

RVA CYBER RESEARCH · SEPTEMBER 13, 2026

AI regulation, *in their own words*

Dario Amodei, Sam Altman, Elon Musk, and Demis Hassabis on who should govern the frontier.

The four men examined here do not divide cleanly into advocates and opponents of AI regulation. Each has assigned government a role in governing the most capable systems. Each has also warned against applying the same burden to every model, company, and use.

Their argument is narrower and more consequential: **When should public authority enter, what evidence should trigger it, who should test a frontier model, and what power should the government have when a system fails?**

This paper traces the public answers given by Dario Amodei, Sam Altman, Elon Musk, and Demis Hassabis. It uses their authored texts, signed statements, formal testimony, full interviews, and official event transcripts. It does not use press summaries or extracted quotation reels. The purpose is not to prove that any speaker is right, consistent, or disinterested. It is to establish what each one put on the record, how that record changed, and where the records meet or conflict.

The record at a glance

- **Dario Amodei** moved from advocating mandatory testing and auditing to a more detailed case for binding, capability-triggered rules. He consistently opposed blanket restrictions on ordinary or low-risk systems.
- **Sam Altman** has supported exceptional regulation for advanced AI since at least 2015. By 2023, he advocated a threshold-based licensing agency and an international authority for superintelligence, while rejecting a fixed calendar pause as the main control.
- **Elon Musk** made the earliest forceful case in this record for proactive regulation. He signed the 2023 call for a 6-month pause, then moved toward a referee model and, by 2026, a short peer-review process with government as the backstop.
- **Demis Hassabis** began with capability benchmarks and government approval, then argued for staged regulation: update existing sector rules now, measure frontier capabilities, and add broader controls as the evidence becomes clearer. In 2026, he proposed a concrete public-private standards body and mandatory pre-release review.

The strongest common ground is not a shared agency chart. It is a shared regulatory shape: concentrate obligations on frontier capabilities, test before release, involve public authority, update the rules as evidence improves, and coordinate internationally.

The clearest direct conflicts concern **timing and control architecture**. In 2023, Musk signed a demand for a fixed 6-month pause; Altman told the Senate that a “calendar-clock pause” was the wrong frame. In 2026, Hassabis proposed an independent standards body with a voluntary period followed by mandatory approval;

Musk responded to that proposal with a different first step, peer laboratories reviewing each other before the government intervenes.

Method: a primary-source record, not a quotation contest

The source set contains 23 dated items: 6 for Amodei, 8 for Altman, 5 for Musk, and 4 for Hassabis. The record begins in 2015 and ends at the research cutoff. Every substantive timeline entry links to the full public source.

An item qualified if the named person wrote it, signed it, delivered it as testimony, or spoke in a full recording or official transcript. A company publication qualified only when it identified the person as author or speaker. The 2023 Future of Life Institute letter is attributed to Musk as a signatory, not as its author. The TIME interview with Hassabis is used as a direct question-and-answer source, with the publication's notice that the interview was condensed and edited. Automated video transcripts were checked against the full recordings and are paraphrased unless a short quotation is necessary.

The analysis uses three labels deliberately:

- **Documented overlap** means 2 or more speakers expressly endorsed the same principle in their own records. It does not imply a negotiated agreement.
- **Direct disagreement** requires a statement that rejects, displaces, or materially contradicts another recorded proposal.
- **Difference** means the speakers emphasized different risks or institutions without establishing a contradiction.

Silence is not disagreement. A change in detail is not automatically a reversal. Company interest can be relevant to a reader's evaluation, but this paper does not infer motive from position.

Dario Amodei: from mandatory measurement to binding controls

2023: measure the danger during training and before release

Amodei's earliest source in this record is his July 2023 written testimony to the U.S. Senate. His core proposal was mandatory testing and auditing for advanced systems. The testing would occur while a model was being trained and again before deployment. He treated the AI supply chain as a possible regulatory surface and argued that legislation should establish standards rather than leave every laboratory to choose its own. He also acknowledged the limiting case: if measured risks became severe enough, a slowdown could be necessary. He pointed to the National Institute of Standards and Technology as a possible technical center for the work. D01

That November, in prepared remarks at the United Kingdom's AI Safety Summit, Amodei explained Anthropic's Responsible Scaling Policy as an "if-then" structure. If a model showed specified dangerous capabilities, the company would not deploy it or train a more powerful successor until stronger safeguards were in place. The important regulatory point came at the end: private scaling policies were "not intended as a substitute for regulation, but rather a prototype for it." Governments, he argued, should turn lessons from company policies into testing and auditing regimes with accountability and oversight. D02

The continuity is already visible. Regulation begins with a measured capability, not a general label called AI. The laboratory can prototype the test, but public authority must eventually make the obligation common.

2025: add transparency, independent measurement, and economic monitoring

At the Paris AI Action Summit in February 2025, Amodei widened the frame. He called for government-enforced transparency about AI practices, third-party measurement of model behavior, and monitoring of economic effects. He also connected AI governance to the distribution of benefits and to a democratic supply chain for advanced technology. D03

This was not a retreat from safety testing. It added 2 layers. First, a regulator needs facts that do not depend entirely on a developer's own description. Second, the public consequences of advanced AI include labor, economic concentration, and access as well as catastrophic misuse.

Early 2026: start with rules that are simple, narrow, and evidential

In *The Adolescence of Technology*, Amodei argued that regulation should be "surgical" and simple enough to avoid crushing beneficial development. Transparency was his preferred first move when evidence remained incomplete. Stronger evidence would justify stronger rules. He also supported exemptions or lighter treatment for small firms that were not operating at the frontier. D04

The essay matters because it states the limiting principle around his earlier proposals. Capability-triggered regulation should not become a pretext for undifferentiated control. Uncertainty calls for measurement and disclosure before prohibition.

Mid-2026: the evidence now warrants binding regulation

In June 2026, Amodei said explicitly that the evidence had changed. In *Policy on the AI Exponential*, he argued that frontier risks now warranted movement beyond transparency to "serious binding regulation." His proposed regime included mandatory third-party testing for cyber, biological, loss-of-control, and automated research-and-development capabilities; incident reporting; stronger security; and government power to block or reverse an unsafe deployment. He also addressed labor displacement and the tax base as policy problems that could require economic adaptation. D05

This is the clearest development in his record. The underlying model remained stable: evidence should trigger obligation. His judgment about the evidence moved. What had been a case for transparency-first rules in January became a case for binding controls by June.

July 2026: regulate dangerous capability, not the open-weights label

Amodei's July statement on open-weights models sharpened the target. He rejected a categorical ban and described non-dangerous open models as a public good. Open weights could increase some misuse risks because a released model cannot be recalled, but that did not make a ban on business use an effective answer. His alternative was mandatory testing for any sufficiently capable model, open or closed, along with controls on advanced chips and large-scale model extraction. D06

He also described Hassabis's newly published standards proposal as close to an emerging consensus. That is one of the few explicit cross-references in this corpus. Amodei did not endorse every institutional detail. He did identify a shared center: common tests, common thresholds, and obligations attached to dangerous capability.

The through line. Amodei has consistently treated measurement as the gateway to regulation. The major change is his assessment of when evidence became strong enough for binding rules. His stated boundary also remained consistent: regulate frontier risk without making ordinary models, small firms, or open development illegal by category.

Sam Altman: exceptional regulation for exceptional capability

2015: less regulation in general, more for advanced AI

Altman's first 2 sources establish the contrast that continues through his later record. In January 2015, he argued broadly for fewer constraints on growth and innovation, but carved out AI for "much more regulation." S01 In March, writing about machine intelligence, he predicted that governments would regulate systems approaching superhuman machine intelligence. He proposed a new agency able to observe capability development, require protections such as air gaps and human authorization, support code review, and fund safety work. The agency should impose as little drag as possible while retaining authority over the exceptional risk. S02

That position predates ChatGPT and the 2023 regulatory debate by 8 years. It joins a generally deregulatory economic view to a special case for high-capability AI.

May 2023: a threshold licensing agency, but not a fixed pause

Before the Senate Judiciary Subcommittee in May 2023, Altman advocated a federal agency that would license systems above a capability or resource threshold. The agency would set safety standards, inspect compliance, require independent audits, and revoke a license when necessary. He repeatedly limited the proposal to the most powerful models. Smaller companies, open-source work below the threshold, and ordinary deployments should not carry the same burden. S03

The hearing also produced the clearest direct disagreement in this paper. Asked about pausing development, Altman said that a pause could be appropriate if tests showed a dangerous capability. He rejected a fixed "calendar-clock pause" as the controlling idea. The right trigger, in his view, was a failed safety evaluation, not the passage of 6 months. S03

Six days later, Altman, Greg Brockman, and Ilya Sutskever published *Governance of Superintelligence*. They proposed coordination among leading developers, an international authority resembling the International Atomic Energy Agency, and a threshold above which the authority could inspect systems, audit safety, test compliance, and restrict rates of capability growth. Systems below that frontier threshold should remain outside the regime. They also argued that stopping all superintelligence work would require a global surveillance system that might be difficult or undesirable to maintain. S04

The domestic licensing agency and the international authority performed the same conceptual job at different scales: identify an exceptional class of system, measure it, and place enforceable obligations above the line.

June 2023: formalize testing, adaptable standards, and distributed liability

Altman's written Senate testimony called for internal and external testing, threshold-based licensing or registration, adaptable standards developed with multiple institutions, and international cooperation. S05 His answers to questions for the record added 2 qualifications. Small firms should receive exemptions or scaled obligations, and liability should follow actual control over the system rather than fall automatically on one participant in the chain. S06

These documents make his scope condition more precise. "AI regulation" does not mean one license for every model or one defendant for every failure. It means a special regime at the frontier, with duties assigned according to capability and control.

2025–2026: connect adoption to public choice and economic resilience

In a July 2025 Federal Reserve conversation, Altman discussed regulation mainly in sectors that already control high-consequence decisions. He argued that officials should weigh the danger of AI adoption against the danger of refusing useful systems. This was not a comprehensive regulatory plan. It placed risk management inside a broader concern for diffusion and economic benefit. S07

In an April 2026 OpenAI forum, Altman argued that public debate should begin before technology locks in social choices. He emphasized broad access to compute, society-wide resilience, and the possibility that tax and transition policy would need to change if AI shifted income from labor toward capital. The event transcript is a first-party publication, but some speaker labels are generic; this paper uses only passages the moderator identifies as Altman's remarks. S08

The through line. Altman's record is stable on the main architecture: strong controls for extraordinary capability, a light burden below the threshold, and both domestic and international institutions. His later remarks broaden the object of governance from model safety to adoption, access, and economic transition. The sharpest boundary is procedural: he permits a pause when evidence demands one, but he does not treat a fixed calendar pause as the default instrument.

Elon Musk: from regulator as gatekeeper to government as backstop

2017: regulate before the public is harmed

Musk offered the earliest extended public argument in this record during a July 2017 conversation with U.S. governors. Ordinary regulation, he said, often follows visible harm. He argued that this sequence was too late for advanced AI. A regulator should first gain insight into what developers were doing, then establish rules, and, if necessary, pause work across the field until safety was demonstrated. He described the government as acting for the public good because no individual company could manage a collective risk alone. He also warned in the same session that overregulation could impede useful activity. E01

The position combined urgency with a familiar Musk distinction. He opposed regulation as accumulated friction in many industries, yet treated advanced AI as an exception because the damage could become irreversible before the normal feedback loop worked.

2018: public oversight for digital superintelligence

At South by Southwest in March 2018, Musk repeated the case for a public body with insight and oversight over digital superintelligence. He separated narrow AI from systems capable of civilization-level consequences. The public authority needed technical visibility before it could govern well. E02

This distinction anticipated the threshold logic the other 3 speakers later used. It did not yet specify a licensing test or international body.

March 2023: sign a fixed pause and a government moratorium

Musk signed the Future of Life Institute's open letter calling on AI laboratories to pause for at least 6 months the training of systems more powerful than GPT-4. The letter asked governments to institute a moratorium if laboratories would not pause voluntarily. It also called for regulators, independent auditors, liability rules, and technical methods to distinguish real content from synthetic output. E03

Because the letter had many signatories, it proves Musk's endorsement of the text, not his authorship of every sentence. Even with that limitation, its calendar requirement was explicit. It is the strongest stopping rule in the 4 records at that date.

November 2023: a referee with independent insight

In a full conversation with U.K. Prime Minister Rishi Sunak after the first AI Safety Summit, Musk described government as a referee. A referee does not play the game, but protects public safety and intervenes when the players create unacceptable risk. He favored independent testing and technical insight. He also argued for alignment among the United States, United Kingdom, and China, and said that even a regulator initially limited to observing and warning the public could add value before it acquired stronger enforcement powers. E04

The image is less absolute than the 2017 gatekeeper or the March 2023 moratorium. Public authority still belongs on the field. Its first powers can be knowledge, testing, and warning rather than a standing stop order.

July 2026: peer review first, government intervention after refusal

In a full interview with *The Economist* in July 2026, Musk said that there was no real stop button for AI development and that society might not choose to press one because the potential benefits were large. He continued to describe the catastrophic risk as nonzero. His immediate proposal was for leading laboratories to test one another's models before release, using a brief review period. If a company refused to respond to a serious finding, the government should serve as the backstop. He cautioned that government often lacked the technical knowledge to be the first evaluator. E05

Musk addressed Hassabis's proposal in the same exchange. He said they had discussed it, regarded it as a useful starting point, and favored an architecture that began with technically capable peers. That creates a direct institutional difference without erasing their common support for pre-release review.

The through line. Musk has remained unusually explicit that advanced AI presents a public-safety problem and that government has a legitimate role. The instrument changed. His record moves from preemptive regulator and possible industry-wide halt, through a fixed moratorium, to a narrower peer-review mechanism with public enforcement in reserve. The later position moderates the stop power; it does not withdraw the claim that the risk is real.

Demis Hassabis: regulate in stages, then make the test mandatory

2023: build benchmarks that can support government approval

In a September 2023 TIME question-and-answer interview, Hassabis argued for rigorous benchmarks that could reveal dangerous model capabilities. Those measurements could support a government approval process: a system that failed should not be released. He preferred capability triggers to simple compute thresholds because the amount of computing power used to train a model was only a proxy for what the model could do. The published interview was condensed and edited, a limitation TIME states on the page.

H01

The sequence matters. First, build a credible test. Then give the government a release decision based on the result.

2024: update sector rules now and defer broad frontier rules until evidence improves

In an official Google DeepMind podcast in August 2024, Hassabis proposed 2 regulatory tracks. Governments should update existing rules in regulated sectors such as health and transportation immediately. For frontier systems, safety institutes should test capabilities and watch the evidence. Broader frontier-specific rules should follow when the risks and benchmarks became clearer. He warned that a detailed AI law written 5 years earlier would have targeted a different technology. The rules therefore had to remain light, nimble, and adaptable. H02

He also called for international cooperation on guardrails and deployment norms. Because AI moves across borders as software, a national rule could not fully contain a frontier risk.

June 2026: self-regulation cannot solve the competitive dilemma

In a full Stanford conversation in June 2026, Hassabis argued that self-regulation was insufficient. Commercial and geopolitical competition created a prisoner's dilemma: even a safety-conscious laboratory could fear losing ground if it slowed alone. Government therefore had to participate. At the same time, he repeated that regulation should be light, dynamic, technically informed, and able to change with the systems. H03

This explains why his staged approach did not end in permanent voluntary control. Measurement buys time and knowledge. It does not remove the collective-action problem.

July 2026: create a standards body and require pre-release approval

In *A Framework for Frontier AI and the Dawning of a New Age*, Hassabis proposed his most complete structure. A federally overseen, industry-funded public-private standards body would develop tests for frontier models. Developers would submit systems for a 30-day review. The process could begin voluntarily but would become mandatory, and a frontier model would need to pass before release. Independent evaluators would use held-out tests, repeated at least quarterly. Open- and closed-weight systems would face the same rule when they met the frontier threshold; non-frontier systems would remain exempt. International standards would reduce regulatory arbitrage, and a coordinated slowdown could become necessary if tests showed unacceptable danger. H04

This architecture connects every stage of his record: capability rather than compute as the trigger, a formal test, government authority, independent review, frontier scope, and international coordination.

The through line. Hassabis did not move from opposition to support. He moved from a principle to a sequence. Existing sector regulators act first. Safety institutes build evidence. A dedicated frontier regime becomes binding once the tests and capabilities justify it. By July 2026, he believed the process was ready to be specified.

Where they explicitly agree

Government has a legitimate frontier role

All 4 expressly assign public authority a role. Amodei asks governments to make testing and auditing accountable. Altman proposes domestic and international licensing bodies. Musk calls government the public's referee and enforcement backstop. Hassabis says self-regulation cannot resolve the competitive dilemma and places federal oversight above his standards body. D02 S03 E04 H03

This is the broadest documented overlap. It does not tell us how intervention begins or how much power the authority should have.

The trigger should be capability or risk, not the word “AI”

Each speaker draws a line around the frontier. Amodei uses dangerous-capability thresholds and exempts ordinary open models. Altman repeatedly excludes systems below a capability or resource threshold. Musk distinguishes digital superintelligence from narrow AI. Hassabis prefers observed capability to compute alone and exempts non-frontier models. D06 S04 E02 H01

Their thresholds are not numerically identical. The common principle is proportionality: the strongest rules attach to systems capable of the strongest harm.

Testing should occur before a frontier model is released

Amodei calls for mandatory third-party testing before deployment. Altman supports internal and external evaluations tied to licensing. Musk's 2026 proposal asks peer laboratories to review a model before release. Hassabis requires independent review and a passing result. D05 S05 E05 H04

The test is the hinge in all 4 systems. The disagreement lies in who owns it and what happens after a failure.

International coordination is necessary

Amodei asks governments to build compatible regimes and addresses models trained outside democratic control. Altman proposes an international authority for superintelligence. Musk calls for alignment that includes the United States, United Kingdom, and China. Hassabis proposes international standards and leaves room for a coordinated slowdown. D02 S04 E04 H04

None supplies a complete answer to enforcement against a nonparticipating state. The overlap is an acknowledgment that frontier capability and model distribution cross national borders.

Rules should adapt without sweeping in low-risk development

All 4 express some concern about blunt control. Amodei favors surgical rules and small-firm exemptions. Altman protects systems and open work below the threshold. Musk warns against the accumulated cost of regulation and later favors technically informed peer review. Hassabis calls for light, fleet-footed rules that follow changing evidence. D04 S06 E01 H02

This is not an agreement on how light the rules should be. It is a shared rejection of one static burden applied to every AI system.

Open versus closed is secondary to dangerous capability

The records converge most clearly by 2026. Amodei rejects a categorical open-weights ban. Altman exempts open-source work below the frontier threshold. Hassabis applies the same review rule to open and closed models once either is frontier-capable. Musk has favored open development in parts of his broader record, but the sources reviewed here do not contain a comparably complete open-weights regulatory test. D06

S04 H04

This is therefore a 3-way explicit overlap and a 4th record that does not contradict it. It should not be upgraded into unanimous agreement.

Where they explicitly disagree

A fixed pause versus an evidence-triggered pause

The March 2023 letter signed by Musk called for a public, verifiable pause of at least 6 months on training systems more powerful than GPT-4, backed by a government moratorium if voluntary action failed. E03

In May, Altman told the Senate that the right trigger was evidence from a safety test, not a “calendar-clock pause.” He allowed that a pause could become necessary if a system crossed a dangerous threshold. S03

This is a direct conflict over the trigger, not over whether development may ever need to stop. Musk endorsed time on the calendar; Altman endorsed a failed evaluation.

An independent standards body versus peer laboratories as the first evaluator

Hassabis's July 2026 framework begins with an independent public-private body. A 30-day review would become mandatory, and a frontier system would need to pass before release. H04

Musk addressed that proposal days later. He preferred leading laboratories to test one another first because they held the relevant technical knowledge. Government would intervene if a developer refused to correct or contain a serious danger. E05

Both require review before release. They disagree about the first institution in the chain and the point at which public compulsion enters.

The regulator's stop power became a disagreement within Musk's own record

Musk's 2017 formulation allowed a regulator to halt the whole field until safety was established. His 2023 signature endorsed a fixed moratorium. In 2026, he said there was no practical stop button and proposed a short peer process with government in reserve. E01 E03 E05

This change does not resolve into a simple contradiction because the contexts and proposed durations differ. It is still a material shift in the role he gives the state: from gatekeeper, to moratorium authority, to backstop.

Transparency first became binding regulation in Amodei's own 2026 record

In January 2026, Amodei presented transparency as the sensible first response when the evidence remained uncertain. By June, he wrote that the risks now justified serious binding rules and government power to prevent or reverse an unsafe deployment. D04 D05

He explains the change as an evidentiary update. The regulatory principle stayed the same; his location on its decision rule moved.

Differences that the record does not establish as disagreements

The speakers use different institutional models. Amodei invokes NIST, testing mandates, and direct public power. Altman proposes a licensing agency and an IAEA-like international authority. Musk uses the referee metaphor and later peer review. Hassabis proposes a federally overseen standards organization. These structures could compete, but they could also occupy different layers of one system. The sources do not prove a direct conflict among all 4.

Their risk emphasis also differs. Amodei gives unusual weight to biological, cyber, alignment, national-security, and labor risks. Altman emphasizes superintelligence, broad access, and economic transition. Musk foregrounds civilization-level danger, public safety, and the technical limits of government. Hassabis stresses rigorous benchmarks, dangerous capabilities, and the competitive dynamics that undermine voluntary restraint. Emphasis is not negation.

The same caution applies to economic policy. Amodei and Altman discuss taxes, labor, or benefit distribution in the later record. Musk and Hassabis say less about redistribution in the sources selected here. Their silence does not amount to rejection.

What changed, and what did not

Across 11 years, the center of this record moved from general warnings toward operational machinery. Early sources ask governments to acquire technical insight, create a specialist agency, or regulate before catastrophe. Later sources specify held-out evaluations, review periods, model thresholds, incident reports, licensing, enforcement, and exemptions.

The people changed too, but not in one direction. Amodei became more willing to bind. Musk became less willing to rely on a broad stop. Hassabis advanced from staged observation to a mandatory process. Altman's frontier-agency architecture remained comparatively stable while his later discussion widened toward access and economic adaptation.

What did not change is the boundary each man draws. None of the 4 records supports treating every use of AI as a civilization-level risk. None supports leaving the most capable systems solely to ordinary product law and private discretion. Between those boundaries lies the actual debate.

Its unresolved questions are practical. Can an evaluation detect a dangerous capability before deployment? Who keeps the test secret enough to prevent gaming and open enough to earn legitimacy? Can a regulator recruit expertise without becoming dependent on the firms it governs? What happens when one laboratory, or one country, refuses the result?

The 4 speakers offer parts of an answer. They do not offer a settled constitution for artificial intelligence. Their clearest common judgment is simpler: when private systems can create public harm at frontier scale, private assurance is not enough. The hard part is deciding who may say **stop**, on what evidence, and for how long.

Primary-source index

Dario Amodei

- **D01 · July 25, 2023.** Written Testimony: Oversight of A.I.—Principles for Regulation. U.S. Senate Committee on the Judiciary.
- **D02 · November 1, 2023.** Prepared remarks at the AI Safety Summit. Anthropic.
- **D03 · February 11, 2025.** Statement on the Paris AI Action Summit. Anthropic.
- **D04 · January 2026.** The Adolescence of Technology. Authored essay.
- **D05 · June 2026.** Policy on the AI Exponential. Authored essay.
- **D06 · July 27, 2026.** Our position on open-weights models. Anthropic; identified as a post by Dario Amodei.

Sam Altman

- **S01 · January 14, 2015.** Policy for Growth and Innovation. Authored blog post.
- **S02 · March 2, 2015.** Machine Intelligence, Part 2. Authored blog post.
- **S03 · May 16, 2023.** Oversight of A.I.—Rules for Artificial Intelligence. Official U.S. Senate hearing transcript.
- **S04 · May 22, 2023.** Governance of Superintelligence. Coauthored with Greg Brockman and Ilya Sutskever.
- **S05 · June 22, 2023.** Written testimony before the U.S. Senate. OpenAI.
- **S06 · June 22, 2023.** Senate Questions for the Record. OpenAI.
- **S07 · July 22, 2025.** Federal Reserve fireside chat. Official transcript.
- **S08 · April 6, 2026.** Sam Altman on Building the Future of AI. OpenAI Forum event replay and transcript.

Elon Musk

- **E01 · July 15, 2017.** Conversation with Elon Musk at the National Governors Association. Full session; official NGA archive page and complete recording.
- **E02 · March 11, 2018.** Elon Musk Answers Your Questions—SXSW 2018. Full SXSW session recording.

- **E03 · March 22, 2023.** Pause Giant AI Experiments: An Open Letter. Future of Life Institute; Musk is a signatory.
- **E04 · November 2, 2023.** Conversation with Prime Minister Rishi Sunak. Full 10 Downing Street conversation.
- **E05 · July 29, 2026.** The full-length interview with Elon Musk. *The Economist* full interview recording.

Demis Hassabis

- **H01 · September 7, 2023.** TIME100 AI Q&A. Direct interview; condensed and edited by TIME.
- **H02 · August 14, 2024.** Unreasonably Effective AI. Google DeepMind Podcast, full episode.
- **H03 · June 2, 2026.** A Conversation with Demis Hassabis. Stanford Graduate School of Business, full conversation.
- **H04 · July 14, 2026.** A Framework for Frontier AI and the Dawning of a New Age. Authored essay.

Source and interpretation limits

This corpus is broad, not exhaustive. It prioritizes sources in which regulation is discussed at enough length to reveal the proposed mechanism and its limits. Public remarks can omit private beliefs, later changes, or details expressed elsewhere. A signature establishes endorsement of a letter, not authorship. An interview transcript can contain transcription errors; the analysis therefore relies on the full recording and avoids making a disputed word carry a conclusion.

The research cutoff is September 13, 2026. Later statements may change the comparison.

Research and editorial production: RVA Cyber **Citation format:** Speaker initial plus chronological item number. **Editorial standard:** Exact. Humane. Consequential.