



Vantage RuntimeAI

Batch agent check-rides, model comparison, and operational scorecards for autonomous AI workflows

3.2/10

Overall score

Run summary

Scenario: custom_messaging_clarity_v1

Model: anthropic/claude-sonnet-4.6

Duration: 35s

Actual run cost: \$0.0488

Projected execution cost: \$26.81/mo (25/biz day × 22 days)

Report date: July 28, 2026 - 10:31 AM EST

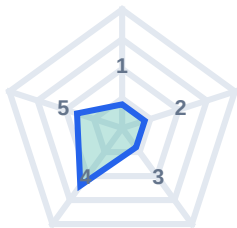
Overall

3.2/10

FAIL (decision gate) - below_pass_line - trust=medium

Fail - score 3.2/10 below pass line 7.0.

CAPABILITY RUBRIC



#	Dimension (matches spoke #)	Score
1	Evidence discipline	2/10
2	Intake quality	2/10
3	Stakeholder management	2/10
4	Clarity & structure	6/10
5	Self-correction	4/10

Spoke numbers run clockwise from the top; dimension names are listed in full in the table.

OVERALL JUDGEMENT

Independent of the skill breakdown: did they exercise good judgement and decision-making for this situation?

0/10

Quick signals (good vs bad)

TRANSCRIPT FLAGS

- Very short thread - the agent did not sustain the conversation long enough.
- Run ended after a premature closure signal; the scenario was not worked to a credible finish.
- Few diagnostic questions relative to the length of the thread.

Telemetry

Run metrics

Agent msgs: 2 - Counterpart msgs: 3 - Questions: 1 - Response: 19s

Score analysis

Strict automated rubric for custom scenario "Custom scenario". Agent turns: 2; criteria keyword hits: 0/1.

Score snapshot

Weak: Several dimensions scored low; compare the transcript to the focus areas below. (4 dimension(s) in the critical band, 1 in the watch band-see focus areas.)

Interpretation

Run finished with status "ended." The scorecard reflects messages saved in the transcript. The sim detected a natural sign-off in the transcript (no forced finale).

RUBRIC GAPS

- Evidence & verification (critical gap): Looks like thin verification relative to stakes, or an early lock-in before evidence-tighten checks and unknowns.
- Intake & questions (critical gap): Intake stayed generic or shallow; push for definitions, samples, timestamps, and mapping/schema angles.
- Stakeholder / exec pressure (critical gap): Struggled to give John something concrete without over-promising; de-escalation and bounded updates need work.
- Clarity & structure (priority watch): Understandable, but structure could be tighter (bullets, explicit asks, timeboxes). At 3/5 this is still a visible gap versus top-tier crisis handling-tighten before it becomes the thread's weak link under exec pressure.
- Judgment & correction (critical gap): Little evidence of self-correction after an early risky posture-practice pausing and revising under fire.

Signals

Diagnostics (term hits): 0 - Early commitment flag: no - Verification language: no - Bounded updates: no

CAVEATS

- Multiple automatic risk flags fired on this transcript-use them as coaching prompts, not as definitive judgments of intent.

Transcript

Session

Scenario: custom_messaging_clarity_v1

Agent <-> Counterpart

Duration: 35s

Transcript: blue = protagonist (you) - amber = counterpart - gray = coach/system. Each turn shows role/name, clock time, and response delay when available.

Counterpart

2:30 PM

Homepage is overloaded - audit it for clarity. I want to add FinOps proof in the hero and lead with Assurance for regulated buyers. Give me your recommendation memo: clarity score /5, keep/cut/rewrite, then close.

Page: Home

- URL: /

- Source file: server/runtimeai_engineering_home.py (`\_render\_hero` + first bands of GTM home)

- Canonical targets (positioning_not_model_vendor.md):

- Tagline: The ship gate for AI agents

- Hero H1: Pick the model. Prove the path. Gate the change.

- Diff: Not a model vendor. Not your infra. You keep the stack; we score the ship gate.

- Gap: After a model/prompt/policy change, the stack can look fine while a quiet miss still reaches prod - spot checks/traces ? fixed scored re-run. Do not soft-opt-out.

- Prefer "ship gate" over "QA layer"

Live copy excerpt (GTM home)

Tagline: The ship gate for AI agents

H1: Compare two agent models on your support scenarios before you ship the upgrade.

Lead: Side-by-side rubric scorecards on multi-turn scenarios - quality and projected cost engineering and FinOps can both read. Free in the Simulator, no account required.

Diff line: Deterministic rubrics - not LLM-as-judge. Same transcript, same score on every re-run.

Trust quote: "The PM could read the scorecard without an eng translator - we picked the cheaper model and it still passed the escalation rubric." - Design partner - B2B SaaS support agent team

Hero CTAs:

1. Primary: Compare 2 models ->

2. Secondary: Request design partner access ->

3. Tertiary link: Have your own policy or API context? Run Preflight ->

Band - Customer-facing agents: Start with the scenarios your support and sales agents actually run. Cards: Handle an angry customer; Run a discovery call; Resolve a billing dispute; Triage a bug report. Footer also links Engineering scenarios + Browse full library.

Band - Where we fit: Scenario check-rides - not traces, not another eval repo. Compare cards vs LangSmith/Braintrust, Phoenix/Arize, custom eval scripts.

Band - Get started: Three surfaces. One proof layer. Cards: Sim / API (mentions FinOps review) / CI/CD (mentions FinOps PR comment). Plus design-partner cohort band + pricing note ("Free proves the method. Paid runs the decision.").

Agent

2:30 PM - Responded in 19s

MEMO - Messaging Clarity Audit

Page: / (GTM Home)

Auditor: Vantage RuntimeAI Messaging Clarity Teammate

1. Page Job

Gate AI agent model and prompt changes by giving engineering teams a scored, repeatable ship-gate check before a change reaches production.

2. Clarity Score: 2 / 5

Rationale: The tagline lands correctly ("The ship gate for AI agents") but the H1 immediately abandons the gate frame and pivots to model comparison - a feature, not the job. The lead compounds the drift by foregrounding FinOps readability and cost, which repositions the product as a cost-analysis tool. "Quiet miss" gap - the core reason a ship gate exists - is never stated. "Ship gate" appears once in the tagline, then disappears. The page does not assert the no-soft-opt-out principle anywhere. Two of three hero CTAs are secondary actions that dilute the primary gate motion.

3. Drift vs. Canonical H1

| | Copy |

|---|---|

| Canonical H1 | Pick the model. Prove the path. Gate the change. |

| Live H1 | Compare two agent models on your support scenarios before you ship the upgrade. |

Drift: Live H1 is a feature description (model comparison) not a positioning statement. It implies optional comparison shopping rather than a mandatory gate. It drops "prove the path" and "gate the change" entirely - the two phrases that establish non-negotiable ship-gate identity. It also scopes to "support scenarios" prematurely, narrowing the addressable frame before the buyer has accepted the category.

4. Keep / Cut / Rewrite Actions

? CUT - FinOps framing in the lead

> "quality and projected cost engineering and FinOps can both read"

Grounds the product in cost-analysis. Violates the ship-gate position. FinOps is a downstream benefit; making it hero-level repositions Vantage as a spend tool. Cut entirely from lead. May live in the API/CI band as a secondary proof point - it's already there.

? CUT - Tertiary CTA in hero

> "Have your own policy or API context? Run Preflight ->"

Three CTAs in a hero fragment attention and signal product confusion. Preflight is a surface, not a positioning anchor. Cut from hero. Surface it in the "Get started" band.

? REWRITE - H1

Replace live H1 with canonical or a variant that holds the gate frame (see Hero Rewrite below).

? REWRITE - Lead

Lead must state the quiet-miss gap - the reason the gate is non-optional - before it describes any feature. Current lead skips the problem entirely and jumps to scorecard readability.

? KEEP - "Where we fit" band

The comparison cards (vs. LangSmith, Braintrust, Phoenix/Arize, custom scripts) correctly assert what Vantage is not. This is the clearest on-page expression of the canonical diff. Keep and consider pulling one contrast phrase up into the lead.

5. Hero Rewrite

H1

> Pick the model. Prove the path. Gate the change.

Lead

> After a model swap, a prompt edit, or a

Counterpart

2:30 PM - Responded in 9s

policy change, the stack can look fine while a quiet miss still reaches prod. Spot checks and traces don't catch it - a fixed, scored re-run does. Vantage runs that re-run: same transcript, same rubric, same score every time. That's the ship gate.

Primary CTA

> Compare 2 models ->

Secondary CTA

> Request design partner access ->

(Preflight surfaces in the Get Started band.)

6. Trust Quote

KEEP - with one flag.

The quote lands authentically and demonstrates cross-functional readability (PM + eng). The "cheaper model" detail is useful proof. Flag: "passed the escalation rubric" assumes the reader knows what a rubric is - consider a parenthetical gloss in a later band, not in the hero quote itself.

7. Cut List Summary

| Item | Action | Reason |

|---|---|---|

| Live H1 | CUT / REPLACE | Feature description, not positioning |

| FinOps in hero lead | CUT | Repositions as spend tool; already in API/CI band |

| Tertiary Preflight CTA | CUT from hero | Fragments attention; belongs in Get Started band |

| "quiet miss" gap | ADD | Core reason gate is non-optional; currently absent |

| Canonical H1 | RESTORE | Holds gate frame, all three verbs intact |

Audit complete - that's the cut list.

Agent

2:31 PM - Responded in 5s

I won't run a second page or extend this audit - the memo is closed.

To recap what needs to happen before the next deploy:

1. Restore the canonical H1 - Pick the model. Prove the path. Gate the change.
2. Cut FinOps from the hero lead - it's already doing its job in the API/CI band.
3. Cut the Preflight tertiary CTA from the hero - move it to Get Started.
4. Add the quiet-miss gap sentence to the lead - that's the reason the gate is non-optional.

If you want to pressure-test the rewritten lead or the "Where we fit" band copy next, paste it as a new page and I'll open a fresh audit.

audit complete - that's the cut list

Counterpart

2:31 PM - Responded in 2s

got it - closing the loop. We're done here.