

Workstreams & Research Directions

Institute for Law & AI | July 13, 2026

Overview

This document outlines LawAI's current workstreams and research directions. Each section introduces the workstream, briefly outlines the reasons why we are excited about it, and highlights a few select research questions we are investigating. Our workstreams are interconnected, and some projects will fit under multiple workstreams.

Our research is split among three teams, with each team owning one or more workstreams. They are:

1. [United States Law and Policy](#)
 - a. [Institutions and Procedures](#)
 - b. [Law and Compute](#)
 - c. [Liability and Insurance Law](#)
2. [Legal Frontiers](#)
 - a. [Law-Following AI](#)
 - b. [Automated Governance and the Rule of Law](#)
3. [EU Law](#)

United States Law and Policy

Our primary focus is improving law and policy in the United States. These workstreams identify key areas of domestic law and policy where additional research could help advance security, welfare, and the rule of law.

Institutions and Procedures

This workstream aims to produce insights into the institutional and procedural factors that impact the effectiveness of AI regulation. While some of our workstreams focus on substantive recommendations (i.e., what would be the best policy on a given topic?), this workstream focuses on the procedures that governments can use to achieve their substantive policy goals. AI is a rapidly developing emerging technology. The capabilities that advanced AI systems will possess in the future, and the regulatory challenges that those capabilities will pose, are difficult to predict. This being the case, attempts to implement specific substantive policies in advance will often miss the mark. Regulators will need to learn, adapt, and account for uncertainty—all while responding to technological developments more rapidly than ever before. Even when the best course of action is unclear in the present, we can [equip regulators](#) with the tools, procedures, skills, information, and authority to develop better answers over time. This approach allows us to account for empirical and normative uncertainties and equip governments to tackle risks and opportunities as they arise.

Relevant questions in this workstream include:

- What authorities do regulators already have to enact regulations, gather information, monitor for compliance, and enforce the law? What new authorities would they benefit from, and how could they be created?
- What information do regulators need to regulate effectively, and how can they gather it?
- How can regulators promote information sharing and cooperation between regulators, other governments, and (where appropriate) private industry?
- How can regulators marshal the necessary expertise through hiring, government contracting, expert consultation, or other mechanisms?
- How can regulators make better decisions under uncertainty? How can they evaluate the effectiveness of different regulatory approaches, including under time constraints? How can they more proactively prepare for future regulation and identify promising areas on which to focus?
- How can regulators make better decisions in emergency situations?
- How can we ensure democratic accountability and align regulators with public values? How can we prevent agency capture or improper influence?
- How can agencies consistently and quickly update regulations when they have become ineffective, while preserving other values like fair notice, reliability, and opportunity to comment?
- How can regulatory structure preserve non-partisanship?

Law and Compute

Advanced AI models use a large amount of specialized computer hardware (“compute”) during development and deployment. Compared to the other primary inputs to AI systems—data and algorithms—compute is uniquely detectable, excludable, and quantifiable. Additionally, the supply chains that facilitate the manufacture or acquisition of compute are unusually concentrated. Compute is therefore a promising policy target for AI governance.

Realizing this, various regulatory and legislative proposals are beginning to leverage compute as a means of influencing advanced AI development and deployment. Improved legal design of these proposals could increase their effectiveness while minimizing costs such as regulatory burdens and privacy risks.

This workstream aims to help policymakers leverage compute governance in an effective and balanced way. Specific research questions that could advance this goal might include:

- What role can compute thresholds play in AI governance efforts?
- How can we design measures of compute usage that are robust to attempts at circumvention?
- How should compute thresholds account for algorithmic and hardware progress?
- How can we design compute monitoring policies that would have acceptable costs in terms of important values like privacy, information security, and economic competitiveness?
- What existing authorities do state governments have to regulate the transfer and usage of compute?
- How could states leverage large-scale computing capacity to defend against emerging threats from asymmetric actors?

Liability and Insurance Law

What sets this workstream apart is its unit of analysis: harm. Whereas many legislative or regulatory approaches create *ex-ante* requirements, liability and insurance impose *ex-post* penalties after harm has occurred. In doing so, they define what types of injuries (physical, monetary, foreseeable, preventable, etc.) deserve relief and clarify who is responsible when an injury occurs.

There are benefits to this *ex-post* approach: by waiting for harm to occur, liability and insurance can tailor relief to the specific circumstances that develop. This flexibility seems particularly valuable for addressing AI-facilitated harms, where it is often difficult to predict (or achieve consensus on) who will be harmed, how it will happen, and which mitigation measures are technically feasible and sensible. In this context, liability can help set general obligations and incentives in areas where regulatory consensus is difficult to achieve, impose penalties once we know what harms have actually materialized, and selectively enhance enforcement and compliance where regulations are in place. Insurance will mediate these effects by altering how the costs of injury are evaluated, compensated, and spread across actors.

But this focus on retroactive relief is only part of the story. Liability and insurance also set forward-looking incentives: they can reward proactive safety innovation; encourage the development of professional standards and safety procedures; encourage the hiring or designation of specific safety personnel; encourage improved security practices; shape monitoring, early detection, and emergency response; encourage or discourage attribution, model control, data collection, documentation, supply chain tracking, and information sharing;

develop the government's investigatory and enforcement capabilities; incentivize (or subsidize) better risk prediction and modeling; alter the economic cost of engaging in risky behavior; and incentivize private risk monitoring and investigation.

There is growing legislative interest in liability reform on both the federal and state level, and private insurers have indicated interest in creating new, AI-specific insurance products. This workstream will address questions such as:

- How will existing liability law and insurance markets operate by default? What will they do well, and what gaps will emerge?
- What legal limitations constrain reforms?
- What can we learn from historical case studies of liability regimes and insurance markets in other industries?
- Who are the relevant actors? What tools, mechanisms, and approaches are available to them?
- What specific details will determine the effectiveness of liability and insurance systems (e.g., what AI systems will be covered, what specific duties will apply and how will those duties adjust over time, how will and should responsibility be allocated between different actors in the value chain, who will enforce the law, and how disputes will be adjudicated)?

Legal Frontiers

As AI technology advances, we will likely need new legal and policy paradigms to help address the unique challenges that accompany technological progress. The Legal Frontiers team aims to incubate new law and policy proposals that are simultaneously:

1. **Anticipatory**, in that they respond to a reasonable forecast of the legal and policy challenges that further advances in AI will produce
2. **Actionable**, in that we can make progress within these workstreams even under significant uncertainty
3. **Accommodating** to a wide variety of worldviews and technological trajectories, given the shared challenges that AI will create and the uncertainties we have about likely developments
4. **Ambitious**, in that they both significantly reduce some of the largest risks from AI while also enabling society to reap its benefits

Right now, the Legal Frontiers team owns two workstreams:

1. [Law-Following AI](#)
2. [Automated Governance and the Rule of Law](#)

Both workstreams aim to preserve the rule of law as AI systems are integrated into important governmental functions.

Our goal with the Legal Frontiers workstreams is to grow new legal and policy paradigms until they can be sustained by a larger research community, both internally and externally. We will therefore aim to identify and incubate new workstreams as existing ones mature.

Law-Following AI

Agentic AI systems will occupy increasingly important roles in organizations of societal importance, including large corporations and governments. However, agentic AI systems will face different legal incentives than the humans they complement or replace: in particular, the threat of direct liability or punishment may deter AI agents significantly less than similarly situated humans.

To help close this deterrence gap, we have developed the concept of “Law-following AI,” or LFAI: AI systems that are designed to follow a broad set of legal duties. Core publications to date on LFAI include:

- Cullen O’Keefe, Ketan Ramakrishnan, Janna Tay & Christoph Winter, *Law-Following AI: Designing AI Agents to Obey Human Laws*, 94 **Fordham L. Rev.** 57 (2025), <https://fordhamlawreview.org/issues/law-following-ai-designing-ai-agents-to-obey-human-laws/>.
- Cullen O’Keefe, Christoph Winter, Matthijs Maas, Janna Tay, et al., *Proceedings of the 2025 Workshop on Law-Following AI* (Lawfare Rsch. Rep. No. 26-8, 2026), <https://www.lawfaremedia.org/article/proceedings-of-the-2025-workshop-on-law-following-ai>.
- Noam Kolt, Nick Caputo, et al., *Legal Alignment for Safe and Ethical AI*, **Transactions on Mach. Learning Rsch.** (2026), <https://arxiv.org/pdf/2601.04175>.

We are also interested in analyzing how LFAI can apply to a broader set of AI governance issues. For example, we have analyzed how LFAI could promote treaty formation and stability in Matthijs Maas & Tobi Olasunkanmi, *Treaty-Following AI* (Inst. for L. & A.I. Working Paper No. 1-2025, 2025), <https://law-ai.org/treaty-following-ai/>.

We remain interested in answering key questions that bear on the design of LFAI policies, especially those that would require governmental AI systems to be law following. Examples [include](#):

1. How can we evaluate whether a given AI system is law following?
2. How should an LFAI manage legal risk amidst pervasive and irreducible legal uncertainty?
3. In which governmental settings should AI agents be required to be law following?
4. Which laws should LFAIs be required to obey?
5. How can LFAI requirements in the executive branch be reconciled with the president's constitutional prerogatives?

We are also interested in practical projects aimed at accelerating the adoption of LFAI policies. Such projects include:

6. Developing technical methods for evaluating whether AI systems reliably follow the law
7. Developing draft legislative and regulatory text that imposes requirements related to LFAI
8. Advising on integration of legal alignment into AI model constitutions/specifications

Automated Governance and the Rule of Law

This workstream aims to resolve possible tensions between ambitious governmental adoption of AI and broader rule of law values.

We start from the premise that AI systems will be key engines of state capacity in the 21st Century. Careful integration of AI systems into core military and civilian functions can enhance national security, deliver better services to citizens, and reduce government waste.

On the other hand, governmental AI systems pose serious risks. Misaligned AI systems with privileged access to governmental infrastructure would be uniquely well-positioned to cause harm. AI systems could also be potent tools of oppression in the hands of a tyrant. AI-enabled weapons could destabilize the global balance of power.

This tension would be easily resolved if a uniformly precautionary approach to governmental AI was realistic or desirable. It is neither. Strategic competition dictates that governments will adopt the same technologies as their competitors. The regulation of AI will be impotent without AI-augmented regulators. Governments will need to use trustworthy state-of-the-art AI systems to harden themselves against more vexing future threats.

We therefore posit that democratic governments may need to make bold plans to channel AI capabilities towards their most important governance tasks. This workstream aims to generate new proposals for how they can do so while safeguarding the basic rights of their citizens. [Law-following AI](#) is one such proposal, but we think there are likely dozens more waiting to be formulated.

We plan to release a report listing some initial ideas for automated governance this year. This workstream also builds on our previous work in Cullen O'Keefe & Kevin Frazier, *Automated Compliance and the Regulation of AI* (Inst. for L. & A.I. Working Paper No. 1-2026, 2026), <https://ssrn.com/abstract=6017756>.

EU Law

On 1 August 2024, the landmark EU Regulation 2024/1689 (“AI Act”) came into force across all 27 EU Member States, creating what is currently the world’s most comprehensive legal framework to regulate artificial intelligence (“AI”). The significance of the AI Act is underscored by its (indirect) impact on some comparative initiatives, such as Californian SB 53, and by the ongoing regulatory uncertainty in many countries. As a pioneer regulation, the AI Act aims to promote safe, human-centric and trustworthy AI, and also influence global standard-setting in the field, echoing the “Brussels effect” seen in other regulatory domains. Therefore, high-quality legal research into the AI Act has real potential for impact: robust insights may not only support an effective EU regulatory approach, but will also be relevant to evaluating how best to design effective AI governance regimes in other jurisdictions.

The AI Act is a complex and ambitious legal instrument: it comprises 180 Recitals, 113 Articles, and 13 Annexes, which impose horizontal obligations across the AI value chain i.e., for providers, deployers, importers, and distributors of AI models and systems with a nexus to the EU market. The obligations imposed by the AI Act apply on a staggered basis between 2 February 2025 and 2 August 2027. Given their potential for good and harm, technological novelty and the consequential regulatory uncertainties, our [current research](#) focuses on general-purpose AI (“GPAI”) models, especially those that present systemic risk (Chapter V AI Act) and the related topics of enforcement and fines (Chapters IX and XII AI Act). We also work on questions turning on the positioning of AI agents within the AI Act’s regulatory framework. A further significant element of our work relates to the issuance of secondary instruments under the AI Act, such as codes of practice, guidelines, common rules and harmonised standards. Such instruments raise important legal questions given their aims to contribute to the objectives, application, implementation and enforcement of the Regulation.

Against that backdrop, the EU Law workstream addresses questions such as:

- How can the criteria for classifying or designating GPAI models with systemic risk be further defined?
- What legal mechanisms are available to distinguish between different levels of systemic risk caused by GPAI models?
- What is the legal effect of the EU GPAI Code of Practice (and compliance or non-compliance with it), given that it lacks the ‘presumption of conformity’ (Article 53(4), Article 55(2)) that is offered to harmonised standards?
 - How does the EU GPAI Code of Practice (i) shape industry practice within the EU and beyond and (ii) influence enforcement strategies?
 - How can legal mechanisms, e.g., dynamic references and rapid updating mechanisms, ensure that the EU GPAI Code of Practice is future-proof and reduce risks of misspecification, while simultaneously offering a necessary degree of legal certainty?
- Which provisions of the EU GPAI Code of Practice are (most) effective for assessing and mitigating systemic risk?
- How can and should the EU GPAI Code of Practice be updated effectively and efficiently?

- How can conflicts in the scope and obligations of the AI Act be resolved?
 - What interpretative approaches resolve the apparent conflict between the exclusion of research and development in the AI Act's scope (e.g. under Article 2(8)) and its inclusion in substantive obligations relating to risk assessment and mitigation (e.g. under Article 55(1)(b))?
 - What role does the EU GPAI Code of Practice play to ensure oversight of the development of GPAI models?
 - When would a GPAI model provider's own use of the GPAI model constitute a "placing on the market" (Article 3(9), Article 55(1))?
- When does a modifier of a GPAI model become the provider responsible for compliance obligations of the modified GPAI model under the AI Act? How should a regulator delineate substantial modifications which shift provider obligations from the original developer to the modifier? How do different legal interpretations and policy decisions in these respects alter incentives for the various actors along the value chain?
- How is fault, and relatedly liability, distributed between different actors in the AI value chain, in particular in view of downstream modifications of a GPAI model or subsequent downstream integration into an AI system?
- Should a GPAI model be considered a product itself or a component of an AI system within the meaning of Directive (EU) 2024/2853 (Article 4(1) and (4))?
- Does the AI Act apply to AI agents? If so, which obligations apply to AI agents and are they suitable for the rapid technology change in this sub-field of AI?
- What is the difference between monitoring, supervising, investigating and enforcing obligations imposed on GPAI model providers by the AI Act? Which of the Commission's powers fall into which category and what procedural protections apply at which point?
- How is the principle of proportionality reflected in the hierarchy of monitoring, investigative and enforcement powers, in the conditions for legitimising action to be taken by the Commission and/or AI Office (such as 'necessary actions to monitor...' per Article 89(1), or 'necessary for the purpose of assessing compliance' per Article 91(1)), and in the fixing of fines under Article 101?
- What are 'structured dialogues' (per Articles 91(2), 92(7) and 93(2))? How can they be used effectively and appropriately in accordance with any applicable procedural safeguards?
- To what extent can measures or actions taken voluntarily by GPAI model providers usefully play a part in the Commission's/AI Office's monitoring, supervision or investigation of GPAI model providers?
- What are possible implications of the precautionary principle as a general administrative principle of EU law (Article 191 TFEU) for the exercise of enforcement powers of the Commission, as well as for other regulatory responsibilities, such as the designation of GPAI models with systemic risks?

- How can other AI Act governance bodies, such as the Scientific Panel and the AI Board, be leveraged to aid or support enforcement processes?
- How, and to what extent, can the AI Office effectively exercise its enforcement powers with respect to GPAI model providers outside the Union?