

Patchwork vs Framework: US and EU AI Regulation Compared, and the Compliance Choices Facing Cross-Border Business

Janine Anthony Bowen
Of Counsel, BakerHostetler

Nils Lölfig*
Senior Counsel, Bird & Bird

☞ Artificial intelligence; Comparative law; Compliance; Consumer protection; Cross-border transactions; EU law; International trade; Product safety; Regulation; United States

Abstract

The United States (US) and European Union (EU) have adopted different approaches to Artificial Intelligence (AI) regulation, creating growing compliance challenges for organisations operating across both markets. While the EU AI Act establishes a comprehensive, risk-based framework rooted in product safety, fundamental rights and market harmonisation, the US relies on a fragmented system of state laws, sector-specific enforcement and voluntary standards, largely focused on consumer protection. This article compares the contrasting approaches to AI regulation in the US and EU and explores the compliance challenges these differences create for transatlantic businesses. It examines how organisations can develop effective governance frameworks that balance regulatory obligations with innovation, arguing that the US primarily regulates marketplace harms while the EU regulates the technology itself. Effective AI governance therefore requires a flexible, risk-based approach capable of addressing both regimes.

1. Introduction

AI regulation has become one of the defining governance challenges of the decade, and the US and EU have taken markedly different paths. The EU has enacted comprehensive, binding, product-safety-style legislation through the AI Act, supplemented by fundamental rights

protections rooted in the Charter of Fundamental Rights of the European Union (“Charter”). The US, by contrast, has pursued a fragmented, innovation-permissive approach: no comprehensive federal statute, a patchwork of state laws, sector-specific federal enforcement, and the non-binding NIST AI Risk Management Framework as a widely referenced technical baseline.

For organisations operating trans-Atlantically, the consequences are practical and immediate. The same AI system may attract entirely different regulatory obligations depending on where it is developed, deployed, or accessed—and on what the system actually does. Building a governance architecture that satisfies both regimes, without over-engineering for jurisdictions where exposure is minimal or under-engineering where it is material, has become a core strategic question for legal, compliance, and product teams alike.

This article does three things. First, it sets out the philosophical and structural foundations of the US and EU approaches. Second, it analyses how consumer protection principles organise USAI regulation, in contrast with the EU’s multi-dimensional framework combining product safety, fundamental rights protection, and internal market harmonisation. Third, it examines the practical implications for businesses operating on both sides of the Atlantic—including how to choose a governance baseline, calibrate compliance investment to actual exposure, and preserve innovation capacity within a structurally divergent regulatory environment.

2. The foundations: two regulatory philosophies

2.1. The US approach: consumer protection within a fragmented architecture

A federal vacuum filled by policy direction and voluntary guidance

There is no horizontal federal AI statute in the US. The vacuum at federal level is filled by three categories of instrument: executive policy direction, sectoral enforcement by existing agencies, and non-binding technical guidance. Executive Order 14179 of January 2025 reversed the Biden-era AI policy and reframed federal posture around innovation acceleration and US technological primacy, with America’s AI Action Plan (July 2025) translating that into pillars on innovation, infrastructure, and international positioning. Sectoral enforcement against AI-driven harms is pursued through existing authorities—notably FTC s.5 powers over unfair or deceptive practices and EEOC oversight of algorithmic employment discrimination—rather than through AI-specific legislation. Technical guidance is supplied

* Special thanks to the following contributors: Pablo Hesse and Harry Qu of Bird & Bird; and Noam Kleinman, Kyle Kennedy, Colleen “Bo” Gair and Jimmy Nguyen of BakerHostetler.

principally by the NIST AI Risk Management Framework, which is voluntary but widely adopted as a practical reference for trustworthy AI design.

State-level experimentation and its consumer-protection centre of gravity

In the absence of federal legislation, the most consequential rule-making has migrated to the states. Texas, Colorado, and California have each enacted statutes that, for all their differences, share a common centre of gravity: the consumer, the marketplace, and the prevention of specific identifiable harms. AI is treated as a transactional input—to employment screening, credit decisions, healthcare interactions, consumer-facing automation—rather than as a technology category requiring *ex ante* governance.¹

The Texas Responsible AI Governance Act (TRAIGA, effective 1 January 2026) targets a discrete set of intentional harms and government uses: it imposes AI-use disclosure on government agencies and healthcare providers; bars AI used to intentionally incite self-harm, harm to others or criminal conduct; prohibits public-sector social scoring producing detrimental outcomes; restricts non-consensual government biometric identification; and prohibits AI developed with intent to discriminate unlawfully, with disparate impact alone insufficient to establish a violation. Enforcement sits exclusively with the Texas Attorney General, civil penalties run from \$10,000 to \$200,000 per violation, and curable violations attract a 60-day cure period.

The Colorado Artificial Intelligence Act (CAIA, effective 30 June 2026) takes a more recognisably risk-based shape. It defines high-risk AI systems by reference to *consequential decisions* in enumerated domains—education, employment, financial services, essential government services, healthcare, housing, insurance, legal services—and imposes a duty of care against algorithmic discrimination, risk management obligations on deployers, annual impact assessments, consumer rights to explanation, appeal and correction, and incident reporting. The statute notably provides an affirmative defence for developers and deployers that identify and cure violations and demonstrate alignment with recognised frameworks such as the NIST AI RMF or ISO/IEC 42001:2023.²

The California Transparency in Frontier AI Act (TFAIA, effective 1 January 2026) shifts the regulatory target upstream, focusing on developers of the most powerful foundation models rather than their downstream deployment. “Frontier model” is defined by reference to a training compute threshold of 10^{26} FLOPs; “large frontier developers”—those with annual gross revenues above \$500 million—must submit quarterly catastrophic

risk assessments to California’s Office of Emergency Services, publish an annual safety framework, file transparency reports when releasing or substantially modifying frontier models, and protect employee whistleblowing. Penalties are capped at \$1 million per violation.

Two features run across all three statutes and characterise the US approach as a whole. First, enforcement is post-market, attorney-general-led, and oriented toward remediation: cure periods, affirmative defences, and capped penalties dominate. Second, federal pre-emption efforts under the current administration—including DOJ task force challenges to state AI regulations—remain a live but as-yet-unsettled overlay.

2.2. The EU approach: product safety meets fundamental rights

A product-safety backbone, layered with Charter rights

The EU AI Act, in force since 1 August 2024, operates from a structurally different premise. Its backbone is product safety regulation: AI systems that fall within its high-risk categories must undergo conformity assessment before being placed on the market, carry CE marking, and be subject to ongoing post-market monitoring and serious incident reporting. This is the long-established EU regulatory technique for goods on the internal market, adapted to AI. Layered onto that backbone is a fundamental rights dimension grounded in the Charter, which underpins the outright prohibition of AI practices deemed incompatible with EU values irrespective of technical safety, and an internal-market harmonisation function that displaces inconsistent national approaches.

A four-tier risk architecture

Unacceptable-risk practices—public-authority social scoring, manipulative AI exploiting vulnerabilities, workplace emotion recognition, biometric categorisation inferring sensitive attributes, among others—are prohibited outright. *High-risk* covers two distinct populations: AI systems that are safety components of (or themselves are) products regulated under EU harmonisation legislation listed in Annex I, and stand-alone systems falling within the Annex III use-case list (biometric identification, critical infrastructure, education, employment, essential services, law enforcement, migration, justice). High-risk obligations span risk management, data governance, technical documentation, logging, transparency, human oversight, accuracy, robustness, and cybersecurity, with conformity

¹ Other states and municipalities have introduced or passed AI-specific legislation. These three are simply illustrative.

² At the time of submission, the Colorado AI Act remained good law. On 14 May 2026, the Colorado legislature passed SB 26-189 (awaiting Governor Polis’s signature at the time of writing), which repeals CAIA and replaces it, effective 1 January 2027, with a significantly narrower regime centred on transparency and consumer disclosure rather than the duty of care, algorithmic impact assessments, and risk management programs originally imposed. Rather than undermining the comparative analysis offered here, this development reinforces it: the retreat from CAIA’s more prescriptive model exemplifies the structural divergence between US and EU approaches to AI regulation that this article identifies.

assessment, CE marking, and post-market monitoring on top. *Limited-risk* systems—chatbots, emotion recognition, biometric categorisation, deepfakes—attract targeted transparency duties without the full high-risk regime. *General-purpose AI models* face a tiered regime of their own: baseline documentation, copyright compliance, and training data summaries for all GPAI models; additional obligations (adversarial testing, systemic risk assessment, cybersecurity, energy efficiency reporting, serious incident reporting) for models with systemic risk above 10^{25} FLOPs or those Commission-designated as such.

Timelines settled by the May trilogue agreement

The application timetable was recalibrated through the Commission's Digital Omnibus proposal of 19 November 2025, on which Council and Parliament negotiators reached a provisional agreement on 7 May 2026. Prohibitions and AI literacy obligations have applied since 2 February 2025, GPAI model obligations since 2 August 2025, and art.50 transparency duties (other than art.50(2)) remain due to apply on 2 August 2026. The trilogue agreement replaces the original high-risk application dates with fixed dates of 2 December 2027 (Annex III stand-alone systems) and 2 August 2028 (Annex I systems embedded in regulated products), dropping the Commission's original mechanism that would have tied application to a Commission decision on standards readiness. Article 50(2) machine-readable marking obligations apply from 2 August 2026, with a transitional period to 2 December 2026 for providers of generative AI systems already on the market. The agreement also adds a new art.5 prohibition targeting AI systems generating child sexual abuse material or non-consensual intimate imagery (applicable from 2 December 2026), extends SME relief measures to small mid-cap enterprises, and resolves the Annex I conformity assessment dispute that broke the previous trilogue. Formal endorsement and adoption are expected before 2 August 2026.

Regulation paired with industrial policy

The EU's regulatory approach is deliberately paired with industrial policy. The Commission's Apply AI Strategy (October 2025) directs approximately €1 billion in EU funding to accelerate AI uptake in healthcare, manufacturing, and energy, drawing on Horizon Europe and Digital Europe; the InvestAI Facility under the broader AI Continent Action Plan stimulates private investment in AI infrastructure, complemented by "AI first" and "buy European" policies for public-sector procurement. This combination—comprehensive regulation with substantial public investment—is the structural counterpart to a US model that relies predominantly on private investment with minimal regulatory intervention.

3. Consumer Protection vs Multi-Dimensional Governance

The most analytically useful way to frame the divergence is by reference to *what each regulator is protecting and from what*. USAI regulation is organised around the consumer as the focal subject and the marketplace as the relevant arena. EUAI regulation is organised around a wider set of protected interests—physical and psychological safety, fundamental rights, democratic values, and the integrity of the internal market—with the consumer being one beneficiary among several. This shapes scope, triggers, enforcement, and remedies in concrete ways.

3.1. Scope and triggers

The consumer protection paradigm tends to trigger regulatory obligations at the point of consumer interaction or consequential decision. CAIA's high-risk category is defined functionally by the existence of a consequential decision in an enumerated set of consumer-facing domains; TRAIGA's transparency duties attach to consumer-facing AI deployed by government agencies and healthcare providers; TFAIA's frontier-model duties target the largest developers but are justified by reference to the downstream consumer and societal risks those models could enable. There is, by design, no horizontal pre-market regulation of AI as a product.

The EU framework, by contrast, regulates AI systems *as products being placed on the market*, well before any individual consumer interaction occurs. High-risk classification under Annex I is product-list-based; classification under Annex III is use-case-based but is not contingent on a specific consumer transaction or harm. Conformity assessment, technical documentation, and CE marking operate *ex ante*. The trigger is market placement or putting into service, not consumer-facing deployment.

3.2. Remedies and enforcement

Consumer protection-style regimes typically rely on enforcement by attorneys general or sector-specific agencies, calibrated to incentivise remediation rather than impose deterrent-level fines. Cure periods (TRAIGA), affirmative defences (CAIA), and capped per-violation penalties (TFAIA at \$1 million) are characteristic. The EU AI Act, by contrast, provides administrative fines of up to €35 million or 7% of global annual turnover for prohibited practices and up to €15 million or 3% for other violations, alongside market surveillance powers and product recalls—penalty levels reflecting the gravity attached to public-good harms rather than individual consumer detriment.

3.3. Substantive coverage

The consumer protection lens also leaves systemic gaps. Manipulative AI is not addressed under TRAIGA in the way art.5 of the AI Act addresses subliminal techniques and exploitation of vulnerabilities. Emotion recognition systems face no specific US state prohibition. Real-time remote biometric identification by law enforcement is regulated extensively in the EU but largely untouched at US state level. Machine-readable marking of synthetic content, mandatory under the AI Act for providers of generative systems, has no US state equivalent. Where the EU regulates AI as a technology with potential to affect rights, the US tends to regulate AI as a transactional input susceptible to specific, demonstrable harm.

3.4. Convergence at the structural level

These differences notwithstanding, structural convergence is more pronounced than is sometimes recognised. Both frameworks employ risk-based tiering; both distinguish between actors in the value chain (providers/deployers in the EU; developers/deployers in CAIA); both rely on transparency as a principal regulatory tool; and both recognise foundation or frontier models as a distinct governance challenge. The table below summarises the principal points of comparison.

Regulatory trend	Key US provisions	Comparison to EU
Prohibited practices	TRAIGA prohibits AI used to intentionally incite self-harm, harm to others or criminal activity; government social scoring causing detrimental treatment; CSAM-generating AI; non-consensual government biometric identification; AI developed with intent to infringe constitutional rights or unlawfully discriminate.	EU art.5 prohibitions apply categorically based on the nature of the practice and Charter rights, generally without requiring proof of intent. Texas's prohibitions are broader on their face but narrower in application, relying on intent thresholds and lacking equivalents for manipulative AI, workplace emotion recognition, or real-time remote biometric identification.
High-risk systems	CAIA defines high-risk functionally by reference to consequential decisions in enumerated domains; imposes duty of care, risk management, impact assessments, consumer rights to explanation/appeal/correction, and incident reporting; affirmative defence for NIST AI RMF or ISO 42001 compliance.	EU defines high-risk by Annex I product integration and Annex III use cases; requires ex ante conformity assessment, technical documentation, CE marking, and post-market monitoring. Both employ provider/deployer-style allocation. Colorado relies on ongoing obligations rather than pre-market approval.

Regulatory trend	Key US provisions	Comparison to EU
Transparency	TRAIGA: clear and conspicuous disclosure of AI use by government and healthcare providers; private sector unaddressed. CAIA: general consumer disclosure unless AI use is obvious.	Providers must design consumer-facing AI to disclose its nature unless obvious, and must machine-readably mark synthetic outputs; deployers must disclose emotion recognition, biometric categorisation, and deepfakes (subject to limited carve-outs). US state laws have no machine-readable marking or deepfake disclosure equivalent.
General-purpose/frontier models	TFAIA regulates frontier models above 10^{26} FLOPs developed by firms with over \$500 million revenue; quarterly catastrophic risk assessments, annual safety framework, transparency reports, whistleblower protections; \$1 million per-violation cap.	EU regulates GPAI in three tiers—baseline obligations for all GPAI models (indicatively from 10^{23} FLOPs); systemic-risk obligations from 10^{25} FLOPs or Commission designation; use-case-specific obligations where GPAI is integrated into high-risk or transparency-triggering systems.

The table confirms a meaningful structural parallel; what it does not capture is the difference in *what these structures are for*. The EU's tiers are gates on market access; the US tiers are triggers for ex post enforcement. That difference is the centre of gravity for everything that follows.

4. Implications for transatlantic business

The third and most operationally significant question is what all of this means for organisations that deploy AI on both sides of the Atlantic. The honest answer is that genuinely transatlantic AI compliance is structurally harder than either jurisdiction alone, and there is no single governance architecture that maps cleanly onto both regimes.

4.1. Three structural challenges

Classification asymmetry: The same AI system may be classified differently across frameworks. An employment screening tool is expressly listed as high-risk under Annex III of the AI Act; under CAIA, equivalent classification turns on whether the system makes or substantially contributes to a “consequential decision”. A consumer-facing chatbot may attract transparency obligations under both regimes but with different triggers and exemptions. A frontier foundation model may be

subject to systemic-risk obligations under the AI Act from 10^{25} FLOPs and to TFAIA from 10^{26} FLOPs. Each system must be assessed independently against each applicable regime; classifications do not travel.

Divergent compliance pathways: The AI Act, in its high-risk tier, is pre-market: conformity assessment, technical documentation, and (in some cases) notified body involvement must be completed before placement on the EU market. US state laws do not provide for pre-deployment approval; obligations are ongoing and largely post-market. A single unified compliance process will not satisfy both, and parallel tracks—with shared inputs where possible—are unavoidable.

Incompatible technical standards ecosystems: EU conformity assessment is structured around harmonised standards developed by CEN-CENELEC, which carry a formal presumption of conformity with AI Act requirements. US frameworks reference the NIST AI RMF and, increasingly, ISO/IEC 42001; some US state regimes reference no specific standards at all. Mapping between the two is possible but partial, and documentation, testing protocols, and quality management systems frequently need to be maintained in two configurations.

4.2. Choosing a governance baseline

There is no universally correct answer to which framework should anchor an organisation's AI governance architecture. The appropriate starting point depends on the organisation's geographic footprint, sectoral context, existing compliance infrastructure, and AI system portfolio. Three broad approaches are available, and most organisations will adopt some form of hybrid.

Approach A: EU AI Act as the primary baseline: Building governance around the AI Act first and layering US state obligations on top is often the more efficient path where the organisation places high-risk AI systems or GPAI models on the EU market, falls within the AI Act's extraterritorial scope, already operates under EU frameworks (GDPR, sectoral financial or medical device regulation) whose infrastructure maps onto AI Act requirements, or anticipates significant EU commercial relationships where AI Act compliance is becoming a procurement consideration. The trade-off is the AI Act's complexity and resource intensity, which can be disproportionate for organisations with limited EU exposure.

Approach B: US state law as the primary baseline: For organisations with limited or no EU market exposure, particularly those in heavily regulated US sectors, it may be more practical to anchor governance in applicable US frameworks. This is sensible where AI systems are deployed exclusively or predominantly in the US, where the organisation operates in a sector already subject to dense federal and state regulation (financial services, healthcare, insurance) imposing meaningful AI-adjacent obligations, or where primary regulatory relationships

are with US agencies, making NIST AI RMF alignment more operationally relevant than CEN-CENELEC harmonised standards.

Approach C: Hybrid, system-by-system mapping: Many transatlantic organisations will find that neither a purely EU-first nor a purely US-first approach fits the portfolio. Where AI systems are spread across jurisdictions and sectors, a system-by-system mapping exercise—assessing each system independently against each applicable framework—is often the most defensible model. It avoids both over-engineering low-risk systems to AI Act standards and under-engineering systems within EU scope or subject to demanding US sector-specific requirements.

Whichever approach is adopted, the operative discipline is to treat classification, documentation, and risk assessment as jurisdiction-specific exercises rather than assuming compliance under one framework translates to compliance under another.

4.3. Calibrating compliance investment to actual exposure

A recurring tension in transatlantic AI governance is the risk of over-engineering compliance relative to genuine regulatory exposure. The AI Act is comprehensive but demanding; US state frameworks are less prescriptive but multiplying; federal sectoral enforcement (FTC s.5, EEOC, sectoral regulators) sits on top of both. Organisations should calibrate investment to actual risk profile, avoiding two equal and opposite traps: building to AI Act standards for systems with no EU nexus, and assuming US-only governance is sufficient where EU exposure exists.

Practical strategies include staged deployment, starting in lighter-regulated jurisdictions before expanding into more demanding ones, particularly for systems with material classification uncertainty; engagement with regulatory sandboxes—both the AI Act's national sandboxes and US state-level pilots—to test higher-risk systems under supervision; compliance-by-design integration of regulatory requirements into AI development from inception, rather than retrofitting once systems are operational; and investment in cross-framework documentation infrastructure (model cards, system cards, impact assessments) capable of feeding both EU technical documentation and US-facing transparency and risk reporting.

4.4. Where cross-framework efficiency is genuinely available

Notwithstanding the structural divergence, meaningful efficiencies are available across frameworks. NIST AI RMF and ISO/IEC 42001 map with reasonable fidelity onto AI Act concepts: govern/map/measure/manage functions translate into risk management, data governance, technical documentation, and post-market monitoring; ISO 42001 management system controls

correspond to many AI Act quality management obligations. Organisations investing in AI Act harmonised-standard compliance can extract significant value for US-facing documentation. Unified impact assessments—combining fundamental rights, data protection, and consumer protection dimensions—can address EU and US requirements within a single process where carefully scoped. Incident response procedures aligned to art.73 AI Act, art.33 GDPR, NIS2, DORA, and sectoral US reporting can run on a single backbone with jurisdiction-specific triggers.

These efficiencies are real but should not be overstated. Classification, conformity pathways, and technical standards remain structurally distinct, and genuine dual-track work will be unavoidable in a number of areas—particularly pre-market conformity assessment, CE marking, machine-readable synthetic content labelling, and the formal documentation that EU notified bodies and market surveillance authorities expect to see.

5. Preserving innovation capacity

The choice of compliance architecture has material consequences for an organisation's capacity to innovate. The risk is not that regulation prevents innovation outright—it does not—but that poorly calibrated compliance imposes friction across the AI development lifecycle, slowing experimentation and raising the threshold at which new use cases become commercially viable. Three practical considerations follow.

First, *governance design should distinguish between AI systems and AI experimentation*. The AI Act's high-risk obligations and US state regulatory triggers attach principally to systems placed on the market or used to make consequential decisions. Internal experimentation, prototyping, and research that does not affect external parties typically falls outside the most demanding obligations, and governance architectures that fail to draw this distinction routinely impose conformity-assessment-level controls on workstreams where they are not legally required.

Second, *human oversight and accountability mechanisms should be designed for the system, not for the audit*. The temptation in transatlantic compliance is to design oversight primarily for documentary defensibility—sign-offs, attestations, recorded reviews. These have value, but only where they reflect substantive review by personnel with the competence, authority, and information to intervene meaningfully. Otherwise they

generate compliance theatre rather than risk reduction, and tend to attract increased regulatory scrutiny over time as authorities develop more sophisticated audit methodologies.

Third, *contractual and procurement frameworks should reflect the actual allocation of regulatory roles*. AI Act roles (provider, deployer, importer, distributor) and the US developer/deployer distinction allocate obligations differently, and contractual frameworks that fail to track those allocations leave gaps in due diligence, warranties, indemnities, audit rights, and incident reporting. Contracting templates have matured substantially over the past 18 months, but the underlying assurance layer—pre-contractual due diligence, ongoing evidence collection, monitoring of model updates and sub-processors—remains uneven, and is now a primary source of regulatory and litigation exposure.

6. Outlook

The regulatory divergence analysed here is unlikely to resolve in the near term. US state-level experimentation will continue, and the prospect of comprehensive federal AI legislation remains uncertain, particularly given the current administration's deregulatory orientation and active federal pre-emption efforts. The EU AI Act is still in its implementation phase, with significant secondary legislation and harmonised standards yet to be finalised.

For transatlantic businesses, the practical implication is that compliance architectures must be capable of *adapting* to new requirements as they emerge—in both jurisdictions—without requiring fundamental redesign each time. The organisations that handle this best will be those that treat AI governance as an ongoing programme rather than a one-off project, calibrate compliance investment to actual exposure rather than worst-case interpretation of either regime, and resist the temptation to treat the choice of governance baseline as a binary decision.

The deeper point is that the US and EU are not merely regulating AI differently; they are regulating different things. The US is regulating the marketplace; the EU is regulating the technology. Until that gap closes—and there is no near-term indication that it will—organisations operating on both sides of the Atlantic will need to maintain governance architectures sophisticated enough to address both, without losing sight of the innovation capacity that makes the technology worth regulating in the first place.