

EMBODIED AND LONG-RUNNING AI SYSTEMS

When AI can act, quality becomes physical.

Robots move, collaborate, remember, observe, and keep running. A small planning error can become a real-world incident.

From Chapter 18: Embodied AI

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

SIMULATION BEFORE REALITY

Run thousands of awkward worlds before touching a real staircase.

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.

PHYSICAL SAFETY IS A SYSTEM PROPERTY

Score task success only after the robot stays inside a **safe envelope**.

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.

RECOVERY IS A FIRST-CLASS CAPABILITY

A robot must know when to retry, ask, stop, undo, or escalate.

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

Approve	refund	after
policy	review	

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

ONE TASK, MANY BOUNDARIES

A “simple” robot job can cross permissions, people, objects, and hazards.

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.

SIMULATION IS A MEASUREMENT TOOL

Virtual worlds expand coverage. They do not replace **reality**.

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.

PEOPLE ARE PART OF THE ENVIRONMENT

Human comfort, consent, privacy, and interruption are **quality requirements**.

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."

LONG-RUNNING SYSTEMS CHANGE OVER TIME

Drift, stale memory, retry loops, and resource growth are **time-travel failures**.

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

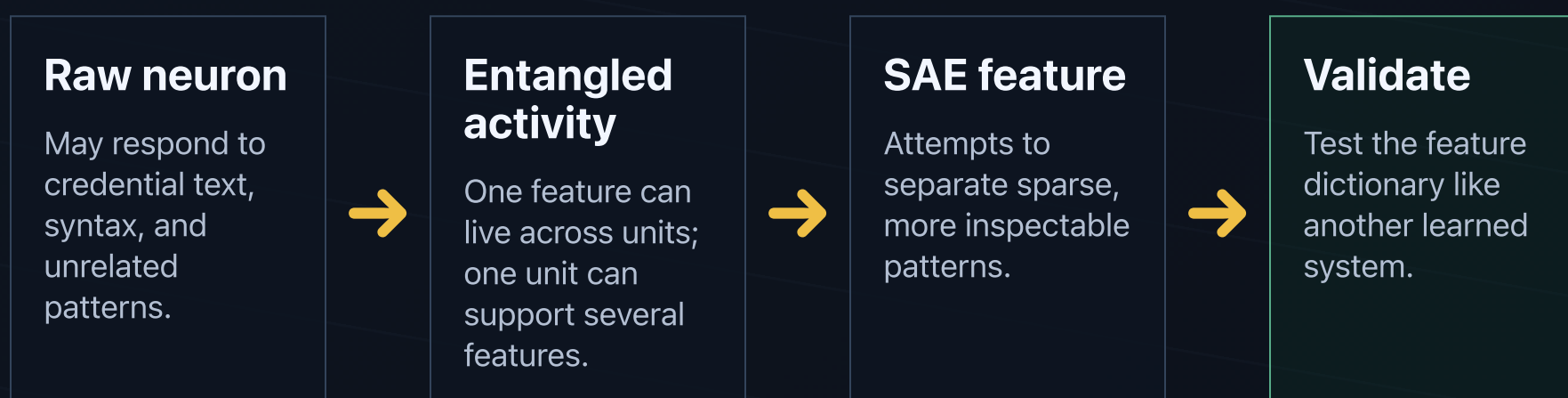
Treat an activation or feature as a classifier that needs validation.

- 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.

PHYSICAL AI NEEDS PRIVACY TOO

A robot can see homes, factories, routines, and personal behavior. Protect that **memory**.



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

ACTIONS REQUIRE INDEPENDENT SAFEGUARDS

Do not rely on the model alone when an interlock, rate limit, geofence, or **stop control can reduce harm.**

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

Test the path, the plan, the recovery, and the real-world consequence.

Embodied AI needs simulation, field evidence, safety envelopes, action traces, and practical recovery design. IcebergQA helps teams make those requirements measurable.

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