

AI and a Global Regulatory Framework

Alfredo M. Ronchi

Politecnico di Milano, Piazza Leonardo da Vinci 32, Milano, 20133, Italy

Tel: +39 393 0629373, Email: alfredo.ronchi@polimi.it, ORCID [0000-0003-1230-4338](https://orcid.org/0000-0003-1230-4338)

Abstract: This paper provides an overview on the key questions posed to policy makers and regulators to define a sounding regulatory framework that will not unintentionally limit the evolution of this technology. How to reduce or even eliminate the range of concerns raised by AI. A selection of the most relevant approached to the problem and already developed regulatory framework will complete the paper.

Keywords: Digital Transformation, Artificial Intelligence, Ethics, Jurisprudence, Human rights

1. AI systems and hallucination

Artificial intelligence has returned to the forefront after a long period of underground development, in an easily accessible and ready-to-use format. Innovative applications like ChatGBT have demonstrated the power of this technology, long considered beyond the reach of end users. In just a few months, ChatGBT has become a daily tool and a topic of discussion among friends and colleagues. Artificial intelligence has equipped a growing number of products and services, and, at the same time, concerns about potential drawbacks and risks have emerged.

Mistakes and hallucinations generated using AI, particularly generative AI, can have very serious consequences. Apart from the nightmares often generated by Sci Fiction movies like, War Games, Eagle Eye or even Lucy¹, this set of technologies has generated the need to regulate both the development and the deployment of its applications. Both experts in the field and governmental bodies are asking to issue specific regulations and laws, one of the most relevant voices in this field was due to San Altman “OpenAI’s Sam Altman Urges A.I. Regulation in Senate Hearing” [May 2023] or a newspaper level “The world wants to regulate AI but does not quite know how” [The Economist - Oct 24th, 2023].

The key questions concerning similar problems are:

- How to achieve precautionary/protective regulatory goals regulating AI design. Use & autonomy without stifling a jurisdiction’s competitive advantage?
- At which level does ought AI be treated “as IF” a morally considerate agent & how would that impact regulatory options?
- Should issues arising from any eventual authentic emerging moral considerateness be integrated into regulatory approaches now?

To better understand the range of concerns and the willingness to regulate a short list of shared and contrasting underlying concerns is:

- Security (homeland/national & defence secrets, military, space, weapons, infrastructure, financial systems, personal information, political processes/elections, commercial secretes, biosecurity),
- Discrimination and BIAS in health, any selection process, procurement, justice, policy, regulation, broadly anything AI may be involved with,

¹ Science fiction movies already proposed similar scenarios e.g. Wargames (1983 American techno-thriller film directed by John Badham) or Eagle Eye (2008 American action-thriller film directed by D. J. Caruso), Lucy (2014 French Sci Fi movie written and directed by Luc Besson)

- AI Moral responsibility – feel released from a personal ethical analysis and related responsibilities,
- Information - Oversight – Misinformation, disinformation, disapproved information filtering, nudging, opinion formation,
- Privacy, human rights & civil liberties protections
- Commercial advantage e.g. promoting national industry – via different methods
- Protection of System and Information Integrity programs from being hijacked
- Consumer protection – liability for harmful products

Some of the ad hoc approaches to the concerns are:

- Mixed frameworks national security, commercial advantage, justice, biosecurity,
- Shared preventative (but generally not extending to precautionary) focus-too late for after the fact liability to be effective.

Possible integrating framework:

- Traditional commercial liability for programme creators (downside v/s preventative effectiveness),
- Environmental law preventative frameworks (downsides v/s competitive or innovative advantage),
- Graduated “AI” focused liability as with other independent actors.

2. Is there a need to create a global AI regulatory framework?

The following paragraphs will provide a short overview on relevant recent developments.

2.1 International Organisations

IEEE published Ethically Aligned Design², a vision for prioritizing human well-being with autonomous and intelligent systems. Similar initiatives were due to GPAI The Global Partnership on Artificial Intelligence, OECD Artificial Intelligence and Robotics, and UNICRI United Nations Interregional Crime and Justice Research publication “Toolkit for Responsible AI Innovation in Law Enforcement³” or “AI for Safer Children”.

2.2 UNESCO

Artificial intelligence and related technologies will contribute to achieving the 2030 Agenda for Sustainable Development. UNESCO is actively contributing to addressing key issues related to the ethics of artificial intelligence, artificial intelligence in education, gender equality, and strengthening the capacity of governments and judiciary. AI’s rapid growth raises ethical concerns, including embedded biases, climate degradation, and human rights threats, exacerbating existing inequalities and harming marginalized groups.

Among other initiatives UNESCO launched the AI Initiative⁴ and published different reports on this topic including UNESCO Generative AI, UNESCO Generative Ai in Education⁵, and more⁶. The Recommendation on the Ethics of Artificial Intelligence was adopted by acclamation by 193 Member States in 2021.

² <https://ieeexplore.ieee.org/stamp/stamp.jsp?tp=&arnumber=9398613>

³ <https://unicri.it/topics/Toolkit-Responsible-AI-for-Law-Enforcement-INTERPOL-UNICRI>

⁴ <https://www.unesco.org/en/artificial-intelligence/recommendation-ethics?hub=32618>

⁵ <https://www.unesco.org/en/articles/generative-artificial-intelligence-education-think-piece-stefania-giannini>

⁶ <https://www.unesco.org/en/artificial-intelligence>

2.3 OECD

The OECD in 2019 published one of the first document providing some basic guidelines to regulate AI technology “Recommendation of the Council on Artificial Intelligence”⁷ that become a key reference globally, including the definition of an AI system provided “Artificial Intelligence (AI) is a general-purpose technology that has the potential to: improve the welfare and well-being of people, contribute to positive sustainable global economic activity, increase innovation and productivity, and help respond to key global challenges. It is deployed in many sectors ranging from production, education, finance and transport to healthcare and security.”

At the end of 2023, the OECD has decided to play its part, leveraging its unique cooperation infrastructure and convening power to support and foster a positive and proactive message about AI and privacy. The OECD active cooperation across borders, sectors, and areas of expertise is evident. The key role of OECD is to act as an observer and an historical partner of Global Privacy Assembly (GPA) that is a reference global network of Privacy Enforcement Authorities (PEAs) together with the Council of Europe (CoE).

In early 2024, the OECD⁸ formally launched the OECD.AI Expert Group on AI, Data, and Privacy, bringing together leading AI and privacy experts worldwide (data protection authorities, policymakers, industry, civil society, and academia). The OECD policy observatory published the 2024 OECD AI principles update⁹. The OECD Ministerial Council Meeting held the same year offered the opportunity to update the principles on AI governance established for the first time in 2019, the already mentioned first intergovernmental standard. The 2024 update considers new technological and policy developments, ensuring they remain robust and fit for purpose. The updated principles now address emerging challenges with an enhanced focus on safety, privacy, intellectual property rights and information integrity.

These principles defined the basis for innovative, trustworthy AI respectfully considering human rights and EU democratic values. Up to now we count 47 countries that adhere to these Principles.

The role of these Principles is to provide recommendations and guidelines to policy makers. This process allows the creation AI risk frameworks, building a foundation for global interoperability between jurisdictions.

The principles are based on shared values and can be summarized:

- Inclusive growth, sustainable development and well-being (Principle 1.1),
- Human-centred values and fairness (Principle 1.2),
- Transparency and explainability (Principle 1.3),
- Robustness, security and safety (Principle 1.4),
- Accountability (Principle 1.5).

Coping with these principles we find recommendations for policy makers such as: Investing in AI research and development (Principle 2.1), Fostering a digital ecosystem for AI (Principle 2.2), Shaping an enabling policy environment for AI (Principle 2.3), Building human capacity and preparing for labour market transformation (Principle 2.4), International co-operation for trustworthy AI (Principle 2.5).

As stated by OECD “The Principles serve as a benchmark for responsible AI development and a critical checklist to address these rapid changes effectively, ensuring that AI continues to benefit society without compromising standards and safety.”

In addition, as already mentioned, there are some core aspects such as interoperability of AI and the definition of AI systems and its lifecycle.

⁷ <https://legalinstruments.oecd.org/en/instruments/OECD-LEGAL-0449>

⁸ OECD's definition of an AI system, OECD ARTIFICIAL INTELLIGENCE PAPERS (2024), OECD March 2024 No. 8

⁹ Care of Juraj Čorba, Audrey Plonk, Karine Perset, Yoichi Iida

2.3.1 Typical Principles life cycle proposed by the US IT Modernisation Centres of Excellence

The typical life cycle is composed by: Design, Development, Deployment¹⁰.

*“The **Design phase** consist in three main steps:*

Understand the problem: To share your team’s understanding of their mission challenge, you first must identify the key project objectives and requirements. Then define the desired outcome from a business perspective. Finally, determining AI will solve this problem. Learn more about this step in the framing AI problems section. (No AI solution will succeed without clearly and precisely understand the business challenge being solved and the desired outcome.)

Data gathering and exploration: This step deals with collecting and evaluating the data required to build the AI solution. This includes discovering available data sets, identifying data quality problems, and deriving initial insights into the data and perspectives on a data plan. (Data is the foundation of any AI solution. Without clearly understanding of the data required and the make-up of that data, a model cannot use it.)

Data wrangling and preparation: This phase covers all activities to construct the working data set from the initial raw data into a format that the model can use. This step can be time consuming and tedious but is critically important to develop a model that achieves the goals established in step 1. (Data preparation is often the hardest and most time-consuming phase of the AI lifecycle.)

*The **Develop phase** consist in two main steps:*

Modelling: This step focuses on experimenting with data to determine the right model. Often during this phase, the team trains, tests, evaluates, and retrains many different models to determine the best model and settings to achieve the desired outcome.

The model training and selection process is interactive. No model achieves best performance the first time it is trained. It is only through iterative fine-tuning that the model is honed to produce the desired outcome. Learn more about types of machine learning and models in Chapter 1. Depending on the amount and type of data being used, this training process may be very computationally expensive meaning it requires special equipment to provide enough computing power and cannot be performed on a normal laptop. See Chapter 5 to learn more about infrastructure to support AI development.

Evaluation: Once one or more models have been built that appear to perform well based on relevant evaluation metrics, test the models on new data to ensure they generalize well and meet the business goals. As this step is particularly critical, we discuss it in much greater detail in the test and evaluation section.

*The **Deploy phase** consist in two main steps:*

Move to production: Once a model has been developed to meet the expected outcome and performs at a level determined ready for use on live data, deploy it into a production environment. In this case, the model will take in new data that was not a part of the training cycle.

Monitor model output: Monitor the model as it processes this live data to ensure that it is adequately able to produce the intended outcome—a process known as generalization, or the model’s ability to adapt properly to new, previously unseen data. In production, models can “drift,” meaning that the performance will change over time. Careful monitoring of drift is important and may require continuous updating of the model.”

¹⁰ Source “Center of Excellence AI Guide for Government: a living and evolving guide to the application of artificial intelligence for the U.S. federal government. <https://coe.gsa.gov/coe/ai-guide-for-government/understanding-managing-ai-lifecycle/>

General remark: “As with any logic and software development, use an agile approach to continually retain and refresh the model. However, AI systems require extra attention. They must undergo rigorous and continuous monitoring and maintenance to ensure they continue to perform as trained, meet the desired outcome, and solve the business challenges.”

As it was demonstrated in different occasions the continuous monitoring and maintenance of AI and ML system is a key aspect to avoid unexpected behaviours and potential drawbacks, even if sometimes the outcome of the system is automatically transferred to operation without a human intervention it is many times wise to enable human involvement or even intervention to block unwanted results.

2.4 Australia

The Australian Government (AGA) uses the OECD's definition of an AI system: “An AI system is a machine-based system that, for explicit or implicit objectives, infers, from the input it receives, how to generate outputs such as predictions, content, recommendations, or decisions that can influence physical or virtual environments. Different AI systems vary in their levels of autonomy and adaptiveness after deployment.”¹¹

The Australian Government has launched consultations on safe and responsible artificial intelligence, finding that the current regulatory system fails to address the specific risks that artificial intelligence poses. Governments, at international level, are introducing new regulations to address the risks of AI, with a focus on creating preventative, risk-based guardrails that apply across the AI supply chain and throughout the AI lifecycle. Australian Government is committed developing a regulatory environment that builds community trust and promotes AI adoption. They opened for consultation 4 September 24, Close 4 October 24) (AI risk categorisation criteria // EU's). We want your views on the proposed guardrails, how we're proposing to define high-risk AI, regulatory options for mandating the guardrails. The proposed approach to use AI safely and responsibly when developing and deploying AI in Australia in high-risk settings. They aim to address risks and harms from AI, build public trust, provide businesses with greater regulatory certainty. Strictly regulate High Risk, free rein to Low Risk to support development (\$ 1.887,5 billion in 24 federal budgets).

Proposed mandatory guardrails for AI in high-risk settings. Proposed Principles are as follows:

- Throughout their lifecycles, AI systems should benefit individuals, society and the environment.
- Throughout their lifecycles, AI systems should respect human rights, diversity & the autonomy of individuals (termed human-centred vision).
- Throughout their lifecycles, AI systems should be inclusive & accessible & should not involve or result in unfair discrimination against individuals, communities or groups (termed Fairness)
- Throughout their lifecycles, AI systems should respect & uphold privacy rights and data protection & ensure the security of data.
- Throughout their lifecycles, AI systems should reliably operate in accordance with their intended purpose.

There should be transparency and responsible disclosure so people can understand when they are being significantly impacted by AI & can find out when an AI system is engaging with them.

When an AI system significantly impacts a person, community, group or environment, there should be a timely process to allow people to challenge the use or outcome of the AI system.

Those responsible for the different phases of the AI system lifecycle should be identifiable and accountable for the outcomes of the AI system, and human oversight should be enabled.

2.5 Canada

There are several definitions we will refer to the Department of National Defence (DND) and Canadian Air Force (CAF) definition: “AI is considered the capability of a computer to do things that are normally associated with human cognition, such as reasoning, learning, and self-improvement.”

¹¹ OECD's definition of an AI system, OECD ARTIFICIAL INTELLIGENCE PAPERS (2024), OECD March 2024 No. 8

With reference to the official documentation published by the Canadian Government, in September 2023, the Minister of Innovation, Science and Industry announced the Voluntary Code of Conduct on the Responsible Development and Management of Advanced Generative AI Systems. Development Focus (\$2.4 billion in 24 federal budget) to encourage capability and infrastructure. Not waiting to “fall further behind”, seeking economic benefit, capitalise on extensive clean energy grid. Artificial Intelligence and Data Act (AIDA) will introduce new requirements for businesses to ensure the safety and fairness of high-impact AI systems every step of the way:

Design: Businesses will be required to identify and address the risks of their AI system about harm and bias and to keep relevant records.

Development: Businesses will be required to assess the intended uses and limitations of their AI system and make sure users understand them.

Deployment: Businesses will be required to put in place appropriate risk mitigation strategies and ensure systems are continually monitored.

The idea is to have a flexible policy, where safety obligations are tailored to the type of AI systems. The more risks are associated with an AI system, the more obligations there will be.

These new regulations would build on existing best practices, with the intent to be interoperable with existing and future regulatory approaches. By drawing on common standards, the government is hoping to ease compliance. Consulting on best approaches with industry players (opened June 24, closed September 6, 2024).

2.6 China: National Level

China invested relevant resources to lead the competition in the AI domain, as a direct consequence China is leading the way in AI regulation releasing new strategies to govern algorithms, chatbots and more. One of the definitions of AI due to Chinese researchers is “*Artificial intelligence (AI) aims to mimic human cognitive functions and execute intellectual activities like that performed by humans dealing with an uncertain environment.*¹²”. The Chinese approach to this technology is in line with Chinese culture and foresees, in recent rules, that AI must always be under human control (since the advent of Generative AI as a precautionary approach). Positive duties for applications to promotes social cohesion, additional regulations prohibitions on algorithms that create division.

Early in 2021/2022 China developed & implemented detailed/binding AI regulation. These regulations foreseen duties on companies regarding content recommendations, rights for users recommended content. “*Worker protections against gig algorithmic scheduling were foreseen as well*”. The Cyberspace Administration of China, the National Development and Reform Commission, the Ministry of Education, the Ministry of Science and Technology, the Ministry of Industry and Information Technology, the Ministry of Public Security, and the National Radio and Television Administration jointly released the Interim Measures for the Management of Generative Artificial Intelligence Services (the “AI Measures”), which is the first administrative regulation on the management of Generative AI services, which came into effect on August 15, 2023.

On September 2024 China has proposed new regulations that would make it mandatory labelling of AI-Generated Content that tries to clamp down on a surge in AI-related fraud.

Rules requiring watermarks identifying deep fakes, protecting people’s “likeness rights” or harm the nation’s image- creating/amending an Algorithm registry.

Regulation created in consultation w/academics, bureaucrats, journalists, researchers & technologies – public debate & advocacy, workshopping etc. factors of policy making.

Draft guidelines were published on September 14 and open for public comment one month as stated by the “*Notice on soliciting opinions on the mandatory national standard (draft for comments) of “Network Security Technology Artificial Intelligence Generation Synthetic Content Identification Method”*”.

¹² Chen Shuang - School of Chinese Medicine, University of Hong Kong, Hong Kong 999077, China

This marks the first time that the Cyberspace Administration of China has proposed specific rules regarding the labelling of AI-generated content.

All relevant units and experts: According to the standard revision plan of the National Standardisation Administration Committee, the Office of the Central Network Security and Information Technology Commission has organised and completed the draft of the national standard of "*Network Security Technology Artificial Intelligence Generation and Synthetic Content Identification Method*", which is now publicly soliciting comments. Opinions were submitted to the drafting department of the organisation before November 13, 2024.

New labelling rules for AI-generated content, effective September 1, 2025, require explicit and implicit labels for all content types. China released three national standards on April 25, 2025, to enhance generative AI security and governance, effective November 1, 2025.

The three standards are:

- Cybersecurity Technology: Generative Artificial Intelligence Data Annotation Security Specification: These standard outlines security requirements for data labelling processes used in training generative AI models.
- Cybersecurity Technology: Security Specification for Generative Artificial Intelligence Pre-training and Fine-tuning Data: It outlines requirements and evaluation criteria for ensuring the security of datasets used in the pre-training and fine-tuning phases of generative AI development.
- Cybersecurity Technology: Basic Security Requirements for Generative Artificial Intelligence Service: This standard establishes security requirements for generative AI services, encompassing user data security assessments, data protection measures, and the safeguarding of training models and datasets.

2.6.1 Transparency rights for citizens whose rights are affected by AI applications

Most recently rules that AI must always be under human control (since the advent of Generative AI- a precautionary approach). Positive duties for applications to promotes social cohesion, additional regulations prohibitions on algorithms that create division. Care taken to encourage production & not discourage development/innovation. Regulation requiring algorithm transparency.

2.7 European Commission & Council of Europe

To better approach the regulatory framework, it is useful to start with the EU definition of AI¹³. “*AI system’ means a machine-based system that is designed to operate with varying levels of autonomy and that may exhibit adaptiveness after deployment, and that, for explicit or implicit objectives, infers, from the input it receives, how to generate outputs such as predictions, content, recommendations, or decisions that can influence physical or virtual environments.*”

The European Commission clearly distinguish between “*simpler traditional software systems or programming approaches ... that are based on the rules defined solely by natural persons to automatically execute operations. A key characteristic of AI systems is their capability to infer*”

The EC join Research Centre explored the opportunity AI-generated synthetic data presents for privacy-safe data use. In such specific field EU Regulations adopt a risk-based approach to regulation. AI should be as neutral as possible to cover techniques that are not yet known/developed.

On 17 May 2024, "*the Council of Europe adopted the first-ever international legally binding treaty*¹⁴ aimed at ensuring the respect of human rights, the rule of law and democracy legal standards in the use of artificial intelligence systems. The treaty, which is also open to non-European countries, sets out a legal framework that covers the entire lifecycle of AI systems and addresses the risks they may pose, while promoting responsible innovation. This document was the outcome of two years of work due to an intergovernmental body, the Committee on Artificial Intelligence (CAI) composed by 46

¹³ Source: Artificial Intelligence Act <https://artificialintelligenceact.eu/article/3/>

¹⁴ [https://search.coe.int/cm/#{%22CoEIdentifier%22: \[%220900001680afb11f%22\],%22sort%22:\[%22CoEValidationDate%20Descending%22\]}](https://search.coe.int/cm/#{%22CoEIdentifier%22: [%220900001680afb11f%22],%22sort%22:[%22CoEValidationDate%20Descending%22]})

CoE members states, 11 non-member states¹⁵, the European union plus the private sector, civil society and academia, who participated as observers. The convention adopts a risk-based approach to the design, development, use, and decommissioning of AI systems, which requires carefully considering any potential negative consequences of using AI systems."

Few days later, on 21 May, the Council of the European Union adopted the EU AI Act, which once published in the EU Official Journal in June, became the first set of AI regulations that has undergone a full legislative approval process.

The EU AI Act is structured in 113 articles and counts 13 annexes, its holistic risk-based approach is suitable for any player in the field of AI from developers to deployers.

The EU vision on AI can be summarised as *"beyond making our lives easier, AI is helping us to solve some of the world's biggest challenges: from treating chronic diseases or reducing fatality rates in traffic accidents to fighting climate change or anticipating cybersecurity threats."*¹⁶ Dealing with a fast evolving technology the key aspect in defining a framework is to find a golden balance between shaping the evolution and leaving to technology the freedom to evolve naturally.

On September 23, 2025, Politico¹⁷ issued an article stating that *"just one year after the European Union adopted a landmark plan to cut risks of artificial intelligence, it's already preparing to put the brakes on."* This decision was mainly due to big company CEOs lobbying and the unavailability of the EU standards defining high risks applications and the consequent impossibility to comply with such standards within the foreseen enforcement date in August 2025. Some EU governments supported the delay of at least twelve months to enforce the AI Act. On September 2025 United Nations General Assembly, the need to regulated AI was raised.

2.8 Finland

There is no established definition of artificial intelligence within Finnish law. In Finland, the EU AI Act is intended to be implemented through supplementary and clarifying provisions, without creating an independent regulatory framework. The current legislative proposals do not include definitions related to artificial intelligence. Since the EU AI Act is directly applicable law in Finland, it is expected that the definitions contained in the EU AI Act will become established in Finnish law as such. The Finnish Centre for Artificial Intelligence (FCAI) provides this vision *"Artificial intelligence is the subject of great hopes but also fears, with demands for greater regulation increasing as we speak. All the while, we don't even have a universal definition for it."*

Split focus (eventually) initially (2017) focused on business opportunities & competitiveness, economic benefits/growth, evolved through 2018 & 19 & 20 to include and supporting a high quality & efficient public service, and ensuring a functioning societal wellbeing of citizens (economic & other kinds of wellbeing).

2.9 India

In India, Artificial Intelligence (AI) refers to *"the development of computer systems that emulate human intelligence, performing tasks like thinking, learning, and decision-making to enhance human life and work"*. The office of the Principal Scientific Adviser to the Government of India introduces AI as *"Artificial Intelligence (AI) offers transformative opportunities to complement human intelligence and enhance how we live and work. With its vast and growing range of applications, AI and machine learning are now deeply embedded in nearly every facet of modern technology."*¹⁸

This definition aligns with global understanding but is often contextualized within India's national strategy, which focuses on applying AI to societal challenges in areas such as healthcare, agriculture,

¹⁵ Argentina, Australia, Canada, Costa Rica, the Holy See, Israel, Japan, Mexico, Peru, the United States of America, and Uruguay

¹⁶ Artificial Intelligence for Europe

¹⁷ <https://www.politico.com>

¹⁸ <https://www.psa.gov.in/ai-mission>

and education to drive inclusive growth and societal transformation, as outlined by NITI Aayog¹⁹ and the Ministry of Electronics and Information Technology (MeitY)²⁰.

2.9.1 Key Aspects of India's Perspective on AI:

India's AI approach focuses on complementing human intelligence, addressing societal needs, and driving development through government initiatives like the National Strategy for Artificial Intelligence²¹.

- Emulation of Human Intelligence: AI is understood as machines that can think, perceive, learn, solve problems, and make decisions like humans.
- Complementing Human Capabilities: A core principle is that AI should work alongside and enhance human intelligence, rather than replace it.
- Focus on Societal Impact: India emphasizes AI's role in addressing societal needs and driving development, particularly in sectors like healthcare, agriculture, and smart cities.
- Strategic National Approach: The Indian government, through initiatives like the National Strategy for Artificial Intelligence and programs by MeitY²², aims to promote AI adoption and innovation for inclusive growth and transformation.
- Data-Driven Learning: AI models in India are developed to learn from examples and contextual information, gathering data and interacting with environments to improve task performance.
- Interdisciplinary Foundation: The concept draws from various fields, including computer science, mathematics, and specialized expertise.

India regulates AI through existing frameworks like the Personal Data Protection Act (DPDP) 2023²³ and IT Act and the Information Technology (IT) Act, 2000²⁴, emphasizing a pro-innovation approach while balancing ethical considerations like fairness, accountability, and transparency for AI systems. An additional reference Act is the Consumer Protection Act 2019²⁵. AI services can fall under the definition of "services" under this act, potentially holding providers liable for harms from biased algorithms or faulty systems. The government also issues advisories, such as the MeitY advisory²⁶, requiring permission for deploying certain AI models and mandating steps against algorithmic discrimination and deepfakes. Specific reference structure, apart from the already mentioned ones is IndiaAI²⁷ and the related India AI Governance Guidelines²⁸. The broader strategy emphasizes a pro-innovation approach, balancing rapid development with ethical considerations like fairness, accountability, and transparency for AI systems.

Government Advisories: The MeitY advisory requires platforms to obtain government permission before deploying under-tested or unreliable AI models. It also mandates measures to prevent algorithmic discrimination and the spread of deepfakes.

Strategic Policies: National Strategy for Artificial Intelligence (2018)²⁹ developed by NITI Aayog, this strategy outlined a vision of "AI for All," emphasizing inclusivity and AI's application in priority

¹⁹ <https://niti.gov.in/>

²⁰ <https://www.meity.gov.in/>

²¹ <https://static.pib.gov.in/WriteReadData/specificdocs/documents/2025/nov/doc2025115685601.pdf>

²² <https://www.meity.gov.in/>

²³ <https://www.meity.gov.in/static/uploads/2024/06/2bflf0e9f04e6fb4f8fef35e82c42aa5.pdf>

²⁴ https://www.indiacode.nic.in/bitstream/123456789/13116/1/it_act_2000_updated.pdf

²⁵ https://ncdr.nic.in/bare_acts/CPA2019.pdf

²⁶ <https://www.meity.gov.in/static/uploads/2024/02/9f6e99572739a3024c9cdaec53a0a0ef.pdf>

²⁷ <https://indiaai.gov.in/indiaaiportal>

²⁸ <https://static.pib.gov.in/WriteReadData/specificdocs/documents/2025/nov/doc2025115685601.pdf>

²⁹ <https://www.niti.gov.in/sites/default/files/2023-03/National-Strategy-for-Artificial-Intelligence.pdf>

sectors. Additional regulations are the Principles for Responsible AI (2021)³⁰ issued by MeitY, these principles establish ethical guidelines for fairness, accountability, and transparency in AI.

A couple of additional initiatives are already planned and in development phase: the Digital India Act³¹ (expected in 2025) a proposed law that aims to regulate India's digital space more effectively probably replacing the IT Act. It seeks to address cybersecurity, online privacy, content moderation, data governance, and the responsibilities of digital platforms, which the old IT Act couldn't fully cover. Lastly the IndiaAI Mission³² a government mission to foster a trustworthy AI ecosystem through initiatives focused on foundational models and AI skills.

2.10 Israel

Israel refers to the updated OECD definition “*An artificial intelligence system is a machine-based system that, for explicit or implicit objectives, infers from the input it receives how to generate outputs such as predictions, content, recommendations, or decisions that can influence physical or virtual environments. AI systems vary in their level of autonomy and adaptability after deployment.*”

Israel's Policy on Artificial Intelligence Regulation and Ethics³³. Israel unveiled its first comprehensive AI policy, led by the Ministry of Innovation, Science and Technology and the Ministry of Justice. The policy aims to foster responsible AI innovation while addressing concerns related to bias, transparency, safety, accountability, and privacy. Ministry of Innovation, Science and Technology and the Ministry of Justice Take Lead in Addressing AI Challenges with Responsible Innovation Policy

Minister of Innovation, Science and Technology Ofir Akunis: " *The revolutionary impact of AI technology is yielding constant improvements in the quality of life of Israel's citizens in countless fields of activity, including in health care, science, technology, agriculture, education, and security. However, together with its many benefits, there are also many risks. These principles we have published facilitate development and responsible innovation, enabling the use of AI, while safeguarding basic rights and the public interest.*"

Meir Levin, Deputy Attorney General (Economic Law), Ministry of Justice: “*In today's rapidly evolving technological landscape, AI presents unprecedented opportunities for our society and economy. Responsibly integrating AI into various aspects of life necessitates adapting our legal and regulatory frameworks. We must ensure that the integration of AI aligns with the public interest and human rights while avoiding unnecessary barriers that could hinder the local industry and deprive the public of innovative products and services. The recommendations in the report represent a significant milestone, providing guidance for the government's future work in AI regulation. I commend the collaborative efforts with the Ministry of Innovation, Science, and Technology, as well as our partners in the Ministry of Justice. I look forward to continuing this important work in the field.*”

Jerusalem 17th December 2023 – Israel has unveiled its first-ever comprehensive policy on Artificial Intelligence (AI) regulation and ethics, demonstrating a commitment to responsible innovation and addressing the complex challenges associated with the widespread use of AI in various sectors. The policy, officially launched today, is the fruit of collaborative efforts led by the Ministry of Innovation, Science and Technology, in conjunction with the Office of Legal Counsel and Legislative Affairs (Economic Law Department) at the Ministry of Justice. As AI systems become integral to industries worldwide, Israel recognizes the need for proactive measures to ensure responsible use and mitigate potential risks. The policy, which builds upon extensive consultations with government departments,

³⁰ <https://www.niti.gov.in/sites/default/files/2021-02/Responsible-AI-22022021.pdf>

³¹ <https://www.delhilawacademy.com/digital-india-act/#:~:text=The%20Digital%20India%20Act%2C%202023,Act%20couldn't%20fully%20cover.>

³² <https://www.india.gov.in/website-indiaai-mission>

³³ <https://aiisrael.org.il>

civil society organizations, academia, and the private sector, establishes a framework to foster innovation while addressing concerns related to bias, transparency, safety, accountability, and privacy.

2.10.1 Key Highlights of Israel's AI Policy:

Israel's AI Policy addresses seven challenges of private sector AI use, including discrimination and privacy, through collaboration with stakeholders and alignment with OECD recommendations. The policy emphasizes "Responsible Innovation," balancing innovation and ethics.

Comprehensive approach: The AI Policy identifies seven main challenges arising from private sector

- (1) AI use: discrimination, human oversight, explainability, disclosure of AI interactions, safety, accountability, and privacy.

- (2) Collaborative development: The AI Policy is formulated through collaboration with various stakeholders, including the Ministry of Innovation, Science and Technology, the Ministry of Justice, the Israel Innovation Authority, the Privacy Protection Authority, the Israeli National Cyber Directorate, the Israel National Digital Agency, leading AI companies in Israel and academia.

- (3) Policy principles: Aligned with the OECD AI Recommendations, Israel's AI Policy outlines common policy principles and practical recommendations to address challenges and foster responsible AI innovation.

- (4) Responsible innovation concept: Emphasizing the concept of "Responsible Innovation," the AI Policy aims to balance innovation and ethics, viewing them as synergistic rather than conflicting goals.

2.10.2 Main Policy Recommendations:

The AI Policy recommends empowering sector-specific regulators, adopting a risk-based approach, and promoting multistakeholder cooperation. It also suggests establishing an AI Policy Coordination Centre and forums for regulators and public participation.

- Government Policy Framework: The AI Policy recommends empowering sector-specific regulators, fostering international interoperability, adopting a risk-based approach, encouraging incremental development, using "soft" regulation, and promoting multistakeholder cooperation.
- Ethical AI Principles: The AI Policy suggests adopting a common set of ethical AI principles based on the OECD AI Principles, focusing on human-centric innovation, equality, transparency, reliability, accountability, and promoting sustainable development.
- AI Policy Coordination Centre: It proposes to establish an expert-based inter-agency body to advise sectoral regulators, promote coordination, update government AI policy, advise on AI regulation, and represent Israel in international forums.
- Mapping AI Uses and Challenges: The AI Policy urges government entities and regulators to identify and map the uses of AI systems and associated challenges in regulated sectors.
- Forums for Regulators and Public Participation: It recommends establishing forums for regulators and multistakeholder participation to promote coordination, coherence, and open discussions.
- International Collaboration: It encourages active involvement in developing international regulation and standards to promote global interoperability.
- Tools for Responsible AI: The policy calls for the development of tools, including a risk management tool, to enable responsible AI development and use, with the AI Policy Coordination Centre leading this effort.

2.10.3 To summarise Israel's AI Policy by topic:

Core Aim

- Ensure responsible and ethical AI innovation.
- Prioritize human-centric development while safeguarding societal values.

- Strengthen Israel's technological leadership through clear regulatory guidance.

Government Commitment

- Adopted by the Minister of Innovation, Science and Technology, Ofir Akunis.
- Policy emphasizes collaboration, transparency, and adaptability to foster a thriving AI ecosystem.
- Plans to formalize the policy through a Government Decision, ensuring regulators have clear instructions for addressing AI-related challenges.

Benefits & Risks

- AI is improving quality of life in healthcare, science, agriculture, education, and security.
- Risks must be managed by keeping humans at the centre of AI development.
- Shared principles will guide responsible innovation while protecting fundamental rights such as privacy and equality.

Legal & Regulatory Framework

- Meir Levin, Deputy Attorney General, highlights the need to adapt legal frameworks to integrate AI responsibly.
- The policy seeks to balance innovation with public interest and human rights.
- Avoids unnecessary barriers that could hinder local industry or limit access to innovative products.

Significance

- Marks a milestone in Israel's approach to AI regulation and ethics.
- Provides guidance for future government work in AI governance.
- Sets a precedent for other nations grappling with similar challenges.

2.11 Japan

Official definition of Artificial Intelligence in Japan³⁴: *"artificial intelligence-related technology" refers to the technology necessary to realize the function of replacing the intellectual ability related to human cognition, reasoning and judgment by artificial methods, and the information entered is processed using the technology, and It refers to the technology related to the information processing system to realize the function of outputting the results.*³⁵

The Japanese approach to AI governance doesn't impose rigid obligations and heavy fines, the AI Promotion Act is based on national objectives, relies on guidance, ad hoc coordination structures, specific laws to ensure responsible AI use. This AI governance model encourages AI adoption while managing risks through existing frameworks, this is rooted in the innovation first model. This law is based on the fact that artificial intelligence-related technologies are the foundation of the development of our country's economy and society, and on the measures to promote the research and development and utilization of artificial intelligence-related technologies. The effective utilization of artificial intelligence-related technologies will bring remarkable efficiency and advancement of administrative affairs and private business activities and the creation of new industries.

Key responsibilities: The government is responsible for formulating and implementing measures related to the promotion of research and development and utilization of artificial intelligence-related technologies in a comprehensive and planned manner. The State shall promote the active use of artificial intelligence-related technologies in the administrative organs of the country to achieve the efficiency and upgrading of administrative affairs. Local public organizations shall, in line with the basic concept and promote the research and development and utilization of artificial intelligence-related technologies make use of the regional characteristics of their local public organizations. It has the responsibility to formulate and implement policies. Universities, R&D corporation and other institutions that carry out research and development of artificial intelligence-related technologies. In addition to the basic philosophy, they can perform the research and development of artificial intelligence-related technologies and the popularization of their results, as well as the cultivation of

³⁴ <https://laws.e-gov.go.jp/law/507AC0000000053>

³⁵ https://laws.e-gov.go.jp/law/507AC0000000053#Mp-Ch_1-At_2

professional and broad knowledgeable personnel. Those who intend to use other artificial intelligence-related technologies in business activities in addition to the basic philosophy and the use of their own active artificial intelligence-related technologies, local public organizations shall strive to improve the efficiency and advancement of business activities and the creation of new industries and comply with the measures implemented by the country.

2.12 Russian Federation

The official definition of Artificial Intelligence in the Russian Federation is “*A set of technological solutions that makes it possible to simulate human cognitive functions*”. In 2019 the President of Russian Federation described Artificial Intelligence as a technology that will allow quick real-time decision-making based on analysing vast amounts of information known as big data, which provides tremendous advantages in terms of quality and performance. In addition, such mechanisms are unparalleled in history in terms of their impact on the economy and productivity, the effectiveness of management, education, healthcare and daily life.

The foundation for Russian Federation National Strategy on the development of artificial intelligence technologies is defined in comprehensive document the key principles can be summarised as follows³⁶:

Fundamental Groundwork

- Priority is to establish new conceptual foundations for AI.
- Develop mathematical methods and operational principles, including those inspired by the human brain.

Global Scientific Collaboration

- Position Russia as a leading platform for solving complex scientific challenges.
- Engage scientists from around the world in collaborative research.

Research Infrastructure

- Open international mathematics centres in Moscow, St Petersburg, and Sochi to complement existing facilities.

Funding & Investment

- Significantly increase funding for AI research.
- Encourage private investment, corporate science, and R&D through incentives.

Workforce & Talent Development

- Strengthen intellectual potential by retaining domestic talent and attracting global professionals.
- Expand AI study programs at universities and colleges, including regional institutions.
- Provide monetary incentives, flexible working conditions (including remote work), and streamlined formalities for citizenship and work permits to encourage experts—especially Russians abroad—to return and contribute.

Legislation & Legal Framework

- Modernize laws to align with technological realities.
- Create flexible regulations supporting AI development and private investment in breakthrough solutions.
- Ensure strong intellectual property protection and efficient patent registration within Russia’s jurisdiction.
- Remove legislative and administrative barriers while safeguarding state security, societal interests, and citizens’ rights.

Data Infrastructure & Regulation

- Establish efficient legal frameworks for data circulation.
- Build advanced infrastructure for data processing and storage.

³⁶ <http://www.en.kremlin.ru/events/president/transcripts/copy/60630>

- Develop balanced solutions enabling data use for AI algorithm training.

Societal Readiness & Education

- Promote mass digital literacy.
- Implement re-training programs, especially for jobs in high demand.
- Prepare society for widespread adoption of AI technologies.

2.13 Saudi Arabia

Saudi Arabia was one of the first country to establish the National Centre for AI (NCAI). Their definition of AI is: *“AI is systems that use technologies capable of making predictions, generating content, providing recommendations, or making decisions with different levels of autonomy.”*

Saudi Arabia aims to implement global best practices for data management and governance to improve performance and support economic transformations. The Saudi Data & AI Authority and the National Data Management Office have created a national data governance framework and standards for adoption by all Public Entities.

The regulatory arrangement of the Saudi Data & AI Authority's Laws and Regulations (SDAIA) were issued by Council of Ministers Resolution No. (292) dated 27/4/1441 AH and were amended by Council of Ministers Resolution No. (195) dated 15/3/1444 AH.

Saudi Arabia aims to implement global best practices for data management and governance to improve performance, productivity, and public services. The Kingdom seeks to leverage data as an economic resource to foster innovation, support economic transformations, and enhance national competitiveness. The Saudi Data & AI Authority and the National Data Management Office have created a national data governance framework, approved by the SDAIA Board of Directors.

The National Data Management and Personal Data Protection Standards, developed by the National Data Management Office, are based on a directive from the Saudi Authority for Data and Artificial Intelligence. These standards, covering 15 domains, are intended for adoption by all Public Entities in the Kingdom.

To support the development of the Data Management and Personal Data Protection standards, a set of international references, internal relevant policies and regulations, and guiding principles were defined. Government Entities must implement the standards, and compliance will be measured yearly to monitor progress and drive efforts towards a successful implementation.

The Data Management and Personal Data Protection Standards include 15 domains, covering data governance, quality, operations, security, and more. The Data Management and Personal Data Protection Standards³⁷ key content are:

- Specifications Prioritization
- Prioritization Details
- Implementation Plan

Kingdom of South Arabia Data Management and Personal Data Protection Standards

1. Data Governance Domain
2. Data Catalog and Metadata Domain
3. Data Quality Domain
4. Data Operations Domain
5. Document and Content Management Domain
6. Data Architecture and Modelling Domain
7. Data Sharing and Interoperability Domain
8. Reference and Master Data Management Domain
9. Business Intelligence and Analytics Domain
10. Data Value Realization Domain

³⁷ National Data Management Office Data Management and Personal Data Protection Standards - Version 1.5 - January 2021 - <https://sdaia.gov.sa/en/SDAIA/about/Documents/DataManagementPersonalDataProtectionStandards.pdf>

11. Open Data Domain
12. Freedom of Information Domain
13. Data Classification Domain
14. Personal Data Protection Domain
15. Data Security and Protection Domain

2.14 South Africa

In South Africa, Artificial Intelligence (AI) is defined as technology that allows machines to sense, comprehend, and act, emulating human intelligence to perform tasks like reasoning, perception, and learning. South African authorities are developing frameworks to address the ethical, social, and economic implications of AI, such as job creation, consumer protection, and accountability for AI-driven decisions.

Core Definition of AI

“AI is a branch of computer science focused on creating machines capable of tasks that typically require human intelligence. This includes capabilities such as reasoning, perception, learning, and even creating new solutions, all driven by pattern recognition in large datasets.”

2.14.1 South Africa's Regulatory Approach

- **Ethical Framework:** South Africa is working on a comprehensive AI-specific legal framework to govern its deployment, with a focus on addressing accountability and ethical concerns.
- **Human-Centred Approach:** The "AI Blueprint" emphasizes technology that enables machines to emulate human capabilities to "sense, comprehend and act," with the goal of making decisions based on absorbed information.
- **Consumer Protection:** A key aspect of the developing regulations is to ensure remedies, consumer protection, and clear lines of responsibility for actions taken by AI systems.

2.14.2 Economic and Societal Considerations

- **Economic Growth and Efficiency:** AI is seen as a tool to drive economic growth and improve efficiencies by enabling machines to "do our jobs better".
- **Skills Development:** The evolving world of work requires South Africans to continuously acquire new knowledge, competencies, and 21st-century skills to adapt to AI's impact on careers.
- **Inclusivity and Access:** There is an emphasis on inclusivity, aiming to create more access to AI technology for everyone and foster collective growth, especially in low-income communities.

2.15 South Korea

South Korean's AI Definition - AI is defined as technology that mimics human-level intelligence, encompassing capabilities like learning, reasoning, and language comprehension. This broad definition helps differentiate between AI types, such as “high-impact AI” and “generative AI,” subject to stricter rules. **Electronic implementation of human intellect:** AI is seen as a technology that mimics or carries out tasks that would normally require human-level intelligence. The definition lists specific capabilities like learning, reasoning, perception, judgment, and language comprehension. This definition is intentionally broad to encompass various AI technologies, from simple algorithms to complex machine learning models. This definition helps the AI Basic Act differentiate between different types of AI, such as "high-impact AI" and "generative AI," which are subject to more stringent rules.

2.15.1 Artificial Intelligence (AI) Basic Act

South Korea's recently enacted Act on the Development of Artificial Intelligence and Establishment of Trust (AI Basic Act)—set to take effect in January 2026, establishes a legal framework for AI

development, prioritizing ethical standards and public trust. The act mandates transparency and safety responsibilities for businesses developing “high impact” AI, including risk assessments and safety measures. It marks a pivotal development for U.S. companies operating in or entering the Korean Artificial Intelligence (AI) market. As the second country after the European Union to adopt a comprehensive AI regulatory framework, Korea is positioning itself as a global leader in trustworthy and innovative AI, while creating both opportunities and new compliance requirements for U.S. companies.

South Korea passed the AI Basic Act in December 2024, intended to provide a legal framework to advance Korea’s national competitiveness in AI while ensuring ethical standards and public trust. The act establishes legal grounds for establishing a national AI control tower, an AI safety institute, and various governmental initiatives in R&D, standardization, and policies. The legislation also mandates separate initiatives to support the national AI infrastructures, including training data and data centres, and fostering SMEs, startups, and talent in the AI field. More importantly, the act assigns transparency and safety responsibilities to businesses that develop and deploy “high impact” AI and generative AI. For example, the law requires businesses to implement AI risk assessment, a set of safety measures, and the designation of a local representative.

The ROK Ministry of Science and ICT is drafting regulations for the AI Basic Act, expected in early 2025. Public discussions and consultations with industry stakeholders may occur.

The ROK Ministry of Science and ICT (MSIT) is currently drafting the subordinate regulations, which are expected to be released in the first half of 2025. Although the Minister of MSIT recently (April 2025) reaffirmed the government’s position is to keep minimum regulations, this law may present both an opportunity and a potential regulatory challenge for American AI businesses operating in South Korea. As detailed guidelines are being developed, public discussions or consultations with industry stakeholders may also take place. If you have any comments regarding the AI Basic Act or the upcoming subordinate regulations, you are encouraged to contact the Ministry of Science and ICT.

2.16 United States: Federal executive, federal legislative & State Level (mixed)

Starting again from the official definition on the legal side we can refer to the “Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence”³⁸

“The term “artificial intelligence” or “AI” has the meaning set forth in 15 U.S.C. 9401(3): a machine-based system that can, for a given set of human-defined objectives, make predictions, recommendations, or decisions influencing real or virtual environments. Artificial intelligence systems use machine- and human-based inputs to perceive real and virtual environments; abstract such perceptions into models through analysis in an automated manner; and use model inference to formulate options for information or action.”

While “the term “AI model” means a component of an information system that implements AI technology and uses computational, statistical, or machine-learning techniques to produce outputs from a given set of inputs.”

At State level: basic requirements are non-discrimination /fairness concerns in States which have AI specific regulations, many states have no AI specific regulations, but ordinary consumer law could be applied in appropriate situations. On 7 September 24 California Governor Gavin Newsom signed three measures to remove deceptive content from large online platforms, increase accountability, and better inform voters. This enacted law criminalising deep fakes which could influence the current election, effective immediately for a period of 120 days before and 60 days after the election. *“Safeguarding the integrity of elections is essential to democracy, and it’s critical that we ensure AI is not deployed to undermine the public’s trust through disinformation – especially in today’s fraught*

³⁸ <https://www.whitehouse.gov/briefing-room/presidential-actions/2023/10/30/executive-order-on-the-safe-secure-and-trustworthy-development-and-use-of-artificial-intelligence/>

political climate. These measures will help to combat the harmful use of deepfakes in political ads and other content, one of several areas in which the state is being proactive to foster transparent and trustworthy AI.” [Governor Gavin Newsom].

This bill, to be known as the Defending Democracy from Deepfake Deception Act of 2024, would require a large online platform, as defined, to block the posting of materially deceptive content related to elections in California, during specified periods before and after an election.

At Federal level: October 20, 2023, Presidential Executive Order - Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence. “*Section 1. Purpose. Artificial intelligence (AI) holds extraordinary potential for both promise and peril. Responsible AI use has the potential to help solve urgent challenges while making our world more prosperous, productive, innovative, and secure*³⁹. *At the same time, irresponsible use could exacerbate societal harms such as fraud, discrimination, bias, and disinformation; displace and disempower workers; stifle competition; and pose risks to national security. Harnessing AI for good and realizing its myriad benefits requires mitigating its substantial risks. This endeavour demands a society-wide effort that includes government, the private sector, academia, and civil society.*”

The executive order identifies eight guiding principles and priorities to advance and govern the development and use of AI:

- [1.] Artificial Intelligence must be safe and secure.
- [2.] Promoting responsible innovation, competition, and collaboration will allow the United States to lead in AI and unlock the technology’s potential to solve some of society’s most difficult challenges.
- [3.] The responsible development and use of AI require a commitment to supporting American workers.
- [4.] Artificial Intelligence policies must be consistent with my Administration’s dedication to advancing equity and civil rights.
- [5.] The interests of Americans who increasingly use, interact with, or purchase AI and AI-enabled products in their daily lives must be protected.
- [6.] Americans’ privacy and civil liberties must be protected as AI continues advancing.
- [7.] It is important to manage the risks from the Federal Government’s own use of AI and increase its internal capacity to regulate, govern, and support responsible use of AI to deliver better results for Americans.
- [8.] The Federal Government should lead the way to global societal, economic, and technological progress, as the United States has in previous eras of disruptive innovation and change.

The executive order refers to different “agencies”⁴⁰ involved in the application of AI technology e.g. Federal law enforcement agency⁴¹, Sector Risk Management Agency⁴².

With specific reference to AI in Critical Infrastructure and Cybersecurity to ensure adequate protection the following action shall be taken:

”Within 90 days of the date of this order, and at least annually thereafter, the head of each agency with relevant regulatory authority over critical infrastructure and the heads of relevant SRMAs, in coordination with the Director of the Cybersecurity and Infrastructure Security Agency within the Department of Homeland Security for consideration of cross-sector risks, shall evaluate and provide to the Secretary of Homeland Security an assessment of potential risks related to the use of AI in critical infrastructure sectors involved, including ways in which deploying AI may make critical infrastructure systems more vulnerable to critical failures, physical attacks, and cyber-attacks, and

³⁹ “Highlights of the 2023 Executive Order on Artificial Intelligence for ...”

⁴⁰ The term “agency” means each agency described in 44 U.S.C. 3502(1), except for the independent regulatory agencies described in 44 U.S.C. 3502(5).

⁴¹ The term “Federal law enforcement agency” has the meaning set forth in section 21(a) of Executive Order 14074 of May 25, 2022

⁴² The term “Sector Risk Management Agency” has the meaning set forth in 6 U.S.C. 650(23).

shall consider ways to mitigate these vulnerabilities. Independent regulatory agencies are encouraged, as they deem appropriate, to contribute to sector-specific risk assessments⁴³”.

This is simply one of the first actions to be taken, within 150 days the Secretary of the Treasury shall issue a public report on best practices for financial institutions to manage AI-specific cybersecurity risks, within 180 days the Secretary of Homeland Security, in coordination with the Secretary of Commerce and with SRMAs and other regulators as determined by the Secretary of Homeland Security, shall incorporate as appropriate the AI Risk Management Framework, NIST AI 100-1, as well as other appropriate security guidance, into relevant safety and security guidelines for use by critical infrastructure owners and operators. Within 240 days of the completion of the guidelines described above. In addition, the Secretary of Homeland Security shall establish an Artificial Intelligence Safety and Security Board as an advisory committee. The Advisory Committee shall include AI experts from the private sector, academia, and government, as appropriate, and provide to the Secretary of Homeland Security and the Federal Government’s critical infrastructure community advice, information, or recommendations for improving security, resilience, and incident response related to AI usage in critical infrastructure.⁴⁴

There are more actions and duties duly described in the Executive Order; the description provided above will describe the approach of the White House in securing the nation against potential threats.

Nevertheless, yet there is no AI specific statute. The 2023 Senate Federal AI Risk Management Bill⁴⁵ is still pending. This bill directs federal agencies to use the Artificial Intelligence Risk Management Framework developed by the National Institute of Standards and Technology (NIST) regarding the use of artificial intelligence (AI).

Exec order declares that AI regulation is a moral duty it:

- Directs application of the Defence Production Act 1950⁴⁶ to require that companies working on an AI model “that poses a serious risk to national security, national economic security, or national public health” to share the results of safety tests with the federal government to ensure those models are not used for malignant purposes.
- Directs that the Department of Commerce will devote “guidance for content authentication and watermarking to clearly label AI-generated content” which federal agencies will be expected to use.
- Instructs federal agencies to screen for bias in model related to housing applications, criminal justice settings and other institutions.

3. Will a global regulatory framework for AI become a reality?

The willingness to create a regulatory framework devoted to Artificial Intelligence is already a reality, several conferences or spokespersons pertaining completely different sectors are concerned about the potential impact of AI on their activities and businesses. Some of them are concerned about the impact on job position and careers other are concerned about human rights or ethics or the possible full control on our lives and more. Regulatory bodies are usually concerned, facing emerging technologies not yet fully developed, to influence their development biasing the potential natural evolution. Typical concerns are related to ethics and human rights, an open discussion is identifying potential concerns like biases, lack of shared values to identify a potential solution. Additional concerns are addressing potential AI hallucinations and the lack of explainability of AI outcomes together with the human supervision of automated decision making. Human supervision many times conflicts with the “real time” decision making process imposed by some activities. Last but already impacting nowadays the unsupervised use of AI by minors establishing an uncontrolled, by adults, trust relation with the

⁴³ Safe, Secure, and Trustworthy Development and Use of Artificial ...

⁴⁴ Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence <https://www.whitehouse.gov/briefing-room/presidential-actions/2023/10/30/executive-order-on-the-safe-secure-and-trustworthy-development-and-use-of-artificial-intelligence/>

⁴⁵ Sponsor Sen. Moran, Jerry [R-KS] (Introduced 11/02/2023)

⁴⁶ https://www.fema.gov/sites/default/files/2020-03/Defense_Production_Act_2018.pdf

system. As it already happened in different digital “global” domains as IPRs, privacy or cybersecurity dealing with global impact the expectation is to release or promulgate a global regulation framework. Currently, each country or Institution is promoting or enacting its own solution. Sometimes approaches to the issue differs other times refer to similar principles. Some research groups propose to create an AI Agency graduated “AI” agency-based liability regulating with an “AS IF” fiction:

- Level Zero – non generative AI – Preventative programming restrictions along the lines of Environmental, Discrimination, Privacy Regulation, penalties based on harm/level of programming malfeasance,
- Level One – Generative or semi-independent AI, or AI as independent actor in the sense of other nonhuman independent actors with less “autonomous agency leeway”. Animal law, preventive programming restrictions & responsiveness (//training’ responsibilities) (NZ Animal Welfare Act 1999 has provision for Governor General to declare no listed entries as “animals”),
- Responsibility for harm lies with creator or custodian of “animal” – penalties lie with both AI and creator/custodians, confinement (redesign) or destruction for the former, financial or custodial penalties for the latter, up through manslaughter and beyond, depending on programming malfeasance,
- Level Two – AI as independent actor in the sense of human independent actors – AI creating new AI – should this trigger morally based responsibilities and processes around control/penalties?

4. Conclusions

The shared willingness among private and public subject to create a regulatory framework devoted to Artificial Intelligence is almost a reality, as reported above several countries and even organisations developed their guidelines or even regulations. Several conferences or spokespersons pertaining completely different sectors are concerned, as it happened on the dawn of new technologies, about the potential impact of AI on their activities and businesses.

The intrinsic nature of “digital” immaterial, cloneable, and borderless suggests finding an agreement on a global regulatory framework addressing AI in its the different phases, design, development and deployment, in any case regulations will carefully consider the different cultural models to do not flatten the outcome to a single cultural model.