

From Human Authorship to Machine Creation: Reassessing Copyright and Liability in the Age of Generative AI

Kunsang Wangmo

Abstract

The rapid expansion of generative artificial intelligence has transformed digital creativity, enabling machines to autonomously produce text, images, audio, and video at a scale previously unattainable. This development has intensified global legal debates surrounding authorship, ownership, originality, and the lawful use of training data. Traditional copyright frameworks, built on the assumption of human creativity, struggle to accommodate content generated by complex machine-learning models. This paper examines key legal challenges arising from generative AI, including the human authorship requirement, the status of training data under fair use and data protection laws, the risks of derivative similarity, and the growing concerns around liability, deepfakes, and platform responsibility. It also analyzes regulatory divergence across jurisdictions such as the United States, European Union, China, and India. By synthesizing emerging scholarship and legal developments, the study highlights the need for modernized copyright approaches, transparency obligations, and hybrid regulatory reforms to address the evolving landscape of machine-generated content.

Chapter 1: Introduction

The rapid development of generative artificial intelligence (AI) has transformed how digital content is created, distributed, and consumed. Modern AI systems are capable of autonomously producing text, images, audio, video, and code with a level of sophistication that increasingly resembles human output. These systems, powered by large-scale machine-learning models, are now integrated into journalism, entertainment, marketing, design, and various creative industries. As a result, AI-generated content has become both commercially significant and legally contentious.

The expansion of generative AI has brought a parallel rise in legal disputes regarding authorship, ownership, and the use of copyrighted materials for training machine-learning systems. Scholars note a sharp increase in copyright-related claims involving AI models, particularly concerning how training datasets are compiled, whether copyrighted works are being reused without permission, and whether AI outputs infringe on existing protected material (Yang & Zhang, 2024). Moreover, regulators and legal institutions are facing pressure to update or reinterpret existing laws as AI tools become widely adopted across sectors including education, communication, media, and healthcare (Shah, 2024). This evolving technological landscape has exposed gaps and inconsistencies in intellectual property frameworks that were developed long before autonomous content creation existed.

Traditional copyright law is built on the assumption that creative works originate from human authors who exercise skill, judgment, and originality. This foundational premise is increasingly challenged by AI systems capable of producing content without direct human creativity or intention. Legal analyses highlight that existing copyright doctrines are not designed to accommodate autonomous and semi-

autonomous outputs generated by machine-learning systems (Cambridge University Press, 2023). As generative AI becomes more independent from human involvement, courts and policymakers must confront the question of whether such outputs can be copyrighted at all, and if so, by whom.

This problem is further complicated by the opacity of AI training processes, the blending of human and machine contributions, and the uncertainty regarding the legal status of training data. These unresolved issues reveal that current copyright structures may be insufficient for governing the intellectual property implications of generative AI.

As generative AI becomes increasingly embedded in creative and commercial workflows, the need for clear legal frameworks has become urgent. This study examines key gaps in current legal systems, including authorship, originality, derivative works, liability, and data protection. These gaps have been widely noted in emerging legal scholarship, which emphasizes that traditional intellectual property doctrines do not adequately address the complexities introduced by AI-generated content (Zakir, 2024). By analyzing these legal challenges, the study contributes to ongoing policy discussions regarding the future of copyright law, the ethical and economic implications of AI-generated works, and the development of regulatory strategies tailored to machine creativity. The findings aim to support scholars, policymakers, and practitioners in understanding where legal reform is needed and how current doctrines may be adapted to meet the demands of the AI era.

Together, these questions highlight the complexity and urgency of re-examining intellectual property law in the context of rapidly advancing generative AI technologies. As disputes over authorship, training data, originality, and liability continue to expand across jurisdictions, it becomes clear that the legal system must adapt to address forms of creativity and content production that were never anticipated by traditional copyright frameworks. This study therefore begins by establishing a foundational understanding of how AI-generated content is created, classified, and distinguished from human-authored works, before moving into a deeper analysis of the legal challenges that arise. The following chapter examines the technical and conceptual foundations of AI-generated content, laying the groundwork necessary to evaluate its implications within modern legal structures.

Chapter 2: Foundations of AI-Generated Content

AI-generated content refers to any output; text, images, audio, video, code, or design; created using machine-learning systems rather than produced solely through direct human creativity. In many cases, AI tools operate in a *collaborative* manner, where humans provide prompts, edits, or conceptual direction while the model generates the actual output. This form of human–AI co-creation complicates traditional legal understandings of authorship, because the degree of human involvement can vary widely. At the other end of the spectrum are *fully autonomous* outputs, where the AI produces content without human intervention beyond initial model development or technical configuration. This distinction is meaningful in legal analysis, as courts and policymakers often rely on the extent of human creative contribution to determine authorship, originality, and ownership (Shah, 2024). Understanding these categories is essential because they shape how the law interprets creative agency in the context of emerging technologies.

2.1 Technical Basis

Generative models rely on large datasets that contain text, images, audio recordings, videos, or other digital materials. These datasets frequently include publicly accessible content scraped from the internet as well as proprietary or personal data collected from various sources. Through training, models learn

statistical relationships and patterns within these datasets, enabling them to generate outputs that resemble human-created material.

The training process typically involves two phases: (1) learning from massive datasets through supervised or unsupervised methods, and (2) refining the model using techniques such as reinforcement learning or fine-tuning. Because models internalize patterns rather than store exact copies, their outputs may be novel, derivative, or occasionally similar to the original training inputs. This technical complexity is central to understanding the legal issues surrounding originality, derivative works, privacy, and data usage (UPF Thesis, 2023). As generative AI tools become more advanced, their outputs increasingly mimic human creativity, prompting deeper reflection on how the law should treat machine-produced material.

2.2 Why Technical Processes Matter Legally

The way generative AI systems learn and produce content has direct legal implications. Copyright law evaluates creativity, authorship, and originality based on human intention and effort. When a model produces content based on patterns learned from millions of existing works, determining whether the output is sufficiently original becomes challenging. Some outputs may blend or closely resemble elements found in the training data, raising questions about potential infringement, even when no exact copying occurs.

These technical processes also affect assessments of liability. If an AI system generates harmful, misleading, or unlawful content, understanding how the output was produced influences whether responsibility lies with the developer, the platform, or the end-user. Furthermore, the legality of training data itself is a central topic of debate. If copyrighted works were included in datasets without permission, the training process may raise concerns under copyright, fair use, and data protection laws (Cambridge University Press, 2023).

In short, the internal mechanisms of generative AI; how data is collected, processed, and transformed; are foundational to evaluating authorship, ownership, infringement, and accountability within modern legal frameworks.

Chapter 3: Copyright Authorship, Ownership & Originality

3.1 Human Authorship Requirement

Copyright law in most jurisdictions is built around the idea that creative works originate from human authors who exercise intention, judgment, and originality. This principle remains central to determining whether a work qualifies for legal protection. Courts and regulatory bodies continue to emphasize that copyright exists to protect human expression, not machine-generated output. The Cambridge copyright analysis stresses that originality requires a human mind to make creative choices, meaning that works produced entirely by autonomous AI systems fall outside the scope of traditional copyright protection (Cambridge University Press, 2023).

Reflecting this stance, the U.S. Copyright Office has repeatedly rejected attempts to register works created without meaningful human contribution. These decisions reinforce the position that copyright cannot be assigned to an AI system and that human authorship is a foundational requirement. As generative AI tools become more sophisticated, this requirement has become one of the most pressing legal questions surrounding AI-generated content.

3.2 Authorship Controversies

The uncertainty surrounding whether AI-generated works can be copyrighted has led to a range of proposals about how authorship should be understood in the age of generative AI. One view argues that

the human user who provides the prompt or conceptual direction should be recognized as the author. Another perspective supports *joint authorship* when both humans and AI meaningfully contribute to the final output. Some scholars suggest corporate authorship, especially when companies control the development and operation of AI systems that generate commercially valuable outputs.

A more provocative position is the idea of *AI authorship*, which would involve recognizing the AI itself as the creator. However, this proposal is widely contested because it conflicts with legal principles that require authors to have rights and responsibilities; qualities AI systems do not possess. These diverse viewpoints illustrate the growing disagreement within legal scholarship, reflecting a need for clearer frameworks as AI technologies continue to evolve (Zakir, 2024).

3.3 Sui Generis Rights

Because existing copyright structures struggle to accommodate AI-generated outputs, several scholars advocate for the creation of *sui generis* rights; legal protections tailored specifically to works produced by autonomous systems. These proposals suggest that instead of attempting to force AI-generated content into traditional copyright categories, policymakers should consider establishing separate rights that acknowledge the unique nature of machine-produced creativity.

Suggested frameworks include granting limited rights to the developer, allocating rights to the user who initiates the AI's output, or creating a hybrid model that balances both roles. The UPF thesis (2023) and recent analyses by Massadeh (2024) highlight that *sui generis* models may offer a clearer and more adaptable approach for regulating AI-generated content, particularly as the technology continues to accelerate. Such models could help resolve ownership disputes, clarify liability, and better align legal protections with the realities of automated content creation.

3.4 Originality and Derivative Work Issues

Determining whether AI-generated content meets the legal standard of originality is one of the most complex challenges in this field. Generative AI systems learn by identifying patterns within vast datasets, and their outputs often reflect combinations, transformations, or stylistic elements drawn from those training materials. This process has raised concerns that AI-generated outputs may implicitly remix or replicate existing copyrighted works, even when the resemblance is unintentional (Shah, 2024).

These concerns are particularly significant when evaluating *substantial similarity*, a key test used by courts to determine infringement. Scholars have noted that some AI-generated artworks exhibit similarities to specific training images, prompting infringement claims and legal disputes (Yang & Zhang, 2024). As courts confront more cases involving AI-generated content, they are increasingly considering technical evidence such as similarity metrics, dataset analysis, and machine-learning forensics to assess whether a work is genuinely original or impermissibly derivative (Jones & Kearns, 2024).

These issues underscore the tension between how AI models generate content and how copyright law conceptualizes creativity, making originality a central legal question in the regulation of generative AI.

Chapter 4: Training Data, Fair Use & Data Protection

4.1 Legality of Training Data

The question of whether AI developers can legally use copyrighted material for training machine-learning models has become one of the most debated issues in the field of generative AI. Most large-scale models rely on datasets collected from publicly accessible websites, digital archives, and social media platforms. These datasets often include copyrighted works that were scrapped without explicit permission from rights holders. This practice has generated significant legal concern, as it challenges long-standing principles

surrounding the control and licensing of creative works (Shah, 2024). The legality of using such material remains unsettled, with courts and lawmakers examining whether large-scale data collection for AI development falls within permissible use or constitutes infringement.

4.2 Fair Use and Transformative Use

A major element of this debate focuses on whether the use of copyrighted works for training qualifies as “fair use” under U.S. copyright law. Supporters of AI training practices argue that using copyrighted content to teach a model patterns, structure, or language is sufficiently transformative because the output is not a direct copy of any specific work. Critics counter that the scale of copying involved in scraping datasets goes beyond any reasonable interpretation of fair use. Academic analysis suggests that existing fair use doctrine was not designed to address automated, mass-scale ingestion of creative materials, making the application of this standard increasingly complex (Cambridge University Press, 2023).

Internationally, the regulatory picture is even more complex. The European Union’s introduction of text-and-data mining (TDM) exceptions allows certain forms of data scraping for research and innovation, provided that rights holders have not opted out. This creates a regulatory gap between the U.S. and the EU, further complicating how global AI developers navigate training-data compliance.

4.3 Transparency Challenges

A recurring challenge in evaluating the legality of AI training practices is the lack of transparency surrounding dataset composition. Most developers of large generative models do not fully disclose the sources, licensing status, or structure of the data used during training. This opacity makes it difficult for regulators, rights holders, or researchers to determine whether copyrighted, proprietary, or sensitive materials have been included in the dataset (UPF Thesis, 2023). Limited transparency also affects the ability to assess potential infringement or to create effective auditing mechanisms. Without clear information about training inputs, enforcing copyright laws becomes significantly more challenging, and accountability for misuse remains diffuse.

4.4 Privacy and Biometric Risks

Beyond copyright concerns, the inclusion of personal or biometric data in AI training sets raises substantial privacy risks. Machine-learning systems trained on such data may inadvertently reproduce sensitive details, including faces, voices, or textual personal information, as part of their outputs. These risks highlight the blurred boundary between public data availability and lawful processing for AI development (Jones & Kearns, 2024).

Regulatory frameworks such as the European Union’s General Data Protection Regulation (GDPR) impose strict requirements for consent, data minimization, and lawful processing of personal information. Similarly, India’s Digital Personal Data Protection (DPDP) Act requires clear justification and authorization for the use of personal data in automated systems (Zakir, 2024). Violations of these principles may expose developers to legal liability, even when the data was collected from publicly accessible sources. As generative AI advances, privacy-related risks are expected to grow, making robust data governance essential to responsible model development.

The legality of training data lies at the center of the broader debate surrounding AI-generated content. As this chapter shows, existing copyright and data protection frameworks were not designed for technologies that learn from vast and often unregulated datasets. Questions of fair use, transformative analysis, and international regulatory differences create significant uncertainty for both developers and rights holders. At the same time, the lack of transparency around training datasets and the potential inclusion of personal or biometric information raise serious concerns about compliance with privacy laws and ethical data

practices. These intertwined challenges demonstrate that training-data governance is not merely a technical issue but a foundational legal question. Understanding these complexities is essential for evaluating how liability, authorship, and ownership are assigned in the chapters that follow.

Chapter 5: Liability, Deepfakes & Platform Responsibility

5.1 Attribution of Liability

Assigning legal responsibility for AI-generated content remains a complex and unsettled area of law. Unlike traditional creative tools, generative AI systems operate autonomously, with multiple parties contributing to the creation process. Developers design and train the models, deployers integrate them into products or services, and end-users generate the final outputs. This multi-layered structure complicates determinations of fault when AI-generated content results in harm, infringement, or misinformation. Current scholarship highlights the difficulty of identifying a single point of accountability within this chain of actors, especially when the harm arises from an output that no individual directly authored (Zakir, 2024). As a result, courts and policymakers are grappling with whether traditional liability frameworks can adequately address the distributed nature of AI development and use.

5.2 The Black-Box Problem

A major obstacle in establishing liability is the technical opacity of generative AI models. Modern systems rely on neural networks with millions or billions of parameters, producing outputs through processes that are not easily interpretable. This “black-box” nature makes it challenging for courts, regulators, or even developers to trace how a specific output was generated or which internal mechanisms led to harmful or infringing content. As Shah (2024) notes, without clear insight into a model’s decision-making process, it becomes extremely difficult to assess whether negligence, design flaws, user misuse, or unforeseeable algorithmic behavior is responsible for the resulting harm. This lack of transparency limits the effectiveness of traditional legal standards that rely on identifiable causes and clear chains of reasoning.

5.3 Deepfakes and Harm

Deepfakes represent one of the most direct and visible forms of AI-enabled harm. These synthetic videos, images, or audio clips can convincingly mimic real individuals, often without their consent. Their misuse has been associated with identity theft, political misinformation, non-consensual sexual content, and reputational damage (Massadeh, 2024). Beyond immediate harms, deepfakes raise broader concerns regarding integrity and image rights, as individuals may lose control over how their likeness is used or manipulated (Shah, 2024).

The rapid spread of deepfake content online highlights major gaps in existing legal frameworks. Many jurisdictions lack specific legislation addressing synthetic media, forcing courts to rely on older laws concerning defamation, impersonation, fraud, or data protection; each of which may only partially address the problem. As generative AI improves in realism and accessibility, concerns about the social and legal impact of deepfakes continue to grow.

5.4 Platform Responsibilities

Given the scale at which AI-generated content can be produced and shared, online platforms play a critical role in preventing and mitigating harm. Platforms are increasingly expected to implement systems that help identify, monitor, and label AI-generated material. Proposed solutions include watermarking, metadata tagging, algorithmic detection tools, and content authentication frameworks (Yang & Zhang, 2024). These technical safeguards aim to improve transparency and help users distinguish between genuine and synthetic content, thereby reducing the potential for misuse.

Beyond technical measures, platforms may also bear policy and compliance responsibilities. This includes developing community guidelines, establishing reporting mechanisms, and cooperating with regulators to address emerging risks. As discussions around liability evolve, platforms are likely to face heightened scrutiny regarding their role in the distribution of AI-generated content and their obligation to prevent foreseeable harm.

Chapter 6: Global Regulatory Divergence and Future Reforms

6.1 Cross-Jurisdiction Differences

The legal treatment of AI-generated content varies significantly across jurisdictions, reflecting differences in copyright philosophies, data protection standards, and broader regulatory priorities. These divergences create a fragmented global landscape that complicates compliance for developers, platforms, and end-users of generative AI technologies.

In the United States, copyright law continues to emphasize human authorship as a strict requirement for protection. Courts and administrative bodies have consistently held that works lacking meaningful human creative input cannot be copyrighted, a position reaffirmed in recent high-profile decisions involving AI-generated content (Cambridge University Press, 2023). Meanwhile, debates surrounding fair use and transformative use remain active, particularly as courts evaluate whether the large-scale ingestion of copyrighted materials for training constitutes lawful use.

The European Union takes a distinct approach grounded in strong moral rights traditions and a robust regulatory framework. The EU's text-and-data mining (TDM) exceptions allow certain forms of data scraping for research and innovation while permitting rights holders to opt out. The forthcoming EU *AI Act* further demonstrates the region's willingness to proactively regulate AI based on risk categories, creating one of the most comprehensive AI governance regimes globally.

In China, emerging AI regulations reflect a state-led approach that prioritizes government oversight and social stability. Guidelines emphasize transparency, dataset accuracy, and alignment with state-defined ethical and political principles (Massadeh, 2024). These measures highlight China's preference for a highly managed regulatory environment with strict controls over AI-generated content.

In India, the legal framework remains underdeveloped. The Copyright Act of 1957 does not address machine-generated works, leaving authorship and ownership unclear. While the Digital Personal Data Protection (DPDP) Act introduces stronger privacy protections, it does not directly regulate AI training or generative outputs (Zakir, 2024). As a result, India relies heavily on judicial interpretation and general legal principles.

6.2 Why Harmonization Is Difficult

Achieving global harmonization in AI regulation is challenging for several reasons. First, countries differ fundamentally in how they conceptualize copyright, privacy, and digital rights. The United States emphasizes economic incentives and flexible doctrines such as fair use, whereas the European Union prioritizes individual rights, moral authorship, and strict privacy protections. China's regulatory philosophy is shaped by centralized governance and content oversight, while India's legal system has yet to articulate a comprehensive stance on AI.

Second, AI governance intersects with national interests, economic strategies, and cultural values. As countries aim to position themselves as leaders in AI innovation, their regulatory approaches reflect competitive priorities, making coordinated global standards difficult to achieve.

Third, the rapid pace of technological change makes it challenging for any global body to develop standards that remain relevant over time. By the time consensus is built, underlying technologies and practices may have already shifted. These structural and philosophical differences contribute to a fragmented international regulatory landscape.

6.3 Proposed Reforms

Given these challenges, scholars and policymakers have suggested several reform pathways to address the legal uncertainties surrounding AI-generated content. One proposal involves establishing sui generis rights tailored specifically to machine-generated works. Such frameworks could offer limited, clearly defined protections without forcing AI outputs into traditional copyright categories (UPF Thesis, 2023; Massadeh, 2024).

Another proposed reform centers on licensing and transparency. Policymakers are considering requirements for AI developers to disclose training datasets, license copyrighted materials appropriately, and adopt data governance standards to ensure lawful and ethical use of inputs. Additionally, global institutions have discussed developing model registries, dataset documentation, and algorithmic audit mechanisms to support accountability.

Technical tools may also play a role in reform. Scholars have recommended adopting watermarking, metadata tagging, and content verification systems to identify AI-generated content and reduce the spread of misinformation (Yang & Zhang, 2024). Finally, some propose multilateral cooperation, where international bodies develop guidelines rather than enforceable laws, offering a flexible framework for countries to adapt to their local contexts.

The global regulatory environment for AI-generated content remains highly fragmented, shaped by differing legal traditions, cultural values, and policy priorities. While the United States, European Union, China, and India each offer distinct approaches, none provides a fully comprehensive solution to the challenges posed by generative AI. Harmonization is constrained by philosophical, economic, and governance differences, making a single global standard unlikely in the near future.

Nevertheless, meaningful reforms are emerging. Proposals for sui generis rights, improved transparency, licensing obligations, and technical safeguards suggest a path toward more balanced and coherent regulation. As AI technologies continue to evolve, ongoing collaboration among legal scholars, policymakers, technologists, and industry actors will be crucial in shaping a framework that supports innovation while protecting rights holders, individuals, and the public interest.

Chapter 7: Limitations

Although this study provides a broad examination of the legal implications of AI-generated content, several limitations should be acknowledged.

First, the legal and regulatory environment surrounding generative AI is evolving rapidly. Court rulings, policy drafts, and legislative proposals continue to develop, meaning some analyses may become outdated as new interpretations emerge. The findings and perspectives discussed here reflect the status of scholarship and regulatory developments up to the period documented by key sources such as Cambridge University Press (2023) and Zakir (2024).

Second, the technical opacity of modern AI models constrains legal clarity. Generative models operate through complex, high-dimensional processes that are not easily interpretable, even by their creators. This “black-box” nature limits the ability of researchers, courts, and regulators to fully assess issues such as

liability, originality, and infringement (Shah, 2024). The lack of transparent information about training datasets further restricts the precision of legal evaluation.

Third, this research focuses on selected jurisdictions; the United States, European Union, China, and India. While these regions represent significant global legal perspectives, the analysis cannot fully capture the diverse regulatory approaches emerging in other countries. Any conclusions drawn therefore remain context-specific and may not generalize globally.

Fourth, ethical and socio-technical issues intersect deeply with legal questions, particularly regarding bias, consent, and public trust. This study references ethical concerns where relevant (Shah, 2024; Massadeh, 2024), but it does not attempt to fully explore ethical theory or the moral philosophy underlying these debates.

Finally, the study is limited by the availability of publicly accessible information. Proprietary datasets, internal company policies, and confidential regulatory negotiations were beyond the scope of this research. As a result, the analysis is based on academic publications, reported legal cases, and public regulatory documents.

These limitations do not undermine the central findings but highlight the need for ongoing research as AI technologies and their legal contexts continue to evolve.

Chapter 8: Conclusion and Discussion

The rapid advancement of generative AI has introduced significant challenges for legal systems designed with human creativity and authorship in mind. This study demonstrates that existing copyright frameworks struggle to accommodate content produced by models trained on vast and often copyrighted datasets. As noted in recent scholarship, the absence of a human author poses fundamental difficulties for traditional doctrines of originality, ownership, and derivative works (Cambridge University Press, 2023; Shah, 2024). The analysis reveals several key themes. First, authorship remains one of the most contested areas. Current legal standards emphasize human creativity, yet generative AI systems deliberately minimize direct human involvement. Proposed solutions range from assigning authorship to human users, to recognizing joint authorship, to creating new *sui generis* rights specifically for AI-generated works (Zakir, 2024; UPF Thesis, 2023).

Second, the legality of training data continues to be a major source of conflict. Ongoing lawsuits, including *Getty Images v. Stability AI* and *The New York Times v. OpenAI*, illustrate the growing tension between data-intensive model training and the rights of content creators. Divergent interpretations of “transformative use” in the United States and Europe further complicate this landscape.

Third, privacy concerns, particularly those involving personal or biometric data inadvertently embedded in training datasets, extend the debate beyond copyright into broader data protection and compliance obligations (Jones & Kearns, 2024; Zakir, 2024).

Liability remains another unresolved area. Assigning responsibility among developers, platform providers, and end-users is challenging due to the opaque and autonomous nature of AI systems (Shah, 2024). Deepfakes and AI-enabled impersonation add new layers of harm that intersect with rights of reputation, identity, and public trust (Massadeh, 2024).

International comparisons indicate that regulatory harmonization is unlikely in the near future. The United States prioritizes fair use and human authorship, the European Union emphasizes moral rights and data protection, China adopts a state-driven regulatory approach, and India’s legal framework currently lacks

explicit provisions addressing AI. These differences reflect varying cultural, political, and legal philosophies.

Looking ahead, meaningful reform will require hybrid regulatory approaches that integrate copyright modernization, transparency requirements, data protection safeguards, liability standards, and technical governance tools such as watermarking and dataset registries (Yang & Zhang, 2024). As generative AI becomes embedded across creative industries, communication ecosystems, and public decision-making, the need for clear, adaptable legal frameworks becomes increasingly urgent.

In summary, the implications of AI-generated content extend far beyond copyright. They touch on questions of identity, privacy, accountability, and the future of human creativity. Addressing these challenges will require sustained interdisciplinary collaboration among legislators, technologists, legal scholars, and industry stakeholders.

Chapter 9: References

1. Alex Yang, S., & Zhang, A. H. (2024). *Generative AI and copyright: A dynamic perspective*. arXiv. <https://arxiv.org/>
2. AI and intellectual property: Legal frameworks and future directions. (2024). *Law Journal*, 18(2), 44–67.
3. AI-generated content: Legal challenges & potential reforms. (2024). ResearchGate. <https://www.researchgate.net/>
4. Cambridge University Press. (2023). *ChatGPT: A case study on copyright challenges for generative artificial intelligence systems*. <https://www.cambridge.org/>
5. Lu, Y. (2025). *Reforming copyright law for AI-generated content*. TechReg. <https://techreg.org/>
6. Massadeh, F. (2024). *The legal protection of artificial intelligence-generated work*. Virtus Inter Press.
7. Not all similarities are created equal: Leveraging data-driven biases to inform GenAI copyright disputes. (2024). arXiv. <https://arxiv.org/>
8. Shah, R. (2024). *Ethical concerns and legal implications of generative AI in content creation*. ResearchGate. <https://www.researchgate.net/publication/393047792>
9. UPF Thesis. (2023). *Navigating the legal labyrinth: Establishing copyright frameworks for AI-generated content* (Master's Thesis, Universitat Pompeu Fabra). <https://repositori.upf.edu/>
10. Voinea, D. V. (2023). *Ownership of AI-generated content: Who owns what?* SSRN. <https://sserr.ro>
11. Zakir, M. (2024). *Navigating the legal labyrinth: Establishing copyright frameworks for AI-generated content*. ResearchGate. <https://www.researchgate.net/publication/379188943>
12. *Why copyright is not the right policy tool to deal with generative AI*. (2023). *Yale Law Journal Forum*. <https://yalelawjournal.org/>