

BUILT WITH AI. GOVERNED BY ME.

I built my own AI intelligence system.
Then I governed it like an enterprise
AI system.

WHAT IT DOES

It watches six lenses on my world, scores what matters, and briefs me weekly.

AI Intelligence

Cybersecurity

Regulation

Jobs

Financial Context

AI Forecasting

Delivered to my phone and a private dashboard. Running autonomously since July.

7,628

items collected

113

automated runs

9

scheduled tasks

~\$13

logged API cost

Documents before code. Code before judgment.

My first commit was six governance documents — risk log, security audit trail, decision journal, governance register. No code until the risk entries existed.

Hard Rule #1

No code is written before the relevant risk entries exist.
12 hard rules govern the system this way.

Where a wrong answer matters, code decides — not the model:

Source trust tiers — fixed per source by config

Legal disclaimers — fixed text, never paraphrased

Forecast dates — validated by code before storage

35 risks assessed

85 decisions recorded

18 governance entries

Humans in the loop, and against my own echo chamber

The system can draft. It can't decide.

LinkedIn assistant

Only drafts. It has no ability to post.

Chatbot actions

A whitelisted set, each needing two-step confirmation.

Forecast verdicts

I mark each one Held, Missed, or Unclear — never the system.

A system that learns my interests could just tell me what I believe.

A guaranteed weekly item that scores low on my current views

A floor so no interest area can be voted to zero

Skeptics of fast AI timelines included deliberately

WHAT BROKE

"Reported success every week for two months. It hadn't produced anything new since July."

A shared budget let a growing scoring backlog spend every call before the weekly briefing got a turn — and the code treated "nothing to report" as success. I only found it by asking the system for its own usage numbers.

Small failures, caught the same way

The "weekly" briefing wasn't weekly — it resurfaced items from months back, not the last 7 days.

A backup never ran — zero errors, and a misleading "last result: success."

The chatbot once claimed a change had succeeded when it hadn't actually run.

STILL OPEN

Not every problem is fixed. That's logged too.

An AI judgment leak between two scores it should keep independent — reduced, not eliminated.

A scoring backlog is currently starving my financial briefing section.

Logged as open risks — not quietly called fixed.

BUILT WITH AI, GOVERNED BY ME

The AI assistant was fast. It was also wrong sometimes.

The governance process is what turned those mistakes into logged findings instead of silent failures.

Speed came from the AI. Trustworthiness came from the process.

TAKEAWAYS

- 1 Monitor outcomes, not exit codes.
- 2 An AI saying "done" is a claim, not a receipt.
- 3 Label uncertainty instead of hiding it.

How are you governing the AI tools you build with?