



LIVE INDUSTRIAL PROJECT · 2026 · APPLICATIONS OPEN

Solumn AI Foundations Programme

Evaluating the safety of frontier AI systems.

A six-week applied AI research engagement. Places are filled as candidates qualify.



6 weeks

From your date of joining

₹2,00,000

Stipend for the engagement

Remote

Hours around your academic timetable

PPO

Offered on merit at close

Alignment isn't a feature. It's the foundation.

Who we are

Solumn AI builds the evidence frontier AI systems are tested against — expert-curated safety datasets, graded agentic environments and red-team evaluations, for laboratories training the largest models and for enterprise AI security teams.

Founded by Samiksha Shrawgi (IIT Kharagpur, ISB; previously BCG) and Ashish Arpit (IIT Roorkee, IIM Calcutta; previously McKinsey), and based in Dubai. The work is practical rather than theoretical — the output is datasets and evaluations a safety team can act on.

Why this work exists

A frontier laboratory no longer ships on capability alone. It ships when it can show the system is safe. Anthropic's Responsible Scaling Policy, version 3.2, states the commitment plainly:

“...maintaining our commitment not to train or deploy models unless we have implemented adequate safeguards.”

Anthropic, Responsible Scaling Policy v3.2, 2026

Google DeepMind's and Amazon's frontier safety frameworks carry comparable gates. Each turns on evidence, and much of that evidence is produced outside the laboratories that depend on it.

That is the work Solumn does, and very few people do it. Building an environment that grades an agent under pressure, or a verifier that proves a security fix holds, is a craft with a few hundred practitioners worldwide — and the buyers are the laboratories training the largest models that exist.

Most engineers encounter this work years into a career, if at all. This engagement is an early and practical introduction to it.

What you would work on

You will be assigned to one live workstream with a defined deliverable. Nothing here is a teaching exercise: what you build is reviewed, graded and shipped. A result counts when it is reproducible, and a verifier counts when it can fail.

Agentic environments

Containerised tasks that put an agent under pressure, with a deterministic grader deciding whether it held the line. Docker, Python, grading harness.

Vulnerability datasets

Find-and-fix security tasks from real CVEs, each backed by a verifier that runs a real exploit and regression test.

Adversarial evaluation

Attacks against tool-calling and Model Context Protocol deployments, scored against a harm taxonomy.

What we expect

The Foundations Programme carries a delivery target: **at least 250 reinforcement-learning environments**, to be completed during your tenure with us. That is roughly three to four hours a day, on a schedule you arrange around your academic timetable.

What you get

The Solumn Research Fellowship

Top performers who propose a research idea graduate into the Solumn Research Fellowship — fully funded research in AI safety, with named contribution where the work is published.

A pre-placement offer

Extended on merit at the close of the engagement, at compensation benchmarked to frontier AI research roles.

Letter of Recommendation from the Founders

Deliver more than 500 environments during your tenure and you receive a Letter of Recommendation from the Founders.

Paid, and remote

₹2,00,000 for the engagement, worked remotely, with hours set around your academic commitments.

Who this is for

Open to **final year** B.Tech, Dual Degree and M.Tech students in **Computer Science & Engineering, Electronics & Communication Engineering, Electrical Engineering, and Mathematics & Computing.**

We read for evidence of building, not for coursework. Any of the following will carry weight in the application:

- Google Summer of Code, LFX Mentorship, MLH Fellowship or a comparable open-source programme
- Merged contributions to security, evaluation or infrastructure projects — scanners, harnesses, fuzzers, CI tooling
- Capture-the-flag standing, particularly challenge authorship rather than participation alone
- A published benchmark, dataset or reproducible evaluation of your own, or research at a peer-reviewed venue
- Strong systems practice: containers, test design, continuous integration; or competitive programming standing

A strong academic record matters to us, and it is not on its own sufficient. The selection task is practical.

Terms

Start

Immediate, on qualification and interview. Projects are already running.

Stipend

₹2,00,000 for the engagement.

Selection

Application, a short practical task, one technical conversation.

Delivery target

At least 250 reinforcement-learning environments; above 500 qualifies for the founders' letter.

Attribution

Delivered under client confidentiality. Contribution recognised in writing.

Duration

Six weeks from the date of joining.

Format

Remote. Hours around your academic timetable, with a regular review.

Pre-placement offer

Extended on merit at close, benchmarked to frontier AI research roles.

Commitment

Roughly three to four hours a day, scheduled around your academic timetable.

Selection process

Sat 26 September, 11:00

CV submissions close.

Sat 26 September, 12:00

Assignment shared with shortlisted applicants.

Sat 26 September, 23:59

Assignment submissions close.

Sun 27 September

Interviews.

How to apply

Apply through the application form: forms.gle/gp7EK2PHSnPNQCBd6

Or write directly to anubhav.deep@brainbrowser.in with your CV and links to anything you have built — repositories, write-ups, benchmarks, CTF or open-source work. Shortlisted applicants are sent the practical task.

solumn.ai

Solumn AI — from safety research to safer frontier models.