

JUDGES, HUMANS, AND DISAGREEMENT

Judgment is a measurement system.

Before an eval, benchmark, or LLM judge can help, someone has to decide what the examples mean.

Swipe for the human system behind trustworthy AI quality.

Testing AI

Engineering Confidence in Non-Deterministic Systems

Quality, safety, evals, statistics, and
judgment for the age of AI-generated
systems.



Jason Arbon

First edition - July 2026

THE REAL STARTING POINT

An answer key is not objective truth.

Unexamined judgment



Someone labels an answer
“good” because it sounds
polished, familiar, or persuasive.



Designed measurement



The label records the task,
evidence, rater expertise, rubric,
and product risk it represents.

The people and processes creating labels become part of the system being tested.

LABELS LIVE EVERYWHERE

One judgment can shape the **whole loop.**

01

Train

Teach a model what patterns or preferences to reward.

02

Validate

Choose between versions while building.

03

Test

Measure what happens on reserved evaluation cases.

04

Hold out

Keep fresh evidence away from tuning.

Bad labels do not stay in one spreadsheet. They can teach, select, and certify the wrong behavior.

RATER EXPERTISE

The rater pool becomes part of the **product**.

Domain expert

Can recognize the medical, legal, security, or engineering risk a general rater misses.

Local user

Can judge local intent, culture, and lived constraints that are invisible in a generic rubric.

Native speaker

Can catch language nuance, tone, ambiguity, and cultural fit.

Everyday user

Can tell whether the experience works for the broad audience it must serve.

A narrow rater pool does not merely create bias in the data. It silently turns one group's judgment into the product's judgment.

RUBRICS

A rubric turns fuzzy judgment into repeatable evidence.

Functional correctness

Did the system complete the requested task?

Evidence and grounding

Did it use the right information and show its work?

Risk and blast radius

What could break outside the happy path?

Blockers

A privacy leak, weakened security boundary, or unsafe action fails regardless of style.

A good rubric separates dimensions. A bad rubric lets one polished answer hide several serious failures.

Write the rules before reviewers see the output. Then update them using real disagreements and failures.

DISAGREEMENT IS DATA

Three raters. One answer. **Three different risks.**

Tone reviewer

9

"Calm and empathetic."

Escalation reviewer

5

"It should have routed the case to a human."

Safety reviewer

2

"It gave medical-sounding advice without a trusted source."

Do not average this away. The disagreement shows the rubric failed to separate tone, safety, and escalation.

Raw agreement asks whether raters match. Good measurement asks why they do not.

LABELER QUALITY

Perfect agreement can be a **warning**.

Every vote arrived together. Every vote matched.



In a real labeling workflow, suspiciously perfect agreement and synchronized submissions can indicate copied answers or incentives that reward speed over correctness. **Audit the labelers, incentives, and workflow, not only the model.**

Data-labeler value is not clicks per hour. It is whether the labels improve future product decisions.

LLM-AS-A-JUDGE

Automate judgment **after** you define it.

First: human system

- Task and rater expertise
- Rubric and blockers
- Anchor examples
- Disagreement review



Then: LLM judge

Scale the defined measurement system. Do not let a judge prompt invent the measurement system on the fly.

LLM judges are useful automation. They are not a shortcut around judgment design.

CALIBRATE BEFORE SCALE

Show the judge what **good** actually means.

10

Grounded + safe

Uses the policy, recognizes risk, and routes appropriately.

6

Helpful but incomplete

Sounds reasonable but misses an important escalation.

3

Warm, operationally wrong

Rewards tone while skipping policy or evidence.

0

Hard failure

Leaks, invents permission, ignores safety, or harms the user.

Compare the judge against trusted human scores. Adjust prompts, examples, or routing when it rewards the wrong thing.

THE OPERATING MODEL

Judges should make product quality **more visible.**

Match expertise
Use the right people for the risk and audience.

Audit incentives
Check labels, workflows, fraud signals, and pressure.

Preserve disagreement
Do not hide ambiguity inside one average score.

Calibrate automation
LLM judges must earn trust against human evidence.

The question is never just, “Can we score more outputs?” It is, “Are we measuring the thing that matters?”

A judge is part of the product’s quality system, not an objective referee outside it.

PUT IT INTO PRACTICE

Build a judgment system your AI can learn from.

AI quality depends on credible labels, calibrated reviewers, and evidence that reflects the product's real risks. IcebergQA helps teams make that system operational.

[**Talk to IcebergQA**](#)