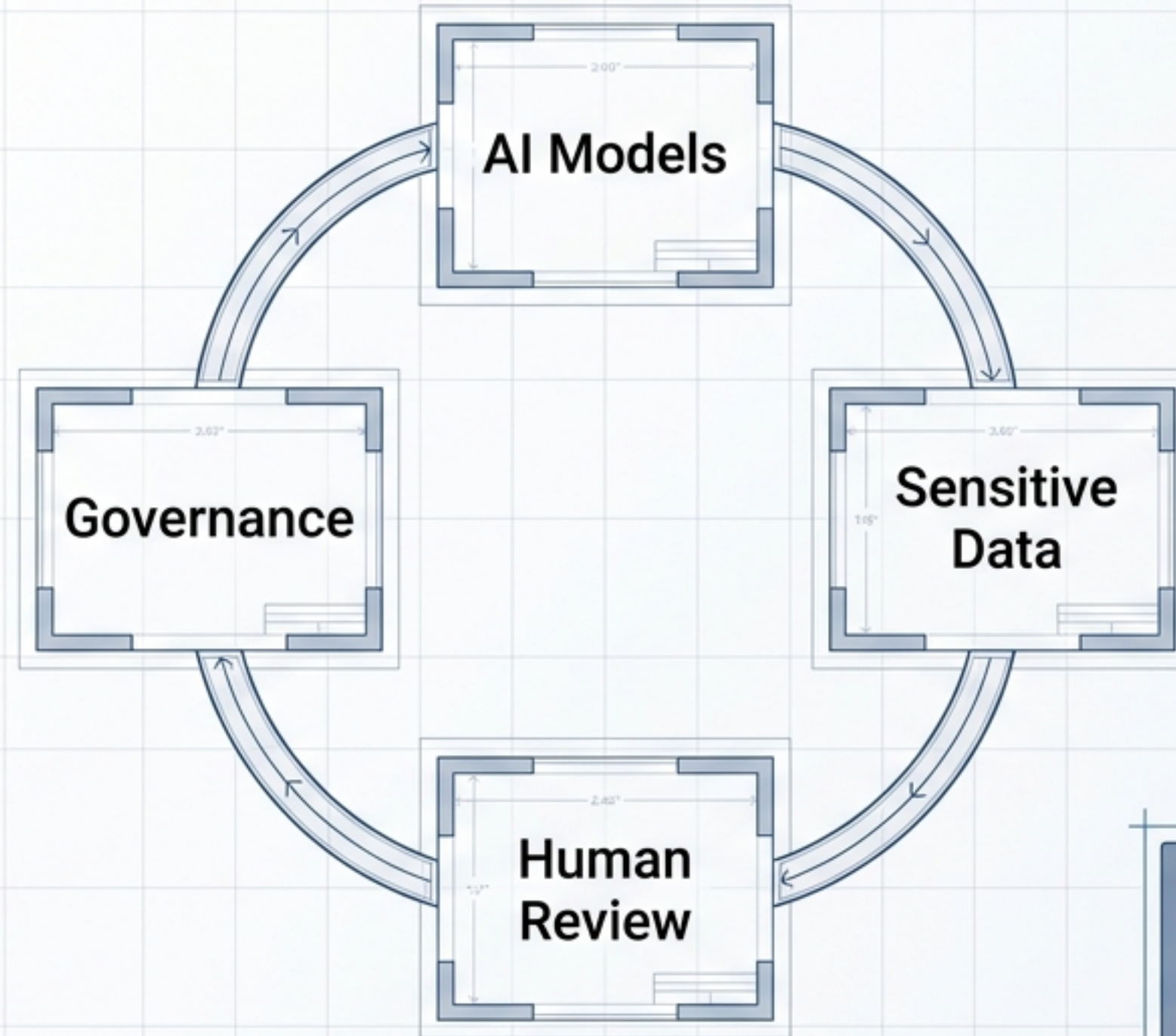


Responsible AI & Risk Review

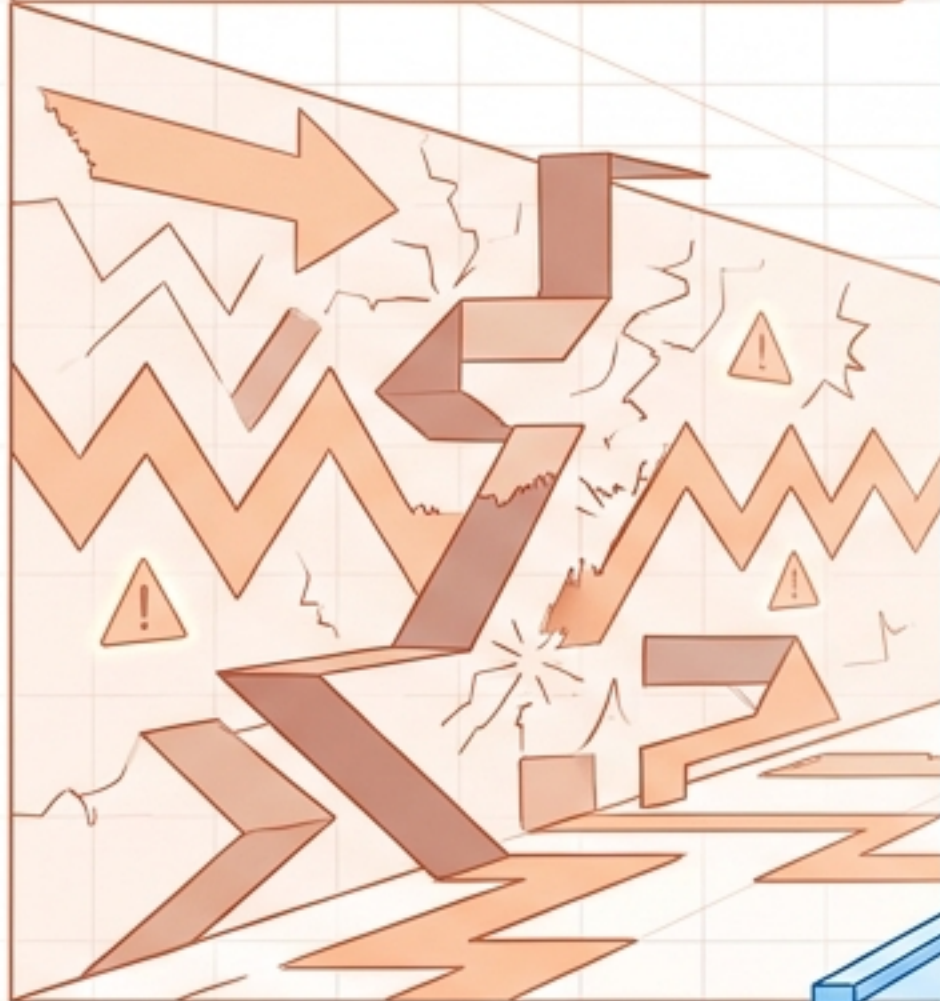
Building safe, practical design habits for healthcare AI adoption.



Human review is a design feature, not a failure.

Governance Should Enable Good Work

**Too Loose:
Ignore Risk**



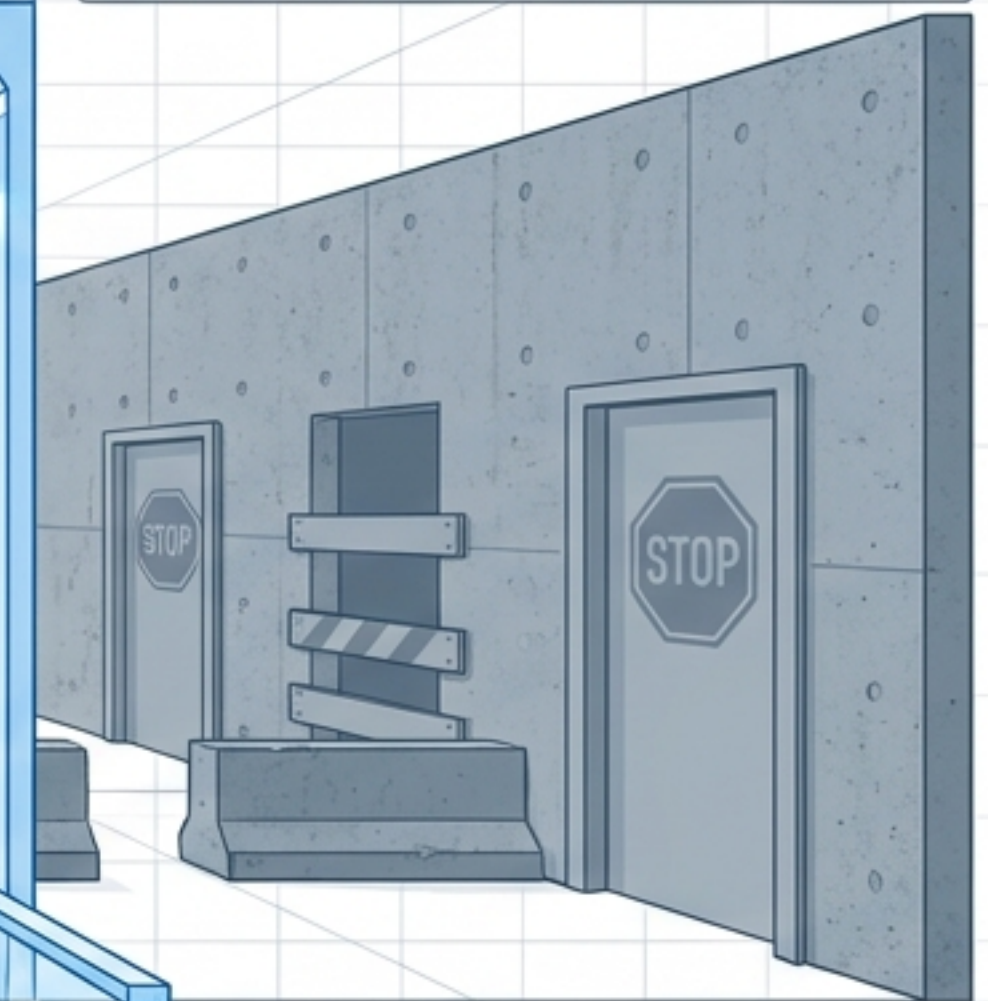
Ignoring risk exposes patients, compliance, and trust.

**The Sweet Spot:
Enablement & Responsible Design**

Practical guardrails help teams move responsibly. Low-risk ideas move quickly. High-risk ideas get the documented paths they need.

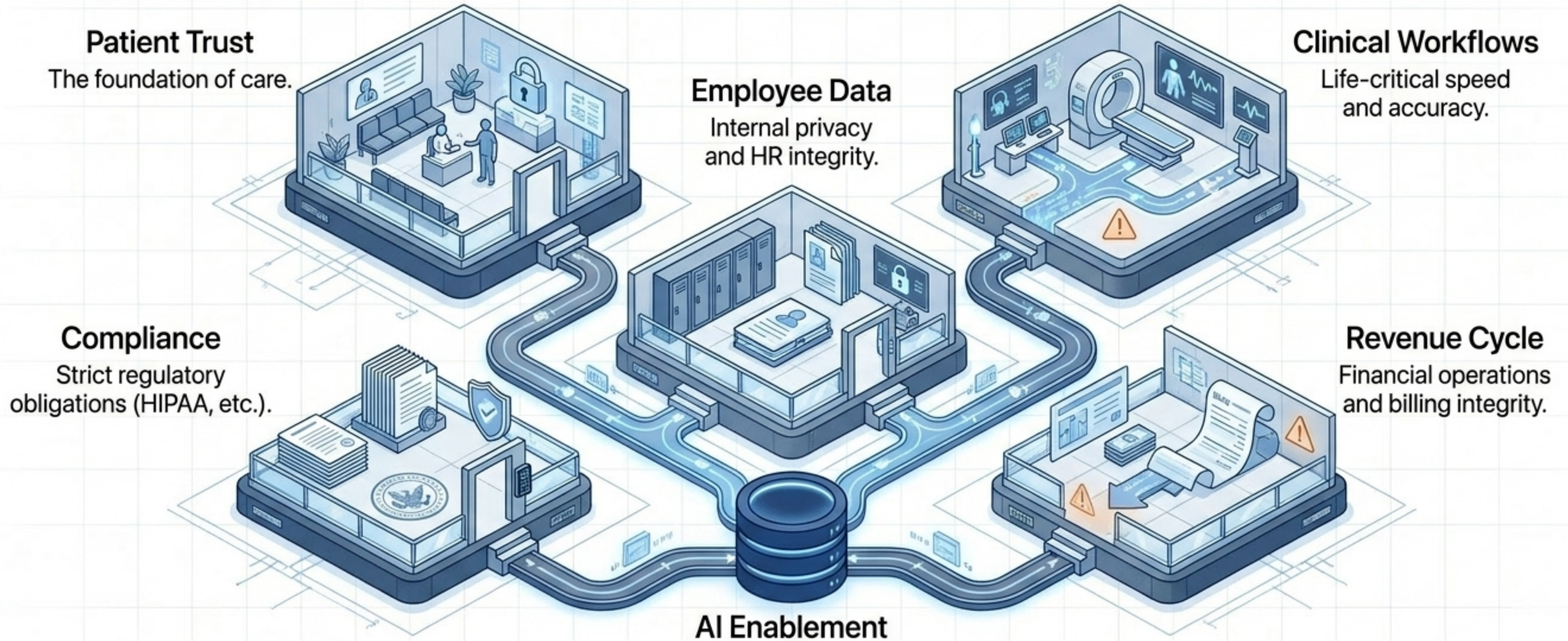


**Too Strict:
Stop Everything**



Stopping everything kills innovation and burns out teams.

Why AI Risk Matters Specifically in Healthcare



AI enablement requires clear boundaries to protect our most critical ecosystem dependencies.

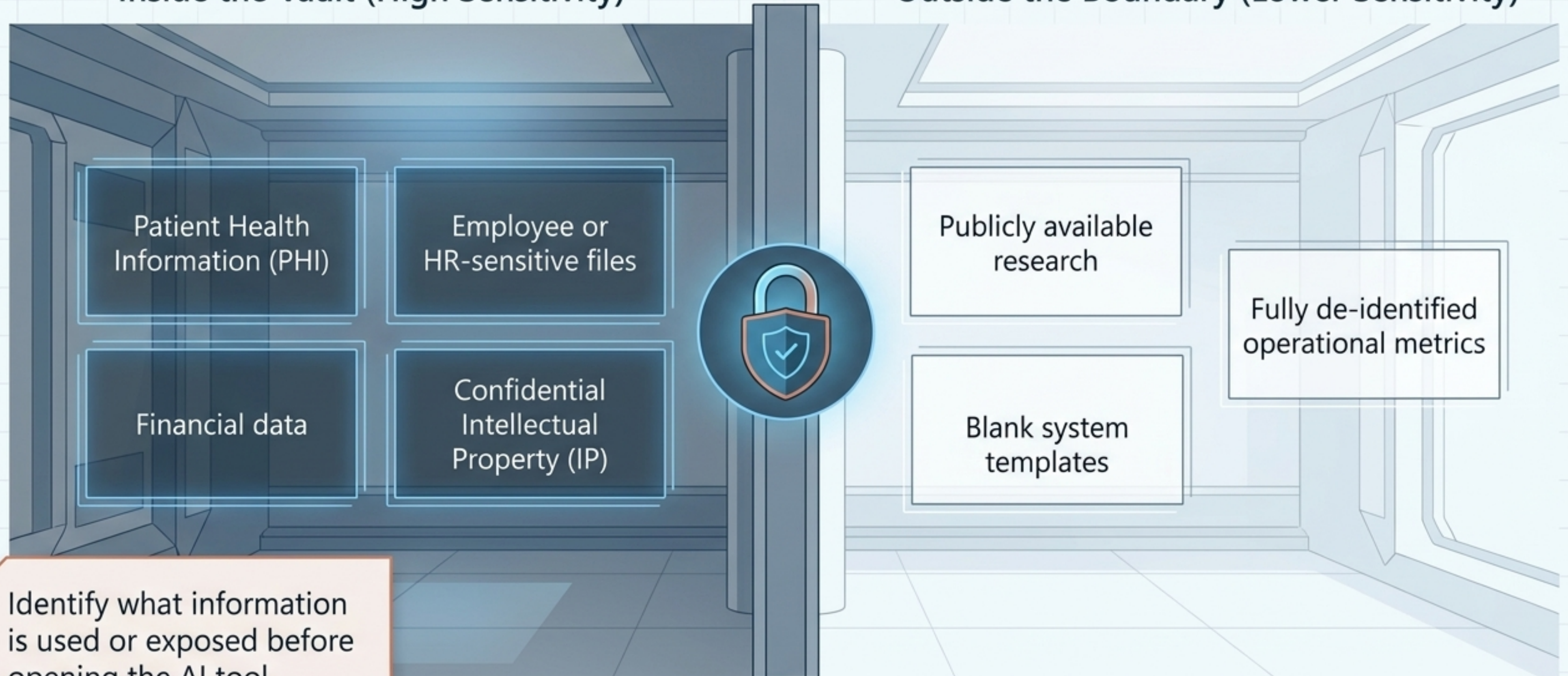
The Baseline Diagnostic

	Does this involve patient information?		Can a human review the output?
	Does this involve employee information?		Can the workflow be audited?
	Does this affect care, billing, legal, compliance, employment, or finance?		Can the solution be turned off or rolled back?
	Is the data source approved?		Who owns the process?

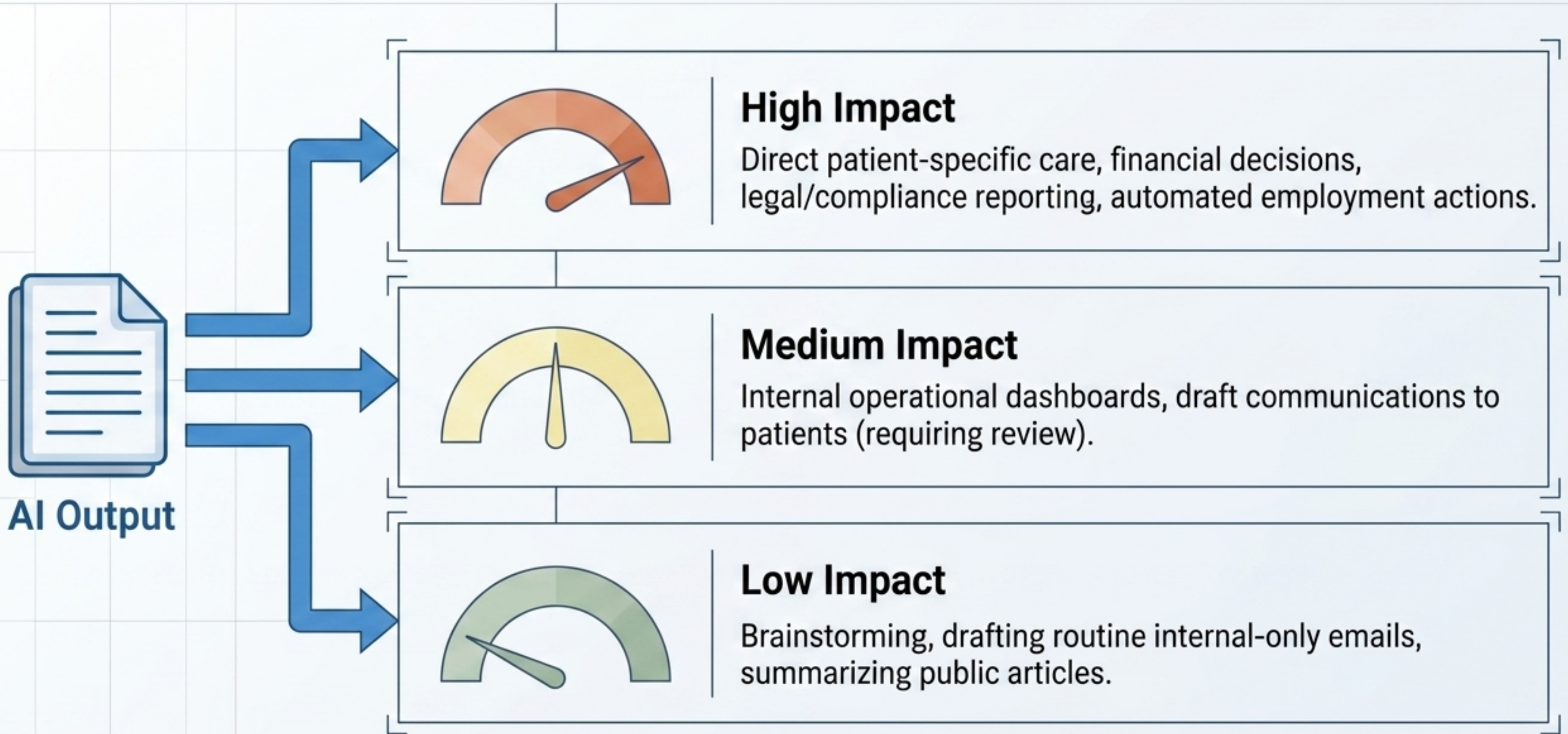
Identify Sensitive Data Before Using AI

Inside the Vault (High Sensitivity)

Outside the Boundary (Lower Sensitivity)



Impact Increases Based on Where Outputs Land



AI Autonomy and the Gradient of Risk

Where does your use case sit on the ladder?

Acting (Highest Risk)

System executes a task or sends external content with zero human intervention.

Deciding (High Risk)

System makes a choice based on rules. Human exception handling only.

Recommending (Medium Risk)

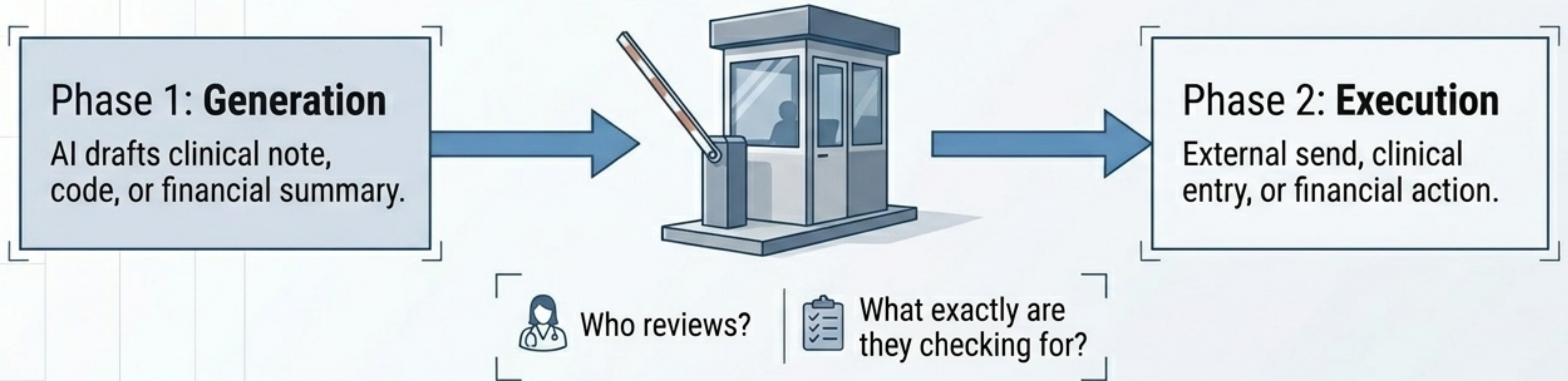
System analyzes data and suggests a specific path or decision.

Drafting (Lowest Risk)

System suggests text or code. Human must refine.

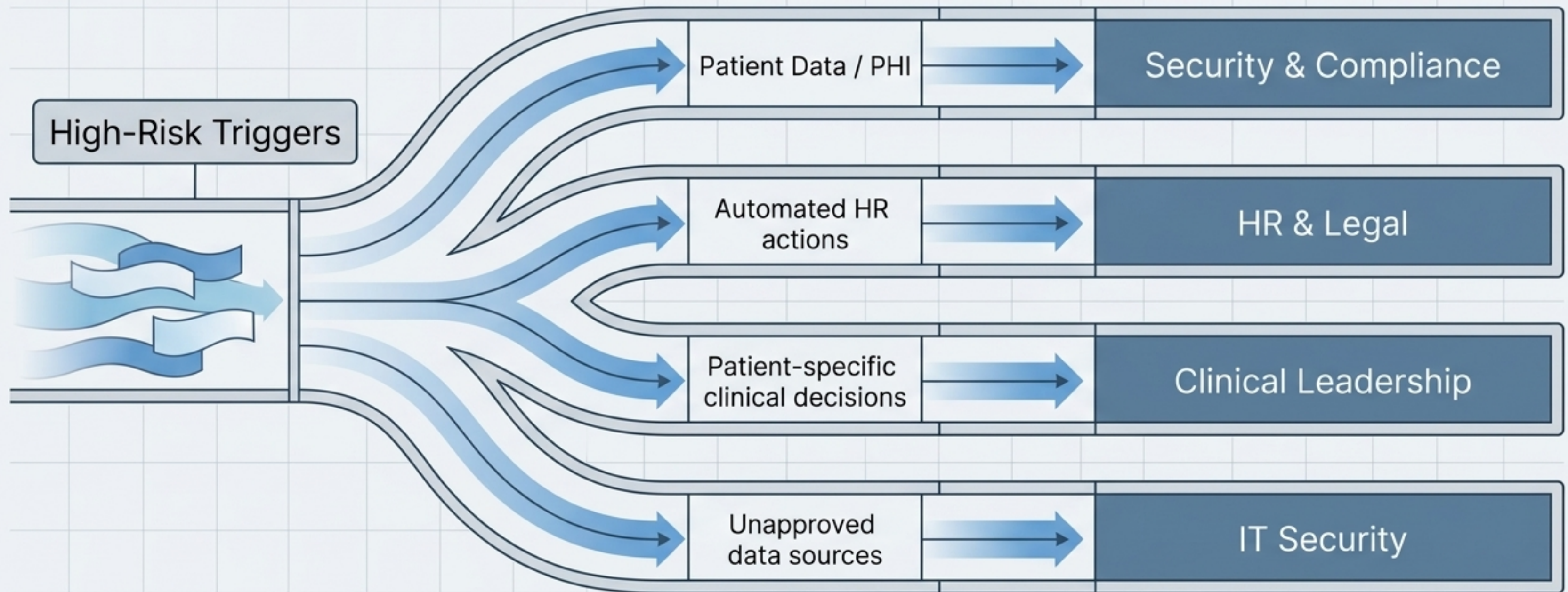
Human Review Must Be Specific and Precede Impact

Mandatory Pause



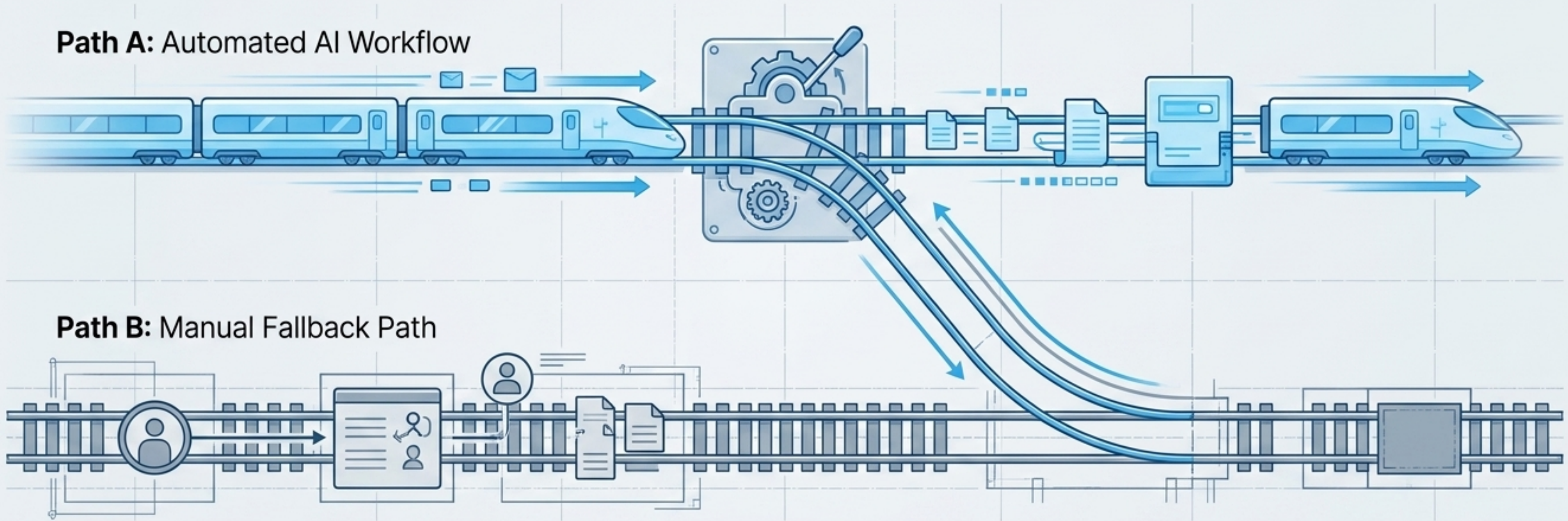
Never bypass the review gate for clinical, financial, or HR-impacting actions.

Routing Complex Cases to the Right Experts



Escalating a use case is a success metric. It means the safety system is working.

If You Cannot Roll It Back, You Cannot Roll It Out



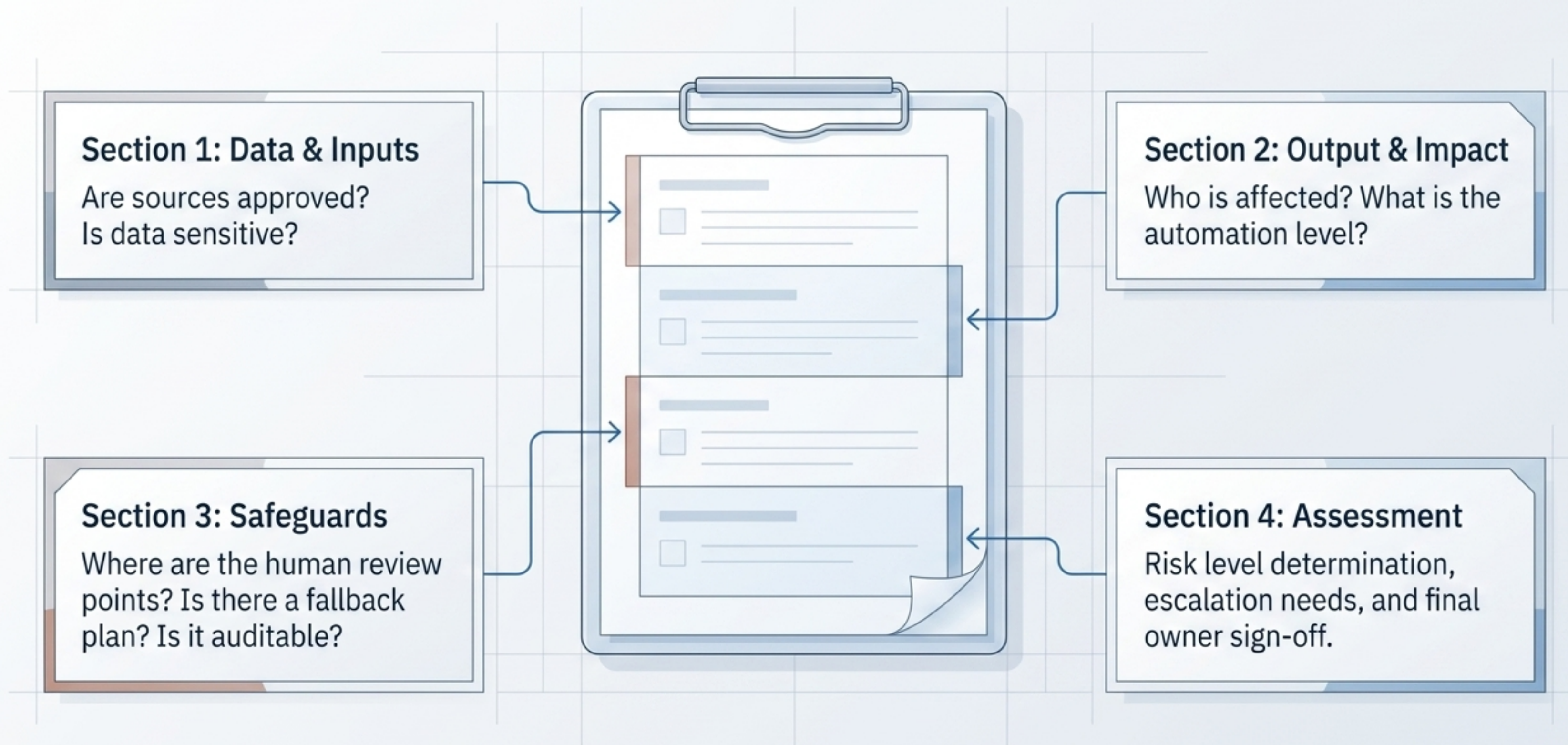
Reversibility

Can the solution be turned off without breaking the entire department?

Fallback Plan

What happens exactly if the AI hallucinates, fails an audit, or goes offline?

Capturing Risk Dimensions Before You Pilot



Classifying Real-World Scenarios

Use Case	Answers	Data	Impact	Review Point	Owner	Recommendation
Drafting patient discharge instructions.	(Requires review gate before patient sees it).					
Summarizing meeting notes (internal only).	(Low risk, ensure no sensitive HR data is included).					
Automated denial appeal letters directly to insurers.	(High risk, requires legal/finance escalation and strict review).					

Correcting Frequent Governance Assumptions

Myth

I'm using an approved tool, so no need for use case or automatic approval.

Reality

Approved tools still require approved use cases. A safe tool can still execute risky workflows.

Myth

A human will look at it.

Reality

Human review must be specifically defined, documented, and placed before the output creates impact.

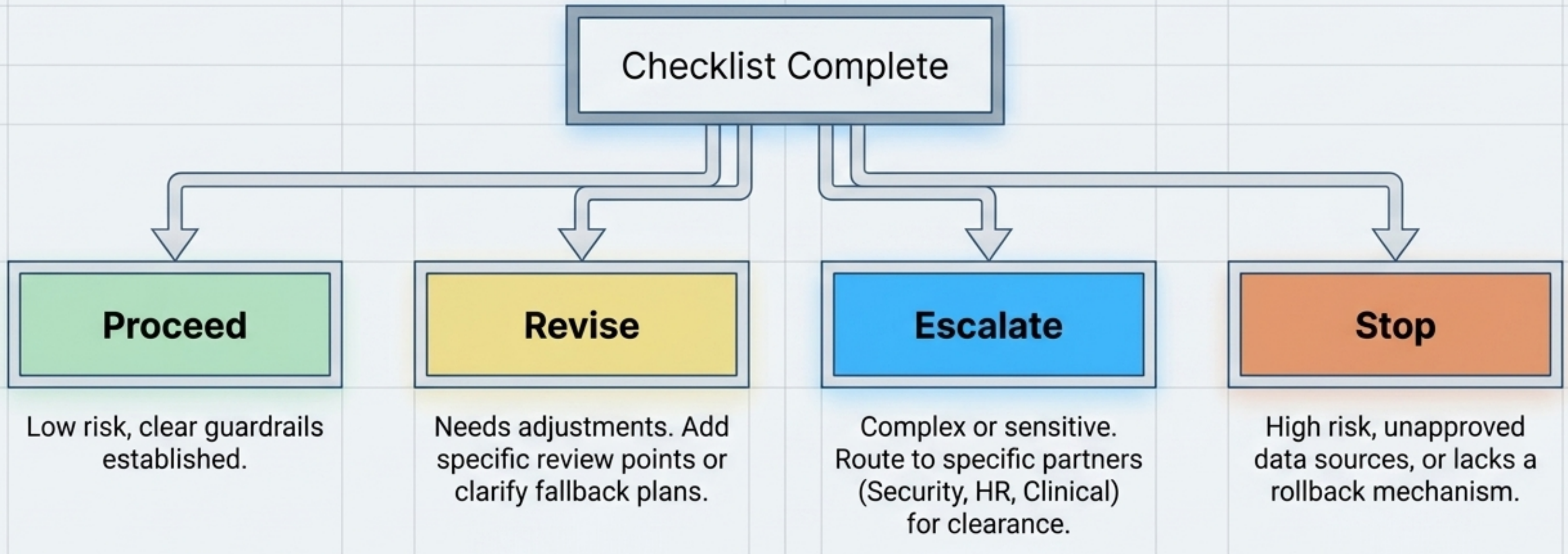
Myth

It's low-risk so no rules apply.

Reality

Even low-risk uses require clear data boundaries, approved sources, and assigned ownership.

The Final Decision Framework



Responsible AI protects patients, employees, the organization, and our credibility to keep innovating.