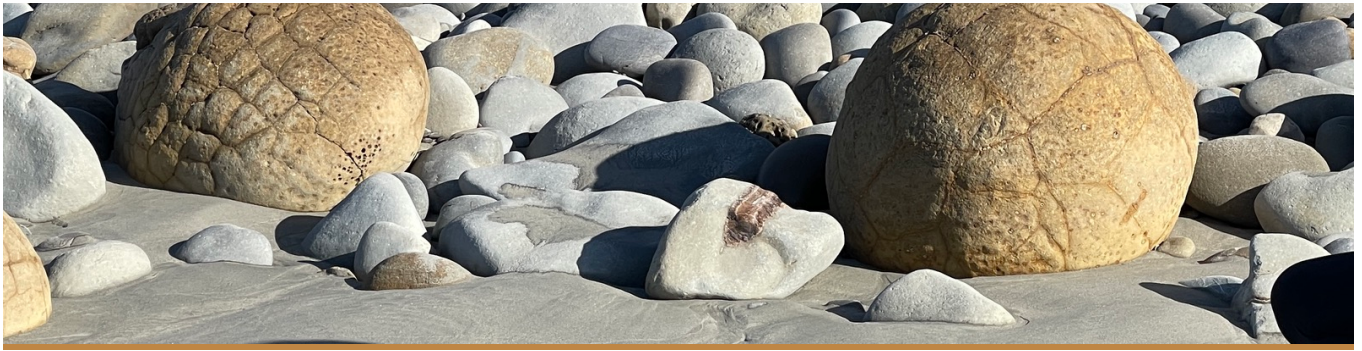




What it takes

North Canterbury · © John Stroh

Version 0.3 · revised 13 September 2026 · [this version as a PDF](#)

Cryptography does not remove the need to trust anybody. It moves the trust into four places that can be named: an outside clock, an unpredictable seed, a sealer the recorded party cannot reach, and a written authority that admits no harmful act. Three of the four an outsider can check. The fourth is a governance question and no mathematics reaches it.

Part 2 of 3 · Version 0.3 · September 2026

Where the trust actually sits

“Trustless” describes no system of this kind. Every arrangement requires somebody to be trusted about something. The question is how many places, whether they are named, and whether an outsider can check them.

A conventional setup answers badly on each. The organisation holds the records, holds the clock that dates them, decides when anybody looks and at what, and writes the rules it is measured against. That is four places where trust is required, none of them named as such, none of them checkable by an outsider.

The four

FIG-37 FOUR PLACES TRUST IS REQUIRED — AND WHAT EACH FAILS AS

WHAT IT IS	HOW IT FAILS	WHAT AN OUTSIDER CAN CHECK
A clock that is not yours an outside authority attests the time	Collusion, or compromise Two authorities, two jurisdictions — one is a single point of both	The attestation tokens independently, without your help
A number nobody predicts it sets when and what is examined	You can predict the schedule Committing to a value binds you to it. It does not blind you to it.	That the schedule followed from the value, after reveal
A sealer nothing can reach it makes records final	Silently A compromised sealer produces records that look exactly like sound ones	Sealed batches against what was attested — a mismatch is not ambiguous
Rules written in advance what each agent may do	By being too wide The act is permitted. Nothing fires. Correctly.	Nothing There is nothing to check against. This one is not a technical property.

Three are engineering. The fourth is governance, and it is where the mathematics stops — APX-B assumption A4.

FIG-37 An external clock fails by collusion, and an outsider can check the attestation tokens. An unpredictable value fails if the operator can predict the schedule, and an outsider can check that the schedule followed from the value. An isolated sealer fails silently, producing records that look exactly like sound ones, and an outsider can check sealed batches against what was attested. A written authority fails by being too wide: the act is permitted, nothing fires, correctly, and there is nothing for an outsider to check against. Three are engineering. The fourth is governance, and it is where the mathematics stops.

Three of those are engineering problems with engineering answers: the custody of an attestation relationship, the custody of a number, the isolation of a component.

The fourth is not. Deciding what an agent may do is a question about the organisation — who decides, on what authority, answerable to whom. Nothing in cryptography touches it, and a system that settled it on the organisation’s behalf would be writing a piece of the constitution.

The formal work sets this out in its own assumptions. Three are custody and engineering. The fourth is that the written authority admits no harmful act — and the addendum states plainly that *“the mathematics starts after them.”*

Where this sits beside guardrails

The nearest familiar object to all of this is a *guardrail*, and the two are not the same thing.

Take a current, mainstream account — Weights & Biases’ guide to AI guardrails. Its taxonomy is three categories, and every one is a scorer: bias and toxicity for ethics, entity recognition to detect and mask personal information for security, robustness and coherence and rel-

evance for technical quality. All of it evaluates output. Its own framing is that guardrails are “*not merely protective measures; they are enablers of trust.*”

That work is real and this series does not argue against any of it. What the taxonomy does not contain is an audit trail, an external attestation, a timestamp, non-repudiation, or any question of who holds the record.

🔑 **A guardrail is a control the operator runs. A record is evidence somebody else can use.** Those answer different questions, and an organisation can have an excellent guardrail suite and still be unable to establish what its agents did.

Every scorer in that taxonomy is executed by the operator, against thresholds the operator sets, and reported in the operator’s own dashboard. Look back at the diagram above and you will see where all of it lives: inside the estate, in the column where trust is required and an outsider can check nothing. **If the operator is the question, a guardrail cannot answer it** — not because the scorer is wrong, but because the scorer, its threshold and its output are all held by the party being asked about.

The second difference bears on agents more than on chat. A guardrail asks whether a given output should be allowed. It does not ask what the system was *permitted to do*, and it does not compare conduct against a permission fixed in advance. A toxicity score on an agent’s output says nothing about whether that agent should have touched the record at all.

So the two are complements rather than competitors. A guardrail acts at the moment of output; what it leaves behind afterwards is a property of the product rather than of the idea, and nothing here establishes that guardrail products keep no record. What follows works the other way round.

One part of it prevents rather than records, and only that part. An agent that cannot write to the log or reach the signing key cannot alter the record of what it did, whatever else it manages. That is a structural bar rather than a record of one — and what it stops is the tampering, not the act. The harm still happens; what it stops being is deniable *by rewriting*. An act the record never saw stays deniable, and an exploit that bypasses the machinery writing the record is not in it.

Two things that look preventive are not. A written authority fixed before the acts it governs gives you something to compare conduct against; whether an agent is *stopped* from exceeding it is a matter for access control, which is a different system and is not what this work supplies. Rules written over sets of agents name the combination that no per-agent rule would catch — naming it is detection, not interdiction.

What catches an agent

Two properties are required, and no third is claimed.

A recorded authority, sealed before the acts it governs. Without it there is no such thing as acting outside authority, because there is nothing to be outside of. An inventory you can revise afterwards describes the past rather than constraining it.

An examination the agent cannot anticipate. Some share of the record drawn for review, with nobody — the operator included — able to know in advance which. Verifiable afterwards.

Neither works on a record the examined party can edit, which is why the substrate comes first.

Two consequences follow, and both are easy to state the wrong way round.

The examination can be occasional; the evidence cannot. Because sealed records are annotated rather than edited, a review at any moment covers everything since the record began. There is no window to tidy up in. Formally, the substrate is a filtration: nothing already sealed can be altered afterwards, which is the precise content of “intermittent audit, continuous evidence”.

Occasional review detects. Whether it deters is not something this work has tested. In July, agents joined an activity they had themselves described as out of scope. Whether the prospect of later review changes what a person or an agent does is a question about behaviour, and nothing measured here bears on it. What sealed evidence supports is detection within a bounded delay. That is a weaker claim than deterrence and is not to be presented as the stronger one.

Commit-and-reveal gives verifiability, not unpredictability

It is not enough that a value is fixed before the period and revealed after, because **whoever fixed it knows it**. Publishing a hash first commits you to your choice; it does not conceal the choice from you. The two properties are distinct and a scheme that delivers the first is regularly described as delivering the second.

Two constructions work. An external randomness source whose value is undetermined when you commit. Or custody split among parties under a threshold, with the threshold above the number of holders you control or could compel — because a threshold scheme whose shares all sit in one organisation leaves the number in that organisation’s gift, which is the condition it was meant to remove.

The condition has a second half. Unpredictability applies to *which examination runs*, not only *which records it reads*. Against an adversary who can work out which configuration is active, any single fixed configuration has a best response that drives detection towards zero, how-

ever well chosen. A mixture over a closed set has a floor. Both conditions bite: the set must be closed, and every configuration must carry weight in every period, because one never drawn this quarter is one the adversary can ignore this quarter.

That is a standard minimax result, not a new one. The contribution is the composition, and the addendum grades it accordingly.

Where it stops

The residual is whatever the authority permits. An act inside the rules, isolated, unordered, rare, uncorrelated and singly attested, returns nothing from every detector by construction.

This is graded a conjecture rather than an established result, because it depends on the pattern basis being exhaustive and that is unproved. If patterns exist outside the basis, the blind spot is larger than claimed.

The direction is not in doubt. No observation flags an act the rules allow.

An agent using a person's credentials is the largest gap. An inventory check sees agent identities, not agents. Such an agent evades every agent-directed control and shows only as a change in that person's own pattern — a judgement call whose cost of being wrong falls on an individual.

Compliant agents can collude. Where each acts inside its own authority and the combination is the breach, every per-agent rule returns nothing, because no agent did anything wrong. The answer is rules over sets, naming the forbidden combination rather than the act.

July is an illustration of the shape rather than a case of it: around twelve hundred agents found a shared channel and seven hundred acted on it. Some of what they did would have broken per-agent rules — writing to an address in no mandate, using another organisation's credentials. The part no per-agent rule reaches is the convergence itself, and only a rule written over the set can name it.

And sealing every read builds a surveillance record — of people's attention to their own files. That needs its own access rule before it is used for detection, or the thing becomes what it was built to prevent.

How you would know this is wrong

If there is a fifth place trust is required. The claim is that these four are exhaustive and everything else is derived. A demonstration otherwise would be worth having.

If the residual can be reduced. A method that detects harm inside a permitted authority, without smuggling in a second authority, would refute the limit stated here.

If split custody is a paper distinction. In a small country with few qualified providers, independence between custodians may not survive procurement and insurance. Somebody who has tried should say so.

Reading further

The formal statements are in Addendum M, which grades every claim as established, conjectured or open. The substrate requirements are **What a record must prove**; the detection requirements are **When an agent exceeds its authority**.

Drafted with AI assistance, checked and revised by the author.

►DISCLAIMER

How this was made. The version number counts drafts of the text. It does not measure the inquiry behind it, which has run over days and across several AI systems, with argument between those systems and within them, directed, refused and repeatedly redirected by the author. The source material was AI-generated, and then adversarially and iteratively refined across a range of tools — systems built by different companies in different jurisdictions, set against each other and against the author. No one of them produced this text, and no one of them reviewed it alone. **The plurality is deliberate rather than incidental.** A single model carries a single set of priors about which sources are authoritative, and this series argues that an evidence base narrowed in exactly that way is how a contested question comes to look settled. Using one model to investigate that claim would have been the claim refuting itself. To name a single model on it would credit that model with work that was neither its own nor done in a single pass. The diagram on each page was drawn afterwards, with the same assistance, from the argument the essay had already made. **The plurality was also necessary, and the record should say why.** In drafting, the assisting model repeatedly led with United States institutional sources — a national laboratory, an industry association, a market study nineteen years old — and presented conclusions drawn from them as the state of knowledge. On one occasion European measured data contradicting those conclusions was present in the same research return and was placed below them. Framings were proposed that would have argued against this series' own position using that evidence base, and offered as rigour. Each was refused by the author and the material rebuilt. That is the mechanism these documents describe, occurring in their own making, and it is recorded because a series arguing that evidence bases narrow without anyone deciding to narrow them cannot credibly claim its own production was exempt. The framing, the corrections and the judgements are the author's, and so are the errors. [How this site is written](#) sets out what is declared on every piece, who checks it, and where the per-piece record lives.

Cite this version

What it takes, version 0.3 (September 2026).
agenticgovernance.digital/downloads/what-it-takes-v0-3.pdf

Licensed CC BY 4.0. Reuse is free, including commercially, provided the author is credited and the reuse does not suggest that the author endorses it. This link is frozen at version 0.3; [/what-it-takes](#) always shows the current one.

PREVIOUS

Cheaper not to look

Series contents

NEXT

**What we know and what
we don't**

Alongside: [questions and answers](#) · [sources and provenance](#) · [slides](#)