



Context Readiness Report — Meridian Systems

SAMPLE · fictionalized engagement · AnswerEcon

1. Executive summary

Meridian Systems is a Level 1 knowledge organization, held there by a four-way tie: Curation, System of Record, Telemetry, and Governance all stop at their first gate. Capture is ahead of the pack at Level 2 — agents write into the workflow, not around it — and Delivery is further ahead still, at Level 3, because a Zendesk AI answer bot has been live since March, drawing on the same Confluence space every human uses. That gap between Delivery and Governance is the finding that matters most in this report: the bot is already answering customers from a corpus I can confirm is 27 percent machine-ready. It is not a capability problem. It is a sequencing problem, and it is fixable in the same 90 days that fixes the rest of Level 1.

Three findings carry the report. First, curated content still ships with duplicates — I found two independently authored fixes for the same VPN error code, filed eleven months apart, neither agent aware the other existed. Second, applicability constraints live in article prose instead of structured fields, so nothing downstream (human or machine) can filter on them reliably. Third, the organization cannot yet compute a defensible cost per answer — not because the math is hard, but because unit costs are attached to two of six live channels, and neither is tied to a resolution outcome. That is not a failure of this report; it is the report's most useful finding, because it tells you exactly what to fix before you spend money on the next channel.

The economic story in one paragraph: Meridian currently manages support economics by self-serve deflection rate, a real number that answers a different question than the one that matters. Deflection tells you how many contacts you avoided; it says nothing about what each resolution cost, however it was resolved. Reaching Level 2 does not hand you a cost-per-answer figure — reaching Telemetry's third gate does, and that is one rung past where this roadmap takes you. What this roadmap buys you is the corpus and the ownership discipline that make that number possible to trust once you do compute it.

The 90-day recommendation is three moves: ratify content standards and a lifecycle policy including an explicit rule that nothing AI-generated enters the record ungated; run the metadata-extraction typing pass across the full corpus, not just this sample; and install per-domain ownership with an orphan sweep, seeding the review-date ledger as you go. None of these require new software.

What I told Meridian not to do yet, and why, deserves its own paragraph because it is the paragraph most reports skip. No AI-drafted candidate articles — Capture's own next gate — until Curation can actually gate what arrives. No further investment in the Zendesk bot's reach until Governance and System of Record catch up to where Delivery already sits; right now every new surface it touches multiplies an existing exposure rather than fixing one. No cost-per-answer claims to leadership until Telemetry is attributing real unit costs to real outcomes. You cannot skip levels in this model, and the fastest way to lose the money already spent on Level 1 is to build Level 3 features on top of it.

2. How I assessed

I conducted four interviews over the engagement's first week: the Support Director, the Knowledge Base owner, a senior support agent from the endpoint-support queue, and the Support Operations Analyst who owns Meridian's Zendesk reporting. Each was walked through the same six-dimension instrument independently before we compared notes, which is what surfaced the gap between how leadership scores the program and how it runs day to day — the Support Director's initial self-score placed Curation and Governance a full level higher than the evidence supported.

Data received: a full Confluence space export (page list and metadata, snapshot dated August 7, 2026), a 90-day Zendesk ticket export, a two-week sample of 40 Zendesk AI bot transcripts, and a partial KB ownership roster. Not received, and marked unverified in this report: a full Slack workspace export and a Google Drive shared-document inventory — both were referenced anecdotally in interviews (see Finding 1) but were out of scope to confirm directly, so I have not counted them in any corpus figure.

I sampled 120 of the 1,840 published Confluence articles, stratified proportionally across the four domains agents named as load-bearing: endpoint-support, billing, platform-admin, and product-defects. Two assessors independently scored a 15-article calibration subset first (93 percent agreement on machine-ready verdicts before reconciliation), then I scored the remaining 105 alone against the same rubric. Engagement ran July 20 through August 7, 2026, three weeks, with this report delivered August 10.

3. Maturity placement

The placement rule is simple and I am applying it exactly as written: an organization's level equals its weakest dimension, full stop, and ties within a dimension break downward. Averaging is not permitted — a program that is excellent at Delivery and absent at Governance is not a "Level 2 program," it is a Level 1 program with an expensive blind spot.

Dimension	Level	Gates passed	First gate not met
Capture	2	C1, C2	C3 — no AI drafting from resolved cases; capture beyond the required workflow step is still manual and after the fact
Curation	1	U1	U2 — review has no written standard to check against
System of Record	1	S1	S2 — templates are inconsistent across the sample
Delivery	3	D1, D2, D3	D4 — the bot does not cite article IDs, filter by validated status, or refuse below a confidence floor
Telemetry	1	T1	T2 — participation and reuse are not tracked against any standard

Dimension	Level	Gates passed	First gate not met
Governance	1	G1	G2 — content health work is reactive, not cadenced
Organization level (lowest row)	1		
Weakest dimensions	Curation, System of Record, Telemetry, Governance (tied at 1)		

The shape here — Delivery at 3 against Governance at 1 — is the configuration this model flags as the most dangerous and, in practice, the most common: a bot answering customers from an ungoverned corpus. It is not a hypothetical for Meridian.

Finding 1 — Delivery: the answer surface is two rungs ahead of the corpus it draws from.

Observed: Meridian enabled Zendesk's AI answer suggestions in March 2026, searching the full Confluence space — published and unreviewed pages alike — with no citation, no applicability filter, and no confidence floor. Evidence: the Support Director, describing the rollout: "We turned on the Zendesk AI answer suggestions in March. It searches everything in the space, published or not." In a spot-check of 40 bot-suggested transcripts from the trailing two weeks, 11 (28 percent) surfaced an article this audit later classified as Partial or Retire-candidate, and zero of the 40 included a citation a customer could verify. Why it matters: a machine reading undifferentiated content at Delivery's D3 gate has no way to know which 27 percent of what it's reading is trustworthy. This is the single highest-leverage risk in the engagement, and it is already live. Dimension: Delivery.

Curation, System of Record, Telemetry, and Governance each stop at their second gate, and the pattern across all four is the same: something exists (a reviewer, a designated system, a usage report, a named owner) but nothing enforces it against a written standard. That is Level 1 by definition — a KB exists; writing and reviewing are voluntary and after the fact.

4. Corpus audit

Meridian's Confluence space holds 1,840 published articles as of the August 7 snapshot, distributed across four domains: endpoint-support 640 (35 percent), billing 310 (17 percent), platform-admin 275 (15 percent), product-defects 195 (11 percent), with the remaining 420 (23 percent) uncategorized by domain entirely. Only 764 articles (42 percent) carry any owner field, and 612 (33 percent) have not been modified in over 18 months.

I sampled 120 articles, stratified to match domain share: 42 endpoint-support, 20 billing, 18 platform-admin, 13 product-defects, 27 uncategorized. Running the metadata-extraction pass (Module 3) against the sample produced this disposition:

Disposition	Count	Share
Keep as-is (validated, no action)	32	27%
Migrate / restructure	51	42%
Retire-candidate	24	20%
Needs SME review (unclear ownership or accuracy)	13	11%

The 20 percent retire-candidate rate matches what this model expects at Level 1 — it is the migration working, not a surprise. Extrapolated across the full corpus, that is roughly 368 articles to queue for retirement review, not deletion.

The headline number, and the one that gates every AI ambition Meridian has: **27 percent of the sampled corpus is machine-ready today** — meeting all four properties (self-contained, explicit applicability, stable identifiers, freshness signals) at once. Notably, this is the same 32 articles the disposition pass above independently flagged "keep as-is" — two different tests, one on structure and one on machine-readiness, converging on the same set is a good sign for the rubric, not a coincidence I'm choosing to read into it.

Property	Pass rate	Articles passing
Self-contained	66%	79 of 120
Explicit applicability	46%	55 of 120
Stable identifiers	84%	101 of 120
Freshness signals	37%	44 of 120
All four (machine-ready)	27%	32 of 120

Finding 2 — System of Record: applicability lives in prose, not metadata. Observed: constraints that should be structured fields — which platform, which plan, which device-management state — are routinely written as a sentence buried in a Notes section instead. Evidence: a sampled endpoint-support article on a VPN reconnect failure states, in prose, "This fix does not apply to personally-owned (non-MDM) devices" — a real applicability boundary with no corresponding metadata field. Only 55 of 120 sampled articles (46 percent) declare applicability as structured metadata at all. Why it matters: a machine reader cannot infer a constraint you didn't state as data, and it will answer anyway — the exact failure mode this instrument's Delivery dimension is designed to catch before it reaches a customer. Dimension: System of Record.

A dedupe and contradiction spot-check across the sample found 6 duplicate or directly conflicting pairs — 12 articles, 10 percent of the sample — concentrated in endpoint-support, where individual agents maintain informal reference pages alongside the official KB.

Finding 3 — Curation: duplicate, conflicting resolutions survive publication. Observed: two articles in the sample resolve the identical VPN error-720 symptom with different repair steps, filed eleven months apart by different agents. Evidence: the senior agent I interviewed, on discovering the older article mid-session: "I didn't know we already had one — I just wrote what worked for my case." Why it matters: nothing in the current review step checks new content against what already exists. Two agents now give customers two different fixes for the same problem, and neither knows the other is wrong or right. Dimension: Curation.

5. Coverage vs. demand

I clustered the trailing 90 days of resolved Zendesk tickets into themes and checked each against the corpus.

Ticket theme	Monthly volume	Coverage verdict	Candidate type
VPN reconnect loop / error 720 (macOS)	74	Covered, but duplicated	Consolidate
Generic sign-in / password reset	61	Covered	None needed
SAML assertion / SSO clock-skew errors	51	Partial — steps exist, troubleshooting thin	Update
CSV export "wrong" timestamps (timezone)	38	Gap — no covering article	New FAQ
Seat count / rollover / billing cycle	33	Covered	None needed
Mobile offline sync duplicate records	29	Covered, open defect, no lifecycle trigger	Monitor / update on fix

The CSV-timezone gap is worth naming specifically: it is exactly the kind of recurring, low-complexity theme an automated recurrence detector would draft a candidate for at Capture's next gate. Meridian isn't there yet — for now it is a manual backlog item for the KB owner, not a system failure.

Finding 4 — Governance: known-error content has no trigger to retire or update itself. Observed: the mobile offline-sync duplicate-records article references an open engineering defect and a documented manual workaround. Nothing in the current process re-checks the article when the defect's status changes. Evidence: the Ops Analyst, on how staleness surfaces today: "We find out a bug's fixed when a customer complains the workaround broke something, not before." Nine of the 120 sampled articles are known-error type, referencing an open defect ID, with no review-date trigger tied to defect status. Why it matters: this is the content most likely to be confidently cited and wrong — the workaround will still read as current the day after the underlying bug ships a fix. Dimension: Governance.

6. Tooling review

Meridian's stack — Zendesk for ticketing and the AI answer surface, Confluence for the knowledge base — already covers the capability classes this model needs at Level 1 and most of Level 2: a designated system of record, a searchable delivery surface, agent-workflow capture hooks, and a machine consumption point. What it lacks is not a tool but configuration and governance layered on the tools it already owns: no metadata-extraction or candidate pipeline (Module 3), no structured applicability fields on the article template, no dedupe pre-check ahead of publish (Curation's U5 gate), no citation or confidence-floor layer on the bot (Delivery's D4 gate), and no cost attribution spanning channels (Telemetry's T3 gate). None of this requires a vendor change; it requires turning on and standardizing what Zendesk and Confluence already support, mapped through Module 2's template and schema guidance.

7. The number that matters, and the answer-economics baseline

Every Meridian answer economics conversation today runs through one number: self-serve deflection rate. It is a real, correctly-measured figure, and it answers a real question — how many contacts did the KB prevent. It does not answer the question this report is built around: what did each resolution actually cost, however it was resolved. Those are different numbers, and collapsing them prematurely produces a figure nobody can defend under scrutiny.

I inventoried six channels through which a Meridian customer gets an answer: self-serve KB search, agent-assisted ticket resolution, phone, email, the Zendesk AI bot, and a largely dormant community forum. Only two — self-serve KB (session cost from web analytics) and agent-assisted (loaded hourly rate times average handle time) — have any unit-cost figure attached at all, and neither is tied to whether the contact actually resolved the customer's problem.

Finding 5 — Telemetry: cost per answer is not computable today, on any channel.

Observed: two of six live answer channels carry a unit-cost figure; zero of six tie that figure to a resolution outcome. Evidence: the Ops Analyst, asked directly whether Meridian can say what a self-serve resolution costs versus an agent-assisted one: "I can tell you page views. I can't tell you if the view solved the case." Why it matters: a blended cost-per-answer figure is the number that should gate every channel investment Meridian makes next — including whether the Zendesk bot is worth expanding — and right now that decision is being made on deflection rate instead, which cannot answer it.

Dimension: Telemetry.

This is not a shortfall in this report. It is the finding: Meridian cannot yet compute a defensible blended cost per answer, and manufacturing one from the two partial figures available would produce a number that looks precise and isn't. The shift-left headroom — how much of today's agent-assisted volume could move to a cheaper channel without a quality loss — is real but not sizeable responsibly until the underlying unit costs are attributed. I have not estimated it in this report for that reason.

8. The roadmap

Three moves, 90 days, all inside Level 1's own next-move list — nothing here climbs past Level 2.

Move 1 (weeks 1–4): Ratify content standards and a lifecycle policy. Module 5. Owner: KB owner, drafting; Support Director, sign-off. Done-when: a written standard covering the five typed templates, a review cadence, and an explicit clause that nothing AI-generated enters the record ungated is signed and circulated to all 45 agents.

Move 2 (weeks 3–8): Run the metadata-extraction typing pass across the full corpus. Module 3. Owner: Ops Analyst executes; KB owner dispositions. Done-when: all 1,840 articles carry type, status, and owner fields; expect roughly 368 articles (20 percent, per the sample rate) flagged retire-candidate and queued — not deleted.

Move 3 (weeks 6–12): Install per-domain ownership, run an orphan sweep, and start the review-date ledger. Modules 5 and 6. Owner: Support Director assigns one accountable owner per domain (team-endpoint, team-billing, team-platform, team-defects); KB owner runs the sweep. Done-when: 100 percent of validated articles carry a named owner, and the review-date ledger is live, seeded with this audit's 120 articles as its first entries.

Six-month line of sight. Finishing the climb to a clean Level 2 means Curation reaching its formal gate (every candidate gets a disposition, a reason code, and a logged edit-distance, held to an SLA), System of Record completing governance metadata on every article (stable ID, type, owner, lifecycle status, applicability), and Telemetry reaching attributed answer events with real unit costs — the precondition for the cost-per-answer number this report deliberately did not manufacture. As an interim control on the existing exposure, I recommend narrowing the Zendesk bot's source set to validated-status articles only within 30 days — not a fix, but a real reduction in what it can get wrong while Governance catches up.

What not to do yet, and why. No AI-drafted candidate articles — Capture's own next gate — until Curation can actually gate what arrives; publishing drafts faster than you can responsibly review them turns backlog into debt. No further expansion of the Zendesk bot's reach, sources, or hand-off scope until Delivery's D4 gate is met and Governance reaches Level 3 — it is already running two levels ahead of where Governance sits, and every additional surface multiplies that exposure rather than fixing it. No cost-per-answer or ROI claims to leadership until Telemetry reaches Level 3; the deflection number you have today is honest but answers a different question. No Module 7 pilot charter yet — a pilot presumes a receiving gate and an owned baseline that don't exist. You cannot skip levels in this model without paying for it later, and Meridian has already found out what that costs once, with the bot.

Re-assessment date: November 8, 2026 (90 days from this report). Re-run the six-dimension instrument with the same four perspectives; the expectation is Capture and Curation both showing progress, Delivery flat by design, and Governance and System of Record moving off their floor.

Appendices

A. Completed scoring grid (raw). See Section 3 table; full gate-by-gate answers (C1–C5, U1–U5, S1–S5, D1–D5, T1–T5, G1–G5) retained in the engagement workbook, assessed 2026-08-07.

B. Data request list. Received: Confluence space export (2026-08-07 snapshot), 90-day Zendesk ticket export, 40-transcript Zendesk bot sample, partial ownership roster. Not received, marked unverified: Slack workspace export, Google Drive shared-document inventory.

C. Interview guide used. The Module 1 six-dimension gate questionnaire, walked identically with all four roles, followed by role-specific follow-up questions on process reality versus stated policy.

D. Glossary. *Three loops* — capture, curation, delivery, the cycles this model scores. *Gate* — a disposition point that must be passed, not implied, for content to move state. *Edit-distance* — how much a curator changes a draft before publishing it, tracked as a quality signal from Curation Level 2 onward. *Cost per answer* — the unit-economics metric this model treats as the successor to deflection rate, requiring attributed answer events with per-channel unit costs (Telemetry Level 3).

E. Attribution and service-mark notice. Contextkeeping is a KCS-aligned operating model. KCS® is a service mark of the Consortium for Service Innovation; this engagement and its materials are independent work, not affiliated with or endorsed by the Consortium. All dollar figures and channel-cost figures in this report are illustrative placeholders.

This is a sample report. Meridian Systems is a fictional company; all data, quotes, and findings are illustrative. Client reports follow this exact structure with real evidence.