

Is there a need to create a global AI 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

Abstract: Starting from the background and the state of the art of the digital landscape, the next paragraph describes some of the main concerns associated with AI and related technologies, this paves the way to deal with the already established or proposed regulatory frameworks in the main countries that invested or are investing huge efforts in these technologies to become leaders world-wide. Closing remarks depicts citizens' standpoints and concerns before approaching possible future trends.

Keywords: Artificial Intelligence, Machine Learning, Digital Transformation. E-Services, Human Rights, Privacy.

1. Background and landscape

The omnipresent digital technology was considered one of the building blocks of this “new global order” [1 – Schwab – 2016] even thanks to the relevant contribution that this technology provided on one of the recent crises: the pandemic. Hence, one of the main vectors of this change was associated to the so-called Digital transformation (DT). The incredibly rapid success of the “Internet”, mainly due to e-commerce, information services, and social media, gave a boost to the globalization trend, a shift toward uniformity, jeopardizing diversities, broadcasting the vision of Cyber Majors [2 – UNESCO-IFAP- 2023]. Social media, global content providers are “training” young generations offering a *pensée unique* approach¹, this will impact future generations and their cultural identity. We remember the first attempt to create a universal “Encyclopaedia” by Microsoft, this project faced rapidly a problem due to the diverse cultural models and backgrounds², far later Wikipedia created “customised by cultural model” description of any entry.

Globalisation and the economic model carried out in the recent past show their limits [3 - Imad A. Moosa, 2023]. Nowadays more often we consider de-globalisation [4 WEF, 2023] as a scenario and the re-discover of local “values” and “identities” [5- ECB, 2023]. We do not like anymore the same sandwich or merchandise all-over the world. In the western world there is an evident lack of values and beliefs [6 - Eckersley -1995], a clear feeling that there is something “wrong” so in such an uncertain environment without clear references young generations are looking for new “gurus” and beliefs. The “cancel culture” movement together with the “politically correct” re-evaluation of history and facts in the light of today’s trends and thoughts [7 - Vogels – 2023, 8 - Dudenhoefer 2020] are some of the examples.

Digital transition can be associated with a progressive switch from traditional procedure and tools toward a cyber mediated set of activities and procedures. Leveraging on laziness and relaxation citizens spend less time outside home, they have shopping online, they buy food and drinks directly delivered on their table, “meet” friends on Zoom or WhatsApp, interact with the “outer environment” though social media and video clips. These aspects are even more evident in young generations that add to the social media the gaming dimension.

¹ *pensée unique* based on a single cultural model

² Telecommunication was invented by G. Bell, A. Meucci or A. Popov?

Technological advances provided ever-improving information processing and communication infrastructures, big data analytics offers the opportunity to identify data patterns that, “invisible” to the humans, can significantly help to better planning and find solutions. Of course, we must carefully consider the benefits and drawbacks due to these changes both in citizens behaviour and the accomplishment of tasks and duties.

2. Artificial Intelligence and typical main concerns

Among the different technologies offered on the shelf of digital transition one of the most promising and concerning is surely Artificial Intelligence (AI) and its rich set of sister applications.

As recalled in a different paper, we don’t know if the concerns often related to AI, considered a potential cyber-Leviathan dominating humanity thanks to its high “intelligence far exceeding any human one, is due to the term that was associated long time ago to the early studies and applications in this field. The term artificial intelligence generates the idea that this technology will compete and exceed the human one thanks to the ability to self-improve even recursively. Apart from the nightmares often generated by Sci Fiction movies like, War Games, Eagle Eye or even Lucy³, 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 Sam 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? (assigning legal personality does not address the issues McDonald’s is a legal person, but is not morally considerate)
- 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 secrets, biosecurity),
- Discrimination and BIAS in health, any selection process, procurement, justice, policy, regulation, broadly anything AI may be involved with,
- AI Moral responsibility – feel released from a personal ethical analysis and related responsibilities,
- Information - Oversight – Misinformation, disinformation, disapproved information filtering, nudging, opinion formation,

³ 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)

- 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 (down sides v/s competitive or innovative advantage),
- Graduated “AI” focused liability as with other independent actors.

3. Artificial Intelligence and regulatory frameworks

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

3.1 UNESCO

UNESCO launched the AI Initiative⁴ and published different reports on this topic including UNESCO Generative AI, UNESCO Generative Ai in Education⁵, and more⁶.

IEEE published Ethically Aligned Design⁷, 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 Enforcement⁸” or “AI for Safer Children”.

3.2 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”⁹ 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

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>

⁷ <https://ieeexplore.ieee.org/stamp/stamp.jsp?tp=&arnumber=9398613>

⁸ <https://unicri.it/topics/Toolkit-Responsible-AI-for-Law-Enforcement-INTERPOL-UNICRI>

⁹ <https://legalinstruments.oecd.org/en/instruments/OECD-LEGAL-0449>

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 OECD 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 update¹⁰. 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.

The typical life cycle is composed by: Design, Development, Deployment¹¹.

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

Understand the problem: To share your team’s understanding of their mission challenge, you first have to 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.)

¹⁰ Care of Juraj Čorba, Audrey Plonk, Karine Perset, Yoichi Iida

¹¹ 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/>

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."

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.

3.3 EUROPEAN COMMISSION & COUNCIL OF EUROPE

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 treaty¹² 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 CoE members states, 11 non-member states¹³, 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."¹⁴

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.

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

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 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].

¹²[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]})

¹³ Argentina, Australia, Canada, Costa Rica, the Holy See, Israel, Japan, Mexico, Peru, the United States of America, and Uruguay

¹⁴ Artificial Intelligence for Europe

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. 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”¹⁵ involved in the application of AI technology e.g. Federal law enforcement agency¹⁶, Sector Risk Management Agency¹⁷.

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,

¹⁵ 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).

¹⁶ The term “Federal law enforcement agency” has the meaning set forth in section 21(a) of Executive Order 14074 of May 25, 2022

¹⁷ The term “Sector Risk Management Agency” has the meaning set forth in 6 U.S.C. 650(23).

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

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.¹⁸

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 Bill¹⁹ 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 1950²⁰ 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.

3.5 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. 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

¹⁸ 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/>

¹⁹ Sponsor Sen. Moran, Jerry [R-KS] (Introduced 11/02/2023)

²⁰ https://www.fema.gov/sites/default/files/2020-03/Defense_Production_Act_2018.pdf

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, workshoping 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".

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. Please feedback your opinions to the drafting department of the organisation before November 13, 2024.

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.

3.6 AUSTRALIA

The Australian Government's consultations on safe and responsible AI have shown that current regulatory system is not fit for purpose to respond to the distinct risks that AI 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.

7 - AI law: emerging trends and challenges

- 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.

3.7 CANADA

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 with regard to 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).

3.8 FINLAND

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).

4. 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. 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. Actually, 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?

5. References

- [1.]Klaus Schwab. “Shaping the Fourth Industrial Revolution”, World Economic Forum 2016
- [2.]V International Conference “Tangible and Intangible Impact of Information and Communication in the Digital Age” to be held in Khanty- Mansiysk, Russian Federation, on 6-8 June 2023 within the XIV International IT Forum with BRICS and SCO participation and the UNESCO Information for All Intergovernmental Programme (IFAP)
- [3.]Imad A. Moosa, 2023, The Western Economic System, In book: The West Versus the Rest and The Myth of Western Exceptionalism, Palgrave Macmillan
- [4.]Deglobalisation: what you need to know, World Economic Forum 2023
<https://www.weforum.org/agenda/2023/01/deglobalisation-what-you-need-to-know-wef23/> and Deglobalisation? The reorganisation of global value chains in a changing world
<https://www.oecd.org/publications/deglobalisation-the-reorganisation-of-global-value-chains-in-a-changing-world-b15b74fe-en.htm>, OCSE
- [5.]Risk or reality? European Central Bank – 2023
<https://www.ecb.europa.eu/press/blog/date/2023/html/ecb.blog230712~085871737a.en.ht>
- [6.]Eckersley Richard, 1995, Values and visions: Youth and the failure of modern Western culture, Youth Studies Australia 14(1):13-21
- [7.]Vogels Emily A. et Al. (2023) Americans and ‘Cancel Culture’: Where Some See Calls for Accountability, Others See Censorship, Punishment, Pew Research Centre -
<https://www.pewresearch.org/internet/2021/05/19/americans-and-cancel-culture-where-some-see-calls-for-accountability-others-see-censorship-punishment/>