

From Biden to Trump: Divergent and Convergent Policies in The Artificial Intelligence (AI) Summer

di Valerio Lubello

Abstract: *Da Biden a Trump: approcci divergenti e convergenti durante nella c.d. “AI Summer”*

- This paper explores the regulatory evolution of artificial intelligence (AI) policy in the United States, focusing on the contrasting yet occasionally convergent approaches of the Biden and Trump administrations during the so-called “AI Summer.” It begins by reconstructing the two central phases of the Biden administration’s strategy: the AI Bill of Rights Phase, centred on the Blueprint for an AI Bill of Rights, and the subsequent To-Do List Phase, characterized by a series of executive orders and agency guidelines aimed at operationalizing ethical principles in real-world AI deployment. These efforts address core rights-related concerns such as privacy, algorithmic discrimination, transparency, and the human oversight of automated systems. The paper highlights the Biden administration’s cross-sectoral focus, which includes national security, labour rights, healthcare, criminal justice, and environmental sustainability—most notably through Executive Order 14141, which tightly links AI infrastructure development to clean energy investments and grid modernization.

The second part of the paper examines the early actions of the Trump administration in 2025, particularly the revocation of Executive Order 14110 and the issuance of new deregulatory measures under the banner of restoring American AI leadership. While this marks a shift toward a more market-driven and innovation-centric approach, the paper notes that several Biden-era instruments—such as the National Security Memorandum and EO 14141—remain in force and continue to shape federal agency activities. Through this comparative lens, the article assesses the extent to which foundational human rights protections, risk-based governance models, and sectoral guidelines can withstand political transitions and contribute to a durable framework for responsible AI governance in the United States.

Keywords: Biden; Trump; Artificial Intelligence; AI; Rights and liberties

1. Introduction

“If misused, AI could threaten United States national security, bolster authoritarianism worldwide, undermine democratic institutions and processes, facilitate human rights abuses, and weaken the rules-based international order”.¹

¹ See the *Memorandum on Advancing the United States’ Leadership in Artificial Intelligence; Harnessing Artificial Intelligence to Fulfill National Security Objectives; and Fostering the Safety, Security, and Trustworthiness of Artificial Intelligence* adopted by President Biden,, October 24, 2024, Sec. 1, (c) available at the following URL: <https://www.whitehouse.gov/briefing-room/presidential-actions/2024/10/24/memorandum-on-advancing-the-united-states-leadership-in->

One of the first demonstrations of the AI capabilities is connected to the epic chess game between Deep Blue and Kasparow. And as well recognised among scholars, the actual s.c. “AI summer” consists mainly in the fact that everybody can develop his/her own AI for his/her own scope and it is not necessary anymore to have a single target, such as playing Chess for Deep Blue. Everybody can design his/her own AI for a different scope.²

From this perspective, we are at the early stages of history considering that Chat GPT (Generative Pre-trained Transformer) arrived on our devices in November 2022.

This acceleration process is not so recent. It exists at least from the moment in which we invented the transistor, and we started to apply its Moore Law, according to which the capacity of our IT infrastructure doubles every eighteen months.

On the other hand, thanks to the proliferation of new AI services and research, different countries around the world feel the need to set new regulatory boundaries, with the ambitious goal of finding the right approach for all possible uses of AI. The AI regulation race has started with different but somehow convergent approaches in different areas of the globe.

Focusing the debate on the field of human rights and freedom it is in some ways self-evident how the AI large scale diffusion is theoretically capable of redefining the burdens of our contemporary bill of rights. Right to life, human dignity, habeas corpus, right to auto-determination, privacy, right to health, labour and justice rights and of course environmental rights are all connected and influenced by the actual level of technologies and for this by AI quickly spread in our daily lives.

Against this evolving backdrop, this essay seeks to highlight the different approaches of the Biden and Trump administrations to AI. Starting with an analysis of the different phases of the Biden Administration (Par. 2), the analysis provides an overview of the first consequences of the Trump policy. As will be described in the following paragraphs, the revocation of E.O. Executive Order Trustworthy Development and Use of Artificial Intelligence of October 30, 2023 is only a part of the whole framework. From this perspective, it is possible to argue that some key Biden legacies seem to be still effective.

2. The Biden approach to the AI and the acts that have a direct impact on rights and liberties: The bill of rights phase and the to do list phase

The AI Biden Administration approach follows a cross sectoral impact, avoiding enthusiastic or despotic scenarios, without refusing – at the same

[artificial-intelligence-harnessing-artificial-intelligence-to-fulfill-national-security-objectives-and-fostering-the-safety-security/](#)

² J. Hawkins, *A Thousand Brains. A New Theory of Intelligence*, New York, 2022. Y. N. Harari, *Nexus, A brief History of Information Network from the Stone Age to AI*, London, 2024. H. Kissinger, E. Schmidt, D. Huttendlocher, *The Age of the AI* London, 2021.

time – to address the global debate toward new directions,³ maintaining a certain continuity with the former Obama⁴ and Trump Administration⁵.

With a specific focus to rights, it is possible to summarise the whole Biden approach in two macro phases.

The first phase, hereinafter the “AI bill of right phase”, is well enshrined by the adoption of the Blueprint for an AI Bill of rights approved in October 2022, and the second one, hereinafter “To do list phase”, enshrined in a bunch of executive orders and guidelines that tried to put in place the next steps of the American democracy in the field of the AI. Namely: the Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence of October 30, 2023, Executive Order on Preventing Access to Americans’ Bulk Sensitive Personal Data and United States Government-Related Data by Countries of Concern February 28, 2024⁶ and, and the (first-ever) National Security Memorandum (NSM) on Artificial Intelligence (AI)⁷ and the Risk

³ For an overview of the evolution of the USA Approach to AI see E. Hine and L. Floridi, *Artificial Intelligence with American Values and Chinese Characteristics: A Comparative Analysis of American and Chinese Governmental AI Policies* (January 11, 2022). AI & Soc (2022). Available at SSRN: <https://ssrn.com/abstract=4006332>. On the Biden Executive Order see N.A. Smuha, *Biden, Bletchley, and the emerging international law of AI*, VerfBlog, 2023/11/15, available at: <https://verfassungsblog.de/biden-bletchley-and-the-emerging-international-law-of-ai/>. For a comparative perspective see M. Wörsdörfer, *Biden’s Executive Order on AI and the E.U.’s AI Act: A Comparative Computer-Ethical Analysis*, Philosophy & Technology, Volume 37, 74 (2024) and A. Engler, *The EU and U.S. diverge on AI regulation: A transatlantic comparison and steps to alignment*, The Brookings Institution, Research Paper available at <https://www.brookings.edu/articles/the-eu-and-us-diverge-on-ai-regulation-a-transatlantic-comparison-and-steps-to-alignment/>. For the EU changing ecosystem see F.P. Levantino, F. Paolucci, *Advancing the Protection of Fundamental Rights Through AI Regulation: How the EU and the Council of Europe are Shaping the Future*, in *European Yearbook on Human Rights 2024*, (ed.) by P. Czech, L. Heschl, K. Lukas, M. Nowak, and G. Oberleitner, Leiden, 2024. For a wider overview, see O. Pollicino, P. Dunn, *Intelligenza Artificiale e Democrazia*, Milano, 2024 and (ed.) G.C. Feroni, E. Raffiotta, C. Fontana, *AI Anthology. Profili giuridici, economici e sociali dell'intelligenza artificiale*, Bologna, 2022.

⁴ See E.O. 13859, 2019 “*Maintaining American Leadership in the Artificial Intelligence*” and *Artificial Intelligence for the American People* 2018, available at <https://trumpwhitehouse.archives.gov/briefings-statements/artificial-intelligence-american-people/> and the E.O. 13960, 2020 “*Promoting the Use of Trustworthy Artificial Intelligence in the Federal Government*”.

⁵ See Executive Order 13960, *Promoting the Use of Trustworthy AI in the Federal Government*, December 2020, the *AI in Government Act*, September 2020 and the Executive Order 13859, *Maintaining American Leadership in AI*, February 2019.

⁶ See the *Executive Order on Preventing Access to Americans’ Bulk Sensitive Personal Data and United States Government-Related Data by Countries of Concern*, adopted by President Biden, February 28, 2024 which is available at the following URL: <https://www.whitehouse.gov/briefing-room/presidential-actions/2024/02/28/executive-order-on-preventing-access-to-americans-bulk-sensitive-personal-data-and-united-states-government-related-data-by-countries-of-concern/>

⁷ See the *Memorandum on Advancing the United States’ Leadership in Artificial Intelligence; Harnessing Artificial Intelligence to Fulfill National Security Objectives; and Fostering the*

Management Profile for Artificial Intelligence and Human Rights adopted in July 2024 by the Bureau of Cyberspace and Digital Policy.⁸

2.1 AI Bill of rights

The Blueprint for an AI Bill of Rights⁹ is a set of guidelines that strongly suggests some principles and best practices in the application of AI systems.

The aim of this pivotal document is to guarantee the spread of the AI systems in accordance with human rights and democratic values, keeping in mind the risks connected to an unethical use of them.

The document, promoted by the Biden Administration together with the White House Office of Science and Technology Policies, is a first tentative to summarise some basic principles in the design, use and deployment of such systems.

It is a very high-level document that should be read together with the initiatives of every single Department, such as for example those activated

Safety, Security, and Trustworthiness of Artificial Intelligence adopted by President Biden, October 24, 2024, Sec. 1, (c).

⁸ See the *Risk Management Profile for Artificial Intelligence and Human Rights* adopted in July 2024 by the Bureau of Cyberspace and Digital Policy which is available at the following URL:

<https://www.state.gov/risk-management-profile-for-ai-and-human-rights/>.

⁹ E. Hine, L. Floridi, *The Blueprint for an AI Bill of Rights: In Search of Enaction, at Risk of Inaction* (November 2, 2022). *Minds and Machines*, 2023., Available at SSRN: <https://ssrn.com/abstract=4279449>. See also A. Oesterling, U. Bhalla, S. Venkatasubramanian, H. Lakkaraju, *Operationalizing the Blueprint for an AI Bill of Rights: Recommendations for Practitioners, Researchers, and Policy Makers*, *rXiv*, available at: <https://arxiv.org/abs/2407.08689>.

by the Department of Energy (DOE),¹⁰ the Department of Defence¹¹ and the U.S. Intelligence Community (IC).¹²

The Blueprint for an AI Bill of Rights sets out five macro principles that represent a new milestone in the debate on AI regulation: 1) Safe and Effective Systems 2) Algorithmic Discrimination Protections 3) Data Privacy 4) Notice and Explanation 5) Human alternatives, Consideration and Fallback.

The first principle, Safe and Effective Systems, affirms in a nutshell that “Automated systems should be developed with consultation from diverse communities, stakeholders, and domain experts to identify concerns, risks, and potential impacts of the system”.

The above-mentioned principle further provides that AI systems “should be designed to proactively protect you from harms stemming from unintended, yet foreseeable, uses or impacts of automated systems”.

It is in somehow a suggestive analogy to Asimov's First Law of Robotics from the 1940s: “A robot may not injure a human being or, through inaction, allow a human being to come to harm” and its complementary but later Zero Law “A robot may not harm humanity, or, by inaction, allow humanity to come to harm”.¹³

¹⁰ The DOE adopted already two versions of the *Department of Energy (DOE) Generative Artificial Intelligence (GenAI) Reference Guide*, the last version has been adopted in July 2024 and it is available at the following URL: <https://www.energy.gov/cio/departments-energy-generative-artificial-intelligence-reference-guide>. The above mentioned document activated the AI Advancement Council to oversee coordination and advise on the implementation of the DOE AI Strategy.

¹¹ The Department of Defence adopted its *Artificial Intelligence Ethical Principles* in 2020. In a nutshell, the document provides five principles: 1) Responsible. «DOD personnel will exercise appropriate levels of judgment and care while remaining responsible for the development, deployment and use of AI capabilities; 2) Equitable. The department will take deliberate steps to minimize unintended bias in AI capabilities. 3) Traceable. The department's AI capabilities will be developed and deployed such that relevant personnel possess an appropriate understanding of the technology, development processes and operational methods applicable to AI capabilities, including with transparent and auditable methodologies, data sources and design procedures and documentation. 4) Reliable. The department's AI capabilities will have explicit, well-defined uses, and the safety, security and effectiveness of such capabilities will be subject to testing and assurance within those defined uses across their entire life cycles. 5) Governable. The department will design and engineer AI capabilities to fulfil their intended functions while possessing the ability to detect and avoid unintended consequences, and the ability to disengage or deactivate deployed systems that demonstrate unintended behaviour». See the following URL: <https://www.defense.gov/News/News-Stories/article/article/2094085/dod-adopts-5-principles-of-artificial-intelligence-ethics/>

¹² The US Intelligence Community since 2020 has developed the *Principles of Artificial Intelligence Ethics for the Intelligence Community* to guide personnel on whether and how to develop and use AI in furtherance of the IC's mission, as well as an *AI Ethics Framework* to help implement these principles. See [https://www.intelligence.gov/images/AI/Principles of AI Ethics for the Intelligence Community.pdf](https://www.intelligence.gov/images/AI/Principles%20of%20AI%20Ethics%20for%20the%20Intelligence%20Community.pdf) and [https://www.intelligence.gov/images/AI/AI Ethics Framework for the Intelligence Community 1.0.pdf](https://www.intelligence.gov/images/AI/AI%20Ethics%20Framework%20for%20the%20Intelligence%20Community%201.0.pdf).

¹³ I. Asimov, *I, Robot*, West Hartford, 1952.

Beyond suggestions, the broad scope of this principle is clearly intended to reduce the risk of the s.c. group fairness implementing some guarantees in the design and training phases, which should be developed and maintained with a clear overview of the stakeholders and the impacted communities.

The second principle, Algorithmic Discrimination Protections, affirms that individuals should not face discrimination by algorithms and systems should be used and designed in an equitable way: “Algorithmic discrimination occurs when automated systems contribute to unjustified different treatment or impacts disavouring people based on their race, colour, ethnicity, sex (including pregnancy, childbirth, and related medical conditions, gender identity, intersex status, and sexual orientation), religion, age, national origin, disability, veteran status, genetic information, or any other classification protected by law”.

As well enshrined in the principle itself “Designers, developers, and deployers of automated systems should take proactive and continuous measures to protect individuals and communities from algorithmic discrimination and to use and design systems in an equitable way”.

As widely recognised, potential cognitive biases represent one of the main risks connected to AI in our days and regulators are trying to limit these prejudices toward a growing attention to the AI design and training stages.¹⁴

The third principle, Privacy, provides that “individuals should be protected from violations of privacy through design choices that ensure such protections are included by default, including ensuring that data collection conforms to reasonable expectations and that only data strictly necessary for the specific context is collected”.

Words that – *mutatis mutandis* – mirror some EU GDPR principles such as privacy by design, privacy by default, transparency and proportionality principles.

Similarly, “Consent should only be used to justify collection of data in cases where it can be appropriately and meaningfully given. Any consent requests should be brief, be understandable in plain language, and give you agency over data collection and the specific context of use”.

Interesting convergences with the EU approach with the aim to solve some privacy dilemmas behind the widespread use of AI systems: first, it is not easy to remove some information from a trained machine; it is possible

¹⁴ E. Loza de Siles, *Artificial Intelligence Bias and Discrimination: Will We Pull the Arc of the Moral Universe Towards Justice?* (December 1, 2021), *J. Int'l & Comp. L.*, Vol. 8, No. 2, 2021, Available at SSRN: <https://ssrn.com/abstract=4002486> and B. Braunschweig, M. Ghallab (ed.), *Reflections on Artificial Intelligence for Humanity*, Springer 2021, and in this Review see C. Casonato, *L'intelligenza artificiale e il diritto pubblico comparato ed europeo*, *DPCE Online* 1-2022 available at <https://www.dpceonline.it/index.php/dpceonline/article/view/1566> and M. Fasan, *I principi costituzionali nella disciplina dell'Intelligenza Artificiale. Nuove prospettive interpretative*, *DPCE Online* 1-2022, available at <https://www.dpceonline.it/index.php/dpceonline/article/view/1567>. For an EU overview see: C. Nardocci, *Artificial Intelligence-based Discrimination: Theoretical and Normative Responses. Perspectives from Europe*, *DPCE Online*, 3-2023, available at the following URL: <https://www.dpceonline.it/index.php/dpceonline/article/view/1981>.

to update a certain dataset, but this does not mean that we permanently erase the first dataset. From this perspective, the concept of synthetic data is becoming obsolete, despite the debate about it, and its regulation seems to be at an early stage. For the same reason, rights to be forgotten or rights to have a correct representation of yourself could be strongly compromised by the spread of such technologies.

Debates about privacy and digital identity are rapidly entering a new dimension in which it is possible, at least in theory, to have a digital copy of certain characteristics of a single natural person.¹⁵

Complete control and self-determination over individuals' digital projections no longer seems possible, at least with the legal and technological tools we have today. Perhaps technologies such as blockchain could help to introduce new ways of effectiveness in the future.¹⁶

The fourth principle is the s.c. Notice and Explanation, according to which “individuals should know that an automated system is being used, and understand how and why it contributes to outcomes that impact you”.

It is something similar to the information principles under the GDPR and EU AI Act in case your data is processed with an automatic system. The ambitious aim is to explain to the stakeholders what is going to happen in the enormous databases trying to calibrate the risk based on the context.

The fifth principle, Human alternatives, Consideration and Fallback provides that “You should be able to opt out, where appropriate, and have access to a person who can quickly consider and remedy problems you encounter”. This principle is composed of two crucial elements that are strongly influencing the global AI debate. The first one is connected to the fact that there is an opt-out right but, once again, it is not easy at all to consider it as effective. The second element is a kind of human touch in the use of the AI which still remains crucial in several applications such as for example in the health, employment and education fields.

2.2 The “To-do list phase”

The second pragmatic phase¹⁷ has been anticipated by the relevant document called From Principle to Practice which provides some design

¹⁵ From a sociological perspective see M. Suleyman, *The Coming Wave, Technology, Power, and the Twenty-first Century's Greatest Dilemma*, New York, 2023.

¹⁶ A.J. Zwitter, O.J. Gstrein, E. Yap, *Digital Identity and the Blockchain: Universal Identity Management and the Concept of the “Self-Sovereign” Individual*, *Frontiers in Blockchain*, 3-2020 available at the following URL:

<https://www.frontiersin.org/journals/blockchain/articles/10.3389/fbloc.2020.00026/full>

For a wider overview see: O. Pollicino, G. De Gregorio (ed.), *Blockchain and Public Law, Global Challenges in the Era of Decentralisation*, Cheltenham, 2021.

¹⁷ For a first reading, M. Bassini, *The Global Race to Regulate AI: Biden's Executive Order Spillover Effects on the EU AI Act*, IEP@BU, available at <https://iep.unibocconi.eu/publications/global-race-regulate-ai-bidens-executive-order-spillover-effects-eu-ai-act>. For a wider analysis see also C. Sbailò, *Governing Artificial Intelligence: Technological Leadership and Regulatory Challenges in an Era of Exponential Growth*, DPCE Online, SP3, Biden Special Issue, available at: <https://www.dpceonline.it/index.php/dpceonline/article/view/2354>

solution and process to respect the above mentioned five principles. The document considers each principle in the Blueprint for an AI Bill of Rights, providing examples and concrete steps for communities, industry, governments, and others to take to build these protections into policy, practice, or the technological design process.¹⁸

This phase furtherly shaped with the pivotal Executive Order Trustworthy Development and Use of Artificial Intelligence of October 30, 2023 and the Executive Order on Preventing Access to Americans' Bulk Sensitive Personal Data and United States Government-Related Data by Countries of Concern of February 28, 2024.

These orders clearly indicate the relevant elements of this “beta phase” with regard to the development and diffusion of AI systems on a large scale¹⁹ and the capability of the AI to be a transversal topic with effects in every aspect of contemporary society.

Following this path, the AI Executive Order's efforts are mainly directed at confronting a new form of the old - but still current - problem of the S.C. digital divide.²⁰

The E.O. is also oriented to well understand the state of the art of the different programs supported by the Artificial intelligence itself, and at the same time there is a clear aim to spread its use in everyday democratic life.²¹

The spectrum of the subjects involved at this scope is wide, and well supported by a consistent flow of data between different private and public subjects, including Agencies, Universities, Health institutions, libraries and ad hoc Task force such as the National AI Research Resource, Patent and Trade Mark Office and Trade Commission.

At the same time, the US is trying to attract talents in this new discipline with the scope to maintain and develop a sort of knowledge leadership in the field. Leadership that has been recognized and affirmed also toward openness with respect to “international allies and partners”.

¹⁸ These are the main directives at stake: 1) Safe and secure AI systems; 2) Unlock technology potential; 3) Support American workers; 4) - Equity and civil rights; 5) - Protection of the consumers; 6) Privacy and civil liberties; 7) Advancing Equity and Civil Rights; 8) Advancing Federal Government Use of AI.

¹⁹ Already under the Trump administration the Executive Order 13960 of 3 December 2020 on *Promoting the Use of Trustworthy Artificial Intelligence in the Federal Government* provided some principles when designing, developing, acquiring, or using AI for purposes other than national security or defence <https://www.federalregister.gov/documents/2020/12/08/2020-27065/promoting-the-use-of-trustworthy-artificial-intelligence-in-the-federal-government>.

These principles — while taking into account the sensitive law enforcement and other contexts in which the federal government may use AI, as opposed to private sector use of A — require that AI is: (a) lawful and respectful of our Nation's values; (b) purposeful and performance-driven; (c) accurate, reliable, and effective; (d) safe, secure, and resilient; (e) understandable; (f) responsible and traceable; (g) regularly monitored; (h) transparent; and, (i) accountable.

²⁰ See Sec. 4.1. *Developing Guidelines, Standards, and Best Practices for AI Safety and Security*.

²¹ Still on this new form of digital divide is also worth mentioning the E.O. *AI Training for the Acquisition Workforce Act*, adopted in October 2022 which mandates the implementation of an AI training program for designated personnel in the federal government.

From the point of view of the AI market, the E.O. aims to avoid any kind of monopolistic scenario, promoting competition among the different players and the different services²². In this light, the Federal Trade Commission has the scope to ensure fair competition in the AI marketplace in order to protect consumers and workers from the harms connected with the use of the AI. Keeping well in mind the ancillary and necessary superconductor industry that should be in some way supported to maintain high level of competitiveness²³.

2.2.1 Access to information

The New York Times case²⁴ demonstrated the necessity of supplementing AI systems with accurate and reliable information to mitigate the risk of generating misleading or erroneous outputs, commonly referred to as "hallucinations."

As well known, the New York Time sued Open AI on the ground that the AI systems have been trained with the journal archives, opening the ongoing copyright conflict between the AI platforms and content creators²⁵.

As a consequence, at this stage of the AI evolution, maintaining high-quality datasets involves significant costs, even though the end-user may not be human.²⁶ Feeding unreliable or distorted data into AI – akin to giving it "magic mushrooms" – would inevitably compromise outputs. Therefore, it is essential to ensure that the libraries and datasets used for AI training are rigorously validated and protected from potential threats, such as malware, that could corrupt the training process.

This concern is directly linked to broader debates around issues like deep fakes and fake news, which fundamentally revolve around the challenge of disinformation. Ensuring the integrity of the information ecosystem is crucial to maintaining trust in AI outputs and preventing harmful misuse of the technology.

This is very clear to the Biden administration which is completely aware about the risks which are at stake for the democracy itself, as well enshrined in the Bulk Data Executive Order²⁷ "These risks may be exacerbated when countries of concern use bulk sensitive personal data to develop AI capabilities and algorithms that, in turn, enable the use of large datasets in increasingly sophisticated and effective ways to the detriment of

²² See par. 5.3

²³ Par. 5.3, Promoting Competition, lett. b)

²⁴ A. Pope, *NYT v. OpenAI: The Times's About-Face*, April 20, 2024, available at: <https://harvardlawreview.org/blog/2024/04/nyt-v-openai-the-times-about-face/>

²⁵ For an overview, O. Pope, *NYT v. OpenAI: The Times's About-Face*, Harvard Law Review Blog, April 10, 2024, available at the following link: <https://harvardlawreview.org/blog/2024/04/nyt-v-openai-the-times-about-face/>

²⁶ <https://www.nytimes.com/2023/12/27/business/media/new-york-times-open-ai-microsoft-lawsuit.html>

²⁷ J.E. Stiglitz, *A Big Defeat for Big Tech, Project Syndicate*, March 18 2024, available at the following URL: <https://www.project-syndicate.org/commentary/big-tech-how-it-blocks-democratic-processes-to-serve-itself-by-joseph-e-stiglitz-2024-03>.

United States national security”. Similarly, “Countries of concern can use AI to target United States persons for espionage or blackmail by, for example, recognizing patterns across multiple unrelated datasets to identify potential individuals whose links to the Federal Government would be otherwise obscured in a single dataset”.

2.2.2 Right to Fair and Decent Work

The right to Fair and Decent Work has been widely covered by the AI Executive Order and the efforts are oriented in the adoption of measures and principles for employees that “could be used to mitigate AI’s potential harms to employees’ well being and maximize its potential benefits” and these guidelines should consider, at least: “(A) job-displacement risks and career opportunities related to AI, including effects on job skills and evaluation of applicants and workers; (B) labor standards and job quality, including issues related to the equity, protected-activity, compensation, health, and safety implications of AI in the workplace; and (C) implications for workers of employers’ AI-related collection and use of data about them, including transparency, engagement, management, and activity protected under worker-protection laws”.²⁸

These principles have been further developed in the Document adopted by the U.S. Department of Labor, Artificial Intelligence and worker well-being - Principles and Best Practices for Developers and Employers.²⁹

This document is in not mandatory and cannot be “intended as a substitute for existing or future federal or state laws and regulations”³⁰ but it introduces a bounce of standards usable at the different levels of AI usage for both developers and employers: 1) Ethically developing AI, according to which “AI systems should be designed, developed, and trained in a way that protects workers”; 2) Establishing AI Governance and human oversight which requires that “Organizations should have clear governance systems, procedures, human oversight, and evaluation processes for AI systems for use in the workplace”; 3) Ensuring transparency in AI use, according to which “Employers should be transparent with workers and job seekers about the AI systems that are being used in the workplace”; 4) Protecting labour and employment rights, according to which “AI systems should not violate or undermine workers’ right to organize, health and safety rights, wage and hour rights, and anti- discrimination and anti-retaliation protections”; 5) The principle Using AI to enable workers provides that “AI systems should assist, complement, and enable workers, and improve job quality”; 6) Supporting workers impacted by AI affirms that employers should “support and upskill workers during job transitions related to AI”; 7) Ensuring responsible use of worker data, according to which “Workers’ data collected,

²⁸ For an overview see A. Seth, G. Racabi, *Varieties of AI Regulations: The United States Perspective*, 77 ILR Review 799 (2024), Available at SSRN: <https://ssrn.com/abstract=4980643>

²⁹ Adopted by the US Department of Labor May 16, 2024 and available at the following URL:

<https://www.dol.gov/sites/dolgov/files/general/ai/AI-Principles-Best-Practices.pdf>

³⁰ Principles and Best Practices for Developers and Employers, 4.

used, or created by AI systems should be limited in scope and location, used only to support legitimate business aims, and protected and handled responsibly”.

2.2.3 Criminal Justice system

Also in the field of the Criminal Justice system there is a certain need to understand the state of the Art so the President asked to the Attorney General to map out all the uses that are at stake in the criminal justice system, such as sentencing, police surveillance, crime forecasting and predictive policing, including the incorporation of historical crime data into AI systems to predict high-density “hot spots”.

In addition, it also considers the use of AI in prison management tools and forensic analysis, providing insights into its current applications and implications for justice, equity and civil liberties³¹.

2.2.4 Healthcare and Human Services

In the sensitive field of Healthcare and Human Services, the Executive Order emphasises the safe, equitable and effective integration of AI to improve delivery, reduce administrative burdens and safeguard patient outcomes.

A key directive is the establishment of an HHS Task Force on AI to develop strategic guidelines for the responsible use of AI in various applications, including quality measurement, programme integrity and patient experience. Priorities in this area include long-term safety and performance monitoring, fair use through bias mitigation, and robust privacy and security standards.

In addition, the Order mandates strategies for AI quality assurance, federal compliance with anti-discrimination laws and a central framework for tracking and analysing AI-related clinical errors. It also lays the groundwork for a regulatory strategy to oversee AI in drug development, ensuring that its use is consistent with public safety and innovation goals.

In this delicate scenario, collaboration with state and local agencies is encouraged to share best practices, while specialised documentation ensures safe implementation in different contexts.

All of these efforts are aimed collectively at harnessing the transformative potential of AI in the HHS sector, with the intention of minimising the risks to patients and caregivers.

³¹ See Sec. 7.1, *Strengthening AI and Civil Rights in the Criminal Justice System*. According to which the Attorney General shall share with the President, among other: the use of AI in the criminal justice system, including any use in: “(A) sentencing; (B) parole, supervised release, and probation; (C) bail, pretrial release, and pretrial detention; (D) risk assessments, including pretrial, earned time, and early release or transfer to home-confinement determinations; (E) police surveillance; (F) crime forecasting and predictive policing, including the ingestion of historical crime data into AI systems to predict high-density “hot spots”; (G) prison-management tools; and (H) forensic analysis”.

2.2.5 National Security

In the context of National Security Systems (NSS), there is the need to strike a balance between AI-enabled national security activities and the protection of human rights, civil rights, civil liberties, privacy, and security.³²

To this end, the above-mentioned memorandum Memorandum on Advancing the United States' Leadership in Artificial Intelligence; Harnessing Artificial Intelligence to Fulfil National Security Objectives; and Fostering the Safety, Security, and Trustworthiness of Artificial Intelligence adopted by President Biden, October 24, 2024³³ – which shapes the AI E.O. directives – assigns specific AI actions to every single Department. To this effect, the Department of Defense (DOD) shall develop, test and integrate AI into national security systems, ensuring its responsible use. Similar actions are assigned to the Department of Commerce, with the AI Safety Institute (AIS) playing a central role in ensuring the safety, security, and trustworthiness of AI systems.

The Department of Homeland Security has a specific role in mitigation of AI risk to critical infrastructure and guides cybersecurity practices for AI systems through the Cybersecurity and Infrastructure Security Agency (CISA).³⁴

Also noteworthy is the growing but symptomatic attention to the connections between the safety and security of the AI systems and the Energy infrastructure. This is a consequence of the digitalization of the electrical infrastructure and the high energy consumption associated with AI.

In this light, the Department of Energy (DOE) coordinates efforts to streamline permitting and approvals for AI-enabling infrastructure, ensuring that these developments are consistent with clean energy production and climate risk management. One of the goals of the EO and the Memorandum is to leverage the Department of Energy's computing capabilities in order to develop new AI models and applications in the areas of energy and climate risks to ensure greater system resilience.³⁵

Moreover, the DOE, through the National Nuclear Security Administration (NNSA), is involved in assessing and mitigating AI-related risks, especially those related to nuclear and radiological threats.

³² See M. Taddeo, D. McNeish, A. Blanchard, E. Edgar *Ethical Principles for Artificial Intelligence in National Defence, Philosophy & Technology*, Vol. 34, pages 1707-1729 (2021), available at

<https://link.springer.com/article/10.1007/s13347-021-00482-3>. For an overview of the principles adopted by different Agencies, see also the *Memorandum for the Heads of Executive of Departments and Agencies - Advancing Governance, Innovation, and Risk Management for Agency Use of Artificial Intelligence* available at the following URL:

<https://www.whitehouse.gov/wp-content/uploads/2024/03/M-24-10-Advancing-Governance-Innovation-and-Risk-Management-for-Agency-Use-of-Artificial-Intelligence.pdf>.

³³ See note 1.

³⁴ See the *Cybersecurity and Infrastructure Security Agency Act* of 2018 (6 U.S.C. 651-674) adopted under the Trump Administration.

³⁵ Sec. *Promoting innovation* 5.2., (g) (iii) of the E.O..

The role of the DOE thus appears to be central to the development of AI infrastructure, ensuring the security of AI applications, and mitigating the risks associated with advanced AI technologies.

It is noteworthy to mention that during the last few days at the White House, January 14, 2025, President Biden adopted The E.O. Advancing United States Leadership in Artificial Intelligence Infrastructure³⁶ which strongly connects AI and Energy issues with the aim of advance the leadership in the “clean energy technologies needed to power the future economy, including geothermal, solar, wind, and nuclear energy; foster a vibrant, competitive, and open technology ecosystem in the United States, in which small companies can compete alongside large ones; maintain low consumer electricity prices; and help ensure that the development of AI infrastructure benefits the workers building it and communities near it”.

The Executive Order 14141 recognises that AI systems require immense computing power and reliable energy sources. The order underlines how the expansion of AI data centres is directly linked to investments in clean energy generation - including geothermal, nuclear, wind and solar.

Therefore, the grid modernisation and improved interconnection are strongly encouraged³⁷ as well as the prioritization of the permitting procedures required for the construction and operation of AI infrastructure.³⁸

2.2.6 Accountability and human rights

61

One of the most interesting aspects of the Biden Administration approach in the field of AI is certainly the AI risk assessment which has been introduced by the AI Risk Management Framework adopted by the National Institute of Standards and Technology - U.S. Department of Commerce³⁹ (AI RMF) and the s.c. “Risk Management Profile for Artificial Intelligence and Human Rights” provided by the U.S. Department of State (AI RMP).⁴⁰

These two complementary documents are an attempt to fulfil the “gap in translating human rights concepts for technologists”.⁴¹

In detail, the both documents have the demanding scope to contribute to human rights due diligence practices. In some ways the Profile complements the Framework: “By referencing universally applicable,

³⁶ The E.O. is available at the following link:

<https://www.federalregister.gov/documents/2025/01/15/2025-00636/framework-for-artificial-intelligence-diffusion>

³⁷ See Sec. 6.

³⁸ See Sec. 7.

³⁹ *The AI Risk Management Framework* has been adopted by the National Institute of Standards and Technology - U.S. Department of Commerce. The Framework is available at the following URL:

<https://nvlpubs.nist.gov/nistpubs/ai/nist.ai.100-1.pdf>

⁴⁰ *The Risk Management Profile for Artificial Intelligence and Human Rights Bureau of Cyberspace and Digital Policy*, July 25, 2024, available at <https://www.state.gov/risk-management-profile-for-ai-and-human-rights/>.

⁴¹ See the Profile Sec 1.

internationally recognized human rights, the Profile provides a globally relevant normative basis for the AI RMF's recommended risk management actions".

The actions for AI designers, developers, deployers and users are summarized in different functions: "1) Govern (set up institutional structures and processes), 2) Map (understand context and identify risks), 3) Measure (assess and monitor risks and impacts), and 4) Manage (prioritize, prevent, and respond to incidents)". These activities "can be applied across applications, stakeholders, and sectors, and throughout the AI lifecycle"⁴².

In this light, both documents provide a pivotal attention on the AI risk due diligence, explaining that guidelines, best practices, risk assessments, remedial and recovering measures, metrics and quantitative indicators of Human Rights risks can be implemented in the different above-mentioned phases as a safeguard against the possible AI risks with respect to Human Rights. Risks that can arise throughout the whole AI lifecycle both as intended and unintended consequences of AI actors' actions.

3. The Trump Revoke and the next steps

As anticipated, one of the first Executive Orders adopted by President Trump in its second mandate has the aim to revoke all the policies adopted under the Biden Administration in the field of AI. As well known, the E.O Initial rescissions of harmful executive orders and actions adopted on January 20, 2025⁴³ expressly revokes the described E.O Executive Order 14110 of October 30, 2023, regarding Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence

Few days later – with the E.O named Removing Barriers to American Leadership in Artificial Intelligence adopted on January 23, 2025⁴⁴ – President Trump made another step forward revoking "certain existing AI policies and directives that act as barriers to American AI innovation, clearing a path for the United States to act decisively to retain global leadership in artificial intelligence". The purpose of the E.O. is a declaration of intent towards the implicit and beneficial market effects: "It is the policy of the United States to sustain and enhance America's global AI dominance in order to promote human flourishing, economic competitiveness, and national security".⁴⁵

For this purpose, the E.O requires the review of the "all policies, directives, regulations, orders, and other actions taken pursuant to the revoked Executive Order 14110 of October 30, 2023 (Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence)". Following

⁴² The Profile, Sec. 1.

⁴³ Available at the following URL:

<https://www.whitehouse.gov/presidential-actions/2025/01/initial-rescissions-of-harmful-executive-orders-and-actions/>

⁴⁴ Available at the following URL:

<https://www.whitehouse.gov/presidential-actions/2025/01/removing-barriers-to-american-leadership-in-artificial-intelligence/>

⁴⁵ See Sec. 2.

this aim, Agencies shall “suspend, revise, or rescind” all the actions that are in some way connected to the Biden E.O.

Moreover, the E.O. has the expressed aim to adopt an Action Plan, demanded to a bunch of competent subjects in the field of the AI and to the heads of the relevant Department and Agencies⁴⁶.

Subsequently, President Trump opened a public debate adopting a s.c. Request for information that closed on 15th March 2025 with 8,755 comments in a plethora of suggested topics.⁴⁷

This *tabula rasa* approach is an openness to the AI industries in the convention that there is no need to limit, at this stage, the growth and the update of the AI tools in different fields.

A *prima facie* “laissez faire” policy that leaves the market to run its own course, without any ex ante bias or regulatory guidance. Neither from the rights nor from an ethical perspective.⁴⁸

The Biden order granularity allowed Agencies and Department to implement their own guidelines and standard, generating a spread of knowledge which seems not possible to delete with a single act. This plethora of acts adopted by the different agencies and executive bodies leave a certain margin of continuity between the two Administrations.

Moreover, it should be noted how some pivotal acts adopted by the former Administration are not directly revoked, such as for example the above mentioned National Security Memorandum and the last minute E.O. on the AI infrastructure that seem to maintain their own effectiveness and capability to orient the USA administrative architecture.

More recently, April 7, 2025, the White House Office of Management and Budget adopted two memoranda that are in somehow symptomatic of the forthcoming policies: the first, Accelerating Federal Use of AI through

⁴⁶ Namely, the Assistant to the President for Science and Technology, the Special Advisor for AI and Crypto, and the Assistant to the President for National Security Affairs, in coordination with the Assistant to the President for Economic Policy, the Assistant to the President for Domestic Policy, the Director of the Office of Management and Budget, and the heads of such executive departments and agencies as the APST and APNSA.

⁴⁷ Such as for example: “hardware and chips, data centers, energy consumption and efficiency, model development, open source development, application and use (either in the private sector or by government), explainability and assurance of AI model outputs, cybersecurity, data privacy and security throughout the lifecycle of AI system development and deployment (to include security against AI model attacks), risks, regulation and governance, technical and safety standards, national security and defense, research and development, education and workforce, innovation and competition, intellectual property, procurement, international collaboration, and export controls” The Request for Information on the Development of an Artificial Intelligence (AI) Action Plan is available at the following link: <https://www.federalregister.gov/documents/2025/02/06/2025-02305/request-for-information-on-the-development-of-an-artificial-intelligence-ai-action-plan>.

⁴⁸ See O. Pollicino, G. Gentile, *How the US threw out any concerns about AI safety within days of Donald Trump coming to office*, *The Conversation*, March 11, 2025, available at: <https://theconversation.com/how-the-us-threw-out-any-concerns-about-ai-safety-within-days-of-donald-trump-coming-to-office-251659>

Innovation, Governance, and Public Trust⁴⁹ and the second concerning Driving Efficient Acquisition of Artificial Intelligence in Government⁵⁰

These acts are oriented to “guidance to agencies on how to innovate and promote the responsible adoption, use, and continued development of AI, while ensuring appropriate safeguards are in place to protect privacy, civil rights, and civil liberties, and to mitigate any unlawful discrimination, consistent with the AI in Government Act”.

At our scope it is noteworthy that the agencies activities should be based on AI risk assessment that should balance the different rights and instances at stake⁵¹ in a certain continuity with the framework introduced by the former Administration.

4. Conclusions

Despotic scenarios – such as the so-called “reserve scenario” in which the use of AI is geographically confined or monopolized by authoritarian regimes – remain, for now, outside the mainstream regulatory debate⁵².

However, what no longer seems remote is the possibility of losing control over the very inputs and outputs that feed and emerge from AI systems. It is no longer possible to understand why a certain pawn is sacrificed in the game of chess, it is somehow no longer possible to understand the whole game and all the variations that the machine knows. .

At the same time, artificial intelligence holds the unprecedented potential to bridge structural inequalities, reduce geographical and social distances, and generate widespread benefits. This duality is already evident in sectors such as healthcare, where well-trained AI models are identifying patterns and correlations that escape even the most skilled human researchers—enabling earlier diagnoses, more personalized treatments, and accelerated scientific discovery.

The proliferation of guidelines, ethical principles, and voluntary frameworks has undoubtedly raised awareness, but it also risks creating a fragmented and inconsistent landscape. On the other hand, the broad embrace of a laissez-faire approach, where innovation is left to evolve without sufficient regulatory anchoring, opens the door to a wide spectrum of AI scenarios, from the utopian to the dystopian.

In both cases the effectiveness of governance mechanisms – be they regulatory, institutional, or technical – becomes crucial in determining the trajectory of AI deployment. Yet, achieving such effectiveness is no simple

⁴⁹ Available at the following URL:

<https://www.whitehouse.gov/wp-content/uploads/2025/02/M-25-21-Accelerating-Federal-Use-of-AI-through-Innovation-Governance-and-Public-Trust.pdf>

⁵⁰ Available at the following URL:

<https://www.whitehouse.gov/wp-content/uploads/2025/02/M-25-22-Driving-Efficient-Acquisition-of-Artificial-Intelligence-in-Government.pdf>

⁵¹ Namely: a) The intended purpose for the AI and its expected benefit; b) The quality and appropriateness of the relevant data and model capability; c) The potential impacts of using AI; d) Reassessment scheduling and procedures; e) related costs analysis; f) Results of independent review; g) Risk acceptance.

⁵² See M. Tegmark, *Life 3.0: Being Human in the Age of Artificial Intelligence*, London, 2017.

task, especially within what increasingly appears to be a permanent “beta phase” of AI.

Following this path, It is not yet clear if the Trump s.c. “revoke” and the new forthcoming AI policies will represent an effective next step in AI Governance.

Some crucial aspects such as principles of transparency, fairness, and risk management seem to be pretty similar in both approaches. It is possible to argue how the apparent policy discontinuity masks a deeper regulatory resilience: the diffusion of AI governance principles across federal agencies and institutions has already created a soft law baseline that cannot be easily undone by a single act or declaration.

Valerio Lubello
Dip. to di Studi giuridici “A. Sraffa”
Università commerciale L. Bocconi
valerio.lubello@unibocconi.it

