

Designing Europe's Future

AI as a Force
for Good

Edited by
Francesco Cappelletti
Maartje Schulz
Eloi Borgne



Published by the European Liberal Forum. Co-funded by the European Parliament.

The views expressed herein are those of the authors alone.

These views do not necessarily reflect those of the European Parliament or the European Liberal Forum.

The European Liberal Forum (ELF) is the official political foundation of the European Liberal Party, the ALDE Party.

Together with over 50 member organisations, we work all over Europe to bring new ideas into the political debate, to provide a platform for discussion, and to empower citizens to make their voices heard. Our work is guided by liberal ideals and a belief in the principle of freedom. We stand for a future-oriented Europe that offers opportunities for every citizen. ELF is engaged on all political levels, from the local to the European. We bring together a diverse network of national foundations, think tanks and other experts. In this role, our forum serves as a space for an open and informed exchange of views between a wide range of different EU stakeholders.

© European Liberal Forum, 2025

Page layout: Altaïs

Cover Image: Freepik

Editors: Francesco Cappelletti, Maartje Schulz, Eloi Borgne

ISBN: 978-2-39067-106-0

ISSN (print): 2791-3880

ISSN (online): 2791-3899

DOI:10.53121/ELFS10

ELF has no responsibility for the persistence or accuracy of URLs for external or third-party internet websites referred to in this publication and does not guarantee that any content on such websites is, or will remain, accurate or appropriate.

Contents

Foreword	3
<i>Svenja Hahn</i>	
Introduction	5
<i>Maartje Schulz</i>	
 <i>Part 1</i>	
AI IN EUROPE	
Building What We Govern: The Public Purpose of AI and the 'Brussels Innovation Effect'	11
<i>Dr Antonios Nestoras</i>	
Is There Still a Brussels Effect for Artificial Intelligence?	21
<i>Charles Mok</i>	
Lists of Good Wishes? The Limits of the 'AI Brussels Effect' in Latin America	27
<i>Claudia del Pozo</i>	
The Fall of Clarity: AI Regulation and the Risk of European Decline	37
<i>Gian Marco Bovenzi</i>	
Building Europe's AI Capacity: A Partnership Made in Brittany	43
<i>Industry note by Philippe Guillotel (InterDigital)</i>	
 <i>Part 2</i>	
AI FOR SOCIETY	
Artificial Intelligence as a Socio-Technical System: What It Is and Where It May Take Us	51
<i>Francesco Cappelletti and Dr Francesco Goretti</i>	
AI for Productivity: Increasing Adoption Levels and Diffusing Technologies in the Workforce	63
<i>Dejan Ravšelj and Aleksander Aristovnik</i>	

AI and Energy Use: How to Wield a Double-Edged Sword in the Direction of Progress <i>Eloi Borgne</i>	73
Creating Opportunities: The Importance of Learning Design for AI Use in Education <i>Sarah K. Howard, Jo Tondeur, Charlotte Neubert, Jonas Hauck, and Richard Böhme</i>	81
CONCLUSION <i>By the editors</i>	93

Foreword

Svenja Hahn

ABOUT

Svenja Hahn is a Member of the European Parliament (Renew Europe Group) and President of the ALDE Party.



Europe stands at a crossroad. For too long, our debates around artificial intelligence have been dominated by fear and regulation. It is time we shift our focus from what we want to restrict to what we want to achieve. AI is not just a technological tool; it is a transformative force that can make our societies more prosperous, sustainable, and free – if we dare to embrace it.

Across the world, innovation is moving fast. Europe cannot afford to stand still while others race ahead. We have the talent, the creativity, and the entrepreneurial spirit to lead in AI, but we also need the right political and economic framework. That means cutting red tape, establishing a genuine single market for digital innovation, improving access to capital, and attracting the brightest minds to stay and build here in Europe.

A liberal vision for AI must be based on trust in people and in innovation, not in bureaucracy. We should empower our innovators instead of overburdening them with complex compliance. The European Union's recent shift towards promoting AI development is a welcome change, but words must now be followed by action.

The future of AI in Europe will not be written in regulation, but in imagination, entrepreneurship, and courage. Let us make Europe a place that does more than simply regulate technology, but one that uses technology to shape and lead the way.

This Study by the European Liberal Forum contributes to this liberal vision by offering concrete policy recommendations that champion innovation, empower entrepreneurs, and promote a human-centric approach to AI development across Europe.

Introduction

Maartje Schulz

ABOUT

Maartje Schulz is a Policy and Research Manager at ELF.



Travelling to the US earlier this year on a study trip with the Friedrich Naumann Foundation, I was struck by the tremendous excitement there around the potential of AI. This excitement came not only from the founders and CEOs of AI companies in San Francisco and in Silicon Valley, but from policymakers on both sides of the political spectrum. Despite their disparate party affiliations and professional backgrounds, they all seemed to be working together to make AI development a success in their country.

I would definitely not argue that the current situation in the US is a guiding light for Europe in every respect. Yet I do wish that my Europe could steal a bit of that appetite for risk-taking and unbridled experimentation and use it to forge a new European model of innovation – not least because our future competitiveness will depend in part on how well we do with advancing our digital economy. The unfortunate reality, however, is that we are lagging far behind the US and China in this area, as Draghi explains in his report.

But this is no reason for Europe to sit back in defeat or succumb to discouragement. We have the foundations, the tools, and the talent in Europe to improve our situation. We just need to get the mindset right first. That is why, in this Study – *Designing Europe's Future: AI as a Force for Good* – we aim to articulate an optimistic, liberal vision for AI. The Study gives voice and provides solutions to the questions that are most pressing at this pivotal moment: how do we go beyond just focusing on AI regulation in Europe? How do we innovate in Europe and steer AI in the right direction so that it can benefit our societies? And how do we have guardrails around AI to protect our values and the way we live – without stifling innovation?

The first chapters in the Study dive into Europe's approach to AI up to the present, reflecting on the EU AI Act from a liberal legal viewpoint and from outsider perspectives. They show that the AI Act, despite its admirable intentions, has introduced unclarity and high reporting burdens, which can make Europe less attractive to investors and innovators. Europe should therefore break away from the age-old 'Brussels Effect' and move towards something more inspiring: the 'Brussels Innovation Effect' – that is, from global rule setter to strategic technological leadership.

With the rise of AI, strategic autonomy will not be defined by market size or regulatory reach alone, but by mastery of such decisive levers of capability as compute, data, and talent. Rather than attempting to outpace private US firms at every turn or take refuge in protectionist defensiveness, Europe should concentrate on sectoral leadership and the intelligent diffusion of technology across society and the workforce. Even if we do not build it all ourselves, we can still try to be the best at diffusing this technology. The present Study highlights numerous paths towards this end.

The second part of the Study focuses on how AI can be a force for good in European society. The restoration of Notre-Dame cathedral is a memorable case: AI was employed to generate a digital twin of the damaged structure, to inform reconstruction, and to facilitate data management. AI can also be used as a supportive tool to increase our productivity, to enable better decision-making, and simplifying routine tasks. In the classroom, in the workplace, and in health services, we need to integrate AI smartly to augment our human experience – not to replace or diminish it.

Let us also build more thriving ecosystems: universities that work with start-ups and with venture capital that is ready to invest and with the support of (local) government. Silicon Valley provides an example, but we can also learn from the Dutch 'Triple Helix' innovation model where the private sector, government, and academia work together to develop new solutions for society.

AI is a challenge: it inspires both fear and hope. It is up to us liberals to own the future and steer AI into direction of positivity, prosperity, and inspiration.

Part 1

AI IN EUROPE

Building What We Govern: The Public Purpose of AI and the ‘Brussels Innovation Effect’

Dr Antonios Nestoras

<https://doi.org/10.53121/ELFS10> • ISSN (print) 2791-3880 • ISSN (online) 2791-3899

ABSTRACT

This chapter contends that Europe must evolve from its historic role as a global rule-setter – the so-called ‘Brussels Effect’ – towards one of strategic technological leadership: the ‘Brussels Innovation Effect’. In the era of AI, sovereignty will be defined not by market size or regulatory reach alone but by mastery of the decisive levers of capability: compute capacity, data collection and analysis, and the cultivation of talent, together with public infrastructures that turn them into durable advantages. Rather than attempt to outpace private US firms on every frontier or retreat into protectionist defensiveness, Europe should concentrate on sectoral leadership anchored in safe, shared, and world-class platforms. Through targeted investment, generative ethics, and purpose-built institutions like Civic Data Trusts, the EU can translate its values into operational systems that others adopt because they function, scale, and inspire trust. The objective is not regulation in isolation, but the construction of a working ecosystem in which advanced intelligence serves the public good and in which Europe leads by building the very systems it governs.

ABOUT THE AUTHOR

Dr Antonios Nestoras is the Founder and CEO of the European Policy Innovation Council (EPIC) and Senior Fellow at the European Liberal Forum.

FROM REFEREES TO ARCHITECTS: EUROPE'S AI INFLECTION POINT

Artificial intelligence is no longer a laboratory curiosity or a novelty on consumer screens. It has become a general-purpose governance technology, a foundational capability that is restructuring the grammar of power, reconfiguring the allocation of resources, and subtly transforming how people relate to the state, the market, and one another. As global AI investment neared \$200 billion in 2024, the emerging technologies have found uses in increasingly diverse contexts – frontier models now underpin applications ranging from clinical trials to national grid balancing.¹

1. Frontier AI models refer to highly advanced foundation models that possess capabilities significant enough to potentially pose severe risks to public safety. These models are at the cutting edge of AI development and are characterised by their ability to perform a wide range of tasks, often with minimal human intervention. McKinsey and Company (2024), *The state of AI in 2024: Investment, innovation, and impact*, McKinsey Global Institute; OECD (2024), *AI policy observatory: Global AI investment trends 2024* (Paris: Organisation for Economic Co-operation and Development).

AI-based systems are already accelerating the design of drugs, allowing new materials to be discovered in months (rather than decades), refining forecasts of energy demand, compressing software development cycles from months into days, and embedding themselves into the interfaces through which societies work, learn, transact, and govern.

The decisive question of this century is no longer whether AI should be used – that is already settled – but what kind of society it will produce if we fail to shape it. In the absence of deliberate design, the organising logic of the technology will be written elsewhere, by actors whose incentives and value systems may diverge sharply from our own.

The decisive question of this century is no longer whether AI should be used – that is already settled – but what kind of society it will produce if we fail to shape it.

Europe is uniquely positioned to take up this challenge. It is the only major political bloc that combines advanced technological capability with the institutional legitimacy needed to align its measures with a conception of the public good. Its layered system of governance, characterised by an interplay between the European Commission's regulatory competence, the Parliament's democratic mandate, and the Council's coordination of national interests, gives the Union a rare ability to transform technological principles into binding political commitments. From the Digital Europe Programme² and Horizon research frameworks³ to the nascent European AI Office,⁴ Europe already possesses institutional instruments that can be repurposed to convert regulatory ambition into strategic action.⁵ Through its combination of democratic accountability and technocratic reach, Europe can ensure that AI serves citizens, not the other way around.

A history of rights-based governance, a tradition of social partnership, and a proven ability to translate norms into enforceable frameworks give Europe an inheritance unlike any other region. Yet influence is not the same as leadership. If Europe confines itself to refereeing technologies invented elsewhere, we risk becoming the authors of the AI rulebook only to discover that others have written the code.

Until now, Europe's claim to digital influence has rested on the 'Brussels Effect' – the capacity to project regulatory standards beyond its borders.⁶ Such influence is real and has guided global practice in many domains. But rules have never been enough to secure technological leadership. What we may term the 'Brussels Innovation Effect' is the necessary next step: not the export of norms alone but the export of working, values-driven systems – the platforms, institutions, and public goods that embody those norms in practice. By capitalising upon the features that make it distinctive, the European Union can deliver technologies and regulations that inspire emulation and international adoption, and which, crucially, embody liberal principles. Only by moving from referee to architect can Europe ensure that its vision of technology for the people becomes a living, global model.

2. European Commission (2023a), *Digital Europe Programme: Work programme 2023–2024* (Brussels: D-G Connect).

3. European Commission (2024a), *Horizon Europe strategic plan 2025–2027* (Brussels: D-G Connect).

4. European Commission (2024b), *Setting up the European AI Office: Mandate and strategic objectives* (Brussels: D-G Connect).

5. European Commission (2024c), *Proposal for a regulation laying down harmonised rules on artificial intelligence (AI Act)*, COM(2021) 206 final (as amended).

6. A. Bradford (2020), *The Brussels Effect: How the European Union Rules the World* (Oxford: Oxford University Press).

WHEN RULES ARE NOT ENOUGH: THE LIMITS OF THE BRUSSELS EFFECT

For more than a decade, the so-called Brussels Effect has been Europe's signature mode of influence in the digital realm: legislate to the highest standard at home and leave it to global firms to implement the same standards abroad. Reliant upon regulatory sophistication and market size, the Brussels Effect has had a significant impact in domains like data privacy. The General Data Protection Regulation (GDPR), for example, not only transformed European practice; it recalibrated corporate behaviour and legal frameworks far beyond our borders. In this way, Europe secured a voice in global technology governance disproportionate to its share of frontier innovation.

But in AI, regulatory export is no longer synonymous with leadership. Rules are downstream of capabilities.⁷ In 2024, Europe accounted for less than 7% of global private AI investment and under 5% of frontier-scale compute capacity, while the United States and China consolidated a near-duopoly over the infrastructure of intelligence.⁸ Unless Europe can design, train, and operate such systems on its own terms (from foundation models to advanced semiconductors to the strategic datasets on which they learn) the power to define their use will steadily erode. Influence becomes conditional, exercised only with the acquiescence of those who control the material substrate.

The deeper danger facing Europe is self-deception. The Brussels Effect can tempt us into mistaking codification for creation, as though inscribing the limits of others' inventions were equivalent to inventing ourselves. Taking refuge in the Brussels Effect, we may neglect the harder demands of technological statecraft: mobilising capital at continental scale, overhauling public procurement, cultivating scarce technical talent, and building institutions capable of deploying AI in service of public purpose. Without such measures, Europe could be remembered not as an architect of the AI order but as its meticulous scribe – or the custodian of a rulebook whose most consequential provisions were authored elsewhere.

To rest on the laurels of the Brussels Effect is no longer sound strategy – we must strive to bring about a Brussels Innovation Effect.⁹ Its test will not be whether Europe can export the *principles* of democratic AI, but whether it can export the *operational realities* – that is, the platforms, infrastructures, and institutional designs through which those principles become lived experience. This is a higher form of leadership rooted not merely in caution, but in the confidence that technology, conceived and governed with democratic purpose, can expand the agency of those who build it and the societies that choose to adopt it.

CAPABILITY AS SOVEREIGNTY IN THE AGE OF INTELLIGENCE

If the Brussels Innovation Effect is to be more than a rhetorical flourish, it must rest on hard capability. In the age of AI, sovereignty will be measured less by the cartography of borders, the tonnage of armies, or even the scale of GDP, and more by command over the decisive levers of machine intelligence itself. Compute capacity, strategic datasets, and advanced technical talent are not peripheral to power; they are its primary instruments. They deter-

In the age of AI, sovereignty will be measured less by the cartography of borders, the tonnage of armies, or even the scale of GDP, and more by command over the decisive levers of machine intelligence itself.

7. M. Leonard and J. I. Torreblanca (2022), *The geopolitics of technology: How the EU can become a global player*, European Council on Foreign Relations (ECFR) Policy Brief.

8. J. Bjerkem and M. Täger (2021), *From the Brussels Effect to the Innovation Effect: Europe's new global role in tech regulation*, European Policy Centre Discussion Paper.

9. P. Butcher and V. van Roy (2023), *AI and Europe's Competitiveness Gap: The Missing Link between Regulation and Capability*, Joint Research Centre (JRC) Technical Report.

mine who can design and operate the systems that will govern energy flows, stabilise financial markets, manage public health, and safeguard national security.

Europe cannot afford to treat these domains as auxiliary concerns. The United States now trains models on compute clusters exceeding 10^{25} floating-point operations; the largest public facilities in Europe remain an order of magnitude smaller. Compute is not simply hardware, it is the commanding terrain of statecraft in the twenty-first century, the vantage point from which all other technological campaigns are fought. Data is not the passive residue of economic life, but the raw material from which learning systems extract their intelligence. Curated, shared, and governed with care, it is a public asset as vital as clean water or reliable electricity. Talent is not a generic labour input; it is the irreplaceable reservoir of human capacity to imagine, engineer, and direct the tools that will shape our collective future.

Without control over these levers, even the most principled regulatory regimes will be hostage to those who possess them. But if Europe secures them, it can anchor an AI paradigm that is competitive without being extractive, open without being naïve, and innovative without compromising the dignity of the people it serves. This conception of capability aligns with the EU's own strategic trajectory so far: from the Digital Decade targets for compute and connectivity to the AI Act's framework for trustworthy intelligence. These policy instruments have the potential to translate sovereignty from an abstract value into a measurable set of technological assets.

The outlines of this potential are already visible. Paris-based Mistral AI, approaching a valuation of \$10 billion in 2025, has demonstrated that European-built foundation models can perform at the global frontier. Berlin's Aleph Alpha is embedding explainable, multilingual AI into the machinery of public administration. Italy's Minerva 7B fuses local linguistic and legal contexts into high-performing systems. Even AlphaFold, though now global, emerged from DeepMind's early work in the United Kingdom and continues to drive breakthroughs in European laboratories.

These are not curiosities; they are signals. They show that when Europe invests strategically in compute, data, and talent, it can set the benchmarks that others will adopt. But isolated victories do not amount to sovereignty. To transform potential into enduring advantage, Europe must build a public AI infrastructure that converts scattered capabilities into a coherent, shared foundation for innovation.

THE COMMONS OF INTELLIGENCE: BUILDING PUBLIC AI INFRASTRUCTURE

If compute, data, and talent are the basic levers of sovereignty, public AI infrastructure is the civic architecture that fixes that sovereignty in place. Without it, Europe's capabilities will remain dispersed and scattered across national projects, corporate ventures, and research enclaves too fragmented to influence the continent's strategic trajectory. With it, Europe can make its leadership deliberate and sectoral: not by joining the indiscriminate race to dominate every AI application, but by constructing safe, open, and trusted platforms in domains where its values and comparative strengths converge.¹⁰

10. European Commission (2023), *High-Performance Computing and Artificial Intelligence: Building Europe's digital commons* (Brussels: D-G Connect).

This is neither a call for a monolithic ‘state AI’ to supplant private enterprise nor an argument for emulating Silicon Valley’s scale-at-all-costs model. It is the recognition that certain infrastructural capabilities, such as high-performance compute, sovereign datasets, and secure model-testing environments, must exist as public goods if innovation is to remain both competitive and aligned with democratic purpose.¹¹ In their absence, European actors will innovate only with the permission, and under the conditions, of foreign or unaccountable gatekeepers.

History offers instructive analogies: the railways and power grids of the nineteenth century, or the highways, universities, and public broadcasting systems of the twentieth. Each was more than a service; it was a generative foundation on which private enterprise, cultural exchange, and political legitimacy could flourish. In the twenty-first century, the equivalent is a commons for intelligence.¹²

Such a commons could take multiple forms:

- **Compute commons:** A federated network of AI-optimised supercomputers and data centres across the Single Market, providing compute credits to SMEs, universities, and public agencies for projects with demonstrable social or economic value.
- **Civic data trusts:** Democratically governed repositories of high-quality, multilingual, domain-specific datasets in sectors such as health, mobility, energy, and culture, operating under strict rights and benefit-sharing frameworks.
- **Open model gardens:** Curated collections of base and fine-tuned models designed for European priorities, audited for safety, transparency, and performance before deployment.
- **Evaluation and compliance sandboxes:** Structured environments in which innovators can test systems against EU safety, security, and rights standards without stalling their path to adoption.

By treating these infrastructures as shared endowments, Europe can enable its researchers, entrepreneurs, and public institutions to operate at the frontier without attempting to mirror – or directly contest – the scale strategies of US and Chinese private giants. This paves the way for a form of leadership that is not measured by who trains the largest model, but by who builds the most trusted and effective systems in the sectors that matter most: health, sustainable industry, multilingual public services, resilient infrastructure.

This is the material foundation of the Brussels Innovation Effect: leadership not by regulatory fiat but by example – that is, by designing and deploying systems that embody European values, proving their worth at scale, and allowing others to adopt them because they have shown themselves to be indispensable.

ETHICS AS THE ENGINE OF RESPONSIBLE SPEED

Public AI infrastructure furnishes the means to act; ethics determines the manner of action. In Europe’s political tradition, ethics is not a decorative afterthought applied once the engineering is done, but a constitutional principle embedded in the very act of building.¹³ This heritage is one of the continent’s greatest strategic assets – yet it carries a latent risk. If ethics ossify into a reflex to defer, to delay, or to insulate ourselves from the burden of decision, they cease to be a compass for innovation and become an excuse for inaction.

11. Joint Research Centre (JRC) (2024), *AI, data and the European commons: Infrastructures for democratic innovation* (Brussels: D-G Connect).

12. E. Morozov (2020), ‘Digital Socialism? The Calculation Debate in the Age of Big Data’, *New Left Review*, 116/117, 33–67, <https://doi.org/10.64590/tt3>.

13. European Commission (2021), *Ethics guidelines for trustworthy AI*. High-Level Expert Group on Artificial Intelligence (Brussels: D-G Connect).

The uncomfortable truth is that in the age of AI, refusing to build can be as damaging as building recklessly. A society that abstains from developing certain capabilities does not prevent their emergence; it simply relinquishes them to actors whose incentives may be indifferent or openly hostile to democratic values. The responsibility, therefore, is not merely to avoid harm, but to pursue good, innovating in ways that validate human dignity, reinforce social trust, and strengthen the resilience of democratic institutions.¹⁴

Europe's moral authority in technology has long come with an internal tension. Ethical leadership has given the Union global legitimacy, but it has also, at times, constrained its industrial competitiveness.¹⁵ The challenge for the AI era is not to abandon this ethical core, but to make it operational and demonstrate that principles and performance can reinforce one another. The EU's framework for trustworthy AI already embodies this balance in theory; the task now is to put it into practice, ensuring that ethical safeguards do not become structural bottlenecks but catalysts for safe, rapid deployment.

This demands a recalibration of standards, a move from the defensive maxim 'first, do no harm' toward the generative imperative 'first, design for dignity and then deploy'. Bias detection, privacy protection, and explainability must be integrated from the outset. Governance should be iterative and adaptive, with deployments monitored and refined transparently. Oversight must be proportionate to risk so that beneficial applications in health, energy, or public services pass from prototype to impact without being stalled by years of attrition.

Seen in this light, ethics do not restrain innovation; they structure it. Speed and safety are not opposing forces but mutually reinforcing disciplines. If Europe can achieve this equilibrium – advancing rapidly while remaining anchored in legitimacy – it will offer the world a living demonstration that advanced intelligence can be governed in the service of the people without yielding either to reckless accelerationism or paralysing caution.

STRATEGIC NICHES: LEADING WHERE VALUES AND CAPACITY CONVERGE

Leadership in the age of AI will not accrue to those who seek to dominate every technological frontier. It will belong to those who carefully choose sectors where capability, values, and strategic interest align, and where success can set de facto standards others are compelled to follow. For Europe, this means resisting the lure of imitating the scale strategies of Silicon Valley or Shenzhen, and instead concentrating on domains where public trust, technical excellence, and governance acumen can be fused into an unrivalled proposition. Four such domains stand out.

Health and life sciences: Europe's integrated public health systems, dense clinical networks, and vast biobank resources form an unmatched foundation for AI-driven diagnostics, personalised medicine, and accelerated drug discovery – all under governance regimes that safeguard patient rights. Sweden's AI-supported breast cancer screening, which in trials cut false positives by 44%, offers a glimpse of what a European-led standard could look like: innovation without the erosion of dignity.

Sustainable industry and mobility: The continent's engineering heritage and manufacturing depth provide fertile ground for AI in predictive maintenance, energy optimisation, and circular logistics. Siemens has already demonstrated that predictive maintenance can reduce costs in European plants by up to 30%. By aligning AI innovation with decarbonisation and industrial resilience, Europe can prove that sustainability and competitiveness are not rival imperatives but mutually reinforcing ones.

14. L. Floridi (2019), 'Establishing the Rules for Building Trustworthy AI', *Nature Machine Intelligence*, 1(6), 261–262. <https://doi.org/10.1038/s42256-019-0055-y>.

15. M. Veale and F. Z. Borgesius (2021), 'Demystifying the Draft EU Artificial Intelligence Act', *Computer Law Review International*, 22(4), 97–112, <https://doi.org/10.48550/arXiv.2107.03721>.

Multilingual public services and culture: With 24 official languages and a cultural heritage spanning centuries, Europe is uniquely qualified to develop AI that supports linguistic diversity, preserves cultural patrimony, and fuels creative industries. The Language Bank of Finland, which is curating more than 1,200 multilingual datasets, exemplifies how such resources can become force multipliers for public services, cultural exports, and democratic participation.

Resilient infrastructure and energy systems: In an era of climate volatility, the ability to keep infrastructure stable is strategic power. European grid operators are already deploying AI to balance renewables and forecast demand; Belgium's Elia has decreased forecasting errors for solar and wind generation by up to 40%. Scaling such capabilities across water systems, transport, and public works could make resilience itself a European export.

In each of these niches, Europe's advantage lies not in scale alone, but in its capacity to embed technology within trusted governance frameworks and safe infrastructure.

GOVERNING INTELLIGENCE: INSTITUTIONS FOR THE AI ERA

If compute, data, and talent are the raw endowments of sovereignty in the age of artificial intelligence, institutions are the constitutional machinery that converts those endowments into lasting capability. They are the repositories of collective intent, the means by which a society determines not only what it can do, but what it ought to do. In the last century, Europe built such machinery in finance, education, and integration: the European Central Bank to anchor a currency, Erasmus to knit together a generation, the Single Market to fuse disparate economies into one. The AI era will demand an act of institutional imagination on a comparable scale.

The challenge is to develop bodies that can govern systems of unprecedented complexity without stifling them. This does not mean multiplying bureaucracies; it means creating entities that are agile, authoritative, and visibly in the service of the public.¹⁶ Early models of this kind are now appearing elsewhere: the United Kingdom's AI Safety Institute, dedicated to testing the most capable models; the United States' AI Safety Institute Consortium, convening public and private expertise on a national scale. Europe should study these examples but adapt them to its own ethos of rights-based governance and social trust.

The architecture could take many forms, but certain functions will be indispensable. An AI Ombudsman for example, designed to be independent, accessible, and empowered to address both individual grievances and systemic risks, would give people a direct line to redress. An AI Safety Agency, with the authority and capacity to evaluate models rigorously and to scan for emerging risks, could define the safety baselines that would become de facto global standards. Civic Data Trusts, curating high-value datasets under democratic governance, could transform data sovereignty from a defensive posture into a source of shared advantage.

These examples are not prescriptions but starting points. Their eventual configuration would depend on political will, technological context, and the evolving expectations of the public. What matters is the underlying philosophy: that in the AI century, legitimacy will rest not only on the capacity to regulate, but on the capacity to build institutions that embody a civilisation's values while commanding its most powerful tools. To govern intelligence, in the end, is to govern ourselves.

16. P. Cihon, M. M. Maas, and L. Kemp (2021), 'Should Artificial Intelligence Governance Be Centralised? Design Lessons from History', *Global Policy*, 12(S5), 20–32, <https://doi.org/10.48550/arXiv.2001.03573>.

FROM NORMS TO CAPABILITIES: TOWARDS THE BRUSSELS INNOVATION EFFECT

For two decades, the so-called Brussels Effect has projected European influence by exporting regulatory norms far beyond our borders. That achievement is real, but it was never, on its own, enough. As AI applications proliferate, the gap between those who set rules and those who build the systems those rules govern will only widen. Rules without the underlying capacity to design, deploy, and steward the technology they address will, over time, lose their force. Sovereignty in this century will belong to those who can pair capability with legitimacy.

The Brussels Innovation Effect demands precisely this pairing: Europe's regulatory reach combined with infrastructure, expertise, and governance frameworks that others adopt voluntarily because they function, scale, and inspire trust. Leadership here does not mean chasing every frontier breakthrough of private US firms, nor erecting protectionist barriers that constrain our own innovators. It means sectoral leadership: anchoring excellence where our comparative advantages are strongest and embracing experimentation as a mode of governance. And ensuring these domains rest on safe, open, and world-class infrastructure. Europe's strength lies not only in codifying norms but in learning from practice, iterating its governance in real time.

EU leadership can be made visible through living institutions: an independent AI Ombudsman; an AI Safety Agency; and Civic Data Trusts, as a start. These can serve as proofs of concept for integrating advanced intelligence into democratic societies without succumbing either to laissez-faire opacity or to centralised overreach. If Europe can institutionalise experimentation by allowing trusted pilots to mature into shared standards, then it will turn governance itself into a tool of innovation. Paired with strategic, sector-specific leadership, a distinctly European model of AI governance can be realised: technology for the people, operationalised at scale. Others would follow not from obligation, but from recognition – because it works, and because it shows that intelligence can serve both prosperity and freedom.

REFERENCES

- Bjerkem, J., & Täger, M. (2021). *From the Brussels Effect to the Innovation Effect: Europe's new global role in tech regulation*. European Policy Centre Discussion Paper.
- Bradford, A. (2020). *The Brussels Effect: How the European Union Rules the World*. Oxford: Oxford University Press.
- Butcher, P., & van Roy, V. (2023). *AI and Europe's Competitiveness Gap: The Missing Link between Regulation and Capability*. Brussels: Joint Research Centre (JRC) Technical Report.
- Cihon, P., Maas, M. M., & Kemp, L. (2021). 'Should Artificial Intelligence Governance Be Centralised? Design Lessons from History'. *Global Policy*, 12(S5), 20–32. <https://doi.org/10.48550/arXiv.2001.03573>.
- European Commission (2021). *Ethics guidelines for trustworthy AI*. High-Level Expert Group on Artificial Intelligence. Brussels: Directorate-General for Communications Networks, Content and Technology (D-G Connect).
- European Commission (2023a). *Digital Europe Programme: Work programme 2023–2024*. Brussels: D-G Connect.
- European Commission (2023b). *High-Performance Computing and Artificial Intelligence: Building Europe's digital commons*. Brussels: D-G Connect.
- European Commission (2024a). *Horizon Europe strategic plan 2025–2027*. Brussels: D-G Connect.
- European Commission (2024b). *Setting up the European AI Office: Mandate and strategic objectives*. Brussels: D-G Connect.
- European Commission (2024c). *Proposal for a regulation laying down harmonised rules on artificial intelligence (AI Act)*. COM(2021) 206 final (as amended).
- Floridi, L. (2019). 'Establishing the Rules for Building Trustworthy AI'. *Nature Machine Intelligence*, 1(6), 261–262. <https://doi.org/10.1038/s42256-019-0055-y>.
- Joint Research Centre (JRC) (2024). *AI, data and the European commons: Infrastructures for democratic innovation*. Brussels: Publications Office of the European Union.
- Leonard, M., & Torreblanca, J. I. (2022). *The geopolitics of technology: How the EU can become a global player*. European Council on Foreign Relations (ECFR) Policy Brief.
- McKinsey & Company (2024). *The state of AI in 2024: Investment, innovation, and impact*. McKinsey Global Institute.
- Morozov, E. (2020). 'Digital socialism? The calculation debate in the age of big data'. *New Left Review*, 116/117, 33–67. <https://doi.org/10.64590/tt3>.
- OECD (2024). *AI policy observatory: Global AI investment trends 2024*. Paris: Organisation for Economic Co-operation and Development.
- Veale, M., & Borgesius, F. Z. (2021). 'Demystifying the Draft EU Artificial Intelligence Act'. *Computer Law Review International*, 22(4), 97–112. <https://doi.org/10.48550/arXiv.2107.03721>.

Is There Still a Brussels Effect for Artificial Intelligence?

Charles Mok

<https://doi.org/10.53121/ELFS10> • ISSN (print) 2791-3880 • ISSN (online) 2791-3899

ABSTRACT

Europe's non-coercive form of global influence on technology governance faces new challenges, and opportunities in the world of artificial intelligence regulations and governance. As the United States and China pursue divergent models of competition and control, Europe must evolve from exporting regulation to exercising genuine governance. The challenge is to transform regulatory strength into strategic capability, while balancing human rights, innovation, and digital sovereignty. By advancing a new Brussels Agenda grounded in values, institutional coherence, and multi-stakeholder collaboration, Europe can reaffirm its global role, demonstrating that ethical governance and technological ambition don't need to be opposing forces in the age of intelligent systems.

ABOUT THE AUTHOR

Charles Mok is a Research Scholar with the Global Digital Policy Incubator, Stanford University, and a former entrepreneur and Legislative Councillor for IT in Hong Kong (2012–2020).

INTRODUCTION

Over the years, global technology governance has come to be seen as a triumvirate, led by the three leading markets and technology powers in the world – the United States, China, and Europe. In her book *Digital Empires: The Global Battle to Regulate Technology* Anu Bradford presented a framework for how these major powers have moulded our digital world order through their different regulatory models and philosophies. This 'triumvirate' framing still matters in 2025, despite growing multipolarity, because these three powers continue to set the pace and define the boundaries of digital governance, each shaping not only their domestic landscapes but also influencing global norms, standards, and expectations in the age of AI.

The United States has advanced a market-driven model, prioritising innovation by autonomous corporations (especially those pertaining to so-called Big Tech) and minimising regulation (to the point of its near non-existence at the federal level). China is driven by the party state, consistently seeking to centralise its singular control, leverage industrial policy support for sectors and firms, and maximise the use of surveillance to protect and preserve its national ruling power.

The European model, on the other hand, is based on respect for human rights and dignity, personal privacy, and democratic accountability. It is backed by robust legal frameworks, typically operating through legislative instruments and directives from the European Union. In a 2012 paper, Bradford described the success and influence of the European model as the 'Brussels Effect', referring to the continent's ability to unilaterally impact global technology regulations through the sheer size

of its market forces, rather than by coercion or diplomacy (e.g., China's affinity for forming multi-lateral bodies).¹ The General Data Protection Regulation (GDPR) served as a prime example of how Europe's market forces alone could make EU regulatory practice the de facto gold standard for the rest of the world: numerous multinational firms opted to extend their European compliance to other parts of the world, and many governments followed Europe's legislative lead, introducing regulations patterned after its data privacy framework.

Then artificial intelligence entered the scene. With the launch of ChatGPT by OpenAI in December 2022, powerful AI tools (mostly in the form of large language machine learning models) were suddenly in the hands of companies, institutions, and individuals. Governments across the world began to grapple with the question of how best to address the implications of AI. In November 2023, the first AI Safety Summit was convened at Bletchley Park, United Kingdom, with government and industrial leaders coming together to discuss the safety of AI and possible regulatory directions. The EU wasted no time in establishing its own legal framework: the Artificial Intelligence Act (AI Act) came into effect in August 2024, largely focused on defining levels of risk and setting transparency requirements for developers.

In this new AI landscape, the strength of the Brussels Effect is not as apparent as before, when the targets of regulation were digital services such as e-commerce and social media platforms. In this chapter, I will first review the governance and regulatory directions taken by United States and China so far and compare them with the European approach, and then examine the new factors and constraints faced by nations in relation to AI. Finally, I will suggest some potential ways for Europe to respond to the current situation.

In this new AI landscape, the strength of the Brussels Effect is not as apparent as before, when the targets of regulation were digital services such as e-commerce and social media platforms.

THE US APPROACH – PIVOTING FROM SAFETY TO 'AMERICA FIRST'

The US government's first AI policies were developed well before 'frontier AI' and 'artificial general intelligence' became household terms. In 2016, the Obama White House released a report on future applications and considerations for AI, emphasising the deployment of AI for social good, fairness, safety, and accountability.² The report favoured adaptive regulations and the building of a skilled AI workforce, supported by federal funds for AI research. While advocacy for AI research continued, the first Trump administration stressed economic competitiveness, technical leadership, national security, and loosened regulation, as exemplified in the executive order 'Maintaining American Leadership in Artificial Intelligence' of February 2019.³

By the time of Biden's executive order on AI (EO 14110), issued in October 2023, the global AI race was already in full swing.⁴ The order can be viewed as both an extension and amalgamation of the approaches espoused by the two previous administrations: it followed Obama's positive, cautionary

1. A. Bradford (2012), 'The Brussels Effect', *Northwestern University Law Review*, 1 December, <https://northwestern-lawreview.org/issues/the-brussels-effect/>.

2. E. Felton and T. Lyons (2016), 'The Administration's Report on the Future of Artificial Intelligence', 12 October, <https://obamawhitehouse.archives.gov/blog/2016/10/12/administrations-report-future-artificial-intelligence>.

3. Executive Order 13859 (2019), 'Maintaining American Leadership in Artificial Intelligence', 11 February, <https://www.federalregister.gov/documents/2019/02/14/2019-02544/maintaining-american-leadership-in-artificial-intelligence>.

4. Executive Order 14110 (2023), 'Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence', 30 October, <https://www.federalregister.gov/documents/2023/11/01/2023-24283/safe-secure-and-trustworthy-development-and-use-of-artificial-intelligence>.

vision for AI (innovative but safe, secure, trusted, and respecting of privacy and civil rights), while echoing the importance Trump placed on global competition, leadership, and standards setting. Even though the order devoted significant attention to implementing safety test reporting requirements, it fell short of proposing any federal legislation on commercial or research AI activities.

On 20 January 2025, the first day of his second term, President Donald Trump revoked Biden's EO 14110 and replaced it with his own EO 14179, 'Removing Barriers to American Leadership in Artificial Intelligence'. Six months later, on 23 July, Trump released his comprehensive strategy for AI, 'Winning the Race: America's AI Action Plan', a 28-page long document with a clear objective: to prevail in the country's competition with China.⁵ This iteration of US strategy is in part a return to Trump's emphasis on innovation, economic dominance, and removing regulatory and permitting barriers for businesses – such as the previous administration's requirements for test reporting high-risk AI models – with an additional 'MAGA Republican' emphasis on 'protecting free speech' and 'American values'. The absence of a binding federal AI law in the US, where the course of innovation is largely left to market forces and corporate discretion, stands in stark contrast to the European Union's legislative approach, which seeks to proactively shape AI development through comprehensive, binding regulation rooted in fundamental rights and risk management.

THE CHINESE APPROACH – 'SHARING' BY MULTILATERALISM

On 26 July, some three days after the US White House put forth its AI action plan, Chinese Premier Li Qiang announced the country's 'Action Plan on Global Governance of Artificial Intelligence' at the 2025 World AI Conference in Shanghai.⁶ Interestingly, the legislation was not presented as an action plan for China alone: it was about *global* governance and the path forward for the international community as a whole, in sharp relief with the 'America First' ethos of its US counterpart. The Chinese plan, containing only 13 points, was shrewdly designed to be succinct and effective in highlighting a willingness to cooperate and share success with other countries. It explicitly avoids stating the country's competitive aspirations in the AI race.

Compared with the US AI action plan's unabashedly self-centred objectives (e.g., to 'meet global demand for AI by exporting [the United States'] full AI technology stack [...] to countries willing to join America's AI alliance') China's narrative is about sharing its technologies to support other countries' AI development, and the use and diffusion of AI for their various industrial sectors. Where the United States' plan talks about eliminating 'bias' references such as 'misinformation', 'diversity, equity, and inclusion', and 'climate change' from the AI Risk Management Framework of the National Institute of Standards and Technology (NIST), China's rhetoric embraces a more 'global' version of shared values (such as tackling algorithmic bias, eliminating discrimination and prejudice, and promoting, protecting, and preserving the diversity of the AI ecosystem).

Of course, in reality, China's 'global' AI action plan still aims to secure the country a leadership role, particularly in the area of AI government. Yet it is about everything that the US plan is not: multilateral collaboration, joint or cooperative technology development, safety, rule/standard setting, and the creation of a responsible, sustainable environmental strategy to deal with the power demands made by AI infrastructure. Here, China is clear about its support for multilateralism in global AI co-operation and governance, basing its plan on the United Nations' Pact for the Future and its Global Digital Compact annex.⁷ It also proposes that the United Nations' International Telecommunication

5. Executive Office of the President (2025), 'Winning the Race: America's AI Action Plan', 23 July, <https://www.white-house.gov/wp-content/uploads/2025/07/Americas-AI-Action-Plan.pdf>.

6. Central People's Government of the People's Republic of China (2025), '人工智能全球治理行动计划 [Action Plan on Global Governance of Artificial Intelligence]', https://www.gov.cn/yaowen/liebiao/202507/content_7033929.htm. English translation: <https://charlesmok.substack.com/p/chinas-action-plan-on-global-governance>.

7. United Nations (2024a), 'What is the pact for the future?', <https://www.un.org/en/summit-of-the-future/pact-for-the-future>; United Nations (2024b), 'Global digital compact', <https://www.un.org/en/summit-of-the-future/global-digital-compact>.

Union (ITU) be given a leadership role in setting AI standards – affirming, in other words, the need for a top-down approach when shaping and governing the technologies of the future, rather than the existing bottom-up multi-stakeholder model, where researchers, the industry, academics, users, and civil societies are empowered and active participants. For years, this has been the vision that China has consistently worked towards.

Unfortunately, the current US administration shows limited interest in pursuing international AI co-operation. Pillar III of US AI action plan, entitled ‘Lead in International AI Diplomacy and Security’, only mentions ‘exporting American AI to allies and partners’. That is, despite what its title implies, the focus is on securing more business and export control for the United States (rather than engaging in any kind of diplomacy, let alone leading global governance). Unlike the Biden administration, which at least expressed verbal support for the multi-stakeholder model in its technology and internet policy initiatives (such as the 2022 Declaration for the Future of the Internet), the Trump administration leaned heavily into criticism of international governance bodies in its AI action plan.⁸ It accused such entities of being subject to ‘Chinese influence’, advocates of burdensome regulations and vague ‘codes of conduct’ promoting cultural agendas that did not align with American values. These claims notwithstanding, American policy has demonstrated a lack of patience or motivation to counter Chinese influence from within the multilateral system. Nor has the US taken any substantive action to prevent the existing multi-stakeholder model from being taken over by multilateralism, led by China.

THE EUROPEAN OPPORTUNITY:

FROM THE BRUSSELS EFFECT TO A NEW BRUSSELS AGENDA

The success of the Brussels Effect for digital services regulation was largely the result of the EU’s large Single Market and the political feasibility of crafting a uniform regulatory framework for all its Member States. But the Brussels Effect, as its name probably implies, is relatively passive. It relies on the ability of the Single Market (by virtue of its size and commercial appeal) to induce global technology and digital services firms – the largest of which are not even based in Europe – to submit to and comply with EU regulations. The European Parliament can pass digital laws, but it is much harder for the Council of Europe to adopt a single industrial policy. Yet in today’s AI race, the United States and China are prioritising competition in their geopolitics and industrial policies, from granting subsidies to exercising export controls. In this more fragmented environment, Europe is increasingly reduced to a collection of individual states, and the Brussels Effect has become less effective in the realm of AI.

Moreover, the global AI industry is very different from that of digital services from past decades (which were dominated mainly by American Big Tech firms and, to a lesser extent, Chinese companies). Across the world, the concept of digital sovereignty has taken hold in debates around AI (whether with respect to artificial general intelligence, the development of large language models (LLMs), or the building of hyper-scaler computing and datacentre infrastructure). Countries want their own LLM models trained on their own languages and cultures, their own national AI unicorns (privately held startups valued at over \$1 billion), and their own data sovereignty restrictions. Even within Europe, concerns about over-regulation are now being voiced by actors in both industry and government, with some (especially from large European economies like France and Germany) contending that the EU AI Act is excessively burdensome on innovative startups and smaller companies.⁹

8. US Department of State (2022), ‘Declaration for the Future of the Internet’, 28 April, <https://www.state.gov/declaration-for-the-future-of-the-internet>.

9. A. Spies (2025), ‘Europe realizes that it is over-regulating AI’, *American–German Institute*, 12 March, <https://americangerman.institute/2025/03/europe-realizes-that-it-is-overregulating-ai/>.

In order for Europe to capture the economic, social, and technological opportunities of the AI revolution, but at the same time extend the sort of soft influence characteristic of the Brussels-Effect, it needs to be proactive.

In order for Europe to capture the economic, social, and technological opportunities of the AI revolution, but at the same time extend the sort of soft influence characteristic of the Brussels Effect, it needs to be proactive. What is needed is not an *effect* but an *agenda* – a Brussels Agenda for AI, representing a continuation of and improvement on the Brussels Effect for digital services. The following three points can form the basis for this forward-looking approach:

Maintaining the Brussels core of values, rights, and rule of law

The essence of the European model is its respect for human rights, personal privacy, and dignity in the digital realms – this must be preserved. In a world where these fundamental values are often said to be ‘incompatible’ with American values, Europe has an important opportunity (and indeed a responsibility) to hold the line for humanity. If China can advocate for creating ‘an inclusive, open, sustainable, fair, safe, and reliable digital and intelligent future for all’, based on the ‘goals and principles of serving the people, respecting sovereignty, development-oriented, safe and controllable, fair and inclusive, and open cooperation’ there is no reason why Europe cannot be an equally viable, or better, alternative for the rest of the world.

The size and potential of the European market, a crucial part of the Brussels Effect, is no different for AI. While the United States and China are usually considered the leaders in AI technology, Europe (especially if the United Kingdom can be enlisted as a ‘free agent’ on the same team) is still a force to be reckoned with, producing critical scientific research and innovative companies. In the AI context, Europe’s market appeal can go beyond its aggregate size. A case can be made that the EU market is more open to entry and cooperation by developers and companies from the rest of the world than China (and nowadays, perhaps the United States), and that it is better protected by the rule of law.

From regulations to governance

The global propagation of the Brussels Effect began with Europe’s regulatory regimes on data privacy and digital services. While the EU should refine its AI regulatory regimes, more focus ought to be placed on *governance* and not merely regulations, which more and more people may see as burdensome and adverse to innovation. Indeed, governance (which includes setting global standards and rules) is now more important than ever due to the fragmentation caused by AI and digital sovereignty. Europe can fill the void left by the US, which is now beset by authoritarian forces seeking to take up and hold onto leadership. Europe must recognise that it may no longer be the preferred model, not just compared to China, but even to the United States. Increasingly, governments and innovative startups, including American ones, are seeking alternatives that demonstrate stronger commitments to human rights and safety.¹⁰

Support for the multi-stakeholder model

The multi-stakeholder model has been a critically important – but often overlooked – factor in the sustainability of digital and Internet technologies over the past several decades. It is increasingly in jeopardy of being undermined by authoritarian governments seeking to seize control from existing rule-setting bodies. As the United States’ attitude towards the global multi-stakeholder model of governance is becoming more dubious, at least for the time being, Europe and other like-minded democracies should take on a leadership role on two fronts: first, in supporting, sustaining, and ex-

10. R. Albergotti (2025), ‘Anthropic irks White House with limits on models’ use’, *Semafor*, 17 September, <https://www.semafor.com/article/09/17/2025/anthropic-irks-white-house-with-limits-on-models-uswhite-house-with-limits-on-models-use>.

panding the multi-stakeholder model from existing digital technologies and services to AI; and second, in confronting and countering the efforts of authoritarian countries to wrest the governance of technology away from entities such as the United Nations and the ITU. It can do both from *within* these organisations. To avoid irreversible backsliding and bolster the intergovernmental agencies now under threat, action is crucial. Europe must step up to lead.

REFERENCES

- Albergotti, R. (2025). 'Anthropic irks White House with limits on models' use', *Semafor*, 17 September, <https://www.semafor.com/article/09/17/2025/anthropic-irks-white-house-with-limits-on-models-uswhite-house-with-limits-on-models-use>.
- Bradford, A. (2012). 'The Brussels Effect', *Northwestern University Law Review*, 1 December. <https://northwesternlawreview.org/issues/the-brussels-effect/>.
- Bradford, A. (2020). *The Brussels Effect: How the European Union Rules the World*. Oxford: Oxford University Press.
- Bradford, A. (2023). *Digital Empires: The Global Battle to Regulate Technology*. Oxford: Oxford University Press.
- Central People's Government of the People's Republic of China (2025). '人工智能全球治理行动计划 [Action Plan on Global Governance of Artificial Intelligence]', https://www.gov.cn/yaowen/liebiao/202507/content_7033929.htm.
- Executive Office of the President (2025). 'Winning the Race: America's AI Action Plan', 23 July, <https://www.whitehouse.gov/wp-content/uploads/2025/07/Americas-AI-Action-Plan.pdf>.
- Executive Order 13859 (2019). 'Maintaining American Leadership in Artificial Intelligence', 11 February, <https://www.federalregister.gov/documents/2019/02/14/2019-02544/maintaining-american-leadership-in-artificial-intelligence>.
- Executive Order 14110 (2023). 'Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence', 30 October, <https://www.federalregister.gov/documents/2023/11/01/2023-24283/safe-secure-and-trustworthy-development-and-use-of-artificial-intelligence>.
- Felton, E. & Lyons T. (2016). 'The Administration's Report on the Future of Artificial Intelligence', 12 October, <https://obamawhitehouse.archives.gov/blog/2016/10/12/administrations-report-future-artificial-intelligence>.
- Mok, C. (2025a). 'Quick take on China's AI Action Plan – It's all about global governance' (Author's Substack), 27 July, <https://charlesmok.substack.com/p/quick-take-on-chinas-ai-action-plan>.
- Mok, C. (2025b). 'In some ways, Trump's new AI Action Plan actually looks familiar', *The Diplomat*, 28 July, <https://thediplomat.com/2025/07/trumps-new-ai-action-plan-looks-awfully-familiar/>.
- Mok, C. (2025c). 'The US aims to win the AI Race, but China wants to win friends first', *Tech Policy Press*, 8 August, <https://www.techpolicy.press/the-us-aims-to-win-the-ai-race-but-china-wants-to-win-friends-first/>.
- Spies, A. (2025). 'Europe realizes that it is over-regulating AI', *American-German Institute*, 12 March, <https://americangerman.institute/2025/03/europe-realizes-that-it-is-overregulating-ai/>.
- United Nations (2024a). 'What is the pact for the future?', <https://www.un.org/en/summit-of-the-future/pact-for-the-future>.
- United Nations (2024b). 'Global digital compact', <https://www.un.org/en/summit-of-the-future/global-digital-compact>.
- US Department of State (2022). 'Declaration for the Future of the Internet', 28 April, <https://www.state.gov/declaration-for-the-future-of-the-internet>.
- White House (2025). 'Winning the Race: America's AI Action Plan', 23 July, <https://www.whitehouse.gov/wp-content/uploads/2025/07/Americas-AI-Action-Plan.pdf>.

Lists of Good Wishes? The Limits of the ‘AI Brussels Effect’ in Latin America

Claudia Del Pozo

<https://doi.org/10.53121/ELFS10> • ISSN (print) 2791-3880 • ISSN (online) 2791-3899

ABSTRACT

This chapter examines the European Union Artificial Intelligence Act (EU AIA) from a Latin American perspective, highlighting both its influence and limitations beyond Europe. The Act has quickly become a global benchmark for responsible AI regulation, with its rights-based and precautionary approach resonating strongly in Latin America. Policy-makers in the region have drawn directly from the EU’s framework, adopting risk-based classifications and ethical language in draft bills and strategies. Yet this influence is often more symbolic than substantive: weak institutions, limited technical capacity, and unstable political environments undermine the enforceability of European-inspired regulation, reducing many initiatives to aspirational declarations. These challenges are compounded by the shifting global context that complicates Latin America’s regulatory choices. The chapter argues that the region’s opportunity lies not in imitation, but in adaptation, using the EU AI Act as a model while building governance that reflects Latin America’s unique institutional and social realities.

ABOUT THE AUTHOR

Claudia Del Pozo is the Founder and Director of Eon Institute, a women-led Mexican think-tank that seeks to future-proof society.

INTRODUCTION

Every new update from the European Commission on the European Union Artificial Intelligence Act (EU AIA) is met with a familiar chorus: accusations of ‘bureaucratic overreach’, complaints about ‘red tape’, and warnings of a looming ‘logistical nightmare’. This growing scepticism stands in stark contrast to the initial wave of international praise that greeted the EU AI Act as a global benchmark for responsible AI regulation.

When it was first introduced in 2021, the AIA was hailed as a pioneering effort to put human rights and democratic values at the core of AI development. Policymakers and civil society leaders around the world commended the Act’s risk-based approach, ethical framing, and ambition to rein in powerful technologies before they became ungovernable. The EU was positioned as the moral compass of AI governance, a ‘first mover’ filling the vacuum of international rules.

Examining the EU AIA from an outsider's perspective offers more than a descriptive account of its influence; it clarifies the actual reach of European regulatory power and the real boundaries of the Brussels Effect.

Yet as generative AI exploded onto the global stage after the release of ChatGPT in 2022, the so-called Brussels Effect has been put to the test and reception has grown more complex. Lawmakers in Brussels were forced to revise key provisions to keep pace while critics abroad, especially from industry, questioned whether the EU's cautious approach could remain viable in a landscape increasingly driven by competitiveness and innovation. At the same time, the United States shifted from Biden's early alignment with the EU toward a deregulatory model under the Trump administration, further eroding the sense of a shared transatlantic vision.

Examining the EU AIA from an outsider's perspective offers more than a descriptive account of its influence; it clarifies the actual reach of European regulatory power and the real boundaries of the Brussels Effect. The asymmetries, dependencies, and translation challenges that arise when a normative model built in Brussels meets regions with different institutional, economic, and political realities are voiced best by observers outside the EU.

Rather than celebrate or dismiss EU leadership, this chapter critically evaluates how far its influence truly extends and where it begins to lose traction. The chapter explores how the EU AI Act is perceived, interpreted, and adapted in Latin America, a region where actors seek to balance emulation of Europe's example with local realities shaped by institutional fragility, socio-economic inequality, and geopolitical dependencies. Against this backdrop, Latin America provides a revealing test case in which the EU's normative leadership is acknowledged, but the assumption that its model can serve as a universal blueprint is challenged.

FROM APPLAUSE TO AMBIVALENCE

With its announcement in 2021, the EU AIA captured international attention. It was immediately lauded as the most important piece of regulation of the AI ecosystem to date, a bold and necessary step in the absence of binding rules elsewhere.¹ For many, it signalled the beginning of a new era of AI governance: one that centred human rights, transparency, and accountability as foundational principles rather than afterthoughts.

The Act reflected Europe's regulatory tradition, marked by a strong precautionary approach and a willingness to anticipate and mitigate harm before it fully materialises. In this spirit, the AI Act sought to regulate not only the technical characteristics of AI systems, but also their structural, political, and social impacts.² It fit neatly, moreover, within the EU's broader product safety framework. As Nicoleta Cherciu, Managing Partner at Cherciu & Co., explained, 'the AI Act aims to be an addition to any part of the EU's current product safety regulation package. For example, toys, medical devices, electronic products, vehicles, computers, etc. are subject to safety requirements and conformity assessments before being placed onto the EU market'.³

The initial international response to the EU AIA was broadly positive. Countries looking to develop their own regulatory frameworks saw the EU's proposal as a model to follow, or at least a reference point that could help reduce legal uncertainty and harmonise international standards. The Act provided something that had been missing: a concrete starting point for national and multilateral debates.

1. R. Csernaton (2025), 'The EU's AI power play: Between deregulation and innovation', *Carnegie Europe*, 20 May, <https://carnegieendowment.org/research/2025/05/the-eus-ai-power-play-between-deregulation-and-innovation?lang=en>.

2. Amnesty International (2021), 'An EU Artificial Intelligence Act for fundamental rights', 30 November, <https://www.amnesty.eu/news/an-eu-artificial-intelligence-act-for-fundamental-rights/>.

3. E. Ghinita (2024), 'The EU AI Act explained by the experts: Will it hinder or enhance innovation?', *The Recursive*, 13 March, <https://therecursive.com/the-eu-ai-act-explained-by-the-experts-will-it-hinder-or-enhance-the-innovation/>.

However, even in this early stage, some cautioned against wholesale adoption. Critics from the worlds of industry and policymaking warned that the Act's complexity and scope could make it difficult to implement or adapt in regions with less regulatory capacity or different political priorities. While the Brussels Effect may have given Europe a first-mover advantage, hesitation remained as to whether such a prescriptive model could be successfully grafted onto very different legal, economic, and cultural environments.

The development of generative AI amplified these doubts. By the time the EU presented the revised AI Act in 2023, other governments were already recalibrating. Some continued to embrace the EU model while others began distancing themselves, questioning whether Europe's slow, bureaucratic approach was the right fit for a technology evolving at breakneck speed.

THE AI BRUSSELS EFFECT IN LATIN AMERICA

Latin America and the Caribbean (LAC) are progressing steadily in their institutional capacity to govern AI, though major structural and technological challenges persist. The region ranks seventh out of nine in Oxford Insights' 2024 Government AI Readiness Index (2024), with many of its constituent nations still in an early phase of AI governance. A small group of countries – namely, Brazil, Chile, and Uruguay – lead regional efforts with updated national AI strategies, strong public sector digital agendas, and an emphasis on ethics and rights-based governance. However, most LAC countries still face significant barriers, ranging from limited infrastructure and weak technology sectors to insufficient technical talent for scaling AI development and regulation.

Despite these constraints, Latin America has been an active participant in global conversations on responsible AI. A defining feature of the region's approach is its prioritisation of governance over innovation, focusing on strategy, ethics, and public sector capacity rather than on competitiveness or technological autonomy. Regional cooperation has been instrumental: the Montevideo Declaration (2024), a follow-up to the Santiago Declaration (2023), adopted at the Ministerial Summit on AI Ethics in Latin America and the Caribbean in October, reflects a collective effort to build a consensus on ethical and inclusive AI governance across Latin America and the Caribbean.⁴

Like the rest of the world, Latin America welcomed the European Union's proposal of the AI Act, applauding it as a groundbreaking and principled attempt to regulate a fast-evolving technology. The EU's emphasis on human rights, transparency, and precaution resonated strongly with Latin American actors from all sectors, but especially civil society and government, given the region's own struggles with inequality, surveillance, and digital asymmetries. As a result, many LAC countries' efforts to regulate AI through bills, national strategies, and public consultations have drawn heavily from the EU model.⁵ In some cases, policymakers have even used the EU AI Act as a direct reference when drafting their own regulatory proposals.⁶

Looking into the five Latin American governments most prepared to take advantage of AI – Brazil, Chile, Colombia, Uruguay, and Peru (see Figure 1) – some trends emerge. Notably, all of them have undertaken AI governance projects patterned after the EU's strategy, be it by prioritising ethics and responsibility or by directly echoing the EU AI Act's risk-based approach. Below is a summary of each country's approach, starting with the most advanced countries in the region.⁷

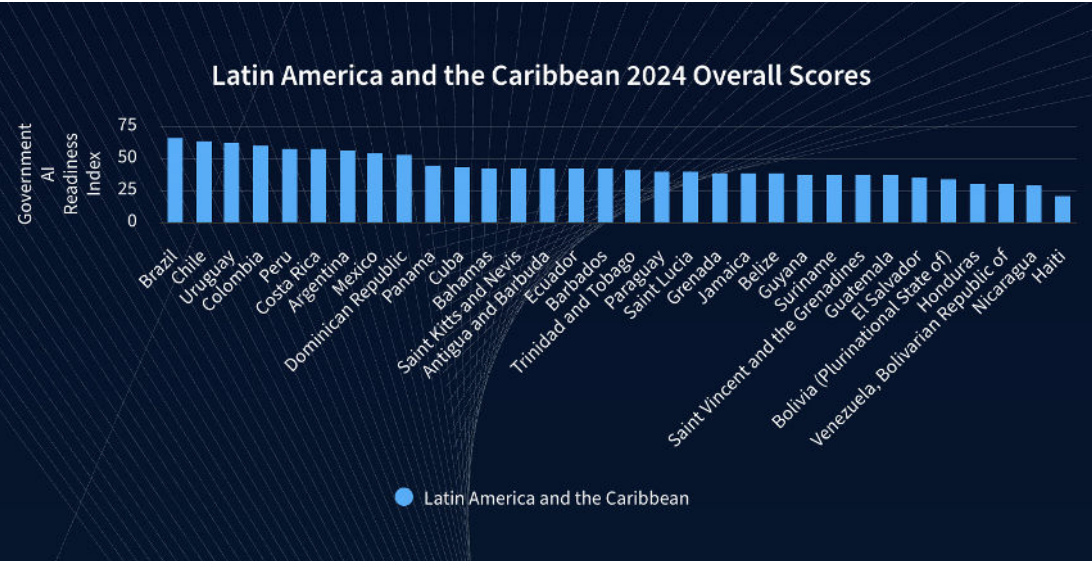
4. Oxford Insights (2024), *Government AI Readiness Index. 2024*, <https://oxfordinsights.com/ai-readiness/ai-readiness-index/?#download-reports>.

5. Access Now (2024), *Regulatory Mapping on Artificial Intelligence in Latin America*, Thomson Reuters Foundation, <https://www.trust.org/resource/regulatory-mapping-on-artificial-intelligence-in-latin-america/>.

6. V. Muñoz (2024), 'Inteligencia artificial: potenciando el futuro de América Latina y el Caribe', *Análisis Carolina*, 18(1), 1–14. https://doi.org/10.33960/AC_18.2024.

7. Oxford Insights, *Government AI Readiness Index*.

FIGURE 1. Latin American and Caribbean countries ranked by level of Government AI Readiness.



Source: Oxford Insights, Government AI Readiness Index 2024, <https://oxfordinsights.com/ai-readiness/ai-readiness-index/?#download-reports>.

Brazil

Together with Chile, Brazil leads Latin America in terms of AI policy, providing a clear example of how the European Union’s regulatory vision informs debates beyond Europe. Brazil is currently advancing Bill No. 2,338/2023, explicitly inspired by the EU AI Act.⁸ In its emulation of European approaches to digital governance, the proposed legislation is in line with other regulations recently put forth in Brazil, such as the GDPR-like data protection law of 2018.

Brazil’s AI Bill, approved by the Senate in December 2024 but awaiting a vote in the House of Representatives and presidential approval, is one of a handful of comprehensive frameworks worldwide aiming to regulate the development, deployment, and use of AI systems. Like the EU AIA, the Bill endeavours to safeguard fundamental rights, ensure secure and reliable AI, and promote human dignity, democratic values, and scientific–technological progress. Central to its design is a tiered risk framework inspired by the EU AIA’s: systems that pose ‘excessive risk’ are prohibited; ‘high-risk’ systems are subject to strict obligations; and all other systems must comply with baseline requirements.⁹

At the same time, Brazil’s AI Bill illustrates the challenges of translating European models into Latin American political and institutional contexts. Despite the Senate’s approval, the bill still requires scrutiny and a vote in the House of Representatives, as well as presidential sanction, leaving it open to amendment and political bargaining. Importantly, no timeline has been set for these next steps, which creates uncertainty regarding when, and in what form, the law will ultimately take effect.¹⁰

8. E. Levy Yeyati (2025), ‘Regulating AI on Latin America’s terms’, *Americas Quarterly*, 30 June, <https://americasquarterly.org/article/regulating-ai-on-latin-americas-terms/>.
9. D. Atanasovska and L. Robeli (2025), ‘Brazil’s AI Act: A new era of AI regulation’, *GDPRLocal*, 26 February, <https://gdprlocal.com/brazils-ai-act-a-new-era-of-ai-regulation/>.
10. White & Case (2025a), ‘AI Watch: Global Regulatory Tracker – Brazil’, 6 June, <https://www.whitecase.com/insight-our-thinking/ai-watch-global-regulatory-tracker-brazil>.

Chile

Chile's legislative answer to the rise of AI was the Artificial Intelligence Regulation Bill (No. 16821-19). This was introduced to the Chamber of Deputies on 7 May 2024 and builds on Chile's National AI Policy, originally published in 2021 and updated in 2024 following UNESCO's recommendations. The bill's framework aims to both regulate and promote the ethical and responsible development of AI, following recommendations put forward by UNESCO.¹¹

The bill focuses on transparency, fairness, and human oversight and seeks to foster the creation and deployment of human-centred AI systems while safeguarding public health and fundamental rights and protecting consumers from harmful applications. It adopts a hybrid model that combines self-regulation with a risk-based framework, classifying AI systems as unacceptable, high, limited, or minimal risk. It takes its guiding principles from internationally recognised standards, particularly those outlined in UNESCO's Recommendation on the Ethics of AI.¹²

Colombia

While Colombia does not have any regulation specifically tied to AI to date, over 20 bills had been proposed to Congress by May of 2025.¹³ One of the most recent of these, put forth by the Ministries of Science and of Information and Communication Technologies (ICT), included a risk classification based on the EU AIA, with prohibited AI systems, high-risk systems that are subject to stringent requirements, and low-risk systems with minimal obligations.¹⁴

Uruguay

Uruguay's Agency for Electronic Government and Digital Society (AGESIC) has recently begun exploring potential AI regulation, signalling a concern with protecting rights, defining permissible uses of AI, and promoting job creation and economic growth in the sector.¹⁵ While details remain scarce, Uruguay's 2023 endorsement of an EU Council declaration on advancing joint AI policy among the EU and LAC suggests that any future framework will likely draw inspiration from the EU AI Act, mirroring the trajectory of its regional counterparts.¹⁶

Peru

Peru is one of the most legislatively active countries in Latin America, especially when it comes to AI. Several of its general regulatory frameworks for AI borrow heavily from the EU AI Act, invoking concepts like 'ethical AI', 'non-discrimination', and 'transparency'. That being said, these principles rarely translate into enforceable obligations and are often inconsistently assigned to users of AI systems instead of developers or deploying institutions, blurring accountability. This has led certain experts to rate Peru's AI legislative activity as impressive in sheer quantity, but lacking in depth.¹⁷

11. UNESCO (2024), 'Chile launches a national AI policy and introduces an AI bill following UNESCO's recommendations', 4 May, <https://www.unesco.org/en/articles/chile-launches-national-ai-policy-and-introduces-ai-bill-following-unescos-recommendations-0>.

12. UNESCO (2021), *Recommendation on the Ethics of Artificial Intelligence*, 23 November, <https://www.unesco.org/en/articles/recommendation-ethics-artificial-intelligence>.

13. F. Fortich (2025), 'Puntos a favor y vacíos del proyecto de ley del Gobierno para regular la IA', *El Espectador*, 26 May, <https://www.elespectador.com/ciencia/nueva-propuesta-de-ley-del-gobierno-de-gustavo-petro-busca-regular-inteligencia-artificial-en-colombia/>.

14. White & Case (2025b), 'AI Watch: Global Regulatory Tracker – Colombia', 26 June, <https://www.whitecase.com/insight-our-thinking/ai-watch-global-regulatory-tracker-colombia>.

15. Montevideo Portal (2025), 'Agesic plantea ley sobre inteligencia artificial: ¿qué contemplaría la regulación?', 18 June, <https://www.montevideo.com.uy/Ciencia-y-Tecnologia/Agesic-plantea-ley-sobre-inteligencia-artificial-que-contemplaria-la-regulacion--uc927330>.

16. Agenciade Gobierno Electrónico y Sociedad de la Información y del Conocimiento (AGESIC) (2023), 'Uruguay adhiere a declaración regulatoria a nivel global sobre Inteligencia Artificial', 20 November, <https://www.gub.uy/agencia-gobierno-electronico-sociedad-informacion-conocimiento/comunicacion/noticias/uruguay-adhiere-declaracion-regulatoria-nivel-global-sobre-inteligencia>.

17. S. Smart and V. M. Montori (2025), 'Peru's AI regulatory boom: Quantity without depth?', *Carr-Ryan Center for Human Rights*, 23 April, <https://www.hks.harvard.edu/centers/carr-ryan/our-work/carr-ryan-commentary/perus-ai-regulatory-boom-quantity-without-depth>.

IMPORTED NORMS, LOCAL REALITIES

While the EU AIA seems to be a touchstone for Latin American countries' AI regulation, differences between the EU and LAC can cause disparities in the impact of legislation.

Latin American countries face challenges like digital illiteracy, strained institutional capacity, regulatory capture and corruption, discriminatory policing, and fragmented public service delivery.¹⁸ Legislation suited to an EU context may be poorly equipped to operate effectively under such circumstances. Moreover, very few Latin American legislative proposals successfully take the region's linguistic diversity and social realities into account, even though these shape how AI operates in practice. Some key obstacles to the implementation of EU-inspired regulatory models in Latin America are outlined below.

Technological literacy gaps. The most notable challenge in LAC countries is the limited technical understanding among many drafters. Certain regulations, for instance, define AI as 'algorithms that people program', or state that 'AI must not lie to a human being' but do not offer a definition of deception in algorithmic systems or how this standard could be implemented.¹⁹

Institutional gaps. Latin America's limited institutional capacity constitutes a major problem – many governments lack the agencies needed to oversee and enforce AI regulations, as well as robust audit mechanisms and redress pathways. Without this scaffolding, European-inspired frameworks cannot functionally mirror the EU model.²⁰ For example, a regulation might mandate audits to detect bias in AI systems without specifying responsible entities or audit procedures. Similar concerns have been raised in Asia and Africa, where experts argue that the EU AI Act cannot simply be transplanted onto diverse legal infrastructures, instead advocating local experimentation through pilot projects and regulatory sandboxes, a practice still rare in Latin America.²¹

Unstable political direction. Shifts in administration often lead to abrupt changes in AI policy, at best altering approaches but more commonly deprioritising them altogether. As a result, efforts to advance regulation may stall after elections, as illustrated by Mexico's abandoned 2018 AI strategy. This volatility stands in stark contrast to the EU's more stable institutional environment and undermines continuity in regulation. Consequently, the effectiveness of AI frameworks in the region often depends less on the quality of the regulation itself than on the political priorities of the government in power, leaving implementation and long-term governance highly vulnerable to electoral cycles.

Geopolitical triangulation. The EU is not the only external influence on Latin America. US companies dominate platforms and cloud services across the region, while China invests heavily in digital infrastructure and surveillance technologies. Latin America's AI governance will be shaped as much by these dependencies as by Brussels' norms.

These challenges have led certain commentators to claim that many Latin American countries seem to prioritise alignment with international norms over substantive local protections, creating 'list[s] of good wishes' without much impact.²²

While the EU AIA seems to be a touchstone for Latin American countries' AI regulation, differences between the EU and LAC can cause disparities in the impact of legislation.

18. Smart and Montori, 'Peru's AI regulatory boom'.

19. Smart and Montori, 'Peru's AI regulatory boom'.

20. Smart and Montori, 'Peru's AI regulatory boom'.

21. Academy of International Affairs NRW (2023), *The EU AI Act and Voices from the Global South Academy of International Affairs NRW (Bonn). AI Policy Workshop: 02 March 2023* (Report), https://www.aia-nrw.org/app/uploads/2023/05/23-05-22_WS_Report-1.pdf.

22. Smart and Montori, 'Peru's AI regulatory boom'; Business & Human Rights Resource Centre (2023), 'Eight coun-

GREAT POWER TENSIONS AND LATIN AMERICA'S AI DILEMMA

If weak institutional capacity complicates the EU's influence in Latin America, the United States' shifting position on AI governance poses an even greater challenge, namely the risk of regulatory fragmentation and uncertainty for Latin American countries seeking external models. Indeed, the course struck by the US may prove to be the decisive external factor in the LAC region's AI regulatory trajectory.

Under the Biden administration, the US seemed briefly to converge with Europe's rights-based vision. Executive Order 14110 on 'Safe, Secure, and Trustworthy Artificial Intelligence' shared the EU's concern with accountability and fairness, and forums like the US–EU Trade and Technology Council projected a spirit of transatlantic partnership. For a time, Europe and the United States appeared poised to jointly set the rules for global AI, offering Latin America a relatively coherent model rooted in shared democratic values.

That moment was fleeting. With Donald Trump's return to power in 2025, the US abandoned its cautious, rights-oriented stance and pivoted sharply toward deregulation. Trump's 'Winning the Race: America's AI Action Plan' reframed governance as an obstacle to innovation, dismantled existing safeguards, and threatened to punish states pursuing rules that could 'slow industry progress'. AI policy was no longer about balancing innovation with rights, but about ensuring US technological dominance at all costs, an ambition further advanced through the American AI Stack. For Brussels, this reversal risks triggering a global regulatory 'race to the bottom'; for Washington, EU rules are dismissed as protectionist and bureaucratic barriers designed to stifle competitiveness.

The impact of this divergence was quickly felt in Latin America. In Brazil, whose original AI Bill (Projeto de Lei No. 2338/2023) closely mirrored the EU's rights-based model, the shift became evident following a March 2024 visit to Washington by Senator Marcos Pontes. He was accompanied by fellow senator Laércio Oliveira and met with US government officials and executives from major tech firms including Amazon, Google, and Microsoft. The delegation sought feedback on Brazil's draft AI legislation, which had proposed strong oversight mechanisms, copyright protections, and safeguards against discriminatory AI systems. Pontes then publicly criticised the bill as 'too strict' and 'based on fear', and subsequently introduced 32 amendments, with Oliveria, that weakened key provisions related to accountability, copyright, and the scope of regulated systems. The result was a significantly diluted version of the legislation, passed by the Senate in December 2024, that reflected growing pressure to align with the US's deregulatory stance under Trump's administration.²³

This episode encapsulates Latin America's central dilemma. While the EU's rights-based approach appeals ideologically to governments wishing to counterbalance US influence (especially amid a wave of left-leaning administrations), the gravitational pull of the US tech sector often dictates what is politically and commercially viable. The result will likely be fragmentation: some governments drifting towards EU-inspired regulation, others embracing US-style deregulation in pursuit of investment, and many defaulting to hybrid or symbolic frameworks that borrow the language of rights without the institutional backing to enforce them.

The US pivot exposes Latin America to a more fundamental vulnerability. Caught between competing global powers, the region risks becoming a passive consumer of external models rather than an active shaper of its own governance. And the equation is no longer only Brussels versus Washington. China has become an increasingly influential actor in the region through investments in digital infrastructure, surveillance systems, and smart city projects.

Peru provides a telling example. In November 2024, Peru and China signed an expanded free-trade agreement underscoring AI cooperation and digital inclusivity, in the wake of President Xi Jinping's

tries in Latin America propose laws to regulate AI, aiming to protect human rights & promote tech innovation in the region', 10 July, <https://www.business-humanrights.org/es/%C3%BAltimas-noticias/eight-countries-in-latin-america-propose-laws-to-regulate-ai-aiming-to-protect-human-rights-promote-tech-innovation-in-the-region/>.

23. Rest of the World staff (2025), 'Brazil's push for comprehensive AI regulation', *Rest of the World*, 31 January, <https://restofworld.org/2025/brazil-ai-regulation-big-tech/>.

visit to Lima for the APEC Forum. Shortly thereafter, the Peruvian government's public messaging began to align with Beijing's diplomatic language on Taiwan, presaging a technological, economic, and ideological convergence.²⁴ In contrast to the EU's normative strategy or Washington's market dominance, Beijing exerts influence through hardware, capital, and turnkey AI solutions, further complicating Latin America's regulatory calculus.

Without a deliberate strategy for adapting external frameworks to its own realities, Latin America's AI governance may be defined not by democratic values or local needs, but by the gravitational pull of Washington, Brussels, and Beijing.

CONCLUSION

The EU AI Act has quickly become a global reference point for responsible AI governance, shaping debates far beyond Europe. In Latin America, its rights-based and precautionary framing resonates with longstanding concerns about inequality, surveillance, and democratic fragility. Policymakers have drawn directly on the Act's language, from Brazil's Senate-approved bill to Chile's ethical frameworks, indicating the appeal of Europe's regulatory approach. Yet due to the factors outlined in this chapter, the influence of the EU AIA is often more symbolic than substantive. The competing forces of Washington and Beijing further attenuate the Brussels Effect – Europe may set aspirational norms, but US and Chinese actors dominate the digital ecosystems in which those norms would need to operate.

Latin America is at a crossroads. Passive adoption of external templates risks turning the region into a regulatory colony of Brussels, Washington, or Beijing, where governance is dictated from abroad. Yet rejecting these models outright would leave countries vulnerable as testing grounds for unregulated AI. The real opportunity lies in adaptation: leveraging the EU AI Act's strengths while designing governance that reflects Latin America's unique institutional weaknesses, social inequalities, and democratic priorities. Whether the region seizes this opportunity will determine its future: either as a credible leader of global AI governance, or a peripheral consumer of rules made elsewhere.

As it witnesses its regulatory influence tested by diverse realities, a broader question emerges for the EU: will the AI Act remain a primarily defensive instrument (focused on preventing harm) or can it evolve into a genuine driver of innovation and inclusion, both within Europe and beyond?

REFERENCES

- Academy of International Affairs NRW (2023). *The EU AI Act and Voices from the Global South Academy of International Affairs NRW (Bonn). AI Policy Workshop: 02 March 2023*. Report.
https://www.aia-nrw.org/app/uploads/2023/05/23-05-22_WS_Report-1.pdf.
- Access Now (2024). *Regulatory Mapping on Artificial Intelligence in Latin America*, Thomson Reuters Foundation,
<https://www.trust.org/resource/regulatory-mapping-on-artificial-intelligence-in-latin-america/>.
- Agenciade Gobierno Electrónico y Sociedad de la Información y del Conocimiento (2023). 'Uruguay adhiere a declaración regulatoria a nivel global sobre Inteligencia Artificial', 20 November,
<https://www.gub.uy/agencia-gobierno-electronico-sociedad-informacion-conocimiento/comunicacion/noticias/uruguay-adhiere-declaracion-regulatoria-nivel-global-sobre-inteligencia>.

24. H. Harsono (2025), 'From belt and road to bits and bytes: China's use of emerging technologies to exercise influence in Peru', *Small Wars Journal*, 2 June, <https://smallwarsjournal.com/2025/06/02/chinas-gray-zone-strategy-in-peru-open-source-tech-and-digital-influence/#>.

- Amnesty International (2021). 'An EU Artificial Intelligence Act for fundamental rights', 30 November, <https://www.amnesty.eu/news/an-eu-artificial-intelligence-act-for-fundamental-rights/>.
- Atanasovska, D. & Robeli, L. (2025). 'Brazil's AI Act: A new era of AI regulation', *GDPRLocal*, 26 February, <https://gdprlocal.com/brazils-ai-act-a-new-era-of-ai-regulation/>.
- Business & Human Rights Resource Centre (2023). 'Eight countries in Latin America propose laws to regulate AI, aiming to protect human rights & promote tech innovation in the region', 10 July, <https://www.business-humanrights.org/es/%C3%BAltimas-noticias/eight-countries-in-latin-america-propose-laws-to-regulate-ai-aiming-to-protect-human-rights-promote-tech-innovation-in-the-region/>.
- Csernaton, R. (2025). 'The EU's AI power play: Between deregulation and innovation', *Carnegie Europe*, 20 May, <https://carnegieendowment.org/research/2025/05/the-eus-ai-power-play-between-deregulation-and-innovation?lang=en>.
- Fortich, F. (2025). 'Puntos a favor y vacíos del proyecto de ley del Gobierno para regular la IA', *El Espectador*, 26 May, <https://www.elespectador.com/ciencia/nueva-propuesta-de-ley-del-gobierno-de-gustavo-petro-busca-regular-inteligencia-artificial-en-colombia/>.
- Ghinita, E. (2024). 'The EU AI Act explained by the experts: Will it hinder or enhance innovation?', *The Recursive*, 13 March, <https://therecursive.com/the-eu-ai-act-explained-by-the-experts-will-it-hinder-or-enhance-the-innovation/>.
- Harsono, H. (2025). 'From belt and road to bits and bytes: China's use of emerging technologies to exercise influence in Peru', *Small Wars Journal*, 2 June, <https://smallwarsjournal.com/2025/06/02/chinas-gray-zone-strategy-in-peru-open-source-tech-and-digital-influence/#>.
- Levy Yeyati, E. (2025). 'Regulating AI on Latin America's terms', *Americas Quarterly*, 30 June, <https://americasquarterly.org/article/regulating-ai-on-latin-americas-terms/>.
- Montevideo Portal (2025). 'Agesic plantea ley sobre inteligencia artificial: ¿qué contemplaría la regulación?', 18 June, <https://www.montevideo.com.uy/Ciencia-y-Tecnologia/Agesic-plantea-ley-sobre-inteligencia-artificial--que-contemplaria-la-regulacion--uc927330>.
- Muñoz, V. (2024). 'Inteligencia artificial: potenciando el futuro de América Latina y el Caribe'. *Análisis Carolina*, 18(1), 1–14. https://doi.org/10.33960/AC_18.2024.
- Oxford Insights (2024). *Government AI Readiness Index. 2024*, <https://oxfordinsights.com/ai-readiness/ai-readiness-index/?#download-reports>.
- Rest of the World staff (2025). 'Brazil's push for comprehensive AI regulation', *Rest of the World*, 31 January, <https://restofworld.org/2025/brazil-ai-regulation-big-tech/>.
- Smart, S. & Montori, V. M. (2025). 'Peru's AI regulatory boom: Quantity without depth?', *Carr-Ryan Center for Human Rights*, 23 April, <https://www.hks.harvard.edu/centers/carr-ryan/our-work/carr-ryan-commentary/perus-ai-regulatory-boom-quantity-without-depth>.
- UNESCO (2024). 'Chile launches a national AI policy and introduces an AI bill following UNESCO's recommendations', 4 May, <https://www.unesco.org/en/articles/chile-launches-national-ai-policy-and-introduces-ai-bill-following-unescos-recommendations-0>.
- White & Case (2025a). 'AI Watch: Global Regulatory Tracker – Brazil', 6 June, <https://www.whitecase.com/insight-our-thinking/ai-watch-global-regulatory-tracker-brazil>.
- White & Case (2025b). 'AI Watch: Global Regulatory Tracker – Colombia', 26 June, <https://www.whitecase.com/insight-our-thinking/ai-watch-global-regulatory-tracker-colombia>.

The Fall of Clarity: AI Regulation and the Risk of European Decline

Gian Marco Bovenzi

<https://doi.org/10.53121/ELFS10> • ISSN (print) 2791-3880 • ISSN (online) 2791-3899

ABSTRACT

This chapter assesses the European Union's AI Act in terms of its legal clarity when applied to companies, businesses, and enterprises. It asserts that the AI Act, with its complex phrasing, endless terminology, and the numerous obligations it places upon recipients, represents an often-unnecessary impediment to the development of technological investments and overall progress. If left unaddressed, these consequences may put the European Union's ability to keep up with its global economic competitors in jeopardy.

ABOUT THE AUTHOR

Gian Marco Bovenzi is a Data Risk Review Expert at Deloitte. He is a PhD Candidate in Strategic and Legal Studies of Innovation for Defence and Security at the Centre for Higher Defence Studies (Italian Ministry of Defence).

INTRODUCTION

In claris non fit interpretatio. This maxim, used in the legal system and jurisprudence of the ancient Romans, can be literally translated as 'in clear things, interpretation is not made'. The principal points to a simple yet fundamental concept: when a text is unambiguous and its meaning evident, interpretive efforts to extract a possible *meaning beyond the meaning* are unnecessary. Its corollary likewise represents a useful syllogism: the more straightforward the formulation of a norm, law, or regulation, the easier it is to understand and apply it – and accordingly, the better the result in terms of legal efficiency.

This principle remains relevant today. While *extensive* or *analogical* interpretations of a norm are often carried out to fill legal gaps and remove potential loopholes (either from counselling or judicial standpoints), strict accordance to the letter of the law remains, where applicable, the ideal. Clearly drafted legislation is a crucial tool for protecting citizens' fundamental rights.

Let us jump back to the future. Surely neither the Romans nor the average twentieth-century individual – besides, perhaps, such authors as Philip Dick, Isaac Asimov, and their peers – could foresee the need for a legal framework regulating artificial intelligence systems. (An AI system can be defined as a 'machine-based system ... designed to operate with varying levels of autonomy and that

may exhibit adaptiveness after deployment, and that ... infers, from the input it receives, how to generate outputs').¹ Conducting a legal overview of such a regulation would have been an exercise in science fiction only a handful of years ago.

PROVIDERS OR DEPLOYERS?

Consider this designation:

Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence and amending Regulations (EC) No 300/2008, (EU) No 167/2013, (EU) No 168/2013, (EU) 2018/858, (EU) 2018/1139 and (EU) 2019/2144 and Directives 2014/90/EU, (EU) 2016/797 and (EU) 2020/1828.

The extended name of the piece of regulation commonly known as AI Act already raises some alarm bells when it comes to clarity. Evaluating its contents, a few elements immediately stand out: the AI Act includes no less than 68 definitions (Article 4); the word provider(s) is mentioned 517 times; and the word obligation(s) (to which providers are subject) appears 228 times. Indeed, the providers of AI systems are the 'most heavily regulated subjects under the AI Act'.² They are defined as natural or legal persons, public authorities, agencies, or other bodies that:

- a) develop an AI system or a general-purpose AI model; or
- b) that have an AI system or a general-purpose AI model developed.

Such systems are either:

- a) placed on the market (within the EEA area); or
- b) put into service under their own name or trademark, whether for payment or free of charge.

Companies, businesses, enterprises, or industries may all qualify as *legal persons*. Not only are such entities subject to obligations as providers, but also as:

- deployers, when using an AI system under their authority (except where the AI system is used during a personal non-professional activity) – Article 4(4);
- authorised representatives, when located or established in the Union and performing or carrying out obligations and procedures on the behalf of a provider – Article 4(5);
- importers, when located or established in the Union and placing on the market an AI system that bears the name or trademark of a natural or legal person established in a third country – Article 4(6);
- distributors, when part of the supply chain, other than the provider or the importer, making an AI system available on the Union market – Article 4(7); and
- operators (a redundant catch-all definition including providers, product manufacturers, deployers, authorised representatives, importers, or distributors) – Article 4(8).

1. European Parliament (2024), *Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence and amending Regulations (EC) No 300/2008, (EU) No 167/2013, (EU) No 168/2013, (EU) 2018/858, (EU) 2018/1139 and (EU) 2019/2144 and Directives 2014/90/EU, (EU) 2016/797 and (EU) 2020/1828*, Article 3(1). Hereafter, for brevity, the legislation will be referred to as the AI Act.

2. M. Roleke (2025), 'The European AI Act', *activeMind.legal*, 3 February, <https://www.activemind.legal/guides/ai-act/>.

The first challenge for a company is to identify the group to which it pertains. While this might seem a simple question to answer, the line between an entity *developing* an AI system and *using it under their authority* is thin and blurred; the distinction depends on several factors. Let us imagine, for instance, an AI tool that is integrated into a given platform and is publicly available upon subscription. Now imagine a company which, after obtaining a licence from a third-party platform, utilises the AI tool in its own environment and fine-tunes it. Fine-tuning may be defined for our purposes as ‘the process of adapting a pre-trained model for specific tasks or use cases ... a subset of the broader technique of *transfer learning*’ which can ‘reduce the amount of expensive computing power and labelled data needed to obtain large models tailored to niche use cases and business needs’ and which ‘plays an important role in the real-world application of machine learning models, helping democratize access to and customization of sophisticated models’.³

The word ‘customisation’ is key to understanding the issue at hand. The question is, does customisation imply the development of a new AI system, or is the company just using the AI system as originally developed (perhaps slightly modified by a third party) in a form suited to their requirements? At first glance, the distinction may seem simple – but how can we precisely define when fine-tuning amounts to the creation of a new AI system? The answer to this question has real-world consequences, affecting the obligations to which providers are subject far more than the obligations placed upon deployers. It may even happen that a company is both a provider and deployer – to which obligations, then, is it subject?

Article 25 of the AI Act offers some elucidation, applicable in cases where a company a) puts its name or trademark on a high-risk AI system already placed on the market or put into service; b) makes a substantial modification to a high-risk AI system that has already been placed on the market or has already been put into service; or c) modifies the intended purpose of an AI system, including a general-purpose AI system, which has not been classified as high-risk and has already been placed on the market or put into service in such a way that the AI system in question becomes a high-risk AI system. In these instances, a deployer, importer, distributor, or other third party is automatically considered to be a provider and must therefore comply with the according obligations.

While this interpretation makes sense in theory, two complications are worth noting. Firstly, the cases enumerated in Article 25 of the AI Act relate exclusively to high-risk systems – general-purpose AI systems appear to fall beyond the scope of the article. Secondly, the recent EU AI Act Compliance Checker (an online tool made available by the European Commission to help clarify the Act’s obligations and requirements) includes the automatic re-qualification of an actor as provider under Article 25 even in cases not involving a high-risk system.⁴ Where, therefore, does the correct interpretation lie?

The example just described is a conundrum companies must deal with daily while carrying out their compliancy checks and assessments. Doing so comes at a cost of time, expense, and investment;

ultimately, it impacts the company’s overall performance in the market. Without an established legal interpretation or linear jurisprudence providing specific guidelines, EU companies bound by the AI Act are at risk of failing to keep pace with their global economic competitors.⁵

EU companies bound by the AI Act are at risk of failing to keep pace with their global economic competitors.

3. D. Bergmann (n.d.), *What is fine-tuning?*, IBM, <https://www.ibm.com/think/topics/fine-tuning>.

4. European Commission (2025), *EU AI Act Compliance Checker*, <https://ai-act-service-desk.ec.europa.eu/en/eu-ai-act-compliance-checker>.

5. See J. Koetsier (2025), ‘Top 10 AI Nations: Global AI Superpowers Ranked In Industry Report’, *Forbes*, 11 September, <https://www.forbes.com/sites/johnkoetsier/2025/09/11/top-10-ai-nations-global-ai-superpowers-ranked/>.

THE CATEGORISATION OF AI SYSTEMS AND RISKS

It is well known that under the AI Act, systems are divided into three distinct categories: prohibited AI systems (Article 5), high-risk AI systems (Articles 6 and following, under Chapter III), and general-purpose AI models (involving fewer associated obligations for providers and no obligations at all for deployers). In the emerging literature, the trend is to divide AI systems/model-related risks into four categories: a) unacceptable risk (associated with prohibited AI systems); b) high-risk systems (as regulated in the AI Act); c) limited risk (systems with simple transparency requirements); and d) minimal risk (unregulated systems/models).⁶

While this categorisation might facilitate risk evaluation by helping to specify companies' obligations under the AI Act, it is absent from the legislation itself. To repeat: the AI Act does not subdivide AI systems into risk-based categories. From an exclusively legal and theoretical standpoint, this does not cause inconvenience; it may, however, raise practical concerns for companies.

When a company assesses an AI system in terms of the AI Act, the first step is understanding its role, as seen above – either provider, deployer, or other (per Article 3). Subsequently, the company must determine how the AI system or model will be used within the business's operations, as well as its components. Understanding whether an AI system is prohibited (cases listed in Article 5) or high-risk may be relatively easy – these criteria are described in Article 6 and Annex III.⁷

However, issues arise in the case of obligations associated with general-purpose systems or models. It is sometimes ambiguous whether these can be simply considered as such or whether the possibility of their being classed as high-risk still stands. For instance, a general-purpose AI *model* is a model trained with a large amount of data using self-supervision at scale and capable of competently performing a wide range of distinct tasks (Article 4(63)). However, such models might embed 'high-impact capabilities' (which match or exceed the capabilities recorded in the most advanced general-purpose AI models) (Article 4(64), with further reference to Article 51(2)) or a 'systemic risk' having a significant impact on the Union market (Article 4(65), further specified under Article 51(1)(a)(b) as well as meeting the criteria listed in Annex XIII); moreover, a general-purpose AI *system* is an AI system that is based on a general-purpose AI model and which has the capability to serve a variety of purposes, both for direct use as well as for integration into other AI systems (Article 4(66)).

As the definitions above demonstrate, the AI Act's description of general-purpose AI and its related features (such as high-impact capabilities or systemic risks) might cause some confusion around the assessment of a system. In pragmatic terms, for a company to make sense of an AI system's legal identity, continuous dialogue between AI developers/computer engineers and lawyers/ethical advisors is unavoidable. Consider, for instance, the requirement stated under Article 51(2), according to which a general-purpose AI model is *presumed* to have high-impact capabilities when the 'cumulative amount of computation used for its training measured in floating point operations is greater than 10²⁵'. On the one hand, the AI developer/trainer must know the cumulative amount of computation and communicate it to the lawyer. On the other hand, the lawyer must still assess the *presumption* that such system has high impact capabilities (as a presumption might be overcome).

On top of this *internal* dialogue, cases in which a business or a company develops or has developed an AI system without interacting with other stakeholders are extremely rare. Hence, introducing an AI system into the market not only implies an assessment from the perspective of a provider,

6. See, for example, B. McElligott (2025), 'EU AI Act. Risk categories', *Explore Legislation Hub*, <https://www.mhc.ie/hubs/the-eu-artificial-intelligence-act/eu-ai-act-risk-categories/>; and Secure Privacy (2025), *EU AI Act: Understanding Risk-Based Classification*, <https://ai-eu-act.eu/blog/eu-ai-act-understanding-risk-based-classification/>.

7. An AI system is considered as a high risk system when '(a) the AI system is intended to be used as a safety component of a product, or the AI system is itself a product, covered by the Union harmonisation legislation listed in Annex I; (b) the product whose safety component pursuant to point (a) is the AI system, or the AI system itself as a product, is required to undergo a third-party conformity assessment, with a view to the placing on the market or the putting into service of that product pursuant to the Union harmonisation legislation listed in Annex I' (Article 6) and in the cases listed in Annex III, which include AI systems used in areas such as biometrics, education, critical infrastructure, employment, or law enforcement.

but also from the viewpoint of external deployers, importers, distributors, and other actors. These interactions undoubtedly tend to delay the entry of an AI system in the market. Numerous back-and-forths between stakeholders may be required to clarify how the AI Act applies or might apply to the system at a later stage.

When conducting an assessment under the AI Act, many factors beyond the AI system itself must also be accounted for: whether the system is hosted on a third-party environment or software, whether it is a software-as-a-service, who is entitled to use it, and whether it is internal-, client-, or external-facing. The numerous evaluations to which a system must be subjected, compounded by the ambiguities in the definitions provided by the AI Act, do more than merely burden companies – they can slow down progress and hamper Europe’s ability to compete globally. Insofar as the AI Act seeks to regulate technology according to liberal values, it is failed by the lack of clarity in its formulation.

CONCLUSIONS

The European Union has long sought to be a safe harbour for the protection of human rights. It cannot be denied that the citizen lies at the very *centre of the village*, the rest orbits around its fundamental rights. With the AI Act, the EU demonstrated its usual commitment to its pivotal line of thought: whenever society evolves, and technology along with it, new threats to fundamental rights may arise, and such threats must be mitigated. However, if we compare the AI Act with the GDPR on data protection and privacy rights, enormous differences emerge. The AI Act’s often unclear legal language stands in contrast with the more straightforward wording of the GDPR. Additionally, the

GDPR placed far fewer burdens and obligations on companies. For reasons like these, the GDPR did not (and does not) constitute a barrier to progress and has become a global benchmark for privacy rights.

The same cannot be said for the AI Act. Its complexity, opacity, and overlap with other regulations put it at odds with the rapidity of enterprise and businesses’ needs. While prohibited and high-risk AI systems perhaps deserve their place in the Act, the treatment of general-purpose AI systems and models is problematic. Their legal formulation, full of exceptions and derogations, results in

The numerous evaluations to which a system must be subjected, compounded by the ambiguities in the definitions provided by the AI Act, do more than merely burden companies – they can slow down progress and hamper Europe’s ability to compete globally.

companies spending an enormous amount of time (and money) to carry out assessments that are often superfluous. This is detrimental to the EU’s role in the global technological landscape, where giants as the United States and China (despite, perhaps, providing lesser levels of legal protection for their citizens) appear way ahead.

In medio stat virtus, to quote the ancient Romans again. Virtue stands in the middle-ground: keeping the citizen at the very heart of the project but freeing companies from ties that impede progress.

REFERENCES

- Bergmann, D. (n.d.). *What is fine-tuning?*, IBM,
<https://www.ibm.com/think/topics/fine-tuning>.
- European Commission (2025). *EU AI Act Compliance Checker*,
<https://ai-act-service-desk.ec.europa.eu/en/eu-ai-act-compliance-checker>.
- European Parliament (2024). *AI Act (Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence and amending Regulations (EC) No 300/2008, (EU) No 167/2013, (EU) No 168/2013, (EU) 2018/858, (EU) 2018/1139 and (EU) 2019/2144 and Directives 2014/90/EU, (EU) 2016/797 and (EU) 2020/1828*,
Regulation - EU - 2024/1689 - EN - EUR-Lex
- Koetsier, J. (2025). 'Top 10 AI Nations: Global AI Superpowers Ranked In Industry Report', *Forbes*, 11 September,
<https://www.forbes.com/sites/johnkoetsier/2025/09/11/top-10-ai-nations-global-ai-superpowers-ranked/>.
- McElligott, B. (2025). 'EU AI Act. Risk categories', *Explore Legislation Hub*,
<https://www.mhc.ie/hubs/the-eu-artificial-intelligence-act/eu-ai-act-risk-categories>.
- Roleke, M. (2025). 'The European AI Act', *activeMind.legal*, 3 February,
<https://www.activemind.legal/guides/ai-act/>.
- Secure Privacy (2025). *EU AI Act: Understanding Risk-Based Classification*,
<https://ai-eu-act.eu/blog/eu-ai-act-understanding-risk-based-classification/>.

Industry Note

Building Europe's AI Capacity: A Partnership Made in Brittany

Philippe Guillotel

<https://doi.org/10.53121/ELFS10> • ISSN (print) 2791-3880 • ISSN (online) 2791-3899

ABOUT THE AUTHOR

Philippe Guillotel is a Distinguished Scientist and Research Manager at InterDigital, specialising in immersive technologies and video innovation.

Artificial intelligence has moved far beyond the realm of hype and is now an embedded force in everyday life – driving efficiency in hospitals, powering autonomous vehicles, enhancing supply chains, and personalising media experiences. Its transformative potential is frequently compared to the dawn of the internet, and, like that earlier revolution, its trajectory will be determined not only by technical ingenuity, but by the strength of the partnerships that nurture it.

The European Commission's ambition is to make Europe the 'AI continent'. The stakes are clear: leading in AI is not just about asserting technological prowess, it is a strategic imperative to ensure sovereignty over data, standards, and innovation, while developing systems that reflect European values such as trustworthiness, inclusivity, and sustainability. Realising that aspiration demands collaboration at scale, between governments, industry leaders, universities, and research institutes.

This industry note discusses several Interdigital projects in which AI has enhanced researchers' skills and knowledge, giving rise to cutting-edge technologies that align with the EU's long-term policy objectives, such as promoting energy efficiency, interoperability, and sustainability.

Europe's AI aims are codified in the AI Continent Action Plan and reinforced by the AI Act, the first major regulatory framework in the world to address the development and deployment of AI. Such legislation articulates the region's vision for its future: to develop and control its own AI ecosystems, ensuring that economic growth, job creation, and technological breakthroughs happen on European soil.

France has aligned closely with this aspiration, launching successive national AI strategies – most recently France 2030 – that seek to double the number of AI-trained graduates, expand AI research

capacity, and fund industry–academia collaborations. The Brittany region, with its dense network of universities, research centres, and technology companies, has emerged as a proving ground for these ambitions.

NEMO.AI: FROM POLICY VISION TO REALITY

Nemo.AI is a concrete product of France’s policy frameworks. It is a joint initiative between InterDigital, a global research and innovation leader in wireless, video, and AI technologies, and Inria, France’s premier public research institute for digital science, as they explore and foster innovative applications of AI.

Established in 2021 and formally launched in June 2022, the initiative was enabled by public–private partnership funding mechanisms championed by the French National Research Agency. These mechanisms recognise that path-breaking AI research often requires both the theoretical depth of academia and the applied expertise of industry to move from concept to commercial impact.

Path-breaking AI research often requires both the theoretical depth of academia and the applied expertise of industry to move from concept to commercial impact.

InterDigital has spent over five decades pioneering foundational wireless and video technologies – contributions that underpin mobile connectivity and media streaming worldwide. In recent years, it has expanded research into AI and machine learning, not as a standalone discipline but as a force multiplier for its core areas of investigation.

Inria operates nine centres across the country, including in Rennes, Brittany. With more than 3,800 scientists working across 220 project teams, often in collaboration with major universities, Inria has a mission that spans fundamental research, open-source development, and the creation of technology start-ups.

InterDigital and Inria’s collaboration for the Nemo.AI Common Lab is built upon shared goals. Their convergent research priorities – immersive media, AI for digital experiences, and energy-efficient technologies – make their combined capabilities greater than the sum of their individual parts. Both InterDigital and Inria operate research facilities in Rennes, which has enabled frequent in-person discussion and joint access to infrastructure, essential ingredients for productive collaboration.

Inside Nemo.AI

More than a research grant, Nemo.AI is a Common Lab in which resources, people, and ideas are fully integrated. Researchers and engineers from InterDigital work alongside PhD students, post-doctoral fellows, and researchers from Inria, brainstorming and tackling projects that have both near-term applications and long-term strategic importance.

The initiative’s name refers to Captain Nemo, the hero of Jules Verne’s *Twenty Thousand Leagues Under the Sea*, reflecting both regional pride and a spirit of exploration. Its mission is to equip Brittany with the scientific and technical capacity to lead in AI-driven media and communications, while feeding into Europe’s broader AI leadership strategy.

That Nemo.AI has its home in Brittany is not incidental – the region’s reputation as a hub for digital technology, maritime innovation, and creative industries is well established. Rennes, in particular, boasts a high concentration of universities, engineering schools, and R&D centres, alongside a thriving start-up scene.

This makes the area fertile ground for AI experimentation, whereby talent from local universities ensures a steady flow of skilled graduates, industry diversity creates opportunities for cross-sector AI applications, and government support at both regional and national levels provides funding stability and strategic alignment.

By anchoring Nemo.AI in Brittany, the partnership bolsters the region's role as both a test bed for emerging technologies and a magnet for global talent. In addition, it ensures that European priorities like privacy, security, and sustainability are embedded from the outset.

Ys.ai – AI for realistic portable avatars

Inspired by the Breton legend of the submerged city of Ys, this project aims to overcome one of the thorniest problems in immersive environments: realistic, portable avatars. Today, users entering a virtual world – whether for gaming, remote work, training, or healthcare – are often obliged to create a new avatar for each platform. Geometry, appearance, motion, and clothing animation are inconsistent, and avatars frequently fail to capture the subtleties of human body language.

Ys.ai uses AI to learn how people move, gesture, and emote in specific contexts. By reading body language and facial expressions, it generates avatars that mirror a user's real-world behaviour with far greater fidelity. This is not just about making virtual interactions 'look nicer' – in fact, realism and consistency in avatars can improve trust, reduce cognitive dissonance, and make cross-platform immersive experiences more seamless.

A critical goal of Ys.ai is standardisation. InterDigital plays an active role in the MPEG consortium, where the groundwork is being laid for global avatar standards. A common specification would mean users could carry a single, consistent digital identity across multiple platforms – a step-change in interoperability similar to the adoption of JPEG for images or MP4 for video.

Nisk.ai-AI for energy-efficient video coding

Niskai is a Celtic water divinity. The choice of this name for the project reflects its concern with sustainability. Launched in January 2025, Nisk.ai addresses the escalating energy footprint of global video consumption, which now accounts for the majority of internet traffic. AI-driven approaches to video compression can deliver equivalent visual quality at lower bitrates, reducing both bandwidth requirements and energy consumption.

Beyond compression, the research explores semantic adaptation – optimising video for different devices, screen sizes, network conditions, and even lighting environments. By preserving the intent and clarity of content while minimising data load, AI-enhanced codecs can make high-quality video experiences more accessible, sustainable, and resilient.

Both Ys.ai and Nisk.ai exemplify how regional research excellence can influence global standards and markets – tying Brittany's innovation ecosystem to worldwide adoption.

STANDARDISATION THROUGH COLLABORATION CREATES COMPETITIVE EDGE

Industry–academia collaborations are not unique to Europe, but the European model places particular emphasis on public–private co-investment, regulatory alignment, and talent development. This allows projects like Nemo.AI to take part in a broader continental strategy.

The benefits are tangible. They include accelerated innovation, as academic research gains practical direction from industry and industry gains early access to breakthrough concepts; talent retention, as researchers work on trailblazing projects close to home; and global influence through standards contributions for technologies that shape international norms.

Standards are the invisible infrastructure of the digital world – without them, interoperability breaks down, markets fragment, and innovation stalls. Europe's leadership in bodies like 3GPP, MPEG, and ETSI has historically been a competitive strength, and Nemo.AI is reinforcing that position.

The work led by Ys.ai in avatar standardisation is a prime example. Once adopted, it could underpin an entire generation of immersive applications, from metaverse platforms to telehealth consultations. Similarly, advances in AI-driven video coding could inform future MPEG video standards, enabling greener, more efficient media delivery worldwide.

The EU is currently revising its legal framework for standardisation. In this context, a firm emphasis needs to be placed on maintaining openness and collaboration. Approaches guided by these values have proven successful in bringing trusted, interoperable, and inclusive cutting-edge technologies to consumers and businesses alike, and can ensure that innovation serves the common good.

COLLABORATION: A BLUEPRINT FOR EUROPE'S AI FUTURE

Nemo.AI offers a replicable model for how Europe can translate policy aspirations into competitive advantage. As Europe moves toward its goal of becoming the 'AI continent', these principles will be essential – not just for flagship projects but for the many smaller collaborations that collectively mould the EU's innovation landscape.

From its base in Brittany, Nemo.AI is helping to meet challenges that are global in scale: how to make immersive experiences more human, how to deliver high-quality video sustainably, and how to make certain that AI evolves in ways that serve people as much as markets. In addition, Nemo.AI offers a space for fostering scientific excellence in innovation, attracting international students in a competitive market and creating a high-level research environment for scientists and engineers.

It also demonstrates a truth that innovators and policymakers alike must remember: technological leadership is built on relationships as much as on research, knowledge, and expertise. When academic insight meets industrial capability – within a supportive policy framework – regions like Brittany can punch far above their weight in the global technology arena.

When academic insight meets industrial capability – within a supportive policy framework – regions like Brittany can punch far above their weight in the global technology arena.

On 16 July 2025, while the European Union was preparing its long-term budget, the European Commission unveiled its proposal for the next Multi-Annual Financial Framework (MFF). A major focus of the MFF is to support research, innovation, and industrial scaling, together with the adoption of critical digital technologies. Part of this new budgetary framework is the establishment of a Competitiveness Fund (ECF), which endeavours to structurally strengthen the EU's competitiveness in strategic technological and industrial sectors. In the context of this new fund, it is important that support of research and innovation remains a key priority of the EU.

Closely linked to the Competitiveness Fund is the new Horizon Europe 2034 research programme. Horizon Europe forms a core component of the ECF, with a separate budget of 451 billion euros. The participation framework of Horizon Europe 2034 facilitates targeted cooperation with third countries. The success of Horizon Europe has historically depended heavily upon international cooperation and the participation of third countries. It is crucial that the EU maintains this open model, indispensable for collaboration with international partners and non-EU researchers, to safeguard the Horizon Europe programme's continued success and global reach.

As we approach this new era of AI opportunity, the lesson is clear: when we collaborate, we win.

Part 2

AI FOR SOCIETY

Artificial Intelligence as a Socio-Technical System: What It Is and Where It May Take Us

Francesco Cappelletti and Dr Francesco Goretti

<https://doi.org/10.53121/ELFS10> • ISSN (print) 2791-3880 • ISSN (online) 2791-3899

ABSTRACT

This chapter offers a discursive and accessible account of artificial intelligence (AI) as a socio-technical phenomenon. The first section introduces technical foundations for non-specialists, explaining key mechanisms such as gradient descent and transformers. The second situates AI in society, law, and governance, with an emphasis on Europe's human-centric model. The third illustrates applications and risks in insurance, cybersecurity, healthcare, and public administration. Finally, the chapter reflects philosophically on AI's meaning and Europe's strategic role in shaping its trajectory. It aims to balance information with accessibility.

ABOUT THE AUTHORS

Francesco Cappelletti is a researcher and Cyber and Data Security Lab & Adjunct Professor at the Brussels School of Governance (Vrije Universiteit Brussels), and a Senior Fellow at the European Liberal Forum.

Dr Francesco Goretti is a biomedical engineer & researcher at the European Laboratory for Nonlinear Spectroscopy, University of Florence.

INTRODUCTION

Artificial intelligence is not simply a new set of tools; it is a reconfiguration of how knowledge is produced, decisions are taken, and public value is created. In practical settings, AI systems already mediate access to credit and insurance, help doctors and nurses to triage patients, and support analysts as they sift through mountains of cybersecurity telemetry in search of weak signals of intrusion. Although it exhibits uneven pacing and frequently overstates its promises, the dissemination of AI from research laboratories into everyday practice is unmistakably evident.

When it comes to political choices, the challenge is to translate technical progress into public value without sleepwalking into unintended concentrations of power or novel forms of exclusion. For readers outside technical disciplines, the challenge is to grasp enough of the essentials to ask the right questions, commission wisely, and design proportionate oversight. This chapter endeavours

to demystify the technical core while showing how institutional and ethical decisions determine whether AI promotes democracy and prosperity or weakens them.

AI must be understood as a socio-technical construct inseparable from the political, economic, and normative structures that both shape and are shaped by its deployment. It operates at the intersection of technical architectures, data governance practices, institutional arrangements, and societal expectations. Large language models (LLMs) exemplify this interplay: intelligent systems trained on immense datasets and optimised algorithms, they are employed to carry out computational tasks. Simultaneously, they alter epistemic authority, transforming how knowledge is created, shared, and secured.

Two conceptual frameworks guide the chapter. The first treats AI as a socio-technical system with capabilities that are inextricable from the data, institutions, labour, and norms that sustain them. In practice, recognising AI as a socio-technical system means that governance cannot focus solely on data or algorithms but must also integrate organisational routines, procurement chains, and human oversight mechanisms that shape how systems perform in a given context.

The second posits that effective governance should be outcome-focused rather than technology-prohibitive. That is, we can protect people and competition by measuring systems against the values we care about – safety, privacy, fairness, and accountability – while remaining flexible about how those outcomes are achieved.

The structure of this chapter is straightforward. First, we offer a discursive technical primer. Next, AI is situated within society and governance and assessed through a European lens. The chapter then turns to high-stakes applications and closes with a philosophical reflection on agency and plural values.

A TECHNICAL PRIMER

AI is just numbers – but understanding this helps to clarify how modern AI works

The process of making a programme clever – that is, creating an AI – primarily relies on a technique known as machine learning (ML). A ‘learning system’ adjusts its internal settings (or decision weights) based on experience, so its results become more accurate, or valuable, as it encounters more examples.

There are three main types of learning: supervised learning, where the system learns from labelled examples that show the correct answer; self- or unsupervised learning, which involves discovering unlabelled patterns or structures in data, akin to how LLMs like GPT understand and generate language; and reinforcement learning, where the system learns to make decisions by trying different actions and learning from the outcomes, similar to trial and error. These methods form the foundation of modern AI.

Every learning pipeline comprises four essential components: data, a model class, an objective (loss) to optimise, and a procedure for optimisation. Think of the ‘model class’ as the core of the system. Different models employ various strategies to generate patterns that enable learning to recognise unseen data. The ‘loss function’ measures how accurately the system’s predictions match actual outcomes (basically, the error). The ‘optimiser’ is the search process that adjusts the model’s parameters, exploring different options within the search space to find a configuration that minimises the error.

For non-experts, it’s useful to think of this process as teaching a system: you provide the data, define what good results look like (loss), select the type of model that can learn from that data, and finally, use an algorithm (the optimiser) to adjust the model until it performs as well as possible.

Throughout this process, data quality is of the utmost importance. Filtering out irrelevant data, removing duplicates, and tracking the origin of the data aren’t just minor steps; they are essential for building trustworthy and reliable AI systems.

Models, losses, optimisation

A model acts like a complex computational system that takes in various inputs (such as data or signals) and produces outputs (such as predictions or classifications). It does this by adjusting its internal settings, known as weights or coefficients, through a process called learning. Tuning a musical instrument to get the right sound is a good analogy. Training the model involves repeatedly changing these settings to make sure it performs well on a specific task.

Training can take several forms. The following table contrasts two main approaches commonly used to training an AI system, reflecting pragmatic trade-offs:

Aspect	Static LLM Deployment	Dynamic Online Learning
Adaptation	Fixed parameters after training; no real-time updates	Continual adjustment based on new data streams
Examples	Most chatbots like GPT series in production	Experimental systems in adaptive environments (e.g., real-time fraud detection)
Pros	Stability and predictability	Flexibility to evolving threats
Cons	May become outdated without retraining	Risk of instability or privacy issues from ongoing data use

It is worth noting that while user feedback on models like OpenAI’s famous ChatGPT may inform future updates by developers, individual conversations do not alter the model in real time, thus maintaining its static nature during use.

Optimisers such as gradient descent can help to reduce error. In short, optimisation adjusts parameters to minimise a task-appropriate loss; the choice of loss and hyperparameters sets practical limits on reliability, documentation, and reproducibility in high-risk settings.

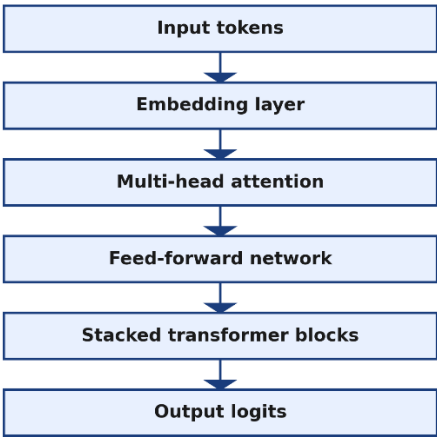
Transfer learning is another strategy widely used to develop AI models. It consists of taking a pre-trained model and fine-tuning it to new instances. A common scenario is to take huge models (both in terms of parameters and pre-training instances) and train them for the same type of task but with different samples; for example, one could adapt a big model trained in recognising dogs and tune it to recognise cats by using just a few samples. This type of approach is very valuable when the available data is not abundant enough to properly train a model from scratch.

Today’s complex networks are an evolution of logical units. A logical unit, which can be imagined as a single neuron of a neural network, computes the weighted sum of all its inputs, passes it to the activation function, and provides an output. Starting from this very basic configuration, thousands and thousands of units can be combined, processing billions of weights and applying complex activation functions: this is the recipe for deep neural networks.

Neural networks can be specialised by choosing and arranging different types of layers (Figure 1). Some layers break down images into basic features, others extract useful information for tasks like classification, and others still help improve training and prevent overfitting. These layers are grouped into blocks; the way they are combined determines how the network behaves. A deep learning expert (or even an AI system today) designs the best neural network structure by adjusting it based on performance.

FIGURE 1: Schematic structure of a neural network.

**Dictionary-trained Neural Network
(schematic mega-structure).**

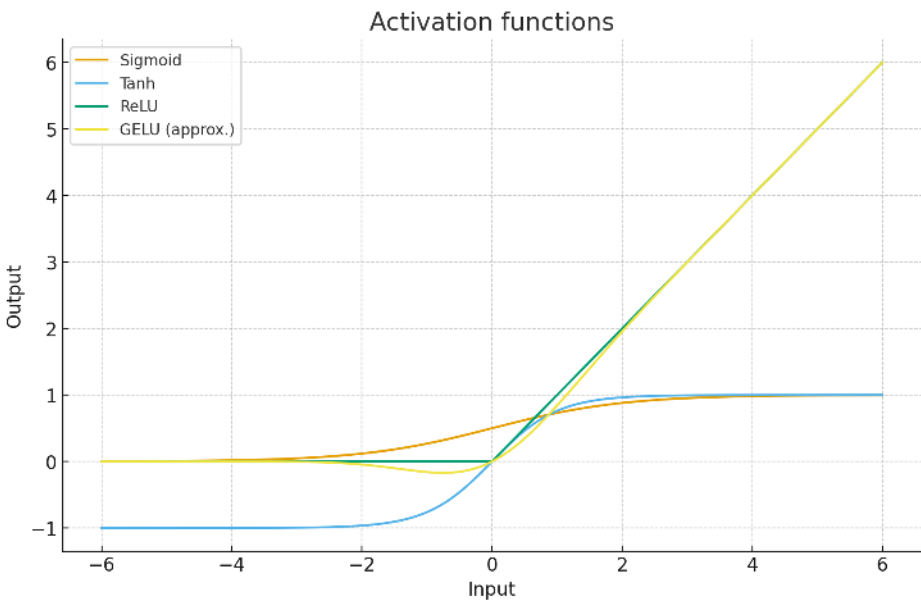


Source: author's elaboration.

Non-linear representation

Non-linear activation functions enable deep neural networks to recognise complex patterns that linear models cannot capture. Typical examples include sigmoid, tanh, ReLU, and GELU, which vary primarily in their smoothness and numerical stability (Figure 2). The fundamental principle is straightforward: these functions introduce non-linearity, transforming weighted sums into expressive decision boundaries while maintaining differentiability to facilitate efficient training.

FIGURE 2: Activation functions (sigmoid, tanh, ReLU, GELU) with distinct shapes and saturation behaviour, showing how non-linearity enables networks to model complex, non-linear relationships in data. These are simple 'switches' inside the AI that decide how strongly each input affects the output, allowing the system to learn curved or irregular patterns instead of just straight lines.



Source: author's elaboration.

Because these activations are differentiable, we can determine how much each layer contributed to an error and adjust it – this is where backpropagation comes in.

Backpropagation leverages the chain rule to propagate errors backwards through each layer, calculating gradients that guide weight adjustments. Contemporary methods such as residual connections and normalisation layers improve the stability and efficiency of training processes.¹ Dependable gradients ease the expansion of exceedingly deep architectures, with transformers currently predominating.

From sequence models to transformers

Transformers are a type of artificial intelligence that understand language by considering the importance of each word or token in relation to others within a sentence or paragraph. This method allows transformers to process information in parallel, unlike older models that operated step by step.² LLMs are initially trained on vast amounts of textual data to predict the next word in a sentence, helping them understand language in a broad, general way. Afterwards, they are fine-tuned for specific tasks or made safer and more helpful through further adjustments such as reinforcement learning from human feedback (RLHF). These steps ensure that models are helpful and safe for real-world applications. The strength of transformers lies in their capacity to grasp the context of a given sentence and their efficiency in computing inferences.³ Even so, what matters is not just training performance but whether the model generalises to new data.

Generalisation and regularisation

If we train a model to learn from one set of data, how can we ensure that it also performs well on new data? Generalisation tackles this exact problem.⁴ It trains the model on one subset, selects parameters on a separate subset, and subsequently reports accurate performance metrics on an independent, held-out test set.

With this in mind, we can skip detailed equations and focus on a fundamental intuition: training reduces error, and evaluation verifies whether a skill transfers. In summary, it can be asserted that contemporary artificial intelligence is devoid of enchantment – it is merely large-scale numerical optimisation.

The physical backbone of AI: RAM, GPUs, and lots of speed

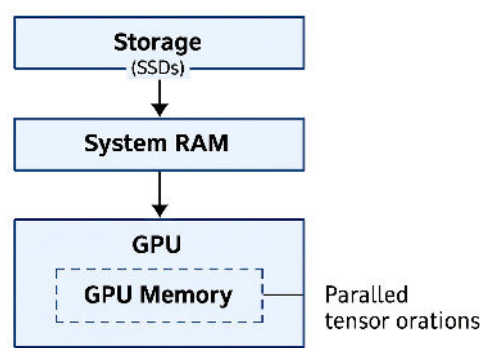
The models under discussion are essentially rapid, large-scale matrix processors; the hardware that runs them is of paramount importance. Graphics Processing Units (GPUs), originally created for graphics rendering, are now vital for training and inference tasks because of their ability to perform many operations at once. This differs from Central Processing Units (CPUs), which are optimised for fewer sequential tasks.⁵ Such parallelism can cut training times from weeks to just hours when the right hardware is used.⁶

-
1. M. P. Deisenroth, A. A. Faisal, and C. S. Ong (2020), *Mathematics for Machine Learning* (Cambridge: Cambridge University Press).
 2. On language models, see chapter 22 of S. J. Russell and P. Norvig (2020), *Artificial Intelligence: A Modern Approach* (4th ed.) (Pearson).
 3. For an interactive visualisation that breaks down the LLM algorithm (as used in ChatGPT), see Bycroft's tool at <https://bbycroft.net/llm>.
 4. This term refers to the capability to perform effectively on novel data, fostered by the application of regularisation techniques (e.g., weight decay, dropout) and the implementation of strict data partitions (training, validation, testing), thereby preventing the leakage of test information during hyperparameter tuning.
 5. S. Mukherjee (2024), 'GPUs vs CPUs explained simply: Parallel computing with CUDA', DigitalOcean, 24 December, <https://www.digitalocean.com/community/tutorials/parallel-computing-gpu-vs-cpu-with-cuda>.
 6. Scale Computing (2025), 'GPU architecture explained: Structure, layers & more', SC//Insights, 16 April, <https://www.scalecomputing.com/resources/understanding-gpu-architecture>.

Generally, data is transferred from storage to system RAM and then to the GPU, where tensors are processed in parallel (Figure 3). GPU memory (vRAM) provides exceptionally high bandwidth (up to ~7.8 TB/s) compared with typical CPU memory (~50 GB/s) – keeping data within the GPU is essential.⁷ When vRAM capacity is exceeded, data swapping via slower pathways reduces performance. On a system level, a large training operation has been measured at approximately 1,287 MWh.⁸ Improving efficiency (through specialised accelerators, optimised scheduling, and advanced cooling methods) is now regarded as a key aspect of responsible infrastructure planning.⁹ Additionally, access to high-vRAM configurations influences who can train large-scale models, reinforcing the need for shared European computing resources.¹⁰

FIGURE 3: Simplified hardware architecture for AI training.

Data Flow During AI Training



Source: author’s elaboration.

From the user’s perspective, however, processes occur seamlessly and imperceptibly: a prompt is converted into numerical tokens, processed through attention and feed-forward layers on the GPU utilising large-scale parallel linear algebra, and subsequently reconstructed into text, which manifests as a straightforward conversational exchange. In actuality, what appears as dialogue is a swift sequence of trained numerical operations – not genuine understanding, but statistically driven prediction – based on patterns derived from extensive human data that the system emulates, inheriting both its virtues and flaws.¹¹

7. Pure Storage (2025), ‘CPU vs. GPU for machine learning’, Purely Technical Blog, 22 February, <https://blog.purestorage.com/purely-technical/cpu-vs-gpu-for-machine-learning>.

8. J. You, J. W. Chung, and M. Chowdhury (2023), ‘Zeus: Understanding and optimizing GPU energy consumption of DNN training’, Proceedings of the 20th USENIX Symposium on Networked Systems Design and Implementation, 119–139. <https://www.usenix.org/system/files/nsdi23-you.pdf>.

9. D. Harris (2024), ‘How AI and accelerated computing are driving energy efficiency’, NVIDIA, 22 July, <https://blogs.nvidia.com/blog/accelerated-ai-energy-efficiency/>; IEA (2025), Energy and AI, Report (Paris: IEA), <https://www.iea.org/reports/energy-and-ai>.

10. IEA, Energy and AI.

11. IEA, Energy and AI.

AI IN SOCIETY AND GOVERNANCE

The historical evolution of AI is marked by an interplay between technical innovation and shifts in governance and societal perception. Early definitions focused on computational logic and symbolic manipulation, treating AI systems as isolated artefacts. Over time, machine learning brought in adaptivity and data-driven methods, which aligned with socio-technical perspectives as digitalisation advanced and networked infrastructures became more prevalent. Concepts like autonomy became central to laws such as the EU AI Act, which links technical outputs to real-world impacts. Evaluative frameworks have evolved from mere performance metrics to include safety and ethics, embedding accountability in certification. Parallel US approaches have emphasised public–private partnerships and security, diverging from Europe’s focus on privacy. Modern discourse views AI as both a tool and a threat, with ISO/IEC standards maturing to integrate ethics into lifecycle management.

In Europe, discussions about AI primarily stress the importance of developing trustworthy, human-centred systems that uphold ethical principles and societal values.

Artificial intelligence presents a complex governance issue requiring careful analysis and strategic policy development. In Europe, discussions about AI primarily stress the importance of developing trustworthy, human-centred systems that uphold ethical principles and societal values. To meet such ideals, it is imperative to implement policies that favour open data sharing, establish interoperability standards across various platforms, and

strengthen international cooperation. These are activities that transcend technical guidelines and can give rise to a comprehensive regulatory framework. Above all, AI must be harnessed in ways that uphold individual rights, promote transparency, and facilitate effective democratic oversight, thereby ensuring that technological progress benefits the broader society.

Network society and algorithmic governance

Today’s internet and its networks heavily rely on algorithms, which act as automated systems that mediate the allocation of various resources, such as information, opportunities, and risks. This mediation process can sometimes lead to a concerning development known as ‘algocracy’, where decision-making and power are increasingly governed by opaque algorithmic models rather than transparent human judgement. To prevent such an outcome and maintain accountability and fairness, European strategies have stressed proportionality (balanced, context-appropriate measures) and contestability (the ability to challenge and scrutinise automated decisions).¹²

The human-centric European approach

As has been stated, the European Union affirms that AI growth must be aligned with safety, legality, and trustworthiness. Initiatives supporting FAIR (Findable, Accessible, Interoperable, Reusable) open data play a crucial role in fostering innovation within the field. However, these initiatives must include safeguards to prevent risky practices and biases.¹³ Institutions should also have recourse mechanisms and ensure that the level of explainability is proportionate to the effects and risks posed by the AI applications (that is, higher-risk systems require greater transparency to enable oversight and accountability).

12. M. Schuilenburg and R. Peeters (2021), *The Algorithmic Society: Technology, Power, and Knowledge* (Abingdon: Routledge).

13. S. Ziesche (2023), *Open data for AI: What now?*, UNESCO, <https://www.unesco.org/en/articles/open-data-ai-what-now>.

Within the European Union, regulations like the GDPR (Regulation (EU) 2016/679) and NIS2 (Directive (EU) 2022/2555) have established fundamental standards for privacy and cybersecurity in the deployment of artificial intelligence. Building upon this legislation, the European Union enacted Regulation (EU) 2024/1689, known as the Artificial Intelligence Act, on 12 July 2024. The Act harmonises AI rules and employs a risk-based supervisory model, rendering it one of the most comprehensive regulatory measures in the field of artificial intelligence.

In contrast with other legislative responses, such as the US AI Initiative or China's AI regulations, the EU's AI Act seeks to establish a coherent legal framework that addresses risks while driving innovation.¹⁴ It categorises AI systems based on their risk levels and sets specific obligations for developers and users accordingly. The Act's horizontal standards vouchsafe baseline consistency across sectors, while its vertical standards cater to domain-specific requirements. The resulting legal framework, provided it remains flexible, can minimise fragmentation and support interoperability.

Global governance

The impacts of AI transcend national borders, requiring coordination at both political and technical levels. The UN High-Level Advisory Body on AI has proposed an international scientific panel, standards exchanges, and capacity development to prevent fragmentation and guarantee equitable access to benefits.¹⁵

A persistent obstacle is the lack of uniform technical standards. Current governance frameworks, including the EU AI Act, rely heavily on conformity assessments, yet many standards remain incomplete or inconsistent. ISO/IEC initiatives, such as ISO/IEC 22989 (AI terminology), ISO/IEC 23053 (framework for AI systems using machine learning), and ISO/IEC TR 24028 (trustworthiness in AI), provide important guidance but stop short of operational criteria. Without reproducible test protocols, shared benchmarks, and lifecycle governance norms, there is a risk of uneven compliance across jurisdictions.

Europe, together with its OECD, G7, and G20 partners, has an opportunity to lead by developing interoperable standards that reconcile horizontal requirements (e.g., fairness, robustness, security) with the needs particular to specific sectors (such as healthcare, mobility, and finance). Creating such standards would allow regulators to verify compliance in consistent, measurable ways while avoiding duplication. A combination of political coordination and technical harmonisation is essential: governance without standards risks vagueness, while standards without governance risk irrelevance.

HIGH-STAKES APPLICATIONS AND RISKS

The practical implementation of AI unveils a spectrum of opportunities and challenges, especially within high-impact sectors such as insurance, cybersecurity, healthcare, and public administration. In these domains, AI offers the promise of increased efficiency; however, it also introduces risks related to bias, errors, and security vulnerabilities. Specifically, in the fields of insurance and healthcare, biased or incomplete training datasets can serve to perpetuate social inequalities or result in misclassification in clinical settings.¹⁶

The practical implementation of AI unveils a spectrum of opportunities and challenges, especially within high-impact sectors such as insurance, cybersecurity, healthcare, and public administration.

14. In the United States, the National Artificial Intelligence Initiative Act of 2020 (enacted as Public Law 116-283, Division E, on 1 January 2021) created the National AI Initiative Office and mandates NIST work on AI standards and risk frameworks.

15. United Nations High-Level Advisory Body on AI (2024), Governing AI for humanity. <https://doi.org/10.18356/9789211067873>.

16. D. Noordhoek (2023), 'Regulation of artificial intelligence in insurance: Balancing consumer protection and inno-

In the field of cybersecurity, the same generative models that enhance detection capabilities also facilitate adversarial attacks and sophisticated phishing schemes, exemplifying the dual-use nature of this technology.¹⁷ European regulations – most notably the EU AI Act – address these concerns through mandates for transparency, accountability, and human oversight, while frameworks such

as NIS2 bolster resilience within critical infrastructure sectors. Furthermore, public-sector deployments, as evidenced by algorithmic decision failures in welfare and education, underscore the necessity for contestability and remedial mechanisms.¹⁸

Collectively, these cross-sector experiences highlight the vital importance of empirical testing and outcome-oriented regulation. However, underlying these questions of governance is a more profound issue: what systems that operate with such proficiency, yet lack understanding, say about our very concepts of intelligence and judgement. It is to this theme that we now turn.

REFLECTIONS AND PHILOSOPHICAL CHALLENGES

The process of optimisation, which seeks to condense a multitude of human values into singular computational objectives, raises urgent questions about the emergence of technocratic governance (especially if such systems operate beyond meaningful political oversight). Opacity in AI systems refers to the difficulty – or often, impossibility – of fully reconstructing the computational steps by which a model arrives at a particular output. This arises from high-dimensional parameter spaces, non-linear transformations, and adaptive mechanisms in modern AI architectures (hence the discussion of AI technicalities at the start of this chapter). In an era characterised by an abundance of quickly accessible information, the more pressing matter is not what we are capable of knowing but rather what we elect to comprehend – and for what reasons. Trust is ultimately a social and moral construct; it cannot be provided solely through optimisation.

Online learning is not static between training cycles, making it an exception in this context. During inference, internal states (such as attention weights and latent embeddings) are computed dynamically; however, these do not modify the underlying model. This complexity poses challenges to interpretability, even for specialists. Nonetheless, research on interpretability has progressed: tools like SHAP, which emphasises feature importance, and LIME, which offers local, model-agnostic explanations, provide partial insights into decision pathways, although their validity is context-dependent.¹⁹

Mechanistic interpretability endeavours to provide causal insights within the layers of models; however, genuine transparency remains elusive, thereby reducing accountability in regulatory contexts and sparking philosophical debates around moral responsibility and agency in AI decision-making. Explainability techniques, whether intrinsic or post hoc (e.g., saliency mapping), assist in narrowing this gap but may lead to oversimplification. Opacity in AI systems affects developers by concealing biases, complicates compliance with the EU AI Act for operators, and diminishes trust among users. Its mitigation requires a multi-layered approach: comprehensive reporting for auditors and accessible explanations for end users, in accordance with international standards.

vention', Geneva Association, 14 September, <https://www.genevaassociation.org/publication/public-policy-regulation/regulation-artificial-intelligence-insurance-balancing>; Russel and Norvig, Artificial Intelligence.

17. A. Kucharay, O. Plancherel, V. Mulder, A. Mermoud, and V. Lenders (eds.) (2024), *Large Language Models in Cybersecurity: Threats, Exposure and Mitigation* (Cham: Springer); R. Islam (2025), *Generative AI, Cybersecurity, and Ethics* (New York: Wiley).
18. A. Kelly (2021), 'A Tale of Two Algorithms: The Appeal and Repeal of Calculated Grades Systems in England and Ireland in 2020', *British Educational Research Journal*, 47(3), 725–741.
19. S. M. Lundberg and S. I. Lee (2017), 'A Unified Approach to Interpreting Model Predictions', *Advances in neural information processing systems*, 30, 1–10. <https://doi.org/10.48550/arXiv.1705.07874>; M. T. Ribeiro, S. Singh, and C. Guestrin (2016), "Why Should I Trust You?" Explaining the Predictions of Any Classifier', *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 1135–1144. <https://doi.org/10.1145/2939672.2939778>.

The following table summarises key ways to address the opacity and explainability risks in high-risk AI systems, aligned with European regulatory priorities:

Category	Mitigation Measures
Operational Integration	Align explanation formats to stakeholder roles (e.g., technical for auditors, narrative for users); test explanations against causal pathways; incorporate adversarial validation to prevent manipulation; document interpretability–performance trade-offs in compliance reports.
Deployment Practices	Map outputs to regulatory explanation depths (e.g., EU AI Act Annex IV); use dual documentation layers (regulator vs. user); conduct audits for reasoning consistency; apply privacy-preserving techniques per GDPR; automate reproducibility checks.
Policy Actions	Mandate certified explainability modules with fidelity testing; standardise multi-tier transparency frameworks under ISO/IEC and AI Act; require disclosure of explanation limitations; incentivise secure interpretability research; run sandbox trials for interpretability–security balance; integrate GDPR in explanation templates; promote case studies on explanation-resilience balances.

Ethical frameworks for AI use must be grounded in both technological realities and socio-political settings. They should articulate principles for design, deployment, and evaluation, weighing benefits against risks to rights and welfare. Ethics require lifecycle integration with continuous monitoring. As a socio-technical construct, value judgements link technical choices to political ones.

Intelligence without understanding

Artificial intelligence systems emulate intelligence through pattern recognition and predictive capabilities without possessing genuine understanding or consciousness – a vital distinction highlighted by most scholars. Nevertheless, these functionalities are already transforming organisational and governance frameworks, necessitating the implementation of comprehensive safeguards such as validation, auditing, contestation, and recourse. These mechanisms serve as pragmatic alternatives to the pursuit of absolute certainty, thereby maintaining public trust, ensuring ethical integrity, and protecting societal interests. Institutions cannot rely on technical assurances – they must create and deploy human-centred procedures.

CONCLUSIONS

Artificial intelligence, particularly LLMs and general-purpose AI systems, functions at the intersection of innovation and socio-political governance, demanding an integrated approach that encompasses technical design, data management, ethical considerations, and regulatory compliance. Transparency and accountability emerge as core principles, requiring detailed documentation of training data provenance, architectural configurations, and operational safeguards to ensure alignment with fundamental rights and safety standards.

Security challenges extend beyond code-level vulnerabilities to include adversarial manipulations, supply chain risks, and infrastructure exposures, necessitating continuous monitoring, adaptive defences, and coordinated incident response mechanisms. The layered regulatory environment in Europe, combining the EU AI Act, GDPR, and NIS2 directives, creates a framework that mandates harmonised compliance efforts, striking a balance between privacy, cybersecurity, and ethical obli-

gations. This regulatory coherence supports the development of standardised evaluation protocols and conformity assessments that can verify system resilience, fairness, and lawful data processing.

Sustainability considerations, particularly in connection with energy consumption and environmental impact, are increasingly recognised as central to responsible AI deployment. Standardised measurement and reporting practices, integrated into security and performance assessments, are imperative. Meanwhile, the burgeoning concentration of data ownership and control raises concerns about equitable access, transparency, and the potential for monopolistic dynamics, underscoring the importance of open standards, third-party audits, and cross-jurisdictional cooperation.

Philosophical and normative discourses acknowledge a rising tension between opacity and explainability, highlighting the need for multi-level transparency that serves diverse stakeholders without compromising security or proprietary interests. Ethical frameworks must be embedded throughout AI lifecycles, ensuring that value alignment, harm prevention, and human rights protections are operationalised through concrete governance checkpoints and technical measures.

Improved AI governance urgently requires harmonised standards that codify lifecycle management, adversarial robustness testing, data provenance verification, and privacy-preserving monitoring. These standards should facilitate interoperability across sectors while accommodating domain-specific nuances, enabling consistent conformity assessments and reducing fragmentation. Coordinated policy efforts and research incentives are essential. This consensus may ultimately need to be sought at a global level.²⁰

The trajectory of AI will depend in large part on governance decisions: regulatory structures, standards, and institutional cultures will shape how technical potential results in societal outcomes. Ultimately, artificial intelligence must be guided by human intelligence. Because legitimacy cannot be automated, durable governance will hinge on the judgements made throughout the AI lifecycle – supported, not replaced, by technical controls. Ongoing oversight, engagement from multiple stakeholders, and flexible governance mechanisms will be crucial to maintaining trust and resilience as AI capabilities grow and expand across various fields.

20. See F. Cappelletti (2023), 'To AI or not to AI? Towards a treaty on Artificial Intelligence', Euractiv, 24 May, <https://www.eureporter.co/uncategorized/2023/05/24/to-ai-or-not-to-ai-towards-a-treaty-on-artificial-intelligence/>.

REFERENCES

- Cappelletti, F. (2023). 'To AI or not to AI? Towards a treaty on Artificial Intelligence', Euractiv, 24 May, <https://www.eureporter.co/uncategorized/2023/05/24/to-ai-or-not-to-ai-towards-a-treaty-on-artificial-intelligence/>.
- Deisenroth, M. P., Faisal, A. A., & Ong, C. S. (2020). *Mathematics for Machine Learning*. Cambridge: Cambridge University Press.
- Harris, D. (2024). 'How AI and accelerated computing are driving energy efficiency', NVIDIA, 22 July, <https://blogs.nvidia.com/blog/accelerated-ai-energy-efficiency/>.
- IEA (2025). *Energy and AI*. Report. Paris: IEA. <https://www.iea.org/reports/energy-and-ai>.
- Islam, R. (2025). *Generative AI, Cybersecurity, and Ethics*. New York: Wiley.
- Kelly, A. (2021). 'A Tale of Two Algorithms: The Appeal and Repeal of Calculated Grades Systems in England and Ireland in 2020'. *British Educational Research Journal*, 47(3), 725–741.
- Kucharavy, A., Plancherel, O., Mulder, V., Mermoud, A., & Lenders, V. (eds.) (2024). *Large Language Models in Cybersecurity: Threats, Exposure and Mitigation*. Cham: Springer.
- Lundberg, S. M. & Lee, S. I. (2017). 'A Unified Approach to Interpreting Model Predictions'. *Advances in neural information processing systems*, 30, 1–10. <https://doi.org/10.48550/arXiv.1705.07874>.
- Mukherjee, S. (2024). 'GPUs vs CPUs explained simply: Parallel computing with CUDA', *DigitalOcean*, 24 December, <https://www.digitalocean.com/community/tutorials/parallel-computing-gpu-vs-cpu-with-cuda>.
- Noordhoek, D. (2023). 'Regulation of artificial intelligence in insurance: Balancing consumer protection and innovation', Geneva Association, 14 September, <https://www.genevaassociation.org/publication/public-policy-regulation/regulation-artificial-intelligence-in-insurance-balancing>.
- Pure Storage (2025). 'CPU vs. GPU for machine learning', *Purely Technical Blog*, 22 February, <https://blog.purestorage.com/purely-technical/cpu-vs-gpu-for-machine-learning>.
- Ribeiro, M. T., Singh, S., & Guestrin, C. (2016). "Why Should I Trust You?" Explaining the Predictions of Any Classifier'. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 1135–1144. <https://doi.org/10.1145/2939672.2939778>.
- Russell, S. J. & Norvig, P. (2020). *Artificial Intelligence: A Modern Approach* (4th ed.). Pearson.
- Scale Computing (2025). 'GPU architecture explained: Structure, layers & more', *SC//Insights*, 16 April, <https://www.scalecomputing.com/resources/understanding-gpu-architecture>.
- Schuienburg, M., & Peeters, R. (2021). *The Algorithmic Society: Technology, Power, and Knowledge*. Abingdon: Routledge.
- Siegel, D. S. (ed.) (2017). *World Scientific Reference on Innovation (vol. 1): University Technology Transfer and Academic Entrepreneurship*. World Scientific.
- United Nations High-Level Advisory Body on AI (2024). *Governing AI for humanity*. <https://doi.org/10.18356/9789211067873>.
- You, J., Chung, J.-W., & Chowdhury, M. (2023). 'Zeus: Understanding and Optimizing GPU Energy Consumption of DNN Training'. *Proceedings of the 20th USENIX Symposium on Networked Systems Design and Implementation*, 119–139. <https://www.usenix.org/system/files/nsdi23-you.pdf>.
- Ziesche, S. (2023). *Open data for AI: What now?*. UNESCO. <https://www.unesco.org/en/articles/open-data-ai-what-now>.

AI for Productivity: Increasing Adoption Levels and Diffusing Technologies in the Workforce

Dejan Ravšelj and Aleksander Aristovnik

<https://doi.org/10.53121/ELFS10> • ISSN (print) 2791-3880 • ISSN (online) 2791-3899

ABSTRACT

This chapter provides an evidence-based assessment of how artificial intelligence (AI) can address the EU's productivity challenges. It explores how AI is redefining human roles in the workforce, examining the rapid but uneven adoption rates across EU countries and occupations. While AI enhances productivity by automating routine tasks and augmenting human skills, its benefits are not uniform, leading to a potential polarisation in the labour market. The analysis demonstrates that the effective diffusion of AI, supported by investments in workforce upskilling and responsible governance, is key to transforming technological progress into sustainable economic growth and long-term competitiveness for the EU.

ABOUT THE AUTHORS

Dejan Ravšelj is an Assistant Professor at the University of Ljubljana, specialising in smart governance, e-government, and public sector economics. He has co-authored studies on AI adoption and digitalisation in public institutions across Europe.

Aleksander Aristovnik is a Professor of Economics and Public Management at the University of Ljubljana. He is a leading researcher in the fields of AI adoption and digital transformation within the public sector.

INTRODUCTION

Artificial intelligence (AI) is now recognised as a major productivity catalyst, sparking dynamic debates across the EU about the implications for employment and work organisation. With the breakthrough emergence of generative AI in 2022, this technology is poised to reshape productivity patterns and redefine the value of knowledge and creativity in the economy. AI, as defined by the European Union (EU) AI Act, refers to machine-based systems capable of operating with varying levels of autonomy and adaptation after deployment. Such systems process input data to generate outputs, including predictions, content, recommendations, or decisions that can influence both

physical and virtual environments.¹ While AI offers powerful tools to boost efficiency and economic performance, it also raises critical questions about job transformation, skills adaptation, and the evolving nature of work in the digital age.² Given its rapid rate of adoption and its impact so far, many now believe that generative AI could become the most transformative technology in decades.³

As AI holds the potential to significantly enhance productivity at both the individual and organisational level, its effects extend beyond firm performance to the economy at large. Productivity remains a central pillar of economic growth and societal wellbeing, and AI may play a vital part in advancing both.⁴ This chapter moves beyond the prevailing hype to provide an evidence-based assessment of how AI can help address the EU's productivity challenges. Specifically, it explores the ways in which AI is redefining human roles in the workforce, examines the pace and scope of AI adoption across EU countries and occupations, and evaluates AI's cumulative impact on workforce and aggregate productivity. In doing so, the chapter argues that the diffusion of AI technologies, supported by inclusive and forward-looking policies, can sustain the EU's competitiveness and long-term prosperity.

REDEFINING HUMAN ROLES IN THE AI-DRIVEN WORKFORCE

AI is fundamentally reconfiguring the modern workforce by changing how tasks are performed, how skills are applied, and how humans engage with technology.⁵ However, the interaction between humans and AI remains complex and the results uneven. In many settings, collaboration between people and intelligent systems enhances efficiency and individual performance but does not necessarily surpass the best results of either working alone. The effectiveness of the partnership depends upon the nature of the task and the abilities each participant can contribute. Creative and open-ended work tends to benefit most from the complementary capabilities of humans and AI, while decision-based tasks often face difficulties related to trust, coordination, and the assignment of responsibility. Achieving real synergy and sustained

AI is fundamentally reconfiguring the modern workforce by changing how tasks are performed, how skills are applied, and how humans engage with technology.

1. Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence and amending Regulations (EC) No. 300/2008, (EU) No. 167/2013, (EU) No. 168/2013, (EU) 2018/858, (EU) 2018/1139 and (EU) 2019/2144 and Directives 2014/90/EU, (EU) 2016/797 and (EU) 2020/1828 (Artificial Intelligence Act), *Official Journal of the European Union*, L 2024/1689, 1–144, <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>.
2. I. González Vázquez, E. Fernández Macías, S. Wright, and D. Villani (2025), *Digital monitoring, algorithmic management and the platformisation of work in Europe* (JRC143072) (Luxembourg: Publications Office of the European Union), <https://data.europa.eu/doi/10.2760/9406086>.
3. M. Chui, R. Roberts, L. Yee, E. Hazan, A. Singla, K. Smaje, A. Sukharevsky, and R. Zemmel (2023), *The economic potential of generative AI: The next productivity frontier* (New York: McKinsey & Company), <https://www.mckinsey.com/capabilities/mckinsey-digital/our-insights/the-economic-potential-of-generative-ai-the-next-productivity-frontier>.
4. F. Filippucci, P. Gal, C. Jona-Lasinio, A. Leandro, and G. Nicoletti (2024), *The impact of artificial intelligence on productivity, distribution and growth: Key mechanisms, initial evidence and policy challenges* (OECD Artificial Intelligence Papers No. 15) (Paris: OECD Publishing), <https://doi.org/10.1787/8d900037-en>.
5. Deloitte AI Institute (2023), *Generative AI and the future of work: Preparing your organization for the boundless potential of AI in the workplace and its impact on jobs* (New York: Deloitte Development LLC), <https://www.deloitte.com/us/en/what-we-do/capabilities/applied-artificial-intelligence/articles/generative-ai-and-the-future-of-work.html>.

productivity gains requires careful task design so that humans and AI can play to their respective strengths. Systems that foster calibrated trust and measure success across multiple performance indicators (such as accuracy, quality, efficiency, and ethical soundness) are likewise vital.⁶

The forms of collaboration between humans and AI systems can be sorted into four main categories. Table 1 illustrates the ways that AI can redistribute responsibilities across the workforce, with outcomes ranging from full automation to the emergence of new forms of human expertise.

TABLE 1: AI impact on human skills and tasks

Category	Description	Examples	Representative Skills / Tasks
Automated tasks: Machines do best	Tasks that are fully executed by AI systems with minimal or no human involvement. These rely on predefined algorithms and automation to ensure speed, accuracy, and consistency.	AI can autonomously generate standardised or repetitive outputs such as customer service responses, reports, or summaries. It can also tailor content to individual users through behavioural and data-driven personalisation.	Machine-only tasks: image generation, text generation, data sorting and categorisation, routine forecasting, language translation, simple graphic design, pattern or trend detection.
Augmented skills: Humans with machines do best	Human performance is enhanced by AI tools that increase efficiency, analytical capacity, and creativity. Collaboration between human intuition and AI computation yields superior outcomes.	Artists can draw inspiration from AI-generated concepts, while researchers and analysts use AI to process complex datasets and make informed, strategic decisions.	Human–AI collaboration: creativity, analytical reasoning, complex problem solving, research and innovation, data visualisation, strategic planning, predictive analytics, rapid prototyping.
New skills: Humans needed	As AI transforms work environments, new competencies are essential for effective collaboration, oversight, and ethical integration of AI systems. Lifelong learning ensures adaptability in AI-driven contexts.	Professionals increasingly manage AI workflows, ensure ethical compliance, and design AI applications aligned with societal and organisational values.	Emerging human competencies: AI ethics and governance, human–AI task coordination, AI system supervision, AI output evaluation and customisation.
Limited-impact tasks: Humans do best	Activities that rely on empathy, moral reasoning, and nuanced judgment remain uniquely human. These require emotional intelligence and contextual understanding that AI cannot replicate.	Leadership, counselling, and ethical decision-making exemplify domains where human presence and interpersonal sensitivity are irreplaceable.	Human-exclusive skills: persuasion and negotiation, motivational leadership, ethical reasoning and integrity, compassion and empathy, relationship building, physical dexterity and care tasks.

Source: Deloitte AI Institute (2023).

In each of these four categories, AI has different implications for workforce productivity. AI-assisted task automation can boost productivity by streamlining repetitive and data-driven processes while improving efficiency, accuracy, and consistency, allowing human workers to focus on higher-val-

6. M. Vaccaro, A. Almaatouq, and T. Malone (2024), 'When Combinations of Humans And AI Are Useful: A Systematic Review and Meta-Analysis', *Nature Human Behaviour*, 8(12), 2293–2303. <https://doi.org/10.1038/s41562-024-02024-1>.

ue activities that require judgment and creativity.⁷ Augmented skills expand human capability – in these applications, AI supports creativity, analysis, and decision making, and productivity gains are reflected by the quality of work, not the speed of its execution.⁸ The development of new skills in the course of human–AI collaboration has made it clear that long-term improvements to productivity require ongoing human learning and adaptability. Organisations that invest in digital literacy, AI ethics, and responsible governance can integrate technologies more effectively and secure sustained performance growth.⁹ Finally, human leadership and ethical reasoning remain crucial for fruitful collaboration in limited-impact tasks – distinctly human competencies continue to underpin meaningful and enduring productivity in an AI-driven workforce.¹⁰

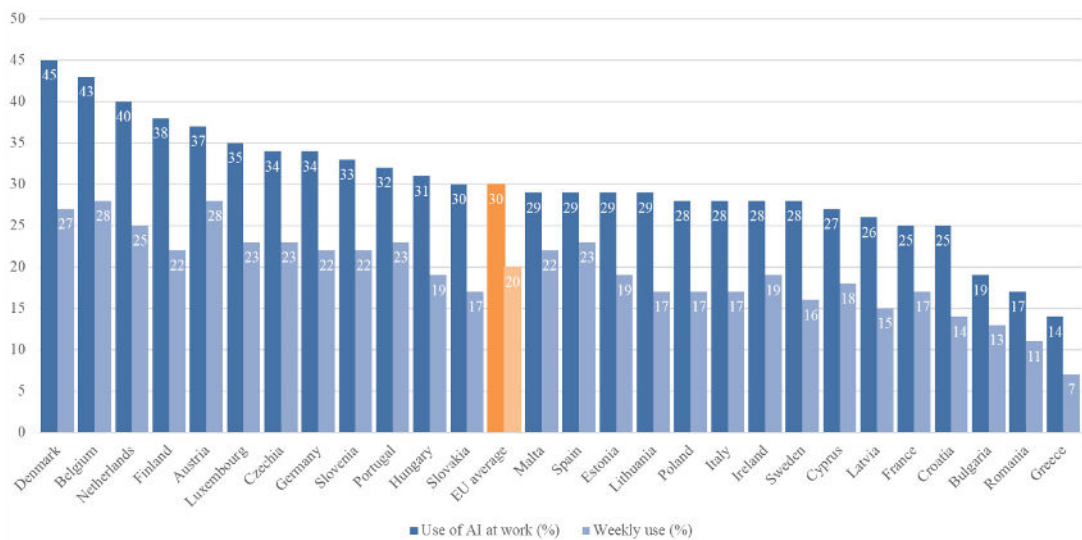
AI ADOPTION IN THE WORKFORCE

The reframing of human roles in the AI era is not a future hypothetical – it is already under way, and Europe is feeling its impact. According to the latest evidence, approximately one third of the EU's labour force has used an AI-powered tool at work at least once in the past 12 months (Figure 1). Considering how recently such tools have become widely available, this is a remarkable figure. The rapid uptake may be explained by the emergence of AI chatbots; generally accessible to the public from late 2022, these swiftly gained popularity worldwide. Some EU countries (including Denmark, Belgium, the Netherlands, Finland, and Austria) report rates of AI use in the workplace that reach or exceed 40%. At the other end of the spectrum are Bulgaria, Romania, and Greece, where fewer than 20% of respondents use AI at work. Still, the variation among EU countries is fairly modest, and Europe's frequency of AI use (compared with the rest of the world) is relatively high.

Looking at this phenomenon in more detail, around 20% of workers in the EU report using AI at least once a week in their main job (once again, with some minor variation across countries; the nations with the lowest proportions of weekly users tend to report the lowest overall prevalence of AI use).¹¹

-
7. B. Y. Kassa and E. K. Worku (2025), 'The Impact of Artificial Intelligence on Organizational Performance: The Mediating Role of Employee Productivity', *Journal of Open Innovation: Technology, Market, and Complexity*, 100474. <https://doi.org/10.1016/j.joitmc.2025.100474>.
 8. X. Sun and Y. Song (2025), 'Unlocking the Synergy: Increasing Productivity through Human-AI Collaboration in the Industry 5.0 Era', *Computers & Industrial Engineering*, 200, 110657. <https://doi.org/10.1016/j.cie.2024.110657>.
 9. C. Y. Ersanlı, F. Çelik, H. Barjesteh, V. Duran, and M. Manoochehrzadeh (2025), 'A Review of Global Reskilling and Upskilling Initiatives in the Age of AI', *AI and Ethics*, 1–10, <https://doi.org/10.1007/s43681-025-00767-9>.
 10. D. A. Spencer (2025), 'AI, Automation and the Lightning of Work', *AI & Society*, 40(3), 1237–1247. <https://doi.org/10.1007/s00146-024-01959-3>.
 11. González Vázquez et al., *Digital monitoring*.

FIGURE 1: Use of AI tools at work in the EU by country



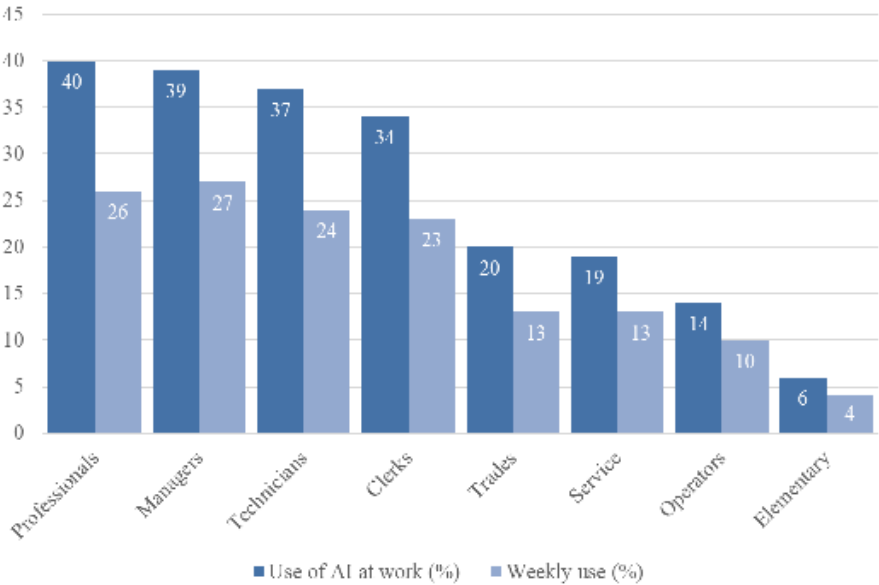
Source: González Vázquez et al. (2025).

More significant differences in AI usage can be seen across occupations (Figure 2). Over one third of white-collar workers (e.g., professionals, managers, technicians, and clerks) use AI tools to perform their duties. In other occupations, AI usage drops to less than one fifth of workers, and to only 6% among those in elementary occupations, with a similar pattern apparent for weekly use.¹² This occupational divide reflects the disparate ways AI affects labour demand and skills. AI tends to automate the routine and repetitive tasks common in lower-skilled or administrative roles while augmenting the complex analytical and creative activities that dominate higher-skilled occupations. As a result, professionals and managers will have more to gain from AI, whereas workers in routine or manual roles face limited opportunities for AI integration and greater risks of task simplification. These structural variations contribute to a widening polarisation in the labour market, with AI expanding opportunities for high-skill workers while reducing demand in lower-skill segments.¹³

12. González Vázquez et al., *Digital monitoring*.

13. W. X. Chen, S. Srinivasan, and S. Zakerinia (2025), *Displacement or Complementarity?: The Labor Market Impact of Generative AI* (Cambridge, MA: Harvard Business School).

FIGURE 2: Use of AI tools at work in the EU by occupation



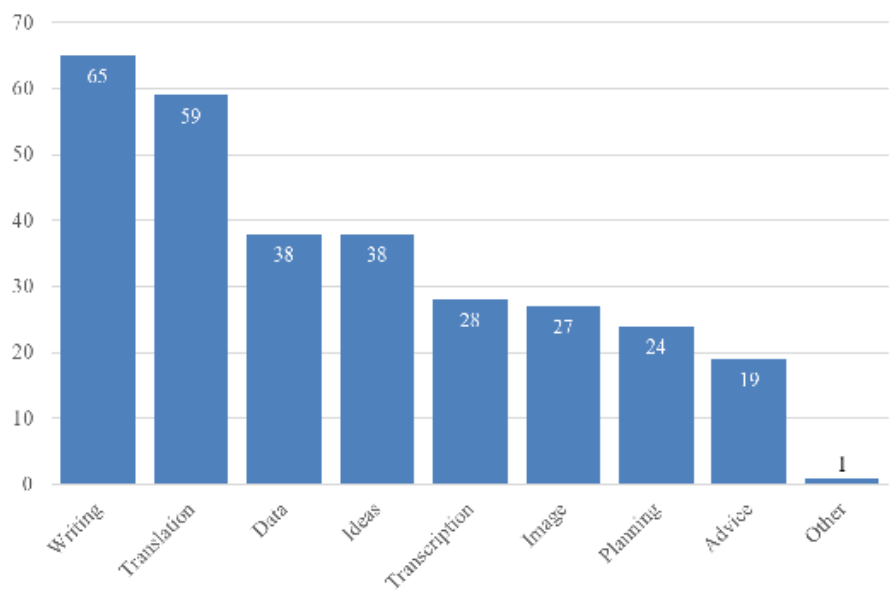
Source: González Vázquez et al. (2025).

IMPACTS ON WORKFORCE PRODUCTIVITY

Across Europe, among workers who use AI, the most common workplace application of it is for writing (65%), followed by translation (59%). Other uses include data processing and idea generation (each 38%), transcription (28%), image generation (27%), planning and scheduling (24%), and customer advice (19%), while only a small share of respondents reported using AI for other purposes (1%) (Figure 3).¹⁴ The fact that AI is most frequently involved in writing and translation tasks suggests that such technology is primarily employed to automate routine cognitive activities, enabling employees to create textual content and reports more efficiently and accurately. Meanwhile, the increasing use of AI for data processing and idea generation signals a shift toward cognitive augmentation, where algorithms assist workers in analysing information and developing new solutions. Such trends show that productivity gains arise not only from faster task execution but also from enhanced creativity, decision quality, and problem-solving capacity.¹⁵ The most significant improvements occur when automation is paired with complementary human competencies such as analytical reasoning and adaptability. In this sense, AI functions as a multiplier of human potential, increasing both the speed and intellectual value of individual output.¹⁶

14. González Vázquez et al., *Digital monitoring*.
15. Cedefop (2025), *Skills empower workers in the AI revolution: First findings from Cedefop’s AI skills survey* (Luxembourg: Publications Office of the European Union), <https://www.cedefop.europa.eu/en/publications/9201>.
16. McKinsey Global Institute (2024), *A new future of work: The race to deploy AI and raise skills in Europe and beyond* (New York: McKinsey Global Institute), <https://www.mckinsey.com/mgi/our-research/a-new-future-of-work-the-race-to-deploy-ai-and-raise-skills-in-europe-and-beyond>.

FIGURE 3: Purpose of AI use in the workplace in the EU



Source: González Vázquez et al. (2025).

Taking a broader view, the cumulative effect of AI-driven task-level efficiencies is becoming evident in measurable improvements to organisational and national productivity across the EU. Recent data suggests that even under moderate AI adoption scenarios, aggregate total factor productivity in the EU could rise by around 1% over five years, with larger gains in countries characterised by higher income levels, stronger digital infrastructure, and greater innovation capacity.¹⁷ Likewise, findings indicate that AI may add between 0.25 and 0.6 percentage points annually to total factor productivity, or 0.4 to 0.9 percentage points to labour productivity, across the EU and other advanced economies over the next decade.¹⁸ These projections suggest that as EU firms embed AI in a wider range of tasks (including writing, translation, and data processing), they may achieve considerable productivity gains. However, progress is not uniform: in areas with robust digital capabilities, advancement proceeds faster than in domains hampered by structural or regulatory barriers. Ultimately, productivity growth in the EU depends upon effective AI diffusion (and its alignment with human skills); stable digital infrastructures; adaptive organisational strategies; and sound governance. Only then can AI fully realise its promise as a powerful multiplier that accelerates output and fortifies long-term economic resilience.

17. IMF (2025), *AI and productivity in Europe* (IMF Working Paper No. 25/64), <https://www.imf.org/en/Publications/WP/Issues/2025/04/04/AI-and-Productivity-in-Europe-565924>.

18. OECD (2025), *Miracle or myth? Assessing the macroeconomic productivity gains from artificial intelligence* (Paris: OECD Publishing), <https://oecd.ai/en/ai-publications/miracle-or-myth-assessing-the-macroeconomic-productivity-gains-from-artificial-intelligence>.

AI AND PRODUCTIVITY IN THE EUROPEAN WORKFORCE

AI is reshaping productivity dynamics across the EU by transforming how work is organised, tasks are performed, and value is created. Its growing presence in the labour force is marked by both progress and imbalance: overall levels of AI adoption are rising quickly, but at uneven rates across sectors, occupations, and regions. Moving forward, the key to harnessing AI's full productivity potential lies in ensuring that its spread across the workforce is equitable and inclusive. Increasing adoption must therefore go hand in hand with policies that promote digital readiness, human adaptability, and responsible innovation.

POLICY RECOMMENDATIONS

Coordinated action at EU and national levels should focus on three interrelated priorities:

1. Policies should promote the diffusion of AI technologies by investing in digital infrastructure, open data frameworks, and innovation networks that facilitate technological transfer and organisational learning. Such initiatives can help scale successful use cases and extend adoption beyond digitally advanced environments.
2. Education and training systems must prioritise digital literacy, analytical thinking, and socio-emotional skills to promote effective collaboration between humans and AI systems. Continuous upskilling and reskilling, alongside a strong understanding of AI ethics and governance, will ensure that productivity gains are sustainable and responsible.
3. Thoughtful, values-informed integration must remain central to AI strategies. The EU AI Act offers a solid foundation for trustworthy adoption, but its impact depends on transparent implementation, clear accountability, and the protection of fair working conditions.

CONCLUSION

AI-driven productivity gains ultimately depend on the harmonisation of innovation, skills, and governance. By supporting widespread adoption and equitable diffusion across the workforce, the EU can convert technological progress into long-term economic resilience and shared prosperity. It is not through proliferation alone that AI systems can best serve society: they must be strategically, inclusively, and ethically integrated into evolving work environments.

REFERENCES

- Cedefop. (2025). 'Skills empower workers in the AI revolution: First findings from Cedefop's AI skills survey'. Brussels: Publications Office of the European Union.
<https://www.cedefop.europa.eu/en/publications/9201>.
- Chen, W. X., Srinivasan, S., & Zakerinia, S. (2025). *Displacement Or Complementarity?: The Labor Market Impact of Generative AI*. Cambridge, MA: Harvard Business School.
- Chui, M., Roberts, R., Yee, L., Hazan, E., Singla, A., Smaje, K., Sukharevsky, A., & Zimmel, R. (2023). *The economic potential of generative AI: The next productivity frontier*. New York: McKinsey & Company.
<https://www.mckinsey.com/capabilities/mckinsey-digital/our-insights/the-economic-potential-of-generative-ai-the-next-productivity-frontier>.
- Deloitte AI Institute. (2023). *Generative AI and the future of work: Preparing your organization for the boundless potential of AI in the workplace and its impact on jobs*. Deloitte Development LLC.
<https://www.deloitte.com/us/en/what-we-do/capabilities/applied-artificial-intelligence/articles/generative-ai-and-the-future-of-work.html>.
- Ersanlı, C. Y., Çelik, F., Barjesteh, H., Duran, V., & Manoochehrzadeh, M. (2025). 'A Review of Global Reskilling and Upskilling Initiatives in the Age of AI'. *AI and Ethics*, 1–10. <https://doi.org/10.1007/s43681-025-00767-9>.
- Filippucci, F., Gal, P., Jona-Lasinio, C., Leandro, A., & Nicoletti, G. (2024). *The impact of artificial intelligence on productivity, distribution and growth: Key mechanisms, initial evidence and policy challenges* (OECD Artificial Intelligence Papers No. 15). Paris: Organisation for Economic Co-operation and Development.
<https://doi.org/10.1787/8d900037-en>.
- González Vázquez, I., Fernández Macías, E., Wright, S., & Villani, D. (2025). *Digital monitoring, algorithmic management and the platformisation of work in Europe* (JRC143072). Luxembourg: Publications Office of the European Union.
<https://data.europa.eu/doi/10.2760/9406086>.
- IMF. (2025). *AI and productivity in Europe*. IMF Working Paper No. 25/64.
<https://www.imf.org/en/Publications/WP/Issues/2025/04/04/AI-and-Productivity-in-Europe-565924>.
- Kassa, B. Y., & Worku, E. K. (2025). 'The Impact of Artificial Intelligence on Organizational Performance: The Mediating Role of Employee Productivity'. *Journal of Open Innovation: Technology, Market, and Complexity*, 100474.
<https://doi.org/10.1016/j.joitmc.2025.100474>.
- McKinsey Global Institute. (2024). 'A new future of work: The race to deploy AI and raise skills in Europe and beyond'. McKinsey Global Institute.
<https://www.mckinsey.com/mgi/our-research/a-new-future-of-work-the-race-to-deploy-ai-and-raise-skills-in-europe-and-beyond>.
- OECD. (2025). *Miracle or myth? Assessing the macroeconomic productivity gains from artificial intelligence*. Paris: Organisation for Economic Co-operation and Development.
<https://oecd.ai/en/ai-publications/miracle-or-myth-assessing-the-macroeconomic-productivity-gains-from-artificial-intelligence>.
- Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence and amending Regulations (EC) No 300/2008, (EU) No 167/2013, (EU) No 168/2013, (EU) 2018/858, (EU) 2018/1139 and (EU) 2019/2144 and Directives 2014/90/EU, (EU) 2016/797 and (EU) 2020/1828 (Artificial Intelligence Act), *Official Journal of the European Union*, L 2024/1689, 1–144.
<https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>.
- Spencer, D. A. (2025). 'AI, Automation and the Lightening of Work'. *AI & Society*, 40 (3), 1237–1247.
<https://doi.org/10.1007/s00146-024-01959-3>.
- Sun, X., & Song, Y. (2025). 'Unlocking the Synergy: Increasing Productivity through Human-AI Collaboration in the Industry 5.0 Era'. *Computers & Industrial Engineering*, 200, 110657.
<https://doi.org/10.1016/j.cie.2024.110657>.
- Vaccaro, M., Almaatouq, A., & Malone, T. (2024). 'When Combinations of Humans and AI are Useful: A Systematic Review And Meta-Analysis'. *Nature Human Behaviour*, 8 (12), 2293–2303.
<https://doi.org/10.1038/s41562-024-02024-1>.

AI and Energy Use: How to Wield a Double-Edged Sword in the Direction of Progress

Eloi Borgne

<https://doi.org/10.53121/ELFS10> • ISSN (print) 2791-3880 • ISSN (online) 2791-3899

ABSTRACT

Data centres currently account for 2–4% of electricity consumption in large economies like the EU. By 2030, it is projected that AI will lead to a doubling of data centre energy consumption. In the same period, AI adoption in Europe could increase total electricity demand across the economy by 4–5%, posing challenges for grid capacity and climate goals. While AI increases energy demand, it also offers transformative potential to optimise energy systems, reduce emissions, and clear the way for breakthrough energy technologies. Europe's AI strategy must therefore balance rapid innovation with sustainability, ensuring AI's energy footprint does not undermine climate neutrality targets.

ABOUT THE AUTHOR

Eloi Borgne is Junior Policy and Research Officer at European Liberal Forum.

INTRODUCTION

Artificial intelligence (AI) is rapidly transforming economies and societies worldwide, with Europe attempting to position itself as a leader in an evolving technological landscape. However, the exponential growth in AI adoption brings with it a major challenge: increased energy consumption, primarily driven by the data centres that power AI systems. This chapter explores the relationship between AI and energy use in Europe, focusing on the significant footprint of AI's energy consumption and the pioneering solutions that could address this issue. It provides an analysis of the current situation, future projections, and potential benefits of AI in energy management. It also examines how Europe can leverage AI to enhance energy security, accelerate innovation, and drive economic growth while mitigating the environmental impact.

THE ENERGY CONSUMPTION OF AI AND DATA CENTRES: CHALLENGES AND PROJECTIONS

The computational muscle behind AI lies in data centres. They host the servers, storage, and networking equipment necessary for developing AI models. Today, data centres account for approximately 2–4% of total electricity consumption in large economies like the European Union.¹ This share is expected to grow significantly as AI adoption gains momentum, with global data centre electricity consumption projected to more than double from 415 terawatt-hours (TWh) in 2024 to around 945 TWh by 2030. In Europe specifically, the electricity required to power data centres could rise by up to 28% by 2030 to reach 98.5 TWh, which would represent 4–5% of total electricity demand (up from 2–3% in 2024).²

This surge in energy demand is driven by the deployment of high-performance accelerated servers specifically constructed for AI workloads, which have greater power density and energy requirements than their conventional counterparts.³ The IEA's Base Case Scenario forecasts a 30% annual growth in electricity consumption from accelerated servers, corresponding to nearly half of the net increase in global data centre electricity consumption.⁴

The rapid growth in AI-related energy demand poses several challenges:

- *Grid capacity and infrastructure:* The pace of grid infrastructure development lags behind that of data centre construction, leading to potential bottlenecks and grid congestion. Some jurisdictions have already imposed moratoriums on new data centre connections due to overwhelming demand. As Laura Cozzi (Chief Energy Modeller at the IEA) notes, this challenge is exacerbated because 'data centres are very concentrated loads. They tend to cluster together ... and it is important and different from other large loads – it tends to be very close to cities, and this means that it's going to be linked to grids that are most likely already congested'.⁵ This concentration means that while the global numbers may seem modest, the localised impact on specific grids is profound and immediate.
- *Environmental impact:* Increased electricity consumption from data centres can lead to higher levels of greenhouse gas emission, with indirect emissions from electricity generation a significant contributor. If this trend is not addressed, Europe's ability to reach its climate targets could be undermined. Additionally, data centres require a huge amount of water for cooling, increasing the potential environmental damage of AI.
- *Energy security:* The concentration of data centre demand in specific regions and the reliance on critical minerals for data centre construction introduce vulnerabilities to supply chain disruptions and geopolitical risks. Cozzi frames the core challenge succinctly: 'Is AI alone a big issue? It is because it is coming as a general purpose technology ... but I think it's important for the energy sector, the fact that this is coming in a situation that was already where we were having a lot of underinvestment, in particular in grids, and tension in parts of the supply chains'.⁶ AI's demand thus further strains a system that is already struggling to cope with considerable pre-existing pressures.

1. T. Spencer and S. Singh (2024), 'What the data centre and AI boom could mean for the energy sector – Analysis', IEA, 18 October, <https://www.iea.org/commentaries/what-the-data-centre-and-ai-boom-could-mean-for-the-energy-sector>.

2. EU Directorate-General for Energy (2023), 'Commission takes first step towards establishing an EU-wide scheme for rating sustainability of data centres', European Commission, 12 December, https://energy.ec.europa.eu/news/commission-takes-first-step-towards-establishing-eu-wide-scheme-rating-sustainability-data-centres-2023-12-12_en.

3. These specialised servers are equipped with GPUs or similar accelerator chips to enhance computing performance for specific tasks.

4. IEA (2025), *Energy and AI* (Paris: IEA), <https://www.iea.org/reports/energy-and-ai>.

5. Jason Bordoff (host) (2025), 'Is AI Friend or Foe to the Clean Energy Transition?' [audiopodcast], in *Columbia Energy Exchange* (Columbia University), 8 July, <https://www.energypolicy.columbia.edu/is-ai-friend-or-foe-to-the-clean-energy-transition/>.

6. Bordoff, 'Is AI Friend or Foe'.

ENVIRONMENTAL IMPACT AND SUSTAINABILITY CHALLENGES

The environmental impact of data centres extends beyond energy consumption. The construction and operation of data centres involve substantial resource use, including water for cooling and rare earth metals for batteries and electronics, contributing to pollution and electronic waste. Diesel generators, commonly used for backup power, can worsen local air quality. The carbon footprint of data centres is in large part contingent upon the energy sources fuelling the grids they draw on, with coal still predominating in some regions.

Europe has made significant strides in addressing these challenges through initiatives such as the Climate Neutral Data Centre Pact. The goal of this pact is for data centres to become climate-neutral by 2030 by improving power usage effectiveness and adopting carbon-free energy sources. To reduce their carbon footprint and mitigate energy price volatility, data centre operators are increasingly investing in projects to promote renewable energy, including solar and wind, through power purchase agreements (PPAs).⁷

Aligning AI adoption with European climate targets requires a combination of technological innovation, strategic planning, and regulatory frameworks.

Further gains in sustainability have been sought through the adoption of more efficient cooling technologies, such as liquid immersion cooling, and the reuse of waste heat in nearby facilities or district heating networks. In 2024, the European Commission established an EU-wide scheme to rate the sustainability of data centres, promoting transparency and best practices in energy efficiency.⁸

Despite these efforts, the precipitous rise in AI-driven energy will demand a more far-reaching and coordinated response. Aligning AI adoption with European climate targets requires a combination of technological innovation, strategic planning, and regulatory frameworks to ensure that AI's energy footprint is minimised.

THE POSITIVE ASPECTS OF AI IN ENERGY MANAGEMENT AND INNOVATION

While AI increases energy demand, it also offers transformative potential to optimise energy systems, reduce emissions, and accelerate innovation across the sector. AI applications in energy management include:

- *Energy efficiency in buildings:* AI can model building energy use and optimise heating, ventilation, and air conditioning (HVAC) systems, reducing energy consumption by at least 8% and lowering carbon emissions significantly.
- *Industrial process optimisation:* AI-driven predictive maintenance augments the efficiency of manufacturing processes, lowering energy use and costs while elevating productivity.
- *Transportation:* AI can improve vehicle efficiency and reduce energy consumption by up to 20% through optimised routing, predictive maintenance, and autonomous driving technologies. The IEA's analysis suggests the potential is even greater – AI optimisation in freight and trucking could offer energy savings 'equivalent to taking away from streets a hundred million cars'.⁹

7. Climate Neutral Data Center (n.d.), *Climate Neutral Data Centre Pact*, <https://www.climateneutraldatacentre.net/>.

8. Directorate-General for Energy (2024), 'Commission adopts EU-wide scheme for rating sustainability of data centres', European Commission, 15 March, https://energy.ec.europa.eu/news/commission-adopts-eu-wide-scheme-rating-sustainability-data-centres-2024-03-15_en Energy.

9. Bordoff, 'Is AI Friend or Foe'.

- *Grid management and renewable integration:* AI enhances grid stability and resilience by anticipating energy demand, managing intermittent renewable sources, and optimising grid operations. AI can help to balance electricity networks as they grow more complex, decentralised, and digitalised. It can likewise facilitate the forecasting and integration of variable renewable energy generation, reducing curtailment and emissions. AI-based detection tools allow swift identification and precise pinpointing of grid faults, reducing outage durations by 30–50%. This supports Europe’s climate goals by enabling higher penetrations of renewable energy. Crucially, AI can also help alleviate the very grid congestion it has so far contributed to. Cozzi points to ‘dynamic line rating’ as a key use case, wherein AI enables existing grid infrastructure ‘to carry more electrons’, with a potential global impact ‘equivalent to what you’ve seen on average built in one year over the past five years of new grid’.¹⁰
- *Energy innovation:* AI has proven value in the development of new energy technologies, including batteries, catalysts for synthetic fuel production, and materials for carbon capture, which are critical for decarbonisation and energy security. There is likewise tremendous untapped potential for AI in research and development (for instance, to streamline the chemical modelling of materials like perovskites, exploited in next-generation solar panels).
- *Cybersecurity:* AI strengthens energy sector cybersecurity by enabling real-time threat detection and automated responses, which are essential as digitalisation increases vulnerability to cyberattacks.¹¹

POTENTIAL SOLUTIONS TO MITIGATE AI’S ENERGY FOOTPRINT

An alarming passage in the recently released US AI Action Plan indicates that ‘the NIST AI Risk Management Framework [is] to eliminate references to ... climate change’.¹² To do so would be a grave error. The increases in energy consumption associated with unchecked AI proliferation should not be ignored, but viewed as an opportunity to develop efficient solutions. Addressing these challenges requires a multifaceted approach with measures that balance innovation with sustainability:

The increases in energy consumption associated with unchecked AI proliferation should not be ignored, but viewed as an opportunity to develop efficient solutions.

- *Improving data centre efficiency:* Investments in energy-efficient hardware, cooling systems, and data centre infrastructure management (DCIM) tools can significantly reduce energy consumption. The adoption of new chip designs, such as NVIDIA’s AI superchips (which offer thirtyfold performance improvements with 25 times less energy use), is a case in point. A more out of the box approach is already under way in China, using experiments with seawater cooling. The scale of demand makes the shift to seawater significant: a 1-megawatt facility, for example, can require up to 26 million litres annually, and total cooling demand across the region could reach 1.7 trillion litres by 2030. Facilities such as Exchange Square in Hong Kong, home to major financial institutions and consulates, al

10. Bordoff, ‘Is AI Friend or Foe’.

11. IEA, *Energy and AI*.

12. The White House (2025), ‘White House Unveils America’s AI Action Plan’, 23 July, <https://www.whitehouse.gov/articles/2025/07/white-house-unveils-americas-ai-action-plan/>.

ready use seawater systems, reflecting a wider trend across Asia toward more sustainable cooling solutions for increasingly resource-intensive data centres.¹³

- *Integrating renewable energy:* Data centres can boost their use of renewable energy through on-site generation (solar panels or wind turbines) and power purchase agreements. The European Data Centre Association aims for 100% renewable energy use by 2030, with current levels at 94%.¹⁴
- *Flexible operations and strategic location planning:* Data centres can leverage such on-site energy resources and backup systems to bolster grid resiliency. Strategic location of data centres in areas with abundant renewable energy and capacity can reduce strain on local grids. While discussions about data centre electricity consumption frequently revolve around terawatt-hours (TWh), it is crucial to note that the electricity system is designed to meet instantaneous peak demand. Over the coming years, data centres are expected to transition from a minimal to a significant portion of instantaneous demand in various regions. There are, however, some drawbacks: high capital intensity in data centres leads to heavy opportunity costs when workloads are curtailed. Additionally, sector fragmentation introduces operational and contractual difficulties. Nevertheless, these issues are manageable. Promising paths towards their solution include recent pilot projects by NVIDIA, EmeraldAI, and their partners, as well as cutting-edge work by Google. To advance these efforts, it is crucial to encourage dialogue between all relevant stakeholders. Only through such collaboration can data centre construction (requiring 1–2 years) and grid expansion (taking 4–8 years in advanced economies) be brought into a healthy alignment, where energy demand does not outstrip capacity.
- *Microgrid and decentralised energy models:* Implementing microgrid models in which data centres operate semi-independently from the grid can enhance resilience and sustainability, enabling the integration of local renewable energy sources and storage technologies.
- *Policy and regulatory frameworks:* The European Union's Climate Neutral Data Centre Pact and the AI Continent Action Plan provide frameworks to promote sustainable AI adoption. The EU AI Act, the world's first comprehensive AI law, ensures safety, fundamental rights, and ethical AI use while fostering innovation.

13. M. Yunus (2025), 'AI-driven data centre boom could drain Asia's rivers without due care', *South China Morning Post*, 16 August, https://www.scmp.com/opinion/asia-opinion/article/3321711/ai-driven-data-centre-boom-could-drain-asias-rivers-without-due-care?module=perpetual_scroll_0&pgtype=article.

14. European Data Centre Association (2023), *Energy Efficiency Directive*, <https://www.eudca.org/energy-efficiency-directive>.

SUMMARY TABLE: Key Data on AI and Energy Use in Europe

Aspect	Impact	Notes
Current data centre electricity use	2–4% of total EU electricity consumption	Accounts for ~96 TWh in Europe (2024)
Projected increase by 2030	140% increase in EU data centre energy use	Reaching 150+ TWh, 4–5% of total electricity demand ¹⁵
Global data centre electricity use	415 TWh (2024) → 945 TWh (2030)	AI-driven demand growing at 15% annually
Environmental impact	~1% of global energy-related GHG emissions	Includes emissions from electricity generation, cooling, and e-waste
Renewable energy integration	94% of European data centres' energy use	Targeting 100% by 2030 via Climate Neutral Data Centre Pact
AI's potential energy savings	8% reduction in building energy use	Through optimisation of HVAC and industrial processes
AI's role in grid management	Enhanced grid stability and renewable integration	Predictive maintenance, demand forecasting, and cybersecurity

CONCLUSION

The rapid growth of AI in Europe is a double-edged sword in terms of energy use. On the one hand, AI's escalating energy demand, driven by data centres, poses major challenges for grid capacity, environmental sustainability, and energy security. The difficulty is magnified by years of underinvestment in grid infrastructure. ¹⁵ On the other hand, AI can deliver powerful tools for addressing longstanding problems.

To overcome the downsides of AI's energy use Europe must adopt a multifaceted strategy that includes improving data centre efficiency, integrating renewable energy, using water more efficiently, and harnessing AI for grid optimisation. The European Union's leadership in AI regulation and sustainability initiatives is a firm foundation on which to build a brighter future in the AI era – in which societal good and economic growth go hand in hand.

15. Bordoff, 'Is AI Friend or Foe'. <https://www.energypolicy.columbia.edu/is-ai-friend-or-foe-to-the-clean-energy-transition/>

By fostering innovation and collaboration across the public and private sectors, Europe can unleash AI's potential to serve humanity while ensuring that its energy footprint supports, rather than undermines, the continent's climate neutrality and energy security goals. This balanced approach will be critical for Europe to maintain its competitive edge and guide the global transition to a sustainable, AI-powered tomorrow.

REFERENCES

- A. Granskog, D. Hernandez Diaz, J. Noffsinger, L. Moavero Milanesi, P. Sachdeva, A. Bhan, & S. von Schantz (2024). 'The role of power in unlocking the European AI revolution', McKinsey & Company, 24 October, <https://www.mckinsey.com/industries/electric-power-and-natural-gas/our-insights/the-role-of-power-in-unlocking-the-european-ai-revolution>.
- Bordoff, Jason (Host) (2025). Is AI Friend or Foe to the Clean Energy Transition? [Audiopodcast]. In *Columbia Energy Exchange*. Columbia University, 8 July. <https://www.energypolicy.columbia.edu/is-ai-friend-or-foe-to-the-clean-energy-transition/>.
- Climate Neutral Data Center (n.d.). *Climate Neutral Data Centre Pact*. <https://www.climateneutraldatacentre.net/>.
- Directorate-General for Energy (2023). 'Commission takes first step towards establishing an EU-wide scheme for rating sustainability of data centres', European Commission, 12 December, https://energy.ec.europa.eu/news/commission-takes-first-step-towards-establishing-eu-wide-scheme-rating-sustainability-data-centres-2023-12-12_en
- Directorate-General for Energy (2024). 'Commission adopts EU-wide scheme for rating sustainability of data centres', European Commission, 15 March, https://energy.ec.europa.eu/news/commission-adopts-eu-wide-scheme-rating-sustainability-data-centres-2024-03-15_en Energy.
- European Data Centre Association (2023). *Energy Efficiency Directive*. <https://www.eudca.org/energy-efficiency-directive>.
- European Union (2023). *Energy Efficiency Directive*. https://energy.ec.europa.eu/topics/energy-efficiency/energy-efficiency-targets-directive-and-rules/energy-efficiency-directive_en.
- International Energy Agency (2025). *Energy and AI – Analysis*. IEA, Paris. <https://www.iea.org/reports/energy-and-ai>.
- Manyika, J. & Woetzel, J. (2024). 'Unlocking the European AI revolution', McKinsey & Company, 24 October, <https://www.mckinsey.com/industries/electric-power-and-natural-gas/our-insights/the-role-of-power-in-unlocking-the-european-ai-revolution> McKinsey & Company.
- Spencer, T. & Singh, S. (2024). 'What the data centre and AI boom could mean for the energy sector – Analysis', IEA, 18 October, <https://www.iea.org/commentaries/what-the-data-centre-and-ai-boom-could-mean-for-the-energy-sector>.
- The White House (2025). 'White House unveils America's AI Action Plan', 23 July, <https://www.whitehouse.gov/articles/2025/07/white-house-unveils-americas-ai-action-plan/>.
- Yunus, M. (2025). 'AI-driven data centre boom could drain Asia's rivers without due care', *South China Morning Post*, 16 August, <https://www.scmp.com/opinion/asia-opinion/article/3321711/ai-driven-data-centre-boom-could-drain-asias-rivers-without-due-care>.

Creating Opportunities: The Importance of Learning Design for AI Use in Education

Sarah K. Howard, Jo Tondeur, Charlotte Neubert,
Jonas Hauck, and Richard Böhme

<https://doi.org/10.53121/ELFS10> • ISSN (print) 2791-3880 • ISSN (online) 2791-3899

ABSTRACT

In this chapter, we present a framework to explore the opportunities of AI tools in learning and the necessary conditions to take advantage of these affordances. A qualitative systematic review of 40 empirical studies of AI use in higher education was conducted. Five synthesised findings were identified, representing opportunities and conditions of AI tool use in teaching and learning: AI competence, human–computer interaction, learning design, AI governance, and future projections. These resulted in an AI Model for Education (AIMed). The model stresses the importance of learning design in creating learning experiences, to support young people to develop AI competence and take advantage of the opportunities provided by AI tools. Practical strategies for learning design for AI tools are discussed, along with policy recommendations to support equitable use of AI tools in learning.

ABOUT THE AUTHORS

Sarah K. Howard is Chair of Digital Education at the University of Leeds and leads the Centre for Research in Digital Education

Jo Tondeur is Professor at the Vrije Universiteit Brussel, specifically in the Brussels Research Institute for Teacher Education

Charlotte Neubert is a Research Associate at the Faculty of Philosophy, Arts, History and Social Sciences at the University of Regensburg

Jonas Hauck is a Research Associate at the Chair of Educational Data Science within the Faculty of Human Sciences at the University of Regensburg

Richard Böhme is a Professor at the Institute for Educational Innovation, Salzburg University of Education Stefan Zweig

INTRODUCTION

The opportunities and risks of using artificial intelligence (AI) in work, education, and our lives are increasing in complexity.¹ However, neither its potential downsides nor the debates around its implications should overshadow the importance of young people learning how to use AI well, being able to make critical decisions about their AI use, and acquiring the competencies that will allow them to navigate a future with such technology.² Mere access to digital tools is not enough for young people to achieve the confidence and skill to deploy them in sophisticated ways.³ In this context, teachers and instructors have a crucial part to play – they can create the environment and design the experiences that build these capacities among learners. This chapter provides a framework – the AI Model for Education (AIMed) – for thinking about how to do so.

It is critical that young people are given opportunities to use AI (and any other digital technologies that may emerge) if they are to avoid systematic 'digital exclusion' from society and the workforce. Research around the digital divide and digital inclusion indicates that educational institutions can play a significant role in developing young people's competency in digital technologies, in particular AI.⁴ While it has been argued that AI tools have a democratising effect,⁵ research has shown that AI has in fact rapidly exacerbated existing inequalities.⁶ It is more important than ever that young people are given the chance to use AI tools, such as chatbots and writing-support programmes, in their learning.

It is critical that young people are given opportunities to use AI (and any other digital technologies that may emerge) if they are to avoid systematic 'digital exclusion' from society and the workforce.

Yet there are significant questions about how to best integrate AI into teaching and learning. Since the advent of ChatGPT and widespread availability of generative AI (genAI), much of the discussion about AI in education has been dominated by concerns about plagiarism, cognitive offloading and dependency, and data privacy issues.⁷ While important, these preoccupations have eclipsed re-

1. Kim, J., S. Yu, R. Detrick, X. Lin, & N. Li (2025), *Designing AI-Powered Learning: Adult Learners' Expectations for Curriculum and Human-AI Interaction*, presented at the Association for Educational Communication & Technology (AECT) 2025 Conference.
2. C. Bajada, P. Kandlbinder, and R. Trayler (2019), 'A General Framework for Cultivating Innovations in Higher Education Curriculum', *Higher Education Research & Development*, 38(3), 465–478. <https://doi.org/10.1080/07294360.2019.1572715>; J. Gläser and G. Laudel (2023), 'The Undercomplexity of Higher Education Policy Innovations', in C. Schubert and I. Schulz-Schaeffer (eds.), *Berlin Keys to the Sociology of Technology* (Wiesbaden: Springer Fachmedien Wiesbaden), pp. 161–182.
3. See, for example, C. Wang, S. C. Boerman, A. C. Kroon, J. Möller, and C. H. de Vreese (2024), 'The Artificial Intelligence Divide: Who Is the Most Vulnerable?', *New Media & Society*, 27(7), 3867–3889. <https://doi.org/10.1177/14614448241232345>.
4. K. Beckman, T. Apps, S. K. Howard, C. Rogerson, A. Rogerson, and J. Tondeur (2025), 'The GenAI Divide Among University Students: A Call for Action', *The Internet and Higher Education*, 101036, 1–9. <https://doi.org/10.1016/j.iheduc.2025.101036>.
5. Taylor, R. R., J. W. Murphy, W. T. Hoston, & S. Senkaiahliyan (2024), 'Democratizing AI in Public Administration: Improving Equity through Maximum Feasible Participation', *AI & Society*, 40, 3653–3662. <https://doi.org/10.1007/s00146-024-02120-w>.
6. T. Bircan and M. F. Özbilgin (2025), 'Unmasking Inequalities of the Code: Disentangling the Nexus of AI and Inequality', *Technological Forecasting and Social Change*, 211, 123925. <https://doi.org/10.1016/j.techfore.2024.123925>.
7. See, for example, W. J. Fassbender (2025), 'Of Teachers and Centaurs: Exploring the Interactions and Intra-actions of Educators on AI Education Platforms', *Learning, Media and Technology*, 50(3), 1–13. <https://doi.org/10.1080/17439884.2024.2447946>; M. Gerlich (2025), 'AI Tools in Society: Impacts on Cognitive Offloading and the Future

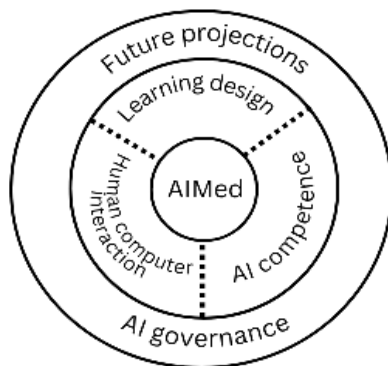
search examining the real-world opportunities of AI for students and how AI tools can be integrated into educational programmes. This has, in turn, limited the information available to instructors, teachers, and lecturers who want to better understand the potential for AI in their pedagogical practice. Ironically, a key strategy for promoting digital inclusion and competence among young people – namely, through learning programmes designed to offer high-quality engagement with cutting-edge technology – points directly to the vital role of these educators.⁸ For example, young people might be given the task of using genAI chatbots to brainstorm for a written assignment. They could then create an outline of the written piece based on the ideas generated and ask the chatbot for feedback. Documentation on how they used suggestions from the chatbot could be elicited, both to support critical engagement with outputs and to make learning processes visible.

AIMed is our framework for thinking about the integration of AI in educational contexts. AIMed was developed from a qualitative systematic review based on 40 empirical studies of AI use in higher education.⁹ While there have been many reviews of the use of AI in education, most have called for a deeper understanding of how AI can be applied to learning and the experiences of students and teachers.¹⁰ AIMed incorporates five synthesised findings based on the systematic review mentioned above, representing opportunities for AI use and conditions under which it is possible: AI competence, human–computer interaction, learning design, AI governance, and future projections.

AIMED

Before unpacking the findings that inform it, we will begin by introducing AIMed, represented schematically in Figure 1.

FIGURE 1. AI model for education



Source: Howard et al. (2025)

of Critical Thinking', *Societies*, 15(1), 6; and C. K. Y. Chan and W. Hu (2023), 'Students' Voices on Generative AI: Perceptions, Benefits, and Challenges in Higher Education', *International Journal of Educational Technology in Higher Education*, 20(43). <https://doi.org/10.1186/s41239-023-00411-8>.

8. F. Gottschalk and C. Weise (2023), 'Digital Equity and Inclusion in Education: An Overview of Practice and Policy in OECD Countries', *OECD Education Working Papers* (Paris, OECD Publishing), no. 299; S. Howard, J. Tondeur, C. Neubert, J. Hauck, and R. Böhme (2025), 'The AIMed model: Creating opportunities for AI in Higher Education', *EdArXiv*. https://doi.org/10.35542/osf.io/pzk9t_v1.

9. For a full report on the study, see Howard et al., 'The AIMed model'.

10. See, for example, F. Ouyang, L. Zheng, and P. Jiao (2022), 'Artificial Intelligence in Online Higher Education: A Systematic Review of Empirical Research from 2011 to 2020', *Education and Information Technologies*, 27(6), 7893–7925. <https://doi.org/10.1007/s10639-022-10925-9>.

The model is a conceptual aid, intended to support teachers and school leaders in thinking about their use of AI tools and the role of these tools in learning. AIMed (and its associated categories and sub-categories) identifies the areas where AI tools can be integrated and clarifies the requisite conditions for doing so. The middle and outer circles of the model contain the five synthesised findings referred to above. At the core is the ‘aim’ of using an AI tool (that is, the desired lesson, experience, or work being evoked by a learning design). This placement is not accidental – the question ‘what is needed to realise a given aim?’ is central to AIMed. The model can operate at multiple levels – from an individual learner, lesson, or course, to a programme or institution. The middle circle comprises learning design, human–computer interaction, and AI competence. Our analysis showed that the learning and teaching opportunities of AI tools are closely tied to interactions. Thus, one may first ask what kind of interaction can support a specific aim. Then, one can turn to necessary conditions – which AI competencies are needed for learners to be able to benefit from these interactions? This can open the door to effective learning designs. Such questions need not be tackled in a specific sequence, but can be addressed according to the priorities linked with a particular aim.

The outer circle deals with AI governance and future projections. AI governance is concerned with the necessary conditions under which all young people can engage with AI tools safely – and in ways that can benefit their learning. Also at this level is the opportunity to think broadly about preparing young people for the future to ensure digital inclusion (i.e., that they will have the competencies to engage in future work and learning). These opportunities are brought about by means of learning experiences that expose young people to good AI practices.

Conceptual models such as this assist practitioners in making decisions about learning designs. The following discussion goes deeper into each of these areas to provide a more granular view of how they are involved in teaching and learning.

OPPORTUNITIES AND CONDITIONS

Each of the following synthesised findings is presented with its associated categories (on the left) and sub-categories (in the middle). The synthesised finding statement, which is a recommendation for using AI tools in learning, is on the right.¹¹ Here, we will present the highlights for each that were found to be particularly important in the literature.

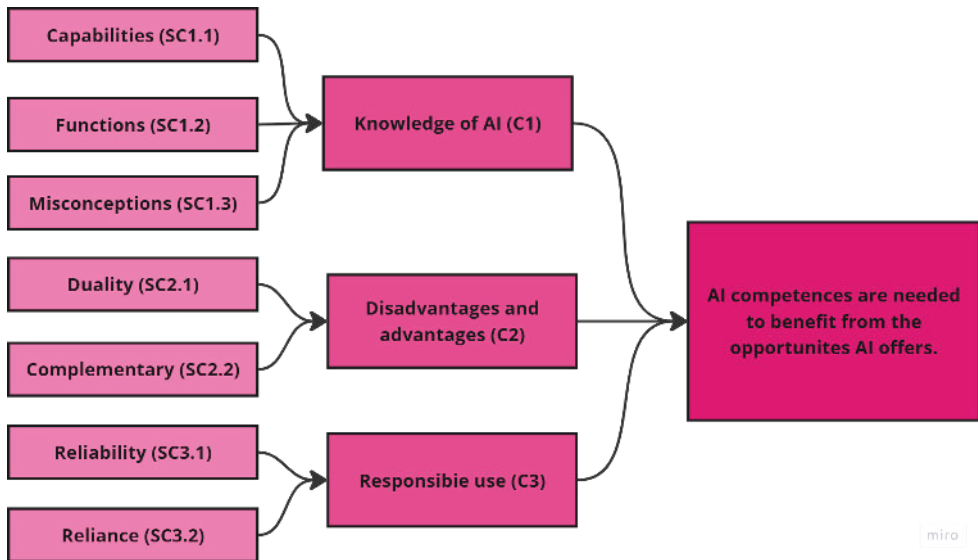
AI competence

To benefit from the opportunities AI can offer, certain competencies are necessary. As stated above, we consider AI competence to be a precondition for using AI tools in learning and teaching. The analysis revealed three categories of competency: knowledge of AI, disadvantages and advantages, and responsible use (Figure 2).

Within the ‘Knowledge of AI’ category, we highlight that to benefit from the use of a tool, an individual must have some understanding of what it is designed to do (functionality) and what it can do (capabilities). This means that both the teacher and learner must be aware of their own misconceptions about AI. A common mistake is in thinking that genAI chatbots, such as ChatGPT, personalise learning. GenAI chatbots may adapt to an individual’s inputs, but this is not personalisation. Knowledge like this is essential if AI tools are to be applied according to their respective strengths and used to create learning opportunities for young people.

11. Again, a full explanation of each of the five synthesised findings can be found Howard et al., ‘The AIMed model’.

FIGURE 2. Factors contributing to AI competence



Source: Howard et al. (2025)

Firmer knowledge of AI can build teachers' and learners' ability to spot the disadvantages and advantages the technology provides.

Firmer knowledge of AI can build teachers' and learners' ability to spot the disadvantages and advantages the technology provides. Specifically, this knowledge allows users to understand where AI can complement human activities (complementarity) and how it can produce both useful and less useful results (duality).¹² Such understanding leads to responsible use, an appreciation of where AI tools are reliable, and the capacity to avoid reliance upon them. While a depth of knowledge about AI tools is

not a prerequisite for young people to begin using them, well-designed learning experiences will contribute to improving their understanding and cultivating good decision making.

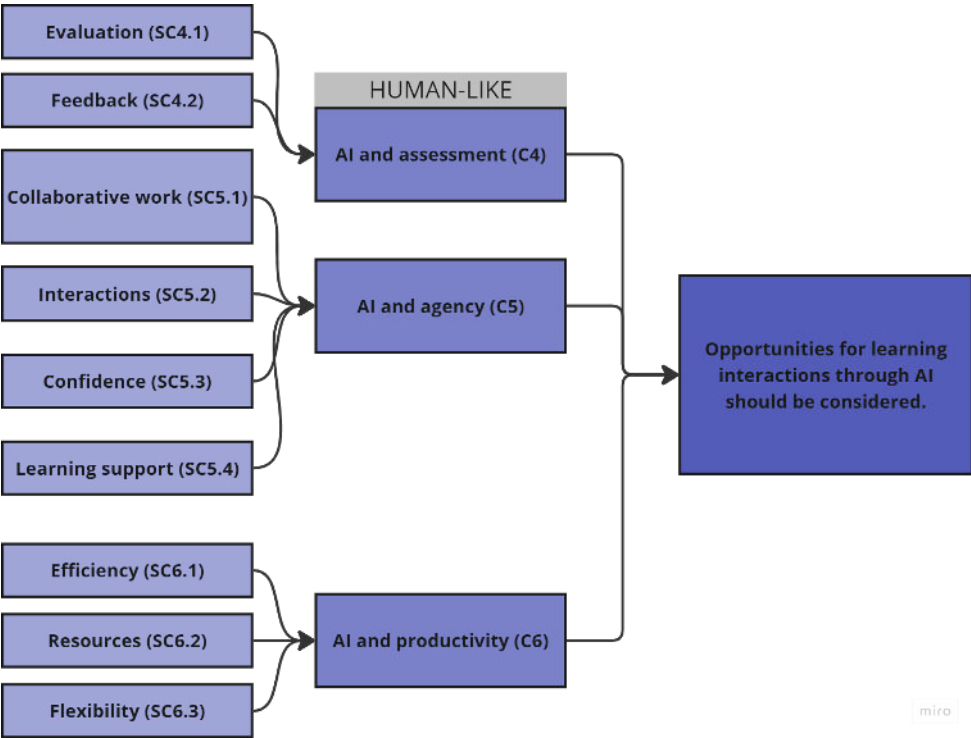
Human–computer interaction

The opportunities for learning interactions between humans and computers through AI should also be considered (Figure 3). Three categories pertain to this topic: 'AI and assessment', 'AI and agency', and 'AI and productivity'. The categories and associated sub-categories summarised in Figure 3 address affordances of AI tools to support a range of interactions in learning and teaching. Notably, 'AI and assessment' and 'AI and agency' were felt to have a 'human-like' quality, but many reported finding them mechanical and potentially unreliable.¹³ This underscores the importance of helping learners understand AI use.

12. S.-H. Jin, K. Im, M. Yoo, I. Roll, and K. Seo (2023), 'Supporting Students' Self-Regulated Learning in Online Learning Using Artificial Intelligence Applications', *International Journal of Educational Technology in Higher Education*, 20(1). <https://doi.org/10.1186/s41239-023-00406-5>.

13. S. Ebadi and A. Amini (2022), 'Examining the Roles of Social Presence and Human-Likeness on Iranian EFL Learners' Motivation Using Artificial Intelligence Technology: A Case of CSIEC Chatbot', *Interactive Learning Environments*, 32(2), 655–673. <https://doi.org/10.1080/10494820.2022.2096638>; S. Y. Liaw, J. Z. Tan, S. Lim, W. Zhou, J. Yap, R. Ratan, R., Chua, W. L. (2023), 'Artificial Intelligence in Virtual Reality Simulation for Interprofessional Communication Training: Mixed Method Study', *Nurse Education Today*, 122, 105718. <https://doi.org/10.1016/j.nedt.2023.105718>.

FIGURE 3. Categories contributing to human–computer interaction)



Source: Howard et al. (2025)

Beyond the human-like dimension of interactions with AI tools, issues of evaluation and feedback frequently came up in the context of ‘AI and Assessment’. One of the most interesting affordances of AI-generated feedback was that learners did not feel judged by AI tools when asking questions or checking their own work or understanding. This had a positive effect on learners’ motivation and their willingness to engage.¹⁴ With regard to AI and productivity, students stated that the flexibility of AI tools increased their learning efficiency by providing advice and guidance at any time.

These findings lead us to the ways of sustaining young people’s engagement in learning – ‘AI and Agency’. Interaction with AI tools brought about an improved sense of confidence among learners, who felt as if they were being supported by a more knowledgeable expert.¹⁵ This is not to say that the use of AI tools is unproblematic, but rather that they can yield opportunities to support young people if appropriately situated within a learning design.

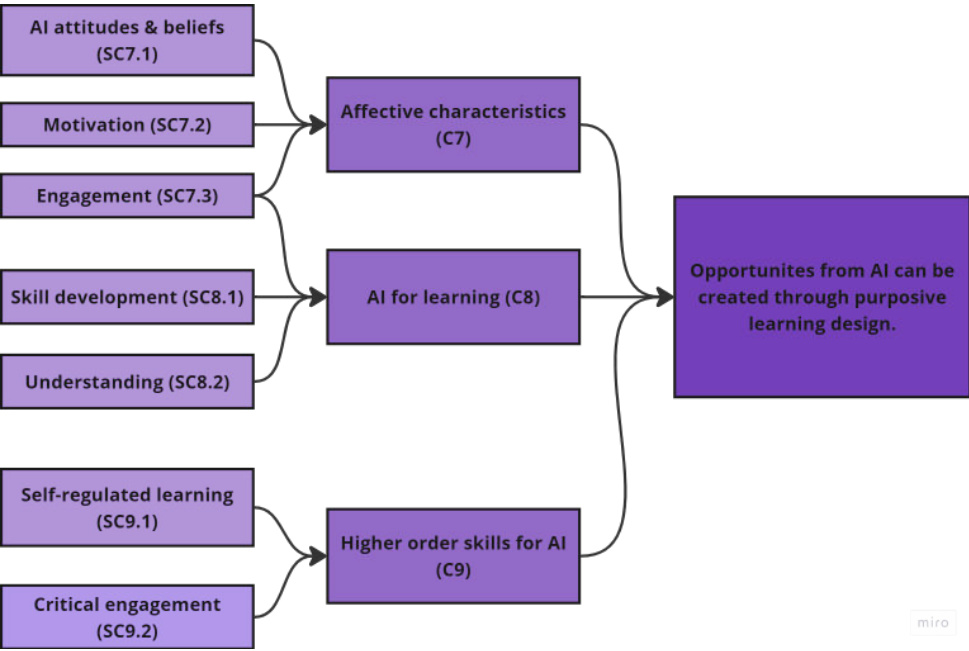
14. See, for example, K. Guo and D. Wang (2024), ‘To Resist It or To Embrace It? Examining ChatGPT’s Potential to Support Teacher Feedback in EFL Writing’, *Education and Information Technologies*, 29(7), 8435–8463. <https://doi.org/10.1007/s10639-023-12146-0>.

15. S. M. Abdelhalim (2024), ‘Using ChatGPT to Promote Research Competency: English as a Foreign Language Undergraduates’ Perceptions and Practices across Varied Metacognitive Awareness Levels’, *Journal of Computer Assisted Learning*, 40(3), 1261–1275. <https://doi.org/10.1111/jcal.12948>

Learning design

We can now turn to the finding that opportunities from AI can be created through purposive learning design (Figure 4). Three categories were identified in this ambit: affective characteristics, AI for learning, and higher-order skills for AI.

FIGURE 4. Categories contributing to learning design



Source: Howard et al. (2025)

The issue of affect is critical in the adoption of AI tools – how a learner feels about the technology has a significant effect on the results of their interaction. Affect is linked to learners’ attitudes and beliefs around AI, which can impact motivation and engagement in learning. For example, learners did not feel judged by chatbots when they asked questions.¹⁶ This substantially contributes to

The issue of affect is critical in the adoption of AI tools – how a learner feels about the technology has a significant effect on the results of their interaction.

the enhanced confidence and motivation observed among learners using AI tools. These are outcomes that good learning design can facilitate. Well-designed tasks are also instrumental in ensuring that learners do not become excessively reliant on AI or slip into cognitive offloading.¹⁷

The most obvious applications for AI in an educational context relate to skill develop-

16. A. Rahman and P. Tomy (2023), ‘Intelligent Personal Assistant – An Interlocutor to Mollify Foreign Language Speaking Anxiety’, *Interactive Learning Environments*, 1–18. <https://doi.org/10.1080/10494820.2023.2204324>.

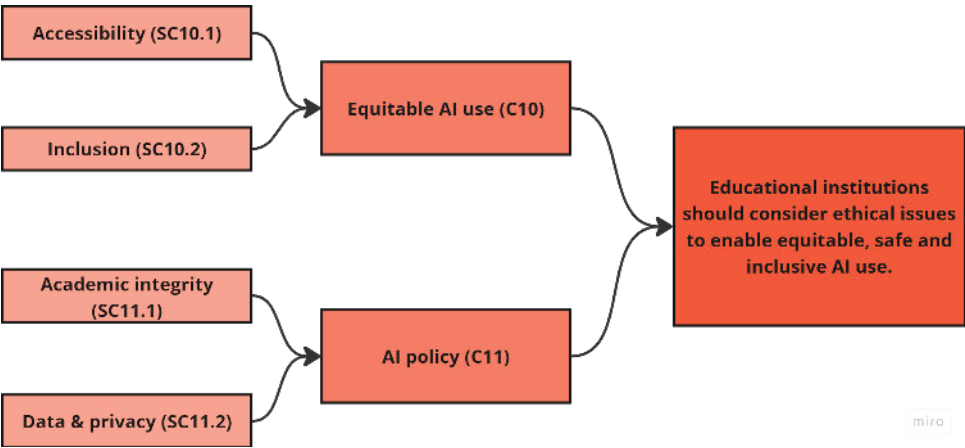
17. Cf. Gerlich, ‘AI Tools in Society’.

ment and understanding, but AI should also be considered in connection to higher-order skills like self-regulated learning and critical engagement. A balance must be struck between learning and using AI tools. While studies have shown that AI can improve writing, revising, and communication, there are concerns about loss of critical thinking.¹⁸ This brings us back to the concept of duality: AI tools come with advantages as well as risks. The latter can be accounted for via careful learning design, rooted in AI literacy and guided by a clear rationale for the technology's use.

AI governance

Figure 5 explores the finding that educational institutions should consider ethical issues to enable equitable, safe, and inclusive AI use. Two factors stand out here: equitable AI use and AI policy.

FIGURE 5. Factors related to AI governance.



Source: Howard et al. (2025)

Equitable AI use can be thought of in terms of accessibility and inclusion. It is common knowledge that the learning and use of technologies presupposes access to technological tools, and AI tools are no different. It is worth emphasising that while some AI tools (such as ChatGPT or Gemini) are free, the most powerful versions of these technologies are not. This creates an access discrepancy, which can drive unequal learning outcomes.¹⁹ Furthermore, by reflecting and reinforcing structural social biases, AI tools can be discriminative and exclusionary for some groups.²⁰ As such, these tools, even if broadly accessible in some form, may continue to widen the gap between ‘haves’ and ‘have-nots’ or between the privileged and vulnerable. This must be kept in mind when bringing AI tools into a learning environment.

18. On improvement in studies, see, for example, Ebadi and Amini, ‘Examining the Roles of Social Presence and Human-Likeness’. On critical thinking, see Chan and Hu, ‘Students’ Voices on Generative AI’.

19. A. Aydınlar, A. Mavi, E. Kütükçü, E. E. Kırımlı, D. Alış, A. Akın, A., and L. Altıntaş (2024), ‘Awareness and Level of Digital Literacy among Students Receiving Health-Based Education, *BMC Medical Education*, 24(1), 38, <https://doi.org/10.1186/s12909-024-05025-w>.

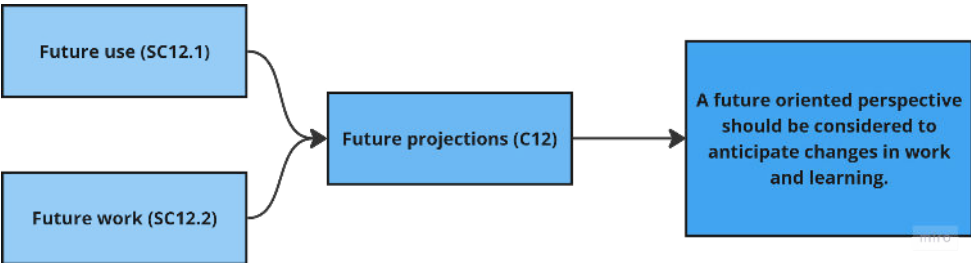
20. See, for example, Chan and Hu, ‘Students’ Voices on Generative AI’.

AI policy is key in mitigating potential harm and ensuring equal access and inclusivity. Such policy must define acceptable AI use and safeguard privacy and security.²¹ Policies should also address academic integrity, which has been a major concern in higher education – how to maintain the integrity of higher education degrees. In many cases, this has less to do with the technology’s causing problems per se, but with inadequate assessments and a lack of clear institutional guidance on good practice.²² Ultimately, for sound AI governance institutional policy that lays out unambiguous expectations of use is irreplaceable, as is modelling good practice. Taken together, these can encourage learners to engage with tools in a safe and beneficial way. Simultaneously, active support for professional learning that fosters AI competence in teachers and instructors is a must.

Future projections

Our analysis indicates that a future-oriented perspective should be adopted to anticipate changes in work and learning (Figure 6).

FIGURE 6. Categories contributing to future projections.)



Source: Howard et al. (2025)

Two categories are highlighted here: future use and future work. Both present opportunities for educational organisations, since quality experiences using AI tools can provide young people with skills for the future.²³ For example, an AI agent could be embedded into a group task to help learners stay focused on the main objectives of the assignment (e.g., by setting goals and monitoring progress).²⁴ Cooperative work along these lines closely resembles the activities common in a variety of workplaces. To be prepared for the future workplace and future use of new technologies, young people must have experiences navigating these new landscapes.

21. Concerns about AI privacy and security at the level of international policy have been expressed by the OECD, https://www.oecd.org/en/publications/ai-data-governance-and-privacy_2476b1a4-en.html, and World Bank, <https://www.worldbank.org/en/programs/accountability/data-privacy>.

22. J. Lodge, S. Howard, M. L. Bearman, P. Dawson, and associates (2023), *Assessment Reform for The Age of Artificial Intelligence*. Tertiary Education Quality and Standards Agency, <https://www.teqsa.gov.au/sites/default/files/2023-09/assessment-reform-age-artificial-intelligence-discussion-paper.pdf>.

23. Chan and Hu, 'Students' Voices on Generative AI'.

24. K. F. Hew, W. Huang, J. Du, and C. Jia (2023), 'Using Chatbots to Support Student Goal Setting and Social Presence in Fully Online Activities: Learner Engagement and Perceptions', *Journal of Computing in Higher Education*, 35(1), 40–68. <https://doi.org/10.1007/s12528-022-09338-x>

TAKE AWAY MESSAGE

This chapter has looked at AI in learning contexts, with particular emphasis on the conditions in which AI can be integrated and the opportunities it can provide. We stressed that the role AI plays in the classroom – and the resulting experiences for students – is largely determined by learning design. The opportunities teachers and instructors introduce into learning environments can be powerful, shaping young people's understanding of how they can work and learn through AI tools. Without meaningful, guided access, however, young people will end up with uneven or insufficient understanding of AI tools and the digital technologies to come.

Teacher training is a case in point: pedagogues beginning their careers need to be instructed in and familiarised with AI tools if they are to pass on such knowledge to learners. With greater competence and confidence using AI, they are more likely to employ it in their teaching practice to support their students. Under the AImed framework, one of their first tasks could be to develop a resource (primary schooling) or a series of lessons (secondary schooling) using a genAI chatbot ('learning design'). The task is tailored to the learners' level of AI competence. It includes a lecture, hands-on prompt-building activities, and 'just-in-time' video modules that guide them in using a chatbot.

In human–computer interaction training, support structures are provided to help teachers critically engage with the outputs of the AI-mediated task and reflect on them. This task, it should be noted, directly prepares learners for future work. While in the hypothetical teaching scenario being described we can assume that guidelines for acceptable use of AI are already in place at the institutional level, teachers would model for students a duty of care in the adoption of the technology. Ultimately, the successful learning design of a task involving AI should both set students up for success and expose them to uses they are likely to encounter in their future professions.

AImed provides a conceptual framework for teachers and instructors as they plan how to design learning for young people. Specifically, it guides them in considering the affordances of AI tools in education and ways to develop young people's competency as users. In several areas, AI can have meaningful benefits for learners – the increased confidence and willingness to ask questions it instils, for example, can have positive effects throughout their education. Still, as is often acknowledged, AI comes with risks, ranging from dependence to widening gaps in access. The greater risk, however, may be 'digital exclusion'. AI competence and AI governance emerge as essential for creating a learning environment that serves all young people in our rapidly evolving digital era.

REFERENCES

- Abdelhalim, S. M. (2024). 'Using ChatGPT to Promote Research Competency: English as a Foreign Language Undergraduates' Perceptions and Practices across Varied Metacognitive Awareness Levels'. *Journal of Computer Assisted Learning*, 40(3), 1261–1275. <https://doi.org/10.1111/jcal.12948>.
- Aydınlı, A., Mavi, A., Kütükçü, E., Kırımlı, E. E., Alış, D., Akın, A., & Altıntaş, L. (2024). 'Awareness and Level of Digital Literacy among Students Receiving Health-Based Education'. *BMC Medical Education*, 24(1), 38. <https://doi.org/10.1186/s12909-024-05025-w>.
- Bajada, C., Kandlbinder, P., & Trayler, R. (2019). 'A General Framework for Cultivating Innovations in Higher Education Curriculum'. *Higher Education Research & Development*, 38(3), 465–478. <https://doi.org/10.1080/07294360.2019.1572715>.
- Beckman, K., Apps, T., Howard, S. K., Rogerson, C., Rogerson, A., & Tondeur, J. (2025). 'The GenAI Divide among University Students: A Call for Action'. *The Internet and Higher Education*, 101036, 1–9. <https://doi.org/10.1016/j.iheduc.2025.101036>.
- Bircan, T., & Özbilgin, M. F. (2025). 'Unmasking Inequalities of the Code: Disentangling the Nexus of AI and Inequality'. *Technological Forecasting and Social Change*, 211, 123925. <https://doi.org/10.1016/j.techfore.2024.123925>.
- Chan, C. K. Y., & Hu, W. (2023). 'Students' Voices on Generative AI: Perceptions, Benefits, and Challenges in Higher Education'. *International Journal of Educational Technology in Higher Education*, 20(43). <https://doi.org/10.1186/s41239-023-00411-8>.

- Ebadi, S., & Amini, A. (2022). 'Examining the Roles of Social Presence and Human-Likeness on Iranian EFL Learners' Motivation Using Artificial Intelligence Technology: A Case of CSIEC Chatbot'. *Interactive Learning Environments*, 32(2), 655–673.
<https://doi.org/10.1080/10494820.2022.2096638>.
- Fassbender, W. J. (2025). 'Of teachers and Centaurs: Exploring the Interactions and Intra-actions of Educators on AI Education Platforms'. *Learning, Media and Technology*, 50(3), 1–13.
<https://doi.org/10.1080/17439884.2024.2447946>.
- Gerlich, M. (2025). 'AI Tools in Society: Impacts on Cognitive Offloading and the Future of Critical Thinking'. *Societies*, 15(1), 6. <https://doi.org/10.3390/soc15010006>.
- Gläser, J., & Laudel, G. (2023). 'The Undercomplexity of Higher Education Policy Innovations'. In C. Schubert & I. Schulz-Schaeffer (eds.), *Berlin Keys to the Sociology of Technology*, pp. 161–182. Wiesbaden: Springer Fachmedien Wiesbaden.
- Gottschalk, F., & Weise, C. (2023). 'Digital Equity and Inclusion in Education: An Overview of Practice and Policy in OECD Countries'. *OECD Education Working Papers*, no. 299. Paris: OECD Publishing.
- Guo, K., & Wang, D. (2024). 'To Resist It or to Embrace It? Examining ChatGPT's Potential to Support Teacher Feedback in EFL Writing'. *Education and Information Technologies*, 29(7), 8435–8463.
<https://doi.org/10.1007/s10639-023-12146-0>.
- Hew, K. F., Huang, W., Du, J., & Jia, C. (2023). 'Using Chatbots to Support Student Goal Setting and Social Presence in Fully Online Activities: Learner Engagement and Perceptions'. *Journal of Computing in Higher Education*, 35(1), 40–68.
<https://doi.org/10.1007/s12528-022-09338-x>.
- Howard, S., Tondeur, J., Neubert, C., Hauck, J., & Böhme, R. (2025). 'The AIMed Model: Creating Opportunities for AI in Higher Education'. *EdArXiv*.
https://doi.org/10.35542/osf.io/pzk9t_v1.
- Jin, S.-H., Im, K., Yoo, M., Roll, I., & Seo, K. (2023). 'Supporting Students' Self-Regulated Learning in Online Learning Using Artificial Intelligence Applications'. *International Journal of Educational Technology in Higher Education*, 20(1).
<https://doi.org/10.1186/s41239-023-00406-5>.
- Kim, J., Yu, S., Detrick, R., Lin, X., & Li, N. (2025). *Designing AI-Powered Learning: Adult Learners' Expectations for Curriculum and Human-AI Interaction*. Presented at the Association for Educational Communication & Technology (AECT) 2025 Conference.
- Liaw, S. Y., Tan, J. Z., Lim, S., Zhou, W., Yap, J., Ratan, R., . . . Chua, W. L. (2023). 'Artificial Intelligence in Virtual Reality Simulation for Interprofessional Communication Training: Mixed Method Study'. *Nurse Education Today*, 122, 105718.
<https://doi.org/10.1016/j.nedt.2023.105718>.
- Lodge, J., Howard, S., Bearman, M. L., Dawson, P., & Associates (2023). 'Assessment Reform for The Age of Artificial Intelligence'. Tertiary Education Quality and Standards Agency.
<https://www.teqsa.gov.au/sites/default/files/2023-09/assessment-reform-age-artificial-intelligence-discussion-paper.pdf>.
- Ouyang, F., Zheng, L., & Jiao, P. (2022). 'Artificial Intelligence in Online Higher Education: A Systematic Review of Empirical Research from 2011 to 2020'. *Education and Information Technologies*, 27(6), 7893–7925.
<https://doi.org/10.1007/s10639-022-10925-9>.
- Rahman, A., & Tomy, P. (2023). 'Intelligent Personal Assistant – An Interlocutor to Mollify Foreign Language Speaking Anxiety'. *Interactive Learning Environments*, 1–18.
<https://doi.org/10.1080/10494820.2023.2204324>.
- Taylor, R. R., Murphy, J. W., Hoston, W. T., & Senkaiahliyan, S. (2024). 'Democratizing AI in Public Administration: Improving Equity through Maximum Feasible Participation'. *AI & Society*, 40, 3653–3662.
<https://doi.org/10.1007/s00146-024-02120-w>.
- Wang, C., Boerman, S. C., Kroon, A. C., Möller, J., & de Vreese, C. H. (2024). 'The Artificial Intelligence Divide: Who Is the Most Vulnerable?'. *New Media & Society*, 27(7), 3867–3889.
<https://doi.org/10.1177/14614448241232345>.

Conclusion

Artificial intelligence is the domain where innovation, sustainability, ethics, and governance can be brought into alignment. Europe's task is not merely to regulate but to craft a model of technological stewardship that unites regulatory authority with the capacity to design, deploy, and sustain these groundbreaking systems. In the twenty-first century, sovereignty will remain with those who can both build the tools and define the terms of their use. Europe's strength lies in turning that duality into a coherent strategy.

The 'Brussels Innovation Effect' should be understood as the fusion of normative influence and practical capability. It entails a shift from the traditional 'Brussels Effect', the export of rules, to strategic technological leadership built on quality, trust, and openness. In the era of AI, true sovereignty will not be achieved through market size or regulatory reach alone, but from mastery of data and nurturing of talent. Europe should not attempt to outpace every engine of private innovation in the US, nor retreat into protectionist defensiveness. Instead, it should lead in sectors where its comparative advantages are strongest, while promoting the ethical diffusion of advanced technologies throughout society and the workforce.

This evolution demands regulation that empowers as much as it restrains. The experience of the GDPR demonstrated that clear and proportionate rules can protect citizens' rights while inspiring global standards. The same cannot yet be said of the AI Act, where overlapping provisions and lack of clarity risk creating uncertainty and slowing progress at the very moment when agility is essential. Guardrails are needed, but they must not become barriers. Europe must ensure that compliance remains rigorous but not overburdening, so that it continues to be a competitive and attractive environment for responsible innovation.

While regulatory complexity can slow deployment, it also reflects Europe's pluralistic commitment to balancing innovation with fundamental rights. The challenge is not simplification alone, but predictability and proportionate enforcement.

Institutions can give this strategy tangible form. An independent AI Ombudsman could provide credible remedies for harms and defend fundamental rights, while an AI Safety Agency could set technically sound baselines for systems that affect public life. Civic Data Trusts would transform abstract claims of data sovereignty into shared, governable assets. These are not symbolic gestures; they are working proofs that artificial intelligence can coexist with democratic accountability, avoiding both opacity and centralised control.

Technically robust lifecycle governance must accompany such institutional progress. Transparency in data provenance, adversarial robustness testing, standardised conformity assessments, and integrated sustainability metrics will be essential to prevent concentration of power and environmental harm while keeping systems auditable and secure. Harmonised standards, developed jointly with industry and civil society, will reduce fragmentation and enhance interoperability across sectors and borders.

AI's environmental implications must be recognised not only as a cause for concern but also as an opportunity. It is now evident that if left unmanaged, AI can drastically increase energy use and emissions. With the right planning, however, it can become an invaluable ally in climate action by optimising grids, accelerating renewable deployment, and transforming urban mobility. Policymakers have a choice: to allow AI to intensify existing environmental pressures or to transform it into a catalyst of long-term sustainability.

Europe's regulatory voice has already shaped global debates, but influence that remains largely symbolic will not suffice. The next phase must translate rights-based frameworks into operational systems that foster inclusion, innovation, and resilience. By championing the multi-stakeholder model and supporting global capacity-building, Europe can offer a viable liberal alternative to both laissez-faire deregulation and authoritarian governance.

If Europe matches ethical clarity with technical prowess, it can demonstrate that liberal democracies are capable not only of harnessing technology but also directing it, making governance itself a creative force that fuses freedom, sustainability, and innovation.

Achieving this ambition requires investment in European AI infrastructure, interoperable data spaces, and shared evaluation facilities that connect research, industry, and regulation. A networked governance model – linking European and national agencies – could ensure that innovation and oversight evolve in tandem.

