



DATA SCIENCE • ADVANCED • THE APPLIED DS/ML CAREER TRACK

LLM Fine-tuning & Deployment

When and how to fine-tune: LoRA/QLoRA, preference tuning, evaluation and cost-effective serving — with GPU labs included.

 4 weeks Live online + GPU labs 4.8 · 180+ trained

PROGRAMME OVERVIEW

About this programme

A precise, engineering-first treatment of model customisation: when fine-tuning beats prompting and RAG (and when it doesn't), parameter-efficient methods, preference optimisation, rigorous evaluation and serving economics.

OBJECTIVES

What this programme sets out to do

- ✓ Build a working command of decision & data
- ✓ Go deep on tuning methods
- ✓ Develop practical skill in eval & serving

AT A GLANCE

Key facts

CATEGORY Data Science	LEVEL Advanced
DURATION 4 weeks · ≈ 24 contact hours	FORMAT Live online + GPU labs
COHORT SIZE Public & private cohorts	RATING 4.8 / 5 (180+ trained)
LANGUAGES 12+ (English, Hindi, Kannada ...)	CERTIFICATE Verifiable digital credential

WHO IT'S DESIGNED & DELIVERED FOR

Cohorts

ML engineers with Python and one shipped model.

ML engineers with Python

one shipped model

Pre-requisites

- Prior hands-on experience in the relevant technical area.
- A laptop for the hands-on segments.
- One or two real use-cases from your work to apply the learning to.

Tentative duration & effort

4 weeks

- ≈ 24 total contact hours
- Live online + GPU labs
- Pacing adapts to your cohort; private cohorts can be re-scoped

WHAT YOU TAKE AWAY

Tangible outputs

- A verifiable, shareable digital certificate
- A role-specific prompt / playbook library
- Qube AI copilot access during the programme
- Session recordings and all working materials
- Entry to the Qubicon alumni community & office hours

MODULE - BY - MODULE

What we cover

MODULE	KEY TOPICS COVERED	HRS*
01 Decision & data	The customisation ladder · Dataset curation · Contamination control · Baseline discipline	8
02 Tuning methods	LoRA/QLoRA hands-on · Preference optimisation · Curriculum strategies · GPU lab	8
03 Eval & serving	Eval harness design · Quantisation trade-offs · Serving stacks · Capstone	8

*Indicative contact hours per module; tuned to your cohort's pace and any private-cohort customisation.

TOOLS & TECHNOLOGIES

What you'll work with

PyTorch

Hugging Face

vLLM

Weights & Biases

PROGRAMME OUTCOMES

What you'll be able to do

- ✓ Decide fine-tune vs prompt vs RAG from evidence
- ✓ Run LoRA/QLoRA and preference-tuning jobs properly
- ✓ Evaluate tuned models against honest baselines
- ✓ Serve efficiently: quantisation, batching, routing

Delivery & certification

- Live, expert-led sessions with hands-on labs and mentoring
- The Qube AI copilot supports every learner between sessions
- Delivered in 12+ languages with domain-accurate translation
- Verifiable digital certificate — verify any code at qubiconai.com/verify

Investment

\$200/hr · ₹7,999/hr

- **Individuals:** enrol & pay online (cards / UPI)
- **Enterprises:** quote → PO → GST invoice → net-30, zero upfront
- Private cohorts, on-site delivery & volume rates on request



For corporates & institutions

- Private cohorts engineered around your stack, data and use-cases
- Curriculum built by principal instructors with the Qube engine
- Capability analytics: engagement, skill-lift and applied-project ROI
- Procurement-friendly: quote → PO → GST invoice → net-30, zero upfront
- MSA, data-security review and on-prem / VPC delivery supported

For individual learners

- Enrol and pay online — cards, UPI or Razorpay
- Live, expert-led sessions (never pre-recorded)
- A verifiable certificate you can share with employers
- The Qube AI copilot supporting you between sessions
- Recordings, materials and an active alumni community

WHY QUBICON.AI

50K+

PROFESSIONALS TRAINED

650+

ENTERPRISE CLIENTS

50+

COUNTRIES SERVED

98%

SATISFACTION SCORE

- Live and expert-led — taught by practitioners who ship AI
- Qube, an adaptive AI copilot, for every single learner

- Verifiable, shareable certification employers trust
- Delivered in 12+ languages, localised to your teams
- Curriculum engineered around your industry and stack



Ready to begin?

Enrol, request a private cohort, or book a free 30-minute consult — our team replies within one business day.

WEBSITE

www.qubiconai.com

Programmes, catalogue & enrolment

EMAIL

contact@qubiconai.com

Corporate & institutional enquiries

OFFICES

Bengaluru · San Francisco

India & USA · delivery in 50+ countries

BOOK A CONSULT

qubiconai.com/book

Free 30-minute capability call

Enrol or request this programme for your team

Reference QBN-LLM-FINE-TUNING-DEPLOY · LLM Fine-tuning & Deployment

qubiconai.com/programs →