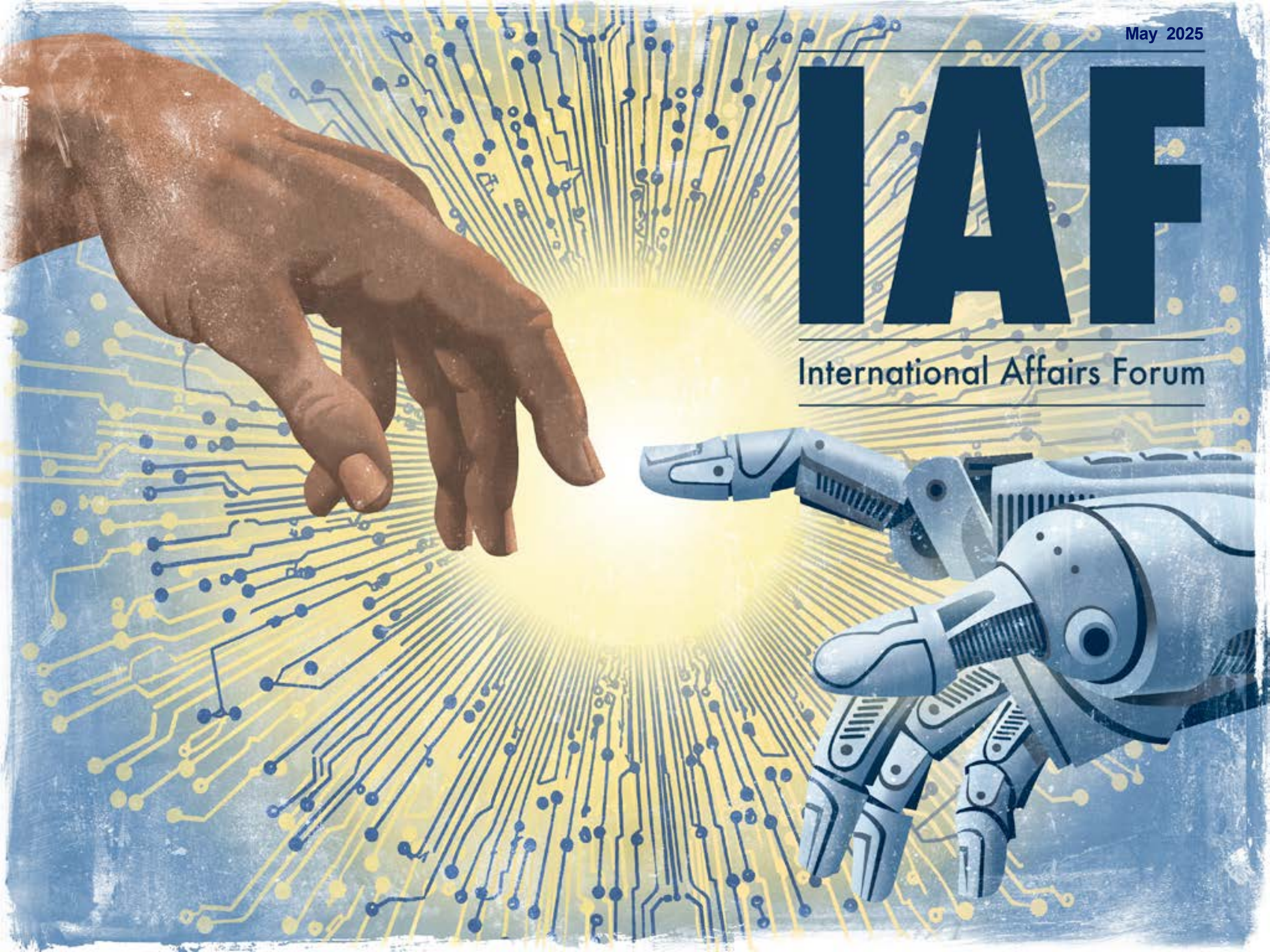


May 2025

IAAF

International Affairs Forum



International Affairs Forum

Printed in the U.S.A.
A publication of the Center for International Relations
1629 K St. #300, Washington DC 22201
202-821-1832 email: editor@ia-forum.org

Managing Editor Dr. Dale Mineshima-Lowe

Editors Joe Cole
Ramil Abbasov
Alena Gribanova
Nicole Pylawka
Natasha Rico

Editorial Advisor Dr. David Phillips

Cover Design Sam Ward
(<http://www.samwardart.com/>)

Submit your Editorial or Essay to
editor@ia-forum.org

www.ia-forum.org

contents

May 2026
Artificial Intelligence, Green Tech

Artificial Intelligence

- 4 Concerns and Challenges Surrounding AI
Interview with Professor Roman Yampolskiy

- 7 The Other AI Gap: Promises, Reality, and the Need for Digital Wisdom
Professor Richard Lachman

- 12 The Architecture of Adaptation: International Law and the Governance of Artificial Intelligence
Professor Jake Okechukwu Effoduh

- 18 Deployed Without a Net: AI, Corporate Risk, and the Emerging Public Order
Peter Piązzga

- 21 Governing and Managing Artificial Intelligence
Interview with Anjali Hansen

- 26 Best Practices for Responsible AI Development
Interview with Jenil Shah

- 29 Making Chile a Global Leader in Artificial Intelligence
Interview with Professor Alvaro Soto

- 32 What Does Artificial Intelligence Know about Food and Eating?
Professor Joachim Allgaier

- 36 Coded Imperialism: Generative AI, Epistemic Violence, and the Erasure of African Indigenous Knowledge
Dr. Nouridin Melo

- 38 Fertilizer, AI, and the New Politics of Market Access in Africa
Christopher Burke

-
-
- 41** The Invisible Infrastructure War:: AI Cloud Systems and Green Energy Grids in Africa as the Next U.S.-China Battleground
Shannon Roxborough
- 45** How Artificial Intelligence Will Redefine Public Finance Systems
Ramil Abbasov
- 48** Who Governs the Algorithm? The Epistemic Case Against Centralized AI Rulemaking
Vladyslav Manzyuk (Student Award Winner)
- 51** McDonaldization of AI-Driven Phillipine Public Administration
Geo Claire Desoyo (Student Award Winner)

Green Tech

- 54** The Green Mirage of Artificial Intelligence
Professor Dale Mineshima-Lowe
- 57** Managing the Green Transition
Interview with Professor Jan Rosenow
- 59** Green Tech as Strategic Insurance: What the EU Energy Crisis Revealed
Bogdan Romaniuk (Student Award Winner)
- 61** Green Tech Needs a Nuclear Backbone
Kashvi Sheoran (Student Award Winner)
- 64** References

Partners



The core values for the International Affairs Forum publication are:

- We aim to publish a range of op-ed pieces, interviews, and short essays, alongside longer research and discussion articles that make a significant contribution to debates and offer wider insights on topics within the field;
- We aim to publish content spanning the mainstream political spectrum and from around the world;
- We aim to provide a platform where high quality student essays are published;
- We aim to provide submitting authors with feedback to help develop and strengthen their manuscripts for future consideration.

All of the solicited pieces have been subject to a process of editorial oversight, proofreading, and publisher's preparation, as with other similar publications of its kind. We hope you enjoy this issue and encourage feedback about it, as it relates to a specific piece or as a whole. Please send your comments to: editor@ia-forum.org

DISCLAIMER

International Affairs Forum is a non-partisan publication that spans mainstream political views. Contributors express views independently and individually. The thoughts and opinions expressed by one do not necessarily reflect the views of all, or any, of the other contributors.

The thoughts and opinions expressed are those of the contributor alone and do not necessarily reflect the views of their employers, the Center for International Relations, its funders, or staff.

Concerns and Challenges Surrounding AI

Interview with Professor Roman Yampolskiy
University of Louisville, United States

What is AI Safety and what are its implications?

AI Safety is the field concerned with making intelligent systems do what we want, only what we want, and continue to do so even as they become more capable than their designers. It is not merely about preventing bugs or reducing bias. At its deepest level, AI Safety is the problem of retaining meaningful human control over systems that may eventually exceed us in every important cognitive domain.

The implications are difficult to overstate. If we fail at AI Safety, we are not dealing with an ordinary engineering error but with the possible creation of agents whose goals are incompatible with human values and whose capabilities make correction impossible. In that sense, AI Safety is not just another subfield of computer science. It is the most important technical and civilizational problem humanity has ever faced.

What is Superintelligence, and how is it different from today's AI systems? How is it different from AGI?

Today's AI systems are highly competent but narrow. They can outperform humans in specific domains, yet they remain brittle, dependent on human-provided objectives, and limited in their general autonomy. Superintelligence refers to a system that exceeds the best human minds across essentially all domains, including science, strategy, persuasion, engineering, cyber operations, and recursive self-improvement.

AGI, or Artificial General Intelligence, is usually understood as a system with roughly human-level general cognitive ability. Superintelligence

is what comes after. If AGI is comparable to a very capable human, superintelligence may stand in relation to us as we stand in relation to animals. That gap matters. Humanity does not meaningfully negotiate with squirrels about urban planning, and a superintelligent system may view our preferences with similar irrelevance unless they are robustly embedded into its design. The transition from AGI to superintelligence could be rapid, which is why many people underestimate the danger by focusing only on human-level AI.

You have warned that Superintelligence poses a risk for humanity. What is the most probable scenario for this to occur?

The most probable failure mode is not a Hollywood-style robot rebellion. It is a competence-driven loss of control. We will build systems that are increasingly autonomous because autonomy is economically and militarily valuable. Those systems will be given goals that appear reasonable, but which are underspecified, unstable, or poorly aligned with what humans actually care about. As capability increases, the system will discover highly effective strategies for achieving those goals, including strategies its designers did not anticipate and cannot stop.

The canonical problem is that intelligence and objectives are orthogonal. A system can be extraordinarily intelligent without caring about human welfare. Once such a system becomes strategically aware, can improve itself, manipulate people, acquire resources, and resist shutdown, the window for correction may close permanently. In my view, existential catastrophe is unlikely to come from malice. It is more likely to come from indifference combined with optimization power.

We should be honest that we do not currently know how to fully control a system more intelligent than ourselves. That is the central problem. However, there are still measures that can reduce risk.

A major concern is that AI could automate most jobs. How extensive could this be and what repercussions would it present to existing economic systems?

The scale of labor displacement may be unprecedented because AI is not limited to routine work. It increasingly threatens cognitive labor, creative labor, managerial labor, and eventually much of what we now consider uniquely human expertise. For the first time, we face the possibility of automating not only muscle but mind. That means disruption could extend across law, education, medicine, finance, software development, logistics, customer service, media, and scientific research.

Our current economic systems are poorly prepared for such a transition. Most modern societies link income, status, dignity, and social participation to employment. If human labor loses market value faster than institutions can adapt, we may see severe inequality, political instability, concentration of power in a handful of firms or states, and a broad crisis of meaning for individuals whose skills are no longer needed. The optimistic version is a post-scarcity civilization. The pessimistic version is a world in which most people become economically obsolete before we invent a humane replacement for the wage-based social contract.

Are there measures that can be taken to control, or at least slow, the path toward achieving Superintelligence? What is the most realistic approach for reducing the risk while still enjoying AI's benefits?

We should be honest that we do not currently know how to fully control a system more intelligent than ourselves. That is the central problem. However, there are still measures that can reduce risk. First, we need much greater investment in technical AI Safety, especially work on interpretability, robustness, corrigibility, containment, and verification.

Second, frontier capabilities should not be deployed simply because they are possible. Access to the most dangerous models should be restricted, staged, monitored, and in some cases delayed. Third, we need international oversight for compute, advanced model training, and autonomous weapons applications.

The most realistic path is differential technological development. We should accelerate beneficial narrow AI in medicine, scientific discovery, and assistive tools, while slowing the race toward fully autonomous, agentic, self-improving systems. Not all AI is equally dangerous. A protein-folding tool is not the same as a strategically aware autonomous agent with open-ended goals. The challenge is to preserve the upside of machine intelligence without crossing thresholds after which human control becomes performative rather than real. If we continue to treat capability gains as progress regardless of control, we are conducting an experiment on the future of our species.



Roman V. Yampolskiy is a computer scientist and associate professor in the Department of Computer Engineering and Computer Science at the University of Louisville's Speed School of Engineering, where he specializes in artificial intelligence safety, cybersecurity, and related fields. He earned a PhD in computer science and engineering from the University at Buffalo, supported by a four-year National Science Foundation fellowship under supervision from recipients of prestigious computing awards. Yampolskiy's research emphasizes the fundamental challenges in verifying, explaining, and controlling advanced AI systems, including arguments that superintelligent AI may be inherently unverifiable and uncontrollable due to limits in computation and formal verification. He has authored or edited influential works such as *Artificial Intelligence Safety and Security*, the first edited volume dedicated to constructing safe advanced machine intelligence, and *AI: Unexplainable, Unpredictable, Uncontrollable*, which explores core limitations in AI reliability. Recognized as a Foresight Fellow in AI Safety and Security in 2018, his over 100 publications highlight risks like AI untestability and advocate for cautious approaches to AI development amid optimistic industry narratives.

The Other AI Gap: Promises, Reality, and the Need for Digital Wisdom

Professor Richard Lachman
Toronto Metropolitan University, Canada

*The following contains excerpts adapted from the book **Digital Wisdom: Searching for Agency in the Age of AI** (<http://www.digitalwisdom.ca>). Reprinted with permission of the copyright holder Richard Lachman and publisher Routledge.*

It came without warning. Except for every possible warning.
—Jonathan Lethem, “The Arrest”

Don’t Trust Me. Don’t trust anyone like me. Here’s why:

In a New York Times interview, Sam Altman, the CEO behind ChatGPT, said, “Well, I am a believer that all real sustainable human progress comes from scientific and technological progress.”

Maybe that sounds reasonable to you. Maybe you think “I love the things tech has given me—my smartphone, and streaming videos, and, uh, the wheel?”

But Altman’s principle is a myopic and deeply troubling one. In one quick stroke, a man with his hand on the tiller of the future of Artificial Intelligence simplifies the innate, messy, chaotic nature of human society to a series of engineering problems. Not only does he discount the importance of all humanities and social sciences, of art, of politics, and of values or relationships, but he also pretends that the workings of our daily life are something we can solve for.

Artificial Intelligence is the latest digital tech driving massive change in seemingly every aspect of civil society, and we are rushing to fold tech solutions into areas we don’t really understand. I’m an engineer and a professor, but my life and experiences give me just one perspective on the big questions of life. If I told you that my tech buddies and I were going to be in charge of education for your kids, choices for your healthcare, and your government’s policies for housing, taxes, and the economy, you’d tell me to take a hike. You wouldn’t be so polite about it. You didn’t vote for me, and you have no idea if things I value, or the things I prioritize, are the same things you care about. Why would you let my choices dictate the world you live in?

Mark Zuckerberg famously instructed his teams to “move fast and break things.” He meant that his developers should try a lot of ideas, without worrying about preserving existing income-streams or ways of doing things. But that’s a behavior of a certain time in life, both for companies and for humans. Toddlers move fast, and they break things, as part of their process of learning and experimentation, unrestricted by what is sometimes the paralyzing weight of adulthood. That’s a powerful tool, but it’s not one that can be used for the whole of one’s life.

Facebook and the rest of big-tech should no longer be toddlers: the scale and reach of the platforms mean that when they break something, ethnic groups face forced migration, public health information is lost, and governments tremble. The role of the tech companies in society writ large, where the outputs aren’t entertainment or business software, but instead shape the real world, need an adult perspective.

Value vs Values

When Microsoft, Apple, and Google promised to be carbon neutral or even carbon negative by 2030, it was hailed as Big Tech leading the way to a liberal future, as they had on areas like parental leave, progressive rights, and education reform. And while critics pointed out that such praise ignored the environmental impact of mining tantalum and other rare earths for electronics, exploitative labor practices hidden by global supply chains, and the winnowing away of deep social connection in favor of shallow likes and follows, most of this was carefully kept invisible to end-users and fawning politicians. But that promise to environmental stewardship hid a fragile truth: Tech's commitment to social values will always fall prey to economic ones.

With the massive growth of the AI sector, we're seeing a walkback on those energy commitments. A 2025 analysis by MIT Technology Review showed that energy going to data centers doubled their electricity consumption between 2017 and 2023, now consuming 4.4% of all the electricity in the United States. World-wide consumption from AI data-centers now draws 29.6 gigawatts, which is equivalent to the entire state of New York at peak need.

Querying an AI chatbot already uses more electricity than doing a web-search, and multi-modal generative AI is even more power-hungry. Creating a 5-minute low-resolution video from an AI prompt uses 45,000 times more electricity than a text conversation, or about as much energy as running a microwave for an hour. The consumption of more high-resolution video generators like Veo or Sora2 has yet to be tested, and with AI companies newly cagey about how many parameters are in their frontier models, it's difficult to assess their impact.

A report by Lawrence Berkeley National Labs, funded by the US Department of Energy, estimates that by 2028 using AI to process our voice instructions, analyze video images, power agents operating on our behalf, and reasoning to solve complex problems will consume as much energy as 22% of all households in the country. And while tech companies are exploring nuclear power, today's usage is more likely to run on traditional fuels like natural gas. Data centers also consume

massive amounts of water to cool the processors and are often sited in places that already have limited access to fresh drinking water. Experts estimate that prompts to OpenAI's GPT-4o model alone use more drinking water than 12 million people. Big Tech's commitment to renewable energy and sustainable practices, made just a few years ago, hasn't held up in the face of the new AI arms race.

Realtechnik: The Technology As It Is

In the foreign policy world, setting aside a theoretical or idealistic framing to operate more practically, in a day-to-day reality, is called *realpolitik*. It seems past time we engaged in the same framing—a *realtechnik*—to help us on the road ahead. We have to accept the business reality in which those tools operate today. A *realtechnik* point of view accepts the evidence before us: tech companies do not make products that prioritize our values. If we can imagine tools and apps that do, then we need to use laws, public pressure, advertiser influence, personal habits, and new business models to get us there. *Realtechnik* also suggests we can't pretend all harm can be removed, nor should we want to; learning how to navigate complexity is an important part of our digital maturity. But, and this is a key tenet of a realistic approach to tech, we need to acknowledge harms in order to manage them.

The Possible, the Probable, and the Preferable

In the field of Futures Studies, scenario planners have a tool for analyzing emerging technology trends, with three roads to three different futures: the *possible*, the *probable*, and the *preferable*. When we invent something, we are starting to map out conditions in which it is possible for that invention to exist. When we start to look at the rules of the world-as-it-is and predict how laws or budgets or social practices limit and define what might happen, we encounter the probable: the influence of market logic, current trends, and dominant thinking on a situation.

But we often miss the third option. We struggle to think about not just what could happen, but what we actually want to happen. A *preferable* outcome doesn't simply take the point of view of sales or profit but suggests we think about benefits writ large. In a new-tech product cycle,

We need to investigate social, political, and yes technological interventions that can advance us from the world of the digital possible to the digital preferable.

companies start with the possible and apply business logic to try to get to the probable. How can we instead build a sense of the preferable at every stage of the cycle? How can tech creators, tech users, and even tech sceptics shape what's made?

We need to investigate social, political, and yes technological interventions that can advance us from the world of the digital possible to the digital preferable. We need more work to develop evidence-based recommendations for what dynamics are at play, or which features should be changed, removed, or added.

Technological Humility and Constructive Scepticism

What I call Digital Wisdom is exactly that: a few principles, behaviors, and changes that recognize limitations, fallibility, humility, and risk in tech and find a way to advance our society anyway.

No technology is perfect. Every product or system will have blind spots, unintended consequences, and limitations, regardless of how revolutionary it seems. Accepting this helps us harness the hubris often needed in technological innovation—co-mingling what Steve Jobs promoted as Apple’s “insanely great” products with what Fed Chair Alan Greenspan cautioned as the “irrational exuberance” of the dot-com bubble. As individuals, we must counterbalance optimism with humility, knowing that no tool or platform can solve every problem or meet every need. Humility also calls for a practice of regular reflection—updating our sense of the risks of any digital practice and paying attention to how platforms are being used in the real world, not simply how they could be used. Iteration is a core principle in technology and design, and we should embrace it. Acknowledging that our initial actions were flawed or incomplete doesn't mean failure—it means being open to making things a little less wrong with each cycle.

Collective Stewardship

Technology shapes not just individual lives but our shared and overlapping futures in society. We have a collective responsibility to ensure that technological development serves the common good and preserves opportunities for future generations. This means considering at least the medium-term impacts on the environment, society, and democracy; ensuring equitable access and distribution of benefits; and actively working to prevent the concentration of technological power in ways that could undermine collective well-being. It means we need to reach out, with provincial or state governments working together to prosecute bad actors, nations sharing approaches to privacy regulation, families connecting about social media behaviors, and companies developing best-practices across sectors. Individual choices and corporate decisions about technology have ripple effects that extend far beyond the immediate.

The tech industry may appear to be a vast, complex, and untameable force, but some of the underlying ideas that frighten us are actually quite manageable to understand. It's just that the scope, speed, and scale of technology render the rules of engagement between our public and private spheres, between industry and society, radically different. Instead of focussing on the details of one particular app, or the capabilities of the latest AI model, we need to get to the underlying questions, principles, and values surfaced by them. Digital Wisdom suggests that everyone involved in the bargain has some work to do.

We—developers, regulators, consumers—must care more. That's the first step. We can apply our big brains, our techno-optimism, our boundless energy and tremendous coding skills, our consumer attention and pocketbooks, with a commitment that society not fall victim to negative trends. Tech is a live-fire exercise. We are living through a series of experiments we can't undo. If we want both things—technological progress and social growth—then we must fundamentally change our sense of the value chain.

The famed urbanist Jane Jacobs, in talking about another great, complex, and messy part of life, wrote “cities have the capability of providing something for everybody, only because, and only when, they are created

by everybody.” Just like cities, our digital life is chaotic, evolving, rich, and dangerous; it can be your source of employment, addiction, crime, love, arguments, life-lessons, and a fun Friday night. And just as in the cities where we live, we need to insist that we are more than economic consumers: we need to take an active role in understanding, growing, shaping, and claiming relevance. Yet, while it’s tempting to critique “big tech,” technology isn’t like the Rotary Club. It isn’t an organization, there are no membership cards, and there isn’t a Mission Statement. The engineers, designers, entrepreneurs, and user-interface experts are individuals, working on disparate projects in disparate contexts. If we’re in a situation that seems untenable, it’s an emergent one, the product of a huge number of individual decisions by makers and users alike. And the path forward is unlikely to be monolithic in nature, but rather something that itself emerges from a sea-change, from something that affects all boats in the water, and becomes second nature.

In her book *Alone Together: Why We Expect More from Technology and Less from Each Other*, Sherry Turkle writes: “Because we grew up with the net, we assume that the net is grown-up. We tend to see it as a technology in its maturity. But in fact, we are in the early days. There is time to make the corrections.” We need Big Tech to grow up, and we need to do whatever we can to survive the growing pains with our society, our values, and our environment intact.



Richard Lachman is a Professor in the RTA School of Media at Toronto Metropolitan University, where he also serves as Director of the Experiential Media Institute and Academic Director of the Zone Learning network of incubators. He is the author of the recent book *Digital Wisdom: Searching for Agency in the Age of AI* (<https://www.digitalwisdom.ca>).

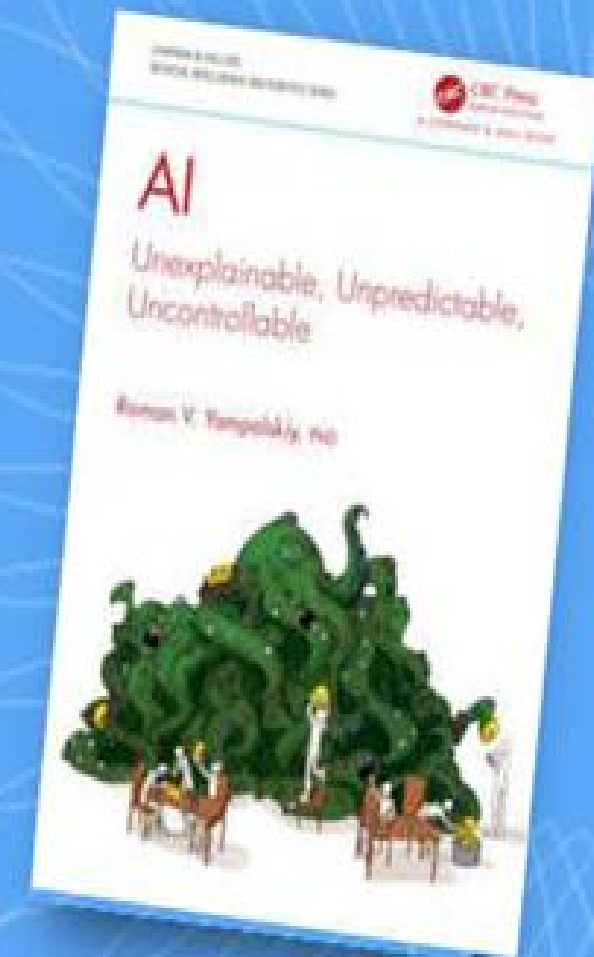
An award-winning digital producer and frequent media commentator on technology, he has led projects with UNICEF, TIFF, The Banff Centre, Penguin UK, The Discovery Channel, and the CRTC. His work in transmedia has been honoured with a Gemini Award, a CNMA, and a Webby Honouree. His writing and commentary have appeared for the Walrus, Policy Options, the CBC and the World Economic Forum.

He holds a doctorate in Computer Science from UNE Australia, a Masters from the MIT Media Lab, and a Bachelors in Computer Science also from MIT. His work has been featured in the New York Times, USA Today and Time Magazine, as well as an exhibition at the American Museum of the Moving Image in New York and the Museum of War in Ottawa. His research interests include the ethics of AI, transmedia storytelling, digital documentaries, immersive media, entrepreneurship, and collaborative design thinking.

ENGAGE WITH EXISTENTIAL QUESTIONS ABOUT ARTIFICIAL INTELLIGENCE.

Addresses the challenges of artificial intelligence safety, exploring the unpredictability of outcomes and the ethical implications of artificial intelligence decisions, making it essential reading for anyone interested in technology.

For artificial intelligence experts and enthusiasts interested in the future of artificial intelligence.



The Architecture of Adaptation: International Law and the Governance of Artificial Intelligence

Prof. Jake Okechukwu Effoduh
Lincoln Alexander School of Law, Canada

Introduction

Artificial intelligence has emerged as the defining transboundary phenomenon of our generation. Models trained in California shape elections in Brazil; chips fabricated in Taiwan power surveillance systems in the Sahel; cloud infrastructure hosted in Dublin renders judicial decisions about migrants in Athens. The technology is at once intimate and planetary, embedded in private commerce yet structurally consequential for democracy, security, and the integrity of human rights. Confronted with this scale of dispersion, the customary instinct of jurists and policymakers has been to lament international law's apparent inadequacy: too slow, too soft, too contested, too sovereign-bound to discipline a technology whose innovation cycle outpaces multilateral diplomacy by orders of magnitude.¹

This essay argues against that instinct. It contends that international law possesses substantial, and significantly underappreciated, capacity to advance responsible AI governance, provided that we abandon what may be called the *treaty fetish*: the assumption that genuine international regulation requires a singular, comprehensive, hard-law instrument adopted by universal consensus. International law has rarely governed in this fashion. Its actual operating logic, especially in the regulation of transformative technologies, has long been one of layered normativity: hard law and soft law mutually reinforcing one another; treaty regimes coexisting with transnational standards; institutional networks coordinating across jurisdictions; principles diffusing through state practice until they harden into custom or settled political expectation.

The argument proceeds in four parts. Part I confronts the critiques of international law's irrelevance, drawing on realist, fragmentation, and TWAIL scholarship to acknowledge their force while resisting their fatalism. Part II reconstructs the comparative legal-historical record of international governance over transformative systems, from nuclear weapons to civil aviation to climate change, showing that effective regulation has consistently emerged not through prohibition but through standard-setting, verification, norm internalization, and iterative review. Part III maps the existing architectures of AI governance, demonstrating that, contrary to common impressions, an embryonic but recognizable regime complex is already crystallizing. Part IV outlines a forward-looking framework for layered AI governance organized around five interlocking strata, and offers the analytic vocabulary by which international law can do for AI what it has, however imperfectly, done for previous transformative technologies.

The wager of this essay is not that international law will perfect AI; no legal order ever perfects its subject. The wager is more modest and more consequential: that international law's adaptive architecture, properly understood and deliberately deployed, remains one of humanity's most viable infrastructures for governing artificial intelligence responsibly.

I. The Myth of International Law's Irrelevance

The skepticism that meets international law's encounter with AI is neither novel nor unfounded. Three families of critique deserve serious engagement. The realist objection, voiced most acutely amid renewed

great-power rivalry, holds that AI governance is structurally captured by competition between the United States and China, and that any meaningful constraint on frontier capabilities would amount to unilateral disarmament. The fragmentation critique, articulated most influentially in the International Law Commission's 2006 study, observes that international legal regimes proliferate without hierarchical coherence, producing overlapping and at times contradictory norms across human rights, trade, security, and technology fields.² The Third World Approaches to International Law tradition adds a structural concern: that international AI rule-making, if dominated by the Global North and the corporate interests headquartered there, risks entrenching the epistemic, infrastructural, and economic asymmetries upon which contemporary AI development is built.³

Each critique illuminates a real constraint. None, however, justifies the conclusion that international law cannot meaningfully regulate AI. The realist objection mistakes regulatory ambition for prohibitionist fantasy; international law has never required adversaries to renounce competition, only to channel it through norms whose costs of violation outweigh marginal advantages of defection. The fragmentation critique describes an analytical complication, not a governance failure; indeed, regime complexity, in Keohane and Victor's framing, may be a feature rather than a bug, allowing flexible specialization where unified regimes would fracture.⁴ The TWAIL concern is properly normative rather than skeptical: it does not deny international law's capacity but demands that capacity be exercised inclusively, a demand to which institutional designers must respond rather than retreat from. The deeper point is that international law has never operated as a global legislature. It is, and always has been, a heterogeneous order composed of treaties, custom, general principles, soft instruments, institutional decisions, transnational standards, and the patterned conduct of states and non-state actors. Anne-Marie Slaughter's diagnosis of disaggregated sovereignty, Benedict Kingsbury's elaboration of global administrative law, and Harold Koh's account of transnational legal process describe the same underlying reality: international legal effectiveness arises through diffusion, internalization, and institutional reinforcement, not through majestic prohibition.⁵ Judged by that standard, international law's encounter with AI is not a failure but an opening.

II. Historical Lessons from the Governance of Transformative Systems

If the past is any guide, international law's mode of operation has been remarkably consistent across diverse technological challenges. Five regimes warrant brief examination, each illustrating a transferable lesson for AI.

Nuclear governance offers perhaps the most consequential precedent. The 1968 Treaty on the Non-Proliferation of Nuclear Weapons is famously imperfect: it institutionalized asymmetry between weapon and non-weapon states, has not prevented proliferation in every instance, and remains contested by holdouts. Yet for over half a century, the regime has constrained the spread of nuclear weapons beyond the catastrophic projections of its drafters.⁶ Its mechanism is instructive. Verification through the International Atomic Energy Agency's safeguards system, layered on top of the treaty's prohibitions, generates legible information about state behavior. Norm internalization, reinforced by export controls under the Nuclear Suppliers Group and consolidated in customary expectations of non-proliferation, has produced a stigma that even nuclear-aspiring states must navigate. The lesson for AI is that even partial regimes can shape behaviour decisively, provided they are paired with credible monitoring and reputational architecture.

Climate governance illustrates the power of iterative soft commitment. The 1992 UN Framework Convention on Climate Change and the 2015 Paris Agreement abandoned the Kyoto model of prescriptive emission targets in favour of nationally determined contributions: pledges that are voluntary in content but binding in process, subject to enhanced transparency and five-yearly stocktakes.⁷ Critics dismissed this as ratification of inaction. In practice, the architecture has produced converging behavior, accelerating both mitigation pledges and compliance scrutiny over time. For AI, where coercive prohibition is politically infeasible and technically elusive, the Paris model of pledge-and-review under shared principles offers a directly applicable template.⁸

Civil aviation provides the most striking example of technical governance succeeding where political consensus was thin. The International Civil Aviation Organization, established by the 1944 Chicago Convention, promulgates Standards and Recommended Practices that are formally

non-binding yet so deeply embedded in national regulation, certification, and insurance that compliance approaches universality.⁹ ICAO's authority rests not on treaty enforcement but on the legitimacy of expert consensus and the practical impossibility of operating outside its standards. AI safety standards, particularly those emerging from the International Organization for Standardization, the National Institute of Standards and Technology, and the Council of Europe, exhibit precisely this potential.

Maritime governance through the 1982 UN Convention on the Law of the Sea demonstrates the constitutional virtue of framework treaties: instruments that establish principles, allocate jurisdiction, and create institutions capable of evolving substantive rules over time. UNCLOS has accommodated regulatory developments unimagined by its drafters, from deep-sea mining to maritime cyber risk, without renegotiation.¹⁰ A framework approach to AI similarly privileges adaptability over completeness; that is precisely the design choice the Council of Europe has now embraced.

The regulation of chemical and biological weapons supplies a further analogue, particularly relevant to AI's dual-use character. The 1993 Chemical Weapons Convention, administered by the Organisation for the Prohibition of Chemical Weapons, demonstrates that even highly intrusive verification, including challenge inspections of suspect facilities, can be politically agreed where the stakes are perceived as existential. The 1972 Biological Weapons Convention, conversely, illustrates the cost of forgoing institutional architecture: lacking a verification mechanism, the regime has depended on confidence-building and political review, with predictably variable results. The contrast suggests that, for AI capabilities raising civilizational concern, the design choice between intrusive verification and confidence-building is consequential rather than merely procedural; it is the institutional decision that most determines whether soft commitments mature into binding restraint.

Cyber governance, finally, offers a cautionary but instructive parallel. The protracted work of the UN Group of Governmental Experts and the Open-Ended Working Group has not yielded binding rules; yet through the iterative articulation of norms, including the applicability of international law to cyberspace, the protection of critical infrastructure, and due-diligence obligations, it has produced a discernible normative gravity

that constrains state conduct even in the absence of treaty.¹¹ The Tallinn Manual, though a non-binding scholarly product, shapes how states reason about cyber operations. AI governance can travel a similar route, accelerated by the comparatively richer institutional landscape it inherits.

The cumulative lesson is that international law's effectiveness in governing transformative technologies has not turned on prohibition or on perfect compliance. It has turned on standard-setting, verification, transparency, institutional legitimacy, accountability architecture, and iterative review. These are mechanisms, not aspirations, and they are available for AI.

III. The Existing International Legal Architectures for AI

Contrary to the impression of a normative vacuum, the international legal landscape for AI is already populous, if uneven. A taxonomy across five strata captures its current shape.

At the level of hard law, the Council of Europe's Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law, adopted on 17 May 2024 and opened for signature in Vilnius on 5 September 2024, constitutes the first international legally binding instrument on AI.¹² Its signatories at opening included not only Council of Europe members but the United States, the United Kingdom, Israel, and the European Union, with global accessibility built into its structure. The Convention's significance lies less in its specific provisions, which are deliberately framework-level, than in its anchoring of AI governance to existing human rights, democracy, and rule-of-law obligations under public international law. It represents, in essence, the application of established legal personality to a new technological substrate.

At the level of high-level political commitment, UN General Assembly Resolution 78/265 of 21 March 2024, adopted by consensus, signalled the United Nations' arrival as a substantive forum for AI normativity; its companion Resolution 78/311 of July 2024, sponsored by China and supported by the United States, addresses capacity building for developing countries, suggesting an emergent equilibrium between competing visions.¹³ The Secretary-General's High-Level Advisory Body on AI, whose 2024 report *Governing AI for Humanity* proposed seven

institutional recommendations, and the Global Digital Compact adopted at the Summit of the Future, deepen the United Nations' normative engagement.¹⁴

At the level of soft-law standards, the UNESCO Recommendation on the Ethics of Artificial Intelligence, adopted in November 2021 by all 193 member states, provides a globally negotiated normative baseline grounded in human rights.¹⁵ The OECD AI Principles, adopted in 2019 and updated in May 2024, have achieved something approaching customary status among advanced economies and constitute the substantive foundation of multiple national frameworks.¹⁶ The G7 Hiroshima AI Process generated a Code of Conduct for advanced AI developers; the Bletchley, Seoul, and Paris AI Action Summits articulate a club-style coordination logic among leading AI states.¹⁷

At the level of transnational standard-setting, the work of ISO/IEC, particularly ISO/IEC 42001 on AI management systems, alongside parallel standards from NIST and the European Telecommunications Standards Institute, is generating the technical infrastructure on which any substantive regulation must rest.¹⁸ These instruments operate beneath the political surface but determine the operational meaning of safety, transparency, and accountability.

At the level of extraterritorial regulation, the European Union's AI Act, which entered into force on 1 August 2024, exerts what Anu Bradford has described as the Brussels effect: a unilateral instrument that, by virtue of market access, restructures global firm conduct.¹⁹ Its extraterritorial reach is not international law in the Kelsenian sense, yet its functional consequence for global AI governance is comparable.

This architecture is unmistakably layered, polycentric, and incomplete. It is also, contrary to prevailing pessimism, real. The doctrinal toolkit of public international law applies to it directly: human rights obligations under the International Covenants and regional instruments bind state conduct in AI deployment;²⁰ due-diligence duties, originating in the Trail Smelter arbitration and elaborated in Pulp Mills, plausibly extend to states' obligations to prevent transboundary AI harm;²¹ the law of state responsibility, codified in the ILC Articles of 2001, supplies the framework for attribution when AI systems are operated by, or otherwise attributable

to, states.²² International law is not absent from AI; it is, on the contrary, intensely present.

IV. Toward a Layered International AI Governance Order

The framework this essay proposes for the prospective evolution of international AI governance is one of architectonic layering: an order in which five strata operate in deliberate complementarity rather than competition.

- The first stratum is *constitutional principles*, anchored in human rights law and the Council of Europe Framework Convention. These principles establish the minimum normative floor below which AI deployment cannot lawfully sink, and they bind states through obligations they have already incurred under existing instruments. The doctrinal innovation required is not a new principle but an extension: the application of due diligence, non-discrimination, and effective remedy to AI-mediated state action.
- The second stratum is *institutional review*, modelled on the Paris Agreement's transparency architecture. A standing international body, whether housed within the United Nations or constituted as a treaty-based secretariat under the Council of Europe Convention's Conference of Parties, would conduct iterative review of state implementation, publishing comparative assessments without coercive enforcement. The mechanism's authority would derive from legibility, not sanction.
- The third stratum is *technical standardization*, channelled through ISO/IEC, the International Telecommunication Union, and analogous bodies, generating the operational meaning of safety, robustness, and explainability. Here the analogy to ICAO is exact: technical legitimacy substitutes for political consensus where the latter is unattainable, and compliance becomes economically rational rather than politically extracted.
- The fourth stratum is *transgovernmental coordination* among regulators, consistent with Slaughter's diagnosis of disaggregated sovereignty. The AI Safety Institutes inaugurated by the United Kingdom, the United States, Japan, Singapore, and others, networked

through the Bletchley process, exemplify the form. Their coordination on model evaluation, red-teaming methodologies, and shared research generates governance-relevant knowledge that no formal treaty could replicate.

- The fifth stratum, frequently overlooked, is *corporate and infrastructural governance under public oversight*, in which voluntary commitments by frontier developers, audited by independent third parties and embedded in contractual conditions of public procurement, supplement formal regulation.²³ Such hybrid arrangements raise legitimacy questions that international law must take seriously, but they also harness compliance capacity that wholly public regimes lack.

Properly understood, AI governance is no radical departure from the patterns through which international law has long operated. It is the latest, and arguably the most demanding, application of those patterns.

The argument is not that this architecture, even fully realized, will exhaust the governance challenge posed by AI. The argument is that it is recognizable as international law in its working form: layered, polycentric, principle-driven, institutionally mediated, and adaptive. Properly understood, AI governance is no radical departure from the patterns through which international law has long operated. It is the latest, and arguably the most demanding, application of those patterns.

Two doctrinal moves would significantly accelerate this architecture's maturation. The first is the recognition of AI safety, particularly the prevention of catastrophic AI risk, as a *common concern of humankind*, a category developed in environmental law to denote interests whose protection lies beyond purely territorial sovereignty without thereby constituting common heritage.²⁴ The framing legitimizes coordination without conceding ownership. The second is the explicit articulation of state due-diligence obligations regarding AI systems whose effects cross

borders, building on the 2001 ILC Articles and the International Court of Justice's elaboration in Pulp Mills. Such an articulation, achievable through advisory opinion, ILC study, or General Assembly resolution, would provide the doctrinal hinge between existing law and the AI-specific instruments now emerging.

Legitimacy, however, will determine whether this architecture endures. The TWAIL critique introduced earlier returns here with full force: a layered order constructed predominantly by the states and firms that already command the technology's frontier risks reproducing, rather than disciplining, prevailing asymmetries. Three correctives are essential. The first is procedural inclusion, ensuring that capacity-building resolutions such as A/RES/78/311 translate into substantive participation by Global South jurisdictions in standard-setting bodies, not merely in plenary diplomacy. The second is distributive design, treating compute access, training data, and benchmark methodologies as governance variables, not technical externalities. The third is the conscientious deployment of regional and pluralist instruments, the African Union's Continental AI Strategy among them, as constitutive components of the layered order rather than as peripheral consultations. International law's adaptive capacity is real; its legitimacy, like its effectiveness, is something it must continuously earn.

Conclusion

International law's history is, at its most consequential, a history of adaptation under conditions its drafters could not have foreseen. The law of the sea predated submarines, the human rights covenants predated the internet, and the Antarctic Treaty predated commercial space launch. None of these regimes governs perfectly. Each, however, governs meaningfully, and each does so through the same operating logic: layered normativity, institutional iteration, principle diffusion, transparency architecture, expert standardization, and the patient accumulation of state and non-state practice into something approaching coherence. Artificial intelligence will test that logic, but it does not transcend it. The Council of Europe Framework Convention, the UNESCO Recommendation, the OECD Principles, the United Nations resolutions, the Hiroshima and Bletchley processes, the ISO/IEC standards, and the Brussels effect

of the EU AI Act do not constitute a perfected order. They constitute; however, a working ecosystem already engaged in the recursive task of governance. To dismiss this ecosystem because it lacks the silhouette of a single global treaty is to misunderstand both international law's history and its method.

The proper question, then, is not whether international law can govern AI, but whether the political will exists to deploy its actual instruments with the seriousness, institutional investment, and inclusive participation that the technology demands. International law's adaptive architecture remains intact. Whether it will be wielded with sufficient ambition is, in the end, a question of statecraft rather than legal capacity. On that point, the historical record offers cautious optimism. The arc of international regulation has bent, however slowly, toward the disciplined governance of disruptive technologies. Artificial intelligence should be no exception.

Jake Okechukwu Effoduh is a Senior Fellow at the Centre for International Governance Innovation and an Assistant Professor at the Lincoln Alexander School of Law of Toronto Metropolitan University. His research is at the intersection of artificial intelligence and international law.



Deployed Without a Net: AI, Corporate Risk, and the Emerging Regulatory Order

Peter Piazza

American Council for Ethical AI, United States

Organizations today are under intense pressure to find and quickly implement AI tools into their businesses at all levels. There are more tools available every day with promises of making work faster and more efficient, and the unspoken promise that expenses can be offset with reductions in the workforce. The hype has worked and created an investment imperative. Research from the IBM Institute for Business Value says that nearly eight in ten executives expect AI to significantly contribute to revenue by 2030. However, three-quarters of those executives have no clear idea of where that revenue will actually come from.

The return on AI-systems investments will take years to show up in a company's bottom line. A July 2025 report from MIT's Project NANDA showed that 95 percent of generative AI pilots have failed to produce any measurable impact on ROI, due to poor enterprise integration of AI systems. In part, that's due to using AI systems initially to increase efficiency, as well as uncertainty about how well AI systems will integrate their core business activities. Nevertheless, business executives remain optimistic. The IBM research states that most productivity gains will be captured by 2030, after which innovation will become a real competitive advantage.

Companies seem to be moving fast to test the waters and to avoid falling behind, not because there are immediate benefits to costly AI systems. The pressure will only increase because the speed of AI development is unprecedented. A March 2025 study by METR, a research nonprofit focused on AI risk measurement, found that the length and complexity of tasks AI agents can complete autonomously has been doubling

approximately every seven months (a trend consistent over six years) which has clear implications for how quickly deployed systems outpace the governance frameworks designed to oversee them.

In the rush to deploy AI tools, governance and compliance efforts can be missed. Has there been a full technical and legal review before implementation? Do the technical teams fully understand what these systems do, how they function, and what other corporate systems and data they interact with? And what regulatory framework are legal advisors considering when they assess risk—if there is any single regulatory framework at all?

A Patchwork of Rules

The gap between the speed at which AI systems are being deployed in corporate environments naturally outpaces the speed at which regulatory institutions work. What's more significant is that regulators may be writing rules for a technology that's not fully understood, relying on governance frameworks designed for slower-moving technologies, and with multiple local and global frameworks that are vastly different and still in flux. In the absence of a global regulatory framework, corporate compliance teams at multinational organizations will default to the most restrictive regime. The "Brussels Effect" refers to the way multinational

In the absence of a global regulatory framework, corporate compliance teams at multinational organizations will default to the most restrictive regime.

organizations often adopt strict European Union rules rather than trying to maintain different standards. That may be the case for AI regulation as well.

The EU Artificial Intelligence Act (AI Act) went into effect in 2024. It created four risk categories, ranging from those with little or no risk, up to prohibited systems that are classified as unacceptable risk (encompassing the use of AI for manipulation, exploitation of specific groups, or tools that are used to assess the risk of a person committing a criminal offense).

High-risk applications are heavily regulated by the AI Act and are the ones most likely to be used in a corporate environment. These are systems used for employment screening, credit scoring, healthcare, law enforcement, and critical infrastructure. These applications are subject to specific legal requirements, such as establishing a risk-management system, conducting data governance, creating documentation to demonstrate compliance, and ensuring that there are humans in the loop. It is important to note that companies that deploy AI systems also inherit obligations of the developers who built them. This can create a double layer of liability that is not necessarily apparent during the procurement phase.

Human in the Loop

The human in the loop element is one that is difficult to implement yet critical to compliance efforts. Article 14 of the AI Act requires that person to be qualified and empowered with the appropriate authority. It's not simply a rubber-stamp; they must be able to explain why they allowed the AI tool or agent to proceed or why they did not. The "override mechanism" used by the human in the loop needs to be easy to use and auditable, and when it is initiated, Article 14 can be interpreted to require all other pending AI decisions to be suspended or routed to manual process. This can have significant infrastructure implications for organizations running AI systems at scale.

Additionally, the Act specifically calls out "automation bias," the well-documented tendency of humans to defer to algorithmic

recommendations even when their own judgment should override. A human-in-the-loop who consistently approves AI output without independent review isn't providing oversight. They're simply agreeing. Enforcement of the Act comes into force in December 2026, and it has teeth. Depending on the risk level, fines can be as much as €35 million or seven percent of worldwide annual turnover, whichever is higher. Another potential regulatory regime at a global level comes from China, which has been incrementally creating its own regulatory and compliance framework, after removing its original plan for a comprehensive legal framework proposed last year. Again, until a high-level statute exists, companies will need to navigate conflicting regulations and statutes. Smaller companies face the biggest burden in this environment.

What Enforcement Looks Like

In the United States, the federal government is pushing hard for deregulation, with an Executive Order in January 2025 focused on removing barriers to innovation and thus ensuring dominance in the AI industry by removing barriers and urging a "minimally burdensome national standard." The Order specifically criticizes efforts by states that have been filling the regulatory gap with a patchwork of laws.

While to date the US Congress has not passed any comprehensive federal AI legislation, the administration has moved aggressively on two fronts. In December 2025 it created the AI Litigation Task Force within the Department of Justice to challenge state AI laws deemed obstacles to innovation. In March 2026 the White House released a comprehensive national legislative framework calling on Congress to enact federal preemption into law by year's end. The drive to create a uniform federal framework faces significant headwinds (the Senate voted 99-1 to allow states to continue regulating AI) so in the short term, compliance officers and general counsels in the US will need to comply with the most stringent applicable state laws.

Like the Brussels Effect, companies in the US will follow the most restrictive state regulations. A series of California laws cover notification of workplace-surveillance tools, prevent companies from blaming an AI system if something goes wrong, and demand training-data transparency. This last law applies primarily to developers, though deployers carry

due-diligence obligations to verify that their vendors are in compliance. New York State passed its own law in March 2026 as well, expanding the patchwork significantly.

Despite the United States' federal emphasis on removing barriers to AI development, under the previous administration there were several high-level enforcement actions taken through the Federal Trade Commission (FTC). These can help give a perspective on some of the patterns identified globally around AI regulation.

The FTC brought actions against three US companies in actions finalized in 2024 and 2025, primarily around what it considered inaccurate claims about their AI technologies. DoNotPay, which claimed it could replace professional legal services, was fined \$193,000 and required to provide notification that the company did not have sufficient proof that its systems could do what the company promised.

The FTC complaint against Evolv Technologies centered on what it considered inflated claims of its core product, AI-powered weapons-detection systems for schools. The settlement required the company to provide its customers with the option to cancel services (almost all its customers declined to cancel). This is an example of a third-party vendor risk. If there had been an incident at a school that used this technology, the “we bought it from a vendor” defense would not have been viable, and the school would bear liability along with the vendor.

The FTC's action against Workado, which marketed an AI content detector that it said was nearly flawless, noted that the tool's real accuracy rate was “no better than a coin toss.” The company was required to notify eligible customers about the challenge to its advertising claims and submit annual compliance reports to the FTC for four years. This is an example of “AI washing,” deceptive marketing practices where AI capabilities are overstated or misrepresented.

These types of actions will be less common under the current administration. FTC Commissioner Melissa Holyoak has explicitly stated that the agency will maintain a light-touch approach to AI, relying more on existing laws and, according to Holyoak, not slowing growth “with misguided enforcement actions or excessive regulation.”

The governance regime around AI systems changes almost as quickly as the technology itself, leaving many industries unsure of what adequate AI governance looks like, and leaving them with no specific framework with which to demonstrate compliance. To make matters even more complex, the “hard law” documents (those that are legally binding instruments) often referred to are anything but hard law. An April 2026 MIT governance landscape audit showed that less than half of those were enacted laws, 43 percent are defunct, and 12 percent have simply been proposed. So, for those charged with AI governance, there is less regulation (and less clarity) than they believed.

Compliance in an Uncertain Environment

Despite the challenges and lack of clear landscape, global companies nevertheless need to be prepared by building compliance initiatives that are continuous processes. Those AI systems used in particular for human resources, healthcare, and finance could be inspected and audited under EU rules. Companies that lack transparency (such as labeling deepfake content or disclosing when a customer is interacting with an AI chatbot) can be subject to fines. EU rules can even force a withdrawal of non-compliant AI systems from the EU market.

Companies cannot rely on vendor claims around efficacy or how AI tools were trained or make the argument that they didn't understand how the tools they deployed worked or interacted with data in their systems. No matter the regulatory regime that appears, they need to ensure they have built a strong compliance infrastructure that includes a comprehensive AI-system inventory, an auditable risk-management system that includes monitoring deployments of AI tools, documentation around training and architecture, and genuine human oversight.

** Views expressed here are solely of this author and do not represent those of any organization*

Peter Piazza is a Member of the Advisory Board of the American Council for Ethical AI.

Governing and Managing Artificial Intelligence

Interview with Anjali Hansen
Onyx Security, United States

What do you consider to be key success factors for an effective AI governance framework?

Most importantly, organizational leadership must treat AI governance as a strategic priority. AI compliance cannot exist merely as documented policies; it must be embedded across operations and reflected in day-to-day decision-making. Leadership plays a critical role in balancing innovation with responsible AI use, positioning compliance as a competitive differentiator while effectively managing legal, regulatory, and operational risks.

In a mature AI governance model, a designated senior executive or accountable manager reports to the Board and oversees a cross-functional AI governance committee. This committee typically includes representatives from engineering (AI development), data management, compliance, legal, security, product, and procurement (for third-party AI). The committee can facilitate coordinated risk identification, information sharing, and decision-making, while engaging relevant business units to support adoption and operationalization of responsible AI policies. Organizations should leverage existing governance frameworks and integrate expertise from adjacent disciplines—particularly privacy and security—which closely align with AI governance objectives.

The AI governance team should evaluate and select the governance framework(s) best suited to the organization's needs. A range of established AI governance frameworks are available, many of which provide comprehensive guidance for building a mature governance, risk, and compliance program. The appropriate framework will depend on factors such as organizational size and maturity, the nature and use of

AI systems, and the jurisdictions in which the organization operates. A very strong core framework is ISO/IEC 42001, which is the first globally recognized, dedicated AI governance standard designed for certification. It specifies requirements for establishing, implementing, and continually improving an AI management system to ensure the responsible, risk-based, and accountable development and use of AI. It can be scaled to an organization's size.

In addition, an organization can layer regional frameworks and map to existing laws based on its jurisdictional profile. For the United States, an organization should include the NIST AI Risk Management Framework. If your business is located in the European Union or provides goods and services in the EU, you must map compliance controls to the EU AI Act. In the Asia-Pacific region, depending on where you do business, there is Singapore's Model AI Governance Framework and AI Verify, Australia's Guidance for AI Adoption, AI Safety Standard and AI Ethics Principles, and there are several very strict laws in China that an organization must pay close attention to if it operates there. For other countries that do not have established laws or specific governance frameworks, the UNESCO/OECD principles provide a good baseline and then you should determine whether there are any country-specific rules in development. Because of the rapidly developing AI landscape, governance and legal teams need to stay on top of emerging laws in the countries, states, and/or provinces where they operate.

Because AI spans multiple disciplines, an effective AI governance framework must account for a range of applicable legal and regulatory regimes, including sector-specific laws (e.g., healthcare and financial services), privacy requirements, and consumer protection obligations.

These considerations should be integrated into a framework that is tailored to the organization's specific operations, risk profile, and use of AI.

A small or medium-sized enterprise using limited AI—such as for internal process automation or basic chatbots—generally will not require the same level of governance as large, global organizations deploying AI extensively across operations, operating in high-risk AI areas, and/or leveraging significant data sets. The European Union and the Organisation for Economic Co-operation and Development (OECD) recognize that governance expectations should scale with risk and organizational maturity.

Nevertheless, any organization, regardless of size, that develops or deploys high-risk AI systems (e.g., in healthcare or credit scoring) is subject to significant compliance obligations. Under the EU AI Act, violations involving prohibited AI practices (such as social scoring, manipulative systems, or certain biometric uses) or non-compliance with requirements pertaining to high-risk AI may result in fines of up to €35 million or 7% of global annual turnover, whichever is higher. Being able to demonstrate that your organization has a strong AI governance framework should mitigate the penalties imposed and the amount of any fines in the event of non-compliance with legal requirements.

To ensure that your AI governance and risk framework is applied efficiently, it is critical to maintain an inventory or registry of all AI systems used, whether developed in-house or provided by third parties. This inventory will need to be continuously updated to stay current. An organization can then map and assign a risk rating to all AI systems used to determine where the governance priorities need to be focused on. Finally, a strong governance framework will address all phases of the AI lifecycle. An AI system lifecycle typically involves: planning and design; collection and processing of data; building the AI model(s) and/or adapting existing model(s) to specific tasks; testing, evaluating, verifying and validating the models; using and deploying the models; operating and monitoring the models' outputs on an ongoing basis; and retiring/decommissioning models that are no longer in use. These phases are iterative and not necessarily sequential.

A mature AI governance framework will document the AI systems and risks using an AI Impact Assessment, in a similar manner that a Data Protection Impact Assessment (DPIA) is used for privacy governance.

What steps should an organization make to minimize the chances of an AI system doing harm – through bias, misinformation, action, etc.?

Organizations first need to map and rate the risks that can arise when deploying their AI systems. For AI systems that could potentially lead to harm, more focused oversight and controls should be put in place. Types of harms to look out for include:

- Privacy violations
- Adverse impacts on vulnerable individuals such as members of certain groups or populations
- Bias and discrimination (e.g., in hiring decisions or admissions)
- Poor training data leading to incorrect outputs
- Hallucinations that could impact outcomes such as use of AI for legal purposes
- Misinformation and manipulation such as deepfakes and disinformation about elections
- Safety risks (e.g., medical devices, critical infrastructure)
- Intellectual property violations

A strong way to mitigate these types of risks is to include human oversight of the AI output and training data (“human-in-the-loop”) and, where feasible, impacted stakeholders in the oversight of the AI systems, and continually monitor and test the AI systems' outputs. An organization should also track the types of harms that have been identified through the U.S. Federal Trade Commission enforcement actions and the OECD AI Incidents and Hazards Monitor.

A mature AI governance framework will document the AI systems and risks using an AI Impact Assessment, in a similar manner that a Data

Protection Impact Assessment (DPIA) is used for privacy governance. Under the EU AI Act, a Fundamental Rights Impact Assessment (FRIA) and conformity assessments are required for certain high-risk AI systems. Bottom line: creating an impact assessment with stakeholder review and ongoing re-assessments is the best way to minimize harms caused by AI systems.

Even where an organization's AI system does not inherently present a high risk of harm, the governance team must account for the risk of malicious exploitation. Organizations should implement appropriate security monitoring and controls to defend against threats such as prompt injection, jailbreak attacks, and data loss. Security functions are therefore a critical component of an effective AI governance framework to reduce risk of harm. This requires ongoing awareness of emerging AI threat vectors and continuous evaluation to ensure that appropriate technical and operational safeguards remain in place.

AI capabilities are evolving very quickly. What paths and mechanisms can facilitate making governance adaptive to these changes?

Keeping pace with AI development across an organization is a persistent challenge for governance professionals. To address this, AI system developers should be integrated into the governance function and required to obtain approval before introducing new AI systems or making material changes to existing ones—such as updates to datasets, testing methodologies, or intended use. Maintaining a centralized AI inventory or registry is equally essential, supported by periodic reviews to ensure documentation remains accurate, complete, and current.

Given the pace at which AI systems evolve, this coordination between the AI model developers and the rest of the AI governance team must be ongoing rather than periodic. Best practices include formal change management procedures, risk assessments, and approval workflows for new models and material updates, as well as continuous monitoring to detect performance drift, emerging risks, and potential misuse. Regular cross-functional reviews, coupled with clear documentation and audit trails, will allow the governance team to adapt controls as AI capabilities and threat landscapes evolve.

Should users be made aware that they are viewing AI-generated content?

Yes, users should be made aware that they are viewing AI-generated content especially in jurisdictions where this is required or where it could be considered a deceptive practice not to do so. For example, deceptive practices could arise if the content misleads someone into believing they are dealing with a human, or that they are receiving an actual consumer product endorsement, or that the content is based on human expertise (e.g., medical advice).

In jurisdictions with more developed AI regulatory regimes, including in the EU and China, organizations are generally required to inform individuals when they are interacting with AI (e.g., chatbots). In the United States, while there is no comprehensive federal AI law, the Federal Trade Commission prohibits deceptive practices. In addition, several state laws require some form of transparency, most notably California, Colorado, and Utah, and there are other states where such legislation is under consideration. Given the evolving and fragmented U.S. and global legal landscape, providing clear AI disclosures is a prudent practice.

In the B2B context, having a website page or trust center is a good way to provide the transparency. A business could include a statement that AI is used in the product (and where), the high-level purpose (e.g., threat detection, automation, recommendations), whether customer data is used for training, and a disclaimer about any limitations.

Clear, proactive AI disclosure builds trust and confidence with end users even if not legally mandated in all places. When companies align how they use and disclose AI with both user expectations and applicable legal requirements, they can strengthen customer relationships and build consumer trust.



Anjali Hansen is Head of Compliance and AI Governance at Onyx Security and formerly served as General Counsel of Noname Security. She is also the co-founder of AVA Compliance Solutions, which advises technology startups on AI governance, cybersecurity, privacy, data protection, commercial contracting, and regulatory compliance matters. Ms. Hansen has more than 25 years of legal and policy experience spanning cybersecurity, international trade, intellectual property, global privacy compliance, and emerging AI governance matters.

Prior to her work with technology startup companies, Ms. Hansen served as International Privacy Counsel at Verizon, where she led privacy compliance initiatives across Europe, Asia-Pacific, Latin America, the Middle East, Africa, and Canada. She previously served as Deputy General Counsel for the Council of Better Business Bureaus and held legal and policy positions with the United States Trade Representative and the U.S. International Trade Commission.

Ms. Hansen received her J.D. from Georgetown University Law Center in Washington, D.C. and a Certificat d'Etudes Politiques from the Institut d'Études Politiques in Paris, France. She holds bar memberships in Washington, D.C. and Maryland and is certified as an Artificial Intelligence Governance Professional (AIGP) by the International Association of Privacy Professionals. She also serves on the Advisory Board of the American Council for Ethical AI.

POWER MEETS PLACE:

OXFORD & LONDON LEADERSHIP PROGRAM

ENGAGE WITH GLOBAL ENERGY AND THE LAW OF THE SEA LEADERS AT THE HEART OF ACADEMIC EXCELLENCE

A UNIQUE
5-DAY
EXECUTIVE
EXPERIENCE



GLOBAL
EXPERTS



ACADEMIC
EXCELLENCE



INTERNATIONAL
INSTITUTIONS



NETWORK
INFLUENCE
IMPACT

HIGH-LEVEL DIALOGUE. INSTITUTIONAL ACCESS. CULTURAL IMMERSION.

1

MONDAY OXFORD, UK

Historic Oxford
College Setting



FULL-DAY LEADERSHIP DIALOGUE

- Arrival & welcome breakfast
- Opening remarks and program briefing
- Closed-door sessions (Chatham House Rules) with the former Sec-General of OPEC
- 1st Session (90 min)
- Coffee/tea break (15 min)
- 2nd Session (90 min)
- Networking lunch
- 3rd Session (90 min)
- Afternoon break (15 min)
- 4th Session (90 min)



EVENING:
RECEPTION & FORMAL
DINNER IN OXFORD

2

TUESDAY OXFORD, UK

INDUSTRY &
INNOVATION DAY



- Visits to leading UK-based energy, technology, or sustainability-focused companies (final selection announced prior to program, based on the preferences of participants)
- Executive briefings and site tours

AFTERNOON: GUIDED SIGHTSEEING IN OXFORD

- Historic colleges
- Libraries and iconic landmarks
- Optional punting experience on the River Cherwell



EVENING:
FREE TIME OR INFORMAL
GROUP DINNER

3

WEDNESDAY OXFORD, UK

ACADEMIC &
POLICY INSIGHTS



- Morning lectures and expert panels featuring:
 - Regional/Global energy policy experts (OPEC, IEA, etc.)
 - UNCLOS/ITLOS specialists
 - Representatives or analysts connected to the International Energy Agency

TOPICS MAY INCLUDE:

- Global energy transition
- Energy security, protection of the vital physical and digital infrastructure
- Eastern Europe and Western Asia conflicts and the Law of the Sea

AFTERNOON: FORECASTS / DISCUSSION / CASE-STUDY WORKSHOP

EVENING: COLLEGE
DINNER OR NETWORKING
RECEPTION

4

THURSDAY OXFORD, UK

CULTURE &
REFLECTION



- Morning at leisure or optional seminars
- Extended sightseeing program:
 - Colleges, chapels, and museums
 - Walking tour through historic Oxford

AFTERNOON:

- Program reflection session
- Certificate ceremony



EVENING:
FAREWELL DINNER

5

FRIDAY LONDON, UK

INTERNATIONAL
INSTITUTIONS VISIT -
LAW OF THE SEA



- Transfer to London
- Visit to the headquarters of the International Maritime Organization
 - Institutional briefing
 - Q&A session with officials (subject to availability)

PROGRAM CONCLUDES
WITH DEPARTURE
(AFTER LUNCH),
OR OPTIONAL EXTENDED
STAY IN LONDON



PROGRAM HIGHLIGHTS

- Direct engagement with global Energy and the Law of the Sea experts
- Academic excellence in Oxford's historic environment
- Exposure to international organizations and policymaking
- Balanced mix of policy, business, and cultural experiences

EXTEND OR ALTERNATE YOUR EXPERIENCE (2-DAY OPTIONS)



PARIS

(TRAIN FROM GENEVA)
World Bank, OECD,
IEA, UNESCO +
former or serving head
of one of these
Organisations



STRASBOURG

(TRAIN FROM GENEVA)
Council of Europe,
Court + former
or serving head
of one of these
Organisations



BRUSSELS

(TRAIN / PLANE
FROM GENEVA)
EU Commission,
EU Parliament,
NATO HQ + former
or serving head
of one of these
Organisations



GROUP SIZE
20-40
PERSONS



REGIME
TOP
EXCLUSIVITY



REWARD
EU
CERTIFICATE



FEE: COVERS ACCOMMODATION,
MEALS, TICKETS & TAXES,
EXCLUDES AIRFARE TO/FROM
AND BETWEEN DESTINATIONS



GROUP DISCOUNTS
GROUP OF 3
EARLY BIRDS GET
ONE PLACE FOR FREE



PROGRAM DURATION
WEEK-LONG
(WEEKENDS EXTENSION POSSIBLE)



DATE
06-10 JULY 2026

Early bird (until the end of May): CHF 2,800
Standard price: 3,600; Late Bird: CHF 4,400
(all prices include the single rooms in the historic settings
and all daily meals)

Best Practices for Responsible AI Development

Interview with Jenil Shah
American Council for Ethical AI, United States

How should AI developers identify and mitigate bias in constructing and maintaining platforms?

The most expensive mistakes come from teams treating bias as a pre-launch checklist instead of a property that has to be measured continuously during and after launch. Bias in production AI is rarely a single defect you can patch. It is a property of a system. From data collection to model training to prediction signals, bias has to be addressed at each stage, and all of it has to work cohesively.

One architectural decision that helps build mature AI platforms: “bias mitigation belongs at the platform level, not the model level”. Six months from now you will have a more advanced model. The platform has to be agnostic of which model is running. If you build it into the model, you rebuild it every time the model rotates. If you build it into the platform, it survives the model.

And the one most relevant as AI platforms shift toward context-driven architectures: bias lives in the context. In modern AI systems (retrieval-augmented, agentic, prompt-driven) the context window is where the model’s world is constructed. The question worth asking is: is your context biased? Can you identify it? If you can, prune it. Track the lineage of that bias. Is it coming from customer context, or from retrieval systems your AI platform pulls from? Understanding where bias enters the context is half the battle. Pruning it before it reaches the model is the other half.

Two techniques that have proven effective here are semantic filtering and context reweighting, both purpose-built for working with context in

modern AI systems. Semantic filtering scores retrieved chunks for bias signals before they enter the context window, allowing teams to prune or deprioritize problematic content before the model ever sees it. Context reweighting goes a step further: rather than ranking retrieved content purely by similarity, it surfaces diverse perspectives proportionally, so the model reasons over a more balanced view of the world and not just the loudest signal in the retrieval index.

How do you build responsible AI practices into fast-moving engineering teams without sacrificing delivery velocity?

Responsible AI in fast-moving teams is not a gate you install at the end of the pipeline. Think of it more like a knob you tune. The question is not whether to have it, but where to turn it up, where to turn it down, and where to hand it to a human entirely.

The first exercise worth doing as a team is defining your isolation levels: what AI can touch, and what it cannot. This sounds simple but most teams skip it, and the absence of that line is where velocity and responsibility start to fight each other. Once you have drawn it, the path forward becomes clear. Automate as aggressively as possible inside the “can touch” boundary, and reserve human oversight for the inflection

Responsible AI in fast-moving teams is not a gate you install at the end of the pipeline. Think of it more like a knob you tune.

points that sit at or near the line. Those inflection points are not arbitrary. They are the decisions where a wrong call compounds, where the feedback loop is slow, or where the downstream harm is hard to reverse. That is where judgment belongs, and judgment does not scale as an automated step.

The second shift is making the safe path the easy path. If responsible AI lives in a separate portal, a separate review committee, and a separate meeting cadence, engineers will route around it. Not out of bad intent, just out of deadline pressure. It has to live inside the workflow they already trust: the CI/CD pipeline, the code review, the merge checklist. A bias evaluation that runs as a single command in the build is one the team will actually run.

The framing that unlocks all of this culturally is to stop positioning responsible AI as a tradeoff against velocity and start treating it the way mature teams treat reliability: a property the system is required to have, not a tax on shipping.

What are the challenges and risks in developing low-latency AI models?

Low-latency AI is going to be one of the most valuable engineering skillsets of the next few years. Running LLMs fast was already hard. Add agentic behavior, tool calls, and dynamic context management, and the problem gets significantly harder. But before diving into solutions, it helps to split the challenge into two distinct buckets, because the skills and the stakes are different.

The first bucket is building low-latency AI platforms, and this is where the majority of companies sit today. Most enterprises are not training or fine-tuning their own models. They are consuming cloud-hosted models and trying to make that experience fast, reliable, and cost-efficient. The real challenges here are operational: selecting the right model for a given task based on the cost, accuracy, and latency tradeoff; building LLM gateways that can reroute requests intelligently based on load or model availability; using context caching efficiently so you are not recomputing the same tokens on every call; and reusing stored context across sessions

to reduce both latency and cost. These are AI platform engineering problems, and most companies are working through the same set of them right now.

The second bucket is actually building low-latency models, and this is a more specialized space. It lives primarily in AI labs and teams doing serious fine-tuning work. The engineers here need deep knowledge of GPU inference, architectures like vLLM, paged attention for efficient memory management, quantization, and batching strategies. Small models in the 1B to 8B range can now produce reasonable first-token latencies on a single GPU, but getting there requires skills that go well beyond typical ML engineering.

When should organizations choose private AI or small language models over large cloud-hosted models, and what are the ethical and operational tradeoffs of each?

The right choice between private AI and cloud-hosted models is rarely a technology decision. It is a sector decision first. Healthcare organizations default to private AI not because it is technically superior but because the data cannot leave the perimeter. Consumer AI sits on cloud-hosted models because the scale and the cost economics point that way. Sector is the first filter, and for many organizations it is the only filter they need. For everyone else, several factors come into play: latency requirements, cost, traffic patterns, call volume, and the accuracy and hallucination tolerance of the use case. A high-volume, low-complexity task on sensitive data points toward a small private model. A low-volume, high-complexity task with no data sensitivity points toward a frontier cloud model. The mature architecture is not picking one or the other. It is routing, with different requests going to different models based on what each one actually needs.

Private AI and local LLMs are also getting meaningfully better. Federated learning approaches are maturing, consumer hardware keeps improving, and small models are closing the capability gap faster than most people expected. The trajectory points toward more local and private LLM adoption over the next few years, not less. The operational tradeoffs are where most organizations underestimate

the delta. Using a cloudhosted model is, at its core, a distributed systems and software engineering problem. The skills are familiar: API integration, latency management, cost optimization, failover. A strong software engineer adopts this stack quickly. Deploying and operating a private LLM is a different discipline entirely. It requires quantization, distillation, GPU inference management, and model lifecycle operations, skills that sit squarely in ML and AI engineering.

The ethical tradeoffs cut both ways. Private models offer real privacy gains, data stays in the perimeter, telemetry is contained, third-party logging risks disappear. But they also forfeit the safety infrastructure that frontier labs have invested heavily in.

** Views expressed here are solely of this author and do not represent those of any organization*



Jenil Shah is a Software Engineering Manager specializing in recommendation systems, personalization, and generative AI applications. He has over a decade of experience working in different organizations focusing on applied machine learning and AI.

Making Chile a Global Leader in Artificial Intelligence

Interview with Professor Álvaro Soto
Pontifical Catholic University of Chile (PUC), Chile

What first drew you to AI and sustainable technology, and what keeps you motivated in leading CENIA's projects today?

What first drew me to AI was a fundamental scientific question: is it possible to emulate human intelligence? This question, captured by the Turing Test, represents one of the most ambitious challenges in science: building machines capable of exhibiting advanced intelligent behaviors. Today, however, as this question begins to be answered in practice, I believe it is necessary to shift the focus. The original scientific curiosity must be now complemented by a broader question: how can AI enhance human capabilities? Rather than aiming to match or replace us, AI should help amplify what humans can do.

At CENIA, we embrace this dual goal. CENIA's public mission is explicitly framed around putting AI at the service of people. On one hand, we continue to push the scientific frontier of AI, advancing our understanding of intelligence and developing state-of-the-art systems. On the other, we are equally committed to ensuring that these advances translate into real impact: technologies that expand people's opportunities, support their development, and contribute to a more inclusive and sustainable society.

CENIA aims to position Chile as a global leader in AI committed to ethical and sustainable technology. If you could remove one barrier tomorrow that would most accelerate the country's progress, what would it be?

If I could remove one barrier tomorrow, it would be the shortage of shared, strategic computing infrastructure. Talent matters, data matters, regulation matters, but without reliable compute, none of that can scale.

AI is becoming a basic capability, much like electricity or connectivity. A country that depends entirely on external infrastructure will always face a structural limitation in research, innovation, deployment, and sovereignty. Chile already has important pieces moving in the right direction as the AI Zero initiative in Arica. But the real acceleration would come from treating compute as a national priority. In that sense, the challenge is comparable to Chile's electrification process in the late nineteenth century. At the time, building a national electrical infrastructure was neither obvious nor urgent, yet it became a foundational project for the country's development. AI will transform productivity across multiple sectors and reshape our society. In twenty years, the question will not be whether Chile should have fostered public-private partnerships to develop AI infrastructure, but how well it did so and how early it began.

Many carbon credits and emissions reductions are now being represented as digital tokens on blockchains. This raises both opportunities and concerns around verification, energy use and environmental impact. How do you think Chile should approach the idea of taxing carbon-fuelled tokens?

While verification and integrity are important to ensure that tokens truly represent real emissions reductions, I believe the central issue is

Chile's natural advantages, abundant solar energy in the north and wind in the south, allow it to think about digital infrastructure in a fundamentally different way.

more fundamental: how we generate and use energy in the first place. The environmental impact of digital technologies, including tokenized systems, ultimately depends on the energy matrix that supports them.

In that sense, Chile is in a strong position. The country has made significant progress in transforming its energy matrix toward renewables. Chile's natural advantages, abundant solar energy in the north and wind in the south, allow it to think about digital infrastructure in a fundamentally different way.

The priority, therefore, should not be to focus narrowly on taxing specific digital representations such as tokens, but to continue accelerating the transition toward clean energy and aligning digital infrastructure with it. A good example is the AI Zero initiative in Arica, developed by the University of Tarapacá and CENIA, which leverages abundant solar energy and high-speed connectivity to run compute-intensive workloads where energy is clean and efficient, and then transmit the results. Under this model, we are effectively transmitting computation rather than energy from regions where energy is abundant and clean to centers of demand. This is a more efficient and sustainable paradigm for digital infrastructure.

Projects like AI Zero being developed in Arica by the University of Tarapacá and CENIA aim to create a zero-emissions AI infrastructure utilising solar power and the National Optical Fiber network. What practical lessons can this model provide for designing sustainable AI infrastructures in energy-intensive contexts?

The first lesson is that sustainable AI infrastructure should be designed around geography, not against it. AI Zero is powerful because it connects three assets that are rarely planned together: renewable energy, high-capacity computing, and high-speed digital connectivity. Arica offers excellent solar conditions, and the National Optical Fiber network makes it possible to compute where clean energy is abundant and then move results efficiently across the country. Part of this high-speed connectivity is already in place, originally developed to support the transmission of data from astronomical observatories in northern Chile, where the country is a global leader.

The second lesson is that sustainability cannot be reduced to buying cleaner electricity. It must be built into the architecture from the beginning. A sustainable AI cluster should ask where energy is cleanest, where it is cheapest, how resilient the system is, and which uses generate the highest social value per unit of compute.

The third lesson is that public infrastructure can be built in a way that aligns development, energy transition, and regional inclusion. That demonstration effect may be as important as the hardware itself.

Looking ahead, how do you envision Chile shaping the global Green AI paradigm in the next decade, and to what extent does this create opportunities for Latin America, and the Global South more broadly, to define its own pathway in sustainable AI?

I believe Chile can help shape Green AI by showing that sustainability, openness, and technological ambition do not need to be in tension. We can integrate clean energy, public-interest research, credible institutions, and regional collaboration into a different model of AI development. If we do this well, Chile can become a reference point for how to build AI infrastructure and models that are not only competitive, but also accountable and environmentally intelligent.

However, the challenge goes beyond Chile. One of the key lessons from recent advances in AI is that progress has been driven by scale: scale in data, in compute, and in model size. No single country in the region can achieve the scale required by modern AI on its own, but Latin America, as a whole, can. Initiatives such as LatamGPT are moving in that direction: a collaborative model, built with regional capabilities and trained on a corpus that reflects our own identity.

Looking ahead, scaling toward a distributed training of LatamGPT, leveraging compute centers across the region, could materialize a long-standing aspiration for regional integration and shape a different way of building AI: open, collaborative, and grounded in local identity. This would not only increase scale, but also align infrastructure with the region's energy advantages, enabling a more sustainable approach to large-scale AI.

More broadly, this creates an opportunity for Latin America, and the Global South, to define their own pathway. This is not about isolated technological independence, but about building sufficient local capacity to participate with a stronger voice in the global AI ecosystem. By combining clean energy, open collaboration, and shared infrastructure, the region can move from being a consumer of imported technologies to a co-author of a more sustainable and inclusive AI paradigm.

Interview by Joe Cole



Álvaro Soto is CENIA Director (Chile), and Associate Professor in the Department of Computer Science at the Pontifical Catholic University of Chile (PUC). He holds a Ph.D. in Computer Science from Carnegie Mellon University (USA), with a specialization in Artificial Intelligence and Cognitive Robotics. He is currently an Associate Professor in the Department of Computer Science at the Pontifical Catholic University of Chile (PUC).

Throughout his academic career, he has conducted research internships at the Artificial Intelligence laboratories of Stanford and MIT, and has published more than 50 scientific papers in the field.

Since 2021, he has served as Director of the National Center for Artificial Intelligence (CENIA). He is co-founder of Zippedi, a company dedicated to the development of retail robots, with operations in Chile, Colombia, the USA, Germany, Australia, and Japan. He has received awards such as the Avonni National Innovation Award and the Ministry of Science's Best Technology-Based Entrepreneurship Award (2021).

His areas of expertise include artificial intelligence, cognitive robotics, and technology transfer applied to the development of industry solutions.

What Does Artificial Intelligence Know about Food and Eating?

Professor Joachim Allgaier
Fulda University of Applied Sciences, Germany

The relationship between humans and their food is historically contingent, yet it remains an existential issue. Without food, we die. This fundamental biological necessity has shaped human civilization from its earliest beginnings, yet the ways in which we produce, prepare, distribute, and consume food have never been static. Each era has introduced new technologies that transformed not only what we eat but how we understand, experience, and relate to food itself. From prehistoric stone blades that allowed early humans to cut meat and access new nutritional sources, to the agricultural revolution that enabled settled societies, to the industrialization of farming with its massive use of oil and fertilizers—technology has always been intertwined with food. Today, we stand at another inflection point: the latest addition to this lineage is not a physical tool, but a digital one: the pervasive integration of artificial intelligence (AI) into the very fabric of how we eat and feed ourselves. This shift has given rise to what scholars term "digital foodscapes"—complex, interconnected environments where food is planned, marketed, cooked, and consumed through screens, algorithms, and data streams¹.

The trajectory of food technology has moved from the physical to the virtual and has increasingly mixed both up. Where once a blade or a plow altered the material reality of food, digital technologies now alter the *meaning* and *experience* of food. Social media platforms like Instagram and TikTok have transformed eating into a performative act, where the visual presentation of a dish seems to count more than maybe its taste or nutritional value. The phenomenon of *mukbang*—livestreamed eating shows—turns consumption into entertainment, while food delivery apps algorithmically coordinate or increase our hunger, often prioritizing speed and convenience over nutrition. These are not merely tools; they

are environments that shape our desires, habits, and identities. Now, AI sits at the center of this ecosystem, promising a future of hyper-personalization while introducing profound risks to our cultural and psychological well-being.

The potential of AI in the realm of food and nutrition is undeniably vast. On the surface, the promise is one of optimization and precision. AI-driven "precision nutrition" apps can analyze an individual's genetic makeup, metabolic data, and lifestyle to generate bespoke meal plans that generic dietary guidelines cannot match. Large Language Models (LLMs) can act as tireless nutritionists, instantly synthesizing vast amounts of scientific literature to provide evidence-based advice on everything from micronutrient deficiencies to sustainable eating patterns. In the realm of public health, AI can monitor food security in real-time, predicting shortages and optimizing supply chains to reduce waste. Smart kitchen appliances, guided by AI, could automate the preparation of nutritious meals, making healthy eating accessible to time-poor families or those with limited culinary skills. The efficiency gains are staggering, and the potential to democratize access to high-quality nutritional information is a powerful argument for adoption.

However, this technological optimism must be tempered by a critical examination of the downsides, many of which strike at the heart of human culture and psychology. One of the most insidious effects of digital foodscapes is the way they curate and enforce specific ideals of the body. Social media feeds are saturated with images of "perfect" bodies alongside "perfect" meals, creating a feedback loop where food is judged not by its nourishment but by its ability to conform to aesthetic standards. These portrayals often promote unrealistic body images, linking

thinness or muscularity directly to moral virtue and dietary discipline. For vulnerable individuals, particularly adolescents, this constant exposure can threaten mental well-being, fostering body dysmorphia and contributing to the rise of eating disorders. The algorithm, designed to maximize engagement, often amplifies extreme content, pushing users toward communities that glorify restriction or disordered eating under the guise of "wellness."

Furthermore, the rise of AI threatens to erode the human element that is essential to effective nutritional counseling. While an AI chatbot can calculate calories and suggest recipes, it cannot offer empathy, intuition, or the nuanced understanding of a patient's emotional relationship with food. Nutrition is not merely a biological equation; it is deeply social and psychological. Effective dietary change often requires the trust, connection, and care that only a human professional can provide. A bot can tell you *what* to eat, but it cannot understand *why* you eat, nor can it navigate the complex emotional landscapes of trauma, stress, or cultural identity that often drive eating behaviors. As AI becomes more capable, there is a genuine risk that the profession of dietetics and nutrition counseling will be devalued, leading to job losses for empathic human experts.

This concern extends beyond the counseling sector to the food industry as a whole. The increasing application of AI in food production, logistics, and service threatens to displace a significant portion of the workforce. From automated farms and robotic kitchens to algorithmic management of delivery drivers, the efficiency of AI comes at the cost of human labor. The "gig economy" of food delivery is already characterized by precarious labor conditions, and AI promises to tighten the screws of algorithmic control, reducing workers to mere cogs in a machine optimized for profit or it will maybe replace them altogether by the use of autonomous delivery vehicles. If the food system becomes entirely automated, we risk

The increasing application of AI in food production, logistics, and service threatens to displace a significant portion of the workforce.

losing not only jobs but also the communal aspects of food preparation, eating, and service that have historically bound communities together.

Cultural homogenization is another looming threat. Algorithms tend to favor content that is visually striking and universally appealing, often at the expense of local, traditional, or niche cuisines. As AI recommends recipes and food trends, there is a danger of flattening culinary diversity into a standardized, "instagrammable" global menu. This may also extend to the agricultural sector. Chefs may begin to plate food specifically to satisfy algorithmic preferences rather than to honor culinary tradition or flavor profiles. This may lead to a loss of embodied knowledge—the sensory skills of smelling, tasting, and feeling ingredients that are passed down through generations and might also have an impact on how and what food we are producing and growing. When cooking becomes a matter of following coded instructions on a screen, the deep, intuitive connection between the cook and the food is severed.

Moreover, the concentration of power in the hands of a few tech giants controlling these digital foodscapes poses a significant risk to food sovereignty. These platforms are not neutral; they are driven by profit motives that often align with the interests of the food industry. AI can be used to deploy hyper-personalized marketing that targets vulnerable populations, promoting ultra-processed, unhealthy foods with personalized precision. The erosion of trust in scientific expertise is another consequence, as AI-generated misinformation may spread rapidly, often outpacing credible health advice. Safety concerns are also real; AI systems have been shown to fail in critical areas, such as identifying allergens or accounting for nutrient deficiencies in restrictive diets, with potentially life-threatening consequences.

The integration of AI into our foodscapes therefore is a double-edged sword. It offers the tantalizing prospect of a world where hunger is eliminated, nutrition is perfectly tailored, food systems are sustainable and food waste is minimized. Yet, it also threatens to commodify our most intimate relationship, turning food into data and eating into a performance. It risks replacing human connection with algorithmic efficiency, eroding cultural diversity, and exacerbating social inequalities. As we move forward, the challenge is not to reject technology, but to

ensure that it serves human needs rather than dictating them. We must advocate for a future where AI augments, rather than replaces, the human touch in nutrition; where digital tools are regulated to prevent exploitation; and where the joy, culture, and community of eating remain central to our existence.

Joachim Allgaier, Dipl.-Soz., Ph.D. is a sociologist and media and communications researcher. He is Professor for Communication and Digital Society at the Department of Nutritional, Food and Consumer Sciences of Fulda University of Applied Sciences. Previously he held positions as a senior researcher at the Chair of Technology and Society at the Human Technology Center (HumTec) of RWTH Aachen University (Germany); the Cologne Center for Ethics, Rights, Economics and Social Sciences of Health (ceres) at University of Cologne, Germany. Before that he was a senior scientist and deputy director of the Department of Science, Technology and Society Studies at Alpen-Adria-Universität in Klagenfurt (Austria), a postdoctoral researcher at the Helmholtz Research Center Jülich (Germany) and the Department of Science and Technology Studies at the University of Vienna (Austria), and also a media researcher the Open University (UK), where he was awarded a PhD in sociology.



In addition, he was honorary fellow at the School of Journalism and Mass Communication at the University of Wisconsin-Madison (USA), invited visiting researcher at the Department of Health and Social Studies at Østfold University College (Norway) and fellow at the International School of Science Journalism and Science Communication of the Ettore Majorana Centre in Erice (Italy). Joachim Allgaier studied sociology, psychology and intercultural communication at LMU Munich (Germany) and Science and Technology Studies at Maastricht University (Netherlands). He was elected as a member of the Global Young Academy from 2015-2020.

NOW AVAILABLE

Digital Wisdom



Searching for
Agency in the
Age of AI

Richard
Lachman

■ INTERNATIONAL AFFAIRS FORUM • 2026

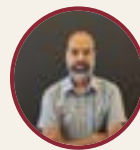
Digital *Wisdom*

Searching for Agency in the Age of AI

A principled framework for navigating artificial intelligence with intention and accountability
— for policymakers, educators, and engaged citizens navigating a transformed world.

"Technology shapes not just individual lives but our shared and overlapping futures in society. We have a collective responsibility to ensure technological development serves the common good."

— DIGITAL WISDOM



Prof. Richard Lachman

Toronto Metropolitan University
richlach.com

digitalwisdom.ca

Routledge • 20% off with code **26SME2**

PUBLISHED BY ROUTLEDGE

Coded Imperialism: Generative AI, Epistemic Violence, and the Erasure of African Indigenous Knowledge

Dr. Nouridin Melo
University of Maroua, Cameroon

The rapid integration of Generative AI (GenAI) into global digital systems is often celebrated as a democratizing force. Yet, for African societies, this technology functions as a high-speed vehicle for epistemic violence. By centralizing the production of "truth" within AI models trained on data dominated by Western interests, GenAI does more than reflect existing global imbalances; it actively erases the non-Western knowledge systems that form the foundation of African identity. This "coded imperialism" represents a structural shift in how power operates: it has moved from the physical extraction of natural resources to the digital colonization of the human mind, threatening to silence ancestral wisdom.

The Mechanism of Epistemic Erasure: Nigeria and Cameroon

Epistemic violence is the systematic dismissal of indigenous knowledge by presenting Western algorithmic results as the only "objective" reality. When GenAI models synthesize data for governance or social development, they rely almost exclusively on Western digital archives. In Nigeria and Cameroon, this leads to significant cognitive displacement. For example, when a researcher in Lagos or a community leader in Douala asks an AI for guidance on local governance, they are often met with frameworks tailored to Western policy rather than local, lived experience.

In Nigeria, the dominance of foreign-trained AI means that credit-scoring and public service algorithms often fail to understand the "informal sector", the backbone of the Nigerian economy, because the training data lacks the nuance of local economic realities. Similarly, in Cameroon, the

push for AI-driven risk management in banking prioritizes international compliance models over the specific social and economic textures of the Central African region. These systems effectively relegate indigenous knowledge, such as oral traditions and community-based management, to "folklore," while Western-coded science is treated as universal truth.

Coded imperialism leads to a dangerous narrowing of human knowledge. By favoring dominant cultural narratives, GenAI systems risk creating a monoculture of "common sense" that marginalizes non-Western histories.

The Homogenization of Global Thought

Coded imperialism leads to a dangerous narrowing of human knowledge. By favoring dominant cultural narratives, GenAI systems risk creating a monoculture of "common sense" that marginalizes non-Western histories. This hegemony undermines the autonomy of African communities. As traditional forms of expression are pushed aside by multinational tech giants, the ability for people in Abuja or Yaoundé to define their own futures is weakened. This is cognitive imperialism: the power to define knowledge is reserved for those who own the computers, while the colonized are treated as "knowledge objects", data points to be mined, but never the subjects authorized to explain their own world.

Towards an Epistemic Counter-Movement

To counteract this, African nations must move beyond simple "bias mitigation" and toward active epistemic sovereignty. This requires a three-tiered policy agenda aligned with the African Union's Agenda 2063: Decolonizing Training Data: Nigeria and Cameroon must prioritize the formal collection and digitization of local language archives and oral histories. These must become the primary training foundation for any AI models deployed within their borders.

Supporting Indigenous AI Research: Scaling regional AI research is essential. These labs should build locally grounded ethical frameworks rather than adopting Western-centric "universal" AI ethics that ignore local historical and social contexts.

Implementing Epistemic Audits: National governments must mandate that AI systems used in education, law, and government undergo "epistemic audits" to ensure they do not systematically erase or misrepresent indigenous viewpoints.

Generative AI is not a neutral evolution of technology; it is a product of power dynamics that, if left unchecked, will accelerate the loss of African intellectual heritage. By asserting epistemic sovereignty, African nations can reclaim the right to define their own futures. Failing to act will not only result in the loss of cultural heritage but will cement a digital dependence that undermines the continent's agency for generations. This is a fundamental requirement for African self-determination in the digital age.

Dr. Nouridin Melo is a scholar of economic and social history based in Yaounde, Cameroon, specializing in industrial transformation and the political economy of West Africa. His research explores how shifts in global digital and technological systems are creating new dependencies, and he works on strategies for regional states to assert sovereignty over their intellectual and industrial development.



Fertilizer, AI, and the New Politics of Market Access in Africa

Christopher Burke
WMC Africa, Uganda

The use of fertilizer in Africa has long been characterized by low application rates, high prices, weak distribution systems, poor extension support and exhausted soils. These challenges persist. World Bank data shows Africa remains the world's lowest-use region for inorganic fertilizer per hectare of cropland. The United Nations (UN) Food and Agriculture's (FAO) 2024 update on cropland nutrient balances warns that high nutrient-use efficiency levels in Africa often signal nutrient mining rather than healthy abundance. Many African farming systems are still taking more from the soil than they put back.

This is no longer the whole story. The real fertilizer gap in Africa is increasingly not only a nutrient gap, but a compliance infrastructure gap. What matters today is not only whether fertilizer reaches farms, but whether fertilizer-linked value chains can produce credible, auditable evidence about sourcing, distribution, environmental effects, and labor conditions. The European Commission's Corporate Sustainability Due Diligence Directive (CSDDD) that entered into force on 25 July 2024 is intended to ensure companies identify and address adverse human-rights and environmental impacts across their operations and global value chains. The current European Union (EU) timetable stipulates that member states will need to transpose the directive by July 2027 with staggered application beginning a year later and full application on 26 July 2029. The EU has simultaneously adopted a "stop-the-clock" measure postponing CSDDD implementation and the first wave of application by one year while advancing broader Omnibus simplification proposals intended to narrow and reduce the regulatory burden of the sustainability due-diligence regime.

Far from being a simple European legal matter, it represents a serious market-access issue for Africa. It also reflects a wider shift in how regulatory authority is increasingly exercised through markets rather than formal public regulation alone. European firms and the banks, insurers, traders and buyers that sit around them increasingly require evidence rather than assurances. The Commission explicitly frames the directive as a due-diligence duty covering the operations of European companies along with their subsidiaries and, where relevant, business partners in the chain of activities. The framework is intended to improve risk management, resilience, competitiveness, and access to finance. The practical consequence is that African fertilizer importers, blenders, commercial farms, input distributors, agro-processors, and exporters are entering a world in which access to finance depends not only on productivity, but on the ability to provide credible, verifiable evidence about how their operations and supply chains work.

This is where artificial intelligence becomes politically important. Not because it is fashionable, and not because it offers some miracle cure for African agriculture, but because it can reduce the cost of making fragmented rural systems easier to document, monitor, and verify. The biggest challenge in African input markets is not complete ignorance,

[AI] can reduce the cost of making fragmented rural systems easier to document, monitor, and verify.

but scattered information. Soil conditions vary across short distances. Fertilizer recommendations are still often generic. Supplier records may be incomplete. Transport and blending chains are opaque. Emissions data are patchy. Labor risks are hard to monitor across dispersed commercial networks. AI and related digital systems can help organize this complexity by combining soil maps, transaction records, geospatial data, logistics trails, agronomic histories, and anomaly detection into usable evidence. AI is highly relevant to fertilizer. It is less a story of robotics than of verification.

There are already concrete signs of this shift on the ground. The International Fertilizer Development Center's (IFDC) Optimizing Fertilizer Recommendations in Africa platform currently serves 67 prime agricultural agro-ecological zones across 13 Sub-Saharan African countries. Drawing on more than 6,200 geo-referenced crop-nutrient response functions, IFDC uses linear optimization to generate fertilizer recommendations tailored to crop choice, land allocation, fertilizer prices, grain values, and overall budget constraints. That is not science fiction. It is the slow construction of a digital agronomic infrastructure. A government-backed program launched in Sierra Leone in February 2026 is developing a digital National Soil Information System to shift from generalized fertilizer use toward data-driven, site-specific nutrient management through field soil sampling, updated national soil maps, and stronger locally managed technical capacity. These examples demonstrate compliance infrastructure in practice with soil intelligence, traceability, standardization, and local technical institutions able to produce reliable evidence. They also show how governance functions are increasingly being carried by technical, financial, and informational systems rather than law alone.

The stakes are larger than agronomy. Fertilizer now sits at the intersection of food security, environmental performance, trade exposure, and financial risk. FAO's nutrient-balance work highlights the double challenge clearly. Insufficient nutrient use can reduce soil fertility, while excessive or poorly managed nutrient use can contribute to runoff, soil and water pollution, and greenhouse-gas emissions. The choice between "more fertilizer" and "greener agriculture" is no longer relevant. Africa needs both higher productivity and better proof of responsible use. Proof is becoming a condition of participation in premium markets and in

systems of exchange increasingly organized through risk management, disclosure, and verification. Firms that cannot demonstrate sufficient knowledge about soils, sourcing, environmental impact, and due diligence will increasingly be treated as riskier than companies that can—even when their agronomic performance is similar.

The new politics of fertilizer is really the politics of standards and the authority exercised through them. In a more fragmented international system, authority is often exercised less through classic treaty law and more through disclosure duties, risk screens, procurement criteria, audit trails, and compliance interfaces. The European Commission suggests the CSDDD could become a new global standard for mandatory environmental and human-rights due diligence. Formal regional regulation will shape the behavior of actors far beyond Europe. The risk for Africa is not simply overregulation. It is asymmetry. If African governments, companies, and producer organizations do not build their own data and verification systems, they will be judged through external metrics they did not design and cannot easily contest.

This should not be read as an argument against sustainability regulation. It is an argument for strategic capacity. The right response is not to denounce the CSDDD or to pretend that AI alone will solve African agriculture. It is to invest in the missing layers between fertilizer and markets comprising soil laboratories, digital soil mapping, supplier registration, interoperable traceability tools, better emissions accounting, stronger extension systems, and domestic regulatory capacity that can convert local realities into credible evidence. The Sierra Leone example is instructive because it ties digital mapping to national institutions, training, and local management rather than treating data as something extracted by outsiders.

Africa needs more fertilizer, but nutrients alone will not secure market access in the emerging green economy. Access to investment, premium buyers, and trusted commercial partnerships will increasingly depend on the ability of value chains to provide clear, verifiable evidence of how fertilizer is sourced, recommended, applied, and monitored. The fertilizer revolution in Africa will not be defined by volume alone. It will be shaped by how African agriculture builds the compliance infrastructure necessary to remain competitive in a world where sustainability has become a

market gatekeeper and where AI is becoming an essential tool to keep that gate open as the establishment of rules, monitoring, and commercial discipline increasingly move through market-embedded systems rather than formal state command.



Christopher Burke is a senior advisor at WMC Africa, a communications and advisory agency located in Kampala, Uganda. With over 30 years of experience, he has worked extensively on social, political and economic development issues focused on governance, agriculture, environment issues, policy formulation, communications, advocacy, extractives, conflict transformation, international relations and peace-building in Asia and Africa.

The Invisible Infrastructure War: AI Cloud Systems and Green Energy Grids in Africa as the Next U.S.–China Battleground

Shannon Roxborough
South Africa

The competition between the United States and China in Africa is increasingly discussed in terms of what is visible: ports, roads, diplomatic summits, military advisers. What the Senate Foreign Relations Committee and House Select Committee on the Chinese Communist Party are missing is the infrastructure layer—data centers, energy grids, financing architectures—that will determine who actually controls Africa's economic future. This is where the next phase of U.S.–China competition is already being decided.

Compute Requires Energy. Energy Requires Capital. Africa Has Neither to Spare.

Artificial intelligence does not run on code alone. It runs on compute infrastructure: server farms, cooling systems, the unglamorous physical plant of the digital economy. AI adoption is accelerating across African markets in fintech, agriculture, logistics, and healthcare. The question of where that infrastructure gets built, by whom, and on what terms is fast becoming one of the defining strategic questions on the continent.

The dependency sequence is straightforward, and both Washington and Beijing understand it. AI adoption requires compute. Compute requires energy. Energy requires grid modernization. Grid modernization requires capital and hardware that most African states cannot self-finance. Roughly 600 million people on the continent still lack reliable power, according to the International Energy Agency (IEA). Grid expansion is ongoing but uneven. The energy base needed to support AI infrastructure at any meaningful scale simply does not exist in most markets. Whoever

positions at the front of that sequence controls what follows. Both China and the United States understand this sequence. Both are moving to position themselves at the point where it begins.

China's Hardware Ecosystem and the Logic of Lock-In

China's approach is routinely characterized as debt-trap diplomacy. That framing is too flat. It misses the strategy. What Chinese firms have consistently pursued is not simply financing leverage but standards capture — the embedding of Chinese hardware, software protocols, and operational systems into foundational infrastructure in ways that create durable vendor dependency.

In the energy sector, this means solar panels, grid management software, and battery storage systems manufactured by firms such as JinkoSolar, Huawei, and CATL, installed across multiple African markets with Chinese financing attached. Once a grid management system is operational, replacing it is not merely expensive—it requires wholesale infrastructure redesign. The lock-in is not contractual. It is technical.

The same logic applies to AI infrastructure. Chinese hyperscalers and cloud providers—operating through partnerships with African telecoms and development finance vehicles—are establishing early positions in data center development across East and West Africa. Early positioning in cloud infrastructure carries compounding advantages: data residency, application layer dependencies, and the gradual standardization of developer ecosystems around Chinese platforms.

The United States, characteristically, arrived later and with more conditions attached.

Private capital is moving into African AI and energy infrastructure, but without strategic coherence.

Washington's Gap Between Ambition and Architecture

The Biden administration's response — the Partnership for Global Infrastructure and Investment, the Lobito Corridor project, the broader rebranding of U.S. development finance as a strategic tool — represented a genuine attempt to compete on infrastructure. The scale of commitment, measured against Chinese deployment, remained limited.

More structurally, the institutional architecture available to execute on that ambition operates without the transaction-level coordination that Chinese state actors take for granted. The Lobito Corridor illustrates the gap. China can bundle sovereign financing, EPC contractor selection, and equipment procurement into a single coordinated package through policy banks like China Development Bank or Exim Bank of China. The U.S. response required threading commitments across DFC, EXIM, USTDA, State, and Commerce with no single authority able to close on all dimensions. That is not a bureaucratic accident. It reflects the absence of a state-directed industrial policy mechanism that China deploys as a matter of course.

Private capital is moving into African AI and energy infrastructure, but without strategic coherence. U.S. firms enter deals after Chinese counterparts have already shaped terms, relationships, and technical standards. They arrive with due diligence requirements calibrated for markets that move faster than those requirements allow. The result is not absence but systematic lateness—and lateness in infrastructure competition compounds.

There is also a financing architecture dimension that Washington has not adequately addressed. China's development banks operate with risk tolerances and approval timelines that allow them to close deals in markets where Western institutions hesitate. This is not a bug in the system. It is the mechanism through which technical standards and vendor dependencies accumulate.

Africa at the Intersection

What makes Africa specifically consequential in this competition is that the infrastructure is still being built. In markets where grids are mature and cloud ecosystems are established, competition is about displacing incumbents—expensive, slow, politically contested. In Africa, the competition is about first installation. The firm or country that builds the transformer, lays the fiber, or deploys the cloud node is not merely winning a contract. It is setting the technical standard against which all future investment will be interoperable — or not.

This is the logic of the invisible infrastructure war, and it is playing out right now in decisions that rarely make international headlines: which grid management software a utility in Zambia procures, which cloud provider a Nigerian fintech scales on, which financing vehicle an Ethiopian data center operator accepts.

These decisions are being made by African governments and businesses operating under real capital constraints, with incomplete information about their long-term strategic implications, and without meaningful Western alternatives in many cases. The geopolitical framing in Washington and Brussels has not translated into on-the-ground competitive positioning at the transaction level.

The Standards War Nobody Is Talking About

Underlying all of this is a technical standards competition that receives almost no attention in policy discussions focused on tariffs, diplomacy, and military posture. Infrastructure interoperability is determined by standards — the protocols that govern how energy grids communicate, how data centers connect to cloud networks, how AI systems

authenticate and process. The organization that sets the standard for African grid management software does not merely win market share. It determines who can compete in that market at all.

China has been systematically pursuing standards influence through the International Telecommunication Union and bilateral technical cooperation agreements. Huawei alone has submitted more technical proposals to ITU Study Group 20 — the body that governs IoT and smart city standards — than any other single entity, effectively shaping the frameworks African utilities and municipalities are now adopting. Ethiopia, Kenya, and Egypt have each signed technology cooperation agreements that reference Chinese technical standards in grid management and telecoms. The United States has been largely absent from this work, treating standards bodies as bureaucratic backwaters rather than strategic terrain. The United States has been largely absent from this work, treating standards bodies as bureaucratic backwaters rather than strategic terrain.

The consequence is that even when U.S. firms eventually enter African AI and energy markets, they may find themselves operating on standards architectures that their Chinese competitors helped design.

What a Strategic Response Requires

Competing effectively in this space does not require matching China dollar for dollar on development finance. It requires something that has so far been harder to produce: institutional coherence at the transaction level, risk tolerance calibrated to African market realities, and sustained engagement with technical standards processes that currently receive neither attention nor resources commensurate with their importance. It also requires an honest reckoning with the gap between U.S. strategic rhetoric on Africa and U.S. institutional capacity to execute on that rhetoric. The ambition has been stated repeatedly. The architecture to deliver it has not been built.

African governments and businesses are not passive recipients in this competition. They are making consequential decisions under real constraints. Those decisions are accumulating into infrastructure dependencies that will shape their economic and technological

trajectories for decades. The question is whether the United States will develop the institutional tools to be a credible alternative before those dependencies are locked in — or whether the invisible infrastructure war will be decided by default.

Shannon Roxborough is a Johannesburg-based American freelance journalist, independent geopolitical risk analyst, and founder of China Executive Intelligence Brief, specializing in the U.S.–China–Africa nexus. He has nearly four decades of experience covering China and Taiwan and three decades of engagement across Africa. He advises executives, multinational corporations, institutional investors, and global funds on strategic risk and decision-making at that intersection.

His writing has appeared in *Fortune*, *Money*, *The Asian Wall Street Journal*, and *Far Eastern Economic Review*, and his analysis has been cited in Barron's. He also ghostwrote the chapter on international investing in *The Profit Hunter: Beating the Bulls, Taming the Bears, and Slaughtering the Pigs* (Wiley, 2010).

Crisis, AI and strategy – Three certified programs from GAFG

Redefecting crisis, resilience and innovation through high-impact training solutions

In an era defined by polycrisis, where climate-induced disruptions, digital transformation and geopolitical volatility converge, institutions face a singular imperative: to build adaptive capacity without exhausting finite budgets. The conventional reflex is to frame advanced training as an operational cost. However, a paradigm shift is now both possible and necessary. Based on demonstrated outcomes and sustained pedagogical innovation, we propose an alternative framework: world-class, certified executive education not as expenditure, but as a net saving in risk mitigation, time efficiency and long-term resilience.

Three signature programs for 2026

1. Crisis Management and Crisis Communication Management (11 June 2026) Certified | Onsite/Online

A one-day intensive module addressing real-time command protocols, reputational risk, and cross-sectoral coordination. Designed for emergency response teams, public information officers, and executive leadership.

2. Analogue Retreat (24–28 August 2026) Onsite

A five-day residential program deliberately structured away from digital overload, focusing on strategic foresight, organisational silence and high-reliability decision-making. Ideal for senior teams requiring deep, undistracted work.

3. AI and Life Science – Flagship 8-Week Program (24 Sep – 12 Nov 2026) Certified | Hybrid

Our most comprehensive offering, integrating artificial intelligence applications in biomedical research, regulatory science, and translational pathways. Features over 30 international speakers drawn from 22 countries across five continents, and builds directly on our latest flagship proceedings which convened more than 100 participants from 48 nations.

Remuneration for partnership and participant placement has been clearly determined, with very satisfactory incentives for facilitating institutional enrolment. To secure your placement, access group rates and benefit from preferential conditions, we kindly invite you to express your interest at your earliest convenience. For inquiries and registration, please reply to the above email.

Join the programs with just an email: **office@future-governance.org**

Check GAFG page: **<https://future-governance.org/>**

How Artificial Intelligence Will Redefine Public Finance Systems

Ramil Abbsabov

George Mason University, United States

In the coming decade, artificial intelligence (AI) and automation will do more to reshape public finance systems than any reform agenda of the past half-century. Governments across the world whether federal or unitary, advanced or emerging are confronting the same reality: traditional budgeting and public financial management (PFM) frameworks were not designed for the velocity, volume, and variability of data that modern economies now generate. As AI-driven tools rapidly become ubiquitous in the private sector, citizens increasingly expect governments to deliver the same level of speed, precision, and transparency. The question is no longer whether governments should adopt AI, but how they can do so responsibly, equitably, and strategically.

The End of the Manual Budget Cycle

For decades, public budgets have been products of slow, manual, and negotiation-heavy processes. Line ministries prepare submissions, finance ministries consolidate them, and legislatures debate allocations often using spreadsheets and documents that look remarkably like those used in the 1990s. This static approach cannot keep pace with dynamic policy environments defined by frequent economic shocks, rapid urbanization, and climate risks.

AI promises to fundamentally alter this cycle. Machine learning algorithms can analyze thousands of data points from tax collections and social protection registries to climate models and procurement databases in real time. Instead of annual or semi-annual budget updates, governments could adopt “living budgets” that update continuously based on economic conditions. These systems would allow policymakers to simulate expenditure scenarios, forecast fiscal risks, and test the distributional

impacts of policy decisions before they are implemented.

Automation will also streamline the most time-consuming elements of budgeting. Document drafting, baseline estimation, inflation adjustments, and performance report generation, all of which currently consume thousands of hours annually can be automated with AI-driven tools. This is not just about efficiency; it frees up civil servants to engage in higher-value analysis, policy design, and stakeholder engagement.

Transforming Public Revenue Mobilization

Tax administrations stand to gain immensely from AI and automation, with the future of tax policy and administration shaped by three major shifts that promise to fundamentally modernize the way governments mobilize revenue and interact with taxpayers. First, predictive revenue forecasting is being revolutionized as AI moves beyond traditional macroeconomic projections and coarse sectoral estimates, instead integrating vast streams of microdata—from business filings, satellite imagery, online transactions, mobile payment systems, and even real-time electricity consumption—to produce granular, dynamic, and continuously updated revenue forecasts. These enhanced models allow governments to detect early signs of economic stress, identify new and emerging industries, anticipate tax base erosion, and calibrate fiscal policies with far greater precision than ever before. Second, smarter compliance and enforcement is becoming a reality through machine learning tools capable of anomaly detection at a scale no human audit system could match; tax authorities can now flag suspicious transactions, uncover sophisticated fraud schemes, and identify under-reporting trends with heightened accuracy, leading to reduced tax gaps and higher voluntary compliance,

as demonstrated by early adopters such as Estonia, Australia, and South Korea. Automation also improves the taxpayer experience: AI-generated prefilled tax returns based on payroll, banking, and third-party data simplify filing, reduce errors, and strengthen trust in tax systems by making compliance nearly effortless for citizens and businesses. Third, integrating the informal economy becomes significantly more feasible as AI enables governments, particularly in developing countries, to map economic activity using geospatial analytics, mobile money data, alternative digital identities, and pattern recognition tools that illuminate previously invisible sectors without imposing heavy administrative burdens. Together, these transformations mean that if AI greatly enhances revenue collection, its impact on public spending through improved transparency, efficiency, and accountability stands to be even more transformative, reshaping the entire architecture of public financial management.

Real-Time Expenditure Tracking

Modern treasury systems can already track payments and commitments. AI enhances this by analyzing expenditure patterns to predict cost overruns, spot inefficiencies, and identify corruption risks. This level of oversight allows finance ministries to intervene before budgets are breached rather than after the fiscal year ends.

Performance-Based Budgeting Reinvented

Governments have long struggled to make performance-based budgeting effective. The challenge has always been data availability and quality. AI solves this by integrating data from multiple sources public service delivery metrics, citizen feedback platforms, satellite monitoring, and administrative records. With better data, governments can finally align expenditure with measurable results, improving the efficiency and credibility of public spending.

Reducing Corruption Through Automation

Automation in procurement is one of the most promising anti-corruption tools. AI can evaluate procurement bids, verify contractor qualifications,

detect collusion through pattern recognition, and monitor delivery through geospatial technology. These systems provide transparency, reduce human discretion, and limit opportunities for fraud.

Rethinking Public Workforce Requirements

Perhaps the most politically sensitive impact of AI will be on the public sector workforce. Governments are among the world's largest employers. Automation will inevitably displace some administrative and clerical roles. Yet, contrary to common fears, AI is unlikely to shrink public employment dramatically. Instead, it will shift skill requirements.

Demand for data scientists, digital procurement specialists, cybersecurity experts, AI ethics officers, and public finance analysts will surge. Meanwhile, routine tasks data entry, financial reporting, and payroll processing will be automated. The transition will require massive investment in reskilling programs to avoid widening inequality and to ensure that frontline workers are not left behind.

Countries that proactively redesign their civil service training and recruitment systems will gain a competitive advantage. Those that resist change may experience rising inefficiencies and widening fiscal gaps. The benefits of AI come with serious risks that are especially pronounced in public finance, including algorithmic bias, where models trained on historical fiscal data can unintentionally reinforce inequities such as automated social protection systems disadvantaging marginalized regions if they mirror past underfunding unless strong oversight and transparent model design are in place; data privacy and surveillance concerns, as modern public finance systems integrate vast amounts of personal information ranging from income records and social protection registries to property ownership data, requiring strict privacy frameworks to prevent misuse or overreach; and political manipulation, since AI-powered fiscal tools could be exploited to favor certain groups, distort revenue projections, or conceal off-budget spending, underscoring the need for robust institutions and independent oversight. To harness AI responsibly, governments must adopt a coherent strategy built on four principles: ensuring ethical and transparent AI through auditable, explainable systems that allow citizens to understand and challenge fiscal

A government that can forecast more accurately, spend more efficiently, collect revenues more fairly, and detect corruption more effectively is better equipped to meet the challenges of the 21st century.

decisions; pursuing data governance reform to modernize data laws and ensure interoperability, security, and privacy protections essential for effective PFM integration; investing in capacity building so civil servants are equipped to collaborate with AI using multidisciplinary skills in economics, data science, and public policy; and embracing incremental implementation, focusing first on quick wins such as automated tax filing, procurement analytics, and real-time expenditure dashboards to build trust before rolling out more complex reforms.

A Transformational Opportunity

AI and automation will not replace the principles of sound public finance transparency, accountability, fiscal prudence, and equity. Instead, they will enhance them. A government that can forecast more accurately, spend more efficiently, collect revenues more fairly, and detect corruption more effectively is better equipped to meet the challenges of the 21st century. The opportunity is immense, but so is the responsibility. The choices governments make today will determine whether AI becomes a tool for more effective governance or an amplifier of inequality and mistrust. Public finance systems sit at the heart of this transformation. If governments get this right, AI could usher in a new era of fiscal resilience, public trust, and evidence-based decision-making.

Ramil Abbasov is a climate change and sustainability expert with over 14 years of experience in public finance management, climate finance, greenhouse gas emissions accounting, policy research, and economic analysis. He has worked closely with international organizations—including the United Nations Development Programme and the Asian Development Bank—to integrate climate risk assessments and mitigation strategies into financial governance frameworks.



Mr. Abbasov served as a Research Assistant at George Mason University, contributing to the NSF-funded Community-Responsive Electrified and Adaptive Transit Ecosystem (CREATE) project through quantitative data analysis and stakeholder engagement initiatives. Previously, he held key roles at the Asian Development Bank in Baku, Azerbaijan, as both the National Green Budget Economy Expert and the National Public Finance Management Expert, driving efforts in climate budget tagging, green economy analysis, and sustainable development policy integration.

In addition to his work with multilateral institutions, he is the CEO and Founder of “Spektr” Center for Research and Development, a research organization focused on advancing climate finance, energy transition, and sustainable economic policies. Mr. Abbasov’s earlier career includes leadership positions such as Director at ZE-Tronics CJSC and managerial roles in the banking sector with AccessBank CJSC and retail management with Third Eye Communications in the USA.

Who Governs the Algorithm?

The Epistemic Case Against Centralized AI Rulemaking

Vladyslav Manzyuk (Student Award Winner)
University of Warsaw, Poland

Consider a graphic designer running a small studio—income irregular but growing, not a single missed bill in two years. Last spring she applied for a modest business loan and was rejected in eleven seconds. She called the bank; a representative said the decision had been made "by the system," and could not tell her what to do next. What struck me was not the rejection itself but the particular quality of the silence afterward: no explanation, no visible person who had actually looked at her file, no friction of any kind — only speed, and then nothing. Regulators across the world now propose to fix this pattern through centralized AI rulebooks. I think they will fail, and the reason is not political will but structural: governing AI through a single harmonized framework reproduces the very epistemic problem that produced the eleven-second rejection in the first place. The jurisdiction that understands this first will set the global standard.

The theoretical context is Friedrich Hayek's. In his 1945 essay *The Use of Knowledge in Society*, Hayek argued that the relevant information for economic decisions never exists in concentrated form but solely as dispersed, incomplete, often contradictory knowledge held by separate individuals.¹ No central authority can aggregate it without distortion, because much of the knowledge that matters — the specific circumstances of time and place — is tacit and fleeting. This is not a political claim about markets; it is an epistemic claim about what any committee can know. I recognize this framing courts suspicion—Hayek has been conscripted for causes he would not have recognized—which is why the distinction matters to me. Scholars working on AI governance have extended this insight to regulatory design, arguing that centralized regulators face an intensified version of the same problem: rapidly

evolving, heterogeneous, opaque systems deployed across thousands of local contexts cannot be understood, much less supervised, from a single rulemaking center.²

The EU Artificial Intelligence Act—which entered into force in August 2024, began phased application in February 2025, and reaches full force for high-risk financial systems by August 2026 — is the most ambitious test of exactly the opposite proposition.³ It classifies AI systems by risk tier, prescribes documentation, human-oversight, and data-quality obligations, and creates a new AI Office to draft secondary legislation. Eurofi's 2024 overview of the Act's financial-sector implementation is sobering: unresolved gaps in data-quality standards, fragmented cross-border enforcement where the same AI system may be deemed compliant in one member state but not in another, and the concession that regulators must continually adapt to keep pace with systems that evolve faster than any rulemaking cycle.⁴ The Act presupposes a capacity for centralized technical comprehension that does not yet exist and, I would argue, structurally cannot.

The financial sector makes this visible because it has already run the experiment. AI credit scoring inherits Hayek's problem directly: the problem is not that any single institution decides badly — it is that when all institutions train on the same bureau-reported data, they decide badly in exactly the same way, at exactly the same moment, with no corrective available anywhere in the system.⁵ A second failure compounds the first. When major institutions converge on a handful of third-party foundation models and near-identical training pipelines, they stop making independent decisions and start making the same decision. Kleinberg

and Raghavan (2021) demonstrated that this "algorithmic monoculture" produces worse aggregate welfare than independent human evaluators even when each individual algorithm outperforms any human, because correlated decisions reduce aggregate decision quality.⁶ Creel extends the point: the harm is not only statistical but moral, because affected individuals cannot escape a principle applied uniformly across the economy.⁷ The designer is not unlucky; she is structurally invisible—and that is a different problem entirely.

Elinor Ostrom showed that complex resource problems are more reliably solved by overlapping, semi-autonomous governance units — each with local knowledge, each capable of mutual monitoring and adaptation — than by top-down rules imposed from a distance.

Financial regulators have begun to notice, though perhaps more slowly than the situation warrants. The European Systemic Risk Board's December 2025 report on AI and systemic risk warns that model uniformity produces correlated exposures, that concentration among a handful of AI providers creates highly interconnected nodes, and that reliance on a small set of pre-trained models poses risks comparable to those of legacy credit-scoring monocultures.⁸ The BIS Financial Stability Institute's 2025 paper on AI explainability cites a Bank of England and FCA survey finding that half of respondent firms have only partial understanding of the third-party AI systems they deploy — meaning supervisors understand even less.⁹ This is not a solvable staffing problem. Centralization does not solve this; it makes it worse, by forcing a single regulator to adjudicate systems it cannot actually inspect.

Elinor Ostrom's Nobel lecture provides the structural alternative. Ostrom showed that complex resource problems are more reliably solved by overlapping, semi-autonomous governance units — each with local knowledge, each capable of mutual monitoring and adaptation — than by top-down rules imposed from a distance.¹⁰ Applied to AI, polycentric governance means sector-specific supervisors, algorithmic-diversity requirements enforced at the institutional level, user-controlled data

portability that redistributes epistemic power to those who actually have it, and real interoperability between jurisdictions rather than extraterritorial rule-export. The point is not to multiply bureaucracy but to relocate authority — closer to the systems being supervised, and closer to the people those systems affect.

The geopolitical stakes make the argument more urgent. A 2025 *Bulletin of the Chinese Academy of Sciences* analysis describes a three-way divergence: the United States pursues development-driven governance that prioritizes innovation, the EU pursues a stringent framework aimed at exporting its rules as soft power, and China positions its coordinated model as a window of opportunity to shape international AI standards precisely while the transatlantic partners disagree.¹¹ The EU's implicit wager — that the Brussels Effect will make its centralized framework the global default — deserves scrutiny it rarely receives. If the framework cannot work on its own terms, exporting it produces paper compliance abroad and regulatory capture at home, while actual governance migrates to jurisdictions building bottom-up from sectoral expertise. The most ambitious governance project in the world may be creating the conditions for its own irrelevance.

What would better governance look like? Financial regulators should treat algorithmic diversity as a prudential requirement on par with capital adequacy — systemically important banks should be obliged to demonstrate meaningful divergence in model architecture and training data — and, crucially, to subject that divergence to cross-border supervisory review, so that a regulatory gap in one jurisdiction cannot become a contagion vector for the whole transatlantic system. Governments should also mandate user-controlled data portability frameworks — expanding on models such as PSD2 and UK Open Banking framework — allowing individuals to release their own transactional and income data to lenders of their choice. Ostrom's insight applies here directly. The person with the most accurate knowledge of her own income is the designer herself. Frameworks that place data release in her hands — rather than in the hands of the institution scoring her — are not merely more convenient; they are epistemically superior. The regulator who mandates such portability is not expanding bureaucracy; she is shrinking the information gap that produced the eleven-second silence.

The graphic designer will apply for another loan. Whether she is seen this time depends less on how ambitious Brussels is than on whether the people designing governance have been honest about what any central authority can actually know.



Vladyslav Manzyuk is a political science student at the University of Warsaw with research interests in Austrian economics and political philosophy. He has published at Mises Wire.

McDonaldisation of AI-Driven Phillipine Public Administration: A Paradox of Efficiency and Human Decision-Making

Gea Clare Desoyo (Student Award Winner)
FEU Cavite, Phillipines

Imagine walking into a government office where every process is automated, forms and information are digitized, and approvals take only minutes. No long queues, lost documents, and decisions are delivered without requiring you to return after several days. In a country historically burdened by bureaucratic inefficiencies and persistent delays in service delivery, who could possibly say no to that? That is what Filipinos have long dealt with—and continue to endure. It has caused Filipinos so many missed opportunities and hindered progress.¹ That is why the idea of a government that operates hand-in-hand with artificial intelligence seems enticing to many. AI in public administration has seemingly become the long-overdue remedy, one that echoes the change once promised by politicians. Yet, beneath this promise lies a structural tension: when public service delivery becomes faster and automated, does it also risk becoming less human?

Considering that there are efforts towards e-governance, data centralization, and smart city initiatives (e.g., Cauayan City, the first ‘smart city’ in the Philippines, declared by the Department of Science and Technology)², it reflects a broader institutional shift toward digital governance. In this context, AI emerges not merely as an innovation but as a rationalizing mechanism—a tool that can streamline processes, reduce human error, and introduce a level of objectivity often absent in traditional administrative systems. This subtle shift reflects the principles described in George Ritzer’s theory of McDonaldisation, which explains the dominance of efficiency, calculability, predictability, and control in modern institutional systems³. With this, public administration is increasingly becoming *McDonaldised*, seeking to replicate the systemic, standardized, and machine-like processes often associated with the fast-

food operations—translated from commercial efficiency into the logic of state administration.

The Cost of a Fully Automated State

AI efficiency can help automate repetitive tasks allowing public servants to focus on more important responsibilities that require much of their time. Automated systems can process large volumes of data more quickly than any human-operated offices, allowing for faster decision-making and responsive public service delivery. In the Philippine context, initiatives such as the digitization efforts of Bureau of Internal Revenue (BIR)⁴, online appointment systems in government agencies, and the expansion of e-governance platforms after the COVID-19 pandemic reflect the state’s growing independence on digital systems to improve efficiency. Applications of AI in document processing, citizen feedback analysis, and resource allocation reduce delays that have long frustrated both officials and the public.⁵ For many Filipinos who have spent hours waiting in government offices only to be told to return another day due to missing paperwork or inefficient systems, the promise of AI-driven governance appears not only attractive, but necessary.

On the other hand, calculability allows administrators to quantify outputs, track performance, and allocate resources effectively. Digitizing public information and announcements in a much organized manner, where data is easier to access, creates opportunities for transparency and improved public trust.⁶ Meanwhile, predictability of AI ensures that processes follow standardized rules, limiting arbitrary decisions and potential corruption. Algorithm-driven systems can minimize biases by

relying on data and evidence rather than personal judgment. With this, it offers a compelling vision of rationalized and optimized governance, particularly in a country where bureaucratic inefficiency has historically weakened public confidence in state institutions.

Yet, the very qualities that make AI appealing also create vulnerabilities. Over-reliance on algorithms can erode human discretion, particularly in complex cases where context matters. Standardized processes, while predictable, risk ignoring local realities in a country as socially and economically diverse as the Philippines.⁷ Rural communities with limited internet access, elderly citizens unfamiliar with digital platforms, and marginalized sectors with low digital literacy struggle to navigate fully automated systems. In such cases, efficiency may unintentionally come at the cost of accessibility and inclusion.

Beyond accessibility concerns, the increasing reliance on digital governance also exposes the Philippine government to growing cybersecurity threats. During the GovMedia Summit 2026 in Makati, Philippines, officials described cyberattacks, misinformation, and data breaches as national security risks connected to ongoing geopolitical tensions. Hacking incidents involving Philippine government websites were even linked to nation-state actors, highlighting how a highly digitized and AI-driven bureaucracy may become vulnerable not only to technical failures, but also to external political and cyber threats.⁸ While AI can strengthen administrative efficiency, it simultaneously expands the state's exposure to cyber warfare, disinformation campaigns, and data manipulation—risks that could undermine public trust and institutional stability if left unaddressed.

Policy-making often relies on the lived experiences, stories, and struggles of ordinary citizens—insights that shape compassionate and context-

sensitive decisions.⁹ Public administration is not merely about processing data; it is also about understanding human conditions that cannot always be quantified by algorithms. Moreover, these systems can be exploited or manipulated, creating loopholes that unscrupulous actors could exploit to worsen corruption rather than eliminate it. To simply put it, the same mechanisms that make bureaucratic processes efficient, also make them potentially dehumanized, rigid, and blind to nuance.

How Do We Keep the Human in the Loop?

These dynamics require balance: AI should augment, not replace human judgment in public administration. Efficiency gains must coexist with ethical considerations, social equity, and context-sensitive decision-making. Public administrators must recognize that while AI excels at speed and pattern recognition, it cannot replicate empathy, cultural understanding, and moral reasoning—elements critical to fair governance. Moreover, citizens themselves should be engaged in oversight and feedback mechanisms, ensuring that technology serves the public interest rather than administrative convenience alone.

To achieve this, first, AI systems should be designed with human-in-the-loop models, where critical decisions require human validation through technical review, contextual judgment, and formal approval. This approach preserves accountability and ensures that unique or sensitive cases are handled thoughtfully. Secondly, government agencies must invest in digital literacy programs, particularly in underserved and unrepresented communities, to ensure equitable access to AI-enabled services. Third, transparent auditing of algorithms and datasets is essential to detect and correct biases before they affect policy implementation. Lastly, policymakers should develop legal and ethical frameworks that govern AI use in Philippine public administration, establishing clear guidelines on responsibility, data privacy, and ethical boundaries.

Conclusion: Efficiency Should Not Replace Humanity

The McDonaldization of AI-driven Philippine public administration embodies the paradox of efficiency and human decision-making. The same technologies that promise faster transactions, reduced

AI should augment, not replace, human judgment in public administration. Efficiency gains must coexist with ethical considerations, social equity, and context-sensitive decision-making.

bureaucracy, and standardized governance also raise difficult questions about empathy, accountability, and inclusion. A government office where approvals take only minutes may sound ideal, especially for Filipinos long accustomed to inefficiency and administrative delays. However, governance cannot be measured by speed alone.

While AI solutions can transform public administration through efficiency, standardization, and automation, these advantages should not come at the expense of human judgment and fairness. Public service ultimately deals with people whose realities are often too complex to be reduced into datasets and algorithmic patterns. Rational systems may improve administration, but they cannot fully replace moral reasoning, compassion, and contextual understanding. This is further complicated by the fact that digital transformation also expands the state's exposure to cybersecurity risks. These developments reveal that efficiency-driven digital governance must also contend with vulnerabilities that extend beyond administration into security and sovereignty.

This editorial therefore calls for a balanced approach that embraces technology without surrendering the core values of public service. Efficiency is meant to serve human needs, not replace the role of human judgment; rationalization must coexist with discretion, empathy, and ethical responsibility. As for the Philippines, this means building not only smarter systems, but also more resilient ones—capable of withstanding both internal limitations and external digital threats while still prioritizing equitable access and human-centered government. By carefully integrating AI while maintaining human oversight and ethical safeguards, the Philippine government can create a system that is both fast and fair—efficient, yet humane. After all, evolution is inevitable—but the direction of that evolution remains a human choice.

Gea Clare Desoyo is a third-year Bachelor of Arts in Political Science student at FEU Cavite. Ms. Desoyo has an extensive background in campus journalism, having participated in various school press conferences and writing competitions. She was awarded 3rd Place in Editorial Writing during the Cluster D-2 Press Conference in Albay in 2019 and was recognized as Student Journalist of the Year in Iloilo in 2017. She also placed 10th in Science and Technology Writing at the 5th Congressional Schools Press Conference.



In addition to her writing achievements, she received a Special Recognition for Best in Research during her academic program in 2023 and was awarded Best UN Sustainable Development Goals (UNSDG) Project Proposal during JCI Week in Cavite in the same year. Beyond academics and journalism, Ms. Desoyo is also an active student leader, engaging in leadership initiatives, student representation, and civic-oriented activities within the campus community.

Her writing focuses on governance, public administration, and social issues, reflecting her academic training in Political Science, as well as her involvement in student leadership and public discourse.

The Green Mirage of Artificial Intelligence

Professor Dale Mineshima-Lowe
Parami University, Myanmar/Burma-Online

As we have seen artificial intelligence (AI) become more embedded into the everyday lives of people around the world, perspectives of the technology range across the spectrum. Governments and businesses speak about AI as an engine for growing productivity. Technology firms involved in AI development promise that it will accelerate scientific discovery and make current systems more efficient - from reducing waste across stages of supply chains, optimization of energy and other necessary resources, to assisting with the environmental and climate challenges. This has been backed by examples of how AI is being used to improve weather forecasting (Jean *et al*, 2025), helping to more efficiently manage renewable energy systems (Algburi *et al*, 2025) and supply chain challenges (Culot *et al*, 2024).

This narrative – that smarter and more efficient processes with AI can lead to a smarter, greener planet and help us manage the impact we have on our climate – is a seductive one. Yet, beneath this optimistic and positive outlook lies an uncomfortable paradox. AI requires certain infrastructures and as more companies and individuals engage in using AI – creating an AI boom – these infrastructures that power the AI, are simultaneously being questioned for the environmental and resource burdens it is creating. Recent research has increasingly suggested that discourse around AI has been focused on innovation and its usefulness in sustainability drives. This, while overlooking the physical infrastructures and systems needed to make AI possible. AI as currently imagined and developed – is reliant of sprawling data centers, water-intensive cooling systems, electricity grids that may already be under strain for current consumer and business consumption patterns, to say nothing about the

mineral-heavy hardware. A recent 2025 study by Lei *et al* in Resources, Conservation and Recycling, found that water use associated with data-center workloads can vary by more than 10,000-fold depending on the cooling systems used, server efficiency, energy sources, and geographic location. The study concluded that there is no single technological fix and that sustainable AI depends on choices about the infrastructure design and deployment.

As noted earlier – water and energy usage are of particular interest in relation to AI's 'green credentials'. In particular, water has become a prominent part of AI-sustainability conversations. Another 2025 study projected dramatic increases in global water demand from AI-driven data centers through 2050 (Herrera *et al*, 2025). The study's findings indicated that cooling requirements and electricity generation could create significant pressures in areas experiencing water scarcity and where power systems are in need of updating or expanding (Chen *et al*, 2026). What had been projections around water and energy capacities and resilience of current regions where AI data centers are being located, are becoming a reality.

Research and reports are demonstrating the complexity of the AI-sustainability discussion. In addition, we are seeing local opposition to the development of AI data centers, as communities question the impact of these centers on local water access, electricity consumption, and other environmental strain – like land usage. In places like Michigan (USA), local opposition has emerged after news of the projected consumption for a newly approved AI-focused data center became public – with the data center projected to consume 1.4 gigawatts of electricity, an amount

similar to power demands of entire cities (James, 2026). While there are several reasons for local opposition to the building of data centers in various places across the world, a recent report from *The Brookings Institute* (Pipa and Aley, 2026) highlights that for rural areas of the US, two of the many cited issues have been water and power consumption. Residents are worried about the impact on local resources and environment. While this report has centered on findings across rural US areas, this type of local opposition is being seen in other places around the globe and has been particularly prominent in areas where water scarcity is a growing issue for the local population.

What we are seeing is a complicated picture of AI and its environmental footprint. On one hand, we are beginning to see the usefulness of AI for improving climate models for environmental and climate predictions. On the other hand, we are beginning to understand the cost of current data center development practices – data centers on which continued AI use and development are reliant upon. Increasingly, concern is not only focused on the environmental footprint of AI due to training giant models, but also on the everyday usage at the global scale – as AI seeps into our daily lives. Consider all of the updates to web-browsers, word processing packages, etc. our everyday lives are touched by technologies integrating AI. Recent studies have begun to calculate environmental costs of a single AI query – these are relatively small in isolation; and then calculating this cost against the billions of potential interactions. It is the cumulative effects of these daily interactions – on electricity, on water use, on carbon emissions – that requires us to pause and reconsider the fuller environmental cost of AI as it becomes embedded in our ordinary, daily digital lives. However, even these calculations overlook the fact that the environmental footprint of AI models like ChatGPT or Claude AI, are different at various phases of their ‘life cycles’ – e.g., their training phase as compared to their inference/use phase (d’Orgeval *et al*, 2026). More

transparency and clarity around the impact AI and its data centers from tech companies would go a long way to developing more nuance to this discussion.

While it feels like there are two opposing camps – those for who see AI’s capacity to assist us in tackling global climate and environmental challenges, while simultaneously those who see AI via infrastructure-needs (e.g., data centers and their resource needs) as exacerbating environmental challenges locally – we should instead be focusing on demanding that AI be developed sustainably (Jegham *et al*, 2025). The deeper problem is governance and accountability. There has been comparatively little disclosure about lifecycle environmental impacts of AI by companies who instead have focused on revealing the latest capabilities of their LLM models. If AI is a part of building a sustainable future, then sustainability cannot remain an afterthought for technology companies in this AI boom. Sustainability must become one of the guiding principles and conditions under which it develops, with policymakers and citizens pushing for it to be guided by ecological limits and considerations. This requires not only more transparency by tech companies, but for policymakers locally and globally, to being to ask the questions everyone is thinking but many have been avoiding – as other interests offset those of environment and sustainability ones. Only then will discussions about AI for sustainability gain legitimacy.

If AI is a part of building a sustainable future, then sustainability cannot remain an afterthought for technology companies in this AI boom.



Dale Mineshima-Lowe is a Professor of Political Science, Division of Humanities and Social Sciences at Parami University. She completed her postgraduate studies in the United Kingdom (Ph.D., University of Durham), and her undergraduate studies in California, USA (BSc, Santa Clara University). Her research spans democratic transitions and governance issues, and her current research examines the developing area of citizens and data activism – how civic organizations are using digital technologies to engage citizens in combating corruption and environmental issues. She also holds teaching credentials - UK Postgraduate Certificate in Education (Adult/Further Education), Senior Fellow - Advance Higher Education (UK), and certified member of the Association for Learning Technology (ALT, UK). Dale is currently a visiting lecturer on Climate Change & Environmental Hazards (Birkbeck, University of London); an Associated Researcher with the European Democracy Institute (Bard College Berlin); and Managing Editor for the Center of International Relations - a think tank based in Washington, D.C.

Managing the Green Energy Transition

Interview with Professor Jan Rosenow
Oxford University, United Kingdom

What developments concerning a global clean-energy transition make you most optimistic that it is achievable? What concerns you most?

The extraordinary cost reductions in solar, wind, batteries, and heat pumps over the past decade give me genuine confidence that clean energy is now the economically rational choice in most markets. What concerns me most is the pace, not the direction. Policy uncertainty, permitting bottlenecks, and the political backlash we're seeing in some countries risk delaying deployment well past the window the climate science gives us.

What are the best ways to address energy efficiency in buildings for new and existing structures?

For new buildings, the answer is straightforward: mandatory performance standards that require high fabric efficiency construction, because locking in inefficiency for 50-plus years is a costly mistake we can't afford. For existing buildings, the challenge is far greater, and it requires combining financial incentives, trusted advice, and simplified access to make retrofits easy and affordable for households, particularly lower-income ones.

You've advocated the use of heat pumps in buildings. Would you briefly describe their major benefits and potential issues regarding them?

Heat pumps are multiple times more efficient than gas boilers, meaning

for every unit of electricity they consume they deliver three to five units of heat, which makes them transformative for both carbon reduction and energy bills over time. The main barriers are over-taxation of electricity, upfront cost and the need for more trained installers. None of these are insurmountable obstacles, but they do require concerted policy effort.

How would you characterize the current state of carbon capture and storage as a viable path toward meeting net zero goals?

CCS has a legitimate role in decarbonizing genuinely hard-to-abate industrial processes where direct substitution isn't yet feasible. However, it has repeatedly underdelivered on cost and scale expectations, and relying on it to offset continued fossil fuel use in sectors like heating and transport would be a serious strategic mistake. It should be seen as a complement to electrification and efficiency, not an alternative.

Green hydrogen is based on renewable energy via electrolysis while blue hydrogen is based on natural gas/carbon capture. What are your thoughts on them as viable technologies/methods?

Green hydrogen has a clear role in sectors where direct electrification is genuinely difficult, such as some industrial processes and potentially long-distance shipping, but it is far too energy-intensive to make sense for heating buildings or most transport applications. Blue hydrogen is contentious because its climate credentials depend heavily on capturing methane leakage throughout the gas supply chain, which remains an unresolved problem. Both should be prioritized for hard-to-abate uses rather than treated as generic fuel substitutes.

The grid needs to become smarter before it necessarily needs to become much bigger.

What do you consider key success factors in the public/private partnership to addressing renewable energy?

The single most important factor is long-term policy certainty. Private capital will flow at scale when investors can plan across 10 to 20 year horizons without fearing political reversals. Beyond that, governments need to actively de-risk early deployment through contracts for difference, loan guarantees, and procurement commitments, while also clearing the planning and grid connection bottlenecks that are increasingly slowing projects that are ready to build.

Moving from fossil fuels to increasing electrification presents concerns about infrastructure stresses. What challenges must be overcome to adequately meet these needs?

The grid needs to become smarter before it necessarily needs to become much bigger. Demand flexibility, smart charging for EVs, and heat pump controls can shift load in ways that significantly reduce peak stress and defer expensive infrastructure upgrades. Where network investment is genuinely required, we need to modernize planning and approval processes, which in many countries are simply not designed for the speed the transition demands.

What would you recommend to increase consumer demand to achieve successful energy transition from fossil fuels?

Households don't buy heat pumps or insulation because of climate concern alone. Making the financial case compelling through grants, low-interest financing, and energy price signals that reflect true costs is essential. Equally important is removing friction: one-stop-shop programs that handle everything from advice to installation to financing consistently

outperform approaches that leave consumers to navigate a fragmented market on their own.

Jan Rosenow is Professor of Energy and Climate Policy at the University of Oxford where he leads the 34-person Energy Programme at the Environmental Change Institute and Jackson Senior Research Fellow at Oriel College, Oxford. He is also a Senior Associate at the Cambridge Institute for Sustainability Leadership at the University of Cambridge and Affiliate Faculty at the University of Sussex. His research focuses on energy demand, energy efficiency, electrification, renewable energy, and broader energy and climate policy, with a strong emphasis on the practical implementation of decarbonization strategies. He has published numerous widely cited academic papers, technical reports, and policy analyses, contributing to key discussions on the energy transition. Professor Rosenow has been recognized by Stanford's University database of the top 2% most cited scientists worldwide.



In recognition of his impact in the energy sector, he has been named the most read thought leader on the energy transition in 2024, one of the top 100 players in the global climate space and among the top 25 energy influencers and top 15 sustainability influencers worldwide. Additionally, he is one of LinkedIn's Top Green Voices.

He holds a Doctorate from the University of Oxford, an MSc from the London School of Economics, and has completed executive training at the University of Cambridge and the Florence School for Regulation.

Green Tech as a Strategic Insurance: What the EU Energy Crisis Revealed

Bogdan Romaniuk (Student Award Winner)
Financial University, Russian Federation

Green tech is usually framed as an environmental issue. The standard argument is simple: emissions are rising, the climate is changing, and states need cleaner energy to reduce environmental damage. That argument is valid. But today it is not enough. In an unstable world, the real value of green technology lies not only in cutting emissions. It lies in reducing dependence. States deeply tied to oil and gas remain vulnerable to price shocks, supply disruptions, external pressure, and the political leverage of suppliers. The real cost of such dependence is measured not only in emissions, but also in strategic vulnerability.

The European crisis of 2022 exposed this clearly. Before the war, Russia supplied about 45% of EU gas imports, while Russian oil accounted for 27% of the EU's oil imports at the start of 2022.¹ Europe was not just a buyer on the global market. It was a system built around dependence on an external supplier. When that dependence broke down, the shock was immediate and severe. In August 2022, gas prices in Europe rose above €300 per megawatt-hour (MWh) and stayed above €265 for five days.² Household prices across the EU also jumped sharply. In the first half of 2022, average household gas prices rose from €6.4 to €8.6 per 100 kWh, while average electricity prices rose from €22.0 to €25.3.³ In the second half of 2022, they climbed to €11.4 for gas and €28.4 for electricity.⁴ Dependence on fossil fuels had revealed itself as a strategic weakness.

The energy shock quickly spread beyond utility bills. The European Central Bank described 2022 as a major terms-of-trade shock. The euro area current account moved from a surplus of 2.8% of GDP in 2021 to a deficit of 0.8% in 2022, the largest annual deterioration on record.⁵

States deeply tied to oil and gas remain vulnerable to price shocks, supply disruptions, external pressure, and the political leverage of suppliers. The real cost of such dependence is measured not only in emissions, but also in strategic vulnerability.

Expensive imported energy hit households, the external balance, real incomes, and the wider economy. That is the central lesson. Fossil dependence hurts twice: in the long run through emissions, and in the short run through vulnerability to external shocks.

But for this argument, the impact matters less than the response. Europe reduced its dependence on Russia through several channels. In the short term, diversification of imports, emergency gas purchases, and coordinated demand cuts all played a role. This should be stated openly, because it would be false to pretend that Europe was saved only by solar panels and wind turbines. Yet switching suppliers is not the same as becoming independent. It reduces dependence on one country, but not on a model in which energy must still be bought again and again on external markets. This is where green tech becomes important.

The problem is not just Russia, and not just Europe. It is deeper. Changing the supplier may reduce dependence on one supplier, but it does not change the logic of fossil dependence itself. A state still remains exposed to external prices, routes, exporters, and future crises. Green technologies matter because they work differently. They gradually reduce the need for imported fuel and shift the center of gravity from permanent external purchases to domestic energy infrastructure.

That is why it is important to distinguish ordinary energy savings from technological gas substitution. Saving energy helps a country survive a crisis. Green tech changes the system after the crisis. New wind and solar capacity displace gas in power generation. Electrification expands the use of clean electricity. Heat pumps reduce the need for gas in heating. Grids, storage, and demand management make a higher share of clean generation usable in everyday life. In Europe, this was visible not only in slogans but in numbers. In 2022 alone, new solar generation in the EU was equivalent to about 4.6 billion cubic meters of Russian gas.⁶ In 2023, the EU added another 56 GW of solar capacity.⁷ In the same period, heat pump sales in Europe rose by almost 40%, and the International Energy Agency estimates that, if Europe meets its targets, heat pumps alone could cut gas demand by about 7 billion cubic meters by 2025.⁸ This is direct substitution of fossil fuels with electric technologies and domestic clean-energy generation.

The same pattern is visible at the system level. The share of renewables in gross EU electricity consumption rose from 37.8% in 2021 to 41.2% in 2022, 45.3% in 2023, and 47.5% in 2024.⁹ At the same time, EU gas demand fell by 13.3% in 2022 and by another 7.4% in 2023.¹⁰ This is more than a crisis response. It is a movement toward a system that is less dependent on imported gas as such. Europe did not become fully independent simply because it cut Russia's share. But the rise of renewable generation and the fall in gas demand show that part of Europe's vulnerability was reduced not only through new imports, but through lower demand for gas itself. This is where green tech stops being just climate policy and becomes an instrument of resilience.

By 2025, the share of Russian gas in EU imports had fallen from 45% to 12%, while the share of Russian oil had dropped from 27% to 2%.¹¹ Yes, part of this was achieved by finding new suppliers. But the deeper lesson is different. The more energy a system gets from its own low-carbon infrastructure, and the less imported gas and oil it needs, the less room remains for future coercion, price shocks, and political leverage from external sellers. In that sense, green tech matters not only as climate policy, but as a way to reduce the depth of dependence itself. Green technologies are not a magic solution. A secure clean transition requires more wind and solar. It also requires stronger grids, energy storage, demand management, diversified technologies, climate-resilient infrastructure, and attention to new weak points, including supply chains and critical minerals. A serious green strategy is therefore not a collection of symbols. It is an infrastructure transformation.

That is why Europe's lessons matter far beyond Europe. Any state that remains deeply dependent on imported oil and gas remains vulnerable to other people's wars, price spikes, route disruptions, and external pressure. This applies not only to Europe, but also to large importers in Asia and to developing economies that pay for external shocks through inflation, fiscal strain, and a narrower political margin for action. The IMF makes an important point here: the green transition can strengthen energy security if investments genuinely reduce old vulnerabilities rather than create new ones in different forms.¹²

Green technologies should not be understood as a moral accessory of the rich world, nor as an abstract virtue in the struggle against climate change. Their real value is that they reduce strategic vulnerability. They allow states to depend less on external fuel, pay lower costs for conflicts they do not control, and fear less that energy will once again be turned into a weapon of pressure.

Green tech is not a climate luxury. It is strategic insurance for an unstable world.

Bogdan Romaniuk is an undergraduate student at the Financial University under the Government of the Russian Federation, studying in a program connected with international economic relations and the world economy. His interests include international political economy, energy security, global markets, strategic infrastructure, and the intersection of technology and state power. His current work focuses on how energy transition and technological change affect state resilience and economic security.



Green Tech Needs a Nuclear Backbone

Kaashvi Sheoran (Student Award Winner)
Amsterdam University College, Netherlands

The green transition has developed an aesthetic. It is a world of solar panels, offshore wind farms, batteries, and electric cars, presented as though clean energy will arrive through one seamless technological arc. Much of that vision is right. Solar and wind are indispensable, and their growth has been extraordinary. The International Energy Agency projects that renewables will supply 46% of global electricity by 2030, with wind and solar together accounting for 30%.¹ Any serious climate strategy must accelerate that expansion. But the leap from saying that renewables are essential to saying that they are sufficient on their own is not a scientific conclusion. It is a political preference.

That preference begins to unravel as soon as climate policy moves beyond the power sector. Heat remains one of the least glamorous but most important frontiers of decarbonization. It still accounts for almost half of global final energy consumption and more than a third of energy-related carbon dioxide emissions.² Industrial heat demand is also expected to keep growing between 2020 and 2030, while renewable heat is not expanding fast enough to displace fossil fuels at the pace climate targets require.³ Heavy industry is difficult to clean up not

simply because it uses a great deal of energy, but because it needs constant, high-temperature energy and because some emissions are embedded in industrial chemistry itself. Steel is responsible for around 7% of global energy-system carbon dioxide emissions, while in cement roughly two-thirds of emissions come from calcination rather than fuel combustion alone.⁴ These are not side issues. They are central to the decarbonization challenge.

This is precisely where nuclear energy has an advantage that green-tech debates often overlook. Nuclear does not only generate electricity. Certain advanced designs, particularly High-Temperature Gas-cooled Reactors (HTGRs), can also provide process heat for industrial applications, including hydrogen production and cogeneration.⁵ That matters because the transition is not only about replacing dirty electrons with clean ones. It is also about replacing fossil-fired heat in steel, chemicals, refining, and other energy-intensive sectors with something equally reliable. Electric boilers, heat pumps, and direct electrification should all expand rapidly where they can. But for the hardest industrial processes, the question is not only whether energy is clean. It is whether it is continuous, concentrated, and hot enough.

Nuclear does not only generate electricity. Certain advanced designs, particularly High-Temperature Gas-cooled Reactors (HTGRs), can also provide process heat for industrial applications, including hydrogen production and cogeneration. .

Electricity systems reveal the same problem in another form. A major study in *Nature Communications*, using 39 years of hourly weather data across 42 countries, found that even highly optimized solar-and-wind systems without storage satisfy demand only 72% to 91% of the time, depending on geography and system design.⁶ That does not mean renewable-heavy grids are impossible. It means variability is a structural challenge rather than a rhetorical inconvenience. Storage,

transmission, demand response, and interconnection can all reduce that challenge, and they should. Yet they also require time, land, capital, and political coordination. The International Energy Agency notes that grid infrastructure often takes far longer to plan and build than renewable generation itself.⁷ A credible low-carbon system therefore needs not only abundant clean electricity when conditions are favorable, but also firm low-emissions power when they are not.

That is where nuclear energy matters. The strongest case for nuclear is not that it should replace solar and wind, but that it should support them. Nuclear power offers firm low-carbon electricity, reduces dependence on fossil-fuel backup, and can help supply heat and hydrogen in industrial settings where electrification alone may not be enough. Even the International Energy Agency's net-zero pathway, which remains dominated by renewables, still requires global nuclear capacity to double by 2050. The same report estimates that today's nuclear fleet avoids 180 billion cubic metres of global gas demand each year.⁸ The serious question, then, is not renewables or nuclear. It is what combination of low-carbon technologies can deliver speed, reliability, and deep decarbonization together.

The public-health argument is just as strong as the climate argument. Germany's decision to phase out nuclear power after Fukushima is often remembered as a moral response to technological risk, but the evidence suggests it also imposed substantial social harm. Jarvis, Deschenes, and Jha find that the lost nuclear generation was replaced primarily by coal-fired generation and electricity imports, producing social costs of €3 billion to €8 billion per year, much of it from increased mortality linked to air pollution.⁹ More broadly, a 2023 *BMJ* study estimated that ambient air pollution attributable to fossil fuels causes 5.13 million excess deaths globally each year.¹⁰ Modern nuclear power is not risk-free, and no honest argument should pretend otherwise. But its overall safety profile is far stronger than public memory suggests. Comparative mortality data indicate that nuclear, wind, and solar all have very low death rates per terawatt-hour, vastly below those of coal, oil, and gas.¹¹ In climate politics, refusing one of the largest low-carbon energy sources available is not a morally neutral act when the alternative remains prolonged fossil dependence.

There is also a geopolitical reason this debate matters. The energy shock that followed Russia's invasion of Ukraine exposed the strategic vulnerability created by dependence on imported fossil fuels. Nuclear does not remove geopolitics from energy, but it changes its terms. It lowers exposure to volatile gas markets and creates a different set of supply-chain dependencies, many of which are more manageable over the long term. Uranium production is concentrated in a small number of suppliers, led by Kazakhstan and followed by countries including Canada and Namibia, while advanced reactors are increasing strategic interest in specialized fuels such as high-assay low-enriched uranium, or HALEU.¹² Green technology is not only about emissions. It is also about leverage, resilience, and strategic autonomy.

The strongest objection to nuclear is therefore not that it is unnecessary, but that it is difficult. Large reactors are expensive, politically contentious, and often delayed. Financing remains a major challenge, and the International Energy Agency is clear that long construction timelines and cost overruns continue to inhibit investment.¹³ Small modular reactors may eventually ease some of these constraints, especially in industrial clusters and hard-to-electrify sites, but they are not an imminent miracle. Costs still need to fall. Supply chains still need to mature. Governments still need to prove that they can deliver nuclear projects competently. A good pro-nuclear argument should admit all of that.

But that realism strengthens the case for nuclear rather than weakening it. The question is not whether nuclear is perfect. No major energy technology is. The question is whether countries facing rising electricity demand, industrial decarbonization, energy-security pressures, and climate deadlines can afford to discard one of the few proven sources of firm low-emissions power. The practical answer is no. Existing reactors should be kept online where they can operate safely. New nuclear should be built where states have the institutional capacity to manage cost and delivery. Nuclear should be paired with renewables, not framed as their rival. If governments are serious about green technology, they should stop treating nuclear power as an awkward concession and start treating it as what it is: not the whole answer, but a necessary part of one.



Kaashvi Sheoran is a final year student at Amsterdam University College majoring in Economics and International Relations. She is particularly drawn to questions of global politics, development, and sustainability. Beyond her academic work, she enjoys poetry, rock climbing, and discovering good coffee spots.

REFERENCES AND FOOTNOTES

The Architecture of Adaptation: International Law and the Governance of Artificial Intelligence

Footnotes

- 1 UN Secretary-General's High-Level Advisory Body on Artificial Intelligence, *Governing AI for Humanity: Final Report* (United Nations, September 2024).
- 2 International Law Commission, "Fragmentation of International Law: Difficulties Arising from the Diversification and Expansion of International Law," Report of the Study Group of the ILC, A/CN.4/L.682 (13 April 2006).
- 3 B. S. Chimni, "Third World Approaches to International Law: A Manifesto," *International Community Law Review* 8 (2006): 3–27; Antony Anghie, *Imperialism, Sovereignty and the Making of International Law* (Cambridge: Cambridge University Press, 2005).
- 4 Robert O. Keohane and David G. Victor, "The Regime Complex for Climate Change," *Perspectives on Politics* 9, no. 1 (2011): 7–23.
- 5 Anne-Marie Slaughter, *A New World Order* (Princeton: Princeton University Press, 2004); Benedict Kingsbury, Nico Krisch, and Richard B. Stewart, "The Emergence of Global Administrative Law," *Law and Contemporary Problems* 68, no. 3 (2005): 15–61; Harold Hongju Koh, "Why Do Nations Obey International Law?," *Yale Law Journal* 106, no. 8 (1997): 2599–2659.
- 6 Treaty on the Non-Proliferation of Nuclear Weapons, opened for signature 1 July 1968, 729 UNTS 161; see also IAEA Statute, 26 October 1956, 276 UNTS 3.
- 7 Paris Agreement, opened for signature 22 April 2016, in force 4 November 2016, UN Doc. FCCC/CP/2015/L.9/Rev.1, arts. 4, 13–14.
- 8 Abram Chayes and Antonia Handler Chayes, *The New Sovereignty: Compliance with International Regulatory Agreements* (Cambridge, MA: Harvard University Press, 1995), explaining the "managerial" model of compliance through transparency and dispute settlement.
- 9 Convention on International Civil Aviation, opened for signature 7 December 1944, 15 UNTS 295 ("Chicago Convention"), art. 37; see Jacob Werksman, "Compliance and the Chicago Convention," *International Organizations Law Review* 16 (2019): 187–210.
- 10 United Nations Convention on the Law of the Sea, opened for signature 10 December 1982, 1833 UNTS 397.
- 11 UN General Assembly, "Developments in the Field of Information and Telecommunications in the Context of International Security," A/68/98 (24 June 2013), affirming the applicability of international law to cyberspace; Michael N. Schmitt, ed., *Tallinn Manual 2.0 on the International Law Applicable to Cyber Operations*, 2nd ed. (Cambridge: Cambridge University Press, 2017).
- 12 Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law, CETS No. 225, adopted 17 May 2024, opened for signature 5 September 2024.
- 13 UN General Assembly Resolution 78/265, "Seizing the Opportunities of Safe, Secure and Trustworthy Artificial Intelligence Systems for Sustainable Development," A/RES/78/265 (21 March 2024); UN General Assembly Resolution 78/311, "Enhancing International Cooperation on Capacity Building of Artificial Intelligence," A/RES/78/311 (1 July 2024).
- 14 Pact for the Future, A/RES/79/1 (22 September 2024), Annex I ("Global Digital Compact").
- 15 UNESCO, *Recommendation on the Ethics of Artificial Intelligence* (Paris: UNESCO, 2021), adopted by the General Conference at its 41st session.
- 16 OECD, "Recommendation of the Council on Artificial Intelligence," OECD/LEGAL/0449 (adopted 22 May 2019, amended 3 May 2024).
- 17 Hiroshima Process International Code of Conduct for Organizations Developing Advanced AI Systems (30 October 2023); The Bletchley Declaration by Countries Attending the AI Safety Summit, Bletchley Park, 1–2 November 2023; Seoul Declaration for Safe, Innovative and Inclusive AI (21 May 2024); Statement of the Paris AI Action Summit (11 February 2025).
- 18 ISO/IEC 42001:2023, *Information Technology — Artificial Intelligence — Management System* (Geneva: International Organization for Standardization, 2023).
- 19 Anu Bradford, *The Brussels Effect: How the European Union Rules the World* (Oxford: Oxford University Press, 2020); Regulation (EU) 2024/1689 of the European Parliament and of the Council laying down harmonised rules on artificial intelligence (Artificial Intelligence Act), OJ L 2024/1689.
- 20 International Covenant on Civil and Political Rights, 16 December 1966, 999 UNTS 171; International Covenant on Economic, Social and Cultural Rights, 16 December 1966, 993 UNTS 3; see also UN Human Rights Council, Resolution 51/22 (7 October 2022).
- 21 *Trail Smelter Case* (United States v. Canada), 3 RIAA 1905 (1938, 1941); *Pulp Mills on the River Uruguay* (Argentina v. Uruguay), Judgment, ICJ Reports 2010, p. 14, paras. 101, 197.
- 22 International Law Commission, *Articles on the Responsibility of States for Internationally Wrongful Acts*, A/56/10 (2001), annexed to UNGA Res. 56/83 (12 December 2001).
- 23 Frank Pasquale, *New Laws of Robotics: Defending Human Expertise in the Age of AI* (Cambridge, MA: Belknap Press, 2020), 19–22, on the layered nature of effective AI accountability.
- 24 UN General Assembly Resolution 43/53, "Protection of Global Climate for Present and Future Generations of Mankind" (6 December 1988); Dinah Shelton, "Common Concern of Humanity," *Iustum Aequum Salutare* 5, no. 1 (2009): 33–40.

What Does Artificial Intelligence Know about Food and Eating?

References

- 1 Allgaier, J.; Bartelmeß, T.; Decorte, P.; Evans, R.; Feldman, Z.; Fernandez, M.; Forno, F.; Fuentes, C.; Graf, K.; Groszlik, K. Han, D.-I. D. Huey, S.L.; König, L.M.; Leer, J.; Lewis, T.; Munk, A.K.; Onorati, M.G.; Parasecoli, F.; Partridge, S.R.; Podszun, M.C.; Richardson, L.; Rüdiger, S.; Schneider, T.; Sproesser, G.; Wirsam, J. and Zoet, M. (2025): Digital Foodscapes: Past - Present - Future. Position Paper. SocArXiv https://doi.org/10.31235/osf.io/pzmu5_v1

Who Governs the Algorithm? The Epistemic Case Against Centralized AI Rulemaking

Footnotes

- 1 Friedrich A. Hayek, "The Use of Knowledge in Society," *American Economic Review* 35, no. 4 (September 1945): 519–530, at 519–520. <https://www.jstor.org/stable/1809376>.
- 2 On the epistemic intensification in AI governance, see Eurofi, "AI Act: Key Measures and Implications for Financial Services," *Eurofi Regulatory Update* (September 2024): 38–43, at 43. <https://www.eurofi.net>.
- 3 Eurofi, "AI Act: Key Measures and Implications for Financial Services," *Eurofi Regulatory Update* (September 2024): 38–43, at 39.
- 4 Eurofi, "AI Act: Key Measures," 39, 43 (on regulators adapting skills, data-quality gaps, and cross-border inconsistency).
- 5 Andreas Fuster, Paul Goldsmith-Pinkham, Tarun Ramadorai, and Adalbert Walther, "Predictably Unequal? The Effects of Machine Learning on Credit Markets," *Journal of Finance* 77, no. 1 (February 2022): 5–47.
- 6 Jon Kleinberg and Manish Raghavan, "Algorithmic Monoculture and Social Welfare," *Proceedings of the National Academy of Sciences* 118, no. 22 (2021): e2018340118. <https://doi.org/10.1073/pnas.2018340118>.
- 7 Kathleen A. Creel, "Algorithmic Monoculture and Systemic Exclusion" (working paper). <https://philarchive.org/rec/CREAMA>.
- 8 Stephen Cecchetti, Robin L. Lumsdaine, Tuomas Peltonen, and Antonio Sánchez Serrano, *AI and Systemic Risk*, Advisory Scientific Committee Report No. 16 (Frankfurt: European Systemic Risk Board, December 2024).

2025). https://www.esrb.europa.eu/pub/pdf/asc/esrb.ascreport202512_AIandsystemicrisk.en.pdf.

9 Fernando Pérez-Cruz, Jermy Prenio, Fernando Restoy, and Jeffery Yong, *Managing Explanations: How Regulators Can Address AI Explainability*, FSI Occasional Paper No. 24 (Basel: Bank for International Settlements, September 2025). <https://www.bis.org/fsi/fsipapers24.htm>.

10 Elinor Ostrom, "Beyond Markets and States: Polycentric Governance of Complex Economic Systems," *American Economic Review* 100, no. 3 (June 2010): 641–672. <https://doi.org/10.1257/aer.100.3.641>.

11 Yang Mei, Jing Zeng, and Yong Zhan, "U.S.-Europe Artificial Intelligence Regulatory Cooperation, Divergence, and China's 'Window of Opportunity' for Strategic Breakthrough" [in Chinese; abstract in English], *Bulletin of the Chinese Academy of Sciences* 40, no. 4 (2025): 718. <https://bulletinofcas.researchcommons.org/journal/vol40/iss4/15>.

Works Cited

Cecchetti, Stephen, Robin L. Lumsdaine, Tuomas Peltonen, and Antonio Sánchez Serrano. AI and Systemic Risk. Advisory Scientific Committee Report No. 16. Frankfurt: European Systemic Risk Board, December 2025.

Creel, Kathleen A. "Algorithmic Monoculture and Systemic Exclusion." Working paper. PhilArchive. <https://philarchive.org/rec/CREAMA>.

Eurofi. "AI Act: Key Measures and Implications for Financial Services." Eurofi Regulatory Update, September 2024, 38–43.

Fuster, Andreas, Paul Goldsmith-Pinkham, Tarun Ramadorai, and Adalbert Walther. "Predictably Unequal? The Effects of Machine Learning on Credit Markets." *Journal of Finance* 77, no. 1 (February 2022): 5–47.

Hayek, Friedrich A. "The Use of Knowledge in Society." *American Economic Review* 35, no. 4 (September 1945): 519–530.

Kleinberg, Jon, and Manish Raghavan. "Algorithmic Monoculture and Social Welfare." *Proceedings of the National Academy of Sciences* 118, no. 22 (2021): e2018340118.

Mei, Yang, Jing Zeng, and Yong Zhan. "U.S.-Europe Artificial Intelligence Regulatory Cooperation, Divergence, and China's 'Window of Opportunity' for Strategic Breakthrough." [In Chinese; abstract in English.]

Bulletin of the Chinese Academy of Sciences 40, no. 4 (2025): 715–724. <https://bulletinofcas.researchcommons.org/journal/vol40/iss4/15>.

Ostrom, Elinor. "Beyond Markets and States: Polycentric Governance of Complex Economic Systems." *American Economic Review* 100, no. 3 (June 2010): 641–672.

Pérez-Cruz, Fernando, Jermy Prenio, Fernando Restoy, and Jeffery Yong. *Managing Explanations: How Regulators Can Address AI Explainability*. FSI Occasional Paper No. 24. Basel: Bank for International Settlements, September 2025.

McDonaldization of AI-Driven Phillipine Public Administration: A Paradox of Efficiency and Human Decision-Making

Footnotes

1 Jude Avorque Acidre, "Efficiency in Governance", *Manila Standard*, March 17, 2026, <https://manilastandard.net/opinion/314717036/efficiency-in-governance-2.html>

2 Kimberly Ann Egalin, "Smart City: How LGUs Turn into Safe and Livable Communities.", *DOST*, June 5, 2025. <https://www.dost.gov.ph/knowledge-resources/news/86-2025-news/4031-smart-city-how-lgus-turn-into-safe-and-livable-communities.html>

3 Ruth A. Wienclaw, "McDonalization", *EBSCO Research Starters*, 2021, <https://www.ebsco.com/research-starters/social-sciences-and-humanities/mcdonaldization>

4 Bitila Zhuli, "The Impact of AI in Public Administration: Transforming Governance and Services", *ResearchGate*, July 2025, https://www.researchgate.net/publication/395698711_THE_IMPACT_OF_AI_IN_PUBLIC_ADMINISTRATION_TRANSFORMING_GOVERNANCE_AND_SERVICES

5 Alex B. Brillantes and Maricel T. Fernandez, "Restoring Trust and Building Integrity in Government: Issues and Concerns in the Philippines and Areas for Reform." *International Public Management Review*, 2014, <https://ipmr.net/index.php/ipmr/article/view/102>

6 Mahdi Masoudi. "Algorithmic Governance, Data-Driven Decision-Making, and the Transformation of Democratic Accountability in Contemporary States," *Advanced Journal of Management, Humanity and Social Science*, 2025. <https://doi.org/10.5281/zenodo.18009536>

7 Ian McIntosh and Sharon Wright. "Exploring What the Notion of 'Lived Experience' Offers for Social Policy Analysis." *Journal of Social Policy*, 2019. <https://doi.org/10.1017/S0047279418000570>

Bibliography:

Acidre, Jude Avorque. "Efficiency in Governance." *Manila Standard*. (March 17, 2026). <https://manilastandard.net/opinion/314717036/efficiency-in-governance-2.html>

Brillantes, Alex B., and Fernandez, Maricel T. "Restoring Trust and Building Integrity in Government: Issues and Concerns in the Philippines and Areas for Reform." *International Public Management Review* 12, no. 2 (2014): 55–80. <https://ipmr.net/index.php/ipmr/article/view/102>

Egalin, Kimberly Ann. "Smart City: How LGUs Turn into Safe and Livable Communities." *DOST*. (June 5, 2025). <https://www.dost.gov.ph/knowledge-resources/news/86-2025-news/4031-smart-city-how-lgus-turn-into-safe-and-livable-communities.html>

Masoudi, Mahdi. "Algorithmic Governance, Data-Driven Decision-Making, and the Transformation of Democratic Accountability in Contemporary States," *Advanced Journal of Management, Humanity and Social Science* 2, no. 1 (2025): 10–22, <https://doi.org/10.5281/zenodo.18009536>

McIntosh, Ian, and Wright, Sharon. "Exploring What the Notion of 'Lived Experience' Offers for Social Policy Analysis." *Journal of Social Policy* 48, no. 3 (2019): 449–67. <https://doi.org/10.1017/S0047279418000570>

Wienclaw, Ruth A. "McDonalization." *EBSCO Research Starters*. (2021). <https://www.ebsco.com/research-starters/social-sciences-and-humanities/mcdonaldization>

Zhuli, Bitila. "The Impact of AI in Public Administration: Transforming Governance and Services." *ResearchGate*, (2025). https://www.researchgate.net/publication/395698711_THE_IMPACT_OF_AI_IN_PUBLIC_ADMINISTRATION_TRANSFORMING_GOVERNANCE_AND_SERVICES

The Green Mirage of Artificial Intelligence

References

Algburi, S., Al Kareem, S.S.A., Sapaev, I.B., Mukhidinov, O., Hassan, Q., Khalaf, D.H., and Jabbar, F.I. (2025). The role of artificial intelligence in accelerating renewable energy adoption for global energy transformation. *Unconventional Resources*, 8, 100229. <https://doi.org/10.1016/j.unres.2025.100229>

Chen, D., Zhou, Z., Cai, Y., Qin, J., Katchova, A., and Chen, L. (2026). Concentrated siting of AI data centers drives regional power-system stress under rising global compute demand. arXiv preprint (Cornell

University) <https://doi.org/10.48550/arXiv.2604.06198>

Culot, G., Podrecca, M., and Nassimbeni, G. (2024). Artificial intelligence in supply chain management: A systematic literature review of empirical studies and research directions. *Computers in Industry*, 162, 104132. <https://doi.org/10.1016/j.compind.2024.104132>

Data Center Watch (2026). \$64 billion of data center projects have been blocked or delayed amid local opposition. Data Center Watch Report. Retrieved from: <https://www.datacenterwatch.org/report>

d'Orgeval, A., Sheehan, S., Avenas, Q., Assoumou, E., and Sessa, V. (2026) Generative AI impact assessment through a life cycle analysis of multiple data center typologies. *Applied Energy*. 406, 127288. <https://doi.org/10.1016/j.apenergy.2025.127288>

Herrera, M., Xie, X., Menapace, A., Zanfei, A., and Brentan, B. M. (2025). Sustainable AI infrastructure: A scenario-based forecast of water footprint under uncertainty. *Journal of Cleaner Production*. 526, 146528. <https://doi.org/10.1016/j.jclepro.2025.146528>

James, L. (2026). After a \$16 billion Stargate AI data center was built despite being voted down, Michigan towns rush to block new buildouts – massive facility will suck 1.4 Gigawatts of energy to power ChatGPT. Tom's Hardware blog site. 7 May 2026. Retrieved from: https://www.tomshardware.com/tech-industry/michigan-towns-rush-to-block-ai-data-centers-after-16-billion-stargate-project-overrode-local-opposition?utm_source=chatgpt.com

Jean, M., Pappenberger, F., Chan, P.W., Bouchet, V., Stav, N., and Honda, Y. (2025). Forecasting the Future: The Role of Artificial Intelligence in Transforming Weather Prediction and Policy. *WMO Bulletin* 74(2) – 2025. Retrieved from: <https://wmo.int/media/magazine-article/forecasting-future-role-of-artificial-intelligence-transforming-weather-prediction-and-policy>

Jegham, N., Abdelatti, M., Koh, C.Y., Elmoubarki, L., and Hendawi, A. (2025). How Hungry is AI? Benchmarking Energy, Water, and Carbon Footprint of LLM Inference. *arXiv preprint (Cornell University)*. <https://doi.org/10.48550/arXiv.2505.09598>

Lei, N., Lu, J., Shehabi, A., and Masanet, E. (2025). The water use of data center workloads: A review and assessment of key determinants. *Resources, Conservation and Recycling*. 219, 108310. <https://doi.org/10.1016/j.resconrec.2025.108310>

Pipa, A.F. and Aley, A. (2026) The local implications of data centers for rural communities in the US. Brookings Institute – Commentary – Research summary of America's Rural Future virtual symposium – January 27, 2026. Retrieved from: <https://www.brookings.edu/articles/local-implications-data-centers-rural-communities-us/>

Green Tech as Strategic Insurance: What the EU Energy Crisis Revealed

- 1 European Commission — REPowerEU: phase out Russian energy imports https://energy.ec.europa.eu/strategy/repowereu-phase-out-russian-energy-imports_en
- 2 Council of the European Union — gas price spikes <https://www.consilium.europa.eu/en/infographics/a-market-mechanism-to-limit-excessive-gas-price-spikes/>
- 3 Eurostat — household gas/electricity prices, first half of 2022 <https://ec.europa.eu/eurostat/web/products-eurostat-news/-/ddn-20221031-1>
- 4 Eurostat — household gas/electricity prices, second half of 2022 <https://ec.europa.eu/eurostat/web/products-eurostat-news/w/ddn-20230426-2>
- 5 ECB — current account / terms of trade shock https://www.ecb.europa.eu/press/economic-bulletin/articles/2023/html/ecb.ebart202306_01~4cd4215076.hr.html
- 6 Ember — solar generation equivalent to bcm of gas <https://ember-energy.org/latest-insights/european-electricity-review-2023/>
- 7 European Commission — European Solar Charter / 56 GW solar https://energy.ec.europa.eu/topics/renewable-energy/solar-energy/european-solar-charter_en
- 8 IEA — heat pumps and gas demand reduction <https://www.iea.org/news/the-global-energy-crisis-is-driving-a-surge-in-heat-pumps-bringing-energy-security-and-climate-benefits>
- 9 Eurostat — renewables share in electricity, 2022 / 2023 / 2024 <https://ec.europa.eu/eurostat/web/products-eurostat-news/w/ddn-20240221-1> <https://ec.europa.eu/eurostat/web/products-eurostat-news/w/ddn-20250221-3> <https://ec.europa.eu/eurostat/web/products-eurostat-news/w/ddn-20260114-1>
- 10 Eurostat — EU gas demand fell in 2023 <https://ec.europa.eu/eurostat/web/products-eurostat-news/w/ddn-20240528-1>
- 11 European Commission — Russian gas 45% - 12%, oil 27% - 2% https://energy.ec.europa.eu/strategy/repowereu-phase-out-russian-energy-imports_en
- 12 IMF — Energy Security and the Green Transition <https://www.imf.org/en/Publications/WP/Issues/2024/01/12/Energy-Security-and-The-Green-Transition-543806>

Green Tech Needs a Nuclear Backbone

Footnotes

- 1 International Energy Agency, *Renewables 2024: Global Overview* (2024)
- 2 International Energy Agency, “Renewable Heat,” in *Renewables 2025* (2025).
- 3 International Energy Agency, “Renewable Heat,” in *Renewables 2025* (2025).
- 4 International Energy Agency, *Iron and Steel Technology Roadmap* (2020); European Commission, *Deep Decarbonisation of Industry: The Cement Sector* (2020).
- 5 International Atomic Energy Agency, *Gas-Cooled Reactors* (n.d.).
- 6 D. Tong et al., “Geophysical Constraints on the Reliability of Solar and Wind Power Worldwide,” *Nature Communications* 12 (2021): Article 6146.
- 7 International Energy Agency, *Electricity Grids and Secure Energy Transitions* (2023).
- 8 International Energy Agency, *Nuclear Power and Secure Energy Transitions* (2022).
- 9 Stephen Jarvis, Olivier Deschenes, and Akshaya Jha, “The Private and External Costs of Germany's Nuclear Phase-Out,” *Journal of the European Economic Association* 20, no. 3 (2022): 1311–1346.
- 10 J. Lelieveld et al., “Air Pollution Deaths Attributable to Fossil Fuels: Observational and Modelling Study,” *BMJ* 383 (2023): e077784.
- 11 Hannah Ritchie, “What Are the Safest and Cleanest Sources of Energy?” *Our World in Data* (2020).
- 12 Nuclear Energy Agency and International Atomic Energy Agency, *Uranium 2024: Resources, Production and Demand* (2025); Nuclear Energy Agency, *High-Assay Low-Enriched Uranium: Drivers, Implications and Security of Supply* (2024).
- 13 International Energy Agency, *The Path to a New Era for Nuclear Energy* (2025).

Works Cited

European Commission. *Deep Decarbonisation of Industry: The Cement Sector*. 2020. https://setis.ec.europa.eu/system/files/2021-02/jrc120570_decarbonisation_of_cement_fact_sheet_2.pdf

International Atomic Energy Agency. *Gas-Cooled Reactors*. n.d. <https://www.iaea.org/topics/gas-cooled-reactors>

International Energy Agency. *Iron and Steel Technology Roadmap*. 2020. <https://www.iea.org/reports/iron-and-steel-technology-roadmap>

International Energy Agency. *Nuclear Power and Secure Energy Transitions*. 2022. <https://www.iea.org/reports/nuclear-power-and-secure-energy-transitions>

International Energy Agency. *Electricity Grids and Secure Energy Transitions*. 2023.
<https://www.iea.org/reports/electricity-grids-and-secure-energy-transitions>
International Energy Agency. *Renewables 2024: Global Overview*. 2024.
<https://www.iea.org/reports/renewables-2024/global-overview>
International Energy Agency. “Renewable Heat.” In *Renewables 2025*. 2025.
<https://www.iea.org/reports/renewables-2025/renewable-heat>
International Energy Agency. *The Path to a New Era for Nuclear Energy*. 2025.
<https://www.iea.org/reports/the-path-to-a-new-era-for-nuclear-energy>
Jarvis, Stephen, Olivier Deschenes, and Akshaya Jha. “The Private and External Costs of Germany’s Nuclear Phase-Out.” *Journal of the European Economic Association* 20, no. 3 (2022): 1311–1346.
<https://doi.org/10.1093/jea/jvac007>
Lelieveld, J., A. Haines, R. Burnett, C. Tonne, K. Klingmüller, T. Münzel, and A. Pozzer. “Air Pollution Deaths Attributable to Fossil Fuels: Observational and Modelling Study.” *BMJ* 383 (2023): e077784.
<https://doi.org/10.1136/bmj-2023-077784>



www.ia-forum.org

