

FULL CURRICULUM • 2026

Large Language Model

Build the chatbots and AI agents that companies are paying a fortune for.

Every module, every lesson — the complete syllabus, end to end.

9

MODULES

128

LESSONS

7

MONTHS

₹33,000

FULL PROGRAMME

What this course covers

1

Level 1 • Beginner

Introduction to Course • Building LLM Products

2 months • 3 modules • 27 lessons • ₹8,000

2

Level 2 • Intermediate

LLM Showdown:Evaluating Models • Mastering RAG

2 months • 3 modules • 47 lessons • ₹11,000

3

Level 3 • Advanced

Fine Tuning Frontier LLMs using LoRA , QLoRA

3 months • 3 modules • 54 lessons • ₹14,000

The whole journey

This is the complete syllabus for Large Language Model — 9 modules and 128 lessons, in the order you will learn them. The level markers show how far each stage carries you, but the curriculum runs as one continuous path from the first lesson to the last.

1

LEVEL 1

Beginner

Introduction to Course · Building LLM Products

2 months · 3 modules · 27 lessons · ₹8,000



1. Introduction to Course

1 lessons

- Introduction to course LLM Engineering



2. Building LLM Products

17 lessons

- Introduction to Large Language model engineering
- Ollama setup - windows
- Ollama setup Mac
- Building LLM Applications & best practises
- Key Skills & Tools for AI development
- Understanding frontier models gpt, claude, open source llms
- Power of local llms - hindi tutor
- Creating AI Powered Web Summarizer using Gemini and BS4
- Comparing LLM models
- Leveraging GPTs canvas feature
- Gemini vs Cohere for analytical tasks
- Evaluating metaAI and perplexity
- Introduction to transformers
- Enhancing LLM reliability through prompt engineering
- Understanding llm parameters
- Understanding gpt tokenization
- Claude 3.5s Alignment and Artifact Creation



3. Build a multimodal chatbot : LLMS, Agents, Gradio UI

9 lessons

- Building MultiModal Chatbots
- Mastering multiple AI APIs
- Building Advanced AI Assistants with LLMs
- Building Multimodal AI Assistants
- Creating Adversarial AI Conversations with OpenAI & Claude
- Introduction to Gradio
- Implementing Streaming Responses from LLM APIs
- Advanced Chatbot Development
- Practical Implementation - Chatbot Creation

Intermediate

LLM Showdown:Evaluating Models · Mastering RAG

2 months · 3 modules · 47 lessons · ₹11,000

4. Open Source Gen AI:Build Autonomous Solutions with Gen AI 14 lessons

- Getting Started with Hugging Face
- Exploring Hugging Face
- Google Colab Free Cloud platform for ML
- Running AI Models with Google Colab & Hugging Face
- Using Google Colab+HuggingFace
- Hugging Face Transformers for AI Tasks
- Hugging Face Pipelines Simplifying AI Tasks
- Tokenization for LLMs
- Tokenization in Modern AI LLAMA 3.1
- Comparing Tokenizers in Open source AI Models
- Efficient Inference with Open Source LLMs on Hugging Face
- Efficiently Running LLMs with Quantisation
- Combining Frontier & Open Source Models for Audio- Text
- Building synthetic data for AI Model Development

5. LLM Showdown:Evaluating Models 22 lessons

- Choosing b w open source & closed source llms
- Scaling Laws, Model Size and Evaluating LLM Performance
- Evaluating Real World LLMs beyond Benchmarks
- Choosing the Best LLM for Coding Projects
- Closed vs Open LLMs
- How to compare Open Source Language Models
- Human Evaluation of LLMs with Chatbot Arena
- How LLMS are transforming industries
- AI Powered Python to C++
- Python to C++ Code Implementation
- Comparing GPT-4 and Claude 3.5 Sonnet for code generation
- GPT vs Claude Code Generation
- Leveraging LLMs for Code Optimisation
- How LLMs handle complex code conversion tasks
- Claude vs Python for Code Generation
- Building LLM powered Code Generation UIs
- Open Source LLMs for Code Generation Hugging Face Endpoints Explored
- Deploying Code Generation Models with Hugging Face Inference Endpoints
- Building Hybrid LLM Systems for Code Conversion
- AI Code Generation Comprehensive Comparison
- Evaluating Code Generation Models
- Evaluating LLM Performance in Real World Applications



6. Mastering RAG

11 lessons

- What is RAG(Retrieval Augmented Generation)
- Building a RAG System from Scratch
- RAG Practical Implementation
- Vector Embeddings & their role in RAG
- Simplifying RAG with LangChain
- Effective Text Splitting for RAG Applications
- Preparing Data for Vector Databases in RAG Applications
- Understanding Vector Embeddings & RAG Systems
- Visualising Vector Embeddings
- RAG Hands On Codes
- Troubleshooting & optimising RAG Systems

3

LEVEL 3

Advanced

Fine Tuning Frontier LLMs using LoRA , QLoRA

3 months · 3 modules · 54 lessons · ₹14,000



7. Fine Tuning Frontier LLMs using LoRA , QLoRA

19 lessons

- Finetuning LLMs - From Inference to Training
- Finding and Crafting Datasets for LLM Fine Tuning Sources & Techniques
- Evaluating LLMs
- LLM Deployment Pipeline
- Prompting, RAG & Fine Tuning When to use which approach
- How to create a balanced dataset for LLM Training
- Building Effective ML Baselines for NLP Tasks
- Feature Scaling in ML
- Linear Regression in ML
- Random Forest Regression
- Bag of Words NLP Implementing Count Vectorizer for Text Analysis in ML
- Support Vector Regression in ML
- Comparing GPT-4 Mini to Traditional AI Models in price estimation
- LLMs vs Traditional ML
- LoRA, QLoRA, PEFT
- Fine Tuning GPT Models with Open API & Monitoring Tools
- Mastering LLM Fine Tuning
- Monitoring & Interpreting LLM Fine tuning process
- Why fine tuning LLMs might fail and how to overcome it



8. Fine-tuned Open Source Models to compete with Frontier Models

20 lessons

- Parameter Efficient Fine Tuning for LLMs
- Fine Tuning LLMs with Limited GPU Memory using QLoRA
- Hyperparameters & best practices in LoRA QLoRA
- Reducing LLM Size without sacrificing performance
- Advanced Quantisation Techniques for reducing LLM Memory Footprint
- Tokenization strategies in LLMs
- Efficient Loading, Tokenization and Fine Tuning of LLaMA 3.1
- Navigating HuggingFace's LLM Leaderboard & selecting base models
- Impact of Quantization on LLM Performance
- Optimising LLAMA to compete with GPT 4
- Epochs and training dynamics in LLM Fine Tuning
- Setting optimal hyperparameters for LLM Fine Tuning
- Monitoring & Optimising LLM Fine Tuning with W&B
- Tracking & visualising LLM Fine tuning Progress
- Visualising Model training with W&B , Hugging Face
- End to End LLM Fine tuning for Business Solutions
- Training Process in LLMs
- Token Prediction & Real Time Monitoring in LLM Fine Tuning
- Fine Tuned Models vs Frontier Models & Performance Optimisation
- Evaluating & Comparing Fine tuned LLMs



9. Build autonomous multi agent systems collaborating with models

15 lessons

- LLMs with Multi Agent & Serverless Deployment
- Efficiently running LLMs in cloud environments
- Deploying AI Models as Serverless APIs in the Cloud
- Building production ready RAG Systems
- Understanding Vector spaces in RAG
- RAG Value proposition & core mechanics
- Hybrid Ensemble architecture for pricing
- Mastering Structured Outputs for LLM Automation
- Mastering structured output with GPT-4
- Defining Agentic AI Hallmarks
- The architecture of Autonomy
- Agents Core Loop Building from scratch
- Gradio AI Interface
- Real Time Monitoring for autonomous AI agents
- Autonomus agent framework

Certificates & what comes next

You receive a certificate at the end of every level you finish. Join for a single level or the full programme — the choice is yours. Learners who stand out may be invited to work on live client projects as interns; these places are merit-based, limited, and depend on projects being available.