

Legal Risks And Governance Approaches Of Generative Artificial Intelligence

Qiufang Zhang^{1*}, Ning Huang²

^{1*}School of Law, Intellectual Property Institute, Zhongyuan University of Technology, Zhengzhou, Henan 450007, China, zhqiufang@zut.edu.cn

²School of Law, Intellectual Property Institute, Zhongyuan University of Technology, Zhengzhou, Henan 450007, China

Abstract

As generative artificial intelligence (AI) penetrates deeply into various fields such as content creation, medical diagnosis, and intelligent customer service, its characteristics—such as automated algorithmic decision-making and large-scale data processing—have raised significant legal risks and regulatory challenges. In the domain of intellectual property, ambiguities in copyright attribution for AI-generated content and potential infringements on original work rights through data scraping during training pose multifaceted legal dilemmas. On the front of data security, issues like the leakage of personal sensitive information from training datasets and discriminatory decisions due to algorithmic black boxes challenge existing data protection frameworks. This paper analyzes the contradictions between the technical logic of generative AI and legal regulations, proposing a synergistic governance framework encompassing "legal regulation + technical compliance + industry self-regulation." It advocates for enhancing intellectual property legislation to clarify ownership of AI-generated works and boundaries of data acquisition while strengthening oversight on algorithm transparency. Establishing data compliance auditing systems and ethical review mechanisms for algorithms are essential steps toward achieving a dynamic balance between innovation and risk under the rule of law.

Keywords: Generative AI; Legal Risks; Governance Approaches



1 Introduction

Generative artificial intelligence (AI) exhibits highly "human-like" characteristics such as mimicking human thought processes, adopting human communication habits, and reflecting value preferences, gradually aligning with ideological frameworks. Its deep learning capabilities, logical reasoning abilities, and programmatic innovation skills demonstrate its competence in performing most tasks traditionally executed by humans. Through continuous input of textual materials, self-searching for additional resources, and integrating these materials autonomously, generative AI remains in a state of rapid advancement. By simulating and extending brain activities, it surpasses human capabilities in computational speed, accuracy, and the execution of procedural tasks, all underpinned by superior computing power and algorithms [1]. As of now, models like GPT-4 Turbo have updated their knowledge base up to April 2023, signifying that AI has completed the assimilation of nearly all existing human knowledge. This indicates that AI technology can aggregate the wisdom of predecessors, effectively standing on the shoulders of global civilization[2].

Currently, generative AI technologies represented by Deepseek, ChatGPT, and others are penetrating various sectors of society at an unprecedented rate. From intelligent customer service and artistic creation to assisting medical diagnoses, their formidable content generation and problem-solving capabilities are reshaping human production and lifestyle patterns. However, technological innovations often come hand-in-hand with the reconfiguration of legal relationships and the emergence of new risks. The intellectual property disputes, data security concerns, and ethical responsibilities arising from data processing, algorithmic decision-making, and content output stages in generative AI have become critical bottlenecks hindering its sustainable development. Establishing a legal governance framework that aligns with the technical characteristics of generative AI, balancing innovation incentives with risk management, has thus become a pivotal issue that both academia and practitioners urgently need to address.

In this context, it is imperative to delve into the inherent contradictions between the technical logic of generative AI and existing legal regulations. The aim is to propose a governance approach that integrates "legal regulation + technical compliance + industry self-regulation." Such an integrated framework not only enhances intellectual property legislation to clarify ownership of AI-generated works and delineate boundaries for data acquisition but also strengthens oversight mechanisms for algorithm transparency. Implementing data compliance auditing systems and establishing ethical review procedures for algorithms are crucial steps toward achieving a dynamic equilibrium between technological innovation and risk mitigation within the confines of the rule of law. This paper endeavors to provide insights and recommendations towards this end, fostering a balanced and sustainable development trajectory for generative AI.

2 Legal Risks in the Application of Generative AI

Generative AI often relies on vast amounts of textual data sourced from the internet, raising significant concerns regarding copyright and intellectual property. During the collection and integration of this data, if proper authorization or permission from the original authors is not obtained, it could constitute an infringement of others' copyrights. Such infringements may involve reproduction rights, information network dissemination rights, adaptation rights, and more, leading to legal disputes. Additionally, during the content generation process, there can be further intellectual property issues. For instance, texts, images, or other creative outputs generated by models might provoke controversies over copyright attribution. If such models are employed for commercial purposes, they could also lead to disputes involving trademarks and patents.

Moreover, generative AI poses privacy and data protection risks. The training process typically involves collecting and processing large volumes of user data. Without explicit user consent or adherence to relevant data protection regulations during data handling, it could result in violations of users' privacy rights. Furthermore, if sensitive user information is leaked by the model, severe legal consequences may ensue. Autonomous learning and content generation by generative AI sometimes produce uncontrolled outcomes, generating harmful or misleading information that could threaten social order, public safety, or even na-

tional security, thereby triggering legal liabilities. Therefore, strict compliance with relevant laws and regulations is essential in ensuring lawful data acquisition and usage, safeguarding user privacy and rights, and preventing the generation of harmful or misleading information. Only through these measures can the healthy and sustainable development of generative AI be ensured.

2.1 Intellectual Property Infringement Risks

While generative AI has undoubtedly revolutionized various sectors by bringing unparalleled convenience and fostering innovation, it also brings to the forefront profound and complex questions surrounding intellectual property. These advanced AI models rely heavily on vast amounts of text resources for their training, which is essential for generating accurate, logical, and coherent content. However, this very process of training often involves accessing and utilizing existing works, which can potentially lead to copyright disputes and legal challenges[3]. For instance, consider a scenario where a generative AI model is trained on a large corpus of text that includes copyrighted novels, news articles, or academic papers. If the model incorporates elements from these works without proper authorization, it could be seen as copying or imitating the original content. According to intellectual property law, any unauthorized use of a work by another party is considered an infringement. Therefore, if generative AI uses unauthorized works during its training phase, it may inadvertently violate the intellectual property rights of the original authors. This issue is particularly pronounced and problematic when dealing with highly creative texts, such as literature, news articles, and academic papers, where the line between inspiration and infringement can be extremely thin.

The generated content by these models might closely resemble or even duplicate existing works. If such content is used commercially or disseminated publicly without the original author's consent, it could also constitute an infringement. For example, imagine a situation where a generative AI model produces a news article that is strikingly similar to one published by a well-known journalist. If this AI-generated article is then published on a news website or used for commercial purposes, it could lead to legal action from the original author, who has the exclusive right to control the use and distribution of their work. Moreover, generative AI faces significant controversies related to the "fair use" principle. Training data often comes from the internet or other public sources, which may themselves have intellectual property issues. For example, content from websites or platforms used without the original author's authorization for model training could indirectly lead to intellectual property infringements. In some cases, even if the model uses copyrighted content, it might not constitute infringement if its purpose aligns with fair use principles like commentary, criticism, or news reporting. However, defining the scope of "fair use" and applying this principle in practice remains a topic requiring deeper exploration and clear guidelines. To illustrate, consider a scenario where an AI model uses a small excerpt from a copyrighted book for the purpose of providing a critical review. This might be considered fair use under certain circumstances. However, if the same model uses a substantial portion of the book to generate a new story that closely resembles the original, it would likely be seen as an infringement. The ambiguity lies in determining what constitutes a "small excerpt" and what is considered "substantial,"[4] as well as the context in which the content is used.

2.2 Information Content Security Risks

The deployment and application of generative AI introduce a host of complex information content security risks that challenge existing regulatory frameworks and societal norms. In China, the Cybersecurity Law, specifically Articles 47 and 48, outlines stringent safety requirements for user-posted information, electronic information, and software. For instance, software is prohibited from "containing" certain types of information that may be deemed harmful or illegal. Additionally, the Provisions on the Governance of Network Information Content Ecology explicitly mandate that "network information content producers" must not create, reproduce, or distribute "illegal information" or "negative information." [5] These regulations aim to maintain a healthy and secure online environment. However, implementing these regulations in the context of generative AI presents significant challenges. Unlike traditional content that directly "contains" prohibit-



ed information, generative AI models generate content dynamically based on user queries. This means that the models themselves do not inherently possess illegal or negative information but can produce it in response to specific prompts. For example, a generative AI model might generate a news article that contains factual errors or misleading information, which appears plausible but is factually incorrect. Identifying and removing such content in real time is extremely difficult due to its dynamic nature.

These risks operate subtly and insidiously, leveraging low transparency, high complexity, and automated algorithms to manipulate users' and the public's psychology and behavior. This phenomenon is often referred to as "hypernudge," where AI subtly influences users' decisions and perceptions without their conscious awareness. Over time, individuals may experience gradual "algorithmic harm," which is both efficient and covert. This makes it challenging to detect and address the harm effectively through regulation and legal redress. For instance, a user might unknowingly be influenced by subtly biased information generated by an AI model, leading to misinformed decisions without realizing the source of the problem. Moreover, generative AI can produce unfair, discriminatory, or harmful outputs. If the training data contains biases or discrimination, the model is likely to inherit these issues and manifest them in the generated texts. For example, if a model is trained on data that reflects gender or racial biases, it may generate content that perpetuates these biases. This not only affects the accuracy and reliability of the model but also exacerbates social inequalities. Worse still, the model might be used to generate false information, malicious attacks, or incite hatred, causing severe societal harm. Imagine a scenario where an AI model is used to generate fake news stories that spread misinformation about a particular community, leading to public unrest or discrimination.

2.3 Data Security Risks

Training generative AI is a highly intricate process that involves handling vast amounts of data, encompassing various types of user information such as personal details, chat records, and other forms of digital interactions. These data sources often contain sensitive information that must be carefully managed to protect user privacy. In today's digital age, safeguarding personal privacy is of utmost importance, and ensuring that user data remains confidential during the training of generative AI models is a critical responsibility. For instance, the 2021 incident involving Clearview AI serves as a stark reminder of the potential dangers. Clearview AI, a facial recognition company, trained its AI models using billions of images scraped from social media platforms without users' consent[6]. This data trove included personal photographs, which, when combined with other publicly available information, could be used to identify individuals, track their movements, and even predict their behavior. The company's actions led to investigations by multiple regulatory bodies, including the Federal Trade Commission (FTC) in the United States and data protection authorities in Europe. Such cases underscore how improper handling of data during AI training can lead to widespread privacy violations and legal consequences.

Another notable example is the Cambridge Analytica scandal[7]. The firm harvested the personal data of millions of Facebook users without their knowledge. This data was then used to train AI-driven models for targeted political advertising. The leaked data included personal information such as names, locations, and browsing histories, which were exploited to manipulate public opinion during elections. As a result, over 87 million users' data was compromised, causing a major global outcry and leading to significant reforms in data protection regulations, such as the General Data Protection Regulation (GDPR) in the European Union[8]. If this data is not properly anonymized or secured, it could be exposed to unauthorized parties. The consequences of such mishandling can be severe. Users' financial security could be compromised if their bank account details or credit card information are leaked. Similarly, identity theft could occur if personal identification numbers or social security numbers are exposed. This could lead to significant financial losses, legal issues, and long-term damage to an individual's credit and reputation.

Moreover, the leakage of sensitive information like chat records can have a profound impact on users' daily lives. In 2023, a popular messaging app faced a data breach that exposed private conversations of

thousands of users. Some of these conversations contained sensitive topics such as job search strategies, medical consultations, and personal disputes. The leaked data was subsequently used for blackmail and harassment, causing emotional distress and social disruption among the affected users. Such incidents highlight that protecting user privacy is not just a moral obligation but a legal necessity for those involved in training generative AI.

Furthermore, the inference processes of generative AI models also pose potential risks to user privacy. These models often analyze user inputs to infer personal characteristics, intentions, or preferences, which can enhance their intelligence and functionality. However, this capability can also lead to privacy infringement. A study by researchers at Carnegie Mellon University demonstrated that AI language models like GPT-3 could accurately infer users' sexual orientation[9], political views, and even mental health conditions based on their writing samples. This information could be misused for discriminatory purposes, such as denying housing, insurance, or employment opportunities. In addition, the potential for misuse of inferred information is a growing concern. For example, in 2022, a travel planning AI service was found to be sharing inferred location data of its users with third-party advertisers without explicit consent[10]. The service used users' chat histories and travel queries to deduce their current and future travel plans, which were then sold to marketers for targeted advertising. This practice not only violated users' privacy but also raised ethical questions about the responsible use of AI-generated inferences.

3 Challenges in Legal Regulation of Generative AI

Currently, existing laws often fall short when addressing new technologies like generative AI, primarily due to limitations in their applicability. This inadequacy results in difficulties in effectively regulating risky behaviors despite some countries and regions having introduced relevant regulations. However, practical enforcement faces challenges due to regulatory limitations and technological constraints, often failing to achieve desired outcomes. Regulatory measures for generative AI mainly encompass the establishment of a relevant legal framework, clarification of corporate and individual responsibilities, and the formulation of regulatory policies. Nonetheless, these measures encounter significant difficulties and challenges in implementation, such as the lagging nature of laws and inadequate regulatory tools. These challenges stem from the rapid pace of technological innovation, the delayed response of legal systems, and the diverse application scenarios.

3.1 The Conflict Between Rapid Technological Advancements and Legal Lag

The iterative speed of generative AI technology is advancing rapidly, whereas the formulation and revision of laws struggle to keep pace. This mismatch renders laws ineffective in addressing new technological risks. Current governance frameworks include many national recommended standards that lack mandatory enforceability, limiting regulatory agencies' ability to effectively constrain relevant entities[11]. Even with the introduction of relevant laws and regulations, they may fail to adapt promptly to new technological characteristics, making it difficult to address the risks posed by generative AI applications. For instance, China's "Regulations on the Management of Algorithmic Recommendation Services" issued in 2022 mandates that algorithms should adhere to mainstream value orientations, actively disseminate positive energy, and not engage in activities harmful to national security, social public interests, economic order, or infringement of others' rights. However, given that models like ChatGPT remain "black box" systems, whether their algorithms can be disclosed adequately and comply with regulatory requirements remains uncertain. The rapid development of AI technology complicates the formation of unified standards and mechanisms for regulation. Legal bodies need to understand AI principles and applications thoroughly to craft effective laws, yet high technical barriers hinder communication and collaboration between the legal and technical communities.



3.2 Liability Determination and Accountability Challenges

The complexity and novelty of AI technology introduce numerous unresolved legal issues and controversies. Existing laws often cannot directly apply to illegal content generated by AI, as traditional legal rules and precedents primarily focus on natural persons or corporations as responsible parties. Determining the legal liability of AI as a technological tool remains contentious, especially concerning copyright attribution, privacy protection, and liability determination in the context of ChatGPT-generated content. For example, a company using ChatGPT to develop an intelligent customer service system encountered issues where the system generated misleading or illegal content, including false advertising, privacy violations, and trade secret leaks. Due to the complexity and diversity of these problems, users received unclear responses and ineffective solutions when lodging complaints. Overlapping legal boundaries across different fields further complicate legal clarity, leaving current legal frameworks unable to provide clear and feasible resolutions. This uncertainty increases regulatory difficulties and potentially threatens the rights and interests of all parties involved. OpenAI's ChatGPT faced lawsuits from writers, comedians, and other professionals who claimed unauthorized use of their works, infringing upon their copyrights[12]. While most cases are still pending, preliminary court decisions suggest mixed outcomes, reflecting the complexity of copyright issues associated with ChatGPT. Cross-border data flows and jurisdictional conflicts further highlight the challenges arising from ambiguous legal boundaries. Handling data through ChatGPT may involve cross-border transfers, subject to varying regulations across countries regarding data privacy and information security. Differences in legal systems and regulatory policies among nations create uncertainties and risks for multinational companies applying AI technologies, leading to ongoing legal and compliance challenges.

4 Governance Pathways for Legal Risks of Generative AI

The governance pathways for legal risks associated with generative AI constitute a multi-dimensional and comprehensive process, involving technological, legal, ethical, and societal aspects. It requires concerted efforts from governments, enterprises, society, and individuals to establish a safe, lawful, and healthy environment for the application of generative AI through measures such as formulating laws and regulations, promoting industry self-regulation, enhancing public legal awareness, and strengthening international cooperation.

4.1 Clarification of Legal Responsibilities

Clarifying the roles and responsibilities of developers, operators, and users in legal liability issues is a crucial step to ensure the healthy and orderly development of generative AI. To define the rights and obligations of each party, relevant laws and regulations must be established to regulate the activities of developers, operators, and users. These laws should specify the attribution of responsibility and accountability mechanisms, ensuring that responsible parties can be accurately identified and held accountable when problems arise[13]. Additionally, laws and regulations should aim to balance various interests to promote the healthy development of generative AI technology.

a) Developers: As creators and designers of generative AI models, developers hold control over the technical architecture and algorithm design. They bear primary responsibility for ensuring the safety, stability, and reliability of the model. Developers must comply with relevant laws and regulations during the development process, avoid infringing on others' rights, and actively take measures to prevent potential risks and issues.

b) Operators: Operators are responsible for the operation and maintenance of generative AI, including providing model usage platforms, ensuring data security, and protecting user privacy. They play a critical role in maintaining the normal operation of the model and safeguarding user rights. Operators need to ensure platform security and stability, prevent data leaks and misuse, and promptly handle user feedback and complaints to protect user rights and interests.

c) Users: Users, as the actual applicators of generative AI, also bear certain legal responsibilities. They must abide by relevant laws, regulations, and platform rules, use models reasonably, and refrain from engaging in illegal or rights-infringing activities. Users should remain vigilant about potential risks and issues during model usage and report them to the platform promptly.

By clarifying the roles and responsibilities of all parties in legal liability issues, we can provide strong legal guarantees and norms for the development and application of generative AI, promoting its healthy and orderly development and better serving society and the people.

4.2 Measures to Address Specific Legal Challenges

4.2.1 Intellectual Property Infringement Risk Management

To mitigate intellectual property infringement risks associated with generative AI, multiple strategies must be employed, including legally acquiring training data, enhancing copyright awareness education, using open-source or authorized datasets, establishing copyright review mechanisms, and collaborating with rights holders.

a) Legal Acquisition of Training Data: Ensure that training data sources are legal and avoid obtaining data from unauthorized channels. Establish partnerships with data providers or use publicly available, legitimate datasets. Respect authors' rights and obtain necessary permissions or licenses when collecting data. Clean and filter collected data to remove potentially copyrighted content. Utilize natural language processing techniques to automatically screen and filter text data to ensure training data does not contain infringing material.

b) Enhancing Intellectual Property Education: Educate developers and users about intellectual property laws, making them aware of regulations they must follow during model development and usage. Through training and education, enhance their ability to identify and avoid intellectual property infringement risks. Prioritize using open-source or already licensed datasets for model training to reduce infringement risks.

c) Establishing Copyright Review Mechanisms: Before generating content, establish copyright review mechanisms to check generated content for copyright infringement. Utilize existing copyright detection tools or algorithms to compare and analyze generated content, ensuring compliance with copyright laws[14]. Collaborate with rights holders to seek authorization for the legitimate use of their works, fostering harmonious development between copyright protection and AI technology.

4.2.2 Information Content Security Risk Management

a) Strengthening Data Quality Management and Supervision: In the training process of generative AI, data quality directly affects the accuracy and reliability of model outputs. Therefore, enhancing data quality management and supervision is essential for mitigating information content security risks. Strictly screen training data, ensure reliable data sources, choose authoritative and credible datasets, preprocess data to remove duplicates, invalid, or low-quality data, and annotate and categorize data for subsequent supervised learning and model training. Verify data authenticity to ensure real and trustworthy training data, avoiding false data leading to model output errors. Review data legality to ensure compliance with laws, regulations, and ethical standards. Pursue data representativeness to cover a wide range of scenarios, enhancing model generalization. Implement supervised learning mechanisms using labeled data for model training and parameter adjustment; monitor model outputs in real-time, manually intervene and correct anomalies. Continuously optimize models through feedback mechanisms to improve harmful content identification and filtering capabilities.

b) Improving Content Review and Management Mechanisms: As generative AI applications expand across various domains, ensuring the safety and compliance of generated content becomes critical. Introduce automated content review technologies using natural language processing and machine learning to



develop automated content review systems, enabling rapid review and filtering of generative AI-generated information. Set corresponding review rules and thresholds based on application scenarios and regulatory requirements to accurately identify and filter out harmful information[15]. Timely detect and remove harmful information: Monitor generative AI outputs in real-time, promptly clear and handle harmful information. Regularly audit and assess generated content to ensure compliance with relevant laws and ethical standards; adjust and optimize models based on assessment results to improve harmful content identification and filtering capabilities. Establish cross-departmental and cross-industry collaborative regulatory mechanisms, strengthen cooperation and exchanges among governments, enterprises, and industry organizations; jointly formulate and refine relevant laws and standards to regulate the application and management of generative AI; enhance supervision and evaluation of generative AI technology applications to ensure alignment with social public interest and national security requirements.

4.2.3 Data Security and Personal Information Protection Risk Management

Given the extensive handling and use of data and personal privacy information by generative AI, strict data protection and privacy security mechanisms must be established.

a) Data Protection Policies: Develop rules for data collection, storage, use, and sharing, clearly defining the purpose, scope, and methods of data collection, ensuring only necessary data related to model applications is gathered and avoiding excessive collection. Employ encryption and other security measures to securely store collected data, preventing leaks or unauthorized access. Limit data access to authorized personnel and strictly control data flow to prevent misuse, ensuring user data security and privacy. Strengthen technical safeguards by establishing security auditing and monitoring systems to continuously monitor and record data processing and usage, promptly detecting and addressing potential security risks. Formulate legal liabilities and penalties for violations of data protection and privacy security regulations, enforcing strict consequences for offenders. Enhance education and awareness to increase user understanding and emphasis on data protection and privacy security, fostering a societal commitment to maintaining data security and privacy.

b) Model Ethics Oversight and Review Mechanisms: Establishing effective ethics oversight and review mechanisms is essential for ensuring the compliance and responsible use of generative AI. First of all, defining regulatory objectives and principles: Clearly define regulatory objectives such as protecting user privacy, preventing discrimination and unfairness, and ensuring content legality. Establish principles like fairness, transparency, and accountability to guide the regulatory process. Develop detailed ethical guidelines covering model training, deployment, and usage, specifying prohibited behaviors and required norms.

Secondly, establishing review institutions and conducting model evaluations: Create specialized institutions with independence, professionalism, and authority to conduct objective and fair reviews.

Pre-launch Review: Conduct thorough governance reviews before model deployment, verifying the legality and ethical compliance of training data, checking for biases or unfairness, and assessing potential risks.

Thirdly, ongoing assessment: Continuously monitor and evaluate deployed models, collecting user feedback and analyzing outputs to ensure performance, stability, and compliance. Regularly test models to identify errors or biases, providing improvement suggestions to developers and users. Focus on user privacy protection and data security to ensure models do not infringe on user rights during usage.

Fourthly, emergency response: Implement rapid response mechanisms to address serious ethical issues, including suspending model usage, emergency repairs, or retraining. Actively collaborate with developers and users to solve problems, ensuring model compliance and safety.

And at last, strengthening user education and feedback mechanisms: Enhance user understanding of generative AI ethics through education and training, encouraging adherence to ethical guidelines and policies. Establish feedback and complaint channels, facilitating user supervision and feedback. Regularly publish user feedback reports to showcase handling processes and improvements, enhancing trust and cooperation between users and institutions.

In conclusion, establishing an effective model ethics oversight and review mechanism requires comprehensive efforts and collaboration. By defining clear regulatory goals and principles, developing governance guidelines, establishing independent review bodies, implementing thorough model evaluations, strengthening user education, and setting up robust feedback mechanisms, we can continuously improve and develop governance and review practices for generative AI. This multi-faceted approach ensures the responsible and compliant use of generative AI technologies, promoting their healthy and sustainable development.

5 Conclusion

The vigorous development of generative AI, coupled with its associated legal risks, poses unprecedented challenges to the social governance system. In the face of this technological wave, only by constructing a comprehensive, multi-dimensional governance system can we achieve a dynamic balance between technological innovation and risk prevention.

At the data governance level, regulating the sources of training data, establishing copyright review mechanisms, and implementing collaborative regulatory systems can effectively cut off the transmission chain of infringement risks. Legal data acquisition pathways and stringent content screening not only protect the rights of original rights holders but also lay the foundation for the legality of model outputs. Cross-departmental and cross-industry collaborative supervision further breaks down data silos, enabling full-process tracking and intervention of risks.

The refinement of ethical oversight mechanisms injects a humanistic dimension into technological development. By clearly defining ethical objectives and governance guidelines, model development can avoid discriminatory or illegal content from the outset. The routine evaluations and transparent reporting of independent review institutions promote the formation of predictable behavioral norms within the industry. Establishing user education and feedback mechanisms transforms passive regulation into proactive governance involving the entire population, thereby building a societal defense against ethical lapses in technology.

The improvement and deepening of international cooperation within the legal system are core supports for ensuring the sustainable development of technology. Clearly delineating the responsibilities of all parties not only compels developers to fulfill compliance obligations but also provides clear judicial criteria for legal practice. Transnational legislative collaboration and standardization help eliminate regulatory arbitrage opportunities, fostering global coordinated governance efforts.

Funding

- [1]Social Science Planning Project of Henan Province (Project No.: 2023BFX027)
- [2]Soft Science Research Program of the Department of Science and Technology of Henan Province (Project No.: 252400411190)
- [3]Intellectual Property Soft Science Research of Henan Province (Project No.: 20250106021)
- [4]Basic Business Fee Project of Zhongyuan University of Technology (Project No.: K2024JJ004)
- [5]Advantage Discipline Enhancement Plan Funded Project of Zhongyuan University of Technology (Project No.: GG202401).

References

- Zhang Linghan. Legal Positioning and Hierarchical Governance of Generative AI.”Modern Law, 2023, 45(04): 126-141.
- Liu Yanhong. Three Major Security Risks and Legal Regulation of Generative AI: A Case Study of ChatGPT. Oriental Law, 2023(04): 29-43.
- Zhang Xin. Algorithmic Governance Challenges and Regulatory Oversight of Generative AI. Modern Law, 2023, 45(03): 108-123.



- Han Weihong. Development and Application of Modern Electronic Information Technology in the Context of ‘Internet+’. *China Equipment Engineering*, 2023(17): 244-246.
- Wang Youmei, Wang Dan, Liang Wei Yi, et al. “Ethical Risks and Avoidance Strategies of ChatGPT in Education.” *Open Education Research*, 2023, 29(02): 26-35.
- Clearview AI Controversy: Data Scraping, GDPR Fines & Ethical Dilemmas | Cruptodamus <https://cryptodamus.io/en/articles/news/clearview-ai-shakeup-scandals-resignations-the-future-of-face-recognition>
- Case Study: The Facebook-Cambridge Analytica Scandal-Cyber Security - AbiEdu <https://abiedu.com/case-study-the-facebook-cambridge-analytica-scandal-cyber-security/>
- [8]Legal framework of EU data protection - European Commission https://commission.europa.eu/law/law-topic/data-protection/legal-framework-eu-data-protection_en.
- Rosenberg M. The Myth of Algorithmic Accountability. *Harvard Law Review*, 2023, 136(4): 947-986.
- Liang Jianchun. Ethical Mechanisms and Guidance Strategies for Data Governance in University Libraries in the Era of Digital Intelligence. *Library and Information Guide*, 2023, 8(07): 1-6.
- Zhao Yue, He Jinwen, Zhu Shenchen. Current Status and Challenges of Generative AI Safety. *Computer Science*, 2024, 51(01): 68-71.
- Deng Jianpeng, Zhu Yicheng. Legal Risks and Countermeasures of the ChatGPT Model. *Journal of Xinjiang Normal University (Philosophical and Social Sciences Edition)*, 2023, 44(05): 91-101+2.
- Zhong Xiangming, Fang Xingdong, Gu Yeye. Governance Challenges and Countermeasures of ChatGPT: The ‘Collingridge Dilemma’ in Intelligent Communication and Breakthrough Paths. *Media Observation*, 2023(03): 25-35.
- Su Yu. Legal Risks and Governance Pathways of Generative AI. *Law Science (Journal of Northwest University of Political Science and Law)*, 2024, 42(01): 76-88.
- Chen Tianyou. Risk Governance of False Information Generated by Generative AI: A Case Study of ChatGPT. *Cultural Journal*, 2024(04): 100-103.