

Teaching a Small LLM Scholastic Voice

Fine-Tuning Qwen 2.5 on the Catechism, Summa, and Augustine via Local MLX

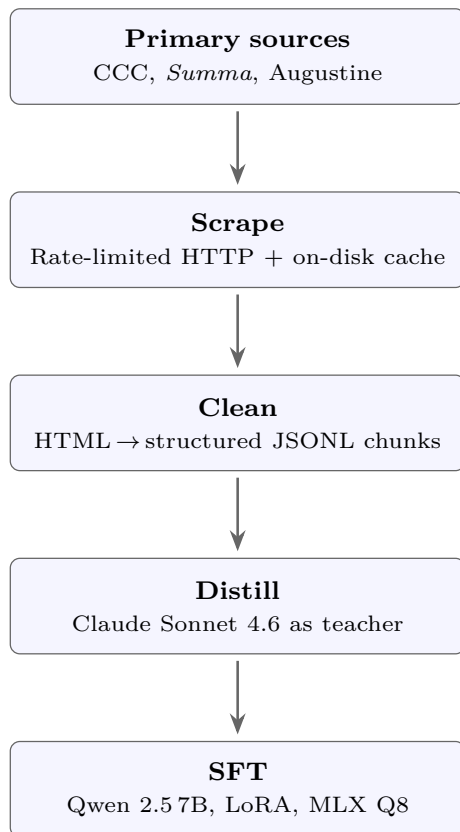
Pablo Leyva · R1 New Jersey Institute of Technology · github.com/pleyva2004/scholastic-llm · [paper](#)

TL;DR. A 7B open-weights model (Qwen 2.5) shifts from flat encyclopedic prose to scholastic / Summa-form prose in ~ 7 minutes on a 48 GB M4 Pro using LoRA + MLX Q8 and just 83 teacher-distilled (question, answer) pairs. A 4-dimension rubric of register, citation grounding, and argument structure rises from **19/120 to 68/120 (+258%)** on 10 held-out philosophical prompts.

Motivation

Instruction-tuned LLMs default to a flat, encyclopedic register. Specialized historical writing — the Summa’s *quaestio* form, Augustine’s autobiographical address, the Catechism’s numbered paragraph citations — has a distinct lexical and structural fingerprint that is absent from base-model output. Can a 7B open-weights model be moved into that register *cheaply*, on a laptop, using only public-domain or fair-use sources?

Pipeline

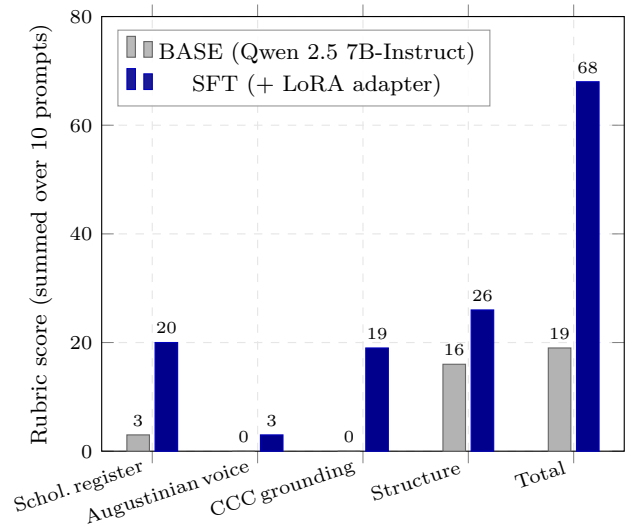


Scrape CCC + Summa + Augustine; clean to 342 JSONL chunks; use Claude Sonnet 4.6 as *teacher* to write 93 scholastic Q&A pairs ($\approx \$0.50$); SFT with LoRA.

Training setup

- Base Qwen2.5-7B-Instruct, MLX 8-bit quantization
- LoRA on top 16 of 28 transformer layers
- AdamW, LR 10^{-5} , batch 1, max seq 2048
- 200 iters, 83 train + 10 valid; wall time ≈ 7 min
- Hardware Apple M4 Pro, 48 GB unified memory
- Peak memory 12.7 GB; trainable params 2.6M / 7.6B (0.034%)

Results



Per-dimension deltas (BASE $\rightarrow$ SFT, max 30) across 10 held-out philosophical prompts not seen in training: scholastic register **3 $\rightarrow$ 20** (+17), CCC grounding **0 $\rightarrow$ 19** (+19), structure **16 $\rightarrow$ 26** (+10), Augustinian voice **0 $\rightarrow$ 3**.

Qualitative example

Q. How do you reconcile divine foreknowledge with free will?

BASE: “The reconciliation of divine foreknowledge with free will is a classic philosophical and theological problem...”

SFT: “... the Angelic Doctor addresses [this] with great precision. **Objection 1.** It would seem that divine foreknowledge destroys free will, for if God knows...”

Phase 2 & DPO negative result

Scaled to 377 pairs with per-source register hints (Aquinas for *Summa*/CCC, Augustinian for *Confessions*/*City of God*), 800 SFT iters; best checkpoint **iter 400**. **Augustinian gap closed** (aug: 3 $\rightarrow$ 7–13); under a balanced rubric ($\max(\text{REG}, \text{AUG}) + \text{CCC} + \text{STR}$), Phase 2 reaches **68/90 vs Phase 1’s 66/90**. A DPO chain (chosen = SFT output, rejected = base) was *inert*: val loss 0.000 from iter 1, accuracy 1.000, margin 35.4 nats. Within-model preference data saturates the policy at initialization.

Takeaway. Small open-weights models move into a niche historical register cheaply when capacity (7B not 1.5B) and quantization (Q8 not Q4) are right. The bottleneck is the *teacher signal*, not the optimization recipe.