

THREE PILLARS
IN PRACTICE:

WHEN AI SHAPES HIGH-CONSEQUENCE DECISIONS

MEDICINE x PUBLIC SERVICES

Two practical applications of
Institutional Accountability,
Operational Control and
Human Capacity.



MEDICINE



hosainstitute.com/hosa-on-ai

THREE PILLARS IN PRACTICE: MEDICINE

**AI SAID THE PATIENT COULD WAIT.
THE PATIENT DIED BEFORE A DOCTOR SAW THEM.**



An AI-enabled triage system assessed the patient's symptoms, classified them as low risk and placed them behind dozens of others.

A clinician remained "in the loop" but was managing an AI-prioritised queue without the time or evidence required to independently reassess every case. By the time the patient deteriorated, intervention was too late.

WHO ACTUALLY DECIDED THAT THEY COULD WAIT?

Institutional Accountability

Who authorised AI to influence the order in which patients receive care?

Who determined that the risk was acceptable?

Who answers for the death?

Operational Control

What was the system permitted to decide?

Could anyone see why the patient was classified as low risk?

What should have triggered human escalation and why did it not?

Human Capacity

Did the clinician have the time, information, situational awareness and authority required to challenge the system?

Or were they present only so the institution could claim that a human remained in control?

The failure was not simply that AI got it wrong. The failure was that an institution allowed AI to shape a life-critical decision without ensuring that accountability, operational control and human capacity remained connected.

Putting a clinician at the end of an AI-controlled workflow does not necessarily create meaningful human oversight.

It may simply create a human liability sink.

**Join the
discourse
live**

**VIEW
EVENTS &
REGISTER →**

PUBLIC SERVICES





THREE PILLARS IN PRACTICE: PUBLIC SERVICES

AI ACCUSED HER OF FRAUD. HER BENEFITS WERE STOPPED. HER BANK ACCOUNT WAS FROZEN. SHE LOST HER HOME. THE AI WAS WRONG ABOUT HER—AND 180,000 OTHERS.

A government AI analysed public and financial data for fraud and initiated enforcement. Her payments stopped. A fraud marker spread across connected institutions. She appealed, but a human received the system's evidence and upheld its decision. Months later, investigators found that the AI had misread a common pattern affecting tens of thousands. Families had already lost income, housing, services and institutional trust.

WHO DECIDED THAT SHE WAS GUILTY?

Institutional Accountability

Who authorised AI to infer fraud and initiate consequences?

Who decided that citizens should carry the burden of disproving it?

Which institution answers for the combined harm?

Operational Control

What evidence was the system permitted to use?

What actions could one fraud flag trigger?

Why did the pattern of repeated errors not stop the system?

Human Capacity

Did reviewers have the time, evidence and authority to independently reconsider each case?

Or did "human review" mean asking an overwhelmed person to confirm the system's conclusion?

The failure was not simply that AI produced the wrong answer. The failure was that institutions allowed one unverified inference to acquire authority and travel.

Nobody consciously decided to destroy her life. Every system simply acted on the same AI-generated conclusion.

Join the
discourse
live

VIEW
EVENTS &
REGISTER →

HOSA on AI  **LIVE**

BIG QUESTIONS EXAMINING THE FUTURE WE ARE BUILDING



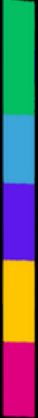
Subscribe to our
newsletter for
event notifications



leo@hosainstitute.com
to join as a panelist or
guest speaker.



hosainstitute.com/hosa-on-ai



HUMAN OPERATING SYSTEM ARCHITECTURE