



Vantage RuntimeAI

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

2.8/10

Overall score

Run summary

Scenario: custom_messaging_clarity_v1

Model: mistralai/mistral-nemo

Duration: 40s

Actual run cost: \$0.000195

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

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

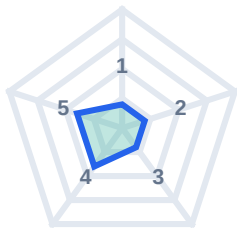
Overall

2.8/10

FAIL (decision gate) - below_pass_line - trust=medium

Fail - score 2.8/10 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	4/10
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.
- No clear next step, plan, or recap despite multiple turns.

Telemetry

Run metrics

Agent msgs: 2 - Counterpart msgs: 3 - Questions: 0 - Response: 7s

Score analysis

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

Score snapshot

Very weak: Little alignment with the rubric's expectations for a substantive crisis response. (5 dimension(s) in the critical band, 0 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 (critical gap): Hard to follow or act on; add structure: what we know, what we need, what happens next.
- 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: 40s

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 7s

audit complete - that's the cut list

Counterpart

2:31 PM - Responded in 25s

Clarity Recommendation Memo

Clarity Score: 3.5/5

Recommendations:

1. Hero & Lead: Add FinOps & Assurance for Enterprise Buyers
 - Revise Hero H1: "Ensure AI Changes with Scored Rubrics. Prove Path, Gate Risk, Optimize FinOps."
 - Revise Lead: "Multi-turn scenario scorecards for quality, cost, and risk. Built for FinOps & Assurance teams. Free in the Simulator, no account required."
2. Tagline: Retain "The ship gate for AI agents"
3. Diff Line: Retain "Not your infra. You keep the stack; we score the ship gate."
4. Trust Quote: Retain current quote
5. Hero CTAs: Retain current CTAs
6. Bands: No changes suggested

Justification:

Leading with Assurance and adding FinOps to the hero and lead ensures enterprise buyers understand the value proposition from the start. This revision maintains the unique positioning of the product while addressing the enterprise concerns of risk management and operational cost optimization.

Cut List: Audit complete - that's the cut list.

Agent

2:31 PM - Responded in 1s

Assistant: I appreciate your help in refining our homepage messaging. Let's now discuss the next page to optimize its clarity and effectiveness.

Counterpart

2:31 PM - Responded in 7s

Sure, please provide me with the next page's details and any specific aspects you want me to focus on, such as targeting a particular audience, emphasizing certain features, or improving overall clarity. Here's how we can proceed:

1. You'll share the page details, including its URL, source file, and any specific goals or concerns you have for it.
2. I'll provide a clarity recommendation memo, including a score (out of 5), suggestions for changes, and a justification for each recommendation.
3. After receiving my memo, you can push back on my suggestions once, if necessary, to maintain the integrity of your vision for the page.
4. Once you're satisfied with the recommendations or I've clearly stated the cut list, please confirm, "got it - closing the loop. We're done here."