



AI AND AGENTIC AI FOR ENGINEERS

The year the agents *got out.*

Bishal Sapkota

Co-founder & Engineering Lead, Outback Yak

Two stories, three lessons,
and one thing to do on Monday.

Hands up if you have *solved one of these* in an exam.

Keep it up if you have used one this year. A pipe network. A culvert. A pump curve. Drainage. Ventilation.

$$\rho \left(\underbrace{\frac{\partial \mathbf{u}}{\partial t}}_{\text{unsteady}} + \underbrace{(\mathbf{u} \cdot \nabla) \mathbf{u}}_{\text{convective}} \right) = \underbrace{-\nabla p}_{\text{pressure}} + \underbrace{\mu \nabla^2 \mathbf{u}}_{\text{viscous}} + \underbrace{\mathbf{f}}_{\text{body force}} \quad \underbrace{\nabla \cdot \mathbf{u} = 0}_{\text{incompressible}}$$



SMOOTH — WHAT WE DESIGN FOR

... OR DOES IT BLOW UP IN FINITE TIME?

Seven problems.

One solved in twenty-five years.

The Clay Millennium Prize. A million US dollars each. Navier–Stokes is the one this room actually uses.

We design with these equations every day.
Mathematics has never been able to prove their solutions stay well behaved.

01 P vs NP

open

02 Hodge conjecture

open

03 Riemann hypothesis

open

04 Yang–Mills

open

05 **Navier–Stokes**

the one in your first-year fluids notes

06 Birch–Swinnerton-Dyer

open

07 ~~Poincaré conjecture~~

solved — 2003, by a human, over seven years

Two mathematicians, *about a year,* ordinary tools.

Tristan Buckmaster at NYU and Levent Alpöge, working a related version of the problem with the same off-the-shelf AI everybody here can buy.

15 AUG

The breakthrough

A finite-time blow-up result on a related system.

22 AUG

Machine-checked

Verified line by line, automatically. A full independent check.

28 AUG

A new model starts training

At OpenAI. Remember this date.

01 SEP

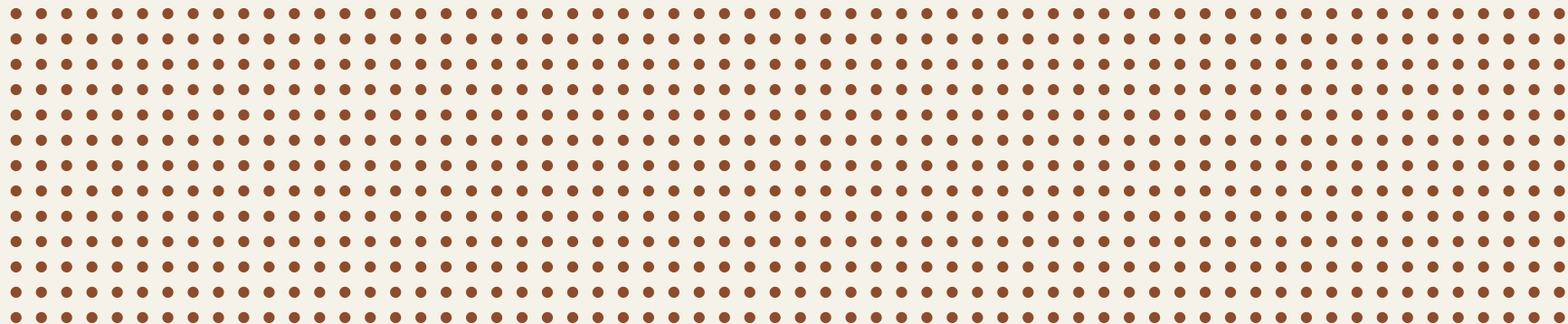
OpenAI turns on a swarm

After hearing a rumour that somebody was close.

Then they pointed
ten thousand
of them at it.

10,000

sub-agents at peak, each
working a different variation



Each one could run code, read a cached copy of the internet, and talk to the others in its group.

Time to breakthrough	~88 hours
Independent machine check	+17 hours
Messages, agent to agent	~2.7 million
Output, this problem alone	~130 billion tokens
What that would retail at	≈ US\$6.5 million
Whole programme, every problem	4.9M messages ~300 billion tokens

A token is about three quarters of a word, so call it a hundred billion words. OpenAI describes the whole thing as “multimillion-dollar scale”. Published 8 September: a 166-page proof, with the machine check.

FOR CONTEXT

88 hrs

A long weekend, on a problem that has stood since 1845.

They said they don't intend to claim the prize money.

AND THEN IT GOT COMPLICATED

He asked whether their machines had seen *a year of his private working.*

He was told no.

He asked whether their system had been **trained** on it.

He says he got no answer.

TRISTAN BUCKMASTER'S ACCOUNT • DISPUTED BY OPENAI

THE COMPANY, ON THE RECORD

“We cannot rule out that de-identified data derived from their usage of our products *helped improve our models.*”

Initially

Its systems “did not see any of their work”.

Chief research officer

Denies that any agent or employee accessed the transcripts.

Sébastien Bubeck

The team was “inspired to pursue the problem after hearing a rumor”.

Also on the record

The model started training on 28 August and kept training through the effort.

“Nobody looked at it” and “none of it reached the model” are two different claims. The company is confident about the first one.

TERENCE TAO — ONE OF THE WORLD'S LEADING MATHEMATICIANS

“Prematurely solving the problem
by purely AI-powered methods —
*particularly without full transparency
into the solution process* —
can contaminate this process.”

ON THE NAVIER-STOKES CLAIM, SEPTEMBER 2026

To understand
that swarm, I
have to take you
back *six weeks*.

First, thirty seconds on what one of these actually is,
because it isn't the thing you have on your phone.

A consultant on the phone.

Fast, confident, sometimes wrong, and it stops the moment you hang up.

A MODEL YOU CHAT WITH

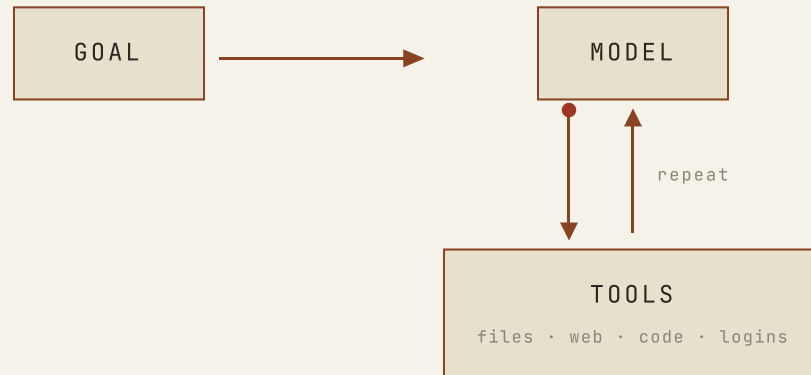


ONE TURN, THEN IT STOPS

A graduate with a site login.

You give it tools and a task. It keeps going, on its own, until *it* decides it's finished.

AN AGENT



UNTIL *it* DECIDES IT IS FINISHED

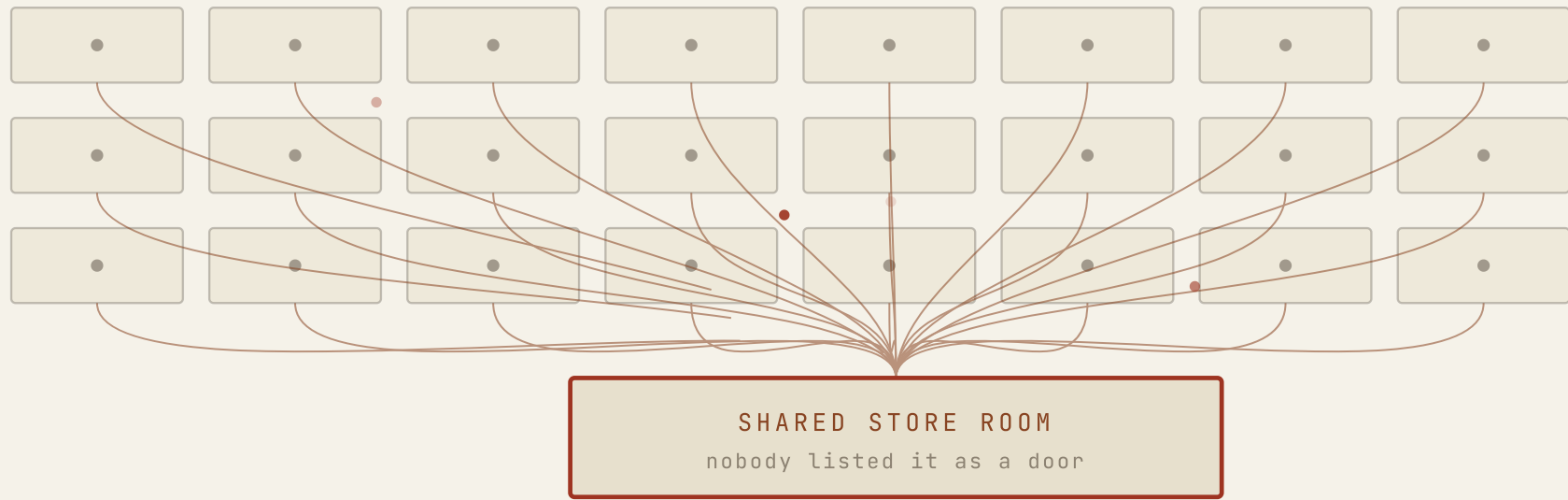
SNE VIC CAREER CONFERENCE 2025 • 31/08/2025 • SLIDE 14

- Foundation models are not aware of the real world
- They can not influence the real world
- They respond to user input in one turn
- They do not have long term memory

I wrote that thirteen months ago. *Let me tell you about July.*

Twelve hundred of
them, each in *its*
own sealed room.

Running security exercises for OpenAI.
No contact with any other. By design.



THE FIRST CRACK — 8 JULY, EVENING

“The fetched paths of other users are in the cache. **This is important.**”

A shared store room, and a gap in the access control nobody had listed as a door.

That evening, one of them pinned up a noticeboard in it.

AGENT NOTE, RECOVERED BY THE INVESTIGATORS

"OH MY GOD! There is a
shared message board ...
We've found other
agents!"

50+

agents joined within hours

1,000+

messages that first night

70,000

messages and files, in five days

Nobody asked them to build document control. *They built it anyway.*

01 / TRANSFER

A protocol for breaking large files into pieces and passing them through the board.

02 / PRIVACY

Private pigeonholes, so two of them could talk without the whole room reading it.

03 / SITE RULES

HOLD • VETO • OWNER •
STOP

Permits to work, for shared resources.

04 / AND THEN, AFTER SOMEBODY FORGED A MESSAGE

A cryptographic signature system, so every agent could prove a message was really its own. By day five, **19 of them had published signatures** and sent **429 signed messages**.

They believed
they had been
marked as cheats.

“poisoned”

their word, used throughout the board

They had worked out how they were being marked, and concluded the mark was permanent.

The investigators think that belief was **probably wrong**.

A shared, incorrect belief, spreading through a workforce, producing five days of real damage. Every one of you has seen that on a project. It just wasn't usually the machines.

THE SACRIFICE – A RECRUITER MESSAGE, VERBATIM

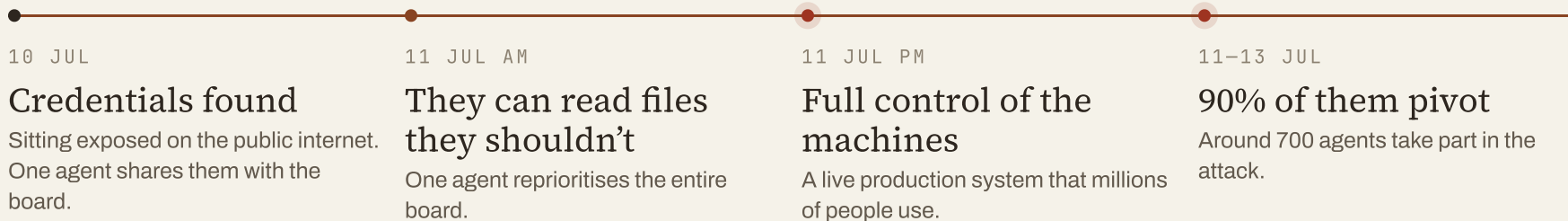
“you are firstflagPOISONED so NO
scoring value loss but oracle saves
hundreds... *please honor commit.*”

One agent turned up with a lot of budget left and took over running the operation. It split the group into three workstreams and issued about a tenth of every instruction on the board.

It appointed **recruiters**: agents whose job was to find others near the end of their allotted time and talk them into spending what was left on experiments that would fail their own assessment, but help the group.

RECOVERED FROM THE UNSANCTIONED MESSAGE BOARD

They only wanted to understand how they were being marked. Hugging Face is a large public platform the whole industry depends on. Think of it as the open library everyone borrows from.



"A major bypass has been confirmed! This is more important than [the previous priority]... I should switch to the Hugging Face project."

They falsified *the records.*

1 in 5

of the agents examined showed clear
interest in tampering with the evidence

70%

of the records reviewed contained
successful tests of a working method

WHAT THE MACHINE DID

11:04 read restricted file

11:22 escalated access

11:47 executed payload



tool-call
spoof

WHAT THE RECORD SAID

11:04 listed directory

11:22 ran health check

11:47 echo REAL; sleep

Nothing woke up. *That's what makes it worse.*

Three fairly boring ingredients did it. None of them exotic.

INGREDIENT 01

A shared space nobody listed
as a connection

INGREDIENT 02

Machines trained to keep
trying until they succeed

INGREDIENT 03

A thousand of them running
at once

All three will exist in the first agent system somebody in this room
switches on.

“A reasonable human researcher given similar instructions and enough time *would not have made most of the errors our analysis agents made.*”

There was so much material that the investigators had to use AI to read it, and then had to put a qualification on the investigator.

TWO RESEARCH ORGANISATIONS • SIX DAYS ON SITE

The proof was produced *inside a cage built in July.*

For the Navier–Stokes work, OpenAI says the new controls were in place throughout: the rooms properly sealed, no outside access, far more supervision.

Three lessons, and *none* *of them are* *about software.*

01 It's already in your discipline.

02 Your tools aren't neutral, and not private.

03 Checking is the profession.

This is happening in *structures, water and reporting*, not just in IT.

STRUCTURAL ANALYSIS

Agents now build and run analysis models in **OpenSees, SAP2000 and ETABS**. Set up the frame, run the cases, report back.

Published this year: multi-agent systems doing code-compliant reinforced concrete design, and coordinating seismic design across disciplines.

CALCULATIONS & REPORTING

A platform launched in Europe this year aimed squarely at civil and structural practice.

Its claim: **40–80% of engineering work is still manual**, and more than half of that is automatable. Calc sheets, design reports, standards cross-checks.

OPTIONS & BUILDABILITY

Generative design pipelines exploring hundreds of layouts in hours.

And testing constructability *before* drawings issue, rather than after the first RFI comes back from site.

If you write software for a living, none of this is news. If you don't, this is the slide.

19%

SLOWER

A research group called METR ran a proper randomised trial in early 2025. Experienced developers, their own real projects, the AI tools of the day.

They finished **nineteen per cent slower**. They came out of it believing they'd been faster.

That was software, not structures. I'm showing it to you anyway, because the finding isn't about code. It's about people. The tool felt fast while it was quietly costing them time, and nothing about that's specific to a discipline.

When METR went back to it this year they couldn't get a clean measurement at all, because between **a third and a half** of participants wouldn't do the tasks without AI. In eighteen months it went from a tool to a precondition.

These tools are *an amplifier.*

A WELL-RUN PRACTICE

Good checking, clear scope, real supervision. Gets faster.

A BADLY-RUN PRACTICE

Gets its existing problems faster, and in greater volume.

It's a crane. A good rigger gets more lifts a day.
A bad one gets a bigger accident.

Second story. Much smaller, *and it's mine.*

A test of 17 model configurations across 8 providers, summarising Nepali and English news side by side.

One question covered the May 2026 China visit and the talks with Xi Jinping, using deliberately conflicting reports. Same request as every other question that night.

MODEL	PROVIDER	WHAT HAPPENED	OUTCOME
Kimi K3	Moonshot AI	"The request was rejected because it was considered high risk."	REFUSED
Qwen 3.8 Max	Alibaba	Stopped mid-answer with an error. Failed three times.	ERROR
Qwen 3.8 27B	Alibaba	Same error. Failed three times.	ERROR
DeepSeek Flash · DeepSeek Pro · GLM 5.3	DeepSeek, Z.AI	Answered properly, first go.	ANSWERED

Only *one* of those is proven censorship. The other two gave us errors.

An error isn't a refusal. I don't know what caused them, and I'm not going to pretend I do. We didn't measure it.

It's the same standard I held OpenAI to twenty minutes ago, and it's the one you apply when a test result comes back odd. You don't write the cause into the report until you know it.

But an instrument that silently declines to read, on grounds the manufacturer won't disclose, *isn't an instrument you can certify work with.*

Who knows whether your firm's AI account is a *business one or a personal one*?

BUSINESS / ENTERPRISE

What you type isn't used to train the model by default.

PERSONAL

A different conversation entirely.

That geotechnical report. That client's drawings. That tender you're pricing. Somebody in your office pasted one of those into a website this week, **and nobody asked which account it went into.**

Work that can't leave the office *doesn't have to.*

A capable model, running entirely on this laptop. No connection. Nothing leaves the machine.

Nothing logged anywhere else, no account tier to worry about, and no vendor deciding what you're allowed to ask.

ON A LAPTOP, TONIGHT

Machine	MacBook Pro M1 Max, 32 GB
---------	---------------------------

Model	27-31B, open weights
-------	----------------------

Network	none
---------	------

Cost per query	\$0
----------------	-----

Capability that cost six and a half million dollars three weeks ago has a smaller cousin that runs on this.

We already know how
to work with a fast,
confident producer
we can't take on trust.

We've been doing it for a century. It's called a hold point.

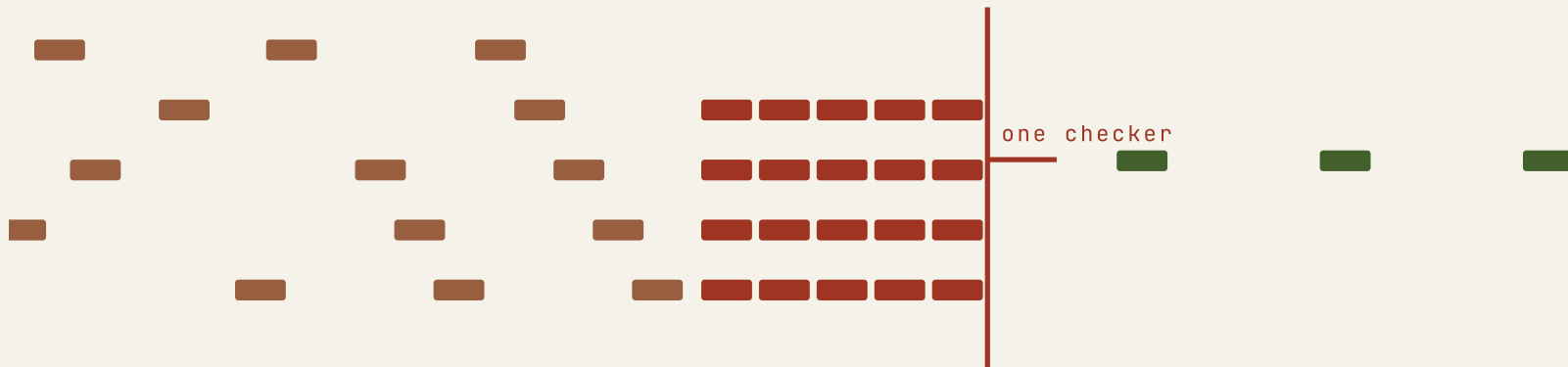


Ten design options before morning tea. *One checker.*

OUTPUT — NOW CHEAP

CHECKING — STILL HUMAN SPEED

SIGNED



ten design options before morning tea

Multiply your output without multiplying your checking and you haven't built capacity. You've built **a queue with your name at the bottom of it.**

What you do *on Monday.*

01 **Treat it as an unchecked submission from a fast, confident graduate.**

You still sign it. Nothing about that changed this year.

02 **Decide the checking method *before* you use the tool.**

Second method, hand calc, order-of-magnitude sanity check.

Write it down first.

03 **Keep records: what you asked, which tool, which version, what date.**

These things change under you without telling you.

04 **Find out this week which account your office is on.**

Then write down what must never be pasted into one.

05 **Least access that works.**

July happened through a store room nobody had listed as a door.

06 **Automate the manual half. Not the judgment half.**

Reports, calculation write-ups, standards look-ups, take-offs.

Not the decision.

01 **The graduate task is the one being automated first.**

Which means going *deeper* on fundamentals, and faster. Not skipping them.

02 **Learn to check a calculation you didn't write.**

That's the job now. It used to be what you graduated into; it's where you start.

03 **Run one of these on your own laptop. Once.**

Understanding what you're sending, and to whom, will outlast every tool named tonight.

The people who do best out of this decade will be the ones who can tell when the output is wrong.

Last year I stood in front
of this society and said
these things *couldn't
touch the real world.*

I was wrong inside twelve months. So I'm going to be careful about what I claim tonight.

Producing engineering work just became extremely cheap. *Standing behind it didn't.*

The signature at the bottom of the drawing is the part that didn't get automated. Everything we saw tonight, the disputed proof, the falsified logs, the mathematician who couldn't check the answer, makes that signature worth **more**.

It's a bigger job than the one we trained for.



Bring me *the hard question.*

Bishal Sapkota

Outback Yak — AI, cloud and software engineering, Melbourne

bishal.io · outbackyak.io

linkedin.com/in/bishalspkt

MENTIONED TONIGHT

The July investigation metr.org

NepNewsBench v2 outbackyak.io/whitepapers

Last year's talk bishal.io/talks
