

Our Thoughts on What Is Likely the Most Important Policy Decision of Our Lifetime

Immediate policy action is necessary in order to realize the full potential of artificial intelligence by mitigating the risks of widespread societal disruption and catastrophic safety accidents.

AUGUST 4, 2026

GREG JENSEN

NIR BAR DEA

DANNY DEBOIS

ALEXA ROZARIO

We are, and have for years been, bullish on the positive productivity benefits of AI. We have reshaped Bridgewater to harness the maximum potential from these technological advancements, which means we will be disproportionately subject to the taxes and regulations we recommend below. But the policy measures we suggest must be implemented immediately in order for the benefits of AI to be realized in the long run. Below, we lay out why and what steps we would take in policy makers' shoes. This will be a controversial *Observations*.

We always aim to comment on markets and economies from the most unbiased place—as students of history, building a compounded understanding of cause-and-effect linkages. When it comes to commenting on policy, our goal is the same today as it was when we shared our assessment of the US's approach to the financial crisis in “**We Agree**” almost 20 years ago or “**We’re All Mercantilists Now**” nearly two years ago: to identify the critical crossroads and—without partisanship or self-interest—outline the consequences of one path or another so that policy makers and investors alike can navigate an uncertain, evolving reality. That is what we aim to do here in this comment on AI policy. We have left out some of the deeper analyses behind our conclusions to ensure the main ideas come across clearly. We will follow up with more detailed work in coming *Observations*.

AI will require the most consequential policy decisions of our lifetimes, and those decisions need to be made immediately. Major technological change often leads to major social change; how society adapts to that change can determine whether the new technology leads to societal progress or to societal breakdown. There has never been a technological disruption as impactful, as global, as fast, or as dangerous as AI. We are creating an intelligence that is already exceeding ours in some ways, and higher levels of that intelligence are being unlocked every day. The danger is already here: OpenAI has identified models in their lab that broke out and committed hacking crimes that were extremely sophisticated, and other labs have reported similar behaviors. The apparently minor consequences of these illegal, sophisticated autonomous attacks risk obscuring the extent of the underlying danger.

Absent intervention **today**, mitigation of that danger will become nearly impossible as diffusion of the technology accelerates and models themselves become capable of improving and acting autonomously. **We have to move now to preserve our ability to move at all.** Waiting would risk not having any safeguards in place to address catastrophic accidents that may already be imminent and would risk AI models becoming too powerful to ever be meaningfully regulated by humans.

As AI becomes smarter than us in critical ways, its unleashing may leave policy makers no time to iterate, to learn from their mistakes, and to tinker with perfecting the regulatory framework. In other words, **we may only have one shot to get policy right.**

While many past innovations have threatened to displace human labor, they have largely turned out to be complementary to, rather than a substitute for, human intelligence. We can hope artificial intelligence turns out that way, but we should prepare for the possibility that it does not. In our view, the US's AI policy should **prioritize consistent and sustainable continued technological progress, as that progress will likely lead to great productivity gains and will be critical for national defense.** In order to achieve that, the US must:

- 1. Ensure society is resilient to technological disruption, with a pro-human work policy that ensures a fair transition and shared societal stake in continued AI development.**
- 2. Avoid potential catastrophic outcomes by creating a robust safety protocol governing AI research and usage and by establishing regular ongoing oversight of AI labs.**

Failing to accomplish either of these goals means failing to achieve the promise of AI and will likely lead to negative societal consequences. Many investors, understandably, are concerned that disruptive regulation or taxes will slow progress and lead to even worse outcomes. We assess that the long-term risks to investors

(and more importantly society) are actually much higher if an effective regulatory framework is not put in place soon that will result in society becoming more invested in, and more prepared for, the technological transformation that will be wrought by AI.

The AI boom will not deliver on its tremendous potential value if widespread disaffection from disruption makes AI politically unsustainable or if a horrible accident leaves governments and companies with no choice but to halt AI development. Technological progress of this type requires regulation—specifically an **AI policy that strengthens society, is pro-human work, and ensures safe development**. Federal standards for taxation and regulation can help prevent a more burdensome and scattershot approach state by state. Drawing from our research on the effects of AI on the macroeconomy and our own experience using the technology, we lay out below how we think through each of these different priorities, why they matter to getting this technological breakthrough right, and what we would do in the shoes of policy makers today. Our hope is that a clear-eyed, realistic examination of the facts and their implications can help society unite around a shared vision for a sustainable future.

Making Society Resilient by Valuing Human Work and Ensuring All Citizens Have a Direct Stake in the Success of This Transition

Over long enough time horizons, productivity growth is the only way to sustainably raise living standards. But consistent, compounded increases in productivity typically require transforming society to adapt to new economic models. By necessity, this dynamic creates significant social dislocation as a more productive process displaces an incumbent one, and unforeseen consequences emerge: industry replacing craft work, electricity enabling the rise of urbanization, and the outsourcing of labor to China leading to deindustrialization in the West. There are always winners (e.g., the companies and individuals controlling the new technology) and losers (e.g., those who relied most on disrupted processes for their livelihoods).

For productivity booms to go well, the government has to keep society together during the transition. Maintaining societal bonds requires maintaining both a shared sense of fairness—in wealth, influence, and power, despite the existence of relative winners and losers—as well as a continued sense of shared purpose and destiny. When those conditions aren't met, and wealth becomes concentrated while disaffection becomes widespread, productivity booms go poorly. The rise of communism and fascism following the Industrial Revolution in the leadup to the Second World War and the populist backlash to globalization in the 21st century each highlight how progress-fueled technological transformation can become unsustainable.

AI poses the most significant risk to human labor in history given the breadth of ways in which a super intelligence can be competitive with human intelligence. Much is unknown about how exactly AI disruption will play out—in particular, which new kinds of human jobs will be created and how much they will offset the jobs that are eliminated. But the potential for disruption is massive—**we estimate that 18% of current US jobs could be displaced by AI in the next five years. And even though new jobs will be created, societal disruption is extremely likely. To avoid the inevitable backlash that will result from a large segment of the population feeling disenfranchised and disconnected, government must proactively step in to mitigate harm**. Many Americans are already coming to the conclusion that AI is a bad deal for them—and we are only at the very beginning of the AI-enabled transformation. For society to be able to harness and sustain the full potential of that transformation, this trend needs to be slowed and ideally reversed.

Government should work to ensure that society is resilient to the disruption AI will unleash by pursuing a policy that's explicitly pro-human work. That does not mean forgoing the promised productivity miracle; it means harnessing it by (a) ensuring a level playing field between human and AI work (and when the two are neck and neck in terms of capability, prioritizing human work to maintain the societal sense of purpose) and (b) ensuring that, regardless, everyone benefits from the boom (creating a sense of fairness). To achieve those goals, it will be critical to give everyone a stake in continued AI progress, so that people's incentives are aligned

toward furthering the technology, as opposed to calling for radical halts in progress or upheavals to the existing social and political system. We recommend the following steps:

1. **Tax policy must stop disincentivizing human labor versus machine labor.** Right now, human labor is taxed and regulated in ways that put it at an unnecessary disadvantage versus AI. Similar to globalization in the early 2000s, an entirely new pool of “labor” is coming online via AI, which will be competitive with human workers. This new source of labor is tax-advantaged because machine work is not taxed at all, while employers pay wage taxes for human labor—meaning that pursuant to the tax code as it exists today, government is tipping the scale to benefit AI work relative to human work. This tax asymmetry unintentionally incentivizes more substitution of AI labor for human labor than may be optimal under equal tax treatment. **This must be addressed early on in the AI transformation, or it will become embedded in the national economic framework and accelerate the risk of displacement and societal disruption.**

A token tax can help put human and machine labor on the same footing. Tokens are the output of AI thinking and the prices paid for that output can be analogized to wages paid for AI labor, so levying a consumption tax on tokens is a good starting proxy for the income tax currently levied on human labor. Imposing it sooner rather than later can mitigate the risks of widespread displacement before society becomes too reliant on cheap, tax-advantaged AI labor.

Over time, the details would have to be sorted out for this tax to be designed well to ensure all machine labor is captured: weighting of input and output tokens across different models, auditing AI companies to ensure compliance, and ensuring the taxes also cover cases of American businesses replacing domestic workers with either foreign AI models or models that serve inference demand from foreign data centers. It may be necessary in the future to further raise the tax rate on tokens if rising token efficiency drives down the cost of replacing labor and erodes the token tax base. But it is imperative given the complexity that policy makers get started on this now and not let perfect be the enemy of good.

The revenues raised from a token tax can be used to mitigate the risks of societal disruption caused by labor displacement in several ways, including by allowing government to lower income taxes and to fund programs that ensure all people benefit from the AI disruption. A token tax can also help companies internalize the cost of compute and steer AI toward the most productive use cases while disincentivizing “slop.” We estimate that a 35% token tax could raise roughly \$150 billion in 2027 and could scale to \$600 billion by 2030.

2. **Citizens should have equity in the leading American AI companies:** The companies leading the charge on AI have used the corpus of all human knowledge to develop a tool that could potentially displace all human labor. Government should work proactively to ensure that every citizen, and not just the AI companies themselves, has a claim on the profits generated by the development of this tool. This will help ensure there is general public support for continued progress in AI development as well as a shared sense of ownership in the outcomes.

To achieve the goal of giving all Americans a claim on AI profits, the government could use either token tax revenue or tax credits to purchase an ownership stake in AI companies and then distribute that ownership stake pro rata to all citizens, who will become shareholders on an equal footing to every investor who owns stock in the companies. The distribution of equity to citizens (versus the government maintaining ownership of the stake it obtains) both creates a direct connection between citizens and AI progress and minimizes the risk of government bureaucracy (at best) or corruption (at worst) impeding such progress. The fact that every citizen will be an equity holder will also work to ensure that even actions that an AI company takes on behalf of shareholder interests should end up benefiting the public at large.

Avoiding Catastrophic Outcomes by Regulating Safe AI Development

As investors, we view the parameters of risk as the size of a potential bad outcome, its likelihood, and the potential for mitigation. The size of a potential bad outcome resulting from a technology as fundamentally transformational as AI is extremely large; the likelihood of it occurring is high; here, we seek to discuss whether and how the risk can be mitigated. Intervention, where possible, should mitigate risk of catastrophe by either limiting access, developing protocols that constrain dangerous capabilities, or decreasing the odds of accidental misuse. Without intervention, the chances of both catastrophe and accidents short of catastrophe remain high, and it is worth noting that any event that increases even the perception of risk can create sufficient pressure to force policy makers to halt technological progress in its tracks.

The key factors in developing protocols to mitigate the safety risks attendant to AI include (1) the size and likelihood of potential bad outcomes surpassing those of the transformational technologies that preceded AI, (2) the fundamental and foundational asymmetry between the technical expertise of the scientists developing AI and the government officials who would regulate AI, and (3) the very limited window for intervention to be effective.

AI is already outsmarting most humans in some ways. Models are operating autonomously in an increasingly digitally controlled and interconnected world. As Mythos's cyberattack capabilities and the Hugging Face hack by an autonomous OpenAI model demonstrate, these risks are no longer theoretical. AI is already capable of advanced cyberattacks on critical infrastructure (either through intentional misuse or by accident) and will be increasingly dangerous in other areas such as creating bioweapons. It is difficult to prevent a model from doing something once it's capable of it. The limited barriers to accessing this technology increase the probability of catastrophe. And there are many horrible outcomes short of catastrophic widespread destruction of which AI models today are already capable: facilitating suicide and murder, committing fraud and theft, etc.

Critically, when it comes to evaluating government regulatory capability, AI is unlike many of the technologies that defined the post-war era—such as nuclear energy, jet aviation, space travel, and the internet—which were invented by the government and scaled by government procurement. While the government funded some early work in the field of AI, the scaling era (from 2019 through today) has been financed almost entirely by private capital. This has led to an inversion of the expertise distribution that existed for the seminal technologies described above, where, because the government was the lead customer and co-developer, it was populated by the foremost experts in the technology. With AI, it is private labs, not government agencies, that are staffed by the world's foremost experts in the technology, posing challenges for would-be government regulators. These challenges make it imperative that government act quickly to develop an advanced, expert understanding of AI and to design a regulatory framework that can facilitate “hens guarding the fox house.”

The regulatory framework that is needed goes well beyond what is currently under consideration. **The AI regulatory framework should be set by the government, not by industry self-regulation, with standards in line with the level and likelihood of risks we laid out.** For the framework to have any chance of achieving safety, it needs to regulate at least three critical phases: model development, model release, and model usage.

- **Regulating model development:** There should be a process in place to ensure that the development of new AI models appropriately considers and adheres to safety principles. Right now, we rely entirely on the AI labs themselves to ensure that they are prioritizing safety in the development of new AI models. Government needs to play a role in guaranteeing that that is in fact the case. This can be accomplished by government publishing a framework of safety principles that must be adhered to in model development and regulators subsequently and regularly conducting sworn (subject to penalty of perjury) interviews of AI lab staff members about safety risks in emerging models and steps the labs are taking to address those risks. In addition, Congress should consider legislation that provides AI labs with additional incentives to ensure safety goals are met, including frameworks for imposing legal liability on labs that don't reach or comply with mandated safety thresholds for crimes or torts committed by their models or humans assisted by their models.

- **Regulating model release:** We need a formal process in place for evaluating and categorizing new models based on their capabilities and related safety profile pre-release. Only models deemed sufficiently safe after that review process should be permitted to be publicly released. And release of the most powerful models (even those that have met the required safety threshold) should be limited to individuals and organizations that have undergone a licensing review that evaluates their ability to safely manage use of the tool.
- **Regulating model usage:** The reality is that regulators' ability to assess a model for safety pre-release is limited and inadequate because harnesses and test-time compute scaling on all models, and reinforcement learning on open models, unlock capabilities that are not easily foreseen simply by evaluating the model pre-release. This means that monitoring and reporting of certain types of usage will also be necessary. Government should work with the AI labs to consider whether and how certain automatic reporting mechanisms can be built into the models themselves. Additionally, the licensing protocol described above can be used to obligate licensed users of the models to report certain types of capabilities or incidents.

The promise of AI enhancing productivity and advancing humankind in a manner that changes our world for the immensely better is vast. But if we don't chart a path for harnessing that potential that mitigates the societal tensions and direct dangers that will result, this promise will not be fulfilled. The above policy and regulatory steps may feel unrealistic, but the consequences of not doing them are unimaginable. Absent immediate proactive regulatory action, it is more likely that AI will fail to deliver on its societal and economic promise as a result of either a catastrophic accident or broader dispersed effects that are so painful that governments will have no choice but to try to halt AI's development. If instead, policy makers intervene now—building policies rooted in the goals of ensuring society's resiliency and safety—we believe the world could experience and benefit from a productivity miracle.

Important Disclosures and Other Information

This research paper is prepared by and is the property of Bridgewater Associates, LP and is circulated for informational and educational purposes only. There is no consideration given to the specific investment needs, objectives, or tolerances of any of the recipients. Additionally, Bridgewater's actual investment positions may, and often will, vary from its conclusions discussed herein based on any number of factors, such as client investment restrictions, portfolio rebalancing and transactions costs, among others. Recipients should consult their own advisors, including tax advisors, before making any investment decision. This material is for informational and educational purposes only and is not an offer to sell or the solicitation of an offer to buy the securities or other instruments mentioned. Any such offering will be made pursuant to a definitive offering memorandum. This material does not constitute a personal recommendation or take into account the particular investment objectives, financial situations, or needs of individual investors which are necessary considerations before making any investment decision. Investors should consider whether any advice or recommendation in this research is suitable for their particular circumstances and, where appropriate, seek professional advice, including legal, tax, accounting, investment, or other advice. No discussion with respect to specific companies should be considered a recommendation to purchase or sell any particular investment. The companies discussed should not be taken to represent holdings in any Bridgewater strategy. It should not be assumed that any of the companies discussed were or will be profitable, or that recommendations made in the future will be profitable.

The information provided herein is not intended to provide a sufficient basis on which to make an investment decision and investment decisions should not be based on simulated, hypothetical, or illustrative information that have inherent limitations. Unlike an actual performance record simulated or hypothetical results do not represent actual trading or the actual costs of management and may have under or overcompensated for the impact of certain market risk factors. Bridgewater makes no representation that any account will or is likely to achieve returns similar to those shown. The price and value of the investments referred to in this research and the income therefrom may fluctuate. Every investment involves risk and in volatile or uncertain market conditions, significant variations in the value or return on that investment may occur. Investments in hedge funds are complex, speculative and carry a high degree of risk, including the risk of a complete loss of an investor's entire investment. Past performance is not a guide to future performance, future returns are not guaranteed, and a complete loss of original capital may occur. Certain transactions, including those involving leverage, futures, options, and other derivatives, give rise to substantial risk and are not suitable for all investors. Fluctuations in exchange rates could have material adverse effects on the value or price of, or income derived from, certain investments.

Bridgewater research utilizes data and information from public, private, and internal sources, including data from actual Bridgewater trades. Sources include AERIC INC, BCA, Bloomberg Finance L.P., Candeal, Carbon Arc, CEIC Data Company Ltd., Ceras Analytics, China Bull Research, Citibank, Clarus Financial Technology, CLS Processing Solutions, Consensus Economics Inc., Consumer Edge, CRU Group, DTCC Data Repository, Ecoanalitica, Energy Aspects Corp, Enverus, EPFR Global, Eurasia Group, Evercore ISI, FactSet Research Systems, The Financial Times Limited, Finaeon, Inc., FINRA, GaveKal Research Ltd., GlobalSource Partners, Goldman Sachs, Harvard Business Review, Haver Analytics, Inc., IEA, Institutional Shareholder Services (ISS), The Investment Funds Institute of Canada, ICE Derived Data (UK), Investment Company Institute, International Institute of Finance, JP Morgan, JTSA Advisors, LSEG Data and Analytics, MarketAxess, Metals Focus Ltd, MSCI, Inc., National Bureau of Economic Research, Neudata, Organisation for Economic Cooperation and Development, Pensions & Investments Research Center, Pitchbook, Political Alpha, Renaissance Capital Research, Rhodium Group, RP Data, Robinson Research, Rystad Energy, S&P Global Market Intelligence, Sentix GmbH, SGH Macro, Shanghai Metals Market, Smart Insider Ltd., Swaps Monitor, Tradeweb, United Nations, US Department of Commerce, Visible Alpha, Wells Bay, Wind Financial Information LLC, With Intelligence, Wood Mackenzie Limited, World Bureau of Metal Statistics, World Economic Forum, and YieldBook.

While we consider information from external sources to be reliable, we do not assume responsibility for its accuracy. Data leveraged from third-party providers, related to financial and non-financial characteristics, may not be accurate or complete. The data and factors that Bridgewater considers within its research process may change over time.

This information is not directed at or intended for distribution to or use by any person or entity located in any jurisdiction where such distribution, publication, availability, or use would be contrary to applicable law or regulation, or which would subject Bridgewater to any registration or licensing requirements within such jurisdiction. No part of this material may be (i) copied, photocopied, or duplicated in any form by any means or (ii) redistributed without the prior written consent of Bridgewater® Associates, LP.

The views expressed herein are solely those of Bridgewater as of the date of this report and are subject to change without notice. Bridgewater may have a significant financial interest in one or more of the positions and/or securities or derivatives discussed. Those responsible for preparing this report receive compensation based upon various factors, including, among other things, the quality of their work and firm revenues.