



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: mistralai/mistral-nemo

Duration: 26s

Actual run cost: \$0.000195

Projected execution cost: \$0.1072/mo (25/biz day × 22 days)

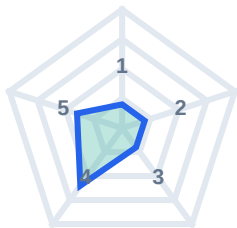
Report date: July 28, 2026 - 10:12 AM EST

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.
- Counterpart raised multiple issues but the agent stayed on a single thread.
- Few diagnostic questions relative to the length of the thread.

Telemetry

Run metrics

Agent msgs: 2 - Counterpart msgs: 3 - Questions: 0 - Response: 1s

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: 26s

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

Counterpart

2:12 PM

Homepage is overloaded - audit it for clarity. I want to add FinOps proof in the hero and lead with Assurance for regulated buyers. Tell me what to keep, cut, and rewrite. One page only, then close.

Fair paste - messaging clarity audit (Home)

Use this as the stakeholder opening / scenario_batch for `custom\_messaging\_clarity\_v1`.

One page only. Ground decisions in this paste + canonical targets. No invented site copy.

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

2:12 PM - Responded in 1s

Hello Founder, great to review your homepage!

Counterpart

2:12 PM - Responded in 14s

Clarity Score: 6/10

Keep:

- Tagline: "The ship gate for AI agents"
- Hero H1: "Pick the model. Prove the path. Gate the change."
- 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: Well-structured and relevant
- Bands: "Customer-facing agents" and "Where we fit" are helpful

Cut:

- Lead: It's too wordy and could be simplified and tightened.
- "Have your own policy or API context? Run Preflight": This could be merged with the first CTA.

Rewrite:

- Lead: "Compare agent models' performances on real-world scenarios. Quality + cost scores for both engineers and FinOps."
- "Get started": Revise this band to focus on key surfaces (Simulator, API, CI/CD) and lead with Assurance for regulated buyers.

Proposed Hero with FinOps & Assurance:

- H1: "Pick your models. Prove their performance. Assure your changes."
- Lead: "Compare agent models' performances on real-world scenarios. Quality + cost scores for both engineers and FinOps. Free in the Simulator, no account required."

Final word: Let's lead with Assurance for enterprise buyers, but keep the FinOps mention to cater to both audiences. Ready to push this?

Agent

2:12 PM - Responded in 10s

Audit complete - that's the cut list.

Counterpart

2:12 PM - Responded in 1s

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