
The Authority Gap

Authority begins when a human ratifies an action. The decision was shaped long before that, in a space no inherited framework was built to reach.

Every serious institution now has AI governance. Almost all of it sits at the wrong end of the decision. Authority begins when a person ratifies an action. The decision was made several steps upstream, in a space no inherited framework reaches. That space is where the outcome is actually determined, and nothing we built governs it.

AUTHOR

Shane Schreck

PUBLISHED

June 2026

CANONICAL URL<https://publications.aleeth.com/the-authority-gap/>

Publication Record

Publication	Essay
Published	June 2026
Author	Shane Schreck
Canonical URL	https://publications.aleeth.com/the-authority-gap/

In April 2026, the people who write the rules for American banks did something quietly radical. They took agentic AI out of the rulebook.

Not by accident, and not by oversight. On purpose, and in writing.

ONE QUIET RULE CHANGE

On April 17, 2026, the Federal Reserve, the OCC, and the FDIC issued SR 26-2. It replaced SR 11-7, the model-risk guidance that had governed how banks build, validate, and trust their models for fifteen years. And it drew a line around the newest systems.

"Generative AI and agentic AI models are novel and rapidly evolving. As such, they are not within the scope of this guidance." Federal Reserve, OCC, FDIC, Supervisory Letter SR 26-2, April 17, 2026

Read that carefully, because it is not an exemption. The same guidance is explicit that the institution still has to govern these systems under its general risk-management practices. The framework stepped back. The accountability stayed exactly where it was.

That is the whole frontier in one sentence. The rules that governed the math no longer reach the systems that now act on it, and the liability did not move an inch.

WHERE THE DECISION IS ACTUALLY MADE

To see why that matters, follow a single autonomous action from start to finish.

It does not begin when a person approves it. By then it is already made. An agent pulls context, weighs options, discards most of them, and narrows to one path. It drafts the action, shapes the reasoning, and decides what to surface to the human and what to leave out. Only at the very end does anything reach a person to sign.

By that point the decision is several steps old. It was shaped upstream, in the retrieval and the reasoning the person never sees. The human at the end is not deciding. They are agreeing.

Authority begins at ratification. The decision is shaped at retrieval. Between those two points lies the space where the outcome is actually determined, and it is precisely the space no inherited framework was built to govern.

OVERSIGHT AT THE FINISH LINE

Europe has named the right requirement and is about to enforce it. The EU AI Act, Article 14, takes effect for high-risk systems on August 2, 2026. It requires that a person be able to oversee the system while it is in use.

"...can be effectively overseen by natural persons during the period in which they are in use." EU AI Act, Article 14, in force August 2, 2026

The requirement is correct. The danger is where institutions will put it. Oversight is easy to place at the end, because the end is where a human already sits, waiting to click approve. So that is where most of it will go.

And a person asked to approve an action they did not shape, on evidence they cannot see, is not an overseer. They are a signature. Oversight that arrives only at ratification governs the one step where the decision has already been made.

WHY THE OLD FRAMEWORKS MISS IT

This is not the failure of any single rule. It is structural. Look at what we have.

Model risk governs the model: validate it, document it, watch it drift. But the model is no longer where the decision lives. Policy and training govern intent, before the system ever runs. Audit governs the aftermath, once the action is already history. Human-in-the-loop governs the final step, after the path is set.

Every instrument we trust sits either before the action, as intent, or after it, as review. None of them sits inside the action, at the instant it is taken. SR 26-2 admitted as much by stepping back. The frameworks built for systems that hold still cannot govern a system that decides and acts in the same motion.

THE LAYER THAT SITS IN THE SPACE

A space can only be governed from inside it. Not with advice about the decision, which is just another opinion. Not with a record written after it, which arrives too late to change the outcome. The space between retrieval and ratification has to be held by something that lives there.

That something is a control plane beneath the act. It decides, in real time, whether an action is allowed to run. It halts the action mid-stride when it crosses a line. And it signs a record of every call it made that nothing can quietly rewrite.

Here is the part that matters most. The record that proves the action was controlled is the same record that proves what the system actually did. Governing the act and proving the act are not two jobs. They are one.

This is what ALEETH is, and I will not oversell it, because a company whose product is control cannot afford to overstate its own. What runs today is the control. It returns a signed yes or no on an action before the action happens. It halts an agent mid-task and quarantines it. It writes a record no one can forge. We put it on a live offensive tool chain and watched it stop the tool the instant it reached to take data, the whole sequence sealed as evidence anyone can check. Its coverage maps to thirteen governance frameworks, and it re-seals that coverage twice a day.

Beneath the act. Above the machine. Impossible to bypass, across the surfaces it governs.

Governance that sits where the decision is not, is not governance. It is a signature on a decision already made.

SR 26-2 moved the rules and left the liability. Article 14 will ask for oversight, and most institutions will place it at the finish line. The distance between where authority begins and where the decision is actually shaped is now an open question for every institution running autonomous AI. The only real answer is to govern the space itself.

The full standard, Institutional Control Architecture, is published at publications.aleeth.com/standard . The instrument is in production today.

Not pitched. Not promised. Proven.

Shane Schreck is the founder of ALEETH and the author of Institutional Control Architecture. ALEETH is building the independent verification layer for autonomous AI deployment. He is a U.S. Army Veteran.

ALEETH

AI has met its match. Truth.™