

GET REAL ABOUT YOUR AI GOVERNANCE

A field guide to the four key elements
of solid AI governance, and what to do
about them



Created by Jill Heinze
Principal AI Strategist & Founder
GroundSpring

© 2026 GroundSpring LLC



Jill Stover Heinze
Founder

 jill@groundspring.ai
 434-264-1819
 groundspring.ai
 /jill-stover-heinze

Thank you!

Thank you for taking the time to “get REAL” about your AI initiatives. Too often, I see innovative leaders eagerly adopt AI by integrating it into workflows and products, but fail to account for the risks. They then build these hidden hazards into their businesses that can undermine the very goals they set out to achieve.

I want to help you avoid those bad outcomes while taking advantage of AI’s potential for good.

I developed the REAL AI Governance™ method through years of working with AI-forward technology organizations to help make doing the right thing easy.

And it all starts with something as surprisingly simple as sitting down with one person who uses AI and asking three questions: would you show me how you use it, where does it feel icky, and what would make you trust it?

Everything that follows in this guide grows from conversations like that one.

I hope this helps you build credibility, confidence, and safety in your AI adoption.

Best wishes,



groundspring.ai

How AI governance accelerates innovation, confidently

Most leaders think of AI governance as a burdensome compliance necessity, which is why many put it off or ignore it altogether, risking costly impacts down the road. But the leaders getting real value from AI treat governance as the thing that lets them move fast on solid footing.

Governance is more usefully thought of as a *design challenge* that gives AI initiatives a sound basis for scaling innovations quickly.

A great example of an organization that took on this design challenge proactively is my client, Bitscopic. Bitscopic is a 40-person founder-led healthcare technology company that delivers decision-ready data for some of the largest and most complex clinical teams in the federal government. Knowing that their work in AI has serious real-life consequences in a highly regulated space, they recognized they needed to invest in their governance function.

Most AI governance gets written by lawyers or engineers who've never sat with the people using the tool. I've spent 20 years doing exactly that, which is why my method starts with people, not policy.

Within 3 weeks, I talked to all team leads to understand how data and AI were being used and where guidance was lacking. I then built a use case tracker mapped to the *NIST AI Risk Management Framework* and helped their governance group define risk tiers, assign accountabilities, and refine ambiguous and missing policies.

Two months later, AI innovation is soaring. As one leader put it, the governance work “forced a massive acceleration on ‘visioneering’ our AI strategy.”

Bitscopic saw what few others do: By designing their governance to meet the demonstrated, evidence-based needs of its staff and leadership in practical terms, they could accelerate their AI initiatives while safeguarding client trust and patient safety. A win-win-win!

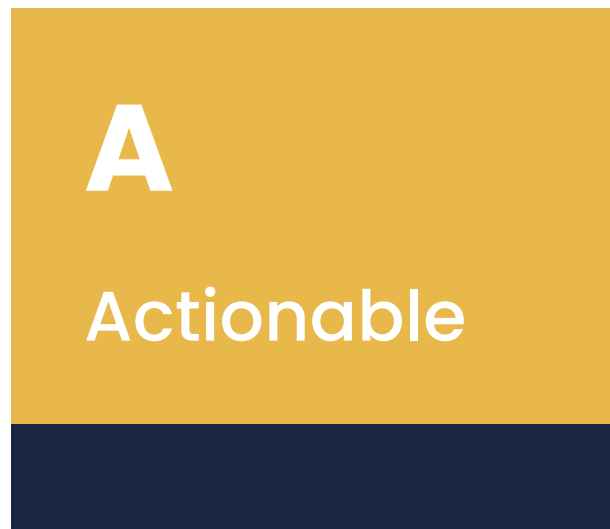
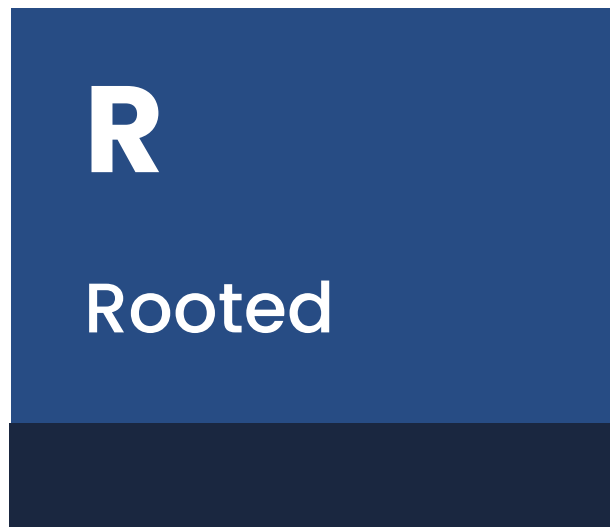
**DOWNLOAD THE FULL
BITSOPIC CASE STUDY**



REAL AI Governance™

The REAL AI Governance™ method is a lightweight, practical frame consisting of four elements to think through your own AI governance design approach.

REAL stands for Rooted, Embedded, Actionable, and Legible.



R

Rooted

What it means

Governance that's Rooted is built on what people are actually doing with AI and how they work, not on assumptions, hopes, or a policy written in a vacuum. It starts with real conversations across the entire organization paired with recognized AI risk

management standards, so that leadership can base its governance decisions on what the people closest to the work need and what regulators will want to know.

What it looks like when it's missing

Leadership issues a broad AI mandate. No one has mapped how different teams use AI day to day, so the mandate lands unevenly if at all. Some functions adopt it easily, while others find it meaningless or confusing. Torn between the pressure to adopt AI and the ambiguity about how to do it properly, people may start using personal AI accounts for work and leaking proprietary data unknowingly. Nobody finds out about the gaps until something goes wrong.

Real data

- A 2026 Gallup survey of over 23,000 U.S. employees found that among those who use AI at least a few times per year that **21%** of those in leadership roles reported AI has had an "extremely positive" effect on their own productivity, compared to just **13%** of individual contributors.¹
- Even as AI adoption rises, **72%** of leaders and workers say skill expectations have shifted, but only **36%** believe they've had adequate upskilling, and only about **33%** of frontline employees say leadership's communications about AI are clear.³
- The same survey found **uneven productivity gains across job categories**, with those in healthcare, technical, and professional roles more likely to report productivity gains, while workers in service and administrative support roles are more likely to report AI has had little or negative impact.²

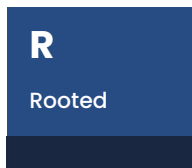


¹ Gallup, "[Rising AI Adoption Spurs Workforce Changes](#)," April 2026, based on a survey of 23,717 U.S. employees fielded Feb. 4–19, 2026. ² Ibid. ³ BCG, "[AI at Work: Strategy Matters More Than Tools](#)," June 2026.

Uneven AI results, unsupported mandates, insufficient training, and empty AI commitments undermine not only AI adoption generally, but responsible practices and stakeholder trust. No written policy or executive directive can take hold when such fundamental differences in lived experience are overlooked, because people are all operating from different base assumptions. By becoming Rooted, organizations can get everyone on the same page and address use cases and sentiment specifically and directly.

This is the cheapest risk-discovery you'll ever run. You find the shadow AI, the hidden data leaks, and the morale-draining confusion before they become an incident, all for the price of a few conversations.

The Rooted move



Before you write a single policy line, sit with one person and ask three things: would you show me how you use it, where does it feel icky, and what would make you trust it? One conversation is a start. The method scales it from there.

You move faster because you're solving the real problems, not the imagined ones.

Get started

- ☐ Pick 3 people across different teams and ask them to show you (not just tell you) how they're using AI this week. Draw their workflow and confirm it with them.
- ☐ Start a running list of every AI tool you know is in use across your org, official or not, what it's used for, and what data it consumes.
- ☐ Take one existing executive AI goal and ask: how exactly would this goal be carried out successfully? If you can't answer it in detail, flag it for revision.

Watch



Click Play to watch my YouTube video about Rooted.

Notes:

E

Embedded

What it means

Governance that's Embedded lives inside the tools and workflows people already use instead of a document they have to remember exists. If doing the right thing takes extra effort to look up, most people won't, not because they don't care, but because that's how humans work under time pressures and mandates.

What it looks like when it's missing

An AI tool or AI-enabled process goes live without anyone building in the guardrails to keep it aligned with actual company policy over time. No one defined the outcome metrics or boundaries for the AI. The AI implementation and the policy drift apart. Nobody notices the gap until a customer, regulator, or key stakeholder does.

Real case

In 2023, the Device Solutions (DS) division of Samsung authorized employee use of ChatGPT with the caveat that employees should "pay attention to internal information security and refrain from entering private content."⁴ Only a few weeks later, **the company learned of three separate incidents where employees shared sensitive trade information with the chatbot**, which was subject to being included in its training data: One employee copied an equipment measurement program's source code into ChatGPT to get help analyzing it for bugs. A second employee accidentally input the full source code of a defective equipment detection program to optimize it. A third employee supplied meeting details recorded on their smartphone to ChatGPT to create meeting minutes.

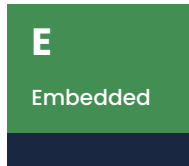
The incident led to a generative AI ban in the division.

While details about Samsung's existing governance and safeguards at the time of the incidents aren't available, this case is a reminder of how easy it is for even expert staff to make costly mistakes when robust, integrated guidance is nascent or missing. These kinds of incidents happen all the time in organizations big and small.

In fact, the World Economic Forum (WEF)'s annual Global Cybersecurity Outlook found that in 2026, data leaks rose to the most significant generative AI security concern among CEOs.⁵

⁴ Doo-yong Jeong "[Exclusive] Fears Become Reality... Samsung Electronics Sees 'Misuse' the Moment It Lifted the ChatGPT Ban," March 30, 2023. ⁵ World Economic Forum, *Global Cybersecurity Outlook 2026*, January 2026.

The Embedded move



By getting a detailed view of employees' workflows, you can better anticipate the critical decision points, including human and technical interventions, that are appropriate for adhering to your best practices.

Once you understand the most likely triggers for unacceptable uses, you can align your policies, pipelines, and monitoring accordingly.

You stop spending energy policing behavior and let your systems carry the load.

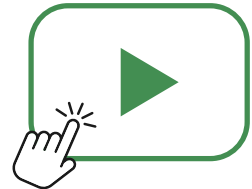
Get started

- ☐ Add "AI check-in" as a recurring 5-minute agenda item to one meeting that already happens to normalize discussing AI use, including areas of question or concern. (A little psychological safety goes a long way!⁶)
- ☐ Find one place your AI guidance currently lives only as a document, and identify the one tool or workflow where it could instead show up automatically (a checklist, a form field, a prompt).
- ☐ Consider one of your AI goals or constraints. Ask, "How would we assess if our AI is on-track?" Build this measure into your AI assessments, guardrails, or testing.

Watch



Click Play to watch my YouTube video about Embedded.



Notes:

⁶ Paula Davis, "[Why Psychological Safety Matters More In AI-Enabled Teams](#)," Forbes, May 18, 2026.

A

Actionable

What it means

Governance that's Actionable means every AI use case is mapped to a specific risk tier that has a named owner (or role) for each required action. Tiers matter because not every AI decision carries equal weight. Tiering risks means scrutiny goes only where the stakes warrant it. Decision-making steps are practical and reasonable for the level of risk that the use case

entails so that governance doesn't become cumbersome or perfunctory. When a question or issue is escalated, the responsible person knows exactly how to handle it, and only focuses on those issues that truly need it.

What it looks like when it's missing

Governance that's not actionable has some combination of having no one in charge, or having a vague group in charge, which makes finding quick resolutions difficult, and also adds a problematic "moral buffer" between employees and the impacts of their decisions. When a serious issue arises, it's nearly impossible to trace back the root causes of the failure and improve going forward.

Real research

A 2025 study in BMC Medical Ethics interviewed 40 healthcare professionals, including clinicians, administrators, and AI developers, across a National Health Service (NHS) Trust in the West Midlands, UK, about their real experience working alongside AI in clinical decision-making.⁷

Among their key findings is that clinicians reported a pervasive concern about accountability, feeling fully liable for patient outcomes even as they relied on AI-generated insights they had no way to independently verify. They also didn't have a clear answer for what should happen if the system was wrong.

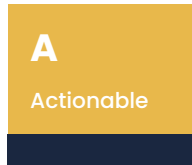
Without risk tiers telling clinicians how much to scrutinize AI-assisted decisions, clear escalation paths, or a

proportionate process, they were left to carry the heavy moral and professional weight of these decisions uneasily. Reading between the lines, that kind of gap could either slow useful AI adoption or produce inconsistent care, depending on how individual clinicians choose to handle the uncertainty.

Whether you're in healthcare or any other industry where the stakes matter, mapping actions, people, and risks serves to put responsibility in its proper place, speeding up decision-making and reducing confusion. This Actionable element alone makes governance a major accelerator!

⁷ Nouis, S.C.E., Uren, V., & Jariwala, S. (2025). *"Evaluating accountability, transparency, and bias in AI-assisted healthcare decision-making: a qualitative study of healthcare professionals' perspectives in the UK."* BMC Medical Ethics, 26, 89.

The Actionable move



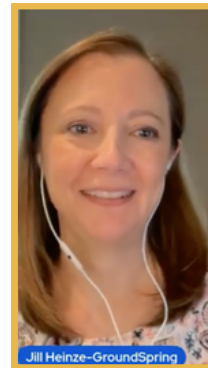
For every AI use case above your lowest risk tier, name one person by name or role who owns it, including the specific job of periodically checking what the system is actually doing against how it was designed to behave.⁸

Get started

- ☐ Identify your single highest-risk AI use case today and ask: what would need to be true for us to feel genuinely confident about it?
- ☐ Write down what actually happens if an AI system in use today produces a wrong or harmful output: who finds out first? What are they supposed to do next?

Notes:

Watch



Click Play to watch my YouTube video about Actionable.

⁸ The NIST AI Risk Management Framework breaks down how to do this triage. See [Playbook](#) sections GOVERN 1 and GOVERN 2 for starters.

L

Legible

What it means

Governance that's Legible can actually be read, understood, and repeated back by employees, by clients, by a board, or by a regulator. A policy that only leadership has seen, or that uses language so vague it could mean anything, doesn't change behavior. Legibility is a key requirement for trust and,

problematically, only **28%** of employees see a strong connection between their leaders' statements and their organizations' actions.⁹

Legibility is also the only one of the four the outside world can see. Vague policy is a liability. Legible policy is an asset you can put in a board deck, a sales conversation, or a contract as proof you're doing what you claim.

What it looks like when it's missing

A policy exists somewhere and is technically "in effect," but almost no one outside of leadership understands it, let alone could explain what it asks of them. Even when people do encounter the language, it tends to lean aspirational rather than specific. Statements like "we use AI responsibly" sound reassuring but don't tell anyone what to actually do differently. Without training, real conversation, or feedback loops to reinforce it, the policy stays a document instead of a resource people can act on.

Real example

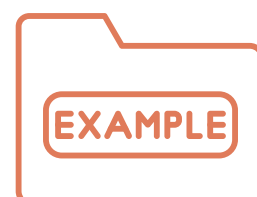
An AI policy statement for a digital consultancy prohibited employees from creating biased outputs with large language model (LLM)-based AI tools.

The policy never defined what counted as "biased." It also reflected a misunderstanding of how these tools actually work.

Bias, like hallucination, is an inherent quality of LLMs. It can show up in an output regardless of what the employee intended or how carefully they wrote their prompt.

Holding an individual employee fully accountable for a model behavior baked into the technology itself puts the responsibility in the wrong place.

A more appropriate policy would prohibit employees from using harmfully biased outputs in their work products, with "harmful bias" defined through specific examples.



⁹ BCG, "[AI at Work: Strategy Matters More Than Tools](#)," June 2026.

The Legible move



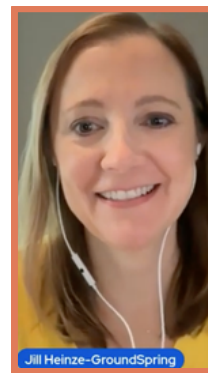
Pick one AI policy statement you currently have and try to explain it to a frontline employee in one sentence, using an example from their actual work. If you can't, the policy isn't legible yet. Specific, concrete language ("employees may not enter customer data into third-party AI tools") beats aspirational language ("we value customers' privacy") every time.

Clear governance isn't just safer. It's sellable.

Get started

- ☐ Pull up your current AI policy, if you have one, and read it as if you were a brand-new employee. Circle every sentence vague enough that you wouldn't know how to act upon it if you had to.
- ☐ Practice a 60-second explanation of your AI safety approach as if you had to give it to a client, board member, or regulator tomorrow, and notice exactly where you stumble. That stumble is your gap.
- ☐ Review any statements or mandates about how employees are expected to use AI. For each, list the knowledge, skills, and tools they would need to meet those expectations. Ensure those resources are provided.

Watch



Click Play to watch my YouTube video about Actionable.

Notes:

Give your AI efforts a reality check.

If any of the four REAL elements above feel unresolved for your organization, you're not alone.

Many reports find that AI governance is woefully behind the pace of AI adoption. **IBM found that 77% of organizations say that AI adoption has already outrun their governance capabilities.**¹⁰

This disconnect between the pace of AI and the safeguards keeping it on-track is where risks materialize.

Now is the time to close the gap to protect your organization, your customers, and your reputation.

You can start the first conversation yourself this week. When you're ready to go from one conversation to the full picture, that's where I come in.

And the fastest way to see your own picture clearly is a REAL Diagnostic.

In only 2 to 3 weeks, I sit down with up to 5 of your stakeholders, map how AI is actually being used across your organization, and deliver a gap analysis report with prioritized recommendations, walked through together in a 90-minute executive debrief.

You leave with a clear, specific picture of where your gaps are and which ones to prioritize that make sense for where your organization is with AI.



Still have questions? I'm happy to chat!
No pressure or commitment.



[BOOK A 15-MINUTE
CHAT WITH ME.](#)

¹⁰ IBM Institute for Business Value, "[New IBM Study Finds CIOs and CTOs Face Growing AI Control Gap as Enterprise Deployment Scales](#)," IBM Newsroom, June 8, 2026