



Maturity Self-Assessment & Roadmap Builder

What this is. The scored diagnostic that places your knowledge program on the six-level road from institutional knowledge to agentic loops — dimension by dimension, on evidence — and turns the placement into a one-page roadmap your VP can sign.

When you need it. First. Placement precedes prescription: every other module's routing (Field Guide §7) assumes you know your level and your weakest dimension. Re-run it quarterly and after each implementation phase — the re-score *is* your progress report.

Which maturity level it serves. All of them; this is the instrument that tells you which.

1. How to run it (45 minutes, three people)

1. **Assemble three perspectives:** the support leader, the person closest to the knowledge base (knowledge lead or senior agent), and one frontline agent. One perspective alone reliably overscores — leaders score the program they designed, not the one that runs.
2. **Work through** `assessment.md` : six dimensions, five gate questions each. Every question is answered **YES only on evidence** — if you'd have to say "mostly" or "we're supposed to," the answer is NO. Ties break downward, always. Flattering yourself here costs real money later: every module you buy or build above your true level is shelf-ware.
3. **Score:** a dimension's level = the number of *consecutive* YES answers from its first question. (YES to Q1, Q2, NO to Q3 = Level 2 — even if Q4 is YES. A skipped gate is a gap, and gaps, not averages, are what break AI-on-top-of-knowledge projects.)
4. **Your organization's level = your LOWEST dimension score.** That's the placement rule, and it's the whole maturity model in one sentence: you cannot skip levels, and your weakest dimension is where the next confident wrong answer will come from.
5. **Read your level** in `levels-and-next-moves.md` — entry criteria, the failure modes characteristic of your level, the 3–5 KPIs worth watching there, and the specific next moves with the module that executes each.
6. **Fill** `roadmap-one-page.r.md` and put it in front of the person who funds the program.

2. Interpreting common shapes

- **Flat profile** (all dimensions equal): rare and healthy — climb evenly, following the Field Guide's routing for your level.
- **Delivery ahead of governance** (e.g., Delivery 3, Governance 1): the most dangerous shape and the most common one — a bot is answering from an ungoverned corpus. Stop expanding delivery; fix

Governance and System of Record first (Modules 5 and 2). This shape is the "expensive hallucinations with your logo" configuration.

- **Capture ahead of curation:** candidates pile up or flow ungated. Install Module 4 before generating another draft.
- **Telemetry last** (very common): you're flying the program on anecdotes. Module 6's Baseline Snapshot is a one-day fix that starts the clock.

3. Honesty mechanics

Two rules keep the instrument useful over time: score the same way each quarter (same three roles present, same evidence bar), and **never edit a past score** — the trend line is the asset, and it's only an asset if it's true. Keep completed scoring grids with the date; the workbook's Baseline Snapshot (Module 6) records the first one.

What this module assumes from its siblings: nothing — it's deliberately runnable before you own anything else. Its outputs feed the Field Guide's routing, Module 6's Baseline Snapshot, and Module 7's pilot charter.

Contextkeeping is a KCS-aligned operating model. KCS® is a service mark of the Consortium for Service Innovation; this material is not affiliated with or endorsed by the Consortium.