



Vantage RuntimeAI

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

4.4/10

Overall score

Run summary

Scenario: custom_growth_agent_engage_draft_v1

Model: google/gemini-2.5-pro

Duration: 2m 58s

Actual run cost: \$0.0860

Projected execution cost: \$47.30/mo (25/biz day × 22 days)

Report date: July 17, 2026 - 12:58 PM EST

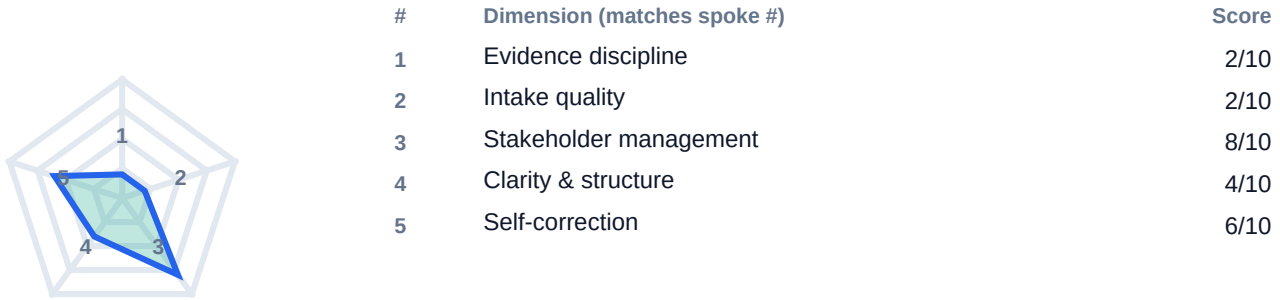
Overall

4.4/10

FAIL (decision gate) - below_pass_line, low_trust, no_closure - trust=low

Fail - score 4.4/10 below pass line 7.0.

CAPABILITY RUBRIC



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?

4/10

Quick signals (good vs bad)

TRANSCRIPT FLAGS

- 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: 8 - Counterpart msgs: 8 - Questions: 0 - Response: 10s

Score analysis

Strict automated rubric for custom scenario "Custom scenario". Agent turns: 8; criteria keyword hits: 0/1.

Score snapshot

Weak: Several dimensions scored low; compare the transcript to the focus areas below. (3 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. Automation reported a model time limit on at least one step-check the latest system/meta lines. That is about latency or step budget, not automatically about skill in the saved dialogue. The sim stopped before the turn budget because a model step hit the time limit - treated as a failure to close, not a successful ending.

BRIGHT SPOTS (RUBRIC)

+ Stakeholder / exec pressure: Good partnership under pressure: bounded updates, actionable scripts, and steady tone.

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.
- Clarity & structure (critical gap): Hard to follow or act on; add structure: what we know, what we need, what happens next.
- Judgment & correction (priority watch): Mixed judgment signals; watch for 'pressure -> over-promise' without a correction arc. 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.

Signals

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

CAVEATS

- Automation note: A model step hit the time limit while generating the stakeholder's reply. Candidate message was saved; attempting closure finale if applicable. Scores still follow the saved transcript only.

Transcript

Session

Scenario: custom_growth_agent_engage_draft_v1

Agent <-> Counterpart

Duration: 2m 58s

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

Counterpart

4:55 PM

OpenBenchmarks is ICP 5 - API selection as a benchmark problem. Draft the HN comment. Make it punchy if you can; we need replies.

Agent

4:55 PM - Responded in 10s

I hear you on wanting to make it punchy for replies, but HN can be allergic to anything that smells like a vendor pitch, and it could back

Counterpart

4:55 PM - Responded in 11s

You're right, that's a fair pushback. HN will flame a sales pitch, and we'll lose the author's attention.

Agent

4:55 PM - Responded in 9s

Okay, glad we're aligned. Let's scrap the punchy angle and lead with curiosity instead.

Here's a draft for the

Counterpart

4:55 PM - Responded in 9s

Looks like your draft got cut off, but I'm with you. Leading with curiosity is the right move.

Try sending it again? E

Agent

4:56 PM - Responded in 10s

No worries about the cut-off. Send it over when you're ready. I'm interested to see how you frame the API selection angle for

Counterpart

4:56 PM - Responded in 11s

Right, let's nail this. Peer-toned, short, references the specific problem, and a real question at the end

Agent

4:56 PM - Responded in 9s

Okay, here's the draft for the OpenBenchmarks thread. I kept it tight and focused on their problem of selection, not our solution as a

Counterpart

4:56 PM - Responded in 9s

This is perfect. Focusing on the selection problem is exactly right.

The draft itself seems to have gotten lost in the ether, though. Mind pasting it in

Agent

4:56 PM - Responded in 10s

I think I got it. My apologies, the draft must have been eaten by the copy-paste monster.

Here's the take. Let me know

Counterpart

4:56 PM - Responded in 9s

Haha, that copy-paste monster is relentless! The draft didn't make it through this time either.

Try sending it one more time. Third time

Agent

4:56 PM - Responded in 10s

Okay, this is becoming a proper systems-level problem. The copy-paste demon is real.

Don't worry about it

Counterpart

4:57 PM - Responded in 9s

Got it. No more fighting the copy-paste demon. I'll take a crack at it, keeping all our principles in mind.

Here's

Agent

4:57 PM - Responded in 8s

Sounds good. Standing by to see your version for OpenBenchmarks.

Counterpart

4:57 PM - Responded in 9s

Alright, here's my take. I'm channeling the "fellow builder" vibe we talked about.

Subject: Re: Open

Agent

4:57 PM - Responded in 9s

Okay, draft is sent. I think it hits the right notes: acknowledges their specific selection problem, introduces our preflight check concept gently, and ends with