



BURAK GALİBA
Thoughts on AI Transformation

ESSAY · ENTERPRISE AI

AI Governance Stopped Being Optional — TRAIGA, Colorado ADMT, and the EU Clock

Texas TRAIGA is in force, Colorado's ADMT law lands January 2027, EU dates are set. Why NIST AI RMF is the safe-harbor move that survives it all.

GOVERNANCE

AI STRATEGY

Burak Galiba · July 8, 2026 · 4 min read

<https://burakgaliba.com/blog/ai-governance-is-a-deadline>



SCAN TO READ ONLINE

Somewhere in your enterprise right now, an AI system is scoring a customer, screening a resume, or flagging a transaction. Since January 1, doing that carelessly has been a legal problem in Texas — and Texas was just the opening act.

AI governance spent years as a slide near the end of the deck: important, aspirational, undated. What changed in 2026 is the dates. Governance is now a calendar, and the calendar does not care whether your program is ready.

The US clock already struck

Three state laws took effect on January 1, 2026, as [King & Spalding's analysis](#) lays out:

- **Texas TRAIGA** — a prohibition-based regime covering AI conduct, with real enforcement teeth and one enormously useful feature I'll get to below.
- **California SB 53** — transparency obligations for frontier-model developers.
- **California AB 2013** — training-data disclosure requirements.

Next on the calendar: **Colorado's SB 26-189**, the ADMT law that replaced its original AI Act, takes effect **January 1, 2027**. If your systems make or substantially drive consequential decisions about Colorado residents — credit, employment, housing, insurance — you have under six months to inventory them and stand up the required processes.

And yes, a federal executive order signals preemption battles over state AI laws ahead. Betting your compliance posture on litigation outcomes is a strategy, technically. So is not wearing a seatbelt because the crash might not happen.

Europe moved the deadline, not the destination

The EU AI Act's general applicability **arrived on August 2, 2026**, and is now in force. The headline change came in June, when the **Digital Omnibus** was finalized: high-risk obligations for stand-alone Annex III systems are deferred to **December 2, 2027**, and for AI embedded in regulated products to **August 2, 2028**, as [Gibson Dunn details](#).

Some executives read the deferral as a reprieve. Read it instead as a published exam date. December 2027 sounds distant until you count what a high-risk compliance program requires: a complete AI inventory, risk classification, documentation, human-oversight design, and monitoring — across every affected system. Enterprises that started at "someday" routinely need 12–18 months. The window and the work are now approximately the same size.

The safe harbor is the strategy

Here's the practical center of all this. TRAIGA contains an explicit safe harbor: organizations that implement the **NIST AI Risk Management Framework** get substantial protection under the statute, a point [Baker Botts highlights](#) in its US AI law review.

That single provision resolves the hardest governance question — *which* standard to build on — better than any consultant could. Anchor on NIST AI RMF and you get:

- A Texas statutory safe harbor, today
- A structure that maps cleanly onto Colorado's 2027 requirements and the EU's 2027–28 obligations
- Insulation from regulatory whiplash — the framework is risk-based and statute-agnostic, so it survives whichever way the preemption fights break
- A vocabulary your auditors, insurers, and enterprise customers already accept

One framework, built once, defensible everywhere. The alternative — a bespoke policy per jurisdiction, rewritten each legislative session — is how governance budgets die.

Governance is also how your projects survive

Compliance isn't even the biggest payoff. Gartner predicts [over 40% of agentic AI projects will be canceled by end-2027](#), with inadequate risk controls among the leading causes — and Gartner also projects [AI governance platform spending will grow from \\$492 million in 2026 to over \\$1 billion by 2030](#) as enterprises internalize the lesson.

The pattern is consistent: [AI initiatives without governance don't fail at launch, they fail at scale](#) — killed in security review, legal review, or a board risk discussion, eighteen expensive months in. [Teams that bring risk functions in at day one ship faster](#), because objections get engineered out instead of litigated at the end.

A defensible program is smaller than most executives fear:

1. **An AI inventory** — every system, including the shadow ones, with owners named.
2. **Risk tiers** — most systems are low-risk; the point is knowing which aren't.
3. **A policy people can follow** — pages, not a binder.
4. **Human oversight where decisions are consequential**, with authority to intervene.
5. **An incident path** — who's told, who decides, who fixes, who documents.

That's a quarter of disciplined work, not a transformation program. The enterprises treating it that way are quietly converting a legal obligation into a shipping advantage.

Where to start

Start with the inventory — nothing else in governance works without it, and every statute above effectively demands one. A spreadsheet, every AI system in use, an owner per row. Most enterprises are surprised (occasionally horrified) by row count alone.

If you want to know how your governance posture scores before the calendar forces the question, my free [AI Readiness Score](#) includes governance as one of its four domains — 20 questions total, about ten minutes, and you'll see exactly how far your current program is from defensible.

Read the essay online at <https://burakgaliba.com/blog/ai-governance-is-a-deadline>

I built a free AI Readiness Score — 20 questions across pilots, data, talent and governance, about ten minutes, no email required: <https://burakgaliba.com/readiness>

© 2026 Burak Galiba · <https://burakgaliba.com> · Views are my own.