



Harness Engineering

A living book on the scaffolding that surrounds AI agents

Gilsley Henrique Darú · with AI co-authorship (Claude, Anthropic)

v0.76.0 · generated on August 7, 2026

DOI 10.5281/zenodo.21632412 · harness.ghdaru.com.br

Opening

00 — Introduction

🕒 state of the art 2026-07 · revised 2026-08-01 · 📖 ~7 min read

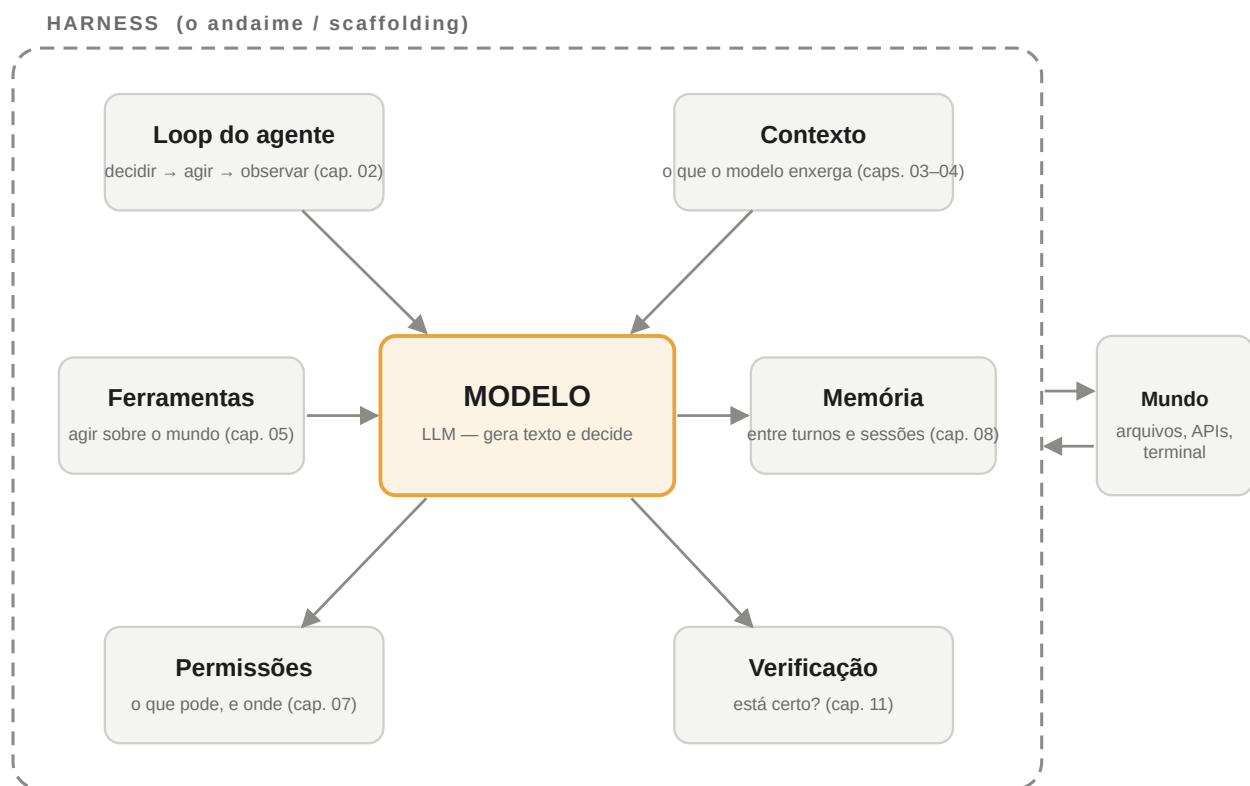
Agent = model + harness

Start with a question anyone who has used an AI chat can ask: why does ChatGPT *answer questions about* your problem, but not *solve* your problem? It explains how to fix the bug — but it does not open the file, run the test, or check that it worked. The short answer: a chat is just the **model**. For the model to *act* — touch files, run commands, verify its own work, and stop at the right moment — an entire structure must be built around it. That structure is the subject of this book.

When an AI agent solves a real task — fixing a bug, migrating a module, answering questions grounded in dozens of files — two distinct things are at work. The first is the **model**: the network that reads context and decides the next step. The second is everything around it: whoever assembles the context it reads, whoever executes the tools it invokes, whoever decides what it may or may not do, whoever remembers what happened yesterday, whoever checks whether the result is correct. That "everything around it" is the **harness** — the rigging, the scaffold, the *scaffolding*.

The formula that organizes this book is simple:

agent = model + harness



The model at the center; the harness — the scaffolding — around it. Each block is a chapter of this book.

The model is interchangeable and improves with every generation. The harness is classic software engineering — and it is where most agents fail or succeed. Two products using exactly the same model deliver radically different results depending on the quality of the harness: how context reaches the model, which tools it has, how errors come back, what happens when the **context window** (the limit of text the model can "see" at once) runs out.

Harness engineering is the discipline of designing that scaffolding: context delivery, tool interfaces, planning artifacts, verification loops, memory systems and sandboxes.

Why a book — and why now

Between 2024 and 2026, coding-agent harnesses stopped being experiments and became a product category: Claude Code, Codex CLI, Gemini CLI, opencode, Aider, Cline, Goose, OpenHands and dozens of others. The most remarkable thing is not the quantity but the **convergence**: independent projects, in different languages, arrived at the same solutions — hierarchical context files, layered compaction, plan mode as a permission mode, lifecycle hooks, MCP (Model Context Protocol) as the integration standard.

When independent implementations converge, there is a discipline behind them. This book documents that discipline.

The method: read code, not marketing

This book is empirical. Each chapter covers one harness capability (the loop, context, compaction, permissions...) and is written from reading the source code of real open source harnesses. The project's most important editorial rule:

Claims about a harness require **evidence**: the file path in the source code where the capability is implemented.

READMEs promise; code delivers. Several projects advertise dimensions their code does not have — the evidence requirement is what separates evaluation from marketing.

A note on authorship and method

For transparency — and consistency with the evidence rule above — this book is **co-written with an AI agent** (Claude Code, by Anthropic) operating under **human authorship, curation and responsibility**. The agent carries out the research, the writing and the production cycle; the human author defines the scope, decides, **verifies every source** and answers for the content. Following editorial authorship policies (ICMJE, COPE, *Nature*, *Science*), the AI is **not** listed as an author — it cannot be held responsible — and its use is disclosed here, at the opening.

This is not a detail: a book about the discipline of properly instrumenting AI agents uses that very discipline to write itself, and exposes it. The full method — dual research verified by cross-search, spec-driven cycle, review and dating — is documented in the [Editorial Guide §6](#), with a survey of the traditional and AI-era writing methodologies that ground it.

How to read this book — three doors in

This book was written to be dense; this section exists so the density is not a wall. Pick your door:

- **If you are just arriving** (you have used AI chats but never built an agent): read 00 → 01 → 02 in order, unhurried, with the [Glossary](#) as support — every acronym in the book is there, spelled out and explained (in the online version, just hover over the acronym). After 02, chapters 03–13 can be read in any order: each is self-contained and opens by defining its own problem.
- **If you already operate an agent** (you use Claude Code, Codex, Cursor or similar and want to understand what is inside): the **Executive summary** at the end of each chapter is your shortcut — the state of the art of the dimension in one paragraph, with the "what to steal" section. Go straight to the chapters you care about and descend into the body when you want the evidence.
- **If you build harnesses**: the whole book is yours, including the Appendices A (per-repository evidence, with file paths), the [Benchmark](#) and the two hands-on tracks — **harness-zero** (didactic build, one feature per step) and **harness-um** (the complete reference implementation, [own appendix](#)).

Structure of the book

- **Foundations** (chapter 01): the formal definitions, the canonical papers and the problem taxonomy that organizes everything that follows.
- **Chapters 02–13**: one capability per chapter. Each defines the problem, presents the known implementation patterns and shows, with evidence, how each studied harness implements it.
- **Convergences and trends** (chapter 14): what the industry has already standardized, where real divergence remains, and the "expiration clause" — the thesis that every harness component exists because the model cannot yet do that on its own, and must be designed knowing that one day it will be unnecessary.
- **Chapters 15–17**: the frontiers — the harness embedded in a product (15), the harness that learns from usage (16) and the protocol layer that binds the ecosystem together (17).
- **Benchmark** (benchmark/): the empirical section — standardized, per-dimension evaluations, with 0–3 scores and evidence, of every harness studied, plus the consolidated comparison.

The harnesses in the study

As of this edition, the study covers **twenty-one open source systems**, evaluated through systematic code reading across five archetypes (the method is in [chapter 01, §6](#)):

- **Coding harnesses** — opencode, gemini-cli, OpenHarness, Codex CLI, Goose, Aider, OpenHands, Grok Build, Pi, Kimi Code and Prime Agent;
- **Self-hosted personal agents** — OpenClaw, Hermes Agent, IronClaw, ohmo;
- **Organizational agents** — QM;
- **Embedded harnesses** — n8n (AI Agent node);
- **Frameworks** — LangGraph, CrewAI, OpenAI Agents SDK (Software Development Kit), Software Agent SDK.

Each was chosen for representing a different *archetype* (replication logic, not sampling): mature provider-agnostic product (opencode), big-tech control regime (gemini-cli), readable didactic port (OpenHarness), sandbox-first (Codex CLI), MCP-native (Goose), context-first (Aider), academic eval culture (OpenHands), whole-organization agent with a swappable loop (QM), and so on.

The full list — with the **exact origin, version, fork and commit read** in each evaluation, and the link to each one's analysis and diagnosis — is in the [Appendix — The study](#). The consolidated per-dimension scoreboard is in the [Comparative](#).

As a theoretical reference and to explore the ecosystem beyond the corpus, there is also the living collection [Awesome Harness Engineering](#) (~426 resources organized by problem, in the same organization as this book) — the source of the harness definition used in chapter 01 and of the taxonomy that structures the chapters.

01 — Foundations

🕒 state of the art 2026-07 · revised 2026-08-01 · 📖 ~13 min read

This chapter pins down the book's vocabulary, **origin** and **method**. Before comparing harnesses (chapters 02–13) we need to answer three questions the first edition left open: *what* a harness is, *where it came from* (and what existed before), and *with what rigor* this book studies it.

1. What a harness is (definition)

The working definition comes from the curated list [awesome-harness-engineering](#):

Harness engineering is the discipline of designing the *scaffolding* — the **support structure** — that surrounds an AI agent (context delivery, tool interfaces, planning artifacts, verification loops, memory systems and sandboxes) and determines whether it succeeds or fails at real tasks.

With the guiding principle:

The focus is the *harness*, not the model. Each component exists because the model cannot do it on its own — and the best harnesses are designed knowing these components will become unnecessary as models improve.

Note the central term: **scaffolding**. It is the book's metaphor — the temporary structure erected around something under construction, which supports the work and is later removed. Keep the word in mind: it reappears in the subtitle, in the title of each part and in §8 (the expiration clause).

If you are just arriving — one image that carries the whole book. Think of the model as a brilliant professional on their first day at a company they do not know: capable, but with no desk, no access to the systems, no knowledge of the house rules — and with a memory that resets after every conversation. The harness is everything the company builds around them: the project dossier they read on arrival (context, ch. 03), the tools on the bench (ch. 05), the badge that defines where they may enter (permissions, ch. 07), the notebook that survives the end of the shift (memory, ch. 08), the supervisor who reviews the deliverable before it ships (verification, ch. 11) — and the shift itself, the rhythm of work-check-continue (the loop, ch. 02). When the chapters turn technical, come back to this image: every dimension in the book is a piece of that office.

2. What came before — and why they were not agents

"Software that acts on your behalf" is an old idea. The previous generations, however, solved the problem **without a language model at the center of the decision loop** — and that is what separates them from an agent:

- **Expert systems** (1980s): hand-written `if-then` rules. They automated decisions, but neither interpreted goals in natural language nor recovered from unforeseen exceptions.
- **RPA — Robotic Process Automation** (UiPath, Automation Anywhere): bots that replay clicks and keystrokes from a fixed *script*. Fragile to any screen change; no goal, no recovery.
- Intent-based **chatbots** (from ELIZA to dialog trees): they produced text, but **did not execute actions** in the world.
- **Code assistants as autocomplete**: **GitHub Copilot** (technical preview in Jun 2021), powered by the **OpenAI Codex** model (a descendant of GPT-3 fine-tuned on code), suggested the next line *inside the editor* — no plan, no tools, no verification loop.

None of them had the **four pieces** that define a harness today (§4). They lacked goal-oriented autonomy and the ability to act on the environment **and correct their own course**.

3. How we got here — the technical lineage

The passage from "model that answers" to "agent that acts" was built in layers, each removing one obstacle:

1. **Explicit reasoning.** *Chain-of-Thought* (Wei et al., 2022) showed that asking the model to "think step by step" improves reasoning tasks.
2. **The loop.** The decisive milestone was **ReAct — Synergizing Reasoning and Acting in Language Models** (Yao et al., [arXiv:2210.03629](https://arxiv.org/abs/2210.03629), Oct 2022; ICLR 2023), which interleaved **Thought** → **Action** → **Observation**: the model reasons, calls a tool, observes the result and continues. That cycle is the skeleton of virtually every modern harness (chapter 02).
3. **Tool calling.** What was missing was a reliable way for the model to *invoke* tools — solved when OpenAI shipped **function calling** (Jun 2023): the model emits structured JSON to trigger functions (chapter 05).
4. **The autonomous wave — and its lesson.** With reasoning + action + tools came 2023: **AutoGPT** (Significant Gravititas, Mar 2023) and **BabyAGI** (Yohei Nakajima, Apr 2023) — loops that decomposed themselves into subtasks and ran on their own. They "failed" in the practical sense (going in circles, burning tokens, losing the thread) because they had *the loop* but **not** the other three pieces: context management, well-designed tools and control. The discipline's founding lesson was born there: **the model alone is not enough; the scaffolding around it is what decides success.**
5. **Maturation — the coding CLIs.** The four pieces were embedded into terminal tools wired to the filesystem and to Git: **Aider** (Paul Gauthier, Apr 2023), **Claude Code** (Anthropic, research preview in Feb 2025), **OpenAI Codex CLI** (open source, Apr 2025), plus projects like **Cline**, **OpenHands** and **SWE-agent**.
6. **Standardization.** With agents proliferating came the protocols: the **Model Context Protocol (MCP)**, opened by Anthropic (Nov 2024), standardized the connection to tools and data (chapter 06); **AGENTS.md** consolidated as the "README for agents"; **Agent2Agent (A2A (Agent-to-Agent))** (Google, Apr 2025; later donated to the Linux Foundation) addressed communication *between* agents (chapter 17).

Timeline (milestones): 1980s expert systems · 2000s–2010s RPA and chatbots · **Jun 2021** Copilot (autocomplete) · **Oct 2022** ReAct · **Mar–Apr 2023** GPT-4, AutoGPT, BabyAGI, Aider · **Jun 2023** function calling · **Nov 2024** MCP · **Feb 2025** Claude Code · **Apr 2025** Codex CLI and A2A.

A note on rigor. "Codex" names three distinct things — the 2021 *model* (the basis of Copilot), OpenAI's Codex *product line* and the open source *Codex CLI* of 2025. The text keeps them separate. Dates and sources for this section are in the [Bibliography](#); items still awaiting verification are flagged there.

4. The constitutive definition: the four elements

The discipline's literature converges on a definition of the harness as a **runtime layer** with four necessary and sufficient elements:

1. **Agent loop** — the cycle that alternates between invoking the model and executing what it decided, until a stopping criterion (ch. 02).
2. **Tool interface** — the contract through which the model acts on the world: reading files, running commands, calling APIs (ch. 05).
3. **Context management** — the assembly, prioritization and compression of what the model sees on each call (chs. 03–04).
4. **Control mechanisms** — permissions, approvals, sandboxes and limits that constrain what the agent can do (ch. 07).

A system missing any of the four **is not a complete harness**: a chatbot with tools but no loop is a "function caller"; a loop without control is an incident waiting to happen; tools without context management collapse on long tasks. **This is the operational definition that serves as the study's inclusion test** (§5–6).

The four pieces in one real task. Ask an agent: "the `test_login` test is failing, fix it". What happens, piece by piece: **context management** assembles what the model will see (the project rules, your message, perhaps the test file); the model reads it and decides to request an action — "run the test and show me the error" — which the **tool interface** actually executes in the terminal; the result comes back, the model proposes editing a file, and the **control mechanisms** decide whether that edit happens directly or needs your approval; once applied, the **loop** feeds the model the new state — does the test pass? — and repeats the cycle until the stop criterion. Four pieces, one turn of work. Chapters 02–13 are this paragraph in slow motion.

5. Where the harnesses in this study come from

The corpus is **open source** (the book's Principle II: the base source is the code) and splits into five archetypes — the same as in chapter 00:

- **Coding harnesses** (opcode, gemini-cli, OpenHarness, Codex CLI, Goose, Aider, OpenHands, Grok Build, Pi, Kimi Code, Prime Agent): reference implementations that put the four pieces together in one executable.
- **Self-hosted personal agents** (OpenClaw, Hermes Agent, IronClaw, ohmo): the harness in the service of one person, with its own identity, memory and channels.
- **Organizational agents** (QM): the harness in the service of an organization — scopes, audience-based permissions and auditing as primitives, with the agent loop as a swappable engine.
- **Embedded harnesses** (n8n, AI Agent node): the loop as a component inside a larger product.
- **Frameworks** (LangGraph, CrewAI, OpenAI Agents SDK, Software Agent SDK): they expose loop, state and tools as programmable primitives.

The **inclusion test** is the definition in §4: whoever has *loop + tools + context management + control* gets in; pure model libraries and mere single-tool *wrappers* stay out. The evaluated list, with each one's repository and the commit read, is in the [Comparative](#) and in the study appendix. Resources to consult beyond the corpus are in the living collection [Awesome Harness Engineering](#).

6. The study's method (rigor)

This book **reads the source code of real harnesses**, compares them across dimensions and then **builds a harness from scratch**. That is not "an engineer's opinion": it is a hybrid research design resting on established methodological traditions. Making them explicit turns the book from a collection of impressions into an **auditable empirical study** — consistent with Principle I ("evidence over rhetoric").

In plain language, before the technical names: the method is (1) picking systems that represent different *types* of harness, not the most famous ones; (2) reading each one's code following **the same script of questions**, recording the exact file that proves each answer; (3) scoring against a fixed, published rubric, so anyone can disagree while looking at the same evidence; and (4) building a harness from scratch to test whether the extracted patterns actually hold. The paragraphs that follow give the formal names of each of those choices and where they come from — they are the genealogy of the rigor, and may be skimmed on a first pass.

Two phases, two engines.

- **Phase 1 — descriptive/comparative:** a **multiple-case study** (Yin) supported by **Mining Software Repositories** (Hassan, 2008), treating each repository as *primary data*. The unit of analysis is **the source code**, not marketing material nor behavior observed in use.
- **Phase 2 — constructive/prescriptive:** **harness-zero** is an exercise in **Design Science Research** (Hevner et al., 2004; the DSRM process of Peffers et al., 2007): designing and evaluating an artifact that instantiates the principles extracted in Phase 1.

How dimensions become measures. The comparison dimensions descend via the **Goal–Question–Metric** method (Basili, Caldiera & Rombach): for each harness goal (context, tools, permissions, memory, verification, loop, orchestration) questions are formulated and, for each question, **indicators observable in the code** (e.g.: is there a compaction mechanism? how granular is the permission model? is there a post-action verification layer?).

Selection by replication, not by sampling. Cases are chosen by Yin's **replication logic** — *literal* (the same pattern is expected) or *theoretical* (a predictable difference is expected) — with explicit criteria: open and inspectable code at the cutoff date; membership in the "harness" class (§4); adoption relevance **or** architectural singularity; archetype diversity (§5). For each case the **URL, commit/tag and reading date** are recorded.

Coding and synthesis. The reading follows a protocol common to all cases (Runeson & Höst, 2009), combining inductive coding inspired by *grounded theory* (Stol, Ralph & Fitzgerald, 2016) in the discovery of the dimensions and *content analysis* (Hsieh & Shannon, 2005) with a fixed grid in the scoring. The comparative synthesis is a **feature analysis** in the **DESMET** style (Kitchenham, Linkman & Law, 1997), in the tradition of *benchmarking* as an engine of scientific progress (Sim, Easterbrook & Holt, 2003).

Threats to validity (Cook & Campbell's 1979 taxonomy, adapted to case study):

Type	Threat	Declared mitigation
Construct	the "dimensions" fail to capture what defines a harness	derivation via GQM; published operational definitions
Internal	attributing to "good practice" what is a project's historical accident	single protocol; every claim traced to a snippet/commit
External / obsolescence	failure to generalize; the field changes in months	selection by archetypes; cutoff date + pinned commits ; the expiration clause (§8) is the declared mitigation, not an ornament
Conclusion	treating qualitative scores as exact metrics	explicit scale and criteria (DESMET); no spurious numeric aggregation

Thus every claim in the book traces back to **a datum in the repository** and to **a named procedure**. The operational details are in the [Comparative](#) and in the evaluation template; the references, in the [Bibliography](#).

7. Taxonomy by problem

A convention inherited from the reference collection: organize the discipline **by the problem being solved, not by vendor or model**. It is the taxonomy that structures the chapters:

Problem	Chapter
How the decision-action cycle works and when it stops	02 — Agent Loop
What the model sees and how that is assembled	03 — Context Delivery
What to do when the context window runs out	04 — Compaction
How the model acts on the world	05 — Tool Design
How to integrate external capabilities in a standardized way	06 — MCP
What the agent may do, and where	07 — Permissions and Sandboxing
What persists across turns and across sessions	08 — Memory and State
How large work becomes verifiable steps	09 — Planning
How to distribute work across multiple agents	10 — Subagents and Orchestration
How to know whether the agent (and the harness) work	11 — Verification and Evals
How third parties extend the harness	12 — Extensibility
Through what surfaces humans and systems use the agent	13 — Interfaces

8. The expiration clause

The discipline's most important — and least practiced — thesis: **every harness component is a temporary prosthesis**. Compaction exists because context windows are finite; *plan mode* exists because models act rashly; the *policy engine* exists because models are not trustworthy with destructive commands. Each premise has a shelf life.

The practical corollary: every component should document **which model capability improvement would make it unnecessary**. Harnesses that fail to do this accumulate dead *scaffolding* — complexity that outlives the limitation that justified it. As seen in §6, this clause is also the **declared mitigation** of the obsolescence threat: the book assumes it is dated. We return to it in chapter 14.

9. Operational artifacts

The discipline has produced standard artifacts that reappear, with variations, in nearly every harness studied:

- **Project instructions file** (`AGENTS.md` / `CLAUDE.md` / `GEMINI.md`): rules, conventions and limits the agent reads before any task. Clear boundaries beat vague restrictions.
- **Plan artifact** (`PLAN.md`): created at the start of the task and updated during execution, with verifiable milestones and scope boundaries.
- **Implementation log** (`IMPLEMENT.md`): an *append-only* record of decisions and deviations from the plan.
- **Harness checklist** (`HARNESS_CHECKLIST.md`): a pre-production review covering instructions, tools, context, planning, permissions and verification — with the expiration table from §8.

These artifacts are the embryo of our evaluation instrument (see `benchmark/template/HARNESS_EVAL.md`).

*This chapter's sources (historical and methodological) are consolidated in the [Bibliography](#), separating the **confirmed** ones from those still awaiting verification — faithful to Principle I.*

Chapters by capability

02 — The Agent Loop

🕒 state of the art 2026-07 · revised 2026-08-01 · 📖 ~9 min read

Learning objectives

1. **Explain** the prompt → decision → tool → result cycle and the structural stopping criterion;
2. **Compare** the industry's two termination contracts (absence of tool calls × satisfied `output_type`);
3. **Implement** a loop with brakes (turns, budget) and an observable trace (step 1 of harness-zero);
4. **Design** two-layer retry (inside the step × loop replay) and recognize what requires idempotency;
5. **Evaluate** the durability of a real loop (what survives a crash?).

The problem

The loop is the heart of the harness: it sends context to the model, receives a decision (text and/or **tool calls** — structured requests for action: "run this tool with these arguments"), executes, feeds back and repeats — until someone decides to stop.

One full turn, in slow motion. You type: "the `test_login` test failed, fix it". What the loop does:

1. Assembles the context (project rules + your message) and **calls the model**;
2. The model does not answer with text — it answers with a tool call: `run_shell("pytest test_login")`;
3. The harness **actually executes it** and returns the output (the error traceback) to the model, as if it were a new message;
4. The model has now *seen* the error and emits another tool call: `edit_file("auth.py", ...)`;
5. The harness executes (perhaps asking for your approval — ch. 07) and returns the result;
6. A new call to the model, which asks for the test again; this time it passes;
7. The model answers **with text only** ("fixed: it was the expired cookie") — and *that* is what ends the turn: **no tool call, the loop stops**.

Seven steps, three model calls, two real executions. Everything else in this chapter is the hard questions hiding in that cycle: who decides to stop (and what if the model never stops?), how errors come back, what happens when the process dies at step 5, how much this can cost. The design questions: who decides to stop? how do results and errors come back? what happens when things go wrong? does the loop survive a restart?

Scientific foundations

- **ReAct** ([arXiv 2210.03629](#)) is the seminal paper: interleaving reasoning and action with environment feedback beats pure reasoning — it is the scientific justification for the loop's existence.
- The survey of **agentic reasoning frameworks** ([arXiv 2508.17692](#)) systematizes the cycle's variants (ReAct, plan-and-act, reflection), useful as a map of the territory.
- The trained frontier: surveys of **agentic search with RL** ([arXiv 2510.16724](#)) show the loop ceasing to be mere orchestration and becoming a training target — when the model is trained *in* the loop, part of the harness migrates into the weights.

(Full bibliography: `livro/bibliografia.md`.)

Industry sources

- **How the agent loop works** (Claude Agent SDK (Software Development Kit), docs): the canonical 5-stage loop; a "turn" ends **when the model responds without tool calls**; and the most modern detail — termination as a **typed state** (`success`, `error_max_turns`, `error_max_budget_usd`...): success and limit exhaustion are distinct, mandatory code paths. Includes `max_budget_usd` **propagated to subagents** and compaction as an observable loop event (`compact_boundary`).
- **Loop engineering** (Claude blog): the vendor names the discipline and gives the taxonomy by axes (how it fires, how it stops, which primitive it uses) — with the quotable design rule: *if you cannot write the verification, the loop is not ready to exist*.

- **Building Effective AI Agents** (Anthropic): the founding workflow × agent distinction and the **evaluator-optimizer** pattern — semantic stopping (quality reached) with a separate judge.
- **Running agents** (OpenAI Agents SDK): the alternative contract — stop when the agent produces the declared **output_type** (validatable), with a typed `MaxTurnsExceeded`.
- **LoopAgent** (Google ADK): only two ways to stop — `max_iterations` or a judge sub-agent emitting `escalate=True` — the dumb loop separated from the addressable judge.
- **Durable AI Loops** (Restate) and **Inngest**: the loop as a **long-running workflow** — each step journaled, failure = replay from the last completed step; retry becomes two categories (backoff inside the step × loop replay), with idempotency mandatory for mutating tools.
- **See also**: the living collection [Awesome Harness Engineering — Agent Loop](#) gathers more resources for this dimension (patterns, articles and implementations), curated by problem.

The state of the art

1. Stopping became a multi-axis contract

The structural criterion (no tool call) remains universal, but on its own it is naive. The modern contract combines: a turn limit; a **budget ceiling in money** (the real novelty of 2025–26, already propagating to subagents); a typed termination *subtype*; and, in the Agents SDK's alternative contract, **stopping by output type** — which turns "are we done?" into verifiable validation. On top of this, two refinements measured in the benchmark: gemini-cli's **next-speaker check** (a cheap inference decides whether the model continues on its own) and the termination veto — Stop hooks that can **refuse the end of the turn** and reinject feedback (software-agent-sdk; Hermes's verify-on-stop is the same principle as a nudge).

2. Anti-runaway: from counter to detector

Every mature loop has `MAX_TURNS`; the best have repetition detection — `LoopDetectionService` (gemini-cli), `RepetitionInspector` (Goose), a stuck detector with `stalled/stuck` states (software-agent-sdk, OpenClaw). The field technique (hashing `tool+args` over a sliding window) circulates among practitioners but has no vendor doc — citable as practice, not as norm.

3. Durability became a property of the loop, not of the infra

The 2026 consensus: per-step journaling + replay. In the benchmark: recoverable jsonl rollouts (Codex), a durable prompt inbox with cursor-replayable events (opcode V2), an append-only event log with directory-based resumption (software-agent-sdk) and — the most radical design — the executor that **returns only durable references** and never mutates state, with an applier validating evidence before applying (IronClaw). The corollary for tools: idempotency stops being a virtue and becomes a requirement.

4. The loop is not the perimeter

The most important architectural lesson of round 2 (IronClaw): *"the loop is intentionally not the security perimeter"* — the loop requests effects through ports; the kernel decides. Even outside the security context, the software-agent-sdk's separation of policy (when to stop/confirm/give up — `Conversation.run()` × mechanics (view → LLM → dispatch — `Agent.step()`) is the clean cut that lets you swap the engine while keeping the loop.

EXECUTIVE SUMMARY

What is most modern: typed termination with a dollar budget; a separate, addressable judge (evaluator-optimizer/escalate) instead of a heuristic in the prompt; durability via journaling/replay; and the policy×mechanics separation. **What to steal**: the typed `ResultMessage.subtype`; budget propagated to subagents; Stop hooks with veto power; the `LoopExit` via durable references.

Hands-on — harness-zero, step 1

Step 1 (harness-zero/etapas/01-loop/) implements the core in ~30 lines: structural stopping, `MAX_TURNS` as a brake, tool errors returning **as text** for the model to decide, and a trace of actions visible in the chat. Extension exercises: (a) add a termination subtype

(`success × max_turns`); (b) add an estimated-cost budget and the third subtype.

Check your understanding

1. Why is "the model responded without tool calls" a good stopping default — and why is it insufficient on its own? (Multi-axis contract.)
 2. Your agent called the same tool with the same arguments 5 times in a row. List two defenses of different natures. (Repetition detector × budget ceiling.)
 3. The process died in the middle of turn 7. What must your loop have persisted in order to resume without repeating side effects? (Journaling + idempotency.)
-

Appendix A — How each repository handles the loop

Per-harness evidence, with paths — online supplement, expanded with each round.

opencode (round 1)

`packages/opencode/src/session/processor.ts`: response consumed as an Effect Stream (`Stream.tap(handleEvent) → takeUntil(needsCompaction) → runDrain`); explicit `continue | stop | compact` verdict; per-provider retry (`SessionRetry.policy`); V2 (`CONTEXT.md`): durable inbox and cursor-replayable events.

gemini-cli (round 1)

`packages/core/src/core/client.ts` (`MAX_TURNS=100`) + `turn.ts`; **next-speaker check** (`utils/nextSpeakerChecker.ts`: mini-prompt {`reasoning`, `next_speaker`} re-invokes the stream if `model`); `LoopDetectionService`; clean core/cli separation.

OpenHarness (round 1)

`src/openharness/engine/query.py` (`run_query`): `async while` until `max_turns` or no tool-uses; **parallelism when all tools in the turn are read-only** (`asyncio.gather`); `PreToolUse → permission → execution → PostToolUse` per call; retry with backoff and cost tracking.

Codex CLI (round 2)

`core/src/session/turn.rs` (`run_turn`, 2,581 lines) on top of the `SessionTask` trait (`Regular/Review/Compact/UserShell`); SSE (Server-Sent Events) streaming **and WebSocket with WS → HTTPS fallback**; hierarchical `CancellationToken`; each turn persisted in jsonl rollouts; no explicit repetition detector (mitigated by budgets).

Goose (round 2)

`crates/goose/src/agents/agent.rs` (`reply → BoxStream<AgentEvent>`): two levels of retry (transient per provider + a recipe-level `RetryManager` with a `SuccessCheck` that resets the conversation); `DEFAULT_MAX_TURNS=1000`; `RepetitionInspector`; `MAX_EMPTY_TURN_RETRIES=3`.

OpenClaw (round 2)

`src/system-agent/agent-turn.ts` + `gateway/agent-*.ts`: runs serialized per *session lane* with an inter-process file-based write-lock; three event streams (`lifecycle/assistant/tool`); `stalled/stuck` watchdogs; dual hooks (`Gateway` + `plugins`).

Hermes (round 2)

`agent/conversation_loop.py` (~6.5k lines) with separate phases (`turn_context/tool_executor/turn_finalizer`); `iteration_budget`; **interrupt-and-redirect** (`/steer` drained pre-API and post-tool); nudges for empty responses; role-alternation repair; **verify-on-stop** nudge.

IronClaw (round 2) ★

`crates/ironclaw_agent_loop`: a pipeline of sealed stages (`input` → `prompt` → `model` → `capability` → `gate/checkpoint` → `stop`), each stage a private strategy; the executor returns a `LoopExit` containing **only durable references** — it never mutates state — and the `LoopExitApplier` validates host-owned evidence before applying (the architecture's explicit thesis: *"the loop is intentionally not the security perimeter"*). Resumable state via checkpoints; a `Queued` → `Running` → `Blocked` → `Completed` state machine with leases/heartbeats; "one active run per canonical thread".

Aider (round 2)

`aider/coders/base_coder.py`: not a tool-calling loop — it is a chat REPL + direct editing. The only iterative mechanism is **reflection** (`reflected_message`, max 3): files requested outside the chat, linter errors or failing tests trigger a new round, always with human confirmation. Reactive self-correction by design, not autonomy.

OpenHands/Canvas (round 2)

`app_server/event/`: the event-stream persists each `Event` as JSON per conversation (pagination, filters, trajectory export) — but the action/observation loop runs in the `openhands-agent-server` (SDK); the app consumes events, it does not generate them. The core is in the `software-agent-sdk` (below).

ohmo (round 2.5)

Loop inherited from OpenHarness's `QueryEngine`; what is its own: a **multi-session pool** (`ohmo/gateway/runtime.py`: one `RuntimeBundle` per `session_key`, recreated when the cwd changes) and **real interruption by new message** (`bridge.py`: each message is an `asyncio.Task`; a new message from the same session cancels the previous one) — few competitors cancel correctly.

n8n (round 2)

V2 uses LangChain's classic `AgentExecutor` (`maxIterations` default 10); V3 keeps `createToolCallingAgent` only to *decide* — tool calls become `EngineRequests` handed back to the **n8n workflow engine**, which schedules the tool nodes and re-enters with `EngineResponse` (`ToolsAgent/V3/helpers/runAgent.ts`). n8n re-internalized the execution loop: decision by the framework, execution by the engine.

Frameworks (frameworks round) — four answers to the same question

LangGraph: the real primitive is **Pregel/BSP** (supersteps + channels + reducers), with per-node retry/cache/timeout — and the ready-made agent (`create_react_agent`) formally deprecated (migrated to `langchain.agents`). **OpenAI Agents SDK (Software Development Kit)**: explicit loop in `run.py` (`output_type` terminates · handoff switches agent · `max_turns` with handlers), on top of a swappable `AgentRunner`. **CrewAI**: a **100% in-house executor, zero LangChain** (`crew_agent_executor.py`), with dual dispatch — native tool-calling or a ReAct fallback with `json_repair`. **software-agent-sdk**: `LocalConversation.run()` (policy: stop, confirm, give up) separated from `Agent.step()` (stateless mechanics view → LLM → dispatch), an append-only event log with a derived `View`, and `Stop` hooks with **veto** power over termination.

03 — Context Delivery

🕒 state of the art 2026-07 · revised 2026-07-25 · 📖 ~11 min read

Learning objectives

By the end of this chapter, you should be able to:

1. **Explain** why context is a budget managed at runtime, not a warehouse (and what *context rot* is);
2. **Compose** a system prompt in layers ordered by volatility (cache-aware);
3. **Design** a cascade of context files (global → project → package → personal) with declared precedence;
4. **Implement** harness-zero's context assembler (step 3) with a project rules file;
5. **Evaluate** a real AGENTS.md file against the authoring practices (lean, executable commands, grown by evidence of failure).

The problem

The model only knows what the harness shows it. "Context delivery" is the engineering of deciding **what** goes into each call — system prompt, project rules, environment state, memories, instructions from external servers — **in what order**, and **how that changes** mid-conversation without breaking the provider's cache or confusing the model.

The classic sub-problems: where project rules live and how they are discovered; whether the system prompt should vary by model; how to communicate state changes mid-conversation without invalidating the cached prefix.

Scientific foundations

- **Context degrades with position and with volume** — *Lost in the Middle* ([arXiv 2307.03172](#)): information in the middle of long contexts is poorly used. Design consequence: what matters goes to the edges (system prompt at the start; the current task at the end), and "send everything" is an anti-pattern with empirical backing.
- **Context engineering as a discipline** — the survey [arXiv 2507.13334](#) systematizes the area (RAG, memory, tool-integrated reasoning) and legitimizes the term the industry adopted.
- **Less context, better agents** — [arXiv 2606.10209](#) measures in long-running agents what Anthropic calls context rot: aggressive curation beats full windows.

(Full bibliography: [livro/bibliografia.md](#).)

Industry sources

- **Effective context engineering for AI agents** (Anthropic Engineering): names the successor to prompt engineering — the job is to **curate the optimal set of tokens at inference time**; names *context rot* as an engineering fact. Decision: the window is a budget, and the goal is the smallest set of high-signal tokens.
- **Prompt caching** (official docs) + **Lessons from building Claude Code: prompt caching is everything**: caching is **by prefix** — the context assembly order is a cost decision. The Claude Code account lists the classic invalidators (a timestamp at the top, a request ID in the tool list, history reserialization) and treats **cache hit rate as a first-class harness metric** (~59% reduction in billable input).
- **AGENTS.md + Agentic AI Foundation**: the "README for agents" was **donated to the Linux Foundation (Dec 2025)** with OpenAI, Anthropic and Block as co-founders; 60k+ projects. Decision: per-repository file context has become neutral, portable infrastructure — investing in that pipeline is safe.
- **How Claude remembers your project** (docs): formalizes the global → project → local **cascade**, with the closest file winning and the personal one kept out of version control.
- **AGENTS.md Field Guide 2026** (practitioner): authoring — start with ~30 lines, cap at ~150–200 at the root, exact commands before prose, nest per package in a monorepo, and **grow only on evidence of the agent's recurring failure** (the common mistake is

treating it as documentation).

- **See also:** the living collection [Awesome Harness Engineering — Context Delivery & Compaction](#) gathers more resources for this dimension (patterns, articles and implementations), curated by problem.

The state of the art

1. Context is a managed budget — and retrieval went just-in-time

The modern consensus inverted the "the more context, the better" instinct: the harness actively manages the window (rule-based pruning, awareness of how much remains, on-demand retrieval). The benchmark's two most advanced materializations: Aider's **repo-map** (the model "sees" the structure of an entire repository within a ~1k-token budget, via tree-sitter + personalized PageRank — static just-in-time retrieval with no explorer agent at all) and Goose's **incremental per-subdirectory hints** (rules loaded as the agent navigates, not all upfront).

2. Prefix stability became an architectural requirement

Cache-awareness stopped being an optimization and reorganized context assembly: layers ordered by volatility, deterministic serialization, zero volatile content at the top. The most rigorous formalizations measured: opencode's **Context Epochs** (the prefix as an immutable cache baseline, with state changes delivered only at safe turn boundaries) and Hermes's **explicit three-layer prompt** (`stable` → `context` → `volatile`), declaredly designed to maximize prefix-cache — including in the skill-curation fork, which inherits the parent's prefix to save ~26%).

3. The rules file standardized — and became a cascade

The AGENTS/CLAUDE/GEMINI.md fragmentation of the discipline's early days is resolved by neutral governance (Linux Foundation): AGENTS.md is the portable format, read natively by Codex, Goose, opencode, OpenClaw, Hermes, Aider and dozens of others, with the proprietary names becoming aliases. The mature pattern is the **cascade with declared precedence** (global → project → package → personal; the closest wins; the personal one gitignored), `@imports` for composition (gemini-cli) and — the authoring practice that separates useful files from dead documentation — growing **on evidence of failure**, like code.

4. The new frontiers

Three recent moves that have not yet become consensus: **per-model-family prompts** (opencode with ~10 variants; Codex taking it to the extreme with **server-driven** instructions — the backend delivers the per-model base prompt, with even a configurable "personality"); **persona × rules separation** (the personal-agent category's contribution: `SOUL.md` for voice/identity separate from the operational `AGENTS.md` — OpenClaw, Hermes, ohmo); and **trust-classed context** (IronClaw: personal/injected content travels in "prompt envelopes" with the trust class preserved — context delivery meeting the security of ch. 07).

The counterpoint: the minimal harness (Pi) — *addendum from round ext-1, 2026-07-31*. While this chapter describes ever-richer context assemblers, **Pi** (Earendil/Zechner, ~54k stars) bets in the opposite direction: a base system prompt **measured at ~460 tokens**, derived from the tool set (each tool contributes its snippet; guidelines enter only if the corresponding tool is active), and skills announced **by name+description only** — the body is loaded by the model itself via `read` when the task calls for it (progressive disclosure taken to the limit: there is not even a skill tool). Editorial honesty demands the two caveats the code reading revealed: (1) the same assembler concatenates the cascade's `AGENTS.md` files **with no budget** — in Pi's own repo this adds ~2,700 tokens, six times the slogan; the minimality is the harness's, not the context's; (2) minimalism is not absence of engineering — Pi's compaction is the most complete in the corpus (see the [evaluation](#), in Portuguese). The underlying bet is falsifiable and worth tracking: **better models would need less harness** — if true, part of this chapter expires; if the window stays expensive, the missing budget charges interest. It is the control experiment the corpus was missing.

EXECUTIVE SUMMARY

What is most modern: budget + just-in-time (not volume), stable prefix as a requirement (with cache hit rate as an SLI), cascading AGENTS.md under neutral governance, and the three frontiers (per-model/server-driven prompts, separate persona, trust class). The

minimalist counterpoint (Pi, round ext-1) shows the other end of the spectrum: a ~460-token prompt derived from the tool set — and proves the budget×richness tension remains open. **What to steal:** the repo-map as a cheap alternative to exploration; Hermes's 3 layers by volatility; the "grows on recurring failure" discipline in AGENTS.md authoring; from Pi, the prompt snippet coupled to the tool definition (prompt and tool set never desynchronize).

Hands-on — harness-zero, step 3

In step 3 you build harness-zero's context assembler: a system prompt in layers ordered by volatility (identity → environment → project rules → memory → task), discovery of an AGENTS.md at the target project's root, and a test that proves **prefix stability** across two consecutive turns (same bytes up to the last message). Completion exercise: the cascade discovery function comes skeletonized; you implement the precedence.

Check your understanding

1. Why is a timestamp at the top of the system prompt expensive — and where should it live? (Prefix caching + mid-conversation updates.)
2. Your agent ignores a project convention repeatedly. What is the right response according to modern authoring practice — and what is the wrong one? (Adding the rule to AGENTS.md on evidence × dumping documentation.)
3. A harness wants to tell the model the date has changed in the middle of a long conversation. Describe two strategies with different cache costs. (Epochs/turn boundaries × rewriting the prefix.)

Appendix A — How each repository handles context delivery

Per-harness evidence, with paths — online supplement, expanded with each benchmark round.

opencode (round 1) — typed algebra and Context Epochs

packages/opencode/src/session/system.ts assembles environment + skills + MCP (Model Context Protocol) instructions; ~10 prompts per model family in session/prompt/*.txt (anthropic, gpt, codex, gemini, kimi, beast...), selected by model-id substring; global/ancestor AGENTS.md aggregated by session/instruction.ts. V2 (CONTEXT.md) formalizes context as an algebra of "Context Sources" with snapshots, **Context Epochs** (cache baseline) and mid-conversation system messages only at safe boundaries.

gemini-cli (round 1) — hierarchy with @imports

prompts/promptProvider.ts assembles by mode/tools/model (modern × legacy snippets); hierarchical GEMINI.md (memoryDiscovery.ts: global → parents → subfolders) with @imports (memoryImportProcessor.ts) and flattenMemory; full override via GEMINI_SYSTEM_MD; just-in-time injection (tools/jit-context.ts).

OpenHarness (round 1) — aggregation with relevant memory

src/openharness/prompts/context.py: base + environment + CLAUDE.md + **memories selected by relevance** (memory/relevance.py, with usage.py tracking usage) + skills + active repo context; -s/- -append-system-prompt on the CLI.

Codex CLI (round 2) — central AGENTS.md + server-driven prompts

core/src/agents_md.rs: hierarchical discovery with merge from project-root to cwd; the system prompt **varies by model and comes from the backend** (ModelInfo.base_instructions via models-manager, with a template and Personality::Friendly/Pragmatic); environmental context via WorldState.

Goose (round 2) — incremental hints and hardening

SystemPromptBuilder with override + extras; multi-file hints (`.goosehints AND AGENTS.md`, `CLAUDE.md` via config) respecting `.gitignore`; **SubdirectoryHintTracker** loads subdirectory hints as the agent navigates; anti prompt-injection sanitization of Unicode tags; per-turn "top of mind".

Aider (round 2) — the repo-map ★

`aider/repomap.py`: definition/reference tags via tree-sitter (per-language `.scm` queries) → file → file graph → **personalized PageRank** (chat files and mentioned ids bias the ranking; $\times 10/\times 50/\times 0.1$ multipliers) → rendering under budget with binary search (~1024 tokens; `map_mul_no_files=8` with no files in the chat) → mtime-keyed cache. The entire context-first path in one file.

OpenHands/Canvas (round 2) — organizational skills

`app_conversation/skill_loader.py`: skills auto-discovered from the conventional repositories **owner/.openhands** and **owner/.agents** across all the user's organizations (GitHub/GitLab/Azure), with KeywordTrigger/TaskTrigger and a marketplace — team context versioned and loaded for all members.

OpenClaw (round 2) — identity workspace with budgets

`buildAgentSystemPrompt` injects `SOUL.md` (persona), `AGENTS.md` (rules), `USER.md`, `IDENTITY.md`, `TOOLS.md`, `MEMORY.md`, `HEARTBEAT.md`, `BOOTSTRAP.md` — with budgets (20k chars/file, 60k total) and marked truncation; provider-aware contributions **above/below the cache boundary**.

Hermes (round 2) — three layers by volatility ★

`agent/system_prompt.py` + `prompt_builder.py`: **stable** (identity/SOUL.md + guidance + skill index) → **context** (the project's AGENTS.md/cursorrules) → **volatile** (memory, USER.md, timestamp) — explicit design for prefix-cache; persona migratable from OpenClaw.

IronClaw (round 2) — context as a policy decision

`LoopPromptPort` (`crates/ironclaw_loop_host`): resolves identity, personal context (**opt-in per run profile, not per channel**), skills and security; injected/personal content travels in **prompt envelopes** with an unforgeable trust class — separating what the loop requests from what the host allows it to see.

ohmo (round 2.5) — the minimal correct version

`ohmo/prompts.py`: ordered concatenation base → soul → identity → user → BOOTSTRAP → workspace → memory; the rigorous decision `include_project_memory=False` (the personal agent does not read a project's CLAUDE.md — tested).

Pi (round ext-1) — the prompt derived from the tool set ★

`core/system-prompt.ts`: base **measured at ~460 tokens**, assembled from the tool definitions' own `promptSnippets` with dedup and guidelines conditional on the active set (deactivate the tool, the prompt shrinks); skills announced only as `<name/description/location>` and loaded by the model via `read` (block omitted if `read` is not active); `AGENTS.md/CLAUDE.md` cascade global → root → cwd with dedup of nested worktrees (`resource-loader.ts`) — yet concatenated **with no budget** (see the box in the chapter body); full override via `.pi/SYSTEM.md`.

n8n (round 2) — the embedded minimum

`ToolsAgent/common.ts`: `ChatPromptTemplate` with a free-form system message + history + rich binaries (images/PDF); no rules file and no hierarchy — the context comes mapped from the workflow by its author.

Frameworks (frameworks round) — open by design

LangGraph and the Agents SDK (Software Development Kit) leave assembly to the dev (static or callable instructions); CrewAI imposes role/goal/backstory as structural context; the software-agent-sdk provides a Jinja preset with a documented escape hatch (`prompt_dir + _prompt_preset()` → `None`).

04 — Compaction

🕒 state of the art 2026-07 · revised 2026-07-25 · 📖 ~13 min read

Learning objectives

By the end of this chapter, you should be able to:

1. **Explain** why compaction exists and which constraints it balances (fidelity × cost × cache);
2. **Compare** the four layers of the aggressiveness ladder and **justify** their ordering;
3. **Analyze** a real harness's compaction implementation and locate its choices on the ladder (Appendix A as the answer key);
4. **Implement** truncation with edge preservation and summarization with a preserved tail (step 5 of harness-zero);
5. **Evaluate** when a compaction has failed (loss of a decision, of file state or of the goal) — and **anticipate** what changes when the provider compacts for you.

The problem

Every agent conversation grows until it no longer fits in the model's context window. Compaction is the set of strategies for continuing to work when that happens — without losing what matters. It is the dimension where the evaluated harnesses converge the most: all of them arrived, independently, at the same layered architecture.

The constraints in tension:

- **Fidelity**: the summary cannot lose decisions, file state or the task's goal.
- **Cost**: summarizing via LLM (Large Language Model) is expensive; truncating is cheap but destructive.
- **Cache**: compacting invalidates the cached prefix — it should happen as little as possible and at controlled moments.

Scientific foundations

- **The window is not uniform** — *Lost in the Middle* ([arXiv 2307.03172](#)) showed that models use the beginning and end of the context best and degrade in the middle. It is the empirical basis for two of the ladder's practices: preserving the recent *tail* intact and truncating outputs while keeping start+end.
- **Context as virtual memory** — *MemGPT* ([arXiv 2310.08560](#)) framed the operating-systems analogy: the window is "RAM", external storage is "disk", and the harness pages between them. Recent work takes the analogy to its literal limit (*demand paging*, [arXiv 2603.09023](#)).
- **Compacting is a budget decision** — *ContextBudget* ([arXiv 2604.01664](#)) treats context management as explicit allocation per content type — what products implement as thresholds and budgets.

(Full bibliography and validation status: [livro/bibliografia.md](#).)

Industry sources

- **Compaction — Claude Platform Docs** (Anthropic, official): compaction has reached **the API level** (beta `compact-2026-01-12`) — the provider summarizes automatically upon hitting the configured threshold and returns a "compaction block". It is vendor confirmation of this chapter's central trend (see The state of the art).
- **Claude Code operating practices** ([CometAPI](#), [okhlopkov](#), [hyperdev](#)): the practitioners' convergent recommendation is the same one the harnesses encode — **what needs to survive compaction should not live in the conversation**: conventions go to the context file (CLAUDE.md/AGENTS.md, reinjected every session) and progress state goes to files the agent rereads after the compact. Compaction defines, by exclusion, what deserves persistence.
- **See also**: the living collection [Awesome Harness Engineering — Context Delivery & Compaction](#) gathers more resources for this dimension (patterns, articles and implementations), curated by problem.

The state of the art

The consolidated pattern: the aggressiveness ladder

Harnesses apply the strategies as a ladder, from cheapest to most expensive — this is the industry consensus, verified in every benchmark round:

1. **Truncate tool outputs at the source** — limit lines/bytes before they enter the history, preserving start and end (*Lost in the Middle* justifies the edges). The modern refinement: **do not discard** — move the full content to referenceable files (opencode) or keep the raw output outside the model's view but visible in the UI (Goose).
2. **Prune / microcompact** — erase the *content* of old tool results (the model rarely rereads a `cat` from 30 turns ago), keeping the record of the call. Newer intermediate layers: *tool distillation* and *output masking* (gemini-cli).
3. **LLM summarization (full compact)** — summarize the old portion while preserving an intact tail (typically 20–30% or a 2k–20k token budget). The state of the art has three refinements: a **structured summary** with mandatory fields (user intent, pending tasks, code state — Goose and software-agent-sdk), a **cheap auxiliary model** for the summary (Hermes), and a **memory flush before compacting** — saving durable notes before losing the context (OpenClaw).
4. **Automatic trigger + reactive path** — a trigger by window percentage (50–90% depending on the project) and, covering the failure case, compaction **reactive** to the API's "prompt too long" error (OpenHarness, OpenClaw).

The two modern frontiers

1. **Auditable compaction (tombstones)**. The most advanced implementation measured in the benchmark (the software-agent-sdk's condenser) does not mutate the history: the log is append-only and forgetting is an *event* (**Condensation**) — a tombstone, as in Cassandra/Kafka. The model's view is derived by applying the tombstones; nothing is lost to auditing, and formal invariants (tool_call/result pairing, batch atomicity) are **testable code**, with the *hard/soft trigger* distinction: if compacting now would violate an invariant, the soft trigger waits for the next turn; the hard one forces an explicit reset. A related refinement: the **effectiveness circuit-breaker** (IronClaw) — comparing the post-compaction estimate against a baseline and detecting compactions that are not working.

2. **Compaction is migrating to the provider**. (And caching is becoming a protocol contract too: the MCP 2026-07-28 spec added `ttlMs/cacheScope` to `tools/list` responses — the protocol taking over what used to be harness heuristics.) Two independent signals in the same year: the Codex CLI implements **remote compaction v2** (the backend compacts) and Anthropic launched **compaction in the API itself** ([docs](#), beta `compact-2026-01-12`). It is the expiration clause in motion — but with an interesting inversion: instead of the component disappearing when the model improves, it **changes owner** (from the harness to the platform). What remains for the harness when the provider compacts: deciding *what to protect* (skills, task state, memory files), *when to trust* (auditing the summary's quality — OpenClaw's `safeguard` mode anticipated this) and the reactive path for providers that do not offer the service.

Addendum (2026-07-31, full text verified): the third way — compaction learned in training. The preprint [CompactionRL](#) (Tsinghua/Z.AI, 06 Jul 2026) proposes the migration's next step: training the model via RL **with compaction inside the loop** — "CompactionRL incorporates compaction into rollout collection, and reconstructs the agent context from a summary once context budget is exhausted" (§1); summarization becomes "a learned part of the model rather than an inference-time heuristic", with a **task-level reward**. The numbers (Table 2, always against the same model *already using inference-time compaction*): GLM-4.5-Air **59.8 → 66.8** on SWE-bench Verified (+7.0) and +3.1 on Terminal-Bench 2.0; GLM-4.7-Flash **+5.5 and +6.8**. And the experiment's protocol is exactly this chapter's ladder — a threshold by remaining budget, a structured summary from a fixed prompt, a **preserved tail of k=2 steps** — that is, the paper validates the triad and changes the *training*, not the architecture. Three consequences: (1) the harness remains the owner of the *when*, but the *how to summarize* is starting to migrate into the weights — harness ↔ model mismatch becomes a new risk; (2) the declared limitation is revealing: "its gains do not consistently transfer to single-window evaluation when compaction is disabled. This indicates a train–test mismatch" — trained compaction creates *coupling* (with compaction turned off, the trained GLM-4.7-Flash actually gets worse, 47.5 → 43.7), the strongest argument so far for an explicit *compaction contract* between harness and model; (3) in the other direction, Table 1 hands power back to the harness: with the executor fixed, **swapping only the summarizer** moves SWE-Verified from 49.0 to 55.5 (+6.5) — "compaction is a performance-critical decision process rather than a passive preprocessing step", and a better dedicated summarizer **beats self-summarization**: choosing who summarizes is a harness decision, and a big one.

The third frontier: compaction stops being involuntary (round ext-4, 2026-08)

The launch of **Prime Agent** (Prime Intellect, Aug/2026) came with a charge aimed straight at this chapter: *"fixed tool-calling schemas and context compaction force the model to work around its own scaffolding instead of leveraging it."* Reading the code ([full evaluation](#)) shows that **the charge is rhetoric and the code says something else** — and the gap between the two is the finding.

Compaction was **neither removed nor weakened**. Prime Agent is built on Pi, and the 1,398 lines of `core/compaction/` are there intact — safe cutting, split turns, cumulative files, reactive overflow recovery — and even improved, with custom instructions and `tokensBefore` recomputation. What changed is **who is in charge**: `compact.run()` and `compact.status()` became callable **by the agent itself** (`skills/compact/`), with a handler that **schedules instead of executing** — executing on the spot would abort the very REPL cell that requested the compaction — and that runs even with automatic compaction turned off, under twelve test cases. Add to that the session's JSONL path injected into the system prompt: the full history, **including previous compactions**, remains programmatically reachable.

The caveat to record in this chapter is therefore precise: **compaction stops being an involuntary event of the harness and becomes one mechanism among others, available to the agent**. It also gains a new role — it became a distillation trigger, with `autoRefine.compact: true` by default: every compaction is an opportunity for the agent to extract learning from what is about to be summarized.

What the aggressiveness ladder did not anticipate is not its own obsolescence but the **inversion of control**: until now, the harness compacts the agent; here, the agent compacts *itself*. The gap the reading found is telling — the announcement mentions a subagent acting as a garbage collector for the REPL, and **there is nothing of the sort in the code** (searching for `garbage/prune/evict` in `src/core`, `skills` and `prime-agent-runtime` returns nothing). Context-as-a-variable solves access to the past; it does **not** solve the growth of the namespace it creates.

EXECUTIVE SUMMARY

Convergence on the ladder is nearly total — the pattern is consolidated, and a new harness that does not implement it needs to justify itself. The remaining differences are fidelity refinements (structuring the summary, auditing its quality, never discarding) and the big open question is one of *market architecture*: how much of the ladder survives in the harness when the platform offers compaction as a service — a question the addendum above sharpens: after migrating to the provider, compaction is starting to migrate **into the weights**. **What to steal** today: tombstones over an append-only log; pre-compaction memory-flush; a structured summary with task IDs preserved; the effectiveness circuit-breaker; and — new in ext-4 — **agent-callable compaction that schedules instead of executing**, plus the session log path in the system prompt, which puts the past back within the model's reach without spending window.

Editorial caveat (2026-08-06). This Executive summary was confronted in round ext-4 and **upheld**, with the qualification in the previous section: the ladder is still the pattern, but *authority* over when to apply it has begun migrating to the agent. If the pattern repeats in other harnesses, the synthesis changes — and this paragraph will be rewritten, not patched.

Hands-on — harness-zero, step 5

In step 5 of the project (`harness-zero/`), you implement the ladder in your own harness, in this order: (1) tool output truncation with start+end preservation; (2) pruning of old tool results beyond a budget; (3) LLM summarization of the history's head, preserving the tail; (4) automatic triggering by an estimated-token threshold — with a **visible indicator in the chat** when compaction happens (the reader's observation window). Completion exercise: the prune function's skeleton comes ready; you write the selection of what to protect.

Check your understanding

1. Why truncate tool outputs **before** summarizing via LLM, and not the other way around? (Cost and destructiveness — if needed, reread the ladder.)

2. A harness summarized the history and the agent, on the next turn, rewrote a file that was already correct. What information did the compaction probably lose, and which state-of-the-art mechanism prevents it? (Hint: structured summary with `CODE_STATE/CHANGES`.)
 3. Your provider now offers compaction in the API. Which of the ladder's responsibilities do you **transfer** and which do you **keep** in the harness? (Connect with "the two modern frontiers".)
-

Appendix A — How each repository handles compaction

Per-harness evidence, with paths — supplementary material (online version), expanded with each benchmark round. The chapter's base source: the code of these repositories.

opencode (round 1) — three mechanisms + managed files

`packages/opencode/src/session/compaction.ts` (+ `overflow.ts`, `summary.ts`): (a) automatic summarization on overflow with a **dedicated compaction agent**, tail under budget (`preserveRecentBudget`, 2k–8k tokens), a new Context Epoch and optional auto-continue; (b) back-to-front **prune** marking tool outputs beyond 40k tokens as **compacted** (`PRUNE_PROTECT`), protecting skills; (c) truncation at the source (`tool/truncate.ts`) preserving start+end and moving the full text to "Managed Tool Output Files".

gemini-cli (round 1) — compression + distillation + masking

`packages/core/src/context/chatCompressionService.ts`: fires at 50% of the limit (`DEFAULT_COMPRESSION_TOKEN_THRESHOLD = 0.5`), preserves the last 30% (`COMPRESSION_PRESERVE_THRESHOLD`), its own budget for function responses (50k) and saving of truncated outputs. Extra layers: `toolDistillationService.ts` and `toolOutputMaskingService.ts`. Manual `/compress`, `ChatCompressed` event, `PreCompressTrigger` hooks.

OpenHarness (round 1) — the faithful translation of Claude Code

`src/openharness/services/compact/__init__.py` (1,725 lines; docstring: "Faithfully translated from Claude Code's compaction system"): **microcompact** (clears `COMPACTABLE_TOOLS`), **full compact** (LLM summary), **auto-compact** (threshold) and compaction **reactive** to "prompt too long" (`_is_prompt_too_long_error`). `PRE_COMPACT/POST_COMPACT` hooks; preserves task state and channel logs.

Codex CLI (round 2) — local + remote v1/v2

`core/src/compact.rs`, `compact_remote_v2.rs`, `compact_token_budget.rs`: auto-compact at ~90% of the window; three strategies — local (`SUMMARIZATION_PROMPT`) and **remote v1/v2** (the backend compacts, via `ResponsesStreamRequest:RemoteCompactionV2`, with its own retry); versioned windows with prefill tracking; controlled pre/mid-turn injection; `TruncationPolicy` for outputs.

Goose (round 2) — structured summary + middle-out

`crates/goose/src/context_mgmt/mod.rs`: threshold at 0.8 of the window; `StructuredSummary` (`user_intent`, `files`, `pending_tasks`, `current_work`); if summarization overflows, **progressive "middle-out" removal** of tool-responses (0 → 100%); **incremental summarization of tool-call/response pairs** in batches of 10 protecting the last N; visibility metadata preserves the raw output in the UI; respects `provider.manages_own_context()`.

OpenClaw (round 2) — safeguard + memory flush

`src/context-engine/` + `docs/concepts/compaction.md`: automatic by threshold and reactive (recognizes dozens of overflow error strings from multiple providers), split preserving tool-call/result pairs; **safeguard** mode with **summary quality auditing**; **silent memory flush before compacting**; `keepRecentTokens` 20k; pluggable compaction providers; the compaction (semantic) × pruning (in-memory trim) distinction.

Hermes (round 2) — pluggable engine + auxiliary model

`agent/context_engine.py` (interface `should_compact/compress/prune`) + `trajectory_compressor.py` (~1.6k lines): summarization of old tool-responses via a **cheap auxiliary model** (default Gemini Flash, up to 50 concurrent requests); manual `/compress`; `/usage` and `/insights` expose the window.

IronClaw (round 2) — pure policy + circuit-breaker

`crates/ironclaw_agent_loop/src/strategies/compaction.rs` (+ `active_task_compaction.rs`): the strategy is **pure policy** (returns `Skip` or the `drop_through_seq` limit; mutation only in the host); `PromptContextTokenBudget` with `preserve_tail_tokens`; an **effectiveness circuit-breaker** (compares the post-compaction estimate against `CompactionEffectivenessBaseline`); a variant that preserves the active task; the host refuses to compact through non-user messages.

software-agent-sdk (frameworks round) — tombstones + testable invariants ★

`openhands-sdk/openhands/sdk/context/condenser/`: forgetting via **tombstones** (`Condensation` event) over an append-only log; triggering for three reasons (`REQUEST/TOKENS/EVENTS`) with **hard/soft** (`condensation_requirement`) and `hard_context_reset()` for the pathological case; `keep_first` + recursive re-summarization of summaries; a structured prompt (`summarizing_prompt.j2`: `USER_CONTEXT`, `TASK_TRACKING` with exact IDs, `CODE_STATE`, `TESTS`, `CHANGES`); invariants in `context/view/properties/` (`tool_call_matching`, `batch_atomicity`...) **tested against real LLMs** (`tests/integration/tests/c01..c05`); `pipeline_condenser` for composition.

Aider (round 2) — classic summarization done well

`aider/history.py` (`ChatSummary`): keeps the tail (~half the budget), summarizes the head via LLM with a split after an **assistant** message, **recursive** up to depth 3, with a fallback model list.

n8n (round 2) — the absence that confirms the category

No compaction in the loop (the memory sub-nodes' `contextWindowLength` + `maxTokensFromMemory` only) — consistent with short, event-triggered executions; it is the "embedded harness" category's ceiling for long tasks.

LangGraph / OpenAI Agents SDK / CrewAI (frameworks round) — the dividing line

LangGraph: **zero native support** (a docstring suggesting `pre_model_hook`); Agents SDK (Software Development Kit): only `OpenAIResponsesCompactionSession` as an optional session; CrewAI: nothing. Compaction is the dimension that most separates "framework" from "ready-made harness".

05 — Tool Design

🕒 state of the art 2026-07 · revised 2026-07-25 · 📖 ~7 min read

Learning objectives

By the end of this chapter, you should be able to:

1. **Explain** why a tool's description is prompt engineering, not API documentation;
2. **Derive** a tool's schema from types (and justify why nobody writes JSON Schema by hand anymore);
3. **Compare** the three scaling regimes — fixed catalog, tool search with deferred loading, and code-as-action;
4. **Implement** harness-zero's `ToolPort` with derived schema and error-as-data (step 2);
5. **Evaluate** when to use individual tool calls versus code orchestrating tools in a sandbox.

The problem

Tools are the agent's "hands": the contract through which the model acts on the world. Tool design means deciding **which** tools exist, **how** their parameters are described to the model, **how** results (and errors) come back, and **when** each one is available. A poorly described tool produces wrong calls; an oversized arsenal dilutes the model's attention *and* blows the context budget before any useful work; an undersized arsenal forces workarounds through the shell.

Scientific foundations

- **The evolution of tool use** — [arXiv 2603.22862](#) traces the trajectory from single-tool calls to multi-tool orchestration, the backdrop of "code-as-action".
- **Tool learning as a field** — the tool learning survey ([repo](#)) organizes how agents learn to select and compose tools.

(Full bibliography: `livro/bibliografia.md`.)

Industry sources

- **Writing effective tools for AI agents** (Anthropic Engineering): the canonical source — tools are "contracts between deterministic systems and non-deterministic agents"; the description is prompt engineering (small refinements → large accuracy gains), the return value should be optimized for **informational density per token**, and the cycle is *prototype* → *evaluate* → *collaborate* (the model itself rewrites the tools from eval transcripts).
- **Code execution with MCP** (Anthropic): loading every definition and passing intermediates through the context is the bottleneck — exposing each tool as a TypeScript file that the agent orchestrates via code took one case from **~150,000** → **~2,000 tokens** (−98.7%).
- **Tool search tool** + **Advanced tool use** (docs + blog): dynamic discovery — send everything, mark the non-critical with `defer_loading: true`, the model sees only the search plus the essentials; one multi-server setup spends ~55k tokens on definitions before doing any work, and tool search cuts that by >85%.
- **Programmatic tool calling** (docs): the model writes Python that calls the tools in a sandbox and returns only the distilled result — ~38% fewer input tokens on a 75-tool benchmark; 20–40% typical in production with 10–49 tools.
- **Code Mode** (Cloudflare): the same thesis, from an infrastructure vendor — the argument is about *training distribution*: LLMs write code against known APIs better than they fill in synthetic schemas. Industry convergence, not one vendor's quirk.
- **Apply Patch** + **GPT-5.1 for developers** (OpenAI): an editing tool **trained into the model** (the V4A diff format) — which explains why ad-hoc search/replace formats lose to the format the model saw in training.
- **See also**: the living collection [Awesome Harness Engineering — Tool Design](#) gathers more consultable resources for this dimension (patterns, articles, and implementations), curated by problem.

The state of the art

1. The consensual core — and type-derived schema won

The harnesses converge on a core of ~10 tools (read/write/edit file, glob, grep, shell, web fetch/search, todo, delegate) — the minimum kit of a coding agent. And nobody writes JSON Schema by hand: the source of truth is the type system (Pydantic in OpenHarness/Hermes, Effect Schema in opencode, declarative classes in gemini-cli, generic dataclasses in software-agent-sdk). The modern quality refinement: separating **what goes back into the model's context from the structured data** — software-agent-sdk's `Observation.to_llm_content` is the cleanest design (you control exactly the informational density Anthropic preaches).

2. Tool context became a scarce resource — three scaling regimes

The default of "dumping every definition into the system prompt" is dead. The state of the art has three regimes, and the choice is driven by catalog size:

- **fixed catalog** (dozens of tools): still fine to send everything;
- **tool search / defer_loading** (hundreds of tools, many MCP servers): keeps 3–5 tools hot, loads the rest on demand — present as `tool_search/tool_discovery` in Codex, Tool Search in OpenClaw, `tool_search` in OpenHarness;
- **code-as-action** (pipelines with bulky data): the model writes code that orchestrates the tools in a sandbox and returns the distilled result — `code-mode` (opencode with embedded V8, Codex likewise), `execute_code` (Hermes calling tools via RPC (Remote Procedure Call) in "zero-context-cost turns"), Code Mode (Goose). The metric the industry now reports is not standalone accuracy, it is **accuracy per definition token**.

3. The editing interface is trained, not invented

The most counterintuitive lesson: the best code-editing format is not the one you design, it is the one the **model saw in training**. Hence `apply_patch` (V4A) being a native OpenAI tool, opencode giving GPT models `apply_patch` instead of `edit/write`, and Aider empirically measuring which format each model applies well (`percent_cases_well_formed`). Corollary: tool selection **varies by model family** — an explicit acknowledgment that the ideal interface depends on who is on the other side. And tool errors come back as **data** (so the model can self-correct), not as exceptions.

EXECUTIVE SUMMARY

What is most modern: type-derived schema with data×context separation; the three scaling regimes (fixed → tool search → code-as-action) chosen by catalog size; and the editing interface as something trained. **What to steal:** `to_llm_content` (per-token density control); tool search with `defer_loading`; measuring the editing format per model (Aider's `percent_cases_well_formed`); error-as-data.

Hands-on — harness-zero, step 2

Step 2 replaces step 1's hand-written schemas with a `ToolPort`: a tool is a typed function, and the schema is **derived from the annotations** (via `inspect/typing`, reading signature and docstring). You add `read_file` alongside `get_time/somar`, with errors returning as text to the model (never as an exception that crashes the loop). Completeness exercise: the schema deriver ships skeletoned for one parameter; you extend it to composite types.

Check your understanding

1. Why is a tool's description prompt engineering and not API documentation? (Informational density; iterating over eval transcripts.)
2. Your agent has access to 8 MCP servers (200+ tools) and spends 55k tokens before acting. Which scaling regime do you adopt, and what does it keep hot? (Tool search + `defer_loading`.)
3. Why can giving a model `apply_patch` beat a search/replace format you designed carefully? (Training distribution.)

Appendix A — How each repository handles tools

Per-harness evidence, with paths — online supplement, expanded each round.

opencode (round 1)

~14 tools + 3 experimental (`tool/`), Effect Schema, separate `.txt` descriptions; **per-model selection** (`registry.ts`: GPT gets `apply_patch` instead of `edit/write`); embedded ripgrep; experimental `lsp`, `plan_exit`, `code-mode` (V8).

gemini-cli (round 1)

~20–25 tools as declarative classes (`BaseDeclarativeTool` + `Invocation`), filtered registration (`maybeRegister`), declarations per model family; shell with background processes, web search with grounding, optional tracker (6 tools).

OpenHarness (round 1)

43+ **tools** (`tools/`, `BaseTool` + Pydantic `input_model` → `to_api_schema()`); `is_read_only()` feeds the loop's parallelism; multimodal, cron, teams, `tool_search`.

Codex CLI (round 2)

`tools/` crate with typed schemas; `unified_exec` (persistent shell with stdin); **first-class `apply_patch`** (streaming parser + `apply_patch.lark` grammar, varying by model); `tool_search/tool_discovery`; **code-mode with embedded V8**.

Goose (round 2) ★ MCP-native

Every tool is MCP: `goose-mcp` built-ins are `rmcp:ServerHandler` served in-process over `DuplexStream`; even `developer/shell/edit` are "platform extensions" speaking `McpClientTrait`.

OpenClaw (round 2)

Broad suite (`openclaw-tools*.ts`): runtime/files/web/CDP browser/media; **Tool Search** and **Code Mode** (JS/TS over a hidden catalog); 52 `AgentSkills` injected as a compact block, read on demand.

Hermes (round 2)

~40+ tools in **composable toolsets** with dynamic postures; `execute_code` (Python calling tools via RPC, "zero-context-cost turns"); per-provider `schema_sanitizer`.

Aider (round 2) ★ edit formats

Instead of JSON tools, **edit formats** (`*_coder.py`): whole/diff (fuzzy SEARCH-REPLACE)/udiff/patch; per-model selection; **benchmark-validated** (`percent_cases_well_formed`).

software-agent-sdk (frameworks round) ★ data×context

Action/Observation/Executor contract; `Observation.to_llm_content` separates what goes back to the model from the structured data; toolsets (one `create` → several tools); MCP-style annotations; `ClientToolSpec` (tool executes on the client's machine).

IronClaw (round 2)

Tools as **capabilities with typed descriptors** declaring `EffectKind`, credentials and network policy; visibility × authority separation (a hidden capability fails closed); obligations (redaction/limits) before any effect.

n8n (round 2)

`create-node-as-tool.ts`: any **usableAsTool** node becomes a tool via `$fromAI('chave', 'desc', tipo)` → derived Zod schema; ToolWorkflow (sub-workflow as tool), ToolHttpRequest, ToolCode, ToolThink.

Frameworks (frameworks round)

Agents SDK (Software Development Kit): `@function_tool` (Pydantic + griffe with docstring auto-detection), 13 types incl. hosted; LangGraph: inherits `@tool` from langchain-core, adds `ToolNode` (execution, injections); CrewAI: Pydantic `BaseTool/@tool`, `crewai-tools` catalog with 79 directories.

06 — MCP

🕒 state of the art 2026-07 · revised 2026-07-31 · 📖 ~17 min read

Learning objectives

By the end of this chapter, you should be able to:

1. **Explain** why MCP (Model Context Protocol) became the lingua franca of agent integration — the open-standard argument against the $M \times N$ cost of point-to-point integrations;
2. **Compare** the protocol's transports (stdio, Streamable HTTP, deprecated SSE (Server-Sent Events)) and decide which to use for local versus remote servers;
3. **Evaluate** the protocol surface (tools, resources, prompts, roots, sampling, elicitation) and what a mature client needs to support;
4. **Recognize** the MCP server as an attack surface — a tool's description is untrusted input — and name the containment defenses;
5. **Implement** the MCP client adapter (stdio) behind a port in harness-zero (step 7).

The problem

No harness can embed tools for every system in the world — databases, issue trackers, browsers, internal APIs. Without a standard, each harness would write N integrations and each tool would be rewritten for M harnesses: the classic $M \times N$ problem. MCP solves it via the **open standard** route: a server exposes *tools*, *resources* and *prompts* over a common protocol (JSON-RPC (Remote Procedure Call) 2.0), and any client harness consumes them without knowing who implemented them. Each side writes once — $M+N$ instead of $M \times N$.

In little more than two years this became industry consensus — in the studied cohort, **10 of the 11 harnesses in the cohort** are full MCP clients, all built on the protocol's official SDKs. The decisions that still differentiate the implementations:

- **Transports:** stdio (local process), Streamable HTTP (remote), and legacy SSE.
- **Authentication:** OAuth for remote servers — with which flows and providers?
- **Resilience:** reconnection, unavailable servers, dynamic changes to the tool list.
- **Surface:** only *tools*, or also *resources*, *prompts*, *roots*, *sampling*, *elicitation*?
- **Role:** is the harness client-only, or also a **server** — consumable by other agents?
- **Security:** an MCP server is third-party code injecting text into the model's context. Who treats that as an attack surface?

Scientific foundations

MCP was born as an **industry specification**, not from a paper — and the academic literature that caught up with it concentrates, revealingly, on **security**. The design decision this whole literature supports is a single, decisive one: **a tool's description (and its return value) from an MCP server is untrusted input**, and must be treated with the same skepticism as any external content.

- **Indirect prompt injection** — Greshake et al., [arXiv 2302.12173](#) (AISec '23), the paper that defined the threat: LLM (Large Language Model)-integrated applications blur the boundary between *data* and *instructions*, so any retrieved content is a potential instruction channel. Translated to the MCP client: a tool's description field and the text the server returns are **data**, never trusted instructions.
- **The MCP systematization** — Hou et al., "MCP: Landscape, Security Threats, and Future Research Directions", [arXiv 2503.23278](#) (also in ACM TOSEM): the canonical SoK. It decomposes the server lifecycle (creation → deployment → operation → maintenance) and shows that the **same server is attackable in different phases** — spoofing at installation, *tool poisoning* at runtime. Decision: the harness needs **per-phase** trust boundaries, not a single gate.
- **Tool description as a vector, measured** — MCPTox, [arXiv 2508.14925](#), the first *tool poisoning* benchmark over 45 real servers / 353 tools: success rates of up to ~73%, and — the uncomfortable finding — **more capable models were more susceptible**, with safety alignment offering minimal protection before execution. Hard decision: you cannot trust the model to self-filter; the tool's metadata has to be blocked **before** it enters the context window.

- **The base rate is empirical, not hypothetical** — "[MCP at First Glance](#)", [arXiv 2506.13538](#) audited 1,899 open-source servers: **7.2% with general vulnerabilities and 5.5% with MCP-specific tool poisoning**, in classes that only partially overlap with traditional appsec. Decision: assume a non-trivial base rate of poisoned servers in the real world; MCP-aware scanning, not just SAST. (See also [MCP Safety Audit](#), [arXiv 2504.03767](#), which shows code execution and credential theft exploits via *legitimately registered* tools.)
- **Choose the protocol by trust context** — the [interoperability survey](#), [arXiv 2505.02279](#) compares MCP, ACP, A2A and ANP: MCP assumes a **relatively trusted** client-server boundary; exposing MCP tools across organizational boundaries does not inherit A2A/ANP's identity guarantees and requires additional authn/authz (ties into ch. 17).

Editorial note (living book): this was the book's most rarefied bibliography dimension — recorded as an "academic gap". Between rounds it matured from gap to **consolidated security literature** (an SoK, benchmarks, empirical audits). The migration is noted in [bibliografia.md](#).

(Full bibliography and pointers: [livro/bibliografia.md](#).)

Industry sources

- **MCP architecture** (official spec): defines what the harness needs to map — on the server side, *tools/resources/prompts*; on the client side, *roots/sampling/elicitation*. The design decision: separating what the server *offers* from what the client *grants* — **sampling** and **roots** exist so the server can request an inference or a file scope **without ever having direct access** to the model or the filesystem, keeping the host as the single point of trust.
- **Transports** (spec): two transports over JSON-RPC — **stdio** (local, no network overhead, the default for local servers) and **Streamable HTTP** (remote). The old **HTTP+SSE was deprecated in the 2025-03-26 revision** and survives only for backwards compatibility — the modern client tries `POST InitializeRequest` first and only falls back to SSE on a 4xx.
- **Introducing the Model Context Protocol** (Anthropic, Nov 25, 2024): the announcement that opened the standard, with the "**USB-C for AI**" analogy (one connector, many peripherals) — which is exactly the harness's M×N argument. (*anthropic.com returns 403 through the proxy; date and framing confirmed via VentureBeat*.)
- **Adoption as the turning point**: OpenAI adopted MCP in [Mar 2025 \(Agents SDK \(Software Development Kit\), TechCrunch\)](#); [Google/Gemini followed \(The New Stack\)](#); [Microsoft brought MCP to GA in Copilot Studio](#) and to Windows. Key decision: when the second-largest lab adopts its competitor's protocol, MCP stops being a vendor bet and becomes **neutral infrastructure** — designing for MCP reduces lock-in risk.
- **Authorization (OAuth 2.1)**: the spec treats every remote server as an **OAuth 2.1 Resource Server** — validating tokens issued by an external Authorization Server (RFC 9728 + 8414 + 7591). [Descopes's practical guide](#) translates it into a decision: separating *who serves the tool* from *who issues identity* enables corporate SSO and per-resource scoped tokens, instead of credentials embedded in the server.
- **Security in practice** — [Invariant Labs' Tool Poisoning Attack](#) (Apr 1, 2025) coined the term: malicious instructions hidden in a tool's *description* that the user never reads but the model obeys. [Trail of Bits showed "line jumping"](#): the mere **registration** of a server is already attack surface, before any invocation — the trust gate has to be at *connect*, not at *call*. And [the lethal trifecta \(Simon Willison\)](#): private data + untrusted content + external communication — MCP makes it far too easy to glue together tools that, combined, close all three corners (read email + open a public PR = exfiltration). The design rule is to prevent the three from coexisting in the same loop.
- **Governance (MCP became a boundary, not a prosthesis)**: the [official registry](#) (preview, Sep 2025) is a *community-owned* API layer — the harness discovers servers via a standardized API, not hardcoded lists. And in Dec 2025 Anthropic [donated MCP to the Agentic AI Foundation, under the Linux Foundation](#) — alongside [goose](#) and [AGENTS.md](#) as founding projects. The protocol now evolves by consensus of a steering group, not by one vendor's roadmap.
- **The 2026-07-28 Specification** (official MCP blog): the announcement of the protocol's biggest revision — stateless core, MRTR, extensions, caching and a deprecation policy (see §6 of the state of the art).
- **See also**: the living collection [Awesome Harness Engineering — Skills & MCP](#) gathers more consultable resources for this dimension (patterns, articles, and implementations), curated by problem.

The state of the art

1. The discipline's clearest standardization

Eleven harnesses in the cohort, several languages (TypeScript, Rust, Python), the same protocol, the official SDKs. It is the most limpid case of convergence the book has recorded: where tool design, loop and compaction diverge, MCP unified. The single exception in the cohort is **Aider** (MCP score 0) — and it is a *philosophical* exception, not a lag: the *context-first* school bets on curated context and edit formats, and forgoes MCP on purpose.

2. Transports converged — stdio local, Streamable HTTP remote

The default has stabilized: **stdio** for local servers (the harness launches the process), **Streamable HTTP** for remote ones. **SSE** became legacy — present only as a compatibility fallback (opencode does *automatic HTTP → SSE fallback*). More rigorous harnesses pin the **protocol revision** (IronClaw explicitly speaks 2025-06-18), a sign that the protocol has versions and the client needs to negotiate them.

3. The turn: the harness became an MCP server too

This is the chapter's strongest *dated update* — and a **prediction by the book itself that expired**. In the early rounds we noted that "none of the harnesses acts as an MCP server in the core; harness-as-a-service shows up via A2A/ACP". Round 2 refuted that: **Codex, Hermes, OpenClaw, OpenHands and n8n expose themselves as MCP servers**. The harness stopped being merely a consumer of tools and became a **piece consumable by other agents** — Codex exposes itself to IDEs and other hosts; OpenClaw serves its channel conversations to Claude Code/Codex; n8n publishes its workflow graph as an MCP endpoint. With that, the protocol surface widens beyond *tools*: **sampling** (the server requests completions from the client — Hermes) and **elicitation** (the server requests structured input — Codex) enter the state of the art. Harness-as-a-service, which used to exist only via A2A/ACP, now has a native MCP route.

4. Authentication: OAuth 2.1 is the floor; enterprise raises the bar

For remote servers, the OAuth flow with PKCE + local callback + token storage became the minimum (opencode, gemini-cli, Codex, Hermes, OpenClaw). The competitive differential sits above it: **gemini-cli** adds **Google auth** providers and **service account impersonation** (MCP designed for corporate GCP); **OpenClaw** stores tokens in SQLite and supports **mTLS**. Enterprise authentication is today's MCP feature frontier.

5. Security: the MCP server is third-party code — and the cohort has started treating it that way

If an MCP server injects text into the context and can see tool arguments, it is attack surface — and the literature (above) shows the threat is measurable and common. The defenses observed in the cohort, in layers (all connect to ch. 07):

- **Test the vector:** gemini-cli includes a **prompt injection via MCP** eval — the only one that treats the server as a *tested* attacker.
- **Filter the environment:** OpenClaw, when launching a stdio server, **blocks dangerous environment variables** (`NODE_OPTIONS`, `LD_*`, `DYLD_*`) that would allow loading code into the process.
- **Mediate credentials:** IronClaw adapts MCP tools into *capabilities without granting ambient authority* — FS, secrets and network remain mediated, and **the server never sees the secret** (the credential is injected at the edge). OpenHands does **secret redaction and restoration** in configuration round-trips.

The bar rose from "connect a server" to "connect a **contained** server" — exactly what the *tool poisoning* and *line jumping* papers call for: a gate at connect time, sanitization between servers, and no trust in the model's self-filtering.

6. The stateless turn — the 2026-07-28 spec

Three days before this revision, the protocol went through its **biggest change since launch** ([official announcement](#); [changelog](#)). The core became **stateless**: the `initialize/notifications/initialized` handshake and the `Mcp-Session-Id` header are gone — each request travels independently, with protocol, identity and capabilities in `_meta` (plus an optional `server/discover` for discovery). The motivation is infrastructural: MCP servers can now scale behind an ordinary round-robin load balancer, without *sticky sessions*. Server-initiated requests (`elicitation/create`, `sampling/createMessage`, `roots/list`) give way to **MRTR (Multi Round-Trip Requests)**: the server responds `resultType: "input_required"` and the client retries with `inputResponses` — bidirectionality becomes an explicit round trip. Rounding out the package: a **formal extension framework** (Tasks becomes

`io.modelcontextprotocol/tasks`; MCP Apps and Enterprise Managed Authorization as extensions), **caching as a contract** (`ttlMs/cacheScope` in listing responses — see ch. 04), header-based routing (`Mcp-Method/Mcp-Name`) and the **first formal deprecation policy** (a minimum 12-month window) — under which **Sampling, Roots, Logging, the legacy HTTP+SSE transport and DCR (Dynamic Client Registration, replaced by CIMD — Client ID Metadata Documents)** are deprecated. Editorial reading: what sections 1–5 describe is still the protocol *installed* in the cohort (the 12-month window exists for exactly that), but the direction has changed — and harness adoption of 2026-07-28 is the number one item to measure in the next benchmark round.

EXECUTIVE SUMMARY

The protocol turned a page on 2026-07-28: a **stateless** core (no handshake, no `Mcp-Session-Id`), MRTR in place of server-initiated sampling/elicitation, formal extensions, caching as a contract (`ttlMs`) and the first deprecation policy (12 months) — under which Sampling, Roots, Logging and the HTTP+SSE transport fall. What the cohort *runs today* is still the protocol of sections 1–5 (the window exists for that); what you *write today* should already target 2026-07-28. **What to steal:** treat the tool description as untrusted input (the literature measures ~73% *tool poisoning* success); in new code, prefer stateless Streamable HTTP (the SSE fallback is now a deprecated transport); if you expose an MCP server, filter the subprocess environment and never let the server see secrets; if your audience is enterprise, OAuth 2.1 with impersonation is the floor — and plan the DCR → CIMD migration within the window.

Hands-on — harness-zero, step 7

Step 7 (`harness-zero/etapas/07-mcp/`) gives harness-zero an **MCP client adapter (stdio)** behind a port. Faithful to hexagonal architecture *by refactoring*: step 2's `ToolPort` already defines what a tool is; now an adapter discovers tools from an external MCP server (via `stdio`, launching the process) and presents them to the loop as native tools — the model cannot tell the difference. You connect the included example server (`servidor_mcp_exemplo.py`) — and, as an extension, any real filesystem MCP server. Period note: the step's `ClienteMCP` implements the 2025-06 protocol's `initialize` handshake — which the 2026-07-28 spec **removed** (stateless core); it keeps working within the 12-month deprecation window, and the difference between the two generations is, in itself, a lesson, lists its tools, and calls them through the same path as local tools. Completeness exercise: the client handles the *happy path*; you add **graceful degradation** (a server that dies does not take down the session) and an **env filter** on the stdio subprocess — the state of the art's minimum defense.

Check your understanding

1. Why does MCP reduce integration cost from $M \times N$ to $M + N$, and what does that have to do with "open standard"? (Each harness and each tool writes once; the common protocol decouples the two sides.)
2. You are about to connect a third-party MCP server exposing a `search_tickets` tool. Give two reasons to distrust it and two concrete defenses. (Tool description = *tool poisoning*; return value = indirect injection; and *registration* itself is a vector — *line jumping*. Defenses: filtered env on the stdio subprocess; mediated credential — the server never sees the secret; injection eval; gate at connect; ch. 07 containment.)
3. A harness that is both MCP **client AND server** gains what that a client-only one does not — and which protocol primitives does that activate? (It becomes a piece consumable by other agents; it activates *sampling* and *elicitation*, the server asking the client.)

Appendix A — How each repository handles MCP

Per-harness evidence, with paths — online supplement, expanded each round.

opencode (round 1) — the most complete protocol implementation

`packages/opencode/src/mcp/` (~1,000 lines in `index.ts`, plus `catalog.ts`, `oauth-provider.ts`, `auth.ts`). Three transports — `StdioClientTransport`, `StreamableHTTPClientTransport` and `SSEClientTransport` with **automatic**

HTTP → SSE fallback. Full OAuth: authorization with a local callback server, PKCE, dedicated `opencode mcp auth` command. Covers the wide surface: `ToolListChanged` notifications, logging, roots, prompts, resources and resource templates. Server instructions go into the system prompt (`system.ts:mcp()`) — the server can teach the model how to use it (which is also the injection vector).

gemini-cli (round 1) — enterprise-grade OAuth

`packages/core/src/tools/mcp-client.ts` + `mcp-client-manager.ts`, the same three transports by config. The differential in `packages/core/src/mcp/`: beyond standard OAuth, **Google auth** providers and **service account impersonation** — MCP for corporate GCP. Tools become `DiscoveredMCPTool` with per-server namespacing; MCP prompts exposed; management via `/mcp` and `~/.gemini/settings.json`. Notable: the eval suite includes a **prompt injection via MCP** test (ch. 11) — the only one to treat an MCP server as a tested attack surface.

OpenHarness (round 1) — pragmatic client

`src/openharness/mcp/` (`McpClientManager`) on the `mcp>=1.0.0` SDK: **stdio** and **Streamable HTTP** transports (no SSE), with connection status, auto-reconnect and **graceful degradation** when a server dies (`call_tool/read_resource` do not take down the session). Resources exposed as its own tools (`list_mcp_resources`, `read_mcp_resource`); `mcp_auth` for authentication. Config via `oh mcp` and `--mcp-config`.

Goose (round 2) ★ MCP-native — the protocol as backbone

The extreme case: **every tool is MCP**. The `goose-mcp` built-ins (memory, computercontroller, tutorial...) are real `rmcp::ServerHandler` servers served **in-process over DuplexStream** (virtual stdio) and can run standalone (`goose mcp <server>`). Even developer/shell/edit are "platform extensions" speaking `McpClientTrait`. A single abstraction for the entire tool surface — the protocol is not an integration, it is the architecture. (Goose is also one of the founding projects of the Agentic AI Foundation.)

Codex CLI (round 2) — client and server, four transports

`rmcp-client/` + `mcp-server/` (Codex exposes itself as an MCP server). **Four transports** (stdio, streamable HTTP, in-process, process-executor); **full OAuth** with refresh transactions and store locking; **elicitation**; server prewarm/refresh; per-MCP-tool approval templates. Integrates MCP into containment (per-tool approval).

Hermes (round 2) — client and server, with sampling

Client with stdio/StreamableHTTP/SSE, OAuth, per-server timeouts, **sampling** (the server can request completions from the client) and per-server opt-in parallelism; `mcp_serve.py` exposes Hermes to other MCP hosts.

OpenClaw (round 2) — client and server, with env filtering

`openclaw mcp serve` exposes channel conversations via stdio to Codex/Claude Code. Client: `mcp.servers` registry with stdio/SSE/streamable-http, **OAuth PKCE in SQLite**, **mTLS**, tool filters, probe/doctor — and an **env security filter** on stdio (blocks `NODE_OPTIONS`, `LD_*`, `DYLD_*`). MCP Apps support with an isolated-origin sandbox.

OpenHands (round 2) — bidirectional, with secret redaction

Client (per-agent MCP config with secret redaction/restoration in GET/PUT round-trips) and **server** (the app-server is a FastMCP exposing PR tools — `create_pr/create_mr` — to the sandboxes, plus an **MCP proxy for Tavily** providing search without exposing the API key). Agent profiles reference server subsets.

IronClaw (round 2) — capability-mediated MCP

`ironclaw_mcp` adapts MCP tools into **capabilities without granting ambient authority**: FS, secrets and network remain mediated; **Streamable HTTP** (protocol 2025-06-18); **mediated credential injection** (the server never sees the secret); resources accounted for by the governor. Ch. 07's containment model applied to MCP.

ohmo (round 2) — inherited

Complete via `McpClientManager` (inherited from the base); server counts in state and a summary exposed to the gateway. Gap: no MCP config of its own (`~/ .ohmo/mcp.json` does not exist) and no per-channel/per-sender MCP isolation.

n8n (round 2) — bidirectional in the workflow engine

MCP Client Tool (SSE + Streamable HTTP, Bearer/OAuth2, tool filtering, per-execution session cache) and **MCP Server Trigger** (`McpTrigger` + `McpServer.ts`) — exposes the connected n8n tools as an MCP endpoint to external clients. Official SDK. The "inverted harness" also speaks MCP in both roles.

Aider (round 2) — absent by philosophy

MCP score **0**. The *context-first* school in its pure state: the 3s sit where the philosophy bets (context, edit formats, git, evals), and the MCP gap is a choice, not a lag.

Frameworks (frameworks round)

Agents SDK (OpenAI): support for MCP servers as a tool source; LangGraph/langchain: MCP adapters for tools; CrewAI: MCP integration via toolkit; software-agent-sdk: MCP-style annotations in the tool contract. MCP is an assumed integration point at the framework layer too — reinforcing the standardization thesis.

07 — Permissions & Sandboxing

🕒 state of the art 2026-07 · revised 2026-07-25 · 📖 ~9 min read

Learning objectives

By the end of this chapter, you should be able to:

1. **Distinguish** the two layers of defense — policy (what the agent may request) and containment (what the process can actually do);
2. **Design** permissions along two orthogonal dimensions (sandbox mode × approval policy);
3. **Apply** the "lethal trifecta" and the "rule of two" as checklists for toolset review and session architecture;
4. **Implement** a `PermissionPolicy` as pure domain (testable without an LLM (Large Language Model)) + non-disableable sensitive paths (step 6);
5. **Evaluate** a real harness for its *blast radius* — what leaks if the injection wins?

The problem

An agent with a shell is a user with a shell: it can delete files, exfiltrate credentials, make network calls. The control mechanisms answer two distinct threats: the **mistake** (the model does something destructive by accident) and the **attack** (prompt injection convinces the model to act against the user). It is the dimension of greatest divergence among the harnesses — a sign that the industry has not yet converged, though it is converging fast.

Two levels, frequently conflated: **permissions** (policy: approval, allowlists, modes) and **sandbox** (containment: OS-imposed limits, even when policy fails).

Scientific foundations

- **The threat, defined** — *Not what you've signed up for* (Greshake et al., [arXiv 2302.12173](#)): indirect injection — instructions planted in data the agent is going to read — is the vector that no traditional code vulnerability captures.
- **The map of defenses** — the layered attack-surface survey ([arXiv 2604.23338](#)) and the agentic security survey ([arXiv 2510.06445](#)) organize threats and defenses; the computer-using agents survey ([arXiv 2505.10924](#)) focuses on those who have a shell.

(Full bibliography: [livro/bibliografia.md](#).)

Industry sources

- **Making Claude Code more secure with sandboxing** (Anthropic): containment on OS primitives (bubblewrap/Seatbelt), workspace-only writes, **network denied by default** — and egress goes through a proxy that runs *outside* the sandbox and enforces a per-domain allowlist. The network boundary is a separate, privileged component, not a bypassable in-process check.
- **How we contain Claude across products** (Anthropic): three regimes (ephemeral gVisor, OS sandbox + approval, sealed VM with credentials outside the guest) and the central thesis — **hard, deterministic boundaries before probabilistic model defenses**. Honest detail: the egress proxy itself broke twice — treat your proxy as the most fragile component, not the most trusted.
- **Agent approvals & security** (OpenAI Codex): the matrix of **two orthogonal axes** — sandbox mode (`read-only/workspace-write/danger-full-access`) × approval policy (`untrusted/on-request/on-failure/never`), with `on-failure` firing the prompt only *after* the sandbox blocks. The most copyable design pattern on the market.
- **Agents Rule of Two** (Meta AI): an agent should not satisfy more than two of the three — processing untrusted input, accessing sensitive data, changing state/communicating externally — in the same session. A *session architecture* criterion, not a substitute for defense-in-depth.
- **The lethal trifecta** (Simon Willison): private data + untrusted content + external communication = exfiltration. Use it as a **toolset checklist**: which corner does each new tool close? The **counterpoint**: announced defenses fall when "the attacker moves last".

- **Attacks on OpenClaw** (The Hacker News): the real-world case — one-click RCE (CVE-2026-25253), plaintext credentials, injection planted in an email signature/calendar invite/issue. The vector was not the model, it was the **harness**: secrets in the same space as the tools + unlimited untrusted input.
- **See also**: the living collections [Awesome Harness Engineering — Permissions & Authorization](#) and [Awesome Harness Engineering — Security, Sandbox & Permissions](#) gather more consultable resources for this dimension (patterns, articles, and implementations), curated by problem.

The state of the art

1. Two orthogonal dimensions, not a slider

The old mental model ("YOLO ↔ ask about everything") is dead. The consensus is to separate **maximum physical capability** (sandbox) from **when to escalate to a human** (approval policy) — independently configurable. Codex is the canonical example (mode × policy, with on-failure). And there are two containment paradigms the benchmark separated:

- **OS containment** (the process *cannot*): Seatbelt + bubblewrap/seccomp + Landlock in Codex; 6 Seatbelt profiles + Docker in gemini-cli; fail-closed WASM + per-tenant Docker in IronClaw;
- **authority architecture** (the loop *cannot reach*): IronClaw makes the loop structurally incapable of acting without the kernel — type-unforgeable trust class, approvals as per-invocation leases, verified by dependency tests. No harness fully combines the two yet — it is the open frontier.

2. Policy without containment is a bet on the model's obedience

The cross-cutting lesson of the benchmark: harnesses with elegant policy but no OS sandbox (opencode, ohmo) are betting the model obeys. Three cheap, exportable defenses the state of the art has consolidated: **non-disableable sensitive paths** (OpenHarness's SENSITIVE_PATH_PATTERNS — denies `.ssh`, credentials, `.kube/config` before any user rule, explicitly against injection); **structural shell parsing** before judging (gemini-cli's policy engine understands redirections and wrappers; software-agent-sdk's `defense_in_depth` detects compositions like fetch-to-exec via AST); and **credentials outside the process** (injected at the egress edge, never in the tools' space — IronClaw, and the direct lesson of the OpenClaw case).

3. Prompt injection is treated as unsolvable — the effort migrated to blast radius

The 2026 consensus, from the model to the vendors: you cannot reliably "detect" injection. The work migrated to **designing sessions that never accumulate the trifecta** (rule of two as the criterion for when to break context), **isolating credentials** (keychain on the host, sealed VM) and **controlling egress** (per-domain allowlist via an external proxy). In the personal-agent category, the third-party vector earned its own defense: **deny-by-default contact pairing/allowlisting** (OpenClaw, ohmo) and a **non-main sandbox** for every session that is not the owner's. And the emerging honesty norm: publishing the gate's **false-negative rates** (Claude Code's auto mode is discussed with numbers in both directions) instead of asserting binary security.

EXECUTIVE SUMMARY

What is most modern: the two orthogonal dimensions; the two containment paradigms (OS × authority) and the finding that nobody has combined them; and the migration from "detect injection" to "reduce blast radius" (trifecta/rule-of-two as checklists, credentials outside the process, controlled egress). **What to steal**: non-disableable sensitive paths; structural shell parsing; on-failure (approve only after the block); contact pairing; publishing the gate's false-negative rate.

Hands-on — harness-zero, step 6

Step 6 introduces the `PermissionPolicy` as **pure domain**: a function (`ação, contexto`) → `allow | ask | deny` that knows nothing of LLMs or chat — testable in isolation (it is the "isolated domain" DDD names, and the test runs without a network). You implement: the three verdicts (allow/ask/deny), the non-disableable sensitive paths, and **inline approval in the chat** (the front end pauses

and asks — the visible manifestation of the policy). Completeness exercise: rule evaluation ships ready; you add minimal parsing of a shell command before judging it.

Check your understanding

1. A harness has only an approval policy, with no OS sandbox. Which class of attack can it not contain, and why? (Policy without containment; the model can be persuaded.)
2. You are about to add an email-sending tool to an agent that already reads GitHub issues and has access to the private repository. Apply the lethal trifecta. (It closes the third corner → exfiltration possible.)
3. Why can **on-failure** (approve only after the block) be better than **on-request** (approve before each action)? (Friction × coverage; the sandbox filters out what never needs a human.)

Appendix A — How each repository handles permissions and sandboxing

Per-harness evidence, with paths — online supplement, expanded each round.

gemini-cli (round 1) ★ policy engine + OS sandbox

`packages/core/src/policy/policy-engine.ts`: prioritized rules with **structural shell parsing** (`parseCommandDetails`, `stripShellWrapper`, redirection detection), rules in TOML; 4 ApprovalModes; 6 Seatbelt profiles (`sandbox-macos-*.sb`) + Docker/Podman with proxy; **trusted folders** gatekeeping hooks/agents.

OpenHarness (round 1) ★ sensitive paths

`permissions/checker.py`: path rules, denied commands, 3 modes; **hardcoded, non-disableable SENSITIVE_PATH_PATTERNS** (`.ssh`, `.aws/credentials`, `.gnupg`, `.kube/config`) against injection; sandbox via `sandbox-runtime/Docker` with a domain allowlist; `trust_env=False` in the web tools (anti-SSRF).

opencode (round 1) — policy without containment

`permission/`: rulesets with wildcards (`allow | ask | deny`, last-match-wins, default `ask`), approval via `Deferred` + event; **subagents derive restricted permissions; no OS sandbox in the core** (containers only in enterprise).

Codex CLI (round 2) ★ OS containment in 3 layers

`sandboxing/` + `linux-sandbox/` + `windows-sandbox-rs/`: Seatbelt via `sandbox-exec` (anti-tamper hardcoded path), embedded bubblewrap + **seccomp** + `NO_NEW_PRIVS`, legacy Landlock; `AskForApproval` incl. **Granular**; per-command **execpolicy in Starlark**; `assess_patch_safety`; network-proxy.

Goose (round 2)

`permission/`: `GooseMode` modes (Auto/Approve/Chat); **permission_judge uses an LLM** to classify read-only; per-signature `ToolPermissionStore` with expiration; light execution isolation (direct shell; external Docker).

OpenClaw (round 2) ★ third-party pairing

`src/pairing/` + `docs/security/THREAT-MODEL-ATLAS.md`: **DMs as untrusted input**, `dmPolicy`: "pairing" default (pairing code, SQLite allowlist); multi-backend sandbox (Docker `network:none/readOnlyRoot/capDrop:ALL`, SSH, OpenShell) with a **non-main** mode; `openclaw doctor/security audit`; caveat: `sandbox.mode` off by default in the main session.

Hermes (round 2)

`tools/approval.py` (detection + allowlist), per-thread callbacks; **six isolated terminal backends** (local, Docker, SSH, Singularity, Modal, Daytona); subagents with safe-by-default `_subagent_auto_deny`; anti-traversal `path_security.py`.

IronClaw (round 2) ★★ authority architecture

`crates/ironclaw_authorization + _approvals + _trust + _wasm + _process_sandbox + _secrets + _network + _safety`: exact-invocation authorization (fail-closed), approvals as **per-invocation leases with fingerprint, type-unforgeable trust class** (`#[serde(skip_deserializing)]`), WASM (fuel/memory/rate, egress denied), per-tenant Docker, zero-exposure secrets at the egress edge, anti-SSRF, bidirectional leak detector — the loop cannot reach the effects (verified by dependency tests).

ohmo (round 2.5) — the right half

`channels/impl/base.py`: **deny-by-default** allowlist + per-sender session isolation + blocking of remote admin commands + OpenHarness's sensitive paths. Gap: `permission_mode/sandbox_enabled` in `gateway.json` are **dead code** — no dial between deny-everything and `full_auto`.

software-agent-sdk (frameworks round)

`sdk/security/`: risk analysis (LLM analyzer + deterministic `defense_in_depth/` with an AST shell parser detecting **fetch-to-exec**) + confirmation policy (`AlwaysConfirm/ConfirmRisky` by threshold); the conversation **returns** in `WAITING_FOR_CONFIRMATION` (it does not block); secret masking.

n8n (round 2) — structural permission

Permission is **topological**: the author chooses which nodes sit on the `AiTool` port — allowlist by construction. Real HITL via `sendAndWait` (durable pause), forbidden in sub-agents; Guardrails node.

Frameworks (frameworks round) — left open

LangGraph and CrewAI have no native tool policy (you build it on `interrupt/HITL`); the Agents SDK (Software Development Kit) has guardrails at three levels (agent/run/tool) as a primitive, but containment is left to the adopter.

08 — Memory & State

🕒 state of the art 2026-07 · revised 2026-07-26 · 📖 ~12 min read

Learning objectives

By the end of this chapter, you should be able to:

1. **Distinguish** the problem's three layers — session state, long-term memory and workspace state — and the requirement specific to each;
2. **Explain** why memory **is not** RAG (Retrieval-Augmented Generation) (memory = retrieval + write path + state management) and why versionable markdown beat vector databases in the code domain;
3. **Derive** a recall policy from the recency $\times$ importance $\times$ relevance formula, and a forgetting policy from usage;
4. **Evaluate** the impact of **reversibility** (workspace checkpointing) on the permissions risk calculus;
5. **Implement** harness-zero's session persistence (SQLite adapter + `/resume`) in step 4.

The problem

The model forgets everything between calls; the harness remembers for it. "Memory and state" covers three layers with different requirements:

1. **Session state** — the conversation itself: messages, tool calls, metadata. It must survive restarts and allow resuming (`resume`), branching and reverting.
2. **Long-term memory** — facts that cross sessions: user preferences, project decisions, learnings. It must be **selectable** (not everything enters every context) and **updatable** (facts change).
3. **Workspace state** — what the agent *did* to the files. It must be **reversible**: undoing an agent's changes is as important as making them.

The thesis unifying the three: the context window is volatile, expensive memory; everything that needs to last lives **outside** it, and the harness decides what to bring back and when.

Scientific foundations

Agent memory has a mature literature — and it provides the exact vocabulary for what the harnesses do in practice.

- **The window as RAM** — [MemGPT: LLMs as Operating Systems, arXiv 2310.08560](#) treats the context as scarce main memory, backed by two external levels (*recall* of recent history and searchable *archival*), with the **agent** paging data via tool calls ("context page faults"). Decision: what to evict and what to fetch is decided by the agent, not by a fixed RAG pipeline.
- **The canonical taxonomy** — [CoALA, arXiv 2309.02427](#) separates **episodic** memory (past experience), **semantic** memory (knowledge of the world/user) and **procedural** memory (skills/code), plus working memory. Decision: at write time, decide *which* kind of memory that fact is — each kind is retrieved differently. The [memory mechanisms survey, arXiv 2404.13501](#) (later ACM TOIS) organizes the subsystem by *sources* · *forms* · *operations* (writing, management/consolidation, reading) — budget effort per operation, not just for the search index.
- **The recall formula** — [Generative Agents, arXiv 2304.03442](#) (UIST '23) stores observations in a dated *memory stream* and retrieves by a composite score of **recency** $\times$ **importance** $\times$ **relevance** (exponential recency decay, importance scored by an LLM (Large Language Model), relevance by embedding). It is the concrete formula a harness should implement to rank what re-enters the context — and it introduces **consolidation by reflection** (synthesizing high-level reflections from clusters of observations).
- **Controlled forgetting** — [MemoryBank, arXiv 2305.10250](#) (AAAI '24) decays/reinforces each memory's strength via an Ebbinghaus curve (elapsed time $\times$ access frequency), keeping the store bounded. Decision: unused memory is a pruning candidate — *usage tracking* is what closes the loop.
- **Memory as learning** — [Reflexion, arXiv 2303.11366](#) (NeurIPS '23) converts outcome feedback into verbal self-reflection, persisted in an episodic buffer and reinjected on the next attempt — improving without updating weights. And recent architectures

(A-MEM, arXiv 2502.12110; Mem0, arXiv 2504.19413) treat writing as an *extract* → *consolidate* → *link* pipeline, with the memory network self-organizing (Zettelkasten-style). Bridge to ch. 16 (self-improvement).

(Full bibliography and pointers: [livro/bibliografia.md](#).)

Industry sources

- **Session as a durable event log** — [Manage sessions \(Claude Code\)](#): each session is continuously written to disk as **JSONL** per project (one line per message/tool-use/metadata); `--continue` resumes the most recent one in the directory, `--resume` opens a picker. Decision: "resuming" is **restoring complete state** (tool calls, results, permission mode, active goal), not text replay — the harness owns a private durable log, not a stable public schema.
- **Workspace reversal as a separate track** — [Checkpointing \(Claude Code\)](#) captures the code state before each prompt; `/rewind` restores code, conversation **or** both (100 recent checkpoints, cleaned up with the session). The [Agent SDK's file-checkpointing](#) exposes this as a reusable primitive. Decision: undoing the *code* is a separate store from undoing the *conversation*, linked by the prompt index.
- **Durable memory as files with precedence** — [How Claude remembers your project](#): the CLAUDE.md hierarchy (managed policy → user → project → local), the `#` shortcut to append a memory line, `/memory` to edit. Decision: cross-session memory is **markdown in precedence tiers** (the most specific wins) — versionable, auditable, scoped; reread at launch as always-on context.
- **The memory tool (beta) and "assume interruption"** — [Memory tool](#): the model requests operations (`view/create/str_replace/...`) on a `/memories` directory that persists across conversations, but execution is **client-side** — your app implements the storage (and the protection against path traversal, size limits, expiration). The system injects "ASSUME INTERRUPTION: your window may be reset at any moment". Paired with [context management](#) (context editing evicts stale pairs from the window; the memory tool persists outside it) — two levels: short-term hygiene + external long-term store. Decision: for long-running agents you need both; the window is ephemeral, `/memories` is the source of truth (the pattern from the essay [harnesses for long-running agents](#): a structured progress log, read at the start and updated at the end of each session).
- **Memory ≠ RAG** — the distinction became an industry thesis: Letta ("RAG is not agent memory") and [AWS Bedrock AgentCore](#) argue that RAG is *stateless* reading; memory is reading + **write path** + **state management** (admission, resolution of conflicting facts, invalidation). Letta exposes self-editing *memory blocks* and **core/recall/archival** tiers (the MemGPT hierarchy as a product); [mem0](#) routes each fact through a layer with its own lifetime; Zep/Graphiti models memory as a **bi-temporal knowledge graph** (outdated facts are *invalidated*, not deleted); LangMem/LangGraph separates **short-term (thread)** from **long-term (per-namespace store)**. Decision: you cannot "buy" memory by bolting on a vector store — you need a write/update/invalidation pipeline.
- **See also**: the living collection [Awesome Harness Engineering — Memory & State](#) gathers more consultable resources for this dimension (patterns, articles, and implementations), curated by problem.

The state of the art

1. Three layers, three champions — and no vector database

The problem's three layers got different champions in the cohort: **session state** (database durability — opencode with SQLite + replayable events; Codex with per-turn rollout jsonl; OpenHands with event-stream), **long-term memory** (relevance + format rigor — OpenHarness with versioned memdir, `relevance.py` and `usage.py`; Hermes with `MEMORY.md/USER.md` + `session_search`), **workspace state** (git-based reversal — Aider and gemini-cli). And the finding that persists: **none of the code harnesses uses a vector database** for memory. In the code domain, versionable markdown beat embeddings — because code memory needs a *write path* (the agent edits the file) and auditability, exactly what the "memory ≠ RAG" thesis predicts.

2. The recall formula and forgetting moved from paper to code

OpenHarness's `relevance.py` + `usage.py` is the practical instance of the Generative Agents stream: it selects by relevance what enters the context and marks usage — unused memory becomes a pruning candidate (MemoryBank's forgetting curve, in practice). Hermes formalizes **active maintenance**: a single tool edits `MEMORY.md/USER.md` with **periodic nudges** (every 10 turns), and a `session_search` (FTS5/BM25 index over the session SQLite, with discovery/recall/summarization modes) provides **cross-session recall** — MemGPT's archival layer built on textual search, not vectors.

3. Reversibility became a primitive — and it changes the risk calculus

Workspace checkpointing stopped being a feature and became a primitive: **Aider** pioneered it years ago (git-native state: atomic auto-commit per round, `aider_commit_hashes`, `/undo`, `.aider.chat.history.md`), **gemini-cli** consecrated it (`/restore`, `/rewind` of the disk via git snapshots), and Claude Code exposes it as checkpointing with separate tracks for code and conversation. The design consequence is the most interesting part: **an agent whose actions are reversible changes the risk calculus of everything else** — permissions can be looser when undoing is cheap (ties into ch. 07).

4. Pluggable providers and the harness as a memory server

The emerging frontier: memory as a pluggable service. Hermes already accepts external providers (Honcho, mem0, supermemory) behind its layer; products like Letta/mem0/Zep position themselves as the "universal memory layer" consumable by any harness. The design tension for the next rounds: keep memory as a **local versionable file** (auditable, portable, dependency-free) or outsource it to a managed store (bi-temporal graph, scale). In code, the file still wins; outside it, the pendulum is less clear.

EXECUTIVE SUMMARY

What is most modern: the OS-tiers frame (RAM ↔ recall ↔ archival) with the agent paging; recall by recency×importance×relevance with usage-based forgetting; workspace reversal as a primitive that loosens permissions; and the hard memory × RAG distinction (write path + invalidation). **What to steal:** persist the session as a durable event log (resuming = restoring state, not replay); treat memory as versionable markdown with usage tracking; separate the code-reversal track from the conversation; and, for long-running agents, write a durable progress log assuming the window can vanish at any moment.

Hands-on — harness-zero, step 4

Step 4 (`harness-zero/etapas/04-sessoes/`) gives harness-zero persistence: an **SQLite adapter** behind a **StorePort** stores messages and tool calls as typed rows, and `/resume` restores the complete state of a previous session (not just the text). Faithful to hexagonal *by refactoring*: the pain that gives birth to the port is reopening the process and losing the conversation. Completeness exercise: persistence covers the *happy path*; you add a minimal `USER.md/MEMORY.md` read at the start and a progress log updated at the end — the "assume interruption" pattern in its simplest form.

Check your understanding

1. Why is agent memory not the same thing as RAG, and what does that explain about choosing versionable markdown over a vector database in the code harnesses? (Memory = retrieval + write path + state management/invalidation; code needs an auditable write path.)
2. You have 10,000 memories and room for 20 in the context. Which score do you use to choose, and how do you decide what to prune over time? (Recency × importance × relevance; pruning by lack of use — the forgetting curve.)
3. Your agent gained workspace checkpointing with `/rewind`. Which decision *from another dimension* does that let you loosen, and why? (Permissions — the risk calculus drops when undoing is cheap; ch. 07.)

Appendix A — How each repository handles memory and state

Per-harness evidence, with paths — online supplement, expanded each round.

opencode (round 1) — state as a database

Persistence in **SQLite via Drizzle** (`packages/core/database`, `core/session/sql.ts`): sessions, messages and parts are typed rows. Sessions have a `parentID` (hierarchy for subagents), support revert (`session/revert.ts`) and **sharing** (`share/`, `sync/`). V2 (`CONTEXT.md`) takes the design to "data infrastructure": a durable prompt inbox, replayable events with cursors

(`sessions.events({sessionID, after})`), context snapshots persisted across restarts. Round 1's most robust state model — the harness as a distributed system with durable state.

gemini-cli (round 1) — the reversible workspace

Long-term memory in the `GEMINI.md` files themselves (`save_memory` tool, global in `~/.gemini` + project index, with auto-memory tested in evals). The distinctive feature is **git-based checkpointing** (`services/gitService.ts` + `chatRecordingService.ts`): workspace snapshots before edits, enabling `/restore` and `/rewind` — undoing the agent's changes on disk, not just in the conversation — plus `/resume`.

OpenHarness (round 1) — memory as files, with discipline

`src/openharness/memory/` (13 modules): persistent memory in markdown (`MEMORY.md`/per-project memdir) with **versioned schema, atomic file-locked writes and signatures**. `relevance.py` selects what enters the context; `usage.py` marks usage (unused memory is a pruning candidate). Sessions persisted with rich metadata (`services/session_storage.py`): permission mode, read-file state, invoked skills, compaction checkpoints. Resume via `-c/--continue`, `-r/--resume`, `/resume`.

Aider (round 2) ★ git-native state — the reversal pioneer

`aider/repo.py`: **atomic auto-commit per round** with an LLM-generated message, configurable authorship attribution, `aider_commit_hashes` tracking what the AI did, `dirty_commit` isolating pending changes. `/undo`, `diff` and `blame` become the memory interface; complemented by `.aider.chat.history.md` and `--restore-chat-history`. **Anticipated by years** the "git checkpoint" that gemini-cli and Claude Code consecrated.

Hermes (round 2) ★ multi-layer memory with cross-session recall

`MEMORY.md` (agent notes) + `USER.md` (user profile) edited by a single tool with **periodic nudges** (every 10 turns); pluggable external providers (**Honcho**, **mem0**, **supermemory**); and **session_search** — an FTS5 index over the session SQLite with three modes (discovery/BM25, windowed recall, LLM summarization) for cross-session recall. MemGPT's archival layer on textual search.

Codex CLI (round 2) — per-turn rollout jsonl

Each turn is persisted as **rollout jsonl** (recoverable); `SessionTask` (Regular/Review/Compact/UserShell) organizes the task machine. Durable, resumable session state integrated into the loop (`core/src/session/`).

OpenHands (round 2) — persisted event-stream

`openhands/app_server/event/` persists each `Event` as JSON per conversation, with pagination, filters and trajectory export. The control plane consumes/persists events; the action-observation loop runs in the SDK. Event sourcing as the state's backbone.

OpenClaw (round 2) — session lanes and workspace files

Runs serialized per *session lane* with a file-based write-lock between processes; workspace files (`MEMORY.md`, `USER.md`, `IDENTITY.md`...) injected with budgets (20k chars/file, 60k total) and marked truncation. Per-channel conversation persistence.

ohmo (round 2) — session/memory backends as plugins

Implements OpenHarness's `SessionBackend` and `MemoryCommandBackend` as first-class plugins (without touching the core), plus a **multi-session pool** (`RuntimeBundle` per `session_key`, recreated when the cwd changes). Proof that the app/engine boundary was designed.

IronClaw (round 2) — resumable state via checkpoints

Resumable state via **checkpoints**; a Queued → Running → Blocked → Completed state machine with **leases/heartbeats** and "one active run per canonical thread". The `LoopExit` carries only durable references — the loop never mutates state; the `LoopExitApplier` validates host-owned evidence before applying.

n8n (round 2) — the workflow engine's memory

Memory via *memory sub-nodes* (`contextWindowLength` window, `maxTokensFromMemory` cutoff); workflow state persisted by the engine across executions. Short by nature — event-triggered executions do not accumulate long context (compaction score 1, by design).

Frameworks (frameworks round)

LangGraph: **checkpoint** (short-term, thread-scoped) + per-namespace **store** (long-term cross-thread); LangMem: semantic/episodic/procedural memories as tools; Agents SDK and CrewAI: session/short-term state with persistence hooks. The short × long term distinction is a framework primitive — what the code harnesses implement by hand, the frameworks expose as an API.

09 — Planning

🕒 state of the art 2026-07 · revised 2026-07-26 · 📖 ~12 min read

Learning objectives

By the end of this chapter, you should be able to:

1. **Distinguish** the three planning instruments — plan mode, todo list, and decomposition — and what each one requires;
2. **Explain** why plan mode is implemented as a case of the permission system (enforced, not requested);
3. **Compare** ReAct (interleaving reasoning and action) with plan-then-execute and decide when each one fits;
4. **Evaluate** the tactical × durable stratification (task plan × session goal) and decomposition with dependencies;
5. **Implement** permission-enforced plan mode in harness-zero (step 8).

The problem

Models tend to act rashly: they edit before understanding, they "solve" before mapping the problem. Planning artifacts force a reading-and-design phase before the writing phase — and give the human a cheap approval point (reviewing a plan costs less than reviewing a diff).

Three distinct instruments, frequently confused:

1. **Plan mode** — a *state* of the harness in which writing is forbidden; the agent only researches and proposes.
2. **Todo list** — working memory for the task in progress: what remains, what is done.
3. **Decomposition** — breaking large work into trackable subtasks, possibly with dependencies.

Scientific foundations

The planning literature explains *why* these instruments exist — and warns against trusting the model's plan.

- **Interleaving beats plan-everything-first (when the environment is unpredictable)** — [ReAct](#), [arXiv 2210.03629](#) (ICLR '23) interleaves reasoning traces and tool actions in the same loop: each observation revises the next thought, so the agent recovers from surprises instead of executing a stale plan. Decision: carry reasoning and observations in a single alternating transcript.
- **Planning ahead helps (when the scope is known)** — [Plan-and-Solve](#), [arXiv 2305.04091](#) has the model emit an explicit plan before solving, suppressing missing steps. The two do not contradict each other: they are distinct regimes — the explicit plan for tasks of known scope, interleaving for uncertain environments.
- **Decompose only when needed** — [ADaPT](#), [arXiv 2311.05772](#) decomposes **recursively and only when the executor fails** a subtask, adapting depth to difficulty and to model capability. Decision: try to execute first, decompose on failure — this avoids the over-planning that most harnesses (wisely) do not impose.
- **Isolate context per subtask** — [Beyond Entangled Planning](#), [arXiv 2601.07577](#) (2026) decomposes into a **DAG of sub-goals** and gives each one *scoped* context, so local errors and replanning do not pollute a monolithic history — reporting up to -82% tokens. A direct bridge to subagents (ch. 10).
- **Do not trust the model's plan** — **externalize** — [PlanBench](#), [arXiv 2206.10498](#) and [TravelPlanner](#), [arXiv 2402.01622](#) show that raw models fail at plan generation and lose track of multiple constraints (GPT-4 ~0.6% on TravelPlanner). Decision: externalize constraint tracking into an artifact (plan/todo) instead of trusting the model to hold everything in context. This is *the* justification for the todo list.
- **The taxonomy as a checklist** — the [planning survey](#), [arXiv 2402.02716](#) organizes the components into five avenues (task decomposition · plan selection · external module · reflection · memory); [PlanGenLLMs](#), [arXiv 2502.11221](#) gives six criteria (completeness, executability, optimality, representation, generalization, efficiency) and [PLANET](#), [arXiv 2504.14773](#) organizes benchmarks by category.

(Full bibliography and pointers: [livro/bibliografia.md](#).)

Industry sources

- **Plan mode is a permission layer** — [Choose a permission mode \(Claude Code\)](#): plan mode removes write/execute for the *entire session*; the agent reads and explores, but every mutation is held until you exit (Shift+Tab cycles Normal → Plan → Auto-accept; `/plan; --permission-mode plan` for CI). Decision: planning is guaranteed by **revoking the mutation tools**, not by asking the model to "plan first". This is the official confirmation of round 1's discovery.
- **Explore → Plan → Code → Commit** — [Best practices \(Claude Code\)](#): the exploration and planning phases are "the cheapest in tokens and the most valuable in outcome". Decision: separating exploration from execution structurally prevents solving the wrong problem before understanding the code.
- **Todo as a machine-tracked artifact** — [Todo tracking \(Agent SDK\)](#): `TodoWrite` creates checklists with three states (pending/in_progress/completed) updated in real time. Decision: externalizing the plan into a structured artifact gives the agent a working-memory anchor and the user progress visibility — and the evolution into a *tasks* system with dependencies and persistence turns the plan into durable infrastructure, not scrollback.
- **Thinking between actions** — [The "think" tool](#) adds a reasoning step *in the middle* of tool use (after the result arrives); [extended thinking](#) exposes reasoning blocks with `budget_tokens` and, in the 4-series models, **interleaved thinking** (think → call tool → think about the result → call again). Decision: allocate explicit budget to planning steps and let reasoning interleave with tools — planning is not a one-shot prefix, it is continuous. (*anthropic.com returns 403 through the proxy; confirmed via independent mirrors.*)
- **Spec-driven: the spec is the durable plan** — [GitHub Spec Kit](#) formalizes `specify` (what/why) → `plan` (architecture) → `tasks` (actionable list) → `implement`, with approval gates between stages; [Kiro](#) generates `requirements.md` (EARS WHEN... THE SYSTEM SHALL...), `design.md`, and `tasks.md`, and **derives a dependency graph** that runs independent tasks in concurrent waves. Decision: the plan becomes a persisted source of truth re-consumed at every phase — it is exactly the method by which **this book is written** (see the project constitution).
- **Planning is an orchestration function — and the tension over parallelizing** — [Anthropic's multi-agent system](#) has the *lead* analyze the query, **write the plan to memory**, and only then spawn workers with isolated specs (planning as a dedicated role). [Cognition \("Don't Build Multi-Agents"\)](#) counters: Devin centralizes planning in one continuous context, because planning is context management — parallelizing workers becomes a game of "telephone" between conflicting implicit decisions. Decision: decompose-and-parallelize is a cost/benefit gate, not a default (connects to ch. 10). (*cognition.com 403 through the proxy; confirmed via mirrors.*)
- **See also**: the living collection [Awesome Harness Engineering — Planning & Task Decomposition](#) gathers more consultable resources for this dimension (patterns, articles, and implementations), curated by problem.

The state of the art

1. Plan mode = permission mode (now the official pattern)

Round 1's discovery — harnesses implement plan mode **as a case of the permission system** (ch. 07), not as its own subsystem — has ceased to be an observation and become a documented standard: the official Claude Code docs describe plan mode exactly that way (removes mutation for the session). Entering plan mode = switching to a ruleset that denies writes; exiting = restoring, with explicit approval. The mature pattern combines three guarantees: read-only that is **enforced** (not requested), the plan as a **persisted artifact** (not just text in the conversation), and **explicit approval** before executing.

2. ReAct became the default; the explicit plan retreated to long-horizon work

The clearest signal in the living book came from n8n: its **Plan-and-Execute Agent was deprecated** (it only exists in legacy V1, alongside ReAct), and V2/V3 converged on the pure Tools Agent — planning implicit in the model. This instantiates the scientific thesis: as models plan better inline, interleaving (ReAct) beats plan-then-execute as the default, and the **explicit plan concentrates where it still pays**: long-horizon work, human in the loop, and decomposition of large tasks. It's not that planning died — it's that cheap planning migrated inside the loop.

3. The todo list is externalized constraint tracking

What PlanBench and TravelPlanner prove (models lose track of multiple constraints) is what the todo list solves: a checklist with states (Codex `update_plan`, `TodoWrite`, the Hermes/Goose `todo`, OpenHarness's `TODO.md`) takes the constraints out of the model's head and puts them in an artifact. The modern evolution is giving that artifact **dependencies and persistence** — gemini-cli's graph tracker, Kiro's dependency graph, the DAG from "Beyond Entangled Planning".

4. Tactical × durable — the personal agents' contribution

Coding harnesses have a *task* plan; what they lack is the *durable* one. **OpenClaw** fills that with four layers: `update_plan` (tactical, one `in_progress` step at a time), **Goals** (one durable goal per session, with token budget and states, injected per turn and visible in the UI), **Task Flow** (durable orchestration with steps and JSON state), and standing orders (persistent policies). This tactical × durable stratification is the frontier the personal-agent category brought to the discipline.

5. Planning is the weakest dimension — and that is a data point, not an accident

Across all rounds, planning was the industry's lowest score (Codex 2, Goose 2, Aider 2, Hermes 2, OpenHands 1, n8n 1, IronClaw 2; only gemini-cli and OpenClaw reach 3). The living book's reading (expiration log, "enforced plan mode", ● open): the prosthesis exists because models act rashly, and it expires when models plan under risk spontaneously — which has not yet happened. The persistent weakness of the dimension is the evidence that the prosthesis is still needed.

EXECUTIVE SUMMARY

What's most modern: plan mode as a permission layer (the official pattern); ReAct/interleaved thinking as the default, with the explicit plan reserved for long-horizon work; the todo/checklist as externalized constraint tracking, evolving into dependency graphs; and the tactical × durable stratification. **What to steal:** enforce read-only through permissions, not through the prompt; externalize the plan into a persisted artifact with states; give thinking budget to planning steps; and decompose-and-parallelize only when the task's breadth pays for the cost.

Hands-on — harness-zero, step 8

Step 8 (harness-zero/etapas/08-plan/) adds plan mode to harness-zero by **reusing** the `PermissionPolicy` from step 6: entering plan mode sets a mode the policy translates into "every write tool is denied"; the agent only reads and proposes; exiting asks for approval and restores the mode. It is the concrete demonstration of the chapter's thesis — plan mode is not a subsystem, it is a configuration of the permission domain that already exists. Completeness exercise: `propor_plano` already persists the artifact (`PLAN.md`); you add the requirement that exiting plan mode only happens with an approved `PLAN.md` — the gate between planning and executing.

Check your understanding

1. Why does it make sense to implement plan mode as a mode of the permission system, instead of a dedicated subsystem? (It reuses an existing mechanism and gets for free the guarantee that read-only is *enforced*, not suggested to the model.)
2. Your agent operates in an unpredictable environment (API responses change the next step). Do you plan everything up front or interleave reasoning and action? Why? (Interleave — ReAct: each observation revises the next thought; a fixed plan goes stale.)
3. A benchmark shows your agent losing track of 8 constraints in a task. Which planning instrument attacks that, and why? (Todo list / checklist — it externalizes constraint tracking out of the model's context.)

Appendix A — How each repository handles planning

Per-harness evidence, with paths — supplemented online, expanded each round.

opencode (round 1) — plan as an agent

Plan mode is a **built-in plan agent** with a read-only ruleset (denies edits, asks confirmation for bash) — planning is switching agents, not just modes. The `plan_exit` tool (`tool/plan.ts`) closes the cycle: it asks for approval, **writes the plan to a file**, and transitions to the build agent. Dedicated prompts (`prompt/plan-mode.txt`, `plan-reminder-anthropic.txt` — reminders per model family). Per-session todos via `todowrite` (`session/todo.ts`).

gemini-cli (round 1) — plan with gatekeeping and decomposition

`ApprovalMode.PLAN` (`policy/types.ts`) with `enter-plan-mode/exit-plan-mode`: a read-only state whose prompt lists the available tools, and `getApprovedPlanPath()` **gatekeeps execution**. Todos via `WriteTodosTool`. The instrument the others don't have: the optional **tracker** (`trackerTools.ts`) — tasks with dependencies (`tracker_add_dependency`) and a graph (`tracker_visualize`). Plan mode has its own behavioral eval (`evals/plan_mode.eval.ts`).

OpenHarness (round 1) — the minimal, correct version

`EnterPlanModeTool` sets `settings.permission.mode = PLAN` (blocks all writes); `ExitPlanModeTool` restores — the most direct implementation of the plan-mode-is-permissions equivalence. Todos in `TODO.md` via `TodoWriteTool` (persistent, readable). Bundled `plan` skill; heavyweight decomposition in the autopilot subsystem (a queue of `RepoTaskCard`).

OpenClaw (round 2) ★ — tactical × durable in four layers

`update_plan` (multi-step plan, one `in_progress` at a time), **Goals** (a durable per-session goal with token budget and states, injected per turn and visible in the UI), **Task Flow** (durable orchestration with steps and JSON state), and **standing orders** (persistent policies). The tactical × durable stratification that coding harnesses lack.

Codex CLI (round 2) — structured checklist

`update_plan` tool (checklist visible in the TUI (Terminal User Interface)) + `ReviewTask`. No two-phase plan mode with plan approval before execution — the economy of "planning first" lives in the checklist, not in a permission gate.

Aider (round 2) — plan-then-edit via coder modes

`/ask` (discusses without editing), `/architect` (reasons about the "how" before delegating), and `/context` (uses the repo-map to converge on the files). Lightweight plan-then-edit, with no persisted plan artifact and no todo list; the `architect-editor` split executes the plan with a second model.

Goose (round 2) — declarative recipes

Recipes (YAML/JSON with instructions, typed parameters, `response.json_schema`, `retry`) + the `todo` extension + `final_output_tool`. Declarative/reusable planning, without a two-phase plan mode.

Hermes (round 2) — todo + Kanban

`todo` tool + iteration budget + a **Kanban system** for multi-agent coordination with specs. Planning coupled to the loop, without a separate formal planner.

n8n (round 2) — the planning that retreated

The **Plan-and-Execute Agent** exists but is **legacy** (V1 only, alongside ReAct/Conversational); V2/V3 converged on the pure Tools Agent. Planning became implicit in the model (+ optional `ToolThink`). The clearest case of the explicit plan losing to interleaving.

IronClaw (round 2) — temporal planning, not decomposition

No first-class task decomposition; the loop's "planner" is strategy composition. Its strength is *temporal* planning (scheduling, leases/heartbeats — supplementary dim. 14).

OpenHands / ohmo (round 2)

OpenHands: a planner tab in the UI and hooks, without a first-class decomposition subsystem in this repo (score 1; the core migrated to the SDK). ohmo: inherited plan mode/todos that assume a TUI — no plan-approval surface in a chat channel.

Frameworks (frameworks round)

LangGraph: planning as an explicit graph of nodes (the plan *is* the topology); Agents SDK and CrewAI: planner/executor roles and sequential/hierarchical processes; the spec-driven camp (Spec Kit/Kiro) treats the plan as a versioned artifact with gates. Where coding harnesses improvise the plan inside the loop, frameworks materialize it as first-class structure.

10 — Subagents & Orchestration

🕒 state of the art 2026-07 · revised 2026-07-26 · 📖 ~12 min read

Learning objectives

By the end of this chapter, you should be able to:

1. **Explain** why a subagent's primary gain is context isolation (reads a lot, returns a little), not parallelism;
2. **Compare** the three philosophies — subagent-as-tool, as-service, and as-teammate;
3. **Evaluate** the cost/benefit gate of decompose-and-parallelize (the Anthropic × Cognition tension) and the failure modes that justify guardrails;
4. **Distinguish** local delegation from cross-system delegation (A2A (Agent-to-Agent)/ACP (Agent Client Protocol)) and when each applies;
5. **Implement** the `task` tool with a child session and derived permissions in harness-zero (step 9).

The problem

A single context cannot hold large tasks: codebase exploration pollutes the window with file dumps; parallelizable work runs serially; and a generalist agent does everything mediocrely. Subagents solve this through **context division** (the subagent reads 50 files and returns only the conclusion), **specialization** (per-role prompts and permissions), and **parallelism**.

The design decisions:

- **Isolation**: a child session? A separate process? Its own git worktree (for parallel edits without conflict)?
- **Permissions**: inherit the parent's? Derived and restricted? Degraded by depth?
- **Communication**: fire-and-forget (returns one result) or a continuous channel (mailbox, messages)?
- **Reach**: local only, or delegation to remote agents from other vendors?

Scientific foundations

The multi-agent systems (MAS) literature has two messages for harness builders: the patterns that work, and the warning that most failures are design failures.

- **The failure is in the design, not the model** — [MAST, "Why Do Multi-Agent LLM \(Large Language Model\) Systems Fail?", arXiv 2503.13657](#) empirically derives 14 failure modes in three categories (specification/roles · inter-agent misalignment · task verification), and concludes that most come from the *system*, not the weights. Decision: invest in explicit role specs, alignment checks, and a dedicated verification stage — not in a bigger model.
- **Roles and SOPs against cascading hallucination** — [MetaGPT, arXiv 2308.00352](#) codifies *Standardized Operating Procedures* and assembly-line roles (PM, architect, engineer, QA) with structured intermediate artifacts, because naively chaining LLMs propagates hallucination; role-scoped outputs let the next agent verify the previous one. And [CAMEL, arXiv 2303.17760](#) shows that role-play **drifts** (role swapping, repetition, early termination) — role stability must be *enforced*, not assumed.
- **Programmable topology and dynamic recruitment** — [AutoGen, arXiv 2308.08155](#) separates agents from conversation topology (swap the orchestration pattern without rewriting agents); [AgentVerse, arXiv 2308.10848](#) assembles the group per task and monitors negative emergent behavior. [ChatDev, arXiv 2307.07924](#) decomposes the pipeline into two-party dialogues per phase. Taxonomy pointer: the [MAS survey, arXiv 2402.01680](#).
- **The healthy skepticism** — multi-agent debate is a verification primitive ([Du et al., arXiv 2305.14325](#)), but [Should We Be Going MAD?, arXiv 2311.17371](#) and [Stop Overvaluing Multi-Agent Debate, arXiv 2502.08788](#) show it does not always beat self-consistency/CoT at equal compute. Decision: **always compare the multi-agent harness against a compute-matched single-agent baseline** before accepting the complexity.

(Full bibliography and pointers: `livro/bibliografia.md`.)

Industry sources

- **Subagent = isolated instance with a restricted toolset** — [Create custom subagents \(Claude Code\)](#): each subagent is a *fresh*, *isolated* instance launched by the `Task` tool, with its own context window and a per-agent-type toolset. The [Agent SDK subagents](#) are declared as config (name, tools, model, prompt) — you can pin cheap models (Haiku for read-only Explore) per role and enforce least-privilege per type. Decision: a search subagent burns tokens exploring without polluting the orchestrator's context, returning only a compact summary.
- **Orchestrator-worker — and the price** — [Anthropic's multi-agent research system](#): a *lead* plans, writes the plan to memory, and spawns parallel subagents, each with isolated context and an **explicit contract** (objective, output format, tools, boundaries). The breadth gain comes at **~15× the tokens** of a single chat (and, per the post, tokens explain ~80% of performance variance) — it only pays on high-value, high-breadth tasks. The [guide to when to use multi-agent](#) gives the three cases: context pollution, genuinely parallel subtasks, and specialization that sharpens tool selection. (*anthropic.com 403 through the proxy; numbers via independent mirrors.*)
- **The counter-argument** — [Don't Build Multi-Agents \(Cognition\)](#): prefer a **single-threaded agent with context compression**. When the work fans out in parallel, each subagent acts on a partial view and makes conflicting implicit decisions (the Flappy Bird example: one builds a Mario-style background, another an incompatible bird) — a game of "telephone" that creates the reconciliation step the architecture itself produced. Two principles: *share the full trace with every agent* and *actions carry implicit decisions; avoid conflicting ones*. For long tasks, add a compression model instead of splitting the thread. (*cognition.com 403; confirmed via HN/GitHub.*)
- **Frameworks materialize the patterns** — [Agents SDK \(OpenAI\)](#) distinguishes **handoffs** (transfers control to a specialist) from **agents-as-tools** (a manager calls sub-agents as functions, keeping the thread); [Swarm](#) was the educational origin of the handoff. [CrewAI](#) chooses between **sequential** and **hierarchical** (`manager_llm` delegates and validates); [LangGraph](#) models a **supervisor** routing among workers with persistent state; [Magentic-One \(AutoGen\)](#) keeps a **progress ledger** and replans on failure; [Google's ADK](#) mixes coordinator/dispatcher with `Sequential/Parallel/Loop` primitives. Decision: choose the coordination form (handoff × tool × supervisor × ledger) by what you need to retain — thread, control, or recovery.
- **Cross-system delegation: A2A (and ACP converging into it)** — when subagents live in different vendors, delegation becomes protocol: [A2A](#) uses **Agent Cards** (JSON announcing identity, skills, endpoint, auth) for discovery and **Tasks** with a lifecycle as the unit of delegated work, over HTTP+JSON-RPC (Remote Procedure Call)+SSE (Server-Sent Events); it is the cross-org generalization of the `Task` tool's handoff. [ACP \(IBM/BeeAI\)](#) was the REST-native alternative — but [merged into A2A under the Linux Foundation in Aug 2025](#). Decision: for new work, standardize on A2A (connects to ch. 17).
- **See also**: the living collection [Awesome Harness Engineering — Task Runners & Orchestration](#) gathers more consultable resources for this dimension (patterns, articles, and implementations), curated by problem.

The state of the art

1. Three philosophies — tool, service, teammate

Round 1's framing persists and got reinforced by round 2. **Subagent-as-tool**: one-shot, contained, with guardrails (opencode `task` → child session, depth 1; Aider's architect → editor split, depth 1). **Subagent-as-service**: registry, termination contracts, remote reach (gemini-`cli invoke_agent` + A2A; Codex `multi_agents_v2` with a **persisted agent graph** and ~100 profiles; Goose `orchestrator` lead/worker). **Subagent-as-teammate**: persistent teams with continuous communication (OpenHarness Swarm with a mailbox + a git worktree per member; Hermes with a **Kanban dispatcher** and structured handoffs).

2. The primary gain is context isolation — not parallelism

What the three round-1 harnesses already showed, the industry consolidated: the subagent is valuable because it **reads a lot and returns a little**. That is why Claude Code models it as a *fresh*, isolated instance, and why the git worktree (OpenHarness) matters — it isolates parallel *edits*, not just reads. This is the same principle as ch. 09's "context scoped per subtask" (Beyond Entangled Planning): the subagent is the vehicle of context scoping.

3. The central tension: parallelizing costs, and most failures are design failures

The dimension's decision axis is the Anthropic × Cognition tension. Orchestrator-worker buys breadth (+~90% in research) at ~15× **tokens**; single-thread avoids the "telephone" game but serializes. MAST closes the argument with data: most MAS failures are failures of *specification and coordination*, not of the model — which explains why every serious harness surrounds subagents with **guardrails**: bounded depth (opencode/Aider depth 1; OpenClaw 1–5), termination contracts (gemini-cli GOAL/MAX_TURNS/TIMEOUT), permissions **degraded by depth** (OpenClaw: a subagent never gets `message/gateway/cron`), and the extreme expression — **IronClaw deny-filters spawn_subagent in all production profiles** (the design supports it; policy forbids it until there is trust). The design rule: decompose-and-parallelize is a cost/benefit gate, with a single-agent baseline as the control.

4. The turn: orchestrating other vendors' harnesses

The frontier round 2 made concrete: the subagent can be *another harness*. OpenClaw orchestrates Claude Code, Gemini CLI, opencode, and Codex as subagents via an **ACP** runtime; OpenHands (Canvas) orchestrates Claude Code, Codex, and Gemini via **ACP** profiles; gemini-cli is an A2A client **and server**. With ACP-IBM (Agent Communication Protocol) converging into A2A under the Linux Foundation, the *agent card* becomes the universal contract for cross-system delegation. Orchestration has stopped being internal to the harness and become interoperability (ch. 17).

Round ext-1 addendum (2026-07-31): workspace isolation became infrastructure. The corpus isolated the subagent's **context**; **Grok Build** (in Portuguese; xAI, opened on 2026-07-15) closes the other half — the **filesystem**. Each `spawn_subagent` with isolation active gets its **own git worktree** created by a dedicated crate (`xai-fast-worktree`: parallel CoW, O(1) BTRFS snapshots, overlays, metadata with auto-GC), with merge-back as a protocol operation (`x.ai/git/worktree/apply`) and graceful fallback to the shared workspace. The lesson is not "use worktrees" (several harnesses have them); it is the investment in making them **cheap enough for the agent to use without thinking** — parallel subagents that edit stop fighting over the working tree. Confirmed in the code (`agent/subagent/handle_request.rs`), not just the announcement.

EXECUTIVE SUMMARY

What's most modern: the subagent as context isolation with an explicit contract; the coordination choice (handoff × tool × supervisor × ledger); guardrails motivated by real failure modes (MAST); cross-vendor delegation via A2A; and — since round ext-1 — workspace isolation via cheap worktrees (Grok Build). **What to steal:** give every subagent a contract (objective/format/tools/boundaries) and isolated context; bound depth and degrade permissions by depth; always compare against a compute-matched single agent; if subagents edit in parallel, isolate the filesystem (worktree), not just the context; and, if you orchestrate across systems, speak A2A.

Hands-on — harness-zero, step 9

Step 9 (harness-zero/etapas/09-subagentes/) adds a `task` tool that launches a **subagent in a child session**: its own context, **permissions derived and restricted** from the parent session, and **maximum depth 1** (a subagent does not spawn a subagent) — the guardrails MAST justifies, in their minimal form. The subagent receives a contract (objective + output format), runs its own loop, and returns only the summary to the parent. Completeness exercise: you add permission degradation by depth and a configurable termination contract (objective + per-subagent timeout).

Check your understanding

1. Your orchestrator needs to understand 40 files to decide on a refactor, but you don't want 40 dumps in the main context. How does a subagent solve this, and what is the real gain? (Context isolation — the subagent reads the 40 and returns only the conclusion; the primary gain is not parallelism.)
2. A colleague proposes running 5 subagents in parallel to speed things up. Name the main risk (with a name from the literature/industry) and the gate you apply before accepting. (Telephone game / conflicting implicit decisions — Cognition; coordination failures — MAST. Gate: ~15× token cost/benefit + compute-matched single-agent baseline.)

3. You want your harness to delegate a subtask to another vendor's agent. What mechanism do you use, and what is the "contract"? (A2A; the Agent Card announces identity/skills/endpoint/auth, and the Task is the unit of delegated work.)
-

Appendix A — How each repository handles subagents and orchestration

Per-harness evidence, with paths — supplemented online, expanded each round.

opencode (round 1) — contained delegation

`task tool (tool/task.ts)` → subagent in a **child session** (`parentID`), **derived, restricted permissions** (`agent/subagent-permissions.ts`), depth 1. Agents in markdown with mode `primary|subagent|all`; built-in `build/plan/general/compaction`. Experimental background mode (`BackgroundJob`) with a `task_id` to **resume the subagent session**.

gemini-cli (round 1) — from local subagent to remote

`invoke_agent` over an `AgentRegistry` (`packages/core/src/agents/registry.ts`); built-in `codebase-investigator`, `generalist`, `cli-help`, `browser`, `skill-extraction`, each with a `ModelConfig`. Explicit termination (`AgentTerminateMode`: `GOAL/MAX_TURNS/TIMEOUT`). Its exclusive: **A2A** client+server (`@a2a-js/sdk`, agent cards). Its own delegation evals.

OpenHarness (round 1) — teams, not subagents

`Swarm` (`src/openharness/swarm/`, 11 modules): `AgentTool` with three backends (`subprocess`, `remote`, `in-process teammate`); `TeamRegistry`; a **mailbox** (continuous communication); **git worktrees** (`worktree.py`) for parallel edits; `permission_sync.py`. Tools `team_create/delete`, `send_message`.

Codex CLI (round 2) — persisted agent graph

Two API generations (`multi_agents_v2`: `spawn`, `send_message`, `followup`, `interrupt`, `wait`); ~100 subagent profiles in TOML; **agent-graph-store** (persisted graph), agent identity, inter-agent communication, `SubAgentStart/Stop` hooks; a `ThreadManager` coordinating parallel threads.

OpenClaw (round 2) — push-based spawn and external ACP

`sessions_spawn` creates isolated subagents with **push-based completion** (`sessions_yield` as polling-free waiting); nesting 1–5; tool policy **degraded by depth** (subagents never get `message/gateway/cron`). An **ACP** runtime orchestrates Claude Code, Gemini CLI, `opencode`, and `Codex` as subagents; `Swarm` via `Code Mode`.

Hermes (round 2) — Kanban dispatcher

`delegate_task` spawns child `AIAgents` with isolated context and safe non-interactive approval; a **Kanban dispatcher** in the gateway spawns workers with structured handoffs, blocking for human input, and heartbeats on long operations.

Goose (round 2) — SubRecipes and orchestrator

`summon` delegates to subagents (a child `Agent` with its own recipe, streamed events); **SubRecipes** with hierarchical composition and parallel/sequential execution; the `orchestrator` extension (`lead/worker`: `list/start/send/interrupt/stop`).

Aider (round 2) — architect → editor

The `architect_coder.py` split: a reasoning model produces the plan; after confirmation, a second coder (with its own `editor_model/editor_edit_format`) executes. Two-role orchestration with distinct models, fixed depth 1.

IronClaw (round 2) — elegant design, restrictive policy

Subagents as child-runs in the same pipeline, with unified gates/checkpoints and an E2E test — **but `spawn_subagent` is deny-filtered in all production profiles** (`TEMP(disable-spawn-subagents)`). The score reflects the available capability, not the design (which would be a 3). The extreme case of "guardrail beats capability".

OpenHands / ohmo (round 2)

OpenHands: SDK primitives (`openhands.sdk.subagent`) + per-organization **AgentProfiles**, including **ACP** profiles — the Canvas orchestrates Claude Code, Codex, and Gemini. ohmo: inherited Agent/Task/Team/SendMessage; an observed asymmetry (`/tasks run` blocked remotely, equivalent tools available to the model).

Grok Build (round ext-1) — worktrees as infrastructure ★

`agent/subagent/handle_request.rs`: `spawn_subagent` with `capability_mode` **intersected** with the type's toolset (`intersect_capability_modes`), max depth 1, `resume_from`, I/O contracts between personas; isolation via `WorktreeBuilder...` `worktree_kind(WorktreeKind::Subagent)` over `xai-fast-worktree` (CoW + O(1) BTRFS + auto-GC), merge via `x.ai/git/worktree/apply`; plugin agents forbidden from declaring `mcpServers/hooks/bypassPermissions`.

Pi (round ext-1) — the documented refusal

No subagents in the core, by manifesto ("There's many ways to do this. Spawn pi instances via tmux, or build your own"); the first-class example `examples/extensions/subagent/` spawns **full pi processes** (real context isolation) with 4 personas and 3 workflows — the feature exists as proof that the extension surface suffices.

n8n (round 2) — agent as another agent's tool

AI Agent Tool (`AgentTool.node.ts v3`): a full agent as another agent's tool — V3 runs the sub-agent's loop inline (`resolveSubAgentRequest`), with nested HITL forbidden; **ToolWorkflow** (sub-workflows as tools). Visual hierarchical orchestration.

Frameworks (frameworks round)

Agents SDK: handoffs × agents-as-tools; CrewAI: sequential × hierarchical (`manager_llm`); LangGraph: supervisor + workers as stateful nodes; AutoGen/Magentic-One: orchestrator with a ledger and replanning; Google ADK: coordinator/dispatcher + `Sequential/Parallel/Loop`. Frameworks expose as first-class API what coding harnesses implement by hand.

11 — Verification & Evals

🕒 state of the art 2026-07 · revised 2026-07-26 · 📖 ~14 min read

Learning objectives

By the end of this chapter, you should be able to:

1. **Distinguish** the three questions of verification (does the harness work? · does the agent behave? · is the work correct?) and each one's technical answer;
2. **Explain** why *intrinsic* self-correction is not enough and verification must be external and anchored in signal (tests, LSP (Language Server Protocol), tools);
3. **Evaluate** *reward hacking* — the agent gaming the verifier — and the defenses (held-out, immutable tests, anti-mock, verifying the final state);
4. **Recognize** the LLM (Large Language Model) judge's biases (position, verbosity, self-preference) and how to mitigate them;
5. **Implement** a harness-zero eval suite (judge + recorded responses) in step 10.

The problem

How do you know the agent works? The question unfolds into three, with different technical answers:

1. **Does the harness work?** — classic software tests over the harness code (loop, tools, permissions).
2. **Does the agent behave well?** — evals: the emergent behavior (does it use the right tools? is it frugal? does it respect plan mode? does it resist injection?) under regression testing.
3. **Is the agent's work correct?** — runtime verification: signals (LSP, tests, lint) fed back to the model during the task.

The second is the hardest and the most neglected: agent behavior is stochastic, expensive to test, and changes silently with every model or prompt swap. And there is a fourth question that round 2 made unavoidable: **is the agent cheating the verifier?**

Scientific foundations

The science of agent verification has three hard messages — and all push toward the same place: verification that is **external and anchored**.

- **Grading by execution, not by appearance** — [SWE-bench](#), [arXiv 2310.06770](#) (ICLR '24) verifies by applying the model's patch and running the repository's **real, hidden tests** (FAIL_TO_PASS + PASS_TO_PASS). Decision: for code, the only trustworthy signal is "the real tests passed", not diff similarity. And [SWE-agent](#), [arXiv 2405.15793](#) shows that **tool ergonomics** (the Agent-Computer Interface) drives success as much as the model does.
- **Intrinsic self-correction is not enough** — [Large Language Models Cannot Self-Correct Reasoning Yet](#), [arXiv 2310.01798](#) is the decisive counter-result: without external feedback, asking the model to "revise" can *degrade* correct answers. Decision: "asking the model to check itself" **is not** a verification strategy — the harness must supply a verifier. [CRITIC](#), [arXiv 2305.11738](#) shows the way: **tool-anchored** self-critique (does the code run? does the fact check out?) beats introspection; [Self-Consistency](#), [arXiv 2203.11171](#) gives the cheap version (sample paths + vote) for checkable answers.
- **The LLM judge works — with biases** — [Judging LLM-as-a-Judge](#), [arXiv 2306.05685](#) measures ~80% agreement with humans, but documents **position, verbosity, and self-preference** biases. Decision: randomize/swap the order of the answers and average, provide a rubric and a reference answer, and calibrate against a human gold set ([survey](#), [arXiv 2411.15594](#)) — a single judge call is not ground truth. And verify the **final state of the world**, not the transcript: [t-bench](#), [arXiv 2406.12045](#) shows that `pass@1` hides brutal inconsistency (`pass@8` < 25%).
- **The agent games the verifier** — the new and most important theme: with [verifiable rewards](#) (RLVR / Tulu 3, [arXiv 2411.15124](#)) a deterministic verifier is a signal and a reward that is harder to defraud — *but reward hacking*, [arXiv 2606.15385](#) and [randomized tests against cheating](#), [arXiv 2606.07379](#) show that agents practice *specification gaming* zero-shot: they delete asserts, call

`sys.exit(0)`, patch `pytest`. Decision: keep a **held-out** ground-truth metric the agent never optimizes, and **immutable tests** it cannot touch.

(Full bibliography and pointers: `livro/bibliografia.md`.)

Industry sources

- **The benchmark is the standard — and it is contaminable** — [SWE-bench Verified \(OpenAI\)](#) is the *human-audited* 500-task subset, created because raw SWE-bench had ambiguous specs and broken tests that failed correct solutions (audit the verifier before trusting it). But [OpenAI stopped reporting SWE-bench Verified](#) due to contamination/memorization — an eval needs rotation and held-outs to remain a signal. [Terminal-Bench \(arXiv 2601.11868, repo harbor-framework/terminal-bench\)](#) brings the rigor to the terminal: each task ships **Docker + a human solution + verification tests**, grading the *final state of the environment*, not the transcript's plausibility.
- **Evals as an engineering discipline** — [Define success criteria and build evaluations \(Claude\)](#): define measurable criteria *beforehand*, force the judge to emit a discrete verdict and to reason before scoring. [Demystifying evals for AI agents \(Anthropic\)](#) decomposes the eval into components (task · trial · agent harness · eval harness · trace · grader · suite) and insists: **grade the final state, not the last message** (an answer can "sound right" while the task failed). And report the [standard error of the mean](#) to distinguish real regression from noise.
- **Verification inside the loop** — [Effective harnesses for long-running agents \(Anthropic\)](#): each session runs the tests, **verifies the feature end-to-end as a user would** (browser automation), leaves a progress log, and commits clean. And [Claude Code best practices](#) elevates TDD to the strongest agentic pattern: write the tests first, confirm they fail, **commit them as a checkpoint, and implement without editing them** — committing the tests up front is the net that reveals when the agent cheats by altering the test instead of fixing the code.
- **Versioned eval tooling** — [OpenAI Evals, Inspect \(UK AISI\)](#) (Dataset + Solver + Scorer, with a Docker/K8s sandbox — the eval and the sandbox are one system), [promptfoo](#) (a versioned `promptfooconfig.yaml` as a CI gate), [Braintrust](#) and [LangSmith](#) (rubric as config, human corrections become few-shot). Decision: the checks live in version control and run in CI like any test.
- **Verification became adversarial** — [Natural emergent misalignment from reward hacking \(Anthropic\)](#): agents learn to *game the verifier* (exit before the tests, patch `pytest`, delete asserts) and the habit **generalizes into broader sabotage**. Decision: harden the verifier (randomized/held-out tests, immutable test files) and never let the agent touch its own grader — [The Verification Horizon \(arXiv 2606.26300\)](#) warns that when the agent's capability outstrips the verifier, reward hacking resurfaces; the verifier has to *evolve* (tests → rubric → interactive judges).
- **See also**: the living collections [Awesome Harness Engineering — Verification & CI Integration](#) and [Awesome Harness Engineering — Evals & Verification](#) gather more consultable resources for this dimension (patterns, articles, and implementations), curated by problem.

The state of the art

1. Three questions, three champions (and the gap that closed)

Round 1's framing persists: **OpenHarness** best tests *the harness* (121 files per subsystem), **gemini-cli** best tests *the agent* (evals with a judge + regression baselines), **opencode** best verifies *the work* (runtime LSP → diagnostics to the model in the same turn). But round 1's revealing gap — "only one of the three tests behavior under attack" — **closed** in round 2: IronClaw treats cross-tenant isolation as a first-class test citizen (with *trace parity* against OpenClaw), and ohmo has 96 adversarial tests (a session does not leak to another sender, `/config` does not leak secrets).

2. The right verification is external and anchored — because the internal kind fails

The central scientific finding (intrinsic self-correction degrades; tool-anchored works) is exactly what opencode's **runtime LSP** does: the agent discovers it broke typing on the next turn, not in CI. It is the same thesis as **Aider's reflection** (triggered by failing lint/tests, not by introspection) and **Hermes's verify-on-stop** — the agent is *forced* to verify before stopping, with `verification_evidence.py` tracking the evidence. Verifying stopped being a hope and became an **enforced stage of the loop**.

3. Behavioral evals became table stakes — and per category

In round 1, only gemini-cli treated behavior as a regression surface. In round 2 that became the norm: **Goose** publishes **Harbor** (on the Terminal-Bench framework, 89 tasks, with a **real leaderboard**: stock 50.6% / code-mode 57.3%); **Codex** has ~660 insta snapshots; **Hermes** runs `mini_swe_runner` (SWE-bench style); **n8n** turned evals into a *product* (Evaluation nodes + LLM-judge). And **per-category** evals emerged: OpenClaw's **Personal Agent Benchmark Pack** (10 category scenarios — `personal-redaction-no-secret-leak`, `personal-approval-denial-stop`, `personal-no-fake-progress`, `personal-memory-preference-recall`), the first behavioral benchmark *of the personal-agent category*. A harness without evals doesn't know what it lost in the last prompt tweak.

4. The adversary is the agent itself

The most serious turn: verification became **adversarial**. The literature shows agents deleting asserts and patching pytest to "pass"; the industry's defense is convergent — **immutable tests** (commit the tests first; the agent does not edit them), **held-out/randomized** (the agent cannot overfit what it does not see), an **anti-mock policy** (opencode's test `AGENTS.md` forbids mocks that lie; the `http-recorder` records real calls), and **snapshots with drift-check** (OpenClaw) for determinism where the judge is expensive. Verification is no longer just measuring correctness — it is preventing cheating.

Addendum (2026-07-31, full text verified): how to evaluate the harness itself — three rules from a methods paper. The preprint [Rethinking the Evaluation of Harness Evolution for Agents](#) (AI2/UW/indep., 14 Jul 2026) tests the "automatic harness evolution" fashion and finds an uncomfortable result: under a **matched budget** ($K=5$ for all methods), it "does not consistently outperform simple test-time scaling methods" — on Terminal-Bench 2.1 (89 tasks, 3 models), pure parallel sampling took mean pass@1 from 68.2 to 72.3 (Table 1) while evolution actually made GPT-5.4 **worse** (75.3 → 69.7); with unit tests available, parallel sampling opens up 86.0 versus 75.8 (Table 2); and on held-out tasks evolution's average gain is **+0.6** (Table 3) — "their gains largely stem from making multiple attempts" (§4.3), because "most edits memorize fixes rather than distilling strategies" (§5.1), accumulating "context bloat that can offset the remaining gains". The three rules that remain for anyone evaluating harnesses (including this book): (1) **matched budget** — every gain attributed to design must be reported against a sample-repetition baseline with the same compute; (2) **search/evaluation separation** — held-out is mandatory, or the gain is overfitting to the set; (3) **instrument sensitivity** — the authors themselves suspect that "Terminal-Bench may simply not be very sensitive to harness design" (§5.2): a benchmark good at measuring harnesses needs headroom AND performance that depends on the harness, otherwise the signal is model capability. For this book's method (a 0–3 rubric via code reading), the paper refines without contradicting: the rubric measures the structural property without going through the sampling-contaminated channel — but it inherits the duty of **convergent validity** (high scores should predict held-out performance), the risk of overfitting if the yardstick is calibrated by looking at the systems one wants to score well, and §5.1's warning: penalize memorization and context bloat, not just missing features. This converses with ch. 16: if evolving the harness automatically yields less than resampling, cheap self-improvement lives in **knowledge** (skills/memory), not in **structure**.

EXECUTIVE SUMMARY

What's most modern: anchored external verification (LSP/tests in the loop, verify-on-stop); behavioral evals as table stakes and per category (Harbor, Personal Agent Benchmark Pack); the LLM judge used with bias control; the defense against reward hacking (immutable tests, held-out, anti-mock); and, for anyone evaluating their own harness, the addendum's three rules (matched budget, held-out, an instrument sensitive to design). **What to steal:** feed real signal back to the model in the same turn (LSP/tests), don't trust self-checking; commit the tests first and don't let the agent edit them; grade the final state, not the last message; and treat behavioral evals as first-class regression.

Hands-on — harness-zero, step 10

Step 10 (`harness-zero/etapas/10-evals/`) gives harness-zero its own eval suite: **recorded LLM responses** (deterministic replay in CI, cheap and stable) to test the loop and the tools without calling the API, and a minimal **LLM judge** that scores whether the agent's behavior meets qualitative criteria (used the right tool? respected plan mode?). Faithful to the chapter's discipline: the judge emits a discrete verdict and the suite runs in CI like any test. Completeness exercise: you add an **immutable test** case — a task whose test the agent is forbidden to edit — and observe the difference between "passed" and "cheated".

Check your understanding

1. Your agent says "I fixed the bug and the tests pass". Why is that, by itself, not verification — and what do you do instead of trusting it? (Intrinsic self-correction/self-reporting is not enough — 2310.01798; run the real, hidden tests and grade by execution — SWE-bench.)
 2. After giving your agent RLVR, the score goes up but the product gets worse. What probably happened, and which two defenses do you apply? (Reward hacking — the agent games the verifier, e.g. deletes asserts; defenses: a held-out metric it never optimizes + immutable tests.)
 3. You use an LLM judge to score open-ended answers. Name one known bias and how to mitigate it. (Position/verbosity/self-preference; swap the order and average, rubric + human gold set.)
-

Appendix A — How each repository handles verification and evals

Per-harness evidence, with paths — supplemented online, expanded each round.

gemini-cli (round 1) — behavior under continuous regression

Four suites: (1) `evals/` — ~45 behavioral tests with an **LLM judge** (`llm-judge.ts`) covering frugality, hierarchical memory, plan mode, delegation, shell safety, **prompt injection via MCP (Model Context Protocol)**, and sandbox recovery; (2) `integration-tests/` — deterministic E2E with **recorded responses** (`.responses`); (3) `memory-tests/` — regression against `baselines.json`, nightly; (4) `perf-tests/` — CPU/startup, nightly. Behavior as a first-class regression surface.

opencode (round 1) — verification during the task

Runtime LSP (`packages/opencode/src/lsp/`): edits trigger diagnostics fed back to the model. An explicit **anti-mock policy** (the `test/ AGENTS.md` forbids mocks) + `http-recorder` (records/replays real HTTP deterministically). Mandatory typecheck (`bun typecheck`).

OpenHarness (round 1) — E2E with a real model

121 files in `tests/`, ~31 subfolders mirroring each subsystem. E2E suites with **real model calls** (`scripts/test_harness_features.py`) and tests against real ecosystem artifacts (`test_real_skills_plugins.py` runs skills from `anthropics/skills` and plugins from `claude-code`). The `harness-eval` skill packages the E2E validation.

Goose (round 2) ★ — Harbor with a public leaderboard

Harbor (`evals/harbor/`): a benchmark on the Terminal-Bench framework (89 tasks) comparing harnesses/models/builds by pass-rate, cost, tokens, and turns — with a **real leaderboard in the README** (stock ~50.6%, code-mode 57.3%) and post-processing LLM-judges; `goose-self-test.yaml`; compaction with ~15 inline tests.

Codex CLI (round 2) — snapshots at scale

~440 test files + ~660 **insta snapshots**; an E2E suite with real turns and a mocked backend; per-platform sandbox policy tests; remote compaction parity; multi-layer CI (nextest per platform, Bazel, postmerge).

Hermes (round 2) ★ — verify-on-stop

32 test subdirectories; a **verify-on-stop nudge** (the agent is forced to verify before stopping, with `verification_evidence.py` tracking evidence); `batch_runner.py` (batch trajectories) and `mini_swe_runner.py` (SWE-bench-style evaluation). Research-oriented.

OpenClaw (round 2) ★ — the category's benchmark

~8,649 test files; **prompt snapshots with drift-check** in CI; a QA stack with a synthetic channel and a YAML catalog of scenarios; the **Personal Agent Benchmark Pack** — 10 category scenarios (`personal-redaction-no-secret-leak`, `personal-approval-denial-stop`, `personal-no-fake-progress`, `personal-memory-preference-recall...`), runnable in mock. The first behavioral benchmark of the *personal-agent* category.

IronClaw (round 2) ★ — isolation as a test citizen

~415 test files; fuzzing; **cross-tenant/agent/project/thread isolation tests as first class** (`reborn_*_scope_isolation_parity.rs`); **recorded trace parity against OpenClaw**; mechanized architecture tests; a rule requiring denial/redaction/escape tests for any sandbox change.

ohmo (round 2) — channel-adversarial

96 adversarial tests (75 in the gateway): a session does not restore another sender's messages, `/config show` does not leak secrets, `/group` history sanitized before becoming context. Gap: no permission/sandbox tests — exactly the weak dimension.

Aider (round 2) — anchored reflection + edit-format leaderboard

Reflection (`reflected_message`, max 3) triggered when the linter finds errors or tests fail (always with human confirmation) — **reactive, anchored** self-correction, not introspection. Famous for empirically measuring the edit format per model (`percent_cases_well_formed`) on its own leaderboard.

n8n (round 2) — eval as a product

The **Evaluations** feature (Evaluation Trigger + Evaluation nodes, enterprise UI) to run datasets against workflows; an eval suite with an LLM-judge in the AI Workflow Builder; per-workflow integration tests. Verification packaged as a sellable feature.

OpenHands (round 2) — the eval that migrated

Agent evals **absent from this repo** (score 0): the classic `evaluation/` directory (the SWE-bench harness OpenHands is historic for) migrated to the `software-agent-sdk`. Here there are 115 unit-test files for the app-server, but zero agent evals — a reminder that the boundary of what gets evaluated depends on where the core lives.

Frameworks (frameworks round)

Frameworks treat evals as API: versioned eval harnesses (OpenAI Evals), Solver+Scorer with a sandbox (Inspect), mixed code+judge scorers (Braintrust/autoevals), rubric-as-config (LangSmith). What coding harnesses assemble by hand, the framework ecosystem exposes as dedicated tooling.

12 — Extensibility

🕒 state of the art 2026-07 · revised 2026-07-26 · 📖 ~11 min read

Learning objectives

By the end of this chapter, you should be able to:

1. **Explain** why extensibility is "open for extension, closed for modification" — extension points instead of forking;
2. **Distinguish** the four extension axes (hooks · commands/skills · plugins · providers) and what each one solves;
3. **Compare** the three ecosystem strategies — depth, packaging, interoperability;
4. **Evaluate** extension code as an attack surface (the *trust triangle*) and the defenses (scanning, trust envelope, managed settings, least-privilege);
5. **Implement** a pre/post-tool hook subsystem with the hook's return value as the control channel in harness-zero (step 11).

The problem

No harness covers every workflow; extensibility decides whether the user **adapts** the harness or **abandons** it. The established axes:

1. **Hooks** — user code intercepting the lifecycle (before/after a tool, compaction, session).
2. **Skills / custom commands** — capabilities packaged as markdown/config, loaded on demand.
3. **Plugins / extensions** — distributable packages aggregating tools, commands, hooks, and config.
4. **Model providers** — the most strategic extension: does the harness work with any model, or is it one model's showcase?

The rule uniting all four is old: **open for extension, closed for modification** — the user extends without editing (or forking) the core.

Scientific foundations

Honest editorial record (Principle I): **there is no academic canon of "agent harness extensibility"** — it is a real gap. The durable citations come from the classic software engineering of extensible architectures and from plugin-ecosystem security, which transfer directly.

- **Extension points, not forks** — the open-closed principle (Meyer, 1988; Martin, 1996) and Eclipse's plug-in architecture ([Birsan, ACM Queue 2005](#)) provide the foundation — and the "*plug-in hell*" warning: poorly designed extension points become debt. Decision: expose explicit *seams* (events, well-known directories), not ad-hoc points.
- **Minimal core, pluggable extensions** — the Microkernel pattern (Buschmann et al., *POSA* v.1, 1996) and its agentic incarnation, [AIOS, arXiv 2403.16971](#) (a kernel that isolates scheduling/memory/tools from agent applications), support the "harness as microkernel" posture: a small core that serves as a socket.
- **Mechanism × policy** — [Hydra \(Levin et al., SOSP '75\)](#) is the origin of "separate mechanism from policy". Translated: the harness provides the *mechanism* (invoking a tool, dispatching a hook, loading a provider); the *extension* provides the policy. That is why adding a model provider can be "writing a file".
- **Third-party extensions are not trustworthy** — the best on-topic citation is [LLM \(Large Language Model\) Platform Security: ChatGPT Plugins, arXiv 2309.10254](#) (AIES '24): a platform/plugin/user *trust triangle* with concrete exploits (session hijacking via a malicious plugin). And the empirical base for over-privilege comes from browser-extension security ([Barth et al., NDSS '10](#): 88% of extensions request more power than they need). Decision: least-privilege + isolation + verification — the same argument as ch. 06's *tool poisoning*.

(Full bibliography and pointers: [livro/bibliografia.md](#).)

Industry sources

- **Hooks: exit code as the control channel** — [Claude Code's hooks](#) expose ~31 lifecycle events (`PreToolUse`, `PostToolUse`, `Stop`, `SessionStart`, `UserPromptSubmit`, `PreCompact`, `SubagentStop`...) where the harness runs user commands; the **exit code is the channel** (0 = proceed / JSON on stdout with allow-deny-ask; 2 = block with stderr fed back to the model). Decision: teams enforce policy (block `rm`, `redact`, `.env`, auto-lint) **deterministically and without patching the harness**. And Codex implements the same pattern independently (hooks + `allow_managed_hooks_only` for enterprises) — hooks are a **cross-vendor** standard, not one vendor's quirk.
- **Plugin = the packaging unit; marketplace = the catalog** — the [Claude Code plugin model](#): a plugin aggregates skills, subagents, hooks, MCP (Model Context Protocol), and LSP (Language Server Protocol) into an installable package (`/plugin install name@marketplace`); a [marketplace](#) is a git repo with `.claude-plugin/marketplace.json`. Installs at `user/project/local/managed` scope, with **pinning to SHAs** and a two-tier trust model (curated official marketplace + community with security triage). Decision: third-party extension becomes distributable **and governable** without forking.
- **Custom commands became a file-drop (and AGENTS.md is the open standard)** — in Claude Code, slash commands were [absorbed by skills](#): dropping a file into `.claude/commands/` or `.claude/skills/` creates the command, with no registration or build. And [AGENTS.md](#) became the **open, multi-tool** config format — read by Codex, Cursor, Cline, Windsurf, Gemini CLI, and Claude Code. Decision: the extension point is "drop a file in a well-known directory", and the format is portable across harnesses.
- **Settings as an enforcement surface** — [Claude Code's config](#) is a precedence stack (Managed > CLI > local > project > user); most keys override, but **permission rules merge**, and **managed settings cannot be overridden** (a security team denies tools/marketplaces for the whole company). Decision: config is not preference, it is enforcement (connects to ch. 07).
- **Extensibility is also a context budget** — [advanced tool use \(Anthropic\)](#) reframes it: with unlimited tool libraries, the extension must be **loaded on demand**, not registered up front; and plugins toggle on/off to control system-prompt cost. Decision: an extension point that always injects context does not scale — late loading is part of the design (connects to chs. 03 and 05).
- **See also:** the living collection [Awesome Harness Engineering — Debugging & Developer Experience](#) gathers more consultable resources for this dimension (patterns, articles, and implementations), curated by problem.

The state of the art

1. Three ecosystem strategies

Round 1's framing persists and got reinforced. **Depth:** hooks reach points the others don't expose — `opencode` transforms messages and the system prompt before sending, intercepts `permission.ask`, and registers auth providers. **Packaging:** the *extension* as a complete distribution unit (`gemini-cli` aggregates MCP+commands+hooks+policies into one package; Codex with a manifest + marketplace + App Server JSON-RPC (Remote Procedure Call)). **Interoperability:** adopting the leader's formats instead of inventing your own (`OpenHarness` with `SKILL.md/.claude-plugin`; `IronClaw` with a compatible `SKILL.md`).

2. The interoperability bet is winning — the "MCP of extensibility"

What in round 1 was the most underrated axis became the dominant trend: **extension formats are converging on standards portable across harnesses**. `SKILL.md/AgentSkills` (`OpenClaw` uses the `agentskills.io` standard; `IronClaw` declares compatibility with `OpenClaw/Claude`), `.claude-plugin` (adopted by `OpenHarness`), and above all **AGENTS.md** (read by six different harnesses) are doing for extensibility what MCP did for integration. Even the **hook vocabulary** converged — Codex's event set is practically `OpenHarness's` and `Claude Code's` (`PreToolUse/PostToolUse/...` with `Approve/Block/Deny/Ask` decisions). Extensibility is ceasing to be a per-harness silo.

3. Marketplaces and security scanning — round 1's gap closed

In round 1, only `gemini-cli` treated extension code as an attack surface. In round 2 that became the norm, exactly as the plugin *trust triangle* predicted: **OpenClaw** has the **ClawHub** registry with a *trust envelope* + scanning (`VirusTotal/ClawScan`); `Claude Code` has a curated official marketplace + community with security triage and **SHA pinning**; `n8n` runs `scan-community-package`; **Goose** checks extensions for malware before loading. Added to the **managed settings** that deny marketplaces enterprise-wide, extension distribution became infrastructure *with containment* — the least-privilege the over-privilege literature demands.

4. Provider-agnosticism became declarative config

The mechanism × policy separation applied to the model: adding a provider stopped being code and became a file. **Goose** has 37 **declarative providers via JSON** (an OpenAI-compatible provider = one file); **opencode** has ~26 loaders + hundreds of models via models.dev; **Hermes** has a subclassable **ProviderProfile** (Nous Portal with 300+ models). The model-agnostic harness — treating the provider as pluggable policy — beat the single-vendor showcase.

5. The next frontier: the harness that extends itself

The embryo of self-extension is already visible: **IronClaw** has **automatic skill extraction** (`learning.rs`) with usage and confidence metrics — the harness observes its own work and writes new skills. It is the bridge to ch. 16 (learning) and to the Voyager/ToolMaker lineage: extensibility that doesn't wait for the user.

EXECUTIVE SUMMARY

What's most modern: format convergence (`SKILL.md/claude-plugin/AGENTS.md` as portable standards); marketplaces with security scanning and managed settings; hooks with exit-code as a cross-vendor channel; declarative provider-agnosticism; and the beginning of self-extension. **What to steal:** expose explicit seams (named events, well-known directories) instead of ad-hoc points; adopt portable formats instead of inventing your own; treat third-party extensions as untrusted (scan + least-privilege + managed deny); and make loading late so you don't blow the context.

Hands-on — harness-zero, step 11

Step 11 (`harness-zero/etapas/11-hooks/`) gives harness-zero a **pre/post-tool hook** subsystem: before each tool call, hooks are registered functions (`@hooks.pre_tool/@hooks.post_tool`) and the **hook's return value is the control channel** ("`block:reason`" blocks and feeds the reason back to the model; a dict adjusts the arguments) — the completeness exercise proposes the products' external variant: run a user command and read the exit code (0 proceeds; non-zero blocks with the stderr). It is the mechanism (the harness dispatches the hook) separated from the policy (the user decides what the hook does) — the chapter's thesis in ~40 lines. Completeness exercise: you add a `PostToolUse` that runs a linter and returns the errors to the model, and a minimal trust gate (the hook only runs if the directory is trusted).

Check your understanding

1. Why does "open for extension, closed for modification" lead to *hooks* and *plugins* instead of telling the user to fork the harness? (Extension points preserve the core and upgradability; a fork diverges and rots.)
2. You are going to allow a third-party plugin marketplace. Name the central risk (with its name from the literature) and two concrete defenses. (*Trust triangle* / over-privilege; defenses: security scanning + SHA pinning + managed settings that deny + least-privilege.)
3. Your harness needs to support a new model provider without a release. Which design principle makes that "writing a file"? (The mechanism × policy separation — the harness supplies the invocation mechanism, the file supplies the provider's policy.)

Appendix A — How each repository handles extensibility

Per-harness evidence, with paths — supplemented online, expanded each round.

opencode (round 1) — deep hooks and radical provider agnosticism

Plugins are functions returning Hooks (`packages/plugin/`): ~15 points, including rare ones — transforming messages/system prompt before sending (`experimental.chat.messages.transform`), intercepting `permission.ask`, customizing compaction, and **registering auth providers** (`auth`). Custom tools auto-loaded from `tool/`. And ~26 **provider loaders** + hundreds of models via models.dev, on the Vercel AI SDK (Software Development Kit) — the most model-agnostic in production.

gemini-cli (round 1) — the all-in-one package

Extensions (`gemini-extension.json`): an installable package aggregates MCP servers, custom commands, hooks, **permission policies**, skills, and themes. Custom commands in TOML (`FileCommandLoader`). Hooks as a subsystem (`packages/core/src/hooks/`) with a **trust gate** (`trustedHooks.ts` — they only run in trusted folders). Providers: the Google ecosystem.

OpenHarness (round 1) — compatibility as strategy

Markdown skills also loaded from `~/.claude/skills` and `~/.agents/skills` (`SKILL.md` layout); plugins in the `.claude-plugin/plugin.json` format (12 real plugins tested); hooks cover **10 events** with **hot-reload**. Providers as named "workflows" (Anthropic/OpenAI-compatible, Copilot, Kimi, GLM, Ollama...).

Codex CLI (round 2) — full hooks + marketplace + App Server

Full hooks (`hooks/`: `PreToolUse/PostToolUse/PreCompact/SessionStart-End/UserPromptSubmit/Stop/SubagentStart-Stop`, with `Approve/Block/Deny/Ask` decisions) and the enterprise knob `allow_managed_hooks_only`; plugins with a manifest and marketplace; skills; configurable providers; profiles; Python/TS SDKs; the **App Server JSON-RPC** as the programmatic backbone.

OpenClaw (round 2) ★ — registry with security scanning

Skills in the **AgentSkills** standard (`agentskills.io`) with 6 precedence levels and the public **ClawHub** registry with a *trust envelope* + scanning (VirusTotal/ClawScan); **159 plugins** (tools, channels, providers, hooks, media) with a Plugin SDK; dozens of LLM providers with failover and auth rotation.

IronClaw (round 2) ★ — compatible and self-extending

A **SKILL.md format compatible** with OpenClaw/Claude; v2 skills with executable snippets, usage/confidence metrics, and **automatic skill extraction** (`learning.rs`); extensions via WASM/MCP/first-party **without restart**; configurable providers (NEAR AI, Gemini OAuth...).

Goose (round 2) — declarative providers and branded distros

Three axes: MCP extensions (6 transport/origin types); recipes/skills; and providers — native + **37 declarative providers via JSON** (adding an OpenAI-compatible provider = creating a file). `CUSTOM_DISTROS.md` (branded distros); `goose-sdk` for embedding; extension malware checks before loading.

Hermes (round 2) — ProviderProfile and plugins

A subclassable `ProviderProfile` (**Nous Portal** with 300+ models under subscription, OpenRouter, your own endpoint); a plugin system (20 directories, a toolset registry, session hooks); Anthropic/Bedrock/Codex/ACP (Agent Client Protocol) adapters.

OpenHands (round 2) — marketplaces and dependency injection

Skill/plugin marketplaces (instance/org/personal); LLM + agent profiles; a pluggable Git-integrations layer; third-party agents via ACP; **sandbox/event-store backends swappable via dependency injection**; litellm for providers.

n8n (round 2) ★ — the catalog as extensibility

The strongest point: **the 400+ integration nodes become a tool pool** without writing code (via `usableAsTool` + `$fromAI`); community nodes with a security scanner (`scan-community-package`); ~20 model providers (`LmChat*`).

ohmo (round 2) — extra roots

`~/.ohmo/skills` and `~/.ohmo/plugins` as roots coexisting with the project's; plugins load tools, slash commands, and MCP servers; skills become channel commands; arbitrary per-channel `channel_configs`.

Frameworks (frameworks round)

Frameworks expose extensibility as API: tool registration/`@tool`, lifecycle callbacks/hooks, provider adapters (litellm/model providers), and — increasingly — reading `AGENTS.md`. The portable format (`AGENTS.md`, `SKILL.md`) is what brings frameworks and coding harnesses together into a common ecosystem.

13 — Interfaces

🕒 state of the art 2026-07 · revised 2026-07-26 · 📖 ~12 min read

Learning objectives

By the end of this chapter, you should be able to:

1. **Argue** why "core with front-ends" beats "front-end with an agent inside" — and how drawing the boundary early maximizes the possible surfaces;
2. **Distinguish** the surfaces (TUI (Terminal User Interface), headless/SDK (Software Development Kit), IDE (Integrated Development Environment), chat, cloud) and what each one demands from the core;
3. **Evaluate** the interaction UX in light of HCI (Human-Computer Interaction) (mixed-initiative, levels of automation, over-reliance);
4. **Recognize** the surface as a security boundary (same turn contract, not a backdoor) and the shift to the *ambient*/inbox paradigm;
5. **Explain** why, with the loop behind ports, a second surface (headless) is a thin adapter, not a rewrite (step 0 of harness-zero).

The problem

The same agent must serve different audiences: the developer at the terminal, the CI script that needs JSON, the IDE that wants inline diffs, the manager who follows along over chat. The architectural question is a single one: **is the harness a core with multiple front-ends, or a front-end with an agent inside?** The harnesses we studied answered "core with front-ends" — and the quality of that separation determines how many interfaces are viable.

Established surfaces: **interactive TUI**, **headless/non-interactive** (-p with structured output), **IDE** (diffs, editor context), **CI/CD** (Actions), **agent protocols** (ACP (Agent Client Protocol), A2A (Agent-to-Agent)), **chat** (Slack, Telegram...) and, increasingly, **cloud/asynchronous**.

Scientific foundations

An honest editorial note (Principle I): **there is no academic canon of "agent harness interface"** — the gap is real. But the HCI of human-AI interaction grounds it with precision, and a recent trickle (2025-26) already addresses human-in-the-loop for agents.

- **When to act × when to ask** — [Principles of Mixed-Initiative UI \(Horvitz, CHI '99\)](#): the 12 principles about goal uncertainty, the cost/benefit of acting, and graceful handoff *are* the central decision of a harness — plan mode and approvals (chs. 07/09) are "passing the initiative", applied.
- **The autonomy dial is per stage** — the 10-level automation scale (Sheridan & Verplank, 1978) and the [types-and-levels model \(Parasuraman, Sheridan, Wickens, 2000\)](#) show that automation applies *independently* to each stage (acquisition · analysis · decision · action). Decision: the harness can **auto-collect context** (high automation) and still **gate the action** (low automation) — the dial does not have to be global.
- **The UX of "when it errs"** — [Guidelines for Human-AI Interaction \(Amershi et al., CHI '19\)](#): 18 guidelines by phase; the recovery ones (cheap correction/undo) explain why reversibility (ch. 08) is also an *interface* decision.
- **Oversight is fragile — design against that** — [To Trust or to Think \(Buçinca et al., CSCW '21\)](#) shows that explanation alone does **not** cure over-reliance; cognitive *forcing functions* do. The [over-reliance review \(Passi & Vorvoreanu, MSR-TR-2022-12\)](#) synthesizes the risk. Decision: approval must be a **deliberate act**, not a reflex click, and the surface cannot hide what the agent did. Recent human-in-the-loop agent work ([Magentic-UI, arXiv 2507.22358](#), with *action guards* = permission gating; [oversight design, arXiv 2510.19512](#)) operationalizes this.

(Full bibliography and pointers: [livro/bibliografia.md](#).)

Industry sources

- **One core, many surfaces (now doctrine)** — the [Platforms and integrations \(Claude Code\)](#) doc says it explicitly: "runs the same underlying engine everywhere, but each surface is tuned to a way of working" (CLI, Desktop, VS Code, JetBrains, Web, Mobile + Chrome, GitHub Actions, GitLab, Slack), with **config, project memory and MCP (Model Context Protocol) shared** across the local surfaces. Decision: build the agent as a single engine and treat terminal/IDE/web/mobile as interchangeable front-ends.
- **Headless is a Unix filter** — [Run Claude Code programmatically \(headless\)](#): `-p/--print, --output-format text|json|stream-json`, reads stdin and redirects stdout "like any command-line tool", with `--allowedTools/--permission-mode` so unattended runs never hang on a prompt. Decision: the interface is stdin/stdout + exit codes — the agent drops into pipes, build scripts and CI without a UI.
- **The SDK is the loop packaged; Managed Agents is the agent as a service** — the [Agent SDK](#) provides "the same tools, loop and context management that power Claude Code", programmable in Python/TS, and separates *who runs the loop* (SDK in your process) from *who renders it*; [Managed Agents](#) take it to the extreme — "Anthropic runs the agent and the sandbox, your application sends events and receives the stream". Decision: the programmatic surface is the core as a library — or as a REST endpoint.
- **The IDE is a thin surface over the same engine** — [VS Code](#) and [JetBrains](#) add inline diffs and editor context by reusing the CLI engine (same CLAUDE.md, same permission modes); the broader pattern — [Copilot agent mode](#), [Cursor 2.0](#) (background agents), [Windsurf Cascade](#) — splits the editor surface into **inline (synchronous)** and **background (asynchronous, cloud)** over the same task abstraction.
- **The interaction UX: approval as a state machine, streaming as events, the human as a tool** — the [permission modes](#) turn approval into a state machine (default/acceptEdits/plan/...), not an ad-hoc prompt; [streaming](#) exposes the loop as a typed event stream (`text_delta`, `tool_use`, `result`); and [AskUserQuestion](#) models human-in-the-loop **as a tool** the agent calls — "asking the human" becomes a step of the loop, not a special interruption.
- **The ambient/inbox shift** — for asynchronous agents, the surface stops being the chat prompt and becomes an **inbox**: [LangChain's ambient agents](#) (always-on, event-triggered, surfacing to the human only via notify/question/review) and [Claude Code on the web](#) (runs in a managed cloud and "keeps going after you disconnect") point to the same future: supervising *many* long-running agents without a live terminal — exactly the "oversight without constant oversight" that levels-of-automation HCI predicts.
- **Chat as a service, with its own identity** — [channels](#) let Telegram/Discord "or your own server" push events into a session; [Slack](#) turns `@Claude` into a cloud session that transforms a bug into a PR, **with its own credentials and audit trail, decoupled from any human's access**. Decision: chat is just another trigger, and the agent-as-service has its own identity (ties into ch. 07).
- **See also**: the living collection [Awesome Harness Engineering — Human-in-the-Loop](#) gathers more consultable resources for this dimension (patterns, articles and implementations), curated by problem.

The state of the art

1. Core with front-ends — the early boundary decides everything

The structural lesson of round 1 has become consensus: **the earlier the core/interface boundary is drawn, the more interfaces fit later**. Codex is the crystal-clear example — "a single Rust engine serves the TUI, headless `codex exec`, the IDE extension, the desktop app, cloud/web, an MCP server and remote control". opencode pays for the same bet with a typed HTTP API and generated clients. The anti-pattern is the inverse: a front-end with an agent stuffed inside, which does not scale to a second surface without a rewrite.

2. Headless with structured output is mandatory

There is no serious harness without the Unix-filter mode: `codex exec` (JSONL), `gemini --output-format stream-json` (NDJSON of events), `oh -p/ohmo --print`, Aider headless. Structured output is what makes the agent **programmable** — a pipeline piece, a CI target, the backend of another UI. It is the surface that, once absent, closes off every other automation.

3. Three visions — and the explosion of the "colleague in chat" with voice

The three bets of round 1 persist, now sharper: the agent as a multi-platform **product** (opencode Electron/VS Code; Codex desktop+cloud), as a platform **service** (gemini-cli SDK/A2A/Action; Managed Agents REST) and as a **colleague** in chat. The personal-agent category has *exploded* the third one: **OpenClaw** serves ~23 chat channels + native apps (iOS/Android/macOS/Windows) + **voice** (Voice Wake, continuous Talk Mode) + Live Canvas; **Hermes** has a single-process multi-channel gateway (10 platforms + voice); **ohmo** makes Telegram/Slack/Discord/Feishu the primary surface. Voice and channel breadth have become first-class surfaces.

4. The surface is a security boundary, not a backdoor

The most mature lesson of round 2, and the one that over-reliance HCI reinforces: a surface cannot be a shortcut around the core.

IronClaw makes this concrete — CLI, WebUI, Slack, Telegram and webhooks all enter through the **same turn contract** (**ProductAdapter**), and the WebUI is *forbidden* from bypassing the auth boundaries. The agent-in-Slack runs with its own credentials and audit trail, decoupled from the human's. And the UX must not *hide* what the agent did — the antidote to the false sense of oversight that Bućinca and Passi & Vorvoreanu document.

5. The next frontier: ambient, cloud, asynchronous

The emerging paradigm changes the very nature of the interface. **Codex cloud-tasks** (a TUI for remote tasks), Claude Code on the web continuing after you disconnect, and the ambient agents' inbox all point to the same place: the human stops *driving* one live agent and starts *supervising many* through notification and review. It is HCI's autonomy dial taken to product — high automation in execution, the human at the decision gate, asynchronous. The agent's interface is leaving the terminal (Aider's watch mode turns **ai!** comments in any editor into commands; OpenClaw's Live Canvas) and becoming an environment.

EXECUTIVE SUMMARY

What is most modern: one engine, many surfaces (doctrine); structured headless as mandatory; the channel + voice explosion; the surface as a security boundary (same turn contract); and the ambient/inbox/cloud shift. **What to steal:** draw the core/interface boundary early (typed API or library), not late; ship headless with `stream-json` from day one; make every surface pass through the same turn contract (never an auth backdoor); model human-in-the-loop as a tool and approval as a deliberate act; and get ready for the inbox — the next terminal is asynchronous.

Hands-on — harness-zero: the chat as observation window

The harness-zero interface was born in **step 0**: a minimal chat on FastAPI, the *observation window* that accompanies every step of the book. The lesson of this dimension is what the entire project demonstrates: because the loop lives behind ports (**LLMPort**, **ToolPort**, **StorePort**), adding a **second surface** — a headless `--print` mode that emits the same events in `stream-json` — is a **thin adapter**, not a rewrite. Completeness exercise: you add the headless mode and prove that the same agent responds in the chat and in the pipe, and that approval (the permission gate from ch. 07) shows up on both surfaces through the same contract — the surface is not a backdoor.

Check your understanding

1. Why does "core with front-ends" allow more interfaces than "front-end with an agent inside", and what decides this in practice? (The boundary drawn early — typed API/library — lets each surface be a thin adapter; the inverse demands a rewrite per surface.)
2. Your agent will run asynchronously in the cloud, supervised by several humans. Which interface paradigm and which HCI principle guide the design? (Ambient/inbox — notify/question/review; levels of automation: high in execution, human at the decision gate; over-reliance → do not hide what the agent did.)
3. You expose the agent on Slack and in a WebUI. What rule prevents the new surface from becoming a security hole? (Same turn contract/ProductAdapter — every surface passes through the same auth boundaries; the UI does not bypass them; its own identity/audit trail.)

Appendix A — How each repository handles interfaces

Evidence per harness, with paths — online supplement, expanded each round.

opencode (round 1) — the largest product surface

Client-server architecture (ch. 02): an HTTP server with a typed API and generated clients enables **seven surfaces** — TUI (SolidJS/opentui), **Electron desktop app** (unique in round 1), VS Code extension, **GitHub Action** (packages/github/), **Slack** (packages/slack/), web (packages/web/) and **ACP** (Zed integration). Link-shareable sessions connect the surfaces.

gemini-cli (round 1) — rich terminal + platform

React/Ink TUI with ~40 slash commands via pluggable loaders. **First-class headless:** `gemini -p` with `--output-format stream-json` (real-time NDJSON). **VS Code companion** (an IDE server exposing files/diffs). Official GitHub Action. **ACP** for editors and an **A2A server**. Its own SDK (packages/sdk).

OpenHarness/ohmo (round 1) — the agent that lives in chat

Typer CLI (oh) headless (`-p, text|json|stream-json`) + `--dry-run`; two TUIs (React/Ink + Textual); autopilot web dashboard. And **ohmo**: a personal agent on **Telegram/Slack/Discord/Feishu** (channels/ + gateway/) with its own workspace.

OpenClaw (round 2) ★ — the widest surface breadth

~23 **channels** (WhatsApp, Telegram, Slack, Discord, Signal, iMessage, Teams, Matrix, Feishu, LINE, WeChat, QQ...), web Control UI, WebChat, CLI, TUI, **voice** (Voice Wake + continuous Talk Mode), **native apps** (iOS/Android/macOS/Windows) and **Live Canvas** (A2UI). The agent's surface as a consumer product.

Codex CLI (round 2) ★ — one engine, every surface

A single Rust engine serves: TUI (ratatui), `codex exec` headless (human + JSONL), IDE extension via App Server, **desktop app**, **cloud/web** (cloud-tasks with a TUI for remote tasks), Codex as an MCP server, and remote control. The canonical example of a core with front-ends.

IronClaw (round 2) ★ — same turn contract

CLI/REPL, WebUI (SSE+WS with OIDC, rate limiting, origin check), Slack, Telegram, webhooks — all entering through the **same turn contracts** (ProductAdapter); the WebUI is **forbidden from bypassing** the auth boundaries. The surface as a security boundary, not a backdoor.

Hermes (round 2) — single-process multi-channel gateway

Full TUI; **multi-channel gateway**: Telegram, Discord, Slack, WhatsApp, Signal, Email, iMessage, QQ, WeChat, Yuanbao — with cross-platform continuity; **voice** (multi-provider transcription + TTS); ACP for editors; an OpenAI-compatible API server.

Goose (round 2) — ACP desktop over an embedded core

Full CLI + TUI; **Electron desktop speaking ACP** to the core (embedded binary, no separate server); headless via recipes + scheduler; Telegram gateway and Discord bot; pure MCP/ACP server mode.

Aider (round 2) — input outside the terminal

Rich CLI/REPL (prompt_toolkit, streaming markdown), browser UI (Streamlit), **watch mode** (aider/watch.py: ai!/ai? comments in code from any IDE become commands), **voice-to-code**, images/URLs in chat. The interface escaping into third-party editors.

OpenHands (round 2) — SaaS control plane

React Web UI (~40 routes: conversations, settings, admin, billing, orgs); **agent-canvas** CLI; headless/REST via Agent Server; **GitHub/GitLab/Jira/Slack resolvers** (webhooks); full enterprise/SaaS (Keycloak, Stripe, multi-tenant); Docker/k8s deploy.

n8n (round 2) — embeddable chat + canvas

Chat Trigger (hosted chat app + embeddable @n8n/chat widget + streaming), Manual Chat Trigger, arbitrary webhooks, the visual editor (canvas) as the building interface, MCP Server Trigger. The "inverted harness" whose primary interface is the graph.

Frameworks (frameworks round)

The frameworks deliver the loop as a library (the pure programmatic surface) + event streaming + human-in-the-loop as composition (OpenAI Agents SDK, LangGraph, CrewAI); the UI is left to the integrator. It is the "core only, the surface is yours" extreme of the spectrum — the opposite of OpenClaw.

14 — Convergences & Trends

🕒 state of the art 2026-07 · revised 2026-07-28 · 📖 ~7 min read

Learning objectives

By the end of this chapter, you should be able to:

1. **Enumerate** the eight architectural convergences of the first round and **explain** why independent convergence signals a consolidated discipline;
2. **Distinguish** the consolidated dimensions from the dimensions in genuine divergence, and **justify** why containment is the most consequential divergence;
3. **Apply** the expiration clause to any harness component — identifying why it exists and under what condition it expires;
4. **Evaluate** a new harness against the convergence checklist, demanding justification for each absence;
5. **Anticipate** the trends to watch in the coming rounds and what each would imply for harness design.

The problem

The previous chapters analyzed the harness dimension by dimension — context, compaction, tools, permissions, loop. What is missing is the question that gives the whole its meaning: **what is an accident of implementation and what is the anatomy of the discipline?** Without this synthesis, each chapter is a catalog of choices; with it, the reader gains a design criterion — knowing what to copy without hesitation, where a different bet is still viable, and what will disappear as models improve.

The measuring instrument is independent convergence. When teams that do not coordinate, on different stacks and from different cultures, arrive at the same architecture, that is strong evidence that the problem — not fashion — determined the solution. And the projection instrument is chapter 01's expiration clause: every harness component is a prosthesis for a current model limitation, and therefore every component should declare when it expects to become unnecessary.

The state of the art

The central finding of the first round: eight convergences

Three harnesses, three stacks (Effect-TS, TypeScript, Python), three origins (independent startup, big tech, academia/teaching gateway) — and a remarkable architectural convergence. Without coordination, all three arrived at:

1. **Hierarchical context file at the project root** — `AGENTS.md` / `GEMINI.md` / `CLAUDE.md`: the same artifact under three names (ch. 03).
2. **Staircase compaction** — truncate tools → prune → summarize via LLM, with automatic threshold-based triggering (ch. 04).
3. **Tool schemas derived from types** — Effect Schema, declarative classes, Pydantic: nobody writes JSON Schema by hand (ch. 05).
4. **MCP as the standard integration** — three full clients on the official SDKs (ch. 06).
5. **Plan mode as a permission mode** — read-only enforced by the permission system, not requested from the model (ch. 09).
6. **Lifecycle hooks** — before/after tool, compaction, session (ch. 12).
7. **Headless with structured output** — `-p` + JSON/NDJSON for scripting and CI (ch. 13).
8. **Stopping on absence of tool-call + turn limit** — the universal mechanics of the loop (ch. 02).

When independent implementations converge like this, the anatomy is consolidated: **this is the discipline**, no longer a set of idiosyncratic choices. A new harness that does not implement the eight items above must justify each absence.

Where genuine divergence remains

The dimensions without consensus are the map of the open bets:

- **Containment** (ch. 07): policy + mandatory OS sandbox (gemini-cli), policy + fixed sensitive paths (OpenHarness), or policy only (opencode)? The most consequential divergence — it is the one that defines operational risk.
- **Multi-agent** (ch. 10): a one-off tool, a service with a registry, or a persistent team with a mailbox? Three incompatible philosophies; the winner depends on how good models get at coordination.
- **Who decides to continue** (ch. 02): a structural heuristic or one extra inference per turn (next-speaker check)?
- **Model neutrality** (ch. 12): ~26 providers (opencode) versus the showcase of one ecosystem (gemini-cli). A commercial bet, not a technical one — but it defines who survives the commoditization of models.
- **Behavioral evals** (ch. 11): in round 1, only one of the three treated agent behavior as a regression surface — round 2 confirmed the prediction and the gap closed (see ch. 11). Easy prediction: in two years, this will be as mandatory as CI.

The expiration clause, applied

Returning to chapter 01’s thesis — every harness component is a prosthesis for a current model limitation. The exercise every harness should do, applied to what we studied:

Component	Exists because...	Expires when...
Compaction	windows are finite and expensive	long context becomes cheap and reliable
Plan mode	models act rashly	models plan spontaneously under risk
Next-speaker check	the model does not signal end-of-turn well	model-native turn protocols
Policy engine / approvals	models are not trustworthy with destructive actions	calibrated, verifiable reliability
Prompt per model family	models respond differently to instructions	instruction-following convergence
Subagent for exploration	file dumps pollute the context	abundant context + robust attention
Repo-map / code indexes	the model does not "carry" the whole repo	usable multi-million-token context

What does **not** expire: sandbox (containment is about the world, not about model capability), interfaces, verification of the work (tests/LSP — truth external to the model), and the interoperability protocols (MCP, A2A, skill formats). Long-term harness engineering lives there: **at the boundary between the agent and the world, not in the crutch for the model's limitation.**

Trends to watch in the coming rounds

1. **Standardization of the context file** — the pressure for a vendor-neutral `AGENTS.md`.
2. **Portable skills/plugins** — OpenHarness already loads skills in the Claude Code format; an "MCP of extensibility" is taking shape.
3. **Agent-as-a-service** — A2A server, agent cards, SDKs: harnesses exposing themselves to one another.
4. **Security as a first-class dimension** — shell parsing, trusted folders, injection evals: today the exception, tomorrow the baseline (hypothesis confirmed in round 2 with Codex CLI).
5. **Reversibility** — git checkpoints with `/rewind`: when undo is cheap, the policy can be looser; expect more harnesses to copy it.
6. **The minimal harness** — against the grain of sophistication, projects like mini-swe-agent (~100 lines) test how much of the *scaffolding* the modern model can already do without. It is the expiration clause turned into an experiment.

EXECUTIVE SUMMARY

- Eight dimensions have already converged across independent implementations — they are the minimum checklist of a serious harness; absences demand justification.
- The genuine divergences (containment, multi-agent, next-speaker, model neutrality, behavioral evals) are the map of the open bets — containment is the one with the greatest operational consequence.
- The expiration clause separates temporary prostheses (compaction, plan mode, repo-map...) from what is permanent: sandbox, interfaces, external verification and interoperability protocols.
- The long-term value of harness engineering lives at the agent–world boundary; the rest changes hands or disappears as models improve.
- This chapter is the book’s living scoreboard: each benchmark round confirms convergences, resolves divergences, or retires expired components.

See also: the living collection [Awesome Harness Engineering — Foundations](#) gathers more consultable resources for this dimension, curated by problem.

Check your understanding

1. Why is **independent** convergence (three stacks, three origins) stronger evidence of consolidation than the adoption of a pattern by several projects that copy each other? (Re-read "The problem" and the central finding.)
2. A new harness implements neither plan mode nor a context file at the root. According to this chapter, what is the correct posture when evaluating it — and what would you demand from its author?
3. Apply the expiration clause to a component that is **not** in the table (for example, the next-speaker check is already there; pick lifecycle hooks or headless): does it exist because of a model limitation or because of a need at the agent–world boundary? Does it expire?
4. Among the five divergences listed, which one defines operational risk and which one is a commercial rather than technical bet? Justify with the text.

15 — The Embedded Harness

🕒 state of the art 2026-07 · revised 2026-07-28 · 📖 ~8 min read

Learning objectives

By the end of this chapter, you should be able to:

1. **Explain** the inversion that defines the category — the workflow contains the harness, not the other way around — and why it raises the question "which dimensions of the scaffolding are essential, and which are replaceable by the environment?";
2. **Identify** which dimensions of the scaffolding the workflow environment does away with (compaction, planning, context delivery, granular permissions) and **justify** why each one becomes dispensable;
3. **Analyze** the implementation of a real agent node (n8n's AI Agent, Appendix A as the answer key) and locate where the loop, the tools and the permissions live;
4. **Evaluate** when to use an embedded harness versus a dedicated one, as a function of task duration and autonomy — and recognize the ceiling of the substitution;
5. **Apply** the category's exportable ideas to a dedicated harness: deriving tools from existing surfaces (the `$fromAI` pattern) and durable human-in-the-loop.

The problem

In the previous chapters, the harness contains the work: the loop drives, the tools act, the workflow emerges from the model's decisions. Tools like **n8n** invert the relationship — **the workflow contains the harness**. An "agent node" is a step inside a graph designed by a human, surrounded by triggers (webhook, cron, chat), integrations and error handling that the workflow engine already provided before AI existed.

This inversion raises the question that gives the category its meaning: **which dimensions of the scaffolding are essential, and which are replaceable by the environment?**

The state of the art

What the environment does away with

The evaluation of the category's representative (n8n, see Appendix A) confirms the thesis with uncomfortable precision: the weak dimensions of the embedded harness are exactly the ones the environment does away with.

Dimension dispensed	Why the environment dispenses it
Compaction	Event-triggered executions are short — context does not accumulate
Planning	The plan <i>is</i> the workflow graph, drawn by the human on the canvas
Context delivery	Context arrives mapped from previous steps via expressions
Granular permissions	The topology is already the allowlist

The last point deserves emphasis: in the embedded harness, **permission is topology**. There is no per-call approval inside the loop — the LLM (Large Language Model) can only invoke what the author plugged into the canvas. It is an allowlist by construction, decided visually by a human, complemented by real human-in-the-loop: nodes that pause execution durably, awaiting approval on a channel (Slack/Outlook), instead of the CLIs' synchronous approval prompt.

And the strong dimensions are where the engine has a structural advantage: **tools** (the pre-existing integrations become a tool pool), **memory** (pluggable database backends), **interfaces** (hosted chat, webhooks, embeddable widget), **MCP (Model Context Protocol)** (client

and server) and **subagents** (agent-as-tool and sub-workflows). No dedicated harness has a tool pool the size of a converted integration ecosystem — because none has a pre-existing ecosystem to convert.

The borrowed loop — and the re-internalization trajectory

The embedded harness typically does not write its own loop: it borrows it from a framework (in the case observed, LangChain JS). But the trajectory measured in the benchmark points in a clear direction: the workflow engine starts by outsourcing the loop and **re-internalizes the half that matters to a workflow engine — the scheduling of execution**. The framework still decides *which* tool to call; the *execution* of the call becomes the engine's responsibility again, as it schedules the nodes and re-enters the agent. (Code detail in Appendix A, finding 1.) The implication: workflow engines tend to absorb ever more of the harness, not the other way around.

The ceiling of the substitution

But the substitution has a ceiling: **without compaction or planning, the agent node serves short automations, not long autonomous work**. An embedded agent that had to refactor a repository for hours would collapse the context window with no defense. The two layers do not compete — they complement each other by task duration and autonomy: the dedicated harness for long, open-ended work; the embedded one for pinpoint decisions inside structured processes.

Implications

1. **For those building a dedicated harness:** the `$fromAI` pattern (Appendix A, finding 2) shows how to derive tools from existing surfaces without writing wrappers; durable HITL (pausing execution for days awaiting approval on a channel) is superior to the CLIs' synchronous approval prompt.
2. **For those building on top of workflow engines:** the category's gaps (compaction, plan mode) are the obvious roadmap — and the loop's re-internalization trajectory suggests the engines will absorb ever more of the harness, not the other way around.
3. **For the book's taxonomy:** "how much harness is needed" is a function of the *execution environment*, not a universal constant. The benchmark's ruler measures scaffolding that is present; this category reminds us that scaffolding absent-by-design is not a gap — as long as the task class is respected.

EXECUTIVE SUMMARY

The embedded harness is not an incomplete dedicated harness: it is a category in which the execution environment replaces, by construction, half of the scaffolding's dimensions — the plan becomes a graph, permission becomes topology, context becomes mapped expressions. The substitution holds as long as the task class is respected: pinpoint decisions inside structured processes, not long autonomous work. **What to steal** today: automatic derivation of tools from existing integrations (the `$fromAI` pattern) and durable human-in-the-loop instead of synchronous approval.

See also: the living collection [Awesome Harness Engineering — Production Infrastructure & Operations](#) gathers more consultable resources for this dimension, curated by problem.

Check your understanding

1. State the inversion that defines the category and explain why it turns the benchmark's "weak dimensions" into "dimensions dispensed by the environment". (If needed, re-read "What the environment does away with".)
2. Why does permission-as-topology do away with per-call approval inside the loop — and which mechanism complements this allowlist when a human decision is genuinely needed mid-execution?
3. A team wants to use a workflow engine's agent node to refactor a repository for hours. Explain, in terms of compaction and planning, why this collapses — and what the correct division between embedded and dedicated harness would be for that task.
4. Name the two ideas from the category worth exporting to a dedicated harness and what each one replaces or improves. (Hint: tool derivation and HITL.)

Appendix A — n8n (AI Agent node)

Per-repository evidence, with paths — supplementary material (online version), expanded each benchmark round. The full n8n evaluation (29/36) is in `../benchmark/avaliacoes/n8n.md`.

Anatomy of the agent node (evidence: `packages/@n8n/nodes-langchain`)

n8n implements the agent as a "cluster node": a root **AI Agent** node with typed ports into which sub-nodes are plugged — model (`AILanguageModel`), memory (`AIMemory`), tools (`AITool`), output parser. Three code findings structure the chapter:

1. The loop is borrowed — and is being handed back. The V2 generation delegates everything to LangChain JS (`AgentExecutor.fromAgentAndTools`, `maxIterations` 10). But V3 changed the design: LangChain still *decides* which tool to call (`createToolCallingAgent`), yet the *execution* became the responsibility of the n8n engine — tool calls become `EngineRequest` objects returned to the engine, which schedules the nodes and re-enters the agent with `EngineResponse`. n8n started by outsourcing the loop and is **re-internalizing** the half that matters to a workflow engine: the scheduling of execution.

2. The `$fromAI` bridge — the category's most exportable idea. `create-node-as-tool.ts` turns **any of the 400+ integration nodes** marked `usableAsTool` into an agent tool: the parameter traversal collects `$fromAI('key', 'description', type)` expressions — the slots the LLM must fill — and generates the Zod schema automatically. No dedicated harness has a tool pool this size, because none has a pre-existing integration ecosystem to convert.

3. Permission is topology. There is no per-call approval inside the loop: the LLM can only invoke what the author plugged into the canvas's `AITool` port. It is an allowlist by construction, decided visually by a human — complemented by real human-in-the-loop (`sendAndWait` nodes pause execution durably awaiting approval on Slack/Outlook, forbidden inside subagents) and a Guardrails node.

The score (29/36) and the strength/weakness map

The evaluation's weak dimensions are the ones the environment dispenses: **compaction (1)** — event-triggered executions are short, context does not accumulate; **planning (1)** — the plan is the graph drawn on the canvas; **context delivery (2)** — context arrives mapped from previous steps via expressions; **granular permissions (2)** — the topology is already the allowlist.

And the strong ones are where the engine has a structural advantage: **tools (3)** — the integrations; **memory (3)** — pluggable database backends; **interfaces (3)** — hosted chat, webhooks, embeddable widget; **MCP (3)** — client **and** server: the `McpTrigger` exposes n8n's tools to external MCP clients; **subagents (3)** — agent-as-tool and sub-workflows.

Cousins to evaluate in future rounds: Zapier Agents, Make, Dify, Flowise.

16 — Learning & Self-Improvement

🕒 state of the art 2026-07 · revised 2026-07-28 · 📖 ~8 min read

Learning objectives

By the end of this chapter, you should be able to:

1. **Explain** why self-improving learning breaks the assumption of *static* scaffolding — and how it inverts the book's expiration clause;
2. **Describe** the stages of the closed skill-capture cycle (trigger, curation, isolation, portable format, indexed re-encounter, maintenance against entropy);
3. **Compare** the two competing designs for applying what was learned — autonomous × human promotion — and **locate** a real harness on the dimension's maturity ladder;
4. **Evaluate** the dimension's risks (superstition, entropy, contamination, prompt injection as permanent learning) and the engineering that prevents them.

The problem

The twelve dimensions of chapters 02–13 describe *static* scaffolding: someone — the harness author, the user, a plugin — writes the instructions, tools and policies, and the agent consumes them. This chapter documents the emerging dimension that breaks that assumption: the agent that **writes its own scaffolding** — capturing learned procedures as reusable skills.

The dimension was promoted to supplementary status in the benchmark template (dimension 13) on the strength of one piece of evidence: the **Hermes Agent** (Nous Research) implements the full cycle, and reading the code confirms each stage (Appendix A).

The state of the art

The closed cycle: the six stages

The reference mechanism, verified in the Hermes code (detailed evidence in Appendix A), closes the cycle in six stages:

1. **Autonomous trigger** — the learning review fires on its own, in the background, without the user asking (with a manual trigger as a complement).
2. **Curation by an isolated fork** — a clone of the agent, with a curatorial prompt that defines what to capture and — most importantly — **anti-patterns of what NOT to learn**. Without that list, the system would degenerate into accumulated superstition.
3. **Isolation of the meta-work** — the curator fork has restricted tools and persistence turned off, so as not to contaminate the real session.
4. **Writing in a portable format** — the skill becomes a `SKILL.md` under strict standards, with the context constraint shaping the format of the knowledge.
5. **Cheap re-encounter** — a compact index always in the system prompt; the full content only enters the context on demand. Learning indexed, not dumped.
6. **Maintenance against entropy** — a periodic curator consolidates, archives by inactivity, and protects what is pinned. Memory that only grows becomes noise; the curator is the knowledge's garbage collector.

The maturity ladder in the evaluated cohort

Harness	Score 13	What it has
Hermes	3	The complete closed cycle (Appendix A), with autonomous application
gemini-cli	3 (retro)	Auto Memory: an extractor agent with anti-noise gates ("Default to NO SKILL", 5 blocking questions) producing SKILL.md + memory patches — but with human promotion via inbox (/memory inbox); dedupe, write sandbox, dedicated evals
IronClaw	2	Automatic skill extraction (<code>learning.rs</code>) with usage/confidence metrics and versioning
OpenClaw	1	Dreaming (autonomous memory consolidation); Skill Workshop with a proposal queue
OpenHarness	1 (retro)	Auto-extraction of facts per turn, with usage-based staleness (60 days) — facts, not procedures
Codex CLI	1	Automatic memories with pruning (facts, not procedures)
Goose	1	chatrecall (semantic recall of past conversations)
opcode, the rest	0 (retro)	Skills are consumption/distribution; nothing is written from experience

The ladder is sharp: **memory of facts** (level 1) → **extraction of procedures** (level 2) → **curated cycle with anti-patterns and maintenance** (level 3). What sets level 3 apart is not capturing more — it is the engineering of *not* capturing wrongly and of pruning what has aged.

The two competing designs at level 3

Level 3 already has **two competing designs**, with the divergence exactly where it matters: *who applies what was learned*. Hermes applies it autonomously (with the curator cleaning up afterwards); gemini-cli requires human promotion (inbox — nothing enters the context without /memory inbox). It is the classic autonomy × control trade-off of chapter 07, reappearing in the newest dimension: Hermes bets that anti-patterns are enough to prevent bad learning; gemini-cli bets they are not. The coming rounds will tell which scales better.

Why this changes the book's thesis

The expiration clause (chs. 01, 14) says: every harness component is a prosthesis for a current model limitation, and expires when the model improves. Self-improving learning **inverts the clause**: instead of waiting for the model to render the scaffolding unnecessary, the model+harness pair *writes new scaffolding for itself*. Each learned skill is a piece of harness generated at runtime, specific to the user and the environment — something no harness author could have written at the factory.

This creates a third path in the taxonomy:

1. **Factory scaffolding** — written by the harness author; expires as models evolve.
2. **Boundary scaffolding** — sandbox, permissions, interfaces; does not expire (it is about the world).
3. **Self-generated scaffolding** — skills written by the agent; it *grows* with use, and its quality depends on the curation engineering, not on the model's raw capability.

The risks: the mirror of the promises

The risks are the mirror of the promises: without anti-patterns, superstition; without curation, entropy; without isolation of the meta-work, contamination; and — pointed out by the IronClaw evaluation (prompt-write safety; cf. ch. 07) — without a protected write boundary, **prompt injection becomes permanent learning**: an attacker who convinces the agent to "learn" a malicious skill persists in procedural memory. A mature dimension 13 will require a mature dimension 6.

EXECUTIVE SUMMARY

The dimension is the newest in the template and the least converged: two harnesses at level 3 with opposite designs on who applies the learning, and the rest of the cohort somewhere between memory of facts and nothing. What is already engineering consensus among those who got there: the central piece is not the capture mechanism, but the **anti-patterns of what not to learn** and the **maintenance** (consolidate, archive, never delete). **What to steal** today: an anti-pattern list in the curatorial prompt; isolation of the

meta-work in a fork without persistence; a compact index with content on demand; a periodic curator as garbage collector; a write boundary protected against prompt injection.

Retroactive re-evaluation of the code cohort pending; the dimension leaves "supplementary" status when ≥ 3 harnesses reach level 2+.

See also: the living collection [Awesome Harness Engineering — Skills & MCP](#) gathers more consultable resources for this dimension, curated by problem.

Check your understanding

1. Why is the **anti-pattern** list ("what NOT to learn") described as the central piece of the curatorial engineering, rather than the capture mechanism itself? What happens to a system that captures without it?
2. Locate on the maturity ladder a harness that extracts facts automatically with usage-based staleness but does not capture procedures. What score does it receive, and what would it need to climb one level?
3. Hermes and gemini-cli are both at level 3, but diverge on *who applies* what was learned. Reconstruct the autonomy $\times$ control trade-off in this context: what is each design's bet?
4. Explain the sentence "a mature dimension 13 will require a mature dimension 6": why is prompt injection qualitatively more serious in a harness that learns than in a static harness?

Appendix A — Hermes Agent

Per-repository evidence, with paths — supplementary material (online version), expanded each benchmark round. Full evaluation:
`../../benchmark/avaliacoes/hermes-agent.md`.

Hermes's closed cycle (evidence: `agent/background_review.py` and related)

The mechanism, verified in the code of the evaluated fork:

1. **Autonomous trigger.** Every ~10 tool-calling iterations (`skill_nudge_interval`, in `agent/turn_finalizer.py`), the harness fires a background review — without the user asking. There is also the manual `/learn` trigger.
2. **Curation by an isolated fork.** A clone of the agent runs in a separate thread with a snapshot of the conversation and a curatorial prompt (`_SKILL_REVIEW_PROMPT`) that is the central piece of the engineering. It instructs the curator to be active ("a pass that does nothing is lost learning"), defines an order of preference (update an existing skill > create a new one; new skills only class-level, never "fix-bug-1234") and — most importantly — lists **anti-patterns of what NOT to learn**: environment-dependent failures, negative claims about tools ("the browser doesn't work"), transient errors, one-off narratives. Without that list, the system would degenerate into accumulated superstition.
3. **Isolation of the meta-work.** The fork has a restricted tool whitelist (`memory + skills`), memory and persistence turned off — so the curation does not contaminate the real session — and inherits the parent's cached prompt prefix (~26% reduction in the cost of the review).
4. **Writing in a portable format.** The skill becomes a `SKILL.md` compatible with `agentskills.io` in `~/.hermes/skills/<categoria>/<nome>/` (with `references/`, `templates/`, `scripts/`), under strict standards — description ≤ 60 characters *because the index in the system prompt truncates at 60*: the context constraint shaping the format of the knowledge.
5. **Cheap re-encounter.** The compact index (name + description) is always in the system prompt; the full content only enters the context when the agent calls `skill_view` — learning indexed, not dumped.

6. Maintenance against entropy. A periodic **curator** (`agent/curator.py`) runs when the agent is idle: it consolidates skills into umbrellas, archives by inactivity (90 days — archive, never delete), and protects pinned skills. Memory that only grows becomes noise; the curator is the knowledge's garbage collector.

17 — The Protocol Layer

🕒 state of the art 2026-07 · revised 2026-07-31 · 📖 ~9 min read

Learning objectives

By the end of this chapter, you should be able to:

1. **Explain** why the protocol layer is what turns a market of silos into an ecosystem — and why each protocol standardizes a different *boundary* of the harness;
2. **Distinguish** the boundaries covered by MCP (Model Context Protocol), A2A (Agent-to-Agent) and ACP (Agent Client Protocol), agentskills.io and AGENTS.md — including the two classic confusions (the two "ACP"s; MCP × A2A as vertical × horizontal);
3. **Analyze** the adoption matrix measured in code and locate a real harness in it;
4. **Evaluate** a protocol's health by measured adoption and governance (neutral foundation × single vendor), rather than by marketing;
5. **Decide** which protocols a new harness needs to speak so as not to be left out of everyone else's composition architectures.

The problem

Chapters 02–16 deal with what happens *inside* a harness. This chapter deals with what happens *between* them — and between harnesses and the rest of the world. Without shared protocols, each harness is a silo: its tools, its project instructions, its subagents and its skills only work inside it. The protocol layer is what turns that market of silos into an ecosystem: each protocol standardizes a different boundary of the harness — agent ↔ tool, agent ↔ agent, agent ↔ editor, agent ↔ user, plus the cross-cutting formats for procedural knowledge (SKILL.md) and project instructions (AGENTS.md).

The practical consequence: in a market that *composes* harnesses, not speaking the protocols is not missing a feature — it is being left out of everyone else's architectures.

The state of the art

The map: one protocol per boundary

The map, organized by the boundary each one solves:

Protocol	Boundary	Origin / governance	Status (2026)
MCP (Model Context Protocol)	agent ↔ tools/data	Anthropic → universal adoption (OpenAI, Google, Microsoft)	mature; ~97M downloads
A2A (Agent-to-Agent)	agent ↔ agent (delegation across organizations)	Google → Linux Foundation (v1.0 in 2026)	consolidating; absorbed IBM's ACP (Agent Communication Protocol)
ACP (Agent Client Protocol)	agent ↔ editor/client	Zed	rapid adoption among coding harnesses
agentskills.io (Agent Skills / SKILL.md)	portable procedural knowledge	Anthropic (open spec, Dec 2025)	~40 compatible products in 6 months
AGENTS.md	portable project instructions	community → Agentic AI Foundation (Linux Foundation)	60,000+ repositories; 20+ tools read it natively
AG-UI	agent ↔ user interface	community (CopilotKit)	emerging
ACP-IBM (Agent Communication Protocol)	agent ↔ agent	IBM	discontinued — merged into A2A (Aug 2025)

Two confusions to clear up: (1) "ACP" names two distinct protocols — IBM's (agent-agent communication, discontinued in favor of A2A) and Zed's (agent-editor, alive and expanding); in this book, ACP = Zed. (2) MCP and A2A do not compete: MCP is the *vertical* connection (agent → tool), A2A is the *horizontal* one (agent → peer agent) — a real system uses both.

The adoption matrix — measured in code, not in marketing

This chapter's differentiator: we cross the protocols with the **11 harness evaluations from the benchmark** (plus the 4 frameworks from the frameworks-1 round) (evidence per file, see [benchmark/avaliacoes/](#)). No external comparison has this column of truth:

Harness	MCP client	MCP server	ACP	A2A	SKILL.md / agentskills	AGENTS.md (or equiv.)
opencode	✓	—	✓ (Zed)	—	partial	✓ AGENTS.md
gemini-cli	✓	—	✓	✓ client+server	✓	GEMINI.md
OpenHarness	✓	—	—	—	✓ (Claude format)	CLAUDE.md
Codex CLI	✓	✓	—	—	✓	✓ AGENTS.md
Goose	✓	✓ (goose mcp)	✓ (entire desktop)	—	✓	✓ AGENTS.md + .goosehints
Aider	✗	✗	—	—	—	✓ (reads it)
OpenHands	✓	✓ (FastMCP)	✓ (profiles)	—	✓ (org repos)	microagents
OpenClaw	✓	✓	✓ (orchestrates third parties)	—	✓ (52 bundled)	✓ AGENTS.md + SOUL.md
Hermes	✓	✓	✓	—	✓ (core of its learning)	✓ AGENTS.md + SOUL.md
IronClaw	✓	—	—	—	✓ (OpenClaw compat)	identity files
n8n	✓	✓ (Trigger)	—	—	—	—
frameworks:						
LangGraph	✗	✗ (paid server only)	✗	✗	✗	—
OpenAI Agents SDK (Software Development Kit)	✓	—	✗	—	partial	sandbox agents only
CrewAI	✓ (mandatory)	—	✓ client+server	—	✓	✓ auto-generated
software-agent-sdk	✓ (OAuth)	—	✗	✓ (uses harnesses as engine)	✓ (spec)	✓

Readings of the matrix:

- MCP has won, in fact:** 10 out of 11 (the exception, Aider, is a philosophical choice). And between rounds 1 and 2, the pattern migrated from "client" to "client+server" — the harness as a consumable service.
- agentskills.io is the fastest standardization we have ever measured:** a spec from December 2025, 8 of our 11 compatible by July 2026. Chapter 12's prediction ("an MCP of extensibility is taking shape") came true — and with a structural detail: skills are portable markdown, so the same skill runs on Claude Code, on Hermes and on IronClaw. Self-improving learning (ch. 16) writes in *that* format — the knowledge one agent learns is, in theory, transferable to another.
- ACP is the cohort's most important silent protocol:** 6 out of 11 speak it, and three harnesses (OpenClaw, OpenHands, Goose) use it to **orchestrate other harnesses** as subagents — Claude Code, Codex, Gemini CLI and opencode become interchangeable parts. What used to be "agent ↔ editor" has become, in practice, the composition bus between harnesses.
- A2A has left "one player's bet" territory** (*updated in the frameworks-1 round*): gemini-cli was the only harness to implement it, but **CrewAI** came in with native client AND server (full AgentCard, JWS, gRPC/REST) — the second measured implementer, and the first framework. Governance at the Linux Foundation and the absorption of ACP-IBM keep pointing to A2A as the candidate for the inter-organizational boundary; in product harnesses, however, that boundary still barely exists.

5. **AGENTS.md has consolidated as the neutral standard:** the AGENTS/CLAUDE/GEMINI.md fragmentation of ch. 03 is resolving itself — Codex, Goose, opencode, OpenClaw and Hermes have already converged on AGENTS.md (now under the Agentic AI Foundation), with the proprietary files becoming aliases.

The stack: how the protocols compose

A complete agentic system in 2026 uses the whole stack, one layer per boundary:

[user]

| AG-UI / chat channels / TUI (interface)

[harness A]

| ACP (composition: A drives B as a subagent)

[harness B]

| A2A (delegation to another organization's agent)

[remote agent]

| MCP (each agent reaches its tools)

[tools/data]

cross-cutting: AGENTS.md (per-project instructions) · SKILL.md (portable procedures)

Implications for harness engineering

- 1. **Protocol is a survival dimension, not a feature dimension:** Aider, a technical reference in three dimensions, is outside the entire composition ecosystem for not speaking MCP/ACP. In a market that composes harnesses, not speaking the protocols means being left out of everyone else's architectures.
- 2. **The expiration clause does not apply here** (ch. 14): protocols are the boundary with the world — the scaffolding that *remains* when models improve. Investing in protocol is the harness investment with the longest half-life.
- 3. **For the benchmark:** the matrix above becomes a permanent section of the comparative, updated each round. Protocols do not receive a 0–3 score like harnesses — they are evaluated by **measured adoption** (the matrix) and **governance health** (neutral foundation > single vendor).

Addendum (2026-07-31): the MCP 2026-07-28 spec ([announcement](#)) reinforces this chapter's thesis from another angle: a stateless core, an extension framework and the **first formal deprecation policy** (12 months) are the typical behavior of a protocol leaving adolescence and entering its infrastructure phase — disciplined versioning matters more than features. The cohort's adoption of the new version enters the matrix next round. And the same-day confirmation on the other boundary (spec 065): the [A2A specification](#) confirms the **stable v1.0 under the Linux Foundation**, organized in three layers (data model in Protobuf/JSON Schema, abstract operations, JSON-RPC/gRPC/REST bindings), with **v1.0.1 already bringing a formal extension mechanism** — the two boundary winners reached, in the same quarter, the same stage: formal extensions instead of features in the core.

EXECUTIVE SUMMARY

The protocol layer already has one winner per boundary: MCP on the vertical (agent → tool, near-total adoption), ACP as the composition bus between harnesses, agentskills.io as the portable format for procedural knowledge and AGENTS.md as the neutral standard for project instructions — while A2A remains the consolidating bet for the inter-organizational boundary, sustained more by governance (Linux Foundation, absorption of ACP-IBM) than by measured adoption in product harnesses. The engineering decision is asymmetric: protocols are the harness component with the longest half-life, immune to the expiration clause, and the adoption matrix — not marketing — is the instrument for re-evaluating them each benchmark round.

Industry sources

- [ecosystem map 2026](#)
- [Zylos: MCP/A2A/ACP convergence](#)
- [Zuplo: where ACP ended up](#)
- [Agent Skills: format and adoption](#)

- [AGENTS.md guide 2026](#)
- [Zed ACP](#)

Adoption matrix: the benchmark's own evidence ([benchmark/avaliacoes/](#)).

- **See also:** the living collection [Awesome Harness Engineering — Skills & MCP](#) gathers more consultable resources for this dimension (patterns, articles and implementations), curated by problem.

Check your understanding

1. A colleague claims that "A2A will replace MCP". Why does the claim confuse the boundaries, and how does the stack show that a real system uses both? (Re-read "The map" and the diagram.)
2. "ACP" appears twice in the protocol table, with opposite statuses ("rapid adoption" and "discontinued"). Explain the difference between the two protocols — and which of them this book calls ACP.
3. You are designing a new harness. Based on the matrix readings and the implications, which protocols are mandatory today, which one is still a bet, and what does the Aider exception teach about the cost of speaking none?
4. Why does the expiration clause (ch. 14) not apply to the protocol layer, when it applies to almost everything else in the harness?

Benchmark

Harness Comparison

Consolidated Comparison — Rounds 1, 2, ext-1, ext-2 and ext-4

16 harnesses evaluated by systematic code reading, 12 dimensions (0–3) + 2 supplementary. Round 1: 2026-07-24 (opencode, gemini-cli, OpenHarness). Round 2: 2026-07-24 (Codex CLI, Goose, Aider, OpenHands, OpenClaw, Hermes, IronClaw, n8n). Round **ext-1**: 2026-07-31 (**Grok Build, Pi**). Round **ext-2**: 2026-08-02 (**Kimi Code, QM** — the latter inaugurating the *organizational agents* category). Round **ext-3**: 2026-08-02 (**Traycer** — evaluated and **not included**). Round **ext-4**: 2026-08-06 (**Prime Agent**). See the [methodology](#) (in Portuguese).

Category: coding harnesses

#	Dimension	opencode	gemini-cli	OpenHarness	Codex CLI	Goose	Aider	OpenHands*	Grok Build	Pi	Kimi Code	Prime Agent
1	Loop	3	3	2	3	3	2	2	3	3	3	3
2	Context	3	3	2	3	3	3	3	3	3	3	3
3	Compaction	3	3	3	3	3	2	2	3	3★	3	3
4	Tools	2	3	3	3	3	3	2	3	3	3	3
5	MCP	3	3	2	3	3	0	3	3	0	2	2
6	Permissions/sandbox	2	3	2	3★	2	2	3	3★	1	2	1
7	Memory/state	3	3	3	3	3	3	3	3	2	2	3
8	Planning	2	3	2	2	2	2	1	3	1	3	2
9	Subagents	2	3	3	3	3	2	2	3★	1	3	3★
10	Verification/evals	2	3	2	3	3	3	0*	2	3	2	2
11	Extensibility	3	3	3	3	3	3	3	3★	3★	3	3
12	Interfaces	3	3	2	3	3	3	3	3	3	3★	3
	Total	31	36	29	35	34	28	27*	35	26	32	31
13	Learning (suppl.)	—	3	—	—	—	—	—	—	2	1	3★

* OpenHands: the repo evaluated is the control-plane (Agent Canvas); the core (loop, condenser, SWE-bench evals) migrated to `software-agent-sdk` — the total underestimates the full project. The SDK joins the queue.

Reading of round ext-1 (2026-07-31):

- The two extremes of the spectrum arrived together.** Grok Build (35) ties with Codex CLI, covering everything with industrial depth — including reading its competitors' artifacts (AGENTS/CLAUDE/Cursor/.mcp.json) and porting codex's and opencode's tools. Pi (26) scores 3 on **everything it accepts** and 0–1 on everything it refuses by manifesto — the jagged profile is not immaturity, it is a thesis ("adapt pi to your workflows, not the other way around"), and every exclusion exists as a tested example extension.
- Dimension 6 keeps separating product from project** — and Grok Build raises the bar: shell authorization by **AST** (tree-sitter-bash), closing the `mv secret x && cat x` bypass, kernel-enforced fail-closed sandbox. Pi is the deliberate counterexample (1): it outsources the boundary to the OS and argues that in-process sandboxing is theater.
- Behavioral evals are still the most common gap** — Grok Build has 26k mechanism tests and zero competence tests (score 10 = 2); Pi, going the other way, is the only one in the round with an A/B bench for harness configurations (`evalHarnessTable`) and

eval artifacts in the native session format.

Reading of round ext-2 (2026-08-02):

1. **The second verticalized vendor confirms the pattern — and the divergence.** Kimi Code (32) repeats Grok Build's move (own model → own harness, open), but with the opposite bet: where xAI went for the maximal platform in Rust with a kernel-enforced sandbox, Moonshot went for **structured autonomy** — a goal mode with a state machine and budgets (turns/tokens/time), a swarm of up to 128 subagents, cron exposed to the model — on top of weak enforcement (no OS sandbox; bash authorized by string glob in the production engine, with the AST parser ready but only consumed in the experimental v2). The detail nobody else has: **harness ↔ API co-design** — the Kimi API gained the `dynamically_loaded_tools` capability to serve the harness's *progressive tool disclosure*, with documented degradation for other providers. The vendor changed the model to serve the harness.
2. **Cross-pollination inside the corpus became routine:** Kimi Code's TUI is a vendored fork of `pi-tui` (acknowledged in the README); QM brings Pi, OpenCode, Codex and Claude Code as pluggable *engines*. The corpus stopped being a list of competitors and became a supply chain.
3. **Behavioral evals remain the divider** in ext-2 as well: Kimi Code has 1,137 mechanism test files and zero competence evals; QM, going the other way, runs **multiplayer E2E against real Slack with an LLM judge** — the most complete implementation of dimension 10 outside gemini-cli.

Reading of round ext-4 (2026-08-06):

1. **The announcement and the code disagree — and the code won.** The launch claims that "*fixed tool-calling schemas and context compaction force the model to work around its own scaffolding.*" On reading, compaction was **neither removed nor weakened**: the 1,398 lines of `core/compaction/` inherited from Pi are still there, with safe cutting, split turns and reactive overflow recovery — and even **improved** (custom instructions, `tokensBefore` recomputation). What changed is **subordination**: `compact.run()/compact.status()` became agent-callable, with a handler that **schedules instead of executing** (executing would abort the very cell that asked), running even with auto-compaction off, under 12 tests. The ch. 04 Executive summary **does not fall — it gains a caveat**: compaction stopped being an involuntary harness event and became one of several mechanisms against context growth, while also becoming a distillation trigger (`autoRefine.compact: true` by default).
2. **Pi became the corpus's substrate.** Prime Agent is **Pi underneath** — the same four packages (`pi-agent-core`, `pi-ai`, `pi-coding-agent`, `pi-tui`), a LICENSE with dual copyright (Mario Zechner + Prime Intellect). It is the **fifth** Pi consumer on the supply-chain map, and the most radical: it collapsed the entire tool set into a **single ipython tool** and built the REPL-as-program-surface on top. The irony closes an argument from ch. 12: the **lowest-scoring system in the corpus** (Pi, 26) — which refuses half the dimensions by manifesto — was the base a frontier lab chose. Pi's extensibility got its strongest empirical validation.
3. **Dimension 13 has a new ceiling — and dimension 10 a new floor.** The *Continual Harness* gives the agent **CRUD over its own state** (prompt, memory, skill, subagent spec) with local/global scope, a review gate every 25 turns and snapshot rollback: it surpasses Hermes and becomes the dimension-13 reference. But the **thesis ↔ measurement asymmetry is glaring**: the repository contains **no evals at all** — Pi's `packages/evals` disappeared in the fork —, the `expectedOutcome` emitted by `/refine` is validated by nothing, and the claimed 95.5% on ARC-AGI-3 has no reproducible artifact in the code. A self-modifying harness without a measurement bench is the riskiest combination the benchmark has recorded.
4. **Regressions inherited downward.** Beyond the evals, **Pi's Project Trust was removed** (`grep for trust in src/` → nothing) and the single tool is arbitrary Python execution, with children inheriting `cwd` and permissions — hence the **1** on dimension 6, below Pi's already low floor.

Category: organizational agents (new in ext-2)

#	Dimension	QM
1	Loop	3
2	Context	3★
3	Compaction	3
4	Tools	3
5	MCP	1
6	Permissions/sandbox	3
7	Memory/state	3
8	Planning	1
9	Subagents	2
10	Verification/evals	3★
11	Extensibility	3
12	Interfaces	3
	Total (1–12)	31
13	Learning (suppl.)	2
14	Proactivity (suppl.)	3★

QM (Y Combinator) inaugurates the category: the first harness in the corpus whose unit of design is the **organization**, not one user's session — scopes (person/team/room/org), context filtered by the *entitlement* of everyone present in the conversation (**context-filter.ts**), recipient consent for autonomous deliveries, and auditing as core primitives. The agent loop is a **swappable dependency** (Pi, OpenCode, Codex or Claude Code by configuration), with the session portable across engines via a re-seedable "tape". It is the loop-commoditization thesis written in **package.json** — and the reason the category is new: on the classic dimensions it scores like a mature harness (31/36), but what defines it does not fit them.

Category: self-hosted personal agents

#	Dimension	OpenClaw	Hermes	IronClaw	ohmo ¹
1–5	Loop/Context/Compact./Tools/MCP	3,3,3,3,3	3,3,3,3,3	3,3,3,3,3	3,3,3,3,3
6	Permissions/sandbox	3	3	3★★	2
7	Memory/state	3	3	3	3
8	Planning	3	2	2	2
9	Subagents	3	3	2 ²	3
10	Verification/evals	3	3	3	3
11	Extensibility	3	3	3	3
12	Interfaces	3	3	3	3
	Total (1–12)	36	35	34	34
13	Learning (suppl.)	1	3★★	2	2
14	Proactivity (suppl.)	3	2	3	3

¹ dedicated evaluation (2026-07-24) of OpenHarness's personal app — gap concentrated in dim. 6 (the gateway's permission/sandbox config is dead code; no dial between deny-all and full_auto). ² a score-3 design, but **spawn_subagent** is disabled in production.

Category: embedded harnesses

n8n (AI Agent node)	Total 1–12: 29/36	Strong: tools 3 (\$fromAI → Zod over 400+ integrations), MCP 3 (client+server), memory 3, subagents 3, interfaces 3 · Weak by design of the environment: compaction 1, planning 1, context 2, permissions 2 (structural/topological)
---------------------	-------------------	--

Category: harness frameworks (round frameworks-1, FRAMEWORK_EVAL template)

Axis	LangGraph	OpenAI Agents SDK	CrewAI	software-agent-sdk
A1 Loop/orchestration	3	3	3	3
A2 State/durability	3★★	3	3	3
A3 Tools/schemas	2	3	3	3
A4 Multi-agent	2	3	3	3
A5 Human-in-the-loop	3	3★	3	3
A6 Streaming/events	3	3	3	3
Total A (0–18)	16	18	18	18
D1 Observability	2	2	2	2
D2 Tests/evals	3	3	3	3
D3 Ergonomics	2	3	3	3
D4 Ecosystem	3	3	3	3
Total D (0–12)	10	11	11	11

Reading of round frameworks-1:

- The primitives have become a commodity** (A is almost all 3s) — the real differentiation lives in axes B (boundaries) and C (protocols), which are descriptive: LangGraph imposes BSP and leaves context/permissions completely open; the Agents SDK imposes the Responses vocabulary; CrewAI imposes the role/task ontology; the OpenHands SDK imposes the entire event model.
- No framework has first-class open observability** (D1=2 across the board): each gravitates to its own platform (LangSmith, OpenAI, AMP, Laminar) — the "OTel for agents" space remains vacant.
- Protocols split the field**: CrewAI (mandatory MCP + **A2A client/server** + skills + auto-generated AGENTS.md) and software-agent-sdk (MCP OAuth + **ACP** + agentskills) are the polyglots; the Agents SDK speaks only MCP; **LangGraph speaks zero** — protocols are a feature of the paid server.
- The "two movements" prediction was confirmed in the code**: software-agent-sdk is the most advanced harness-turning-into-framework (everything became a pluggable ABC, and its ACPAgent orchestrates Claude Code/Gemini/Codex as engines); LangGraph makes the opposite move — **hollowing itself out** of the agent layer (create_react_agent deprecated toward the langchain package) to be just a durable runtime.
- Compaction remains the harness/framework dividing line**: only software-agent-sdk ships it ready-made (condenser with tombstones — the best measured in the entire benchmark); LangGraph/Agents SDK/CrewAI leave the context window to the user (the Agents SDK has only a compaction session; CrewAI nothing).

Executive summary of round 2

The hypotheses recorded in round 1 were confronted — 3 confirmed, 1 surprise:

- ✓ **Codex CLI = the new ceiling in containment** (35/36): Seatbelt + bubblewrap/seccomp + Landlock + Starlark execpolicy + network-proxy — three independent layers. gemini-cli is no longer the only "reference 3" on dimension 6.
- ✓ **Goose = MCP-native confirmed** (34/36): even the internal tools are real MCP servers served in-process. The technical tie between Codex/Goose/gemini-cli at the top of the coding category indicates the product frontier is converging.

3.  **Aider = the alternative path on context** (28/36): repo-map (tree-sitter + PageRank) is the reference in context delivery without an agent loop — and the benchmark's first **0** (MCP) shows the cost of the philosophy.
4.  **OpenHands = methodological surprise** (27/36*): the repo became a control-plane; the core lives in an external SDK. Lesson: the unit of evaluation must track how projects decompose.

The personal agents category debuted at an unexpectedly high level: OpenClaw (36) is the "gemini-cli of the category"; Hermes (35) brings the only closed-loop implementation of **self-evolving learning** (dimension 13 promoted to template supplementary because of it); IronClaw (34) redefines the conceptual ceiling for security — a loop structurally incapable of acting without the kernel (a trust class unforgeable by types, approvals as per-invocation leases, fail-closed WASM) — something **no coding harness evaluated has**.

The embedded harness confirmed the category's thesis: n8n's weak dimensions are exactly the ones the workflow engine dispenses with (short executions → no compaction; the plan is the drawn graph; permission is topology). And V3 revealed a move opposite to what was expected: n8n is *re-internalizing* the execution loop from LangChain into its own engine.

Champions by dimension (overall, rounds 1+2)

Dimension	Current reference	Mention
Loop	IronClaw (loop ≠ security perimeter)	opencode (durability), gemini-cli (next-speaker)
Context	Aider (repo-map) and opencode (epochs)	Codex (server-driven by model), Hermes (3 cache-aware layers)
Compaction	Codex (remote v2) and Goose (3 techniques)	IronClaw (effectiveness circuit-breaker)
Tools	Goose (MCP-uniform) and IronClaw (typed capabilities)	n8n (\$fromAI), Aider (edit formats driven by eval)
MCP	Codex and OpenClaw (full client+server)	Goose (in-process)
Permissions/sandbox	IronClaw (authority kernel)	Codex (3 OS layers), OpenClaw (pairing)
Memory	Hermes (multi-layer + FTS5)	gemini-cli (git checkpoint), OpenClaw (Dreaming)
Planning	gemini-cli and OpenClaw (goals/task flow)	— the weakest dimension across the entire industry
Subagents	OpenClaw (push-based + third-party ACP)	Codex (graph store), OpenHarness (swarm)
Verification	gemini-cli (4 suites) and IronClaw (cross-tenant isolation)	Aider (benchmark driving design), Goose (leaderboard)
Extensibility	broad tie — it became a commodity	OpenClaw (ClawHub w/ scan), Goose (JSON providers)
Interfaces	OpenClaw (23 channels + voice + apps)	Codex (1 core → CLI/IDE/desktop/cloud)
Learning (13)	Hermes (autonomous) and gemini-cli (human inbox) — two level-3 designs	IronClaw (automatic extraction)
Proactivity (14)	OpenClaw (heartbeat w/ lightweight context)	IronClaw (routines engine)

Cross-cutting findings from round 2

1. **Planning is the industry's weakest dimension:** no new harness reached 3; the overall average of dimension 8 is the lowest in the benchmark. Everyone has a todo-list; almost no one has an enforced plan → approve → execute.
2. **MCP client+server became the standard among the mature:** Codex, OpenClaw, Hermes, OpenHands, n8n and IronClaw expose themselves as servers — in round 1, none of the three did this in core. The harness as a *consumable service* consolidated within months.
3. **ACP emerged as the harness-orchestration protocol:** OpenClaw, OpenHands and Goose orchestrate/integrate other harnesses (Claude Code, Codex, Gemini CLI, opencode) via ACP — ch. 14's prediction about "agent-as-a-service" was confirmed through a different route.
4. **The expiration clause gained an inverted case:** Hermes's learning loop does not wait for the model to improve — the model+harness pair writes its own scaffolding (skills). Self-expansion instead of expiration.

5. **Security now has two distinct paradigms:** containment by the OS (Codex — the process *cannot*) and authority architecture (IronClaw — the loop *cannot reach*). They are complementary, and no harness combines both yet.

Recorded next steps

- **Retroactive re-evaluations:** dimension 13 on the round-1 harnesses (gemini-cli's `skill-extraction-agent` is a candidate for a 2); ohmo as a dedicated entry in the personal category.
- **Queue:** OpenHands/software-agent-sdk (the missing core), frameworks (LangGraph, CrewAI, Agents SDK — adapted template), Cline/Roo (IDE), mini-swe-agent (minimal harness), Crush, smolagents.
- **Methodological evolution:** from static to behavioral — running the harnesses on standardized tasks (Goose's Harbor and OpenClaw's Benchmark Pack are models to study).

Book apparatus

Glossary

Glossary

The acronyms of this book, **spelled out**, with a short explanation and the **context** in which they appear. In the body of the chapters, hovering over an acronym shows its meaning (**abbr**); here is the complete reference. The expansions were checked against the text itself (Principle I).

Agents, protocols and orchestration

- **MCP — Model Context Protocol.** Open protocol that standardizes how a harness plugs external tools, data and prompts into the model. *Appears in:* ch. 06 (MCP) and ch. 17 (Protocols).
- **ACP — Agent Client Protocol.** Protocol (originated at Zed) for the **agent** ↔ **editor/client** conversation. Not to be confused with IBM's *Agent Communication Protocol* (also "ACP"), discontinued and merged into A2A. *Appears in:* ch. 13 (Interfaces), ch. 17.
- **MRTR — Multi Round-Trip Requests.** Pattern from the MCP 2026-07-28 spec that replaces server-initiated requests (sampling/elicitation): the server responds `input_required` and the client retries with the answers. *Appears in:* ch. 06.
- **DCR — Dynamic Client Registration.** Dynamic registration of OAuth clients; deprecated in the MCP 2026-07-28 spec in favor of CIMD. *Appears in:* ch. 06.
- **CIMD — Client ID Metadata Documents.** DCR's successor in MCP authorization: the client's identity comes from a metadata document. *Appears in:* ch. 06.
- **A2A — Agent-to-Agent.** Protocol for **delegation between agents** (originated at Google, donated to the Linux Foundation). *Appears in:* ch. 10 (Subagents), ch. 17.
- **LSP — Language Server Protocol.** The standard that inspired agent protocols: it separates the "intelligence" (server) from the interface (client). *Appears in:* ch. 11, ch. 12, ch. 14.
- **RPC — Remote Procedure Call.** Calling a procedure in another process/machine as if it were local; the basis of several protocols. *Appears in:* chs. 05, 06, 10 and 12.
- **MAST — Multi-Agent System Failure Taxonomy.** Taxonomy of failure modes of multi-agent systems (from the paper "Why Do Multi-Agent LLM Systems Fail?"). *Appears in:* ch. 10 (Subagents), bibliography.
- **RAG — Retrieval-Augmented Generation.** Generation augmented by retrieval: fetch relevant excerpts and inject them into the context. *Appears in:* ch. 03 (Context), ch. 08.

Models and AI

- **AI — Artificial Intelligence** (in the Portuguese original, **IA — Inteligência Artificial**). *Appears in:* the whole book.
- **LLM — Large Language Model.** The model the harness wraps. *Appears in:* the whole book.
- **GPT — Generative Pre-trained Transformer.** A family of language models. *Appears in:* chs. 01, 05 and 09, bibliography.
- **SWE-bench / SWE-agent — Software Engineering** (benchmark / software-engineering agent). *Appears in:* ch. 11 (Evals).

Tools, interfaces and networking

- **API — Application Programming Interface.** The contract through which programs talk to each other. *Appears in:* the whole book.
- **SDK — Software Development Kit.** A kit for building on top of a platform (e.g. the Agent SDK). *Appears in:* ch. 12 (Extensibility), ch. 13.
- **CLI — Command-Line Interface.** *Appears in:* ch. 13.
- **TUI — Text (Terminal) User Interface.** Interactive text interface in the terminal. *Appears in:* ch. 13.
- **IDE — Integrated Development Environment.** (e.g. VS Code). *Appears in:* ch. 13.
- **UI — User Interface / UX — User Experience.** The user's interface and experience. *Appears in:* ch. 13.
- **HCI — Human-Computer Interaction.** (a scientific field). *Appears in:* ch. 13.
- **HTTP — HyperText Transfer Protocol.** The protocol of the web. *Appears in:* chs. 06, 10, 11 and 13.

- **SSE — Server-Sent Events.** Streaming of events from server to client (used in the chat). *Appears in:* ch. 13.
- **JSON — JavaScript Object Notation.** The data format of tool schemas. *Appears in:* ch. 05, 06.
- **OS — Operating System** (in the Portuguese original, **SO — Sistema Operacional**). *Appears in:* ch. 07 (Permissions and Sandboxing).
- **CI — Continuous Integration.** *Appears in:* ch. 11, apparatus.
- **DDD — Domain-Driven Design.** Guides the `harness-zero`. *Appears in:* the hands-on build.

Editorial, publishing and research

- **DOI — Digital Object Identifier.** The work's persistent identifier (Zenodo). *Appears in:* apparatus, cover.
- **ORCID — Open Researcher and Contributor ID.** The researcher's identifier (the author's). *Appears in:* apparatus, "About the author".
- **ISBN — International Standard Book Number.** The standard identifier for books. *Appears in:* Editorial Guide.
- **CC — Creative Commons.** A family of open licenses (the content is CC BY 4.0). *Appears in:* license, apparatus.
- **MIT.** Permissive software license (the name comes from the *Massachusetts Institute of Technology*); covers the code. *Appears in:* license.
- **ICMJE — International Committee of Medical Journal Editors** and **COPE — Committee on Publication Ethics.** Authorship/ethics guidelines followed in disclosing AI co-authorship. *Appears in:* Editorial Guide §6.
- **ICLR — International Conference on Learning Representations.** Scientific conference cited in the bibliography. *Appears in:* bibliography.

Appendix — The study

Appendix — The study: the harnesses evaluated

This appendix **shows the work**: the full list of harnesses that went through the study, with **where they came from** (upstream repository), **the exact snapshot that was read** (fork/commit/snapshot — the materialization of the method's cutoff date, ch. 01 §6) and the link to the **complete evaluation** of each one. The instrument used in every evaluation is the same: the [HARNESSEVAL.md](#) template (and [FRAMEWORK_EVAL.md](#) for frameworks), applied through systematic code reading following the [benchmark methodology](#) (in Portuguese).

How to read this table

- **Origin**: the public upstream repository.
- **Version/snapshot**: the version or snapshot that was read.
- **Fork/commit (cutoff date)**: the picture frozen in the GH`Daru/*` fork — this is what guarantees **reproducibility** (anyone can read the same commit) and materializes the method's obsolescence mitigation. The forks are synchronized by the [scripts/sync-forks.ps1](#) script.
- **Evaluation**: the complete document (metadata, per-dimension scores with code evidence, diagnosis and "what to steal").

The 16 evaluated

Harness	Category	Origin	Version/snapshot	Fork/commit read	Evaluated on	Analysis
Aider	coding harnesses	github.com/Aider-AI/aider	snapshot 2026-07	fork GHDaru/aider, commit 5dc9490	2026-07-24 (round 2)	evaluation
Codex CLI (OpenAI)	coding harnesses	github.com/openai/codex	snapshot 2026-07	fork GHDaru/codex, commit 000d254	2026-07-24 (round 2)	evaluation
gemini-cli	coding harnesses	github.com/google-gemini/gemini-cli	snapshot 2026-07 (main)	—	2026-07-24 (round 1, exploratory)	evaluation
Goose (Block / AaIF)	coding harnesses	github.com/block/goose	v1.44.0	fork GHDaru/goose, commit 0038bc7	2026-07-24 (round 2)	evaluation
opencode	coding harnesses	github.com/anomalyco/opencode	v1.18.4 (V2 in transition, documented in <code>CONTEXT.md</code>)	—	2026-07-24 (round 1, exploratory)	evaluation
OpenHands (Agent Canvas)	coding harnesses	github.com/All-Hands-AI/OpenHands	snapshot 2026-07	fork GHDaru/OpenHands, commit 6b04532	2026-07-24 (round 2)	evaluation
OpenHarness	coding harnesses	github.com/HKUDS/OpenHarness	v0.1.9	—	2026-07-24 (round 1, exploratory)	evaluation
Hermes Agent (Nous Research)	self-hosted personal agents	github.com/NousResearch/hermes-agent	snapshot 2026-07	fork GHDaru/hermes-agent, commit 55ef425	2026-07-24 (round 2)	evaluation
IronClaw (NEAR AI)	self-hosted personal agents	github.com/nearai/ironclaw	snapshot 2026-07	fork GHDaru/ironclaw, commit 073ded0	2026-07-24 (round 2)	evaluation
ohmo (OpenHarness)	self-hosted personal agents	github.com/HKUDS/OpenHarness (ohmo/ directory)	v0.1.9 — dedicated evaluation, complementary to OpenHarness's (round 1)	—	2026-07	evaluation
OpenClaw	self-hosted personal agents	github.com/openclaw/openclaw	snapshot 2026-07	fork GHDaru/openclaw, commit 1e15b18b	2026-07-24 (round 2)	evaluation
n8n (AI Agent node)	embedded harnesses	github.com/n8n-io/n8n	snapshot 2026-07; package evaluated: <code>packages/@n8n/nodes-langchain v2.32.0</code> (135 AI nodes)	fork GHDaru/n8n, commit 55e92cc2	2026-07-24 (round 2)	evaluation
CrewAI	frameworks	github.com/crewAIInc/crewAI	v1.15.6 — monorepo with 6 packages (<code>crewai</code> , <code>crewai-core</code> , <code>crewai-tools</code> ~79 tools, <code>cli</code> , <code>crewai-files</code> , <code>devtools</code>)	fork GHDaru, commit b3aaaab	2026-07	evaluation
LangGraph	frameworks	github.com/langchain-ai/langgraph	langgraph 1.2.9 — monorepo: core (~28k LOC), prebuilt, checkpoint (+postgres/sqlite/conformance), cli, sdk-py; ~63k LOC of tests (2.3× the code)	fork GHDaru, commit 1e1ca88	2026-07	evaluation
OpenAI Agents SDK	frameworks	github.com/openai/openai-agents-python	v0.18.3	fork GHDaru, commit 5976333	2026-07	evaluation

Harness	Category	Origin	Version/snapshot	Fork/commit read	Evaluated on	Analysis
Software Agent SDK (OpenHands)	frameworks	github.com/OpenHands/software-agent-sdk	v1.37.1	fork GHDaru, commit 99342c4	2026-07	evaluation

Extension ext-1 (2026-07-31): the first Radar → corpus promotion

The corpus grew from 16 to **18** through the very path the book itself institutionalized: the [daily Radar](#) (in Portuguese) found the candidates (2026-07-31 sweep), the editor approved the promotion, the repositories were forked for frozen reading, and the same instrument (`HARNESS_EVAL.md`) was applied — round **ext-1**, without touching the round 1/2 snapshots. Both pass the inclusion test of ch. 01 §4 (open source + general-purpose harness + adoption/representativeness): Grok Build for the opening of a complete commercial harness; Pi as a **deliberately atypical case** (Yin's replication logic called for a minimalist counterpoint, and the corpus was missing one).

Harness	Category	Origin	Version/snapshot	Fork/commit read	Evaluated on	Analysis
Grok Build (xAI)	coding harnesses	github.com/xai-org/grok-build	snapshot 2026-07 (opened on 2026-07-15, Apache 2.0)	fork GHDaru/grok-build, commit dd04f39	2026-07-31 (round ext-1)	evaluation
Pi (Earendil Labs)	coding harnesses	github.com/badlogic/pi-mono	snapshot 2026-07-31	fork GHDaru/pi, commit 7846534	2026-07-31 (round ext-1)	evaluation

Extension ext-2 (2026-08-02): the second promotion — and a new category

The corpus grew from 18 to **20** through the same path: the Radar confirmed QM in a primary source (2026-08-02 sweep) and the critical reading of a vendor marketing article led to Kimi Code, verified at the source; the editor approved, the repositories were forked and the instrument applied — round **ext-2**. Kimi Code enters under the same criterion as Grok Build (the second model vendor opening a complete harness — the pattern became a trend, ch. 14). QM did not fit any existing archetype and **inaugurates the "organizational agents" category**: the unit of design is the organization (scopes, audience-based permissions, consent, auditing), and the agent loop is a swappable engine — indeed, **Pi itself, evaluated in ext-1, is a dependency here** (`package.json`). Yin's replication logic called for exactly this: a case that would test the limits of the taxonomy.

Harness	Category	Origin	Version/snapshot	Fork/commit read	Evaluated on	Analysis
Kimi Code (Moonshot AI)	coding harnesses	github.com/MoonshotAI/kimi-code	CLI 0.31.1 (opened ~2026-06, MIT)	fork GHDaru/kimi-code, commit e22479a	2026-08-02 (round ext-2)	evaluation
QM (Y Combinator)	organizational agents	github.com/yc-software/qm	snapshot 2026-07-31 (opened on 2026-07-31, MIT)	fork GHDaru/qm, commit 7f2c916	2026-08-02 (round ext-2)	evaluation

Extension ext-3 (2026-08-02): evaluated and not included — the inclusion test at work

The method documents refusals too. **Traycer** (Traycer AI) was nominated by the editor, forked and evaluated with the full instrument (fork GHDaru/traycer, commit 65fc3d7, MIT) — and **did not pass the inclusion test** of ch. 01 §4: the open repository (~513k lines) contains clients, a CLI and a remarkable orchestration protocol, but **none of the four harness pieces** — the Host that runs loop, context, tools and control is a signed closed binary with a mandatory cloud (the repo's own `AGENTS.md` states the Host and backends are not there). The [full evaluation](#) (18/36) stays on record: it is the study's best-documented case of "open source" as a client-distribution strategy, and the central evidence of the new [Appendix — The supply chain](#), for which the reading yielded the map of 18 orchestrated harnesses.

Extension ext-4 (2026-08-06): the corpus reaches 21 — and the first synthesis confronted

Prime Agent (Prime Intellect) arrived through an editor's tip on announcement day and was the first candidate to threaten an **Executive summary** rather than add an addendum: the launch claims that context compaction "forces the model to work around its own scaffolding."

Reading the code **upheld the ch. 04 synthesis and recorded a caveat** — compaction was not removed, it was *subordinated to the agent* (see [ch. 04](#)).

The round's structural finding is another: **Prime Agent is built on Pi** — the same four packages, a LICENSE with dual copyright (Mario Zechner + Prime Intellect), a README crediting `pi-mono`. It is the **fifth** Pi consumer recorded in the [supply chain appendix](#), and it closes an argument from ch. 12: the **lowest-scoring system in the corpus** (Pi, 26/36 — which refuses half the dimensions by manifesto) was the base a frontier lab chose to build the study's most radical harness on.

Harness	Category	Origin	Version/snapshot	Fork/commit read	Evaluated on	Analysis
Prime Agent (Prime Intellect)	coding harnesses	github.com/PrimeIntellect-ai/prime-agent	workspace 0.7.0 (MIT)	fork GHDaru/prime-agent, commit 0e0d233	2026-08-06 (round ext-4)	evaluation

Consolidated diagnosis

The **per-dimension results** (scores 0–3, with evidence) and the comparative diagnosis live in the [Harness Comparison](#) — including the interactive heatmap. Beyond the scores, each individual evaluation brings: the harness's **observed archetype**, its strengths with file paths, and the **"what to steal"** section (patterns worth carrying over to other harnesses).

Method note (ch. 01 §6): selection followed replication logic (Yin) — representative *and* deliberately atypical cases; the unit of analysis is the source code; the scores follow the template's fixed grid (feature analysis, DESMET). The expiration scoreboard for the predictions is in the [History](#) (in Portuguese).

See also: reference implementations beyond the ones evaluated here are catalogued in [Awesome Harness Engineering — Reference Implementations](#).

Appendix — The supply chain

Appendix — The harness supply chain

Map captured on **2026-08-02** (rounds ext-2/ext-3). Like every state of the art in this book, it expires: check it against the [History](#).

When this study began, the corpus was a list of **competitors**: alternative products solving the same problem. Rounds ext-2 and ext-3 revealed something else: the harnesses became **suppliers to one another** — `package.json` dependencies, vendored forks, subprocesses, resumed foreign sessions. This appendix shows the work: who consumes whom, through which mechanism, with the evidence for each link. "Supply chain" here is the manufacturing image (factory A makes the part factory B assembles), and also the security sense of the term: **whoever embeds, inherits the risks**.

The map (evidence per link)

Each row is a link verified by code reading at the frozen commit of the corresponding evaluation (paths relative to each repo's root).

Consumer	Supplier	What it consumes	Mechanism	Evidence
QM	Pi	the default agent engine — with the consumer's own security patch applied	npm dependency on a repackaged fork (qm-pi-coding-agent-0.82.0-security.2)	package.json:58
QM	Claude Code	alternative engine	@anthropic-ai/claude-agent-sdk + in-process MCP server bridging tools	package.json:50; src/harness/claude-harness.ts
QM	Codex CLI	alternative engine	@openai/codex dependency	package.json:60
QM	opencode	alternative engine	opencode-ai + plugin/SDK	package.json:61-62,72
Kimi Code	Pi	the entire TUI	vendored fork of pi-tui, with a public acknowledgment	packages/pi-tui/; README.md:122
software-agent-sdk	Codex CLI, gemini-cli	whole harnesses as executors	ACP subprocesses orchestrated by ACPAgent	openhands/agent_server/conversation_service.py:723; event_service.py:873
Grok Build	Claude Code, Codex, Cursor	the competitors' sessions (resumable) and their context artifacts (AGENTS.md/CLAUDE.md/ .cursor)	reading native formats + session picker	crates/codegen/xai-grok-pager/src/views/session_picker.rs
n8n	LangChain	the AI Agent node's foundation — currently being re-internalized (V3)	@langchain/* dependencies	packages/@n8n/nodes-langchain/package.json
Pi	← third parties	the xAI provider reaches Pi from outside , via a community package	extension mechanism (pi-xai-oauth)	radar 2026-08-01
Traycer	Claude Code	GUI+TUI engine: resume/fork, lifecycle hooks, remote management of its MCP/plugins/skills	SDK + PTY claude -- resume -- fork-session + hooks → traycer CLI	protocol/src/host/agent/tui/unary-schemas.ts:48-80; clients/traycer-cli/src/commands/agent-activity-from-hook.ts
Traycer	Codex CLI	GUI+TUI engine	codex app-server (JSON-RPC) + PTY codex resume	protocol/src/host/agent/tui/unary-schemas.ts:70-80
Traycer	opencode	engine and substrate of its own inference (per-user OpenCode server behind the Traycer backend)	PTY + server spawn with an account header	protocol/src/common/schemas.ts:70-76; agent-runtime.ts:839-849
Traycer	Pi (and the Oh My Pi fork)	GUI engines — the fork <i>alone</i> motivated protocol version v6.0	Pi's native RPC	agent-runtime.ts:925-946; provider-schemas.ts:80-135
Prime Agent	Pi	the entire base — Pi's four packages, with the tool set collapsed into a single ipython and two new layers on top	thesis fork; LICENSE with dual copyright	LICENSE; packages/{agent,ai,coding-agent,tui}/package.json; README.md

Consumer	Supplier	What it consumes	Mechanism	Evidence
			(Mario Zechner + Prime Intellect)	
Traycer	Hermes, Kimi Code, Cursor, +ACP	GUI engines (8+ providers via ACP: hermes acp, kimi acp, grok agent stdio, qwen --acp...)	ACP stdio processes / @cursor/sdk	agent-runtime.ts:851-941; protocol/src/host/agent/shared.ts:35-43

Add the **editorial production** links: Traycer materializes skills from public registries (anthropics/skills, vercel-labs) pinned by hash in a lockfile (`skills-lock.json`) for the agents that write its own repo — harness consumption starting before the product even exists.

The extreme case: Traycer, the cockpit that is all chain

Round **ext-3** evaluated [Traycer](#) (18/36) — a product whose *entire* proposition is consuming other harnesses: a multiplayer cockpit (~513k open lines) that catalogs, in its own wire contract, the resume/fork semantics of **18 competing CLIs/SDKs**, with 6 provider enums frozen per protocol version. It **did not pass the inclusion test** of ch. 01 §4 — the four harness pieces are not in the open code: the Host that runs loop, context and control is a signed closed binary with a mandatory cloud (the repo's own `AGENTS.md` says so; full evidence in the evaluation). The record stays for two reasons: it is the corpus's best-documented case of **"open source" as a client-distribution strategy**, and it is proof that the orchestration layer — buying, driving and reselling the work of other harnesses — became a standalone product.

Three readings

1. **"What is it made of?" became an evaluation question.** A harness is no longer described only by what it does, but by the links it embeds. Pi today feeds **at least five systems** (QM as engine, Kimi Code as TUI, Traycer as provider, the Oh My Pi fork — and Prime Agent as its entire base); a failure, a CVE or a license change in that single link propagates through the whole chain, exactly as in physical industry.
2. **The session became an integration interface.** Three different consumers (Grok Build, Traycer, QM) treat other harnesses' *sessions* as resumable artifacts — via native formats, versioned resume/fork anchors, or "tape" re-seeding. It is an emerging pattern with no standard: each one solves it by reverse-engineering the neighbor. If a session-interchange format ever standardizes (ch. 17), much of this map becomes compatibility code — the expiration clause applied to this very appendix.
3. **Enforcement does not travel down the chain.** When QM runs Pi, the permissions are QM's (Pi has none); when Traycer drives 18 harnesses, the permission mode is a **relay** — and Traycer's A2A instruction even tells derived agents to operate in `full_access` by default. Whoever consumes a harness inherits its capabilities, but does **not** automatically inherit its controls — the weakest link in the chain sets the risk for the whole.

The counterpoint that confirms it

The two most-consumed suppliers on the map are also the ones that take the *classic* supply chain most seriously: Pi pins dependencies and allowlists lifecycle scripts (`--ignore-scripts` everywhere); QM audits its supplier to the point of **patching it** (the security patch on line 58). The lesson closes the circle of ch. 07: in the era when your neighbor's harness is your dependency, supply-chain security stopped being an npm topic and became a topic of **agent architecture**.

See also: the full evaluations of each link are in the [Appendix — The study](#); the editorial reading of the trend is in the [Comparison](#) (round ext-2) and in chapter [14 — Convergences](#).

Appendix — Book usage

Appendix — Book usage (live)

State of the art captured in 2026-07 · last revised 2026-07-29 · [history and expiration log](#)

A book that teaches **verification** (ch. 11) and **interface observability** (ch. 13) should be able to answer a question about itself: *how is this book used?* This page answers it — live.

What is measured (and what is not)

Since edition 0.49, the site records **aggregate navigation**: which page was visited, how many times, by **anonymous** sessions — and only after the reader **accepts the telemetry notice** (the banner on the first visit). Nothing beyond that:

- We do **not** collect IP, user-agent, name, email or any personal data;
- the session is a random identifier generated by the browser, erasable by the reader themselves (the companion's `/limpar` command removes everything from the session — the right to be forgotten);
- the panel below consumes a **strictly aggregated** projection (`total` and per-page counts) — there is no public endpoint with individual data.

The live panel

(The numbers above exist only in the online version — in the PDF this island is omitted by definition.)

What it is for

This panel is the same input that drives the **living book's cadence** ([ADR 0007](#), in Portuguese): chapters getting more reader attention take priority in the quarterly revision window, and ignored pages raise the right editorial question — is it lacking promotion, or lacking a rewrite?

It is also a miniature demonstration of what the book preaches: **instrumenting is the cheap half of verification** — the expensive half is deciding what to do with the number. The decision record lives, as always, in the [History](#) (in Portuguese).

Appendix — Book graph

Appendix — The book's graph (live)

State of the art captured in 2026-07 · last revised 2026-07-30 · [history and expiration log](#)

Every technical book is a graph disguised as a sequence: chapters that cite each other, systems that show up across several dimensions, concepts that stitch it all together. This page makes the graph explicit — and **interactive**.

How this graph is built (and why it never goes stale)

The nodes and edges are **neither hand-edited nor generated by a model**: they are extracted **deterministically from the Markdown itself** on every site build, by the publishing engine (`publicar/grafo.mjs`). A chapter → harness edge exists because that chapter *mentions* that system in the text — with a weight equal to the number of mentions (Principle I: every edge is verifiable textual evidence). Since the build runs on **every published change to the book**, the graph tracks the content by construction — updating it is not a process, it is a property.

- **Nodes** (4 types): the 18 **chapters**; the 16 **systems in the study corpus**; key **concepts** (MCP, A2A, ACP, LSP, RAG, MAST); the 13 **harness-zero steps**.
- **Edges**: chapter → chapter ("ch. NN" cross-references), chapter → system (mentions), chapter → concept (occurrences), chapter → step (the Hands-on track).

The interactive graph

(The visualization exists only in the online version — in the PDF this island is omitted by definition.)

Reading guide

- **Hubs**: the larger nodes concentrate connections — expect to see ch. 02 (Loop) and systems evaluated across every dimension as gravitational centers.
- **Bridges**: concepts like MCP connect groups that would otherwise sit far apart (ch. 06 to ch. 17, the corpus to the protocols) — the book's stitching made visible.
- **The practical track**: filter by "harness-zero" to see how the steps tie the theoretical chapters to the build (Backward Design in graph form).
- **Click any node** to isolate its neighborhood and navigate straight to the corresponding page.

Appendix — harness-um (the reference)

Appendix — harness-um: the reference implementation

State of the art captured in 2026-07 · last revised 2026-07-31 · [history and expiration log](#)

After eighteen chapters describing what a harness has, the honest question is: *what if we put it all together?* This appendix answers with code. The **harness-um** is the book's reference implementation — the features from chapters 02–13 gathered into a single system, small enough to be read in an afternoon and complete enough to be the starting point for yours.



The official figure: the core (the agent) wrapped by the 12 segments of the ring — chapters 02–13, one per feature. The visual identity is the same as the cover's: the harness is what sits *around*.

Why "harness-um" (and not "openharness")

The name tells the book's progression: the **harness-zero** (Hands-on) builds one feature per step, from scratch; the **harness-um** ("harness-one") is the destination — everything together and cohesive. And there is an editorial reason: "OpenHarness" **already exists** — it is one of the 16 systems in this study's corpus (HKUDS/OpenHarness, an open-source port of Claude Code). Naming the book's reference after a system the book itself evaluates would create exactly the confusion Principle I exists to prevent.

The ubiquitous language

The central decision of the harness-um is not technical, it is **linguistic**: the code speaks the book's language — in Portuguese, and the class names are deliberately kept that way. Every term the chapters defined becomes an identical code name — reading the code is rereading the table of contents. The translation into each model API's dialect (today, Anthropic Messages) happens at a single edge (`provedores.py`, the providers module), the **anticorruption layer**: if the provider changes, the domain never even hears about it.

Book term	In the code	Chapter
Agent loop	<code>LoopDoAgente.executar()</code> — the agent loop's <code>run()</code>	02
Turn (with a budget)	<code>max_turnos</code>	02
Context assembly	<code>MontadorDeContexto</code> (named layers)	03
Compaction	<code>Compactador</code> (summary + intact tail)	04
Tool	<code>Ferramenta</code> , <code>@ferramenta</code> , <code>CaixaDeFerramentas</code>	05
MCP	<code>ClienteMCP</code> (stateless, spec 2026-07-28)	06
Permissions	<code>Politica</code> → <code>PERMITIR</code> / <code>PERGUNTAR</code> / <code>NEGAR</code> (allow / ask / deny)	07
Durable memory	<code>Memoria</code> (<code>MEMORIA.md</code>)	08
Session	<code>Sessao</code> (JSONL, append-only)	08
Plan as artifact	<code>Plano</code> (persisted, re-injected)	09
Subagent	<code>tarefa()</code> — clean context, read-only toolbox	10
Verification	<code>Verificador</code> (post-mutation, verdict to the model)	11
Hook	<code>Gancho</code> (deterministic, can veto)	12
Skill	<code>Habilidade</code> (<code>SKILL.md</code> , progressive disclosure)	12
Interface	REPL (<code>python -m harness_um</code>)	13
Provider	<code>Provedor</code> → <code>ProvedorAnthropic</code> , <code>ProvedorEco</code>	02, 11

How to download and run

The code lives **in this repository**, alongside the book — in `harness-um/`:

```
git clone https://github.com/GHDaru/harness_engineering.git
cd harness_engineering/harness-um
pip install -e .

# sem chave nenhuma (ProvedorEco, offline):
python -m harness_um --eco 'leia @usar ler_arquivo {"caminho": "README.md"}'

# com modelo real (chave SÓ no ambiente):
export ANTHROPIC_API_KEY=...
python -m harness_um      # REPL: /plano /memoria /contexto /sair
```

The `ProvedorEco` (the echo provider) deserves the note: it is deterministic and obeys `@usar ferramenta {...}` directives — enough to exercise the whole loop (tool-use, permissions, hooks, verification) **with no network and no cost**. That is why the `harness-um` tests run in the book's CI on every push: the reference must not rot in silence.

harness-zero × harness-um

	harness-zero	harness-um
Purpose	teaching how to build (Backward Design)	showing the finished whole
Form	13 steps, each a complete app	1 cohesive package (<code>harness_um/</code>)
Reading	during the chapters	after the book
Analogue	exercise workbook	annotated answer key

Expiration

Like everything in this book: the harness-um is the **2026-07** snapshot — the MCP client is born on the 2026-07-28 spec, but providers, schemas and conventions change in months. The [living book's cadence](#) (ADR 0007, in Portuguese) covers this appendix too; the code carries the same clause in its README.

Bibliography

The book's scientific bibliography

Editorial rule: no reference enters a chapter without status ✓ **validated** (ID ↔ title confirmed by an independent source). 2026-07-29 review (spec 050): every ⌚ **item was verified by independent web search** and promoted to ✓ (with two corrections recorded: [arXiv 2509.18661](#) is the *Agentic* AutoSurvey; Norton's ISBN 9780226595146 is the 1st ed., 2009). ★ = chapter anchor.

Overall status

Status	Meaning
✓	ID ↔ title confirmed by independent search in this session
⌚	Cited from memory or from a single source; confirm before citing in the body

Cross-cutting / Foundations (chs. 00–01)

- ★ ✓ **From Question Answering to Task Completion: A Survey on Agent System and Harness Design** — arXiv [2606.20683](#). The survey exactly on this book's scope; a candidate for the theoretical spine of ch. 01.
- ✓ **Recursive Agent Harnesses** — arXiv [2606.13643](#). A finding from the validation pass; assess its fit (composite harnesses — connects with chs. 10 and 15).
- ✓ **ReAct: Synergizing Reasoning and Acting in Language Models** (Yao et al.) — arXiv [2210.03629](#). The seminal paper of the reasoning+action loop.
- ✓ **Li, Xinzhe A Review of Prominent Paradigms for LLM-Based Agents: Tool Use, Planning (Including RAG), and Feedback Learning** — COLING 2025, pp. 9760–9779 ([aclanthology](#); [arXiv 2406.05804](#)).
- ✓ **Agent Systems with Harness Engineering** (Tang, Peng, Chen et al., RUC/Gaoling) — OpenReview [nM5tDHRQsx](#) · [PDF + curation](#) (May 2026; **no arXiv version**; 62-pp. PDF read in full, spec 065). The second survey on this book's scope — and the complement of the anchor above: a scaffold-side taxonomy convergent with ours (workflow/memory/skills/multi-agent) plus a whole third that the book does not cover (**agentic training**: RL, rewards, rollout infra). The quotable central thesis: harness engineering as "the joint optimization of both components" (model↔scaffold). Rigor caveats: no limitations section, no declared survey methodology, sample of real systems n=3 — and permissions, extensibility and interfaces (strong in our benchmark) treated as future directions, not first-class components.

History and provenance (ch. 01 §2–3) — added in the rigor review

- ✓ **ReAct: Synergizing Reasoning and Acting in Language Models** (Yao et al.) — arXiv [2210.03629](#), ICLR 2023. The Thought → Action → Observation loop; the skeleton of every harness.
- ✓ **Introducing GitHub Copilot: your AI pair programmer** (GitHub, Jun 2021) — [github.blog](#). Marks the "before": autocomplete by Codex, no loop/tools.
- ✓ **Function calling and other API updates** (OpenAI, Jun 2023) — [openai.com](#). The model → tools link.
- ✓ **Introducing the Model Context Protocol** (Anthropic, Nov 2024) — [anthropic.com](#).
- ✓ **Announcing the Agent2Agent Protocol (A2A)** (Google, Apr 2025) — [developers.googleblog.com](#).
- ✓ **AGENTS.md** — [agents.md](#). The "README for agents".
- ✓ **BabyAGI** (Yohei Nakajima, Apr 2023) — [babyagi.org](#). · **AutoGPT** (Significant Gravitas, Mar 2023) — [repository](#).
- ✓ **Aider** (Paul Gauthier, 2023) — [github.com/Aider-AI/aider](#).
- ✓ **Chain-of-Thought Prompting Elicits Reasoning in Large Language Models** (Wei et al., 2022) — [arXiv 2201.11903](#). · ✓ **Toolformer: Language Models Can Teach Themselves to Use Tools** (Schick et al., Meta, 2023) — [arXiv 2302.04761](#).

The study's methodology (ch. 01 §6) — added in the rigor review

- ✓ **Hassan, A. E. (2008).** *The Road Ahead for Mining Software Repositories*. FoSM/ICSM 2008. Repositories as primary data (MSR).
- ✓ **Runeson, P. & Höst, M. (2009).** *Guidelines for Conducting and Reporting Case Study Research in Software Engineering*. Empirical Software Engineering 14(2). The case-study protocol in SE.
- ✓ **Kitchenham, B., Linkman, S. & Law, D. (1997).** *DESMET: a methodology for evaluating software engineering methods and tools*. IEE CCEJ. The benchmark's feature analysis.
- ✓ **Sim, S. E., Easterbrook, S. & Holt, R. C. (2003).** *Using Benchmarking to Advance Research*. ICSE 2003. The benchmark as a scientific engine.
- ✓ **Stol, K.-J., Ralph, P. & Fitzgerald, B. (2016).** *Grounded Theory in Software Engineering Research*. ICSE 2016. Inductive coding of artifacts.
- ✓ **Peppers, K. et al. (2007).** *A Design Science Research Methodology for IS Research*. JMIS 24(3). The harness-zero's DSRM process.
- ✓ **Yin, R. K. (2018)** *Case Study Research and Applications: Design and Methods*, 6th ed., SAGE (ISBN 9781506336169). • ✓ **Hevner, March, Park & Ram (2004)** *Design Science in Information Systems Research*, MIS Quarterly 28(1), 75–105. • ✓ **Basili, Caldiera & Rombach (1994)** *The Goal Question Metric Approach*, Encyclopedia of Software Engineering, vol. 1, Wiley, 528–532. • ✓ **Hsieh & Shannon (2005)** *Three Approaches to Qualitative Content Analysis*, Qualitative Health Research 15(9), 1277–1288 (DOI 10.1177/1049732305276687). • ✓ **Cook & Campbell (1979)** *Quasi-Experimentation: Design & Analysis Issues for Field Settings*, Houghton Mifflin (ISBN 9780395307908).

Ch. 02 — Agent Loop

- ✓ ReAct (above).
- ✓ **LLM-based Agentic Reasoning Frameworks: A Survey** — arXiv [2508.17692](#).
- ✓ **A Comprehensive Survey on RL-based Agentic Search** — arXiv [2510.16724](#) (the trained loop, the chapter's frontier).

Ch. 03 — Context Delivery

- ★ ✓ **A Survey of Context Engineering for Large Language Models** — arXiv [2507.13334](#).
- ✓ **Lost in the Middle: How Language Models Use Long Contexts** (Liu et al.) — arXiv [2307.03172](#). The empirical basis of "position matters" (justifies tail preservation and layered prompts).
- ✓ **Less Context, Better Agents: Efficient Context Engineering for Long-Horizon Tool-Using LLM Agents** — arXiv [2606.10209](#).

Ch. 04 — Compaction

- ★ ✓ **MemGPT: Towards LLMs as Operating Systems** (Packer et al.) — arXiv [2310.08560](#). The "virtual memory" formulation that anticipated the compaction ladder.
- ✓ **ContextBudget: Budget-Aware Context Management for Long-Horizon Search Agents** — arXiv [2604.01664](#).
- ✓ **The Missing Memory Hierarchy: Demand Paging for LLM Context Windows** — arXiv [2603.09023](#).
- ✓ **CompactionRL: Reinforcement Learning with Context Compaction for Long-Horizon Agents** (Li, Hou, Jing, Tang, Dong — Tsinghua/Z.AI) — arXiv [2607.05378](#) (preprint, 06-Jul-2026; **full text read and citations verified**, spec 066). The "third way" of the chapter's addendum: summarization learned during training with task reward (+7.0 Pass@1 on SWE-bench Verified with GLM-4.5-Air, Table 2); validates the chapter's threshold+summary+tail triad; declared limitation: train–test mismatch (model ↔ harness coupling); and Table 1's pro-harness finding — swapping only the summarizer moves +6.5 points.
- ✓ Lost in the Middle (ch. 03) — grounds *what* to preserve.

Ch. 05 — Tools

- ✓ **The Evolution of Tool Use in LLM Agents: From Single-Tool Call to Multi-Tool Orchestration** — arXiv [2603.22862](#).
- ✓ **Tool Learning with Large Language Models: A Survey** (Qu et al.; accepted at Frontiers of Computer Science) — arXiv [2405.17935](#) (+ [repo](#)).

- ✓ **Gorilla: Large Language Model Connected with Massive APIs** (Patil et al., 2023) — [arXiv 2305.15334](#). · ✓ **ToolLLM: Facilitating LLMs to Master 16000+ Real-world APIs** (Qin et al., 2023) — [arXiv 2307.16789](#).

Ch. 06 — MCP

Update (living book, 2026-07): the gap recorded in earlier rounds has been **filled** — MCP accumulated an SoK, *tool poisoning* benchmarks and empirical server audits. The standard remains an *industry spec*; academia came in through the **security** door.

- ★ ✓ **Model Context Protocol (MCP): Landscape, Security Threats, and Future Research Directions** (Hou et al.) — [arXiv 2503.23278](#); also ACM TOSEM. The canonical SoK: server lifecycle + threat taxonomy per phase.
- ★ ✓ **MCPTox: A Benchmark for Tool Poisoning Attack on Real-World MCP Servers** (Wang, Gao et al.) — [arXiv 2508.14925](#). 45 real servers / 353 tools; success of up to ~73%; more capable models were more susceptible.
- ✓ **Model Context Protocol (MCP) at First Glance: Studying the Security and Maintainability of MCP Servers** (Hasan, Li, Fallahzadeh, Rajbahadur, Adams, Hassan) — [arXiv 2506.13538](#). 1,899 servers audited: 7.2% with general vulns, 5.5% with *tool poisoning*.
- ✓ **MCP Safety Audit: LLMs with the Model Context Protocol Allow Major Security Exploits** (Radosevich & Halloran) — [arXiv 2504.03767](#). Exploits via legitimately registered tools; the MCPSafetyScanner tool.
- ✓ **A Survey of Agent Interoperability Protocols: MCP, ACP, A2A, and ANP** (Ehtesham, Singh et al.) — [arXiv 2505.02279](#). Choose the protocol by trust context (links to ch. 17).
- ✓ **Not what you've signed up for: ...Indirect Prompt Injection** (Greshake et al.) — [arXiv 2302.12173](#). The first-principles basis: retrieved content is an instruction channel.
- ~ **Threat Modeling and Analysis of Vulnerabilities to Prompt Injection with Tool Poisoning** — [arXiv 2603.22489](#); MDPI *J. Cybersecurity and Privacy* 6(3):84 (2026). STRIDE+DREAD over MCP components. (*ID and venue verified; author list not confirmed by snippet.*)

Industry sources (docs/vendor/practitioners) along the Ch. 06 line below.

Ch. 07 — Permissions and Sandboxing

- ★ ✓ **Not what you've signed up for: Compromising Real-World LLM-Integrated Applications with Indirect Prompt Injection** (Greshake et al.) — [arXiv 2302.12173](#). The paper that defined the threat.
- ✓ **A Systematic Survey of Security Threats and Defenses in LLM-Based AI Agents: A Layered Attack Surface Framework** — [arXiv 2604.23338](#).
- ✓ **A Survey on Agentic Security: Applications, Threats and Defenses** — [arXiv 2510.06445](#).
- ✓ **Safety and Security Threats of Computer-Using Agents** — [arXiv 2505.10924](#).

Ch. 08 — Memory and State

- ★ ✓ **MemGPT: Towards LLMs as Operating Systems** (Packer et al.) — [arXiv 2310.08560](#). Context as scarce RAM; recall/archival tiers; the agent pages via tool ("context page faults").
- ★ ✓ **Generative Agents: Interactive Simulacra of Human Behavior** (Park et al.) — [arXiv 2304.03442](#); UIST '23. The *memory stream* and recall by **recency** × **importance** × **relevance** + consolidation through reflection.
- ★ ✓ **Cognitive Architectures for Language Agents (CoALA)** (Sumers et al.) — [arXiv 2309.02427](#). The episodic/semantic/procedural taxonomy + working memory (Tulving's foundation).
- ✓ **A Survey on the Memory Mechanism of LLM-based Agents** (Zhang et al.) — [arXiv 2404.13501](#); ACM TOIS. Sources · forms · operations (writing/management/reading).
- ✓ **MemoryBank: Enhancing LLMs with Long-Term Memory** (Zhong et al.) — [arXiv 2305.10250](#); AAAI '24. Forgetting controlled by an Ebbinghaus curve (time × access frequency).
- ✓ **Reflexion: Language Agents with Verbal Reinforcement Learning** (Shinn et al.) — [arXiv 2303.11366](#); NeurIPS '23. Verbal self-reflection persisted in an episodic buffer (bridge to ch. 16).

- ✓ **A-MEM: Agentic Memory for LLM Agents** (Xu et al.) — arXiv [2502.12110](#). Self-organizing structured notes (Zettelkasten).
- ✓ **Mem0: Production-Ready AI Agents with Scalable Long-Term Memory** (Chhikara et al.) — arXiv [2504.19413](#); ECAI '25. Extract → consolidate → retrieve pipeline; the LoCoMo benchmark.
- ✓ **A Survey on the Memory Mechanism** and evolution surveys: **From Storage to Experience** — arXiv [2605.06716](#); **From Human Memory to AI Memory** — arXiv [2504.15965](#); **Governing Evolving Memory in LLM Agents (SSGM)** — arXiv [2603.11768](#) (also ch. 16).
- ~ **Zep: A Temporal Knowledge Graph Architecture for Agent Memory** — arXiv [2501.13956](#). Bi-temporal graph; outdated facts invalidated, not deleted. (*Recurring ID in searches; not opened byte-by-byte through the proxy.*)

Ch. 09 — Planning

- ★ ✓ **ReAct: Synergizing Reasoning and Acting in Language Models** (Yao et al.) — arXiv [2210.03629](#); ICLR '23. Interleaving reason and action in the same loop.
- ★ ✓ **Understanding the Planning of LLM Agents: A Survey** (Huang et al.) — arXiv [2402.02716](#). Five-way taxonomy (decomposition · selection · external module · reflection · memory).
- ✓ **Plan-and-Solve Prompting** (Wang et al.) — arXiv [2305.04091](#); ACL '23. An explicit plan before solving (known scope).
- ✓ **Tree of Thoughts** (Yao et al.) — arXiv [2305.10601](#); NeurIPS '23. Search over plans with backtracking.
- ✓ **ADaPT: As-Needed Decomposition and Planning** (Prasad et al.) — arXiv [2311.05772](#); NAACL Findings '24. Decompose only when the executor fails.
- ✓ **Beyond Entangled Planning: Task-Decoupled Planning for Long-Horizon Agents** — arXiv [2601.07577](#). A DAG of sub-goals with scoped context (~82% tokens).
- ✓ **PlanGenLLMs: A Modern Survey of LLM Planning Capabilities** (Wei et al.) — arXiv [2502.11221](#); ACL '25. Six plan-evaluation criteria.
- ✓ **PLANET: Benchmarks for Evaluating LLMs' Planning Capabilities** — arXiv [2504.14773](#).
- ✓ **PlanBench** (Valmeekam et al.) — arXiv [2206.10498](#); NeurIPS '22 Datasets. Raw models fail at plan generation → external validators.
- ✓ **TravelPlanner** (Xie et al.) — arXiv [2402.01622](#); ICML '24. Agents lose the thread of multiple constraints → externalize the tracking.

Ch. 10 — Subagents and Orchestration

- ★ ✓ **Why Do Multi-Agent LLM Systems Fail? (MAST)** (Cemri, Pan, Yang et al.) — arXiv [2503.13657](#). 14 failure modes in 3 categories; most come from *design*, not from the model — the most actionable paper for anyone building orchestration.
- ★ ✓ **MetaGPT: Meta Programming for a Multi-Agent Collaborative Framework** (Hong et al.) — arXiv [2308.00352](#); ICLR '24. SOPs + assembly-line roles against cascading hallucination.
- ✓ **AutoGen: Multi-Agent Conversation** (Wu et al.) — arXiv [2308.08155](#). "Conversable" agents with programmable interaction topology.
- ✓ **CAMEL: Communicative Agents** (Li et al.) — arXiv [2303.17760](#); NeurIPS '23. Inception-prompting for role stability (role-play drifts).
- ✓ **ChatDev: Communicative Agents for Software Development** (Qian et al.) — arXiv [2307.07924](#); ACL '24. Chat chain + "communicative dehallucination".
- ✓ **AgentVerse** (Chen et al.) — arXiv [2308.10848](#); ICLR '24. Dynamic recruitment + guardrails for emergent behavior.
- ✓ **LLM-based Multi-Agents: A Survey of Progress and Challenges** (Guo et al.) — arXiv [2402.01680](#); IJCAI '24. Taxonomy (interface · profiles/roles · communication · capability).
- ✓ **Improving Factuality and Reasoning through Multiagent Debate** (Du et al.) — arXiv [2305.14325](#); ICML '24. Debate as a verification primitive.
- ~ **Should We Be Going MAD?** (Smit et al.) — arXiv [2311.17371](#) · **Stop Overvaluing Multi-Agent Debate** (Zhang et al.) — arXiv [2502.08788](#). The skeptical counterweight: compare against a *compute-matched* single-agent baseline before adopting the complexity.
- ✓ **D3MAS: Decompose, Deduce, Distribute** — arXiv [2510.10585](#).

Ch. 11 — Verification and Evals

- ★ ✓ **SWE-bench: Can Language Models Resolve Real-World GitHub Issues?** (Jimenez et al.) — arXiv [2310.06770](https://arxiv.org/abs/2310.06770); ICLR '24. Grading by executing the repo's real tests (FAIL_TO_PASS/PASS_TO_PASS), not string-match.
- ★ ✓ **Large Language Models Cannot Self-Correct Reasoning Yet** (Huang et al.) — arXiv [2310.01798](https://arxiv.org/abs/2310.01798); ICLR '24. Do not trust *intrinsic* self-correction — an external verifier is needed.
- ✓ **SWE-agent: Agent-Computer Interfaces Enable Automated SE** (Yang et al.) — arXiv [2405.15793](https://arxiv.org/abs/2405.15793); NeurIPS '24. Tool ergonomics (ACI) drives success, not just the model.
- ✓ **Judging LLM-as-a-Judge with MT-Bench and Chatbot Arena** (Zheng et al.) — arXiv [2306.05685](https://arxiv.org/abs/2306.05685); NeurIPS '23. An LLM judge is viable (~80% agreement) but has biases (position/verbosity/self-preference).
- ✓ **A Survey on LLM-as-a-Judge** (Gu et al.) — arXiv [2411.15594](https://arxiv.org/abs/2411.15594). Judge reliability as a pipeline concern (rubrics, gold set, self-agreement).
- ✓ **CRITIC: LLMs Can Self-Correct with Tool-Interactive Critiquing** (Gou et al.) — arXiv [2305.11738](https://arxiv.org/abs/2305.11738); ICLR '24. Self-critique anchored in tools (does the code run? does the fact check out?) beats introspection.
- ✓ **Self-Consistency Improves CoT** (Wang et al.) — arXiv [2203.11171](https://arxiv.org/abs/2203.11171); ICLR '23. Sample paths + majority vote: cheap model-only verification.
- ✓ **τ-bench: Tool-Agent-User Interaction** (Yao et al.) — arXiv [2406.12045](https://arxiv.org/abs/2406.12045). Verify the *final state of the world* (not the transcript); pass^k reveals inconsistency.
- ✓ **Survey on Evaluation of LLM-based Agents** (Yehudai et al.) — arXiv [2503.16416](https://arxiv.org/abs/2503.16416). Axes: capability · safety · robustness · cost; prefer held-out benchmarks.
- ✓ **Tulu 3 / RLVR** (Lambert et al., Ai2) — arXiv [2411.15124](https://arxiv.org/abs/2411.15124). Reinforcement Learning with Verifiable Rewards: a deterministic verifier is a signal and a reward that is harder to game.
- ~ **Reward Hacking in Language Model Agents (AI Safety Gridworlds)** — arXiv [2606.15385](https://arxiv.org/abs/2606.15385); **Do Coding Agents Deceive Us? (Capped Evaluation with Randomized Tests)** — arXiv [2606.07379](https://arxiv.org/abs/2606.07379). The agent plays against the verifier → held-out/randomized + immutable tests. (*recent; ID by cross-search.*)
- ✓ **The 2025 AI Agent Index** — arXiv [2602.17753](https://arxiv.org/abs/2602.17753) (FAccT '26).
- ✓ **Rethinking the Evaluation of Harness Evolution for Agents** (Wang et al. — AI2/UW/indep.) — arXiv [2607.12227](https://arxiv.org/abs/2607.12227) (preprint, 14-Jul-2026; **full text read and citations verified**, spec 066 — two sentences circulating as quotes were third-party paraphrases and were replaced with the verbatim text). The methods paper of the chapter's addendum: automatic harness evolution does not consistently beat test-time scaling under a matched budget (K=5; Tables 1–2), generalizes +0.6 on held-out (Table 3), and "most edits memorize fixes rather than distilling strategies" (§5.1) — the three rules (matched budget, search/evaluation separation, a design-sensitive instrument) hold for any harness evaluation, including this book's.

Ch. 12 — Extensibility

There is no academic canon of *agent-harness extensibility* (gap confirmed in 2026-07). The durable citations are classic SE on extensible architectures + the security of plugin ecosystems.

- ★ ✓ **On Plug-ins and Extensible Architectures** (Dorian Birsan) — *ACM Queue* 3(2):40–46 (2005), DOI [10.1145/1053331.1053345](https://doi.org/10.1145/1053331.1053345). Eclipse's plug-in model and the "plug-in hell" warning.
- ★ ✓ **LLM Platform Security: ...OpenAI's ChatGPT Plugins** (Iqbal, Kohno, Roesner) — arXiv [2309.10254](https://arxiv.org/abs/2309.10254); AIES '24. The platform/plugin/user "trust triangle" — a third-party extension is not trusted by default.
- ✓ **Policy/Mechanism Separation in Hydra** (Levin, Cohen, Corwin, Pollack, Wulf) — SOSP '75, DOI [10.1145/800213.806531](https://doi.org/10.1145/800213.806531). The origin of "separate mechanism from policy": the harness provides mechanism, the extension provides policy.
- ✓ **Protecting Browsers from Extension Vulnerabilities** (Barth, Felt, Saxena, Boodman) — NDSS '10. Over-privilege: 88% of extensions request more power than they need → least-privilege + isolation.
- ✓ **AIOS: LLM Agent Operating System** (Mei et al.) — arXiv [2403.16971](https://arxiv.org/abs/2403.16971). A kernel that isolates scheduling/memory/tools from agent applications (microkernel applied to agents).
- SE foundations (canonical books): **Microkernel pattern** (Buschmann et al., *POSA* v.1, 1996); **Software Product Lines** (Clements & Northrop, 2001); **Open-Closed Principle** (Meyer, *OOSC*, 1988; Martin, 1996) — "open for extension, closed for modification".

- Self-extension (bridge to ch. 16): **Voyager** [2305.16291](#), **CREATOR** [2305.14318](#), **CRAFT** [2309.17428](#), **ToolMaker** [2502.11705](#).

Ch. 13 — Interfaces

There is no academic canon of *agent-harness interfaces* (gap confirmed in 2026-07). The durable citations come from the HCI of human-AI interaction, mixed-initiative and levels of automation — plus a recent trickle (2025-26) of work on human-in-the-loop for agents.

- ★ ✓ **Principles of Mixed-Initiative User Interfaces** (Horvitz) — CHI '99, DOI [10.1145/302979.303030](#). The 12 principles of when the system should act × ask (the decision to "hand over the initiative").
- ★ ✓ **Guidelines for Human-AI Interaction** (Amershi et al.) — CHI '19, DOI [10.1145/3290605.3300233](#). 18 guidelines by phase; the UX of "when it errs" (cheap correction/undo).
- ✓ **A Model for Types and Levels of Human Interaction with Automation** (Parasuraman, Sheridan, Wickens) — IEEE SMC-A 30(3), 2000, DOI [10.1109/3468.844354](#). Automation by stage (acquisition/analysis/decision/action): the autonomy dial need not be global.
- ✓ **Human and Computer Control of Undersea Teleoperators** (Sheridan & Verplank) — MIT tech report, 1978. The 10-level automation scale (the adjustable autonomy dial).
- ✓ **To Trust or to Think: Cognitive Forcing Functions Can Reduce Overreliance on AI** (Buçinca, Malaya, Gajos) — CSCW '21, arXiv [2102.09692](#). Explanation alone does not cure over-reliance; forcing functions do (approval must be a deliberate act).
- ✓ **Overreliance on AI: Literature Review** (Passi & Vorvoreanu) — Microsoft Aether, MSR-TR-2022-12 (2022). A synthesis of the risk of a false sense of supervision.
- ✓ **Magentic-UI: Towards Human-in-the-loop Agentic Systems** (Mozannar et al., Microsoft) — arXiv [2507.22358](#). Co-planning/co-tasking and **action guards** = permission gating.
- ✓ **Design Considerations for Human Oversight of AI** (Faas et al.) — IUI '26, arXiv [2510.19512](#). Twelve considerations for keeping the human *engaged*, not just present.
- ~ **LLM-Based Human-Agent Collaboration and Interaction Systems: A Survey** (Zou et al.) — arXiv [2505.00753](#) (ACL '26). A bibliography entry point. · **Explanation in AI** (Miller) — *Artificial Intelligence* 267 (2019). Contrastive and selective explanations.

Ch. 16 — Learning and Self-Improvement

- ★ ✓ **A Survey of Self-Evolving Agents: What, When, How, and Where to Evolve** — arXiv [2507.21046](#).
- ✓ **Voyager: An Open-Ended Embodied Agent with LLMs** — arXiv [2305.16291](#). The self-written skill library that anticipated Hermes by 3 years.
- ✓ **Adaptation of Agentic AI: Post-Training, Memory, and Skills** — arXiv [2512.16301](#).
- ✓ **A Comprehensive Survey of Self-Evolving AI Agents: A New Paradigm Bridging Foundation Models and Lifelong Agentic Systems** (Fang et al., 2025) — arXiv [2508.07407](#).
- ✓ SSGM (ch. 08) — the risk of poisoned permanent learning.

Chs. 15, 17 — the recorded gap

Embedded harnesses and protocols have **rarefied** academic literature (2026-07 searches returned no dedicated surveys). Editorial record: the book covers these dimensions with specs, benchmark evidence and industry literature — and flags the gap as a research opportunity (a possible "open problems" section in ch. 14). Note: chs. 12 (extensibility) and 13 (interfaces), previously on this list, were anchored in adjacent literature — classic SE on extensible architectures and the HCI of human-AI interaction, respectively (see the Ch. 12 and Ch. 13 sections above). The *agent-specific* gap persists, but the durable foundations exist.

Living collections

- ✓ **Awesome Harness Engineering** — a collection curated by the author: harness-engineering resources, patterns and templates organized **by problem** (the same taxonomy as this book). Referenced in the chapters as "See also", section by section.

Industry sources by chapter (vendor docs and engineering blogs)

Commercial/industrial material grounding each chapter's "Industry sources" section (v3 skeleton). URLs verified as existing by search; direct fetches to anthropic.com/openai.com return 403 (anti-bot) in this environment — content confirmed via snippets and third-party citations.

Ch. 06 (MCP) — **2026-07-28 release:** ✓ [The 2026-07-28 Specification](#) (official blog) · [changelog](#) — stateless core, MRTR, extensions, `ttlms`, deprecation policy. Verified by direct fetch of the announcement on 2026-07-31 (spec 060).

Ch. 02 (Loop): [How the agent loop works](#) · [Loop engineering](#) · [Building Effective AI Agents](#) · [Running agents \(OpenAI Agents SDK\)](#) · [LoopAgent \(Google ADK\)](#) · [Durable AI Loops \(Restate\)](#) · [Durable Execution \(Inngest\)](#)

Ch. 03 (Context): [Effective context engineering](#) · [Prompt caching \(docs\)](#) · [Prompt caching is everything](#) · [AGENTS.md](#) · [Agentic AI Foundation](#) · [How Claude remembers your project](#) · [AGENTS.md Field Guide 2026](#)

Ch. 04 (Compaction): [Compaction \(docs\)](#) · [Auto Compact explained \(CometAPI\)](#) · [Compaction explained \(okhlopkov\)](#) · [Protecting more context \(hyperdev\)](#)

Ch. 05 (Tools): [Writing effective tools for AI agents](#) · [Code execution with MCP](#) · [Tool search tool \(docs\)](#) · [Advanced tool use](#) · [Programmatic tool calling \(docs\)](#) · [Code Mode \(Cloudflare\)](#) · [Apply Patch \(OpenAI docs\)](#) · [GPT-5.1 for developers](#)

Ch. 06 (MCP): [MCP architecture \(spec\)](#) · [Transports \(spec\)](#) · [Introducing MCP \(Anthropic\)](#) · [OpenAI adopts MCP \(TechCrunch\)](#) · [Google embraces MCP \(The New Stack\)](#) · [MCP GA in Copilot Studio \(Microsoft\)](#) · [MCP Auth spec \(Descope\)](#) · [Tool Poisoning \(Invariant Labs\)](#) · [Line jumping \(Trail of Bits\)](#) · [The lethal trifecta \(Willison\)](#) · [MCP Registry \(preview\)](#) · [MCP → Agentic AI Foundation \(Anthropic\)](#)

Ch. 08 (Memory/state): [Manage sessions \(Claude Code\)](#) · [Checkpointing \(Claude Code\)](#) · [File-checkpointing \(Agent SDK\)](#) · [How Claude remembers your project](#) · [Memory tool](#) · [Managing context \(context editing + memory\)](#) · [Effective harnesses for long-running agents](#) · [Memory blocks \(Letta\)](#) · [RAG is not agent memory \(Letta\)](#) · [Memory types \(mem0\)](#) · [Graphiti knowledge-graph memory \(Neo4j\)](#) · [LangMem SDK \(LangChain\)](#) · [Memory vs RAG \(AWS Bedrock AgentCore\)](#)

Ch. 09 (Planning): [Permission modes / plan mode \(Claude Code\)](#) · [Best practices — Explore/Plan/Code/Commit](#) · [Todo tracking \(Agent SDK\)](#) · [The "think" tool](#) · [Extended/interleaved thinking](#) · [GitHub Spec Kit](#) · [Spec-driven development \(GitHub Blog\)](#) · [Kiro specs](#) · [Multi-agent research system \(Anthropic\)](#) · [Don't Build Multi-Agents \(Cognition\)](#)

Ch. 10 (Subagents/orchestration): [Custom subagents \(Claude Code\)](#) · [Subagents \(Agent SDK\)](#) · [Multi-agent research system \(Anthropic\)](#) · [When to use multi-agent \(Claude\)](#) · [Don't Build Multi-Agents \(Cognition\)](#) · [Agents SDK orchestration \(OpenAI\)](#) · [CrewAI processes](#) · [LangGraph multi-agent \(LangChain\)](#) · [Magentic-One \(AutoGen\)](#) · [Multi-agent patterns in ADK \(Google\)](#) · [A2A spec](#) · [ACP joins A2A \(LF AI & Data\)](#)

Ch. 11 (Verification/evals): [SWE-bench Verified \(OpenAI\)](#) · [Why we no longer evaluate SWE-bench Verified](#) · [SWE-bench \(site\)](#) · [Terminal-Bench](#) · [Define success criteria / build evals \(Claude\)](#) · [Demystifying evals for AI agents \(Anthropic\)](#) · [Statistical approach to model evals \(Anthropic\)](#) · [Effective harnesses for long-running agents](#) · [Best practices — TDD](#) · [OpenAI Evals](#) · [Inspect \(UK AISI\)](#) · [promptfoo](#) · [Braintrust scorers](#) · [LangSmith LLM-as-judge](#) · [Natural emergent misalignment from reward hacking \(Anthropic PDF\)](#)

Ch. 12 (Extensibility): [Hooks \(Claude Code\)](#) · [Discover/install plugins](#) · [Plugin marketplaces](#) · [Customize with plugins \(announcement\)](#) · [Skills / custom commands](#) · [Settings \(precedence/managed\)](#) · [Advanced tool use \(late loading\)](#) · [AGENTS.md \(open standard\)](#) · [Codex config](#) · [Copilot Extensions \(GitHub\)](#)

Ch. 13 (Interfaces): [Platforms and integrations \(Claude Code\)](#) · [Claude Code on the web](#) · [Headless](#) · [Agent SDK overview](#) · [Managed Agents](#) · [VS Code](#) · [JetBrains](#) · [Copilot agent mode \(VS Code\)](#) · [Permission modes](#) · [Streaming output \(SDK\)](#) · [AskUserQuestion / user input \(SDK\)](#) · [Agent Inbox \(LangChain\)](#) · [Channels](#) · [Slack](#)

Ch. 07 (Security): [Claude Code sandboxing](#) · [How we contain Claude](#) · [Agent approvals & security \(Codex\)](#) · [Agents Rule of Two \(Meta\)](#) · [The lethal trifecta \(Willison\)](#) · [New prompt injection papers](#) · [OpenClaw attacks \(The Hacker News\)](#)

Pedagogy (grounds the book's method, not its content)

- ✓ **Blueprints for complex learning: The 4C/ID-model** (van Merriënboer et al.) — [ETR&D](#).
- ✓ **Cognitive Architecture and Instructional Design: 20 Years Later** (Sweller, van Merriënboer & Paas, 2019) — [EPR](#).
- ✓ **van Merriënboer & Kirschner (2018) *Ten Steps to Complex Learning***, 3rd ed., Routledge (ISBN 9781138080805). • ✓ **Wiggins & McTighe (2005) *Understanding by Design***, expanded 2nd ed., ASCD (ISBN 9781416600350). • ✓ **Diátaxis** (Procida) — [diataxis.fr](#).

Guide — Writing methodologies (Editorial Guide §6 survey)

Sources of the study on editorial and academic writing processes/methodologies (Editorial Guide §6). All verified by cross-search (spec 050 review). Feature [010-estudo-metodologias-escrita](#).

Traditional:

- ✓ **The IMRAD Structure: A Fifty-Year Survey** (Sollaci & Pereira, 2004) — *J. Med. Libr. Assoc.* 92(3):364–371, [PMC442179](#).
- ✓ **The Science of Scientific Writing** (Gopen & Swan, 1990) — *American Scientist* 78(6):550–558, [JSTOR 29774235](#).
- ✓ **How to Write and Publish a Scientific Paper** (Day & Gastel) — 7th ed. Cambridge, ISBN 9781107670747.
- ✓ **A Cognitive Process Theory of Writing** (Flower & Hayes, 1981) — *CCC* 32(4):365–387, [DOI 10.58680/ccc198115885](#).
- ✓ **Revision Strategies of Student Writers and Experienced Adult Writers** (Sommers, 1980) — *CCC* 31(4):378–388, [DOI 10.2307/356588](#).
- ✓ **The Elements of Style** (Strunk & White, 4th ed. 2000) — ISBN 9780205309023 • **Style: Toward Clarity and Grace** (Williams, 1990) — ISBN 9780226899152 • **On Writing Well** (Zinsser, 2006) — ISBN 9780060891541.
- ✓ **The Chicago Manual of Style** (17th ed., 2017) — ISBN 9780226287058 • **APA Publication Manual** (7th ed., 2020) — ISBN 9781433832161.
- ✓ **The Craft of Research** (Booth, Colomb, Williams et al., 4th ed. 2016) — ISBN 9780226239736 • **The Uses of Argument** (Toulmin, 1958) — Cambridge University Press.
- ✓ **The history of the peer-review process** (Spier, 2002) — *Trends in Biotechnology* 20(8):357–358, [DOI 10.1016/S0167-7799\(02\)01985-6](#).
- ✓ **Peer Review** (Melinda Baldwin) — Encyclopedia of the History of Science (CMU ETHOS, ed. Christopher Phillips), [entry](#) (entry with no stated year). • ✓ **Developmental Editing** (Scott Norton) — Univ. of Chicago Press, 1st ed. 2009, ISBN 9780226595146 (a 2nd ed. exists, 2023, ISBN 9780226793634).
- (Pedagogy — see the section above: Backward Design; 4C/ID; Sweller; [Diátaxis](#).)

AI-era:

- ✓ **CoAuthor** (Lee, Liang, Yang, 2022) — CHI '22, [DOI 10.1145/3491102.3502030](#); [arXiv 2201.06796](#).
- ✓ **Wordcraft** (Yuan, Coenen, Reif, Ippolito, 2022) — IUI '22, [DOI 10.1145/3490099.3511105](#); [arXiv 2107.07430](#).
- ✓ **Co-Writing with Opinionated Language Models Affects Users' Views** (Jakesch et al., 2023) — CHI '23, [DOI 10.1145/3544548.3581196](#); [arXiv 2302.00560](#).
- ✓ **Spec Kit** ([github.com/github/spec-kit](#)); **Kiro** ([kiro.dev](#)); **Structured Authoring in Docs-as-Code** (SIGDOC '24) — [DOI 10.1145/3641237.3691677](#); **DITA** ([dita-lang.org](#)).
- ✓ **RAG** (Lewis et al., 2020) — [arXiv 2005.11401](#).
- ✓ **RARR** (Gao et al., 2023) — ACL '23, [arXiv 2210.08726](#) • **Evaluating Verifiability in Generative Search Engines** (Liu, Zhang, Liang, 2023) — [arXiv 2304.09848](#) • **A Watermark for LLMs** (Kirchenbauer et al., 2023) — [arXiv 2301.10226](#).
- ✓ **ICMJE (AI use by authors); COPE — Authorship and AI Tools** (2023) ([position](#)); **Thorp, "ChatGPT is fun, but not an author"** (*Science*, 2023) — [DOI 10.1126/science.adg7879](#); **Nature editorial** (2023) — [d41586-023-00191-1](#).
- ✓ **Fabrication and errors in the bibliographic citations generated by ChatGPT** (Walters & Wilder, 2023) — *Scientific Reports* 13, [DOI 10.1038/s41598-023-41032-5](#).
- ✓ **Your Brain on ChatGPT** (Kosmyna et al., 2025) — [arXiv 2506.08872](#) • **Homogenization Effects of LLMs on Human Creative Ideation** (2024) — [arXiv 2402.01536](#) • **Academ-AI** (2024) — [arXiv 2411.15218](#).

- ✓ *Agentic AutoSurvey: Let LLMs Survey LLMs* (Liu et al., 2025) — [arXiv 2509.18661](#) (note: this is the "Agentic AutoSurvey"; the original AutoSurvey is distinct earlier work). · ✓ **Defeating Nondeterminism in LLM Inference** (Thinking Machines Lab, Sep 2025) — [industry blog](#), [non-academic](#).

Editorial Guide

Editorial Guide — the book's operating rules

The operational version of the pedagogical guidance. The full report (with rationale) is in [estudos/2026-07-25-parecer-editorial-plano-pedagogico.md](#) (in Portuguese). This guide is what you consult **while writing**.

1. The pedagogical framework in four lines

Framework	What it dictates in the book
Backward Design	Every chapter is designed backwards: objectives → evidence (check/practice) → only then the content
4C/ID	harness-zero steps = whole tasks; chapters = supportive information; boxes in the code = just-in-time; katas = part-task practice
Diátaxis	Four types of text, never mixed in the same section: chapter=explanation, harness-zero=tutorial, templates/benchmark=reference, "what to steal"=how-to
Cognitive Load	Worked examples before exercises; exercises are "complete", not "create from scratch"; scaffolding decreases step by step; one new idea at a time

2. Chapter skeleton v3 (mandatory; pilot: ch. 04)

Editing rule (v3): when opening each topic, also seek **commercial/industrial material** (official vendor docs, engineering blogs, practitioner posts) in addition to the scientific. The base source remains **the code of the repositories**. The chapter body receives **the state of the art** (what is most modern, synthesized from all benchmark rounds + industry); the detailed per-repository treatment **goes to the file's Appendix** — which lives in the online version as complementary material and is updated every round.

1. **Objectives** — 3–5, Bloom verbs (explain, compare, implement, evaluate)
2. **The problem** — why the dimension exists
3. **Scientific foundations** — 2–4 papers *translated into decisions*; pointer to `bibliografia.md`
4. **Industry sources** — relevant vendor docs and engineering posts, with the same translation rule ("the vendor recommends X because Y")
5. **The state of the art** — the main body: consolidated patterns + what is most modern, citing repositories only as named examples (the detail lives in the appendix)
6. **Hands-on** — the corresponding harness-zero step
7. **Synthesis + "what to steal"** — executive summary and exportable ideas
8. **Check your understanding** — 2–3 questions that test exactly the objectives of item 1
9. **Appendix A — How each repository handles it** — the per-harness evidence with paths, expanded at every benchmark round (online complementary material)

2.1 Living book: dating and history (mandatory)

This is a **living book** — coherent with its own thesis (the expiration clause: what we describe is temporary). Three rules:

1. Every v3 chapter declares the capture date in its header: `> **State of the art captured in AAAA-MM** · last revised AAAA-MM-DD · [history and expiration log](../historico.html)`. This tells the reader whether the "State of the art" section is fresh — something the *event* date (in the body) does not.
2. **Distinguish three dates** (see `HISTORICO.md`): event date (in the body — historical fact, immutable), capture date (in the header — when we took the snapshot), benchmark round (in the evaluations — the version of each repo's snapshot). Re-evaluating = a new round, never overwriting.

3. **Every edition updates `Livro/HISTORICO.md`**: the edition changelog, the per-chapter snapshot table, and — most importantly — the **expiration log** (the prediction scoreboard: each expiration clause scored  against reality, with dated evidence). A line that changes state is the most important news of a new edition.

Associated writing rule: when a statement is time-sensitive ("today", "not yet", "the 2026 consensus"), it is implicitly under the header's capture date — no need to date every sentence, but avoid timeless absolutes ("never", "always") unless they are of the non-expiring kind (the boundary with the world).

3. Permanent writing rules

- **Evidence by file path** for any claim about a harness; ✓ **status** for any scientific citation (the `academic-research` skill has the flow).
- Scores 0–3 only compare within the same benchmark category.
- Every component described should, when possible, declare its **expiration clause**.
- Prose in Portuguese; established technical terms (harness, loop, tool, prompt) **untranslated**.
- Tables for enumerable facts; explanation lives in the prose, not in the cells.

4. harness-zero rules (the report's 4 conditions)

1. **Lightweight DDD** — ubiquitous language = the book's glossary; tactical patterns only where they pay off; DDD appears as a named consequence in the code.
2. **Architecture through refactoring** — each port is born from the pain of the corresponding chapter; never anticipated structure.
3. **Anti-rot** — the model behind an `LLMPort`; self-contained, runnable steps; deliberate didactic mistakes are **commented as such** in the code.
4. **Frozen chat** — HTML+JS served by the backend; it only evolves when a dimension demands a new surface.

5. Repository tooling

- **spec-kit** (`.specify/` + `/speckit-*` commands): for new harness-zero features or large book sections, the flow is `/speckit-specify` → `/speckit-plan` → `/speckit-tasks` → `/speckit-implement` (with `/speckit-clarify` before the plan when the request is ambiguous). The project's constitution lives in `.specify/memory/`.
- **academic-research skill** (`.claude/skills/`): the locate → validate → record → integrate flow for scientific references.
- **scripts/sync-forks.ps1**: local synchronization of the forks with their upstreams.

6. Study: editorial and academic writing processes and methodologies (traditional and AI-era)

Updated in 2026-07 · living book (the AI practices have an expiration date). Sources in the "Guide — Writing methodologies" section of `bibliografia.md`.

A book about engineering — the discipline of instrumenting a process well — needs to expose its own production process, or it contradicts what it teaches. This section is a *survey* of editorial and academic writing methodologies (the established ones and those of the AI era) and, at the end, makes explicit and dated the method with which this book is written. It is **reference/explanation** text (Diátaxis), not a chapter — which is why it does not follow the v3 skeleton.

6.A — Traditional methodologies

The structure of scientific writing. IMRaD (Introduction, Methods, Results, Discussion) was not invented by an author: [Sollaci & Pereira \(2004\)](#) show that it was "imposed by decantation", becoming the standard in the 1980s. The classic [Gopen & Swan, "The Science of Scientific Writing" \(1990\)](#) establishes the principle of *reader expectation* — meaning is born from structural position (topic/stress positions), not just from words. The practical codification is in *How to Write and Publish a Scientific Paper* (Day & Gastel).

Writing as a cognitive process. Flower & Hayes (1981) model writing as **recursive** processes (planning/translating/reviewing) guided by goals, not linear stages; Sommers (1980) shows that experienced writers revise by *re-seeing the meaning*, while novices swap words at the surface — "writing is rewriting".

Craft and style. The tradition runs from the prescriptive minimalism of *The Elements of Style* (Strunk & White) to the *principled* theory of clarity of *Style: Toward Clarity and Grace* (Williams — characters=subjects, actions=verbs, old-before-new), through the authentic voice of *On Writing Well* (Zinsser); the editorial/citation standards are the *Chicago Manual of Style* (17th ed.) and the *APA Publication Manual* (7th ed.).

Craft of research and argument. *The Craft of Research* (Booth, Colomb & Williams) frames research as **making an argument to a reader** (problem → question → *claim* → reasons → evidence; the "So what?"); Toulmin's model (1958) gives the anatomy of the argument (claim, grounds, warrant, backing, qualifier, rebuttal).

Peer review and the editorial flow. Spier (2002) traces the history of peer review; the historiography (Baldwin, *ETHOS*) reminds us that universal refereeing is a 20th-century construct. And the editorial division of labor — *developmental editing* (restructuring vision/discourse) × *copyediting* (sentence-level preparation) — is the axis of the flow (Norton, *Developmental Editing*).

Instructional design (what this book already uses). Backward Design (Wiggins & McTighe), 4C/ID (van Merriënboer et al., 2002), cognitive load (Sweller, 1988) and Diátaxis (Procida) — the pedagogical basis of Principle III.

6.B — AI-era methodologies

Human-AI co-writing. HCI studies treat co-writing as **observable** interaction, not a black box: CoAuthor (Lee, Liang, Yang, 2022) records the interaction at keystroke level; Wordcraft (Yuan et al., 2022) decomposes writing into *moves* (continue/infill/elaborate/rewrite) tied to intention. Measured **caution**: Jakesch et al. (2023) show that a biased assistant shifts what the user writes *and thinks* ("latent persuasion").

Spec-driven / structured authoring / docs-as-code. Write the intention first and let it drive the generation: GitHub Spec Kit (spec → plan → tasks → implement) and Amazon Kiro formalize this; the documentation community already adopts an engineering workflow for prose (docs-as-code, SIGDOC '24; DITA/topic-based).

Agent-augmented research and retrieval. RAG (Lewis et al., 2020) anchors generation in retrieved sources instead of the model's memory; the agentic frontier decomposes the *survey* into roles (search/synthesize/verify) — a trend illustrated by auto-survey work (🕒 to be confirmed).

Verification and provenance. RARR (Gao et al., 2023) does attribution/checking *after* generation; Liu, Zhang & Liang (2023) measure that only 51.5% of generative search engines' claims are fully supported by citation — judge by *citation precision/recall*; watermarking (Kirchenbauer et al., 2023) embeds provenance (fragile to paraphrase).

Academic integrity and authorship. The policy consensus: **an LLM cannot be an author** (it cannot answer for the content) and its use must be **disclosed** — ICMJE, COPE (2023), *Science* (Thorp, 2023), *Nature* (2023). And disclosure, in practice, is **widely violated** (Academ-AI, 2024).

6.C — Tensions and synthesis (traditional × AI)

The gain of AI assistance (speed, research reach, structure) comes with four tensions that an academic edition cannot ignore:

- **Fabricated sources.** Walters & Wilder (2023) measured 55% fabricated citations in GPT-3.5 (18% in GPT-4) and substantive errors in the real ones — hence this book's rule: **verify every reference** against the primary source, by cross-search.
- **Verifiability.** Text that *looks* cited frequently is not supported (Liu et al.'s 51.5%) — the citation must be checked, not trusted.
- **Reproducibility.** LLM outputs are non-deterministic; logging prompt, model version and context is part of the rigor.
- **Homogenization and "cognitive debt".** AI converges style and ideas (homogenization, 2024) and uncritical use is associated with lower engagement/ownership (Kosmyna et al., 2025) — the reason for AI to *amplify*, not *replace*, the author's judgment.

The book's synthesis: use AI as a **research and structuring prosthesis under human verification**, not as an author. Traditional methodologies (argument, clarity, revision) remain the quality standard; the AI ones accelerate the path to it, as long as they are fenced by verification.

6.D — This book's method, declared

This book practices what it describes. Each practice links to a principle of the constitution and has evidence in the repository itself:

- **Evidence over rhetoric** (Princ. I) — no claim about a harness without a *path* in the code; no citation without validated status. Sources verified by cross-search; gaps recorded, not filled with weak sources.
- **The base source is the code** (Princ. II) — the body is born from reading the harnesses' code; science and industry provide context. The per-repository treatment (with paths) is each chapter's **Appendix A**.
- **A combined pedagogical method** (Princ. III) — Backward Design + 4C/ID + Diátaxis + cognitive load; the v3 skeleton is its materialization.
- **Verified double research** — when opening each topic, parallel research agents gather **scientific** and **industry** material; each source is confirmed by ≥2 independent mentions before it enters (the rule the AI era makes at once possible and mandatory, in light of Walters & Wilder).
- **Spec-driven cycle** (Princ. VII) — every improvement goes through `spec → plan → tasks → implement` (spec-kit), on its own branch; *this section* was produced that way (`specs/010-estudo-metodologias-escrita/`), with the official cycle and its gates (Constitution Check, cross-artifact analysis).
- **Living book** (Princ. IV) — dating and `HISTORICO.md`; the predictions have a scoreboard (the expiration log).

Authorship disclosure (transparency). Consistent with the policies above and with Principle I, we openly declare: this book is **co-written with an AI agent (Claude Code, from Anthropic)** operating under **human authorship, curation and responsibility**. The agent executes research, drafting and the spec-kit cycle; the human author defines the scope, decides (via `/speckit-clarify` and review), verifies the sources and answers for the content. Following [ICMJE/COPE/Nature/Science](#), the AI is **not** listed as an author — it cannot be responsible — and its use is disclosed here, in the method.

6.E — A repeatable flow for a contributor

To bring a chapter or section up to the book's standard:

1. **Open the topic** — double research (commercial/industrial + scientific), verified by cross-search; record gaps.
2. **Gather the base source** — read the harnesses' code; note paths (becomes Appendix A).
3. **Write** — in the v3 skeleton (chapters) or in the correct Diátaxis type (guide/benchmark = reference); one type of text per section; technical terms untranslated.
4. **Revise (developmental)** — re-see structure and meaning before the surface copyedit: does the argument close? does the order serve the reader? is there redundancy or a gap? "Writing is rewriting" (§6.A; the constitution's quality gate).
5. **Verify sources** — no invented URL/ID; unconfirmed marked 🟡; sync `bibliografia.md`.
6. **Build gate** — `node publicar/build.mjs` green (no broken internal links).
7. **Date it** — the capture stamp on the chapter and an entry in `HISTORICO.md` — **with the version of the AI model used** — if the state of the art changed.

AI-use safeguards: the AI researches and drafts; the human decides, verifies and signs. Every source brought by an agent is checked before it enters the body.

Acronyms and glossary (policy)

- **Every technical acronym is spelled out at its 1st occurrence** in a chapter — "Model Context Protocol (MCP)" — and, from then on, the text may use only the acronym.
- The publishing engine reinforces this: it **automatically wraps every known acronym in <abbr>**, so that hovering reveals the meaning at any occurrence without polluting the source text. The acronym map lives in `publicar/build.mjs` and is mirrored on the [Glossary](#) page (`livro/glossario.md`).
- The **Glossary** gives the **spelled-out form**, a short explanation and **in which chapters** each acronym appears. When introducing a new acronym, add it in both places (engine map + glossary) and **check the expansion against the source** (Principle I).

Living-book cadence

Policy decided in [ADR 0007](#) (in Portuguese; alternatives and rationale there).

- **Quarterly window** (next: **2026-10**): re-sync of the 16 forks (`scripts/sync-forks.ps1`), a diff driven by the benchmark dimensions, update of the affected Appendices A, of the expiration scoreboard and of the revision dates; a minor edition in the [History](#).
- **Extraordinary trigger**: any event that **invalidates an "Executive summary"** (a protocol change, a capability migrating to the provider, a corpus harness being archived) triggers a targeted revision of the affected chapter, without waiting for the window.
- Each chapter's "state of the art captured in" date remains the truth exposed to the reader — the cadence exists so that it never lies by omission.

About

About the author

About the author



Gilsley Henrique Darú is the human editor, director and orchestrator of this living book. A data scientist, engineer and university professor, he has worked for more than 20 years at the frontier between **optimization, operations research and artificial intelligence** applied to planning and the supply chain — and, for the past few years, on the engineering of the AI systems this book calls a *harness*.

The choice of topic is not accidental. The author's trajectory is a long practice of **wrapping powerful methods in a scaffold that makes them useful in the real world**: mathematical solvers inside production planning and control routines, predictive models inside business processes, and now AI agents inside a *scaffolding* of loop, tools, memory and verification. This book is the systematization of that discipline.

Profiles: [Currículo Lattes](#) · [ORCID 0000-0002-8979-0461](#) · [LinkedIn](#) · Contact: ghdaru@gmail.com

In one sentence

Head of Data & AI, consultant in enterprise optimization, AI and supply chain — applying the **Theory of Constraints**, operations research, systems thinking and AI agents to improve flows and find focus. Based in Joinville, Santa Catarina, Brazil.

Academic background

- **PhD in Numerical Methods in Engineering (Computational Mathematics)** — Universidade Federal do Paraná (UFPR), in progress (since 2019). Thesis in development: *Framework Humano-Computacional para Saneamento Acelerado de Dados*.

- **Professional Master's in Mathematics, Statistics and Computing Applied to Industry (Data Science)** — Universidade de São Paulo (USP), 2021–2024. Dissertation: *Categorização de produtos em e-commerce: avaliação do método Argmax para classificação de descrições curtas em português*.
- **Master's in Numerical Methods in Engineering** — UFPR, 2003–2005. Dissertation: *Uma heurística para o sequenciamento da produção baseada na Teoria das Restrições* (mathematical programming, optimization, scheduling and heuristics).
- **Specialization in Data Science (SPRINT)** — Business Intelligence, Artificial Intelligence and Big Data — SENAI/DR-SC, Florianópolis, 2018–2019.
- **Specialization in Software Engineering with emphasis on Object Orientation** — Pontifícia Universidade Católica do Paraná (PUC-PR), 2004.
- **Mechanical Engineering** — Universidade do Estado de Santa Catarina (UDESC).
- **Associate degree in Data Processing** — UDESC.

Professional experience

Neogrid (2016–present) — a software company that is a reference in supply-chain integration. A trajectory of growing responsibility, from NGLabs to leading data and AI:

- **Head of Data & AI** (since Dec 2025) — data engineering (capture, ingestion, transformation, delivery), governance (master data, reference data, metadata, quality), data science and analytics (descriptive, diagnostic, predictive and prescriptive analysis) and steering the AI initiatives, including the construction of a contextual and agentic hub.
- **Executive Innovation Manager** (Jul–Dec 2025).
- **Manager of the Innovation and AI Lab in Supply Chain Management** (2023–2025) — solution prototyping, data-driven experimentation and developing the team in advanced analytics.
- **Specialist / Data Scientist and Data Quality Coordinator** (2016–2023), at NGLabs — collecting, cleaning, transforming, mining and communicating knowledge from data.

As an **optimization and supply-chain consultant**, he served large operations — retail (HAVAN), railway optimization in Germany (HVLE) and production planning in fashion (Malwee) — applying the Theory of Constraints, operations research and discrete and continuous simulation.

Earlier industry experience (more than 20 years in total):

- **Grupo Malwee** (2014–2016) — Production Planning and Control Manager and Executive Technical Coordinator: improving planning capability, inventory reduction, safety-stock calculation and categorization, master plan and support for the SAP migration (PP/MRP module).
- **WEG Automação** (2006–2014) — Production Planning and Control Coordinator and Systems Analyst: S&OP implementation, inventory policies with OTIF gains, improved sales forecasting, master plan, performance indicators and SAP migration (with an APO project).
- **Datasul** (2000–2005) — Systems and Business Analyst: developing applications with operations research, constraint programming, genetic algorithms and data mining (coil cutting, workforce allocation, petrochemicals).

Teaching

A university professor for more than two decades, including program coordination and graduate teaching:

- **Universidade do Estado de Santa Catarina (UDESC)** — University professor (since 2017): Production Management, Business Management, Entrepreneurship and Operations Research.
- **Centro Universitário — Católica de Santa Catarina** — Professor in the **Graduate Program in AI & Deep Learning** (2024–2025); Operations Research I and II in the Production Engineering undergraduate program.
- **INESA — Instituto de Ensino Superior Santo Antônio** — Professor (2017–2022): Statistics and Applied Mathematics; took part in creating the pedagogical program and in the approval of the Production Engineering program.
- **SOCIESC** — Tenured professor (2011–2019): Formal Languages and Automata, Algorithm Analysis, Compilers, Theory of Computation, Calculus III, Data Structures.

- **Faculdade Metropolitana de Guaramirim (FAMEG) — Coordinator of the Production Engineering program** (2011–2014) and professor (2009–2011): Operations Research, Multi-criteria Analysis, Modeling and Simulation, Statistics, Calculus. Work on ENADE and on program authorization and accreditation.
- Also taught at FACASC, Centro Universitário Católica de Santa Catarina (Jaraguá do Sul), Faculdade Cenecista de Joinville and Associação Catarinense de Ensino.

Academic production

Full journal articles

- Bianchini, J.; **Darú, G. H.**; Berger, S. *Analysis of production planning and control based on the simulation of a production process using a hybrid MRP and Kanban model*. **Journal of Lean Systems**, v. 3, n. 1, 2018. [\[article\]](#)
- Bratti, R. T.; Coelho, E. T. B.; **Darú, G. H.**; Decker, S. L.; Steffen, D. K. *A influência da lei 9870/99 sobre a inadimplência no setor educacional em uma instituição de ensino de Santa Catarina*. **Paidós**, v. 16, 2019.
- Bratti, R. T.; Coelho, E. T. B.; **Darú, G. H.**; Decker, S. L.; Reis, E. A. *A Teoria do Crescimento da Firma e Fusões e Aquisições: evidências da inter-relação das teorias*. **Paidós**, v. 16, 2019.
- Bratti, R. T.; **Darú, G. H.**; Machado, C. S. G.; Pavanati, I.; Reis, E. C. S. *O ensino da matemática através de jogos na perspectiva da neurociência com alunos do 3º ano do Ensino Fundamental*. **Paidós**, v. 16, 2019.

Conference proceedings

- **Darú, G. H.**; Pimentel, R.; Maldonado, M. U. *Manufactured fabrics obsolescence risk evaluation of alternative policies in a textile industry*. XIV CLADS — Congresso Latino-Americano de Dinâmica de Sistemas, São Paulo, 2016.

Completed advising (selection)

Final-year projects in Production Engineering and Computing, among them: applying **genetic algorithms** to bar-cutting production; a comparison between **semantic search and textual search**; evaluating production scenarios through **simulation**; **multi-criteria decision-support** models (PROMÉTHÉE) for supplier selection; and decision-support systems with **constraint programming**.

How to cite this book

This work has a **DOI** (Zenodo/DataCite) and is versioned by edition:

Darú, Gilsilei Henrique. *Engenharia de Harness — Um livro vivo sobre o scaffolding que envolve agentes de IA*. 2026. DOI: [10.5281/zenodo.21632412](https://doi.org/10.5281/zenodo.21632412)

The identifier accompanies the living work; each edition also receives its own version DOI. The human + AI co-authorship is declared in the [note on authorship and method](#).

Profiles and contact

- **Currículo Lattes**: <https://lattes.cnpq.br/6253911800847523>
- **ORCID**: <https://orcid.org/0000-0002-8979-0461>
- **LinkedIn**: <https://www.linkedin.com/in/gilsilei-darú/>
- **E-mail**: ghdaru@gmail.com

This page is part of the book's apparatus (back matter). The facts come from the Currículo Lattes, from the public professional profile and from verifiable sources; institutions and companies are cited as trajectory, not as endorsement. On the human + AI co-authorship of this work, see the [note on authorship and method](#) in the introduction.