



Cross-Session Handoff for AI Agents: A Failure Report

What we built, exactly how it broke, and where the next person should start.

Summary

We tried to give a Claude Code session a memory that survives death. The idea: before a long session gets reset — by `/clear`, by a quit-and-restart, or by automatic compaction — it writes down who it is and what it was doing, and a fresh session automatically reads that back and picks up where the last one left off. No lost context, no re-explaining, no re-deriving what was already figured out.

We built most of it. A `/handoff` skill that writes durable state to disk, a `SessionStart` hook that loads it back, a trust model to keep a hostile repo from injecting itself as "your predecessor," an auto-nudge that watches for context degradation, and an emergency writer that fires just before compaction. Two rounds of adversarial review, all high-severity findings closed.

Then we tested it on a real session with real work at stake. It failed twice, in two different ways, for two reasons that had nothing to do with our logic and everything to do with the platform underneath us:

1. A running Claude Code process does not re-read `settings.json`. Our hook was installed on disk but invisible to the process that needed it, so the reset wiped the session and nothing loaded it back.
2. Hook output above roughly 10 KB is truncated — the full text is spilled to a file and only a pre-view is injected. Our whole point was a *comprehensive* handoff, which is by definition large. The successor got a 2 KB preview and a dead end.

We stopped. This document is what's left, and it is more useful than the code was. Every failure below is one that any comparable system will hit — the platform facts don't care whose code sits on top of them. So this is written two ways at once: as an honest post-mortem, and as a starting kit for the next person or company that wants to actually make cross-session handoff work. The design that held up is here to copy. The failures are here so you don't rediscover them the expensive way.

1. Why this is worth solving

A long-running LLM session degrades well before it hits the context limit, and it degrades in ways the token counter doesn't show you.

- **Persona drift.** Wang et al. (arXiv 2402.10962) measured persona stability dropping ~44% by turn 8 of multi-turn dialogue — pure attention decay on the original instructions, no adversarial input required.
- **Context rot.** Chroma's 2025 study of 18 frontier models found accuracy degrading with context length across retrieval, reasoning, and instruction-following. Anthropic's own docs name the effect.
- **Reasoning fails before retrieval does.** RULER (arXiv 2404.06654) and follow-ups show reasoning quality falling long before the model loses the ability to fetch a fact. A coding session — which is reasoning on top of retrieval — hits the reasoning cliff first.

NoLiMa (Modarressi et al., ICML 2025) put a number on it: Claude 3.5 Sonnet's *effective* non-lexical context was ~4K tokens against a nominal 200K window — two percent. Newer models are better, not by orders of magnitude. Anthropic's own Claude Code guidance says to start a new session when you start a new task, and notes the model is "at its least intelligent point when compacting."

So the operator is stuck with a bad choice: keep working in a session that's quietly getting worse, or reset and lose everything the session learned — the workspace understanding, the dead ends already ruled out, the "here's what we tried" that never made it into a file. The existing escape hatches all leak. Built-in `/compact` is a lossy summary produced by a model that's already impaired. Manual notes need discipline and can't auto-fire. Restart-from-scratch is what most people actually do, and it throws away the most expensive thing in the room.

What was missing is a primitive: *same agent, same machine, save state before a reset, load it after, with a trust boundary and a staleness rule*. That's the gap we went after. (The closest published analog is Google's A2A protocol, now under the Linux Foundation — but A2A is agent-to-agent over a network, not same-agent-across-time inside one process. The problems rhyme; the mechanism doesn't transfer.)

A working version would have to clear three bars before anything else counted: a reset loses no state; parallel agents don't clobber each other; and cloning a random repository can't inject arbitrary content as trusted context. Everything else — auto-triggering, compaction safety, autonomous runs — is amplification on top of those three.

2. What we built

Enough of this worked that it's worth handing over intact. If you build cross-session handoff, this is the part to steal.

2.1 Two artifacts, not one

The load-bearing decision: split identity from state, because they decay at different rates.

A **persona file** (`~/.claude/personas/<slug>.md` , global, one per agent) holds *who the agent is* — role, domain, explicit in/out scope, loaded skills, user preferences, product context, sibling agents. It changes on the order of weeks and is reloaded verbatim at every session start. Never summarized.

A **state file** (`<project>/.claude/handoffs/<slug>/state-<ts>.md` , per project, per handoff) holds *what the agent was doing* — objective, progress anchored to file paths and commit SHAs, the paths already tried and why they failed, current hypothesis, next concrete step, verification commands, live background processes, booby traps. It's transient; one per handoff.

Merge the two and you rewrite identity every time the work changes, which feeds the exact drift you're trying to prevent. Splitting them lets each reload on its own clock. LangGraph's checkpointers/BaseStore split and CrewAI's always-loaded `backstory` field independently landed on the same shape, which is some evidence it's right.

The state schema forces a field most summarizers drop and every framework's users complain about losing:

```
# Objective                - one sentence
# Progress                 - anchored to file paths, SHAs, test names
# Attempted paths that failed and why - required; the successor's anti-loop
# Current hypothesis
# Load-bearing tool outputs - excerpts, not dumps
# For architecture questions, read - pointer list into the actual code
# Next concrete step       - one action
# Verification commands
# Live background processes - pids/ports, or "None."
# Cautions
```

"Attempted paths that failed and why" is the single most-cited handoff failure in the literature: successors re-tread dead ends because summarizers throw away negative results. Make the field required and refuse to write a handoff without it.

2.2 Trust tiers, so a repo can't impersonate your predecessor

A hook that blindly injects any `.claude/handoffs/` content from the current directory is a prompt-injection vector — clone a hostile repo and its "handoff" becomes your ground truth. So loading is gated three ways:

- **Ground truth.** A `session_id → slug` breadcrumb at `~/.claude/.session-slug-map/<session_id>`, written by the predecessor's own `/handoff`. Same session id means it's genuinely us. Injected as "trust after a sanity check."
- **Advisory.** No breadcrumb (e.g. after a full restart, which mints a new session id), but the slug has a persona file that was created *on this machine*. That proves local provenance, not same-session identity. Injected as "verify each claim against reality before acting."
- **Blocked.** Neither. The hook skips injection. This is what closes the hostile-repo hole — a cloned repo can carry a state file but not a locally-authored persona.

2.3 Auto-nudge, because handoff-before-degradation beats handoff-after

A bad session writes a bad handoff, so the trigger has to fire while the session is still sharp. Three canaries evaluated on every `UserPromptSubmit`:

- **Context load** (primary) — tokens as a fraction of the window, read from a side file the statusline writes. 40% soft / 60% hard on 200K models; 30/45 on 1M.
- **Turn count** (secondary) — 20 soft / 35 hard.
- **Tool symptoms** (tertiary) — same call 3x in a row, or the same file read 5+ times. 1 symptom soft, 2+ hard.

Any canary hard, or two soft, fires a hard nudge. Token count alone misses drift in long chat-heavy sessions where every turn is cheap; turn count alone over-fires in code sessions with a few enormous turns. You need both, plus a cooldown log so it doesn't nag every turn once a threshold is crossed.

2.4 A compaction safety net

Even with a perfect nudge, a user ignores it and auto-compact fires. The `PreCompact` hook runs synchronously just before that. It *can* block compaction; we chose not to, because a wrong block leaves the harness unable to compact and the context fills forever. Instead it writes a cheap emergency snapshot (frontmatter + git state + last-N transcript lines), and `SessionStart` folds that in as a "recent activity since last checkpoint" section if it's newer than the last real handoff. The model-authored handoff is high-quality but stale; the shell-authored snapshot is low-quality but fresh. The successor gets both.

3. What broke

This is the part worth reading. Everything here was observed, root-caused, and — where noted — has a fix that we either shipped or know how to ship. The unshipped fixes are the leads: they're where a successor starts.

3.1 The two that killed it

The session that got wiped with nothing to load

What happened. After the core shipped, we tested with a live coding agent mid-task on real work. It ran `/handoff` — wrote every file correctly to disk — then ran `/clear` per our own documented flow. The context wiped. Nothing loaded back. The session came up empty.

Root cause. Claude Code reads `settings.json` at process startup, not on `/clear`. We'd registered the `SessionStart` hook *while that process was already running*. The hook existed on disk and was invisible to the one process that needed it. `/clear` did exactly what it always does; nothing was there to reverse it.

Why we missed it. Our docs said "after install, run `/clear` to activate." We conflated "written to the filesystem" with "loaded in this process." Those are different states, and the gap between them is exactly the first-run experience.

The fix, and why documentation can't be it. The mechanical fix is trivial — quit and restart, don't `/clear`, until the hook was present at process start. But a warning buried in the skill file is only visible *after* you've installed and started using the thing, and the moment you get burned is that very first test. Any system that needs a process restart to activate has a rollout hazard that docs can't close. The real fix is structural: either hard-gate on a restart before *any* use, or design the primary flow to degrade gracefully with no hooks at all (skill-only, manual load) so a missing hook costs convenience, not the session. Build for the second. Recovery here was only possible because the files were on disk — we loaded them into the next session by hand.

The 25 KB handoff that got truncated

What happened. With the core shipped, I wrote myself a comprehensive handoff before degrading: a 19 KB state file dense with decisions and dead ends, plus a 6 KB persona. Testing it by running the hook directly produced a perfect 262-line injection. Then the user did a real quit-and-restart, and the successor came up with no visible pickup at all.

Root cause. Claude Code caps hook stdout at roughly 10 KB of injected context. Above that, it writes the full output to a `tool-results/` file and injects a short preview plus the path — *with no instruction to read the file*. Our 25 KB payload became a 2 KB preview and a dangling pointer. The

core promise — a fresh session auto-loads the full handoff — is exactly what the budget forbids, because a handoff worth writing is large.

Why we missed it. We tested by invoking the hook directly in a shell (`bash session-start.sh`), which has no injection budget and prints everything. The truncation only exists inside the real hook harness. **We tested around the harness instead of through it** — the single most expensive mistake in the project.

The fix (unshipped — start here). Split the artifact and load on demand. Keep the persona under ~3 KB. Make the state file a small pointer document (under ~3 KB) that references detail files — `state-details/progress.md` , `.../attempted.md` , and so on. The hook injects only the pointer and instructs the successor to `Read` details as needed. It's a layer of indirection, but it respects a budget you don't control. Discover that budget empirically *before* you design the schema, not after.

3.2 The one that made it look broken even when it half-worked

Resume is not startup

What happened. After the truncation, the user tried again — from their seat, "started a new session in this tab." The telemetry showed `SessionStart` firing with source `resume` , not `startup` : Claude Code had re-attached to the *same* session id, not spawned a fresh one. Our hook injected the handoff anyway — redundantly, into a context that already had it, and truncated on top. There was no "successor restates and waits" moment because there was no successor. It was the same session, talking to itself.

Root cause. "Quit and restart" from a terminal doesn't reliably produce a fresh session — it depends on terminal behavior, the exact re-invocation, and Claude Code's session persistence. Our hook treated all four `SessionStart` sources the same. Our tests always used a novel `session_id` , so we only ever simulated a fresh start and never saw `resume` .

The fix (unshipped). Branch on `source` . If `source=resume` and the session id already has a breadcrumb, the handoff is already in context — skip injection entirely. The four-value matcher (`startup|resume|clear|compact`) exists for a reason; handle each case explicitly instead of collapsing them.

3.3 The small cuts — shell and platform footguns

None of these killed the project, but every one cost hours, and every one will bite the next person writing hooks in bash. They're worth more as a checklist than as war stories.

- `set -e` + a function that ends in a short-circuit. `find_latest_precompact` ended with `[[ -n "$best" ]] && echo "$best"` . When `$best` was empty (the normal case), the `[[` re-

turned non-zero, `echo` never ran, and the function's exit code was that non-zero — which, called inside `$(...)` under `set -e`, killed the whole hook silently. It passed every earlier test because those tests happened to have a file present. **Fix:** every bash function whose last command can fail needs an explicit `return 0`, and every function needs its "nothing to do" branch tested.

- **Env-var propagation across a pipe.**

`FIELD_NAME="$field" printf '%s' "$INPUT" | python3 -c '...'` sets the variable for `printf`, not for `python3`. Python's `os.environ["FIELD_NAME"]` raised `KeyError`, got swallowed by a bare `except`, and returned an empty string — so every field parsed as empty, every canary read zero, and no nudge ever fired. The tests "passed" because most of them expected no nudge. **Fix:** put the prefix on the correct side of the pipe (`... | FIELD_NAME="$field" python3 -c`), and never let a test accept the right answer for the wrong reason — assert the mechanism, not just the outcome.

- **Python inlined in a bash heredoc.** `python3 - <<'EOF'` worked in a fresh subshell and silently failed in the accumulated outer test shell. We never fully isolated why (some interaction of inherited `set -e` state, stdin routing, and function scope), but the pattern was reproducibly fragile. **Fix:** extract Python to a dedicated `.py` file called with args and env vars. This one change removed a whole class of ghost bugs.
- **Shell values interpolated into `python3 -c` source.** Any value spliced into the Python source string is an injection sink — a session id containing a quote executes arbitrary Python. **Fix:** pass values through `os.environ` or `sys.argv`, never string interpolation.
- **Every path built from disk content.** A slug or filename read from a file and turned into a path is a traversal risk (`../../etc/passwd` as a "slug"). **Fix:** regex-gate everything — `^[a-z0-9][a-z0-9-]{0,63}$` for slugs, `^state-[0-9T-]+Z?.md$` for state filenames.
- **Every file a hook reads.** Write `.tmp` then `mv`. A hook can fire mid-write and read a torn file otherwise. Applies to state files, the `current.md` pointer, breadcrumbs, side files, logs — all of them.
- **Smaller ones, for completeness:** Python <3.12 forbids backslashes inside an f-string expression (extract the value to a variable first); `set -u` plus an empty glob crashes on `${ARR[0]}` (guard with `${#ARR[@]}` or use `find`); macOS ships bash 3.x, so no associative arrays, no `${var,,}`, no `readarray` if you want it to run out of the box.

4. Why we stopped

Three things converged.

The failures were real, not theoretical. The wiped session had hours of real work in it, recoverable only because the files happened to be on disk before `/clear`. Had the `/handoff` itself failed, that work was gone. Two live tests against the bar "good enough to hand to a real AI company, and low-friction" produced two nukes. That's honest signal, and it said the foundation wasn't there yet.

The fragility compounded. Every individual failure was small and fixable. But each component added surface area, and the next planned phases were autonomous mode (the highest blast radius of anything in the design) plus a pile of extras — more surface stacked on a base that hadn't yet passed "a fresh session loads the handoff." Adding features to an unreliable foundation just concentrates the unreliability.

The remaining bugs were the platform's, not ours. Settings-vs-process, the injection budget, resume-vs-startup — these aren't logic errors a user-space tool can robustly fix. You can work around each one, but you're building on facts the platform doesn't guarantee and can change under you. Past a certain point that's not a coding problem, it's a "this belongs in the platform" problem.

Stopping was the right call. A half-shipped memory system is worse than none: it creates an expectation that state is safe, and then breaks at the exact moment the state matters. We took the code down and kept the lessons.

5. If you want to make this work

Concrete leads for the next person, roughly in the order I'd tackle them.

Nail the platform facts first, empirically, before writing a line of schema. Three numbers/behaviors decide your whole design, and none are guaranteed to hold at the version you're on:

- What is the hook-stdout injection budget, exactly? Measure it by injecting increasing payloads through the *rea/* harness until you see truncation. Design the artifact to live under it.
- What does each `SessionStart` source (`startup|resume|clear|compact`) actually do to context, and which of them already has your state in-window? Test with a fixed, reused `session_id`, not a fresh one each time.
- When does a running process pick up `settings.json` changes? Assume never-without-restart until you prove otherwise.

Then ship the three fixes we didn't. They're specified above and they're the difference between "demo" and "works": split-and-load-on-demand artifacts (§3.1), `source`-aware injection (§3.2), and a restart hard-gate or a hooks-optional degrade path (§3.1).

Test through the harness and on a session you'd be sad to lose. Mock tests validate the inputs you thought of; the fatal bugs lived in inputs we didn't — harness behavior, session lifecycle, size limits. Direct-invoking the hook proved nothing. Dogfood on real work early, because that's the only place the real failures appear.

Keep the bash hygiene as non-negotiable defaults. Every function returns 0 explicitly; every hook runs `trap 'exit 0' ERR` so a failure exits clean instead of injecting garbage; all Python lives in `.py` files, fed by args and env vars, never inlined and never string-interpolated; every disk-derived path is regex-gated; every written file goes `.tmp → mv`. These aren't polish — each one maps to a specific bug above.

Run adversarial review, and don't skip it under deadline. The two review rounds we ran caught the hostile-repo injection, the slug-hijack, the Python-injection sink, and the path-traversal holes before they shipped — design-level bugs that no unit test was going to surface. The one plan we didn't review is the one that shipped the `set -e` bug. The review is cheap next to the incident.

Where this really belongs. Every unfixed failure is platform-level, which is the honest conclusion: a robust version of this is a native primitive, not a hook-and-skill stack. If you have platform influence, the highest-leverage moves are small and specific — publish the injection budget so systems fail loudly at write time instead of silently at read time; document the `SessionStart` source semantics; commit to a stable transcript format so tooling can parse turn counts and tool calls without reverse-engineering an internal file. Those three unblock everyone building in user-space, us included.

Appendix: schemas to copy

The reusable contracts, verbatim, so a successor doesn't reinvent them.

Persona file — `~/.claude/personas/<slug>.md` :

```
---
persona_slug: <string, matches filename>
persona_version: <int, bump on material change>
autonomous: false                                # explicit opt-in; safe default
created_at: <ISO8601>
updated_at: <ISO8601>
---
# Role
# Domain
# Explicit scope                                (in / out – the out list matters most)
# Loaded skills
# Mannerisms
# User preferences
```

# Product context	(macro picture – needed to strategize, not just act)
# Sibling agents	(other agents in play, their scope vs. mine)

State file — `<project>/.claude/handoffs/<slug>/state-<ts>.md` :

```
---
handoff_version: 1
created_at: <ISO8601>
git_sha: <hash>
predecessor_session_id: <string>
predecessor_turn_count: <int>
predecessor_context_pct: <0-100 or "unknown">
trigger_reason: <manual|checkpoint|nudge_response|precompact_emergency>
persona_slug: <string>
---
# Objective
# Progress
# Attempted paths that failed and why      (required)
# Current hypothesis
# Load-bearing tool outputs
# For architecture questions, read
# Next concrete step
# Verification commands
# Live background processes
# Cautions
```

Hook registration — `~/.claude/settings.json` . Note the `bash $HOME/...` prefix, which side-steps a `chmod +x` fragility, and the four-source `SessionStart` matcher you must then branch on:

```
{
  "hooks": {
    "SessionStart": [{
      "matcher": "clear|compact|startup|resume",
      "hooks": [{"type": "command", "command": "bash $HOME/.claude/hooks/session-
start.sh"}]
    },
    "UserPromptSubmit": [{
      "matcher": "",
      "hooks": [{"type": "command", "command": "bash $HOME/.claude/hooks/user-
prompt-submit.sh"}]
    },
    "PreCompact": [{
      "matcher": "auto|manual",
      "hooks": [{"type": "command", "command": "bash $HOME/.claude/hooks/pre-com-
pact.sh"}]
    }
  ]
}
```

```
}]
}
}
```

Trust check — the core of `session-start.sh`, the ground-truth tier. The `resume -skip` from §3.2 belongs right after the slug resolves:

```
# Ground-truth trust: a breadcrumb the predecessor's /handoff wrote
if [[ -n "$SESSION_ID" ]] && [[ -f "$SESSION_MAP_DIR/$SESSION_ID" ]]; then
    RAW_SLUG=$(head -1 "$SESSION_MAP_DIR/$SESSION_ID" | tr -d '\r' \
        | sed 's/^[[:space:]]*//; s/[[:space:]]*$//')
    if [[ "$RAW_SLUG" =~ ^[a-z0-9][a-z0-9-]{0,63}$ ]]; then # traversal gate
        PREFERRED_SLUG="$RAW_SLUG"
        TRUST_LEVEL="predecessor"
    fi
fi
# Advisory trust otherwise: scan handoffs/ for a slug that has a LOCAL
# persona file, pick the freshest state. No local persona → skip (blocked).
```

The rest of the surface a builder inherits: a statusline processor writing a per-session JSON side file (session id, context %, token counts, pid, timestamp) that the canaries read; append-only `log.jsonl` audit trails per persona; a per-level nudge-cooldown log; and the atomic-write discipline on every one of them.

This system did not ship, and taking it down was correct. It produced a map of what breaks when you try to solve cross-session state for AI agents in user-space — with a root cause and a fix beside each failure. The hard part of the design isn't the code; it's knowing what not to assume. That part is done. Start from section 3.