

MAZE.AI: A Claim-Bounded Answer Loop with Machine-Checked Provenance

J. Konstapel Leiden, 20-9-2026.

An AI that can be trusted is not one whose every behaviour is proven — that is impossible for open language — but one whose every claim is proven. This paper describes a running system built on that criterion, the theorems that bound it, and the tests that tie the theorems to the code in production.

J. Konstapel, Leiden — with computational assistance from Claude (Anthropic) · 19 September 2026 · live: maze.swarp.nl/ai · code: [maze/ai](#), [maze/lean](#) (Lean 4.34, no Mathlib)

Abstract. Large language models produce fluent assertions with no bound on what they assert. We argue that the only certifiable relation between an answer and the world is provenance: that a piece of content was retrieved from a specific, certified store at a specific, computed address, along a derivation that is written down. MAZE.AI is an answer loop whose output type admits exactly four forms — derived, attested, vacancy, proposal — each carrying its address, and whose emission path is single. The three invariants of the loop (derived $\Rightarrow$ the pair is in the store; vacancy $\Rightarrow$ the address is empty; attested $\Rightarrow$ a ledger entry settled this address for this canonical question, by a named person, on a date, and the content is in the store) are proven in Lean 4 against the running definitions, in both a list-

store and a finite-map-store model where the first invariant becomes an equivalence. The address kernel (balanced-ternary bijection, range, uniqueness, prefix-stable truncation) is proven and differentially tested against the production JavaScript kernel (30,572 addresses, 0 differences). A separate verifier module re-reads store and ledger before any emission and refuses rather than rewrites. The ledger is append-only, hash-chained and signed; the canonicalisation of questions is a versioned module with a determinism test. A pre-registered twenty-question placement test scored 15/20 against a threshold of 14. The language model is retained only as a proposer: it may phrase, it may not assert. We know of no other system that combines these elements in one output path with machine-checked invariants; we claim novelty only for the combination, and state the limits — placement quality, absence of personal keys, and the semantic gap between "found on an image of the address" and "found at the address" — explicitly.

1. The problem: assertion without bound

A language model is a function from text to text. Nothing in that function distinguishes a sentence that reports something from a sentence that merely continues something. Retrieval-augmented generation reduces the frequency of invention but does not change the type of the output: the model still emits an assertion, now decorated with citations that it also chose. The person reading it cannot tell, from the artefact alone, which parts are held by the sources and which are the model's continuation.

The obvious repair — "prove the model correct" — is blocked in principle. Correctness of an assertion about the world requires a truth

relation between sentences and states of affairs; for an open language that relation is not a total, certifiable object (Tarski's undefinability theorem gives the sharp form). So "every behaviour proven" is not a research programme with a long horizon; it is not a programme at all.

2. The criterion: every claim proven

What can be certified is narrower and, we argue, sufficient: the relation between an output and a store. Define a certified store as a set of (address, content) pairs whose membership is decidable. Then an answer system satisfies criterion P if every output it can emit is of a type whose meaning is a decidable statement about the store, and each such statement is proven for the emitting function. Under P the system never says "X is true"; it says "X is what stands at address a, reached by derivation d" — and that sentence is either provable or the output cannot be constructed.

P moves the burden of meaning to where it can be borne: to the person who settles that a given address answers a given question, and to a ledger that records that settlement as an event. The machine proves provenance; the human settles meaning; neither pretends to do the other's work.

3. The construction

MAZE.AI runs as a context under the MAZE, a self-addressing ternary knowledge net of 8.19 million pieces (Wikipedia, Mathlib theorems, EPO patents, art, music, open questions) in which every piece has an address computed from its content. Five components, each small:

1Kernel (C1). Addresses are lists of trits ($-1, 0, +1$) in balanced ternary. The kernel is a bijection between addresses of depth n and integers in $[-(3^n-1)/2, (3^n-1)/2]$; truncation to a prefix is idempotent and stable under deepening. Proven in Lean (Trit.lean, Kern1.lean, Brug.lean).

2Store (C2). The production store is PostgreSQL (address, content, receipt of origin). Two Lean models: a list of pairs, and a finite map — a list with a proof that each address carries at most one content (Winkel.lean), matching the MAZE invariant that an address is never rewritten after admission.

3Loop (C3). One function, one output type with four constructors, each carrying an address: inductive Uitvoer

```
4 | afgeleid   (addr : Int) (content : String)           -- derived
5 | voorstel  (addr : Int) (bron : String)              -- proposal
6 | vacature   (addr : Int)                             -- vacancy
7 | geattesteerd (addr : Int) (content : String) (who : String) (date : String)
```

The loop reads the ledger for the canonical question first; a settlement whose content is still in the store yields attested. Otherwise it reads the store at the computed address: content yields derived, nothing yields vacancy. A proposal is what a language model may supply — a phrasing with the addresses it cites; it enters the loop and leaves it only as derived or vacancy.

8Verifier (C4). A separate module through which every output passes before emission. It does not trust the loop: it re-reads the store and the ledger and checks the invariants on the output as it stands. If it refuses, nothing is emitted (an error with reason) — never a silent retyping.

9Ledger (C5). Append-only rows (canon version, canonical question, material id, address, who, via, date, previous hash, hash, signature, key fingerprint); hash = sha256 over previous hash and fields; signature = HMAC over the hash with the key that admitted the settlement. A chain check walks the whole list. Two settlements of the same question both remain; the latest is operative.

Two input layers feed the loop: exact names (a question that is the name of an inhabitant resolves by identity), and open language (embedding of the question through the same gate as material $\rightarrow$ 24-trit address $\rightarrow$ deepest inhabited prefix $\geq 6 \rightarrow$ cosine ranking within that neighbourhood; a name bridge with a gate test rejects name hits that lie far from the question's meaning). The canonicalisation `canon(q)` — NFKC, lower case, letters/digits/spaces only, single spaces, trimmed, 400 characters — is its own versioned module with a determinism test (frozen examples, idempotence, 5,000 random inputs).

One output path. The MAZE's main conversational page also goes through the loop: the language model's answer is delivered as a proposal, the loop's output is shown above it, and both pass the verifier. There is no debug path, and no second path "for now".

4. The theorems

All statements below are checked by Lean 4.34 with no additional axioms beyond `propext`, `Classical.choice` and `Quot.sound`, and no sorry. Names are as in the source.

KERNEL

waarde_bereik: $|\text{value}(a)| \leq (3^{|a|}-1)/2$. waarde_inj: equal depth and equal value $\Rightarrow$ equal address. waarde_ofInt: $\text{value}(\text{ofInt } n \ v) = v$ on the range.

trunc_idem, trunc_prefix: truncation is idempotent and stable under deepening. waarde_eq_decode_reverse: the two kernel formulations (deepest-first, shallowest-first) compute the same integer.

Proof sketch: induction on the address list with Int.pow_succ ; the range bound by generalising 3^n and case analysis on the head trit (ω); injectivity by contradiction from the bound on tails.

LOOP, LIST STORE (Lus.lean)

lus_afgeleid_store: $\text{lus } s \ g \ q \ \text{addr} = \text{afgeleid } a \ c \Rightarrow (a, c) \in s$.

lus_vacature_leeg: $\text{lus } s \ g \ q \ \text{addr} = \text{vacature } a \Rightarrow a = \text{addr} \wedge \forall c, (a, c) \notin s$.

lus_geattesteerd_grootboek: $\text{lus } s \ g \ q \ \text{addr} = \text{geattesteerd } a \ c \ \text{who } \text{date} \Rightarrow (\exists e \in g, e.\text{canonQ} = q \wedge e.\text{addr} = a \wedge e.\text{who} = \text{who} \wedge e.\text{date} = \text{date}) \wedge (a, c) \in s$.

Proof sketch: unfold the loop; case analysis on List.find? over ledger and store; membership from $\text{List.mem_of_find?_eq_some}$, emptiness from $\text{List.find?_eq_none}$.

LOOP, FINITE-MAP STORE (Winkel.lean)

Winkel.vind_iff: $\text{lookup } s \ a = \text{some } c \Leftrightarrow (a, c) \in s$.

antwoordW_iff: $\text{answer } s \ \text{addr} = \text{derived } \text{addr } c \Leftrightarrow (a, c) \in s$ — the first invariant becomes an equivalence: the loop emits exactly what the store holds.

antwoordW_vacature_iff: $\text{vacancy} \Leftrightarrow$ the address is empty.

Winkel.voeg_bewoond: inserting on an inhabited address changes nothing (write-once as a theorem); vind_voeg_leeg, vind_voeg_andere: insertion on an empty address is visible there and nowhere else.

The three loop invariants are re-proven over the map (`lusW_*`).

Proof sketch: the uniqueness field of the structure closes the converse direction; the write-once property follows from a dependent if on the lookup.

LEDGER (`Keten.lean`)

`chain_deterministic`, `totality`, `chain_snoc` (prefix property of the chain over a list), and `chain_last_changes`: altering the last entry changes the final hash — proven only under the explicit hypothesis that the step function is injective. Without collision-freeness it is false, and Lean does not let it be stated otherwise.

What is not proven, and not provable: that any derived content is a good answer to the question. That is the semantic boundary of §1, and the system does not cross it. It is measured instead (§6).

5. Binding the theorems to the running system

A proof about a model is worthless if the production code is a different function. Three bindings:

- Differential kernel test. The Lean executable prints 30,572 addresses ($-100,000 \dots 100,000$ at depth 12; 2,000 pseudo-random integers up to 3^{42} at depth 42) in both formulations; the production JavaScript kernel prints the same series; diff reports 0 differences.
- Production-path test. The real loop runs on the real store for a fixed question list; every output is passed through the verifier (C4), which re-reads store and ledger; the ledger chain is walked. Any refusal is red.
- One control command. `node maze/ai/control.mjs` rebuilds the Lean project, prints the axiom report and rejects anything beyond the three standard axioms, runs the differential test, the canon

determinism test and — with a database — the production-path test. It stops at the first red step.

The residual gap is stated rather than hidden: production "derived" uses the deepest inhabited prefix of the computed address as an image and ranks within it by cosine; the Lean loop models exact lookup. The verifier therefore checks store membership and address provenance (the inhabitant's address, or the image as its prefix), and a *zwak* (weak) flag marks a top inhabitant with $\text{cosine} < 0.32$ to the question. The equivalence of §4 holds for exact addresses; on images it holds as implication.

6. Measurements

Twenty-question placement test (plan hashed before the run, sha256 8bc2f1fe...; threshold $\geq 14/20$ fixed in advance). Twenty open questions in Dutch and English with the expected inhabitant named beforehand. Result: 15/20, of which 2 weak; 5 misses. Reading of the misses, made after the run and changing nothing: three are the name bridge choosing a lemma that is too general ("Lithium" for lithium-ion battery; "Theory" for MOND); two are names that are not lemmas in the net ("Sardex", "Grokking") returned as a distant weak neighbour — where the honest output would have been a vacancy. The weak flag caught four of the five misses. Latency 2–11 s per question, dominated by the name bridge's external lookups.

A separate, failed engineering track, reported for completeness. Before the claim-bounded loop, a byte-level language model with a computed address per byte, prefix-routed experts and native ternary weights was trained (NumPy, CPU) under two hashed plans. Against a

conventional small model at 1.39 bits/byte, the computed-address model reached 5.54 (v0) and 2.96 (v2). Two of three pre-registered proofs passed to the letter but not in spirit (prefix routing equalled hash routing; the union criterion was trivially satisfiable). That track is not MAZE.AI; it is booked in [maze/ai/uitkomst.md](#) and kept as a candidate cheap host, not as the core. The archive-level union, measured on the net itself, is clean: of 4.08 million distinct addresses, 0.06% are shared by two independently built sources.

7. What is claimed, and what is not

- Claimed: every output of the deployed loop is of one of four types, carries an address, and satisfies a proven statement about the store; a proposal never leaves the loop as an assertion; settlement is an auditable event; the proofs are bound to the code by differential and production-path tests under one command.
- Not claimed: that answers are true; that placement is reliable (15/20 is a first measurement, not a guarantee); that the ledger is unforgeable against the key holder (signatures are HMAC with the admitting key; personal keys per human do not yet exist, so who is a name next to a key fingerprint); that sha256 is collision-free (assumed, stated).
- Novelty, carefully: proof assistants, retrieval with citations, transparency logs and human-in-the-loop review each exist. We know of no system in which they are combined into a single output path whose output type is claim-complete by construction and whose invariants are machine-checked against the running definitions and tested against production. We claim the combination, not the parts.

8. Relation to existing work

Retrieval-augmented generation (Lewis et al. 2020) and citation-conditioned answering (e.g. Nakano et al. 2021) attach sources to generated text but keep the generation as the answer. Verified software (seL4, Klein et al. 2009; CompCert, Leroy 2009) proves programs against specifications and is the methodological model for §4–5, with one difference of scope: we prove a small loop and a kernel, not a compiler, and we bind them to an unverified production language by differential testing rather than by extraction. Transparency logs (Laurie 2014, Certificate Transparency) are the model for the ledger: append-only, hash-chained, publicly checkable, with the meaning of an entry left to the parties. The MAZE's address-from-content principle descends from content-addressable storage, but the balanced-ternary kernel gives addresses a metric structure (prefix = coarser class) that a hash lacks, and the truncation theorems are what make "an image of the address" a well-defined notion. On the impossibility side, Tarski (1936) gives the boundary that motivates P.

9. Reproduction

```
git clone ... && cd swarp
node maze/ai/controle.mjs      # Lean rebuild + axiom report + differential test +
                                canon test
DATABASE_URL=... node maze/ai/controle.mjs # + production-path test on the
live store
node maze/ai/twintig-toets.mjs  # twenty-question test (plan hash must match)
curl -X POST https://maze.swarp.nl/api/ai/vraag -H 'content-type: application/json' -d
'{"vraag":"balanced ternary"}'
```

Lean sources: maze/lean/MazeAI/{Trit,Kern1,Brug,Keten,Lus,Winkel}.lean (≈ 740 lines, no Mathlib). Production loop: maze/ai/{canon,lus,verifier,route}.mjs. Ledger table ai_attest.

Annotated references

- Tarski, A. (1936). Der Wahrheitsbegriff in den formalisierten Sprachen. *Studia Philosophica* 1. Why read: the theorem that a sufficiently rich language cannot define its own truth predicate — the formal reason why "prove the model correct about the world" is not a programme.
- Klein, G. et al. (2009). seL4: formal verification of an OS kernel. *SOSP*. Why read: the reference case for proving a running system against its specification, and for being explicit about what the proof does not cover.
- Leroy, X. (2009). Formal verification of a realistic compiler. *CACM* 52(7). Why read: the model for tying a proof to executable code; we substitute differential testing for extraction and say so.
- Laurie, B. (2014). Certificate Transparency. *CACM* 57(10). Why read: append-only, hash-chained logs whose entries are facts about events, not judgements — the design our ledger follows.
- Lewis, P. et al. (2020). Retrieval-augmented generation for knowledge-intensive NLP tasks. *NeurIPS*. Why read: the standard way to attach sources to a language model, and the point where MAZE.AI departs — the retrieved content, not the generation, is the output.
- Nakano, R. et al. (2021). WebGPT: browser-assisted question-answering with human feedback. *arXiv:2112.09332*. Why read: citations chosen by the model, which is exactly what leaves the assertion type unchanged.
- de Moura, L., Ullrich, S. (2021). The Lean 4 theorem prover and programming language. *CADE*. Why read: the tool in which

the theorems of §4 are checked; statement and program are the same object.

- Knuth, D. E. (1997). The Art of Computer Programming, vol. 2, §4.1. Why read: balanced ternary, its range and uniqueness — the arithmetic of the address kernel.

- Konstapel, J. (2026). Provable AI. and maze.ai/1 — Specification. Why read: the criterion P, the boundary argument and the work programme (a–j) this paper reports on; every item is now either built or listed in §7 as open.