



Implementing Your Own Data Concierge

A scoping guide for collapsing fragmented operational reporting into one governed model with an AI agent grounded on top of it — what the engagement actually involves, phase by phase, and a rough read on your own scope before a scoping call.

PART 1

The methodology: 5 phases, in order

Each phase depends on the one before it. Skipping ahead is exactly how these projects usually fail — an AI agent bolted onto ungoverned data, a dashboard with no shared definitions underneath it.

PHASE 1

One Clock, One Shape

Every domain gets unified onto a single model with one time anchor and one calendar. This is the unglamorous foundation: until "last 30 days" means the same thing to every team and every report, nothing built on top of it can be trusted. Legacy reports and tribal knowledge aren't discarded — they're mined as a business-logic reference, then rebuilt cleanly instead of patched.

Delivers: a single governed model every domain reads from — no new platform, no new data feed.

PHASE 2

Cross-Domain Enrichment

This is where the model starts answering questions no single-domain dashboard ever could: linking a change to the incident wave it caused, a recurring problem to the asset that keeps failing, an outage to the business unit it actually hurt. The hard rule: a partial cross-domain link is labeled partial, never presented as complete.

Delivers: the cross-domain view — usually where the first genuinely new insight shows up.

PHASE 3

The Model Explains Itself

A self-documenting data dictionary — plain-English definitions, synonyms, grain warnings — embedded directly into the model, visible on hover to a report author and readable by an AI agent. This is what makes the model survive contact with a large organization.

Delivers: the semantic layer the AI agent will eventually stand on, plus an immediate drop in "which number is right" debates.

PHASE 4

The Insight Engine + Governance Handback

A small set of business rules (a handful of high-signal ones, not dozens) run automatically and surface findings, not just metrics. Ownership of these rules moves to the business teams that own the definitions — not permanently to the technical team that built them — with a full change history behind every edit.

Delivers: proactive findings instead of passive dashboards, and governance that doesn't bottleneck on the build team forever.

PHASE 5

The Agent

Only now does the AI agent get connected — grounded in the semantic layer from Phase 3, so a plain-English question gets a real answer: what the number is, why it's that number, and what caveats apply. An agent connected to raw tables without Phases 1–3 underneath it is a liability, not an accelerant.

Delivers: the "ask it a question" layer — the most visible part of the system, and the part that should ship last.

PART 2

The rollout is staged, not a big-bang

A properly run implementation moves through these stages, each gated on the last one actually working.

STAGE	WHAT HAPPENS	WHY IT’S SEPARATE
Build	Phases 1–5 built and validated against a small, bounded scope first	Prove the model is right before anyone depends on it
Validation & Hardening	Every output reconciled to source; becomes a permanent, ongoing process	New sources/metrics get baselined and proven before they're trusted, forever
Closed Beta	First reports migrate for a small group of real users	Proves the migration path itself; nothing legacy retires until its replacement is validated
General Availability	Opens to the full organization; disagreeing legacy reports retire	Adoption is actively driven, not just switched on and left
Agent, validated release	The AI layer gets its own separate validation track	Speed of a wrong answer is worse than no answer — deliberately the last gate

The point of laying it out this way: this is not a rip-and-replace project. Nothing is retired until its replacement has already earned trust — that's what makes it fundable in stages rather than needing one big budget approval up front.

PART 3

Scope yourself: what makes this bigger or smaller

Before a scoping call, a rough self-assessment. None of these need precise answers — directional is enough.

1. How many source domains/systems need to agree on one set of definitions?

One system, one team → smallest scope (mostly Phase 1). 3–6 domains each with their own reporting today → the typical case this methodology was built for. 6+ domains, or domains with no shared common key → largest scope, since Phase 2's linkage work grows with the number of pairs that need to connect.

2. Does more than one team currently disagree about the same metric?

If yes — that's the actual trigger for this methodology, and a strong signal Phase 1 alone will deliver visible value fast.

3. Do you need the cross-domain view, or is one unified dashboard per domain enough?

If leadership's real pain is "the expensive problems live between our systems and nobody can see them," Phase 2 is not optional — it's the point.

4. What's your current governance maturity?

No existing data dictionary or shared definitions anywhere → Phase 3 is a bigger, more foundational lift, but also where the most goodwill gets built fastest. Some documentation exists already → Phase 3 becomes consolidation, not creation.

5. Do you actually want an AI agent, or just the governed reporting?

Phases 1–4 stand alone as a complete, valuable engagement with no agent at all. Phase 5 is additive, not required — and should never be pulled forward ahead of the phases that ground it.

6. Who owns the rules once they're built — IT, or the business?

If the honest answer is "IT will own it forever," Phase 4's governance handback needs to be scoped explicitly as change management, not just as a technical deliverable — usually the harder half of that phase.

What “done right” actually looks like

Three patterns show up reliably when this is built with real discipline, not just clever engineering.

Normalized metrics reveal the truth raw counts hide

A metric that looks like it's getting worse in raw terms can be revealed as actually improving once it's normalized to the right denominator — a model built with this discipline is what stops leadership from acting on the wrong read of a bigger raw number.

The cross-domain gap is usually much bigger than anyone assumes

Structural blind spots that "everyone kind of knew about" tend to be worse in practice than any single team's dashboard could show — quantifying it for the first time is often the single most convincing moment in the whole rollout.

A model built to validate itself catches its own defects

The build process itself — reconciling every changed number back to a known baseline — routinely surfaces real, pre-existing data bugs nobody had found, because nobody had ever validated that rigorously before.

None of this requires new platforms or new data feeds. It requires the discipline to build it once, correctly, on tools already licensed and data already governed.

PART 5

Why hire this out instead of building it internally

This methodology is straightforward to *describe*. It is not straightforward to *execute* without the scar tissue of having done it before — every phase above has a failure mode that looks reasonable in the moment (bolting the agent on early, treating a partial cross-domain link as complete, letting IT own governance forever) and only shows up as a problem months later. That's the actual value of hiring this out: not the framework itself, but the judgment about where the real risk hides in each phase, and the validation discipline to catch it before it ships.

NEXT STEP

Talk through your scope

Bring your rough answers from Part 3 — techstylesolutions.studio/contact