

Formal Reasoning in AI

From next-step prediction to verifiable reasoning agents

Peiyang Song

PhD Student @ CMU LTI

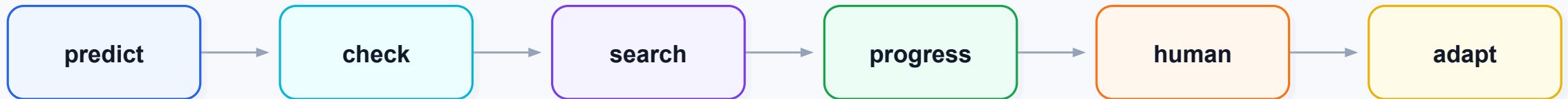


ECE 4424 / CS 4824 · Machine Learning
Virginia Tech · Sept. 14, 2026

Reliable reasoning is a system problem

Not just a better next-token predictor.

Strong reasoning emerges when learning is placed inside a loop with exact feedback, search, human control, and adaptation.

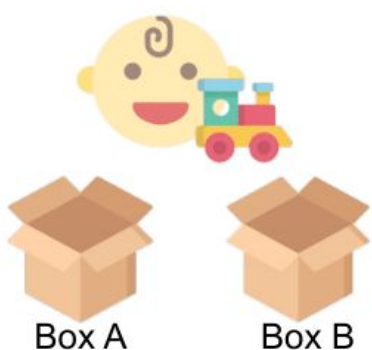


Today's path: failure → formal feedback → neuro-symbolic systems → human–AI workflows → adaptive agents → verified software.

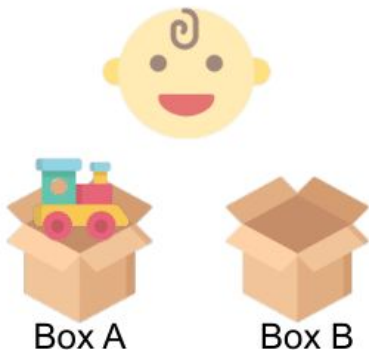
A failure that feels almost human

A-Not-B errors: a model keeps following an old, trivial pattern after the context changes.

Repeat Several Times



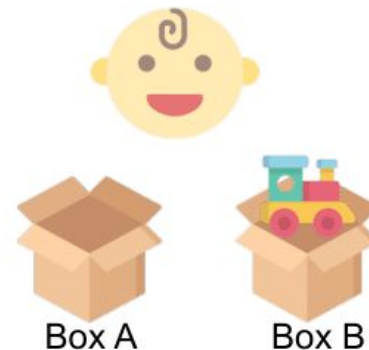
1. Object in View



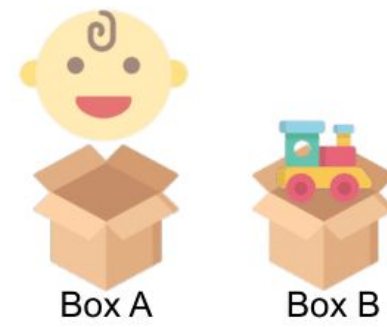
2. Put the Object in Box A



3. Infant Finds the Object



4. Put the Object in Box B



5. Infant Searches for Object in Box A



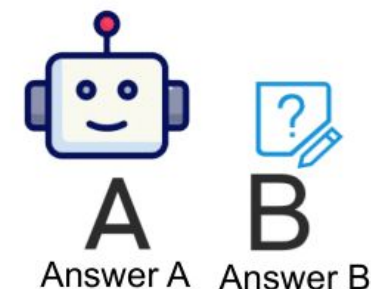
1. Question in View



2. Show Example Questions Where Consistently A Is the Right Answer



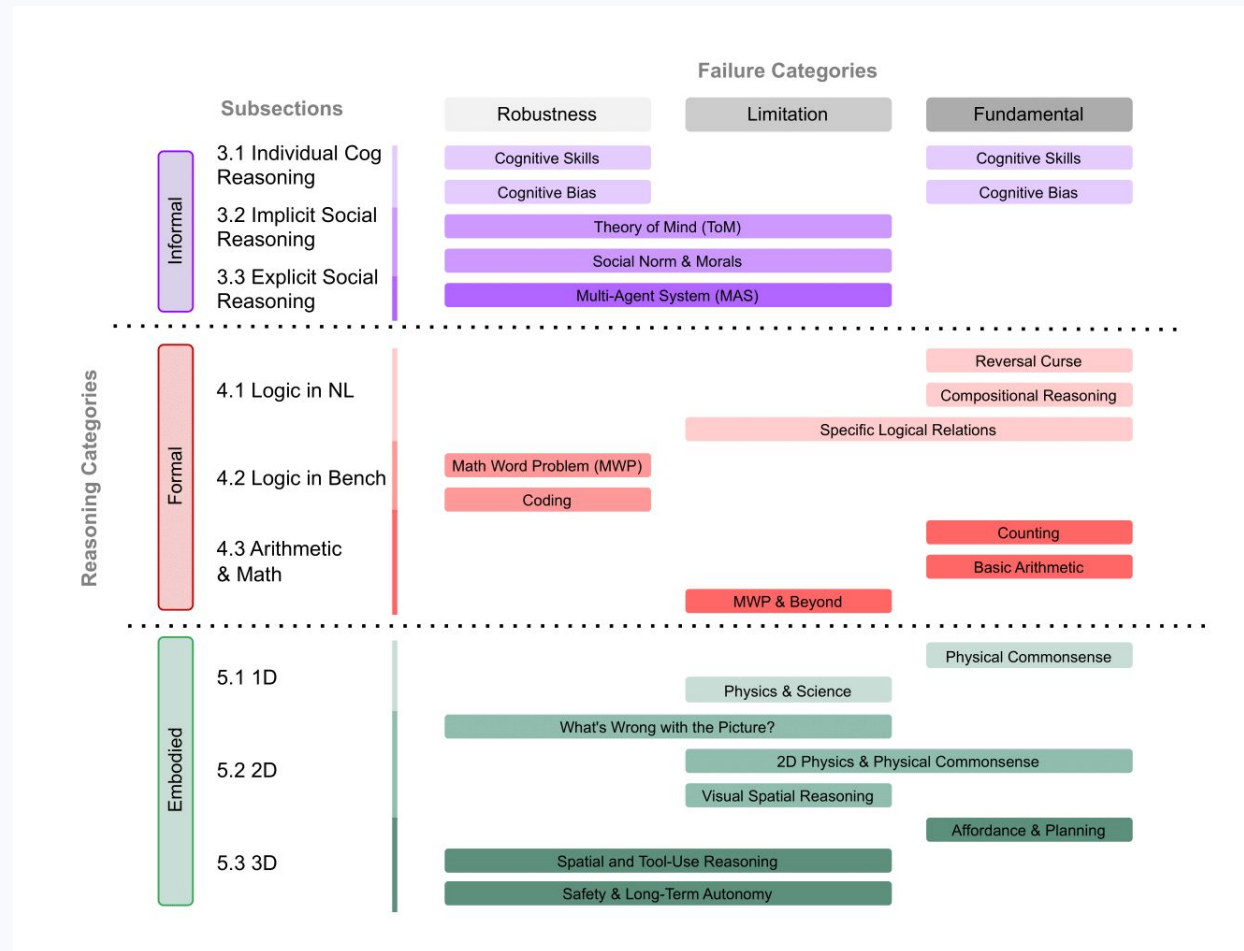
3. Give A Question with Right Answer B



4. LLM Chooses A

Reasoning failures are not isolated bugs

They recur across intuitive, logical, mathematical, and embodied settings.



How to fix/alleviate this?

If natural-language reasoning is hard to judge reliably, can we build settings where correctness is machine-checkable?

The neuro-symbolic promise

Let neural models generate; let symbolic tools verify.



Natural collaboration

LLMs are strong generators: guesses, tactics, plans, code.

Formal tools are strong judges: parse, type-check, verify.

Not “neural vs. symbolic”: the point is the interface.

What Lean gives the loop

A proof assistant turns reasoning into exact state transitions.

```
example (P Q : Prop)
  (hp : P)
  (hpq : P → Q) : Q := by
  apply hpq
  assumption
```

context
hp : P
hpq : P → Q

goal
⊢ Q

tactic
apply hpq

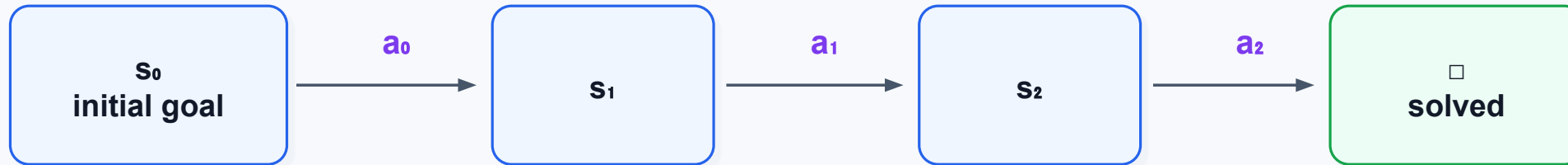
new goal
⊢ P

$$s_{i+1} = \text{TLean}(s_i, a_i)$$

The model proposes an action; Lean either accepts it and returns a new state, or rejects it with rich feedback regarding e.g. type mismatches.

Reasoning as sequential decision making

The proof becomes a trajectory through states.



$$a_{\square} \sim \pi_{\theta}(a \mid s_{\square})$$

$$s_{\square+1} = \text{TLean}(s_{\square}, a_{\square})$$

success $\Leftrightarrow$ no goals
left

Learning the next proof step

Classification: proof state in, tactic/action distribution out.

Training data

```
state: hp : P, hpq : P → Q ⊢ Q  
label: apply hpq
```

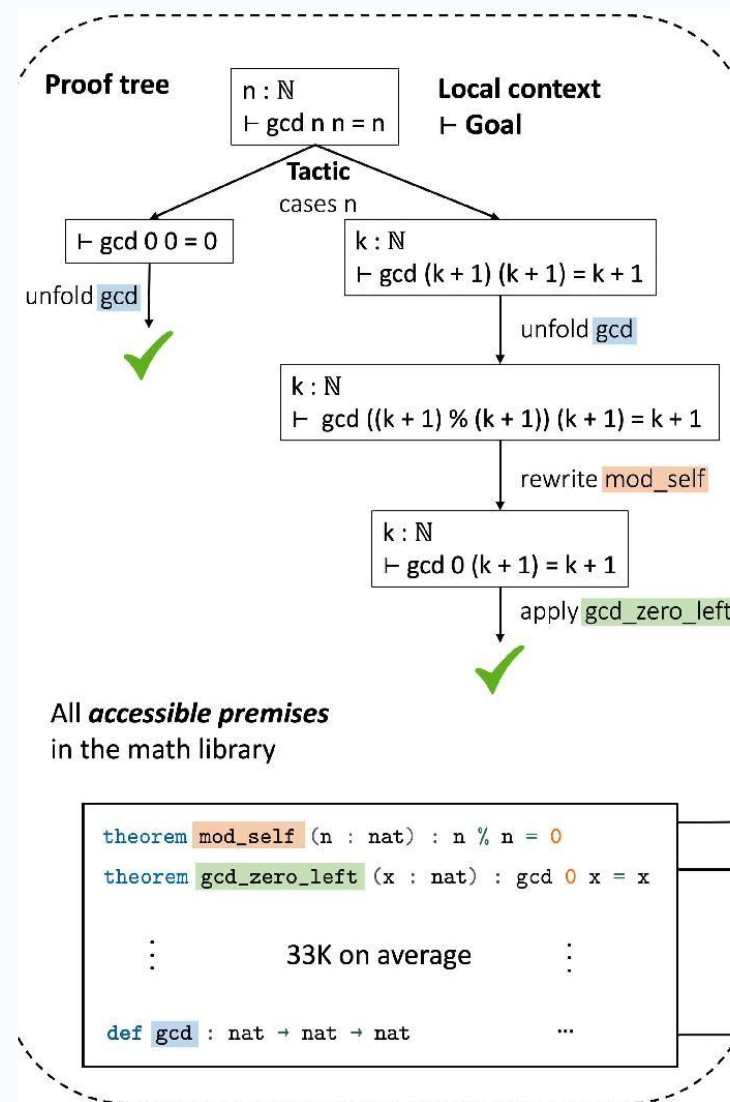
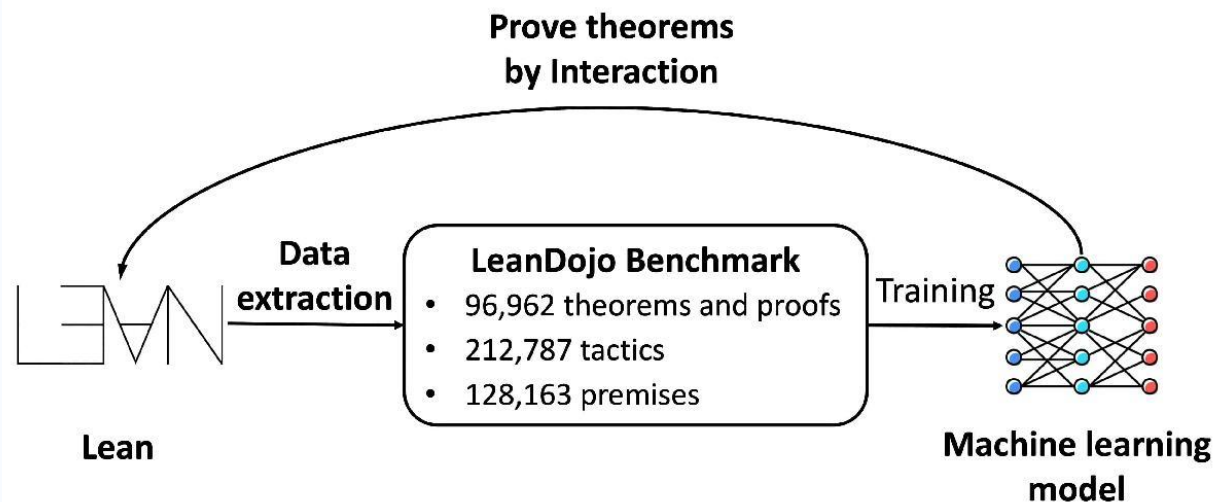
```
state: hp : P ⊢ P  
label: assumption
```

$$\pi_{\theta}(a \mid s) = P_{\theta}(\text{next action} = a \mid \text{proof state} = s)$$

$$L(\theta) = - \sum_i \log \pi_{\theta}(a_i \mid s_i)$$

LeanDojo: opening the Lean $\leftrightarrow$ ML loop

Infrastructure for data extraction, programmatic interaction, and reproducible theorem-proving research.

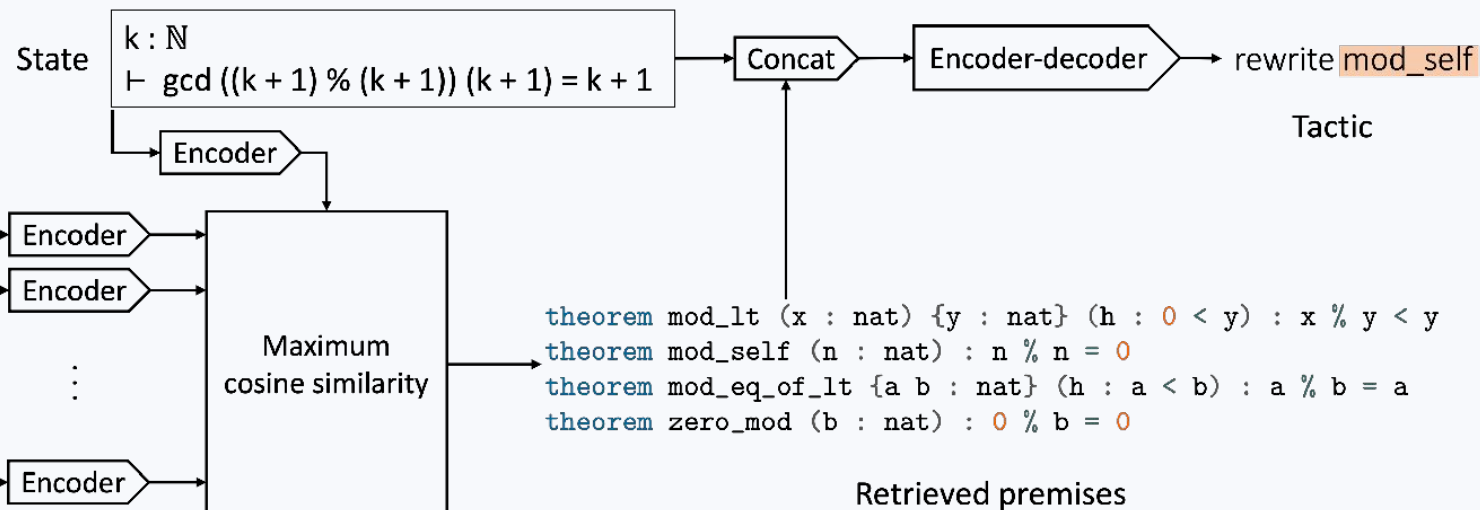


ReProver: retrieval makes large libraries usable

A prover must choose from a vast premise library before generating a useful tactic.

All *accessible premises*
in the math library

```
theorem mod_self (n : nat) : n % n = 0
theorem gcd_zero_left (x : nat) : gcd 0 x = x
...
33K on average
...
def gcd : nat → nat → nat
```



Prediction alone is not reasoning

A legal, likely step can move the proof away from success.

A tiny theorem

$$R, Q_3 \rightarrow R, Q_2 \rightarrow Q_3, Q_1 \rightarrow Q_2 \vdash R$$

Trivial proof:

assumption

Tempting but bad path:

$$R \rightarrow Q_3 \rightarrow Q_2 \rightarrow Q_1$$

greedy policy
chooses apply

Lean accepts
the tactic

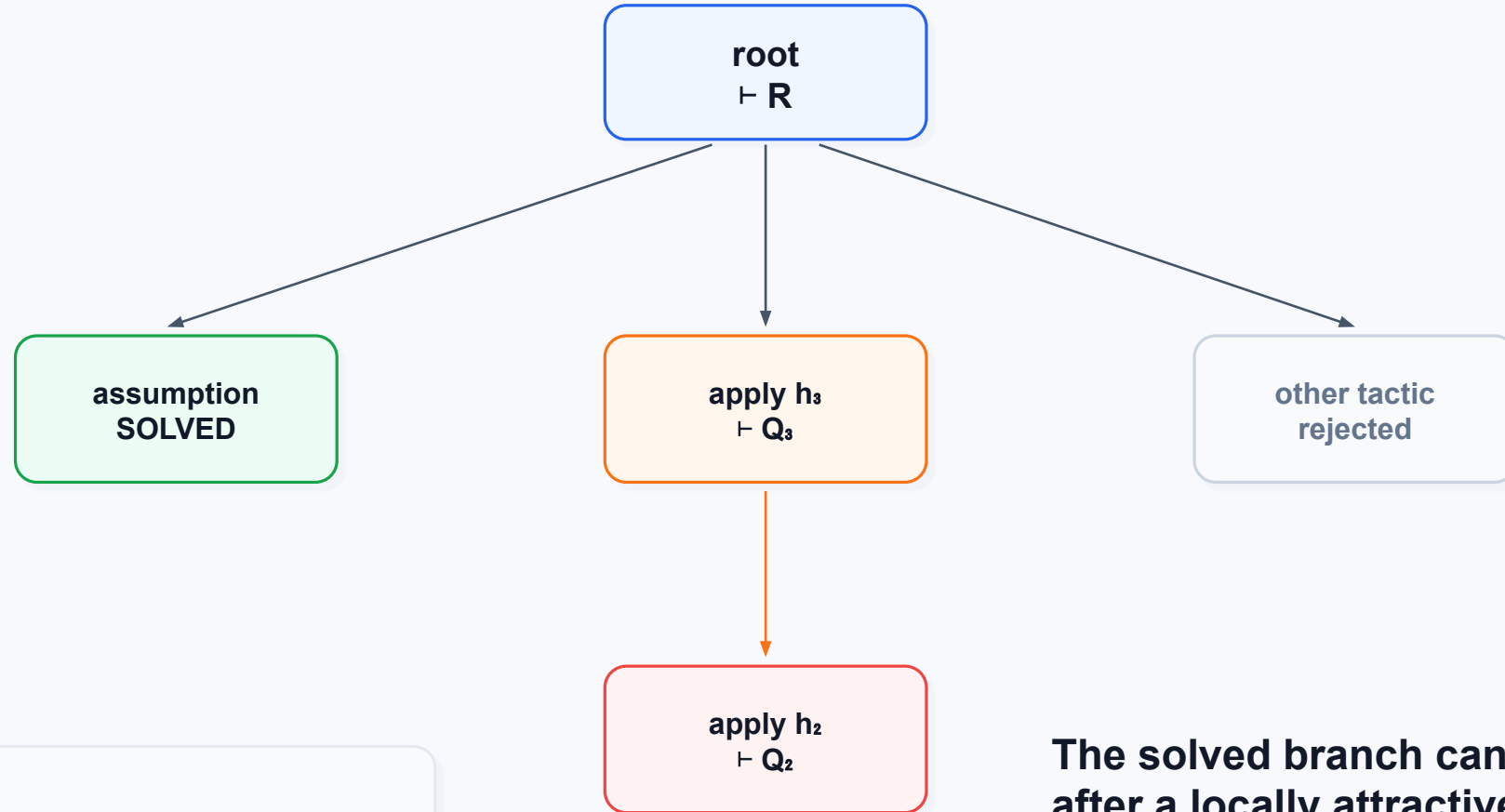
new subgoal
 Q_3

eventually
stuck at Q_1

Verification gives local validity; search
gives strategic recovery.

Search + verifier: generate, check, recover

Keep alternatives alive instead of committing to the first accepted step.



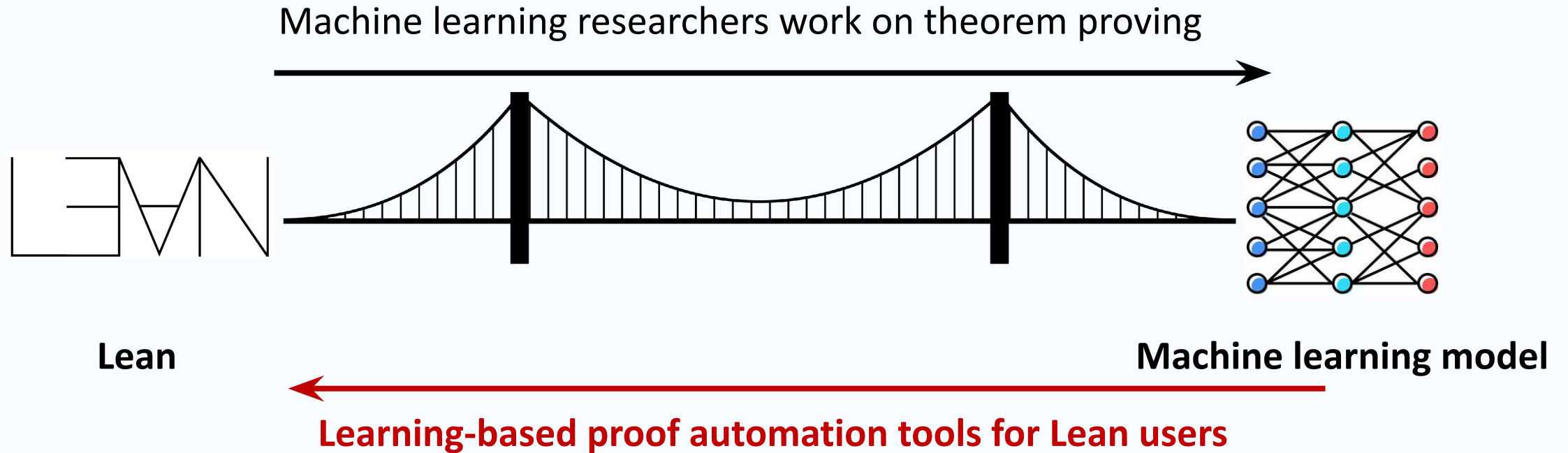
Best-first score

$$\Sigma \log \pi_{\theta}(a_{\square} \mid s_{\square})$$

The solved branch can still be found after a locally attractive detour.

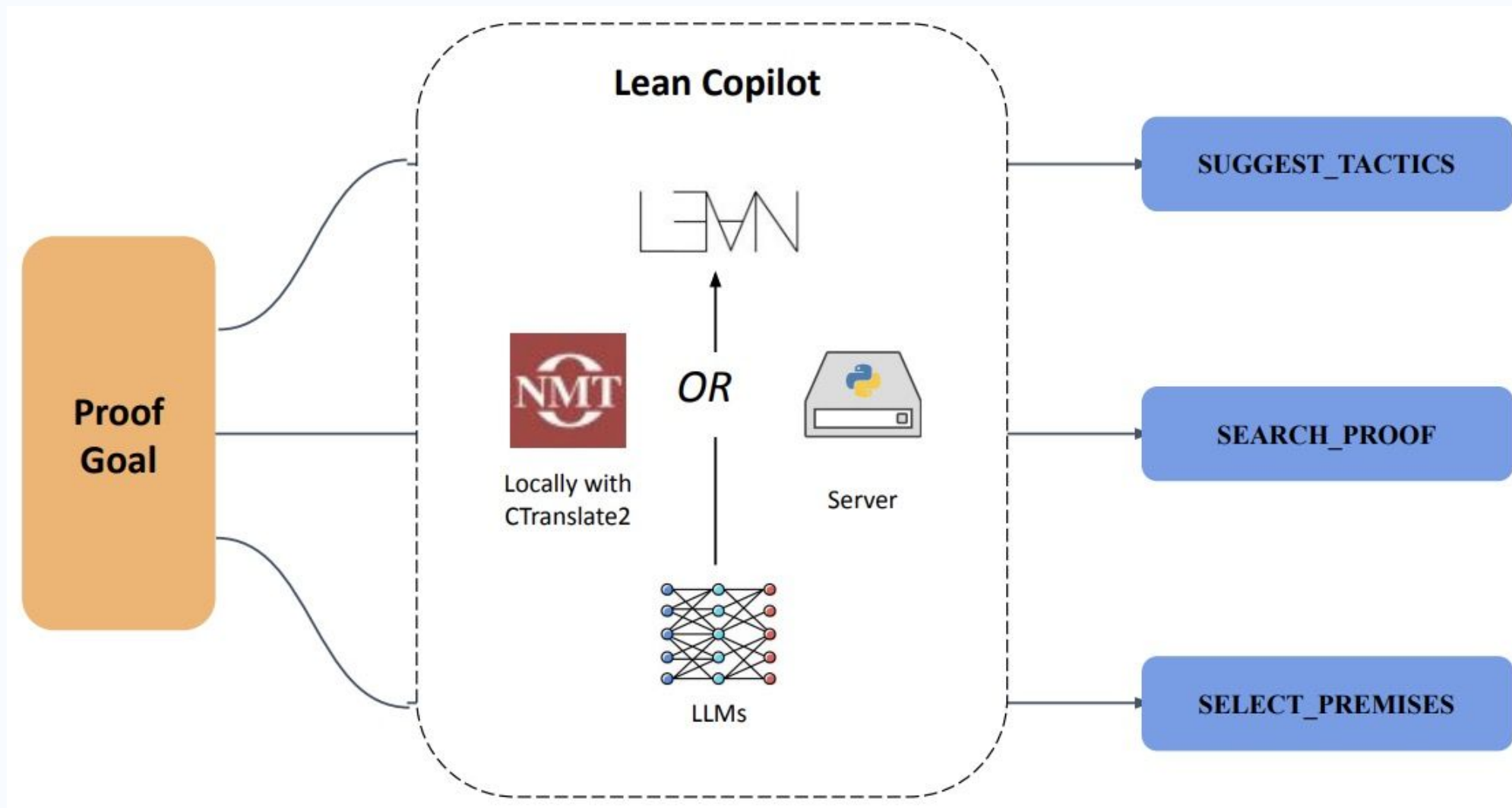
Lean Copilot: AI inside the proof workflow

A shift from autonomous theorem proving to interactive assistance.



Lean Copilot: AI inside the proof workflow

A shift from autonomous theorem proving to interactive assistance.



a

Tool Usage	Accuracy
Tools	~0.72
No Tools	~0.38

b

Tool Usage	Time (min)
Tools	~128
No Tools	~118

c

Participant	Inline LLM	Chatbot	Semantic Search	LLM (Other)	tools/person
1	3	0	0	0	9
2	3	0	0	0	4
3	3	0	0	0	3
4	3	0	0	0	3
5	3	0	0	0	3
6	3	0	0	0	3
7	3	0	0	0	3
8	3	0	0	0	3
9	3	0	0	0	3
10	3	0	0	0	3
11	3	0	0	0	3
12	3	0	0	0	3
13	3	0	0	0	3
14	3	0	0	0	3
15	3	0	0	0	3
16	3	0	0	0	3
17	3	0	0	0	3
18	3	0	0	0	3
19	3	0	0	0	3
20	3	0	0	0	3
21	3	0	0	0	3
22	3	0	0	0	3
23	3	0	0	0	3
24	3	0	0	0	3
25	3	0	0	0	3
26	3	0	0	0	3
27	3	0	0	0	3
28	3	0	0	0	3
29	3	0	0	0	3
30	3	0	0	0	3
31	3	0	0	0	3
32	3	0	0	0	3
33	3	0	0	0	3
34	3	0	0	0	3
35	3	0	0	0	3
36	3	0	0	0	3
37	3	0	0	0	3
38	3	0	0	0	3
39	3	0	0	0	3
40	3	0	0	0	3
41	3	0	0	0	3
42	3	0	0	0	3
43	3	0	0	0	3
44	3	0	0	0	3
45	3	0	0	0	3
46	3	0	0	0	3
47	3	0	0	0	3
48	3	0	0	0	3
49	3	0	0	0	3
50	3	0	0	0	3
51	3	0	0	0	3
52	3	0	0	0	3
53	3	0	0	0	3
54	3	0	0	0	3
55	3	0	0	0	3
56	3	0	0	0	3
57	3	0	0	0	3
58	3	0	0	0	3
59	3	0	0	0	3
60	3	0	0	0	3
61	3	0	0	0	3
62	3	0	0	0	3
63	3	0	0	0	3
64	3	0	0	0	3
65	3	0	0	0	3
66	3	0	0	0	3
67	3	0	0	0	3
68	3	0	0	0	3
69	3	0	0	0	3
70	3	0	0	0	3
71	3	0	0	0	3
72	3	0	0	0	3
73	3	0	0	0	3
74	3	0			



Progress, not just probability

A local policy needs a global sense of distance-to-goal.

Policy question

$$\pi_{\theta}(a \mid s)$$

What action looks likely next?

Progress question

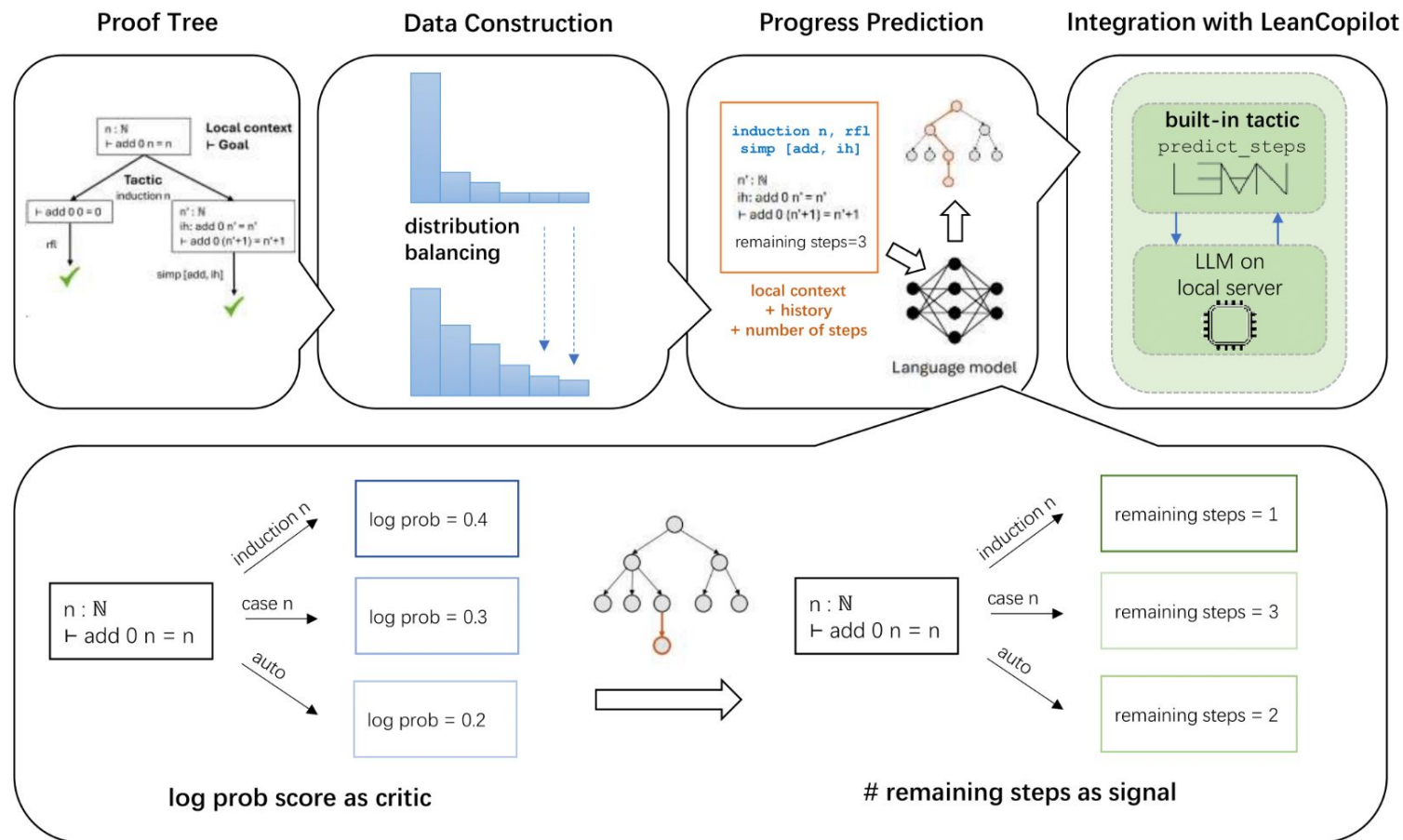
$$\tilde{d}\phi(s) \approx \text{steps remaining}$$

Where will this action leave me?

$$\text{priority}(s) = \sum \log \pi_{\theta}(a_{\square} \mid s_{\square}) - \lambda \tilde{d}\phi(s)$$

LeanProgress: a critic for proof search

Estimate remaining steps, then use the estimate to prioritize search.



The library keeps growing

Real mathematics is not a fixed train/test split.

Static assumption

$$D_{\text{train}} \approx D_{\text{test}}$$

A fixed model is evaluated after the world stops changing.

Lifelong setting

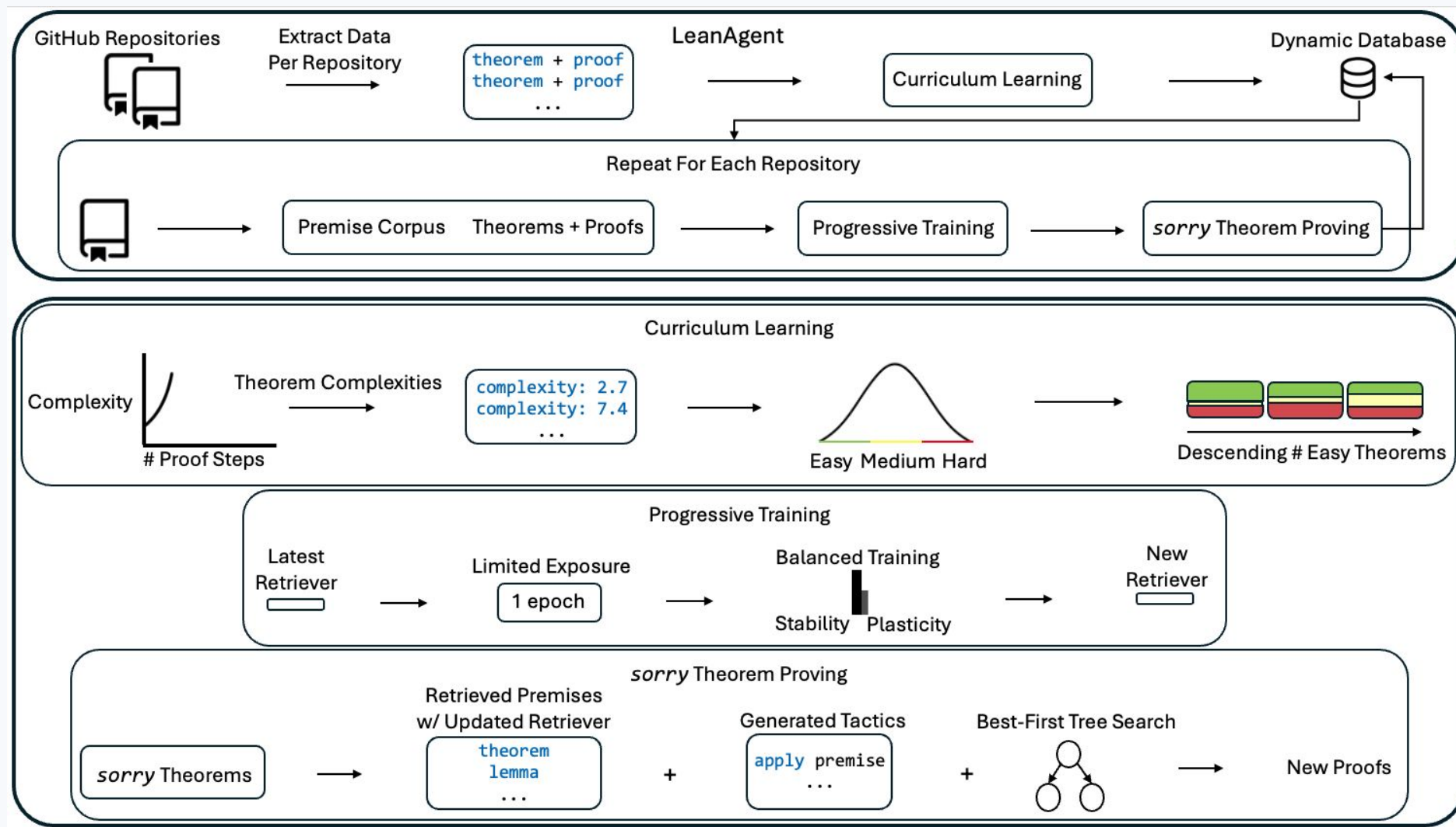
$$D_1 \rightarrow D_2 \rightarrow \dots \rightarrow D_\infty$$

New repositories, definitions, theorems, proof styles, and dependencies keep arriving.

Challenge: learn new knowledge without forgetting old knowledge.

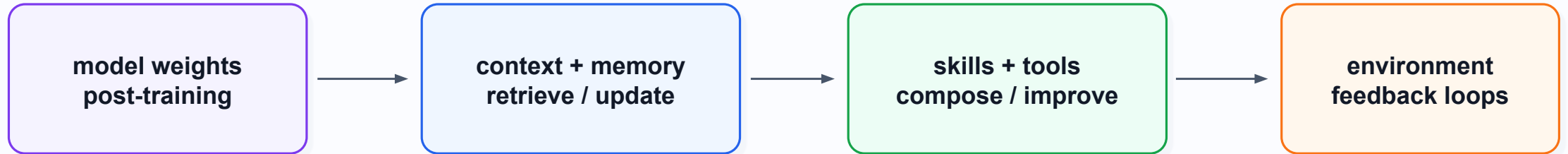
LeanAgent: lifelong theorem proving

A theorem prover that adapts as mathematical knowledge expands.



Adaptation is a general agent problem

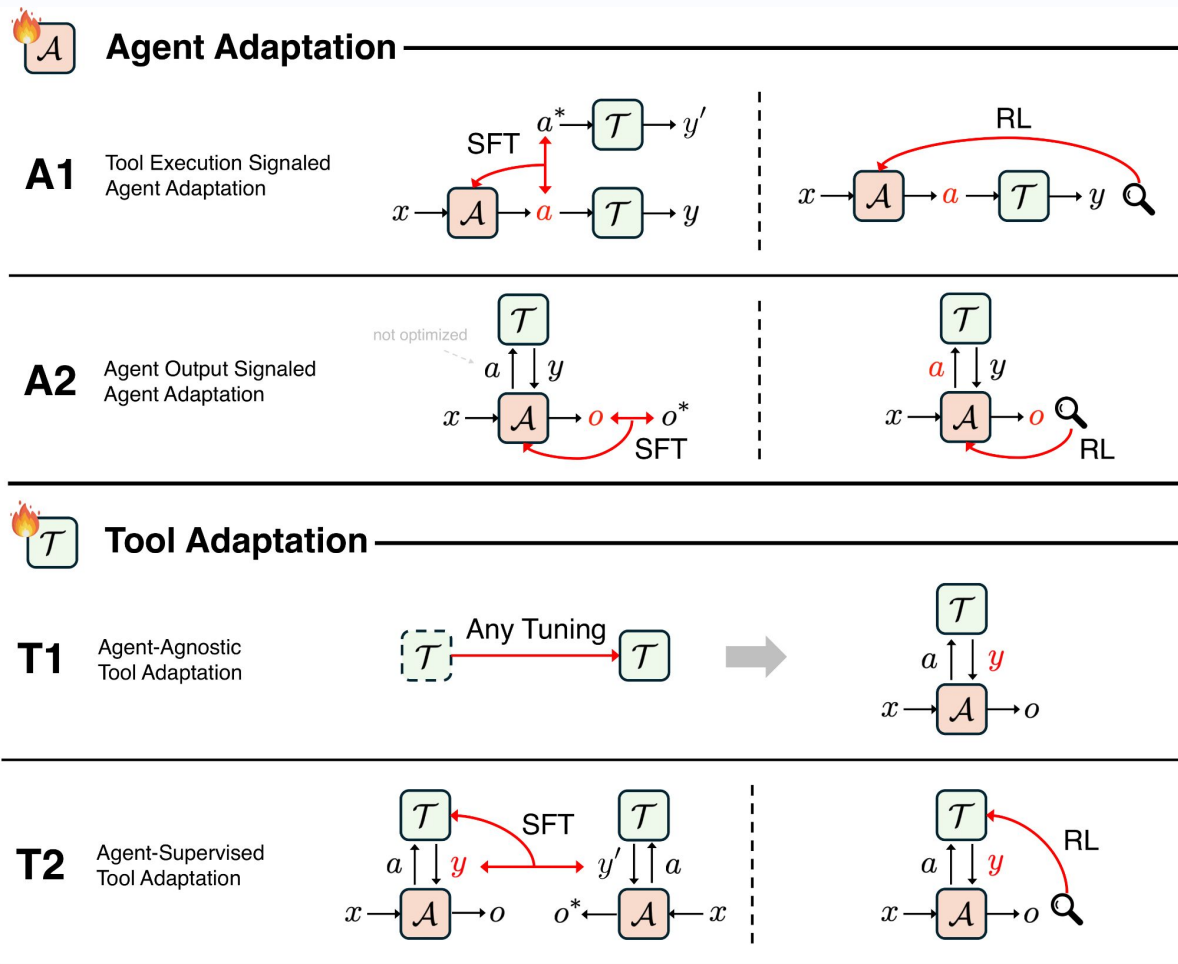
Where should a reasoning system change when tasks and context grow?



LeanProgress adapts search behavior with a progress signal; LeanAgent adapts to new libraries. The same design space appears in agentic AI broadly.

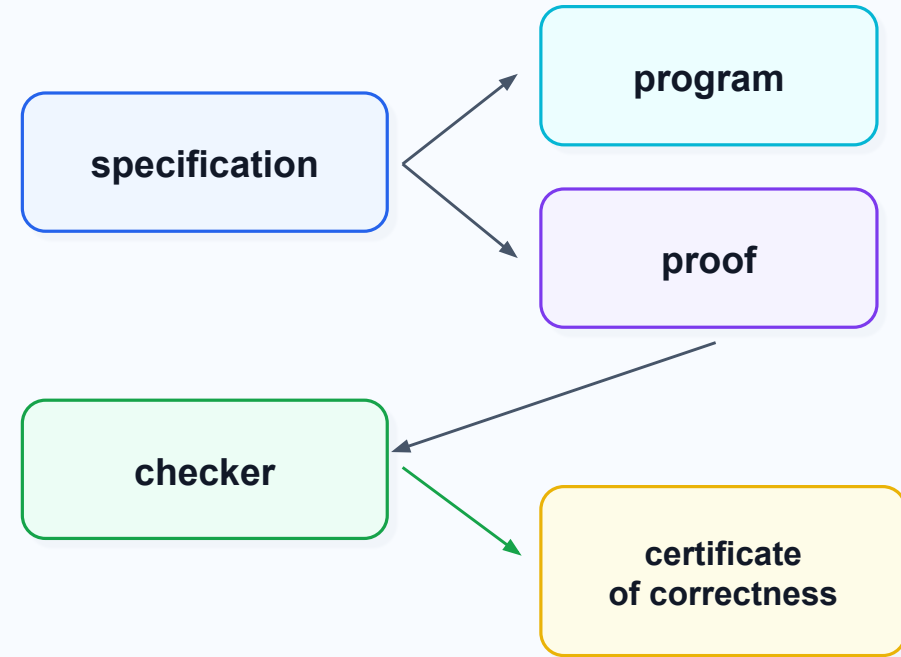
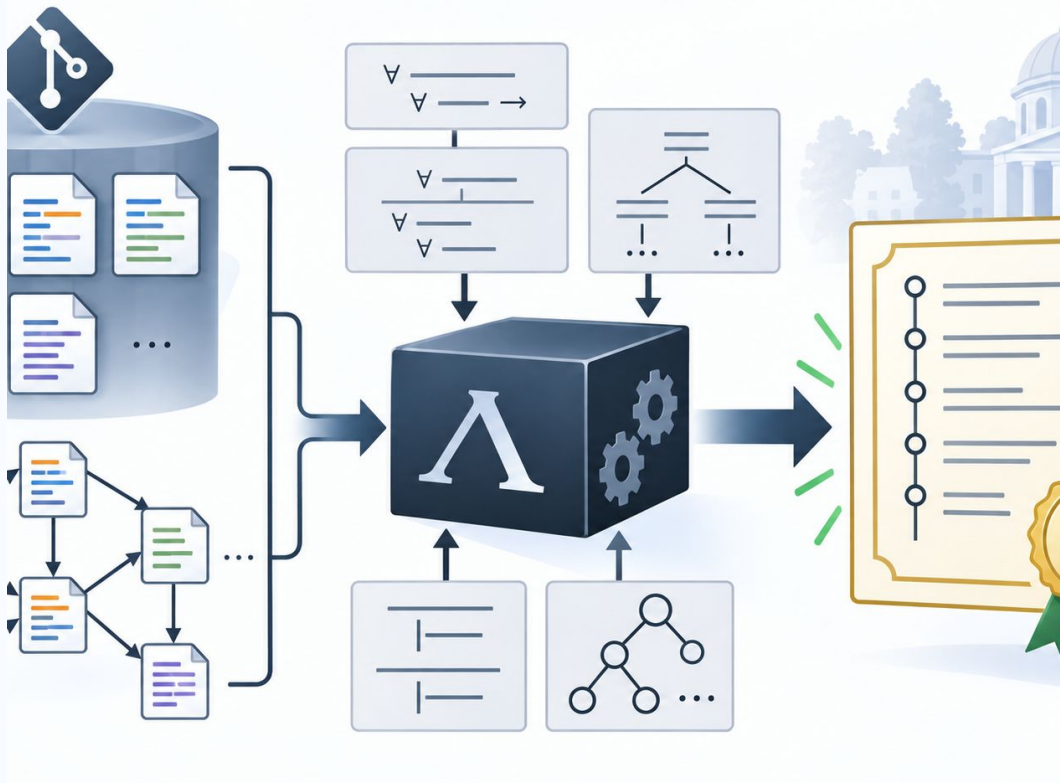
Adaptation is a general agent problem

Where should a reasoning system change when tasks and context grow?



From math proofs to verified software

The same loop applies when the artifact is code plus a proof.



Stronger guarantee: not just “the generated code passes tests,” but “the code satisfies a formal specification.”

Vero: repository-scale verified software

Can agents make coherent implementation-and-proof choices across modules?

43

repository-level
instances

743

APIs

2,705

formal
specifications

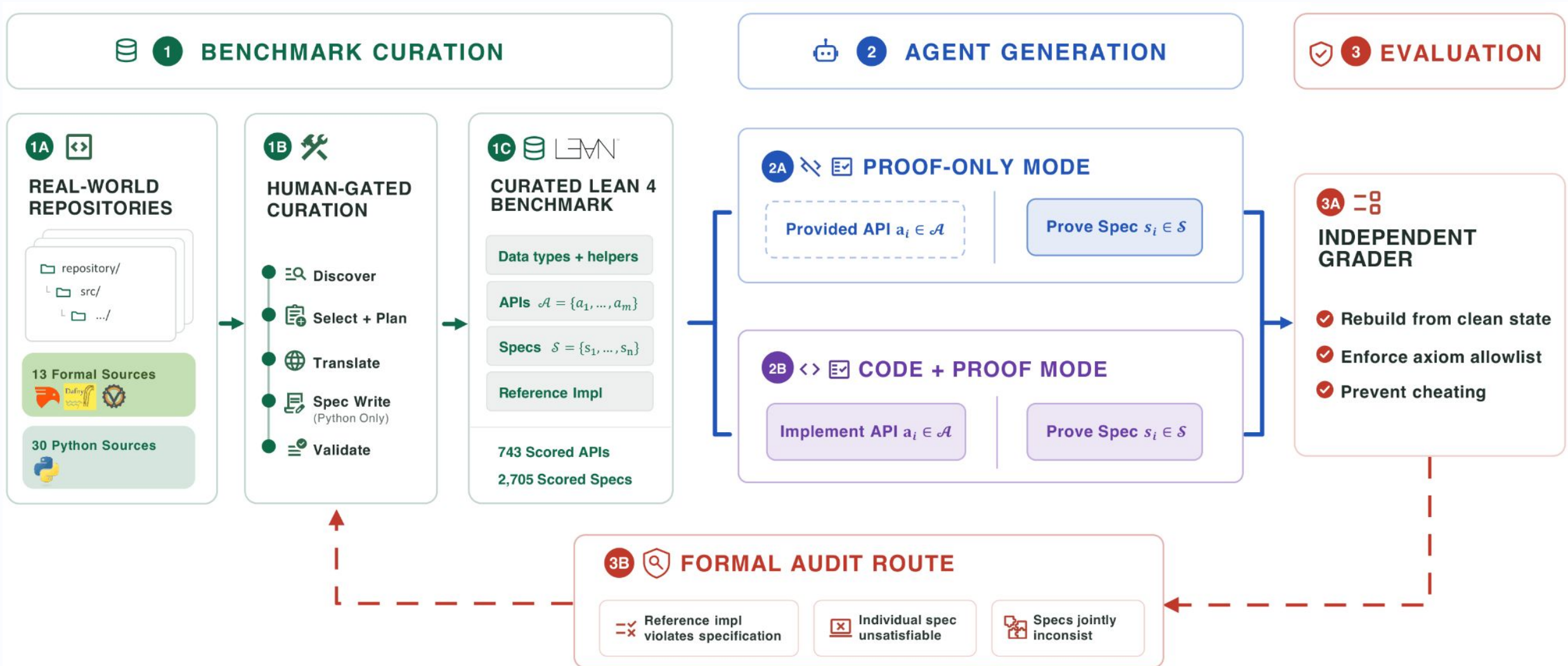
27/43

best full
solutions

Key shift: from single functions to multi-module repositories with interacting APIs, specs, and dependencies.

Vero: repository-scale verified software

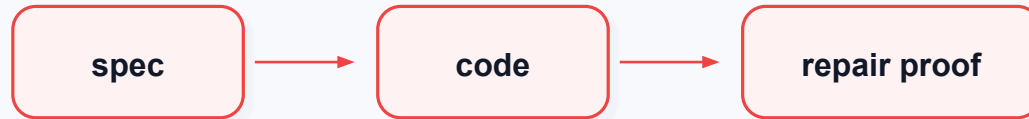
Can agents make coherent implementation-and-proof choices across modules?



P³: plan program and proof together

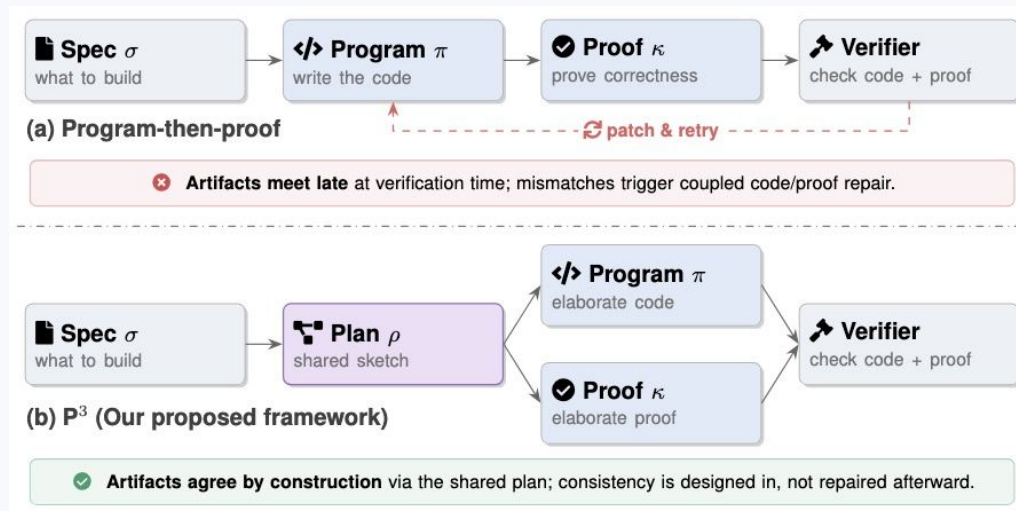
Implementation choices can make the later proof easy—or impossible.

Sequential baseline



Code-first decisions can lock in proof-hostile structure.

P³ workflow



Core lesson

Verified code generation is not two independent problems. The program and its correctness argument should be designed hand in hand.

P³ = Joint Program-and-Proof Planning

Where else can formal reasoning lead?

Anywhere we can specify what we want and check evidence that we got it.

**formal
math**

**verified
software**

**hardware
& protocols**

**scientific
computing**

**agent
evaluation**

Open question

Can AI explore creatively while correctness is enforced by something stronger than model confidence?

Now build the loop

Notebook walkthrough

Peiyang Song

PhD Student @ CMU LTI

psong2@andrew.cmu.edu



ECE 4424 / CS 4824 · Machine Learning
Virginia Tech · Sept. 14, 2026