clear escalation path, because triage is decision *support*, not decision replacement. The playbook includes knowing what the model must never decide alone.

Scenario D: a consultancy's research assistant. A strategy consultancy wants an internal tool to help analysts explore unfamiliar industries. The task is open-ended by design: every engagement is a new domain, and "correct" is a matter of insight, not a rubric. No stable distribution, no repeatable behavior to install, and the "data" would be the analysts' past decks — brilliant, but each one

unique. The flowchart stops at Question 1: not narrow $\times \rightarrow$ prompt the generalist, possibly with RAG over the firm's knowledge base. Fine-tuning here would be expensive overfitting to yesterday's engagements — the model would learn the *style* of old decks while the new engagement needs fresh thinking. This is the scenario where Chapter 9's generalist wins outright, and choosing it is not a failure of ambition. It is the playbook working.

The industry recipe: NVIDIA's migration algorithm

Chapter 7 cited NVIDIA Research's position paper arguing that small models are the future of agentic AI. The same paper sketches a six-step recipe for migrating an existing system off the big generalist onto fine-tuned small models — and it's worth laying out here, because it turns this chapter's flowchart into an operational plan:

1. **Log everything.** Instrument the system to record real production calls — inputs, outputs, tool calls, latencies. Your future training data is today's traffic exhaust, and if you're already running on the generalist, step 1 is free.
2. **Curate 10,000–100,000 examples.** Clean the logs: strip personal data, deduplicate, filter for quality. The paper's rule of thumb is that tens of thousands of examples suffice to fine-tune a small model — Chapter 6 says the same thing from a different angle.
3. **Cluster into tasks.** Group the logged calls by pattern — intent recognition here, field extraction there, summarization of one document type, code generation against one toolset. Each recurring cluster is a candidate for its own specialist.
4. **Pick the small model per task.** Match each cluster to a candidate: check its instruction-following and reasoning on the relevant benchmarks, its license, and its memory footprint on your hardware.
5. **Fine-tune each specialist.** Train each task's dataset into its model with LoRA or QLoRA — or distill the big model's behavior on that task into the small one. The paper explicitly blesses both of this book's core mechanisms.
6. **Keep the loop running.** Retrain periodically on fresh traffic. Tasks drift; the specialists must drift with them.

Notice what the recipe assumes — and what it doesn't. It doesn't ask you to fine-tune everything at once. You migrate task by task, starting with the highest-volume clusters, where the savings are largest and the evaluation is easiest. That's the playbook's Question 1 made concrete — find the calls that are narrow, stable, and repeated, which are usually the majority of your traffic — applied in Question 4's order: convert the expensive, latency-hungry clusters first. The generalist stays for the long tail and the genuinely open-ended work: exactly the heterogeneous architecture Chapter 9 describes as the sane end state.

The data checklist

Before any fine-tuning project, you should be able to check every box:

- **Real inputs.** Examples drawn from the actual task, not invented. Synthetic data is fine as a supplement — the phi and Orca results prove that curated synthetic data can be transformative — but the core should be the real distribution, because that's what you'll be graded on.
- **Correct outputs.** Every target answer verified. A thousand correct examples beat a hundred thousand sloppy ones — the LIMA lesson (Zhou et al., 2023). Budget for verification; it is the highest-leverage spending in the project.
- **Consistent format.** Same structure, same style, same conventions throughout. Inconsistency teaches the model that anything goes, and "anything goes" is the opposite of specialization.
- **Long-tail coverage.** The common cases *and* the awkward ones: ambiguous inputs, adversarial phrasings, near-misses between categories. Specialists earn their keep on the tail; a dataset of only easy cases produces a model that is confident exactly where it should be careful.
- **A held-out test set, locked away.** A few hundred examples the model never trains on, drawn from the same distribution, reserved for honest evaluation. If the test set leaks into training — even indirectly, through data that resembles it — your numbers are fiction. Lock it before training starts.

Evaluate like a skeptic

Chapter 8 showed how evaluations lie. The playbook's evaluation rules:

1. **Task-aligned evals.** Measure what you will actually ship: the real task, the real inputs, the metric that matches the business goal — not a generic benchmark that happens to be convenient. A 2-point gain on a public leaderboard means nothing if your users' metric doesn't move.
2. **Contamination hygiene.** If your test examples could have appeared in any model's training data, the scores are suspect. The MMLU-CF result — every model dropping 13–17 points on a contamination-free variant (GPT-4o: 88.0 $\rightarrow$ 73.4) — is the reason, alongside direct evidence that models can reproduce masked benchmark answers from memory (Deng et al., 2024). Prefer fresh, private test sets; treat public benchmark wins as directional, not decisive.
3. **Judge-dependence awareness.** If you use a model to judge models, remember the Alpaca story: 90 wins to 89 against the baseline by human raters, but only a 26.5% win rate by the automatic judge (Taori et al., 2023). Judges favor their own family and their own style. Pair judge scores with ground-truth benchmarks wherever possible, and distrust any evaluation where the judge and the student share a lineage.
4. **Measure the tax.** Fine-tuning can erode general abilities and safety guardrails (Chapter 8). Re-test the base model's broad capabilities after tuning, and re-test safety specifically — including with adversarial prompts. The specialist should gain the skill without losing its mind or its manners.
5. **A/B against the incumbent.** The final exam is not a benchmark — it is the current system. Run the fine-tuned model against whatever you use today (prompted generalist, rules engine, human team), on live traffic or a faithful replay, and measure the business metric: resolution rate, handling time, error rate, cost per task. Ship the winner, not the story.

Cost modeling without fantasy

Chapter 7 gave September-2026 prices — date any numbers you use, because they move fast. The playbook's costing has three parts:

- **One-time costs:** data curation (usually the biggest line item — human time), compute for the fine-tuning run (a QLoRA run on a 7B model fits a single 24GB consumer GPU; a 65B run fits under 48GB), and evaluation labor. For most teams, the GPUs are the cheapest part and the data work the most expensive — plan accordingly.
- **Recurring costs:** inference. API pricing is per-token and never goes to zero; a self-hosted small model costs hardware or GPU rental and then *nothing per token*. At sustained volume, the fine-tuned model usually wins within months — Chapter 7's worked example showed ~\$96/month versus ~\$3,200/month for 100,000 monthly conversations on identical traffic, a 33× gap. Latency compounds the win: in agentic workloads where the model is called dozens of times per task, a 10× latency gap becomes a 10× cost gap.
- **The break-even formula, not a number:** $(\text{monthly API cost} - \text{monthly self-host cost}) \times \text{months} > \text{one-time fine-tuning cost}$. Plug in today's prices and your volume. If the inequality holds within your planning horizon, fine-tune. If your volume is tiny, the API is cheaper — and that is fine. Print the formula in your proposal, not a price; prices shift too fast for single numbers to survive contact with next quarter.

Also cost the *risks*, which both options carry. A fine-tuned model needs re-tuning when the task drifts — budget for a data-and-retrain loop, not a one-off project. An API model's price, behavior, and availability can change under you with a vendor's changelog. Neither option is free of maintenance. The honest comparison prices both.

When NOT to fine-tune

Tape this to the wall:

- The task is open-ended, creative, or different every time.
- The gap is missing *knowledge*, not missing *behavior* — use RAG.

- You cannot assemble a few hundred high-quality examples.
- The task distribution shifts faster than you can re-tune.
- You need broad, unfamiliar reasoning or open-ended agentic behavior (Chapter 9).
- Your volume is so low that API costs are rounding error — don't build infrastructure for a problem you don't have.
- You are fine-tuning to *imitate* a big model's style rather than to *master* a task — style is not capability (Chapter 8).
- Legal or contractual terms forbid training on the data you have — check before you start, not after.

Revisit on a schedule

The boundary moves. Today's fine-tuning win can become tomorrow's commodity (when the next base model handles your task zero-shot), and today's "too open-ended to tune" can become tunable (when you accumulate enough data to define it). Put a review on the calendar: every six to twelve months, re-run the flowchart for each AI feature you ship. The playbook is not a one-time decision — it is a standing question, and the teams that keep asking it are the ones that keep winning on cost and quality while everyone else argues about model sizes.

The thesis, as strategy

Here is the whole book in one paragraph, rewritten as something you can act on.

Default to the generalist for exploration and the unfamiliar. The moment a task becomes narrow, stable, and repeated — the moment you can describe exactly what "correct" looks like and collect examples of it — stop renting intelligence by the token and start owning it: fine-tune a small model on your data, evaluate it like a skeptic, host it where your constraints demand, and pocket the difference in cost, latency, and control. Revisit the decision whenever the task changes shape, because the boundary moves. Small models win —

conditionally, which is to say, *strategically*: not everywhere, but exactly where it counts.

Key takeaways

- Four options, one flowchart: prompt the generalist, retrieve knowledge with RAG, fine-tune behavior into a small model, or buy — chosen by asking whether the task is narrow and stable, whether the gap is knowledge or behavior, whether good data exists, and what latency, cost, and privacy demand.
- Fine-tune for behavior, retrieve for knowledge: RAG consistently beats fine-tuning at injecting facts (Ovadia et al., 2023); fine-tuning wins at installing reliable behavior. The cheat-sheet test tells you which gap you have.
- Data quality beats data quantity: ~1,000 excellent examples can rival models trained with far more (LIMA; Zhou et al., 2023) — but you need real, consistent, long-tail-covering examples plus a locked-away test set.
- Evaluate like a skeptic: task-aligned metrics, contamination-free test sets, judge-bias awareness, post-tuning safety re-tests, and A/B against the incumbent system on the business metric.
- Cost it as a break-even formula with dated prices, and know when NOT to fine-tune: open-ended tasks, knowledge gaps, thin data, fast drift, tiny volume, or mere style imitation. Revisit the decision every 6–12 months.

Sources

- Ovadia et al., "Fine-Tuning or Retrieval? Comparing Knowledge Injection in LLMs," 2023 (preprint). <https://arxiv.org/abs/2312.05934>
- Zhou et al., "LIMA: Less Is More for Alignment," 2023 (preprint). <https://arxiv.org/abs/2305.11206>
- Deng et al., on MMLU contamination evidence, 2024. (See Chapter 8 sources.)
- Taori et al., "Alpaca: A Strong, Replicable Instruction-Following Model," Stanford CRFM, 2023. <https://crfm.stanford.edu/2023/03/13/alpaca.html>

- Asai et al., "Self-RAG: Learning to Retrieve, Generate, and Critique through Self-Reflection," 2023 (preprint). <https://arxiv.org/html/2310.11511v1>
- Dettmers et al., "QLoRA: Efficient Finetuning of Quantized LLMs," 2023 (preprint). <https://arxiv.org/abs/2305.14314>
- Belcak et al., NVIDIA Research / Georgia Tech, "Small Language Models are the Future of Agentic AI," 2025 (position paper, preprint). <https://arxiv.org/abs/2506.02153> — source of the six-step LLM → SLM migration recipe.

Appendix A: Glossary

Plain-language definitions of every technical term used in this book, in alphabetical order.

Agentic. Describes a model that can pursue a multi-step goal on its own: calling tools, reading results, adjusting its plan, and continuing without a human in the loop. An agentic coding assistant doesn't just suggest code — it edits files, runs tests, and fixes failures.

Alignment. Shaping a model's behavior to match what humans want: helpful, honest, and safe. Alignment is usually done after pretraining, with techniques like instruction tuning, RLHF, or DPO.

Alignment tax. The performance cost of alignment: optimizing a model for human preferences can slightly degrade its scores on standard benchmarks or its precision on narrow factual tasks. Named in the InstructGPT work (Ouyang et al., 2022).

API. A service interface that lets your software call a model running on someone else's servers. You send text, you get text back, and you pay per token. Convenient, but your data leaves your building and the price never goes to zero.

Base model. A model fresh out of pretraining: it has absorbed vast knowledge but hasn't been taught to follow instructions or behave like an assistant. Base models complete text; instruction-tuned models answer you.

Benchmark. A standardized test for models: a fixed set of questions or tasks with right answers, so different models can be compared on the same yardstick. MMLU, HumanEval, and GSM8K are benchmarks used throughout this book.

Benchmark gaming. Optimizing for a benchmark's score instead of the real capability it was meant to measure — after which the score stops meaning anything. Also known as Goodhart's law: when a measure becomes a target, it ceases to be a good measure.

Bitter lesson. Richard Sutton's 2019 observation that, across decades of AI history, general methods that scale with compute have always beaten hand-

engineered, knowledge-heavy approaches. The book's reply: the bitter lesson is about how to *build* base models, not about whether to specialize them.

Catastrophic forgetting. When fine-tuning on new data degrades skills the model previously had — the new training overwrites old knowledge. Mitigated with low learning rates, few epochs, mixing in old data, or PEFT methods that freeze most of the model.

Chain-of-thought. A prompting or training technique where the model writes out its reasoning step by step before giving the final answer. Distillation often works by training small models on a big model's reasoning traces, not just its answers.

Compute-optimal. The most efficient split of a fixed training budget between model size and training data. The Chinchilla finding: about 20 tokens of data per parameter.

Contamination. When a benchmark's test questions leak into a model's training data, inflating its score. The model isn't solving the test — it's remembering it. A major reason to prefer fresh, private evaluations.

Dense retriever. In RAG systems, the component that searches a document collection for passages relevant to a query, using vector similarity rather than keyword matching. "Dense" refers to the dense vector embeddings it compares.

Distillation. Training a small "student" model to mimic a large "teacher" model — its outputs, its reasoning traces, or its internal probabilities. The student ends up far smaller and faster while keeping much of the teacher's skill on the distilled tasks.

Distribution. The statistical pattern of the data: what kinds of inputs show up and how often. *In-distribution* means new inputs that resemble the training data; *out-of-distribution* means inputs from a different pattern, where specialized models tend to stumble.

DPO (Direct Preference Optimization). A simpler alternative to RLHF: instead of training a separate reward model and running reinforcement learning, DPO turns human preference pairs ("answer A was better than answer B") directly into a classification-style training loss (Rafailov et al., 2023).

Embedding. A list of numbers (a vector) that represents a piece of text, image, or other data, positioned so that similar meanings sit near each other. The foundation of retrieval, similarity search, and much of modern NLP.

Emergence. Abilities that appear suddenly in models past a certain scale — like multi-step tool use — rather than improving gradually. Some researchers argue emergence is partly a measurement artifact; the practical gap at small scale is real either way.

Few-shot. Giving a model a few examples of the task inside the prompt, so it can infer the pattern. The model learns the format from the examples without any training.

Fine-tuning. Continuing to train an already-trained model on additional, task-specific data, so it specializes. The central technique of this book.

Foundation model. A large model trained on broad data that can be adapted to many downstream tasks — the starting point that fine-tuning, distillation, and prompting all build on.

GPU. The graphics processing unit — the chip that does the massive parallel arithmetic of training and running neural networks. Model sizes are often described by what GPU memory they need.

Guardrails. Rules or filters that constrain a model's behavior: refusing certain requests, avoiding toxic outputs, staying on topic. Fine-tuning can accidentally erode them — one reason to re-test safety after tuning.

Hallucination. When a model generates fluent, confident statements that are false. A key reason RAG exists: grounding answers in retrieved documents reduces invention.

Human preference. Judgments by people about which of two model responses is better. The raw material of RLHF and DPO — and of benchmarks like LMArena.

In-context learning. A model's ability to pick up a task from examples or instructions given in the prompt, without any weight updates. Few-shot prompting is the classic case.

Inference. Running a trained model to get outputs — as opposed to training it. Inference cost and speed are what make small models economically dominant at scale.

Instruction tuning. Fine-tuning a base model on tasks phrased as natural-language instructions, which teaches it to follow directions and dramatically improves zero-shot behavior (Wei et al., 2022).

Jailbreak. A prompt crafted to bypass a model's guardrails and get it to do something it was trained to refuse. Benign fine-tuning has been shown to make models more jailbreakable — guardrails are part of what can be forgotten.

Knowledge cutoff. The date up to which a model's training data extends. Anything that happened after the cutoff is unknown to the model unless supplied via RAG or a newer model.

Latency. How long you wait for a response. Two flavors matter: time to the *first* token (TTFT) and overall generation speed (throughput). Small models win both, and latency is often the difference between shippable and unshippable.

LLM-as-a-judge. Using a large model to score or compare other models' outputs. Cheap and scalable, but biased: judges tend to favor models from their own family, so judge scores should be paired with ground-truth benchmarks.

Long tail. The vast number of rare, unusual cases that individually occur seldom but collectively matter enormously. Specialists earn their keep on the tail — if your data covers only common cases, the model will fail exactly where humans need help most.

LoRA (Low-Rank Adaptation). A PEFT method that freezes the model's weights and trains only a tiny low-rank "adapter" update — about 8.4 million parameters (0.12%) when adapting a 7-billion-parameter model (Hu et al., 2021). The adapter can be merged back into the model afterward with zero inference overhead.

NPU. A neural processing unit — an AI accelerator chip in phones and laptops that runs small models locally, enabling on-device AI without the cloud.

On-device. Running a model locally on a phone, laptop, or other device rather than on a server. Private (data never leaves), offline-capable, and free per query — but limited to models small enough to fit.

Out-of-distribution. See *distribution*. Inputs that differ systematically from the training data — the zone where specialized models lose their edge and generalists reassert theirs.

Overfitting. When a model memorizes its training data instead of learning the underlying pattern, performing well on training examples and poorly on new ones. Narrow fine-tuning on thin data is especially prone to it.

Parameters. The adjustable numbers inside a neural network — its weights. "7B" means 7 billion parameters. More parameters generally mean more capacity, more memory, and slower, costlier inference.

PEFT (Parameter-Efficient Fine-Tuning). The family of methods — LoRA and QLoRA foremost — that adapt a model by training a tiny fraction of its parameters instead of all of them. What makes fine-tuning cheap enough for small teams.

Per-token pricing. How API providers charge: a price per million input tokens and per million output tokens. Output tokens cost several times more than input tokens at every provider.

PPO. Proximal Policy Optimization — the reinforcement-learning algorithm originally used in RLHF to optimize a model against a reward model. Powerful but fiddly; DPO was invented partly to avoid it.

Pretraining. The initial, massive training phase where a model learns from vast text (or other data) to predict what comes next. Nearly all of a model's knowledge is acquired here; fine-tuning mostly teaches it how to *use* that knowledge.

Prompt. The text you give a model to steer it: instructions, context, examples. Prompting is the zero-training way to get behavior out of a model.

Prompt engineering. The craft of designing prompts — wording, structure, examples — to get better outputs. Often the first thing to try before fine-tuning, and sometimes all you need.

QLoRA. LoRA combined with 4-bit quantization of the frozen base model, shrinking fine-tuning memory so dramatically that a 65-billion-parameter model can be fine-tuned on a single 48GB GPU (Dettmers et al., 2023).

Quantization. Storing a model's weights in fewer bits (e.g., 4-bit instead of 16-bit), shrinking memory use and speeding inference with minimal quality loss. What makes big models servable and small models phone-sized.

RAG (Retrieval-Augmented Generation). Connecting a model to external knowledge: at query time, relevant documents are retrieved and placed in the

prompt, so the model answers from your data instead of its memory. The right tool for knowledge problems; the wrong tool for behavior problems.

Reranker. In a RAG pipeline, a second-stage model that re-scores the retrieved passages and keeps only the best before generation. Retrieval gets candidates; the reranker picks winners.

Reward model. In RLHF, a model trained on human preference comparisons to predict which response a human would prefer. The language model is then optimized to score highly with the reward model.

RLHF (Reinforcement Learning from Human Feedback). The three-stage recipe — supervised fine-tuning, reward-model training on human preferences, then PPO optimization — that taught models to follow instructions and human values (Ouyang et al., 2022). The heaviest post-training method, now often replaced by DPO for smaller projects.

Scaling laws. Empirical rules describing how model performance improves as a smooth power-law function of parameters, data, and compute (Kaplan et al., 2020). Their Chinchilla correction (Hoffmann et al., 2022) showed data and parameters should scale equally.

Self-hosting. Running a model on your own servers or cloud GPUs instead of calling a vendor's API. Fixed infrastructure cost, zero per-token cost, full data control — the economic endgame of the small-model strategy.

Supervised fine-tuning (SFT). Fine-tuning on labeled input-output examples with a standard training loss. The workhorse stage of every modern post-training pipeline, including LIMA's thousand-example demonstration.

Synthetic data. Training examples generated by a model (often a large one) rather than collected from humans. Powers distillation and data-specialization strategies like phi — quality-filtered synthetic textbooks beat raw web scrapes.

System message. A privileged instruction placed before the conversation that sets the model's role, tone, and constraints. Fine-tuning can bake the system message's behavior into the model itself.

Temperature. A generation setting controlling randomness: low temperature makes outputs deterministic and conservative; high temperature makes them varied and creative. Also used in distillation, where a high "softmax temperature" reveals the teacher's informative near-miss probabilities.

Throughput. How many tokens a model generates per second once it's going. Small models reach hundreds or thousands of tokens per second; large models tens.

Token. The unit models read and write — roughly a word fragment. "Fine-tuning small models" is about six tokens. API pricing, context limits, and speed are all measured in tokens.

Tool use. A model calling external functions — search, calculators, code execution, APIs — as part of answering. The foundation of agentic behavior, and an area where broad general models still lead on unfamiliar tasks.

TPU. Tensor Processing Unit — Google's AI accelerator chip, the hardware behind much large-scale training.

TTFT (time to first token). How long before the model starts responding. Critical for interactive products: users feel TTFT directly, and small models cluster at the fast end.

Weights. See *parameters*. The learned numbers that define what a model does. "Freezing" weights means leaving them unchanged during fine-tuning.

Zero-shot. Asking a model to do a task with no examples, just instructions. Instruction tuning is what made zero-shot behavior reliable — and a fine-tuned specialist often needs no examples at all on its home task.