

ACCEPTED POSTER

THE 3RD AI FOR MATH WORKSHOP · ICML 2026

Typed Inference Dynamics

Auditable State Transitions for Mathematical Reasoning

Behnood Mashayekhi · Amesha · behnood@amesha.co

Seoul · COEX Hall D1 · 11 Jul 2026 · presented remotely

One claim.

Seven objections.

An honest verdict on each.

Don't take the claim on trust. Try to break it.

The rest of the talk is the seven objections a referee raises against typed inference dynamics — and where each one lands.

1	“Isn't this just structured prompting?”	PENDING
2	“Does it actually help anything?”	PENDING
3	“Aren't you just leaking the gold answer?”	PENDING
4	“Isn't it just a longer prompt?”	PENDING
5	“Isn't it just more compute?”	PENDING
6	“Is the object even real, or inference-time theatre?”	PENDING
7	“Does any of this generalize?”	PENDING

Watch the ledger, top-right. By the end: **three answered**, **two bounded**, **two conceded** — and the concessions are the point.

Reasoning is judged through objects that are too **flat**

FINAL ANSWER

Right / wrong

Says whether a path **ended** correctly — not which subpath to retire.

CHAIN-OF-THOUGHT

A line of prose

A fluent line of text — not a persistent field of **active and inactive candidates**.

PROCESS REWARD

A scalar score

Scores steps — but a number can't say what **state should survive** the next solve.

Between these lies a substrate: a **structured state whose transitions can be inspected, ablated, credited, and learned**. That object is what the rest of the talk defends.

We build reasoning to be **scrutinized**, not admired

Four commitments shape the whole project — and the third is why this talk invites its own objections.

1 The transition is the unit

Study **how state changes** — paths retired, constraints deposited, frames reframed — not just where a solve ends.

2 Everything is logged — including failure

A rejected operation is **evidence, not noise**. The object is auditable by construction, not by afterthought.

3 Boundaries are a feature

We report the nulls and reversals. A method that **always wins can't be trusted**; one with known regimes can.

4 One object, many uses

The same transition record serves **inference, audit, ablation, and training** — write it once, reuse it four ways.

One reasoning step becomes an **auditable, trainable object**

THE TRANSITION IS THE UNIT OF STUDY

$$(\mathcal{F}_t, \tau) \mapsto (e, r, \Delta\mathcal{F}) \quad \mathcal{F}_{t+1} = \mathcal{F}_t \oplus \Delta\mathcal{F}$$

A helper proposes typed ops over a persistent field; a **deterministic reducer** accepts, rejects, and logs them; the accepted state is rendered to a **frozen solver**; the outcome supplies credit.

State

$\mathcal{F}_t$

paths · constraints ·
subgoals · frame · values

Typed actions

$\Delta\mathcal{F}$

retire · suspend · reactivate ·
reframe · constrain

Reducer

accept / reject
local validity + audit log

Observation

render

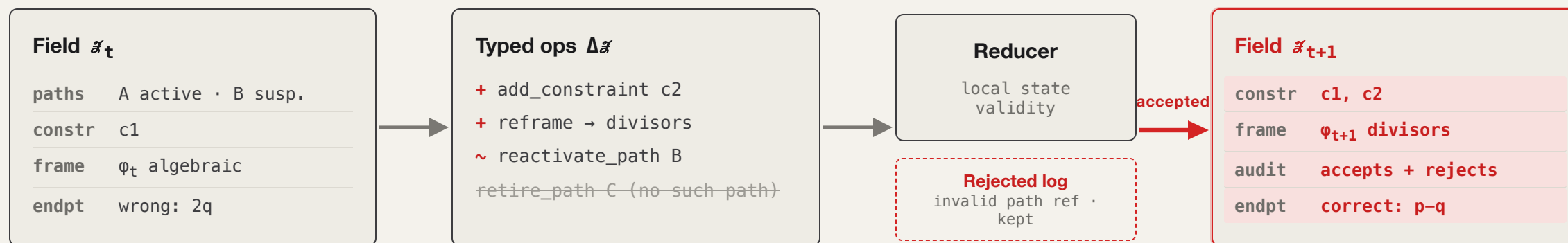
accepted field to a frozen
solver

Credit

outcome

solved / unsolved + trace

The helper proposes; the reducer **accepts, rejects, and logs**



The reducer checks **local form and state invariants, not mathematical truth**. A failed op is kept as evidence — this is the object every objection will now attack.

“Isn’t this just **structured prompting** with extra steps?”

- **Typed mutations** over paths · constraints · subgoals · frames — not free-form prose.
- **Deterministic reduction**: every op is accepted or rejected against invariants — not a scalar score.
- **Rejected ops are kept** as provenance. A prose rationale blurs a failed path into fluent text; the reducer preserves it.
- **Replayable & solver-facing**: the accepted state is the next prompt, and the whole trace can be replayed or ablated.

verdict —

ANSWERED

A structured prompt has **no memory** of what it rejected. This does.

“Fine — but does it **actually help** a solver?”

BENCHMARK	CARRIER	N	COUNT	LIFT
MATH-500 L5	7B	134	35→53	+13.4
MATH-500 L5	14B-AWQ	134	36→59	+17.2
OlympiadBench	14B-AWQ	29	4→11	+24.1
AIME 24–25	14B-AWQ	30	1→1	+0.0

Both carriers move positively on **full** MATH-500 L5 (n=134, not a cherry-picked subset). OlympiadBench transfers.

95% paired-bootstrap CIs exclude zero for MATH: 7B [+7.5, +19.4], 14B [+11.2, +23.9].

verdict — **ANSWERED** Real downstream force on frozen solvers — **and AIME is already flagged null**, which we’ll return to.

“You’re just **leaking the gold answer** into the prompt.”

+23.3

R2 gold-guard lift, pp
CI [+10.0, +40.0] — **excludes 0**

R2 adds a **direct gold-substring guard** to every operation string. The hard-subset lift **survives** at +23.3 pp.

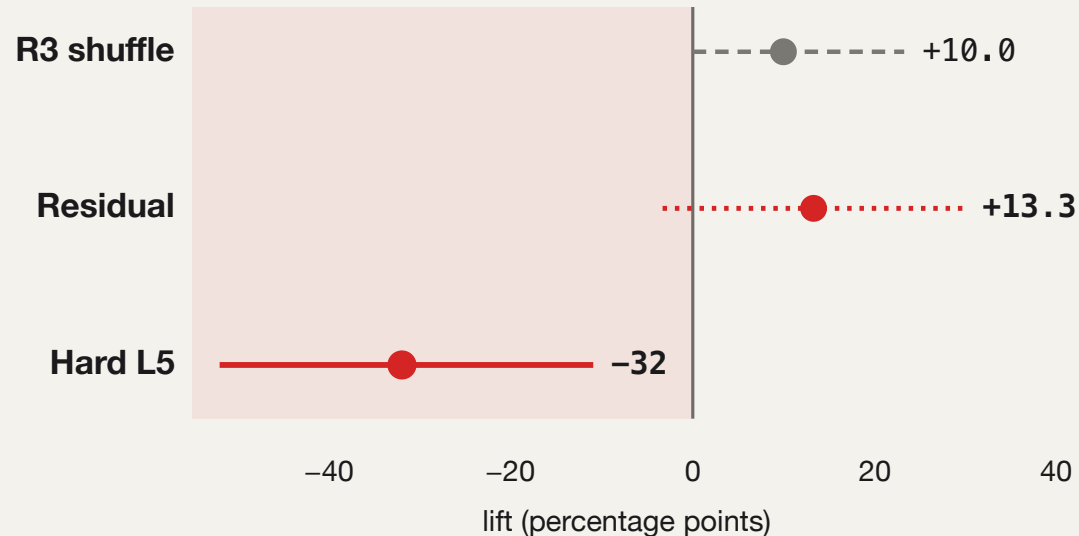
So verbatim answer leakage is **not load-bearing** in the guarded run.

verdict —

BOUNDED

Honest limit: R2 rules out **direct** leakage, not a theorem over **all** possible leakage. So — bounded, not closed.

“It’s just **more text** — any scaffold would help.”



R3 shuffle replays cross-item op sequences: scaffold alone gives **+10.0** — so structure **is** real (conceded).

But the content-specific **residual is +13.3**, CI crosses 0. And on a hard subset the sign **reverses to -32 pp**.

verdict —

BOUNDED

“More text always helps” predicts monotone gains. A **sign reversal** falsifies that — it’s an intervention, not padding.

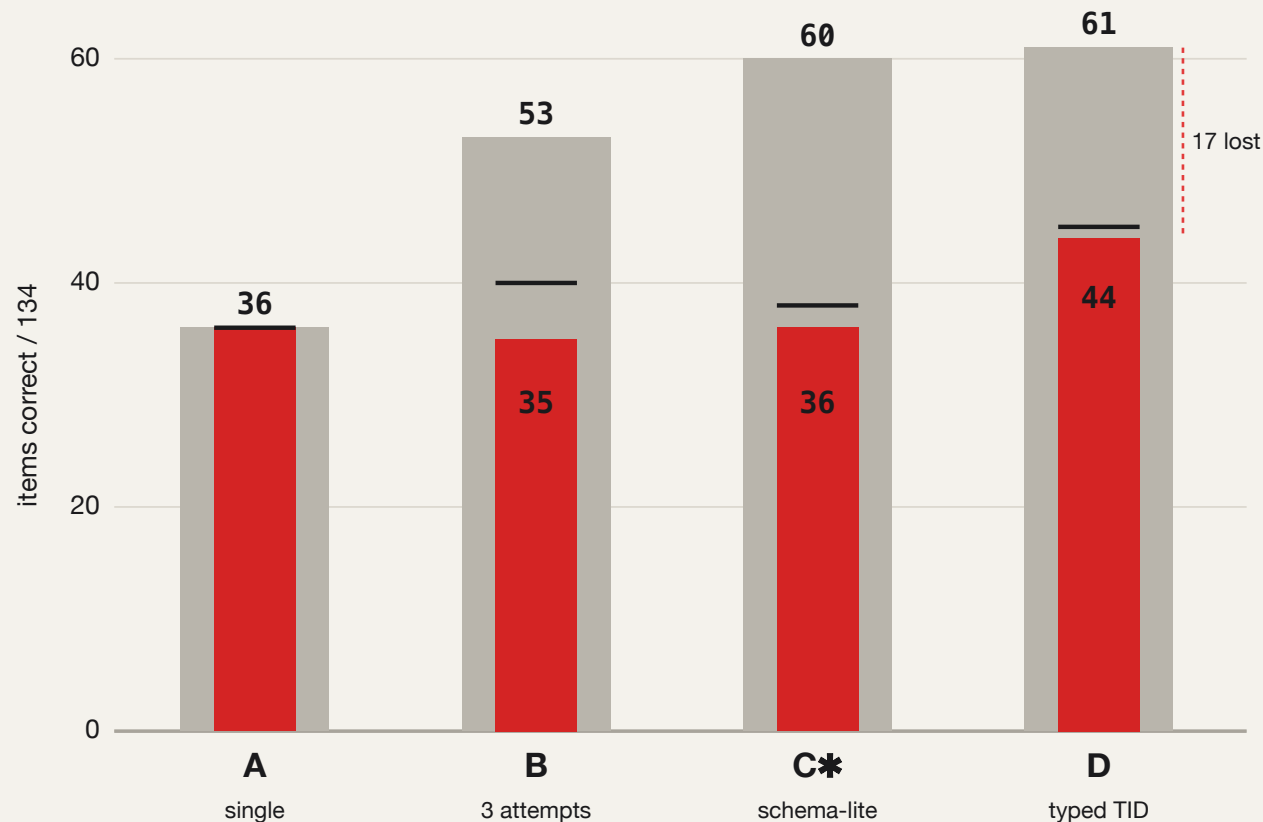
“You just spend **more compute** — more attempts, bigger prompts.”



Conceded: D costs **~2× B**. So we ran a **fresh fixed-horizon comparison** — same 134 items, same 14B family, **no gold-conditioned stopping** — to ask whether the coverage is just more tries.

The honest answer is on the next slide — and it's **not what we expected**.

D found 61. D kept 44.



Against schema-lite: **+0.8 pp, p=1.000. No net coverage gain.**

verdict — **CONCEDED**

But look **where** D differs: it discovers a correct answer on **61** items and only keeps 44. The gap is the finding.

any-turn discovery terminal retained frozen selector

The contribution isn't the lift — it's the **object** that made the gap visible

WHERE D'S 61 DISCOVERIES LAND AT THE TERMINAL TURN



■ persistent-correct ■ rescued **retained 44** ■ corrupted ■ dropped after T1 **lost 17**

THE WIN

16 rescued

items the frozen solver missed alone,
recovered by the typed state.

THE LEAK

17 lost

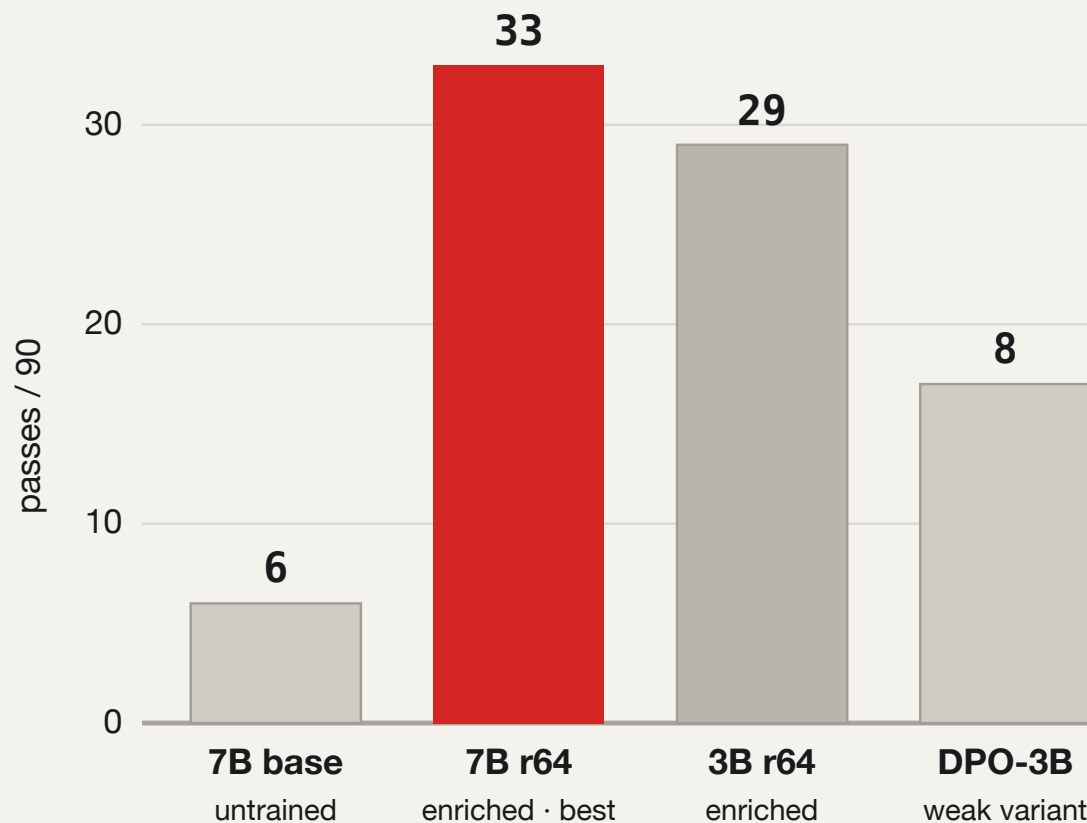
8 corrupted + 9 discovered then dropped
before the endpoint.

ONLY POSSIBLE BECAUSE

the trace is auditable

no flat transcript can tell you **where**
discovery isn't bankable.

“Is the object **real**, or just inference-time theatre?”



A **5,000-example** transition corpus lifts a 7B LoRA from **6** to **33 / 90** on a hard harness. The trace is a real training signal.

Weak variants stay low (DPO-3B **8**) — so it is **not** “any fine-tuning helps.”

verdict — **ANSWERED** The same transition object used for inference and audit is **distillable** — it carries learnable structure.

“Does any of this generalize?”

AIME C3

Exactly **null** at the registered T2 endpoint: **1→1** . No spin.

verdict —

CONCEDED

Hard stratum

Directionality **reverses**: **-32 pp** on a hard L5 subset.

Carrier scope

All headline carriers are one **Qwen2.5** family. Non-Qwen is future work.

Proof corpus

200 trajectories, **local SymPy** checks — **not** Lean / Coq / Isabelle validity.

Bounded and diagnostic. **The nulls and reversals are what make it falsifiable** — a real intervention has regimes.

Seven objections. Here's where they **actually** land.

3 ANSWERED

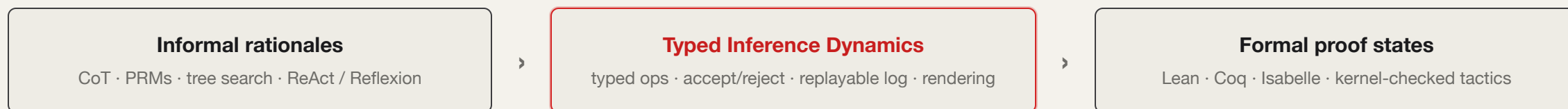
2 BOUNDED

2 CONCEDED

1	Structured prompting?	ANSWERED
2	Does it help?	ANSWERED
3	Answer leakage?	BOUNDED
4	Just a longer prompt?	BOUNDED
5	Just more compute?	CONCEDED
6	Is the object real?	ANSWERED
7	Does it generalize?	CONCEDED

A weaker talk shows you only the three that survive. **This one shows all seven** — which is exactly why you can **trust the three.**

A **pre-formal** layer — after prose, before the proof kernel



A failed route can deposit a constraint, a frame can change, a subgoal can open, and malformed state motion is rejected — **before** a problem is ever translated into a proof language. The state, action, transition, and credit interface is exactly what later credit assignment or RL would optimize.

TAKEAWAY

Typed inference dynamics makes a reasoning step an explicit, **auditable and trainable** object — one that survives cross-examination precisely because it **reports its own boundaries**.

Helpful in some regimes, null or harmful in others — always an auditable record. **Retention and state-authority calibration are the next bottlenecks.**

Behnood Mashayekhi · Amesha

behnood@amesha.co

The 3rd AI for Math Workshop · ICML 2026

