



■ LUNCH & LEARN • BRINGING AI ABOARD

Towing the Fleet Forward

Bringing AI and LLMs aboard, under governance, starting now.



■ RESOLVING THE PARADOX



We can't out-wait it, so we tow

The trap holds only if hull and engine are welded together. **Towing splits them.**



- 1 Keep the hull** — our people, operations, knowledge — the vessel we've built.
- 2 Swap the engine** — hitch to the newest; never bet the fleet on one.
- 3 The flip** — the faster ship becomes our next engine — waiting loses.

■ WE TOW WHEN — A TRIGGER, NOT A DATE

$$\text{value}_{\text{now}} > \text{friction}_{\text{switch later}}$$

where **value** = what starting today earns us · **friction** = the cost of re-hitching when a better engine arrives

Capability clears the bar · governance is ready · cost beats manual.

(A framework built around the uncertainty itself — the way a qubit harnesses it.)

The paradox never closes — towing means it doesn't have to. **We clear all three triggers today, so we start. That's the case for this class.**



Two AIs on the table, and a gap that’s closing

Governed enterprise AI

OUR LANE — EXAMPLES OF THE CLASS, NOT COMMITMENTS

- ✓ Compliant by default — runs inside a governed compliance boundary
- ✓ Our data never leaves the boundary
- \$ Seat-licensed — predictable cost, already in our stack
- *Constrained, for now*

Copilot (gov cloud)

ChatGPT wrappers

IBM watsonx

Frontier models

AVAILABLE NOW — NOT IN OUR ENVIRONMENT

- ✓ Fast and highly capable
- ✓ State-of-the-art reasoning
- \$ Metered tokens — pay-per-use, easy to start, easy to overrun
- *External and unbounded*

Claude

GPT-5

Gemini



The gap is closing. Private cloud can bring whole frontier models inside governed walls — or split them down (distillation, SLMs) when it fits — on whichever compliant platform we ride. The “safe vs. powerful” trade-off is dissolving.



Tow the fleet

Elevate what already works, no rip-and-replace. Three lines, one cable — **and today we pull hard on exactly one of them.**

- 01

Software

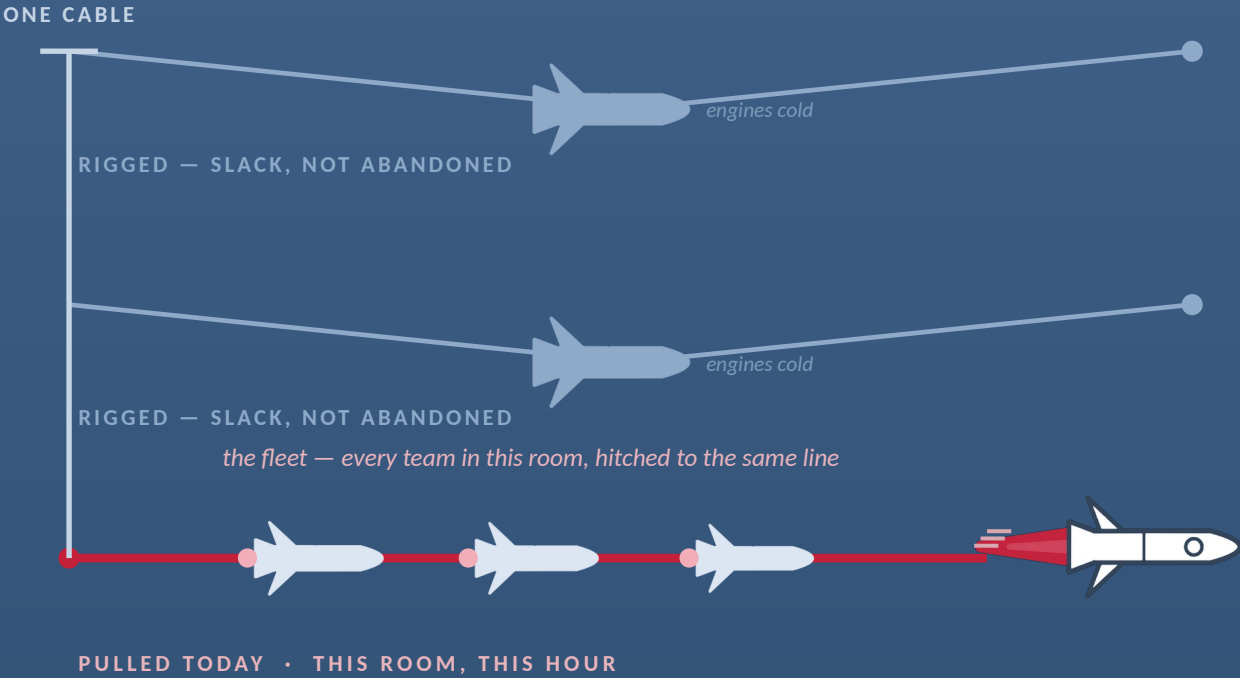
Wrap legacy systems in AI layers. Modernize the intelligence, keep the proven core.
- 02

Operations

Accelerate proven workflows inside policy-compliant guardrails. Speed within governance.
- 03

People

Arm our people rather than replace them. This is where the real change happens — and it's the only line we pull today.

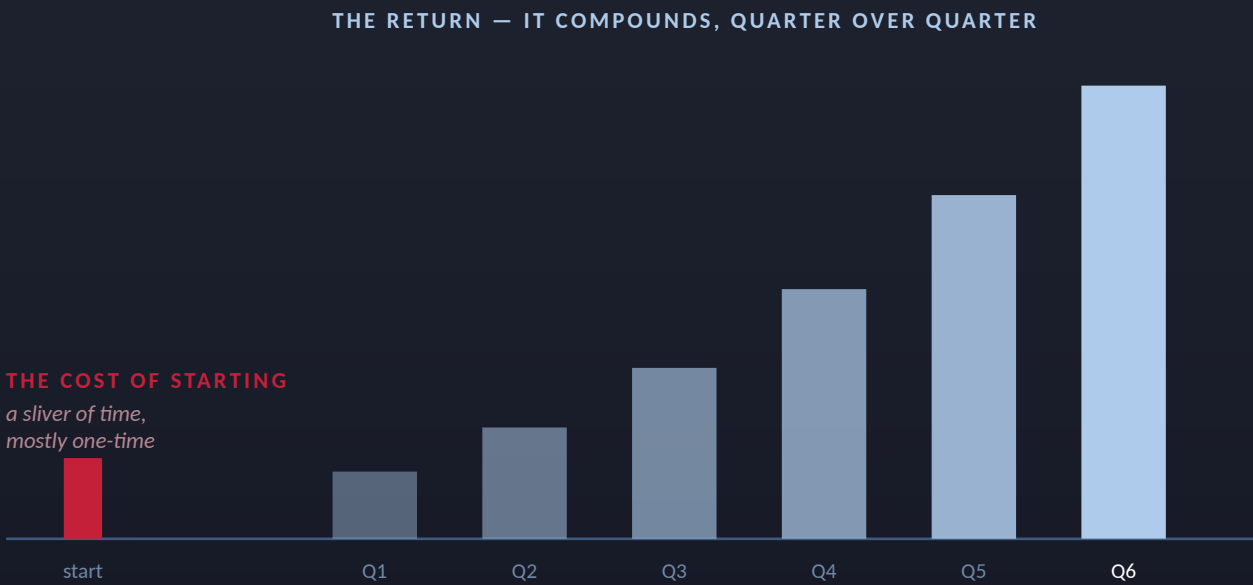


Governance is the cable's tension, not a brake on it. *Software and Ops are already rigged — they're the next pull, not the forgotten ones.*



Small cost now. Compounding return.

Drawn to scale — and the advantage it buys is the part you can't purchase later.



The model was never the moat.

- Your data — organized so AI can actually use it
- Your workflows — rebuilt around the new speed
- Your people — fluent from thousands of reps

$$\text{value}_{\text{now}} > \text{friction of switching later}$$

The trigger from earlier — and it's already cleared.

HOW IT COMPOUNDS — THE FLYWHEEL, AND EVERY LAP ADDS TO THE STACK

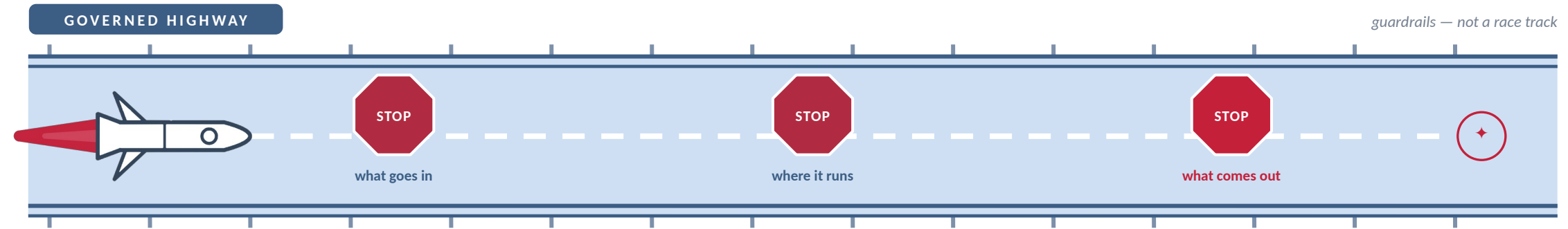


You can buy the model. You can't buy two years of your team having used it. So the expensive choice isn't starting — it's waiting. **That's why I'm here.**



Governance is what lets us go fast

The part you asked me to focus on — and the part that actually lets us say yes.



1 What goes in

Know your data class before you paste. Approved data, approved tools — all inside our governed environment. If it couldn't leave the boundary in an email, it doesn't leave in a prompt.

2 Where it runs

Governed tools only. Sensitive information never touches open, consumer models — the boundary is what keeps our clearances and contracts intact.

3 What comes out

A person owns every output. Verify before it ships — the work carries your name, not the AI's. Rigor on the workflow, always.

WHEN IN DOUBT:

Ask before you paste

Default to the governed tool

Verify, then send

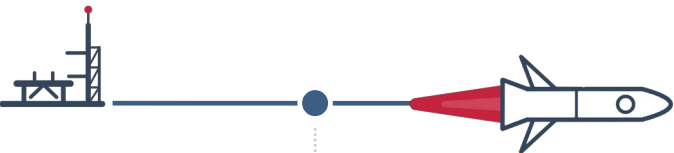
Guardrails don't slow the tow — they make the speed safe, and safe speed is what lets us commit.



Start with a pilot you can measure

Not a transformation program — one first launch, with a course you can chart.

launch pad — today



under way — week 1

the first measured win — day 30



1 Pick one routine task

WEEK 1

Low-risk, easy to double-check, done every week — a task where a win is obvious.

status emails • meeting notes • report digests

2 Set the measure up front

BEFORE LAUNCH

Before you launch, write down the one number that will prove the win:

hours saved

cycle time

error rate

3 Prove it, then share it

DAY 30

Run it weekly, verify every output, and say the number out loud at day 30.

the win funds tow #2 →

It's not only about ROI. The first pilot is how a team gets comfortable, breaks the uncertainty wall, and starts the compounding now.

■ THE QUALITY GUARDRAIL



Slop is real.

AI is a probabilistic yes-man: it predicts the next likely word, so it hands you confident, wrong, unverified output. It looks real, and it won't be flawless for years.

■ OUR JOB NOW



Hallucinations look exactly like good work:

- A source or citation that was never written
- A number that's confidently, subtly wrong
- A summary that quietly invents a detail

Without governance and your judgment, these slip through, **and “AI-assisted” becomes slop with your name on it.**

“We'll adopt once it's perfect” is the perfect-ship trap again.

There's always a cleaner version coming. Govern the flaws. Tow anyway.



Your part starts this week

Three days. Three moves. Every prompt below is printed on your worksheet, word for word — type it exactly, then swap one noun for your own work.

MONDAY

Pick it, then draft it

One routine thing you already do every week. Hand it over and read the draft — don't touch it.

■ TYPE THIS

Draft the weekly status email I send my director: what shipped, what slipped, what I need, what's next.

WEDNESDAY

Critique — don't edit

Say what's wrong in your own words and make it go again. Three rounds. No hand-polishing.

■ TYPE THIS

Too corporate and too long. Cut it by half. You buried the ask — put it in the first two lines.

FRIDAY

Meta-prompt — the one nobody tries

Make it improve your prompt, then run its version. This is the move that changes how you work.

■ TYPE THIS

Rewrite my original prompt so it gets a better first draft on the first try - then run it.

Keep the rigor. Verify, think critically, own the result. That's how we get the speed without the slop.



Five habits that do most of the work

Not tricks. These five show up in almost every prompt I write — and they're what built this deck.

1 Batch it. Number it.

One message, twelve numbered notes — not twelve messages. It answers as one coherent pass instead of twelve patches that fight each other.

2 Say what's wrong. Don't fix it.

"I don't understand at all" beat any design spec I could have written. Name the symptom; let the engine find the cure.

3 Give the constraint, not the solution.

"Make it look like Google made it." One constraint carries further than a paragraph of instructions — and leaves room to be surprised.

4 Bring your own material.

My hand-drawn ships. My role, my voice, my gecko. The output is only as original as what you feed it — AI recreates, it shouldn't invent you.

5 Make it show its work.

Sources re-checked, files hashed, every claim listed so I can verify it. Speed is worthless if you can't put your name on the result.

■ A REAL PROMPT FROM BUILDING THIS DECK • VERBATIM

`"1) i dont like the left vertical story block bar accent, make it look like google made it • 2) "interrupt me" take out that whole call to action, its obnoxious • 3) youre right its not 12-18 months, close to 6-12 months • ...and nine more"`
why it worked: twelve numbered notes in one message — blunt, batched, and not one of them told it how to fix anything.

None of this is prompt engineering. It's the same thing you already do with a sharp new hire: be clear, be honest, and check the work.