

CASE AUTHORS V. ANTHROPIC PBC 2025 (S.U.A.) – BETWEEN COPYRIGHT INFRINGEMENT AND „FAIR USE OF AI”

Ileana Ioana BÎLBĂ (POP) *

Abstract

The case Authors v. Anthropic is one of the most relevant cases in the US regarding the use and training of AI with authors' works protected by copyright. This case highlights the enduring tension between technological innovation and creators' rights. In the United States, where judicial rulings carry precedential authority, outcomes such as Anthropic v. Authors are likely to shape the trajectory of copyright law and the ethical use of creative works in AI training, marking a pivotal moment in the enforcement of intellectual property in the era of rapidly evolving technology.

The June 23, 2025, ruling in Bartz v. Anthropic by Judge William Alsup marks a landmark U.S. decision on copyright law as applied to large language model (LLM) training. The court distinguished between the acquisition of copyrighted works and their use in AI training, treating each as a separate „use” under the fair use doctrine. While the mass acquisition of pirated books was deemed „inherently, irredeemably infringing,” the subsequent training of LLMs on lawfully obtained works was held to be highly transformative and generally fair use.

Judge Alsup emphasized that LLMs generate new outputs rather than reproduce original works, and limited memorization does not constitute infringement absent traceable outputs. The court also upheld digital format-shifting of lawfully acquired books for internal AI training.

This ruling provides critical guidance for AI developers, affirming that transformative use is permissible under copyright law while underscoring the legal risks of unlawful acquisition, establishing a pivotal precedent at the intersection of AI and intellectual property.

Keywords: Artificial Intelligence, „fair use” of AI, Authors vs. Anthropic, Copyright.

1. Introduction

The advancement of AI-based technologies has begun to gain momentum, so that tensions have begun to arise regarding the protection of intellectual property rights, and discussions about training algorithms on the works of human authors have reached the courts in the United States. There are serious concerns in academia about the impact of using AI, which are very realistic: „it will compete with humans and even surpass humans, but also that this AI could become more powerful than we would like it to be, with catastrophic consequences for humanity.”¹

In this article, I will analyze to what extent the use of large language models (LLM) qualifies for training artificial intelligence within the limits of current legislation, in the absence of a license or authorization from human authors, using Authors v. Anthropic (2025) as a case study. For comparison, we analyzed the application of the American „fair use” doctrine² which offers greater flexibility in use through legal interpretation, in contrast to the European framework of Directive EU 2019/790,³ which provides greater legal clarity but requires complementary mechanisms, such as collective licensing.

* PhD Candidate in Intellectual Property Law, „Nicolae Titulescu” University of Bucharest, Master's student in Business Administration (MBA) at the UNESCO department, within the University of Bucharest, SEMP scholarship (2nd semester, 2024) at the University of Geneva, Switzerland.

¹ Viorel Roș, Andreea Livădaru, Ciprian-Raul Romițan, *Dreptul proprietății intelectuale. Dreptul de autor și drepturile conexe și drepturile de proprietate industrială*, editura C.H. Beck, București, 2025, p. 10.

² 17 U.S. Code § 107 - Limitations on exclusive rights: Fair use: “the fair use of a copyrighted work, including such use by reproduction in copies or phonorecords or by any other means specified by that section, for purposes such as criticism, comment, news reporting, teaching (including multiple copies for classroom use), scholarship, or research, is not an infringement of copyright.” https://www.law.cornell.edu/uscode/text/17/107#:~:text=%20Title%2017.%20*%20CHAPTER%201.%20*%20C2%A7%20107, accessed on 2nd of march 2026.

³ Directive (EU) 2019/790 of the European Parliament and of the Council of 17 April 2019 on copyright and related rights in the Digital Single Market and amending Directives 96/9/EC and 2001/29/EC (Text with EEA relevance.) <https://eur-lex.europa.eu/eli/dir/2019/790/oj?locale=ro>, accessed on 2nd of march 2026.

Systems that use large linguistic models (LLM) rely on huge amounts of data, whether textual, visual or audiovisual, in order to generate results that are as relevant and close to expectations as possible, but this process raises big questions about the protection of intellectual property rights, especially when these systems are trained on copyrighted works.

Some situations have generated a wave of litigation in the United States, and courts have been called upon to rule on whether the use of large linguistic models (LLMs) and the training of artificial intelligence algorithms on copyrighted works violate the American doctrine of fair use. One such case is *New York Times v. OpenAI* ⁴ in which the famous newspaper accused OpenAI of having trained its algorithms on millions of newspaper articles without a license, and more than that, having memorized and reproduced parts of the articles, and in many cases the reproductions were wrong, affecting both the work of journalists and the reputation of the newspaper, because OpenAI, by reproducing NYT articles, enters into direct competition with the newspaper, but without paying for these reproductions.⁵

In the case of *Getty Images v. Stability AI*⁶, a UK judge ruled in November 2025 that Getty had partially succeeded against Stability AI in trademark infringement claims regarding artificial intelligence results that displayed distorted Getty Images watermarks. This case represented a significant development in emerging jurisprudence in the intersection of artificial intelligence and intellectual property law. The judgment did not provide a definitive resolution to the complex legal debates regarding AI training and copyrights infringements but underlines the structural limitations of current intellectual property law frameworks when applied to machine learning structures trained on transnational dataset with unprecedented scale and complexity. Terms like authorship, ownership, originality, creative labour and technological mediation becomes evolutionary concepts and which will need to be redefined in the future and contain broader frameworks to allow for the integration of technological developments.

2. Case Authors v. Anthropic PBC 2025 (S.U.A.)

„Authors vs. Anthropic” is a class action lawsuit and landmark settlement from 2025 regarding the use of copyrighted books to train generative artificial intelligence models. The case, filed by authors Andrea Bartz, Charles Graeber, and Kirk Wallace Johnson, highlighted the legal, ethical, and financial conflicts between artificial intelligence developers and content creators. The authors alleged that Anthropic violated Copyright Law by downloading pirated copies of their books from “shadow libraries” (LibGen, PiLiMi) to train its artificial intelligence system Claude, without permission or compensation.

The plaintiffs, a group of published authors, filed a class action lawsuit against Anthropic, alleging that their works were illegally copied for use in the teaching of Anthropic, an LLM, alleging three counts: that copyrighted works were used to teach LLMs, that books that were legally purchased in print were digitized, and that pirated digital copies of the books were purchased and retained to build a digital library. Anthropic argued that its uses constituted fair use under Section 107 of the U.S. Copyright Act.⁷

The court ruled on June 23, 2025 and found three conclusions relevant to the case of „Authors vs. Anthropic”:

- Using copyrighted works to train AI algorithms constitutes „fair use” because it is „transformative”, with the results obtained from using the AI being different from reproducing the work. The Court drew an analogy to the human learning process, and Copyright Law does not extend to methods, concepts, or principles contained in a work.
- The digitization of books legally purchased by Anthropic is fair use, even if some works were destroyed during scanning, the digital result replaces the printed work, being transformative in the sense of fair use, because it facilitates the storage of the work, without increasing the number of duplicates and without having been distributed outside the Anthropic company.

⁴ *New York Times v. OpenAI* https://nytco-assets.nytimes.com/2023/12/NYT_Complaint_Dec2023.pdf, accessed on 5th of march 2026.

⁵ A deeper analysis in Harvard Law Review: *NYT v. OpenAI: The Times’s About-Face* by Audrey Pope: <https://harvardlawreview.org/blog/2024/04/nyt-v-openai-the-times-about-face/> accessed on 5th of march 2026.

⁶ *Getty Images v. Stability AI* <https://www.judiciary.uk/wp-content/uploads/2025/11/Getty-Images-v-Stability-AI.pdf>, accessed on 5th of march 2026.

⁷ Regina Sam Penti, *From Books to Bots: Key Takeaways from the Anthropic Fair Use Decision for AI Developers and Copyright Holders*, <https://www.ropesgray.com/en/insights/alerts/2025/06/from-books-to-bots-key-takeaways-from-the-anthropic-fair-use-decision-for-ai-developers>. accessed on 9th of march 2026.

- The acquisition and storage of pirated works does not constitute fair use, because they were not reasonably necessary for LLM training, especially since copyright-compliant alternatives existed.⁸

The Court's decision states that authors should be more proactive in protecting their works and take action when pirated digital libraries appear in order to better protect their works and prevent language programs from doing unauthorized training. Three authors have filed a class action lawsuit, which will benefit over half a million authors in the US, who will receive over \$1.5 billion, a quarter of which will be in attorneys' fees.⁹

3. Technical elements of generative artificial intelligence

Generative artificial intelligence models are trained on large data sets through a process that involves tokenization, statistical learning, and optimization of model parameters. During training, the system processes large amounts of data to identify patterns and relationships between words, images, or other forms of content. Tokenization refers to raw data (text, images) being converted into smaller units called "tokens" (words, sub-words, or pixels), which can be numerically processed by a computer.

The model analyzes the probabilities of co-occurrence of these tokens. It learns that a particular word or pixel tends to follow another in a given context, identifying syntactic, semantic, or visual structures. During training, the model continuously adjusts its internal parameters (weights) to minimize the difference between its prediction and the actual data. This process is repeated billions of times until the model "understands" the data representations.¹⁰

From a legal perspective, a crucial distinction must be made between input data and output content. This distinction is essential for determining liability, as training involves copying data, while output raises issues of reproduction or derivative works. Although AI systems are designed to generalize, studies have shown that models can occasionally reproduce their training data verbatim, raising concerns about infringement.¹¹

3.1. Technical element in the case Getty Images v. Stability AI (2025)

In the case Getty Images v. Stability AI¹², there is one of we find one of the first legal analyses that evaluates the technical structure of a generative artificial intelligence system from the perspective of current United Kingdom legislation. Stability AI developed a technical artificial intelligence system called Stable Diffusion and trained it using public databases of photos on the internet. The court paid particular attention to the explanation of the mechanism of deep learning and image generation.

Stable Diffusion was presented as a neural-network architecture operating through probabilistic modeling. During the training of the artificial intelligence system, it processed billions of image-text pairs extracted from datasets. The dataset contained URLs, metadata, captions, and associated image references obtained through large-scale web scraping and crawling processes.¹³ Getty Images alleged that millions of copyrighted visual assets originating from its databases were included in these datasets and subsequently materialized during model training.¹⁴ Technically, the judgment distinguished between "training" and "inference." Training involved repeated exposure of the neural network to visual data in order to optimize "model weights," namely the adjustable parameters controlling the system's internal statistical representations. By contrast, inference referred to the operational stage in which a trained model generates synthetic images from textual prompts by iteratively removing noise from an initial random image.¹⁵ The court emphasized that Stable Diffusion does not

⁸ *Idem*.

⁹ Dave Hansen, The Bartz v. Anthropic Settlement: Understanding America's Largest Copyright Settlement, <https://legalblogs.wolterskluwer.com/copyright-blog/the-bartz-v-anthropic-settlement-understanding-americas-largest-copyright-settlement/> accessed on 15th of March 2026.

¹⁰ Stanford University: Reexamining "Fair Use" in the Age of AI, <https://hai.stanford.edu/news/reexamining-fair-use-age-ai>, accessed on 7th of march 2026.

¹¹ *Idem*.

¹² Getty Images v. Stability AI <https://www.judiciary.uk/wp-content/uploads/2025/11/Getty-Images-v-Stability-AI.pdf>, accessed on 7th of march 2026.

¹³ Britannica, *Stable Diffusion – generative AI*, <https://www.britannica.com/technology/Stable-Diffusion>, accessed on 7th of march 2026.

¹⁴ Getty Images v. Stability AI <https://www.judiciary.uk/wp-content/uploads/2025/11/Getty-Images-v-Stability-AI.pdf>, accessed on 7th of march 2026.

¹⁵ *Idem*.

retain literal copies of training images within the model itself; however, its functional capacities remain indirectly shaped by the training data embedded within the learned parameters.

A further technical issue concerned the appearance of synthetic Getty Images watermarks in AI-generated outputs. Getty argued that these outputs demonstrated the internalization of protected visual indicators and trade marks during training. Stability AI accepted that certain prompts could generate images containing distorted Getty watermarks but contended that such outputs resulted from user manipulation rather than autonomous conduct attributable to the company itself.

The litigation also addressed distributed cloud-computing infrastructures. Evidence demonstrated that training processes occurred primarily on Amazon Web Services (AWS) clusters located outside the United Kingdom. This technical-geographical dimension proved legally significant because Getty ultimately abandoned parts of its copyright infringement claims after acknowledging insufficient evidence that model training occurred within UK territorial jurisdiction.

3.2. Technical element in the case *Authors v. Anthropic PBC* 2025

The litigation in *Authors v Anthropic PBC* is a significant judicial examination of the technical infrastructure underlying large language models and their interaction with United States copyright law. The dispute focused on the development and training of Anthropic's AI systems, particularly the use of extensive literary datasets allegedly acquired from both lawful and unauthorized digital repositories.¹⁶

From a technical perspective, the proceedings addressed the operational architecture of transformer-based language models. Anthropic's systems were trained through large-scale computational ingestion of digitized books contained in datasets originating from repositories such as LibGen and PiLiMi. The court distinguished between: the acquisition and storage of copyrighted works, the preprocessing and organization of datasets; and the subsequent machine-learning training process.¹⁷ The judgment emphasized that the defendant maintained a centralized digital "library" composed of millions of books copied into machine-readable formats. Technically, this repository enabled: text extraction, tokenization, dataset filtering, deduplication, indexing, and repeated retraining operations for LLM optimization.

The court examined how transformer models process textual corpora through probabilistic computation. During training, textual material is converted into numerical tokens and processed iteratively across neural-network layers to optimize statistical parameters governing linguistic prediction. Similar to diffusion-image systems, the model does not retain directly accessible copies of complete books; instead, it develops generalized semantic and syntactic relationships through parameter weighting and predictive optimization. A central technical issue concerned the distinction between: permanent storage of pirated literary works, and transformative computational use during model training.

Judge William Alsup recognized that reproductions created specifically for machine-learning training could constitute transformative uses because the system utilized textual material to develop generalized linguistic capabilities rather than to reproduce expressive literary content verbatim.¹⁸ However, the court treated differently Anthropic's large-scale downloading and retention of unauthorized digital books from pirate repositories. The maintenance of a persistent internal library was characterized as a distinct act of reproduction potentially constituting copyright infringement independently of the later AI-training process.¹⁹

Another technical aspect involved evidentiary questions surrounding "memorization" and output generation. Plaintiffs alleged that LLMs may reproduce protected passages when prompted. Anthropic argued that outputs are generated probabilistically through learned statistical relationships rather than through retrieval of stored textual copies. The litigation therefore highlighted the growing doctrinal distinction between: latent statistical representation, memorization, and direct reproduction of copyrighted expression. The settlement

¹⁶ Bartz et al v. Anthropic PBC, No. 3:2024cv05417 - Document 437 (N.D. Cal. 2025), <https://law.justia.com/cases/federal/district-courts/california/candce/3:2024cv05417/434709/437/>, accessed on 9th of March 2026.

¹⁷ *Idem*.

¹⁸ Regina Sam Penti, From Books to Bots: Key Takeaways from the *Anthropic* Fair Use Decision for AI Developers and Copyright Holders, <https://www.ropesgray.com/en/insights/alerts/2025/06/from-books-to-bots-key-takeaways-from-the-anthropic-fair-use-decision-for-ai-developers>, accessed on 9th of March 2026.

¹⁹ Bartz et al v. Anthropic PBC, No. 3:2024cv05417 - Document 437 (N.D. Cal. 2025), <https://law.justia.com/cases/federal/district-courts/california/candce/3:2024cv05417/434709/437/>, accessed on 9th of March 2026.

framework also reflected the scale of AI data infrastructures. The proceedings involved nearly half a million literary works and extensive digital evidence, including terabytes of computational and dataset information.

The Court's Functional Understanding of AI Training

A central aspect of the court's reasoning was its distinction between: the storage of copyrighted works, and the transformative computational use of those works during AI training.²⁰

Judge William Alsup adopted a functional interpretation of machine-learning processes. The court recognized that transformer-based LLMs do not operate as digital archives reproducing books in a conventional sense. Instead, the system processes textual corpora through tokenization, parameter optimization, and probabilistic modelling to generate generalized linguistic capabilities.²¹

Technically, this reasoning demonstrates judicial acceptance of the distinction between: expressive content; and statistical abstraction. The court implicitly acknowledged that model weights and learned parameters constitute mathematical representations derived from training data rather than literal textual reproductions. This reflects a sophisticated understanding of neural-network architecture rarely observed in earlier copyright jurisprudence.

3.2.1. Separation Between Dataset Piracy and AI Training

One of the most important technical distinctions in the judgment concerns the differentiation between: unlawful acquisition of datasets; and the subsequent machine-learning training process. The court treated Anthropic's alleged downloading of books from repositories such as LibGen and PiLiMi as technically and legally distinct from the later use of those texts in model training.²²

This distinction is doctrinally important because the court effectively fragmented the AI pipeline into separate legally assessable stages: ingestion, storage, preprocessing, training, and output generation. By isolating these stages, the decision avoids collapsing the entire AI lifecycle into a single act of infringement. Instead, liability becomes dependent upon the specific computational function being performed at a given stage. From a technological perspective, the court demonstrated awareness that modern AI systems depend upon persistent internal repositories used for: deduplication, indexing, filtering, retraining, and dataset management.²³

3.2.2. Judicial Recognition of Transformative Computation

A particularly innovative element of the reasoning lies in the court's treatment of AI training as potentially "transformative"²⁴ under U.S. fair use doctrine. The court accepted that LLM training serves a fundamentally different function from the original literary works. Rather than communicating expressive narratives to readers, the computational process extracts statistical relationships, semantic patterns, and linguistic structures to develop predictive text-generation capabilities.

This reasoning extends earlier jurisprudence concerning: search indexing, digital scanning, and computational analysis, particularly the logic developed in *Authors Guild v Google*. However, the court did not fully immunize AI systems under fair use. Instead, it imposed a conceptual limit: transformative computational use may exist even where the underlying acquisition of training data remains unlawful. This creates a dual analytical framework: lawful transformation, coexisting with potentially unlawful procurement.

4. Conclusions

The rapid expansion of generative artificial intelligence has exposed a fundamental structural tension within copyright law: the need for widespread access to data to enable technological innovation versus the imperative to protect the economic and moral interests of authors. The case of *Authors v. Anthropic* illustrates

²⁰ *Idem*.

²¹ *Idem*.

²² Hansen, Dave, *The Bartz v. Anthropic Settlement: Understanding America's Largest Copyright Settlement*, <https://legalblogs.wolterskluwer.com/copyright-blog/the-bartz-v-anthropic-settlement-understanding-americas-largest-copyright-settlement/> accessed on 15th of march 2026.

²³ *Bartz et al v. Anthropic PBC*, No. 3:2024cv05417 - Document 437 (N.D. Cal. 2025), <https://law.justia.com/cases/federal/district-courts/california/candce/3:2024cv05417/434709/437/>, accessed on 9th of march 2026.

²⁴ *Idem*.

this tension in its most distilled form, focusing specifically on whether training AI systems constitutes unlawful reproduction or permissible transformative use.

In this article, we have shown that, in the US legal framework, the fair use doctrine offers a flexible and adaptable tool, capable – at least in principle – of accommodating AI training practices. In particular, the concept of transformativity allows courts to interpret machine learning as an analytical, non-expressive process, rather than as a direct exploitation of copyrighted works. However, this flexibility comes at the cost of legal uncertainty, especially in cases involving systematic, large-scale data ingestion and potential market substitution effects.

Instead, the European Union takes a more structured and predictable approach through the text and data mining exceptions introduced by Directive (EU) 2019/790. While this framework increases legal certainty, it remains limited by its formalistic nature and reliance on opt-out mechanisms, which may prove insufficient in practice given the scale and opacity of artificial intelligence (AI) training processes.

A comparative analysis of recent litigation—including *Authors v. Anthropic*, *New York Times v. OpenAI*, and *Getty Images v. Stability AI*—reveals the emergence of a continuum of liability, ranging from lawful data analysis to unlawful reproduction. This continuum underscores the inadequacy of binary legal classifications in the context of generative AI and highlights the need for more nuanced regulatory solutions.

Accordingly, I argue that neither the US nor the EU framework, in their current forms, is fully capable of addressing the systemic challenges posed by AI. A purely flexible approach risks undermining authors' rights, while an overly rigid system may stifle innovation. The most viable way forward lies in developing a hybrid model that combines the interpretative flexibility of fair use with the structural clarity of EU-style regulation.

Such a model should include, at a minimum, (i) the introduction of collective or compulsory licensing mechanisms to facilitate legal access to training data, (ii) increased transparency obligations for AI developers regarding the composition of training datasets, and (iii) clearer standards to distinguish between legal training and infringing production.

Ultimately, the evolution of copyright law in the age of artificial intelligence will depend not only on judicial interpretation but also on legislative intervention and international coordination. As generative AI continues to reshape the creative industries, the challenge for legal systems will be to ensure that innovation and authorship remain mutually reinforcing, rather than fundamentally incompatible.

References

- *** U.S. Congressional Research Service – Generative AI and Copyright Law, <https://crsreports.congress.gov/product/pdf/LSB/LSB10922>, accessed on 5th of march 2026.
- Avarvarei, Simona Catrinel, "Inteligența artificială, noul Golem al unui viitor care a început deja?", (2024), 2, Revista română de dreptul proprietății intelectuale.
- Bartz et al v. Anthropic PBC, No. 3:2024cv05417 - Document 437 (N.D. Cal. 2025), <https://law.justia.com/cases/federal/district-courts/california/candce/3:2024cv05417/434709/437/>, accessed on 9th of march 2026.
- Britannica, Stable Diffusion – generative AI, <https://www.britannica.com/technology/Stable-Diffusion>, accessed on 9th of march 2026.
- Budileanu, Cristiana, „Inteligența artificială și cadrul legislativ actual al dreptului de autor. Studiu de caz CHATGPT,” (2024), 2, Revista română de dreptul proprietății intelectuale.
- Coulter, Matthew, Aiming for Fairness: an Exploration into Getty Images v. Stability AI and its Importance in the Landscape of Modern Copyright Law, 34 DePaul J. Art, Tech. & Intell. Prop. L. 124 (2024) Available at: <https://via.library.depaul.edu/jatip/vol34/iss1/4>, accessed on 5th of march 2026.
- Dumitrașcu, Ramona „Drepturi de proprietate intelectuală pentru dezvoltarea tehnologiilor cu inteligență artificială din perspectiva rezoluției Parlamentului European din 20 octombrie 2020,” (2023), 3, Revista română de dreptul proprietății intelectuale.
- Getty Images v. Stability AI <https://www.judiciary.uk/wp-content/uploads/2025/11/Getty-Images-v-Stability-AI.pdf>, accessed on 5th of march 2026.
- Gervais, Daniel J., The Machine As Author (March 25, 2019). Iowa Law Review, Vol. 105, 2019, 2053-2106, Vanderbilt Law Research Paper No. 19-35, Available at SSRN: <https://ssrn.com/abstract=3359524>, accessed on 5th of march 2026.

- Hansen, Dave, The Bartz v. Anthropic Settlement: Understanding America's Largest Copyright Settlement, <https://legalblogs.wolterskluwer.com/copyright-blog/the-bartz-v-anthropic-settlement-understanding-americas-largest-copyright-settlement/> accessed on 15th of march 2026.
- Kissinger, Henry; Eric Schmidt, Daniel Huttenlocher, The Age of AI, John Murray, 2021.
- El Mouttaqui, Hajar and Salama, Feras, Organizational Form and Audit Pricing (June 1, 2023). Available at SSRN: <https://ssrn.com/abstract=4487339> or <http://dx.doi.org/10.2139/ssrn.4487339>, accessed on 5th of march 2026.
- Lemley, M. A., & Casey, B. (2021). Fair learning. Texas Law Review, 99(4), 743-785. <https://texaslawreview.org/fair-learning/>, accessed on 5th of march 2026.
- New York Times v. OpenAI https://nytco-assets.nytimes.com/2023/12/NYT_Complaint_Dec2023.pdf, accessed on 5th of march 2026
- Pentti, Regina Sam, From Books to Bots: Key Takeaways from the Anthropic Fair Use Decision for AI Developers and Copyright Holders, <https://www.ropesgray.com/en/insights/alerts/2025/06/from-books-to-bots-key-takeaways-from-the-anthropic-fair-use-decision-for-ai-developers>. accessed on 9th of march 2026.
- Roș, Viorel; Livădariu, Andreea, Drepturile morale de autor în era inteligenței artificiale. De la inteligența naturală creatoare, la inteligența artificială, în Viorel Roș, Ciprian Raul Romițan, Provocările dreptului de autor. La 160 de ani de la prima lor reglementare legală în România, editura Hamangiu, București, 2022.
- Roș, Viorel; Livădariu, Andreea; Romițan, Ciprian-Raul; Dreptul proprietății intelectuale. Dreptul de autor și drepturile conexe și drepturile de proprietate industrială, editura C.H. Beck, București, 2025.
- Spineanu-Matei, Octavia, "Strategia curții de justiție a Uniunii Europene cu privire la utilizarea inteligenței artificiale," (2024), 3, Revista română de dreptul proprietății intelectuale.
- Stan, Oana-Gabriela, "Considerente cu privire la posibilitatea protejării prin drept de autor a unei creații realizate de inteligența artificială generativă" (2025), 1, Revista română de dreptul proprietății intelectuale.
- Stanciu, Marius „O imagine pixelată asupra încălcării drepturilor de autor în utilizarea inteligenței artificiale,” (2023), 4, Revista română de dreptul proprietății intelectuale.
- Vișoiu, Ruxandra, "Inteligența artificială și drepturile de autor – ce ne rezervă viitorul ?," (2024), 2, Revista română de dreptul proprietății intelectuale.