



How to manage the gap between enterprise AI use and AI regulation

By Jon Polenberg

Two years ago, almost no states regulated AI. Today, [more than 30 do](#). Companies operating across multiple states often build to the strictest requirement they face, but only when the requirements differ by degree. When they differ by kind, no single bar exists. One state's rule can push a company to normalize a skewed applicant pool. Another can push the company to follow raw performance data. Compliance in one jurisdiction can create exposure in the next, forcing the enterprise to choose which law to break.

CIOs must make AI governance decisions based on [today's patchwork of state laws](#), not under the federal framework they hope Congress eventually adopts. The longer that gap persists, the more enterprise AI governance turns into legal fragmentation management rather than technology implementation. The problem lands on CIOs because AI now runs through HR, marketing, customer data, cybersecurity and legal compliance simultaneously, not just one department's tech stack.

The hiring challenge

[Hiring usually shows the problem most clearly](#) because no other task or function runs AI against as many people as often. While volume doesn't make hiring the first place trouble appears, it does make hiring the place where trouble becomes hardest to miss.

A tool that reduces 100,000 resumes to a shortlist of 10 makes judgment calls about commute times, work history and career gaps. No one wrote those judgments into policy. No one reviewed them before the shortlist reached a hiring manager. LinkedIn funnels candidates who might have never applied. Data providers scrape the open internet to identify and solicit people who never put themselves forward for anything.

None of this requires bad intent. A commute-time filter can quietly exclude applicants from underserved neighborhoods without any decision-maker choosing to discriminate. An applicant pool that skews 70%-30% along gender lines raises a hard question: Should the company normalize the ratio, leave it alone or follow whatever the underlying performance data shows? The third option is the only one that looks neutral. It assumes unbiased

historical performance data, which is the same assumption that fails in the retail promotion context discussed below. A tool trained on decades of history [inherits that history's patterns](#).

There is no clean answer, and the states that have tried to write one down do more than set different bars; some point in opposite directions. The legal exposure companies face comes from disparate impact, which occurs when a neutral policy or practice disproportionately harms members of a protected group. This exposure comes from how AI performs once deployed at scale against real people, not from anyone's intent or from whether the company disclosed the practice.

Development vs. deployment: Where's the true exposure?

Vendors build AI systems, but the companies using them decide where to deploy them, what role they will play in consequential decisions and whether that use complies with each state's law. The statutes don't agree on the legal hook. Some reach the developer. Some reach the deployer. Some reach both.

The allocation shifts from one framework to the next. But one fact doesn't shift: the company operating the tool against real people remains present in every jurisdiction at once. [Deployment, not development, concentrates exposure](#) regardless of how any single statute assigns responsibility. A vendor's safety testing claims or federal review status doesn't move that exposure off the deployer's books.

This dynamic extends beyond hiring. Retailers that disclose their cameras and offer an opt-out can still deliver advertising along lines that track race, gender or age. Disclosure and opt-out solve a notice problem. They do nothing about impact because the underlying model selects for engagement, not for any demographic anyone chose. The retailer can act transparently and still produce a skewed outcome. That's the point: Consent mechanisms don't cure disparate impact. Promotion decisions carry the same risk. A tool trained on decades of performance data inherits that history's patterns and can produce an outcome nobody at the company would have chosen.

Is a single federal standard the answer?

A federal framework with one set of standards, rather than 30 separate state ones, sounds like the obvious fix. The White House has tried twice in the last year. First, it created a [task force to challenge state AI laws in court](#). Then it created a [voluntary security review window for AI developers](#). Neither substitutes for legislation. An executive order can't preempt state law without Congressional action, and a voluntary review imposes no compliance obligation on vendors that can opt in or out.

However, even a working mechanism wouldn't answer the harder question: What level of algorithmic bias can the law tolerate? The White House's efforts fail procedurally. The policy question remains unresolved for a different reason. Zero bias is impossible. No human being is bias-free, and neither is the data humans produce.

No federal lawmakers, left or right, have been willing to draw that line. The problem isn't the bias itself. Anti-discrimination law has tolerated imperfect, biased human decision-making for decades without demanding zero. An algorithm changes the politics by making residual bias legible. It can measure, report and audit that bias after the fact. A statute would have to name a tolerance and defend it. That means conceding that the approved system remains biased and saying by exactly how much. Doing that is harder politically than tolerating the same bias when it spreads across a thousand human managers, and no one has to sign the number.

A single federal standard remains the better end state than 30 conflicting ones. But one risk on the other side deserves acknowledgment, even if it doesn't change the conclusion. Technology is moving fast enough that some of today's hardest questions, especially bias normalization, may have technical answers no one has built yet. A federal standard written today could lock in today's assumptions and foreclose tomorrow's solutions. That risk calls for careful drafting. It doesn't justify the patchwork of state standards, which imposes its costs now and every day it persists.

What to do today

None of this changes [what companies can do today](#). CIOs should map AI by use case and jurisdiction, so they know which laws govern each deployment. They should then build governance around practices that should exist regardless of the regulatory landscape, including clear terms of service, reliable data provenance and human review where the consequences matter most.

Congress, courts and future administrations will keep shaping AI regulation, but enterprise AI won't wait for them. Companies already deploy systems under laws that states can enforce today. CIOs must govern within the legal landscape that exists now, not the federal framework they hope will eventually arrive.

Jon Polenberg is a shareholder and vice chair of Becker's Business Litigation Practice. As a business trial lawyer in Florida, he's known for his advocacy and strategic insight in complex commercial litigation. Jon has decades of experience representing businesses in high-stakes cases and handling intricate legal disputes with precision and a commitment to excellence. His practice spans a range of industries, assisting companies in resolving matters that affect their operations, reputations and financial interests.