

Chapter 24. Eighteen pilots

Introduction

The pilots are not experiments to see whether an idea might work. They are working programmes designed to make the idea work in practice. We will build them, use them, learn from what fails, correct them and keep going until they deliver.

Learning is part of the work. When something fails, the response is not “the pilot failed”. The response is: “That did not work. Why? What do we change now?” Each failure gives us information for the next move.

Ignition matters just as much. Once a pilot works, it must be cheap enough, visible enough, simple enough and useful enough that other people can pick it up and use it themselves.

The standard is therefore demanding. Each pilot must produce a working mechanism, improve through repeated use and become easier to reproduce. We are not looking for proof of a theory. We are building practical capability.

This chapter stays at programme level. It explains what each pilot must achieve, where it sits in the sequence and what we will build through it. The Appendix contains a standardised pilot format for each of the eighteen pilots, detailed enough for a team to pick up the template and start work immediately.

How this chapter unfolds

The eighteen pilots are grouped and sequenced so that the programme begins with a common orientation, then builds individual and group capability, develops shared visibility, diagnostic, control and safety tools, and finally applies those capabilities to problem solving, paradox settlement, international cooperation and democratic decision-making. Many pilots can run in parallel, but the sequence makes clear what we are building and how the pieces connect.

Group A — Orientation and shared language

1. **More Mores for More — orientation and promulgation.** Build More Mores for More into a simple, usable public purpose for AI and spread it widely enough to orient the more specific aims that sit underneath it.
2. **Primitives Literacy.** Build simple, level-appropriate ways for people to understand the core primitives well enough to use the later tools intelligently.

Group B — Agency, self-understanding and mirrors

3. **Permissions.** Build practical ways for individuals to control what their AI can know about them and what it is permitted to do with that information.
4. **Biomarkers.** Add inexpensive physiological information where it genuinely improves the usefulness of a Personal Visibility Mirror, while keeping clear limits around what the data can tell us.
5. **The Personal Visibility Mirror.** Give individuals better visibility of their own patterns, inconsistencies and outcomes so they can reflect and choose differently.
6. **The Group Visibility Mirror.** Extend visibility to groups so that participation, disagreement, reciprocity, cooperation and problem-solving patterns can be seen and improved.
7. **National Visibility Mirrors.** Build national and, later, civilisational visibility without averaging away disagreement, minorities or changing conditions.

Group C — Diagnosis, truth, control and safety

8. **Periodic Tables.** Build usable maps of the relatively invariant parts of systems so that what is absent, weak or operating at the wrong level becomes visible.
9. **The Abacus.** Make paradox positions visible so that people can see where they agree, where they differ and what settlement is currently being held.
10. **Truth verification.** Build an AI-supported truth-seeking process that improves the handling of claims, evidence, uncertainty and correction without turning AI into the authority on truth.
11. **Control and correction using the 8+1 concept.** Build control systems that sense deviation, compare reality with intention and correct human systems faster.
12. **Early warning system for human failure modes.** Build a visibility system that detects recurring human failure patterns early enough for correction before they become crises.
13. **Pilot must-not-do and early warning.** Put explicit boundaries and dynamic early-warning mechanisms inside every pilot so that emerging risk is detected and corrected while the work continues.
14. **AI civilisation must-not-do.** Build a credible process for identifying, constraining and continually revising boundaries around AI capabilities whose consequences could be catastrophic.

Group D — Solving, settling and deciding together

15. **The ten-step solvable problem process.** Use disciplined AI-assisted problem solving to produce better solutions faster, while recognising when the issue is a paradox rather than a soluble problem.

16. **Paradox settlement.** Use truth-seeking, the Abacus, deliberation and AI to reach workable current settlements where no permanent objectively correct answer exists.

17. **USA-China cooperative paradox pilot.** Use shared visibility to move the USA-China AI relationship from slogans and winning-versus-losing toward specific paradoxes that the two countries must manage together.

18. **Turbo-charging democracy.** Put richer citizen allocation, the Abacus, factual information, AI support and Visibility Mirrors into a working democratic process that improves the information flowing between citizens and political decision-makers.

Group A — Orientation and shared language

1. More Mores for More — orientation and promulgation

More Mores for More has been the common aim running through the book. Now we put it to work. AI will otherwise be driven by the separate aims of laboratories, governments, companies, investors, users and increasingly autonomous systems. Some aims will align. Others will conflict. We need a simple human purpose strong enough to sit above them: more physical and relational abundance for more people over time.

This pilot turns that purpose into usable language. We will build stories, visual explanations, public communication, education, organisational statements and AI-supported conversations until people can understand the idea quickly, argue with it intelligently and use it when judging real AI choices. The aim is not propaganda. The aim is a clear public orientation that people can actually use.

The learning loop is straightforward. We use the phrase, watch where it works and where it does not, change the language, remove confusion and use it again. We keep refining it until citizens, builders, governments and organisations can ask a practical question of new AI uses: “Does this create More Mores for More?” Once people use that question without us prompting them, ignition has begun.

2. Primitives literacy

We will build short, level-appropriate learning modules around the basic primitives: the brain as a survival and problem-solving machine; its gifts and frailties;

predispositions and repurposing; holons and levels; solvable problems versus paradoxes; truth-seeking; cybernetic control and correction; More Mores for More; and what AI can and cannot appropriately do. This is not a philosophy course. A child, manager, politician and AI researcher need different depths of explanation while sharing the same underlying ideas.

The work is to make these ideas usable. People should begin recognising what is happening in real time: threat, status, felt truth, wrong-level problem solving, hidden paradoxes, poor incentives and failed correction. If the first version of the material does not produce that recognition, we simplify it, improve it and try again.

The result we are driving toward is practical literacy. People should be able to say: “I am defending my position rather than examining it.” “We are treating a paradox as a problem.” “That rule changes the incentives.” The modules must become simple enough to travel into schools, companies, governments, communities and online learning, because the rest of the pilot programme works better when people share this language.

Group B — Agency, self-understanding and mirrors

3. Permissions

Permissions provide the practical architecture through which individuals decide what their AI can know about them and what it can do with that information. They can progressively cover documents, communications, behaviour, biomarker information and more intimate information voluntarily supplied by the individual.

We will make the boundary between usefulness and intrusion explicit. The system must show how much information the mirror actually needs, what permission has been given, what benefit that permission creates and how easily it can be withdrawn.

The required result is a permissions bargain people understand and control: greater access produces visible usefulness while the person retains agency. Once that works reliably, it becomes reusable infrastructure for thousands of later AI applications.

4. Biomarkers

We begin with information already obtainable cheaply from watches and similar devices, including sleep, heart rate, activity and other physiological measures. With explicit permission, relevant information can feed the Personal Visibility Mirror.

The work is to find the physiological signals that genuinely improve judgement and reflection. We are not claiming that a watch can read a mind. We will use the data, compare it with behaviour and outcomes, discard what adds noise and keep what repeatedly adds useful context.

The result we want is simple: the mirror occasionally sees something useful that the individual did not see—for example, that particular decisions or conflicts repeatedly coincide with fatigue or physiological stress. If that happens reliably, existing mass-market hardware makes the capability cheap and scalable.

5. The personal visibility mirror

The Personal Visibility Mirror is one of the pivotal pilots. With permission, AI progressively remembers what I say, what I do, what I say I want, what subsequently happens and any relevant physiological information, and then reflects the patterns back to me.

The mechanism is straightforward. Humans are good at seeing inconsistency in others and much less reliable at seeing it in themselves. Memory is selective. We rationalise. We defend status and identity. We forget predictions that proved wrong. The mirror changes the information available to us. It does not need to order us to behave differently; it makes it harder not to see ourselves.

We will improve the mirror until better self-visibility produces useful reflection without becoming irritating, manipulative or intrusive. Observation must remain distinct from judgement. The target is the moment when a person can say: “I can see something about myself that I could not see before, and because I can see it, I can choose differently.” That is the capability we are building.

6. The group visibility mirror

The Group Mirror applies the same principle to families, teams, companies, communities and deliberative groups. Subject to permissions, AI observes the process and makes patterns visible: who speaks, who does not, repeated disagreements, areas of hidden agreement, inconsistencies, unresolved assumptions and perhaps movement over time.

We will use the Group Visibility Mirror until groups can see themselves more accurately and use that visibility to improve domination, exclusion, listening, reciprocity, cooperation and problem solving. Behaviour and consequence become visible together, so the group can correct itself rather than merely discuss what went wrong afterward.

The result we want is better participation, less repetition, earlier recognition of real disagreement and a stronger ability to reach settlements without suppressing minority positions. Once ordinary teams and communities can demonstrate that repeatedly, the mechanism will be easy for others to copy.

7. National visibility mirrors

The national architectural Visibility Mirrors extend the same principle beyond the group. They use AI to represent what a nation—or ultimately humanity—is saying, doing and experiencing: agreements, disagreements, contradictions, changing values, factual conditions and longer-term consequences.

The job is to make very large populations visible without collapsing millions of voices into a meaningless average. Minorities must remain visible. Competing truths and paradox positions must remain visible. Where the first representation distorts reality, we correct the representation and improve it.

The result is a national picture that citizens and leaders can use because it is richer than polling, political rhetoric or social media alone. Once one country can see itself more clearly in this way, other governments and populations will have a strong reason to build their own versions.

Group C — Diagnosis, truth, control and safety

8. Periodic Tables

Periodic Tables are structured maps of the relatively invariant components—the nouns—of systems whose actual processes—the verbs—vary by level, culture, institution and technology. We already have candidates for governance, truth-seeking, deliberation and corporations.

We will use the Periodic Tables against real systems, challenge the nouns, remove weak categories, add what is missing and keep refining the levels until the tables reliably show where a system is absent, weak or operating at the wrong level. In corporations, the same method can expose capability gaps against competitors, industries and functions.

The required result is immediate visibility: a user can look at the table and say, “Here is the part of our system that is absent, weak or operating at the wrong level.” Publishing the tables openly and improving them through use is part of the work, not an optional academic exercise.

9. The Abacus

The Abacus is a way of making paradox positions visible along spectrums rather than forcing them into binary choices. It can be applied to nations, corporations, leaders and particular decisions.

We will use the Abacus until people can distinguish factual disagreement from legitimate value disagreement and can see where each side actually sits rather than hiding the dispute inside left/right or yes/no labels.

The result is a much better conversation: “We disagree, but now we can see precisely where we disagree and the settlement we are currently prepared to accept.” The Abacus only becomes real when people use it on live problems, so worked demonstrations are part of the pilot from the start.

10. Truth verification

Truth Verification would use an AI-supported process that asks what the claim is, what evidence supports it, where the evidence came from, what contradicts it, how confident we should be and what would change our conclusion.

Truth Verification builds an AI-supported process around a simple discipline: what is the claim, what evidence supports it, where did that evidence come from, what contradicts it, how confident should we be and what would change our conclusion?

We will improve the process until people can separate fact, inference, uncertainty, opinion and manipulation more reliably and correct themselves faster when evidence changes. AI assists the search for truth; it does not become the oracle.

11. Control and correction

The Control and Correction pilot applies a general-purpose control and correction architecture to ask whether a human system actually knows what it is trying to achieve, can sense what is happening, compare reality with intention and correct itself.

We will apply Control and Correction to real programmes, companies and pilots until deviations are seen earlier and corrected faster. The work is practical: define the aim, establish measures, sense what is happening, compare it with intention, assign authority to correct, and learn from the result. AI makes continuous sensing and comparison much cheaper.

The required result is a system that discovers drift while there is still time to act. We will use the architecture across very different settings and keep improving it until the same control logic proves useful repeatedly.

12. Early warning system for human failure modes

Human systems tend to fail in recurring ways. Threat can narrow time horizons. Status conflict can become domination. Belonging can harden into tribalism and exclusion. Competition can become predation. Anger can become revenge. Felt truths can harden while truth-seeking weakens. Reciprocity can break down and short-term advantage can overwhelm longer-term consequence. The purpose of this pilot is to see whether those movements can be detected early enough to act before they become crises.

We will combine AI with the Personal, Group, and National Visibility Mirrors, the Abacus, Truth Verification and Control and Correction. We will look for changes in patterns rather than pretend to read individual minds: rising threat language, collapsing dialogue, widening paradox positions, repeated truth failures, loss of reciprocation, concentration of power and other agreed indicators relevant to the system being observed. We can begin inside a company, community or public institution where the history and outcomes are visible.

The target is an early-warning process that sees meaningful deterioration before ordinary management or political processes do, triggers review and correction, then learns from whether the warning was useful or false. False alarms are not a reason to stop. They are data for improving the warning system. The purpose is earlier visibility of known human failure modes while there is still room to act.

13. Pilot must-not-do and early warning

Every pilot should carry its own explicit boundaries and monitoring system. The Must-Not-Do component specifies prohibited actions, while the Early-Warning component watches trajectories before a boundary is actually crossed.

We will make safety dynamic. Every pilot will carry explicit boundaries and an early-warning system that watches the trajectory before a boundary is crossed.

When an emerging problem appears, the pilot slows, changes or stops the relevant part, records what happened, corrects the design and continues. Safety therefore becomes part of making the pilot work, not a reason for avoiding the work.

14. AI civilisation must-not-do

This pilot builds the process by which societies identify AI capabilities whose consequences could be catastrophic and therefore require strong boundaries—for example, autonomous lethal systems or dangerously democratised biological engineering capabilities.

We will define the boundaries, build ways to monitor them, create processes for international agreement and keep updating them as technology changes. Conventional regulation is too slow for this task, so the control process itself has to learn.

The required result is not a permanent list. It is a working international process that can identify catastrophic risks, agree restrictions, detect violations and revise the boundaries as capability changes. We keep building it until enough governments, laboratories, companies and citizens recognise that preventing shared catastrophe is a common interest.

Group D — Solving, settling and deciding together

15. The ten-step solvable problem process

This pilot takes real problems for which objective improvement can be identified and applies the disciplined AI-assisted problem-solving process.

We will use AI to accelerate diagnosis, option generation, modelling, implementation and correction. Just as importantly, the process must recognise when the apparent problem is actually a paradox and switch methods rather than force a false solution.

The target is better solutions, produced faster and more cheaply, followed by measurement and correction. We choose conspicuous real problems where results can be measured, make the process work, learn from every failure and keep improving it until the advantage is obvious.

16. Paradox settlement

For problems without a permanent objectively correct answer—freedom versus order, present versus future, competition versus cooperation and so forth—we use truth-seeking, the Abacus, deliberation and AI to reach a workable current settlement.

We will improve the process until people can disagree more intelligently, understand each other's positions, see consequences and identify a settlement without handing the value decision to AI.

The target is not unanimity. It is a legitimate settlement that participants understand, can live with, can monitor and know can later be changed. We keep working the process until disagreement no longer has to mean either false consensus or paralysis.

17. USA-China cooperative paradox pilot

The United States and China are competitors, but they are also exposed to many of the same AI paradoxes. Capability versus restraint, competition versus cooperation, secrecy versus transparency, national advantage versus global responsibility and present advantage versus future consequence do not disappear because the countries have different political systems. We are going to make those shared tensions visible at enough granularity to create room for cooperation where cooperation is necessary.

The pilot uses the same ten-row Abacus and the National Visibility Mirrors developed earlier in the book. Participants from both sides place their own country and the other country independently, then compare the results. Truth verification

separates factual disagreement from value disagreement. We then work the small number of paradoxes where both sides can see the same failure risk and where neither country can create a safe settlement alone.

The first win does not need to be a grand bargain. It needs to be one bounded cooperative settlement that both sides can explain and monitor: a shared threshold, reciprocal disclosure rule, incident-communication process, testing protocol or review mechanism around a clearly shared risk. We make one work, learn from it, strengthen the process and move to the next. The purpose is to change the conversation from slogans and winning-versus-losing into specific tensions that can actually be managed together.

18. Turbo-Charging Democracy

This is where several streams come together. We will put positive and negative allocation voting, including the \$10,000 mechanism, alongside the Abacus, factual information, AI support and Visibility Mirrors. Citizens will do more than choose between political packages every few years; they will reveal the intensity and direction of their preferences and confront the trade-offs themselves.

We will run the process, compare citizens' allocations with politicians' choices and actual expenditure, make the differences visible and require explanation. Then we do it again the next year. Citizens learn. Politicians learn. Bureaucrats learn. The democratic system becomes a visible correction loop rather than a largely hidden argument between elections.

The target is better information in both directions and better decisions over repeated cycles. We start at a scale where the mechanism can be built within months, make it work, show the result and let political competition drive adoption. A successful local demonstration should not remain local for long.

And finally: Ignition

The first pilot establishes the orientation; the remaining seventeen turn that orientation into working mechanisms. Ignition begins when a mechanism works well enough that other people want it.

Ignition is therefore a design requirement alongside low cost, high speed, high visibility and rapid sharing. We are not waiting for one institution to build the new civilisation. We are building working sparks that other people can pick up and improve.

A Personal Visibility Mirror works. Somebody copies it. An early-warning system catches a human failure pattern before it becomes a crisis. Another organisation adopts it. A town makes Turbo Democracy work. Another town improves it. A

company uses an Abacus and competitors adopt variants. A school teaches brain literacy and other schools take the material. A USA-China cooperative pilot produces one reciprocal settlement that works, and the same method moves to the next shared problem.

Some attempts will fail. Good. We learn why, change them and go again. Some will combine. Some will become much better than the first version. None of that changes the direction of travel.

The end state is not a shelf full of pilot reports. It is a growing field of working mechanisms that people are using, improving and spreading faster than any central programme could direct.

The programme therefore has a simple discipline: Make it work. Learn when it does not. Correct it. Make it work better. Then make it easy for others to use. Repeat that often enough and the pilots become a civilisation-building process.