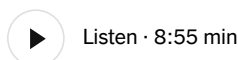


GUEST ESSAY

We Returned From China. We Realized Our Century's Biggest Challenge.

July 11, 2026

**By Eric Schmidt and Selina Xu**

Mr. Schmidt is the chief executive of Relativity Space and a former chief executive of Google. Ms. Xu is a China and technology analyst.

See more of our coverage in your search results.

[Add The New York Times on Google ↗](#)

Humanity is inching along a precarious tightrope.

Our world is in the midst of deciding how the artificial intelligence revolution will unfold and what limits should be drawn. Too much caution could waste A.I.'s promise of faster economic growth, greater scientific discovery and more prosperity. Too little caution could unleash labor-market chaos and social disorder. Balance matters. If we get the balance between control and growth wrong at any point, we'll fall into the gorge below. And no country is guaranteed to get to the opposite side.

This is our century's biggest survival challenge.

The Chinese tech executives we met on a trip last month were optimists about the technology, but they were much more conservative than Silicon Valley on how fast the growth of A.I. ought to be. Some were wary about replacing their employees too aggressively with A.I. Others worried that if A.I. became destabilizing, China's entire industry might be reined in by the government. One chief executive even said point-blank that growing more slowly was preferable to going full speed ahead, so as to avoid a repeat of the Luddite backlash toward the Industrial Revolution, when 19th-century English workers tried to halt the spread of mechanization by storming factories and destroying power looms.

Chinese policymakers are also visibly wrestling with how to encourage growth while preserving social and regime stability, which is the top priority for the Chinese Communist Party. In April, China stopped issuing new licenses for autonomous vehicles after dozens of robotaxis abruptly stranded passengers on the streets of Wuhan. That same month, a Chinese court ruled that companies cannot terminate employees just to replace them with A.I. systems. In June, a new

employment five-year plan pledged to prevent large-scale unemployment risks and to use A.I. to promote job creation. A ban on A.I. companions for minors in China is set to take effect in July.

Although China is no paragon, such actions have helped buoy the population's feelings toward A.I. Around 84 percent of respondents in China said they were excited about A.I., according to a recent Stanford report. Meanwhile, American skepticism toward the technology remains high. Communities across the United States are fighting the construction of data centers. Parents worry about children forming unhealthy attachments to A.I. companions. Workers fear being replaced. Policymakers warn of national security vulnerabilities. Researchers debate catastrophic risks. These fears are not irrational. People are asking a simple question: Will A.I. make my life better, or make it worse?

The question is hard to answer because the benefits of the technology are unevenly distributed and can often feel overhyped. One reason is what some have called "jagged intelligence," where A.I. fails at incredibly simple tasks even as it excels in specific areas. So when some in Silicon Valley predict an impending jobpocalypse, while others view the industry's claims as overblown, it's hardly surprising that the public mood around A.I. is darkening. In a Gallup survey from last year, 80 percent of American adults thought the government should regulate A.I., even if doing so means slower progress, a view that was shared by Democrats and Republicans alike. Even young Americans who historically embraced new technologies are angrier and more skeptical.

Sign up for the Opinion Today newsletter Get expert analysis of the news and a guide to the big ideas shaping the world every weekday morning. [Get it sent to your inbox.](#)

Dismissing A.I. entirely would be a mistake. The technology can help diagnose diseases, predict protein folding, improve farming, forecast disasters better, design new materials, accelerate scientific and drug discovery and power robots in dangerous environments to improve human safety (such as in space, firefighting and minefields).

Yet how technology spreads is never inevitable. If A.I. is viewed as benefiting the few at the expense of the majority, then the public will rage against the machine. And A.I. won't be able to make our lives better in the long run if it cannot survive in the short term. The real challenge, then, isn't whether the United States or China will build an overwhelming, insurmountable advantage over the other. It's whether either can figure out how to realize the benefits of A.I. without ripping apart its social fabric. Neither has found the answer yet.

Silicon Valley tech executives and policymakers across the country are waking up to that fact. States have introduced dozens of bills this year to put safety and privacy guardrails around A.I. The Trump administration issued a new executive order that seeks to give the government more oversight over new models before they're released to the public. Companies like OpenAI are starting to come around to the idea that strong safety rules can help reverse growing public opposition. But there is still no clear consensus on how the United States should move forward. We know big disruptions are coming, and they're coming fast as A.I. capabilities advance faster than policy response. We need bigger ideas to fix what may soon break.

To get to the other end of the tightrope, we need radical and incremental solutions alike. Here's one to start: a populist A.I. agenda that treats the technology as a public project. Just as NASA made space a national mission rather than a private one, the government should now do the same for A.I. to ensure that its benefits reach the public, not just the companies building it. Here are some concrete ways to do that:

One, treat some A.I. profits as a shared resource and distribute them directly to citizens. This can be in the form of a sovereign wealth fund, seeded by contributions from A.I. companies in the form of stock or cash. Prior models for this include Singapore's Temasek, seeded in 1974 with equity in state-linked companies, and Australia's Future Fund, which turned budget surpluses and government shares in a privatized telecom into a permanent endowment for future generations. Another format would be to reinvest A.I. profits into the younger generation, whose members risk being displaced before their careers even begin, and teach them how to use these tools. After all, the fruits of the age of A.I. are not the outcome of individual companies alone, but are built on the knowledge that society has accumulated over centuries.

Two, support public-interest A.I. models that the private sector might not be incentivized to build. Examples include models to help citizens navigate government services and benefits, to help them gain legal aid or to help educate their children. We strongly believe in the importance of open-source A.I., which anyone can freely use and modify. That would allow local governments, libraries, researchers, small businesses, nonprofits and independent developers to participate in building, owning, inspecting and adapting models for their community and privacy needs.

For the government and companies, this means expanding shared A.I. computing infrastructure, teaching the public how better to understand and use A.I. for its personal well-being, and funding genuinely open-source and safe American models. Companies working alone won't get us there: the enormous cost of computing power pushes them toward serving customers who can pay the most. Some technologies need to become public infrastructure through expanded access and price reductions. Britain once nationalized electricity; the United States regulated railroads. We've done it before and can do it again.

Three, regulate one area that attracts bipartisan agreement: the way children and teens interact with A.I. Several bipartisan bills have been introduced that propose safeguards for minors, parental controls and penalties for platform violations. Let's advance those. Absent regulations, industry is unlikely to police itself, as we've seen before in the age of social media.

Our society has navigated and absorbed major technological upheavals before, from the Industrial Revolution to the computer age. Those transitions endured politically because reformers and lawmakers built labor protections, social insurance, public schooling and antitrust law that helped broaden technology's gains. This time around, many Americans have no faith that the benefits of A.I. will be distributed. It is up to policymakers and the companies building A.I. to change that.

A.I. must help people thrive, not merely enrich a handful of companies and leave others behind. If that doesn't happen, the country will resist it, slow it and fight over it. That is how America loses its balance on the tightrope.

Eric Schmidt, a former chief executive and chairman of Google, is the chairman and chief executive of Relativity Space.

Selina Xu leads China and A.I. policy research in the office of Eric Schmidt.

The Times is committed to publishing a diversity of letters to the editor. We'd like to hear what you think about this or any of our articles. Here are some tips. And here's our email: letters@nytimes.com.

Follow the New York Times Opinion section on Facebook, Instagram, TikTok, Bluesky, WhatsApp and Threads.

A version of this article appears in print on , Section A, Page 19 of the New York edition with the headline: America Needs a Populist Vision for A.I.