

Executive Accountability in the Age of AI

Governance is not what you claim. It's what you can prove.

For the executive accountable for AI they did not build, and cannot personally inspect.

THE QUESTION THAT COUNTS

“On what do you base that this was responsible?”

Not: how does the technology work. A simpler, and harder, question. And it always comes: from an auditor, a regulator, a counterparty, a board member.

The trap is substitution

- ◆ The machine's output takes the place of a decision no one can defend.
- ◆ Robodebt: debts calculated before the evidence was sought. A\$2,754 on average, hundreds of thousands of people.
- ◆ The technology did not fail first. The burden of proof was reversed.

The same mistake, six times

Exhibit 2.1. Where one thing stood in for another

In each of these cases, something usable was present. In none of them had anyone established that it could carry the next step.

AVAILABLE	USED AS	WHAT HAD NOT BEEN ESTABLISHED
Annual income	fortnightly income	<i>actual income in the relevant periods</i>
Chatbot answer	a correct statement of the policy	<i>agreement with the applicable policy</i>
Model score	fitness for use	<i>fitness for this application</i>
Vendor attestation	sufficient proof of safety	<i>what had actually been tested and verified</i>
Prior approval	permission to proceed	<i>whether the original assumptions still held</i>
Logged ticket	evidence of follow-up	<i>which decision had actually been made</i>

You are bound by what your AI says

- ◆ Air Canada was held to what its own chatbot promised a customer.
- ◆ “It was the chatbot” was no defense.
- ◆ Your name is on it, even if you did not write the code.

THE INSTRUMENT

One instrument. One question.

- ◆ The AI Control Index is not a checklist, but one question.
- ◆ What can you prove, and to whom, when they ask?
- ◆ Every chapter turns on a real, public-record case. Every chapter ends with five questions you can ask on Monday.

THE INSTRUMENT

The AI Control Index

LAGEN

L1 Strategy & Accountability

L2 Ethics & Fairness

L3 People, Skills & Operating Model

L4 Applications & Agents

L5 AI Engineering

L6 Data, Context & Governance

L7 Systems & Sources

L8 Infrastructure

SCHILDEN

S1 GRC

S2 AI Security & IR

S3 Supply Chain

S4 Observability

S5 FinOps

Eight layers, five shields, no checklist. One question: what can you prove, and to whom, when they ask?



PART II

The Decision

What your organization actually decides when it deploys AI.

Behind every model sits a decision

- ◆ Behind every model sits a decision someone must sign for. This part finds it, names it, and asks: who carries it?
- ◆ “There’s a human in the loop.” Which follow-up question actually counts?
- ◆ Not whether there is a human, but whether that human can act, and in time.



PART III

Intervention

Can your people carry an AI outcome and intervene in time?

Who moves when the machine acts?

- ◆ Uber and the NTSB: when the system acts, who moves? And is there still time?
- ◆ Oversight is not a box to tick. It is the ability to stop.
- ◆ An oversight loop that does not close is not oversight. It is a signature after the fact.

PART III · INTERVENTION

5.6 seconds. Three layers of intervention.

Exhibit 7.1. 5.6 seconds. Three intervention layers

The collision arose in seconds. What the reader must see is that the ability to intervene was not squandered in those seconds, but designed much earlier.



What could no longer be recovered in 5.6 seconds had already been organized before the drive.



PART IV

The Stack

What to demand of each layer of the technology stack.

Every layer hides a decision

- ◆ From data to model to deployment: not how the technology works, but which decision is hidden in each layer, and what evidence that layer must produce.
- ◆ “The model scores 94%.” Who decided in advance what good enough meant?
- ◆ A sepsis model: 0.76 in the lab, 0.63 on real patients. The number hid the decision.



PART V

The Five Shields

Noticing when reality drifts from what you approved.

Monitoring is a tripwire

- ◆ The five shields notice when reality drifts from what you signed off.
- ◆ Monitoring is not a dashboard to watch. It is a tripwire.
- ◆ The mechanism that tells you: the model that is running is no longer the model you approved.

PART VI

Forensic Exposure

How to prove it, as the world changes.

From observation to validity

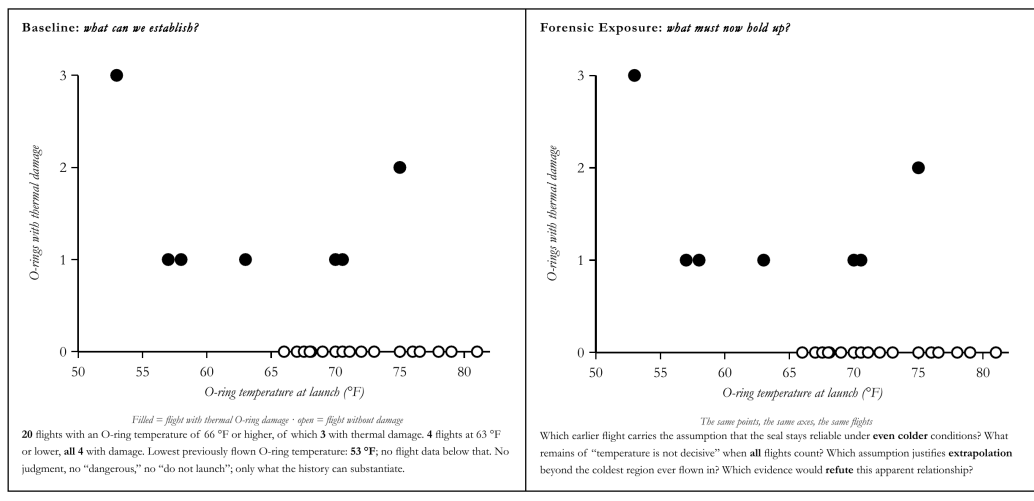
- ◆ How an observation becomes a finding, a finding becomes a decision, and a decision keeps its validity as the world changes.
- ◆ The same map, read two ways: baseline and adversarial.
- ◆ The baseline and the forensic reading show not different facts, but a different burden of proof.

PART VI · FORENSIC EXPOSURE

The same flight history, two readings

Exhibit 19.1. The same flight history, two readings

One and the same chart, twice: the full flight history, all flights with and without thermal O-ring damage, plotted against the O-ring temperature at launch. No point, no flight, no axis changes. Only the question the map has to answer.



No point changes, no flight changes, no fact changes. The points are the same. The question is not.

CHALLENGER · ROGERS COMMISSION

**“The map does not change.
The burden of proof does.”**

No point changes, no flight changes, no fact changes. What looked like a pattern in the baseline became, read adversarially, a warning that had been ignored.

THE METHOD

Five questions for Monday

- ◆ Every chapter ends with five questions you can put to your own organization on Monday morning.
- ◆ The point is not to admire the framework. The point is to use it.
- ◆ You don't need to build the technology to ask the right question.

REAL, NOT HYPOTHETICAL

The cases

- ◆ Robodebt · Air Canada · the sepsis model · Zillow · Clearview · Knight Capital.
- ◆ The Grenfell fire test that outlived the certified product claim.
- ◆ The New York City chatbot that told businesses to break the law.
- ◆ Public record. No invented scenarios.

WHO THIS IS FOR

Who signs for this

- ◆ Executive accountable officers, the named person on the governance register.
- ◆ Chief Data & AI Officers, who translate board ambition into operational AI.
- ◆ Risk, compliance & audit: they know the gap between policy and practice.
- ◆ Strategy & consulting leads, who want a governance backbone for their recommendations.

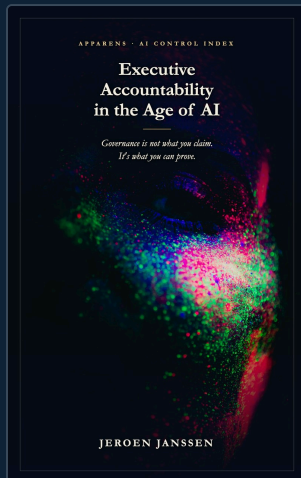
THE CORE

What you don't need. What you must know.

- ◆ You don't need to build the technology.
- ◆ You must know which decision is hidden inside it, who must sign for it, and what you must see before you believe it.
- ◆ That is where accountability begins.

THE BOOK

Executive Accountability in the Age of AI



English

E-BOOK · HARDCOVER



Nederlands

E-BOOK · HARDCOVER

Six parts, from the trap to the proof. Available as e-book and hardcover, in English and Dutch.

THE APP

Not just a book, a workspace

- ◆ The AI Control Index app applies the method to your own organization.
- ◆ “Decision Ready”: evidence reviewed, risks identified, gaps prioritized.
- ◆ A personal workspace for the executive who carries the accountability.



THE FOUNDATION

Published research, not marketing

From Battlefield to Boardroom

Strategic Red Teaming as an epistemic governance instrument

ARXIV:2607.01913

From Record to Finding

Why tamper-proof logs cannot establish legal oversight of agentic AI

ARXIV:2607.00941

A Supervisory-Evidence Ontology

for agentic AI under EU law

ZENODO 10.5281/ZENODO.19758441

From Battlefield to Boardroom: Strategic Red Teaming as an Epistemic Governance Instrument in the Age of AI

Jeroen Janssen, February 2026

Abstract

Strategic planning in contemporary organizations proceeds, almost universally, on the assumption that the approval of a strategic initiative represents a sufficient test of its merit. It does not. The strategy approval process is characterized by advocacy rather than adversarial inquiry, and the assumptions that hold a strategy together routinely go untested until contact with operational reality forces correction at high cost. This essay argues that red teaming, an adversarial methodology originating in military doctrine and subsequently institutionalized in cybersecurity

Open and testable on arXiv and Zenodo. The method rests on peer-reviewable work.

KNOWLEDGE SHARING

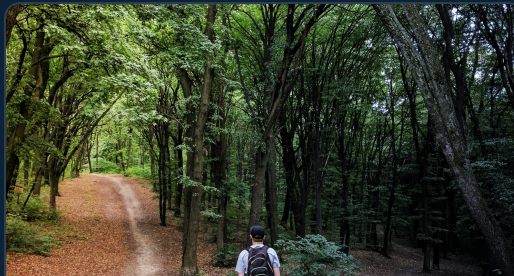
A series of sharp analyses on apparens.nl



PART 1

Your AI Strategy Has No Contract With Reality

Jan 2026



PART 2

Your AI Just Started Making Decisions. Who Told It To?

Jan 2026



PART 3

Your AI Is Scaling. Your Understanding Isn't.

Feb 2026

Ongoing, on AI governance, red teaming and the EU AI Act.

**Governance is not what you claim.
It's what you can prove.**

[**apparens.nl/book**](https://apparens.nl/book)

Jeroen Janssen · Apparens