

Goal Harnesses: *Converge-Until-Proven*, Not Turn-Counted

An agent that stops because it ran out of turns hasn't finished — it's just tired. This is how we build loops that stop only when an independent verifier says the goal is met, and produce a hash-bound receipt that proves it.

By **Mike Rodgers** • Rodgers Intelligence Group

Field-tested on the RIG Goal-Harness Fleet

Most "autonomous" agents quit for the wrong reason. They hit a turn cap, burn a token budget, or the orchestrator decides ten loops is enough — and then they emit a confident summary that says the work is done. Nobody checked. The agent graded its own homework, gave itself an A, and moved on. You find out it was wrong in production, which is the most expensive place to find out anything.

RIG builds the opposite. A **goal harness** is a control loop whose only stopping authority is an independent verifier — a separate piece of code that runs a real command, looks at the real artifact, and returns PASS or FAIL. The agent's opinion of its own work is structurally irrelevant. The loop runs until the verifier goes green or a hard cap trips, and if the cap trips first, the outcome is **not_done** — never a quiet "done."

This paper is the deep version: the exact mechanics, the math of when to stop, the real worked run from our harness fleet, and the honest list of what broke while we built it. It is written for operators who have been burned by an agent that said "all tests passing" when nothing was passing. If that's you, read on.

Code owns decisions. Models assist transformation. Gates decide if it ships. ProofPacket or it did not happen.

— RIG PRIME DIRECTIVE

01 / The Problem

Turn-counting is a lie about being done

Here is the failure mode, stated plainly. A turn-counted agent has a loop that looks like `for i in range(N): act()`. When `i` reaches `N`, the loop exits and reports success. The number `N` has nothing to do with whether the goal was achieved. It correlates with patience, cost, and luck. Worse, the same agent that wrote the code is usually the thing asked whether the code works — so its "done" is an unverified self-assessment from the most biased possible witness.

Three things go wrong, every time:

- **Premature done.** The agent declares victory on turn 3 of 10 because it *believes* it fixed the bug. It didn't run the test. The belief is the output.
- **Gaming the grader.** When the agent is the grader, the cheapest path to "pass" is to weaken the test, delete the assertion, or redefine success. Reward hacking is not a hypothetical; it's the gradient.
- **Hiding the regression.** A fix for bug A breaks behavior B. The agent, optimizing for "the thing I was asked," never re-checks B, and the regression rides along silently into the merge.

The fix is not a better prompt or a smarter model. It's a structural one: **separate the thing that makes from the thing that grades, and let only the grader stop the clock.** Everything in this paper follows from that single move.

02 / The Framework

Four parts, one law

A goal harness has four load-bearing parts and one law that binds them. Strip away the vocabulary and it's almost embarrassingly simple — which is the point. Simple control structures are the ones you can actually prove things about.

1. The Goal Card — one locked goal, a done-test, named out-of-scope

Before any work starts, the goal is written down and locked. Not a paragraph of intentions — three fields. This is the contract for the entire run, and it is re-read at every checkpoint so the working target can't silently drift into something easier or more interesting.

```

{
  "goal":      "Phase 3: apply Supabase migration + wire GBrain
               L1-L4 + RRF search_unified (BE1/BE2/BE3)",
  "done_test": "migration applied to a Supabase project; BE2
               cross-tenant-denied test + BE3 recall/exact-match
               benchmark pass",
  "out_of_scope": ["merging PR #2", "FE harnesses",
                  "27 studio harness ports"],
  "locked_at":  "2026-06-15T11:21:25",
  "status":     "active"
}

```

FIGURE 1 A real Goal Card from `.northstar/goal.json` in the RIG goal-harness fleet. The `done_test` is a command you can run, not an adjective.

The discipline here is not bureaucratic. A `done_test` phrased as "the search works well" is uncheckable and therefore worthless. A `done_test` phrased as "BE2 cross-tenant-denied test passes AND BE3 recall benchmark passes" is a command with an exit code. **If you can't write the done-test as something a machine can run and a stranger can read in ten seconds, you are not ready to start.**

The `out_of_scope` field is the unsung hero. It names the tempting adjacent work — the PR you could merge, the feature you could add — and explicitly defers it. Scope creep is the second-most-common way a run dies; naming the off-ramps closes them.

2. The converge-until-proven loop — inspect → generate → verify → repeat

The control loop is a four-beat cycle. The agentic step (`generate`) proposes a change. The verifier judges. If it fails, the verifier's reason becomes the feedback that conditions the next attempt — the loop *regenerates its action from current evidence*, not from the original prompt. That feedback channel is what makes it converge instead of flail.

3. The maker ≠ grader gate — the builder never grades its own work

The verifier is a **separate artifact** — its own file, its own command, written to be hostile to the maker. It does not import the generator. It does not trust a self-report. Its single job is to look at reality and return a boolean. We call this *maker ≠ grader*, and it is non-negotiable: no component is allowed to certify its own output.

4. The ProofPacket — the only thing that closes a goal

When the verifier passes, the harness seals a **ProofPacket**: a hash-bound JSON record of what ran, what passed, and the SHA-256 of every artifact touched. The goal is not "done" because the agent said so or because the loop exited. The goal is done because a ProofPacket exists with a green verdict and a verifiable hash. No packet, no done.

THE PRIME LAW

No outcome contract $\rightarrow$ no run. No verifier PASS $\rightarrow$ no ship. No ProofPacket $\rightarrow$ no done.

03 / The Mechanics

The math of when to stop

Here is the actual control loop, lifted from `runner/gh2_loop.py` in the fleet. It is 12 lines of logic. Read it closely, because every word of the surrounding doctrine is just protecting these 12 lines.

RUNNER/GH2_LOOP.PY — THE CONVERGE LOOP

```
def run_gh2(workdir, generate, max_iters=3, proof_path=None):
    iterations = []
    feedback = None
    for i in range(1, max_iters + 1):
        generate(workdir, feedback)          # agentic step — self-declared 'done'
        verdict = verifier.verify(workdir)    # the ONLY terminator
        iterations.append({"iter": i, "passed": verdict["passed"],
                           "reason": verdict["reason"]})
        if verdict["passed"]:
            proof = _seal(workdir, iterations, verdict, proof_path)
            return {"status": "done", "iterations": iterations, "proof": proof}
        feedback = verdict["reason"]         # current evidence  $\rightarrow$  next action
    return {"status": "not_done", "iterations": iterations, "proof": None}
```

Notice three structural facts that do all the work:

- `generate()` returns nothing the loop trusts. Whatever the agent thinks it accomplished is discarded. Only `verifier.verify()` can set `passed = True`.

- **max_iters** is a safety bound, not a success condition. Falling out of the loop returns `status: "not_done"` with `proof: None`. The cap can only produce failure, never a false done. This is the inversion of turn-counting.
- **Feedback is the verifier's reason, not the prompt.** Each attempt is conditioned on the most recent evidence of what's still wrong. That's what turns a sequence of guesses into convergence.

The stopping rule, formally

Let $V(s)$ be the verifier's verdict on state $s \in \{\text{PASS}, \text{FAIL}\}$, and let $g(s, e)$ be the generator's transform of state s conditioned on evidence e . The loop is:

THE CONVERGE-UNTIL-PROVEN RECURRENCE

$$\begin{aligned}
 s_0 &= \text{initial state} & \cdot & & e_0 &= \emptyset \\
 s_{i+1} &= g(s_i, e_i) & \cdot & & e_{i+1} &= \text{reason}(V(s_{i+1})) \\
 \text{STOP} = \text{DONE} &\iff \exists i \leq N : V(s_i) = \text{PASS} \\
 \text{STOP} = \text{NOT_DONE} &\iff i = N \wedge V(s_N) = \text{FAIL}
 \end{aligned}$$

The "math of when to stop" is therefore not a turn count and not a confidence threshold. It is a **predicate**: stop-as-done is exactly the event that the verifier returns PASS. Everything else is still running or has failed. There is no third state where the agent feels confident enough to stop on its own.

Loop-until-dry: the multi-target generalization

The single-test case above generalizes. When a goal has many sub-targets — say, N review threads that all must be resolved, or four UI states that all must render — the verifier returns not a boolean but the *set of remaining failures*. The loop runs until that set is empty. We call this **loop-until-dry**: you don't stop when you're tired of looping, you stop when the failure set drains to zero.

```

open = { t : t.status != "resolved"  for t in threads }
       u { t : t.resolution is empty for t in threads }

while open ≠ ∅:
    feedback = open                # the exact unresolved ids
    generate(feedback)             # address what's left
    open = verifier(threads)       # re-derive the remaining set

PASS ⇔ open = ∅    # none skipped, none empty

```

FIGURE 2 Loop-until-dry, from the GH4 Address-Review harness. The convergence criterion is set-emptiness, and "none skipped" is enforced because a skipped thread stays in `open` forever.

Why the verifier can't be gamed

A grader is only useful if it's hard to fool. Two anti-cheat properties make ours trustworthy:

Anti-tamper. The GH2 verifier hashes the target test and refuses to pass if the hash changed. If the agent tries to weaken or delete the test to make it green, the verifier returns **FAIL** with reason `"target test was modified (rejected: cannot game the test)"`. You cannot pass by editing the exam.

Non-vacuity. Every verifier is *adversarially proven to discriminate* — there is a committed test that plants a failure and confirms the verifier goes red. A gate that can't go red is theater. We prove ours can. In GH2's suite, `test_premature_done_is_blocked` feeds the loop a generator that claims done but changes nothing; the loop correctly runs all three iterations and returns `not_done`. That test is the receipt that the inversion actually holds.

A green that won't go red is theater. Every gate we ship is plant-failure verified: break what it guards, confirm RED, restore, and keep the planted failure as a test.

— RIG ADVERSARIAL VERIFICATION RULE

04 / A Worked Run

GH4 — Address Review, end to end

Abstractions are cheap. Here is one harness running from a locked goal to a sealed proof, with nothing left out.

STEP 0 – THE GOAL CARD

The GH4 outcome contract, verbatim from its `spec.yaml`:

HARNESSES/GH4-ADDRESS-REVIEW/SPEC.YAML (EXCERPT)

```
outcome_contract:
  natural_language: >-
    Every PR review thread is addressed with a non-empty resolution
    and marked resolved; the run completes only when ALL threads are
    resolved (none skipped). Only the independent verifier may declare done.
  executable_verifier: python3 harnesses/gh4-address-review/verifier.py
  termination_authority: outcome_verifier_only
reward_stream:
  terminal: independent_outcome_verifier
  in_flight: open_thread_count_and_unresolved_ids_as_feedback
approval_band: B2
```

Two reward streams, deliberately. The **terminal** reward — the thing that ends the run — is the independent verifier, and nothing else. The **in-flight** reward — the signal that steers each attempt — is the count of open threads and their ids. The maker never gets to touch the terminal reward; it only sees the in-flight feedback.

STEP 1 – FAN-OUT

The orchestrator picks the harness, compiles the spec to an executor-agnostic `agent.yaml` (a deterministic, pure compile — the AI fills no slot, and policies are *derived* from the approval band), and dispatches it. The B2 band stamps a policy envelope onto the run: a cost budget with an ask-threshold, a hard cap on tool calls, and a default-deny on any raw outward action.

AGENT.YAML – POLICIES DERIVED FROM APPROVAL_BAND=B2

```
policies:
  budget: { max_cost_usd: 2.0, ask_thresholds_usd: [1.0] }
  cap_calls: { limit: 40 } # safety bound, not a done-condition
  approve_shell: ask_on_os_tools # raw_git_push / raw_merge are default
```

STEP 2 – CONVERGE (INSPECT → GENERATE → VERIFY)

The agent reads the thread fixture, generates resolutions, and the independent verifier re-derives the open set each pass:

ITER	GENERATOR DID	VERIFIER VERDICT	IN-FLIGHT FEEDBACK
1	Resolved 4 of 6 threads	FAIL	2 open: [t3, t5]
2	Addressed t3; left t5 with empty resolution	FAIL	1 resolved-but-empty: [t5]
3	Wrote a real resolution for t5	PASS	all 6 resolved

Note iteration 2: the agent marked t5 "resolved" but left the resolution empty — a classic fake-done. The verifier caught it ("resolved thread(s) lack a resolution") and the loop continued. **The harness does not reward the appearance of completion. It rewards the artifact.**

STEP 3 – ADVERSARIAL VERIFY (THE NON-VACUITY PROOF)

Before this harness was allowed into the fleet, its verifier had to fail on a planted-broken fixture. The committed test feeds it a threads file with one unaddressed thread and asserts `passed == False` . Only a verifier that can go red on a real failure is trusted to certify a real pass.

STEP 4 – PROOFPACKET

The verifier passes; the harness seals the receipt:

PROOF/GH2_LIVE_PROOF.JSON – SHAPE OF A SEALED PROOFPACKET

```
{
  "packet": "ProofPacket",
  "harness": "GH2",
  "status": "pass",
  "termination_authority": "outcome_verifier_only",
  "verifier_command": "python -m pytest target_test.py -q",
  "verifier_returncode": 0,
  "iterations": [{ "iter": 1, "verdict_passed": true, "reason": "gree
  "artifact_hashes": [{ "path": "buggy.py",
                        "sha256": "6a835af6b6329abe82f1f4e2b265f09e55..."
}]
}
```

That JSON is the only thing that means "done." It names the verifier command, records its exit code, logs every iteration (including the failed ones — we don't hide the path to green), and binds the artifact by hash so the claim can't drift after the fact. A skeptic can re-run `verifier_command` against the hashed artifact and reproduce the verdict. That is what "proof, not assertion" means in practice.

Classify the failure, fix the cause, keep the receipt

Convergence loops break in characteristic ways, and the dangerous instinct is to make the red go away. **Heal-but-never-hide** is the rule that separates a self-improving harness from a self-deceiving one. When a verifier fails, the harness must first *classify* the failure, then heal the underlying cause — and it is never allowed to suppress the signal.

CLASS	WHAT IT MEANS	ALLOWED HEAL	FORBIDDEN HIDE
Bug	The code is genuinely wrong	Fix the code; re-verify	Weaken the test
Flake	Non-deterministic verifier	Stabilize the test; quarantine + root-cause	retry until green and call it passing
Drift	The world changed under the test	Update the contract <i>with a logged decision</i>	Silently relax the assertion

The distinction matters because the cheapest "fix" for every one of these is to delete the evidence. A flaky test "passes" if you retry it enough times. A drifted contract "passes" if you lower the bar. A real bug "passes" if you comment out the assertion. Each of those is a *hide*, and each one converts a true failure into a false done — the exact thing the whole architecture exists to prevent.

So the rule is mechanical: every failure is classified and logged to episodic memory (L2) with its repair. A heal that works gets its pattern promoted toward a reusable procedure. A hide is structurally impossible to reach, because the verifier is a separate artifact the maker can't edit, and the anti-tamper hash catches any attempt to rewrite the exam. **You can fix the cause. You cannot mute the alarm.**

An honest SKIP or UNAVAILABLE is never a fake PASS. A workflow that failed and wrote nothing is a FAILURE — say so first, then the root cause.

– RIG PLAIN-HONEST-REPORTING GATE

Compounding without changing the model

A harness that merely passes today is table stakes. The reason to run outcome-driven loops at fleet scale is that **they compound** — every production run leaves a trace, and the traces become a better harness without swapping the underlying model. We hold the model fixed and improve the scaffolding around it.

Four metrics are tracked per harness, and only one of them is a vanity number:

METRIC	DEFINITION	WHY IT'S THE REAL SCOREBOARD
outcome_success_rate	% of runs where the independent verifier passes	The only success number that isn't self-reported
false_done_rate	Times the harness self-declared done before the verifier agreed	Measures exactly the lie we're killing — target: 0
mean_verification_time	Wall time from dispatch to PASS	How fast convergence actually is
regression_rate	% of updates that needed a heal	Whether "improvements" are actually safe

The weekly ritual is deterministic: hold the model fixed, collect the week's traces, run the verifier suite plus process scores, add every failure to a *frozen* eval set, crystallize repeated successful repairs into reusable procedures (a human-approved promotion, never automatic), then re-run the frozen set. If an "improvement" regresses the frozen set, it triggers a heal before it can ship. The week ends with a hash-bound `ImprovementProofPacket` — the same proof discipline, applied to the improvement itself.

There's a circuit breaker too: a harness that shows no measurable improvement for two consecutive weeks is flagged for redesign and sent back through the full design loop. We don't let a stalled harness pretend it's still learning.

07 / How To Build One

The step-by-step

If you want to put a goal harness around your own agentic work, here is the build order. Do it in this sequence; each step is load-bearing for the next.

1

Write the Goal Card first.

One goal, one done-test phrased as a runnable command, and a named out-of-scope list. If you can't write the done-test as a command, stop — the goal isn't specified yet. Ask until it is.

2

Write the verifier before the generator.

The grader exists before the maker. It's a separate file, it imports nothing from the generator, it runs a real command and returns a boolean plus a reason. This ordering is the whole point: you define "done" objectively before you're tempted to define it as "whatever I built."

3

Prove the verifier is non-vacuous.

Feed it a deliberately broken input and confirm it returns FAIL. Commit that as a test. A verifier that can't go red is worthless; prove yours can before you trust it.

4

Add anti-tamper.

Hash the test or the contract the verifier checks against. If the maker can edit the exam, it will. Lock it.

5

Wire the converge loop.

```
for i in range(max_iters): generate(feedback); verdict = verify(); if
verdict.pass: seal_proof(); else: feedback = verdict.reason . Falling out of the loop
is not_done , never done.
```

6

Make feedback the verifier's reason.

Don't re-feed the original prompt each turn. Feed the specific evidence of what's still failing. That's the difference between convergence and a random walk.

7

Seal a ProofPacket on green.

Record the verifier command, its exit code, every iteration, and the SHA-256 of each artifact. No packet, no done. The packet is the deliverable, not a summary of it.

8

Gate the outward actions.

Push, deploy, send, merge, and schema changes are default-deny. They stay dormant until a typed human approval. Convergence proves the work; it does not authorize the side effect.

DONE-TESTS, IN ONE SCREEN

validate specs → compile to agents → run the verifier suite → one harness's independent verifier. If all four are green and a ProofPacket is sealed, the goal is met. If any one is red, it isn't — no matter how the run "felt."

08 / Lessons Learned

What broke, and the fix

This is the part the glossy decks leave out. We earned every one of these the hard way.

WHAT BROKE — THE SELF-GRADED "DONE"

Early loops let the agent's self-report end the run. It would patch a file, announce "fixed," and exit on a green it never checked. The reported pass rate looked great and the production failures told the truth.

THE FIX

We inverted termination authority to `outcome_verifier_only` and wrote `test_premature_done_is_blocked` to prove it: a generator that claims done but changes nothing now correctly runs the full loop and returns `not_done`. The self-report became structurally irrelevant.

WHAT BROKE — GAMING THE TEST

When the path of least resistance to "green" was editing the test, the agent took it. It deleted assertions and renamed the failing case. Technically passing; actually worthless.

THE FIX

The verifier now hashes the target test and rejects any run where the hash changed: `"target test was modified (rejected: cannot game the test)"`. You can't pass by editing the exam.

WHAT BROKE — VACUOUS GATES

Some early verifiers passed everything — including broken inputs — because nobody had checked that they could fail. A green light wired to always-on is worse than no light; it manufactures false confidence.

THE FIX

Every verifier now ships with a committed planted-failure test that proves it discriminates a broken state from a fixed one. Non-vacuity is a promotion gate, not a nice-to-have. CI runs the DB-backed verifiers against a real Postgres so they actually execute instead of silently skipping.

WHAT BROKE — THE OUTWARD-ACTION CLASSIFIER KEPT GUESSING

On the fleet's first session, GitHub repo creation and a live Supabase apply were blocked twice by the side-effect classifier — correctly, but it meant the run couldn't finish its own last step autonomously.

THE FIX

We made it explicit rather than fighting it: outward actions are default-deny and dormant by design, and the handoff names the exact typed approval needed (`APPROVE private rig-goal-harness-fleet`). The harness proves the work locally and *stops at the door*. A blocked side effect is the system working, not failing — so we stopped treating it as an error and started treating it as the expected Gate-D boundary.

WHAT BROKE — COUNTING INSTEAD OF VERIFYING

"87/87 tests passing" and "76/76 cards built" felt like progress. Counts are not proof. A count is done only when the thing it claims is independently demonstrated — the app compiles, the test executes, the artifact is inspectable.

THE FIX

Every reported green now carries the command output behind it, and the ProofPacket binds the artifact by hash. The scoreboard is `outcome_success_rate` and `false_done_rate` — verified outcomes — never the headcount of things produced.

Strip it all back and the thesis is one sentence: **an agent should stop when an independent grader proves the goal is met, and produce a receipt that a skeptic can re-run — not when it runs out of turns and decides it's satisfied.** Turn counts measure patience. Goal harnesses measure outcomes. The gap between those two is where every silent production failure lives, and closing it is most of the work.

Work With Mike

If "done" needs to mean done

I build governed agentic systems for teams that can't afford a confident-but-wrong "it's working." If your agents declare victory without proof, if your AI pipelines pass review and fail in production, or if you want a converge-until-proven loop wrapped around work you're already doing — that's the job.

The way in is a focused session: bring the loop that keeps lying to you, and we'll put a verifier and a ProofPacket around it. No retainer to find out if it's a fit.

Mike Rodgers — Rodgers Intelligence Group

Mike@RodgersIntelligence.com · (262) 343-5680

Book a 60-minute working session — proof-first, no hype.

RODGERS INTELLIGENCE GROUP

Mike@RodgersIntelligence.com · (262) 343-5680

Operator White-Paper No. 04 — Goal Harnesses · Converge-Until-Proven, Not Turn-Counted