

RISHABH SABHARWAL

+44 7486 045605 | rishabh.sabharwal.app@gmail.com

<https://sabharwalrishabh.github.io/> | <https://www.linkedin.com/in/sabharwalrishabh/>

EDUCATION

MSc in Artificial Intelligence, University of Edinburgh, UK (August, 2026) | **Grade: Distinction***

Coursework: Accelerated NLP, Advanced NLP, ML Systems (Inference Optimization), Probabilistic Modeling & Reasoning

B.Tech in ECE (AI & ML), NSUT, India (2024) | **GPA: 3.68/4** (CVSPK 2024 Fellow, **Top 5%**)

SELECTED PUBLICATIONS

Rishabh Sabharwal, H. Wang, A. Storkey, J. Pan, “Multi-Turn Evaluation of Deep Research Agents Under Process-level Feedback,” in SCALE Workshop, **ICML 2026**. (Oral - Top 6%) ([Link](#))

A.Mazumder, **Rishabh Sabharwal**, ..., P. Rathore, “On Convergence of Adam with Data Dependent Stepsize” in **IEEE Transactions on Artificial Intelligence** (Q1 Journal, Impact Factor: 4.8). ([Link](#))

PROJECTS

Multi-Turn Evaluation of Deep Research Agents (DRAs) ([Code](#))

- Designed a multi-turn evaluation framework and evaluated 3 LangChain-scaffolded DRAs using LLM-as-judge.
- Implemented observability and tracing via LangSmith to monitor tool usage, costs, and latency across runs.
- Exposed a reliability flaw:** the full-rewrite architecture of LangChain deep research causes models to regress on up to **20%** of previously satisfied criteria. Designed a ‘resume path’ for DRAs, which decreased regressions by 4x while consuming 78% fewer tokens and half the latency. **Overall costs reduced by 70%.**

Multi-Step Agentic RAG System for Multi-Hop QA ([Code](#))

- Built a RAG pipeline for multi-hop QA over 270K facts using a **FAISS** vector index and HuggingFace models
- Deployed via **FastAPI** with **vLLM**-served Qwen3-8B for low-latency inference, alongside GPT-4.1-mini. Containerized and orchestrated with **Docker** Compose.
- Designed a decoupled agentic RAG system with multi-step retrieval via LLM tool calling, improving EM from 0.56 (single-step RAG) to 0.67. Achieved roughly 0.17s retrieval p99 and 3.1s end-to-end p99.

LLM Post-Training on Small-Scale Models ([Code](#))

- Fine-tuned Qwen2.5-0.5B on GSM8K via SFT and RLHF; mitigated reward over-optimization using process reward models assigning dense rewards to CoTs, improving accuracy from 31% to 37% on test set
- Employed a latent-space knowledge distillation framework, transferring K latent reasoning tokens from Qwen2.5-7B directly into the student’s representation space, achieving 41% accuracy on the test set.
- Experimenting with Qwen2.5-32B / 72B as teacher model trained via DeepSpeed ZeRO-2 distributed training.

EXPERIENCE

Research Engineer

Jan 2024 - Jul 2025

Indian Institute of Science

Bengaluru, India

- Collaborated to design a parameter-free module for spectral GNNs, reducing pre-training memory footprint **50x**, GFLOPs **100x**, and latency **6x** while achieving state-of-the-art results across 20 real-world datasets.
- Led a team to design a plug-and-play graph-diffusion-inspired attention module for ConvNets. Our module **adds just 20 parameters to ResNet-50 yet outperforms ResNet-101** (~20M more parameters); generalizes across ConvNeXt-V2 and MogaNet backbones. Tested deployment on edge device (Jetson).

SKILLS

- Programming & Frameworks:** Python, Bash, SQL, Pandas, NumPy, Scikit-Learn, PyTorch, FastAPI
- ML/LLM Tooling:** HuggingFace, PEFT (LoRA, QLoRA), vLLM, Weights & Biases, AWS, Vertex AI
- Generative & Agentic AI:** LangChain, LangGraph, Prompt Engineering
- DevOps/Systems:** Git, Linux, CI/CD, Docker, Distributed Training (PyTorch DDP)