

Navigating the Risks of Generative AI: A Comparative Analysis of International Regulatory Approaches

Szacując ryzyka związane z generatywną Sztuczną Inteligencją: analiza porównawcza międzynarodowych podejść regulacyjnych

Huang Xinbo¹, Liu Guo^{2*}

¹*Nanchang Institute of Technology, School of Humanities and Law, Nanchang 330044, China
E-mail: huang_xinbo@gsgsg.uum.edu.my, ORCID:0000-0002-0905-1803*

²*China University of Political Science and Law, School of Law, Beijing 102249, China.
E-mail (Corresponding Author): CU212028@cupl.edu.cn

Abstract

The rapid evolution of artificial intelligence (AI) offers unprecedented opportunities for technological innovation and societal advancement, while also posing significant risks and challenges. Generative Artificial Intelligence (Generative AI) – exemplified by models such as ChatGPT – relies heavily on deep learning architectures, resulting in increased technical complexity and reduced interpretability. As a result, Generative AI exacerbates a variety of risks common to traditional AI systems, including data security vulnerabilities, privacy concerns, copyright infringements, and algorithmic biases. This study critically assesses the adequacy of current regulatory frameworks in addressing these risks. The findings suggest that traditional legal instruments are insufficient for managing the unique challenges posed by Generative AI, increasing the likelihood of regulatory failure. By conducting a comparative analysis of regulatory approaches in the European Union, the United States, the United Kingdom, and China, the research reveals a growing global consensus on the necessity of developing responsible and trustworthy AI that aligns with Sustainable development goals. The paper advocates for the establishment of a sustainable and adaptive regulatory framework that aims to enhance governmental capacity in addressing the evolving risks associated with Generative AI, while simultaneously promoting its safe, accountable, and sustainable development.

Key words: Artificial Intelligence (AI), generative AI, regulatory approaches, sustainable development

Streszczenie

Szybka ewolucja sztucznej inteligencji (SI) oferuje bezprecedensowe możliwości innowacji technologicznych i postępu społecznego, jednocześnie stanowiąc znaczne ryzyko i wyzwanie. Generatywna sztuczna inteligencja (Generative AI) – z modelami takimi jak ChatGPT – w dużym stopniu opiera się na architekturach głębokiego uczenia, co skutkuje zwiększoną złożonością techniczną i zmniejszoną interpretowalnością. W rezultacie Generative SI nasila szereg ryzyk typowych dla tradycyjnych systemów SI, w tym luki w zabezpieczeniach danych, obawy dotyczące prywatności, naruszenia praw autorskich i uprzedzenia algorytmiczne. Niniejsze badanie krytycznie ocenia adekwatność obecnych ram regulacyjnych w zakresie radzenia sobie z tymi ryzykami. Wyniki sugerują, że tradycyjne instrumenty prawne są niewystarczające do radzenia sobie z wyjątkowymi wyzwaniami stawianymi przez Generatywną SI, zwiększając prawdopodobieństwo niepowodzenia regulacyjnego. Przeprowadzając analizę porównawczą podejść regulacyjnych w Unii Europejskiej, Stanach Zjednoczonych, Wielkiej Brytanii i Chinach, badanie ujawnia rosnący globalny konsensus co do konieczności rozwijania odpowiedzialnej i godnej zaufania SI, która jest zgodna z Celami zrównoważonego rozwoju. W dokumencie postuluje się ustanowienie zrównoważonych i elastycznych ram regulacyjnych, których celem będzie zwiększenie zdolności rządowych do radzenia sobie ze zmieniającymi się ryzykami związanymi z generatywną sztuczną inteligencją, przy jednoczesnym promowaniu jej bezpiecznego, odpowiedzialnego i zrównoważonego rozwoju.

Słowa kluczowe: Sztuczna Inteligencja (SI), generatywna SI, podejście regulacyjne, zrównoważony rozwój

1. Introduction

Generative Artificial Intelligence (Generative AI) refers to a category of intelligent technologies capable of autonomously producing diverse forms of content, including text, images, audio, and video. By leveraging large-scale datasets, these models generate novel outputs with distinctive features of creativity and innovation. Representative models of Generative AI include Variational Autoencoders (VAEs), Generative Adversarial Networks (GANs), Diffusion Models, Generative Pre-trained Transformers (GPT), and Neural Radiance Fields (NeRFs) (Dubey & Singh, 2024; Singh & Ogunfunmi, 2021; Yang et al., 2023). Owing to its wide range of applications, Generative AI is already driving substantial transformations in key sectors such as healthcare, education, and climate action.

Importantly, the role of Generative AI in advancing the United Nations Sustainable Development Goals (SDGs) has attracted growing international attention. Established at the 2015 UN Sustainable Development Summit, the 17 SDGs aim to provide a comprehensive framework to tackle global challenges such as poverty, inequality, and environmental degradation, and are widely recognized as a blueprint for a peaceful, prosperous, and sustainable future (United Nations, 2015). The term ‘sustainability,’ used interchangeably with the term ‘sustainable development’ in the academic literature has started gaining prominence since its appearance in the Brundtland Report in 1987 (Brundtland 1987 cited in Pahuja et al., 2025). Generative AI is increasingly viewed as a key enabler of these goals, with applications that promote green economic transitions, improve public services and education access, enhance global cooperation, and facilitate open and inclusive innovation (Vinuesa et al., 2020). The academic community generally acknowledges that Generative AI presents both opportunities and risks for sustainable development. While it holds promises for addressing complex global challenges, it simultaneously raises concerns regarding the digital divide, resource consumption, and ethical dilemmas. For example, in the context of climate change, Generative AI can support environmental monitoring and the design of sustainable solutions. At the same time, however, its large-scale computational demands pose environmental risks. According to The New Yorker, training OpenAI’s GPT-3 consumed approximately 1.287 million kilowatt-hours of electricity. After deployment, with an estimated 200 million daily requests, ChatGPT’s inference service may consume up to 500,000 kilowatt-hours per day, further highlighting the ecological impact of AI infrastructure (Strubell, Ganesh, & McCallum, 2019).

In addition, the rapid progress of Generative AI is accompanied by significant safety and socio-ethical concerns. Generative AI rely heavily on extensive datasets to enhance their capabilities in understanding, using, and generating human language. This reliance often involves utilizing copyrighted works as training data, which has sparked significant controversy and legal disputes. The incorporation of such materials without proper consent has led to numerous copyright infringement claims. Artists, writers, and programmers, aggrieved by the perceived exploitation of their work for profit by tech companies, have initiated class-action lawsuits against major entities like OpenAI and Microsoft, alleging unauthorized use of their copyrighted content (Zhang & Kamel Boulos, 2023). Moreover, Generative AI deep forgery technology has the capability to produce highly realistic synthetic images, videos, and even human voice simulations. The misapplication of this technology introduces significant risks, including the potential for severe privacy breaches, identity theft, and the propagation of disinformation with the capacity to disrupt social stability and influence political processes. As the underlying algorithms and models become more sophisticated, the distinction between authentic and manipulated content becomes increasingly obscure, presenting profound challenges for cybersecurity, information integrity, and public trust in digital communications (Amoozadeh et al., 2024).

Furthermore, Generative AI technologies, including chatbots and image generators, frequently perpetuate stereotypes by depicting engineers predominantly as white and male and nurses as white and female. Additionally, facial recognition systems employed by law enforcement are prone to higher rates of misidentification for Black individuals, which can result in wrongful arrests. Similarly, hiring algorithms often exhibit a gender bias, favouring male candidates over female candidates. These phenomena illustrate systemic biases embedded in Generative AI, raising significant concerns about their impact on perpetuating and amplifying existing social inequalities. Addressing these biases is crucial for the equitable deployment and ethical use of AI technologies (Zhou et al., 2024).

David Collingridge’s control dilemma highlights the challenge of regulating emerging technologies: implementing regulations too early may stifle innovation and hinder technological progress while delaying regulation can lead to harmful social effects that become difficult to manage once they surface (Collingridge, 1982). This dilemma is particularly relevant to Generative AI. For instance, if the unauthorized use of copyrighted materials for training data is deemed a copyright infringement, it could pose significant legal challenges for the AI industry. Conversely, if AI companies are allowed to freely use copyrighted works, it could not only undermine the rights of human creators but also lead to ethical issues. As a result, the debate over whether and how AI should be regulated has intensified. Given these regulatory tensions, scholars have increasingly emphasized the need for a proactive and coordinated global response. In particular, there is growing recognition that AI governance must not only address immediate legal and ethical concerns but also align with broader societal objectives, such as the pursuit of sustainable development. Truby (2020) argues that regulatory frameworks for AI should be designed to support the achievement of the Sustainable Development Goals (SDGs), while Djeflal (2020) advocates for integrating *sustainable development* as both a normative foundation and a strategic objective in AI governance. He proposes a three-tiered framework for sustainable AI development – comprising technical, societal, and governance dimensions – where the technical layer draws on Mulder’s concept of *sustainable technology*, the societal layer assesses the broader social implications of AI, and the governance layer aims explicitly at fulfilling the SDGs.

This study aims to systematically analyze the risks and challenges associated with Generative AI. Through qualitative content analysis, it provides an in-depth examination of the regulatory approaches adopted by the European Union (EU), the United States (US), the United Kingdom (UK), and China in response to these emerging issues. To achieve this aim, the study explores the following research questions: What specific risks are associated with Generative AI? How do current regulatory frameworks in the EU, US, UK, and China respond to Generative AI risks? What are the common trends in the regulatory approaches to Generative AI across various countries and regions? Accordingly, the research objectives are threefold: (1) to identify the unique risks introduced by Generative AI; (2) to examine how existing regulatory texts in the EU, US, UK, and China address Generative AI risks; and (3) to analyse the global regulatory trends in Generative AI. The research framework is structured as

follows. First, the paper identifies and categorizes the specific risks inherent in Generative AI models. Second, it conducts a comprehensive analysis of existing legal frameworks governing AI, using theoretical examination and critical evaluation of relevant legal texts to uncover gaps in their capacity to manage the evolving risks of Generative AI. This comparative inquiry, based on document analysis and policy comparison, reveals divergent approaches: while the EU and China emphasize risk-based regulation, the US and UK largely adhere to market-driven models. The insights drawn from this analysis serve as the foundation for recommending improvements to regulatory practice. Finally, building on the findings, the study proposes a flexible, hybrid regulatory framework designed to promote responsible innovation while ensuring effective oversight.

2. Risks and challenges of generative AI

Compared to the security issues of traditional AI systems, the automatic content generation enabled by the big new data-driven approach in Generative AI introduces new risks and challenges, including security and privacy risks, copyright infringement, algorithmic bias, and discrimination risks.

2.1. Data security and privacy risks

The rapid advancement of Generative AI has necessitated urgent global attention to data security and privacy concerns. With the exponential growth in the capabilities of AI models, including those developed by OpenAI, there has been a corresponding increase in the volume of data required for training these systems. This data often includes sensitive personal information, raising significant privacy risks. On March 31, 2023, the Italian Personal Data Protection Authority imposed a temporary ban on ChatGPT, citing concerns over data breaches involving user conversations and payment information (Arcila, 2023). Subsequently, the U.S. Department of Commerce issued an AI Accountability Policy, emphasizing the importance of data security and privacy protection for AI accountability (NTIA, 2024). These regulatory moves underscore a global consensus on the critical importance of data security and privacy protection in the development of Generative AI.

Generative AI necessitate the ingestion of vast datasets to achieve their advanced capabilities. For instance, GPT-1 incorporated 117 million parameters, GPT-2 expanded to 1.5 billion, GPT-3 reached 175 billion, and GPT-4 is anticipated to encompass over 100 trillion parameters (OpenAI, 2023). The data employed in training these models often includes personal and sensitive information, such as user identities, preferences, and behavioral patterns. Insufficient protection of this data poses significant risks of privacy breaches. Despite assurances from Generative AI service providers regarding the security of user data—typically achieved through methods like anonymization and encryption, as outlined in privacy policies and user agreements—evidence from numerous public reports indicates that the risk of data security remains significant. For instance, on March 24, 2023, OpenAI issued an official statement acknowledging that approximately 1.2% of ChatGPT Plus user data was vulnerable to a breach. The compromised data included sensitive information such as users' names, snippets of chat history, email addresses, and payment details. In one notable case, developers at Samsung inadvertently exposed corporate confidential data, including new program source code, while attempting to utilize ChatGPT for code patching. In addition, even anonymized or de-identified datasets are not immune to privacy risks. Generative AI possess the capability to infer personal identifying information and basic characteristics from such datasets, potentially leading to re-identification and subsequent privacy violations. The process of data aggregation and analysis by Generative AI can reveal latent correlations that compromise individual privacy. The training of Generative AI frequently requires access to external data sources, including social media platforms, medical records, and other repositories of personal information. This access introduces additional privacy risks, as data sharing with third-party entities amplifies the likelihood of unauthorized use or disclosure of sensitive information. OpenAI's privacy policy explicitly acknowledges the collection of extensive user information, including account details, conversation content, and various other private data, which may be shared with third-party entities. At Black Hat 2024, cybersecurity researcher Michael Bargury exposed critical vulnerabilities in Microsoft's AI-powered assistant, Copilot. Bargury demonstrated how malicious actors could exploit these vulnerabilities to manipulate Copilot into stealing sensitive data. His research further revealed that Copilot could be co-opted into becoming an effective phishing tool, capable of crafting deceptive messages that could trick users into divulging confidential information. Beyond individual privacy violations, the global nature of AI data collection and processing raises concerns about national data security. The cross-border flow of data, integral to the training of Generative AI models, could inadvertently expose sensitive national information, posing a threat to political and economic stability.

2.2. Copyright issues

Generative AI, such as ChatGPT, have revolutionized content creation by producing text, images, and other media that mimic human output. However, these advancements come with significant legal and ethical challenges, particularly regarding copyright infringement. Training Generative AI requires vast amounts of data, which often includes copyrighted works such as books, images, music, and other intellectual property. If AI developers use this data without permission from the copyright holders, it can lead to copyright infringement (Smits & Borghuis, 2022). Copyright law grants creators' exclusive rights over their works, including the right to control reproduction, distribution, and adaptation. Copyright laws are designed to protect creators' rights, and unauthorized use of these works during AI training can undermine these protections. The reliance on copyrighted content for AI training raises questions about fairness to original creators and the potential erosion of intellectual property rights. A notable example is the lawsuit filed by Getty Images against Stability AI, the developer of the AI painting tool Stable Diffusion. Getty Images accused Stability AI of using millions of images from its website without permission, illustrating the risks associated with using copyrighted material in AI training (James, 2023). Additionally, several artists have sued AI developers for allegedly using billions of images scraped from the internet without consent, claiming that this practice violates their rights and devalues their original works (Sophia, 2024).

The content generated by Generative AI, based on their training data, can infringe on existing copyrights by producing works that resemble or derive from the original copyrighted material. A derivative work is one that is based on or incorporates substantial elements of a pre-existing work, such as adaptations or transformations. If AI-generated content (AIGC) closely mirrors

or directly incorporates elements of copyrighted works, it could be considered a derivative work, leading to copyright infringement. For instance, using ChatGPT to generate summaries of copyrighted books can pose a significant copyright risk. These summaries could serve as substitutes for the original works. Another complex issue arises when AI models generate content in the style of specific authors or artists. While copyright law generally does not protect artistic styles, the creation of works that closely imitate the original could still raise legal concerns, especially if the generated content competes with the original in the marketplace. The line between permissible influence and copyright infringement becomes blurred in such cases, particularly when AIGC is substantially like the copyrighted work.

ChatGPT exhibit creative capabilities by producing new content, raising the question of whether such content can be recognized as legally protectable work. However, the originality requirement in copyright law – which mandates that a work must possess a minimal degree of creativity – presents challenges for AIGC. While these works are technically new creations, their originality often stems from the recombination of existing material, making their eligibility for copyright protection a contentious issue. In November 2023, the Beijing Internet Court ruled in China's first AIGC case, determining that Generative AI are tools that assist users in the creative process. The court found that the user's contributions – from conceptualizing the image including designing characters, choosing prompts, and setting parameters to selecting the final image – constituted original work. Therefore, the images were protected by copyright law. Conversely, the United States Copyright Office has consistently refused to grant copyright protection to content generated using Generative AI, stating that users cannot claim creative contributions in the process, regardless of the complexity of prompts or modifications.

2.3. Biases and discrimination risks

The training of Generative AI, such as ChatGPT, on vast and predominantly unlabeled datasets inherently carries the risk of embedding and perpetuating biases, discrimination, and harmful content. Examples of harmful content in training data include insult, hate speech, violence, pornography, and other forms of toxic material (e.g., offensive, or threatening information). As these models generate content based on patterns learned from such data, they inevitably reproduce these biases, posing significant risks in their outputs. Among the most critical of these risks are prejudice and discrimination, which are systemic issues that can manifest in AI outputs as biased and unjust treatment of specific groups. The process of manual data annotation introduces biases as human annotators may unconsciously project their own prejudices during the interpretation and labeling stages. Moreover, during data processing, the model's conclusions are sometimes adjusted to align with societal expectations, inadvertently reinforcing existing biases.

The National Institute of Standards and Technology (NIST) categorizes AI bias into three main types: systemic bias, statistical and computational bias, and human bias. Systemic bias stems from the cultural, institutional, and societal norms that influence data collection and interpretation. Statistical and computational bias occurs due to the underrepresentation of certain demographic groups within the training datasets, resulting in skewed and often inaccurate outputs. Human bias arises from the cognitive errors and subjective judgments made by individuals involved in the data annotation and model development processes. Empirical studies have highlighted these biases in generative models like Stable Diffusion (Schramowski et al., 2023). For instance, when asked to generate images of a *CEO*, these models predominantly produce images of men in business attire, reflecting both systemic and statistical biases inherent in the training data. Systemic bias is particularly concerning when training data is sourced from a specific language or cultural context, as it embeds the dominant cultural traditions, values, and ideologies of that context into the model. This not only risks propagating these biases on a global scale but also raises the potential for cultural and ideological conflicts when such models are deployed in diverse environments, thus posing ideological security risks (Struppek et al., 2023). Furthermore, the integration of *machine learning* with *manual annotation* to enhance the intelligence and accuracy of Generative AI models exacerbates the problem of algorithmic bias. This approach allows for the infusion of human subjective judgments and preferences into the model, making these biases more entrenched and challenging to detect or mitigate.

Generative AI such as ChatGPT rely on reinforcement learning from human feedback to optimize their language generation processes. While the primary goal of these models is to align AI behavior with human intentions, the outcomes often deviate from this objective. Problems such as the inability of moral algorithms to filter out inherent biases in foundational text data, the amplification of implicit biases by reward algorithms, and the reinforcement of algorithmic biases through repeated training cycles are prevalent. These biases significantly influence the model's outputs, raising serious concerns about gender discrimination, racial discrimination, and other forms of social inequity. Gender discrimination, exacerbated by algorithmic bias, is particularly concerning. Historical social and biological differences between men and women have led to distinct roles and expectations, fostering gender biases that result in systemic inequities and unfair treatment of women across various sectors. When Generative AI algorithms internalize these biases, they can amplify sexism in their outputs (Bansal et al., 2022). For example, a UNESCO study published on March 7, 2024, reported that Generative AI models tend to exacerbate gender biases, with women depicted as performing housework four times more often than men in descriptions generated by large language models (Hacker, Mittelstadt, Borgesius, & Wachter, 2024). The risk of racial discrimination is similarly troubling. AI's historical development has been marred by persistent racial biases, often concealed under the guise of scientific objectivity. For instance, Microsoft's Tay chatbot was taken offline within 24 hours of its launch after it began generating racist and offensive language learned from interacting with users (Wolf, Miller, & Grodzinsky, 2017). These risks are not merely theoretical but have profound implications for social justice and equity, underscoring the urgent need for comprehensive strategies to identify, mitigate, and ultimately prevent bias in AI systems.

3. Current regulations of generative AI in the EU, US, UK, and China

Major global entities – including the EU, the US, the UK, and China – have actively initiated the development of regulatory frameworks in response to the rapid advancement of Generative AI technologies. Table 1 provides a summary of key policy and regulatory documents issued by these four jurisdictions between 2022 and 2024. The data illustrate the evolving landscape

of Generative AI governance and highlight the increasing regulatory attention devoted to this emerging technology. The findings suggest that these jurisdictions are progressively strengthening their regulatory efforts, aiming to establish frameworks that not only mitigate potential risks but also promote the responsible and sustainable development of Generative AI.

3.1. European Union

The European Union has made considerable progress with the *European Union Artificial Intelligence Act (AIA)*, which adopts a risk-based regulatory approach, mandating fundamental rights impact assessments, transparency requirements, and continuous oversight for high-risk AI systems. The European Union's (AIA) employs a risk-based approach to regulate AI systems, guided by principles of human oversight, technological robustness, privacy, transparency, diversity, and social well-being. It categorizes AI risks into four levels: unacceptable risks, which are banned; high risks, which face stringent regulation; limited risks, which require minimal regulation with user awareness; and minimal risks, subject to basic compliance (European Commission, 2024). Currently, the AIA introduces specific regulations for underlying models – AI systems trained on extensive datasets that can be adapted for various applications – without categorizing Generative AI as high-risk. While Generative AI drive advancements in general AI, the disruptive effects on industrial structures and social relations are not yet fully apparent, making it challenging to evaluate their potential harm. Therefore, it is premature to make a conclusive judgment on their risk level. Instead, a thorough investigation of the specific risks posed by Generative AI and an evaluation of the existing regulatory framework's capacity to manage these risks are essential.

Table 1. The development of regulatory frameworks of generative AI between 2022 and 2024, own elaboration

Jurisdiction	Document Title	Issuing Authority	Key words	Release Date
EU	EDPS Guidelines on Generative AI and Personal Data Protection	European Data Protection Supervisor (EDPS)	GDPR compliance, data protection, privacy	Jun 2024
EU	Artificial Intelligence Act	European Commission	Risk-based approach, mandatory compliance, transparency, accountability	Jul 2024
US	Blueprint for an AI Bill of Rights	White House Office of Science and Technology Policy (OSTP)	Rights protection, data privacy, algorithmic transparency, fairness	Oct 2022
US	Principles for the Development, Deployment, and Use of Generative AI Technologies	Association for Computing Machinery's global Technology Policy Council (ACM TPC)	Generative AI, development principles, accountability	Jul 2023
US	Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence	White House	Safety, trust, federal AI strategy, national security	Oct 2023
US	Generative AI Copyright Disclosure Act	U.S. Congress (Proposed)	Copyright disclosure, transparency, creator rights	Apr 2024
US	The Content Origin Protection and Integrity from Edited and Deepfaked Media (COPIED) Act	U.S. Congress (Proposed)	Originality, AI-generated content, copyright protection	Jul 2024
UK	Establishing a Pro-Innovation Approach to Regulating AI.	UK Government	flexibility, proportionality, and a context-specific approach	Jul 2022
UK	A Pro-Innovation Approach to AI Regulation	UK Government	Pro-innovation, flexible regulation, regulatory coordination	Mar 2023
China	Administrative Measures for Generative AI Services (Draft for Public Comment)	Cyberspace Administration of China (CAC)	Safety, algorithm accountability, content supervision	Apr 2023
China	Interim Measures for the Management of Generative AI Services	Cyberspace Administration of China (CAC)	Data compliance, user protection, platform responsibility	Jul 2023
China	Basic Requirements for the Security of Generative AI Services	National Information Security Standardization Technical Committee	Technical security standards, system protection, risk control	Oct 2023

The EU has taken a significant step toward mitigating the potential risks of Generative AI by enacting the AIA, which establishes comprehensive obligations for providers of foundational AI models. It outlines three key obligations for providers of these underlying models. First, a risk management obligation requires comprehensive measures to ensure the model's safety and stability, including risk analysis, data source management, performance evaluation, energy consumption control, and quality management. Second, an information guarantee obligation mandates that providers supply clear technical documentation and usage instructions to aid downstream users in the compliant development and application of high-risk AI systems. Third,

additional requirements for Generative AI include ensuring that generated content does not contain illegal material, does not infringe on fundamental rights, and provides summaries of the use of copyrighted training data (European Commission, 2024). This legislation outlines requirements related to transparency, safety, and traceability, and advocates for full lifecycle governance of AI systems. Notably, it emphasizes the involvement of independent experts in compliance assessments and technical evaluations, thereby reinforcing the objectivity and credibility of the regulatory process. In parallel, the EDPS released the *Guidelines on Generative AI and Personal Data Protection* – the first EU-level document to specifically address data protection issues associated with Generative AI. These guidelines articulate the legal principles and standards that must be followed when processing personal data using Generative AI technologies. Key requirements include lawfulness, transparency, data minimization, purpose limitation, and the safeguarding of data subject rights. The EDPS stresses that public and private actors must ensure that the deployment of Generative AI systems respects fundamental rights and data protection obligations, promoting responsible innovation within a legally and ethically sound framework.

3.2. The United States

In October 2022, the United States introduced the *Blueprint for an AI Bill of Rights* as a foundational framework for its approach to AI governance. Issued by the OSTP, this document outlines five core principles designed to safeguard the rights of individuals in the development and deployment of automated systems: safe and effective systems, algorithmic discrimination protections, data privacy, notice and explanation, and human alternatives, consideration, and fallback. Although not legally binding, the blueprint serves as a policy guide for promoting responsible and rights-oriented AI development across federal agencies and the private sector. In July 2023, the ACM TPC published its *Principles for the Development, Deployment, and Use of Generative AI Technologies*. This document offers a set of globally applicable principles emphasizing transparency, accountability, fairness, robustness, and respect for human autonomy in the context of Generative AI. It underscores the importance of interdisciplinary collaboration and ethical foresight in guiding the rapidly evolving Generative AI landscape.

The United States has embraced a market-oriented approach to AI regulation, exemplified by President Biden's *Executive Order on Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence* (AIEO) issued in October 2023. This order emphasizes transparency and accountability, mandating that AI companies disclose safety test results to federal authorities. It encompasses eight policy domains: safety and security, innovation, labour support, AI bias and civil rights, consumer protection, privacy, federal AI usage, and international leadership. The AIEO directs over 50 federal agencies to undertake more than 100 specific tasks and establishes a White House Council on AI composed of leaders from 28 federal departments (AIEO, 2023). Subsequently, the U.S. National Institute of Standards and Technology (NIST) released four draft guidance documents on April 29, 2024. These documents focus on mitigating risks related to Generative AI, safeguarding data used in AI training, managing synthetic content risks, and fostering international collaboration on AI standards (U.S. Department of Commerce, 2024).

In April and July 2024, the United States introduced two notable legislative proposals aimed at addressing the regulatory challenges posed by Generative AI. The first, the *Generative AI Copyright Disclosure Act*, seeks to enhance transparency by requiring developers of Generative AI systems to disclose the copyrighted materials used in training their models. This proposed legislation reflects growing concerns about intellectual property rights in the age of AI, particularly regarding the unauthorized use of creative works in large-scale datasets. Subsequently, in July 2024, U.S. lawmakers proposed the COPIED Act, which aims to combat misinformation and manipulation by mandating clear labelling of AI-generated content and establishing penalties for the malicious creation and distribution of deepfakes. Federal Trade Commissioner Alvaro Bedoya has indicated that Generative AI is regulated under Section 5 of the Federal Trade Commission Act, which addresses unfair and deceptive practices. Additionally, US Senate Majority Leader Chuck Schumer noted that the US legislative framework for AI is evolving to create adaptable regulations that balance safety, accountability, and transparency with the promotion of innovation. The US prioritizes technological advancement over stringent regulation, aiming to sustain its leadership in the AI sector. In 2023, prominent tech companies, including Google, Microsoft, Amazon, Meta, OpenAI, Anthropic, and Inflection, entered into voluntary agreements with the White House to enhance responsible AI development. These agreements involve independent expert evaluations prior to public release, comprehensive studies on societal impacts such as bias and privacy, and third-party audits to identify and mitigate vulnerabilities. This approach highlights how the US strategically balances innovation with ethical considerations through industry-led self-regulation.

However, despite these initiatives, relatively few AI-related legislative proposals in the U.S. have been successfully enacted into binding law. This pattern reveals a growing awareness within the U.S. government of the strategic importance of AI governance and a willingness to engage with its implications. At the same time, it also highlights structural and political challenges in the U.S. policy-making process, including stakeholder disagreements, concerns over overregulation stifling innovation, and the complexity of achieving bipartisan consensus on emerging technology issues.

3.3. The United Kingdom

In July 2022, the UK government – through the Department for Digital, Culture, Media and Sport (DCMS) and the Office for AI – published a policy paper titled *Establishing a Pro-Innovation Approach to Regulating AI*, which outlined a flexible, context-specific regulatory strategy designed to support innovation while managing emerging risks. The framework emphasized empowering existing regulators, ensuring proportionality, and promoting principles such as safety, transparency, and accountability (Department for Digital, Culture, Media & Sport, 2022). This laid the groundwork for the more comprehensive *A pro-innovation approach to AI regulation* released in March 2023, which further developed the UK's pro-innovation regulatory strategy. It sets out five cross-sectoral regulatory principles: safety, security and robustness; appropriate transparency and explainability; fairness; accountability and governance; and contestability and redress (Department for Science, Innovation and Technology, 2023). This document seeks to advance responsible AI innovation and reinforce the UK's position as a global leader in AI, aligning with OECD principles on safety, security, robustness, transparency, explainability, fairness, accounta-

bility, governance, competition, and resilience. However, the AI supply chain – including algorithms, datasets, parameter designs, and model training architectures – often lacks transparency, posing challenges to effective risk regulation and accountability in practice.

The UK government plans to integrate existing regulatory frameworks, such as the Equality Act 2010 and the UK General Data Protection Regulation (UK GDPR), into its AI oversight strategy. Regulatory duties are distributed among various agencies, including the Health and Safety Executive, the Equality and Human Rights Commission, and the Competition and Markets Authority, each adopting a tailored approach to managing AI. Additionally, the UK has established a central monitoring and evaluation framework to assess and enhance the effectiveness of its regulatory measures. This flexible governance model is designed to address the rapid pace of technological advancement, advocating for adaptable regulations that support innovation while managing risks. This approach contrasts with rigid, one-size-fits-all legislation, promoting an industry-led strategy that facilitates nuanced risk management across sectors. Nonetheless, this model's reliance on collaboration among government, regulators, and the public presents challenges in achieving effective coordination and cooperation.

3.4. China

In China, the government has prioritized the advancement of General AI within a framework that fosters innovation while maintaining a focus on risk management. The inclusion of the Draft Artificial Intelligence Law in the State Council's legislative work plan in May 2023, followed by the formulation of the *Interim Measures for the Management of Generative AI Services* (Interim Measures) in July 2023, underscores China's commitment to balancing technological development with safety. The Interim Measures, jointly issued by seven Chinese government departments – including the National Development and Reform Commission, Ministry of Education, and Ministry of Public Security – govern services that use Generative AI to create and distribute content such as text, images, audio, and video to the public. The Interim Measures emphasize regulatory flexibility and shared responsibility. First, they limit their scope to the service level, exempting the research and design stage to foster innovation. The measures adopt a dual regulatory approach: classification-based supervision, which tailors regulatory intensity to specific application scenarios and risk levels, and industry-specific supervision, which allows sectors to apply rules relevant to their operational context. Second, the Measures impose obligations on both providers and users of Generative AI services. Providers must assume responsibility for generated content and comply with cybersecurity requirements, including issuing warnings, restricting or terminating services, and keeping operational records. They must also establish service agreements with users that clarify mutual rights and responsibilities. While users bear legal liability for illegal use – such as disseminating fake news via Generative AI – providers are expected to implement corrective actions to mitigate harm and ensure compliance with national security standards.

As of 2024, under the regulatory framework of the Interim Measures, two major filing mechanisms have been established in the domain of Generative AI services: algorithm filing and large model deployment filing. The algorithm filing mechanism is based on Article 17 of the Interim Measures, which requires service providers fulfilling filing, modification, and deregistration procedures. Meanwhile, the filing requirement for large model deployment refers to the obligation to conduct security assessments *in accordance with relevant national regulations*, as also stipulated in Article 17, and is specifically targeted at safety evaluations in the AIGC (AI-Generated Content) sector. To support the implementation of the Interim Measures, the National Information Security Standardization Technical Committee released the *Basic Requirements for the Security of Generative AI Services*. This document defines fundamental safety requirements for Generative AI services, including corpus security, model security, security measures, and content governance, and provides a basis for conducting security assessments. It clarifies the standards for the security assessments required by Article 17 of the Interim Measures, obliging service providers to thoroughly evaluate these four core areas. This domain-specific evaluation approach enables service providers to systematically identify and address potential security risks, thereby improving the efficiency and effectiveness of their security management systems. By adopting these standards, providers can better ensure compliance of generated content, safeguard user privacy, and uphold cybersecurity.

3.5. Comparative analysis of Generative AI regulatory approaches

Similarities: First, the EU, the US, the UK, and China have all recognized that Generative AI introduces distinctive risks compared to traditional AI systems and are actively working to establish dedicated regulatory frameworks. These initiatives reflect a shared understanding that existing AI governance models are insufficient to fully address the unique challenges posed by Generative AI, such as misinformation, bias amplification, and intellectual property infringement. Second, all jurisdictions aim to ensure that Generative AI technologies are developed and deployed in a safe, reliable, and sustainable manner. They seek to promote technological innovation and the widespread adoption of Generative AI while simultaneously implementing risk mitigation measures to ensure that innovation remains controllable and socially beneficial. Third, strengthening international cooperation and dialogue has become a common priority. Recognizing that Generative AI technologies transcend national borders, major regulatory authorities have begun emphasizing the need for harmonized standards, shared best practices, and participation in global initiatives to address cross-border challenges.

Differences: The United States prioritizes fostering innovation leadership in Generative AI. Its regulatory philosophy emphasizes minimizing government intervention to avoid stifling technological progress. U.S. policies favor a sector-specific, flexible market governance approach, relying heavily on voluntary industry standards and self-regulation. Although privacy, intellectual property protection, and national security concerns are acknowledged, enforcement is often decentralized, and concrete legislation remains limited. In contrast, the European Union emphasizes the establishment of a comprehensive, legally binding framework, exemplified by the AI Act, which imposes stringent obligations on developers and deployers of foundation models. The EU adopts a risk-based classification system, requiring rigorous oversight throughout the Generative AI lifecycle – from development to deployment and use – with particular emphasis on transparency, accountability, and fundamental rights protection. The United Kingdom promotes a pro-innovation regulatory model grounded in five core principles: safety, transparency, fairness, accountability, and contestability. Rather than legislating immediately, the UK has issued regulatory guidance encouraging sector-specific regulators to adapt existing laws while maintaining flexibility. Its approach balances innovation support

with safeguarding public trust. China combines government-led oversight with the establishment of technical standards and regulatory instruments. Through measures such as the Interim Measures and supporting technical guidelines, China implements a classified and graded supervision model that differentiates oversight based on service scenarios and risk profiles. The emphasis is on balancing innovation with societal stability, requiring providers to undergo security assessments and comply with content governance rules. Furthermore, China underscores inter-agency coordination and dynamic risk management to adapt regulatory strategies as the technology evolves.

Table 2. Comparative features of generative AI regulation in the EU, US, UK, and China, own elaboration

Jurisdiction	Regulatory Approach	Main Focus	Key Characteristics
EU	Comprehensive legal framework (AI Act)	Risk-based classification; lifecycle regulation	Strict legal obligations; strong fundamental rights protection; transparency; accountability across development, deployment, and use stages
US	Sector-specific, market-driven governance	Encourage innovation; protect privacy and IP	Minimal legislative intervention; industry self-regulation; decentralized enforcement; emphasis on national leadership in AI
UK	Pro-innovation flexible model	Sector-specific adaptation; regulatory agility	Five principles (safety, transparency, fairness, accountability, contestability); guidance over immediate legislation; emphasis on responsible innovation
China	Government-led regulation with technical standards	Security risk management; content governance	Classified and graded supervision; mandatory security assessments; inter-agency coordination; dynamic regulatory updates

Table 2 summarizes the key features of Generative AI regulation across the EU, the US, the UK, and China. It highlights that while all four jurisdictions recognize the need to address the unique risks of Generative AI and aim to ensure its safe, reliable, and sustainable development, they adopt distinct regulatory strategies. The EU emphasizes a comprehensive legal framework with risk-based classification and lifecycle regulation. The US favors a market-driven, sector-specific approach that promotes innovation with limited legislative intervention. The UK adopts a pro-innovation, principle-based strategy, focusing on regulatory flexibility and accountability. China, meanwhile, combines government-led oversight with technical standards, focusing on classified supervision and comprehensive security risk management. These differences reflect each jurisdiction's policy priorities and governance traditions.

4. Sustainable regulatory trends in generative AI

Through a comparative study of the regulatory policies on Generative AI in the EU, the US, the UK, and China, a clear trend emerges: all four regions emphasize the importance of developing specialized regulatory frameworks to address the unique risks associated with Generative AI, which differ from those of traditional AI. This shared commitment to ensuring the safe, reliable, and sustainable development of Generative AI is evident in the regulatory approaches of each jurisdiction. On one hand, these countries aim to foster technological innovation and encourage the widespread adoption of Generative AI. On the other hand, they also prioritize managing the potential risks associated with these emerging technologies.

The regulatory evolution across these regions reflects a significant shift towards dynamic and adaptive regulatory strategies. For example, the EU, through its AI Act, has outlined comprehensive regulations for AI, including generative models, by classifying AI systems based on risk levels and tailoring obligations accordingly. The US, while initially focused on minimizing regulatory burdens to promote innovation, has increasingly incorporated privacy and intellectual property protections into its framework, with specific regulations now being proposed for Generative AI technologies. The UK's approach, centered around five guiding principles, advocates for a balance between encouraging innovation and ensuring accountability and governance. In China, the government has emphasized a hybrid regulatory model, combining government oversight with technical standards and promoting sector-specific collaboration through safety assessments and algorithm registration requirements. The development of these regulatory frameworks aligns with theories of adaptive governance and risk governance. Adaptive governance emphasizes the need for regulatory systems to remain flexible and responsive to the rapidly evolving nature of technology (Folke et al., 2005). Risk governance, on the other hand, advocates for a comprehensive approach to managing complex and uncertain risks through interdisciplinary collaboration (Renn, 2008). These theoretical foundations suggest that regulatory frameworks for Generative AI must not only address current risks but also be adaptable to future challenges, ensuring that regulations evolve in tandem with technological advancements.

In conclusion, the trend toward sustainable and adaptive regulation is becoming a central theme in the global governance of Generative AI. Going forward, the development of regulatory frameworks that are flexible, risk-sensitive, and capable of evolving with the technology will be critical to ensuring the safe and responsible growth of Generative AI across jurisdictions.

5. Conclusion

This paper highlights the complex risk landscape posed by Generative AI, emphasizing the inadequacy of traditional regulatory frameworks in addressing their unique challenges. The study identifies critical risks, including security and privacy risks, copyright infringement, algorithmic bias, and discrimination risks. These risks are compounded by the rapid pace of technological advancements, which current legal frameworks struggle to keep up with. Specifically, regulations related to data protection and liability are often too general and fail to address the specificities of Generative AI, demonstrating an urgent need for more agile, forward-thinking regulatory approaches.

The comparative analysis of regulatory strategies in selected jurisdictions—China, the EU, the UK, and the US—reveals diverse approaches with distinct emphases. Both China and the EU adopt risk-based models that focus on stringent risk mitigation,

emphasizing comprehensive governance mechanisms to manage the inherent complexities of Generative AI. These jurisdictions have established detailed regulatory structures, particularly regarding data protection, algorithmic transparency, and AI accountability. Conversely, the UK and the US prioritize market-driven models that emphasize fostering innovation and technological advancement, often at the expense of comprehensive regulatory oversight. This divergence highlights a critical gap in balancing innovation with regulatory caution, underscoring the need for a more flexible governance framework capable of reconciling rapid technological progress with the protection of public welfare. Based on these findings, it is clear that the future of Generative AI regulation requires a proactive and adaptive strategy. This approach must not only address current risks but also anticipate emerging challenges. A sustainable regulatory framework should be capable of evolving with the technology, ensuring that AI innovations continue to develop within a system that safeguards societal interests while fostering technological progress.

However, this study has certain limitations that should be acknowledged. Firstly, the comparative analysis was based on a limited selection of jurisdictions, and further research could extend this analysis to other countries with differing regulatory approaches, such as Japan, Canada, and Australia. Secondly, the rapid development of Generative AI means that new risks and challenges may emerge, which were not covered within the scope of this research. Future studies should therefore adopt longitudinal approaches to track the effectiveness of existing frameworks over time and evaluate how they adapt to new technological shifts. Additionally, research into the role of private sector self-regulation and industry standards could provide deeper insights into complementing public regulatory efforts.

References

1. AMOOZADEH M., DANIELS D., NAM D., KUMAR A., CHEN S., HILTON M. et al., 2024, Trust in Generative AI among students: An exploratory study, *Proceedings of the 55th ACM Technical Symposium on Computer Science Education V. 1*, 67–73.
2. ARCILA B.B., 2023, Is It a Platform? Is It a Search Engine? It's ChatGPT! The European Liability Regime for Large Language Models, *J. Free Speech L.*, 3, 455.
3. BANSAL H., YIN D., MONAJATIPOOR M., CHANG K.W., 2022, How well can text-to-image generative models understand ethical natural language interventions? *arXiv preprint*, arXiv:2210.15230.
4. CHANG Y., WANG X., WANG J., WU Y., YANG L., ZHU K., et al., 2024, A survey on evaluation of large language models, *ACM Transactions on Intelligent Systems and Technology*, 15(3), 1–45.
5. COLLINGRIDGE D., 1982, *The social control of technology*, Francis Pinter, London.
6. DEPARTMENT FOR DIGITAL, CULTURE, MEDIA & SPORT, 2022, *Establishing a pro-innovation approach to regulating AI*, UK Government, <https://www.gov.uk/government/publications/establishing-a-pro-innovation-approach-to-regulating-ai/establishing-a-pro-innovation-approach-to-regulating-ai-policy-statement>.
7. DEPARTMENT FOR SCIENCE, INNOVATION AND TECHNOLOGY, 2023, *A pro-innovation approach to AI regulation (CP 796)*, GOV.UK, <https://www.gov.uk/government/publications/ai-regulation-a-pro-innovation-approach/white-paper>.
8. DJEFFAL C., 2020, Sustainable AI development (SAID): On the road to more access to justice, *Technology, innovation and access to justice: Dialogues on the future of law*, eds. De Souza S.P., Sphorm M., *Edinburgh University Press*, 112–130. <https://doi.org/10.1515/9781474473880-013>.
9. DUBEY S.R., SINGH S.K., 2024, Transformer-based generative adversarial networks in computer vision: A comprehensive survey, *IEEE Transactions on Artificial Intelligence*.
10. EUROPEAN COMMISSION, 2024, *Artificial Intelligence Act*, https://ec.europa.eu/digital-strategy/our-policies/artificial-intelligence_en.
11. EXECUTIVE OFFICE OF THE PRESIDENT, 2023, *Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence*, <https://www.whitehouse.gov/briefing-room/statements-releases/2023/07/27/fact-sheet-executive-order-on-artificial-intelligence/>.
12. EXECUTIVE ORDER NO. 14110, 2023, 88 Fed. Reg. 75191 [hereinafter AI EO], at Sec. 2.
13. FOLKE C., HAHN T., OLSSON P., NORBERG J., 2005, Adaptive governance of social-ecological systems, *Annual Review of Environment and Resources*, 30: 441–473, <https://doi.org/10.1146/annurev.energy.30.050504.144511>.
14. GENERATIVE MODELS: VAEs, GANs, diffusion, transformers, NeRFs. <https://www.techtarget.com/searchenterpriseai/tip/Generative-models-VAEs-GANs-diffusion-transformers-NeRFs>.
15. HACKER P., ENGEL A., MAUER M., 2023, Regulating ChatGPT and other large generative AI models, *Proceedings of the 2023 ACM Conference on Fairness, Accountability, and Transparency*, 1112–1123, <https://doi.org/10.48550/arXiv.2302.02337>.
16. HACKER P., MITTELSTADT B., BORGESIU F.Z., WACHTER S., 2024, Generative Discrimination: What Happens When Generative AI Exhibits Bias, and What Can Be Done About It. *arXiv preprint*, arXiv:2407.10329.
17. JAMES W., 2023, *AI art copyright lawsuit: Getty Images vs. Stable Diffusion*, <https://www.theverge.com/2023/2/6/23587393/ai-art-copyright-lawsuit-getty-images-stable-diffusion>.
18. KASNECI E., SEßLER K., KÜCHEMANN S., BANNERT M., DEMENTIEVA D., FISCHER F. et al., 2023, ChatGPT for good? On opportunities and challenges of large language models for education, *Learning and Individual Differences*, 103: 102274.
19. LIN J.C., YOUNESSI D.N., KURAPATI S.S., TANG O.Y., SCOTT I.U., 2023, Comparison of GPT-3.5, GPT-4, and human user performance on a practice ophthalmology written examination, *Eye*, 37(17), 3694–3695.
20. NATIONAL TELECOMMUNICATIONS AND INFORMATION ADMINISTRATION (NTIA), 2024, *AI accountability policy report: Overview*, <https://www.ntia.gov/issues/artificial-intelligence/ai-accountability-policy-report/overview>.
21. OPENAI, 2023, *GPT-4 Technical Report*, <https://cdn.openai.com/papers/gpt-4.pdf>.

22. PAHUJA A., HASSAN R., CHANDEL A., KAUSHIK N., BHANOT N., FAREED A., 2025. Domain Development and Future Research Agendas in Corporate Governance and Sustainability Research: A Review. *Sustainable Development*, 0:1–22 <https://doi.org/10.1002/sd.3559>
23. RENN O., 2008, Risk governance: Coping with uncertainty in a complex world (2nd ed.), *Earthscan*.
24. SCHRAMOWSKI P., BRACK M., DEISEROTH B., KERSTING K., 2023, Safe latent diffusion: Mitigating inappropriate degeneration in diffusion models, *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 22522–22531.
25. SINGH A., OGUNFUNMI T., 2021, An overview of variational autoencoders for source separation, finance, and bio-signal applications, *Entropy*, 24(1), 55.
26. SMITS J., BORGHUIS T., 2022, *Generative AI and intellectual property rights, Law and Artificial Intelligence: Regulating AI and Applying AI in Legal Practice*, T.M.C. Asser Press, Hague, 323–344.
27. SOPHIA W., 2024, Artificial intelligence and artists' intellectual property: Unpacking copyright infringement allegations in Andersen v. Stability AI Ltd., <https://itsartlaw.org/2024/02/26/artificial-intelligence-and-artists-intellectual-property-unpacking-copyright-infringement-allegations-in-andersen-v-stability-ai-ltd/>.
28. STATISTA MARKET INSIGHTS, 2024, Generative AI – worldwide. Statista, <https://www.statista.com/outlook/tmo/artificial-intelligence/generative-ai/worldwide>.
29. STRUBELL E., GANESH A., MCCALLUM A., 2019, Energy and policy considerations for deep learning in natural language processing, *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, 3645–3650, <https://doi.org/10.18653/v1/P19-1355>.
30. STRUPPEK L., HINTERSDORF D., FRIEDRICH F., SCHRAMOWSKI P., KERSTING K., 2023, Exploiting cultural biases via homoglyphs in text-to-image synthesis, *Journal of Artificial Intelligence Research*, 78, 1017–1068.
31. TORTORA L., 2024, Beyond Discrimination: Generative AI Applications and Ethical Challenges in Forensic Psychiatry, *Frontiers in Psychiatry*, 15, 1346059.
32. TRUBY J., 2020, Governing artificial intelligence to benefit the UN Sustainable Development Goals, *Sustainable Development*, 28(4), 946–959, <https://doi.org/10.1002/sd.2048>.
33. U.S. DEPARTMENT OF COMMERCE, 2024, *Department of Commerce announces new actions to implement President Biden's strategy*, <https://www.commerce.gov/news/press-releases/2024/04/department-commerce-announces-new-actions-implement-president-bidens>.
34. UNITED NATIONS, 2015, *The 2030 Agenda for Sustainable Development*, <https://sdgs.un.org/2030agenda>.
35. VINUESA R., AZIZPOUR H., LEITE I., et al., 2020, The role of artificial intelligence in achieving the Sustainable Development Goals, *Nature Communications* 11(1), 233, <https://doi.org/10.1038/s41467-019-14108-y>.
36. WOLF M.J., MILLER K., GRODZINSKY F.S., 2017, Why we should have seen that coming: comments on Microsoft's Tay experiment and wider implications, *ACM SIGCAS Computers and Society* 47(3), 54–64.
37. YANG L., ZHANG Z., SONG Y., HONG S., XU R., ZHAO Y. et al., 2023, Diffusion models: A comprehensive survey of methods and applications, *ACM Computing Surveys* 56(4), 1–39.
38. ZHANG P., KAMEL BOULOS M.N., 2023, Generative AI in medicine and healthcare: promises, opportunities and challenges, *Future Internet* 15(9), 286.
39. ZHOU C., LI Q., LI C., YU J., LIU Y., WANG G. et al., 2023, A comprehensive survey on pretrained foundation models: A history from BERT to ChatGPT, *arXiv preprint*, arXiv:2302.09419.
40. ZHOU M., ABHISHEK V., DERDINGER T., KIM J., SRINIVASAN K., 2024, Bias in generative AI, *arXiv preprint*, arXiv:2403.02726.