



## OFFICE OF THE GOVERNOR

**SEP 29 2024**

To the Members of the California State Senate:

I am returning Senate Bill 1047 without my signature.

This bill would require developers of large artificial intelligence (AI) models, and those providing the computing power to train such models, to put certain safeguards and policies in place to prevent catastrophic harm. The bill would also establish the Board of Frontier Models – a state entity – to oversee the development of these models.

California is home to 32 of the world's 50 leading AI companies, pioneers in one of the most significant technological advances in modern history. We lead in this space because of our research and education institutions, our diverse and motivated workforce, and our free-spirited cultivation of intellectual freedom. As stewards and innovators of the future, I take seriously the responsibility to regulate this industry.

This year, the Legislature sent me several thoughtful proposals to regulate AI companies in response to current, rapidly evolving risks – including threats to our democratic process, the spread of misinformation and deepfakes, risks to online privacy, threats to critical infrastructure, and disruptions in the workforce. These bills, and actions by my Administration, are guided by principles of accountability, fairness, and transparency of AI systems and deployment of AI technology in California.



SB 1047 magnified the conversation about threats that could emerge from the deployment of AI. Key to the debate is whether the threshold for regulation should be based on the cost and number of computations needed to develop an AI model, or whether we should evaluate the system's actual risks regardless of these factors. This global discussion is occurring as the capabilities of AI continue to scale at an impressive pace. At the same time, the strategies and solutions for addressing the risk of catastrophic harm are rapidly evolving.

By focusing only on the most expensive and large-scale models, SB 1047 establishes a regulatory framework that could give the public a false sense of security about controlling this fast-moving technology. Smaller, specialized models may emerge as equally or even more dangerous than the models targeted by SB 1047 – at the potential expense of curtailing the very innovation that fuels advancement in favor of the public good.

Adaptability is critical as we race to regulate a technology still in its infancy. This will require a delicate balance. While well-intentioned, SB 1047 does not take into account whether an AI system is deployed in high-risk environments, involves critical decision-making or the use of sensitive data. Instead, the bill applies stringent standards to even the most basic functions — so long as a large system deploys it. I do not believe this is the best approach to protecting the public from real threats posed by the technology.

Let me be clear – I agree with the author – we cannot afford to wait for a major catastrophe to occur before taking action to protect the public. California will not abandon its responsibility. Safety protocols must be adopted. Proactive guardrails should be implemented, and severe consequences for bad actors must be clear and enforceable. I do not agree, however, that to keep the public safe, we must settle for a solution that is not informed by an empirical trajectory analysis of AI systems and capabilities. Ultimately, any framework for effectively regulating AI needs to keep pace with the technology itself.

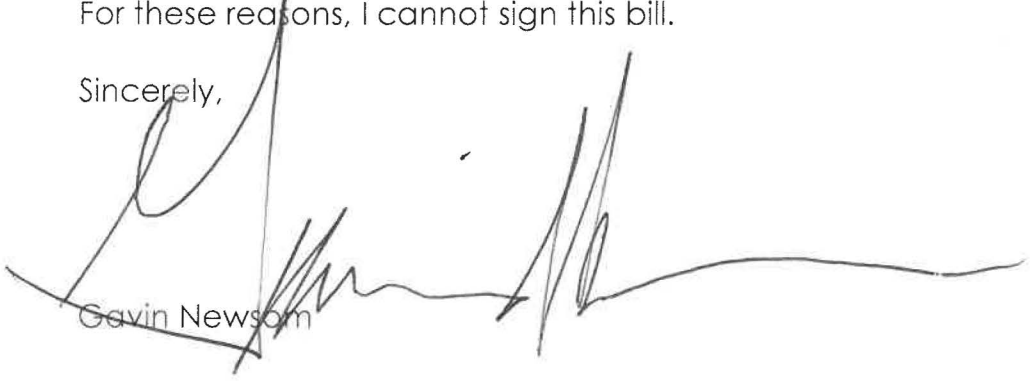
To those who say there's no problem here to solve, or that California does not have a role in regulating potential national security implications of this technology, I disagree. A California-only approach may well be warranted – especially absent federal action by Congress – but it must be based on empirical evidence and science. The U.S. AI Safety Institute, under the National Institute of Science and Technology, is developing guidance on national

security risks, informed by evidence-based approaches, to guard against demonstrable risks to public safety. Under an Executive Order I issued in September 2023, agencies within my Administration are performing risk analyses of the potential threats and vulnerabilities to California's critical infrastructure using AI. These are just a few examples of the many endeavors underway, led by experts, to inform policymakers on AI risk management practices that are rooted in science and fact. And endeavors like these have led to the introduction of over a dozen bills regulating specific, known risks posed by AI, that I have signed in the last 30 days.

I am committed to working with the Legislature, federal partners, technology experts, ethicists, and academia, to find the appropriate path forward, including legislation and regulation. Given the stakes – protecting against actual threats without unnecessarily thwarting the promise of this technology to advance the public good – we must get this right.

For these reasons, I cannot sign this bill.

Sincerely,



Gavin Newsom