

Fine-Tuning Transformers for Medical Reasoning with LoRA and Hugging Face Trainer

Friday, 17 July 2026 12:00 (1 hour)

This talk demonstrates a practical, end-to-end notebook for fine-tuning a reasoning-capable transformer model for medical question answering. Participants will learn how to use Hugging Face Datasets and Trainer to handle the training workflow, from loading and cleaning data to tokenization, checkpointing, evaluation, and inference. The session demonstrates parameter-efficient fine-tuning with LoRA, showing how a 3B-class Mistral reasoning model can be adapted on a single 16 GB GPU by training only small adapter weights instead of the full model. The notebook combines MedReason and medical-o1 reasoning datasets into a unified question, chain-of-thought, and answer format, then trains and evaluates the model on a small demo subset. By the end, attendees will understand the key engineering choices behind efficient LLM fine-tuning and see a side-by-side comparison of base and fine-tuned model behavior on medical reasoning tasks, including practical notes on GPU setup, mixed precision, and resource cleanup for reproducible classroom demos.

Presenter: DOLGOPOLYI, Roman (GRNET)