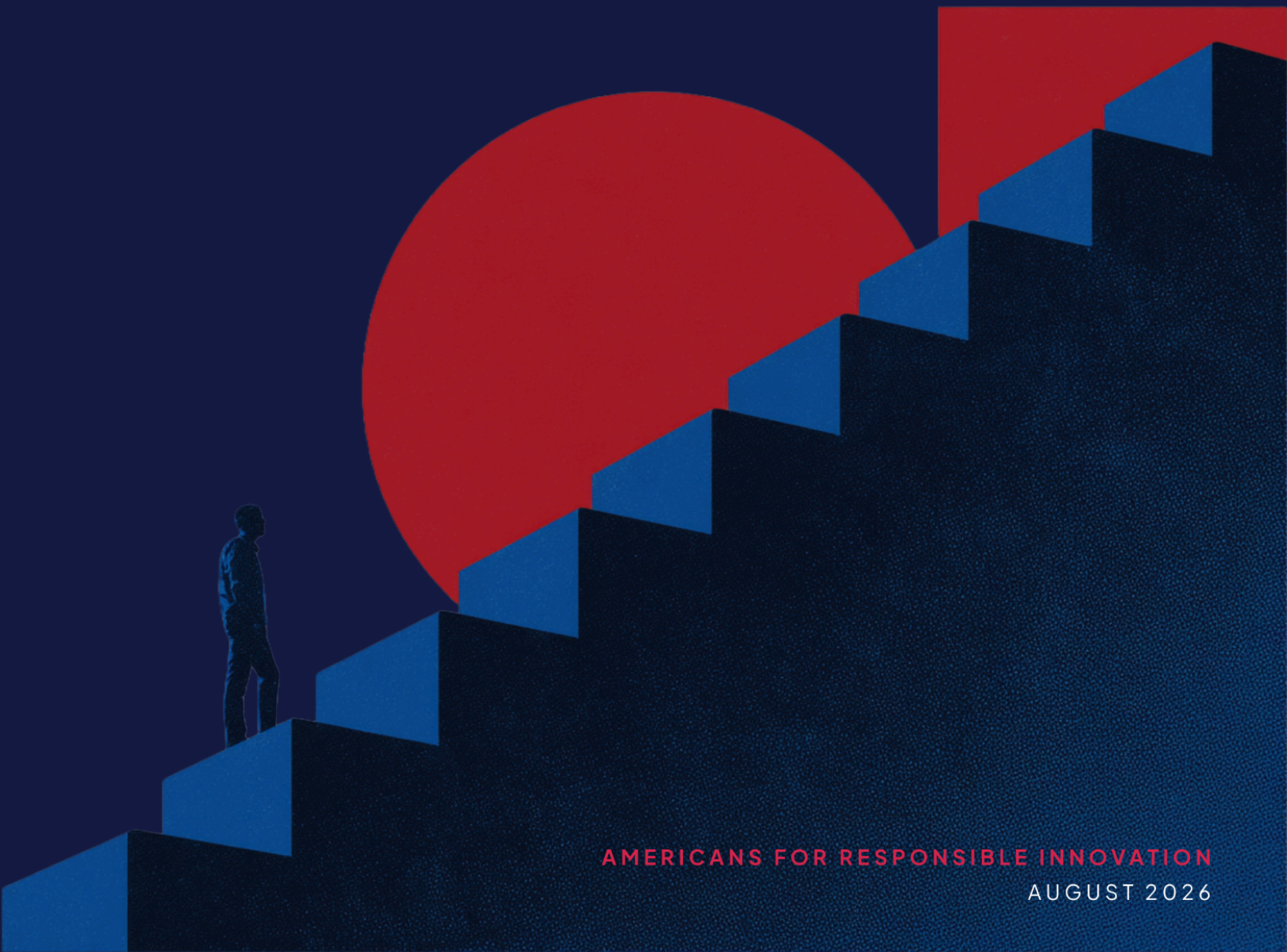




RESPONSIBLE INNOVATION AT THE FRONTIER:

A Blueprint for Federal AI Governance



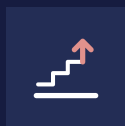
AMERICANS FOR RESPONSIBLE INNOVATION

AUGUST 2026

EXECUTIVE SUMMARY

Over the last two years, frontier AI has advanced at a pace that has forced policymakers to reckon with a difficult reality: the technology is too consequential to be governed by a laissez-faire approach. This recognition has prompted a growing push for government action. Right now, the Trump administration is regularly preventing frontier model releases based on unspecified criteria; three states have enacted frontier safety framework statutes; leading developers (including Anthropic, Google, and OpenAI) have each published a blueprint for federal legislation; and Congress is weighing a litany of its own measures. But, whatever final form frontier governance takes, it must ensure that developers' safety practices adequately address the risks of frontier AI.

ARI's blueprint for federal AI governance is designed to promote safe frontier AI development in America. The blueprint is built around three governance functions any federal proposal should incorporate. To set adequate safety *standards*, the federal government should hold every covered developer's published safety framework to minimum federal standards that define adequacy. To verify compliance with those standards, there should be independent *assurance*. To secure *transparency* into frontier AI development, there should be federal visibility into both internally deployed frontier models and the growing automation of AI research and development.



01

Standards

Minimally Adequate Federal Safety Standards In Publicly Available Frameworks.



02

Assurance

Independent Assurance From The Government That Federal Standards Are Being Met.



03

Transparency

Regular, Sustained, And Confidential Federal Visibility Into The Automation Of AI Research And Development.

ARI's blueprint describes a model for federal governance of AI that fits the moment. The blueprint is intentionally limited to a handful of well-resourced frontier developers, where the most salient risks are concentrated. Federal compliance satisfies state requirements that are equivalent or less demanding, leaving states free to legislate above the federal floor. The whole plan sunsets in nine years unless Congress reauthorizes it. In return, this blueprint delivers a binding and adaptable response to the challenge of ensuring adequate frontier safety practices.

A BLUEPRINT FOR FEDERAL AI GOVERNANCE

Responsible Innovation at the Frontier

01 Standards

MINIMALLY ADEQUATE FEDERAL SAFETY STANDARDS IN PUBLICLY AVAILABLE FRAMEWORKS.

The regulator sets standards, integrating input from industry practice, government technical analysis, and independent research. Standards never fall below a level already set and reflect best practices as supported by leading safety science. When an imminent threat outruns the standards in force, the government may intervene under a residual emergency authority, subject to prompt judicial review.

WITHOUT IT

If companies are allowed to define safety on their own terms, there are no standards.

02 Assurance

INDEPENDENT ASSURANCE FROM THE GOVERNMENT THAT FEDERAL STANDARDS ARE BEING MET.

The government examines each covered developer from the outset, testing models and inspecting safety practices to confirm that a developer follows its framework and that the framework is adequate. Although assurance examination may be augmented by a future independent verification market accredited by the government, the government always serves as a backstop, and it alone imposes every binding consequence.

WITHOUT IT

Without assurance, the floor for frontier AI safety is based on developer-made promises.

03 Transparency

REGULAR, SUSTAINED, AND CONFIDENTIAL FEDERAL VISIBILITY INTO THE AUTOMATION OF AI RESEARCH AND DEVELOPMENT.

Covered developers file quarterly disclosures on the capability most likely to cause frontier development to outrun any fixed standards review cycle. Disclosures center on auditable metrics and shed light on the state of human oversight of automated AI R&D, and any material-change disclosures are filed within four business days.

WITHOUT IT

When the frontier labs are the only ones that can determine whether standards and assurance are still meaningfully measured against the very edges of the frontier, there is no transparency.

Governing Innovation at the Frontier

ARI's blueprint combines three key governing principles – standards, assurance, and transparency – into a single comprehensive and coherent vision.

Binding safety standards that are set by the federal government for each developer's framework are the keystone of this vision. The standards supporting this vision are not set in a vacuum. They are developed collaboratively and cyclically, based on best industry practice, government technical analysis, and independent expertise.

Even adequate standards that go unchecked are meaningless. Once the government establishes these binding standards, government examiners assure that every developer adheres to the standards in its framework. Eventually, and only when a robust market that can resist capture exists, government-accredited independent verifiers assist with this assurance function. Even in a fully mature market, final adjudication of whether a developer has met its obligations to the published minimal standards rests with the government.

ROLES AND RESPONSIBILITIES

Who does what: standards, assurance, and transparency

The role of the government in overseeing AI development must, at times, be augmented or assisted, but the government must never be replaced as the central regulatory authority.

PARTY	Standards	Assurance	Transparency
Government	▪ Sets binding adequacy standards on a fixed cycle ▪ Issues penalties for noncompliance ▪ Issues residual emergency orders subject to prompt judicial review	▪ Conducts the primary assurance function ▪ Accredits and supervises verifiers ▪ Assigns one verifier per covered risk domain in the final phase	▪ Receives confidential filings ▪ Publishes annual aggregate report ▪ Directs for-cause examination by qualified accredited verifiers
Covered developers	▪ Publish and maintain safety frameworks meeting federal standards	▪ Undergo government examination ▪ Choose accredited verifiers and receive government-assigned audits	▪ File quarterly and material-change disclosures on AI R&D automation
Accredited private verifiers	▪ Contribute information and assessments to the next standards cycle	▪ Verify framework compliance under government supervision	▪ Conduct examinations directed by the government

Finally, each covered developer must provide the government with the needed transparency into the automation of its own research and development. This transparency is provided via confidential federal disclosure to help ensure that the standards-setting and review cycles do not miss critical advances on the frontier that are occurring as the result of increasing automation.

If these three mutually reinforcing principles still fail to prevent serious threats to the public from materializing, the government may take action through the courts to block or halt system development, internal deployment, or external usage. In some parts of ARI's blueprint the role of the government is augmented or assisted, but the government is never replaced as the central regulatory authority.

SCOPE

A covered developer

The blueprint is intentionally limited to a handful of well-resourced frontier developers, where the most salient risks are concentrated.



COMPUTE THRESHOLD

10^{26} FLOP

Has trained or initiated the training of a frontier model, a model trained using, or undergoing training intended to use, at least 10^{26} floating-point operations of compute; and



SPEND THRESHOLD

AND

\$100 million

Has spent at least \$100 million in the past year on aggregate training of AI models.

Only the largest training runs today reach or exceed 10^{26} floating-point operations. However, compute is an imperfect proxy for capability, because algorithmic efficiency gains mean that models trained with less compute may eventually match the performance of current frontier models. As such, the regulator must review the compute threshold on a fixed cycle and adjust it as the proxy degrades, keeping coverage focused on the frontier itself. In time, and with technical assistance from the Center for AI Standards and Innovation (CAISI), the regulator should adopt a workable capabilities-based definition of a frontier model.

To promote responsible frontier innovation, federal frontier governance should cover AI used internally by a developer in addition to any public or external offerings; there is already evidence that dangerous capability thresholds can be reached by internally deployed frontier models before any of those models are available to the public. As such, the term “deployment” in ARI's blueprint refers to both internal and external uses or deployments of frontier AI. Obligations relating to standards and assurance apply with respect to covered AI models, regardless of who is using them or where they are being used. This is a critical clarification of scope in ARI's blueprint, as covered risk domains already include harms arising from internal deployment, and obligations relating to automated AI R&D inherently implicate internal-facing technologies and use cases.

Standards: An adaptable and rising floor

Current state laws that require developers to publish safety frameworks were necessary steps toward frontier governance. However, these laws do not guarantee that frameworks meet a standard of adequacy. Since the deepest technical expertise sits with the developers, they should have real input into the content of any standards. However, input is different from authorship, and standards-setting must always rest with the government. As part of ARI's framework, each developer publishes and maintains a safety framework addressing myriad risks with direct relevance to either public safety or national security. ARI developed its risk domains with current developer testing capabilities in mind.

STANDARDS

Covered risk domains

Each framework must address its model's capabilities with respect to each of the following:

- 01  chemical, biological, radiological, nuclear, and explosive threats (CBRNE);
- 02  offensive cyber capability;
- 03  automated AI R&D;
- 04  harmful manipulation, including large-scale influence operations, particularly those directed by foreign states; and
- 05  autonomy, which includes misalignment and loss of control.

For each covered risk domain, the developer must include certain details as part of its framework, such as:

- how each risk is evaluated (including specific capabilities testing and red-teaming) and why those techniques are sufficient;
- thresholds at which the developer must respond or enact mitigations, expressed in terms of model behavior or evaluation results; and
- at what point the developer notifies federal authorities.

While new risk domains may be added, regulators may never summarily remove a covered risk domain since each is codified in statute.

Whether a threat emerges from one risk domain or several, the published standards cover it. All standards are rooted in model capabilities and operational conduct, not hypothetical or supposed capabilities. While this initial set of risk domains serves as the baseline set in ARI's proposal, over time the government may identify new risk domains as new threats from models emerge. While new risk domains may be added, regulators may never summarily remove a covered risk domain since each is codified in statute. For each covered risk domain, the developer must include certain details as part of its framework, such as:

- how each risk is evaluated (including specific capabilities testing and red-teaming) and why those techniques are sufficient;
- thresholds at which the developer must respond or enact mitigations, expressed in terms of model behavior or evaluation results; and
- at what point the developer notifies federal authorities.

To help provide direct accountability, each developer must name an officer responsible for the framework's maintenance and implementation. Developers must file frameworks with the federal regulator, which publishes every framework in a public catalog. For publicly released models, information about model evaluation and results must be made publicly available concurrently with model release through the publication of certain transparency documents (e.g., system cards). In ARI's blueprint, evaluations and risk assessments should not end with commercial deployment, and risk reports must be published on a regular basis.

Framework publication is only the first step in ARI's proposal. After the initial framework filing cycle, the regulator periodically reviews and updates minimum standards for every covered framework. These updated standards draw on developers' published commitments and research, technical analysis from CAISI, independent safety research, and the aggregated findings of assurance examinations. While standards should continually evolve and adhere to the best practices supported by safety science, they should never promote developer laxity. The floor set during each cycle ensures that developers cannot backslide on safety commitments or deploy systems that exhibit previously prohibited behavior. Advances in AI safety may ultimately permit a relaxation of a framework's technical standards, but ARI's blueprint ensures that the public is never exposed to lower safety margins than those set in the previous cycle. This resolves the instability currently at the heart of self-regulation, where developers can revise their own safety commitments at will.

Ultimately, all standards-setting rests with the regulator's independent judgment rather than averaging industry custom or defaulting to industry best practices. Proven feasibility informs the floor but does not cap it, so a developer that demonstrates stronger measures can prompt the government to strengthen the subsequent safety requirements in future review cycles. The standards derivation process is informed by a blend of industry, government, and independent inputs, so no single party other than the regulator can set the floor.

Given the unpredictable advances in the frontier, there may be moments when fixed standards cannot anticipate a particular danger and the standards-setting process might be too slow to adequately prevent those dangers. In these extreme cases, the government should have emergency authority to temporarily halt dangerous developer activities, including development, internal deployment, or external release. This authority should include ordering specific mitigations and recalling a deployed model from service if it risks causing or has caused severe harms. This

emergency authority should run through clear, rules-based protocols. In ARI's blueprint, such authority is designated in statute, granted only to senior federal officials in named offices, and informed by the relevant technical advisors to those officials. If an emergency order is issued, it takes effect immediately and lapses after 72 hours unless the government petitions a court for an extension within that window. Once an extension has been sought, the order remains in effect while courts rule on an expedited basis, but lapses after a fixed period without a ruling. For a court to uphold an emergency order, it should require substantial evidence of present or imminent severe harm and an affirmative ruling should be subject to the normal appellate process.

This emergency authority lays out a system of foreseeable consequences, constrained by statutory criteria and judicial review. While an emergency authority is necessary, it should rarely be needed: in ARI's blueprint, the conditions for stopping development or deployment must be included in developer safety frameworks like any other requirement. This framework requirement helps reduce regulatory friction surrounding development or deployment intervention. For example, where a given mitigation's adequacy is disputed, the response should be to halt further development or deployment until the dispute is resolved.

Assurance: Continuous verification under government authority

Assurance is what makes any federal minimum safety standard meaningful. In ARI's blueprint, every covered developer is subject to two kinds of continuous examination that occur across each of the specified risk domains:

- ***Compliance assurance***, which assesses whether the developer is adhering to its published frontier safety framework; and
- ***Adequacy assurance***, which assesses whether the developer's framework meets the minimum federal safety standards currently in force.

Importantly, these forms of assurance are not interchangeable. Compliance assurance is a technical determination made against a fixed and published standard; it may be conducted by the government or by a delegated authority. Adequacy assurance is a determination that directly affects public policy decisions related to future rulemaking. This kind of assurance determination interrogates whether a developer's framework is sufficient and helps inform what the next set of safety standards should require. In ARI's blueprint, that judgment is never delegated. Irrespective of the market maturity of independent verification organizations (IVOs), the government never fully relinquishes its inherent public safety responsibility to conduct adequacy assurance.

Since no healthy IVO market exists today, both assurance functions must be conducted by the government, with the dedicated funding and hiring authority to do so. These assurance functions could be housed in a single entity (e.g., using the special hiring and pay authorities of a national lab) or spread across multiple agencies (e.g., led by the Department of Commerce and supplemented by agencies from other cabinet-level departments). In conducting these assurance functions, federal evaluation teams test frontier systems against the safety standards currently in force and against the capability thresholds listed in each developer's framework. At the same time, a dedicated oversight team, modeled after the Federal Reserve's standing supervision of the largest banking institutions, examines each developer's safety practices. Comprehensive assurance verification requires both evaluations and examinations: model evaluations can show that

capability thresholds were crossed without any corresponding actions by the developer, and examinations can show that evaluations were not conducted.

Regardless of who conducts the examination, these tests occur at critical times in the development and deployment lifecycle. Scheduled review is the principal assurance mechanism, but it is not the only one. As part of this blueprint, covered developers must also promptly report safety incidents to the government, and whistleblowers who disclose violations are protected in statute. Either of those channels can prompt a for-cause examination. Examiners must also immediately report any imminent risks of severe harm they observe that neither the standards in force nor the developer's framework addresses; findings of that kind can supply the basis for the residual emergency authority described earlier.

As frontier capabilities advance, the covered population is likely to grow; there may be more developers, more instances of significant fine-tuning, more agentic deployments, more risk domains, and more points in a lifecycle that require examination. The government can and must sustain the deep expertise required to maintain its assurance primacy, but it probably cannot sustain the throughput required for continuous examination of every covered developer and domain. The constraint most likely to require augmenting the government's assurance function with IVOs is one of volume, not competence. Even if a viable IVO market emerges, the government must permanently retain some deep and narrow examination functions. These functions should include, at a minimum, audits of compliance examinations, adjudication of examiner discrepancies, development of standards for novel and contested capability domains, continuous inspection of IVO work, for-cause examination capacity, and adequacy assurance in every domain.

ASSURANCE

The three-phase shift

If a shift of compliance assurance responsibility occurs, ARI's blueprint envisions it occurring in three phases. Examination work shifts only in domains where the regulator certifies, against fixed criteria, that sufficient accredited capacity exists.

PHASE ONE

The government conducts all assurance examinations of covered developers.

Government
Conducts

No second channel



PHASE TWO

Each developer must be examined by the government and one accredited IVO, selected by the developer, for each covered risk domain.

Government
Conducts

Accredited IVO
Developer-selected



PHASE THREE

Each developer selects an accredited IVO for each risk domain, and the government assigns the other.

Accredited IVO
Developer-selected

Accredited IVO
Government-assigned

CAPABILITY FLOOR

In every domain, cycle, and phase, the government conducts a share of the total examinations, so federal capacity remains a permanent core capability rather than a bridging function that sunsets as the IVO market emerges.

In ARI's blueprint, government capacity is a permanent and core capability, not a bridging function that sunsets as a new IVO market emerges. To help ensure the government is prepared to act as the assurance backstop ARI envisions, the government should conduct a share of the total examinations in every domain, cycle, and phase. Maintaining that kind of capability floor helps ensure federal examiners keep pace with the frontier, preserves the regulator's ability to reproduce IVO work, and provides a comparison basis by which IVO findings can be assessed.

It is not clear how long it will take for a robust IVO market to emerge, or whether it will emerge at all. Once a statutorily designated official (e.g., the Secretary of Commerce) certifies, based on fixed criteria, that initial IVO testing and examination capacity exists, IVOs may start to augment compliance assurance testing under the government's direction and by its authority. IVO accreditation is granted by domain and always rests with the government alone. Because accreditation occurs by domain, the examination configuration is also defined per domain. For example, a developer may need several IVOs to cover every risk domain. Modeled after the Public Company Accounting Oversight Board's supervision of public company auditors, the government regulator monitors each IVO's performance and inspects its work through its own retained backstop capacity. As compliance responsibility shifts toward a mature private market, examination work for compliance assurance shifts only in domains where the regulator certifies, against fixed criteria, that sufficient accredited capacity exists.

If the verification burden shifts, the liability faced by IVOs should also change. If the government is conducting examinations in the same domain as an IVO, it may be possible to relax IVO liability. In that instance, the government bears the burden of assurance since it provides the primary assurance for public safety. Initial, limited liability protection may also help catalyze an IVO market. Once compliance assurance examinations are principally driven by IVOs, that liability protection diminishes. Generally speaking, IVOs should not enjoy liability protection if their examinations are load-bearing; there could be nothing worse for frontier AI safety than a robust IVO market that enjoys immunity from its own negligent conduct. Modern regulatory history is replete with examples of that kind of failure. For developers, no examination outcome, whether by the government or IVOs, ever creates immunity, safe harbor, or a presumption against liability.

ARI's blueprint also includes market safeguards that help maintain IVO impartiality and independence, if and once that market exists. All examination fees, paid by developers and distributed to IVOs, pass through a regulator-administered fund. There is also a cap on the percentage of revenue that an IVO can receive from any single developer, calibrated for the market's size. Periodic rotation of IVOs and government examination teams is mandated, with revolving-door restrictions on personnel movement between both IVO and government examiners and the developers they examine and evaluate. Additionally, and irrespective of IVO market maturity, any marked difference between two examination assessments draws scrutiny. This additional scrutiny prompts review, potentially triggers a for-cause examination, and may affect the accreditation of the IVO involved if a bad faith determination was made.

Even with the government retaining adequacy assurance functions as part of this blueprint, the introduction of any kind of delegated examination invites comparisons to self-regulation. In ARI's blueprint, IVOs are always accredited contractors operating under government supervision;

they do not constitute some kind of self-regulating guild. IVOs have no rulemaking authority, no direct role in setting or interpreting the federal safety standards set by the regulator, and no separate governance body composed of other IVOs. Accreditation rests with the government alone and may be rescinded only by the government regulator. Even if a mature IVO market exists, the government issues every binding consequence. In short, ARI's blueprint is deliberately structured so that it does not depend on a robust IVO market ever emerging. In this way, the public interest is preserved by the government's guaranteed examination authority and talent acquisition efforts, rather than by whether the IVO market matures, stalls, or fails.

Transparency: Structured visibility into AI R&D automation

Automated AI R&D is a covered risk domain like any other, so every framework must address it by detailing evaluation methods, capability thresholds, and corresponding responses or mitigations in the event a threshold is crossed. It is also the one domain that receives a unique disclosure program, for three reasons. First, AI R&D has the potential to exacerbate other threats. Automation of the research pipeline can accelerate advances in every covered risk domain. In the long run, these advances could outrun the standards-setting and assurance cycles, which work only if frontier capabilities advance at a pace the cycles can match. Second, human oversight of AI R&D diminishes precisely as automation grows, so the window for establishing structured disclosure narrows as the need for it rises. Third, disclosure provides evidence regarding the ongoing dispute about AI's recursive self-improvement potential. The evidence that indicates whether automating AI R&D produces runaway acceleration or a stalling loop should accumulate where the government can see it. Today it does not, because it sits in internal deployments that developers are incentivized to keep private. As part of ARI's blueprint, disclosure must therefore be mandatory and uniform.

To secure transparency into AI R&D automation, ARI's blueprint adapts the SEC's 10-Q and 8-K reporting regime to the research pipeline. Each covered developer files quarterly disclosures on the automation of its AI R&D activities, plus an immediate material-change filing within four business days after discovery of a material shift in AI R&D automation between quarterly filings.

Because the categories describe broad functions, this program survives even if there are major changes in the way frontier AI is developed. These broad categories are set in statute, but specifics such as subcategories, measurement rules, and filing deadlines are delegated to rulemaking. Disclosures anchor on auditable measures, including the share of compute devoted to autonomous work in each category; measurements that developers estimate themselves (e.g., researcher-equivalent hours) act only as a supplement. Each filing must also report the state of human review against a predefined test of meaningful oversight, paired with auditable review-sampling rates (e.g., the share of autonomous outputs examined by a human). Meaningful human oversight of AI R&D is measured and defined by conditions including:

- whether humans can and do intervene in an AI system's outputs or actions before they are integrated into broader R&D workflows;
- whether a supervisor faces a volume and velocity of AI outputs that still permit substantive assessment or verification of each distinct automated R&D action;
- whether the system's human supervisor has sufficient information and resources to verify the system's safety; and
- whether review protocols ensure material human control over autonomous decision making.

AI R&D Disclosure Filings

TYPE OF FILING	TIMELINE	CONTENT
Regular	Quarterly	Auditable measures across the four disclosure categories; self-estimated figures as supplements
Material Change	Within four business days after a material change in R&D automation	Auditable measures of material R&D automation gains; any framework responses activated by the change

Filings under this part of ARI's blueprint are confidential and held under security controls commensurate with their sensitivity, and any trade secrets or proprietary information shared with the government are protected by Freedom of Information Act exemptions.

The blueprint requires public summaries only in the safety, evaluation, and oversight category, because disclosures in this category are less commercially sensitive and more useful for the public than others.

Developers must always maintain internal records, which the government will incorporate into periodic aggregated reports on the state of automated AI R&D. If a filing warrants further examination, the regulator may either examine the developer or contract an accredited IVO to conduct the examination under government authority, giving the IVO appropriate and secure access to those internal records.

Findings can tighten reporting obligations, bear on a developer's compliance with its own framework, and feed into the next standards-setting cycle. If a developer makes knowingly false statements or fails to retain the mandated records, then penalties result. However, a developer that identifies and corrects its own filing errors in good faith is protected. Under ARI's blueprint, no developer is ever penalized for a truthful disclosure alone.

Bounded by design

The standards, assurance, and transparency functions may be consolidated within a single authority (e.g., the Department of Commerce) or divided among several federal bodies. No matter the distribution of regulatory authority among agencies, the authority to penalize a developer that fails to adequately evaluate and address frontier risks is conferred by statute.

A TARGETED FEDERAL FLOOR

What the blueprint does not do

Agency discretion is bounded throughout: standards-setting, emergency authority, and enforcement are tied to statutory criteria.

PREEMPTION	Preemption operates by compliance equivalence rather than broad displacement of state law. Under this blueprint, federal compliance satisfies a state requirement only where the federal obligation covers the same conduct, and is at least as protective while other state laws remain untouched.
NO IMMUNITY	For covered developers, compliance creates no immunity, safe harbor, or presumption against liability under otherwise applicable state law.
STATE AUTHORITY	States remain free to legislate above the federal floor.
SUNSET	The whole regime sunsets nine years after enactment unless Congress reauthorizes it, with function-specific efficacy reviews staggered throughout the latter years of implementation, including a review of shifts in compliance assurance responsibilities, and a Government Accountability Office review of each function before sunset.

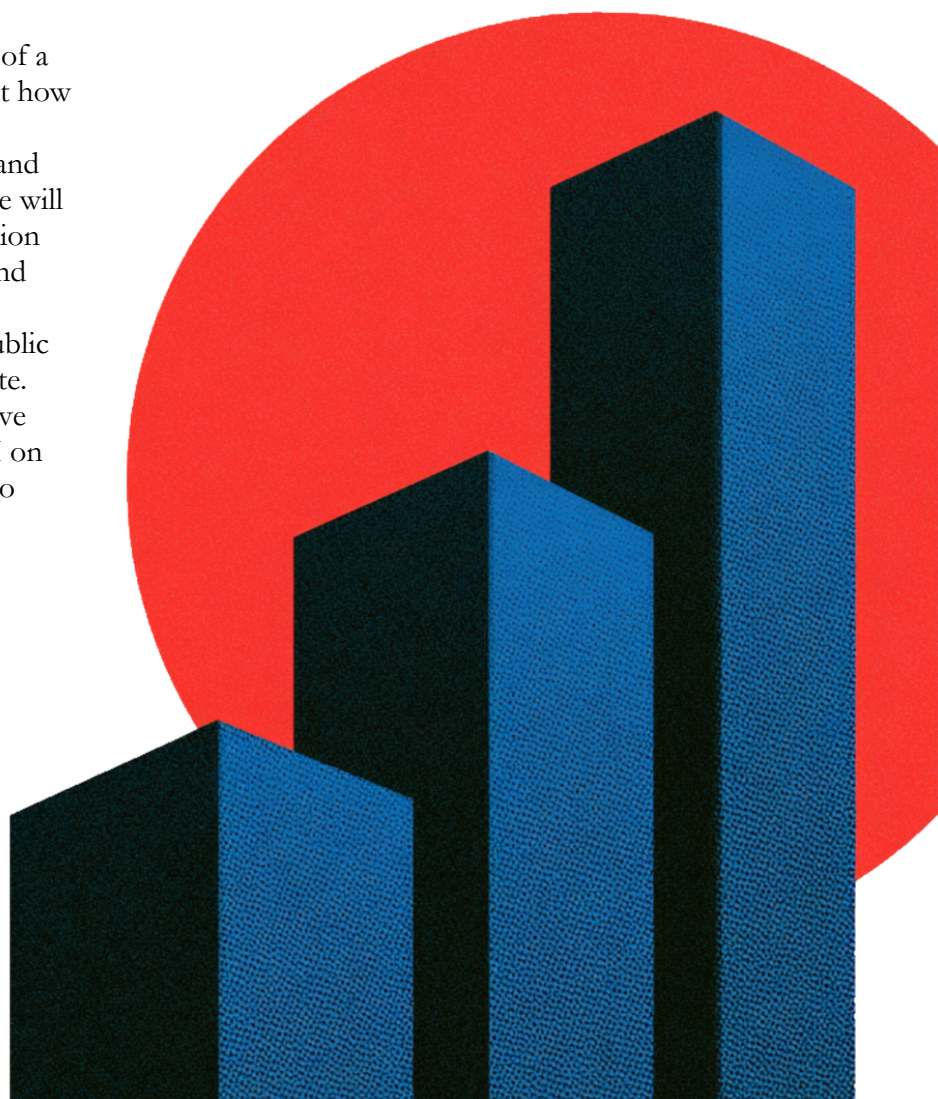
ARI’s proposed approach binds tightly where it must and lightly wherever it can.

RESPONSIBLE INNOVATION THAT MEETS THE MOMENT

The three pillars described in this blueprint are inseparable. If companies are allowed to define safety on their own terms, there are no standards. Without assurance, the floor for frontier AI safety is based on developer-made promises. When the frontier labs are the only ones that can determine whether standards and assurance are still meaningfully measured against the very edges of the frontier, there is no transparency. This blueprint unifies all three functions in order to mitigate the most serious risks that might arise from the development of frontier AI while still supporting responsible innovation by industry in a technological domain in which America cannot afford to fall behind.

ARI's blueprint offers a clear path toward responsible innovation at the frontier. It affects only the developers building the world's most powerful and cutting-edge models. Wherever possible, it defers authority to the states on all issues not directly related to its own scope of regulation. States remain, in most cases, free to raise standards wherever they feel the need while protecting the general public—wherever they live—with a clear federal standard for conduct. This clarity also helps the developers of this technology. ARI's blueprint provides a regulatory system that replaces ad hoc intervention and unclear safety standards with clear and impartial rules administered by technical experts and enforced by elected and appointed officials.

Finally, ARI offers this blueprint as part of a broader, ongoing national dialogue about how to regulate this critical technology. This blueprint signals ARI's legislative intent and aspiration, but specific statutory language will always matter and determine ARI's position on any final legislative proposals. We stand ready to engage with lawmakers, policy experts, civil society, industry, and the public to see this blueprint translated into statute. The opportunity to build a comprehensive federal framework governing frontier AI on the right terms is here, but the window to build it may be quickly closing.



ACKNOWLEDGMENTS & CONTRIBUTIONS

This blueprint is the product of months of dialogue and research, both among individuals and collectively by ARI's Policy Department. It was not created in a vacuum, and it is not the product of isolated thinking. Although many contributed to the concepts included in this blueprint, ARI acknowledges the following individuals for their contributions to its final form and substance:

AUTHORS

- Iskandar Haykel, Acting Policy Director (*Principal Author*)
- Bradford Kimball, Policy Intern (*Assisting Author*)
- Morgan Plummer, Vice President of Policy Design & Delivery (*Assisting Author*)

CONTRIBUTORS

- Dr. Justin Bullock, Vice President of Policy
- Q Chaghtai, Digital Director
- David Robusto, Associate Policy Director for Policy Strategy
- Giovanni Rocco, Communications Manager