



RESEARCH
ARTICLE



OPEN
ACCESS



PEER
REVIEWED

How Brussels reproduces Silicon Valley technosolutionism: Sociotechnical imaginaries in the EU's regulatory approach to AI (2019–2025)

Alvaro Oleart *Université libre de Bruxelles* alvaro.oleart@ulb.be

Alejandro Flores Moleón *Universidad Autónoma de Madrid*
alejandro.floresm@uam.es

DOI: <https://doi.org/10.14763/2026.3.2111>

Published: 21 August 2026

Received: 21 November 2025 **Accepted:** 21 May 2026

Funding: The authors did not receive any funding for this research.

Competing Interests: The author has declared that no competing interests exist that have influenced the text.

Licence: This is an open-access article distributed under the terms of the Creative Commons Attribution 3.0 License (Germany) which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited. <https://creativecommons.org/licenses/by/3.0/de/deed.en>
Copyright remains with the author(s).

Citation: Oleart, A., & Flores Moleón, A. (2026). How Brussels reproduces Silicon Valley technosolutionism: Sociotechnical imaginaries in the EU's regulatory approach to AI (2019–2025). *Internet Policy Review*, 15(3). <https://doi.org/10.14763/2026.3.2111>

Keywords: Sociotechnical imaginaries, European Union, AI Act, Big tech, Artificial intelligence

Abstract: The European Union has produced the world's first comprehensive regulation on Artificial Intelligence (AI). The two main regulatory instruments mobilised by the EU are the binding 2024 AI Act and the voluntary General-Purpose AI (GPAI) Code of Practice, a combination that reflects a structural tension between the regulatory oversight exercised by EU regulators and the growing role of technology platforms as co-stewards of regulation. Drawing on the concept of sociotechnical imaginaries, understood as shared visions of desirable futures and risks to be avoided, we ask: What sociotechnical imaginaries are embedded in the EU's approach to AI? To explore this question, we empirically compare the imaginaries embedded in the 2024 AI Act, the 2025 GPAI Code of Practice and the preceding 2019 High Level Expert Group on AI report 'Ethics guidelines for trustworthy AI' as a window to trace the imaginaries upon which the EU has built its approach to AI. We examine how the EU articulates the promises and risks of AI, which actors are legitimised to govern safety, what solutions are promoted, and what type of regulatory leadership emerges. We show how responsibilities are distributed between EU institutions and Big Tech, highlight the risks of regulatory capture, technosolutionism, and reliance on private infrastructure, and point towards democratic alternatives.

1. Introduction: The EU's sociotechnical imaginary when regulating AI

Artificial intelligence (AI) has become a central political issue in Europe and beyond, with direct implications for electoral integrity, democracy, the protection of critical infrastructure, the environment, the functioning of essential services, or the management of systemic risks, among other domains. The European Union (EU) has produced the world's first comprehensive regulation on AI, a twofold regulatory response: on the one hand, the 2024 AI Act, which has already entered into force, established a risk-based approach with binding rules; and on the other, the 2025 General-Purpose AI Code of Practice (hereinafter GPAI Code of Practice), a voluntary instrument aimed at guiding compliance with general-purpose AI models through commitments made by industry players. This combination of binding legislation with voluntary industry commitments reflects a structural tension between the regulatory oversight exercised by EU regulators and the growing role of Big Tech platforms as co-stewards of tech governance in general, and AI in particular.

These instruments also reflect a dual imaginary of promise and risk. The AI Act argues that "AI is a fast evolving family of technologies that contributes to a wide array of economic, environmental and societal benefits across the entire spectrum of industries and social activities" by "improving prediction, optimising operations and resource allocation, and personalising digital solutions" (recital 4). In fact, AI "should serve as a tool for people, with the ultimate aim of increasing human well-being" (recital 6). Almost all sectors of society appear as beneficiaries of AI, even if the next article argues that "depending on the circumstances regarding its specific application, use, and level of technological development, AI may generate risks and cause harm to public interests and fundamental rights that are protected by Union law" (recital 5). This ambivalent framework places AI at the heart of EU governance not only as a regulated technology, but also as a terrain where the relationship between public authorities and Big Tech companies is negotiated.

Drawing on the concept of sociotechnical imaginaries (Jasanoff & Kim, 2009), understood as shared visions of desirable futures and risks to be avoided, in this article we ask: What sociotechnical imaginaries are embedded in the EU's approach to AI? To explore this question, we empirically compare the imaginaries embedded in the 2024 AI Act, the 2025 GPAI Code of Practice and the 2019 High Level Expert Group on AI (HLEG on AI) report 'Ethics guidelines for trustworthy AI' as a window to trace the imaginaries upon which the EU has built its approach to AI. Crucially, we examine how the EU articulates the promises and risks of AI, which actors are

legitimised to govern safety (AI Office, Big Tech, civil society, etc.), what solutions are promoted (bans, standards, voluntary codes), and what type of regulatory leadership emerges. This framework allows us to map how the AI Act and the GPAI Code of Practice articulate the division of responsibilities between the EU and Big Tech.

We show how responsibilities are distributed between EU institutions and Big Tech, highlight the risks of regulatory capture, technosolutionism, and reliance on private infrastructure. We argue that the combination of mandatory and voluntary instruments reconfigures European regulatory leadership toward a hybrid model, where Big Tech is not only an object of regulation but also a co-producer of it in sensitive areas such as systemic risk management. This shift connects directly with the debates on the privatisation of security, public-private relations, and the consolidation of Trust & Safety as a professional field within the technology industry. . We argue that the EU reproduces the technosolutionist logic stemming from Silicon Valley's Big Tech companies, corporations that overwhelmingly rely on a surveillance capitalist business model that claims "human experience as free raw material for translation into behavioural data" (Zuboff, 2019, p. 8). In the field of AI, we encompass within the label of 'Big Tech' the dominant corporations, most of them based in California's Silicon Valley in the United States (US): Google/Alphabet (Gemini, DeepMind), OpenAI (ChatGPT), Anthropic (Claude), Microsoft (Copilot), Grok (SpaceX), Amazon, Nvidia or Apple (although Apple has invested less than its competitors on AI as their revenues are coming mostly from their hardware sales). While there are other relevant companies in the field, such as French Mistral or Chinese High-Flyer (DeepSeek) and Alibaba Cloud (Qwen AI), the central AI players (particularly in the EU) remain the US-owned ones.

After this introduction, the article develops a theoretical framework based on the idea of sociotechnical imaginaries. The third section introduces the methodology that we have mobilised to uncover the sociotechnical imaginaries upon which the 2024 AI Act, the 2025 GPAI Code of Practice and the 2019 HLEG on AI report 'Ethics guidelines for trustworthy AI' are built. Fourth, we develop the empirical analysis of the EU's regulatory initiatives on AI. Last, the conclusion explains the main contributions of the article and puts forward alternative paths to democratise the relation of society with technology by pointing towards alternative non-technosolutionist imaginaries that could operate as democratic counterweights to the EU's approach.

2. Sociotechnical imaginaries and digital futures: Technocracy, risk and technosolutionism

The regulation of AI in Europe cannot be understood without reference to the imaginaries that underpin it. As Jasanoff and Kim (2009, p. 120) argue, socio-technical imaginaries are “collectively shared, institutionally stabilised and publicly represented visions of desirable futures” that are made achievable through science and technology (see also Jasanoff & Kim, 2015). Far from being neutral, these imaginaries serve as normative horizons: they prescribe not only how technologies should be developed, but also how they should be governed. As a result, these imaginaries not only reflect on and offer sociotechnical trajectories but, at the same time, coproduce the installment of these futures and, thus, yield a performative function. This performativity explains why imaginaries have become so important in debates about AI, “as they guide activities, provide structure and legitimisation, attract interest and foster investment” (Bareis & Katzenbach, p. 860). Imaginaries are not mere symbolic frameworks, but material and discursive configurations that guide collective action and define what is possible within a given social order. As Castoriadis (1987) pointed out, imaginaries constitute the creative source of social institutions, while Taylor (2004) emphasises that they are not abstract ideas, but a creation of meaning that shapes the shared moral background enabling legitimacy and coordinated action. This reading allows us to understand imaginaries not only as ideas, but as living structures that mediate between culture, politics and technology. However, as Markham points out, this power to shape the future also has an exclusionary dimension through a “discursive closure”, by naturalising, neutralising and legitimising particular values and infrastructures, closing off discussion of alternatives that might counter current hegemonic power (Markham, 2021, p. 382).

In the specific case of the EU, AI governance has been structured primarily around risk management. In this regard, complex technological issues are framed as matters that can be anticipated, classified, and addressed through the design of regulations. For instance, the 2024 AI Act embodies this logic by introducing a pyramid of criticality that differentiates systems according to their levels of risk (unacceptable risks, high risk, limited risk, minimal risk) (Chamberlain, 2023, p. 5). Through this architecture, the EU presents itself as a guardian of safety and responsibility, in line with what Bradford (2020, p. 16) identifies as the precautionary principle, which legitimises regulatory action even in the absence of absolute certainty.

However, critics have highlighted the tensions inherent in this model. Coeckelbergh (2024, pp. 1491-1497) warns of a “democratic deficit” when conflicts over

values are displaced by technical solutions. In this sense, signifiers such as “*human-centred AI*” or “*trustworthy AI*” tend to function less as substantive democratic commitments than as devices for legitimising regulatory authority. This argument is situated within debates on “technosolutionism” (Morozov, 2013, p. 415) which frames social challenges as technical problems, obscuring structural causes. Indeed, citizens are often framed as subjects who must be protected through automated solutions. Rather than empowering individuals as democratic agents, this model reproduces a paternalistic logic where technology acts as a mediator between risk and public action. Thus, technological solutionism allows us to understand that the European technocratic imaginary is not limited to risk management, but reproduces a paternalistic vision of digital government where AI appears as both a threat and a solution.

This technosolutionist mindset does not operate in a vacuum, but is built upon a structural dependence on major technology platforms, whose integration into the provision of public services reinforces their ability to shape the very regulatory frameworks that are supposed to oversee them. In this sense, it is important to highlight the structural dependence on these platforms, since “the more entrenched intermediaries are in the provision of public (infrastructural) services, the more indispensable they may appear and the better they may be able to impose their preferences in institutional reforms” (Kausche & Weiss, 2025, p. 292). Moreover, this power develops through non-binding codes of conduct that reinforce co-regulation (Oleart & Bouza, 2025).

Likewise, it is also necessary to highlight the discourse and epistemic influence (Birhane et. al, 2026). Thanks to their expertise and information advantages, intermediaries often occupy the driver's seat: “their expertise, information advantages and experience place them in a better position than the regulator... [so] intermediaries become the leaders of regulation, with the apparent legislators following them” (Kausche & Weiss, 2025, p. 293). Intermediaries “seek to influence the beliefs of others by promoting their own ideas”, often presenting their preferred solution as the neutral option (Kausche & Weiss, 2025, p. 292). The more they present themselves as responsible experts deploying (seemingly) neutral technosolutionist arguments, the better they can shape reform agendas and public policies. Technosolutionism also operates as a grammar of legitimation: companies “strongly rely on solutionist ideas that describe security problems as technological problems and posit a compatibility between the economic interests of private actors and the broader political or normative interests of public actors” (Obendiek & Seidl, 2023, p. 1307). By presenting their technical arrangements as neutral and necessary, they

normalise their role as co-governors and persuade decision-makers to accept and adopt their views. Conceptually, this fits with governance as public-private hybridisation: public authorities increasingly turn to so-called “experts” that are conceived as neutral actors in order to articulate public-private partnerships. In consequence, the regulation of technology is often shaped by complex multi-level institutional politics, geopolitical competition and strong corporate interests (see Calcara et al., 2020).

Tracing the socio-technical imaginaries (Jasanoff & Kim, 2009) on which the EU has built its regulatory response to the rise of AI is not only an ideological endeavour, but also a material one. In this regard, there is a wealth of literature documenting the many ways in which large technology companies have lobbied the EU to influence policy (Bank et al., 2021; Bouza García & Oleart, 2024; Bouza et al., 2025). In fact, no other economic sector currently spends as much money in lobbying as the tech sector in Brussels, with Meta, Microsoft, Apple, Amazon and Google among the top spenders (CEO, 2025). The large amount of resources of tech lobbyists in Brussels helps to explain how Big Tech companies have skilfully navigated the EU arena on multiple policy issues, such as the Digital Services Act (Oleart & Bouza, 2025) and the AI Act. This is not entirely surprising, as business lobbying in the EU has historically been more powerful than that of civil society or trade unions (Coen et al., 2021).

Finally, this regulatory capture by Big Tech companies cannot be separated from the EU’s own self-perception as a global regulatory superpower, an ambition that is taking on an increasingly geopolitical dimension. This is particularly relevant considering the increasing geopolitical framing of tech regulatory debates. Indeed, Anu Bradford (2023) has put forward the thesis of the ongoing global competition between three so-called ‘Digital Empires’: the free-market US model, the authoritarian and state-led Chinese model, and the rights-based model of the EU. The EU has entirely assumed this framing, as illustrated by the Commission’s 2025 AI Continent Action Plan where there are explicit references to a perceived “AI race” according to which the EU risks falling behind China and the United States (European Commission, 2025a). However, this narrative illustrates a colonial mindset. In this sense, there is an increasing corpus of literature that looks at technology (and particularly AI) through an anti-colonial lens, as illustrated by Kate Crawford’s ‘Atlas of AI’ or Karen Hao’s ‘Empire of AI’. Indeed, Couldry and Mejías put forward a crucial question to analyse the relationship between Big Tech and democracy from the perspective of intertwining the articulation of capitalism and colonialism: “what if we interpret data and technology today in terms of not just historic, but

contemporary and evolving, relations between colonialism and capitalism?” (Couldry & Mejías, 2023, pp. 787-788). From this perspective, the EU is competing with China and the US for global leadership in AI, naturalising the hierarchy between the Global North and the Global South. As Bradford argued in her popular ‘Brussels Effect’, the normative justification for the EU to have a global regulatory influence is that it allows countries with less resources to “outsource their regulatory pursuits to a more resourceful and experienced agency” (Bradford, 2020, p. 253). In doing so, it further entrenches global injustices and political inequalities. Looking at the EU’s role in regulating technology from a decolonial perspective is relevant considering the past and present of European colonialism (Rodney, 2018 [1972]), as well as the specific colonial origins of European integration (Hansen & Jonsson, 2014). In spite of calls to (un)learn Europe and the EU through a decolonial lens (Oleart et al. 2025), the dominant narrative in the field of technology is heavily shaped by Bradford’s work, situating the EU as a progressive “rights-based” actor and global regulatory leader.

That said, the EU is certainly aware of normative concerns around its own regulatory approach to technology. It is in this sense that we can make sense of the combination of geopolitical ‘hard power’ with a shift towards ‘ethics’. An interesting example of this could be how ethics functioned as a hybrid instrument. In this regard, the 2019 High-Level Expert Group on AI Guidelines were criticised as industrial “ethics washing” and for being “watered down” as a result of compromises between the public interests of sustainability and “the common good” and the tech industry’s interests of boosting industrial capacity and growth (Carlsson and Rönnblom, 2022, p. 2). This shift ‘towards ethics’ and away from politics has been attributed to industrial pressure to avoid binding rules. Relatedly, there is evidence that the EU mobilised democratic innovations to legitimise tech policies such as the AI Act in what constituted a form of ‘citizenwashing’ (Petit & Oleart, 2025).

In short, these dynamics illustrate a “regulatory decoy” driven by a techno-legal solutionism, where institutional design and ideational dominance reinforce each other, ensuring that regulation serves incumbent benefits rather than the public interest (Vertesi et al., 2026, p. 11). Hence the democratic risks: vagueness about data ownership and control and avoidance of limits on extraction; greater control of public affairs by the technology industry (Carlsson & Rönnblom, 2022, p. 8); and a contradiction between democracy and Big Tech. This hybrid regime also reconfigures the governor’s dilemma: the dependence on private companies to improve their competence and the need to control it (Obendiek & Seidl, 2023, pp. 1310–1311). Beliefs about competence and controllability are ideationally mediat-

ed, so firms invest in a reputation for expertise and effectiveness aligned with public values. Given that this authority is not only objective capacity but also ideational power, it is understandable that companies are not only regulated but also integrated as co-governors. In doing so, they contribute to define standards, provide infrastructure, and co-monitor compliance, while the EU maintains its image as a firm, rules-based regulator.

Overall, the combination of the AI Act— a binding regulatory framework—and the GPAI Code of Practice —a voluntary framework co-produced with industry—signals the rise of a hybrid governance model. The GPAI Code of Practice signals how EU regulatory authority is intertwined with the self-regulatory commitments of Big Tech companies, shifting part of the responsibility for safety and risk control to the private intermediaries themselves. This architecture is not merely complementary, but constitutive of the way in which the EU exercises its regulatory leadership: while the AI Act projects technocratic determination, the GPAI Code of Practice institutionalises dependence on corporate infrastructures to implement and oversee fundamental commitments.

3. Methods and data: Tracing socio-technical imaginaries of AI in EU public policies

The analysis combined deductive coding, guided by the theoretical framework of sociotechnical imaginaries, with inductive coding to capture new and unexpected patterns in how AI is problematised and legitimised across EU policy instruments. This mixed strategy allowed the conceptual framework to remain sensitive to the empirical material, connecting the normative structure of the texts with their discursive dynamics and institutional context. In practical terms, it enabled comparison between binding regulation (*AI Act*) and voluntary frameworks (*General-Purpose AI Code of Practice*) with the 2019 HLEG's Ethical Guidelines serving as a normative precursor. The three documents are not exactly equivalent nor equally important, as the AI Act is clearly the most relevant and impactful one, particularly in legal terms. However, the three documents are helpful to trace the underlying socio-technical imaginaries upon which the EU has constructed its vision and desired future for AI. Our research approach is similar to Gernot Rieder's (2018) critical analysis of the European Commission's socio-technical imaginaries of Big Data. Interestingly, Rieder argues that 'Big Data' is not simply an analytical concept to define an established reality, but a metaphor that operates as a legitimising device that is key to situate Big Data as a crucial element for the EU's global competitiveness and a more prosperous digital future. This type of imaginaries are in direct

tension with democracy, as the ever growing need for data collection and surveillance is adopted by both corporate actors as well as public authorities (Waller, 2020).

The coding scheme was organised around eight analytical dimensions specifically adapted to the study of sociotechnical imaginaries: (1) desired future, referring to the visions of the social order or desirable outcomes that AI is expected to bring about; (2) geopolitics, encompassing references to global competition, strategic autonomy or the EU's regulatory leadership; (3) narrative justification, capturing the values invoked to legitimise intervention, such as security, democracy, resilience, competitiveness or trust; (4) responsible actor, identifying the entities portrayed as legitimate to govern or implement AI (EU institutions, the AI Office, Big Tech, civil society, users...); (5) AI as problem or solution, distinguishing whether AI is depicted primarily as a source of risk, a governance instrument or a dual object; (6) issues related, indicating the specific policy domains linked to AI, such as elections, content moderation, systemic risk management, compliance or innovation; (7) risk, referring to the types of threats identified (fundamental rights, systemic, geopolitical, social...); and (8) function attributed to AI, capturing what AI is expected to do in practice, from ensuring transparency and detecting harmful content to enforcing compliance and strengthening security.

These dimensions are inspired in the framework of sociotechnical imaginaries. The desired future, risks and issues related captures their *normative orientation* (Jasanoff & Kim, 2009) ; geopolitics reflects the EU's external projection of imaginaries as global regulatory power; narrative justification traces the *social imaginary signification* that sustains legitimacy (Castoriadis, 1987); responsible actor reveals how authority and accountability are stabilised within a strategic action field (Fligstein & McAdam, 2012, p. 256); and function attributed to AI express the performativity that structure EU governance (Bareis & Katzenbach, 2021).

For each document, a close reading was conducted to identify and code relevant passages. In the 2024 *AI Act* (European Commission (2024)), analysis focused on the recitals and articles that define risk hierarchies, conformity obligations and governance structures; in the 2025 *GPAI Code of Practice* (European Commission (2025b)) attention was given to the commitments, measures and reporting mechanisms addressing model safety, transparency and lifecycle security (the code itself is divided in three chapters: transparency, copyright and Safety and Security); and in the 2019 *HLEG's Ethical Guidelines* (European Commission 2019) the focus was on ethical principles, the assessment list, and indicators of participation and contestability. Although this structure facilitated comparison, the interpretation was integra-

tive rather than fragmentary, aimed at capturing how AI functions as a *background imaginary*: a structuring presence shaping institutional expectations, governance practices and legitimising narratives, even when not explicitly named.

Given the institutional and normative nature of the corpus, language is predominantly formal and technical, with limited metaphorical density. Accordingly, linguistic features were secondary to the analysis of functions, actors and justifications. To ensure traceability, every finding refers to identifiable text segments (instrument, section) and to corresponding codebook rules. Finally, the analysis identified moments where policy discourse tended to dematerialise AI by overlooking its labour, infrastructural or environmental dimensions.

4. The technosolutionist imaginary of AI present in EU regulations and democratic alternatives: human-centrism as political bullshit

The EU launched the European AI Strategy in April 2018, after which the Commission put forward the High-Level Expert Group on Artificial Intelligence, composed by a set of experts, including from the industry. Contemporaneously, in December 2018 the European Commission's Joint Research Centre published a big report on AI where it was already making the case that, given the ongoing global competition between US, China and the EU, the EU "should embrace the opportunities afforded by AI", even if not "uncritically" (European Commission 2018). In parallel, the EU also launched the European AI Alliance as a wide multi-stakeholder forum that would provide feedback on the work of the High Level Expert Group on AI as well as EU policy-making on AI. Notably, these two key forums (High Level Expert Group on AI and AI Alliance) brought together a broad range of stakeholders, including major industry actors (Smuha, 2024).

The AI Act is conceived by the EU as a crucial piece of legislation that will reinforce a future where AI technology plays a central role, "promoting the European human-centric approach to AI and being a global leader in the development of secure, trustworthy and ethical AI" (recital, 8). In doing so, the EU conceives of itself as a global "leader" in the management of a "human-centric AI". While actors such as Russia or China are not directly mentioned, the AI Act refers to the need to "address the risks of undue external interference" (recital, 62), which highlights the important geopolitical component with which the legislation has been designed. This conception aligns with the 2019 Ethical Guidelines of the High-Level Expert Group on AI, which had already articulated European leadership in explicitly ethical terms. In its 2019 guidelines, the group defined "trustworthy AI" as one

grounded in fundamental rights and supported by four principles: respect for human autonomy, prevention of harm, fairness, and explicability. The High-Level Expert Group on AI states a trustworthy approach is key to enabling 'responsible competitiveness', enabling producers to gain competitive advantage while positioning Europe as "the home and leader of cutting-edge and ethical technology" (European Commission, 2019, p. 4). Indeed, the AI Act explicitly acknowledges that its imaginary builds on the guidelines drawn years earlier by the High Level Expert Group on AI : "While the risk-based approach is the basis for a proportionate and effective set of binding rules, it is important to recall the 2019 Ethics guidelines for trustworthy AI developed by the independent High-Level Expert Group on AI appointed by the Commission" (recital 27).

The responsible actors included in the 2024 AI Act are mostly "providers and deployers" of AI. Interestingly, Big Tech companies are not directly mentioned. The only companies that are referred to are Small and Medium Enterprises (SMEs), actors that are situated as potential beneficiaries of AI. In fact, the AI Act includes a reference whereby some copyright law might be nuanced to ensure these companies can benefit from AI usages: "Without prejudice to Union copyright law, compliance with those obligations should take due account of the size of the provider and allow simplified ways of compliance for SMEs, including start-ups, that should not represent an excessive cost and not discourage the use of such models" (recital, 109). The GPAI Code of Practice reinforces the centrality of providers by positioning the signatories as directly responsible for governance and safety. They must "create, implement, and update" the Safety and Systemic Risk Framework (Safety & Security, Commitment 1) and define "systemic risk acceptance criteria" to guide model deployment (Safety & Security, Commitment 4). Internally, the Code requires "clear responsibilities" and "appropriate resources" (Safety & Security, Commitment 8), including executive roles such as the Chief Risk Officer, thereby consolidating an organisational architecture of co-regulation under the supervision of the AI Office. The 2019 High-Level Expert Group on AI expands this vision by promoting a shared, multi-actor responsibility that involves both European institutions and civil society. Its ethical guidelines assign specific roles to developers, executives, regulators, and citizens, seeking to embed ethics into everyday practices.

The generic references to "providers" underlines the flaws of the very concept of "human-centric". Such a concept is not defined, but rather mentioned as a broad reference to counter the idea that AI will function autonomously from humans. But which humans ought to be situated in the centre of "human-centrism"? Is it hu-

mans such as Mark Zuckerberg, Jeff Bezos or Elon Musk, or are content moderators and data annotators being exploited in the Global South for the purpose of building AI systems? The perceived lack of trade-offs between the companies building AI systems and the broader society depoliticises AI, and makes it possible to conceive it as potentially positive for everyone in society. The GPAI Code of Practice shifts the “human-centered” ideal toward a form of administrative humanism, where ethical aspirations are translated into technical procedures and standardised documentation. Through Commitment 7 and Measures 7.1 and 7.3 (Safety & Security), trust in AI is grounded in its *documentability*: signatories must produce an updated Model Report, track version changes, and record evaluation results such as benchmarks or red-teaming outcomes. Within this framework, trust moves away from public deliberation and toward procedural legibility—an AI system becomes “trustworthy” to the extent that it can be described, measured, and audited within a bureaucratic regime of controlled transparency. In contrast, the High-Level Expert Group on AI warns that ethics cannot be reduced to norms or technical codes, emphasising that “specific ethics code – however consistent, developed and fine-grained future versions of it may be – can never function as a substitute for ethical reasoning itself, which must always remain sensitive to contextual details that cannot be captured in general Guidelines” (European Commission, 2019, p. 9).

The framing of “human-centrism” in the AI Act and the GPAI Code of Practice can be conceptualised as a form of political bullshit, a type of discourse oriented to deceive the public and avoid addressing the inescapable trade-offs that the development of AI entails. Therefore, the political economy of AI is entirely absent from the AI Act, in which AI is perceived to evolve as a technology that can “provide key competitive advantages to undertakings and support socially and environmentally beneficial outcomes” in almost every societal area, ranging from healthcare and agriculture to climate change mitigation and public services (AI Act, recital 4). There are some risks related to AI, which explains the “risk-based approach” of the AI Act. The AI Act explains that this approach “should tailor the type and content of such rules to the intensity and scope of the risks that AI systems can generate. It is therefore necessary to prohibit certain unacceptable AI practices, to lay down requirements for high-risk AI systems and obligations for the relevant operators, and to lay down transparency obligations for certain AI systems” (recital 26). However, in spite of these risks, the imaginary upon which the AI Act is constructed is based on the idea that such risks can be overcome in order to deploy technology in a beneficial way for society. The GPAI Code of Practice crystallises this anticipatory logic in a streamlined governance architecture. It mandates a state-of-the-art Safety and Security Framework to keep systemic risks “acceptable” (Safety & Secu-

rity, Commitment 1), using transparency as a preventive tool for pre-deployment assessment. Providers must follow a continuous monitoring cycle—model evaluations, risk estimation, and post-market monitoring—deploying models only when residual risk meets the acceptability threshold (Safety & Security, Measures 3.1–3.5). In parallel, the GPAI Code of Practice advances a defense-in-depth approach, requiring mitigations “robust under adversarial pressure” (Safety & Security, Measure 5.1), lifecycle-wide cybersecurity (Safety & Security, Commitment 6), and attention to “non-state external threats” (Safety & Security, Measure 6.1). Together, these provisions consolidate a proportional relationship between model capability and security safeguards. The High Level Expert Group on AI framed AI governance as an ongoing and inherently open-ended endeavour, noting that the Ethical Guidelines “should be seen as a living document to be reviewed and updated over time to ensure their continuous relevance as the technology, our social environments, and our knowledge evolve” (European Commission, 2019, p.3). This orientation was reinforced by the acknowledgement that “AI systems are continuously evolving and acting in a dynamic environment,” making “the realisation of Trustworthy AI therefore a continuous process” (European Commission 2019, p.20). At the same time, the document sought to balance the promotion of innovation with harm prevention, stating that the goal is “seeking to maximise the benefits of AI systems while at the same time preventing and minimising their risks” (European Commission, 2019, p. 4).

Throughout the AI Act, there is no acknowledgement of the direct material consequences of the development of AI systems, including the massive environmental impact of AI systems or the work performed by millions of workers (often in the Global South) in doing data annotation and other AI trainings. The AI Act does include a reference to “critical infrastructure”, in which examples “of safety components of such critical infrastructure may include systems for monitoring water pressure or fire alarm controlling systems in cloud computing centres” (recital 55). In spite of this reference, most of the time AI is conceived as a “software” and mostly fails to incorporate the important material dimension of it. The GPAI Code of Practice deepens this dematerialisation of AI by focusing solely on the technical management of risk, without addressing the environmental impact or the human labor that underpins model development. AI is portrayed as a *software* system governable through documentation and standardised procedures, consolidating an imaginary of purely technical control, detached from the ecological, labor, and material dimensions of the digital ecosystem. In contrast, the High-Level Expert Group on AI places social and environmental well-being among its core evaluation pillars, urging the assessment of environmental impact, attention to bias, and the

establishment of safeguards against surveillance and manipulation (European Commission, 2019 p. 10, p. 19). In doing so, it reconnects AI governance with its material and distributive consequences, situating ethics within broader social and ecological accountability.

In coherence with previous co-regulatory frameworks such as the 2025 Code of Conduct on Disinformation included within the framework of the DSA, the AI Act also attempts to build cooperative relations with companies and “stakeholders”. In fact, the AI Act argues that “In cooperation with the relevant stakeholders, the Commission and the Member States should facilitate the drawing up of voluntary codes of conduct to advance AI literacy among persons dealing with the development, operation and use of AI” (AI Act, recital 20). Relatedly, much of the responsibility for the well functioning of AI is placed on individuals. Rather than relying on a more structural approach that situates Big Tech companies liable and overwhelmingly responsible, the notion of “AI literacy” (see Oleart & Flores Moleón 2026): “In order to obtain the greatest benefits from AI systems while protecting fundamental rights, health and safety and to enable democratic control, AI literacy should equip providers, deployers and affected persons with the necessary notions to make informed decisions regarding AI systems” (AI Act, recital 20). The GPAI Code of Practice takes this logic of co-regulation a step further by institutionalising a bureaucracy of continuous transparency. Rather than promoting literacy or public dialogue, it requires signatories to prepare and regularly update a technical model report before and during commercialization, documenting the model’s architecture, risks, evaluation results, and external reviews. Cooperation thus becomes a permanent documentation regime, where trust is built through technical traceability and procedural oversight rather than public deliberation. In contrast, the High-Level Expert Group on AI offers a more open and participatory interpretation of this cooperative ideal. It proposes combining qualitative pilot projects in selected organizations with a public, quantitative registry open to all stakeholders, ensuring that AI ethics operates as a collective and revisable process, rather than as an internal exercise in regulatory compliance (European Commission, 2019, p. 24).

TABLE
1:
Dimensions
of
the
sociotechnical
imaginary
of
AI
in
the
AI
Act,
the
GPAI
Code
of
Practice
and
the
High-
Level
Expert
Group
on
AI

DIMENSIONS OF SOCIOTECHNICAL IMAGINARY	AI ACT	GPAI CODE OF PRACTICE	HLEG ETHICAL GUIDELINES
DESIRED FUTURE	“AI as a human-centric technology that should serve as a tool for people, with the ultimate aim of increasing human well-being” (recital 6)	A continuously updated ecosystem in which providers must “create, implement and update” the Safety & Security Frameworks (Safety & Security, Commitment 1), produce Model Reports (Safety & Security, Commitment 7), deploy models only once systemic risk acceptance criteria are met (Safety & Security, Commitment 4).	A Europe that is “the home and leader of cutting-edge and ethical technology” (pp. 4-5), enabling “responsible competitiveness” and positioning Europe as a global model of ethical governance.
GEOPOLITICS	EU as a global leader in the uptake of trustworthy AI.	EU regulatory sovereignty centred on the AI Office, and alignment with international standards (Safety & Security, Measure 7.3). Attention to “non-external threats” (Safety & Security, Measure 6.1)	Europe as a leader through “responsible competitiveness” (pp. 4-5) and a model of governance based on fundamental

DIMENSIONS OF SOCIOTECHNICAL IMAGINARY	AI ACT	GPAI CODE OF PRACTICE	HLEG ETHICAL GUIDELINES
			rights, as an implicit alternative to other global models.
NARRATIVE JUSTIFICATION	Promotion of the uptake of human centric and trustworthy AI and support innovation	Legitimacy stems from state-of-the-art methods (Safety & Security, Commitment 1), continuous reporting (Safety & Security, Commitment 7) and procedural traceability.	Trustworthiness as a necessary condition for society to adopt AI and realise its benefits; ethics as a legitimate competitive advantage.
RESPONSIBLE ACTORS	Mostly providers and deployers of AI (Big Tech companies are absent) and individuals (AI literacy).	Signatories (GPAI providers) design, monitor and report under AI Office oversight. Internal accountability is reinforced through "clear responsibilities" and executives roles such as the Chief Risk Officer (Safety & Security, Commitment 8)	Shared, multi-actor responsibility that involves both European institutions and civil society. Its guidelines assign specific roles to developers, executives, regulators, and citizens.
AI: PROBLEM OR SOLUTION?	While bearing some risks, AI can bring benefits to society in almost every area.	AI seen as a potential systemic risk manageable through continuous monitoring (Safety & Security, Measures 3.1–3.5)	Ambivalent, seeks to maximise the benefits of AI systems while at the same time preventing and minimising their risks.
ISSUES RELATED	Healthcare, education, climate, innovation, culture, security, energy, public services, justice, agriculture, food safety, media, sports...	Comprehensive risk pipeline (Safety & Security, Measures 3.1–3.5), copyright compliance (1.1–1.5) and lifecycle-wide cybersecurity (Safety & Security, Commitment 6).	Social and environmental well-being, bias, and safeguards against surveillance and manipulation as core governance issues (p. 10, p. 19)
RISK	The risk-based approach is oriented to manage	Systemic and technically manageable through a "defence-in-depth approach" requiring	Primary concerns include/manipulation, unjustified surveillance, and black-box opacity

DIMENSIONS OF SOCIOTECHNICAL IMAGINARY	AI ACT	GPAI CODE OF PRACTICE	HLEG ETHICAL GUIDELINES
	the potential harmful risks and harness the potential of AI.	lifecycle-wide cybersecurity (Safety & Security, Commitment 6) and a continuous assessment cycle (Safety & Security, Measure 4.2).	and inequality. Proposes safeguards, human oversight, and ethics as a political and deliberative practice.
FUNCTION ATTRIBUTED TO AI	Enable “a European ecosystem of public and private actors creating AI systems in line with Union values and unlocking the potential of the digital transformation across all regions of the Union” (recital 8).	General-purpose infrastructure governed through continuous documentation – model evaluations (Safety & Security, Measures 3.2), and Model Reports (Safety & Security, Measure 7.1)	A means—not an end— to serve the common good and human development, provided it remains subject to participatory, context-sensitive ethical oversight and human rights constraints.

The findings of this study align with Griffin’s (2023, p. 10) argument that the EU’s approach reinforces “the image of platforms as benevolent stewards of the public interest, rather than companies pursuing private gain”. In doing so, lobbyists for Big Tech companies have succeeded in framing policy debates and regulation as a ‘technical’ question, and they are conceived as neutral “third-party-experts” (Kausche & Weiss, 2025, p. 303) rather than as actors with an evident conflict of interest. In spite of being recognised as a threat to democracy (Schaae, 2024), Big Tech companies continue to be situated as a reliable ‘partner’ by the EU in its regulatory logic. AI poses problems conceived as *technical* that can be ‘solved’ through the cooperation of different stakeholders. Similarly, the GPAI Code of Practice embodies a distinctly technical approach to AI governance. Its commitments focus on exhaustive documentation—“ensuring the quality, security, and integrity of the documented information” (Copyright, Measure 1.2, 1.3)—and on procedural safeguards such as “not circumventing effective technological measures” Collectively, these provisions shift the ethical and political debate about AI’s purposes and structural inequalities into a realm of technical compliance and documentary verification, where legitimacy derives from procedural precision rather than public deliberation. In contrast, the High-Level Expert Group on AI explicitly rejects this technocratic reduction: it warns it is not about ticking boxes, but about continuously identifying and implementing requirements, evaluating solutions, ensuring improved outcomes throughout the AI system’s lifecycle, and involving stakeholders

in this (European Commission, 2019, p. 3). Whereas the GPAI Code of Practice translates ethics into protocols and audits, the HLEG envisions it as a political and deliberative practice, one meant to question—rather than reproduce—the power structures shaping AI development.

Our comparative analysis emphasises that there is a continuous thread across the AI Act, GPAI Code of Practice and High-Level Expert Group on AI. However, the High-Level Expert Group on AI was more explicit about questions related to ethics beyond ticking boxes, whereas the AI Act and the GPAI Code of Practice are much more industry-oriented documents. Hence, the comparison reveals a gradual evolution rather than a rupture. The High-Level Expert Group on AI promotes an adaptive and precautionary form of AI governance that balances innovation with social protection and explicitly includes social and environmental well-being among its ethical principles. In contrast, the GPAI Code of Practice focuses almost exclusively on the technical management of risk, overlooking the material, labour, and ecological dimensions of AI. While the High-Level Expert Group on AI calls for participatory and context-sensitive ethics, the GPAI Code of Practice institutionalises a bureaucratic regime of continuous transparency based on reports, audits, and procedural traceability. In this sense, the High-Level Expert Group on AI and the GPAI Code of Practice do not embody opposing visions but rather two stages of the same trajectory: the former lays the ethical and reflective foundations, while the latter formalises and bureaucratises them through technical and measurable procedures. The AI Act culminates this process by enforcing a risk-based regulation oriented towards minimising the risks that AI brings, yet still conceiving AI as a type of technology that has the potential to “solve” many “problems” in society. The AI Act is heavily in support of “innovation”, even if it comes at the expense of democracy and the environment. In doing so, the EU’s approach normalises AI as an ‘inevitable’ technology that is generally ‘good’, reproducing what has been Vertesi et al. (2026) called the “inevitability decoy” (Vertesi et al., 2026, pp. 6-8): a discursive operation that frames AI as a predetermined future, foreclosing alternative trajectories, and rushes certain (AI-powered) futures while preventing others from taking hold. Such a technosolutionist imaginary not only contributes to entrench Big Tech companies in society, but situates them in a powerful regulatory role, as the enforcement of the AI Act and the GPAI Code of Practice is reliant on these companies cooperating.

5. Conclusion: How the EU continues empowering Big Tech through its AI (co)regulation

While there is a growing literature critical of Big Tech, AI and new technologies (Crawford, 2021; Merchant, 2023; Adams, 2024; Hao, 2024; Bender & Hanna, 2025), we still know relatively little about the specific sociotechnical imaginary upon which EU public policies on AI are built. This article has made a comparative analysis of the sociotechnical imaginary present in the 2024 AI Act, the 2025 GPAI CoP and the 2019 High Level Expert Group on AI report ‘Ethics guidelines for trustworthy AI’. We paid particular attention to the distribution of power and responsibilities between regulators and platforms. Our article is particularly relevant in the context of a deregulatory push by the second von der Leyen Commission (2024-2029). Indeed, in November 2025, within the broader proliferation of Omnibus proposals, the Commission put forward the Digital Omnibus package (Omnibus VII), composed of two proposed regulations: the Digital Omnibus (Data Omnibus), and the Digital Omnibus on AI. Both components of the Digital Omnibus package are oriented towards deregulating existing EU digital laws. The Data Omnibus is closely related to AI, given that it proposes a change in GDPR that would facilitate the exploitation of personal data by Big Tech to train AI models and AI systems (EDRI, 2025a; Shrishak, 2025). The Digital Omnibus on AI focuses primarily on watering down some elements of the AI Act, undermining crucial oversight and transparency mechanisms (EDRI, 2025b). While the Digital Omnibus package has not yet been adopted, civil society and trade unions have collectively argued that it would represent “the biggest rollback of digital fundamental rights in EU history. It is being done under the radar, using rushed and opaque processes designed to avoid democratic oversight” (European Civic Forum, 2025). Relatedly, in the Commission’s 2025 AI Continent Action Plan, there is a reference to ‘Artificial General Intelligence (AGI)’, a longstanding myth mobilised by Silicon Valley leaders that refers to the possibility of AI reaching or even surpassing human capabilities. According to the Commission, AI models will become increasingly complex, “with the next generation of frontier AI models expected to unlock a leap in capabilities, towards Artificial General Intelligence (AGI) capable of tackling highly complex and diverse tasks, matching human capabilities.” (European Commission, 2025, p. 7). Hence, contemporary developments in EU regulation of AI are therefore heavily marked by attempts to develop AI at all costs and rooted in Silicon Valley AI technosolutionism, even if they are justified on the grounds of building European digital sovereignty.

The deregulatory Digital Omnibus package has been interpreted as a betrayal of

the EU's longstanding championing of digital rights. For instance, journalists at Politico Europe claimed that "Brussels is done being the world's digital policeman" (Haeck & O'Regan, 2025). However, our article emphasises something different: the Digital Omnibus package should be interpreted as an acceleration of the pre-existent technosolutionist approach to AI taken by the EU: a 'technosolutionism on steroids' (Flores Moleón & Oleart, 2026) that takes the larger pro-Big Tech trajectory visible in the present paper in an even more radical direction. Hence, the Digital Omnibus reflects more continuity than a rupture with the previous EU regulatory approach to AI, driven by the logic of 'innovation' and masked under the idea 'human-centrism', which we conceive as a form of political bullshit. In other words, the EU has never actually challenged the surveillance capitalist business model upon which Big Tech companies have been built. Indeed, our article shows how the AI Act and the GPAI Code of Practice crystallise new forms of privatising AI governance, where the relationship between EU regulators and the tech industry is closely entangled not only through public-private partnerships but also through the underlying sociotechnical imaginary, rooted in technosolutionism as an ideational paradigm. These partnerships materialise in practices such as the co-design of codes of conduct, privileged access to regulatory forums, and the ability of platforms to set standards in the form of "best practices", definitions, or technical processes. Hence, the article highlights how Big Tech companies consolidate themselves in EU regulatory approaches to AI, which in turn reduces the space for democratic intervention into AI (Garibay-Petersen et al., 2025) beyond 'technical' arguments. In our analysis, we observe how Big Tech companies have embedded themselves in sensitive areas such as crisis response, content moderation, the implementation of international sanctions, and the configuration of the emerging professional field of Trust & Safety.

However, it is not only relevant to analyse what we see in EU regulation, but also what is absent. A key dimension unaddressed by EU regulation (albeit mentioned in the Ethical Guidelines of the High-Level Expert Group on AI) is the materiality of AI. This is an inescapable dimension because part of the rationale for supporting AI is its capacity to enhance energy efficiency. Unfortunately, the reality is the exact opposite due to the physical and environmental needs of AI, as it requires vast amounts of servers hosted in gigantic data centres across the globe which have considerable energy and water demands (Rone, 2022). As the EU aims to triple the amount of data centres, and invest in AI Gigafactories, through the proposed Cloud and AI Development Act - which is part of the European Commission's (2026) "European Technological Sovereignty Package" -, the environmental impact and energy demands will only increase.

Overall, our analysis indicates how the EU regulatory approaches to AI mostly align with Silicon Valley technosolutionist imaginaries. In doing so, the EU is reproducing the epistemic and political power of Big Tech companies rather than challenging it. This is normatively problematic by itself, because it does not question the surveillance capitalist business model of these companies, and AI is the ultimate illustration of a model built on mass data extractivism. But the EU's approach is even more problematic due to the global aspirations of the AI Act and the increasing geopolitical logic that is driving EU tech policy (Bradford 2023; Farand 2026). Even if the AI Act is unlikely to become a global framework for AI (Crum 2025), attempts at conceiving the EU as a 'regulatory leader' are inherently problematic, particularly considering the "hidden human labor powering AI" (Cant et al., 2024), which is mostly located in the Global South. There is evidence that Big Tech companies exploit workers across the world - particularly from the Global South, as the case of Kenya illustrates, where workers operated in precarious conditions and earned less than two dollars per hour to train OpenAI's ChatGPT (Perri-go, 2023) - to perform tasks without which AI simply cannot exist, such as resource extraction, data annotation or content moderation. We should then think about AI from a decolonial perspective, one that accounts for the close intertwining of capitalism and new forms of colonialism when making sense of technology (see Couldry & Mejías, 2023), a perspective that is particularly relevant in the EU context considering the past and present of European colonialism and its close intertwinement with capitalism. It is therefore necessary to think of technology and democracy otherwise, a reflection that needs to start by countering the surveillance capitalist business model of Big Tech corporations and reversing its privatisation (Oleart & Rone, 2025). This is the starting point from which we can begin a critical reflection on the risks of technological capture and dependence on private corporations in order to build democratic and decolonial counterweights to Big Tech not just in the EU, but globally.

ACKNOWLEDGEMENTS

We are grateful to the three reviewers for this paper, Mathias Berguber, Paul Waller and Péter Mezei, as well as the Internet Policy Review editor Frédéric Dubois. We are grateful to them, as the paper has greatly benefitted from their excellent comments and suggestions.

References

Bareis, J., & Katzenbach, C. (2022). Talking AI into being: The narratives and imaginaries of national AI strategies and their performative politics. *Science, Technology, & Human Values*, 47(5), 855–881. <https://doi.org/10.1177/01622439211030007>

Bender, E. M., & Hanna, A. (2025). *The AI con: How to fight big tech's hype and create the future we want*. Random House.

Birhane, A., Angius, R., Agnew, W., Pandit, H. J., Mitra, B., Dobbe, R., & Talat, Z. (2026). Big AI's regulatory capture: Mapping industry interference and government complicity. *Proceedings of the 2026 ACM Conference on Fairness, Accountability, and Transparency*, 2586–2607. <https://doi.org/10.1145/3805689.3806740>

Bouza García, L., & Oleart, A. (2024). Regulating disinformation and big tech in the EU: A research agenda on the institutional strategies, public spheres and analytical challenges. *Journal of Common Market Studies*, 62(5), 1395–1407. <https://doi.org/10.1111/jcms.13548>

Bouza, L., Roch, J., & Oleart, Á. (2025). Civil society coalitions in the EU media regulation: When field theory meets network analysis. *European Politics and Society*, 26(2), 304–320. <https://doi.org/10.1080/23745118.2024.2420229>

Bradford, A. (2020). *The Brussels effect: How the European Union rules the world*. Oxford University Press.

Bradford, A. (2023). *Digital empires: The global battle to regulate technology*. Oxford University Press.

Calcara, A., Csernaton, R., & Lavallée, C. (Eds). (2020). *Emerging security technologies and EU governance: Actors, practices and processes* (1st edn). Routledge. <https://doi.org/10.4324/9780429351846>

Cant, C., Muldoon, J., & Graham, M. (2024). *Feeding the machine: The hidden human labor powering AI*. Bloomsbury Publishing USA.

Carlsson, V., & Rönblom, M. (2022). From politics to ethics: Transformations in EU policies on digital technology. *Technology in Society*, 71, 102145. <https://doi.org/10.1016/j.techsoc.2022.102145>

Castoriadis, C. (1987). *The imaginary institution of society*. MIT Press.

Chamberlain, J. (2023). The risk-based approach of the European Union's proposed artificial intelligence regulation: Some comments from a tort law perspective. *European Journal of Risk Regulation*, 14(1), 1–13. <https://doi.org/10.1017/err.2022.38>

Coeckelbergh, M. (2024). *Why AI undermines democracy and what to do about it*. John Wiley & Sons.

Coen, D., Katsaitis, A., & Vannoni, M. (2021). *Business lobbying in the European Union*. Oxford University Press.

Corporate Europe Observatory. (2025, October 29). Big tech lobby budgets hit record levels. *Corporate Europe Observatory*. <https://corporateeurope.org/en/2025/10/big-tech-lobby-budgets-hit-record-levels>

Couldry, N., & Mejias, U. A. (2023). The decolonial turn in data and technology research: What is at

stake and where is it heading? *Information, Communication & Society*, 26(4), 786–802. <https://doi.org/10.1080/1369118X.2021.1986102>

Crawford, K. (2021). *The atlas of AI: Power, politics, and the planetary costs of artificial intelligence*. Yale University Press.

Crum, B. (2025). Brussels effect or experimentalism? The EU AI Act and global standard-setting. *Internet Policy Review*, 14(3). <https://doi.org/10.14763/2025.3.2032>

EDRi. (2025a). *Press release: Commission's Digital Omnibus is a major rollback of EU digital protections*. EDRi. <https://edri.org/our-work/commissions-digital-omnibus-is-a-major-rollback-of-eu-digital-protections/>

EDRi. (2025b, November 19). *Why the Digital Omnibus puts GDPR and eprivacy at risk*. EDRi. <https://edri.org/our-work/why-the-digital-omnibus-puts-gdpr-and-eprivacy-at-risk/>

European Civic Forum. (2025). *Joint letter: The EU must uphold hard-won protections for digital human rights*. European Civic Forum. <https://civic-forum.eu/publications/open-letter/joint-letter-the-eu-must-uphold-hard-won-protections-for-digital-human-rights>

European Commission. (2018). *Artificial intelligence: A European perspective*. European Commission. https://ai-watch.ec.europa.eu/publications/artificial-intelligence-european-perspective_en

European Commission. (2019). *Ethics guidelines for trustworthy AI*. European Commission. <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>

European Commission. (2024a). *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) (Text with EEA relevance)*. European Commission. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>

European Commission. (2024b, September 30). *The general-purpose AI code of practice*. European Commission. <https://digital-strategy.ec.europa.eu/en/policies/contents-code-gpai>

European Commission. (2025). *AI continent action plan*. European Commission. <https://digital-strategy.ec.europa.eu/en/library/ai-continent-action-plan>

European Commission. (2026). *Communication on European tech sovereignty, accompanied by an EU open source strategy*. European Commission. <https://digital-strategy.ec.europa.eu/en/library/communication-european-tech-sovereignty-accompanied-eu-open-source-strategy>

Farrand, B. (2026). *Geopolitical union: Europe's attempt to take back control of technology regulation* (1st edn). Cambridge University Press. <https://doi.org/10.1017/9781009691093>

Flores Moleón, A. (2025). Socio-technical imaginaries of AI's role in the strengthened EU Code of Practice on Disinformation. *International Journal of Communication*, 19(2025), 1–21.

Flores Moleón, A., & Oleart, A. (2026). When AI-first becomes democracy-last: The European Commission's Digital Omnibus and its technosolutionism on steroids. *European Journal of Risk Regulation*, 1–15. <https://doi.org/10.1017/err.2026.10123>

Garibay-Petersen, C., Lorimer, M., & Menzat, B. (2025). Creating certainty where there is none: Artificial intelligence as political concept. *Big Data and Society*, 12(4). <https://orca.cardiff.ac.uk/id/eprint/181677/>

- Griffin, R. (2023). Public and private power in social media governance: Multistakeholderism, the rule of law and democratic accountability. *Transnational Legal Theory*, 14(1), 46–89. <https://doi.org/10.1080/20414005.2023.2203538>
- Haeck, P., & O'Regan, E. (2025, November 25). *Brussels is done being the world's digital policeman*. <https://www.politico.eu/article/brussels-police-world-digital-tech-us-china-regulations/>
- Hansen, P., & Jonsson, S. (2014). *Eurafrica: The untold history of European integration and colonialism*. Bloomsbury Academic.
- Hao, K. (2025). *Empire of AI: Dreams and nightmares in Sam Altman's OpenAI*. Penguin Press.
- Heermann, M. (2024). Undermining lobbying coalitions: The interest group politics of EU copyright reform. *Journal of European Public Policy*, 31(8), 2287–2315.
- Jasanoff, S., & Kim, S.-H. (2009). Containing the atom: Sociotechnical imaginaries and nuclear power in the United States and South Korea. *Minerva*, 47(2), 119–146. <https://doi.org/10.1007/s11024-009-9124-4>
- Jasanoff, S., & Kim, S.-H. (Eds.). (2015). *Dreamscapes of modernity: Sociotechnical imaginaries and the fabrication of power*. University of Chicago Press.
- Kausche, K., & Weiss, M. (2025). Platform power and regulatory capture in digital governance. *Business and Politics*, 27(2), 284–308. <https://doi.org/10.1017/bap.2024.33>
- Markham, A. (2021). The limits of the imaginary: Challenges to intervening in future speculations of memory, data, and algorithms. *New Media & Society*, 23(2), 382–405. <https://doi.org/10.1177/1461444820929322>
- Merchant, B. (2023). *Blood in the machine: The origins of the rebellion against big tech*. Hachette UK.
- Morozov, E. (2013). *To save everything, click here: The folly of technological solutionism*. PublicAffairs.
- Muldoon, J. (2022). *Platform socialism. How to reclaim our digital future from big tech*. Pluto Press.
- Navarro, J. T., García, L. B., & Oleart, A. (2025). How the EU counters disinformation: Journalistic and regulatory responses. *Media and Communication*, 13. <https://doi.org/10.17645/mac.10551>
- Obendiek, A. S., & Seidl, T. (2023). The (false) promise of solutionism: Ideational business power and the construction of epistemic authority in digital security governance. *Journal of European Public Policy*, 30(7), 1305–1329. <https://doi.org/10.1080/13501763.2023.2172060>
- Oleart, A., & Flores Moleón, A. (2026). Why the EU's technosolutionist focus on AI and media 'literacy' empowers big tech: Centering structural approaches to counter the undemocratic political economy of surveillance capitalism. *Social Sciences*, 15(7), 430. <https://doi.org/10.3390/socsci15070430>
- Oleart, A., & García, L. B. (2025). From self to co-regulation in the EU's approach to disinformation: The framing power of big tech business lobbies in the lead to the Digital Services Act. *International Review of Public Policy*, 7(2), 235–256. <https://doi.org/10.4000/14rse>
- Oleart, A., & Rone, J. (2025). Reversing the privatisation of the public sphere: Democratic alternatives to the EU's regulation of disinformation. *Media and Communication*, 13. <https://doi.org/10.17645/mac.9496>
- Oleart, A., Van Der Waal, M., & Van Weyenberg, A. (2025). (Un)learning 'Europe' as decolonial practice. *Politique Européenne*, 88(2), 6–28. <https://doi.org/10.3917/e.poeu.088.0006>

Perrigo, B. (2023, January 18). Exclusive: OpenAI used Kenyan workers on less than \$2 per hour to make ChatGPT less toxic. *Time*. <https://time.com/6247678/openai-chatgpt-kenya-workers/>

Petit, P., & Oleart, A. (2026). Citizenwashing EU tech policy: EU deliberative mini-publics on virtual worlds and artificial intelligence. *Politics and Governance*, 14, 10468. <https://doi.org/10.17645/pag.10468>

Popiel, P. (2018). The tech lobby: Tracing the contours of new media elite lobbying power. *Communication, Culture and Critique*, 11(4), 566–585. <https://doi.org/10.1093/ccc/tcy027>

Rieder, G. (2018). Tracing big data imaginaries through public policy: The case of the European Commission. In A. R. Sætnan, I. Schneider, & N. Green (Eds), *The politics and policies of big data* (pp. 89–109). Routledge.

Rodney, W. (2018). *How Europe underdeveloped Africa*. Verso Books.

Rone, J. (2022). The politics of data infrastructures contestation: Perspectives for future research. *Journal of Environmental Media*, 3(2), 207–214. https://doi.org/10.1386/jem_00086_1

Rone, J. (2024). The shape of the cloud: Contesting data centre construction in North Holland. *New Media & Society*, 26(10), 5999–6018. <https://doi.org/10.1177/14614448221145928>

Schaake, M. (2024). *The tech coup: How to save democracy from Silicon Valley*. Princeton University Press.

Shrishak, K. (2025, November 10). Von der Leyen caught in AI hype trap while tech bros cash in. *EU Observer*. <https://euobserver.com/digital/ar68689659>

Smuha, N. A. (2024). *The work of the high-level expert group on AI as the precursor of the AI Act*. SSRN. <https://dx.doi.org/10.2139/ssrn.5012626>

Taylor, C. (2020). What is a “social imaginary”? In D. P. Gaonkar, J. Kramer, B. Lee, & M. Warner (Eds), *Modern social imaginaries* (pp. 23–30). Duke University Press. <https://doi.org/10.1515/9780822385806-004>

Vertesi, J., boyd, danah, Taylor, A., & Shestakofsky, B. (2026). *Reckoning with the political economy of AI: Avoiding decoys in pursuit of accountability* (Version 2). arXiv. <https://doi.org/10.48550/ARXIV.2604.16106>

Waller, P. (2025). *A critique of government policies built on sociotechnical imaginaries—Nightmare of the imaginaries 2*. SSRN. <https://doi.org/10.2139/ssrn.5230465>

Zuboff, S. (2019). *The age of surveillance capitalism: The fight for a human future at the new frontier of power*. PublicAffairs.

Published by



ALEXANDER VON HUMBOLDT
INSTITUTE FOR INTERNET
AND SOCIETY



RESEARCH
FOR THE
DIGITAL AGE

in cooperation with



CREATE



centre
— internet
et société



R&I
IN3
Internet
interdisciplinary
Institute
Universitat Oberta de Catalunya



UNIVERSITY OF TARTU
Johan Skytte Institute of
Political Studies