

ARCPOINT CONSULTING · WHITE PAPER · AUGUST 2026

The AI Control Layer

Why every AI-enabled organization ends up building one, and what it costs to build it late

ArcPoint Consulting, LLC · Washington, DC · August 2026

Jermaine Leonard, AIGP · Principal

Executive summary

Most organizations have responded to AI risk by producing documents: an acceptable use policy, a set of principles, a governance committee with a charter and a quarterly meeting. None of that is a control. A control is a mechanism that can intervene in a live system, and the practical test of whether an organization has one is simple. Can anything in your organization stop an AI system from operating, over the objection of the team that owns it, and produce a record of having done so?

For most organizations the honest answer is no, and that answer is becoming expensive. In April, one major regulator held a deadline and agencies lost the use of live systems. In July, another moved its own deadline sixteen months to the right, with six days to spare. The calendar is unstable in both directions, so building to the calendar is not a strategy. The organizations that came through both events intact were the ones that had built controls rather than a compliance schedule.

What survives the calendar is the layer itself: an oversight structure that sits above AI systems rather than inside them, with four functions that either exist or do not. Visibility. Validation. Control. Auditability. This paper works through what each one requires, why the financial and cybersecurity analogues tell us how this gets adopted, and the ways organizations conclude they have a layer when they do not. It also makes the case for waiting, which is stronger than the consulting literature admits, before saying why we think it loses.

I. The claim

A control layer is defined by what it can stop.

Everything else an organization produces in the name of AI governance is commentary on control rather than control itself. A policy describes intended behavior. A principle describes intended values. A committee describes intended attention. None of them acts. The layer is the part that acts, and the reason to name it separately is that organizations consistently mistake the commentary for the thing.

Call it the stop test. Name the mechanism that can halt a consequential AI system over the objection of the team that owns it, and name the last time it was used. An organization that cannot answer the first half does not have a control layer. An organization that answers the first half and not the second has an untested one, which is better than nothing and is not the same as a working one.

One objection lands immediately, and answering it describes the architecture. Not every control stops something. An audit trail stops nothing. Neither does a disclosure, or an appeal right. But each of them stands in a fixed relationship to the stop. Visibility tells you there is something there to halt. Validation tells you when halting is warranted. Auditability is what lets you defend the decision after you have made it, and that function does more than it looks like it does, because an intervention nobody can reconstruct six months later gets reversed, and the authority that made it erodes along with it. Halting a system once is an incident. Retaining the standing to halt the next one is the layer.

The second half of the claim is structural. A control layer sits above the system it governs, not inside it.

That distinction does real work. A model vendor's built-in safety filter is a setting, configured by the same team that ships the product, adjustable by that team, reported on by that team. It may be an excellent setting. It is not a control, because it has no independent reporting line and no authority the product team does not grant it. The same goes for a governance feature in an MLOps platform, a guardrail library in the application code, and a risk register kept by the group whose deliverables it evaluates.

This is not a novel idea. It is the oldest idea in institutional risk management, applied to a new object. Both of the disciplines that AI governance is most often compared to arrived at the same architecture, and both arrived at it the same way: after a failure large enough to make the argument for them.

II. Why the question changed in the last sixteen months

What changed is easiest to see in two dates.

A deadline arrived with teeth. OMB Memorandum M-25-21, issued April 3, 2025, required federal agencies to apply seven minimum risk management practices to any high-impact AI use case, among them pre-deployment testing, an AI impact assessment, ongoing monitoring, human oversight with intervention and accountability, and a consistent remedy or appeal process. The memo gave agencies 365 days and stated the consequence plainly: if a high-impact use case is not compliant with the minimum practices, "the agency must safely discontinue use of the AI functionality." That deadline was April 3, 2026. FedScoop, canvassing 28 agencies in the days after it, found agencies that met it, agencies that reclassified use cases out of scope, and agencies that pulled systems rather than defend them. At Veterans Affairs, four high-impact use cases were rolled back from deployed to pre-deployment or retired, which

exempted them from the requirement. Whatever one thinks of the memo, this is the material fact: the absence of a control layer became a shutdown event, on a date, for real systems.

A deadline moved, and the move became law with six days to spare. The European Union's AI Act was to apply its high-risk obligations to Annex III systems from August 2, 2026. It does not. Regulation (EU) 2026/1744, the Digital Omnibus on AI, was published in the Official Journal on July 24, 2026 and entered into force on July 27, six days before the date it displaced. Annex III high-risk obligations now apply from December 2, 2027, and Annex I systems from August 2, 2028. The deferral reaches Sections 1 to 3 of Chapter III and stops there: the Article 5 prohibitions, the general-purpose AI obligations in Chapter V, and the Article 50 transparency requirements all held their original schedule. So an organization that built its program around August 2, 2026 found the target sixteen months further out, and an organization that deferred everything to that date found that a meaningful part of the regime had never moved at all. Both were wrong, in opposite directions, about the same statute.

Put the two together and you have the argument. One regulator's date held and cost organizations the use of live systems. The other's slipped and cost them the coherence of their plans. An organization that had built controls was fine in both jurisdictions. An organization that had built a compliance calendar was wrong in both.

Underneath both dates, conditions were getting harder. Stanford's 2026 AI Index records 362 AI incidents in 2025, against 233 the year before, on a base that stayed under 100 per year until 2022. The same chapter shows organizations recognizing risks well ahead of mitigating them: 63 percent considered personal privacy a relevant risk while 40 percent actively mitigated it, and 72 percent named cybersecurity while 61 percent mitigated it. On agentic deployment specifically, security and risk concerns were named as the primary obstacle to scaling by 62 percent of respondents, ahead of technical limitations and regulatory uncertainty at 38 percent each.

That last number is the interesting one, because it inverts the usual framing. The constraint on agentic AI is no longer capability, and no longer, primarily, the law. It is that organizations cannot supervise what they would be deploying. The control layer is not a tax on AI adoption. At the current frontier it is the precondition for it.

III. What the layer is, and what it is not

The layer has four functions. They are not phases, and they are not maturity stages. An organization either has each one or does not.

Visibility. A current inventory of every AI system in operation, including the ones acquired inside other software and the ones staff adopted without asking, with a named owner, the data each touches, and the decisions each influences. Visibility is where most programs fail first, and it fails quietly, because an inventory that was accurate when it was built looks identical to one that is accurate now.

The system list is the easier half and the less useful one. What governs is the decision list: for each consequential decision the organization makes, what feeds it, how far that influence runs, whether the resulting action can be undone without a human, and who answers for it. The two lists diverge quickly. Twenty tools that draft text and summarize meetings fill twenty rows on a system inventory and none on a decision inventory, which does not make them harmless, because they remain a real data question, but attention paid to them does not govern a decision. One system that sets a price, screens an applicant, or releases a payment is a single row on the first list and the only row that matters on the second. A board that asks how many AI systems the organization runs has asked for the first list. The number it wanted was on the second.

Validation. Testing that happens before deployment and again on a defined cadence, against criteria written down in advance, producing an artifact someone outside the building team can read. The distinction that matters here is between a benchmark and a validation. A benchmark tells you about a model. A validation tells you about your deployment of it: your prompt, your retrieval corpus, your business rules, your filters, your version of every layer as of a specific date.

Control. Named human accountability for each consequential decision, an enforcement point where policy actually binds rather than advises, and the authority to use it against a system in production. This is the function most often present on paper and absent in fact. The test again: who can stop this, without asking the owner, and has that ever happened?

Control is also set at a level, and the level is a delegation the organization chose rather than a property of the technology. A system that informs leaves the decision with a person. One that recommends puts a default in front of that person, which is not the same thing, because a default that is accepted almost every time is a decision the system is making with a human signature attached. One that decides settles the outcome inside stated parameters. One that acts carries the outcome into another system, and that is the line worth drawing, because it is where reversal stops being cheap. What sorts the four is not model capability. It is whether the action can be undone without a human, and if it cannot, who holds the authority to reverse it and how long that takes. The assistant that drafts the refund email and the agent that issues the refund can run on the same model. Only one of them needs a brake.

Auditability. Evidence produced as a by-product of the system operating, not assembled after someone asks for it. If your evidence is generated by the request for evidence, you do not have

auditability. You have a research project with a deadline. It is also, for the reason given in Section I, what keeps the other three from decaying.

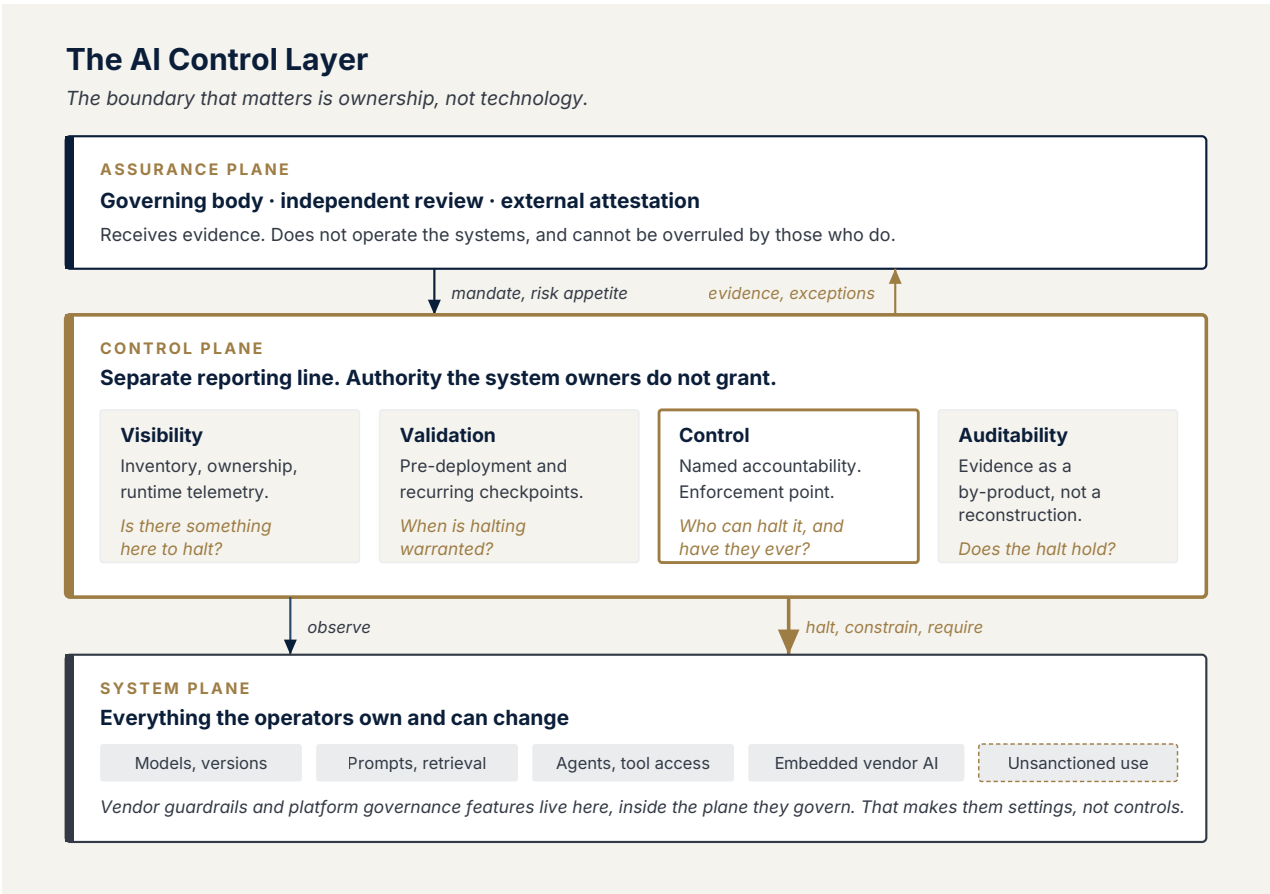


Figure 1. The boundary that matters is ownership, not technology. Vendor guardrails and platform governance features sit inside the plane they govern, which is precisely what makes them settings rather than controls.

The layer is not a committee. A committee is a venue in which control can be exercised, which is useful and not sufficient. A governance board with no enforcement point beneath it produces minutes.

Nor is it a document, which is the failure we see most. It has a distinctive signature: an organization can produce a policy for every requirement and cannot produce a single instance of a policy changing an outcome.

IV. What the analogues actually teach

The comparison to financial controls and cybersecurity is usually made loosely, as a way of saying that AI governance is important. Made precisely, both analogues predict specific things about how this will go.

Financial oversight. Sarbanes-Oxley Section 404, effective August 2003, requires management to report annually on the effectiveness of internal control over financial reporting, evaluated against a recognized framework and attested to by the external auditor. The instructive part is the architecture, not the requirement. SOX did not improve accounting software. It built a structure above the accounting function, with a separate reporting line and an outside attestation, on the premise that the function could not credibly assess itself. The Institute of Internal Auditors formalized the same premise in the Three Lines Model, updated in July 2020, which preserves internal audit's independence from management because that independence is, in its words, "critical to its objectivity, authority, and credibility."

Cybersecurity. When NIST published Cybersecurity Framework 2.0 on February 26, 2024, the substantive change was the addition of a sixth function, GOVERN, defined as the function under which "the organization's cybersecurity risk management strategy, expectations, and policy are established, communicated, and monitored." NIST placed it at the center of the wheel because it "informs how an organization will implement the other five Functions." The NIST AI Risk Management Framework had already reached the same conclusion, describing GOVERN as cross-cutting and "infused throughout AI risk management," with governance requiring "continual and intrinsic" attention across a system's lifespan.

What carries over is this.

The layer is defined by separation. In both disciplines, the operative move was to give oversight a reporting line the operators do not control. Nothing about AI changes this, and everything about AI's current organizational placement, typically inside the function being asked to show returns on it, argues that it is harder here.

The layer gets built after the failure, at a multiple of the cost. Neither SOX nor the modern security operations function was a voluntary innovation. Both were retrofits, priced accordingly. The organizations that came through those transitions cheaply were the ones that had already built most of the structure for their own reasons.

The layer becomes an entry requirement, not a differentiator. No enterprise wins business by having internal controls over financial reporting, and none is permitted to operate a public company without them. The same transition is already visible in AI procurement.

V. The worked example: what a buyer has already specified

The District of Columbia is a useful case, not because it is the largest market but because its requirements are published, specific, and further along than most.

Mayor's Order 2024-028 establishes six AI values and attaches a per-deployment duty: every AI deployment by a District agency carries a documented obligation, and §III.C assigns human accountability for every action or determination. That order does not sit alone. The OCTO AI/ML Governance Policy is binding, not advisory, and carries obligations the six values do not cover. The AI Procurement Handbook adds a set of standard contract clauses, of which clause E.2 requires the contractor to demonstrate that the AI system and its usage, deployment, and maintenance "do not conflict with Mayor's Order 2024-028, the District's AI/ML Governance Policy, or the District's Handbook for AI Values Alignment." Conformance with one of the three is not conformance.

Clause E.4 is the part worth reading closely, because it is a control layer written as a contract term. It imposes ten operational requirements on the contractor, among them a means of performance monitoring, evidence of bias management, a human override capability, explainability documentation, ongoing monitoring, staff training, and audit rights permitting the District to observe or audit relevant work processes at no additional cost.

Read that list against the four functions. Visibility, validation, control, auditability, all four, specified by a buyer, in a clause a vendor signs. The District did not ask its vendors for a policy. It asked for the layer, and it reserved the right to come look at it.

Sophisticated buyers are already past the policy stage, and the gap between what is contractually required and what vendors have actually built is where the commercial risk now sits. A signature on clause E.4 is a representation about an operating capability. Most vendors signing it are representing an intention.

One point is DC's alone, and it carries a date. The AI Taskforce and Advisory Group both sunset December 31, 2026, under §VI.J and §IV.J of the same order. The per-deployment mandate in §III does not sunset. It is permanent. So the support structure agencies have leaned on for interpretation goes away while the obligation stays, which means that from January 1, 2027, the layer is either internal to the agency or supplied under contract. There is no third option short of not deploying, and the window for choosing is measured in months.

VI. Five ways organizations conclude they have a layer when they do not

These are patterns we see in assessment work, stated generically. Each is a way of passing a document review and failing an operating review.

The future-tense commitment presented as an existing control. A governance document states that notice “will be” provided, that an incident-reporting address “will be” established, that limits “will be” recorded in an agreement when finalized, and the surrounding structure treats the statement as satisfying the requirement. A commitment is not a control. The diagnostic that makes this concrete is arithmetic: compute the elapsed time between the system going live and the document being published, and state it. A control promised in a document published a year after go-live has had a year not to exist.

The capability stated where a trigger is required. “We are able to deactivate the system” is a statement about potential. A control specifies the condition under which deactivation happens, who decides, and on what timeline. The two failures pair, and they frequently appear in the same document.

Baseline compliance presented as exceptional control. An organization completes the mandatory training, files the required inventory, adopts the published framework, and presents this as evidence of a mature program. But the floor is the floor. The question worth asking is what the organization does that nobody required of it, and whether any of that would have caught something the mandatory work would have missed.

Contract rights held but never exercised. The audit right is in the agreement. The reporting obligation is in the agreement. Neither has ever been invoked. An unexercised right is indistinguishable, in operation and at audit, from a right that does not exist, and vendors calibrate to what buyers actually do.

Complete against one instrument, silent on the others. A governance artifact is worked through carefully against the framework everyone knows about, and says nothing at all about the binding policy sitting beside it: approval authority, governance board review, pre-use cybersecurity and business risk assessment, how exceptions get granted and by whom. The document is not weak. It is complete, which is exactly what makes the omission hard to see, because a reviewer checking it against the instrument it was written for will find nothing wrong with it. This is the most common failure we find in governance work that has already been done by someone else, and it follows directly from assessing the headline order instead of the whole published obligation set.

If more than one of these describes your program, the useful next step is not more policy. It is to find one enforcement point and make it real.

VII. The case for waiting, which is better than it sounds

The honest objection to everything above is that the European Union just moved a deadline sixteen months to the right, and organizations that spent 2025 building toward it are now sixteen months early with capital they could have deployed elsewhere. In a field where the tooling gets cheaper and more capable every quarter, being early is not obviously a virtue. Waiting is a defensible use of money, and treating it as negligence is the sort of thing consultants say when urgency is the product. We sell this work. Weigh the argument accordingly.

Here is the answer, with a qualification and a concession attached.

The two errors are not the same size. Building early costs carrying cost, and carrying cost is knowable in advance and payable in installments. Building late costs the use of the system, on a date somebody else chooses, with no notice period in which to negotiate. April was not hypothetical. The asymmetry is the whole argument. One error is a line item. The other is an outage.

The deferral was narrower than the headline. Three sections of Chapter III moved. The prohibitions did not. The general-purpose AI obligations did not. The Article 50 transparency requirements did not. An organization that read “the AI Act is delayed” and stood down was wrong about its own regulator, in the same month, on the same statute. Reading a deferral as a reprieve requires knowing precisely what was deferred.

Most of the layer pays for itself outside compliance. The clearest evidence is the agentic figure from Stanford: security and risk concerns, not capability and not regulation, were the primary obstacle to scaling for 62 percent of respondents. An organization that cannot see, validate, halt, and evidence its AI is an organization that cannot safely expand it, which makes the layer not a tax on adoption but the constraint currently binding it. The same capability answers a customer’s security review, a partner’s diligence questionnaire, and an insurer’s underwriting question, none of which appear on any regulator’s calendar.

If you do decide to spend now, know what the money can and cannot buy. A market is forming: Gartner expects spending on AI governance to reach \$492 million in 2026, driven by regulatory fragmentation it predicts will quadruple and extend to 75 percent of the world’s economies by 2030. Tooling does real work in visibility and auditability, which are instrumentation problems,

and doing those by hand does not scale past a handful of systems. We are vendor-neutral, and none of this argues against buying. It argues about limits. A platform can tell you what your systems are doing. It cannot tell you who is permitted to stop one, and it cannot create the authority to do so. Those are organizational facts, and no procurement produces them. Buying a governance platform without settling decision authority means fitting an instrument panel to a vehicle with no brake pedal. The instruments will be accurate. That is the problem: they will accurately display a situation nobody has the standing to change.

And the concession. If nothing you run with AI touches a person's access to money, safety, employment, benefits, housing, or legal standing, this paper is premature for you, and the proportionate response is an inventory, one named person who can switch things off, and a note in the calendar to revisit when that changes. A control layer scales down further than most of this literature admits. A forty-person firm's version is a spreadsheet, a named owner, and a quarterly review that produces a written record, which is a real layer even though it would not impress anyone. But one part of it does not scale down at all. There is no size of organization at which it is acceptable that nobody has the standing to say no.

VIII. Where to start, if you are starting now

Three questions, in this order. They are diagnostic rather than procedural, and any organization can answer them in a week.

One. Name every AI system currently influencing a consequential decision, and the person accountable for each. Not the systems you procured. The systems operating. The gap between those two lists is your visibility gap, and it is usually larger than the team expects, because it includes AI embedded in software bought for other reasons and tools staff adopted independently.

Two. Apply the stop test to the most consequential of them. Who can halt it, and when did they last do so? If the answer is the team that built it, you have a setting. If the answer is a named person outside that team who has never used the authority, you have an untested control.

Three. Pick one obligation you are already under and try to produce the evidence today. A contract clause, a regulatory requirement, an internal commitment. Set a two-hour limit. What you can produce in two hours is your auditability. What takes three weeks to assemble is a reconstruction, and reconstructions do not hold up when the person asking has a reason to doubt you.

Those three answers map which of the four functions you actually have. In our experience the results are uneven rather than uniformly bad, and the unevenness is the useful finding, because the weakest function determines what the layer can do. A program with excellent monitoring and no enforcement authority will observe its own incident in high resolution.

IX. What happens next

The regulatory calendar will keep moving in both directions, and there is no reason to expect the next three years to be more orderly than the last three. Building a compliance program indexed to a date means rebuilding it every time the date changes, and it means an organization's actual risk posture becomes a side effect of legislative scheduling.

The layer does not have that property. Visibility, validation, control, and auditability are worth having on the merits, they satisfy most of what any of these regimes ask for, and they are indifferent to when the asking happens. Build to the layer, not to the deadline.

Inside the organization, the layer changes managerial work before it changes anything else, and it is usually discovered rather than designed. Approval shrinks. A supervisor who signed individual transactions ends up setting the thresholds at which transactions route for signature, watching what arrives, and judging when the routing rule itself was wrong. That is a different job, performed against different evidence, and organizations go on measuring the old one. It is also where a working layer quietly stops working, because thresholds drift toward whatever the reviewing population has time for, and nobody files an incident report about a threshold that moved.

The transition to watch is the one both analogues already went through. Internal financial controls and security operations each started as a differentiator, became an expectation, and ended as a condition of being allowed to operate at all. AI is moving along that path faster, and buyers are ahead of the regulators: the District wrote the control layer into a contract clause before anyone was required to have one.

So the question that decides who is still running consequential AI in regulated environments at the end of this decade is not whose models are best. Everyone's will be good. It is who can turn one off, on the record, and show you why.

About ArcPoint

ArcPoint Consulting is a senior-led AI governance, risk, and control advisory practice in Washington, DC. We do not help organizations adopt AI. We help them govern it.

We build the AI Control Layer: the evidence, oversight, and documentation that make an AI deployment visible, validated, controlled, and auditable. It is produced by ArcPoint's Weftline Framework, which assesses five dimensions together, Strategic Intent, Decision Authority, Operational Integration, Workforce Enablement, and Technical Infrastructure, with AI Influence Mapping running across all five, building on the NIST AI Risk Management Framework and adding structured ethical value elicitation drawn from IEEE 7000 and a defined bias-testing protocol. Jurisdiction-specific requirements layer on top.

Every engagement is delivered and signed by the principal.

Jermaine Leonard, AIGP · ArcPoint Consulting, LLC · Washington, DC Certified Business Enterprise, District of Columbia, six designations · UEI JE2DLNX4BVQ3

This paper is analysis and general professional commentary. It is not a compliance assessment of any organization, and it is not legal advice.

Sources and verification

/verify pass run 2026-08-21. Every citation below was re-checked live against the primary source on that date, not against the search result or the secondary analysis that first surfaced it. Four corrections resulted and are noted in the table. Status key: **P** = primary source read directly · **S** = secondary, attributed as reporting in the text.

#	CLAIM	PRIMARY SOURCE	STATUS	RESULT
1	M-25-21 issued April 3, 2025; high-impact AI definition; seven minimum risk management practices; 365-day deadline; "the agency must safely discontinue use of the AI functionality"	OMB Memorandum M-25-21, whitehouse.gov	P	Corrected. Draft said "six" practices (seven) and quoted "must be discontinued," which is not language in the memo. Both fixed; the quotation is now verbatim
2	April 3, 2026 deadline; 28 agencies canvassed; VA rolled four high-impact use cases back from deployed to pre-deployment or retired, exempting them	FedScoop, Lindsey Wilkinson, April 9, 2026	S	Corrected and attributed. Department now named (public reporting); the "exempted them" mechanism added, since reclassification is the point

#	CLAIM	PRIMARY SOURCE	STATUS	RESULT
3	Regulation (EU) 2026/1744 of 8 July 2026 (Digital Omnibus on AI); OJ 24 July 2026; in force on the third day after publication; Annex III to 2 December 2027, Annex I to 2 August 2028; deferral confined to Sections 1, 2 and 3 of Chapter III	Regulation (EU) 2026/1744, EUR-Lex, full text	P	Corrected. Draft said the deferral reaches "Chapter III"; it reaches Sections 1 to 3 only. Now cited to the OJ text rather than law-firm analyses
4	362 AI incidents in 2025 vs 233 in 2024, base under 100/year until 2022; privacy 63/40 and cybersecurity 72/61 recognition-vs-mitigation gaps; 62% name security and risk as the primary obstacle to scaling agentic AI, vs 38% technical and 38% regulatory	Stanford HAI, 2026 <i>AI Index Report</i> , Chapter 3 (chapter PDF)	P	Confirmed
5	NIST AI RMF: GOVERN is cross-cutting, "infused throughout AI risk management," requiring "continual and intrinsic" attention	NIST AI RMF 1.0 Core, airc.nist.gov	P	Confirmed
6	NIST CSF 2.0 published February 26, 2024; six Functions; GOVERN definition; GOVERN "informs how an organization will implement the other five Functions"	NIST CSWP 29 (PDF)	P	Confirmed
7	SOX §404 effective August 14, 2003; management assessment of ICFR against a suitable recognized framework; auditor attestation	SEC final rule, <i>Management's Report on Internal Control Over Financial Reporting</i>	P	Confirmed
8	Three Lines Model published July 2020; internal audit independence "critical to its objectivity, authority, and credibility"	Institute of Internal Auditors, <i>The IIA's Three Lines Model</i> (PDF)	P	Confirmed
9	AI Procurement Handbook clause E.2 conformance language (quoted verbatim); clause E.4 ten contractor requirements including performance monitoring, bias evidence, human override, explainability, ongoing monitoring, training, and audit rights at no additional cost	DC OCTO, <i>AI Procurement Handbook</i> , techplan.dc.gov (PDF)	P	Confirmed
10	MO 2024-028 §III per-deployment duty; §III.C "human accountability for every action or determination"; §IV.J and §VI.J sunset of the Advisory Group and AI Taskforce on December 31, 2026	Mayor's Order 2024-028, techplan.dc.gov	P	Confirmed by live re-fetch. Supersedes the vault's 2026-07-06 / 07-31 provenance, which carried confidence: medium pending owner review

#	CLAIM	PRIMARY SOURCE	STATUS	RESULT
11	OCTO AI/ML Governance Policy is binding, not advisory: authority under DC Official Code § 1-1401 et seq. to "write and enforce IT policies"; reaches providers and third parties; §4.1.3 Agency Director written approval; §4.1.7 cyber and business risk assessment; exceptions escalate to the OCTO CISO; violators "may be denied access... and may be subject to other disciplinary action"	DC OCTO AI/ML Governance Policy, octo.dc.gov	P	Confirmed by live re-fetch. Supersedes vault provenance
12	Spending on AI governance expected to reach \$492 million in 2026; fragmented AI regulation predicted to quadruple and extend to 75% of the world's economies by 2030	Gartner press release, February 17, 2026	P	Corrected. Draft conflated the \$492M governance figure with a separate "\$1 billion in total compliance spend" figure whose horizon is not stated unambiguously in the release. The \$1B claim is now dropped rather than hedged. The release's "3.4x more effective" claim remains deliberately unused: the effectiveness construct is self-reported and undefined

Finding patterns in Section VI are ArcPoint's own, harvested from engagement work and recorded in the practice's observed-practice baseline and finding components: future-tense commitment (baseline B-03), capability where a trigger is required (B-05), complete against one instrument and silent on the others (B-08), plus the baseline-compliance and unexercised-contract-rights components. All five are stated as generic patterns and contain no client-identifying material.