

Competitive Differentiation

Questions this doc answers

- How should Reflect be categorized vs Mem0, Supermemory, Notion AI, or built-in ChatGPT/Claude memory?
- Why do we call this infrastructure rather than another AI app?

Category, not feature checklist

Reflect Memory is an **infrastructure layer below AI tools**, not a chat app and not a memory feature locked inside one vendor.

Buyers typically compare us to:

- **Built-in vendor memory** (ChatGPT / Claude memory) — convenient inside one product; does not travel when the team switches tools.
- **App-layer memory products** — wrappers or assistants that sit on top of a single workflow; usually weaker on private deploy, audit, and multi-tool ownership.

Reflect's pitch for diligence:

Give every AI tool one provider-neutral organizational memory that can live in your network, with explicit visibility and audit controls.

What matters to regulated buyers

- Private / isolated / self-host options
- Ambient Mode that is toggleable and audited (personal-only until shared)
- Vendor-neutral access via REST + MCP
- Docs and prompts designed for async AI review before a live call

We do **not** publish a public "moat teardown," algorithm inventory, or competitor kill-sheet. Side-by-side technical comparisons for active evaluations are handled under NDA.

How to use this in evaluation

1. Read `/diligence/architecture` and `/diligence/deployment-architecture` for shape and options.
2. Read `/diligence/security-compliance` for control themes.
3. Run the hub prompt at `/diligence` through your own AI.
4. Ask for the NDA pack if you need deeper implementation or competitive Q&A.