

OUR PRINCIPLES FOR A HUMAN-CENTERED AI FUTURE.

Five principles, the policy actions behind them, and the laws already pointing the way. September 2026.

■ INTRODUCTION

AI is the defining issue of our time. It is already reshaping how we work, how students learn, and how some of the most impactful decisions are made. The choices that we make now will affect our society for generations, and the stakes could not be larger. The potential upside that AI provides is extraordinarily large, while the potential downside is equally as steep. During this pivotal time, Forward Progress USA holds the crucial role of ensuring the youth have a say in the discussions that will affect our future.

AI is a nonpartisan issue. We are a nonpartisan organization. Workers, students and voters from across the political spectrum understand that the changes from Artificial Intelligence are coming, and they are coming fast. We are not affiliated with nor do we endorse any political party. We also do not explicitly endorse any vendor or corporation. Successful AI policy can and will cross ideological boundaries to ensure a better future for all.

We know that an AI-centric future is coming. Our goal is not to hinder the rapid developments seen from Artificial Intelligence, but rather prepare the world for the future. Society must invest in AI literacy and training to ensure that workers and students are prepared for what comes next. We must ensure safety and transparency in regards to development of frontier AI models, while also understanding the importance of innovation. We oppose heavy-handed regulations and overly bureaucratic compliance requirements. We understand that competition in this era is crucial, and hope to protect development while simultaneously guaranteeing transparency and safety from the largest firms in the AI sector.

This memo lays out the 5 core principles that guide our organization: People First, Safety by Design, Broad Prosperity, Freedom to Innovate and Public Readiness. Together, they outline the AI future that we think is possible, and the changes we must make to secure it for the next generation.

01 PEOPLE FIRST

■ WHAT IT MEANS

AI should be built for the sole purpose of augmenting and improving human potential, not as a replacement for human agency or creativity. Innovation in the AI sector should have the clear goal of improving the universal human experience. Artificial intelligence can not and will not serve as a replacement for humanity.

■ WHY IT MATTERS

When AI makes a decision without human oversight, real people bear the consequences. From denied loans, to rejected insurance claims, the negative effects of automation show up in our lives everyday. As AI rapidly embeds itself into our common systems, the risk is clear. Human agency and accountability may quietly disappear. Every AI related invention, progression and innovation should have a clear benefit to our society as a whole. The change is coming, but we still have time to mitigate the risks and ensure that the benefits of automation assist us rather than replace us.

■ POLICY ACTIONS

Require meaningful human review for high-stakes automated decisions (hiring rejections, benefits denials, medical diagnoses)

Require disclosure when a person is interacting with or being evaluated by an AI system

Preserve a clear right to appeal an AI-driven decision to a human

Require companies/political campaigns to disclose when AI is used to generate creative content (writing, art, music)

Protect workers rights to organize and bargain over the introduction of AI tools that could affect job security, workload or working conditions

■ EXAMPLES

Colorado AI Act (2026) → General state AI legislation that prioritizes human decision making in high-risk areas. It also preserves the consumer's right to appeal for human review.¹

AI Transparency in Elections Act (2026) → This legislation would mandate political campaigns to include a disclaimer over any political advertisement that utilized AI within 120 days of the election. Democracy is a fundamental right, and preserving honesty and integrity in voting matters is crucial.²

■ SOURCES

1. Colorado Senate Bill 26-189, Colorado General Assembly, 2026. <https://leg.colorado.gov/bills/sb26-189>

2. AI Transparency in Elections Act, S. 3875, 118th Congress. <https://www.congress.gov/bill/118th-congress/senate-bill/3875>

02 SAFETY BY DESIGN

■ WHAT IT MEANS

No AI system should be implemented in a high-risk environment without extensive prior research and testing. We believe that any scenario that involves human safety, health or livelihood must be taken seriously, and cannot be rushed to market. AI implementation into government services and military branches must be thoroughly researched and transparent.

■ WHY IT MATTERS

The impact of Artificial Intelligence in these high-risk situations could not be more prevalent. A survey of 272 AI experts revealed that they estimate a one-in-five chance that AI will gain dangerous weapon capabilities and cause mass harm by 2030.³ 63% of Americans think that AI could someday destroy mankind.⁴ All while more and more experts in the field critique the negligence for safety that the largest most powerful AI corporations demonstrate. The risks are real, and we must take them seriously before it's too late.

■ POLICY ACTIONS

Require pre-deployment safety testing and third-party audits for systems used in critical infrastructure.

Prohibit deployment of autonomous weapon systems without a human in the decision-making loop.

Establish a federal body to review and certify AI systems before their usage in high-risk government functions.

Mandate phased rollout processes for AI systems in public services.

Require incident reporting systems for both the government and corporations, so dangerous incidents or failures are tracked and investigated.

Prohibit the use of experimental AI models in active military operations or emergency response systems.

■ EXAMPLES

HALO Act (2026) → This act, proposed by Senator Adam Schiff, requires a human to hold the ultimate discretion over the use of force with autonomous and semi-autonomous weapons. It aims to create strict safeguards and human oversight requirements for Department of Defense systems.

The RAISE Act (2026) → This legislation, passed in New York State, requires developers of high-risk frontier AI models to create written safety protocols before deployment, go through independent third-party audits and report critical safety incidents. It includes heavy civil penalties for violations and creates an oversight office within the Department of Financial Services to monitor compliance.

■ SOURCES

3. Herb Scribner, "These 5 AI Risks Have the Highest Potential for Catastrophe," Axios, July 27, 2026.

<https://www.axios.com/2026/07/27/mit-study-ai-risks-urgent-weapons-cyber-attacks>

4. The Hill, poll on American attitudes toward AI and existential risk, 2026. <https://thehill.com/policy/technology/6092503-poll-ai-humanity-risks/>

03 BROAD PROSPERITY

■ WHAT IT MEANS

The productivity gains created by AI should benefit workers and communities, not just a small number of powerful companies and CEO's. One of the largest AI-related risks facing society today is mass unemployment. As Artificial Intelligence systems are implemented more and more into corporations, the risk to workers becomes even larger. We support a future where workers, employees and families are the primary beneficiaries of the AI-centric future.

■ WHY IT MATTERS

Right now, the power and wealth that comes with Artificial Intelligence is held in the hands of a very small number of corporations and people. 70% of the AI revenue in the United States comes from just two companies.⁵ Students, workers and unions are being left out of the conversations that decide their future everyday, and the risks for us could not be larger. If we do not act now, we risk locking ourselves into a future where a handful of companies control the economic destiny of millions of American workers.

■ POLICY ACTIONS

Require companies above a certain size to conduct and disclose workplace impact statements before large-scale deployment.

Support shorter standard work weeks or reduced hours at full pay for workers in industries that have become far more productive as a result of AI.

Encourage employee ownership models at companies undergoing AI-driven transformation so that workers also reap the benefits of increased productivity.

Fund economic development programs targeted at communities most vulnerable to AI-related job loss.

■ EXAMPLES

Alex Bores' AI Dividend Plan → This plan proposes a direct-payment program to employees, activated when AI significantly displaces American workers. It aims to serve as a way for the masses to reap the benefits of increased AI-era productivity, and would be funded by digital age levies such as a token tax and equity stakes in major AI companies.⁶

California SB 951 → Requires a 90-day advance notice for layoffs that hit 25 people or 25% of staff.⁷

New York LOADinG Act → Protects public employees across New York State from displacement due to Artificial Intelligence. Preserves civil service protections and collective bargaining rights.⁸

■ SOURCES

5. Investor's Business Daily, analysis of AI industry revenue concentration and market risk. <https://www.investors.com/news/technology/ai-bubble-ed-zitron-money-will-run-out-bear-case/>
6. "Exclusive: Alex Bores Rolls Out 'AI Dividend' Plan to Share AI Wealth," Axios, April 20, 2026. <https://www.axios.com/2026/04/20/alex-bores-ai-dividend-plan-wealth>
7. California Senate Bill 951, California State Legislature, 2025–2026 session. https://leginfo.ca.gov/faces/billNavClient.xhtml?bill_id=202520260SB951
8. New York Senate Bill S7543 (LOADinG Act), New York State Senate. <https://www.nysenate.gov/legislation/bills/2023/S7543/amendment/A>

04 FREEDOM TO INNOVATE

■ WHAT IT MEANS

We understand that at this point in time, it is absolutely crucial to put necessary guardrails and regulations on AI. However, we also believe that the freedom to innovate in this unprecedented time period is equally as important. Policymakers should encourage competition, rather than bogging down companies with inessential compliances and overregulated bureaucratic measures. Good AI policy should keep people and systems safe, while also encouraging firms to strive for breakthroughs that benefit everyone.

■ WHY IT MATTERS

Overregulation doesn't just slow down Big Tech companies, it also fractures and hurts the smaller firms and startups that promote competition and prevent monopolies in the market. In fact, most of the time Big Tech corporations have resources to combat and respond to regulations that smaller innovators don't. Freedom to innovate is not an excuse to bypass crucial safety measures, but rather a way to invite more leaders to the table, and break up the concentration of power that the largest AI companies hold today.

■ POLICY ACTIONS

Establish regulatory sandboxes to allow startups and researchers to test new AI applications in controlled environments before full compliance requirements apply.

Set a clear federal baseline for AI regulation to avoid fractured state laws and prevent innovators from having to follow 50 different sets of rules.

Increase funding for safe, secure open-source AI research to ensure innovation is open to all.

Scale compliance laws to company size and risk levels to encourage startup innovation and avoid slowing down safe implementation.

Streamline permitting and approving safe, harmless AI-infrastructure (Closed-loop cooling data centers, energy projects).

Expand high-skill immigration pathways for AI researchers and engineers to keep innovation and talent in the U.S.

■ EXAMPLES

TRAIGA (Texas Responsible Artificial Intelligence Governance Act) → While some aspects of the law are heavily debated (such as the reporting mechanisms for AI biases), TRAIGA's contains an extremely beneficial regulatory sandbox program which promotes innovation and avoids compliance concerns.⁹

■ SOURCES

9. Texas Responsible Artificial Intelligence Governance Act (TRAIGA), Office of the Texas Attorney General.
<https://www.texasattorneygeneral.gov/consumer-protection/file-consumer-complaint/consumer-ai-rights>

05 PUBLIC READINESS

■ WHAT IT MEANS

Society must invest in AI literacy, education and workforce training to guarantee a smooth transition towards an AI-centric world. This means giving students the skills to understand and critically evaluate AI systems, giving workers real pathways to retrain as their industries change and giving the government the expertise to actually understand the technologies that they are regulating.

■ WHY IT MATTERS

AI is a technology unlike anything we have ever seen before, and its pace of change is outrunning institutions' ability to prepare. It is absolutely crucial that starting now, we work towards investing in equitable programs and economic development initiatives that guide us all towards a better future. With this level of unprecedented change, the risk of certain communities falling behind increases. Public readiness guarantees that we all benefit and thrive from what comes next.

■ POLICY ACTIONS

Fund AI literacy standards and curricula in K-12 schools, covering how AI systems work, their limitations, and how to use them responsibly.

Require AI ethics and literacy training as part of a teacher's professional development.

Support Community College and vocational partnerships that offer accessible, low-cost AI skills training.

Prioritize public readiness funding for rural and lower-income communities with an increased risk of falling behind.

■ EXAMPLES

California AB 2876 (2026) → Signed by Governor Newsom, incorporates AI literacy and media literacy into K-12 curriculum frameworks and instructional materials.¹⁰

Hawaii SB 2212 (2026) → Requires a statewide AI literacy curriculum for 11th and 12th grade public school students, starting in the 2027–2028 school year.¹¹

■ SOURCES

10. California Assembly Bill 2876, LegiScan, 2026. <https://legiscan.com/CA/text/AB2876/id/2977688>

11. Hawaii Senate Bill 2212, Hawaii State Legislature, 33rd Legislature, 2026.
https://data.capitol.hawaii.gov/sessions/session2026/bills/SB2212_.HTM

■ CONCLUSION

Taken together, these principles represent what we believe to be the best path forward towards a positive AI future. But these principles don't represent policy. That work falls on us, the people, to act. The future of Artificial Intelligence is not yet predetermined. The way that we, as a society, react to this technology now will change how we live forever. That's why at this pivotal turning point, it is extremely valuable that we have leaders in government who are willing to tackle the problems caused by Artificial Intelligence head on. At this current moment, the most practical and valuable change we can make is through advocating and bolstering these leaders. We will continue to support and follow anyone who we believe to be strong fighters that align with the values outlined in this memo.