

Towards legal regulation of artificial intelligence*

Hacia la regulación jurídica de la inteligencia artificial

Moisés Barrio Andrés**

ABSTRACT

This paper addresses the current state of AI regulation and discusses the feasibility and need for regulation through legally binding instruments, beyond the field of ethics. History shows what can happen if the State withdraws and allows private companies to set their own exclusive regulatory standards. We have a unique opportunity to create laws and principles governing AI on a common basis and with an indispensable public-private partnership that should preferably be international in scope with the leadership of the UN or, failing that, at European level led by EU institutions, as they had already achieved in the areas of privacy and data protection.

KEYWORDS

Artificial intelligence regulation, AI regulation, AI laws, robotics, regulation, Digital law, ICT law

RESUMEN

Este artículo aborda el estado actual de la regulación de la IA y analiza la viabilidad y la necesidad de regularla mediante instrumentos jurídicamente vinculantes, más allá del ámbito de la ética. La historia demuestra lo que puede ocurrir si el Estado se retira y permite que las empresas privadas establezcan sus propias normas reguladoras. Tenemos una oportunidad única para crear leyes y principios jurídicos que regulen la IA sobre una base común y con una imprescindible colaboración público-privada que debería ser preferentemente de ámbito internacional con el liderazgo de la ONU o, en su defecto, a nivel europeo liderado por las instituciones de la UE, como ya se consiguió en los ámbitos de la privacidad y la protección de datos.

PALABRAS CLAVE

Regulación de la inteligencia artificial, regulación de la IA, leyes de la IA, robótica, regulación, Derecho digital, Derecho de las TIC, Derecho de las nuevas tecnologías

*Artículo de Investigación postulado el 24 de enero de 2020 y aceptado el 9 de noviembre de 2020

**Letrado del Consejo de Estado, España. (moises.barrio@consejo-estado.es) orcid.org/0000-0002-2877-5890

SUMMARY

1. Introduction
2. National Regulatory Frameworks For AI
 - a) European Union
 - b) Usa
 - c) Japan
 - d) China
3. Concluding Remarks
4. References

1. Introduction

Artificial intelligence (AI)¹ is a *vintage* technology. Its inception began in the wake of the Second World War. We are currently witnessing AI's boom due to a significant reduction in hardware costs, the capabilities and availability of big data and the cloud, and the "Artificial Intelligence as a Service" (AIaaS), which allows software developers to use components from, for instance, IBM, Google or Microsoft instead of having to program all aspects of the application from scratch. AI is *gradually* becoming ubiquitous: in our domestic devices of the Internet of Things (IoT)², in services and digital platforms, in robots, in the streets of smart cities, in offices, in factories, in hospitals, etc. And finally, albeit very timidly for the time being, in law firms and courts.

In the present day, of course, the legal world was not going to be left impervious to this disruptive technology. States are beginning to legislate on the technology, albeit unhurriedly. This paper discusses the attempts to create a legally binding regulations for AI.

2. National Regulatory Frameworks for AI

As Jacob Turner³ has pointed out, States' public policies regulating AI generally fall into at least one of the following three categories: (i) promoting the growth of a local AI industry; (ii) ethics and regulation of 1AI; and (iii) tackling the problem of unemployment caused by AI. These categories may sometimes be

¹ About the concept and legal challenges of AI I will refer to my book Barrio Andrés, Moisés, *Manual de Derecho digital*, Valencia, Tirant lo Blanch, 2020.

² Weber, Rolf H., *Internet of Things*, Berlin, Springer, 2010. See also, Barrio Andrés, Moisés, *Internet de las Cosas*, Madrid, Reus, 2020, 2a edition.

³ Turner, Jacob, *Robot Rules*, London, Palgrave Macmillan, 2018, p. 225 et. seq.

in tension and, at other times, can be mutually supportive. This paper focuses on regulatory initiatives rather than economic or technological ones, which are analysed in other contributions.⁴

The following exposition is merely a brief summary and does not intend to comprehensively examine all laws and government initiatives concerning the regulation of AI. Public policies are clearly developing fast and any exhaustive study would soon become outdated. Instead, our intention is to capture an array of general regulatory approaches with a view toward establishing the general tendency in various of the foremost jurisdictions involved in the AI industry.

a) European Union

The European Union has launched several initiatives aimed at developing a comprehensive AI strategy, including its regulation. The three key documents in this regard are the General Data Protection Regulation⁵ (GDPR), the European Parliament's Resolution of February 2017 on Civil Laws for Robotics⁶, and the Ethics guidelines for trustworthy AI produced by the European Commission's High-Level Expert Group on Artificial Intelligence⁷ (AI HLEG), a final version of which was presented in April 2019.

Although the GDPR was not aimed specifically at AI, its provisions nevertheless appear likely to have fairly drastic effects on the industry even beyond what its drafters might have intended.⁸ The GDPR extends the scope of EU data protection law to all foreign companies processing data of EU residents. The GDPR is intertwined with AI for several reasons, including that it requires a certain amount of explanation, which can be challenging with “black box” AI systems. Article 22 GDPR stipulates that: *“In particular, the controller must allow for a human intervention and the right for individuals to express their point of*

⁴ See also, Calo, Ryan, Froomkin, A. Michael and Kerr, Ian (eds.), *Robot Law*, Edward Elgar Publishing, New Work, 2016, and Barrio Andres, Moises (ed.), *Derecho de los Robots*, Madrid, Wolters Kluwer, 2019, 2nd edition.

⁵ Regulation (EU) 2016/679 of the European Parliament and of the Council of 27 April 2016 on the protection of natural persons with regard to the processing of personal data and on the free movement of such data, and repealing Directive 95/46/EC (General Data Protection Regulation).

⁶ European Parliament resolution of 16 February 2017 with recommendations to the Commission on Civil Law Rules on Robotics (2015/2103(INL)), [Accessed 23 January 2020] Available at: <http://www.europarl.europa.eu/sides/getDoc.do?pubRef=-//EP//TEXT+TA+P8-TA-2017-0051+0+DOC+XML+V0//EN>

⁷ [Accessed 23 January 2020] Available at: <https://ec.europa.eu/digital-single-market/en/news/draft-ethics-guidelines-trustworthy-ai>

⁸ Lopez Calvo, Jose (ed.), *La adaptación al nuevo marco de protección de datos tras el RGPD y la LOPDGD*, Madrid, Wolters Kluwer, 2019.

view, to obtain further information about the decision that has been reached on the basis of this automated processing, and the right to contest this decision.”

The European Parliament’s Resolution of February 2017 on Civil Laws for Robotics includes thought-provoking content, but it has not yet created binding law; instead it was merely a recommendation to the Commission for future action. Specifically, it is the preparatory document for the drafting of a Directive concerning civil-law rules on robotics. The results and a summary of this consultation were made available in a later report published in October 2017. Meanwhile, the European Parliament voted on the resolution in February 2017 to regulate the development of AI and robotics throughout the European Union. The Joint Declaration on the EU’s legislative priorities for 2018-2019 also named data protection, digital rights, and ethical standards in artificial intelligence and robotics as priorities.

Taking up the European Parliament’s appeal to create binding legislation, the European Commission issued a call in March 2018 for a High-Level Expert Group on Artificial Intelligence (AI HLEG), which, according to the Commission, “*will serve as the steering group for the European AI Alliance’s work, interact with other initiatives, help stimulate a multi-stakeholder dialogue, gather participants’ views and reflect them in its analysis and reports.*”

The work of the AI HLEG includes “*propos[ing] to the Commission AI ethics guidelines, covering issues such as fairness, safety, transparency, the future of work, democracy and more broadly the impact on the application of the Charter of Fundamental Rights, including privacy and personal data protection, dignity, consumer protection and non-discrimination.*”⁹ An initial version of the guidelines was published on 18 December 2018 and the experts presented their final version¹⁰ to the Commission in April 2019 after extensive consultation throughout the European AI Alliance. Based on fundamental rights and ethical principles, the document lists seven key requirements that relevant systems should meet in order to be trustworthy. Aiming to operationalise these requirements, an assessment list is presented to provide guidance on the requirements’ practical implementation. This assessment list will undergo a piloting process to which all interested stakeholders can participate. The objective is to then bring Europe’s ethical approach to the global stage. The

⁹ European Commission, Call for a High-Level Expert Group on Artificial Intelligence, website of the European Commission, [Accessed 23 January 2020], Available at: <https://ec.europa.eu/digital-single-market/en/high-level-expert-group-artificial-intelligence>, accessed 23 January 2020.

¹⁰ [Accessed 23 January 2020], Available at: https://ec.europa.eu/newsroom/dae/document.cfm?doc_id=58477

Commission is opening up cooperation to all non-EU countries that are willing to share the same values.

In April 2018, 25 EU countries signed a joint declaration of cooperation on AI, the terms of which included a commitment to “[e]xchange views on ethical and legal frameworks related to AI in order to ensure responsible AI deployment”¹¹. Subsequently, on 25 April 2018, the European Commission adopted a Communication on Artificial Intelligence for Europe¹² laying down the European approach to take utmost advantage of the opportunities offered by AI and address the new corresponding challenges of AI. In 2019, the Commission developed and made available the Guidance on the interpretation of the Product Liability Directive to prepare its reform in 2020.

Also, on 7 December 2018, the Commission submitted a Communication to the European Parliament, the European Council, the Council, the European Economic and Social Committee and the Committee of the Regions entitled *Coordinated Plan on Artificial Intelligence*¹³, accompanied by the *Coordinated Plan on the Development and Use of Artificial Intelligence Made in Europe - 2018* prepared by Member States (as part of the Group on digitizing European industry and Artificial Intelligence), Norway, Switzerland, and the Commission.

Next, on 19 February 2020, the European Commission published a White Paper¹⁴ aiming to foster a European ecosystem of excellence and trust in AI and a Report on the safety and liability aspects of AI.¹⁵ The White Paper proposes measures that will streamline research, foster collaboration between Member States and increase investment into AI development and deployment. It also proposes policy options for an imminent EU regulatory framework that would determine the types of legal requirements that would apply to relevant actors, with a particular focus on high-risk applications.

Finally, on 20 October 2020, the European Parliament has approved three resolutions that analyze how the European Union understands that different matters affected by Artificial Intelligence (AI) can be regulated, while promoting innovation, ethical standards and confidence in this technology. These resolutions analyze (i) the intellectual property rights for the development of AI technologies (supporting an effective system to guarantee them and to

¹¹ EU Member States Sign Up to Cooperate on Artificial Intelligence, website of the European Commission, 10 April 2018, [Accessed 23 January 2020], Available at: <https://ec.europa.eu/digital-single-market/en/news/eu-member-states-sign-cooperate-artificial-intelligence>

¹² COM(2018) 237 final.

¹³ COM(2018) 795 final.

¹⁴ COM(2020) 65 final.

¹⁵ COM(2020) 64 final.

safeguard the European patent standards); (ii) the civil liability regime in matters of AI (a regulatory framework is proposed to guarantee the strict liability of the operators of “high risk” AI systems in case of damages); and (iii) certain ethical aspects of AI, robotics and related technologies (having as a key guiding principle the human control and centrality).

Despite these encouraging signs and laudable intentions, the EU’s regulatory agenda remains at a nascent stage.

b) USA

In its final months, the Obama administration produced a major report on the Future of Artificial Intelligence, along with an accompanying strategy document.¹⁶ Although these documents focussed primarily on the economic impact of AI, they also briefly covered topics such as “AI and Regulation” and “Fairness, Safety and Governance in AI”.¹⁷ In late 2016, a large group of US universities sponsored by the National Science Foundation (NSF) published *A Roadmap for US Robotics: From Internet to Robotics*, a 109-page document edited by Ryan Calo.¹⁸ The report included calls for further work on AI ethics, safety, and liability.

Although the subsequent Trump administration initially appeared to have abandoned the topic as a major priority, as of the beginning of 2019 stands out to have changed course. In a 31 July 2018 memo from the Executive Office of the President, leadership in AI (along with “*quantum information sciences and strategic computing*”) is listed as the second-highest R&D priority for the fiscal year 2020, just after the security of the Americans. And on 7 September 2018, the U.S. Department of Defense (DoD) announced that it would invest up to \$USD 2 billion over the following five years in the advancement of AI. That amount would be in addition to existing government spending on AI R&D, which totalled more than \$USD 2 billion in 2017 alone, an amount that

¹⁶ The Administration’s Report on the Future of Artificial Intelligence, website of the Obama White House, 12 October 2016, Available at: <https://obamawhitehouse.archives.gov/blog/2016/10/12/administrations-report-future-artificial-intelligence>, accessed 1 March 2019. For the reports themselves, Available at: https://obamawhitehouse.archives.gov/sites/default/files/whitehouse_files/microsites/ostp/NSTC/preparing_for_the_future_of_ai.pdf, and also, Available at: https://obamawhitehouse.archives.gov/sites/default/files/whitehouse_files/microsites/ostp/NSTC/national_ai_rd_strategic_plan.pdf, [Accessed 23 January 2020]

¹⁷ Preparing for the Future of Artificial Intelligence, Executive Office of the President National Science and Technology Council Committee on Technology, 17–18 and 30–32 October 2016, [Accessed 23 January 2020], Available at: https://obamawhitehouse.archives.gov/sites/default/files/whitehouse_files/microsites/ostp/NSTC/preparing_for_the_future_of_ai.pdf

¹⁸ A Roadmap for US Robotics: From Internet to Robotics, 31 October 2016, [Accessed 23 January 2020], Available at: <http://jacobsschool.ucsd.edu/contextualrobotics/docs/rm3-final-rs.pdf>

includes only unclassified programs and not any amounts under the Pentagon and intelligence agencies' budgets. Existing funding has already propelled more than 20 active programs under the Defense Advanced Research Projects Agency (DARPA) exploring diverse aspects and uses of AI, and dozens of new projects have now been promised.

This funding follows the announcement in August 2018¹⁹ of a National Security Commission on Artificial Intelligence that was subsequently made official with President Trump's signing of the *2019 National Defense Authorization Act* (NDAA). The Commission will include 15 members selected by various government officials in the coming months. The Commission will assess the national-security implications of AI, including the ethical considerations of AI in defense. The DoD also established²⁰ a Joint AI Center (JAIC) in July 2018 to explore the agency's use of AI, although the contours of the JAIC's mission have yet to be defined. The JAIC will ostensibly work on AI National Mission Initiatives, improve collaboration with the private sector, academia, and military allies; attract AI talent and establish an ethical framework for AI in defence; and aid the National Defense Strategy. The DoD may also soon publish an AI Strategy.

In May 2018, President Trump and the White House held a Summit on Artificial Intelligence for American Industry with the participation of key technology companies. The White House also released a Fact Sheet, entitled *Artificial Intelligence for the American People*²¹, highlighting the Trump administration's priorities for AI. Trump declared his intention for the US to be the global leader in AI, pointing out that "[t]o the greatest degree possible, we will allow scientists and technologists to freely develop their next great inventions right here in the United States." Any attention to job losses, the impact of immigration policies on the technology sector, privacy, cybersecurity, and the impact on vulnerable groups was apparently minimal. Instead, the priorities discussed were funding AI research, removing regulatory barriers to the deployment of AI-powered technologies, training the future American workforce, achieving strategic military advantage, leveraging AI for government services, and working with allies to promote AI R&D. The White House announced plans to help provide US companies with new data sources and to establish a Select Committee on Artificial

¹⁹ [Accessed 23 January 2020], Available at: <https://www.executivegov.com/2018/08/fy-2019-ndaa-to-authorize-10m-for-ai-national-security-commission/>

²⁰ [Accessed 23 January 2020], Available at: <https://www.fedscoop.com/dod-joint-ai-center-established/>

²¹ [Accessed 23 January 2020], Available at: <https://www.whitehouse.gov/briefings-statements/artificial-intelligence-american-people/>

Intelligence to help government agencies contemplate and use the technology, as well as consider partnerships with industry and academia.

What is more, President Trump specifically named artificial intelligence as an Administration R&D priority in his 2019 Budget Request to Congress. AI was also featured for the first time in the National Security Strategy in relation to its role in helping the US lead in technological innovation, as well as AI's role in information statecraft, weaponisation, and surveillance. AI also appears for the first time in the National Defense Strategy, where it is described as one of the technologies that will change the character of war and afford increasingly sophisticated capabilities to our adversaries, including non-State actors. Moreover, autonomous systems that include AI and machine learning (ML) are described as one of the primary areas in which modernisation of key capabilities is desirable.

President Trump issued an Executive Order launching *the American AI Initiative*²² on 11 February, 2019. The Executive Order explained that the Federal Government plays an important role not only in facilitating AI R&D, but also in promoting trust, training people for a changing workforce, and protecting national interests, security, and values. And while the Executive Order emphasizes American leadership in AI, it is stressed that this requires enhancing collaboration with foreign partners and allies. The initiative is guided by five principles, which include (in a summarized form) the following: 1. Driving technological breakthroughs, 2. Driving the development of appropriate technical standards, 3. Training workers with the skills to develop and apply AI technologies, 4. Protecting the American values including civil liberties and privacy and fostering public trust and confidence in AI technologies, and 5. Protecting the US technological advantage in AI, while promoting an international environment that supports innovation. The day after the Executive Order was released, the US Department of Defense followed up with the release of an unclassified summary of its own Artificial Intelligence Strategy. The U.S. Air Force released an Annex to this strategy to share its own 2019 Artificial Intelligence Strategy in September 2019.

42

Numerous bills have also been introduced in Congress that either refer to or focus on artificial intelligence. There are at least nine bills that relate to autonomous driving, including *The SELF DRIVE Act*, now called *AV START Act*, which passed the House in September 2017.²³ The bill charges the Department

²² On March 19, 2019, the US federal government released <http://AI.gov> to make it easier to access all of the governmental AI initiatives currently underway. The site is the best single resource from which to gain a better understanding of US AI strategy.

²³ See, [Accessed 23 January 2020], Available at: <https://www.congress.gov/bills/115th-congress/house-bill/3388/>

of Transportation (DoT) with undertaking research on the best way to inform consumers about the capabilities and limitations of highly automated vehicles. But without a doubt, the most relevant regulatory instrument is the *Algorithmic Accountability Act*²⁴ presented last April 2019 in the Senate, which, if finally approved, requires companies that apply automated decision-making techniques to audit their machine-learning systems for bias and discrimination and to take corrective action in a timely manner if such issues were identified. It would also require those companies to audit all processes beyond machine learning involving sensitive data for privacy and security risks. Should it pass, the bill would place regulatory power in the hands of the US Federal Trade Commission (FTC), the agency in charge of consumer protections and antitrust regulation.

Several other AI-related bills are being introduced at state and local levels. For example, in August 2018, the California State Senate passed a resolution in support of the Asilomar AI Principles²⁵ (a set of twenty-three guidelines for the safe and beneficial development and use of AI). Likewise, the New York City Council passed an algorithmic accountability bill in 2017 that established the New York Algorithm Monitoring Task Force; the group studies how municipal agencies employ algorithms to make decisions that affect the lives of New Yorkers. In December 2017, Supervisor David Canepa introduced a resolution in California's San Mateo County that called on Congress and the United Nations to restrict the development and use of lethal autonomous weapons. Elsewhere in California, San Francisco Supervisor Jane Kim created an initiative in 2017 called the Jobs of the Future Fund to help prepare for the likelihood of job losses due to automation.²⁶

At the very least, it appears that the US Federal Government aspires to regain lost ground and has been attempting to position itself among the leading AI nations.

c) Japan

Industry in Japan has for nearly half a century placed a particular focus on automation and robotics.²⁷ The Japanese Government has generated various

²⁴ See, [Accessed 23 January 2020], Available at: <https://www.wyden.senate.gov/imo/media/doc/Algorithmic%20Accountability%20Act%20of%2019%20Bill%20Text.pdf>

²⁵ See, [Accessed 23 January 2020], Available at: <https://futureoflife.org/ai-principles/>

²⁶ See, [Accessed 23 January 2020], Available at: <https://www.jobsofthefuturefund.com/>

²⁷ Shimpo, Fumio, "The Principal Japanese AI and Robot Strategy and Research Toward Establishing Basic Principles", *Journal of Law and Information Systems*, Vol. 3 (May 2018).

strategy and policy papers with a view toward maintaining this position. For instance, in its 5th Science and Technology Basic Plan (2016–2020), the Japanese Government declared its aim to “*guide and mobilize action in science, technology, and innovation to achieve a prosperous, sustainable, and inclusive future that is, within the context of ever-growing digitalization and connectivity, empowered by the advancement of AI.*”²⁸

In line with these goals, the Japanese Government’s Cabinet Office convened an Advisory Board on Artificial Intelligence and Human Society in May 2016 under the initiative of the Minister of State for Science and Technology Policy “*with the aim to assess different societal issues that could possibly be raised by the development and deployment of AI and to discuss its implication for society.*” The Advisory Board published a report in March 2017 that recommended further work on issues such as ethics, law, economics, education, social impact and R&D.²⁹

The Japanese Government’s proactive approach, driven by its national industrial strategy and aided by a strong public discourse on AI, provides an excellent model for how governments can foster discussion nationally and internationally. The challenge for Japan will be to sustain this early *momentum*, something that will be maintained if other countries follow its approach. Pending the submission of binding legislation, Japan has so far produced only a number of ethical recommendations.

d) China

In July 2017, the AI 2.0 proposal from the China Academy of Engineering triggered the launch of a fifteen-year New Generation Artificial Intelligence Development Plan. The plan is focused on a forward-looking blueprint for basic theories and common key technologies, including big-data intelligence, swarm intelligence, cross-media intelligence, hybrid-enhanced intelligence, and autonomous systems, and their applications in manufacturing, urbanisation, healthcare, and agriculture, as well as AI hardware and software platforms, policies and regulations, and ethical concerns. Another R&D project related to AI is the *Brain Science and Brain-Inspired Research*, comparable to Europe’s Human

²⁸ Report on Artificial Intelligence and Human Society, Japan Advisory Board on Artificial Intelligence and Human Society, 24 March 2017, Preface, [Accessed 23 January 2020], Available at: http://www8.cao.go.jp/cstp/tyousakai/ai/summary/aisociety_en.pdf

²⁹ Report on Artificial Intelligence and Human Society, Japan Advisory Board on Artificial Intelligence and Human Society, *op. cit.*

Brain Project, the BRAIN Initiative in the US, and other State-level projects. It is expected to be approved this year and should run for fifteen years.

Also in July 2017, China's State Council issued the *Next Generation Artificial Intelligence Development Plan* (新一代人工智能发展规划).³⁰ The policy plan outlines China's strategy to build a domestic AI industry worth nearly \$USD 150 billion over the next few years and to become the leading AI country by 2030. This document officially marked the development of the AI sector as a national priority and was included in President Xi Jinping's "grand vision" for China. Although this represented the first time that AI had been specifically mentioned in a work report of the Communist Party of China, the sentiment is seen more broadly as a continuation of the 13th Five-Year Plan and the State-driven industrial plan *Made in China 2025*. The Next Generation Artificial Intelligence Development Plan was described by two experienced analysts of Chinese digital technology as "[o]ne of the most significant developments in the artificial intelligence (AI) world" that year.

Although its main focus was on fostering economic growth through AI technology, the Plan also provided that "[b]y 2025 China will have seen the initial establishment of AI laws and regulations, ethical norms and policy systems, and the formation of AI security assessment and control capabilities." As Jeffrey Ding³¹ points out, "[n]o further specifics were given, which fits in with what some have called opaque nature of Chinese discussion about the limits of ethical AI research."

The Ministry of Science and Technology (MOST), as well as a new office named the AI Plan Promotion Office, are responsible for the implementation and coordination of the emergent AI-related projects, which are driven primarily by government-led subsidies. An AI Strategy Advisory Committee was also established in November 2017 to conduct research on strategic issues related to AI and to make recommendations. Furthermore, an AI Industry Development Alliance was established; the Alliance is co-sponsored by more than 200 enterprises and agencies nationwide and focuses on building a public-service platform for the development of China's AI industry in order to integrate resources and accelerate growth.

³⁰ Available in English translation from the New America Institute: A Next Generation Artificial Intelligence Development Plan, China State Council, translated by Rogier Creemers, Leiden Asia Centre; Graham Webster, Yale Law School Paul Tsai China Center; Paul Triolo, Eurasia Group; and Elsa Kania, 20 July 2017, [Accessed 23 January 2020], Available at: <https://na-production.s3.amazonaws.com/documents/translation-fulltext-8.1.17.pdf>

³¹ Ding, Jeffrey, "Deciphering China's AI Dream", in Governance of AI Program. Future of Humanity Institute (Oxford: Future of Humanity Institute, March 2018), 30, [Accessed 23 January 2020], Available at: https://www.fhi.ox.ac.uk/wp-content/uploads/Deciphering_Chinas_AI-Dream.pdf

In November 2017, Tencent Research, an institute within one of China's largest technology companies, and the China Academy of Information and Communications Technology (CAICT) produced a book of 482 pages, the title of which roughly translates to *A National Strategic Initiative for Artificial Intelligence*. Topics covered include law, governance, and the morality of machines.

In a paper entitled *Deciphering China's AI Dream*³², Ding hypothesises that "AI may be the first technology domain in which China successfully becomes the international standard setter." The report points out that the book National Strategic Initiative for Artificial Intelligence identified Chinese leadership on AI ethics and safety as a way for China to seize the strategic high ground. Ding notes that the book emphasises that "China should also actively construct the guidelines of AI ethics, play a leading role in promoting inclusive and beneficial development of AI. In addition, we should actively explore ways to go from being a follower to being a leader in areas such as AI legislation and regulation, education and personnel training, and responding to issues with AI."³³

Ding³⁴ observes further:

One important indicator of China's ambitions in shaping AI standards is the case of the International Organization for Standardization [...] Joint Technical Committee (JTC), one of the largest and most prolific technical committees in the international standardization, which recently formed a special committee on AI [SC 42]. The chair of this new committee is Wael Diab, a senior director at [Chinese multinational company] Huawei, and the committee's first meeting will be held in April 2018 in Beijing, China - both the chair position and first meeting were hotly contested affairs that ultimately went China's way.

In furtherance of its policies, China established a national AI-standardisation group and a national AI expert-advisory group in January 2018.³⁵ At the launch event for these groups, a division of China's Ministry of Industry and

³² Ding, Jeffrey, "Deciphering China's AI Dream", *op. cit.*

³³ *Ibid.*

³⁴ *Ibid.*

³⁵ Triolo, Paul and Goodrich, Jimmy, "From Riding a Wave to Full Steam Ahead As China's Government Mobilizes for AI Leadership, Some Challenges Will Be Tougher Than Others", *New America*, 28 February 2018, [Accessed 23 January 2020], Available at: <https://www.newamerica.org/cybersecurity-initiative/digichina/blog/riding-wave-full-steam-ahead/>

Information Technology released a 98-page White Paper on AI standardization.³⁶ The White Paper noted that AI raised challenges in terms of legal liability, ethics and safety, stating:

[...] considering that the current regulations on artificial intelligence management in various countries in the world are not uniform and relevant standards are still in a blank state, participants in the same AI technology may come from different countries which have not signed a shared contract for artificial intelligence. To this end, China should strengthen international cooperation and promote the formulation of a set of universal regulatory principles and standards to ensure the safety of artificial intelligence technology.

China's goal of becoming a leader in the regulation of AI may be one of the motivations behind its call in April 2018 to United Nations Group of Governmental Experts on lethal autonomous weapons systems *"to negotiate and conclude a succinct protocol to ban the use of fully autonomous weapon systems."*³⁷ In so doing, China for the first time adopted a different approach regarding autonomous weapons than that of the US. The Campaign to Stop Killer Robots announced that China had joined twenty-five other nations in calling for such a ban.³⁸

Triolo and Goodrich³⁹ have pointed out that *"[a]s in many other areas, Chinese government leadership on AI at least nominally comes from the top. Xi has identified AI and other key technologies as critical to his goal of transforming China from a 'large cyber power' to a 'strong cyber power' (also translated as 'cyber superpower')"*. This approach seems to originate from the White Paper.

³⁶ White Paper on Standardization in AI, National Standardization Management Committee, Second Ministry of Industry, 18 January 2018, [Accessed 23 January 2020], Available at: <http://www.sgcc.gov.cn/upload/f1ca3511-05f2-43a0-8235-eeb0934db8c7/20180122/5371516606048992.pdf>. Contributors to the white paper included the China Electronics Standardization Institute, Institute of Automation, the Chinese Academy of Sciences, the Beijing Institute of Technology, the Tsinghua University, the Peking University and the Renmin University, as well as private companies such as Huawei, Tencent, Alibaba, Baidu, Intel (China) and Panasonic (formerly Matsushita Electric) (China) Co., Ltd.

³⁷ The original recording of the Chinese delegation's statement is available on the UN Digital Recordings Portal website, [Accessed 23 January 2020], Available at: https://conf.unog.ch/digitalrecordings/index.html?guid=public/61.0500/E91311E5-E287-4286-92C6-D47864662A2C_10h14&position=1197

³⁸ Convergence on Retaining Human Control of Weapons Systems, Campaign to Stop Killer Robots, 13 April 2018, [Accessed 23 January 2020], Available at: <https://www.stopkillerrobots.org/2018/04/convergence/>

³⁹ Triolo, Paul and Goodrich, Jimmy, "From Riding a Wave to Full Steam Ahead As China's Government Mobilizes for AI Leadership, Some Challenges Will Be Tougher Than Others", *op. cit.*

In May 2019, the Beijing AI Principles were released by a multistakeholder coalition including the Beijing Academy of Artificial Intelligence (BAAI), Peking University, Tsinghua University, Institute of Automation and Institute of Computing Technology in Chinese Academy of Sciences, and an AI industrial league involving firms like Baidu, Alibaba and Tencent. The 15 Principles call for “the construction of a human community with a shared future, and the realization of beneficial AI for humankind and nature.”

The Principles are separated into three sections: Research and Development, Use, and Governance. They include focus on benefitting all of humanity and the environment; serving human values such as privacy, dignity, freedom, autonomy, and rights; continuous focus on AI safety and security; inclusivity; openness; supporting international cooperation and avoiding a “malicious AI race”; and long-term planning for more advanced AI systems, among others.

Finally, there are also local government AI policy initiatives throughout China. For example, the Shanghai government issued its own implementation plan for new-generation AI in November 2017; Beijing announced a major new AI-focused industrial park to be constructed in Mentougou District in June 2018; Guangzhou launched an International Institute of AI; and many other districts have committed funds for AI research.

3. Concluding Remarks

In the absence of an international treaty or mandatory EU or national legislation to regulate AI, private companies have begun to act unilaterally. For instance, in 2016, six major technology companies –Amazon, Apple, Google, Facebook, IBM and Microsoft– formed the Partnership on Artificial Intelligence to Benefit People and Society⁴⁰ to *„study and formulate best practices on AI technologies, to advance the public’s understanding of AI, and to serve as an open platform for discussion and engagement about AI and its influences on people and society.“* Similarly, in October 2017, DeepMind, one of the world’s leading AI companies acquired by Google in 2014, created a new ethics committee, DeepMind Ethics & Society⁴¹, *„to help technologists put ethics into practice, and to help society anticipate and address the impact of AI in a way that works for the benefit of all.“*

⁴⁰ See, [Accessed 23 January 2020], Available at: <https://www.partnershiponai.org/>

⁴¹ See, [Accessed 23 January 2020], Available at: <https://deepmind.com/applied/deepmind-ethics-society/>

These initiatives are highly positive and valuable but are not sufficient. They lack the legitimacy that the State can provide. In fact, they leave out the myriad small to medium-sized enterprises that are also developing AI. Nevertheless, it is also imperative that the State ensures compliance with the legal systems and the fundamental principles and rights enshrined in national constitutions and the Charter of Fundamental Rights of the European Union.

But why do we need *binding* global legal regulation? In other studies⁴² we have dealt at length with new issues that are not covered by current laws. For example, it is not entirely clear who should be held liable if AI causes damage (for example, in an accident with an autonomous car or by incorrect application of an algorithm): the original designer, the manufacturer, the owner, the user or even the AI itself. We also discuss whether an autonomous electronic personality should be recognized for the most advanced systems that directly assigns them rights and obligations. There are even moral dilemmas about how AI should make specific important decisions even if there would be decisions involved on which it should not have the last word. If we apply solutions on a case-by-case basis, we risk uncertainty and confusion. As Oliver Wendell Holmes⁴³ said, „*hard cases make bad law*“, referring to which an extreme case is a poor basis for a general law that would cover a wider range of less extreme cases. Lack of regulation also increases the likelihood of hasty, instinctive, or even anger-fed reactions.

Moreover, a problem closely connected with AI regulation is that of data quality. One of the key elements of any AI system is the acquisition and preparation of data sets. They usually come from different sources, so they have to be integrated, cleaned, filtered and converted into a convenient format (normalized) to be processed by the available machine learning tools. Some working data sets are used for training the learning algorithms in order to create models. These models must be validated to ensure that they are doing the right pattern matching (validation) and that they have certain desirable properties such as coherence, consistency, etc. (verification). The best performing model is chosen for production normally undergoing a prior test session with another data set.

These issues are not merely theoretical concerns to entertain academics. AI systems already have the ability to make difficult decisions that have until now been based on human intuition or the laws and the practice of courts. Those decisions range from questions of life and death, such as the use of autonomous

⁴² Barrio Andrés, Moisés (ed.), *Derecho de los Robots*, Madrid, Wolters Kluwer, 2018, 2a edition, p. 87 and Barrio Andrés, Moisés, "Hacia una personalidad electrónica para los robots", *Revista de Derecho Privado*, N° 2/2018.

⁴³ *Northern Securities Co. v. United States*, 193 U.S. 197 (1904).

killer robots in the military, to issues of economic and social importance, such as how to avoid algorithmic biases when artificial intelligence decides, for example, whether to award a scholarship to a student or when to grant parole to an inmate. If a human being were to make these decisions, the human would always be subject to a legal rule and must accompany the decision with a legal motivation, i.e. to explain the *rationale* for the decision under the law. There are at present no such rules for AI.

The regulation of AI is currently presided over by corporate interests and is promoted from an ethical approach; this is not always desirable. A mere glance at the global financial crisis of 2008 illustrates the result of a self-regulated industry careening out of control. While the State has intervened to require banks to hold better assets to back their loans, the global economy continues to suffer the repercussions of a framework that was previously fundamentally self-regulatory.

That is not to say that progress is not being made. DeepMind has hired leading public analysts, including transhumanist philosopher Nick Bostrom and economist Jeffrey Sachs, as members of its ethics committee, and the list of the Partnership on AI members now includes non-profit organisations such as the American Civil Liberties Union, Human Rights Watch, and UNICEF. By early 2020, however, the Partnership on AI only has representatives from thirteen countries.

Nevertheless, ethical frameworks differ notably from legal frameworks given that legal frameworks can only be developed by international or State legislatures. Furthermore, ethical rules are only binding in the internal forum and entail, in cases of non-compliance, sin and potential eternal punishment, while legal rules are binding in the external forum, with their non-compliance entailing liability, sanctions, fines or even prison sentences.

For the time being, all States are still trying to catch up with Silicon Valley regarding AI regulation; the longer they wait, the more difficult it will be to properly manage the future of AI. Earlier we noted that the European Commission had launched a group of experts in June 2018 to examine the challenges posed by the development of artificial intelligence and its impact on the fundamental rights of the European Union (the High-Level Expert Group on Artificial Intelligence, AI HLEG). On 8 April 2019, this group had presented the final version of ethical guidelines for the development and use of artificial intelligence.

While it is a difficult achievement, it is not an impossible one. At the national level, States already oversee many other complex technologies including nuclear power and cloning. At the international level, the European Medicines

Agency (EMA) sets pharmaceutical standards for twenty-eight countries and ICANN regulates key parts of the entire Internet.

It is important to have an imperative –yet prudent and thoughtful– body of laws. Self-regulation is insufficient. If standards remain purely voluntary, some technology companies will decide to ignore any rules that do not benefit them, giving some organisations advantages over others. For example, none of the major Chinese AI companies, such as Alibaba, Tencent or Baidu, has announced that they will set up ethics committees or that they intend to join the Partnership on AI. Nor is it easy for a company to establish an ethics committee. The difficulties faced by Google in this field are very enlightening.

In addition, without a unified framework, too many private ethics committees could also lead to too many sets of rules. It would be chaotic and dangerous for every large company to have its own code for AI, just as it would be if every private citizen could establish his or her own legal statutes. Only the State has the power and the mandate to ensure a fair system that imposes this type of compliance in all areas. That is why all States are sovereign, typically have parliaments and a judiciary. In short, they have the backing of democratic legitimacy.

Therefore, when rules are drafted for AI systems, the voices of companies must remain contributors; yet, while highly relevant, they should not be akin to legislators. Technology companies may be well-positioned to design rules because of their experience in the field, but industry actors are rarely in the best position to adequately assess democratic, moral, ethical, and legal risks.

History shows what can happen if the State withdraws and allows private companies to set their own exclusive regulatory standards. Allowing this to happen in the case of AI would be not only reckless but also exceedingly dangerous. However, States have not yet reached definitive positions on how AI should be governed. We have a unique opportunity to create laws and principles governing AI on a common basis and with an indispensable public-private partnership that should preferably be international in scope with the leadership of the UN or, failing that, at European level led by EU institutions, as they had already achieved in the areas of privacy and data protection.

Reflecting our view, the European Commission drafted a White Paper, the final version of which was published in February 2020, which sets out the key pillars of the forthcoming regulatory framework for AI and actions to facilitate access to data through legally binding legislation in the EU law. And on 20 October 2020, the European Parliament adopted three reports outlining how the EU can best regulate AI while boosting innovation, ethical standards and trust

in technology. These three reports support an ethics framework and legal obligations, a civil liability system for breaches of the law through AI means, and an intellectual property system focused on granting rights only to humans. Ultimately, the need for legally binding rules governing AI is being consolidated.

4. References

- Barrio Andrés, Moisés (ed.), *Derecho de los Robots*, Madrid, Wolters Kluwer, 2019, 2nd edition.
- Barrio Andrés, Moisés, „Hacia una personalidad electrónica para los robots“, *Revista de Derecho Privado*, N° 2/2018.
- Barrio Andrés, Moisés, “Alastria as a case of Internet governance”, at Ibañez Jimenez, Javier (ed.), *Alastria mission and vision: A multidisciplinary research*, Reus, Madrid, 2020.
- Barrio Andrés, Moisés, *Ciberdelitos 2.0: amenazas criminales del ciberespacio*, Buenos Aires, Astrea, 2020, 2nd edition.
- Barrio Andrés, Moisés, *Internet de las Cosas*, Madrid, Reus, 2020, 2nd edition.
- Barrio Andrés, Moisés, *Manual de Derecho digital*, Valencia, Tirant lo Blanch, 2020.
- Benedikt, Carl Frey and Osborne, Michael, “The Future of Employment: How Susceptible Are Jobs to Computerisation?” Oxford Martin Programme on the Impacts of Future Technology Working Paper, September 2013.
- Boden, Margaret, *AI: Its nature and Future*, Oxford, Oxford University Press, 2016.
- Bostrom, Nick, *Superintelligence*, Oxford, Oxford University Press, 2014.
- Calo, Ryan, Froomkin, A. Michael and Kerr, Ian (eds.), *Robot Law*, Edward Elgar Publishing, New Work, 2016.
- Ding, Jeffrey, “Deciphering China’s AI Dream”, in *Governance of AI Program. Future of Humanity Institute* (Oxford: Future of Humanity Institute, March 2018).
- Geraci, Robert M., *Apocalyptic AI: Visions of Heaven in Robotics, Artificial Intelligence, and Virtual Reality*, New York, Oxford University Press, 2010.
- Hernandez-Orallo, Jose, “Beyond the Turing Test”, *Journal of Logic, Language and Information*, Vol. 9, No. 4 (2000).
- Kaplan, Jerry, *Artificial Intelligence: What Everyone Needs to Know*, New York, Oxford University Press, 2016.
- Kurzweil, Ray, *The Singularity Is Near: When Humans Transcend Biology*, New York, Viking Press, 2005.
- Lopez Calvo, Jose (ed.), *La adaptación al nuevo marco de protección de datos tras el RGPD y la LOPDGD*, Madrid, Wolters Kluwer, 2019.
- Nilsson, Nils J., *The Quest for Artificial Intelligence: A History of Ideas and Achievements*, Cambridge, Cambridge University Press, 2010.
- Rolf H., *Internet of Things*, Berlin, Springer, 2010.

- Russell, Stuart and Norvig, Peter, *Artificial Intelligence: International Version: A Modern Approach*, Englewood Cliffs, Prentice Hall, 2010.
- Shimpo, Fumio, "The Principal Japanese AI and Robot Strategy and Research Toward Establishing Basic Principles", *Journal of Law and Information Systems*, Vol. 3 (May 2018).
- Triolo, Paul and Goodrich, Jimmy, "From Riding a Wave to Full Steam Ahead As China's Government Mobilizes for AI Leadership, Some Challenges Will Be Tougher Than Others", *New America*, 28 February 2018.
- Turing, Alan M., "Computing Machinery and Intelligence", *Mind: A Quarterly Review of Psychology and Philosophy*, Vol. 59, No. 236 (1950).
- Turner, Jacob, *Robot Rules*, London, Palgrave Macmillan, 2018.
- Wallach, Wendell and Allen, Colin, *Moral Machines: Teaching Robots Right from Wrong*, Oxford, Oxford University Press, 2009.
- Weinbaum, David and Veitas, Viktoras, "Open Ended Intelligence: The Individuation of Intelligent Agents", *Journal of Experimental & Theoretical Artificial Intelligence*, Vol. 29, No. 2 (2017).