



The Case For Mandatory AI Ethics Audits In Government Utilization

Jack Skager Berg

Attorney at Law, Certified AI Leader, USA.

***Corresponding Author:** Jack Skager Berg, Attorney at Law, Certified AI Leader, USA.

Submitted: 29 August 2026

Accepted: 05 September 2026

Published: 12 September 2026

Citation: Berg, J. S. (2026). The Case For Mandatory AI Ethics Audits In Government Utilization. J of Research and Reviews in Economics and Management. 2(5), 01-11.

Abstract

This white paper examines the need for mandatory AI ethics audits in government decision-making. It proposes a framework based on transparency, fairness, accountability, human oversight, and continuous review to ensure responsible and ethical use of AI while protecting individual rights and democratic values.

Keywords: Artificial Intelligence (AI); AI Ethics, Ethics Audits, Government AI, Algorithmic Governance, AI Accountability, Transparency, Fairness, Human Oversight, Due Process.

Executive Summary

In all of my years in the law, forty-five to be exact, I have seen and used all of the technological advancements available, including the manual typewriter, the IBM Selectrix with the little jumping typing ball, the word processor, the computer, electronic court systems and even the mobile phone. My first one weighed about eight pounds with its own case and a curly Q pigtail cord attached to the handset. You get the picture. I am an old man with a young heart (thank you Dr. Hoenicke) who is still trying to attain wisdom.

With all that being said, I have never witnessed technologically anything like the advancement of artificial intelligence. With its roots going back to the 1940's it is now completely out of the womb and existing amongst us in nearly every facet of our lives. For instance, from a practical standpoint, the algorithmic systems that are determining whether veterans receive benefits, whether small business owners qualify for federal contracts, whether defendants are granted bail, and whether families are flagged by child welfare agencies all have AI systems running right now, 24 hours a day. Today. And in too many cases, the people those systems are judging do not know how they were evaluated, cannot get a plain-language explanation, and have no meaningful path to contest the outcome. I have sat across the table from those people mostly in a pro bono legal capacity. But I have also seen large corporations and government agency clients face the same problem. I have witnessed the frustration and

dismay. It is not abstract to me.

This white paper is my attempt as a practicing attorney and as a Certified AI Leader to make the case for mandatory AI ethics audits in government utilization. Not voluntary guidelines. Not aspirational frameworks. Mandatory, independent, publicly disclosed audits with real accountability behind them. The philosophical case for this is strong. The constitutional case is clear. The practical case is urgent. And the bipartisan coalition willing to act is closer to a majority than Washington's current inertia suggests.

The governing principle of everything that follows is this: Keep the Humane in the AI Loop. I mean that not as a feelgood slogan but as a structural constitutional requirement. One that flows directly from the Fifth and Fourteenth Amendments, from the founding compact of this Republic, and from more than two thousand years of Western political philosophy. When government exercises power over persons, those persons are entitled to an explanation, an appeal, and a human being in the equation who takes responsibility. That is the American deal. It does not expire because the decision-maker happens to be code. It requires ethics, ethical considerations, inherent moral principles, ethic audits and, most of all, ethical intelligence or what I refer to as "EI." AI comes from the mind. EI comes from the heart.

This Paper Makes Four Core Claims, and I will State Them

Plainly Up Front

First, government AI systems that make consequential decisions without transparency, appeal rights, or human accountability are not just bad policy, they are constitutionally suspect. The Fifth Amendment's due process guarantee and the Fourteenth Amendment's equal protection clause do not contain carve-outs for algorithmic decisions. If a government system can deny your benefits without giving you a reason you can understand or contest, that system has a due process problem. It is that simple.

Second, this is not a novel insight. The philosophical tradition of the West, from Plato's insistence that the ruler must give an account, to Aquinas's natural law requirement that authority serve the common good, to Leo XIII's subsidiarity principle, to Will Rogers's plain-spoken defense of the ordinary American's right to understand what his government is doing, all of it converges on the same point. Government power exercised over persons must be made accountable to those persons. Every age has had to rediscover this truth in the specific form its own technologies require. Ours is no different.

Third, the concern about unaccountable algorithmic government is not a left issue or a right issue. It's an American issue. Conservatives who have spent generations fighting arbitrary bureaucratic power should be more alarmed about algorithmic governance than anyone, at least a human bureaucrat can be questioned, fired, or hauled before Congress. Progressives who have documented algorithmic discrimination in recidivism scoring, resume screening, and healthcare allocation know that the most vulnerable communities are bearing the cost of this opacity. The common ground here is wider than Washington would have you believe.

Fourth, mandatory AI ethics audits are not just theoretically desirable, they are realistically achievable. The EU AI Act has demonstrated that risk-tiered governance works without destroying innovation. The NIST AI Risk Management Framework has given us a rigorous voluntary foundation that simply needs the teeth of legal obligation behind it. States including Colorado, Illinois, and California are already moving. The policy infrastructure exists. What has been missing is the federal will.

The audit framework I propose rests on five pillars: Transparency: all government AI must be explainable to the citizens it affects. Fairness: audits must test for disparate impact across protected characteristics. Accountability: a qualified, trained and named human official must be legally responsible for every AI-assisted decision. Redress: any adversely affected citizen must have a meaningful right to appeal to a human decision-maker. And Continuous Review: ethics audits (internal and external) must be ongoing, not one-time events. These are not radical demands. They are the same standards we apply to every other consequential exercise of governmental power. Think of the non-AI sections of the IRS.

To make these pillars operational, I propose three institutional structures: an Office of AI Ethics Accountability (OAEA) within

GAO or OMB; an independent AI Ethics Audit Corps (AEAC) modeled on the inspector general community; and, a Public AI Audit Registry (AAR) making all audit findings available to any citizen, journalist, or researcher who needs them. These institutions are not bureaucratic overhead. This is democracy at work.

Why am I writing this? Because I have sat with clients who received a code where they deserved an explanation. Because I have read the philosophical literature and the constitutional history, and I am convinced the argument is sound. And because I believe genuinely, not performatively, that America has risen to challenges like this before. We brought the administrative state under the rule of law in 1946. We extended civil rights protections in 1964. We built FOIA in 1966. The accountability structures we build now, or fail to build, will define this generation's relationship with its government. I think we can get this right. But we have to decide to do so.

Introduction

A few years ago, I sat across from a friend and pro bono client, The Weathervane (breast cancer) Foundation, located in Houston the area, which had just received a denial letter from a federal procurement system. The letter gave my friend, Fran, a code. No explanation. No name. No face. She slid the paper across the desk to me and asked, in a voice that was equal parts frustration and genuine bewilderment, what had she done wrong? I did not know what to tell her, because neither of us could understand the system that had judged her. We could not identify what variable had tripped the denial. We could not find a human being to call. The system had rendered its verdict, and it was not interested in questions.

That moment stuck with me. Not just as a professional frustration, though it was certainly that, but as a signal. Something was wrong, structurally wrong, about a system of government that could make a consequential decision affecting an altruistic charity's funding and not be required to explain itself in any language a human being could understand or challenge. What I saw in that office was not a technical glitch. It was a governance failure.

That experience is not unique to my practice, and it is not unique to procurement. AI systems are already embedded in American government at every level. They are determining who receives Social Security disability benefits and who gets flagged for review. They are generating risk scores that influence bail and sentencing decisions. They are screening job applicants for federal positions. They are flagging families for child protective services investigations. They are scoring contractors, ranking grant applications, and allocating healthcare resources in public programs. The AI-in-government story is not coming, it's here. It arrived while we were still debating whether it was a good idea.¹ And here is the accountability gap that concerns me most: the legal and institutional structures that democratic governance demands have not kept pace. The Administrative Procedure Act was written for a world of human administrators making human decisions. The Freedom of Information Act assumes documents that can be read and understood. The due process rights enshrined in

the Fifth and Fourteenth Amendments presuppose that a person who has been adversely affected can, at some level, understand what happened to them and contest it. When the decision-maker is an algorithm, whose logic is either proprietary, mathematically opaque, or buried in training data no one has audited for bias, all of those structural guarantees begin to hollow out.

This paper makes four core arguments. First: government AI without meaningful ethics audits is constitutionally vulnerable. Second: two and a half millennia of Western philosophical tradition from Plato to Aquinas to Leo XIII to Will Rogers converge on the principle that power over persons must be accountable to those persons. Third: concern about unaccountable algorithmic governance is not a partisan position. It is an American one. And fourth: mandatory AI ethics audits are not a theoretical aspiration, they are a practical, achievable, urgently needed reform.

The governing principle that animates everything I argue here is simple enough to say in six words: Keep the Humane in the Loop. I mean it structurally. I mean it constitutionally. I mean it as both attorney and as someone who has earned AI leadership recognition. Someone who has studied these systems closely enough to understand their genuine promise and seen their failures up close enough to know the cost when we deploy them without accountability.

What I have not done is become a pessimist about technology. I think AI can make government more accurate, faster, and more equitable if we govern it right. The question this paper addresses is not whether to use AI in government. That ship has sailed. The question is whether we are going to build the accountability infrastructure those systems require before the harms are so widespread and the systems so entrenched that meaningful reform takes a generation.

America has risen to that challenge before. We have brought power to heel. administrative, corporate, governmental, when it threatened the constitutional compact that makes self-governance possible [1]. This paper argues, with some urgency, that we must rise to it again.

Section: 1

The American Covenant and Algorithmic Power

Start with the Preamble. "We the People of the United States, in Order to form a more perfect Union, establish Justice, ensure domestic Tranquility, provide for the common defence, promote the general Welfare, and secure the Blessings of Liberty to ourselves and our Posterity, do ordain and establish this Constitution for the United States of America." I've read that sentence hundreds of times. What strikes me now, in the context of algorithmic governance, is the first word: We. Not the executive. Not the legislature [1].

Not an algorithm. We. The Constitution is a covenant, a compact between a people and the government they create. And the central term of that compact is accountability. Government derives its legitimacy from the people, answers to the people, and

must be comprehensible to the people it serves. That covenant does not dissolve because the decision-maker is code rather than a civil servant. And yet, in agency after agency, algorithmic systems are making consequential determinations: who gets benefits, who gets a contract, who gets flagged for additional scrutiny, with no requirement that the logic be explained, the outputs be audited, or any named human take responsibility for the result. That is not a covenant. That is a machine operating in a constitutional vacuum.

The Constitutional Bedrock

Please allow me to be direct about the legal framework I return to in practice because it matters here. The Fifth Amendment guarantees that no person shall "be deprived of life, liberty, or property, without due process of law [2]. The Fourteenth Amendment extends that guarantee and adds equal protection of the law to the states. These are not abstract principles. In practice, they mean that when government takes an adverse action against you, you are entitled to notice, an opportunity to be heard, and a decision grounded in articulable reasons [3].

Here is the question that keeps me up at night: how can anyone meaningfully appeal a decision they can't understand? If a federal system denies your disability claim or any claim and gives you a code, not an explanation, have you received the constitutional process you are due? If a recidivism algorithm scores you as high-risk and the variables that drove that score are a trade secret legally protected from disclosure, can you meaningfully contest it? If a procurement ranking system flags your business as noncompliant based on weighted inputs no one will reveal, what does your right to appeal actually mean?

These are not hypotheticals. These are the operational realities of algorithmic governance as it is currently practiced. And my answer, grounded in the constitutional text and the case law that has built up around it is that due process and equal protection rights are constitutionally suspect when the decision-making mechanism is opaque by design.

Consent Requires Comprehension

The Declaration of Independence grounds the legitimacy of government in the consent of the governed. And consent, meaningful consent, not the fictional kind, requires comprehension [4]. A citizen cannot meaningfully consent to a power she cannot understand or contest. When government operates through systems whose logic is proprietary or mathematically inaccessible, the consent of the governed becomes a legal fiction. You can vote for the legislators who funded the agency, and the agency can deploy a system that affects your life in ways no elected official has examined, and the loop between citizen and government, the accountability loop, is broken. That is not a technical problem. It is a constitutional one.

America Has Done This Before

I want to push back, gently but firmly, on the thought that treating AI accountability as an urgent legal priority is somehow unprecedented or radical. It is not. America has a long, hard-won

tradition of bringing new forms of power under the rule of law.

The Administrative Procedure Act of 1946 was enacted precisely because the administrative state had grown powerful enough to affect millions of lives without adequate standards for fairness, transparency, or appeal. Congress looked at what the administrative agencies had become and said: this is not good enough. There have to be rules. Decisions have to be reasoned. Citizens have to be able to contest them [5].

The Civil Rights Act of 1964 confronted a different form of systemic harm, one embedded in institutions and practices that had the appearance of neutrality but produced discriminatory outcomes. Sound familiar? The documented disparate impact of algorithmic systems in recidivism scoring, facial recognition, and healthcare allocation is the civil rights challenge of our particular moment [6].

The Freedom of Information Act of 1966 extended the principle of government transparency to the documentary record of executive action. A searchable, publicly accessible registry of AI ethics audit reports is, at its core, the FOIA principle applied to the algorithmic age [7].

Each of these was, in its moment, contested. Each was called burdensome, impractical, or premature. Each became a pillar of accountable governance we would never want to give up. Mandatory AI ethics audits belong in that tradition, not as a departure from American values, but as their extension.

Section: 2

The Philosophical Inheritance: From Athens to Washington

I will admit that when I first started thinking seriously about AI governance, not just the regulatory mechanics, but the underlying why, I didn't expect to find myself reaching for Plato. It seemed like a long walk from fourth-century Athens to a federal procurement database in 2026. But the more I worked through the problem, the more I kept arriving at the same insight: the thinkers who have wrestled most seriously with the relationship between power and accountability are more useful here than almost anything written in the last twenty years. The questions they asked are exactly our questions. They just did not have the specific technology in front of them.

What follows is not a philosophy lecture. It is an argument. Each of these thinkers, Plato, Aquinas, Leo XIII, and, yes, Will Rogers, gives us a distinct angle on the same underlying problem: when does authority over persons become legitimate, and what conditions must hold for it to remain so? I find all four persuasives, and I think together they make a case that no reasonable person, of any political persuasion, can simply wave away.

Plato: The Obligation to Give an Account

Here is the thing about Plato's Republic: it contains one of the most precise descriptions of the accountability problem we face today, and it does so without ever mentioning a computer.⁸ Plato's philosopher-king is not legitimate merely because

he produces good outcomes. He must possess episteme, genuine knowledge, not just successful output, and he must be able to give an account of his decisions. He must be able to say: here is why, here is the reasoning, here is how this serves the common good. Absent that capacity, the appearance of wise governance is merely the exercise of power dressed in borrowed clothes.

An AI system that produces an output it cannot explain, that denies a benefit or assigns a risk score based on statistical patterns in training data, with no capacity to articulate the reasoning in terms of values, facts about the individual, or the common good fails Plato's test completely. It does not matter how often the output is "correct" in some aggregate statistical sense. A system that cannot give an account is, in Plato's framework, not a legitimate governor. It is, at best, a sophisticated guesser. At worst, it is a tyrant dressed in the efficiency of mathematics.

The explainability requirement I argue for throughout this paper, the demand that AI systems affecting citizens' lives be capable of producing plain-language explanations of their determinations, is, at its root, a Platonic demand. It is the demand that power be grounded in reason, not just in results.

Aquinas: Law Must Serve Persons, Not Merely Functions

Thomas Aquinas gave us one of the most durable tests in the Western legal tradition, and I come back to it constantly in my practice: *lex iniusta non est lex* [9]. An unjust law is no law at all. But what makes a law just? For Aquinas, human law has binding force, genuine moral authority, only when it meets two conditions. First, it must be ordered toward the common good. Second, it must respect the rational dignity of the persons it governs.

Apply those two conditions to algorithmic governance. An AI system that denies benefits without explanation fails on both counts. It does not serve the common good when it produces discriminatory outcomes against identifiable protected classes without any mechanism for detection or correction. And it fails to respect rational dignity when it treats a person as a data point, a set of weighted inputs to be scored, without ever engaging with her as an individual whose circumstances deserve genuine consideration.

What I find most useful in Aquinas is not just the normative standard, it's the insight that procedural legitimacy is not separable from substantive justice. You do not get to call a process fair simply because it was consistently applied if the thing consistently applied is an unjust standard. That is a lesson the designers of algorithmic systems and the agencies that deploy them need to hear clearly.

Leo XIII: Subsidiarity and the Irreducibility of the Human Person

Pope Leo XIII's 1891 encyclical *Rerum Novarum* was written about industrial capitalism, not artificial intelligence [10]. But the principle at its core, subsidiarity, applies here with striking force. Decisions that affect individual human lives should be

made at the most human level possible. Higher-order systems like governmental, institutional, mechanical may assist; they may not supplant the irreducibly personal judgment that the dignity of each human being demands.

Pope Leo was alarmed by the tendency of industrial systems to reduce workers to inputs; to treat persons as fungible units of productive capacity rather than as ends in themselves. The parallel to algorithmic governance is not subtle. When a benefits system processes a claimant as a feature vector — a bundle of data inputs to be scored against a model, rather than as a person whose individual circumstances deserve genuine human evaluation, it commits exactly the category error Leo diagnosed. The machine may inform the decision. It may not be the decision.

That is not a rejection of AI as a tool. It is an insistence that the tool remains subordinate to the human judgment it's meant to assist. Bishop Robert Barron, a revered acquaintance of mine through a mutual friend stated it best when explain AI in Pope Leo XIV (the current Pope Leo) 's Magnifica Humanitas: "AI and its attendant technologies are good in the measure that they function as tools in the hands of responsible agents acting for a moral purpose; they are problematic in the measure that they come to dominate both thought and action, bending the properly human in the direction of the machine." For these reasons, "Keep the Humane in the Loop" is, in this reading, a restatement of subsidiarity for the digital age.

Will Rogers: Common Sense from the People's Corner

Now, I have to be honest, this is the section of the philosophical argument I enjoy most. Not because Rogers is less serious than the others, but because he gets to the same truth from a completely different angle, and he does it in a way that no amount of academic philosophy can quite replicate [11].

Will Rogers understood ordinary Americans better than almost any public figure of his era. He championed the common man's right to understand his government. Will was not anti-intellectual, but he did recognize that complexity is frequently weaponized. When government hides behind the complexity of its operations, it has not become more sophisticated. It's just become less accountable. Rogers would have had a field day with algorithmic governance. I can almost hear him: "I never met a government algorithm I fully trusted, and I've also met a few that nobody understood, including the people who built them."

Rogers understood something that policy wonks often miss that complexity is not wisdom, and a government that hides behind the complexity of its tools has forgotten who it works for. The ordinary Oklahoman (Will Rogers was from Oklahoma) who had been denied a benefit by a system he couldn't understand, Rogers would have considered that a scandal. Not a technical one. A democratic one.

What I take from Rogers, beyond the humor, is this: the philosophical argument for accountability does not depend on being

educated in philosophy. It depends only on believing that the people this government serves deserve to be treated like adults who can understand what is being done to them. That is a conviction that crosses every ideological line I know of and one that I think most Americans, when it is put plainly, simply agree with always.

Plato, Aquinas, Leo XIII and XIV, and Will Rogers may or may not agree on much. But they agree on this: power exercised over persons must be accountable to those persons. It must be explainable. It must serve genuine human goods. Good morals. It must respect individual dignity. And it must be capable of being understood and contested by the people it affects. That four-part consensus, spanning twenty-five centuries and wildly different intellectual traditions, is not a coincidence. It is a convergence on something true.

Section: 3

The Bipartisan Moral Common Ground

My business and legal careers have shown me that the concern about unaccountable algorithmic government doesn't belong exclusively to any political party. It belongs to Americans as a whole. The conjecture that positions AI accountability as either a "left-wing" regulatory overreach or conversely, as a "right-wing" anti-technology panic is wrong on both counts and honestly, I think it's being promoted by people who benefit from the current absence of accountability, not by people who've actually thought through the interests at stake. Let me make both arguments, seriously and fairly, because I think they both deserve to be heard.

The Conservative Case

If you believe in the rule of law in the substantive sense, not just as a slogan, then you should be deeply uncomfortable with algorithmic governance as currently practiced. Rule of law means that governmental authority is exercised according to known, articulable rules that apply consistently and that can be challenged when they do not. An algorithm whose decision logic is proprietary, whose training data has never been audited, and whose outputs can't be explained to the person they affect does not meet that standard. It fails the rule of law on its own terms.

The conservative tradition has long been suspicious of arbitrary bureaucratic power. That suspicion is, I think, entirely healthy. But here is the thing: that suspicion should be more acute when the agent is a machine, not less. At least a human bureaucrat can be summoned before a congressional oversight committee. At least a human official can be questioned about his or her reasoning, challenged on her facts, and held accountable if he or she violated the rules. An algorithm cannot be deposed. It cannot be cross-examined. It cannot be fired for cause. If you are worried about unaccountable government power, you should want more accountability over AI systems, not less.

And individual liberty, the cornerstone of conservative governance philosophy, demands a meaningful right to contest government determinations. Not a nominal right buried in administrative procedure that is practically impossible to exercise. A real

right: to know what the system decided, why it decided it, and to have a human being with actual authority, review the individual's circumstances and make a genuine judgment. That is not regulatory overreach. That is constitutional minimum.

The Progressive Case

The progressive case is, in some ways, more empirically grounded and more urgently documented. Algorithmic discrimination is not a hypothetical. It has been demonstrated in recidivism scoring tools that produce racially disparate risk assessments. It has been documented in resume-screening software that disadvantages women. It has been found in facial recognition systems that misidentify people of color at dramatically higher rates. It has been identified in healthcare allocation algorithms that systematically underestimate the needs of Black patients. Unfortunately, and hopefully more by accident than purpose, this is true.

Opacity is an equal protection problem. When a system produces disparate outcomes across protected characteristics, and the people designing, deploying, or overseeing that system can't tell you what's driving those disparities because no one has audited the training data, the decision logic, or the real-world impact data, you don't just have a technical failure. You have a potential constitutional violation that no one is in a position to detect, much less remedy.

And who bears the cost? Consistently, it is the most vulnerable communities, those with the least resources to hire lawyers, challenge administrative decisions, or navigate appeal processes that weren't designed with them in mind. It is a structural feature of algorithmic systems trained on historical data produced by unequal systems. Mandatory fairness auditing is the remedy that government equal protection doctrine points toward.

The Common Ground

Here is where these two arguments converge, and I want to be precise about this because it matters for the politics: the accountability requirements I am proposing: transparency, fairness auditing, human accountability, meaningful redress, continuous review, serve both sets of values simultaneously. They protect individual liberty by ensuring government cannot make uncontested determinations. And they protect equal protection by requiring that disparate impacts be detected and addressed.

The evidence of an AI bipartisan movement is already here. Colorado's SB 205 (2024), which requires algorithmic impact assessments for high-risk automated decisions, passed with support from both sides of the aisle [12]. Illinois enacted AI interview screening requirements years ago, recognizing that automated hiring tools deserve scrutiny regardless of which party controls the statehouse. Illinois is now passing AI audit laws even though they don't exactly know what those audits should look like yet [13]. The NIST AI Risk Management Framework developed through an extensive public process and refined across two presidential administrations reflects genuine cross-partisan consensus

that AI governance is a national priority [14].

Is there genuine disagreement about the precise form accountability should take, federal versus state, prescriptive versus principles-based? Yes, and those are debates worth having. But the premise that government AI must be subject to meaningful accountability? That is not a partisan position. I have never met a constituent of any ideology who, when they actually understood what was at stake, said: "Yes, I'd like a machine to make consequential decisions about my life without any obligation to explain itself or give me an appeal." That is not a political preference anyone actually holds. The debate is downstream of shared values that, if we named them honestly, most Americans already agree on.

Section: 4

The Case for Mandatory Ethics Audits: What They are and why They Work

So, what does accountability actually look like in practice? That is the question I get most often from clients, from colleagues, and from policy folks who have heard the philosophical argument and want to know what I am actually asking. It is the right question. Let me be direct about what a mandatory AI ethics audit is and what it is not, because there is a lot of confusion in this space, and some of it is not accidental.

Defining the Ethics Audit

An AI ethics audit is an independent, third-party review of an AI system's design, training data, decision logic, outputs, and real-world impact. That independence is not a detail; it is the entire point. A vendor self-assessment is not an audit. A cybersecurity certification is not an audit. A privacy impact assessment is not an audit, though it might be a component of one. The distinction I draw in practice is this: a technical compliance check asks, "Does it work?" An ethics audit asks, "Should it work this way, and does it treat people fairly?"

Those are fundamentally different questions, and the fact that the field has sometimes conflated them is one reason we are in the situation we are in. A system can be technically functional, cybersecure, and fully compliant with its contractual specifications while simultaneously producing racially disparate outcomes, denying benefits to eligible applicants at elevated rates, and operating with logic that no affected citizen could possibly understand or contest. A technical compliance check will not catch those problems. An ethics audit is designed to note those problems.

The Five-Pillar Framework

The ethics audit framework I propose is built on five pillars. Each pillar reflects a constitutional value and generates specific, measurable audit requirements. Some of these are already being discussed by other AI professionals. It is the hottest topic in AI, I believe. But I want to present these in structured form because the structure matters. It is the difference between principles and a workable compliance framework:

Pillar	Core Requirement	What Auditors Must Assess
1 — Transparency	All government AI systems must be explainable to the citizens they affect.	Can the system's decision logic be described in plain language? Are the key variables and their weights disclosed? Can an affected individual receive a meaningful explanation of any decision made about them?
2 — Fairness	Audits must test for disparate impact across protected characteristics.	Does the system produce materially different outcomes on the basis of race, gender, income, disability, geography, or other protected characteristics? Is the training data examined for historical bias?
3 — Accountability	A named human official must be legally responsible for every AI-assisted government decision.	Is there a designated human official with legal accountability for the system's outputs? Are decision logs maintained? Is the chain of responsibility documented and publicly disclosed?
4 — Redress	Any citizen adversely affected by an AI-assisted decision must have a meaningful right of appeal to a human decision-maker.	Does the agency maintain a functional appeal process in which a human reviews the individual's circumstances? Is the appeal process accessible, timely, and genuinely capable of overriding the algorithmic determination?
5 — Continuous Review	Ethics audits must be ongoing, not one-time events.	Is the system reviewed on a defined schedule as it evolves? Are post-deployment impacts monitored? Are audit findings publicly disclosed? Is there a mechanism for affected communities to raise concerns between audit cycles?

Each of these pillar's map to something constitutionally significant. Transparency maps to due process. Fairness maps to equal protection. Accountability maps to the basic constitutional principle that governmental authority must be traceable to a responsible actor. Redress maps to the right to be heard. Continuous review maps to the recognition that living systems require ongoing democratic oversight, not just initial approval.

Existing Models: Proof of Concept

I want to address the objection that this framework is novel or untested because it is not. The EU AI Act, now in force as Regulation (EU) 2024/1689, establishes a risk-tiered governance framework with mandatory conformity assessments, human oversight requirements for high-risk AI applications (including benefits administration and law enforcement tools), and post-market monitoring obligations. Is the EU system perfect? No. Is it proof that sophisticated, risk-calibrated AI governance on scale is achievable without destroying innovation? Yes, unambiguously [15].

The NIST AI Risk Management Framework provides a rigorous, four-function structure, Govern, Map, Measure, Manage, that has already been widely adopted by responsible AI deployers across the public and private sectors [16]. Its limitation is not its rigor; it is its voluntary character. Voluntary frameworks are valuable. They establish norms, develop expertise, and build the professional culture of responsible deployment. But voluntary frameworks cannot substitute for binding legal requirements when constitutional rights are at stake. An agency that decides, without legal compulsion, that it does not need to audit its benefits adjudication AI is making a choice that will cost real people real things. Voluntary compliance does not change that calculus.

At the state level, Colorado, Illinois, and California are developing models that the federal government can learn from and, where appropriate, build upon. The regulatory infrastructure is not starting from zero [17].

Addressing the Three Objections

I hear three objections regularly to mandatory ethics audits and I want to address them honestly rather than dismissively, because some of the people raising them are raising them in good faith. Objection One: Cost: Yes, audits cost money. Independent, rigorous audits cost more than no audits. That is true. But the relevant comparison is not "cost of auditing versus zero." It is "cost of auditing versus cost of undetected algorithmic harm." When a flawed benefits system denies eligible claimants for months or years, the human cost is enormous and the eventual legal and remedial cost to government is substantial. A risk-tiered approach, calibrating audit intensity to the stakes of the decision, allows resources to be focused where they are most needed. High-stakes decisions (benefits, liberty, housing) warrant comprehensive auditing. Lower-stakes administrative applications warrant lighter review. That is proportionate, not burdensome.

Objection Two: Innovation Slowdown: This one frustrates me, because it gets the risk analysis exactly backwards. The real threat to sustainable AI innovation in government is not accountability. It is loss of public trust. We have already seen algorithmic systems pulled from deployment not because auditors found problems, but because the problems became public scandals. Proactive ethics auditing protects innovation by building the public trust that allows these technologies to be deployed at scale without a political backlash that sets the whole enterprise back years.

Objection Three: Technical Complexity: I will grant that algorithmic auditing requires specialized expertise. But algorithmic auditing is a developed field with established methodologies, credentialed professionals, and a growing literature. Auditors don't need to replicate the mathematics of a machine learning model in their heads. They need to assess design choices, training data quality and representativeness, output patterns across demographic groups, and the adequacy of appeal mechanisms. That is sophisticated work, but it's no more technically demanding than financial auditing and we've managed to build a functional financial auditing profession. I understand this. I have been directly involved as the project manager and liquidator of 114 FSLIC, RTC and FDIC financial institutions during my career.

Section: 5

A Vision For Governance: The Ethics Audit as Democratic Institution

I want to refute the notion that ethics audits are "red tape." They are not. That idea gets it exactly backwards, and I think it is worth taking a moment to correct it, because the concept shapes the debate in ways that obscure what's actually at stake.

"Red tape" is bureaucratic process that impedes legitimate activity without producing commensurate benefit. That is not what I'm describing. What I am describing is democratic accountability, the same function performed, in their respective domains, by the Government Accountability Office, the Inspector General community, and the environmental review process. These institutions are not red tape. They are the mechanisms by which a self-governing people maintains oversight of the power it has delegated to its government. Ethics audits for AI belong in exactly that tradition.

The Accountability Analogy

Consider the GAO. Its mandate is to ensure that federal programs operate legally, effectively, and in accordance with congressional intent [18]. It does not prevent government from acting; it ensures that government acts accountably. Nobody seriously argues that the GAO is "red tape." It is a pillar of democratic governance.

Or consider the Inspector General community. The Inspector General Act of 1978 created an independent oversight function within federal agencies precisely because internal self-oversight is insufficient [19]. IGs investigate waste, fraud, and abuse. They issue public reports. They create accountability where it otherwise would not exist. An AI ethics audit corps is that same principle applied to the specific challenge of algorithmic governance.

Or consider the National Environmental Policy Act, which requires environmental impact assessments before major federal actions that significantly affect the environment [20]. NEPA does not prevent infrastructure. It ensures that the consequences of governmental action are examined before the fact, not discov-

ered after. A mandatory AI ethics audit requirement is the same logic: examine the impact before deployment, not after the harm is already distributed across a population. These institutions represent democracy at work. They exist because self-governance requires the tools to govern itself. Mandatory AI ethics audits belong in that same lineage.

Three Institutional Structures

Let me be certain about what I am proposing institutionally, because the five-pillar framework is only as strong as the institutional architecture that enforces it. This is a synthesis of many thoughts on the subject generated not only by me but by others who have given this subject deep thought.

Office of AI Ethics Accountability (OAEA)

I propose establishing an Office of AI Ethics Accountability within either the Government Accountability Office or the Office of Management and Budget. The OAEA would develop and regularly update the mandatory audit framework, certify third-party auditors who meet professional standards, receive and publish audit reports on a searchable public registry, and refer material violations to appropriate oversight and enforcement bodies. Its staff would be genuinely interdisciplinary. Law, ethics, computer science, social science, and civil rights expertise are all essential to this function, and no single discipline is sufficient alone.

AI Ethics Audit Corps (AEAC)

The AEAC would be an independent auditing body modeled on the inspector general community and the Public Company Accounting Oversight Board, institutions that have demonstrated the viability of independent, publicly accountable professional audit functions. Membership would be drawn from civil society, academic institutions, and technical professionals on a rotating basis, with robust conflict-of-interest requirements and publicly disclosed membership rosters. The independence of the AEAC from both the agencies being audited and the vendors supplying the systems is not a design preference. It is a structural necessity.

Public AI Audit Registry

Every audit report produced under this framework would be publicly disclosed through a searchable registry maintained by the OAEA. Core findings such as identified disparities, explainability assessments, accountability chain documentation, recommended remediations would be available to any citizen, researcher, journalist, or advocacy organization that wants them. This is the FOIA principle applied to algorithmic governance. Transparency is the mechanism that gives the framework its teeth. Justice Brandeis understood this better than anyone [21].

Brandeis wrote that quote about financial transparency. The principle applies with equal force to algorithmic transparency. A system that must submit to independent audit, whose findings will be publicly disclosed and will be designed differently before the fact. The prospect of public disclosure changes the design conversation. It shifts the question from "What can we get away with?" to "What can we defend?" That shift matters enormously.

This, too, is what "Keep the Humane in the Loop" means. Not just at the moment of a specific decision made about a specific individual. Throughout the entire lifecycle, design, deployment, review, renewal, or retirement. The loop is not a circuit breaker that activates in exceptional cases. It is a continuous connection between governmental power and the democratic accountability that legitimizes it [22-24].

Section 6

Conclusion: The Covenant Renewed

Let me come back to my client example: the breast cancer 501c3 charity operated by my cancer-survivor friend, Fran, mentioned at the start. After we worked through her situation together, after we finally got a human reviewer to look at her file with fresh eyes, the outcome changed. Not because the facts were different. The facts were the same as they had always been. The outcome changed because a person looked at the whole picture instead of a model scoring a profile. A person asked the right questions, understood the context, exercised judgment, and arrived at a different conclusion than the algorithm had [25-27].

That is not a criticism of AI. It's an argument for keeping humans in the loop. The algorithm was not wrong because algorithms are bad. It was wrong because no algorithm, however sophisticated, has yet acquired the capacity to see a whole human being to understand context, to weigh competing values, to recognize when a rule's application produces a result its drafters would never have intended. That capacity is what human judgment offers. And it is what gets lost when we remove humans from consequential decisions without putting any accountability structure in their place. May I state this plainly: There May be Artificial Intelligence but there will Never be Artificial Ethics. Humanity, not machines, has the lock on ethics. Believe me. Ethics are and will always be the guardrails for AI [28-29].

The Philosophical Synthesis

Plato, Aquinas, Leo XIII & XIV, and Will Rogers do not share a tradition, a century, or even a cup of coffee. But I have come to think of them as unlikely co-counsel in this argument, each pressing a different facet of the same essential claim.

Plato insisted that legitimate governance requires the capacity to give an accounting, to explain, in terms of reasons and values, why a decision was made. I find that demand exactly right. An algorithm that produces outputs it cannot explain does not meet the standard of legitimate governance; it meets the standard of a sophisticated coinflip.

Aquinas gave us the natural law framework: authority is binding only when it serves the common good and respects the rational dignity of persons. I keep coming back to that second condition, rational dignity. What does it mean to respect the rational dignity of a person denied a contract by a system he or she cannot understand and cannot contest? It means giving them an explanation they can evaluate and an appeal they can meaningfully and mercifully pursue. Anything less is, in Aquinas's terms, not law. It is power pretending to be law [30].

Leo XIII's subsidiarity principle and Leo XIV human dignity proclamation that decisions affecting human lives should be made at the most human level possible is the most operationally precise of the four. The machine may inform but it may not decide alone. The humane in the loop is the constitutional minimum supporting those universal proclamations.

And Will Rogers, the "Cherokee Kid," bless him, just tells us the truth plainly: a government that hides behind the complexity of its tools has forgotten who it works for. I have not forgotten. I do not intend to. I am listening to Will. So did several of our Presidents [31].

A Call to Action

I want to close by speaking directly to four audiences, because I think each of them has an important distinct role in getting this right.

To Lawmakers: You have the authority right now. The framework exists. The bipartisan will is forming. The moral case has been made, philosophically, constitutionally, and practically. The choice before you are whether to act before the harms are so widespread and the systems so entrenched that reform requires a generation of corrective effort. I have seen what entrenchment looks like in other areas of contemporary (my generation) law, especially with consumer protection and especially in mortgage redlining. It is expensive, slow, and painful. The window for proactive governance is open. Act before it closes. That moment is now.

To Procurement Officers: You are not just administrators. Every standard you write, every audit clause you insist on, every accountability requirement you demand from a vendor as a condition of contract is democratic governance. That matters, and more than you may realize. The standards built into procurement today become the expectations of the industry tomorrow. You have more power to shape this than anyone in government, and I'd ask you to use it.

To Technology Leaders: Build systems that can be audited. That means keeping records, maintaining documentation, designing for explainability from the beginning rather than trying to retrofit it after the fact. Disclose what auditors need. Treat ethics review as a professional obligation, not a regulatory obstacle to be minimized. Here is the honest truth: the industry's long-term credibility, your credibility, depends on whether the public comes to believe that AI systems in government can be trusted. The fastest path to that trust is accountability. The fastest path to a generation-long regulatory backlash is the opposite.

To Citizens: This fight did not start with lawyers or legislators. It started with people, people like my client, Fran, insisting that their government owes them an explanation. That insistence is as valid in the age of AI as it ever was, and it's, as necessary. Ask where your representatives stand. Support organizations that are challenging algorithmic harm, support positive solutions. And

when an automated system makes a decision about your life, about your benefits, your housing, your liberty, your livelihood, insist on your right to an explanation. Insist on your right to a human being who can look you in the eye and take responsibility. Do not accept a code where you're owed a reason.

A Note of Hope

I want to close on a genuinely positive note. I am not a pessimist about AI or AI ethics or even AI audits. I have seen what America does when it decides that a form of power has exceeded the bounds of democratic accountability. As a teenager, I even watched every minute of the Watergate hearings. Accountability acts. Imperfectly, often slowly, sometimes only after the harms have become undeniable but it acts. The Administrative Procedure Act. The Civil Rights Act. The Freedom of Information Act. The Inspector General Act. Each was, in its moment, a recognition that the existing accountability structures were insufficient for the power they were meant to govern, and a decision to build something better.

My patriotism, legal background and even my Jesuit education makes me believe we can do this again. The philosophy is clear. The law is ready. The policy framework exists. The bipartisan will is forming. What remains is the act and the conviction that the people this technology affects deserve nothing less than what we've always promised them: a government that answers to them, that can explain itself to them, and that remains, in some genuinely human sense, theirs.

About the Author

Jack Skagerberg is a Texas-and Wisconsin licensed (45 years) attorney admitted to the USSC and various other federal courts. Attorney Skagerberg is a Certified AI Leader with a graduate AI certificate from the University of Texas at Austin McCombs Graduate School of Business. His practice spans business law, technology policy, and government procurement matters, where he has worked directly with clients and financial services agencies navigating automated decision-making systems. (www.linkedin.com/in/jack-skagerberg-b9a54ab).

Jack's work at the intersection of law, ethics, and emerging technology reflects a conviction he returns to again and again in his career, to wit: that the tools of governance must serve the people they govern and that accountability is not a luxury but a constitutional necessity. For Jack, this maxim is an absolute truth with Artificial Intelligence ("AI"). He writes, advises, and speaks on AI governance, ethics, legal technology, and the philosophy of law in the digital age.

References

1. U.S. Constitution. (1787). Preamble; Amendment V; Amendment XIV, § 1.
2. The Declaration of Independence. (1776). United States.
3. Administrative Procedure Act, 5 U.S.C. §§ 551–559, 701–706. (1946).
4. Civil Rights Act of 1964, Pub. L. No. 88-352, 78 Stat. 241. (1964).
5. Freedom of Information Act, 5 U.S.C. § 552. (1966).
6. National Environmental Policy Act, 42 U.S.C. §§ 4321–4347. (1970).
7. Inspector General Act of 1978, Pub. L. No. 95-452, 5 U.S.C. app. §§ 1–13. (1978).
8. Government Accountability Office Enabling Act, 31 U.S.C. §§ 702–720. (2006).
9. National Institute of Standards and Technology. (2023). Artificial intelligence risk management framework (AI RMF 1.0) (NIST AI 100-1).
10. Executive Office of the President, Office of Management and Budget. (2024). Advancing governance, innovation, and risk management for agency use of artificial intelligence (Memorandum M-24-10).
11. Artificial Intelligence Video Interview Act, 820 Ill. Comp. Stat. 42. (2019).
12. Colorado General Assembly. (2024). Consumer protections for artificial intelligence (S.B. 24-205).
13. European Parliament & Council of the European Union. (2024). Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). Official Journal of the European Union, L 1689, 1.
14. Plato. (1992). *The Republic* (G. M. A. Grube, Trans.; C. D. C. Reeve, Rev.). Hackett Publishing Company. (Original work published ca. 375 BCE)
15. Plato. (1970). *The Laws* (T. J. Saunders, Trans.). Penguin Classics. (Original work published ca. 348 BCE)
16. Aquinas, T. (1947). *Summa theologica* (Fathers of the English Dominican Province, Trans.; I–II, Q. 90–97). Benziger Brothers. (Original work published ca. 1265–1274)
17. Leo XIII. (1891). *Rerum novarum*: Encyclical on capital and labor. Holy See. Vatican website
18. Brandeis, L. D. (1914). *Other people's money and how the bankers use it*. Frederick A. Stokes Company.
19. Rogers, W. (1993). *The wit and wisdom of Will Rogers* (A. Ayres, Ed.). Meridian.
20. Barocas, S., Hardt, M., Narayanan, A. (2023). *Fairness and machine learning: Limitations and opportunities*. MIT Press.
21. Diakopoulos, N. (2019). *Automating the news: How algorithms are rewriting the media*. Harvard University Press.
22. Eubanks, V. (2018). *Automating inequality: How high-tech tools profile, police, and punish the poor*. St. Martin's Press.
23. Floridi, L., Cowls, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., ... & Vayena, E. (2018). An ethical framework for a good AI society: Opportunities, risks, principles, and recommendations. *Minds and Machines*, 28, 689–707.
24. Kearns, M., Roth, A. (2019). *The ethical algorithm: The science of socially aware algorithm design*. Oxford University Press.
25. Mitchell, M. (2019). Model cards for model reporting. In *Proceedings of the Conference on Fairness, Accountability, and Transparency* (pp. 220–229). ACM.
26. Noble, S. U. (2018). *Algorithms of oppression: How search engines reinforce racism*. New York University Press.
27. O'Neil, C. (2016). *Weapons of math destruction: How big*

-
- data increases inequality and threatens democracy. Crown Publishers.
28. Regan, P. M., Jesse, J. (2019). Ethical challenges of edtech, big data and personalized learning: Twenty-first century student sorting and tracking. *Ethics and Information Technology*, 21, 167–178.
29. Selbst, A. D. (2019). Fairness and abstraction in sociotechnical systems. In *Proceedings of the ACM Conference on Fairness, Accountability, and Transparency* (pp. 59–68). ACM.
30. Sunstein, C. R. (2005). *Laws of fear: Beyond the precautionary principle*. Cambridge University Press.
31. Wachter, S., Mittelstadt, B., Floridi, L. (2017). Why a right to explanation of automated decision-making does not exist in the General Data Protection Regulation. *International Data Privacy Law*, 7(2), 76–99

Copyright: ©2026 Jack Skager Berg This is an open-access article distributed under the terms of the Creative Commons Attribution License, which permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited.