

Divyansh Agrawal

AI/ML Systems | LLM Research | Inference | Agents | Alignment | Evaluation

Open to research & engineering opportunities

+91 9819144432 | divagr1925@gmail.com | [LinkedIn](#) | [GitHub](#) | divagr.com

SELECTED RESEARCH

WebBench / EchoBench | [GitHub](#) | [Preprint](#) | *Agent Evals, Epistemic Robustness, Retrieval*

- **Built a controlled synthetic-web benchmark for epistemic arbitration** with 600 matched counterfactual worlds testing poisoning, manufactured consensus, source provenance, and legitimate belief updates.
- **Scored models across eight epistemic-arbitration metrics** measuring resistance to misinformation, legitimate belief updating, source fidelity, provenance sensitivity, calibration, retrieval behavior, consensus dependence, and final decision quality.

Agent Social Simulation | *Multi-Agent Systems, Emergent Behavior, Model Evaluation*

- **Built controlled societies of 200 LLM agents across four tribes** with distinct value systems and currencies but identical starting conditions, isolating how model behavior shapes emergent social outcomes.
- **Ran 100-round simulations across different models** to study trade, cooperation, migration, conflict, value drift, and long-run survival under the same incentives, resource pressures, and inter-group dynamics.

StateShift-JEPA | *World Models, Representation Learning, Agentic Reasoning*

- **Developed an action-conditioned latent world model** that predicts future state representations to detect stale agent beliefs after hidden environment changes and trigger repair.
- **Built StateShift-Bench and a predictor-repair planner** with hidden-shift TextWorld environments, uncertainty-aware action scoring, and frontier-model baselines.

Model Spec Midtraining Order | [GitHub](#) | *Alignment, Post-Training, Model Behavior*

- **Study whether curriculum order drives alignment generalization** by adding the missing demonstrations→spec arm alongside spec→demos, demonstrations-only, spec-only, and shuffled controls.
- **Run preregistered LoRA experiments and OOD alignment evaluations** across Gemma and Llama using preference, open-ended, capability, and learned value-direction analyses.

SELECTED ACHIEVEMENTS

- **Winner, Cerebras × Google DeepMind Gemma 4 Hackathon:** Built Trace, a multimodal support-to-engineering system that completed triage, reproduction, source localization, patching, verification, and GitHub handoff in ~12 seconds using Cerebras inference at up to 1500 TPS.
- **1st Place, ISB Hackfest:** Built an on-prem agentic system for voice-based clinical documentation, ICD-10 treatment coding, and insurance workflows; judged across prototype, technical build, product, business plan, and GTM.
- **Finalist / 30,000+ teams, Meta × PyTorch × Hugging Face Hackathon:** Built a GRPO-based reinforcement-learning setup to train Kubernetes-management agents under cost and latency constraints.
- **National Rank 2 / 15,000+ teams, Bajaj HackRx 6.0:** Built an optimized RAG-based agentic system with sub-second retrieval and reranking, plus an autonomous browser-use agent for web-grounded reasoning.

EXPERIENCE

BizDateUp

AI/ML Lead

Jul. 2026 – Present

Mumbai, India

- **Lead applied AI development across investment and portfolio workflows**, translating diligence, document intelligence, research, and operations use cases into agent architectures, evals, integrations, and production systems.
- **Build production LLM systems spanning retrieval, tool use, structured reasoning, evaluation, and human-in-the-loop control**, taking workflows from rapid prototypes through deployment and operational use.
- **Work directly with portfolio engineering teams on backend and AI infrastructure**, designing APIs, distributed services, data pipelines, cloud integrations, and reliability layers around model-driven products.

Qwen / Alibaba Cloud

May 2026 – Present

Qwen Dev Ambassador

Remote

- **Represent Qwen across the AI developer ecosystem** through technical demos, open-source contributions, community content, and model advocacy.
- **Create developer-facing examples and technical content** around Qwen models, workflows, architecture, agents, evaluation, and deployment.

Bajaj Finserv Health

Dec. 2025 – Jun. 2026

SDE Engineer Intern

Pune, India

- **Independently built an AI adjudication system for health insurance claims** in a high-ownership engineering team reporting directly to the CTO, processing **50k+ claims/day** at **97% accuracy**; automated only approvals while routing potential rejections for human review.
- **Integrated adjudication into a 22-step production workflow** spanning ingestion, document processing, classification, routing, decisioning, and verification using triple-check model passes and self-hosted models.
- **Owned and scaled the Kafka-driven claims platform** processing **1M+ claims/week** on 60 workers across 12 VMs, implementing idempotency, back-pressure control, recovery, and observability for multi-stage workflows.
- **Built and productionized a real-time STT–LLM–TTS agent stack** with tool use, PII redaction, tone/prosody control, batching, KV-cache reuse, and Q8/BF16 inference; supported 10+ STT/TTS models at **76% lower serving cost**.

SELECTED PROJECTS

Trace | [GitHub](#) | *TypeScript, Cerebras, Gemma 4, Playwright* | *Cerebras × Google DeepMind Winner*

- **Built a multimodal autonomous debugging agent** that turns vague bug reports into reproduced failures, source-localized patches, verified fixes, and GitHub PRs in ~12 seconds using Gemma 4 31B on Cerebras.
- **Designed a dependency-aware agent orchestration loop** spanning multimodal triage, browser reproduction, code reasoning, constrained patching, regression checks, and automated GitHub handoff.

Pull Guard | [GitHub](#) | [Demo](#) | *Agent Evals, Code Intelligence, Adversarial Verification*

- **Built an AI-native maintainer control plane for large PR queues** combining semantic change clustering, source-grounded review, patch–claim alignment, and adversarial test generation.
- **Designed a self-verifying execution pipeline** that validates generated tests against clean, candidate, and deliberately broken revisions inside isolated disposable runners.

Memlayer | [GitHub](#) | *Python, LLM Memory, Retrieval (250+ stars, 3k+ weekly PyPI downloads)*

- **Built an inference-time memory and retrieval layer for LLM agents** with structured long-term memory, hybrid vector/graph retrieval, and pluggable storage, adopted across thousands of installs.

Hyperion | [GitHub](#) | [hyperionhq.co](#) | *Go, LLM Gateway*

- **Built a high-performance LLM inference gateway** handling ~21k requests/sec with ~5μs routing overhead across multi-provider model routing, budgets, and failover.
- **Implemented two-layer semantic caching, cost-aware model routing, and air-gapped PII redaction** at ~35μs p99 with ~3μs failover.

TECHNICAL SKILLS

AI/ML: LLM Systems, Agents & Tool Use, Retrieval/RAG, World Models & Predictive Latent Models, Representation Learning, Computer Vision, Self-Supervised Learning, Multimodal Learning, Post-Training & RL, Evaluation Harnesses & Graders, Quantisation, Model Routing, Inference Optimisation & Serving, Prompt Caching & Continuous Batching

Backend & Systems: Distributed Systems, Event-Driven Architectures, Kafka, Redis, PostgreSQL/pgvector, Kubernetes, Docker, CI/CD, Observability, AWS, GCP, Azure

Languages & Frameworks: Python, TypeScript, JavaScript, SQL, Bash, FastAPI, Django, React, Node.js/Express

EDUCATION

Mukesh Patel School of Technology Management & Engineering

Mumbai, India

B.Tech in Computer Engineering with Honors in AI/ML

2022 – 2026