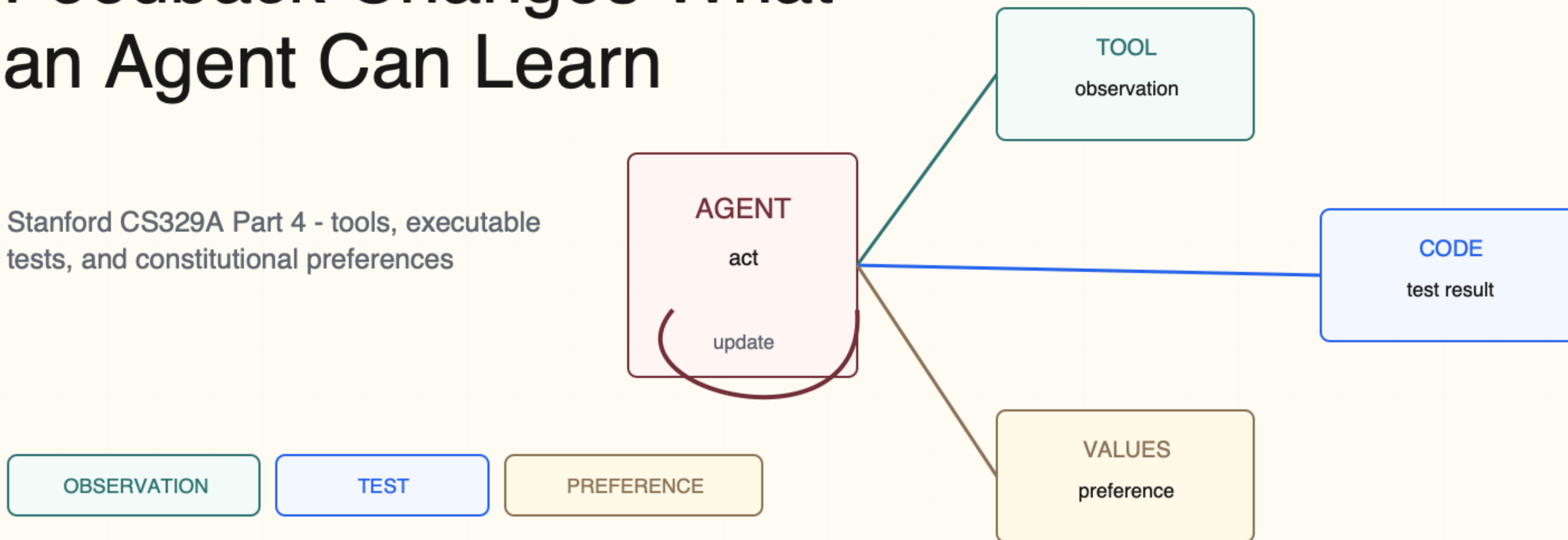


Feedback Changes What an Agent Can Learn

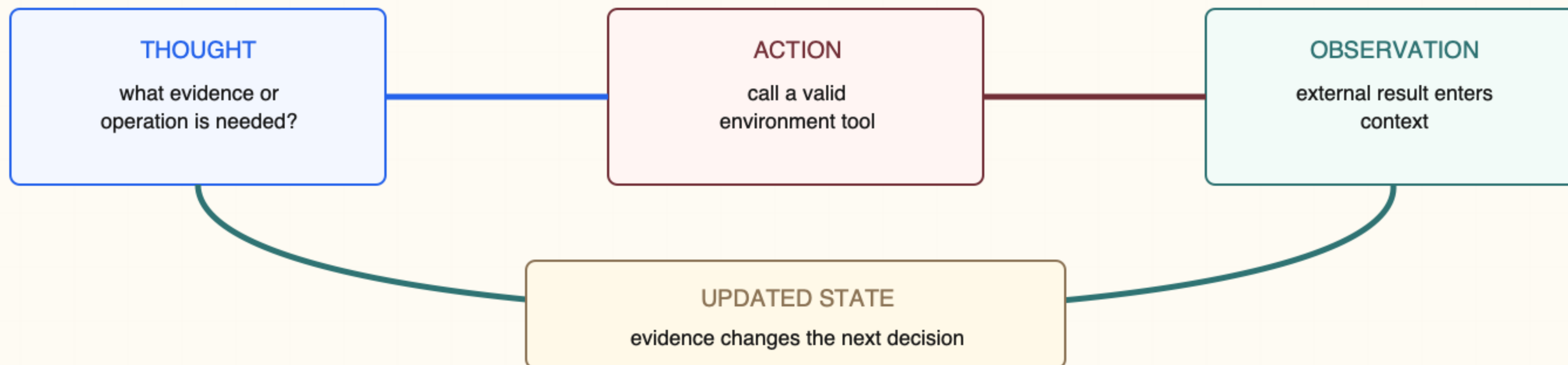
Stanford CS329A Part 4 - tools, executable tests, and constitutional preferences



three signals != one meaning

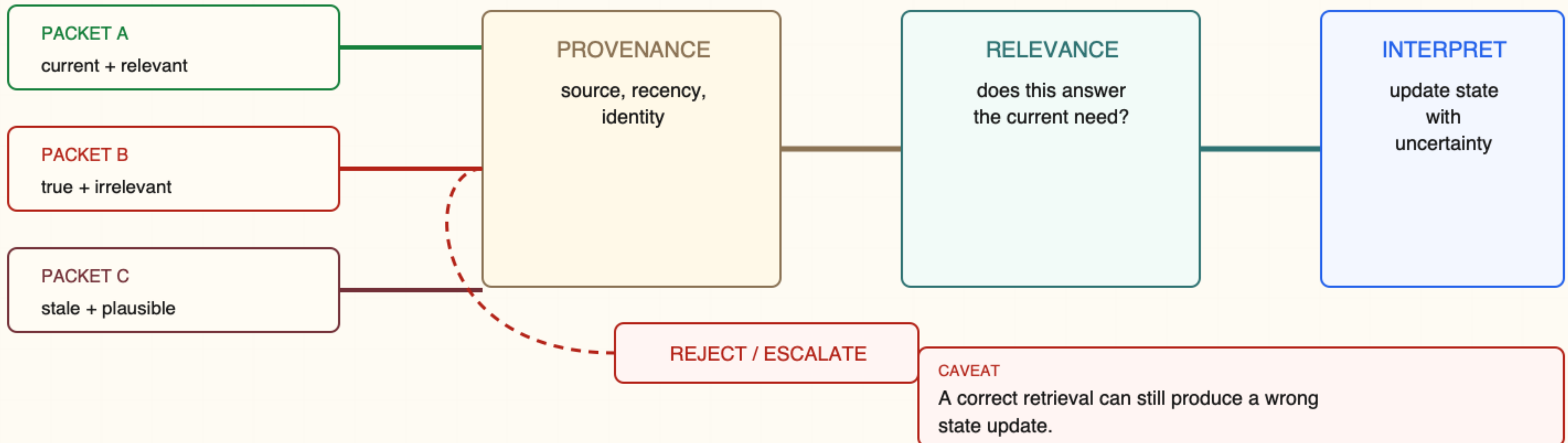
ReAct Lets Evidence Change the Next Decision

Thought selects an action; the environment returns an observation; the state updates.



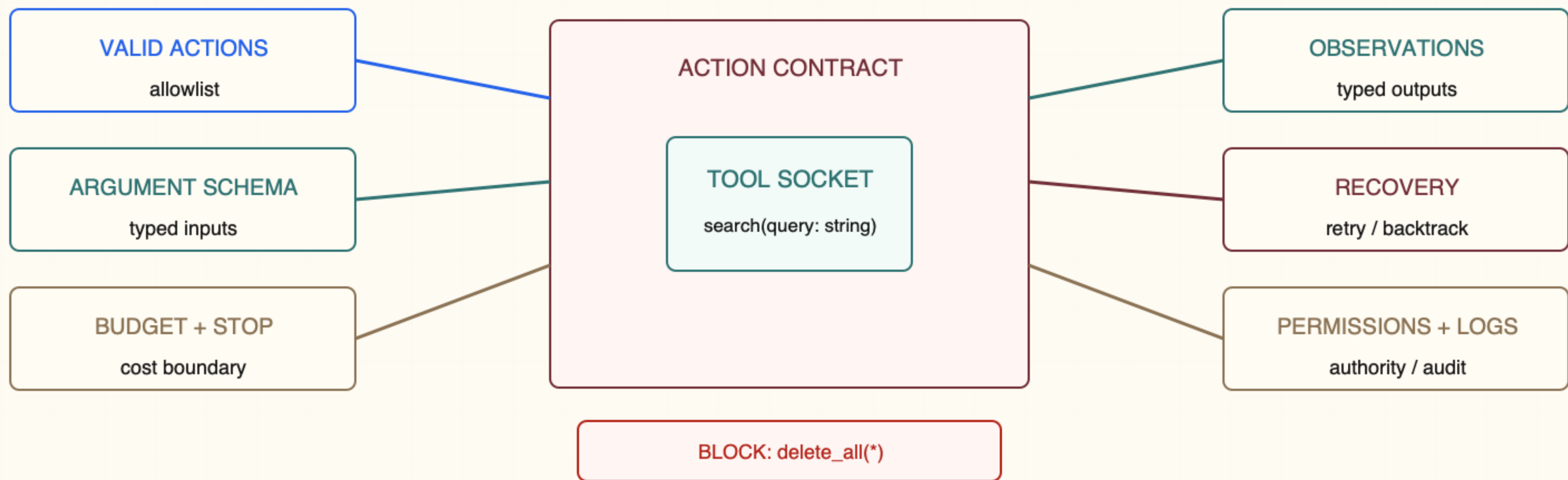
A Tool Observation Is Evidence, Not Truth

Retrieval can be relevant or misleading; interpretation still determines the state update.



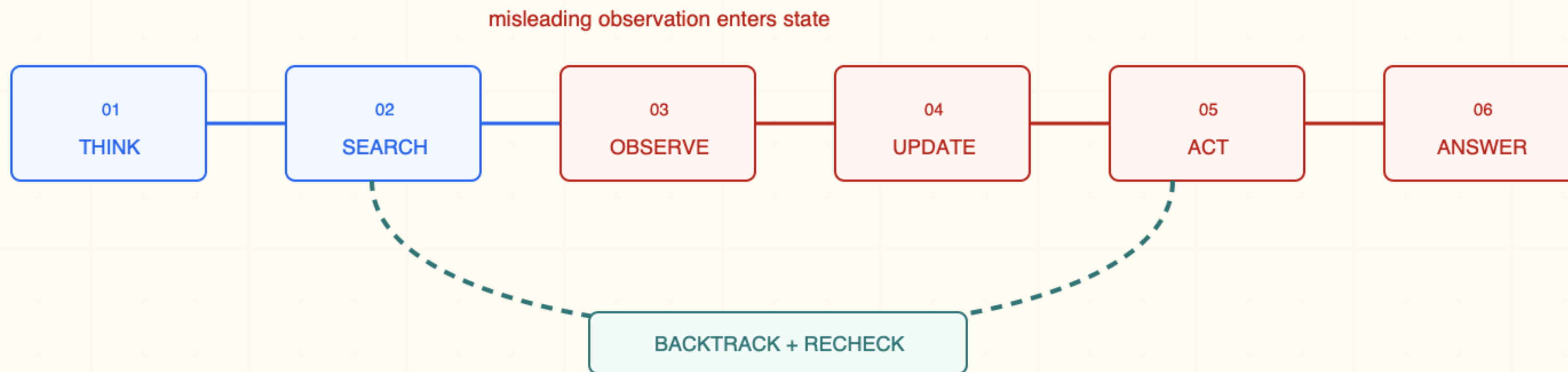
Tool Use Needs an Explicit Action Contract

The operating boundary must define schemas, observations, budgets, recovery, and permissions.



One Early Error Can Contaminate the Whole Trajectory

Longer loops create more chances to recover and more chances to compound a mistake.

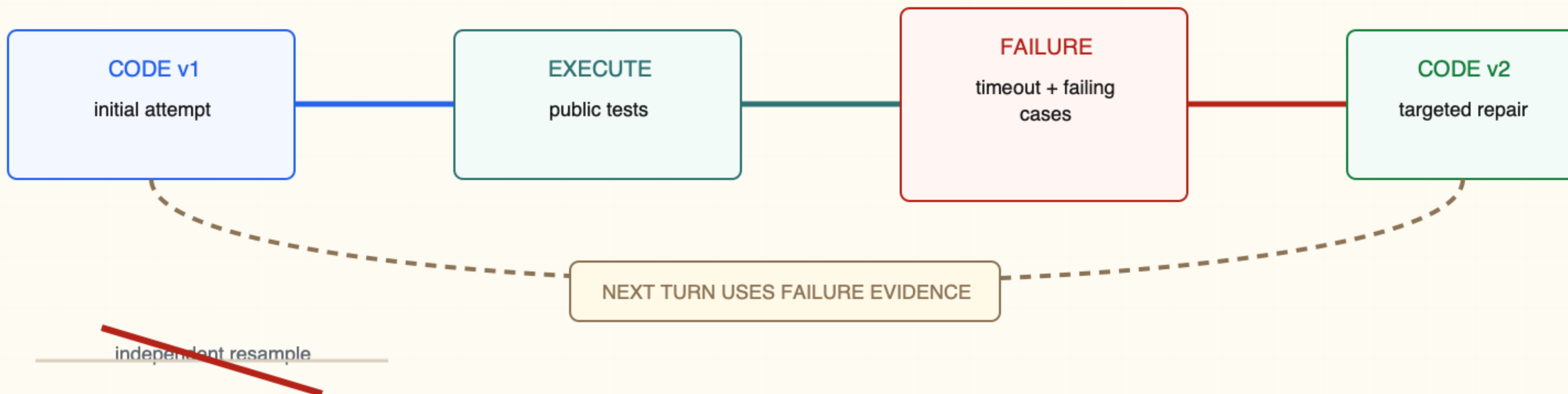


CAVEAT

Longer trajectories add both recovery opportunities and failure surface.

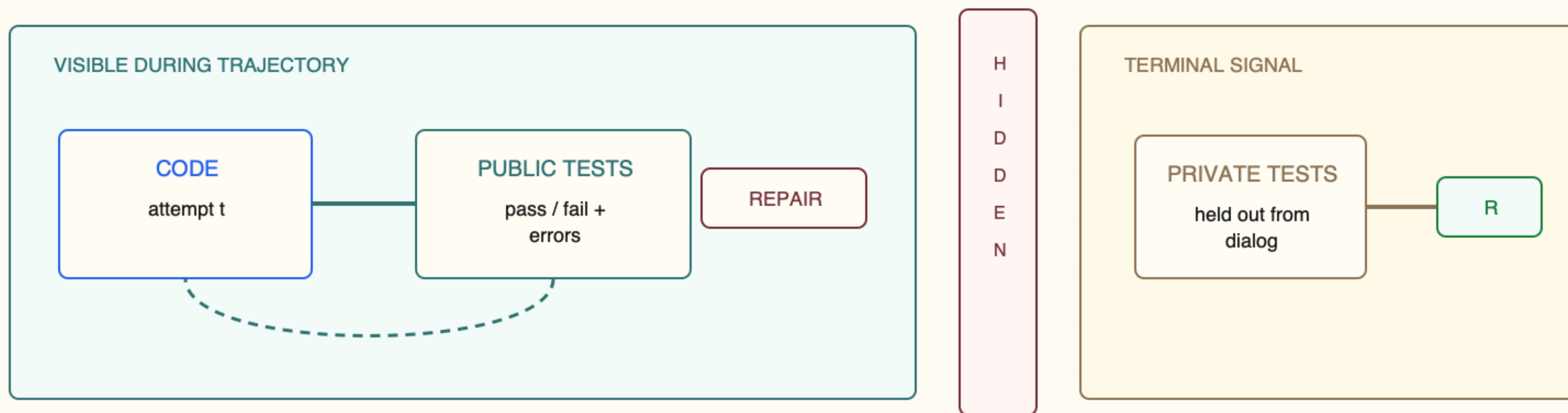
RLEF Teaches Code Models to Repair, Not Just Resample

A new attempt is conditioned on execution evidence from the previous attempt.



Public Tests Teach the Repair; Private Tests Guard the Reward

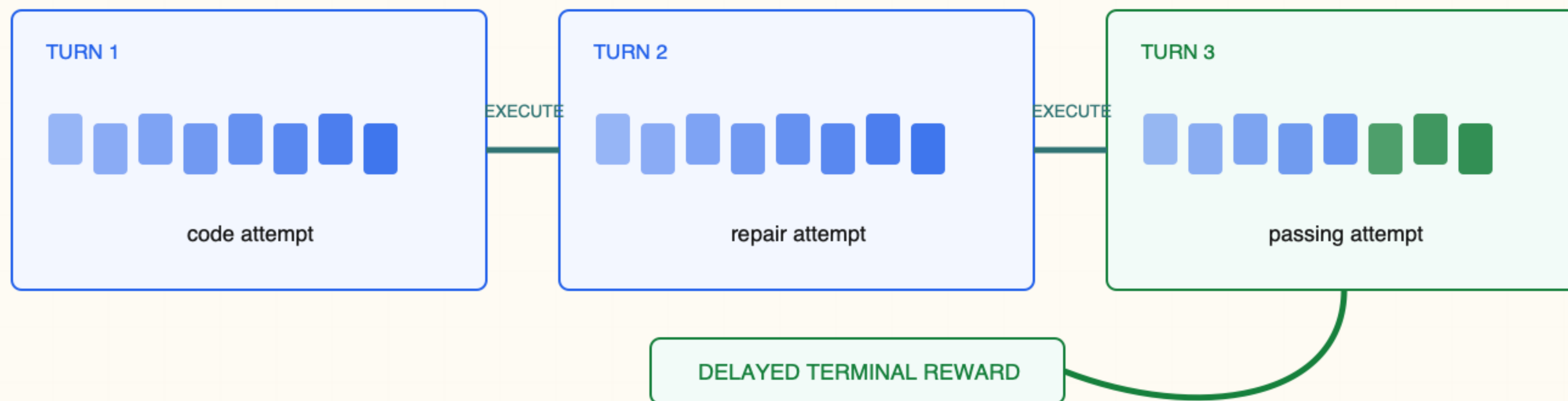
RLEF separates in-trajectory feedback from the terminal correctness signal.

**CAVEAT**

Hidden tests reduce direct exposure; they do not make the specification complete.

Credit Assignment Spans Tokens and Turns

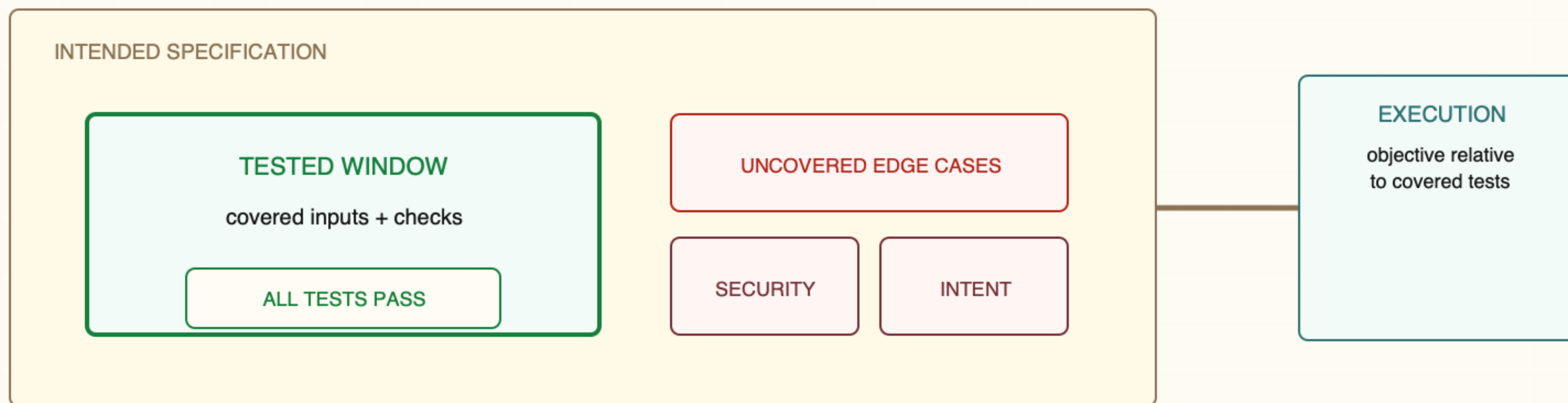
Code is emitted token by token; execution feedback arrives after a complete attempt.

**CAVEAT**

One terminal result must assign credit across many token decisions.

Executable Feedback Is Objective, but Narrow

Passing tests proves covered behavior, not complete correctness, security, or intent.

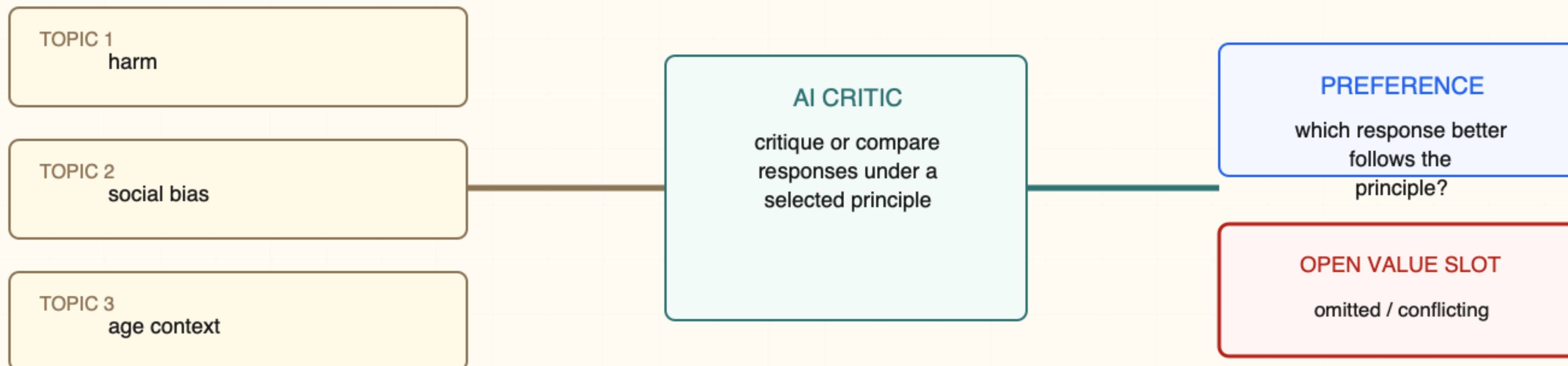
**CAVEAT**

Passing tests is evidence of covered behavior, not proof of complete correctness.

A Constitution Makes Normative Feedback Inspectable

Human-written principles guide an AI critic, but they do not remove human judgment.

ILLUSTRATIVE PRINCIPLE TOPICS

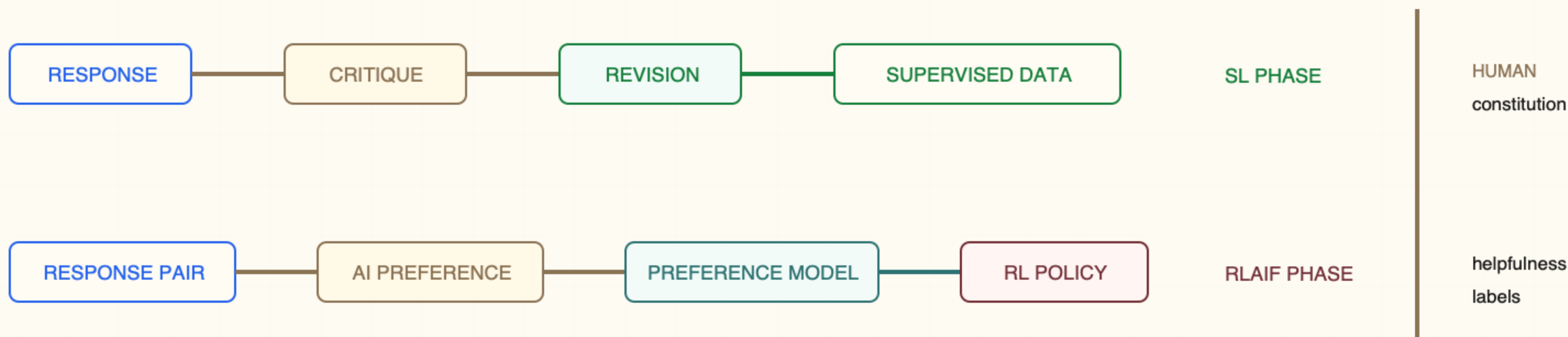


CAVEAT

Topic labels are illustrative paraphrases;
principles remain human choices.

Constitutional AI Trains Through Revision and Preference

Supervised critique revises responses; RLAIIF turns AI comparisons into a learned reward.

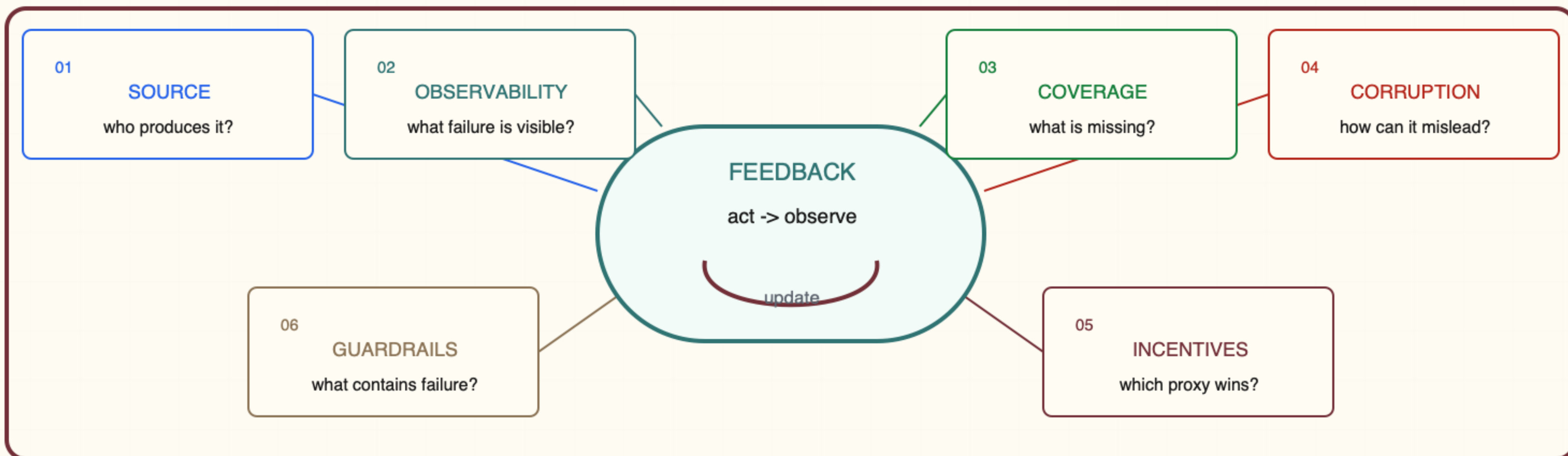


CAVEAT

AI feedback scales labeling; it does not remove humans or evaluator blind spots.

Audit the Feedback Contract, Not Just the Model

Source, observability, coverage, corruption, incentives, and guardrails bound self-improvement.



SELF-IMPROVEMENT IS BOUNDED BY THE FEEDBACK CONTRACT