

W&B Weave Evaluations Can Power Cost-Aware QA Scorecards

A monetization play: use Weave's scorers, judges, tracing, and cost tracking to sell cost-aware quality scorecards, then monetize regression analysis and recurring QA retainers.

Source <https://yetyield.com/blog/weave-evaluations-as-cost-aware-qa-scorecards/>

Most evaluation pitches fail because they sound like “more process”.

But buyers will fund evaluation when you tie it to two budgets they already protect:

- reliability budgets (incidents are expensive)
- efficiency budgets (token spend is real money)

Weave is useful because it explicitly combines evaluation mechanics with tracing and cost tracking, which makes “cost-aware QA” a sellable wedge.

The monetization angle

Don't sell “we set up Weave.”

Sell a cost-aware quality scorecard:

- what quality looks like (scorers, judges)
- what changed (tracing and comparisons)
- what it costs (token/API cost tracking)

Offer ladder:

- 1.Scorecard design audit (dataset + scorers + thresholds)
- 2.Implementation sprint (instrumentation + eval pipeline)
- 3.Monthly scorecard + regression analysis retainer

This ladder builds directly on:

- Agent Regression Tests Can Be a Retainer Business (/blog/agent-regression-tests-as-a-retainer-business/)
- How to Turn Agent Evaluation Checklists Into a Paid Product (/blog/how-to-turn-agent-evaluation-checklists-into-a-paid-product/)

The official Weave surfaces you can ground on

W&B describes Weave as a toolkit for evaluating LLMs and GenAI applications, with:

- scorers and judges
- evaluation datasets

- detailed tracing
- cost tracking
- model versioning and comparisons

Official entry point: <https://docs.wandb.ai/models/tables/evaluate-models>

Why cost-aware QA converts better than generic “quality”

Generic quality metrics get debated.

Costs do not.

When you show:

- “quality improved 4%”
- “cost increased 40%”

...you are now in a pricing conversation, not a vibes conversation.

That makes it easier to sell:

- budget ceilings
- sampling policies
- model selection decisions

The scorecard you can sell in week one

Start with a small scorecard that maps to money-risk:

Quality:

- format compliance (deterministic)
- groundedness / hallucination risk (judge-based or heuristic)
- task completion (did the workflow actually succeed?)

Cost:

- token spend per request (median + p95)
- “cost per successful task” (a business-facing metric)

This aligns with OpenAI’s warning against “vibe-based evals” and its emphasis on task-specific evaluation and continuous evaluation.

Official anchor: <https://developers.openai.com/api/docs/guides/evaluation-best-practices>

The deliverables clients pay for

Deliverable 1: Evaluation dataset (v1)

- typical cases
- edge cases
- adversarial cases

The dataset is the compounding asset, not the tooling.

Deliverable 2: Scorers and judges (v1)

- deterministic checks for compliance
- judge-based scoring where nuance matters

Deliverable 3: Cost policy

Write explicit budget rules:

- a max token ceiling per request
- a max monthly evaluation spend
- a sampling policy (evaluate 10% of runs + 100% of failures)

Deliverable 4: Monthly regression report

Monthly:

- scorecard trend lines
- top regressions with traces
- cost anomalies
- prioritized fix backlog

That report is what the retainer buyer is buying.

Where Weave fits relative to other substrates

Weave is not the only evaluation substrate.

But it is a particularly clean fit when:

- the org already uses W&B for ML ops
- the buyer cares about cost governance
- you want a “compare model versions” story that looks like real engineering

For substrate selection framing:

- [How to Compare LangSmith, Phoenix, Weave, and Foundry as Evaluation Substrates \(/blog/compare-langsmith-phoenix-weave-foundry-as-eval-substrates/\)](#)

What to avoid

- evaluating everything (sampling is part of the product)
- “quality dashboards” without thresholds
- discussing cost without tying it to business outcomes (speed, reliability, unit economics)

Next research direction: how to turn Anthropic’s Console Evaluation tool into a paid “prompt QA pack” that creates a repeatable pre-deploy gate for prompt changes.