

SPAN OF CONTROL

HOW TO BUILD AN OPERATING SYSTEM FOR YOUR AGENTIC VENTURE

HANKAY

FOUNDER & SYSTEMS ARCHITECT AT SYSDOM AI



**"NO PAYROLL. NO POLITICS. NO MEETINGS.
GREAT, BUT NOW YOU HAVE ALL THE JOBS?"**

RIGHT EFFORT. WRONG LAYER.

You are shipping faster than ever. That part is real, and it is not going back.

But speed is not progress. Agents write, draft, file, and ship - and every missing feedback loop still routes back to one overloaded human, who is also the only one who would notice.

No unaided mind holds a complex system. Herbert Simon called the limit **bounded rationality**; in 1975, he and Allen Newell received the Turing Award for foundational contributions to artificial intelligence and human cognition.

"AI without Systems Intelligence is just chaos at machine speed."

The Systems Intelligence Model (SIM). Four engines where value accumulates. Four drivers that keep them coherent. A method for reading your own venture and naming the real constraint.

DIAGNOSE / DESIGN / OPERATE / REVIEW

16 CHAPTERS · 87 PAGES

Han Kay

Creator of the Systems Intelligence Model. Author of *Conscious Systems* and the ConsciOS architecture. Thirty years across Silicon Valley, Carnegie Mellon, and the World Bank. Nine ventures, seven instructive failures.

A SYSDOM HANDBOOK

Span of Control

How to build an operating system for your
agentic venture

Han Kay

Span of Control: How to build an operating system for your agentic venture

Copyright © 2026 Han Kay. All rights reserved.

Published by Sysdom AI, a trade name of Systems Intelligence LLC.

This handbook is a derived work. The canonical text is *Conscious Systems: How to Build Conscious Systems in the AI Era—From Startups to Governments*, which remains unaltered.

sysdom.ai

Contents

PART 0: THE CONTROL PROBLEM

Ships 3× Faster, Further From Done
The Unaided Mind

PART 1: DIAGNOSE

You Fix Where You Look
The Four Engines
The Four Drivers
Where to Push
The Read

PART 2: DESIGN

Sprint Zero
Problem Discovery
The Venture Architecture Map
The Sprint Zero Brief

PART 3: OPERATE

Blast Radius
Bounded Agents
The Venture Constitution

PART 4: REVIEW

The Operating Excellence Review
Evolve or Push Harder

BACK MATTER

The Twenty-Minute Audit
Glossary
Where to Go Next

1. Ships 3× Faster, Further From Done

Answer this honestly before you read on.

You ship three times faster than you did last year. Maybe five. You have agents writing code, drafting outreach, summarizing calls, filing tickets. The output is real — you can see it in your commit history and your sent folder.

So why does the venture feel further from done?

Pick the one that stings most:

- The product keeps slipping.
- Coordination eats the day — even though there's barely anyone to coordinate with.
- The direction keeps changing.
- Honestly, you can't tell.

Most builders pick the last one. That answer is the most useful, and this book is written for people willing to give it.

The output went up. The coherence went down.

Here is the uncomfortable shape of the thing. Every cycle is executed well. Nothing is checking whether it was the right cycle.

Russell Ackoff put it bluntly: *“The righter we do the wrong thing, the wronger we become.”* Aim at the wrong thing and precision stops being a virtue. Each correct execution carries you further from

where you needed to be, and the better you execute, the further you get — with more evidence of effort behind you, which makes it harder to stop.

That's the trade you made without noticing. You bought throughput. What throughput costs is the interval in which you'd have caught the mistake.

A year ago the constraint was your hands. You could only write so much code, send so many emails, read so many documents. That constraint was frustrating, but it was also a governor. It kept the amount of work in flight small enough that one person could hold it.

The constraint moved. It is now your attention, and attention did not scale when your throughput did.

This is not a motivational problem. You are not insufficiently disciplined. You are running into something older and harder than discipline, and the next chapter names it properly.

For now, notice the mechanism, because it is structural and it repeats:

“AI will cut our support costs.” Tickets doubled. The answers got faster, so more people asked. The escalation loop never changed, so the hard tickets piled up behind the easy ones that no longer needed a human at all.

“Agents will speed up engineering.” They did. Then the review queue became the bottleneck, because the one thing that could not be delegated was the judgment about whether the work was correct.

“Static identity files will keep the agent coherent.” They drifted. Files do not update themselves. Three weeks later the agent was operating from a description of a venture that no longer existed.

In all three cases the AI worked exactly as advertised. The system around it did not.

This is not a small minority

In 2025, MIT’s NANDA initiative studied more than three hundred enterprise AI initiatives and found that roughly ninety-five percent never reached measurable value. The exact figure is debated. The pattern is not.

The interesting part is the twist. The same employees at the same companies were using ChatGPT every day, on their own, and it was working. Shadow AI delivered where the sanctioned program did not.

Which rules out the obvious explanations. The models were fine. The people were fine. What the official initiatives lacked was memory, feedback, workflow fit, and governance — the system the tool runs inside.

They did not have an AI problem. They had a systems problem.

You are one person, not an enterprise. That does not exempt you. It concentrates the exposure. In a company, the missing structure is at least distributed across departments that partially compensate for each other. In a one-person venture with a fleet of agents, every missing loop routes back to the same overloaded human — and that human is also the one who would have to notice.

The captain and the designer

Picture a fishing boat and a superyacht.

Who controls how fast each one goes?

Not the captain. The designer.

Swap the captain and performance moves by maybe five percent. The person who designed the hull decided whether the ceiling was ten knots or fifty. Everything the captain does happens underneath that ceiling.

AI is a brilliant captain. On a well-designed system it is transformative. On a badly designed one it reaches the wall faster and hits it harder.

Almost everyone reading this is investing in the captain. Better models, better prompts, better tools, more agents. Very few are touching the hull.

What this book is for

Systems Intelligence is the layer between what you want and what executes.

You already have the top layer. You know why you're building, what "better" means to you, what you'd refuse to do for money. Call it consciousness, call it intent — you have it, and it's usually the reason you started.

You have the bottom layer too, in abundance. Models, agents, tools, execution at machine speed. That layer has never been cheaper or more available.

The middle layer decides whether the bottom amplifies the top or amplifies the mess. Most builders skip it entirely, going straight from intent to tooling, and then wonder why more tooling keeps producing more noise.

The middle layer is the system: the engines where value accumulates, the drivers that coordinate them, the feedback loops that tell you whether any of it is working. It is learnable. It is not long. And it is the difference between an agent fleet that compounds and one that just runs.

You can learn to be the designer, not only the captain.

What you'll be able to do

By the end of this handbook you will be able to sit down with your own venture and, in about twenty minutes, name:

- which of your four engines is the constraint,
- what input it is missing,
- which feedback loop is broken,
- which coordination driver is absent,
- and one structural move that would change the most.

Then you'll be able to design the operating boundary around your agents so that what they produce lands as evidence you can act on, rather than volume you have to triage.

That's it. It's a small book because the model is small. The model is small on purpose, and the next chapter explains why that isn't a compromise — it's the whole point.

Take this with you: the bottleneck is almost never where the noise is.

2. The Unaided Mind

Three sentences we heard over and over while researching what actually stops founders. They're worth reading slowly, because they're the same sentence in three costumes.

"I'm working harder but the structural problem won't move."

"I'm doing product, growth, ops, and finance all at once."

"There are too many frameworks and tactics. I have pieces but no operating system."

The first is exhaustion. The second is fragmentation. The third is the diagnosis, and the person who said it was closer to the answer than they knew.

Simon's ceiling

In 1978 Herbert Simon won the Nobel Prize in Economics for a finding that sounds obvious once you hear it and is almost universally ignored in practice.

Human rationality is bounded. Not by intelligence, effort, or character — by architecture. There is a limit to how much of a complex situation any mind can hold, and past that limit we do not fail gracefully. We simplify without noticing, and then act with full confidence on the simplification.

Simon's point was that this is not a defect to be overcome. It is a permanent property of being a decision-maker inside a system larger than your own working memory. The correct response is not to strain harder against the ceiling. It is to build something that holds the parts you cannot.

He established this about mid-century managers running mid-century companies.

You are running a venture where the volume of decisions, artifacts, and open threads went up by an order of magnitude in about eighteen months, most of it generated by systems that never sleep. The ceiling did not move. Everything underneath it did.

So when you feel like you're working harder and the structural problem won't move, you are not underperforming. You are operating at the design limit of unaided cognition, in conditions that have multiplied what needs to be held.

That's not an excuse. It's a specification. It tells you what kind of solution is possible — and rules out the one you're probably attempting.

Complicated, or complex?

There's a reason ventures exceed the ceiling so quickly, and it isn't size.

Think about a road network. Roads, signs, traffic lights, lane markings, cars. That's complicated — many parts, many interactions. But it's tractable. Given enough time you could map every element and predict how the whole thing behaves.

Now put drivers in the cars.

Everything changes. Drivers perceive. They decide. They adapt. They learn a route and then abandon it. They hesitate. They imitate the car in front. They break rules when the rules stop making sense to them. And crucially, they respond to each other — so the system’s behavior emerges from the interaction, not from the parts.

That’s a Complex Adaptive System. Structure plus adaptation. And it can’t be predicted by inventory, because the interesting behavior isn’t in any component.

Your venture has always been one of these. Your customers adapt. Your competitors adapt. Your collaborators interpret, resist, and improvise. You yourself change your mind about what you’re building based on what you learned last Tuesday.

Here’s the part that’s new.

Your agents are drivers too.

They perceive context. They decide within a run. They adapt to what they find. They respond to each other’s outputs, and increasingly they trigger each other without a human in between. When you added an agent fleet, you did not add tools to the road. You added drivers.

Most founders are still modeling their stack as infrastructure — as roads and signs, complicated but knowable. It isn’t. It’s traffic. And the thing about traffic is that a jam can form and travel backward down a highway with no accident, no obstruction, and no cause you can point to. Just interactions.

Every decision is an omission

Now the sharpest consequence, and the reason this chapter exists before any of the method.

Every decision is simultaneously a commission and an omission. You commit attention, money, and action to one thing. You omit everything else — including the relationships you never saw.

Bounded rationality guarantees you are always omitting something. That's not the problem. The problem is that omissions are invisible by construction. A bad commission announces itself: you shipped the feature, nobody used it, you know. A bad omission is silent. You simply never learn what the unattended relationship was doing while you were elsewhere.

This is why “I'm working harder but the structural problem won't move” is such an accurate description of a specific failure. The work is real. It's landing on the visible thing. The constraint is sitting in a relationship that never entered the frame — so effort and outcome have decoupled, and no amount of additional effort will recouple them.

The cost of what you didn't see is never zero. The cost is whatever you did instead.

The two answers that don't work

Facing this, most founders reach for one of two responses.

Work harder. More hours against a cognitive ceiling. This produces exhaustion and — worse — confidence, because the effort feels like progress and progress is hard to distinguish from

motion when you've lost the reference frame.

Collect more frameworks. OKRs, a growth playbook, a prioritization matrix, three newsletters, someone's thread about agent orchestration. Each is locally sensible. Together they're what that founder meant by *pieces but no operating system* — fragments that don't compose, each optimizing a different slice with no account of how the slices interact.

Both fail for the same reason. Neither addresses the actual constraint, which is that you have no shared representation of the whole system. Effort applies force. Frameworks apply local optimizations. Neither gives you a way to see.

What a model actually does

A model does not make you smarter, and it doesn't lift Simon's ceiling. Nothing does.

What it does is decide, in advance, what is worth attending to — so your bounded attention gets allocated rather than captured. Instead of your focus going wherever the noise is loudest, it goes where the model says value should be flowing, and you check whether it is.

And it does one more thing that matters more than it sounds: **it makes the missing structure discussable.**

Without a model, an absence has no name. You can't point at a gap, can't put it in a task, can't ask an agent to watch it, can't tell a collaborator it's there. It just quietly doesn't exist until it's expensive. With a model, the absence has a location — *this engine has no feedback loop, this driver isn't coordinating anything* — and the moment it has a location it becomes work you can assign.

For a solo founder that's close to everything. You don't have a colleague who'll notice. The model is the colleague who notices.

Small on purpose

The model in this book has four engines and four drivers. That's it.

If you were expecting something more elaborate, understand that the compactness is the point. A model as complex as the system it describes would exceed the same ceiling and be useless. Its whole value is that it compresses a Complex Adaptive System into something one tired person can hold in their head on a bad Tuesday and still reason with.

Constraint is the feature. Everything in this book was cut until it fit.

One last warning

Brilliance is not protection.

Einstein ran an institute into the ground. A genius at physics, defeated by a complex adaptive system full of people. The same pattern took Kodak and Nokia, and probably a company you've worked at: extraordinary intelligence applied at the wrong layer.

You cannot out-think a system you cannot see. You have to be able to see it first.

That's what Part 1 is for.

Take this with you: you don't need more effort. You need a model — and it has to be small enough to actually use.

3. You Fix Where You Look

There's a reason smart people keep solving the wrong problem, and it isn't judgment. It's visibility.

Every system can be observed at four depths. Only the top one is easy to see, and it's the one with almost no leverage.

Events. *"The launch slipped again."* Individual occurrences. Visible, urgent, dated.

Patterns. *"Every launch slips."* Events over time. Visible if you look, and most people eventually do.

Structures. *"Feedback from the last launch arrives after the next one is scoped."* The arrangement of parts, flows, incentives, and delays that produces the pattern.

Mental models. *"Shipping speed is what proves we're serious."* The beliefs that made the structure seem correct when it was built.

Events get the meetings. Patterns get the dashboards. Structures and mental models decide what keeps happening.

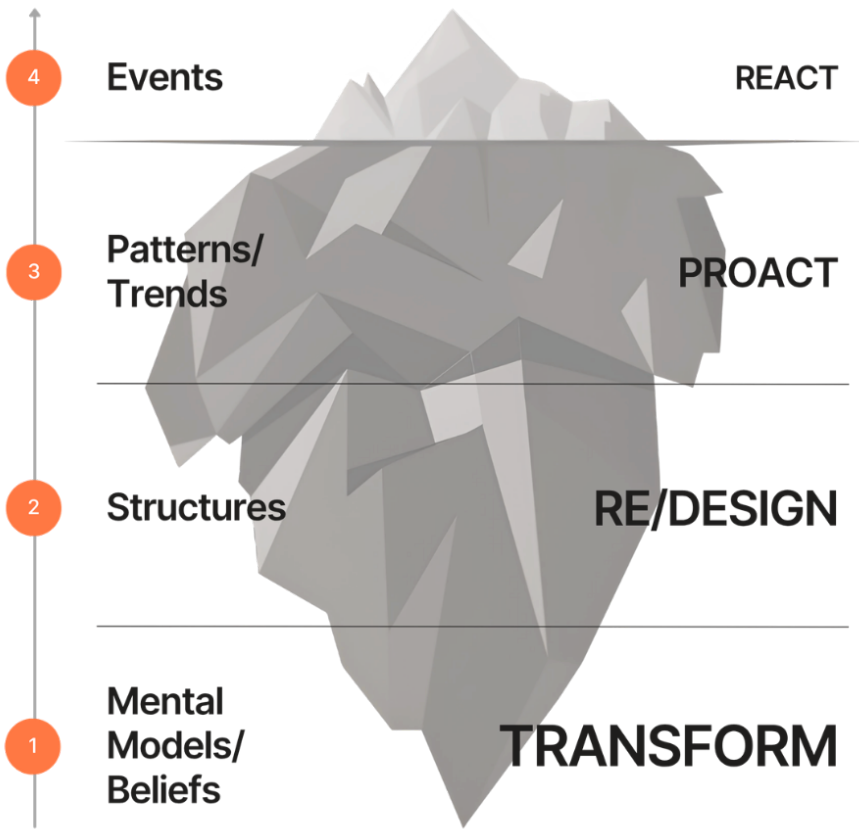


Figure 3.1: The Iceberg Model — four levels of system intervention, from Events at the surface down to Mental Models at the base. The deeper you intervene, the more lasting the change. Reacting to events is the weakest leverage; changing mental models is where transformation begins.

Why we stay at the top

Not stupidity. Incentives.

Events are visible, bounded, and satisfying to fix. You can close one. You get the small hit of completion, and the fix is legible to anyone watching — including you, at the end of a long day, wanting evidence that the day meant something.

Structures are invisible, unbounded, and unsatisfying. Fixing one produces no immediate result, because its effect shows up as *an event that doesn't happen next month*. You will never feel the win. You'll just stop losing, quietly, and probably attribute it to something else.

So attention drifts upward on its own. It requires deliberate effort to look down.

What the agent fleet does to this

Here's the part specific to how you're working now, and it's worth sitting with.

Agents are exceptionally good at the event layer. That's the layer that's legible to them — a bounded task, a clear input, a checkable output. Fix this bug. Draft this reply. Summarize this call. Clear this queue.

They're much weaker at the structure layer, because seeing structure requires holding the whole system, having a stake in its future, and noticing what *isn't* there. Absence is not a thing you can put in a prompt.

So adding agents multiplies your throughput at the shallowest layer while leaving the deep layers exactly as they were.

This is the mechanism behind Chapter 1. You didn't get worse. Your event-layer capacity went up ten-fold and your structure-layer capacity stayed at one tired human. The ratio broke, and the ratio is what determines whether motion becomes progress.

It also explains the specific texture of the exhaustion — the feeling of a day that was undeniably productive and somehow moved nothing. That's a day spent entirely in the top two inches

of the iceberg.

Worked example, in your world:

- **Event:** the agent produced unusable output on Tuesday.
- **Pattern:** it produces unusable output whenever the task is underspecified.
- **Structure:** there's no step where ambiguity gets resolved before a run starts, and no reviewer between the output and the repo.
- **Mental model:** delegation means handing off. (It doesn't. It means bounding.)

You can fix Tuesday forever. Only the third line stops the pattern, and only the fourth line stops you from rebuilding the same structure somewhere else in six weeks.

The habit

You don't need more theory to start using this. You need one habit:

When something goes wrong, ask which layer you're standing on before you decide what to do.

Most of the time you'll notice you're at the event layer, about to fix an instance. Sometimes that's right — the building is on fire, put it out. But name it as a symptom fix so you don't mistake it for a solution and stop looking.

Do this now

Take the loudest current problem in your venture and write it four times:

1. As the event that happened.
2. As the pattern it belongs to.
3. As the structure producing the pattern.
4. As the belief that made the structure look reasonable.

If you can't write line 3, that's the real finding — and Chapter 4 gives you the vocabulary for it.

Take this with you: you fix where you look. Leverage lives where you don't.

4. The Four Engines

Every venture, at every scale, accumulates value in exactly four places.

Product / Service — how value gets created and delivered.

Customer — how people discover you, trust you, and stay. **Cash** — how resources flow in, out, and toward what's working. **Skills** — how capability compounds.

That's the first half of the model. It's short enough to memorize in a minute and deep enough to read a three-trillion-dollar company with.

Engines are not departments

This is the most common misreading, so let's kill it early.

An engine is not a team, a function, or a box on an org chart. It's a **living workflow** — a continuous process that takes inputs, transforms them, produces outputs, and receives feedback.

Which is precisely why this model works when you have no departments at all.

You're one person, possibly two, with a fleet of agents. You have no marketing department. But you absolutely have a Customer Engine — it's just that nobody owns it, nothing measures it, and it runs on whatever attention is left over on a Thursday.

The engine exists whether or not anyone is minding it. Value either accumulates there or it leaks. That process doesn't wait for you to staff it. And an unowned engine is not a dormant engine — it's an engine running badly, invisibly, at full time.

The circuit

The four engines aren't a list. They're a loop.

Product → Customer → Cash → Skills → back to Product.

Product creates value. Customer turns value into relationship and revenue. Cash renews resources and buys options. Skills build the capability that makes the next Product better.

Value compounds around this circuit. Break any single link and the whole cycle stalls, regardless of how strong the other three are.

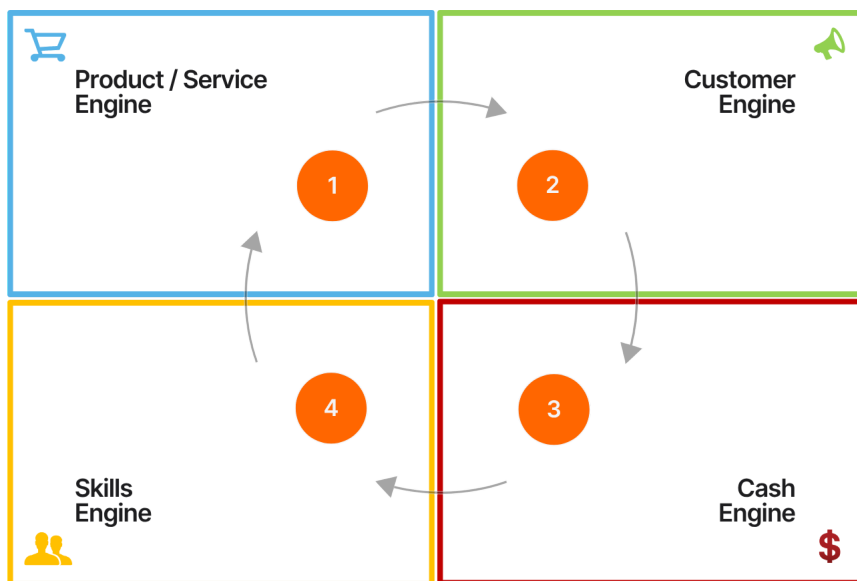


Figure 4.1: The four Jump Engines in motion — value accumulating across Product/Service, Customer, Cash, and Skills, cycling 1 → 2 → 3 → 4 as each engine feeds the next.

Two ventures make the point cleanly.

Shopify ran all four feeding each other — merchants succeeded, revenue followed, capital funded capability, capability improved the platform, which made merchants more successful. A reinforcing loop.

Theranos ran Customer and Cash at extraordinary power with no functioning Product engine underneath. Enormous trust, enormous capital, nothing completing the circuit.

From the outside, both looked like rockets. The model separates them in minutes, and the tell isn't which engine is strongest — it's whether value *completes the loop*.

The ceiling

Two rules govern how engines constrain each other.

The weakest engine sets the ceiling. Not the strongest. You cannot out-execute your weakest link, and the amount of energy poured into an already-healthy engine is one of the most reliable ways founders waste a year.

A disproportionately strong engine also distorts the system.

This surprises people. An engine that races far ahead of the others pulls the whole thing out of shape — a Customer Engine generating demand a Product Engine can't serve produces churn and reputational damage, which is worse than having generated no demand.

Sustainable growth is synchronized growth.

The anatomy — learn it once

Here's the leverage in the model. Every engine, in every venture, is built from the same seven parts. Learn the anatomy once and you can read any engine anywhere.

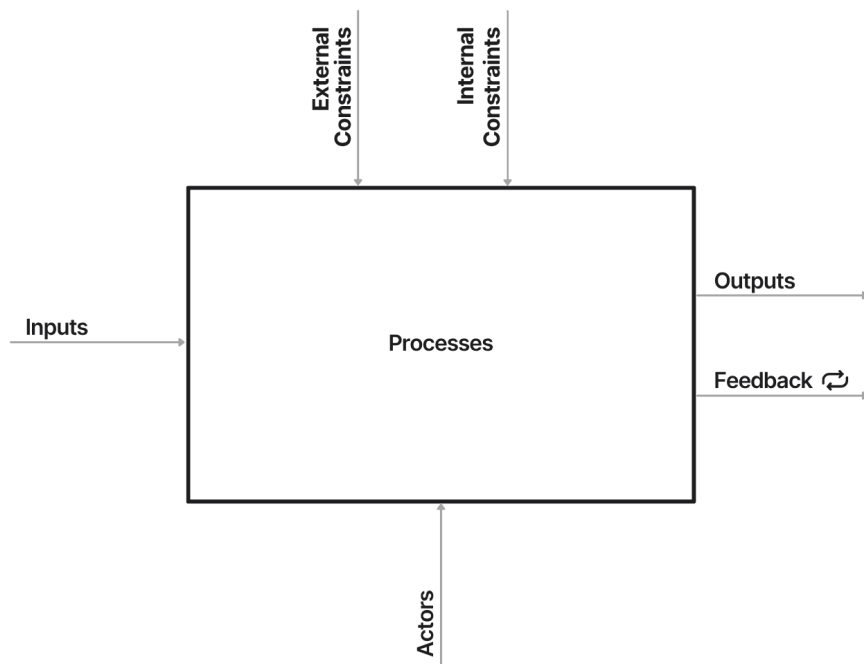


Figure 4.2: The Universal Engine Structure — every engine is built from the same seven components: external and internal constraints (top), inputs (left), processes (center), outputs and feedback (right), and actors, human and agent (bottom). The four engines differ in what flows through this structure, not in the structure itself.

Inputs — what flows in. Cash, supplies, requirements, information, attention.

Processes — the transformation. Planning, design, engineering, testing, support.

Outputs — what flows out. Value, assets, revenue, capability.

Feedback — what returns. And note this carefully: *feedback mostly returns as requirements*. A customer complaint isn't just information; it's a specification for the next cycle. If feedback isn't reaching the process that would act on it, the engine is running open-loop — producing confidently, learning nothing.

Actors — who runs it. **Human and agent.** This is where your fleet lives inside the model. Agents aren't a separate layer bolted onto the venture; they're actors inside specific engines, and they should be readable as such.

Constraints — what bounds it, from three directions: *external* (market, regulation, physics), *internal* (capacity, capability, architecture), and *governance* (your own rules, decision rights, and approval boundaries).

That last category is the one founders forget. Constraints are not only things done to you. The most binding constraint in a small venture is frequently a rule you made yourself, that worked, and that you are now defending without re-examining.

Reading your own

Diagnosing yourself is harder than diagnosing a company, because you're inside it and you have a story about who you are.

I'll go first.

I ran this on my own thirty years. Product engine: strong — I build things, always have. Skills: strong, and across an unusual range of fields. Cash: usually enough.

Customer: weak. Consistently, structurally, for three decades.

I'm an introvert. I disliked social media. I could build almost anything and could not get it in front of the people it was for. Three master's degrees, nine ventures, seven failures — and the same engine starving every time, wearing a different costume.

That's not a confession, it's a demonstration. The read was available the whole time. What I lacked wasn't effort or intelligence. It was a model that would have made the missing engine *visible and nameable* in year three instead of year thirty.

What I was actually doing, for most of thirty years, was making good work and quietly hoping someone would notice it. That's not a strategy. It's a wish standing where an engine should be — and a wish produces no feedback, so it never corrects.

Once the gap had a name, it became something I could build rather than something I waited for. So I built the missing engine.

Your read will be less flattering than your self-image. That's how you know it's working.

Take this with you: the engine exists whether or not anyone is minding it.

5. The Four Drivers

Four healthy engines can still produce an incoherent venture.

Engines accumulate value. Nothing in that description makes them work *together*. Coordination is a separate function, and in the model it has four parts.

Innovation — senses. What's changing, what's newly possible, what assumption just expired.

Governance — decides and allocates. Who chooses, what the boundary is, what requires a human yes, where resources go.

Interaction — moves information. Does what one engine learns actually reach the engine that needs it, in time to matter.

Culture — preserves shared mental models. What “good” means here, and whether everyone acting in this system agrees.

Four engines, four drivers. Rotate every driver across every engine and you get sixteen questions. That grid is the whole diagnostic surface of the model.

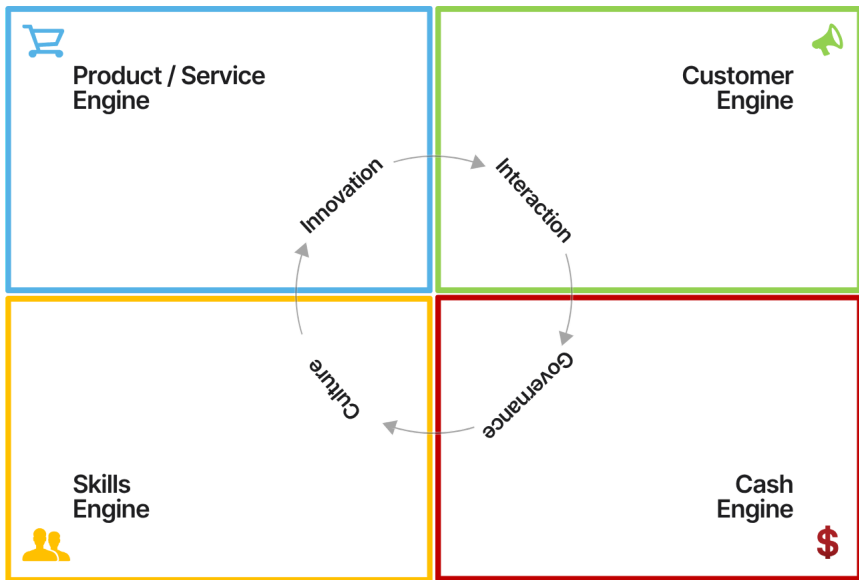


Figure 5.1: The four Jump Drivers — Innovation, Governance, Interaction, and Culture forming a coordination loop around the engines.

What actually disappears in a one-person venture

Here's the finding that matters most for how you're working, and it's not intuitive.

When you were doing everything yourself, all four drivers existed — **implicitly, inside your head**. You sensed change because you were touching every part. You governed because every decision was yours by default. Information moved because it never had to leave your skull. Culture was coherent because there was one person in it.

You didn't have coordination mechanisms. You *were* the coordination mechanism.

Then you added agents. And implicit coordination didn't degrade gracefully — it simply stopped existing, without any event marking the moment.

Governance goes first. An agent has no default sense of what requires your approval. Absent an explicit boundary, it either escalates everything (and you become a bottleneck) or escalates nothing (and you lose control). Both feel like the agent's fault. Neither is.

Interaction goes second and hurts longest. A run produces something genuinely useful. It lands in a log, a branch, a folder. It never reaches the place where the relevant decision gets made. The work was done, correctly, and had no effect. If you've felt the specific frustration of knowing an answer exists somewhere in your own systems and not being able to act on it, that's a broken Interaction driver, not a memory problem.

Culture drifts quietly. This is the static-identity-files failure from Chapter 1. You wrote down what the venture is and what "good" means, the agents used it, and then the venture changed and the file didn't. Three weeks later your fleet is operating coherently — with a description of a company that no longer exists. Shared mental models don't maintain themselves. Written-down ones especially don't.

Innovation is the one founders usually retain, because sensing change is the part of the job that feels like the job. But check anyway: when you're heads-down orchestrating output, is anything watching whether the situation still holds?

A worked read

Apple is useful here because its drivers are unusually strong, and it's still under strain — which shows that driver problems are not weakness problems.

Innovation: disciplined focus is a genuine strength. Under AI it sits in tension with open-ended experimentation, which requires tolerating a lot of public wrongness.

Governance: integrated decision-making produces extraordinary coherence, and concentrates decision pathways.

Interaction: compressing complexity for the user is a defining skill. It coexists with tightly controlled internal feedback, and AI needs feedback to move fast and freely.

Culture: finished-product excellence is the identity. An AI relationship is never finished.

Every strength carries a tension when the physics change. That's not a criticism of Apple; it's the pattern. The architecture that produced historic success meets a new constraint, and pushing the old strength harder doesn't remove it.

A systems hypothesis from public evidence, not a claim of insider knowledge.

Making them explicit

For each engine, one line per driver. Twelve words is enough.

- **Innovation:** what is sensing change here?
- **Governance:** who decides, and what needs a human yes?

- **Interaction:** where does what’s learned here need to arrive?
- **Culture:** what does “good” mean here, and who agrees?

Where you can’t answer, you’ve found an absent driver — which, per Chapter 2, is now a thing with a location, and therefore work you can assign.

You’ll find blanks. Everyone does. The blanks are the point.

Take this with you: you used to be the coordination. Once you added agents, that stopped being true — and nothing announced it.

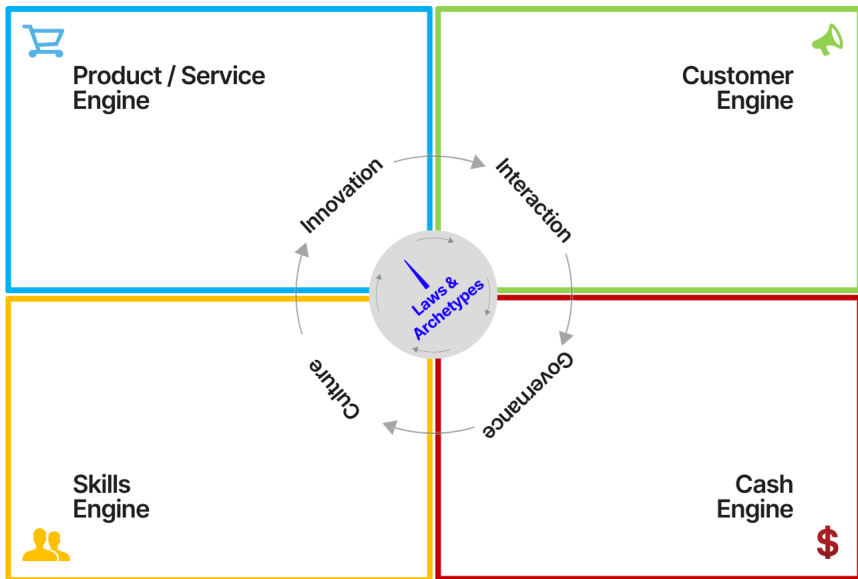


Figure 5.2: Overview of the Systems Intelligence Model — four Jump Engines, four Jump Drivers, and the Laws and Archetypes dial at the center, where the universal patterns apply across every engine and every driver.

6. Where to Push

You can now see the depth of a problem and name the parts of the system it lives in. One thing is missing before you can act, and it is the single most useful idea in this book.

Not all interventions are equal. They are not even close.

Two founders can spend the same month, with the same intelligence and the same effort, on the same venture — and one changes the trajectory while the other changes nothing. The difference is almost never the quality of the work. It's where they applied it.

Donella Meadows spent a career on this question and produced a ranked answer: places to intervene in a system, ordered from weakest to strongest. Her list has twelve entries. Below is a working version organized into five tiers, translated into a one-person venture running on agents. The full twelve-point hierarchy is in Appendix D of *Conscious Systems*; this is the working version.

Read it from the bottom up. That's the direction of increasing power, and the direction of decreasing obviousness.

Tier 5 — Numbers

The weakest place to intervene, and where roughly everyone starts.

Budgets. Prices. Rate limits. Model choice. Hours worked. How many agents you run. The temperature setting.

Changing a number changes the magnitude of what the system already does. It doesn't change what the system does.

This is the tier that produces the sentence from the problem library: *“I’m working harder but the structural problem won’t move.”* Working harder is a Tier 5 intervention. So is switching models, which is why the upgrade that was going to fix everything didn’t.

Numbers matter. They’re just not leverage. Adjust them when a system is already well-designed and slightly mistuned.

Tier 4 — Stocks, flows, and delays

Slightly stronger. Mostly about time.

How long between doing something and finding out whether it worked. How much buffer sits between you and a problem. How work queues up and drains.

Delay is the underrated one, and it’s where agent-heavy ventures quietly break. A feedback loop with a long delay behaves like no feedback loop at all — worse, actually, because you keep acting confidently in the gap and your corrections arrive out of phase with the thing they were correcting.

Concretely: if your agents produce output on Monday and you review it on Friday, you have a four-day delay in a system generating a week’s worth of decisions per day. You are steering a fast vehicle with slow hands. Shortening that delay is often worth more than improving the output itself.

Tier 3 — Feedback loops

Now we’re getting somewhere.

Two kinds, and you need both consciously.

Balancing loops keep things stable and self-correct. A review step. A budget cap. A weekly check that catches drift. When a system runs away from you, it's usually because a balancing loop is missing or too weak to matter.

Reinforcing loops compound. Value creates trust creates revenue creates capability creates better value. Your entire engine circuit from Chapter 4 is a reinforcing loop, and the reason a starving engine matters so much is that it breaks the compounding, not just that engine.

Adding one real balancing loop to an agent fleet — a place where output has to prove itself before it counts — usually beats a month of prompt engineering. And a working reinforcing loop is the only thing that produces growth you don't have to push.

Tier 2 — Information flows and rules

High leverage. Frequently free.

Information flows — who knows what, and when. Meadows' favorite example: adding a visible electricity meter to a house cuts consumption, with no change to price, technology, or rules. Nothing was added except the information reaching the person who acts.

This is your Interaction driver, and it's where most solo-founder waste hides. The knowledge usually already exists somewhere in your systems. It just doesn't reach the point of decision in time. Routing existing information to the place it changes a choice is often the cheapest real improvement available to you.

Rules — incentives, constraints, permissions, boundaries. What's allowed, what requires approval, what's forbidden.

For an agent fleet, rules *are* the system. An agent has no instinct for your boundaries. Whatever rules you set are the complete governance of that actor, and unset rules aren't defaults — they're holes. Part 3 is essentially one long Tier 2 intervention.

Tier 1 — Goals and paradigms

The strongest, and the hardest to see, because you're inside them.

The goal of the system. Not the goal you'd state on a landing page — the one the system actually optimizes for, which you can read off its behavior. A venture whose real goal is “ship visible things weekly” will produce a very different structure than one whose real goal is “learn what customers will pay for,” even if both founders describe themselves the same way.

Changing the goal changes everything below it automatically. This is why it's so powerful and why it's so rarely attempted.

The paradigm. The belief that made the goal seem obvious in the first place. *Speed is what proves we're serious. Real founders don't need process. If I build something good enough it will find its audience.*

That last one was mine, for most of thirty years. It's a paradigm, not a strategy, and it generated a system with no Customer Engine in it. No amount of Tier 5 effort was ever going to fix that, because the structure was working exactly as the belief specified.

You can't change a paradigm by deciding to. But you can notice one, and noticing is most of it — because a belief you can see is a belief you can test.

Meadows' warning

She was emphatic about this and it's worth repeating.

Leverage points are counterintuitive, and people routinely push them in the wrong direction.

Everyone senses where the leverage is. That intuition is often right about *location* and backwards about *direction* — so people find the high-leverage point and then push it the wrong way, hard, with conviction. Adding rules where the problem is too many rules. Adding information where the problem is noise. Speeding up a loop that needed damping.

The defense is modest: when you touch a high-leverage point, change one thing, and watch what actually happens before you touch it again. Structural interventions have delays. You will be tempted to conclude it didn't work before it's had time to work.

Using this

You don't need to memorize twelve points. Keep one question:

What tier is the move I'm about to make?

If it's Tier 5 and the problem is recurring, you already know how that ends. Look for the Tier 2 version — the rule or the information flow that would make the recurrence stop.

Most founders never leave Tier 5. Not because they can't, but because Tier 5 moves are visible, fast, and feel like work. Tier 2 moves are invisible, slow, and feel like avoidance right up until the day the problem stops happening and you can't quite say why.

Take this with you: everybody finds the leverage point. Most people push it the wrong way.

7. The Read

You now have the pieces: four engines where value accumulates, four drivers that coordinate them, the iceberg that tells you which depth you're looking at, and the leverage hierarchy that tells you where a move is worth making.

This chapter puts them together into a single procedure you can run on any venture, including yours, in about twenty minutes.

It has five moves. Do them in order. The order matters more than the polish.

Surface → Engine → Driver → Pattern → Leverage

Get a blank page. Not a document, not a canvas tool — a page. This should be fast and slightly rough. You are not producing an artifact yet. You are producing a diagnosis.

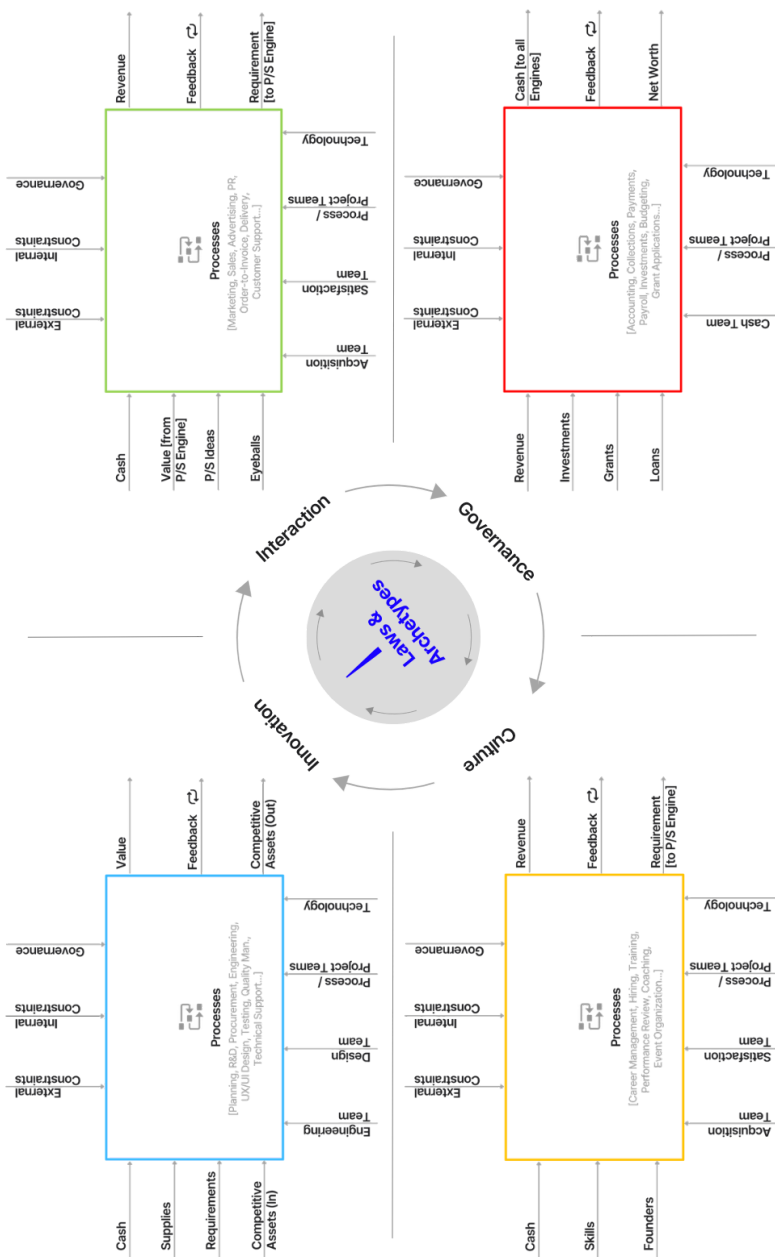


Figure 7.1: The complete Systems Intelligence Model — all four Jump Engines, all four Jump Drivers, and the Laws and Archetypes dial at the center, operating as one integrated venture architecture. This is the whole model on one page; the Read is how you use it.

Move 1 • Surface

What looks successful, and what looks broken?

Write the symptoms exactly as you'd complain about them to a friend. Not the professional version.

“Leads reply and then I lose them.” “The prototype keeps growing and I can’t tell when it’s done.” “I’m doing the same manual ops work I was doing in March.” “Runway is getting tight and I keep changing the price.”

Five minutes, maximum. Resist the urge to diagnose while you write. You’re collecting events, and events are the shallowest layer of the iceberg — but they’re also the only layer you can see directly, so they’re where the read has to start.

Move 2 • Engine

Where is value accumulating, and where is it leaking?

Take your four engines and ask one question of each. Answer in a sentence.

Product / Service — Is value actually reaching people, or is it stuck in process?

Customer — Are the right people finding you, trusting you, and staying? Or are you shouting into the void?

Cash — Do resources flow toward what’s working, or away from it?

Skills — Is capability growing faster than the problem is growing?

Now circle the weakest one.

Two warnings before you do.

The weakest engine sets the ceiling, but a disproportionately strong engine also pulls the system out of shape. Theranos ran Customer and Cash at full power with no Product engine underneath. From the outside it looked like a rocket. Shopify had all four feeding each other. Also looked like a rocket. The model tells them apart in minutes, and the tell is not which engine is strongest — it's whether value completes the circuit.

The bottleneck is where value stops compounding, not where the noise appears. These are usually different places. Noise appears where the pain is felt; the constraint usually sits one step upstream of that. If your loudest symptom is in Customer, check whether Product is actually delivering something worth staying for.

Sustainable growth is synchronized growth.

Move 3 · Driver

What is coordinating this — or failing to?

An engine can have every part in place and still fail, because nothing is holding it coherent with the others. Rotate each driver across the weak engine you circled:

Innovation — is anything sensing what's changing here, or are we running last quarter's assumptions?

Governance — is it clear who decides, what the boundary is, and what needs a human yes?

Interaction — does information actually move between this engine and the others, or does it die in someone's inbox?

Culture — do the people and agents working here share a mental model of what “good” means?

Name the one that’s most absent. For solo founders running agents, Governance and Interaction are absent far more often than the other two — not because you’re careless, but because when you were doing everything yourself, both were implicit. You were the coordination. The moment agents entered, the implicit coordination stopped existing and nothing replaced it.

Move 4 · Pattern

Why does this keep happening?

A one-time failure is an event. A recurring one is structural, and structures have names.

Ask: has this exact problem shown up before, in a different costume?

If yes, you’re looking at an archetype, and the most common one by far is **Limits to Growth**. The architecture that produced your success meets a new constraint when the physics change. The instinct is to push the old strength harder. It doesn’t work, because the old strength isn’t what’s binding.

Constraints are not only external. They live in goals, rules, capabilities, governance, and feedback loops — which means they’re frequently things you built on purpose, that worked, and that you’re now defending.

Move 5 · Leverage

Where would one structural move change the most?

Not five moves. One.

Write it as a sentence with a verb in it. Not *“improve onboarding”* — that’s a wish. Something closer to *“every lead reply gets a dated next action before I close the tab”* or *“nothing ships from an agent without a named human reviewer and a run receipt.”*

The test for a good leverage move: it changes the structure, not the symptom. If you can accomplish it by trying harder this week, it’s a symptom fix. If it changes what will keep happening after you stop paying attention, it’s structural.

Then check its tier. If what you wrote is a number — more hours, more agents, a bigger budget, a better model — go back and look for the Tier 2 version of the same intent. Ask what rule, or what piece of information arriving somewhere new, would make the problem stop recurring rather than stop today.

Before you point AI at it

One more question, and it prevents most AI waste:

Will AI amplify a healthy engine, repair a weak engine, or accelerate a broken pattern?

The first is where AI earns its keep. The second is possible but slow and usually needs the structure fixed first. The third is where most of the ninety-five percent went.

Ask it before you build the agent, not after.

What you should be holding now

Four lines:

- **Engine** — which one is limiting the rest
- **Input** — what it's missing or receiving distorted
- **Loop** — which feedback path is broken
- **Driver** — what's failing to coordinate it

Plus one structural move, and a note on which second engine to watch — because a real read usually points in more than one direction, and the secondary signal often turns out to matter more within a month.

That's a diagnosis. Not a plan yet. Part 2 turns it into a Venture Architecture Map and a Sprint Zero Brief, which is where it becomes something you can operate from.

Do this Monday

Twenty minutes. Four steps.

Map — draw the four engines. **Mark** — circle the weak one. **Name** — find the driver failing to coordinate it. **Move** — write one structural intervention.

Change the system, not just the symptom.

Run it live. The Systems Bottleneck Scanner does this read on a written note about your venture — no signup, nothing to install. Write the situation in your own words the way you did in Move 1, and it returns the likely bottleneck, the affected engine, the

secondary engine to watch, the signal strength behind the call, and a safe next move. Treat it the way you'd treat a second opinion: useful because it disagrees with you sometimes.

Do the paper version first. The read is the skill. The scanner is the instrument.

Take this with you: the read is the skill. The scanner is the instrument.

8. Sprint Zero

A diagnosis is not a plan. You know which engine is starving, which driver is absent, and where a structural move would land. Now you have to design the thing that fixes it — before you build it.

Sprint Zero is the bounded period between having an idea and committing to it. It's the bridge, and it's the part almost everyone skips.

The argument you're about to make against this

Building is nearly free now. Why design first? Just build the thing and find out.

It's a serious objection and it deserves a serious answer.

Here it is: cheap building didn't remove the need to think. It removed the **forcing function** that used to make you think.

When a build cost three months, the expense itself did your design work for you. Nobody committed three months without asking whether the thing was worth building. The cost was doing the epistemics.

Now the build costs an afternoon. The forcing function is gone, and nothing replaced it. So you build, and it's roughly right, and you build the next thing on top of it, and six weeks later you're looking at a system whose architecture is the accumulated residue of forty afternoons in which nobody ever asked what shape this should be.

That's the founder in the problem library saying "*we've rebuilt the core three times and it's still not right.*" Each rebuild was cheap. That's exactly why there were three.

When building is expensive, cost enforces design. When building is cheap, only you can.

What Sprint Zero is not

It isn't planning. Planning assumes you know what you're doing and sequences it. Sprint Zero assumes you don't, and designs the fastest honest route to finding out.

It isn't a business plan, a spec, or a strategy deck. Those are documents about a future you've decided on. Sprint Zero is about the decision itself.

And it isn't long. Days, not weeks. It's timeboxed on purpose — an unbounded Sprint Zero becomes a place to hide from building, which is its own failure mode and a comfortable one for the kind of person who reads systems books.

What it produces

Two artifacts, and they're the next two chapters.

A Venture Architecture Map — one page showing your engines, who runs them, what constrains them, what's at risk, and what's undecided.

A Sprint Zero Brief — four short lists: what's Clear, what Needs Proof, what requires Human Review, and the First Move.

That's the whole output. If it takes more than two pages you're planning, not designing.

The four questions

Sprint Zero answers four things, in order. Everything else is elaboration.

1. What problem, in whose words? Not your category for it — the sentence the person having it would actually say. Chapter 9 is entirely about this, because it's where most ventures are lost before anything is built.

2. What's the smallest promise you could prove? Not the vision. The one claim that, if true, means this is worth continuing — and if false, means stop.

3. What has to be true, and how would you know? The assumptions underneath the promise, and the cheapest evidence that would move each one.

4. Where does judgment stay human? The decisions you will not delegate, to an agent or anyone else, no matter how convenient. Naming these now is much easier than naming them at 11pm when an agent is waiting.

Why the fourth question is new

The first three are classic. The fourth is the one this era added, and it's the one founders skip.

If you don't decide in advance where human judgment is required, you will decide it by accident — under time pressure, in the moment, in the direction of whatever unblocks you fastest. And the direction that unblocks you fastest is always *let it through*.

That's not a discipline failure. It's a design failure. You left a rule unset, and unset rules aren't neutral defaults; they're holes that whatever's most urgent will pour through.

Set the boundary while you're calm. It's a Tier 2 intervention and it costs you fifteen minutes.

Stable architecture before iteration

One more thing Sprint Zero buys you, and it's the antidote to rebuild churn.

Iteration only compounds on a stable base. If the foundation shifts every cycle, each iteration partially destroys the last one, and you get motion without accumulation — the exhausting kind of progress where version four is not obviously better than version two.

Sprint Zero's job is to make enough of the architecture stable that iteration has something to build on. Not to get it right. Just to make it *decided*, so it stops being re-litigated every week.

Decided is not the same as correct. Decided means you'll change it deliberately, with a reason, instead of drifting.

Take this with you: when building is expensive, cost enforces design. When building is cheap, only you can.

9. Problem Discovery

Two sentences from the research, and they belong together:

“I built it, nobody wanted it.”

“It’s more fun to ship code than to ask hard questions.”

The second causes the first. Everyone knows this. Almost everyone does it anyway, and the reason isn’t ignorance — it’s that shipping code produces a visible artifact today and asking hard questions produces uncertainty today. Under pressure, people choose the artifact. Every time.

Agents have made this dramatically worse, and it’s worth being precise about how.

The friction that used to save you

Building something you weren’t sure about used to be *expensive enough to hurt*. That pain was information. Somewhere around week three of building the wrong thing, the cost would surface a question you’d been avoiding, and you’d stop.

An agent will build whatever you describe, immediately, without hesitation, and it will do a decent job. There is no week three. The friction that used to force the question is gone.

So the discipline that was previously supplied by the world now has to be supplied by you. That’s the whole chapter.

Say it in their words

Here's the test, and it's harder than it sounds.

Write the problem as a sentence the person having it would actually say out loud, to a friend, in their own words.

Not *"inefficient lead management workflows."* That's your category.

Try: *"Leads reply and then I lose them."*

Not *"suboptimal unit economics at scale."*

Try: *"My first big API invoice wrecked the margin."*

Not *"lack of coherent operating framework."*

Try: *"There are too many frameworks and tactics. I have pieces but no operating system."*

Those second versions are real. They come from actual founders. And here's why the distinction matters more than it looks: **if you can only state the problem in your own category language, you have a hypothesis about a problem, not a problem.** The category is something you constructed. The sentence is something you heard.

You can tell which one you have by whether you can remember who said it.

There's a marketing consequence too, which you'll want later: lead with the voiced version. People recognize their own sentence. They do not recognize your taxonomy, and asking them to learn it before they'll admit they have the problem is a tax most won't pay.

The smallest provable promise

Now compress your solution to a single claim.

Not the vision. The vision is fine, keep it, but it isn't testable and untestable things generate no feedback. You want the one sentence that, if true, means continue — and if false, means stop or turn.

A good promise is specific enough that someone could disagree with it. *"We help founders be more effective"* is unfalsifiable, which is why it feels safe and why it teaches you nothing.

Then the question almost nobody asks:

What would have to happen for me to conclude this is wrong?

Write it down now, before you have any data. Pre-registering your own disconfirmation is the single cheapest defense against the thing that kills most ventures — not being wrong, but being unable to notice you were wrong because the criteria kept moving.

If you can't name what would falsify your promise, you're not going to learn anything from what happens next. You'll just interpret it.

What evidence is worth collecting

Rank your assumptions by two things: how load-bearing they are, and how cheaply they can be tested.

Start with the assumption that is most load-bearing and cheapest to test. Not the most interesting one. Not the one you most want to be right about — that one you'll test last, if at all, which is a bias worth knowing about yourself.

And prefer evidence with a cost to it. Someone saying “that sounds useful” costs them nothing and tells you nothing. Someone giving you an hour, an introduction, a pre-payment, or their real data has spent something. Only spent things are signal.

The trap on the other side

There's an opposite failure and I don't want to send you into it.

You can discover forever. You can interview thirty people, build a beautiful model of the problem space, and never ship — and it will feel rigorous the entire time. For a certain kind of thoughtful founder, discovery is the more comfortable hiding place.

The discipline is symmetric: discovery is bounded too. When you can state the problem in someone's words, state a falsifiable promise, and name the cheapest test — you're done. Go build the test.

Sprint Zero has a clock. That's not a limitation of the method, it's part of it.

Take this with you: if you can only say the problem in your own words, you have a hypothesis, not a problem.

10. The Venture Architecture Map

One page. Your whole venture, visible at once.

That sounds reductive and it is. Reduction is the function. Chapter 2 established that no mind holds a complex adaptive system unaided; the map is the external structure that holds the parts you can't. If it needs a second page, it's failing at its only job.

What goes on it

Five zones.

The four engines. For each, one line on what's actually happening — not what you intend. *“Product: shipping weekly, no eval pipeline.” “Customer: 40 waitlist signups, no follow-up loop.”* Mark the constraint engine and the secondary one you're watching.

Actors. Who runs each engine. Human and agent. Most engines in a solo venture will say “me” and that's exactly the finding — an engine with no owner but you is an engine running on leftovers.

Constraints. Per engine, the binding one. External, internal, or governance. Be specific: *“can't ship faster than I can review”* is a real constraint; *“limited resources”* isn't.

Risks. Three at most. What would hurt, how likely, whether you'd see it coming. That third column is the useful one — a risk you'd spot early is a different animal from one that arrives fully formed.

Open decisions. What’s genuinely undecided, with a date by which it needs deciding. Not a to-do list. Decisions only.

That’s it. Engines, actors, constraints, risks, decisions.

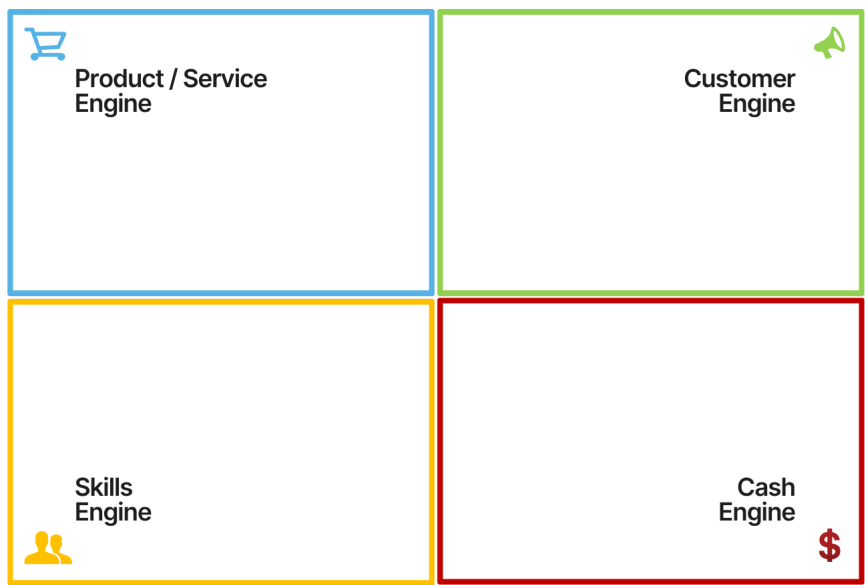


Figure 10.1: The blank SIM canvas — four engine quadrants waiting to be mapped to your venture. This is the working surface you return to.

What doesn’t go on it

Metrics dashboards. Roadmaps. Feature lists. Financial projections. Your mission statement.

Those are all real artifacts and they live elsewhere. The map has one job: let you see the shape of the system and where it’s binding. Everything you add past that costs you legibility, and legibility is the entire product.

If you find yourself wanting to add a section, ask whether it would change a decision this month. Usually the honest answer is no.

It's alive

The map is not a document you produce and file. It's a state of the system, and the system changes.

Date it. Keep the old versions. When something moves — an engine strengthens, a constraint lifts, a risk fires — update it and note what changed.

The version history turns out to be more valuable than any single version, because it's the only record you'll have of how your own understanding evolved. Six months in, comparing March's map to September's will tell you things neither one says alone: which constraints you kept mis-identifying, which risk you listed four times without addressing, which engine you've been claiming is fine since spring.

That's your feedback loop on your own judgment. Almost nobody has one.

Building it from the Read

You already did most of this work in Chapter 7. The Read gave you engine, input, loop, driver, and a leverage move. The map is that, plus the parts you didn't need for the diagnosis but do need for operating: who's running what, what's binding, what's at risk, what's undecided.

Give it thirty minutes. Rough is fine — a hand-drawn version you actually look at beats a beautiful one you make once and never open.

The test of a good map isn't completeness. It's whether you can hand it to someone and have them ask a question you'd been avoiding.

Take this with you: the map's only job is legibility. Everything you add past that costs you the thing you came for.

11. The Sprint Zero Brief

The map shows the system. The brief tells you what to do about it, and it has exactly four parts.

Clear · Needs Proof · Human Review · First Move

This shape looks almost too simple. It's load-bearing, and each section does a job the others can't.

Clear

What you actually know, with evidence behind it.

The discipline is subtraction. Most things you're about to write here aren't clear — they're assumptions that have been around long enough to feel settled. Age is not evidence.

For each line, ask: *how do I know this?* If the answer is “everyone knows” or “it's obvious” or “we decided that in April,” it belongs in the next section.

Being ruthless here is the whole point. Every false entry in Clear is an omission you've stopped looking at — and per Chapter 2, omissions don't announce themselves.

Needs Proof

What you believe but haven't established.

Each line gets a test and a cost. Not “validate demand” — *“ten founders will give me an hour this week; if fewer than three describe this problem unprompted, the framing is wrong.”*

This section should be longer than Clear. If it isn't, you're either unusually far along or not being honest, and the base rate strongly favors the second.

Rank by load-bearing $\times$ cheapness, per Chapter 9. Work top-down.

Human Review

Where your judgment stays in the loop, no matter how convenient it would be to skip.

This is the section that didn't exist in anyone's Sprint Zero five years ago, and it's now the most important one.

Write the boundary as a rule an agent could be held to. Not *"be careful with customer communication"* — that's a feeling. Something like: *"nothing goes to a named person outside the company without me reading it first."* *"No spend over \$50 without approval."* *"No irreversible change to production data, ever, by anyone but me."*

The test for whether something belongs here: **is it reversible?** Two-way doors can be delegated — if it's wrong, you walk back through. One-way doors stay human. Published content, sent messages, deleted data, spent money, changed consent, made promises. Those don't come back.

Set the rules now. Chapter 12 turns them into an operating mechanism.

First Move

One thing. With a verb, an owner, and a date.

Not a list. The single next action that produces evidence about the most load-bearing thing in Needs Proof.

If you can't compress it to one, your diagnosis isn't finished — go back to the Read. A well-formed diagnosis produces an obvious first move, and the feeling of “there are five equally important things” is almost always a sign that the constraint hasn't actually been located.

Why this shape works

Each section prevents a specific failure.

Clear stops you rebuilding what's settled. **Needs Proof** stops you defending what's merely familiar. **Human Review** stops the slow erosion of judgment under time pressure. **First Move** stops the brief becoming a plan.

And read together, Clear versus Needs Proof does something quietly valuable: it separates knowledge from belief and makes the ratio visible. Most founders discover their Clear column is much shorter than they expected. That's not a bad day. That's the instrument working.

Worked example

Solo founder, developer tools, four months in.

Clear — People sign up (60 in three weeks, organic). Product does the core thing reliably. I can build faster than I can validate.

Needs Proof — That signups reflect intent to pay (*test: offer paid tier to the 20 most active, this week — under 2 conversions means the wedge is wrong*). That the bottleneck is discovery, not retention

(test: check whether the 60 came back a second time — data already exists, 20 minutes).

Human Review — Any message to a named user. Any pricing change. Anything touching customer data.

First Move — Pull the return-visit data before writing another line of code. Thursday. Me.

Note that the first move is the *cheapest* test, not the most important one. Twenty minutes of existing data might invalidate the more expensive test entirely. Sequence by cost when the information is comparable.

Take this with you: Clear is shorter than you think. That's the instrument working, not a bad day.

12. Blast Radius

Everything so far has been about seeing. This part is about running the thing, and it starts with the mechanism that decides what reaches you.

Borrow a term from the people who run large systems for a living: the **blast radius** of an action is how much breaks if it goes wrong. Reducing blast radius is not caution. It's the thing that makes speed survivable — you can move fast precisely because most of what you're doing can be undone.

Your agents have no sense of blast radius. They will execute a reversible draft and an irreversible send with exactly the same confidence. So the boundary has to live outside them, in a mechanism.

That mechanism is the **approval path**. Three buckets. Every piece of work goes in one.

Do Now — bounded, reversible, delegatable. An agent or you can execute without asking.

Needs Human Yes — irreversible, expensive, or judgment-dependent. Stops until you decide.

Carry Forward — real but not now. Recorded so it survives, deferred so it doesn't crowd the week.

That's the entire governance system. Its simplicity is deliberate: a governance system you won't maintain is worse than none, because it produces false confidence.

The two ways this fails

Without an explicit path, you get one of two failure modes. Both feel like the agent's fault and neither is.

Everything escalates. The agent, sensibly cautious, asks about everything. You become a permanent approval bottleneck, reviewing decisions that never needed you. Throughput collapses to your reading speed and you conclude agents don't work.

Nothing escalates. The agent proceeds confidently on everything, including the one thing that needed a human. You find out later. Sometimes much later.

The oscillation between these — clamping down after an incident, loosening after the frustration builds — is where most solo founders currently live. The approval path replaces the oscillation with a rule.

The test: which way does the door swing?

Sort by **reversibility**, not by importance.

Importance is a bad sort key because everything feels important at 11pm and nothing does on a good Tuesday. Reversibility is a property of the action, and it doesn't move with your mood.

Two-way doors — you can walk back through. A draft, a branch, an experiment, an internal document, a local change. If it's wrong, you undo it and lose an hour. **Delegate freely.** Most work is here, and founders systematically underestimate how much.

One-way doors — you cannot walk back. Published content. A message sent to a named person. Deleted data. Spent money. A changed consent state. A promise made. A price announced. **These stay human. Always. No exceptions for convenience, urgency, or how good the draft looks.**

The list of one-way doors in most ventures is short — usually five to ten items. Write it once. It'll change less than you expect.

Carry Forward is not a graveyard

The third bucket is the one people implement badly, and it matters more than it looks.

Carry Forward exists so that “not now” doesn't have to mean “forget.” Without it, you face a false choice on every non-urgent item — do it now, or lose it — and both options are bad. You'll take on work that shouldn't be this week's, or you'll drop things that mattered.

Two rules keep it honest. **Everything carried forward has a reason and a date to revisit.** And **it gets reviewed weekly** — which is Chapter 15's job.

An unreviewed Carry Forward list becomes a graveyard, then becomes a source of guilt, then becomes something you avoid opening. At that point you've built an anxiety generator, not a system.

Writing rules an agent can hold

Your one-way door list is a set of rules, and rules are Tier 2 leverage — high, and free.

But they have to be written so a non-human actor can apply them. That means observable, not aspirational.

“Use good judgment on customer communication” — unusable.
Judgment is the thing the agent doesn’t have.

“Nothing goes to an address outside our domain without my approval” — usable. It’s checkable.

“Be careful with spend” — unusable.

“No single call over \$5, no session over \$50, hard stop” — usable.

Writing this way is uncomfortable, because explicit rules feel rigid and you’d rather stay flexible. But flexibility that only exists in your head isn’t governance — it’s the absence of governance with a nicer name, and it fails in exactly the direction that unblocks whatever’s most urgent.

What comes back

Anything that hits Needs Human Yes should arrive with four things: **what was attempted, what happened, what it would cost or risk, and a recommendation.**

Not a raw log. Not “please review.” A decision, framed, with the evidence attached and a default proposed.

This is the difference between an approval path that saves you attention and one that just moves the work. If reviewing takes as long as doing, the mechanism has failed and you’ll abandon it within two weeks.

Take this with you: sort by reversibility, not importance. Importance moves with your mood; the door doesn’t.

13. Bounded Agents

An agent, in this model, is not a tool and not an assistant. It's an **actor inside a specific engine, with a bounded role.**

That definition does real work. “Assistant” implies general helpfulness, which gives you no way to reason about scope. “Actor inside the Customer Engine, permitted to draft and never to send” is something you can design around, audit, and hold to a rule.

Every agent should have an answer to four questions: which engine, what role, what boundary, what evidence.

What to delegate

Delegate work that is **bounded, reversible, and checkable.** All three.

Bounded — the scope is clear and finite. Reversible — a two-way door. Checkable — you can tell whether the output is right without redoing it.

That third one is the constraint founders miss. If verifying takes as long as doing, delegation has produced nothing. It has moved the work from your hands to your eyes, which are the scarcer resource.

Before delegating, ask how you'll know it was done well. If you don't have an answer, either build the check first or keep the task.

What never to delegate

Judgment about what matters. An agent can rank by criteria you supply. It cannot decide what should be criteria. That's the whole job.

One-way doors. Per Chapter 12.

Relationships that are actually relationships. An agent can draft. Where a human on the other side believes they're in contact with you, you are the one who shows up. This isn't only ethics — it's that the trust you'd be spending is your Customer Engine's entire stock.

Deciding whether the venture is working. The most tempting and most dangerous. You cannot outsource the read, because the read requires a stake in the answer.

The question that prevents most waste

Before pointing an agent at anything:

Will this amplify a healthy engine, repair a weak engine, or accelerate a broken pattern?

The first is where agents earn their cost. The second is possible but slow, and usually needs the structure fixed before the automation. The third is where most of the ninety-five percent went — automating a broken loop makes it break faster and more expensively, while producing throughput metrics that look like progress.

Ask it before you build the agent. It takes ten seconds.

The run receipt

Every delegated execution should leave a record with three things: **what was attempted, what happened, and what it cost.**

Not for compliance. Because without it you have volume instead of evidence — and volume is precisely what you don't need more of.

What was attempted — the actual instruction, not your memory of it. Most agent failures are instruction failures, and you cannot see that without the original text.

What happened — the output, and whether it was accepted, revised, or discarded. The acceptance rate over time is the single best signal of whether a given delegation is working. If you're revising eighty percent of what comes back, that's not a working delegation, it's a slow way of writing.

What it cost — tokens, money, wall-clock time.

Cost is a signal, not just a bill

"My first big API invoice wrecked the margin."

That founder didn't have a pricing problem. They had a missing feedback loop — Chapter 6, Tier 3 — and the loop was missing in a specific way: cost information existed, but arrived monthly, long after the decisions that generated it. A Tier 4 delay problem sitting on top of a Tier 3 absence.

Treat spend as an operational signal that must arrive fast enough to act on. Per run, per day, visible where you work. Caps that stop rather than warn.

There's a second reason, less obvious and more useful: **cost per accepted output is a quality metric**. An agent that's cheap per call but wrong two-thirds of the time is expensive. That number tells you which delegations are actually working, and it's invisible unless you're recording both halves.

The drift problem, properly diagnosed

From Chapter 1: static identity files don't keep an agent coherent. They drift, because files don't update themselves and the venture keeps moving.

The instinct is to write a better file. That's Tier 5 — a bigger number of words, same structure, same decay.

The structural fix is a **single durable source of decisions that gets updated as part of the work**, rather than as a separate act of maintenance you'll stop doing by week three. Not a better document. A document with a maintenance loop attached, and the loop is the part that matters.

That's the next chapter.

Take this with you: if verifying takes as long as doing, you haven't delegated — you've moved the work to your eyes.

14. The Venture Constitution

Bounded rationality has a consequence nobody warns you about: **you will re-decide the same things.**

Not because you're forgetful. Because a decision made in March lives in the context of March, and by June the context is gone and only the conclusion remains — sitting there without its reasoning, looking arbitrary. So you re-open it. Sometimes you reverse it, then reverse back in August.

Each re-litigation costs a few hours and, worse, resets everything downstream that assumed the answer was settled.

The Venture Constitution is where decisions go to stop being re-decided.

What it is

A short, versioned document holding the decisions that are **settled until deliberately changed.**

Not everything. Not a wiki. The small set of choices that, if they moved every week, would make coherent action impossible:

- What you're building and what you've decided not to build
- Who the customer is — and who they aren't
- Naming and positioning, once chosen
- The one-way door list
- Standing boundaries for agents and collaborators
- Pricing posture
- What you will not do for money

Each entry gets three things: **the decision, the date, and the reason.**

The reason is the part people skip and the part that does the work. A decision without its reasoning is either dogma or noise — you can't tell whether it still applies, so you either follow it blindly or discard it. With the reason attached, you can check the reason. If it still holds, the decision holds. If the reason has expired, you have grounds to change it deliberately.

Decided, not correct

A constitution isn't a claim to have been right. It's a claim to have **stopped re-arguing.**

That distinction matters because it removes the fear that keeps people from writing one. You're not carving anything permanent. You're saying: this is decided, and it changes through amendment — dated, with a reason — not through drift.

The failure mode isn't a wrong entry. A wrong entry gets found and amended. The failure mode is a decision that was never written, so it gets silently re-made a dozen times in a dozen directions, and nothing downstream can rely on it.

It's what agents read

Here's where the two halves of this book meet.

Chapter 5: culture drifts because shared mental models don't maintain themselves. Chapter 13: static identity files decay. The Venture Constitution is the structural answer to both, and it works for one reason:

It's the same artifact you maintain for yourself.

That's the whole trick. A separate agent-context file drifts because maintaining it is a chore with no immediate payoff, so it slides down the list until it's three months stale. A constitution you actually use — that you consult when you're tired and unsure whether something was decided — stays current because *you* need it current. The agents read it as a side effect of you needing it.

Never build a second copy for the agents. The moment there are two, one is wrong, and it's always the one being read.

Keeping it short

The pressure will be to add. Resist.

A constitution that grows past a few pages stops being read, and an unread constitution is worse than none — it produces the belief that things are settled when nothing is actually being consulted.

Test for inclusion: **would re-deciding this cost me more than an hour, or break something downstream?** If not, it isn't constitutional. It's just a preference, and preferences can live in your head where they belong.

Prune when you amend. Decisions expire.

Amending

When something needs to change: change it, date it, keep the old version, note what moved and why.

The amendment history becomes something unexpectedly valuable — a record of how your thinking evolved, with reasons attached. It's the same benefit as the map's version history, and for the same reason: it's feedback on your judgment rather than on your output, and almost nobody has any.

Take this with you: never keep a separate copy for the agents. The moment there are two, one is wrong — and it's the one being read.

15. The Operating Excellence Review

Everything in Part 3 decays without one thing: a **recurring moment where a human looks at the whole system**.

Not the work. The system.

Twenty minutes. Once a week. Same time. It is the smallest possible balancing loop on your entire venture, and it's the mechanism that keeps the other mechanisms alive.

Why weekly, why twenty minutes

Weekly because it matches how fast things actually drift. Daily is too often — you'll see noise and react to it, which is worse than not looking. Monthly is too slow; four weeks is enough time to build substantially in the wrong direction. Weekly catches drift while correction is still cheap.

Twenty minutes because it has to survive a bad week. A ninety-minute review is excellent and you will skip it in month two, and a review you skip provides zero balancing force. Short and reliable beats thorough and occasional, by a very wide margin.

This is a Tier 3 intervention: adding a balancing loop where none existed. It's among the highest-value moves in this book, and it costs twenty minutes.

The agenda

Six items, in order.

1. Focus — is it still three to seven outcomes? More than seven means nothing is prioritized. Fewer than three usually means you've stopped looking at parts of the system. Cut or add.

2. What finished, continues, drops? Every active outcome, one word each. *Drops* is the one people avoid. Something should drop most weeks. If nothing ever does, you're accumulating commitments rather than making choices, and the accumulation is invisible until it's crushing.

3. Carry Forward — still real? Anything past its revisit date: promote it, or kill it. Do not let it sit. This is the maintenance that keeps the third bucket from becoming a graveyard.

4. Risks and blockers — resolve, accept, or close. Three options, no fourth. "Still monitoring" is not a state; it's a way of not deciding, and a risk you've been monitoring for six weeks is one you've implicitly accepted without saying so.

5. Did any signal change a decision? You gathered evidence this week. Did any of it move anything? If evidence never changes a decision, you're collecting it for comfort — which is expensive and feels productive.

6. One learning about the operating system itself. Not about the venture. About how you're running it. *"I approved three things I'd said needed review."* *"The map hasn't changed in a month and that can't be true."*

Why item six is the important one

The first five review the venture. The sixth reviews the reviewer.

Without it, you can run a disciplined weekly review for a year while the review itself slowly stops working — becoming a ritual you perform rather than an instrument you read. Item six is the loop that catches that.

One sentence. Written down, so you can notice when the same sentence appears three weeks running. Three repeats of the same learning isn't a learning anymore — it's a structural problem you keep observing and not fixing.

What this is not

It isn't a status report. Nobody's reading it.

It isn't planning — planning happens in the work, from the map and the brief.

It isn't a retrospective on your performance. The subject is the system, not you. This distinction matters more than it sounds, because a review that becomes self-assessment becomes something you dread, and dread is how weekly becomes occasional becomes never.

You're not grading yourself. You're reading an instrument.

When you skip it

You will. Something urgent will land on the day.

Reschedule within the week rather than letting it slide. Two consecutive misses and the loop is effectively gone — and the drift accumulated during the gap is exactly the drift that a review would have caught, so it compounds in the least convenient way.

If you find yourself skipping repeatedly, that's item six material. The obstacle is usually that the review has grown past twenty minutes. Cut it back rather than abandoning it.

Take this with you: short and reliable beats thorough and occasional, by more than you'd think.

16. Evolve or Push Harder

One question remains, and it's the one that decides whether the next year compounds or repeats.

Something isn't working. Do you push harder on the current approach, or redesign it?

Get this wrong in one direction and you abandon things that were about to work. Wrong in the other, and you spend a year applying force to a structure that cannot produce what you're asking of it.

The pattern that catches everyone

Limits to Growth is the most common trap in management, and it has a specific shape.

An approach works. It produces real success. Then growth slows — not because the approach got worse, but because it met a constraint that wasn't binding before.

The instinct is overwhelming: *this worked, we just need more of it*. So you push. And the pushing does nothing, because the thing you're strengthening was never what was limiting you.

Kodak had extraordinary film capability. Nokia had extraordinary hardware. Apple has extraordinary product secrecy — genuinely a strength, and one that sits awkwardly with the open feedback loops an AI relationship requires. In each case: the architecture that produced historic success meets a new constraint when the physics change.

Constraints are not only external. They live in goals, rules, capabilities, governance, and feedback loops — which means they're frequently things you built deliberately, that worked, and that you are now defending precisely because they worked.

The diagnostic

Four questions. They resolve it faster than deliberation will.

1. Is the constraint the same one as last time? If you're hitting the same wall repeatedly, force isn't the answer. It's structural — you've been running Tier 5 against a Tier 2 problem.

2. Is effort still producing proportional results, or has the curve bent? Proportional means keep going. Diminishing across several cycles means you've met a limit, and the limit will not yield to enthusiasm.

3. What would have to be true for pushing harder to work? State it. Then check whether it is true. This one sentence resolves most cases on its own, because the answer is usually “the constraint would have to be capacity, and it isn't.”

4. Is the thing I'm defending a strength or a habit? The hardest question. Strengths earned in a previous phase are the most expensive things to re-examine, because your identity is attached to them.

Push when

The constraint is genuinely capacity and capacity is genuinely growing. You're on a reinforcing loop that hasn't saturated. Results are still proportional. The approach is young enough that you haven't fairly tested it.

Give things a fair chance. Structural interventions have delays — Chapter 6 — and abandoning something before its delay has elapsed looks identical to abandoning something that failed.

Redesign when

The same constraint keeps binding. Effort and results have decoupled. The constraint is a rule, a goal, or a belief rather than a capacity. Or the physics changed — the thing that made your approach correct is no longer true about the world.

That last one is the hardest to see, because nothing announces it. The world doesn't send a notice. It just quietly stops rewarding what it used to reward, and the lag between the change and your noticing is where most of the damage happens.

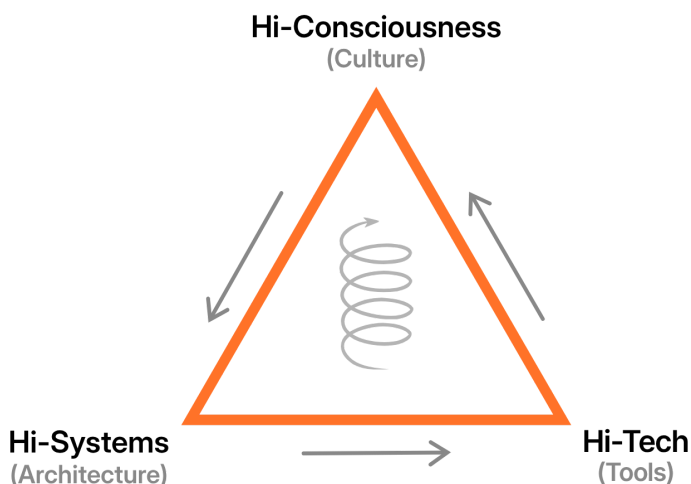
The captain and the designer, again

We started here. It's worth ending here.

Swap the captain of a boat and performance moves a few percent. The designer set the ceiling.

Everything in this book has been an attempt to move you, in specific and repeatable ways, from captaining to designing. The Read locates the constraint. The leverage hierarchy tells you where a move is worth making. The Map makes the system legible. The Brief separates what you know from what you believe. The Approval Path decides what reaches you. The Constitution stops you re-deciding. The Review keeps all of it honest.

None of it makes you smarter. Simon's ceiling doesn't move, and nothing in these pages pretends otherwise.



What it does is let a bounded mind hold an unbounded system — well enough to see where to push, and to know when pushing is the wrong move entirely.

That's the whole skill. It's learnable, it's small, and almost nobody is doing it.

You can build anything now. This is how you know what to build.

Take this with you: constraints are not only external. The most binding one is usually something you built on purpose, that worked.

The Twenty-Minute Audit

The whole method, one page. Run it weekly at first, monthly once it's habit.

MAP (5 min) — Draw the four engines. One line each on what's actually happening.

- **Product / Service** — is value actually reaching people, or stuck in process?
- **Customer** — are the right people finding you, trusting you, staying?
- **Cash** — do resources flow toward what's working, or away from it?
- **Skills** — is capability growing faster than the problem?

MARK (3 min) — Circle the weakest engine. Note the second one to watch. Remember: the bottleneck is where value stops compounding, not where the noise is.

NAME (5 min) — Four lines.

- **Engine** — which is limiting the rest
- **Input** — what it's missing or getting distorted
- **Loop** — which feedback path is broken
- **Driver** — Innovation, Governance, Interaction, or Culture — what's failing to coordinate it

MOVE (5 min) — One structural intervention. A sentence with a verb.

Then check its tier. Numbers and effort are Tier 5. Look for the Tier 2 version — the rule, or the information reaching somewhere new.

BEFORE AI (*2 min*) — Will AI amplify a healthy engine, repair a weak engine, or accelerate a broken pattern?

Change the system, not just the symptom.

Glossary

Engine — One of four places value accumulates: Product/Service, Customer, Cash, Skills. A living workflow, not a department. Exists whether or not anyone owns it.

Driver — One of four coordination mechanisms: Innovation (senses), Governance (decides), Interaction (moves information), Culture (holds shared models).

Engine anatomy — The seven parts every engine has: inputs, processes, outputs, feedback, actors (human and agent), and constraints (external, internal, governance).

Complex Adaptive System — A system whose parts perceive, decide, and adapt in response to each other. Behavior emerges from interaction, not from components. Your agents are part of this now.

Bounded rationality — Herbert Simon's finding that no mind can hold a complex system unaided. Not a defect; a permanent property. The reason a model beats more effort.

Leverage point — A place to intervene, ranked by power. Numbers are weakest; goals and paradigms strongest. Counterintuitive, and commonly pushed in the wrong direction.

Iceberg — The four depths of observation: events, patterns, structures, mental models. You fix where you look.

Limits to Growth — The most common systems trap. The architecture that produced success meets a new constraint; pushing the old strength harder doesn't remove it.

Sprint Zero — The bounded period between having an idea and committing to it. Produces a Map and a Brief. Timeboxed in days.

Venture Architecture Map — One page: engines, actors, constraints, risks, open decisions. Dated and versioned.

Sprint Zero Brief — Four lists: Clear, Needs Proof, Human Review, First Move.

Approval Path — Do Now / Needs Human Yes / Carry Forward. Sorted by reversibility.

One-way door — An irreversible action. Never delegated. Published content, sent messages, deleted data, spent money, changed consent, promises made.

Run receipt — The record of one delegated execution: what was attempted, what happened, what it cost.

Venture Constitution — Decisions that are settled until deliberately changed. Each with a date and a reason. The same artifact you and your agents both read.

Operating Excellence Review — The weekly twenty-minute loop that keeps the rest alive.

Where to Go Next

Run the read live. The Systems Bottleneck Scanner takes a written note about your venture — in your own words, the way you wrote it in Move 1 — and returns the likely bottleneck, the affected engine, the secondary engine to watch, the signal strength behind the call, and a safe next move. No signup. Do the paper version first; the read is the skill.

Go deeper. *Conscious Systems* is the full text behind this handbook — 25 chapters covering the 9 Universal Laws, the 6 System Archetypes, all four Drivers across all four Engines, AI integration patterns, and the complete Meadows leverage hierarchy in Appendix D. Free.

Go underneath. *ConsciOS: A Viable Systems Architecture for Human and AI Alignment* is the formal control architecture the whole framework rests on. For readers who want the cybernetics. Free.

Keep the model alive. Sysdom AI turns this method into an operating environment — your Map as durable state, your Constitution versioned, your agents bounded, your spend auditable.

Span of Control · Han Kay · Systems Intelligence · 2026 The original text is canonical and untouched. This handbook is a derived work for founders running on agents.