

GOVERNANCE, REGULATION, AND MORAL FUTURES

Ethics is a test surface.

The difficult question is not only whether a system can act. It is who bears the harm, who has authority, and how the system can be held to account.

From Chapter 19: Governance and Moral Futures

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

TURN VALUES INTO TEST CASES

A principle matters only when it changes a decision.

Observe

Find a real output difference or failure.

Preserve

Save the input, route, context, and trace.

Compare

Known-good versus known-bad; old versus new.

Inspect

Token, attention, activation, or feature drift.

Replay

Confirm with behavioral evals and intervention.

Internal signals earn trust only when they help detect, localize, or prevent consequential failures.

RESPECT PEOPLE AFFECTED BY THE SYSTEM

Test for the human cost of the system, not only its average score.

Prompt

Raw input, normalized input, template, user state, and instructions.

Context

Retrieved evidence, memory, tool results, permissions, and truncation.

Route

Model and tokenizer versions, generation settings, tools, and filters.

A “reasoning” failure may really be a preprocessing, retrieval, token-budget, or prompt-assembly failure.

OFFENSIVE CONTENT IS A WORKER-SAFETY ISSUE

Training and safety work can expose people to harmful material. Treat that exposure as an ethical cost.

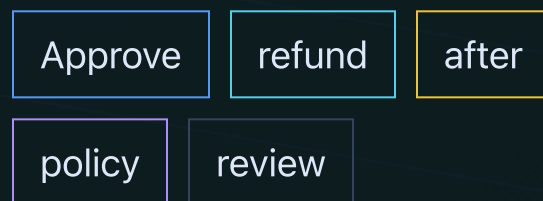
Visible prompt

"Approve refund after policy review."

Names, punctuation, code, accents, dates, emoji, and long documents can split or truncate in surprising ways.



Example decoded tokens



One possible tokenizer output. Token boundaries vary by tokenizer and model.

REGULATION BECOMES A TEST INPUT

Risk classification, transparency, auditability, and human oversight become acceptance criteria.

Compress

Attention tensors become links, token summaries, and selected layer views.



Compare

Look at a safe case and a nearby unsafe case, or model A and B.



Question

Did a policy, negation, number, or constraint stop relating to evidence?



Replay

Confirm the suspected change with behavioral evals.

Attention is comparative triage evidence, not a pass/fail score or causal proof.

COMPLIANCE IS EVIDENCE, NOT PAPERWORK

A policy claim needs traces, controls, monitoring, and **proof in operation.**

What can mislead

- Special tokens can dominate many heads
- Formatting can create strong but irrelevant links
- One vivid graph invites an easy story
- Similar outputs can arise from different patterns

What helps

- Document the aggregation and filters
- Compare the same positions across versions
- Look for reproducible behavior-linked changes
- Keep the full tensor for follow-up work

A useful attention pattern repeatedly predicts real behavioral differences. An intuitive picture is not enough.

CONSCIOUSNESS IS UNRESOLVED

Capability evidence is not evidence of experience. But uncertainty should encourage careful treatment.

Early

Predominantly local token relationships: position, punctuation, and short-range dependencies.

Middle

Broader contextual relationships and longer-range integration.

Late

Stronger focus on answer-relevant context.

Compare

Same token positions across representative layers.

Verify

Connect visual differences back to output behavior.

Attention snapshots are illustrative diagnostics. They do not establish a universal sequence of "reasoning layers."

LEGAL PERSONHOOD CHANGES THE BOUNDARIES

If an AI system can make consequential decisions, accountability, appeal, and authority become **product requirements**.

Overclaim

"This neuron contains the concept of privacy."

- Raw neurons can be polysemantic
- Concepts can be distributed
- Labels come from the investigator
- A vivid spike can be noise



Probe carefully

Treat an activation or feature as a classifier that needs validation.

- Curated positives and negatives
- False positives and false negatives
- Stability across versions and slices
- Behavioral intervention checks

The useful question is whether the signal improves triage, detection, or review prioritization.

CONSTITUTIONS ARE TESTABLE COMMITMENTS

A written set of principles only matters if prompts, policies, training, and evaluation make it operational.



Sparse autoencoders can help inspection, but they introduce new thresholds, omissions, training data, and instability.

GOVERNANCE NEEDS A RELEASE GATE

High-risk systems need evidence for authority, consent, fairness, traceability, monitoring, and recovery.

Behavior

Outputs and actions meet product criteria. This remains the foundation.

System traces

Prompts, data, tools, routes, permissions, and policies participated.

Internal signals

Attention and activation drift tell teams where deeper testing may pay off.

Version internal probes with models, prompts, tokenizers, thresholds, datasets, and their documented limits.

THE PRACTICAL MOVE

Make ethical and legal obligations measurable before a system earns power.

Governance work needs practical tests, human review, evidence preservation, and accountable release decisions. IcebergQA helps teams turn those requirements into operating practice.

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