

Concept Paper

Generative Artificial Intelligence and Regulations: Can We Plan a Resilient Journey Toward the Safe Application of Generative Artificial Intelligence?

Matteo Bodini 

Department of Economics, Management and Quantitative Methods, University of Milan,
Via Conservatorio 7, 20122 Milan, Italy; matteo.bodini@unimi.it

Abstract: The rapid advancements of Generative Artificial Intelligence (GenAI) technologies, such as the well-known OpenAI ChatGPT and Microsoft Copilot, have sparked significant societal, economic, and regulatory challenges. Indeed, while the latter technologies promise unprecedented productivity gains, they also raise several concerns, such as job loss and displacement, deepfakes, and intellectual property violations. The present article aims to explore the present regulatory landscape of GenAI across the major global players, highlighting the divergent approaches adopted by the United States, United Kingdom, China, and the European Union. By drawing parallels with other complex global issues such as climate change and nuclear proliferation, this paper argues that the available traditional regulatory frameworks may be insufficient to address the unique challenges posed by GenAI. As a result, this article introduces a resilience-focused regulatory approach that emphasizes aspects such as adaptability, swift incident response, and recovery mechanisms to mitigate potential harm. By analyzing the existing regulations and suggesting potential future directions, the present article aims to contribute to the ongoing discourse on how to effectively govern GenAI technologies in a rapidly evolving regulatory landscape.

Keywords: GenAI; GenAI regulations; GenAI regulatory framework; resilience



Citation: Bodini, M. Generative Artificial Intelligence and Regulations: Can We Plan a Resilient Journey Toward the Safe Application of Generative Artificial Intelligence? *Societies* **2024**, *14*, 268. <https://doi.org/10.3390/soc14120268>

Academic Editor: Theodora Saridou,
Charalampos Dimoulas

Received: 3 November 2024

Revised: 9 December 2024

Accepted: 13 December 2024

Published: 18 December 2024



Copyright: © 2024 by the author. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

1. Introduction

The latest advancements of Generative Artificial Intelligence (GenAI) technologies, such as ChatGPT, Microsoft Copilot, Google Gemini, and many others have led to the creation of conversations capable of miming human interactions, generating realistic images, engineering of advanced computer code, and execution of several other human-related tasks [1]. However, the latter progress has raised several concerns, such as potential issues related to future job loss and displacement [2], the rise of deepfakes [3], and violations of intellectual property rights [4]. On the other hand, GenAI technologies could significantly boost productivity, potentially driving growth for businesses and economies [5]. For instance, Bank of America predicted that by the year 2030, GenAI could bring USD 15 trillion into the global economy [6]. Given the above-mentioned implications, it is no surprise that regulatory agencies are trying to take actions against GenAI. Indeed, as per the Stanford Artificial Intelligence (AI) Index Report, over 140 laws related to AI were enacted by the year 2023 across world states (refer to Figure 1) [7]. Despite the introduced laws and regulations, a major question still remains, i.e., what form of regulations are the most effective in addressing the recent challenges posed by GenAI technologies, across multiple sectors? The present concept paper adopts a cross-sectoral perspective, recognizing that regulatory challenges and demands may vary significantly across multiple domains, such as education, healthcare, politics, and the military. Indeed, the resilience-based research design provided herein is aimed to overcome domain-related discrepancies, thereby delivering a flexible yet robust approach applicable across multiple sectors.

Cumulative AI-related bills passed into law since 2016, as of 2023

Bills passed into law by national legislative bodies (e.g., congress, parliament) with the keyword "artificial intelligence" (translated to the respective languages) in the title or body of the bill.

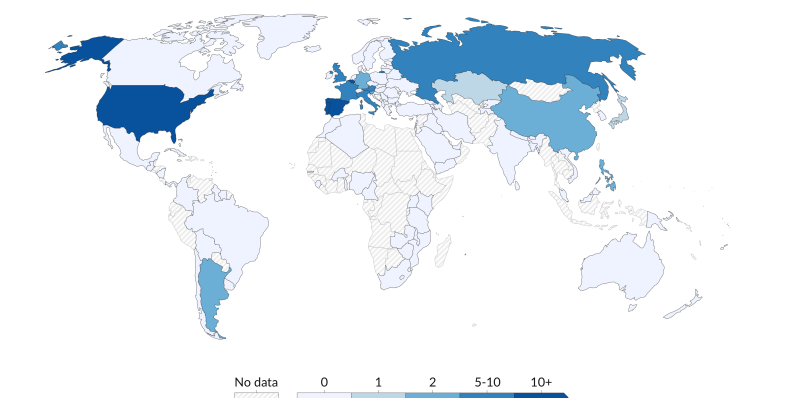


Figure 1. The figure reports the number of bills passed into law by national legislative bodies (e.g., congress and parliament) with the keyword “artificial intelligence” (translated to the respective languages) in the title or body of the bill. The data reported in the figure were collected from the 2023 Stanford AI Index Report [7] (p. 378). The top five world countries for passed AI-related bills (whose number of bills is here, respectively, reported for each state within parenthesis) are the United States (US) (23), Portugal (15), Belgium (12), Spain (11), and South Korea (10). The figure was adapted from Roser [8], with minor processing, from the 2023 Stanford AI Index Report [7] (p. 376), under the terms of the Creative Commons Attribution License—CC BY 4.0 (<https://creativecommons.org/licenses/by/4.0/>, accessed on 6 December 2024).

Regarding the above-mentioned term of “resilience”, within the context of GenAI technologies, such a term refers to the capability of systems to withstand, adapt to, and recover from disruptions while maintaining their functionality. Drawing from the “Ethics Guidelines for Trustworthy AI” provided by the High-Level Expert Group on AI, resilience is a crucial component of technically robust and safe systems, encompassing features such as fault tolerance, security against adversarial attacks, and recovery mechanisms [9]. The latter concept also aligns with the insights of the European Union Agency for Cybersecurity (ENISA) into AI cybersecurity, emphasizing the importance of preemptive risk identification, incident mitigation, and post-incident recovery [10]. Therefore, a resilient GenAI system must not only integrate preventive safeguards but also ensure dynamic adaptability to evolving challenges, safeguarding its reliability across diverse applications.

The complexity of regulating GenAI lies in its multiple available applications and the range of risks it introduces. Current regulatory efforts have primarily focused on issues such as data privacy, security, and transparency. For instance, the European Union’s (EU) proposed AI Act sets forth a risk-based framework that classifies AI systems according to the potential harm they could cause, ranging from minimal to unacceptable risk [11,12]. High-risk systems, such as those used in critical infrastructure or law enforcement, would face stricter regulatory scrutiny, including requirements for transparency, oversight, and accountability. However, are the already available regulations enough to address the challenges specific to the latest advancements brought by GenAI? It must be acknowledged that, while the present work predominantly examines novel regulatory initiatives, not all the challenges introduced by GenAI are entirely new. Indeed, depending on the considered domain, pre-existing regulatory frameworks, such as those related to the International Humanitarian Law (IHL), have shown remarkable adaptability in addressing emerging technological challenges. For instance, the IHL framework, as discussed in Davison [13], demonstrated how established principles, such as distinction and proportionality, have been applied successfully to regulate autonomous weapon systems which, similarly to GenAI, are characterized by rapid technological advancements. Drawing from such an example, certain aspects of GenAI regulation could benefit from leveraging existing laws

and principles, thereby reducing the need for entirely new frameworks. However, such a reflection even underscores the dual necessity of both evaluating the sufficiency of existing frameworks while considering the unique attributes of GenAI that demand innovative regulatory responses. Thus, while the present work acknowledges the remarkable insights offered by established frameworks, it points out that the latter ones alone may not be sufficient to address all the GenAI-related challenges. As a final consequence, the latter point underscores the necessity for the development of novel and comprehensive frameworks specifically designed to meet all the unique challenges posed by Generative AI.

In the past years, several researchers attempted to categorize the designed regulatory AI and GenAI frameworks within different world countries relying on the degree of legally binding and non-binding regulations, and whether such regulations were sector-specific or could be applied to the entire considered state economy [14]. For instance, the Executive Order on AI issued by President J. Biden in the US in October 2023 represents a federal initiative to establish safety and security standards, potentially guiding the US toward more rigorous restrictions [12,15]: 14 US states as of 2024 approved some sort of AI regulation mainly focused on consumer and/or data privacy, e.g., the Illinois AI Video Interview Act [16]. The government of the United Kingdom defined several principles (safety, security and robustness, appropriate transparency and explainability, fairness, accountability and governance, and contestability and redress) and a central function to assess regulation and monitor AI and GenAI technologies [17,18]. The People's Republic of China introduced regulations on recommendation algorithms [19] (in the year 2022), deep synthesis [20] (in the year 2023), and GenAI [21] (in the year 2023) [22]. The previously mentioned EU AI Act (released in the year 2024) states requirements based on risks induced by AI and GenAI systems. In particular, applications denoted with high risks require an assessment before release [11,12]. When considering the balance of interests, it can be observed that the US leans toward fostering innovation while the EU prioritizes citizen welfare and China emphasizes full control over the state [23]. As a final point, it must be finally noted that the majority of countries of all the world continents (except Africa) developed, or are developing, national AI strategies (refer to Figure 2). For the sake of readability, Table 1 summarizes the key aspects of the above-described regulations.

Countries with national artificial intelligence strategies, 2023

An AI strategy is a policy document that communicates the objective of supporting the development of AI while also maximizing the benefits of AI for society.

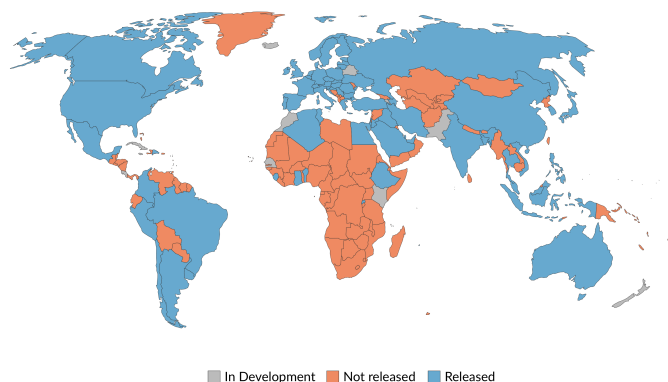


Figure 2. The figure reports world countries with national AI strategies (in light blue color), countries where AI strategies are under development (in grey color), and countries where AI strategies are not released or not already planned (in orange color). The data reported in the figure were collected from the 2023 Stanford AI Index Report [7] (p. 392). The figure was adapted from Roser [8], with minor processing, from the 2023 Stanford AI Index Report [7] (p. 391), under the terms of the Creative Commons Attribution License—CC BY 4.0 (<https://creativecommons.org/licenses/by/4.0/>), accessed on 6 December 2024).

Table 1. Overview of AI and GenAI regulatory frameworks for the analyzed countries or world regions, categorized by the degree of stringency and focus. The table outlines key regulations and taken actions and highlights each country's or region's main regulatory focus.

Country/Region	Regulatory Approach	Key Regulations	Main Focus
US	Legally non-binding, evolving	Executive Order on AI (2023) [15], state-level regulations (e.g., Illinois AI Video Interview Act) [16]	Innovation, safety, and security
United Kingdom	Legally non-binding, with oversight	Principles for AI regulation: safety, transparency, fairness, accountability, and monitoring [17]	Balancing innovation with safety
People's Republic of China	Legally binding, sector-specific	Regulations on algorithms (2022) [19], deep synthesis [20] (2023), and GenAI [21] (2023)	State control and oversight
EU	Legally binding, risk-based	EU AI Act (2024), requiring risk assessments for high-risk AI applications [11]	Citizen welfare and protection

2. Research Methodology

This article relied on a qualitative research methodology to explore the regulatory landscape of GenAI and, within the forthcoming sections, it introduces a resilience-focused regulatory framework. In particular, the designed research process involved the following key steps:

- **Literature review:** A comprehensive review of the existing literature on GenAI technologies, their societal and economic impacts, and current regulatory approaches was conducted. The latter one included academic papers, industry reports, policy documents, and regulatory guidelines from major global players such as the US, United Kingdom, China, and the EU.
- **Comparative analysis:** This study compared the regulatory frameworks of different countries and regions to identify commonalities and differences in their approaches to GenAI governance. Such a comparative analysis helped in understanding the strengths and weaknesses of various regulatory models and informed the development of the proposed resilience-focused framework.
- **Thematic analysis:** Key themes related to GenAI regulation, such as data privacy, security, transparency, and resilience, were identified and analyzed. The thematic analysis provided insights into the critical aspects that need to be addressed in a robust regulatory framework.
- **Case studies:** Specific case studies of regulatory responses to other complex global challenges, such as climate change, nuclear proliferation, and space debris, were examined. The latter case studies provided valuable lessons and analogies that were applied to the context of GenAI regulation.
- **Framework development:** Based on the insights gained from the literature review, comparative analysis, thematic analysis, and case studies, a resilience-focused regulatory framework for GenAI was developed. Such a framework emphasizes adaptability, swift incident response, and recovery mechanisms to mitigate potential harm associated with GenAI technologies.

Regarding the above-described research methodology, it must be noted that the present study focused on the regulatory strategies of the EU, China, the US, and the UK to illustrate divergent global approaches to AI governance. In particular, such regions were selected

based on their distinct regulatory priorities and methodologies, including risk-based frameworks (EU), innovation-driven initiatives (the US and UK), and state-controlled policies (China). Additional countries, such as Russia, Argentina, Japan, and India, are acknowledged for their active AI legislative landscapes; however, they were not included in the current analysis due to the scope of this study. Indeed, the focus was on selecting world regions with the highest legislative activity, as identified in the Stanford AI Index 2023, while also aiming to highlight different approaches to emphasize the current contrasts in regulatory frameworks and underlying policy objectives.

Moreover, while the present concept paper is focused on the regulation of GenAI, it is important to notice that GenAI is a subset of the broader category of AI [24]. Indeed, AI encompasses various forms, including narrow AI, general AI, and super AI. However, given the recent emergence of GenAI, most global entities have not yet developed specific regulations exclusively for GenAI [25]. Therefore, the regulatory frameworks discussed in the current article primarily pertain to AI in general. Such broader focus is necessary to provide a comprehensive understanding of the current regulatory landscape, as GenAI is inherently included within the scope of AI regulations. The references to AI regulations highlight foundational principles and approaches that can be adapted and applied to the regulation of GenAI. Indeed, as the field of GenAI continues to evolve, it is anticipated that more specific regulatory measures will be developed to address its unique challenges and implications. Last, but not least, the present concept paper adopts a cross-sectoral perspective, recognizing that regulatory challenges and demands may vary significantly across domains such as education, healthcare, politics, and the military. Thus, the resilience-based research design outlined in the present article aims to address domain-specific discrepancies, providing both a robust and adaptable framework that can be applied across multiple sectors.

3. Comparative Analysis of GenAI and Other Global Challenges

Although the categorization presented in the Introduction Section helps toward the comprehension of regulatory AI and GenAI laws and frameworks [14], it does not pinpoint which regulations are the most capable of addressing the specific attributes of GenAI. A more effective approach might be to compare GenAI technologies with other complex issues tackled by world governments for several years such as climate change, nuclear weapons proliferation, and space debris. Such a comparison could rely on the approach suggested by Roberts [26], which suggests measuring the level of agreement in addressing the considered problem and the capability to effectively manage it. Indeed, measuring the people's consensus and the effectiveness of resolution could help in determining the gravity of the challenge and the necessary regulations. Regarding the issue of climate change, multiple global regulations, e.g., the Paris Agreement of 2016 have been signed, despite the several competing economic interests [27]. As a result, unfortunately, it is difficult to promote the reduction in emissions. Global politics is mostly aligned regarding the proliferation of nuclear energy as only nine countries currently have nuclear weapons and a few governments can actually produce nuclear weapons through the availability of enriched uranium and related technology [28]. Regarding the issue of space debris, currently, a few states and space agencies can send rockets into space. However, there is limited consensus on preventing and/or removing the existing debris, and the several agencies proceed on their own [29]. Finally, regarding GenAI, several competing interests are nowadays faced due to the wide economic benefits for companies. Furthermore, it is difficult to limit the usage of GenAI systems since they are mostly freely available, open source, and the employed technology to implement them is changing rapidly [5,7,30]. As a final result, regarding the low consensus toward implementing GenAI regulations and the current limited control over its applicability, GenAI technologies nowadays present a particularly complex problem, even potentially worse with respect to the pressing issues of climate change, the nuclear proliferation of weapons, and the limitation of space debris

release within space. A summary comparison of GenAI technologies with the analyzed global challenges is reported in Table 2.

It must be noted that the methodology provided by Roberts [26] was selected in the context of the present study due to its comprehensive framework, capable of evaluating complex global challenges. Such a methodology considers both the level of consensus in addressing a specific problem and the capability to effectively manage it, thereby providing a thorough approach to regulatory analysis. The latter dual focus is particularly related to the regulatory landscape of GenAI, which involves multiple stakeholders with different interests and capabilities. Indeed, by applying the Roberts methodology, in the previous paragraphs, it was possible to effectively compare the novel challenge of GenAI with other global challenges such as climate change, nuclear proliferation, and space debris. The latter comparison both remarks the unique regulatory needs of GenAI technologies and underscores similarities and differences with the other considered global challenges.

Table 2. Comparison of GenAI technologies with the analyzed global challenges, focusing on consensus and management capabilities and relying on the approach proposed by Roberts [26]. The last column of the reported table highlights the distinct regulatory priorities of the analyzed entities, i.e., the US, UK, China, and the EU.

Global Challenge	Level of Agreement	Management Capability	Key Issues	Country/Region Highlights
Climate change	High	Limited	Competing economic interests, difficulty in reducing emissions despite global agreements [27]	EU: strong environmental focus; US: varied state-level policies
Nuclear proliferation	High	Strong	A few countries have the capability to produce nuclear weapons [28]	US/UK: focus on global deterrence; China: strategic military considerations
Space debris	Low	Strong	Limited consensus on debris prevention, independent actions by agencies [29]	EU: collaborative efforts (e.g., ESA); China: rapid expansion of space programs
GenAI technologies	Low	Limited	Rapidly changing technology, low consensus on regulations [5,7,30]	US: innovation-driven; UK: balanced innovation/safety; China: state control; EU: citizen welfare

In a similar way to the above-presented issues, GenAI evolves swiftly; exerts a wide-ranging influence on the economy, society, and governance; demands substantial expertise for comprehensive understanding; and presents an unpredictable balance of risks and rewards [5,31]. However, unlike the other presented issues, there is still a lack of consensus: the multitude of conflicting interests among governments and corporations obstructs the identification of a shared problem to address, let alone a collective solution [31]. The latter point is not unexpected, as the potential benefits of AI are so substantial that concerns about potential risks have not significantly slowed its development [5]. Furthermore, the pace of AI and GenAI evolution is so rapid that it challenges regulatory bodies to stay abreast. Indeed, even though there are still numerous tasks where human intuition, creativity, and emotional intelligence remain unparalleled by AI and GenAI, the advancements in these technologies over the past decade have been nothing short of remarkable. In just ten years, AI and GenAI methods have evolved from being relatively underwhelming to outperform-

ing humans in several domains traditionally dominated by human expertise. Such domains include, but are not limited to, reading comprehension [32], image recognition [33,34], and language understanding [35]. The studies reported by Roser [8] provide a comprehensive overview of these advancements. For instance, AI systems now excel in reading comprehension tasks, often achieving higher accuracy rates than human counterparts. Similarly, in the field of image recognition, AI algorithms have reached a level of precision that surpasses human capabilities, enabling applications ranging from medical diagnostics to autonomous driving. Language understanding, another critical area, has seen AI models like OpenAI ChatGPT and its successors demonstrate an impressive ability to generate and comprehend human-like text. For a more detailed comparison between human and AI capabilities, refer to Figure 3. The latter figure illustrates the relative performance of AI systems compared to humans across various tasks, highlighting the areas where AI has not only caught up but also exceeded human performance.

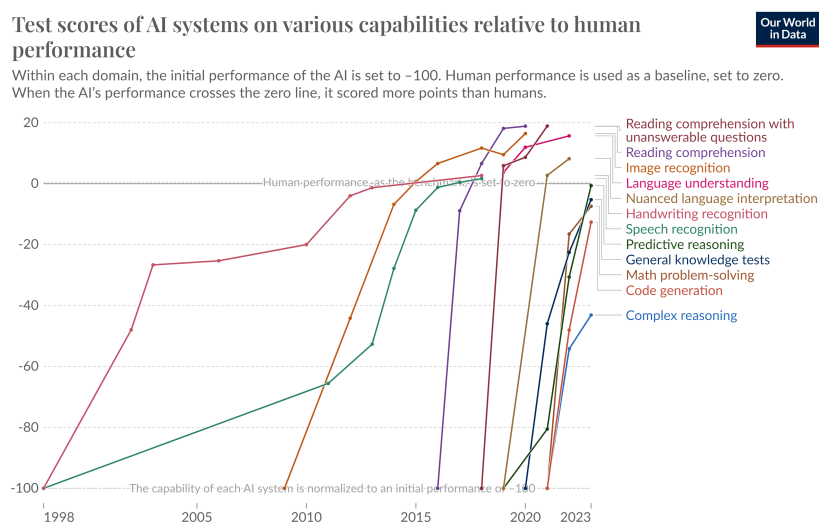


Figure 3. The figure reports the scores of several AI and GenAI systems on several capabilities related to human performance. Several capabilities were tested on the models (refer to the legend on the right of the figure), and the human performance was used as a baseline by setting it to the value of zero. As a result, AI and GenAI models outperformed humans in several of the analyzed tasks, for instance, in reading comprehension (18.85), image recognition (16.45), and language understanding (15.67). On the other hand, they also underperformed compared to humans in many tasks, such as complex reasoning (−43.12), code generation (−12.64), and math problem solving (−7.44). The data displayed in the figure were collected from Kiela et al. [36], and figure was adapted from Roser [8] under the terms of the Creative Commons Attribution License—CC BY 4.0 (<https://creativecommons.org/licenses/by/4.0/>), accessed on 6 December 2024).

AI and GenAI technologies present multiple opportunities for proper regulatory interventions that must be taken immediately in the next forthcoming years. Indeed, the creation of state-of-the-art GenAI models necessitates specialized knowledge, high-end hardware such as Graphics Processing Units (GPUs), substantial computational power, and financial resources that could exceed USD 100 million, for instance, in the case of the development of top-tier GenAI models, such as the OpenAI GPT-4 [37]. However, such a landscape is evolving and the high expenses associated with the development and operation of even the most advanced AI and GenAI models are becoming less, thus enfostering the development of novel models [8,30]. Furthermore, AI models and applications are increasingly being made available as open source [38]. For instance, meta Llama 3 was made available to developers in open-source format, but even other models, such as the Google Gemma model, are also available [39]. Moreover, AI repositories are now leveraged by AI developers to develop and deploy AI and GenAI models. For instance, the well-known AI model repository Hugging Face (<https://huggingface.co>) currently houses

nearly 800,000 AI models available for download and customization in its transformer library as of August 2024. The most currently downloaded model, a speech recognition model, has been downloaded over 292 million times by users within the community [40].

4. How to Safely Apply GenAI Technologies with Resilience?

Differently from the other mentioned problems reported in Section 1, issues with limited consensus and low control, such as the ones related to GenAI, necessitate prompt regulations centered on the concept of resilience. Regarding the issue of climate change, soft laws were applied in the past years where several voluntary agreements were made within and between many world states, for instance, regarding standards, codes of conduct, and certifications [27]. Regarding the proliferation of nuclear weapons, historically preventive agreements and regulations were made to restrict the usage of technologies and actions [28]. Finally, toward the issue of space debris, the civil or criminal liability of space companies was regulated [29]. However, regarding GenAI, the presence of multiple conflicting interests between world states and companies and the swift pace of technological advancements makes it challenging to develop preventive measures and soft regulations. Furthermore, regulations merely focused on the concept of liability may demand an excessive degree of control for establishment and enforcement. In contrast, regulations focused on the concept of resilience may represent a valid alternative, since such regulations could adopt a permissive approach, thus allowing GenAI technologies to progress and proliferate while strengthening institutions to formulate principles, standards, procedures, and capabilities to preemptively address potential issues which may arise due to their usage. Such regulations may also adopt a retrospective approach, implying they react promptly to issues and facilitate recovery following any potential incident.

GenAI regulations focused on the concept of resilience necessitate the involvement of competent and prompt institutions and stakeholders, such as international bodies, national governments, corporations, civil society, and individuals. Such entities are required to demonstrate the following seven key resilient features:

- Strong institutions capable of preempting incidents: the incorporation of principles, standards, procedures, and capabilities that encourage the development, utilization, and testing of GenAI systems for proactive risk identification and mitigation forms the bedrock for averting AI-related incidents.
- Swift action to mitigate risks and halt incidents: Events such as intellectual property theft, the disruption of critical infrastructure, election manipulation through deepfakes, or other unforeseen occurrences may happen. In such cases, institutions should be equipped with the necessary authority, procedures, and capabilities to promptly detect and halt these incidents.
- International collaborations toward unified AI standards: International cooperation is crucial for developing and enforcing unified standards for GenAI. The latter involves creating a global framework that ensures consistent regulation, especially in cross-border applications of AI technologies, to prevent deficiencies in regulations that could be exploited by bad actors.
- Investments in public awareness and education: Institutions should focus on raising public awareness about GenAI technologies and their associated risks. By educating the public and fostering digital literacy, society can better understand GenAI's impact, which is essential for resilience, particularly in recognizing and mitigating risks associated with GenAI misuses.
- Fostering innovation with built-in safety mechanisms: Encourage the development of GenAI systems that inherently include safety features such as AI explainability [41,42], robustness against adversarial attacks [43], and the ability to deactivate or revert decisions [44]. This proactive approach ensures that GenAI advancements are aligned with resilience goals and minimize potential harm from the outset.
- Implementing continuous monitoring and auditing: Establish ongoing monitoring systems to track the performance and impact of GenAI technologies in real time.

Regular audits can identify emerging risks early, allowing for swift intervention and the adjustment of practices to prevent potential harm.

- Bouncing back from risks by reducing harm and providing compensation to victims: After an incident, it is crucial to ensure that victims, businesses, organizations, and individuals receive appropriate compensation. Such an approach will maintain trust in GenAI usage and prevent demands for halting GenAI utilization or implementing other severe measures.

The actualization of the above-described requirements need the mentioned institutions to implement certain regulatory measures. At the international level, efforts should focus on establishing binding GenAI principles and technical standards to create a minimum threshold for GenAI regulations that can be universally monitored by bodies such as the G7 and China, or a potential new International Atomic Energy Agency (IAEA) equivalent for GenAI [45]. In the event of an incident, GenAI risk and incident early warning systems should be in place to detect and alert the international community about potential threats. Post-incident, there should be a robust system for tracking and analyzing the aftermath of GenAI-related incidents to continuously improve global safety standards. Enhancing international collaboration for unified standards is crucial, as seen in the success of the Montreal Protocol in phasing out ozone-depleting substances. Investing in public awareness and education, similar to public health campaigns for disease prevention, can help society understand GenAI's impact and recognize risks. Fostering innovation with built-in safety mechanisms, akin to the international automotive industry's crash testing protocols, ensures GenAI advancements align with resilience goals. Continuous monitoring and auditing, as practiced in environmental regulations to prevent pollution, can identify emerging risks early [46]. Such measures, combined with swift action to mitigate risks and halt incidents, will strengthen institutions and ensure a resilient approach to GenAI.

It must be noted that existing multilateral cooperation efforts on AI-related issues provide a strong foundation for such above-presented initiatives. For instance, the European Commission's Expert Group on AI [47], the ENISA [48], and the United Nations' AI Advisory Body [49] are already working toward establishing robust AI governance frameworks. Additionally, the Organisation for Economic Co-operation and Development, i.e., OECD's Global Partnership on AI [50] and its Network of Experts [51], as well as initiatives by the United Nations Interregional Crime and Justice Research Institute (UNICRI) [52] and the AI for Good initiative by the International Telecommunication Union (ITU) [53], are fundamental in fostering international collaboration and setting global AI standards.

National governments have a critical role in structuring and enforcing technical standards within their jurisdictions. This includes licensing GenAI companies, auditing AI models, establishing post-release controls, conducting red team testing, and implementing Know Your Customer (KYC) procedures for APIs [54]. Governments should also offer incentives for companies to exceed these regulatory minimums and invest in developing more resilient GenAI systems. In response to incidents, there should be an emergency authority equipped to act swiftly to contain and mitigate the damage. Furthermore, national governments must establish clear liability and compensation standards to address the consequences of GenAI incidents. Regarding the available regulations, it must be noted that the EU's AI Act outlines specific obligations for member states, including the designation of national competent authorities for AI oversight [11]. Moreover, Spain was the first EU country to fulfill the latter obligations by establishing the Spanish Agency for the Supervision of Artificial Intelligence (AESIA) [55]. Such an agency is responsible for supervising high-risk AI systems and ensuring compliance with the AI Act.

Enhancing international collaboration for unified standards is crucial, as seen in the success of the Kyoto Protocol in reducing emissions. Investing in public awareness and education, similar to national financial literacy programs, can help society understand AI's impact and recognize risks. Fostering innovation with built-in safety mechanisms, akin to the pharmaceutical industry's clinical trials, ensures AI advancements align with resilience goals. Continuous monitoring, as practiced in food safety regulations, can identify

emerging risks early. By taking quick steps to address risks and prevent incidents, such measures will fortify institutions and promote a robust strategy for GenAI regulation.

Companies must take proactive steps to encourage transparency, share their expertise, and drive the development of technical standards for GenAI. They should provide data and models for testing, create watermarks for GenAI-generated content, and develop use cases for mitigating risks. Internally, companies need to develop controls, processes, and compliance mechanisms to ensure responsible GenAI practices. If incidents occur, companies should have safety breaks in place, such as requiring human oversight for critical decisions involving GenAI. To manage the aftermath, companies should contribute to an insurance fund that can compensate for damages resulting from AI incidents. Enhancing international collaboration for unified standards is crucial, as seen in the Basel Accords in banking regulation [56]. Investing in public awareness and education, similar to environmental conservation campaigns, can help society understand AI's impact and recognize risks [46]. Fostering innovation with built-in safety mechanisms, akin to the company-level nuclear industry's safety protocols, ensures GenAI advancements align with resilience goals. Continuous monitoring and auditing, as practiced in aviation safety, can identify emerging risks early. Such actions, together with prompt efforts to manage risks and stop incidents, will enhance institutions and foster a resilient framework for regulating GenAI.

Civil society has the responsibility to act as a watchdog, identifying best regulatory practices and advocating for society's interests in AI governance. This involves conducting research to identify system risks and developing professional standards for AI practitioners. When responding to incidents, civil society should work to identify victims, document the impact of incidents, and push for accountability. In the recovery phase, civil society should provide independent assessments of the damage caused by AI incidents to ensure fair and transparent evaluations. Enhancing international collaboration for unified standards is crucial, as seen in the success of the Stockholm Convention in eliminating persistent organic pollutants [57]. Investing in public awareness and education, similar to road safety campaigns, can help society understand AI's impact and recognize risks. Continuous monitoring and auditing, as practiced in occupational health and safety, can identify emerging risks early. Swift actions to mitigate risks and stop incidents, alongside the latter measures, will bolster institutions and ensure a resilient regulatory framework for GenAI.

Finally, the role of people in shaping the future of AI cannot be understated. Individuals should have access to AI models, data, and computational resources, and there should be incentives, such as "bug bounties," to encourage users to provide feedback and report potential issues. During the response phase, people can contribute by crowdsourcing warnings and relevant information about AI-related risks. In the aftermath of incidents, people should be involved in reporting damages and collecting compensation to ensure their concerns are addressed and remedied. Enhancing international collaboration for unified standards is crucial, as seen in the success of the Minamata Convention in reducing mercury pollution [58]. Investing in public awareness and education, similar to digital literacy programs, can help society understand GenAI's impact and recognize risks. Continuous monitoring and auditing, as practiced in consumer protection regulations, can identify emerging risks early. Combining the latter measures with timely interventions to reduce risks and prevent occurrences will reinforce institutions and create a solid approach to GenAI regulation.

To clearly present the conclusions of the present concept paper, Table 3 finally summarizes the key regulatory measures and actions proposed throughout this research. In particular, it includes the suggested recommendations and outlines the institutional roles necessary to ensure a resilient approach to the governance of GenAI. Such a table serves as a concise summary of the presented conclusions, integrating the multiple aspects of the resilience-focused regulatory framework discussed in the present article.

Table 3. Overview of GenAI regulatory measures across institutions, focusing on preemptive actions, risk mitigation, international collaboration, public education, innovation, monitoring, and compensation.

Institution/Entity	Preempt Incidents	Mitigate Risks	International Collaboration	Public Awareness	Innovation with Safety	Monitoring and Auditing	Compensation Post-Incident
International bodies	Establish binding GenAI principles and standards	Implement early warning systems	Create a unified global framework	Promote awareness akin to public health campaigns	Encourage GenAI safety similar to automotive crash testing protocols	Monitor global GenAI practices	Track and analyze incidents, improve standards
National governments	License GenAI companies, audit models	Equip emergency authorities for swift response	Collaborate on unified standards, e.g., Kyoto Protocol	Run public awareness campaigns	Incentivize resilient GenAI systems	Conduct ongoing audits	Establish liability and compensation systems
Companies and industry	Develop internal controls and compliance mechanisms	Implement safety breaks, human oversight	Share expertise, contribute to standard development	Engage in awareness and transparency initiatives	Build GenAI with inherent safety features	Monitor GenAI performance, adjust practices	Contribute to an insurance fund for damages
Society	Advocate for best regulatory practices	Identify victims, document impacts	Push for unified standards	Research and educate the public on GenAI risks	Promote responsible GenAI use among practitioners	Conduct independent assessments, identify risks	Advocate for fair compensation, document damage
Individuals	Provide feedback, report issues (e.g., bug bounties)	Crowdsource warnings, share information	Support global GenAI regulations	Participate in education and awareness programs	Encourage transparent and safe GenAI practices	Report emerging risks	Report damages, seek compensation

5. Future Directions Toward Resilient GenAI Regulations

AI regulations focused on the concept of resilience are already slightly beginning to surface. The OECD has established an AI incident observatory and categorizes AI systems based on risk. The above-mentioned EU's risk classification for AI systems, the US 2023 Executive Order on AI which encourages collaborations with AI firms to report red team findings, and China's algorithm registry along with corporate security self-assessments all incorporate elements of resilience regulations to varying degrees. Indeed, corporations such as OpenAI, Microsoft, and Google have defined technical standards and have even made them public, an example being Microsoft's Responsible AI Standard [59].

The rise in reported AI incidents, as shown in Figure 4, highlights the growing challenges in ensuring responsible AI deployment and the urgency for resilience measures. Notable examples of AI incidents are related to deepfake videos [60]. Indeed, the recent misuse of deepfake technology has led to significant legal actions in various world countries. For instance, in the US, the DEEPFAKES Accountability Act was introduced to address the malicious use of deepfakes [61]. Such legislation aims to hold individuals accountable for creating and distributing deepfake content with the intent to deceive or harm others. Moreover, within the UK, a recent notable deepfake case involved the CEO of a UK-based energy firm who was tricked into transferring EUR 200,000 to a fraudster using deepfake audio technology to mimic the voice of his superior. Despite the fraudulent nature being recognized, the company faced challenges in tracing the origins of the deepfake, highlighting the need for robust legal frameworks to address such incidents [62]. Although advised that legal action might be challenging due to existing legislation gaps, such a case underscored the urgent need for laws to protect individuals from non-consensual deepfake content [63]. Similarly, the EU Code of Practice on Disinformation calls for measures to deal with the spread of manipulative deepfake content, particularly in political contexts [64]. Indeed, a recent notable example of the misuse of deepfake technology occurred during the Russia–Ukraine conflict, where a deepfake video falsely depicting Ukrainian President Volodymyr Zelenskyy surrendering was circulated. Such a video was quickly debunked, but it highlighted the potential for deepfakes to be used in disinformation campaigns [65]. Last but not least, AI technology has been used to monitor inmate calls in US prisons, aiming to detect and prevent criminal activities. Such technology has been employed to uncover threats, drug smuggling operations, and other illicit activities within the prison system [66].

There are six further measures that could enhance resilience in the context of GenAI. Such additional measures were derived from the sources processed during the research, particularly through the comprehensive analysis of the existing literature, case studies, and thematic analysis conducted as part of this study, i.e., (1) establish mandatory principles and baseline technical standards for GenAI, which could be overseen by an organization akin to the IAEA. (2) Ensure greater accessibility and the provision of resources for civil society and individuals to utilize data training sets and computational power for model development and testing. (3) Formulate structured recovery regulations linked to resilience. This could be the most complex task as it might suggest accountability for incidents. However, in the absence of explicit liability and compensation norms, GenAI firms might hesitate to disclose information due to potential future liability concerns. (4) Implement the continuous scenario-based stress testing of GenAI systems to identify potential vulnerabilities before they can cause harm. (5) Develop an international framework for AI incident reporting to standardize the process across borders, ensuring the timely and transparent communication of risks. (6) Encourage cross-sector collaborations between GenAI developers, regulators, and industries that heavily rely on GenAI, fostering a shared responsibility in addressing potential risks and strengthening resilience.

Regarding the previously reported sixth point, the importance of international co-operation and the establishment of global standards for AI governance was even recently emphasized by the United Nations High-Level Advisory Body on AI [67] which focused on the following:

- Ethical AI development: establish ethical guidelines to ensure AI technologies are developed and used in ways that uphold human rights and promote the social good.
- Inclusive AI: promote inclusivity in AI development, ensuring equitable access to AI benefits across diverse regions and communities.
- Transparency and accountability: ensure transparency in AI systems and establish accountability mechanisms to prevent misuse and address potential negative consequences.
- Capacity building: focus on building AI expertise and infrastructure in developing countries, thus enabling their participation in and benefit from AI advancements.
- Alignment with Sustainable Development Goals: align AI governance with the United Nations Sustainable Development Goals, hence promoting AI applications that contribute to global sustainability objectives.

Annual reported artificial intelligence incidents and controversies, World

Notable incidents include a "deepfake" video of Ukrainian President Volodymyr Zelenskyy surrendering, and U.S. prisons using AI to monitor their inmates' calls.

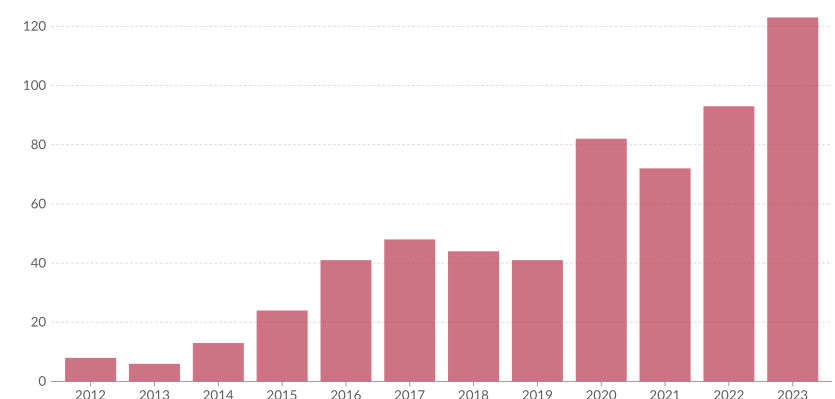


Figure 4. Annual reported AI incidents and controversies worldwide. The increase in incidents over the years underscores the growing challenges in managing AI risks. Notable examples include deepfake videos, such as one falsely depicting Ukrainian President Volodymyr Zelenskyy surrendering, and the use of AI to monitor inmate calls in US prisons. Data source: AI Incident Database via AI Index (2024) [7]. Figure was adapted from Roser [8] under the terms of the Creative Commons Attribution License—CC BY 4.0 (<https://creativecommons.org/licenses/by/4.0/>), accessed on 6 December 2024).

As a final remark, regulators could enhance consensus and control to shift away from resilience-focused regulations. The latter one is a logical approach: just as armies strive to alter power dynamics to achieve victory in warfare, corporations aim to secure monopoly power. However, modifying such parameters is challenging. It has been forty-four years since climate change was first discussed at the UN's inaugural Earth Summit in 1972, leading to the establishment of binding targets in the 2016 Paris Agreement through the efforts of advocates and scientists [27]. But the world cannot afford to wait that long: if GenAI triggers a banking crisis, disrupts critical infrastructure, or inflicts other forms of widespread societal damage, governments might be compelled to enforce legally binding authoritarian regulations, irrespective of their practicality or the potential benefits that might be forfeited. The time to act is now, and the forthcoming AI Safety Summits in France in 2025 could offer platforms for regulators to more decisively adopt a resilience-oriented approach, such as the one presented in the present article.

Funding: This research received no external funding.

Institutional Review Board Statement: Not applicable.

Informed Consent Statement: Not applicable.

Data Availability Statement: No new data were produced.

Conflicts of Interest: The author declares no conflicts of interest.

References

1. Gupta, P.; Ding, B.; Guan, C.; Ding, D. Generative AI: A systematic review using topic modelling techniques. *Data Inf. Manag.* **2024**, *8*, 100066. [CrossRef]
2. Idrisi, M.J.; Geteye, D.; Shanmugasundaram, P. Modeling the Complex Interplay: Dynamics of Job Displacement and Evolution of Artificial Intelligence in a Socio-Economic Landscape. *Int. J. Networked Distrib. Comput.* **2024**, *12*, 185–194. [CrossRef]
3. Gambín, Á.F.; Yazidi, A.; Vasilakos, A.; Haugerud, H.; Djenouri, Y. Deepfakes: Current and future trends. *Artif. Intell. Rev.* **2024**, *57*, 64. [CrossRef]
4. Chesterman, S. Good models borrow, great models steal: Intellectual property rights and generative AI. *Policy Soc.* **2024**, puae006. [CrossRef]
5. Al Naqbi, H.; Bahroun, Z.; Ahmed, V. Enhancing Work Productivity through Generative Artificial Intelligence: A Comprehensive Literature Review. *Sustainability* **2024**, *16*, 1166. [CrossRef]
6. Adigwe, C.S.; Olaniyi, O.O.; Olabanji, S.O.; Okunleye, O.J.; Mayeke, N.R.; Ajayi, S.A. Forecasting the Future: The Interplay of Artificial Intelligence, Innovation, and Competitiveness and its Effect on the Global Economy. *Asian J. Econ. Bus. Account.* **2024**, *24*, 126–146. [CrossRef]
7. Maslej, N.; Fattorini, L.; Perrault, R.; Parli, V.; Reuel, A.; Brynjolfsson, E.; Etchemendy, J.; Ligett, K.; Lyons, T.; Manyika, J.; et al. The AI Index 2024 Annual Report. AI Index Steering Committee, Institute for Human-Centered AI, Stanford University, Stanford, CA, April 2024. Available online: <https://www.dirittoue.info/?p=4682> (accessed on 6 December 2024).
8. Roser, M. The Brief History of Artificial Intelligence: The world has changed fast—What might be next? *Our World Data* **2022**. Available online: <https://ourworldindata.org/brief-history-of-ai> (accessed on 6 December 2024).
9. High-Level Expert Group on Artificial Intelligence. Ethics Guidelines for Trustworthy AI. 2019. Available online: <https://www.aepd.es/sites/default/files/2019-12/ai-definition.pdf> (accessed on 6 December 2024).
10. European Union Agency for Cybersecurity (ENISA). AI Cybersecurity Challenges—Threat Landscape for Artificial Intelligence. 2021. Available online: <https://protecciondata.es/wp-content/uploads/2021/06/Panorama-de-amenazas-de-ENISA-sobre-inteligencia-artificial-informe-de-2020.pdf> (accessed on 6 December 2024).
11. The European Commission. Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024. *Off. J. Eur. Union* **2024**. Available online: <http://data.europa.eu/eli/reg/2024/1689/oj> (accessed on 6 December 2024).
12. Wörsdörfer, M. Biden's Executive Order on AI and the E.U.'s AI Act: A Comparative Computer-Ethical Analysis. *Philos. Technol.* **2024**, *37*, 74. [CrossRef]
13. Davison, N. A Legal Perspective: Autonomous Weapon Systems under International Humanitarian Law. *UNODA Occas. Pap.* **2017**, *30*, 5–18. Available online: <https://www.icrc.org/en/document/autonomous-weapon-systems-under-international-humanitarian-law> (accessed on 27 November 2022).
14. Bryer, A.R. Understanding Regulation: Theory, Strategy, and Practice. *Account. Eur.* **2013**, *10*, 279–282. [CrossRef]
15. The White House. Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence. 2023. Available online: <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/> (accessed on 6 December 2024).
16. Illinois General Assembly. Artificial Intelligence Video Interview Act. 2019. Available online: <https://www.ilga.gov/legislation/publicacts/fulltext.asp?Name=101-0260> (accessed on 6 December 2024).
17. Government of the United Kingdom—Department for Science, Innovation & Technology. Implementing the UK's AI Regulatory Principles: Initial Guidance for Regulators. 2024. Available online: <https://www.gov.uk/government/publications/implementing-the-uks-ai-regulatory-principles-initial-guidance-for-regulators/102aa401-60f6-46e8-95dc-b48150eba7dd> (accessed on 6 December 2024).
18. Roberts, H.; Babuta, A.; Morley, J.; Thomas, C.; Taddeo, M.; Floridi, L. Artificial intelligence regulation in the United Kingdom: A path to good governance and global leadership? *Internet Policy Rev.* **2023**, *12*, 1–31. [CrossRef]
19. Cyberspace Administration of China. Regulations on the Administration of Internet Information Service Recommendation Algorithms. 2022. Available online: https://www.gov.cn/zhengce/zhengceku/2022-01/04/content_5666429.htm (accessed on 6 December 2024).
20. Cyberspace Administration of China. Provisions on the Administration of Deep Synthesis in Internet-Based Information Services. 2023. Available online: http://www.cac.gov.cn/2022-01/28/c_1644970458520968.htm (accessed on 6 December 2024).
21. Cyberspace Administration of China. Interim Measures for the Management of Generative Artificial Intelligence Services. 2023. Available online: https://www.cac.gov.cn/2023-07/13/c_1690898327029107.htm (accessed on 6 December 2024).
22. Franks, E.; Lee, B.; Xu, H. Report: China's New AI Regulations. *Glob. Priv. Law Rev.* **2024**, *5*, 43–49. [CrossRef]
23. Fung, P.; Etienne, H. Confucius, cyberpunk and Mr. Science: Comparing AI ethics principles between China and the EU. *AI Ethics* **2023**, *3*, 505–511. [CrossRef] [PubMed]
24. Damar, M.; Özen, A.; Çakmak, Ü.E.; Özoğuz, E.; Safa Erenay, F. Super AI, Generative AI, Narrow AI and Chatbots: An Assessment of Artificial Intelligence Technologies for The Public Sector and Public Administration. *J. AI* **2024**, *8*, 83–106. [CrossRef]
25. Huang, K.; Joshi, A.; Dun, S.; Hamilton, N. AI Regulations. In *Generative AI Security. Future of Business and Finance*; Huang, K., Wang, Y., Goertzel, B., Li, Y., Wright, S., Ponnappalli, J., Eds; Springer: Cham, Switzerland, 2024. [CrossRef]

26. Roberts, N. Wicked Problems and Network Approaches to Resolution. *Int. Public Manag. Rev.* **2000**, *1*, 1–19. Available online: <https://ipmr.net/index.php/ipmr/article/view/175> (accessed on 6 December 2024).
27. Bodansky, D. The legal character of the Paris agreement. *Rev. Eur. Comp. Int. Environ. Law* **2016**, *25*, 142–150. [CrossRef]
28. Ghoshal, D. Historical and Contemporary Missile Development: Nuclear Weapon States, Regional Powers and Other Powers. In *Role of Ballistic and Cruise Missiles in International Security*; Palgrave Macmillan: Cham, Switzerland, 2023; pp. 69–143. [CrossRef]
29. Klimburg-Witjes, N. A Rocket to Protect? Sociotechnical Imaginaries of Strategic Autonomy in Controversies About the European Rocket Program. *Geopolitics* **2023**, *29*, 821–848. [CrossRef]
30. Albayrak Ünal, Ö.; Erkeyman, B.; Usanmaz, B. Applications of Artificial Intelligence in Inventory Management: A Systematic Review of the Literature. *Arch. Comput. Methods Eng.* **2023**, *30*, 2605–2625. [CrossRef]
31. Saheb, T.; Saheb, T. Topical review of artificial intelligence national policies: A mixed method analysis. *Technol. Soc.* **2023**, *74*, 102316. [CrossRef]
32. Xiao, C.; Xu, S.X.; Zhang, K.; Wang, Y.; Xia, L. Evaluating Reading Comprehension Exercises Generated by LLMs: A Showcase of ChatGPT in Education Applications. In Proceedings of the 18th Workshop on Innovative Use of NLP for Building Educational Applications (BEA 2023), Toronto, Canada, 13 July 2023; pp. 610–625. [CrossRef]
33. Bodini, M. Will the Machine Like Your Image? Automatic Assessment of Beauty in Images with Machine Learning Techniques. *Inventions* **2019**, *4*, 34. [CrossRef]
34. Bodini, M. A Review of Facial Landmark Extraction in 2D Images and Videos Using Deep Learning. *Big Data Cogn. Comput.* **2019**, *3*, 14. [CrossRef]
35. Karanikolas, N.; Manga, E.; Samaridi, N.; Tousidou, E.; Vassilakopoulos, M. Large Language Models versus Natural Language Understanding and Generation. In Proceedings of the 27th Pan-Hellenic Conference on Progress in Computing and Informatics (PCI '23), Lamia, Greece, 24–26 November 2023; pp. 278–290. [CrossRef]
36. Kiela, D.; Thrush, T.; Ethayarajh, K.; Singh, A. Plotting Progress in AI. *Contextual AI Blog* **2023**. Available online: <https://contextual.ai/blog/plotting-progress> (accessed on 6 December 2024).
37. Telenti, A.; Auli, M.; Hie, B.L.; Maher, C.; Saria, S.; Ioannidis, J.P.A. Large language models for science and medicine. *Eur. J. Clin. Investig.* **2024**, *54*, e14183. [CrossRef] [PubMed]
38. Hagg, A.; Kirschner, K.N. Open-Source Machine Learning in Computational Chemistry. *J. Chem. Inf. Model.* **2023**, *63*, 4505–4532. [CrossRef] [PubMed]
39. Chang, Y.; Wang, X.; Wang, J.; Wu, Y.; Yang, L.; Zhu, K.; Chen, H.; Yi, X.; Wang, C.; Wang, Y.; et al. A Survey on Evaluation of Large Language Models. *ACM Trans. Intell. Syst. Technol.* **2024**, *15*, 39. [CrossRef]
40. Gong, Y.; Chung, Y.A.; Glass, J. Ast: Audio spectrogram transformer. In Proceedings of the Annual Conference of the International Speech Communication Association, INTERSPEECH, Brno, Czechia, 30 August–3 September 2021; pp. 571–575. [CrossRef]
41. Bodini, M.; Rivolta, M.W.; Sassi, R. Interpretability Analysis of Machine Learning Algorithms in the Detection of ST-Elevation Myocardial Infarction. *Comput. Cardiol.* **2020**, *47*, 1–4. [CrossRef]
42. Bodini, M.; Rivolta, M.W.; Sassi, R. Opening the Black Box: Interpretability of Machine Learning Algorithms in Electrocardiography. *Philos. Trans. R. Soc. A* **2021**, *379*, 20200253. [CrossRef] [PubMed]
43. Dong, H.; Dong, J.; Wan, S.; Yuan, S.; Guan, Z. Transferable Adversarial Distribution Learning: Query-Efficient Adversarial Attack Against Large Language Models. *Comput. Secur.* **2023**, *135*, 103482. [CrossRef]
44. Chiang, C.-W.; Lu, Z.; Li, Z.; Yin, M. Enhancing AI-Assisted Group Decision Making through LLM-Powered Devil's Advocate. In Proceedings of the 29th International Conference on Intelligent User Interfaces (IUI '24), Greenville, SC, USA, 18–21 March 2024; pp. 103–119. [CrossRef]
45. De Castro, B. L.; Gracia, F. J.; Peiró, J. M.; Pietrantoni, L.; Hernández, A. Testing the Validity of the International Atomic Energy Agency (IAEA) Safety Culture Model. *Accid. Anal. Prev.* **2013**, *60*, 231–244. [CrossRef] [PubMed]
46. Bodini, M. Charting the Future of Conservation in Arizona: Innovative Strategies for Preserving Its Natural Resources. *Conservation* **2024**, *4*, 402–434. [CrossRef]
47. European Commission. Expert Group on AI. Available online: <https://digital-strategy.ec.europa.eu/en/policies/expert-group-ai> (accessed on 6 December 2024).
48. European Union Agency for Cybersecurity (ENISA). Available online: <https://www.enisa.europa.eu/> (accessed on 6 December 2024).
49. United Nations. AI Advisory Body. Available online: <https://www.un.org/techenvoy/ai-advisory-body> (accessed on 6 December 2024).
50. OECD. Global Partnership on AI. Available online: <https://www.oecd.org/en/about/programmes/global-partnership-on-artificial-intelligence.html> (accessed on 6 December 2024).
51. OECD. Network of Experts. Available online: <https://oecd.ai/en/network-of-experts> (accessed on 6 December 2024).
52. United Nations Interregional Crime and Justice Research Institute (UNICRI). Accessed on: https://unicri.it/topics/ai_robotics/ (accessed on 6 December 2024).
53. International Telecommunication Union (ITU). AI for Good. Available online: <https://aiforgood.itu.int/> (accessed on 6 December 2024).
54. Kapsoulis, N.; Psychas, A.; Palaiokrassas, G.; Marinakis, A.; Litke, A.; Varvarigou, T. Know Your Customer (KYC) Implementation with Smart Contracts on a Privacy-Oriented Decentralized Architecture. *Future Internet* **2020**, *12*, 41. [CrossRef]

55. Spanish Agency for the Supervision of Artificial Intelligence (AESIA). Available online: <https://www.lamoncloa.gob.es/lang/en/gobierno/news/Paginas/2024/20240619-ai-oversight-agency.aspx> (accessed on 6 December 2024).
56. ElBannan, M. A. The Financial Crisis, Basel Accords and Bank Regulations: An Overview. *Int. J. Account. Financ. Report.* **2017**, *7*, 225–275. [CrossRef]
57. Hagen, P. E.; Walls, M. P. The Stockholm Convention on Persistent Organic Pollutants. *Nat. Resour. Environ.* **2005**, *19*, 49–52. Available online: <https://www.jstor.org/stable/40924611> (accessed on 6 December 2024).
58. Evers, D.C.; Keane, S.E.; Basu, N.; Buck, D. Evaluating the Effectiveness of the Minamata Convention on Mercury: Principles and Recommendations for Next Steps. *Sci. Total Environ.* **2016**, *569*, 888–903. [CrossRef] [PubMed]
59. Camilleri, M.A. Artificial intelligence governance: Ethical considerations and implications for social responsibility. *Expert Syst.* **2024**, *41*, e13406. [CrossRef]
60. De Rancourt-Raymond, A.; Smaili, N. The unethical use of deepfakes. *J. Financ. Crime* **2023**, *30*, 1066–1077. [CrossRef]
61. DEEPFAKES Accountability Act (H.R. 3230). 116th Congress (2019–2020). Available online: <https://www.congress.gov/bill/116th-congress/house-bill/3230> (accessed on 6 December 2024).
62. Quirk, C. The High Stakes of Deepfakes: The Growing Necessity of Federal Legislation to Regulate This Rapidly Evolving Technology. *Princet. Leg. J.* **2023**. Available online: <https://legaljournal.princeton.edu/the-high-stakes-of-deepfakes-the-growing-necessity-of-federal-legislation-to-regulate-this-rapidly-evolving-technology/> (accessed on 6 December 2024).
63. HyperVerge. Are Deepfakes Illegal? 2023. Available online: <https://hyperverge.co/blog/are-deepfakes-illegal/> (accessed on 6 December 2024).
64. Borz, G.; De Francesco, F.; Montgomerie, T.L.; Bellis, M.P. The EU soft regulation of digital campaigning: Regulatory effectiveness through platform compliance to the code of practice on disinformation. *Policy Stud.* **2024**, *45*, 709–729. [CrossRef]
65. Sky News. Ukraine War: Deepfake Video of Zelenskyy Telling Ukrainians to ‘Lay Down Arms’ Debunked. 2022. Available online: <https://news.sky.com/story/ukraine-war-deepfake-video-of-zelenskyy-telling-ukrainians-to-lay-down-arms-debunked-12567789> (accessed on 6 December 2024).
66. ABC News. US Prisons and Jails Using AI to Mass-Monitor Millions of Inmate Calls. 2019. Available online: <https://abcnews.go.com/Technology/us-prisons-jails-ai-mass-monitor-millions-inmate/story?id=66370244> (accessed on 6 December 2024).
67. United Nations. Governing AI for Humanity: Final Report of the UN High-Level Advisory Body on Artificial Intelligence. 2024. Available online: https://www.un.org/sites/un2.un.org/files/governing_ai_for_humanity_final_report_en.pdf (accessed on 6 December 2024).

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.